



Universidad Michoacana de San Nicolás de Hidalgo

Facultad de Ciencias Físico-Matemáticas

**Métodos de Alta Resolución para la solución
numérica de leyes de conservación en
hidrodinámica clásica**

Tesis
que para obtener el título de

Licenciado en Ciencias Físico-Matemáticas

presenta:

Néstor Enrique Ortiz Madrigal

Asesor:
Dr. Francisco S. Guzmán Murillo
Instituto de Física y Matemáticas, UMSNH.

Morelia, Michoacán, Mayo de 2008

Dedicatoria

Este trabajo está dedicado a toda mi familia y a todos mis amigos. En el más alto nivel jerárquico, como siempre y en todo, están mis papás y mi hermana.

Con gusto comparto mi tesis con mis amigos y compañeros de la licenciatura, aquellos con quienes a lo largo de tantas horas de trabajo, hombro con hombro, nos hicimos humildes ante el conocimiento de la naturaleza y maduramos razonamientos a fuerza de rigor matemático, apasionados de su armonía y profunda belleza.

Agradecimientos

Agradezco en su justa proporción a quienes me han formado como científico y a quienes siguen haciéndolo. La influencia más considerable ha sido de los admirables Físicos:

Alberto Mendoza, Rafael G. Campos, y Francisco Guzmán.

Por sus valiosos aportes a este trabajo agradezco a los Doctores: José Antonio González, Rafael G. Campos, Ricardo Becerril, Francisco D. Mota y Mario César Suárez. A mi asesor, gracias por facilitarme los medios y una cantidad importante de conocimientos técnicos y científicos.

Contenido

1	Introducción	7
2	Discretización del espacio-tiempo en Volúmenes Finitos	11
2.1	La discretización en una dimensión	11
2.2	La formulación en Diferencia de Flujos	14
2.3	Reconstrucción de la función primitiva	16
2.4	La noción de Alta Resolución	17
2.5	Reconstrucción lineal	17
3	Métodos de Alta Resolución	19
3.1	Oscilaciones	19
3.2	La condición de Variación Total Decreciente	20
3.3	Limitadores de pendientes	22
3.4	Condiciones de frontera	23
3.5	Convergencia y estabilidad	25
4	Implementación y pruebas con ecuaciones escalares	29
4.1	La ecuación escalar de advección	29
4.1.1	Pseudocódigo para resolver la ecuación escalar de advección	32
4.2	La ecuación de Burgers	35
5	Sistemas hiperbólicos lineales de leyes de conservación	41
5.1	Desacomplamiento del sistema de coeficientes constantes	42
5.2	Generalización del Método de Alta Resolución para sistemas de coeficientes constantes	43
5.2.1	Reformulación del método de Alta Resolución en términos de limitadores de flujos	43
5.2.2	Generalización	46
5.2.3	Implementación	49
5.3	Pruebas con la Ecuación de Onda	49
5.3.1	Desacoplando el sistema	51

5.3.2	Aplicando la generalización del método de Alta Resolución	55
5.4	Sistemas lineales de coeficientes variables	57
5.4.1	Pruebas con la Ecuación de Onda	61
6	Sistemas hiperbólicos no-lineales de leyes de conservación	65
6.1	Discontinuidades en la solución	66
6.2	Cálculo de las velocidades de propagación	66
7	Las Ecuaciones de Euler	69
7.1	Deducción de las leyes de conservación	69
7.2	La forma conservativa del sistema	73
7.3	El Tubo de Choques de Sod	74
7.4	Implementación del método de Alta Resolución y resultados	75
8	Conclusiones	79

Capítulo 1

Introducción

En años recientes el desarrollo de algunas áreas de la física ha sido crucialmente dependiente de los avances en el campo de la computación, el análisis numérico y los métodos que de éste se desprenden para dar solución a los sistemas de ecuaciones diferenciales que modelan ciertos fenómenos naturales.

Una parte fundamental de la física es el modelado teórico de la naturaleza. Los modelos matemáticos que describen, por ejemplo, el trasfondo espacial y temporal donde ocurren los fenómenos físicos, las cantidades ahí medibles o los medios continuos -los sólidos o fluidos-, naturalmente no apelan a la Teoría de Números ni a las Matemáticas Discretas, sino al Análisis Real, lo que permite montar la física y sus cantidades sobre Espacios Métricos, estudiar su Topología y su Geometría. Los modelos llegan a complicarse y generalizarse introduciendo Espacios Métricos tan abstractos como las Variedades Diferenciables, por ejemplo. La idea clásica y fundamental es hacer física en espacios continuos y *suaves* para facilitar el estudio de razones de cambio infinitesimales, ya sea de cantidades propias de la variedad, ya de las cantidades que ahí se miden. Como es bien sabido, lo anterior conduce a la *formulación matemática* de la naturaleza en términos de sistemas de ecuaciones diferenciales, cuyas soluciones describen los fenómenos para todo tiempo y en todo el espacio -al menos idealmente- dados ciertos valores iniciales y condiciones físicas de frontera. Así se lleva a la física a su fin último: entender, modelar y predecir los fenómenos naturales objetivamente.

Formular las ecuaciones diferenciales que describen un fenómeno físico no supone de antemano que la ecuación resultante será soluble de manera exacta, de hecho en la mayoría de los casos se recurre a los llamados Métodos Numéricos: algoritmos computacionales que calculan soluciones aproximadas a las exactas y que resultan ser herramientas útiles para los físicos en su tarea de explorar soluciones de ecuaciones diferenciales bajo particulares situaciones.

En esta Tesis describo la formulación de un tipo especial de Métodos Numéricos, conocidos como Métodos de Volumen Finito. El propósito es mostrar una manera de resolver problemas generales que involucren sistemas no lineales de ecuaciones diferenciales par-

ciales acopladas. El ejemplo paradigmático de este tipo de sistemas es el que rige el comportamiento de fluidos compresibles no viscosos, el cual consiste en cuatro ecuaciones, tres son de balance (conservación de masa, energía y momento lineal) y la cuarta es una ecuación termodinámica de estado. Más adelante se deducirán las leyes de conservación y se justificará la ecuación de estado. Este sistema se conoce como *Ecuaciones de Euler*, en honor al célebre físico y matemático suizo Leonard Euler (1707-1783), quien las dedujo en el siglo XVIII. Se encuentra de inmediato que la solución general de dicho sistema no se ha obtenido por métodos analíticos, y peor aún, que la no-linealidad del sistema provoca que las soluciones desarrollen discontinuidades en un tiempo finito, lo que supone una construcción muy cuidadosa de los métodos numéricos que se aplican.

Valga mencionar que el modelo general y completo de la hidrodinámica no son las Ecuaciones de Euler, sino una versión aumentada (principalmente considerando esfuerzos de corte y fuerzas externas actuando sobre el fluido), llamada *Ecuaciones de Navier-Stokes*, en honor a Claude-Louis Navier (1785-1836) y George Gabriel Stokes (1819-1903), francés e irlandés, respectivamente. Las Ecuaciones de Navier-Stokes son mucho más complicadas que las de Euler, e implican un trabajo numérico todavía más cuidadoso debido a la cantidad de términos acoplados que involucran y a las no-linealidades que presentan. Incluso hay que decir que existe un problema abierto muy importante en torno a este sistema: probar que tiene una solución de flujo laminar para todo tiempo dadas condiciones iniciales de flujo laminar. Este enunciado forma parte de los *siete problemas del milenio*, los cuales han sido elegidos por el *Instituto Clay de Matemáticas* en Cambridge, Massachusetts, que premiará con un millón de dólares a quien logre resolver al menos uno de ellos. De tal magnitud es la consideración que se tiene a las ecuaciones que gobiernan el movimiento de los fluidos sin carga.

La principal complicación al tratar de resolver numéricamente las Ecuaciones de Euler es la formación de discontinuidades en la solución, incluso para datos iniciales suaves, esto debido a la no linealidad del sistema. Este problema complica mucho la implementación de métodos basados en diferencias finitas porque cerca de las discontinuidades, las aproximaciones usuales a las derivadas de las funciones divergen y provocan oscilaciones espurias. Sin embargo se puede afrontar el problema usando diferencias finitas, añadiendo un término que suavice la solución cuando ésta tienda a diverger. El problema es que se trata de un artificio que físicamente significaría trabajar con fluidos viscosos. El método es conocido como *diferencias finitas con viscosidad artificial* y aunque es relativamente fácil de implementar, no siempre es satisfactorio.

Lo que se busca es resolver el problema de las discontinuidades sin la necesidad de agregar términos a las ecuaciones, evadir el uso de viscosidad artificial e implementar un método que a pesar de las dificultades, devuelva la solución físicamente correcta. Actualmente se ha desarrollado una importante cantidad de métodos altamente complejos y que tienen las bondades requeridas. Generalmente están basados en una discretización del espacio-tiempo en pequeñas celdas de volumen finito (de ahí el nombre citado anteriormente), lo que implica profundos cambios conceptuales respecto a la discretización clásica usada en

diferencias finitas, pretendiendo ahora resolver un sistema de ecuaciones integrales y no diferenciales, hecho que tiene importantes ventajas, pues la formulación integral de las ecuaciones es menos rigurosa en cuanto a propiedades de continuidad y diferenciabilidad de las soluciones.

El tipo de métodos que tienen la propiedad de considerar con mayor cuidado las vecindades de las discontinuidades, mientras que las partes suaves son tratadas de manera *convencional*, son propiamente llamados *de Alta Resolución*, y hay varios tipos de ellos. Puntualizando, en esta tesis nos dedicaremos a la formulación de un Método de Alta Resolución particular que servirá para resolver sistemas de leyes de conservación no lineales y acoplados, sin tener que recurrir a la disipación numérica.

En el camino hacia las Ecuaciones de Euler se estudian diversas ecuaciones más sencillas. Primero se discretiza del espacio-tiempo en volúmenes finitos (Capítulo 2), enseguida se presenta la formulación del método (capítulo 3) y su implementación para la solución de ecuaciones escalares, tomando los casos particulares de la ecuación de advección y el problema no-lineal de la Ecuación de Burgers (Capítulo 4). En el quinto capítulo se abordan sistemas lineales escritos en forma conservativa, probando el método con la ecuación de onda. En el sexto capítulo se generaliza la teoría para sistemas no-lineales y se prueba otra vez con la ecuación de onda. En el séptimo capítulo se deducen las Ecuaciones de Euler y se resuelven para datos iniciales discontinuos. Al final se dan algunas conclusiones.

Capítulo 2

Discretización del espacio-tiempo en Volúmenes Finitos

2.1 La discretización en una dimensión

La discretización en volúmenes finitos que a continuación se describe, en conjunto con los métodos de Alta Resolución, son las partes fundamentales de la estrategia a seguir para resolver sistemas no-lineales de leyes de conservación. Veremos enseguida que los métodos de volumen finito se basan en la formulación integral de las ecuaciones, y por lo tanto guardan una comunión más cercana con la física que describen. Más adelante se discutirá ese hecho.

Es oportuno mencionar que en este trabajo sólo se tratan casos unidimensionales. Primero describo la discretización clásica que se utiliza para trabajar con métodos de diferencias finitas, en base a ello por fin defino los harto anunciados volúmenes finitos. Sea D un dominio espacio-temporal finito con una coordenada espacial x y una temporal t . Denótese el dominio espacial por D_x . Por cuestiones computacionales es imposible tener un espacio-tiempo continuo, por lo que debe *discretizarse* en el sentido de que sólo tendremos un número finito N_x de puntos en el espacio y N_t puntos en el tiempo, es decir, $Card(D_x) = N_x$ y $Card(D_t) = N_t$, donde a un tiempo dado, una función arbitraria $q : D_x \rightarrow \mathbb{R}$ tomará sus valores. Cada punto de esta *mallla espacial* está separado de sus vecinos por una distancia espacial Δx , de tal manera que el i -ésimo punto se localiza a una distancia $x_i = i\Delta x$ de la frontera izquierda del dominio. Las fronteras se definen como los puntos x_0 a la izquierda y x_N a la derecha. Si de principio se imponen las fronteras y se pide un número arbitrario de puntos N_x en la malla, fácilmente se calcula la llamada *resolución* de ésta: $\Delta x = \frac{x_N - x_0}{N_x}$. La discretización de la coordenada temporal t se hace de la misma manera, tomando N_t instantes en el tiempo, separados por un intervalo Δt de tal manera que el n -ésimo instante está dado por $t^n = n\Delta t$.

Lo que realmente no hay que perder de vista es que para un tiempo dado, la función q

asociada a una ecuación diferencial y definida en D_x , solamente puede tomar valores en los N_x puntos que han sido definidos. Cuando q evoluciona en el tiempo (al satisfacer una ecuación diferencial que contiene derivadas temporales, por ejemplo), un instante después, sus valores en cada punto del espacio habrán cambiado, pero sigue estando definida sólo en los puntos de la malla. Por supuesto que el hecho de que el tiempo esté discretizado implica que la evolución se dará en pequeños saltos. Los valores que la función toma en un tiempo t^n en un punto x_i serán denotados por q_i^n .

Por otro lado, hay que subrayar que la uniformidad de la malla, es decir, que los puntos en el espacio y el tiempo estén separados entre ellos por una distancia constante, implica que $\Delta x = x_{i+1} - x_i$ y $\Delta t = t^{n+1} - t^n$, para cualesquiera n e i , son constantes.

Hasta ahora se ha construido una malla que se usa comunmente para trabajar con métodos de diferencias finitas. Para obtener una discretización en volúmenes finitos se consideran parejas de los valores que toma la función q en puntos vecinos de la malla, luego se toma su promedio de la siguiente forma:

$$Q_i^n \equiv \frac{1}{\Delta x} \int_{x_i}^{x_{i+1}} q(x, t^n) dx. \quad (2.1)$$

Es claro que sobre un dominio discreto como D_x , dicho valor es:

$$Q_i^n = \frac{1}{\Delta x} \left(\frac{\Delta x q_i^n + \Delta x q_{i+1}^n}{2} \right) = \frac{q_i^n + q_{i+1}^n}{2}. \quad (2.2)$$

El concepto importante a introducir es el de *celda de volumen finito*, comunmente llamada solamente *celda* y definida por el producto cartesiano (ver Figura 2.1):

$$C_i^n \equiv [x_i, x_{i+1}] \times [t_n, t_{n+1}].$$

La parte espacial de la celda se define por $C_i = [x_i, x_{i+1}]$, a la cual buscamos asociar su correspondiente promedio Q_i^n a un tiempo dado t^n mediante $Q^n(C_i) = Q_i^n$, lo que se ilustra en la Figura 2.1 y explícitamente se escribe:

$$Q^n(C_i) = \frac{1}{\Delta x} \int_{C_i} q(x, t^n) dx. \quad (2.3)$$

En principio se tenía un mapeo en el espacio $x - q$ y después se estableció otro en $C - Q$. En la práctica es preferible trabajar en el espacio $x - Q$, empresa que no resulta del todo complicada puesto que si Q_i^n es el promedio de cantidades definidas en x_i y x_{i+1} , basta con asociarle una posición justo en $x_{i+\frac{1}{2}}$, como lo muestra implícitamente la Figura 2.1 y explícitamente la Figura 2.2, que de paso exhibe puntos x y celdas en una misma recta, representando cada celda como un intervalo semiabierto. Más adelante se hará uso de esta interpretación. Por otro lado es importante notar que no se tienen N_x celdas, sino $(N_x - 1)$, lo cual jugará un papel muy importante cuando se impongan condiciones de frontera.

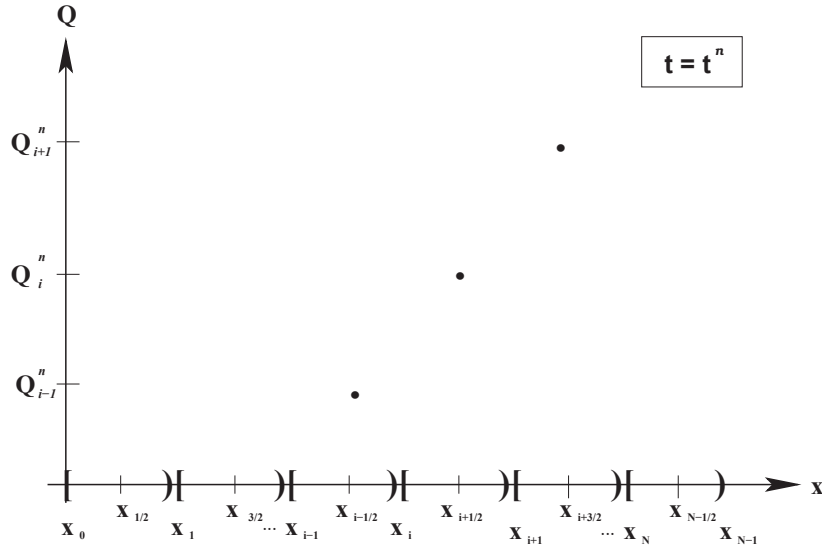


Figura 2.2: En el espacio $x - Q$ el promedio Q_i^n se asocia a la posición $x_{i+1/2}$.

2.2 La formulación en Diferencia de Flujos

Para mostrar la utilidad de la construcción que se acaba de hacer, considérese una ecuación escalar general para q (en una dimensión) de la forma:

$$\partial_t q + \partial_x f(q) = 0, \quad (2.4)$$

donde $f(q)$ es una función real de q , conocida como *función de flujo*, o simplemente *el flujo*. El nombre se ha heredado de términos usados comúnmente en varias ramas de la física, como el electromagnetismo, la termodinámica o la hidrodinámica, donde se obtienen ecuaciones de este tipo (en general no-homogéneas) para describir fenómenos de transporte de carga, de calor, de masa, etcétera, lo que de antemano supone un balance y conservación de la cantidad transportada. Un ejemplo clásico es la ecuación de balance de masa en una dimensión, también conocida como *ecuación de continuidad*, que en ausencia de fuentes o sumideros de materia se escribe como:

$$\partial_t \rho + \partial_x (\rho u) = 0, \quad (2.5)$$

donde ρ denota la densidad volumétrica de masa y u la velocidad de transporte a través de una superficie transversal al movimiento. La cantidad ρu tiene un sentido físico fundamental: es un flujo superficial y temporal de masa y tiene unidades de masa por unidad de área por unidad de tiempo. Luego la ecuación (2.5) tiene la misma forma que (2.4) con $q = \rho$ y

$f(\rho) = \rho u$. Esto ilustra el nombre de la función f .

De hecho la ecuación (2.5) es una de las Ecuaciones de Euler, las cuales se deducirán formalmente más adelante. Por ahora valga decir que la anterior es una *ley de conservación* en el sentido de que establece una rigurosa identidad entre la razón temporal de cambio de la densidad de masa y la razón espacial de cambio del flujo de la misma, es decir, sin fuentes ni sumideros, cualquier cambio en el tiempo de la densidad debe ser estrictamente compensado por un cambio del flujo en el espacio, lo que alude a la conservación de la materia.

Todas las ecuaciones de conservación tienen la forma de (2.4), es por ello que ésta se llama *ecuación escalar general escrita en forma conservativa*. Es importante decir que tales expresiones no se obtienen directamente de la física, sino que resultan de oportunas aplicaciones del Teorema de Stokes (que en alguna de sus versiones redundan en el Teorema de la Divergencia, por ejemplo) sobre ecuaciones integrales que surgen de primeros principios aplicados sobre volúmenes fijos (llamados *volúmenes de control*) que encierran pequeñas porciones del medio continuo.

Para pasar de las ecuaciones integrales a las diferenciales es necesario asumir que las funciones son continuas y diferenciables en todo volumen de control, para todo tiempo. Entonces es claro que la forma integral presenta una conexión más básica con la física y exige menos propiedades matemáticas a la solución, por lo que resulta ser una alternativa muy atractiva para abordar el problema de encontrar numéricamente la solución de ecuaciones cuyas soluciones presentan discontinuidades.

Volviendo a la ecuación (2.4) e integrándola sobre la i -ésima celda se obtiene:

$$\begin{aligned} & \int_{C_i^n} (\partial_t q + \partial_x f(q)) dx dt \\ &= \int_{t^n}^{t^{n+1}} \int_{C_i} \partial_t q dx dt + \int_{t^n}^{t^{n+1}} \int_{C_i} \partial_x f(q) dx dt \\ &= \int_{C_i} [q(x, t^{n+1}) - q(x, t^n)] dx + \int_{t^n}^{t^{n+1}} [f(q(x_{i+1}, t)) - f(q(x_i, t))] dt = 0. \end{aligned} \quad (2.6)$$

Introduciendo (2.1) y definiendo la *función de flujo numérico* como el promedio temporal del flujo a través de la superficie $x = x_i$ como:

$$F_i^n = \frac{1}{\Delta t} \int_{t^n}^{t^{n+1}} f(q(x_i, t)) dt, \quad (2.7)$$

se tiene $\Delta x (Q_i^{n+1} - Q_i^n) + \Delta t (F_{i+1}^n - F_i^n) = 0$ y así escribimos una de las fórmulas más importantes y útiles de esta exposición: la llamada *forma en diferencia de flujos* de (2.4):

$$Q_i^{n+1} = Q_i^n - \frac{\Delta t}{\Delta x} (F_{i+1}^n - F_i^n) . \quad (2.8)$$

El hecho de haber partido de una ecuación escrita en forma conservativa deviene en una importante propiedad del método: si sumamos $\Delta x Q_i^{n+1}$ sobre cualquier conjunto de celdas desde I hasta J :

$$\Delta x \sum_{i=I}^J Q_i^{n+1} = \Delta x \sum_{i=I}^J Q_i^n - \Delta t (F_{J+1}^n - F_I^n) , \quad (2.9)$$

la suma sobre las diferencias de flujos se anula en todo el espacio excepto en las fronteras, lo que nos lleva a sospechar que tendremos una conservación exacta de flujos en todo el dominio si ponemos condiciones de frontera adecuadas. Con ello en mente se dice que el método es conservativo, lo cual se tendrá siempre que la ecuación se escriba en la forma (2.4). Sin embargo hay una razón más importante para escribir las ecuaciones en forma conservativa, y es que si en cambio las escribimos en forma ligeramente modificada (cuasilineal), derivando el flujo previamente para obtener algo como: $\partial_t q + g(q)\partial_x q$, por ejemplo, se producen errores en la solución cerca de las discontinuidades. Esto es así porque las soluciones son equivalentes para ambas formas de escribir la ecuación pero exclusivamente en las regiones suaves y no en puntos donde hay discontinuidades.

2.3 Reconstrucción de la función primitiva

Consideremos nuevamente la ecuación (2.8). Las particularidades del método resultante dependen de la manera en la que se calculan los flujos numéricos F_i^n y F_{i+1}^n . Estos flujos están definidos justo en las interfaces de celdas vecinas. Irónicamente ahora no tenemos información de la cual echar mano para calcular los flujos en los puntos donde originalmente había estado definida la función q . El siguiente objetivo es entonces reconstruir dicha función q en cada punto. Una suposición razonable es que sus valores deben depender de los datos que sí conocemos, es decir, de los promedios Q^n cercanos a la i -ésima interface, donde F^n está definido.

La estrategia a seguir para reconstruir la información necesaria es la siguiente: considérese la parte espacial de la i -ésima celda C_i como un intervalo real y definamos ahí una función $\tilde{q}_i^n$ en términos de al menos Q_{i-1}^n , Q_i^n y Q_{i+1}^n a un tiempo dado t_n . A sabiendas de que Q^n depende de x , escribimos: $\tilde{q}_i^n = \tilde{q}_i^n(x, t_n) \quad \forall x \in C_i$. De esta manera la frontera izquierda de la celda tendrá asociado un valor numérico. Si enseguida *pegamos* cada una de las $\tilde{q}_i^n$ tendremos una función continua a trozos $\tilde{q}^n$ definida en todo el eje x , que naturalmente se llama *función reconstruida*.

El tipo de reconstrucción que se elija caracterizará el método numérico. De entre todas las posibles opciones para $\tilde{q}_i^n$, la más sencilla es hacerla una función constante con valor Q_i^n a

lo largo de toda la celda C_i , es decir, $\tilde{q}_i^n = Q_i^n \ \forall x \in C_i$. Entonces resolver el problema de valores iniciales se logra usando *alguno* de los métodos existentes para resolver un Problema de Riemann (ver [2]) en cada interface entre celdas vecinas y en cada paso temporal, y así calcular los flujos numéricos y después hacer la evolución con algún método integrador en el tiempo, por ejemplo el de Runge-Kutta con una discretización de Método de Líneas. Esta idea fue propuesta originalmente por Godunov en 1959 [4]. El llamado *Método de Godunov* tiene precisión a primer orden y no será discutido aquí en detalle, pues resulta ser un caso muy particular del método que se expone enseguida.

2.4 La noción de Alta Resolución

Decimos que un método es de *Alta Resolución* cuando es capaz de identificar las partes no-suaves de la solución y usar métodos especiales mayor en la región cercana, procurando devolver la solución física en todo tiempo sin la presencia de alguna oscilación espuria. Lo anterior tiene un costo: la precisión del método generalmente pierde un orden en esas regiones (y sólo ahí), lo cual se considera un buen precio a pagar a cambio de un adecuado y consistente tratamiento de las regiones problemáticas que con diferencias finitas, por ejemplo, son infranqueables.

Un adjetivo que comunmente se asocia a los métodos de Alta Resolución es el de *Capturador de Choques* (Shock-Capturing, en inglés), haciendo alusión a la propiedad que tienen de identificar y *atrapar* los *frentes de choque*, que es como en el ámbito de la hidrodinámica se conoce a las discontinuidades en general. En el capítulo sexto se profundiza al respecto y se hace una clasificación de los frentes de choque. Por ahora sólo es necesario decir que usualmente la literatura moderna se refiere a estos métodos como High Resolution Shock Capturing (HRSC) Methods.

También es necesario mencionar que hay distintas clases de métodos de Alta Resolución, siendo algunos de los más populares el *Resolvedor de Riemann de Roe* [7] o la *Fórmula para el Flujo de Marquina* [7], [9]. Los anteriores se basan en estrategias para resolver el Problema de Riemann en cada celda, mientras que el método que se discute en esta Tesis ataca el problema desde otra perspectiva, la cual se introduce en la siguiente sección.

2.5 Reconstrucción lineal

Nuestro método toma una vertiente distinta a la de los resolvedores de Riemann al dar mayor importancia a la manera en que se elige la función que reconstruye los datos: $\tilde{q}_i^n$. Otra elección particular, relativamente sencilla, y que además da resultados más exactos (a segundo orden) que los del Método de Godunov es una función lineal en x , es decir, una recta con alguna pendiente σ_i^n en términos de los datos conocidos. Considérese lo

siguiente:

$$\tilde{q}_i^n = Q_i^n + \sigma_i^n(x - \bar{x}_i), \quad x_i \leq x < x_{i+1}, \quad (2.10)$$

donde

$$\bar{x}_i = \frac{1}{2}(x_i + x_{i+1}) = \frac{1}{2}(2x_i + x_{i+1} - x_i) = x_i + \frac{\Delta x}{2} \equiv x_{i+\frac{1}{2}}.$$

Es fundamental que el promedio sobre C_i de la función reconstruida sea exactamente Q_i^n , además al evaluar $\tilde{q}_i^n$ a la mitad del intervalo entre x_i y x_{i+1} debiéramos obtener también Q_i^n . Ambas propiedades son de esperarse debido a la construcción que hemos hecho:

$$\tilde{q}_i^n(\bar{x}_i, t^n) = Q_i^n = \frac{1}{\Delta x} \int_{C_i} \tilde{q}^n(x, t^n) dx.$$

Analísese ahora la pendiente: al elegir $\sigma_i^n = 0$, la función $\tilde{q}_i^n$ se reduce a una constante a trozos (Método de Godunov). En cambio para tener una pendiente que aproxime la derivada $\partial_x q$ en cada celda, las opciones son:

$$\begin{aligned} \sigma_i^n &= \frac{Q_{i+1}^n - Q_{i-1}^n}{2\Delta x}, \\ \sigma_i^n &= \frac{Q_i^n - Q_{i-1}^n}{\Delta x}, \\ \sigma_i^n &= \frac{Q_{i+1}^n - Q_i^n}{\Delta x}. \end{aligned} \quad (2.11)$$

La pendiente (2.11) corresponde a una diferencia centrada y da lugar al *Método de Fromm*; (2.12) es una diferencia adelantada y esa pendiente corresponde al *Método de Beam-Warming*, mientras que (2.12), una diferencia atrasada genera un esquema llamado *Método de Lax-Wendroff*. Un análisis de las propiedades de cada uno de los métodos se encuentra en [1]. En el siguiente capítulo se muestra que la elección particular de alguna de las pendientes anteriores introduce oscilaciones indeseadas. Esto es causado porque cerca de las discontinuidades o en cualquier lugar donde haya gradientes pronunciados, la derivada puede crecer desmesuradamente y por lo tanto también la pendiente que aproxima esa derivada, de tal manera que debemos ser muy cuidadosos en la manera de elegir a σ_i^n en esas regiones delicadas. En ese sentido necesitamos implementar un algoritmo que *limite* las pendientes específicamente donde pudieran producirse las oscilaciones. Así tendremos una función *limitadora de pendientes*, que será la esencia del Método de Alta Resolución que estamos a punto de describir.

Capítulo 3

Métodos de Alta Resolución

3.1 Oscilaciones

La elección de una pendiente particular de cualquiera de las propuestas (2.11) genera oscilaciones cerca de una discontinuidad. Para mostrarlo, considérese una situación sencilla en la que se tienen los siguientes datos iniciales:

$$Q_i^0 = \begin{cases} 1 & \text{si } i \leq 0 \\ 0 & \text{si } i > 0. \end{cases}$$

Tómese una pendiente fija de Lax-Wendroff para la función reconstruida $\tilde{q}_i^0$. Haciendo así, se tiene que todas las pendientes son cero, excepto σ_0^0 que es $-1/\Delta x$. Por lo tanto $\tilde{q}_0^0 = 1 - \frac{1}{\Delta x}(x - \frac{\Delta x}{2}) = -\frac{x}{\Delta x} + \frac{3}{2}$. Notar que $\tilde{q}_0^0(0) = 1.5$ y $\tilde{q}_0^0(\Delta x) = 0.5$ son independientes del valor de Δx .

Supongamos que en el primer paso temporal Δt la función global reconstruida $\tilde{q}^0$ evoluciona y se mueve a la derecha una distancia $u\Delta t$ tal que $0 < u\Delta t < \Delta x$, donde u es la velocidad a la que se propaga la información. Para apreciar claramente el efecto, elíjase $u\Delta t = \Delta x/2$. En ese caso:

$$\tilde{q}_0^1(x) = \begin{cases} 1 & \text{si } x \leq \Delta x/2 \\ \tilde{q}_0^0(x - \Delta x/2) = -\frac{x}{\Delta x} + 2 & \text{si } x > \Delta x/2. \end{cases}$$

Después del primer paso temporal en la celda C_0 se tiene por definición:

$$\begin{aligned} Q_0^1 &= \frac{1}{\Delta x} \int_0^{\Delta x} \tilde{q}_0^1(x) dx \\ &= \frac{1}{\Delta x} \left[\int_0^{\Delta x/2} dx + \int_{\Delta x/2}^{\Delta x} \left(-\frac{x}{\Delta x} + 2 \right) dx \right] = \frac{9}{8} > 1, \end{aligned}$$

mientras que en C_1 :

$$Q_1^1 = \frac{1}{\Delta x} \int_{\Delta x}^{3\Delta x/2} \left(-\frac{x}{\Delta x} + 2 \right) dx = \frac{3}{8} > 0.$$

El problema es que nunca se espera que los datos tengan valores distintos a 1 y a 0 después de la evolución. De hecho, puesto que ahora habrá una pendiente positiva en la celda a la izquierda de $x = 0$, el resultado será un $Q_{-1}^2 < 1$. Igualmente, una pendiente negativa a la derecha de $x = \Delta x$ origina $Q_1^2 < 0$, dejando ver claramente el comienzo de una oscilación que crecerá con el tiempo. Ver la Figura (3.1).

Para evitar el fenómeno oscilatorio descrito arriba, una estrategia consiste en anular la pendiente para la celda C_0 , lo cual es razonable, pues las pendientes fueron propuestas para soluciones suaves y no habría por qué suponer que serán bien comportadas en regiones de discontinuidades. El coste de hacer las pendientes iguales a cero es la pérdida de un orden en la precisión, por eso se busca implementar un método de segundo orden de precisión en regiones suaves, es decir, que use pendientes de Lax-Wendroff o Beam-Warming, mientras que en las vecindades de las discontinuidades imponga pendientes cero.

3.2 La condición de Variación Total Decreciente

Ahora la principal tarea consiste en encontrar la manera de impedir que las pendientes de las funciones reconstruidas $\tilde{q}_i^n$ tiendan a crecer en las regiones de altos gradientes en la solución. Es aquí donde toma sentido la propiedad de *Alta Resolución* de método, pues éste será capaz de elegir la pendiente σ_i^n más apropiada para cada celda, en cada paso temporal, asegurándose de tomar aquella de mínima magnitud como sea posible cerca de las discontinuidades.

Antes de mostrar explícitamente la manera de limitar las pendientes, se debe imponer una condición que el método deberá cumplir necesariamente, garantizando que la magnitud de las variaciones de la solución numérica en torno a la solución verdadera sea decreciente durante la evolución. Un parámetro que estima cuánto varía la solución a un tiempo dado es la llamada *Variación Total* (TV, por sus siglas en inglés):

$$TV(Q^n) = \sum_{j=1}^{N_x} |Q_j^n - Q_{j-1}^n|.$$

Para un Q_i^n en particular:

$$TV(Q_i^n) = \sum_{j=1}^i |Q_j^n - Q_{j-1}^n|.$$

Se dice que un método tiene una *Variación Total Decreciente* cuando sucede que a lo largo

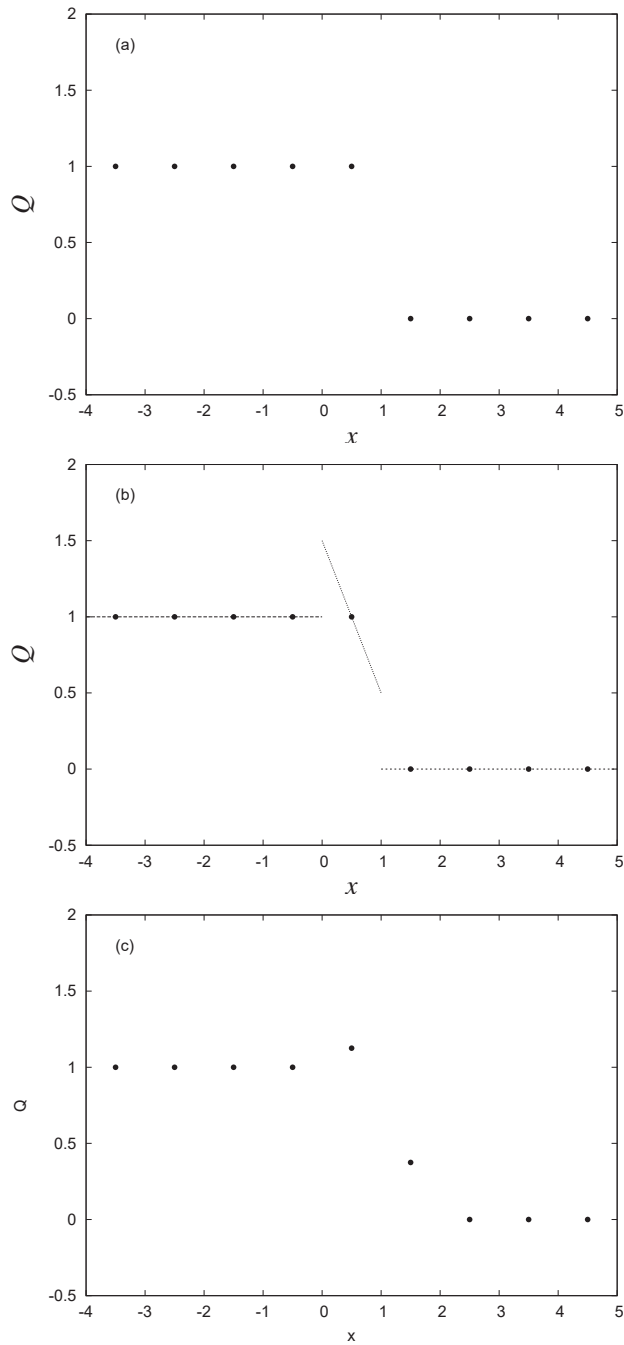


Figura 3.1: En (a) se muestran los datos iniciales antes de la evolución. En (b), la forma de la función reconstruida global usando una pendiente fija de Lax-Wendroff (línea punteada), antes de comenzar la evolución. En (c), la función reconstruida ha evolucionado después del primer paso temporal: se mueve a la derecha una distancia $\Delta x/2$. En la figura se muestran los promedios Q_i^1 . Ello hace evidente el comienzo de una oscilación que crecerá con el tiempo.

de toda la evolución:

$$TV(Q^{n+1}) \leq TV(Q^n). \quad (3.1)$$

Se sabe que la solución exacta de una ley de conservación tiene una Variación Total no-creciente [1], así que es razonable pedir que el método tenga una Variación Total Decreciente. Todos los métodos aquí discutidos cumplirán con ello.

3.3 Limitadores de pendientes

Retomando la función reconstruida (2.10) y la misión de evitar que las pendientes crezcan desmesuradamente, se da una primera alternativa de *función limitadora de pendientes*, llamada *función minmod*, que por supuesto genera una Variación Total Decreciente al tiempo que ofrece una convergencia a segundo orden en partes suaves de la solución (y a primer orden en zonas *difíciles*). Las pendientes quedan en los siguientes términos:

$$\sigma_i^n = \minmod\left(\frac{Q_i^n - Q_{i-1}^n}{\Delta x}, \frac{Q_{i+1}^n - Q_i^n}{\Delta x}\right), \quad (3.2)$$

con

$$\minmod(a, b) = \begin{cases} a & \text{si } |a| < |b| \text{ y } ab > 0 \\ b & \text{si } |b| < |a| \text{ y } ab > 0 \\ 0 & \text{si } ab \leq 0. \end{cases}$$

Hay que notar que por ejemplo, si ambas pendientes $\frac{Q_i^n - Q_{i-1}^n}{\Delta x}$ y $\frac{Q_{i+1}^n - Q_i^n}{\Delta x}$ son de signos distintos, entonces debe tratarse de la vecindad de un máximo o un mínimo local, así que para evitar oscilaciones y satisfacer (3.1) es necesario hacer $\sigma_i^n = 0$, lo que nos lleva al caso del método de Godunov y se reduce el orden de precisión. Así tenemos un método robusto que evita oscilaciones cerca de las discontinuidades, pagando el precio de perder precisión a cambio del manejo consistente de esas zonas, mientras que en partes suaves usa la pendiente de Lax-Wendroff o Beam-Warming, eligiendo la de menor magnitud.

Otra propuesta de limitador, que reduce menos drásticamente las pendientes cerca de las discontinuidades, es el *limitador MC* (de *Monotized Central difference*) propuesto en 1977 por van Leer [10]:

$$\sigma_i^n = \minmod\left[\left(\frac{Q_{i+1}^n - Q_{i-1}^n}{2\Delta x}\right), 2\left(\frac{Q_i^n - Q_{i-1}^n}{\Delta x}\right), 2\left(\frac{Q_{i+1}^n - Q_i^n}{\Delta x}\right)\right], \quad (3.3)$$

usando:

$$\minmod(a, b, c) = \begin{cases} a & \text{si } |a| < |b| \text{ , } |a| < |c| \text{ y } bc > 0 \\ b & \text{si } |b| < |a| \text{ , } |b| < |c| \text{ y } bc > 0 \\ c & \text{si } |c| < |a| \text{ , } |c| < |b| \text{ y } bc > 0 \\ 0 & \text{si } bc \leq 0 . \end{cases}$$

El limitador MC *selecciona* la mejor pendiente σ_i^n comparando la diferencia centrada de Fromm con el doble de las pendientes cargadas hacia adelante o hacia atrás a cada lado, de tal manera que en regiones suaves se tiene la pendiente centrada del método de Fromm, que ofrece precisión de segundo orden, mientras que cerca de las discontinuidades se sigue teniendo precisión de primer orden pero con una mejor resolución que *minmod* y cumpliendo siempre con la condición de Variación Total Decreciente.

En el siguiente capítulo se presentan algunos ejemplos con la intención de especificar cómo funcionan los limitadores de pendientes. Antes es necesario introducir algunos otros ingredientes importantes en la solución numérica de ecuaciones diferenciales.

3.4 Condiciones de frontera

El método numérico y las variantes que se han descrito hasta ahora se formulan suponiendo que en cada paso temporal, para calcular Q_i^{n+1} , en general siempre se tiene información de celdas vecinas: al menos Q_{i-1}^n y Q_{i+1}^n . La misma suposición se hace cuando se escribe la función reconstruida $\tilde{q}_i^n$ en términos de una pendiente que depende de Q_{i-1}^n , Q_i^n y Q_{i+1}^n . Sin embargo hay que tener en cuenta que por restricciones computacionales, el dominio es finito cuando se considera una malla uniforme y no se le puede extender arbitrariamente para tener siempre disponible la información en celdas vecinas cuando sea necesaria. Forzosamente se debe incluir en el algoritmo una manera especial de calcular las funciones en los extremos del dominio. A ello se llama imponer *condiciones de frontera*. Al modelar un fluido confinado a un recipiente de paredes rígidas, por ejemplo, se requiere que los cálculos en las fronteras del dominio sean tales que se logre un efecto de *reflexión* íntegra de la información que las alcance; en ese caso se tienen *condiciones de frontera reflejantes*. Si en cambio se está estudiando sólo una porción de un flujo estacionario que se extiende ampliamente, es razonable *cerrar* el dominio uniendo los extremos para conseguir un efecto de *periodicidad*, es decir, la información que sale por un lado, entra por el otro, y ésto se repite periódicamente, aludiendo a la estacionariedad ideal del flujo; en ese caso se tienen *condiciones de frontera periódicas*. En problemas de Relatividad General, por dar otro ejemplo, con frecuencia se tienen sistemas aislados y fenómenos que no admiten periodicidad, por eso es necesario pedir que las fronteras *dejen salir* la información e impidan cualquier reflexión; más particularmente, al reproducir la expansión de la onda de choque de una Supernova en una porción finita del espacio-tiempo, es natural imponer este tipo de condiciones, llamadas *absorbentes*. Queda claro que la improvisación de condiciones de

frontera es poco conveniente. En lugar de ello hay que pensar con cuidado la manera de imponerlas.

Al implementar condiciones de frontera absorbentes, formalmente, para calcular las pendientes en x_0 y x_{N_x-1} se extrapola a la izquierda y derecha de C_0 y C_{N_x-1} respectivamente, para obtener dos números: Q_{-1}^n y $Q_{N_x}^n$, los cuales corresponden a celdas virtuales C_{-1} y C_{N_x} (llamadas *celdas fantasma*), y luego se calculan las pendientes σ_0^n en términos de Q_{-1}^n , Q_0^n y Q_1^n , mientras que $\sigma_{N_x-1}^n$ estará en términos de $Q_{N_x-2}^n$, $Q_{N_x-1}^n$ y $Q_{N_x}^n$. Si la extrapolación es lineal (como se ha hecho para todos los ejemplos en este trabajo) no se requiere del cálculo explícito de Q_{-1}^n y $Q_{N_x}^n$; a continuación se discuten varios casos:

i) Si se usan pendientes fijas de Beam-Warming (diferencias adelantadas), entonces no hay problemas en la frontera derecha; en la izquierda basta con darse cuenta de que la pendiente de la extrapolación es justo la que corresponde a la función reconstruida en C_0 y también es la misma que la pendiente en C_1 . Por lo tanto $\sigma_0^n = \sigma_1^n$.

ii) Si en cambio se usan pendientes fijas de Lax-Wendroff (diferencias atrasadas), no hay problemas en la frontera izquierda; en la derecha se aplica el mismo razonamiento que en i) y se concluye que $\sigma_{N_x-1}^n = \sigma_{N_x-2}^n$.

iii) Cuando se usan pendientes fijas de Fromm (diferencias centradas), en la frontera izquierda $\sigma_0^n = (Q_1^n - Q_0^n)/\Delta x$ y en la derecha $\sigma_{N_x-1}^n = (Q_{N_x-1}^n - Q_{N_x-2}^n)/\Delta x$. Lo anterior no es obvio y por eso a continuación se da una prueba para el lado izquierdo: si se denota con m a la pendiente de la extrapolación: $m = (Q_1^n - Q_0^n)/\Delta x = (Q_0^n - Q_{-1}^n)/\Delta x$. Entonces $2m = (Q_1^n - Q_0^n)/\Delta x + (Q_0^n - Q_{-1}^n)/\Delta x = (Q_1^n - Q_{-1}^n)/\Delta x = 2(Q_1^n - Q_{-1}^n)/2\Delta x = 2\sigma_0^n$. Por lo tanto $\sigma_0^n = (Q_1^n - Q_0^n)/\Delta x$. Para el lado derecho la comprobación es análoga.

iv) Si en lugar de una pendiente fija se usa el limitador *minmod*, el razonamiento en la frontera izquierda es el siguiente: *minmod* debe comparar y elegir la menor entre las diferencias $(Q_1^n - Q_0^n)/\Delta x$ y $(Q_0^n - Q_{-1}^n)/\Delta x$, pero éstas son exactamente las mismas debido a que la extrapolación para calcular Q_{-1}^n es w lineal de pendiente $m = (Q_0^n - Q_{-1}^n)/\Delta x = (Q_1^n - Q_0^n)/\Delta x$. Por lo tanto: $\sigma_0^n = (Q_1^n - Q_0^n)/\Delta x$ y $\sigma_{N_x-1}^n = (Q_{N_x-1}^n - Q_{N_x-2}^n)/\Delta x$.

v) Con razonamientos similares se concluye que aun usando el limitador MC, las pendientes en las fronteras siguen siendo: $\sigma_0^n = (Q_1^n - Q_0^n)/\Delta x$ y $\sigma_{N_x-1}^n = (Q_{N_x-1}^n - Q_{N_x-2}^n)/\Delta x$. Al implementar condiciones de frontera absorbentes la tarea es más sencilla: basta con hacer una identificación de la celda del extremo derecho del dominio con la del extremo izquierdo, lo que implica que: $\sigma_0^n = \sigma_{N_x-2}^n$ y $\sigma_{N_x-1}^n = \sigma_1^n$.

Ahora las condiciones de frontera para Q_i^n : la implementación más sencilla es la de las condiciones periódicas, simplemente se *cierra* el dominio usando dos valores virtuales: $Q_{-1}^n = Q_{N_x-1}^n$ y $Q_{N_x}^n = Q_0^n$. En general, cuando Q_i^{n+1} está en términos de más de dos celdas vecinas, basta con imponer la misma condición de periodicidad añadiendo las celdas fantasma necesarias para calcular Q_0^{n+1} y $Q_{N_x-1}^{n+1}$.

Al imponer condiciones absorbentes se busca que la información que alcance las fronteras salga y no regrese, alejándose hacia $\pm\infty$. Sería catastrófico que al ser evolucionados dentro del dominio, los datos se *contaminaran* por información proveniente de las fronteras. Se debe poner cuidado en ello. La idea es usar información virtual de nuevo: si extrapolamos

linealmente hacia afuera de la frontera izquierda obtenemos que en la celda fantasma C_{-1} , $Q_{-1}^n = 2Q_0^n - Q_1^n$, lo que después se usará en el cálculo de Q_0^{n+1} . Análogamente, para calcular $Q_{N_x-1}^{n+1}$ hará falta conocer $Q_{N_x}^n$, que se calcula extrapolando linealmente sobre la celda fantasma C_{N_x} mediante: $Q_{N_x}^n = 2Q_{N_x-1}^n - Q_{N_x-2}^n$.

En realidad no es difícil implementar las condiciones descritas arriba. Los códigos usados para obtener los resultados de este trabajo fácilmente se adaptan para tener condiciones de frontera periódicas y absorbentes. En particular se usarán estas últimas al resolver las ecuaciones de Euler y veremos que la elección queda justificada por la naturaleza del fenómeno que se modela.

3.5 Convergencia y estabilidad

Convergencia

No hay que perder de vista que el objetivo primordial es encontrar la solución de algún sistema de ecuaciones diferenciales para todo tiempo. Sucede que la solución que se obtiene numéricamente en un dominio discretizado es esencialmente distinta de la solución exacta en el espacio-tiempo continuo: la que se busca. No es ninguna calamidad descubrir a estas alturas que los métodos numéricos nunca darán la solución exacta de las ecuaciones. Es natural, porque un espacio-tiempo continuo es un concepto (y una computadora no puede reproducirlo) y uno discreto es otro: un subconjunto del continuo. A lo más que puede aspirarse es a lograr que si se aumenta la resolución del dominio discretizado, entonces la solución numérica se aproxime uniformemente, es decir, *converja*, a la solución exacta. No se trata de una modesta pretensión, de hecho es muy ambiciosa, y cuando se haya logrado, el trabajo habrá terminado exitosamente. Pero las garantías de éxito no se tienen a priori, y las dificultades pueden ser muchas, todas provenientes de la formulación e implementación del método.

Hay más dificultades todavía cuando el sistema a resolver es tan complicado que ni siquiera conocemos la solución exacta; es aquí donde se necesita un razonamiento redentor: si $\Delta x, \Delta t \rightarrow 0$, entonces $N_x, N_t \rightarrow \infty$, lo que define un espacio continuo. Trasladando esto a la práctica: si al hacer Δx y Δt tan pequeños como lo permita la computadora, la solución converge a *cierta solución*, entonces ¡esa debe ser la solución en el continuo! Es la parte linda de los métodos numéricos. Los detalles técnicos van a continuación.

Al calcular una solución numérica inevitablemente se comete un error, el cual se define como la diferencia en cada punto entre la solución exacta en el continuo y la numérica. La magnitud de los errores generalmente depende del método. Decimos que una solución es precisa a orden (n, m) si el error que se comete al calcularla es del orden de $O(\Delta x^n) + O(\Delta t^m)$.

Sea $Q_k(x)$ una función que toma valores en un dominio discreto y que proviene de un método que devuelve resultados precisos a segundo orden. No hay que confundir el subíndi-

ce: aquí es una etiqueta para la función y no hace referencia al valor que ésta toma en la k -ésima celda. Q_k puede escribirse como $Q_k(x) = Q_0(x) + E(x)(\Delta x^2) + O(\Delta x^3)$, donde $Q_0(x)$ denota a la solución exacta en el continuo y E (desconocido) es el error que se comete inevitablemente al calcular Q_k . Si conocemos $Q_0(x)$ y se tiene un par de funciones Q_1 y Q_2 calculadas con diferentes resoluciones, Δx y $\Delta x/2$, respectivamente, entonces:

$$\frac{Q_1 - Q_0}{Q_2 - Q_0} = \frac{\Delta x^2 + O(\Delta x^3)}{(\Delta x/2)^2 + O(\Delta x^3)} = 4 + O(\Delta x^3).$$

Esta fórmula dicta la manera práctica de comprobar que un método calcula soluciones convergentes, precisas a segundo orden. Cuando se consigue ese anhelado 4 (llamado *factor de convergencia*) se dice que se tiene *convergencia a segundo orden*. De hecho 4 es el factor entre las dos resoluciones (que en este caso es 2) elevado a la segunda potencia, que es el orden de convergencia.

En la práctica normalmente no se conoce Q_0 . En ese caso se hace lo que se llama una *prueba de convergencia de Cauchy*, usando tres funciones: Q_1 , Q_2 y Q_3 , calculadas con resoluciones diferentes Δx , $\Delta x/2$ y $\Delta x/4$, respectivamente. Así:

$$\begin{aligned} \frac{Q_1 - Q_2}{Q_2 - Q_3} &= \frac{\Delta x^2 - \frac{1}{4}\Delta x^2 + O(\Delta x^3)}{\frac{1}{4}\Delta x^2 - \frac{1}{16}\Delta x^2 + O(\Delta x^3)} \\ &= \frac{\Delta x^2 + O(\Delta x^3)}{\frac{1}{4}\Delta x^2 + O(\Delta x^3)} = 4 + O(\Delta x^3). \end{aligned} \quad (3.4)$$

Una vez más el 4 es el factor de convergencia de la implementación del método numérico. Un comentario importante al respecto de la convergencia a segundo orden es que ésta puede verse reducida progresivamente en el tiempo a convergencia de primer orden, esto como consecuencia del uso de extrapolaciones lineales en las fronteras, pues siempre que haya propagación de información de éstas hacia el interior del dominio, al mismo tiempo se propagará un error de primer orden, y siempre hay que tenerlo en consideración.

Estabilidad

Otro aspecto importante en el análisis numérico es la *teoría de la estabilidad*. Cada algoritmo computacional que pretenda resolver numéricamente un sistema de ecuaciones diferenciales amerita un análisis teórico que proporcione un criterio o una condición para que no sea el mismo algoritmo el que provoque la divergencia de la solución; en tal caso se dice que el método es inestable: las evidencias pueden ser disipación u oscilaciones indeseadas.

Resulta que para una gran cantidad de métodos, el exhaustivo análisis de estabilidad [2]

puede resumirse en un enunciado que involucra las resoluciones espacial y temporal: Δx y Δt , las cuales no pueden ser independientes y deben cumplir con la llamada *condición de Courant-Friedrichs-Lewy (CFL)*:

$$0 < \lambda \leq \frac{1}{u},$$

donde $\lambda \equiv \frac{\Delta t}{\Delta x}$ se llama *Factor de Courant*, y u es en general la velocidad de propagación de la información. Esta condición será necesaria para la estabilidad del método (y todas sus variantes) que se discute en esta Tesis. En cada ejemplo se procura cumplirla.

Capítulo 4

Implementación y pruebas con ecuaciones escalares

4.1 La ecuación escalar de advección

Cuando se modela el transporte de una sustancia contenida en un fluido durante el movimiento de éste, el fenómeno es referido como *trasporte advectivo*. Si por ejemplo se considera un contaminante siendo transportado en un flujo unidimensional a velocidad fija u , entonces la densidad del contaminante como función del espacio y del tiempo satisface la llamada *ecuación escalar de advección*. Se trata de una ecuación escrita en forma conservativa con un coeficiente real y constante u :

$$\partial_t q + u \partial_x q = 0, \quad u \in \mathbb{R}, \quad (4.1)$$

u se llama velocidad de advección de la cantidad conservada $q = q(x, t)$. En este caso se conoce la solución exacta de la ecuación: es muy sencillo verificar por sustitución directa que cualquier función de la forma $q(x, t) = q_0(x - ut)$ la satisface.

La solución numérica de esta ecuación usando una discretización en volúmenes finitos comienza por definir Q_i^n para cada celda C_i^n como en (2.2). Notar que en este ejemplo $f(q) = uq$. Enseguida se pretende usar la fórmula (2.8), que es la más relevante hasta ahora, así que primero hay que calcular los flujos:

$$\begin{aligned} F_i^n &= \frac{1}{\Delta t} \int_{t_n}^{t_{n+1}} f(\tilde{q}_i(x_i, t)) dt \\ &= \frac{1}{\Delta t} \int_{t_n}^{t_{n+1}} u \tilde{q}_i(x_i, t) dt \\ &= \frac{u}{\Delta t} \left[\frac{\Delta t \tilde{q}_i^{n+1}(x_i, t^{n+1}) + \Delta t \tilde{q}_i^n(x_i, t^n)}{2} \right]. \end{aligned}$$

Ahora el problema es calcular $\tilde{q}_i^{n+1}(x_i, t^{n+1})$, y la manera de hacerlo consiste en resolver (4.1) en cada celda, tomando (2.10) como la función reconstruida y teniendo en cuenta que:

$$\begin{aligned}\tilde{q}_i^n(x_i, t^n) &= Q_i^n + \sigma_i^n \left(-\frac{\Delta x}{2} \right) \\ \tilde{q}_i^n(\bar{x}_i, t^n) &= Q_i^n.\end{aligned}\tag{4.2}$$

Por otro lado la versión discreta (a primer orden en el espacio y en el tiempo) de (4.1) se escribe:

$$\frac{\tilde{q}_i^{n+1}(x_i) - \tilde{q}_i^n(x_i)}{\Delta t} + u \left[\frac{\tilde{q}_i^n(x_{i+\frac{1}{2}}) - \tilde{q}_i^n(x_i)}{\frac{\Delta x}{2}} \right] = 0.$$

Sustituyendo (4.2) en la ecuación anterior:

$$\begin{aligned}\tilde{q}_i^{n+1} &= \tilde{q}_i^n - u\Delta t \left[\frac{Q_i^n - Q_i^n + \sigma_i^n \frac{\Delta x}{2}}{\frac{\Delta x}{2}} \right] \\ &= \tilde{q}_i^n - u\Delta t \sigma_i^n \\ &= Q_i^n + \sigma_i^n (x - u\Delta t - \bar{x}_i) \\ &= \tilde{q}_i^n(x - u\Delta t, t^n),\end{aligned}$$

donde también se ha usado (2.10).

Antes de calcular el flujo numérico es importante notar que:

$$\begin{aligned}\tilde{q}_i^{n+1} + \tilde{q}_i^n &= \tilde{q}_i^n(x - u\Delta t, t^n) + \tilde{q}_i^n(x, t^n) \\ &= Q_i^n + \sigma_i^n (x - u\Delta t - \bar{x}_i) \\ &\quad + Q_i^n + \sigma_i^n (x - \bar{x}_i) \\ &= 2 \left[Q_i^n + \sigma_i^n \left(x - \frac{u\Delta t}{2} - \bar{x}_i \right) \right] \\ &= 2 \left[\tilde{q}_i^n \left(x - \frac{u\Delta t}{2}, t^n \right) \right].\end{aligned}\tag{4.3}$$

Finalmente se calcula el flujo numérico como está indicado por (2.7). Considerando que $\tilde{q}_i^n$ no está definido para todo $t \in [t^n, t^{n+1}]$, la integral para F_i^n es un promedio que se transforma en:

$$F_i^n = \frac{1}{\Delta t} \left[\frac{\Delta t f(\tilde{q}_i^{n+1}(x_i, t^{n+1})) + \Delta t f(\tilde{q}_i^n(x_i, t^n))}{2} \right], \quad (4.4)$$

y usando los resultados que se han obtenido hasta (4.3), el flujo numérico es:

$$F_i^n = u \left[\tilde{q}_i^n \left(x_i - \frac{u \Delta t}{2}, t^n \right) \right].$$

Si $u \geq 0$ entonces la cantidad entre corchetes no está definida sobre la celda C_i , sino en C_{i-1} , así que:

$$\begin{aligned} F_i^n &= u \left[Q_{i-1}^n + \sigma_{i-1}^n \left(x_i - \frac{u \Delta t}{2} - x_{i-1} - \frac{\Delta x}{2} \right) \right] \\ &= u Q_{i-1}^n + \frac{u}{2} \sigma_{i-1}^n (\Delta x - u \Delta t), \quad u \geq 0. \end{aligned} \quad (4.5)$$

En el caso en que $u < 0$ el flujo resulta ser:

$$F_i^n = u Q_i^n - \frac{u}{2} \sigma_i^n (\Delta x + u \Delta t), \quad u < 0. \quad (4.6)$$

El paso final es calcular Q_i^{n+1} usando la forma en diferencia de flujos (2.8):

$$\begin{aligned} Q_i^{n+1} &= Q_i^n - u \frac{\Delta t}{\Delta x} \left[(Q_i^n - Q_{i-1}^n) + \frac{1}{2} (\Delta x - u \Delta t) (\sigma_i^n - \sigma_{i-1}^n) \right], \quad u \geq 0, \\ Q_i^{n+1} &= Q_i^n - u \frac{\Delta t}{\Delta x} \left[(Q_{i+1}^n - Q_i^n) - \frac{1}{2} (\Delta x + u \Delta t) (\sigma_{i+1}^n - \sigma_i^n) \right], \quad u < 0. \end{aligned} \quad (4.7)$$

Para el primer ejemplo concreto de la implementación del método, se elige $u = 1$, entonces:

$$Q_i^{n+1} = Q_i^n - \frac{\Delta t}{\Delta x} \left[(Q_i^n - Q_{i-1}^n) + \frac{1}{2} (\Delta x - \Delta t) (\sigma_i^n - \sigma_{i-1}^n) \right]. \quad (4.8)$$

Hay consistencia en el hecho de que si las pendientes son cero, entonces el método se reduce al estándar de diferencias finitas adelantadas. Por otro lado, como se ha dicho antes, tener las pendientes nulas equivale a implementar el Método de Godunov, por lo tanto, para la ecuación de advección con $u = 1$, el Método de Godunov se reduce a un esquema de diferencias finitas adelantadas.

Los datos iniciales para el ejemplo particular son los siguientes, definidos para $x \in [0, 1]$:

$$q^0(x) = \begin{cases} 2 \exp \left[-\frac{(x-0.35)^2}{(0.1)^2} \right] & 0 \leq x \leq 0.7 \\ 1.5 & 0.7 < x \leq 0.9 \\ 0 & \text{en otro caso} \end{cases}.$$

Así se tiene una función continua y suave a trozos, para experimentar y comparar los métodos en varias situaciones.

Durante la evolución se espera que los datos iniciales sean advectados a la derecha, pues $u = 1 > 0$. Si además se imponen condiciones de frontera periódicas, la señal que salga por la derecha entrará por la izquierda y se repetirá el ciclo, permitiendo monitorear cualquier efecto disipativo (decremento con el tiempo de la amplitud de los datos), al mismo tiempo que se evalúa la tendencia del método a suavizar las discontinuidades y a generar oscilaciones (en caso de no implementar un limitador de pendientes).

Discretizando la malla con 1000 puntos se tiene $\Delta x = 0.001$, y tomando el factor de Courant $\lambda \equiv \frac{\Delta t}{\Delta x} = 0.5$, se calcula $\Delta t = \frac{\Delta x}{2} = 0.0005$. Como la extensión del dominio es 1, los datos lo habrán recorrido completo una vez cuando hayan avanzado una distancia $1 = 1000\Delta x$, pero $\Delta x = 2\Delta t$, entonces se requiere de dos mil iteraciones para lograrlo. La razón de imponer condiciones de frontera periódicas viene del hecho de que $q^0(x - ut)$ es una solución de (4.1), y en particular para $u = 1$ y $t = 1$, $q^0(x - 1)$ es solución de la ecuación, y como el dominio mide exactamente 1, se espera que a $t = 2000\Delta t = 1$ la solución numérica se encuentre justo en la misma posición donde estaban los datos iniciales; así puede observarse fácilmente cualquier falla del método: oscilaciones, disipación, etcétera.

Utilizando un integrador de Runge-Kutta de tercer orden [3], se obtienen los resultados mostrados en la Figura 4.1, donde se comparan las soluciones obtenidas a $t = 1$ usando pendientes nulas, pendientes fijas de Lax-Wendroff y los limitadores *minmod* y MC. Se observa que al incrementar la complejidad del método, la disipación y las oscilaciones cerca de las discontinuidades tienden a desaparecer. El limitador MC devuelve los mejores resultados, y teóricamente converge a segundo orden en regiones suaves, lo cual se comprueba en la Figura 4.2).

4.1.1 Pseudocódigo para resolver la ecuación escalar de advección

Resulta conveniente dar a conocer cada paso a seguir en la programación. Los ingredientes están expuestos, la receta queda completa al describir el procedimiento:

0. Inicia el algoritmo.
1. En un archivo de entrada de parámetros se especifican:
 - * valores numéricos de las fronteras del dominio.
 - * número de puntos que tendrá la malla.
 - * factor de Courant.
 - * número de iteraciones.

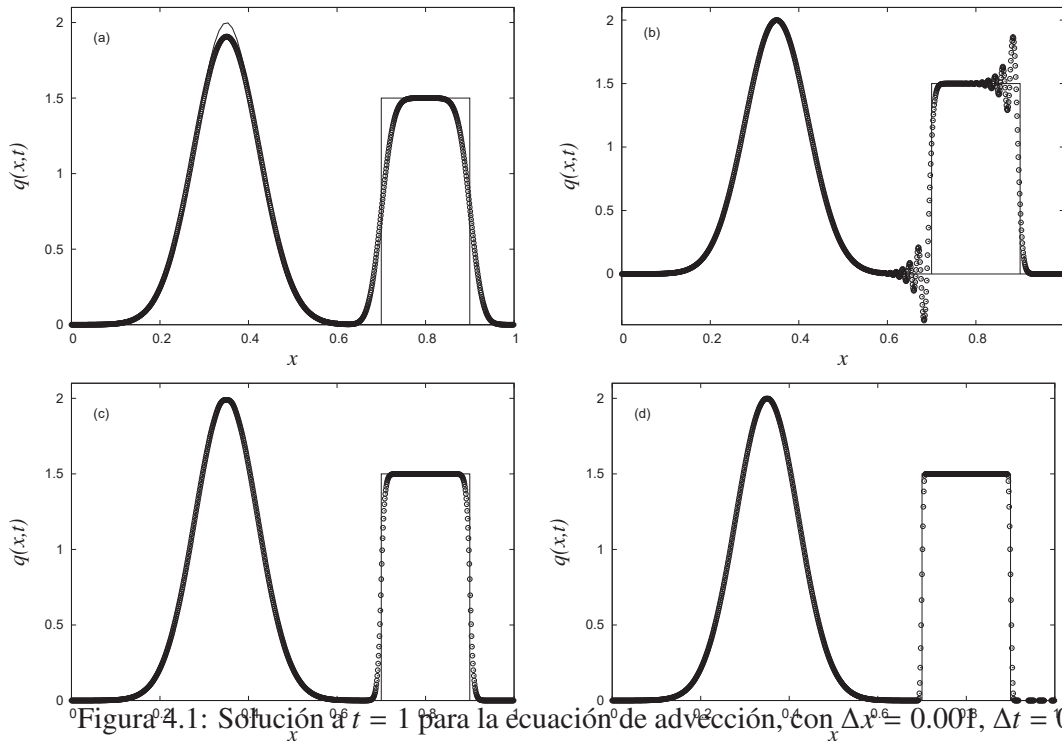


Figura 4.1: Solución a $t = 1$ para la ecuación de advección, con $\Delta x = 0.001$, $\Delta t = 0.0005$ y condiciones de frontera periódicas. La línea continua es la solución exacta a $t = 1$. (a) Usar pendientes nulas arroja resultados muy modestos: hay una disipación muy notoria y las discontinuidades se suavizan rápidamente. (b) Usando pendientes fijas de Lax-Wendroff (atrasadas) hay menos disipación, pero se ven oscilaciones espurias cerca de las discontinuidades. (c) Las diferencias son apreciables con el limitador de pendientes *minmod*. (d) Limitador MC.

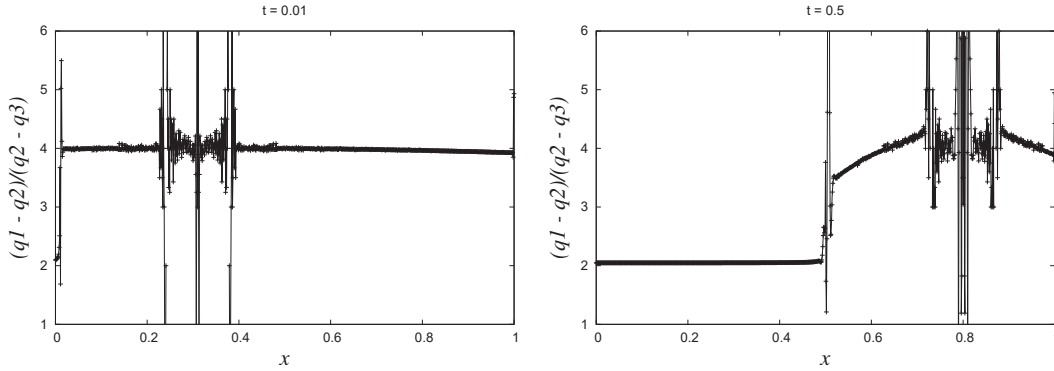


Figura 4.2: Usando el limitador MC, condiciones de frontera absorbente y un perfil gaussiano centrado en $x = 0.3$ como dato inicial, se calcula la solución numérica para la ecuación de advección ($u = 1$) en su forma conservativa para tres resoluciones diferentes: $\Delta x = 0.001 \equiv \Delta x_0$ (q_1), $\Delta x = \Delta x_0/2$ (q_2) y $\Delta x = \Delta x_0/4$ (q_4). Se comparan las soluciones según (3.4). A la izquierda, $t = 0.1$; a la derecha, $t = 0.5$. Notar que la convergencia a segundo orden es inestable en la región del máximo del perfil; la convergencia a primer orden se debe a la propagación de un error de primer orden desde las fronteras

- * datos iniciales positivos.
- * velocidad de advección.
- * tipo de pendiente.

Se calculan Δx y Δt .

2. Se definen los datos iniciales q_i^0 en la malla recién construida.
3. Se calculan los promedios Q_i^0 así: $Q_i^0 = (q_i^0 + q_{i+1}^0)/2$.
4. Se registran estos valores como datos iniciales.
5. Comienza la evolución temporal:

- Una subrutina calcula las pendientes σ_i^n en cada celda (en términos de Q_{i-1}^n , Q_i^n y Q_{i+1}^n), usando ya sea una pendiente fija, o algún limitador.
- Tomando en cuenta el signo de la velocidad de advección u (que en el ejemplo fue $u = 1$), el programa decide cuál de las dos opciones de (4.7) debe usar y entonces calcula el *lado derecho* de la fórmula elegida, usando las pendientes previamente obtenidas. En este mismo paso se aplican las condiciones de frontera.
- Se integra con el algoritmo Runge-Kutta de tercer orden para obtener Q_i^{n+1} .

- Se guardan los datos. Para comenzar una nueva iteración se hace $Q_i^n = Q_i^{n+1}$.

6. Fin del algoritmo.

4.2 La ecuación de Burgers

Uno de los ejemplos más sencillos de ecuaciones escalares no-lineales es la llamada ecuación de Burgers inviscida, que tiene la siguiente forma:

$$\partial_t q + q \partial_x q = 0. \quad (4.9)$$

El principal problema de esta ecuación proviene de la no-linealidad, lo que implica la formación de una discontinuidad en la solución a un tiempo finito. Sin importar si los datos iniciales son suaves, la discontinuidad aparece y es difícil de tratar consistentemente usando la aproximación de diferencias finitas. Por esta razón, la implementación de un método de diferencias finitas en principio fracasará, pues los gradientes pronunciados hacen diverger las aproximaciones de las derivadas. Incluso se debe poner atención a la manera en que se escribe la ecuación, pues si (4.9) no se escribe en su forma conservativa y en cambio se discretiza en su forma cuasilineal, el esquema en diferencias finitas sería:

$$q_i^{n+1} = q_i^n - \frac{\Delta t}{\Delta x} q_i^n (q_i^n - q_{i-1}^n), \quad (4.10)$$

al poner datos iniciales discontinuos se obtiene una solución que se mueve con una velocidad incorrecta, lo cual se muestra más adelante en base a resultados numéricos; pero si se escribe la ecuación de Burgers en su forma conservativa:

$$\partial_t q + \partial_x \left(\frac{1}{2} q^2 \right) = 0, \quad (4.11)$$

inmediatamente se identifica la función de flujo como $f(q) = \frac{1}{2} q^2$ y la ecuación se discretiza así:

$$q_i^{n+1} = q_i^n - \frac{\Delta t}{\Delta x} \left[\frac{1}{2} (q_i^n)^2 - \frac{1}{2} (q_{i-1}^n)^2 \right], \quad (4.12)$$

y se obtienen resultados notoriamente mejores (ver Figura 4.3). Es importante notar que en las regiones donde las soluciones para (4.9) y (4.11) son suaves, éstas son equivalentes, pero dejan de serlo cerca de las discontinuidades. Esto muestra la importancia de escribir siempre las ecuaciones en su forma conservativa, pues de otro modo se obtienen soluciones no convergentes.

Hay una alternativa *sencilla* para conseguir mejores resultados. Se considera la ecuación de Burgers con viscosidad:

$$\partial_t q + q \partial_x q = \epsilon \partial_{xx} q, \quad \epsilon > 0. \quad (4.13)$$

El método suele llamarse *de viscosidad artificial*, y consiste en agregar un término disipativo (el lado derecho de (4.13)), llamado *término de viscosidad artificial*, para reducir el comportamiento erróneo de la solución en las regiones cercanas a los altos gradientes. La idea es la siguiente: si ϵ es muy pequeño, el término $\epsilon \partial_{xx} q$ es despreciable respecto a los demás en regiones donde los datos son suaves, pero conforme comienza a formarse la discontinuidad, este término crece mucho más rápido que $\partial_x q$, por ser proporcional a la segunda derivada, y en algún momento será comparable a los otros términos, generando disipación y manteniendo estable la solución. Cuando $\epsilon \rightarrow 0$, la zona donde el término de viscosidad se vuelve importante queda mejor localizada y la solución numérica se aproxima mejor a la solución real. Incluso del nombre el término proviene de las ecuaciones de la hidrodinámica, pues los términos de viscosidad provocan el mismo efecto disipativo que ϵ en (4.13).

En primera instancia puede pensarse que es un error asumir que la solución de (4.9) es la misma que la de una ecuación tan distinta como (4.13), sin embargo en [2] se da una demostración rigurosa de que las soluciones coinciden en el límite cuando ϵ tiende a cero.

El problema con el método de viscosidad artificial es que ϵ no es físico (se agrega a mano, y su valor se convierte en un parámetro extra de la solución numérica), y es preferible abordar el problema prescindiendo de cualquier parámetro sin sentido físico, es decir: lo más razonable es tratar de evitar coeficientes de viscosidad cuando se está modelando un fluido invíscido. Por otro lado, el valor de ϵ no es arbitrario y depende de la ecuación que se pretende resolver. Algunos criterios formales para optimizarlo son discutidos en [2].

Ahora se muestra la implementación de un Método de Alta Resolución. Como con la ecuación de advección, la idea es usar:

$$Q_i^{n+1} = Q_i^n - \frac{\Delta t}{\Delta x} [F_{i+1}^n - F_i^n],$$

con F_i^n dados por (4.4). Si se elige una reconstrucción lineal $\tilde{q}_i^n = Q_i^n + \sigma_i^n(x - \bar{x}_i)$, entonces:

$$F_i^n = \frac{1}{\Delta t} \left[\frac{\frac{\Delta t}{2} (\tilde{q}_i^{n+1}(x_i, t^{n+1}))^2 + \frac{\Delta t}{2} (\tilde{q}_i^n(x_i, t^n))^2}{2} \right]. \quad (4.14)$$

donde $\tilde{q}_i^{n+1}(x_i, t^{n+1})$ debe calcularse resolviendo (4.11) y usando (4.2). La discretización a primer orden de (4.11) se ve:

$$\frac{\tilde{q}_i^{n+1}(x_i) - \tilde{q}_i^n(x_i)}{\Delta t} + \left[\frac{(\tilde{q}_i^n(x_{i+\frac{1}{2}}))^2 - \frac{1}{2}(\tilde{q}_i^n(x_i))^2}{\frac{\Delta x}{2}} \right] = 0,$$

luego,

$$\begin{aligned} \tilde{q}_i^{n+1} &= \tilde{q}_i^n - \frac{\Delta t}{\Delta x} \left[(Q_i^n)^2 - \left(Q_i^n - \frac{\Delta x}{2} \sigma_i^n \right)^2 \right] \\ &= \left(Q_i^n + \frac{\Delta x}{2} \sigma_i^n \right) (1 - \Delta t \sigma_i^n) + \frac{\Delta t \Delta x}{4} (\sigma_i^n)^2. \end{aligned}$$

Ahora se calcula el flujo numérico a partir de (4.14):

$$\begin{aligned} F_i^n &= \frac{1}{4} \left[(\tilde{q}_i^{n+1}(x_i, t^{n+1}))^2 + (\tilde{q}_i^n(x_i, t^n))^2 \right] \\ &= \frac{1}{4} \left[\left(Q_i^n + \frac{\Delta x}{2} \sigma_i^n \right)^2 (1 + [1 - \Delta t \sigma_i^n]^2) \right] \\ &\quad + \frac{1}{4} \left[\frac{\Delta t \Delta x}{2} \left(Q_i^n + \frac{\Delta x}{2} \sigma_i^n \right) (1 - \Delta t \sigma_i^n) (\sigma_i^n)^2 \right] \\ &\quad + \frac{1}{4} \left[\left(\frac{\Delta t \Delta x}{4} \right)^2 (\sigma_i^n)^4 \right]. \end{aligned}$$

Y finalmente la forma en diferencia de flujos es:

$$\begin{aligned} Q_i^{n+1} &= Q_i^n - \frac{\Delta t}{4\Delta x} \left[\left(Q_{i+1}^n + \frac{\Delta x}{2} \sigma_{i+1}^n \right)^2 (1 + [1 - \Delta t \sigma_{i+1}^n]^2) \right] \\ &\quad + \frac{\Delta t}{4\Delta x} \left[\left(Q_i^n + \frac{\Delta x}{2} \sigma_i^n \right)^2 (1 + [1 - \Delta t \sigma_i^n]^2) \right] \\ &\quad - \frac{\Delta t}{4\Delta x} \frac{\Delta t \Delta x}{2} \left[\left(Q_{i+1}^n + \frac{\Delta x}{2} \sigma_{i+1}^n \right) (1 - \Delta t \sigma_{i+1}^n) (\sigma_{i+1}^n)^2 \right] \\ &\quad + \frac{\Delta t}{4\Delta x} \frac{\Delta t \Delta x}{2} \left[\left(Q_i^n + \frac{\Delta x}{2} \sigma_i^n \right) (1 - \Delta t \sigma_i^n) (\sigma_i^n)^2 \right] \\ &\quad - \frac{\Delta t}{4\Delta x} \left(\frac{\Delta t \Delta x}{4} \right)^2 \left[(\sigma_{i+1}^n)^4 - (\sigma_i^n)^4 \right]. \end{aligned} \tag{4.15}$$

Los resultados se presentan en la Figura 4.3. Compárese la solución cuando se usa la

discretización cuasilineal (4.10), disipación numérica con diferencias finitas centradas, y usando el método de Alta Resolución (4.15). Claramente la velocidad de propagación es muy distinta en cada caso. La solución más satisfactoria por su convergencia, al menos a primer orden, es la que resulta de aplicar el limitador MC a un perfil inicial suave (Figura 4.4).

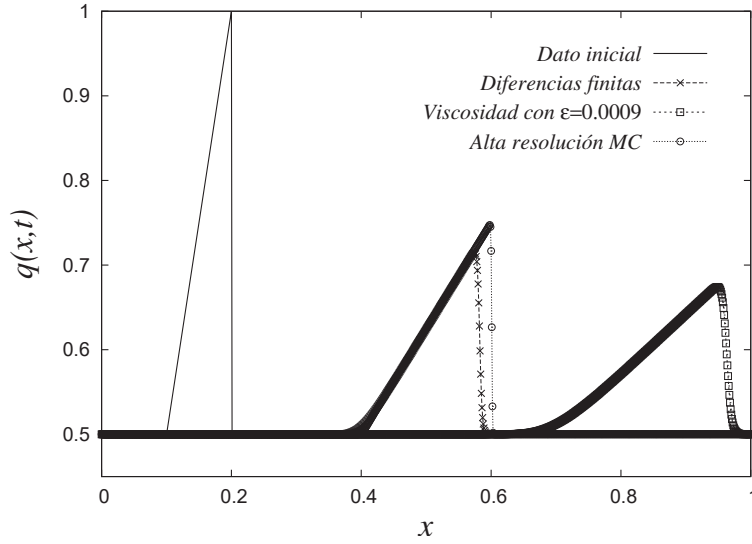


Figura 4.3: Solución a $t = 0.6$ para la Ecuación de Burgers con $\Delta x = 0.001$, $\Delta t = 0.0005$ y condiciones de frontera absorbente. El dato inicial es un segmento de recta de pendiente pronunciada (línea continua). ($\times$) es la solución a la forma cuasilineal (4.9) usando diferencias finitas. ($\square$) es la solución a (4.13) con $\epsilon = 0.0009$ con diferencias finitas también. ($\odot$) es la solución a (4.11) usando el método de Alta Resolución con un limitador MC.

Pseudocódigo para resolver la ecuación de Burgers

Es exactamente el mismo que aquel presentado en la sección (4.1.1) para resolver la ecuación de advección, excepto por un cambio en el segundo punto del quinto paso, el cual debe decir: *Se calcula el lado derecho de (4.15) usando las pendientes previamente obtenidas. En este mismo paso se aplican las condiciones de frontera.*

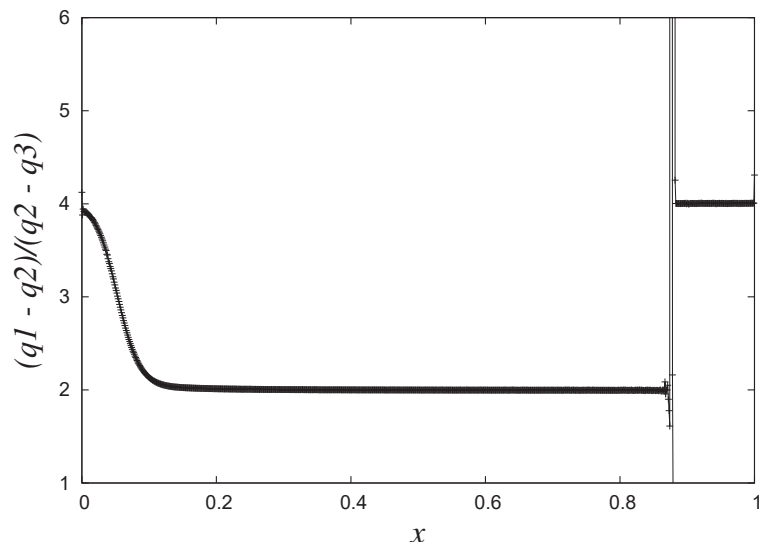


Figura 4.4: Usando el limitador MC, condiciones de frontera absorbente y un perfil gaussiano centrado en $x = 0.35$ como dato inicial, se calcula la solución numérica para la Ecuación de Burgers en su forma conservativa para tres resoluciones diferentes: $\Delta x = 0.001 \equiv \Delta x_0$ (q_1), $\Delta x = \Delta x_0/2$ (q_2) y $\Delta x = \Delta x_0/4$ (q_4). Se comparan las soluciones a $t = 0.6$ según (3.4) y se obtiene un factor de convergencia de segundo orden en regiones suaves. Notar que a partir de la frontera izquierda se propaga un error de primer orden debido a las extrapolaciones lineales.

Capítulo 5

Sistemas hiperbólicos lineales de leyes de conservación

Existen varias clasificaciones de los sistemas de ecuaciones diferenciales escritos en forma conservativa, una en particular es:

- *Sistemas lineales*, cuyo flujo es un vector de componentes lineales en q .
- *Sistemas no-lineales*, cuyo flujo es un vector de funciones no-lineales en q .

Dentro de los sistemas lineales se puede hacer otra distinción: sistemas de coeficientes constantes o de coeficientes variables, es decir, las entradas de la matriz del sistema pueden depender o no, de la posición x . En los sistemas no-lineales no puede hacerse esta subclasificación, pues al depender el flujo de las variables conservadas, necesariamente la matriz tiene cantidades que dependen implícitamente de x . En las tres primeras secciones de este capítulo se aplican métodos de Alta Resolución para resolver exclusivamente sistemas lineales de coeficientes constantes escritos en forma conservativa.

En la primera sección se presenta una estrategia *sencilla* que consiste en *desacoplar* las ecuaciones del sistema y resolverlas por separado, utilizando los métodos descritos en los capítulos anteriores. En la segunda sección se hace una reformulación del método de Alta Resolución presentado para la ecuación de advección, que luego se generaliza para sistemas lineales de coeficientes constantes, sin necesidad de desacoplarlos; ello será la base para formular el método que resuelva sistemas no-lineales: la intención más ambiciosa de este trabajo.

Los métodos descritos en la primera y segunda secciones, ambos son de Alta Resolución, basados en discretizaciones en volúmenes finitos. La justificación de invertir tiempo en implementar el primero, es la posibilidad de comparar los resultados con las nuevas reformulaciones y generalizaciones que están por exponerse, y así enfrentar a los sistemas no-lineales con un buen precedente.

5.1 Desacomplamiento del sistema de coeficientes constantes

Considérese el sistema de m ecuaciones escalares:

$$\partial_t \mathbf{q} + \partial_x (A \mathbf{q}) = \mathbf{0} . \quad (5.1)$$

$\mathbf{q}$ es el vector de las variables conservadas y tiene m entradas, naturalmente. A es una matriz de tamaño $m \times m$. Por su analogía con la ecuación (2.4), se dice que el sistema está escrito en *forma conservativa*, donde $A \mathbf{q}$ se llama *vector de flujo*.

Asumiendo que el sistema es *hiperbólico* (lo cual no es una suposición débil), es decir, que A es diagonalizable con eigenvalores reales, entonces A puede descomponerse como el producto: $A = P \Lambda P^{-1}$ donde $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_m)$, con λ_k el k -ésimo eigenvalor de A , y P una matriz cuyas columnas son los eigenvectores de A . El sistema se llama *estrictamente hiperbólico* si $\lambda_i \neq \lambda_j$ cada vez que $i \neq j$ [3].

Entonces el sistema en (5.1) se puede escribir como:

$$\partial_t \mathbf{q} + \partial_x (P \Lambda P^{-1} \mathbf{q}) = \mathbf{0} .$$

Si se define $\mathbf{w} \equiv P^{-1} \mathbf{q}$:

$$\begin{aligned} \partial_t \mathbf{q} + \partial_x (P \Lambda \mathbf{w}) &= \mathbf{0} \\ \partial_t \mathbf{q} + (\partial_x P) \Lambda \mathbf{w} + P (\partial_x \Lambda) \mathbf{w} + P \Lambda (\partial_x \mathbf{w}) &= \mathbf{0} , \end{aligned} \quad (5.2)$$

multiplicando por P^{-1} a la izquierda de cada término:

$$\begin{aligned} P^{-1} \partial_t \mathbf{q} + P^{-1} (\partial_x P) \Lambda \mathbf{w} + (\partial_x \Lambda) \mathbf{w} + \Lambda (\partial_x \mathbf{w}) &= \mathbf{0} \\ P^{-1} \partial_t \mathbf{q} + \Lambda (\partial_x \mathbf{w}) &= -(\partial_x \Lambda) \mathbf{w} - P^{-1} (\partial_x P) \Lambda \mathbf{w} . \end{aligned}$$

En el caso de coeficientes constantes, los eigenvectores de A y las entradas de Λ son independientes de x y de t , así que lo anterior se reduce a:

$$\partial_t \mathbf{w} + \Lambda \partial_x \mathbf{w} = \mathbf{0} , \quad (5.3)$$

que gracias a que Λ es diagonal, es un sistema desacoplado de m ecuaciones de advección para cada una de las componentes de $\mathbf{w}$, las cuales se llaman *variables características* del sistema (notar que son proyecciones del vector advectado en los eigenvectores de A).

Cada ecuación debe resolverse por separado usando alguna de las dos opciones en (4.7),

dependiendo del signo de la velocidad de advección (que en este caso serán los eigenvalores de A), e implementando algún limitador de pendientes de los expuestos en esta Tesis para tener un método de Alta Resolución.

5.2 Generalización del Método de Alta Resolución para sistemas de coeficientes constantes

Una manera alternativa de abordar el sistema (5.1), es escribirlo desde el principio como:

$$\partial_t \mathbf{q} + A(\partial_x \mathbf{q}) = -(\partial_x A) \mathbf{q} , \quad (5.4)$$

que cuando A es constante, se reduce a:

$$\partial_t \mathbf{q} + A(\partial_x \mathbf{q}) = \mathbf{0} . \quad (5.5)$$

Para resolverlo sin tener que desacoplarlo, es preciso generalizar el método de Alta Resolución mostrado para la ecuación de advección. Antes es conveniente reformular este método de tal manera que los flujos no se escriban más en términos de las pendientes, sino de un parámetro más conveniente, el cual ofrece importantes ventajas al resolver sistemas de ecuaciones.

5.2.1 Reformulación del método de Alta Resolución en términos de limitadores de flujos

Considérese el flujo numérico para la ecuación escalar de advección (4.5) cuando $u \geq 0$; tómese una pendiente fija atrasada, es decir, $\sigma_{i-1}^n = (Q_i^n - Q_{i-1}^n)/\Delta x$. Sea $Q_i^n - Q_{i-1}^n \equiv \Delta Q_i^n$. Entonces el flujo puede reescribirse así:

$$F_i^n = u Q_{i-1}^n + \frac{u}{2} \left(1 - \frac{\Delta t}{\Delta x} u \right) \Delta Q_i^n , \quad u \geq 0 . \quad (5.6)$$

Sea ϕ una función escalar tal que $\phi = \phi(\theta_i^n)$, donde $\theta_i^n = \left(\frac{\Delta Q_{i-1}^n}{\Delta Q_i^n} \right)$. ϕ se llama *función limitadora de flujo*. En general este flujo numérico se escribe como:

$$F_i^n = u Q_{i-1}^n + \frac{u}{2} \left(1 - \frac{\Delta t}{\Delta x} u \right) \phi(\theta_i^n) \Delta Q_i^n , \quad u \geq 0 . \quad (5.7)$$

Ciertamente el parámetro θ_i^n podría tener problemas cuando su denominador sea cero, sin embargo siempre se procura que ese denominador se cancele al multiplicar θ_i^n con algún otro factor.

Mientras tanto en la ecuación anterior se tiene la libertad de elegir $\phi(\theta_i^n)$, como antes podían elegirse las pendientes σ_i^n , por ejemplo, elegir $\phi(\theta_i^n) \equiv 0$ equivale a tomar pendientes cero en (4.5) y el método se reduce a las clásicas diferencias finitas adelantadas. Otra posibilidad evidente es tomar $\phi(\theta_i^n) \equiv 1$, lo que reduce (5.7) a (5.6), que es lo mismo que (4.5) con pendiente fija atrasada $\sigma_{i-1}^n = (Q_i^n - Q_{i-1}^n)/\Delta x$.

Otra alternativa es definir $\phi(\theta_i^n) \equiv \theta_i^n$ para que (5.7) se vea:

$$F_i^n = uQ_{i-1}^n + \frac{u}{2} \left(1 - \frac{\Delta t}{\Delta x} u \right) \Delta Q_{i-1}^n ,$$

(notar que el denominador de θ_i^n , ΔQ_i^n , efectivamente se cancela y no hay problemas en la implementación numérica), lo que es equivalente a tomar (4.5) con una pendiente adelantada $\sigma_{i-1}^n = (Q_{i-1}^n - Q_{i-2}^n)/\Delta x$.

Una opción menos obvia es tomar $\phi(\theta_i^n) \equiv \frac{1}{2}[1 + \theta_i^n]$, que devuelve:

$$F_i^n = uQ_{i-1}^n + \frac{u}{2} \left(1 - \frac{\Delta t}{\Delta x} u \right) \frac{1}{2} (Q_i^n - Q_{i-2}^n) ,$$

y que equivale a usar pendientes centradas de Fromm en (4.5). A todo esto, lo que se busca es limitar el flujo F_i^n , y eso se logra limitando $\phi(\theta_i^n)$: una manera sencilla es escribiendo el limitador de Alta Resolución *minmod* (3.2) en términos de θ_i^n , esto es:

$$\phi(\theta_i^n) = \minmod(1, \theta_i^n) .$$

El limitador MC adoptaría la forma:

$$\phi(\theta_i^n) = \minmod\left(\frac{1 + \theta_i^n}{2}, 2, 2\theta_i^n\right) = \max\left(0, \min\left[\frac{1 + \theta_i^n}{2}, 2, 2\theta_i^n\right]\right) .$$

Ha habido varias propuestas, algunas más finas, resultantes de meticolosos análisis, pero a fin de cuentas sus resultados no se alejan mucho de los obtenidos con el limitador MC. Un ejemplo muy popular es el limitador *Superbee* de Roe [7]:

$$\phi(\theta_i^n) = \max(0, [\min[1, 2\theta_i^n], \min[2, \theta_i^n]]) ,$$

o el limitador de van Leer:

$$\phi(\theta_i^n) = \frac{\theta_i^n + |\theta_i^n|}{1 + |\theta_i^n|} .$$

En la literatura de Análisis Numérico ([1], [2], [7]), puede encontrarse una amplia discusión sobre las propiedades algebraicas que estos limitadores deben tener para que los métodos que los usan cumplan con la condición de Variación Total Decreciente y sean estables y

precisos a segundo orden. En general se aceptan los limitadores MC y de van Leer como las mejores opciones.

Después de todo, vale la pena aclarar que se trata de una *reformulación del método*, o si se prefiere, de una *reinterpretación*, y que por lo tanto no debe de haber ninguna diferencia en los resultados numéricos. No hay pasos oscuros: una vez que se tiene la función lineal reconstruida $\tilde{q}_i^n$, se introduce σ_i^n de manera natural, y se calculan los flujos en términos de ellas: $F_i^n = F_i^n(\sigma_i^n)$. Entonces hay dos opciones: limitar σ ó escribir el flujo en términos de ϕ , $F_i^n = F_i^n(\phi)$ y luego limitarlo.

El camino a seguir para reformular el método en términos de limitadores de flujos para la ecuación de advección con $u < 0$, es el mismo que seguimos cuando $u \geq 0$, sólo que en este caso hay que asegurarse de empezar con (4.6) en lugar de (4.5). Primero se toma una pendiente fija atrasada, pero es crucial entender que cuando la velocidad tiene signo negativo, la dirección en la que se propaga la información es de derecha a izquierda, por lo tanto los conceptos *atrasado* y *adelantado* se invierten en significado, siendo ahora la pendiente atrasada en la i -ésima celda: $\sigma_i^n = (Q_i^n - Q_{i-1}^n)/\Delta x$. En analogía a (5.6) se obtiene:

$$F_i^n = uQ_i^n - \frac{u}{2} \left(1 + \frac{\Delta t}{\Delta x} u \right) \Delta Q_i^n, \quad u < 0. \quad (5.8)$$

Para tener resultados consistentes (en particular con los casos analizados arriba para $\phi(\theta_i^n) \equiv \theta_i^n$ y $\phi(\theta_i^n) \equiv \frac{1}{2}[1 + \theta_i^n]$) en este caso hay que definir $\theta_i^n = \frac{\Delta Q_{i+1}^n}{\Delta Q_i^n}$ y escribir:

$$F_i^n = uQ_i^n - \frac{u}{2} \left(1 + \frac{\Delta t}{\Delta x} u \right) \phi(\theta_i^n) \Delta Q_i^n, \quad u < 0. \quad (5.9)$$

El inconveniente práctico de tener que considerar siempre dos casos por separado ($u \geq 0$ y $u < 0$), debe ser solucionado antes de pretender generalizar el método para sistemas de ecuaciones. Por lo pronto, para simplificar la escritura de los resultados más relevantes, se introduce un par de cantidades: $\delta_{i[+]}^n$ y $\delta_{i[-]}^n$, ambas iguales a $\phi(\theta_i^n) \Delta Q_i^n$, pero cuyo subíndice + y - indica si el numerador de θ_i^n es ΔQ_{i-1}^n ó ΔQ_{i+1}^n , respectivamente, lo cual depende directamente del signo de u . De esta manera se tiene que el flujo numérico para la ecuación de advección en términos de limitadores de flujo es en resumen:

$$F_i^n = \begin{cases} uQ_{i-1}^n + \frac{u}{2} \left(1 - \frac{\Delta t}{\Delta x} u \right) \delta_{i[+]}^n, & \text{si } u \geq 0 \\ uQ_i^n - \frac{u}{2} \left(1 + \frac{\Delta t}{\Delta x} u \right) \delta_{i[-]}^n, & \text{si } u < 0. \end{cases} \quad (5.10)$$

Para lo sucesivo es necesario escribir el flujo anterior en una sola expresión, haciendo previamente un par de definiciones convenientes:

$$\begin{aligned} u^+ &= \max(u, 0) \\ u^- &= \min(u, 0) . \end{aligned}$$

Notar que: $u = u^+ + u^-$ y $|u| = u^+ - u^-$.

En términos de u^+ y u^- , el flujo más general para la ecuación escalar de advección es:

$$\begin{aligned} F_i^n &= u^+ Q_{i-1}^n + u^- Q_i^n + \frac{1}{2} \left[u^+ \left(1 - u^+ \frac{\Delta t}{\Delta x} \right) \delta_{i[+]}^n - u^- \left(1 + u^- \frac{\Delta t}{\Delta x} \right) \delta_{i[-]}^n \right] \\ &= u^+ Q_{i-1}^n + u^- Q_i^n + \frac{1}{2} \left[u^+ \left(1 - u^+ \frac{\Delta t}{\Delta x} \right) - u^- \left(1 + u^- \frac{\Delta t}{\Delta x} \right) \right] \delta_{i[sgn]}^n \\ &= u^+ Q_{i-1}^n + u^- Q_i^n + \frac{1}{2} \left[(u^+ - u^-) - (u^{+2} + u^{-2}) \frac{\Delta t}{\Delta x} \right] \delta_{i[sgn]}^n \\ &= u^+ Q_{i-1}^n + u^- Q_i^n + \frac{1}{2} \left[|u| - |u|^2 \frac{\Delta t}{\Delta x} \right] \delta_{i[sgn]}^n \\ &= u^+ Q_{i-1}^n + u^- Q_i^n + \frac{1}{2} |u| \left(1 - |u| \frac{\Delta t}{\Delta x} \right) \delta_{i[sgn]}^n , \end{aligned} \quad (5.11)$$

donde

$$sgn = \begin{cases} +, & \text{si } u \geq 0 \\ -, & \text{si } u < 0 . \end{cases}$$

Este importante resultado será el punto de partida para el trabajo con sistemas de ecuaciones.

Finalmente, la forma en diferencia de flujos (2.8) toma la forma:

$$Q_i^{n+1} = Q_i^n - \frac{\Delta t}{\Delta x} (u^+ \Delta Q_i^n + u^- \Delta Q_{i+1}^n) - \frac{\Delta t}{\Delta x} \frac{|u|}{2} \left(1 - |u| \frac{\Delta t}{\Delta x} \right) (\delta_{(i+1)[sgn]}^n - \delta_{i[sgn]}^n) . \quad (5.12)$$

5.2.2 Generalización

El objetivo ahora es generalizar (5.12) a una expresión vectorial que corresponda al sistema en (5.5), donde es claro que A tiene el papel que u tenía en (4.1). El primer paso es escribir un flujo numérico vectorial que generalice la expresión (5.11), para eso es necesario definir las cantidades matriciales análogas a los escalares u^+ , u^- y $|u|$, a saber: A^+ , A^- y $|A|$ (que a priori no tiene por qué llamarse o entenderse como *valor absoluto de A*, pues éste no es un concepto con un sentido claro). De antemano se sabe que el sistema en cuestión puede

desacoplarse en m ecuaciones de advección cuyas velocidades son las entradas de Λ , por lo tanto, la principal consideración en la tarea de definir A^+ , A^- y $|A|$, debe ser el signo de los eigenvalores de A . Sean Λ^+ y Λ^- matrices diagonales que contienen los eigenvalores positivos y negativos de A , respectivamente. Entonces $\Lambda = \Lambda^+ + \Lambda^-$, luego la definición de A^+ y A^- es natural:

$$A = P\Lambda P^{-1} = P(\Lambda^+ + \Lambda^-)P^{-1} = P\Lambda^+P^{-1} + P\Lambda^-P^{-1} \equiv A^+ + A^-.$$

Análogamente a $|u| = u^+ - u^-$, se define la matriz $|A| \equiv A^+ - A^-$, y en estos términos puede escribirse una expresión general para el flujo numérico del sistema, a partir de (5.11):

$$\mathbf{F}_i^n = A^+ \mathbf{Q}_{i-1}^n + A^- \mathbf{Q}_i^n + \frac{1}{2}|A| \left(Id - |A| \frac{\Delta t}{\Delta x} \right) \delta_{i[sgn]}^n, \quad (5.13)$$

donde Id es la identidad en el espacio de las matrices $m \times m$.

En la expresión anterior la cantidad vectorial $\delta_{i[sgn]}^n$ carece de un sentido claro, pues en analogía con (5.11), ésta debe ser proporcional al salto $\Delta \mathbf{Q}_i^n$ en la dirección de propagación de la información:

$$\delta_{i[sgn]}^n \sim \Delta \mathbf{Q}_i^n.$$

Pero (5.13) es un flujo generalizado para un sistema de ecuaciones acoplado, y por lo tanto la dirección en la que se propaga la información (y que es precisamente la dirección del vector $\delta_{i[sgn]}^n$) es más bien una suma vectorial de las direcciones características de propagación, por eso se escribe:

$$\delta_{i[sgn]}^n \sim \Delta \mathbf{Q}_i^n = \sum_{k=1}^m \alpha_{ki}^n \mathbf{r}_k, \quad (5.14)$$

donde los $\mathbf{r}_k$ son los eigenvectores de la matriz A (y columnas de P , por supuesto). Este *desacoplamiento* de la información en *eigencomponentes* nos permite retomar el hecho de que el k -ésimo eigenvalor λ_k es la velocidad de advección de la k -ésima componente de la información propagada, cada una de las cuales debe ser limitada por separado, y para hacerlo se dejan m parámetros libres en $\delta_{i[sgn]}^n$, así que se define:

$$\delta_{i[sgn]}^n = \sum_{k=1}^m \phi(\theta_{ki}^n) \alpha_{ki}^n \mathbf{r}_k,$$

con

$$\theta_{ki}^n = \frac{\alpha_{kI}^n}{\alpha_{ki}^n} \quad \text{donde} \quad I = \begin{cases} i-1 & \text{si } \lambda_k \geq 0, \\ i+1 & \text{si } \lambda_k < 0. \end{cases}$$

ϕ puede ser cualquiera de las funciones limitadoras expuestas en el apartado anterior.

Finalmente (5.13) puede escribirse como sigue:

$$\mathbf{F}_i^n = A^+ \mathbf{Q}_{i-1}^n + A^- \mathbf{Q}_i^n + \tilde{\mathbf{F}}_i^n, \quad (5.15)$$

donde

$$\tilde{\mathbf{F}}_i^n = \frac{1}{2}|A| \left(Id - |A| \frac{\Delta t}{\Delta x} \right) \sum_{k=1}^m \phi(\theta_{ki}^n) \alpha_{ki}^n \mathbf{r}_k,$$

y la forma en diferencia de flujos (2.8) resulta en:

$$\mathbf{Q}_i^{n+1} = \mathbf{Q}_i^n - \frac{\Delta t}{\Delta x} (A^+ \Delta \mathbf{Q}_i^n + A^- \Delta \mathbf{Q}_{i+1}^n) - \frac{\Delta t}{\Delta x} (\tilde{\mathbf{F}}_{i+1}^n - \tilde{\mathbf{F}}_i^n), \quad (5.16)$$

que por su generalidad y alcance, representa la segunda fórmula, junto con (5.12), más importante en este trabajo, después de (2.8).

A estas alturas es muy reconfortante (y necesario) comprobar que si $A = u$, entonces el único eigenvector es $\mathbf{r}_1 = 1$ y el único eigenvalor es precisamente $\lambda_1 = u$. α_{1i}^n se calcula mediante:

$$\Delta Q_i^n = \Delta \mathbf{Q}_i^n = \sum_{k=1}^1 \alpha_{ki}^n \mathbf{r}_k = \alpha_{1i}^n,$$

luego:

$$\begin{aligned} \tilde{\mathbf{F}}_i^n &= \frac{1}{2}|u| \left(1 - |u| \frac{\Delta t}{\Delta x} \right) \sum_{k=1}^1 \phi(\theta_{ki}^n) \alpha_{ki}^n \mathbf{r}_k \\ &= \frac{1}{2}|u| \left(1 - |u| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{1i}^n) \alpha_{1i}^n \\ &= \frac{1}{2}|u| \left(1 - |u| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{1i}^n) \Delta Q_i^n \\ &= \frac{1}{2}|u| \left(1 - |u| \frac{\Delta t}{\Delta x} \right) \delta_{i[\text{sgn}]}^n. \end{aligned}$$

Y como era de esperarse, con en esta breve y elegante maniobra se comprueba que (5.16) se reduce a (5.12).

5.2.3 Implementación

Para propósitos prácticos, el resultado en (5.16) resulta muy incómodo, por lo que se aprovechan algunas relaciones presentadas a continuación. Antes defínanse, para $k = 1, 2, \dots, m$:

$$\begin{aligned}\lambda_k^+ &= \max(\lambda_k, 0) \\ \lambda_k^- &= \min(\lambda_k, 0) .\end{aligned}$$

Como se estableció en (5.14), $\Delta \mathbf{Q}_i^n = \sum_{k=1}^m \alpha_{ki}^n \mathbf{r}_k$, entonces:

$$A^+ \Delta \mathbf{Q}_i^n = A^+ \sum_{k=1}^m \alpha_{ki}^n \mathbf{r}_k = \sum_{k=1}^m \alpha_{ki}^n A^+ \mathbf{r}_k = \sum_{k=1}^m \alpha_{ki}^n \lambda_k^+ \mathbf{r}_k = \sum_{k=1}^m \lambda_k^+ \alpha_{ki}^n \mathbf{r}_k .$$

De la misma manera:

$$A^- \Delta \mathbf{Q}_{i+1}^n = \sum_{k=1}^m \lambda_k^- \alpha_{k(i+1)}^n \mathbf{r}_k ,$$

y

$$\tilde{\mathbf{F}}_i^n = \frac{1}{2} \sum_{k=1}^m |\lambda_k| \left(1 - |\lambda_k| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{ki}^n) \alpha_{ki}^n \mathbf{r}_k ,$$

donde las k α 's se calculan de $\Delta \mathbf{Q}_i^n = \mathbf{Q}_i^n - \mathbf{Q}_{i-1}^n = \sum_{k=1}^m \alpha_{ki}^n \mathbf{r}_k$. De un bonito resultado del álgebra lineal [6], cuando A es simétrica, los eigenvectores son ortogonales y la fórmula es muy directa:

$$\alpha_{ki}^n = \frac{1}{|\mathbf{r}_k|^2} (\Delta \mathbf{Q}_i^n \cdot \mathbf{r}_k) , \quad (5.17)$$

en otro caso hay que trabajar un poco más y resolver el sistema para α_{ki}^n usando algún método algebraico.

5.3 Pruebas con la Ecuación de Onda

La solución numérica de la ecuación de onda es un excelente ejercicio como preámbulo al trabajo en Relatividad Numérica, sobre todo porque además de ejemplificar la implementación de métodos clásicos y muy útiles como el de diferencias finitas para resolver un problema de datos iniciales con condiciones de frontera, permite introducir conceptos fundamentales de la descomposición 3+1 del espacio-tiempo en Relatividad General, en particular los de *función de lapso* y *vector de corrimiento*, así se tiene una completa li-

bertad de norma y se puede experimentar con varias situaciones comunes en Relatividad Numérica. Una amplia discusión acerca de la importancia de la solución de la ecuación de onda, así como de las técnicas más usuales para hacerlo puede encontrarse en [11]. En esta Tesis sólo se dirá lo justo acerca del problema a resolver y en cambio se hará una discusión a detalle sobre la implementación del Método de Alta Resolución que se ha presentado.

Lo justo es anotar que la ecuación de onda en general se escribe en términos del operador diferencial D'Alambertiano, que actúa sobre campos escalares ϕ , y que dadas las componentes de una métrica $g_{\mu\nu}$ en un espacio-tiempo arbitrario se define como $\square\phi = \frac{1}{\sqrt{-g}}\partial_\mu[\sqrt{-g}g^{\mu\nu}\partial_\nu\phi]$.

La ecuación $\square\phi = 0$ se llama *ecuación de onda* porque los datos iniciales para el campo ϕ que la satisfaga, se tomarán como una perturbación que se propagará como una onda viajera en el espacio-tiempo. Si tal ecuación fuera no-homogénea, el término extra representaría una fuente que genera continuamente las perturbaciones que se propagan como ondas. Aquí se resuelve sólo la ecuación homogénea.

En un espacio-tiempo en 1+1 dimensiones, la representación matricial de la métrica más general es:

$$g_{\mu\nu} = \begin{pmatrix} (-a^2 + b^2) & b \\ b & 1 \end{pmatrix},$$

donde a se llama *función de lapso* y b , *vector de corrimiento*, que en este caso es un vector con una sola componente, o sea, un escalar. El determinante de la métrica se calcula mentalmente: $g = -a^2$, y entonces la ecuación de onda general en 1+1 se escribe explícitamente como (ver detalles del álgebra en [3]):

$$\begin{aligned} 0 &= \square\phi \\ &= -\partial_t \left[\frac{1}{a} \partial_t \phi - \frac{b}{a} \partial_x \phi \right] + \partial_x \left[\frac{b}{a} \partial_t \phi + a \left(1 - \frac{b^2}{a^2} \right) \partial_x \phi \right]. \end{aligned} \quad (5.18)$$

Esta ecuación de segundo orden se puede escribir como un sistema de dos ecuaciones de primer orden; basta definir un par de nuevas variables:

$$\begin{aligned} \psi &:= \partial_x \phi \\ \pi &:= \frac{1}{a}(\partial_t \phi - b \partial_x \phi) = \frac{1}{a}(\partial_t \phi - b \psi), \end{aligned}$$

notar que π es el argumento de la derivada temporal en (5.18), y de esa misma ecuación se tiene $\partial_t \pi = \partial_x \left[\frac{b}{a} \partial_t \phi + a \left(1 - \frac{b^2}{a^2} \right) \psi \right]$, que rápidamente se reduce a $\partial_t \pi = \partial_x (a\psi + b\pi)$. Si se despeja $\partial_t \phi$ de la definición de π , y asumiendo que ϕ tiene derivadas continuas ($\partial_t \psi =$

$\partial_x(\partial_t \phi)$), entonces $\partial_t \psi = \partial_x(a\pi + b\psi)$. En síntesis el sistema para ψ y π es:

$$\begin{aligned}\partial_t \pi &= \partial_x(a\psi + b\pi) \\ \partial_t \psi &= \partial_x(a\pi + b\psi) .\end{aligned}\tag{5.19}$$

Una vez resuelto, la función original ϕ se recupera usando la definición de π .

Para los fines de este trabajo es importante escribir el sistema (5.19) en la forma matricial:

$$\partial_t \mathbf{u} + A \partial_x \mathbf{u} = -(\partial_x A) \mathbf{u} ,\tag{5.20}$$

donde $\mathbf{u} = (\pi, \psi)^T$ y:

$$A = -\begin{pmatrix} b & a \\ a & b \end{pmatrix} .$$

5.3.1 Desacoplando el sistema

A continuación se resuelve el sistema (5.20) desacopándolo en dos ecuaciones de advección, para ello se comienza por calcular los eigenvalores y eigenvectores de A para escribir $A = P \Lambda P^{-1}$. Sean λ_π y λ_ψ dichos eigenvalores; no es difícil encontrar que:

$$\begin{aligned}\lambda_\pi &= a - b \\ \lambda_\psi &= -(a + b) ,\end{aligned}$$

y

$$P = \begin{pmatrix} 1 & 1 \\ -1 & 1 \end{pmatrix} , \quad P^{-1} = \frac{1}{2} \begin{pmatrix} 1 & -1 \\ 1 & 1 \end{pmatrix} ,$$

así que (5.20) ahora se ve como:

$$\partial_t \mathbf{w} + \Lambda \partial_x \mathbf{w} = -(\partial_x \Lambda) \mathbf{w}\tag{5.21}$$

con

$$\mathbf{w} = P^{-1} \mathbf{u} = \frac{1}{2}(\pi - \psi, \pi + \psi)^T \equiv (R, L)^T ,$$

R y L son las llamadas variables características del sistema, y corresponden a los dos *modos* (componentes de onda) en los que se divide el perfil de los datos iniciales. Cada uno de

éstos se propaga como una onda viajera independiente, uno a la derecha y otro a la izquierda del dominio, respectivamente.

Si asumimos una matriz constante (es decir, a y b constantes), entonces (5.21) se reduce a:

$$\partial_t \mathbf{w} + \partial_x (\Lambda \mathbf{w}) = 0 ,$$

que explícitamente puede escribirse como el siguiente par de ecuaciones escalares de advección desacopladas, cuyas velocidades son precisamente los eigenvalores de A :

$$\begin{aligned} \partial_t R + (-b + a) \partial_x R &= 0 \\ \partial_t L + (-b - a) \partial_x L &= 0 . \end{aligned} \quad (5.22)$$

Éstas se resuelven por separado usando (4.7), que es un método formulado para limitar las pendientes de las funciones reconstruidas; o bien pueden resolverse (también por separado, claro) usando (5.12), que es un método formulado para limitar flujos y no pendientes. Los resultados son los mismos.

Para calcular π y ψ hay que resolver de antemano el sistema algebraico:

$$\begin{aligned} \frac{1}{2}[\pi - \psi] &= R \\ \frac{1}{2}[\pi + \psi] &= L , \end{aligned}$$

así, una vez resuelto el sistema en (5.22) para R y L , basta hacer un cálculo sencillo:

$$\begin{aligned} \pi &= L + R \\ \psi &= L - R , \end{aligned}$$

Finalmente, ϕ se recupera resolviendo: $\partial_t \phi = a\pi + b\psi$.

Implementación y resultados

Dado el perfil inicial ϕ_i^0 en cada punto de la malla, se calcula ψ_i^0 usando diferencias finitas centradas: $\psi_i^0 = \partial_x \phi_i^0$ (las derivadas en los extremos del dominio están dadas como se establece en [3]). Como a $t = 0$ la derivada temporal de ϕ es cero, entonces $\pi_i^0 \equiv 0$.

Enseguida se calculan $R_i^0 = \frac{1}{2}(\pi_i^0 - \psi_i^0)$ y $L_i^0 = \frac{1}{2}(\pi_i^0 + \psi_i^0)$, luego los promedios en cada celda: $\bar{R}_i^0 = \frac{1}{2}(R_i^0 + R_{i+1}^0)$ y $\bar{L}_i^0 = \frac{1}{2}(L_i^0 + L_{i+1}^0)$.

Para $\bar{R}$, utilizando (4.7), se itera usando el integrador de Runge-Kutta de tercer orden:

$$\begin{aligned}\bar{R}_i^{n+1} &= \bar{R}_i^n - (-b + a) \frac{\Delta t}{\Delta x} \left[(\bar{R}_i^n - \bar{R}_{i-1}^n) + \frac{1}{2}(\Delta x - (-b + a)\Delta t)(\sigma_i^n - \sigma_{i-1}^n) \right], \quad (-b + a) \geq 0, \\ \bar{R}_i^{n+1} &= \bar{R}_i^n - (-b + a) \frac{\Delta t}{\Delta x} \left[(\bar{R}_{i+1}^n - \bar{R}_i^n) - \frac{1}{2}(\Delta x + (-b + a)\Delta t)(\sigma_{i+1}^n - \sigma_i^n) \right], \quad (-b + a) < 0,\end{aligned}$$

usando el limitador de pendientes MC, según (3.3):

$$\sigma_i^n = \minmod \left[\left(\frac{\bar{R}_{i+1}^n - \bar{R}_{i-1}^n}{2\Delta x} \right), 2 \left(\frac{\bar{R}_i^n - \bar{R}_{i-1}^n}{\Delta x} \right), 2 \left(\frac{\bar{R}_{i+1}^n - \bar{R}_i^n}{\Delta x} \right) \right],$$

así mismo, para $\bar{L}$:

$$\begin{aligned}\bar{L}_i^{n+1} &= \bar{L}_i^n - (-b - a) \frac{\Delta t}{\Delta x} \left[(\bar{L}_i^n - \bar{L}_{i-1}^n) + \frac{1}{2}(\Delta x - (-b - a)\Delta t)(\sigma_i^n - \sigma_{i-1}^n) \right], \quad (-b - a) \geq 0, \\ \bar{L}_i^{n+1} &= \bar{L}_i^n - (-b - a) \frac{\Delta t}{\Delta x} \left[(\bar{L}_{i+1}^n - \bar{L}_i^n) - \frac{1}{2}(\Delta x + (-b - a)\Delta t)(\sigma_{i+1}^n - \sigma_i^n) \right], \quad (-b - a) < 0,\end{aligned}$$

con

$$\sigma_i^n = \minmod \left[\left(\frac{\bar{L}_{i+1}^n - \bar{L}_{i-1}^n}{2\Delta x} \right), 2 \left(\frac{\bar{L}_i^n - \bar{L}_{i-1}^n}{\Delta x} \right), 2 \left(\frac{\bar{L}_{i+1}^n - \bar{L}_i^n}{\Delta x} \right) \right].$$

Las condiciones de frontera, ya sean periódicas o absorbentes (en el caso de la ecuación de onda se llaman de onda saliente), se imponen en cada paso temporal. Al final de cada paso también deben recuperarse R_i^n y L_i^n , es decir, los valores que corresponden a cada punto de la malla original, haciendo una interpolación lineal: $R_i^n = \bar{R}_{i-1}^n + \frac{1}{2}(\bar{R}_i^n - \bar{R}_{i-1}^n)$ y $L_i^n = \bar{L}_{i-1}^n + \frac{1}{2}(\bar{L}_i^n - \bar{L}_{i-1}^n)$. Luego se recuperan las variables características:

$$\begin{aligned}\pi_i^n &= \frac{1}{2}(L_i^n + R_i^n) \\ \psi_i^n &= \frac{1}{2}(L_i^n - R_i^n),\end{aligned}$$

y finalmente se calcula ϕ_i^n de $\partial_t \phi = a\pi + b\psi$, usando nuevamente el integrador de Runge-Kutte de tercer orden.

Se hicieron los cálculos para $a = 0.7$, $b = 0.2$, $\Delta x = 0.01$, $\Delta t = 0.005$, tomando como dato inicial un perfil gaussiano centrado en $x = 0.0$, de amplitud 1.0 y 0.1 de desviación. La Figura 5.1 muestra los resultados: notar que como era de esperarse, el perfil inicial se divide en dos modos independientes, de la misma amplitud, y que se propagan en direcciones contrarias y con velocidades distintas.

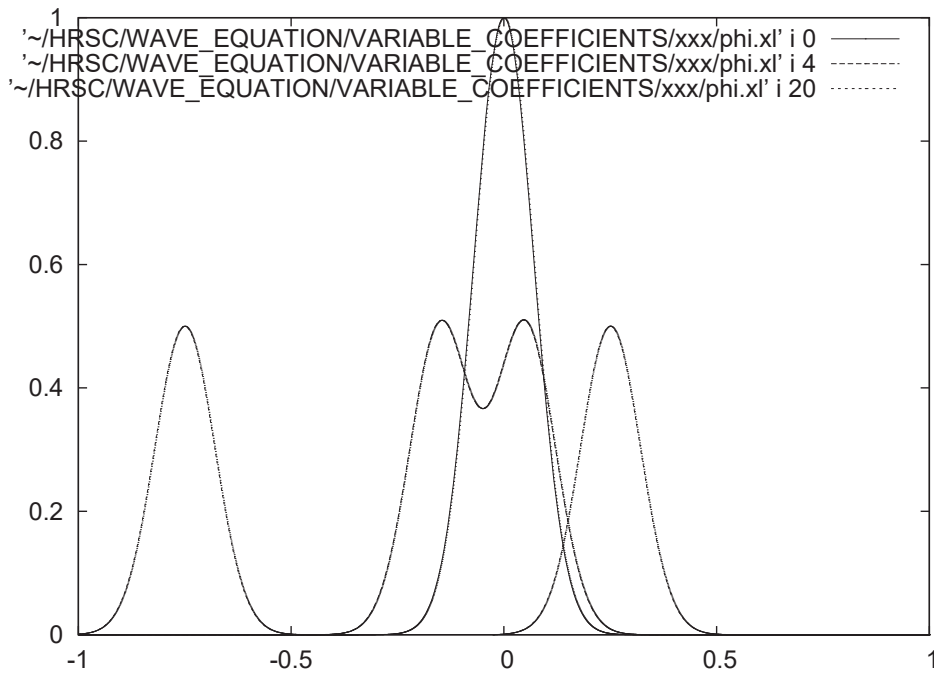


Figura 5.1: Solución de la Ecuación de Onda usando el limitador de pendientes MC, previamente desacoplado el sistema en un par de ecuaciones de advección. $\Delta x = 0.01$, $\Delta t = 0.005$. La línea continua son los datos iniciales; la de puntos largos es la solución a $t = 0.25$ y la de puntos cortos a $t = 0.6$.

5.3.2 Aplicando la generalización del método de Alta Resolución

Ahora se resuelve el sistema (5.20) sin necesidad de desacoplarlo, mediante el método de Alta Resolución generalizado. La implementación es directa: se comienza por dar los datos iniciales ϕ_i^0 , se calcula ψ_i^0 usando diferencias finitas centradas: $\psi_i^0 = \partial_x \phi_i^0$, y $\pi_i^0 \equiv 0$ por las razones antes mencionadas.

Enseguida se calculan los promedios de las variables características:

$$\begin{aligned}\bar{\pi}_i^0 &= \frac{1}{2}(\pi_i^0 + \pi_{i+1}^0) \\ \bar{\psi}_i^0 &= \frac{1}{2}(\psi_i^0 + \psi_{i+1}^0).\end{aligned}\tag{5.23}$$

Los eigenvalores de la matriz son $\lambda_\pi = (-b + a)$ y $\lambda_\psi = (-b - a)$. Los eigenvectores: $r_\pi = (1, -1)^T$ y $r_\psi = (1, 1)^T$.

Utilizando el integrador de Runge-Kutta de tercer orden, se itera según la forma vectorial general en diferencia de flujos (5.16), cuyas componentes (1) y (2) se calculan por separado:

$$\begin{aligned}\bar{\pi}_i^{n+1} &= \bar{\pi}_i^n - \frac{\Delta t}{\Delta x}(A^+ \Delta \bar{\pi}_i^{n(1)} + A^- \Delta \bar{\pi}_{i+1}^{n(1)}) - \frac{\Delta t}{\Delta x}(\tilde{\mathbf{F}}_{i+1}^{n(1)} - \tilde{\mathbf{F}}_i^{n(1)}) \\ \bar{\psi}_i^{n+1} &= \bar{\psi}_i^n - \frac{\Delta t}{\Delta x}(A^+ \Delta \bar{\psi}_i^{n(2)} + A^- \Delta \bar{\psi}_{i+1}^{n(2)}) - \frac{\Delta t}{\Delta x}(\tilde{\mathbf{F}}_{i+1}^{n(2)} - \tilde{\mathbf{F}}_i^{n(2)}),\end{aligned}\tag{5.24}$$

donde, siguiendo las fórmulas de la subsección (5.2.3):

$$\begin{aligned}A^+ \Delta \bar{\pi}_i^n &= \sum_k \lambda_k^+ \alpha_{ki}^n \mathbf{r}_k = \lambda_\pi^+ \alpha_{\pi i}^n \mathbf{r}_\pi + \lambda_\psi^+ \alpha_{\psi i}^n \mathbf{r}_\psi \\ &= \lambda_\pi^+ \alpha_{\pi i}^n (1, -1)^T + \lambda_\psi^+ \alpha_{\psi i}^n (1, 1)^T \\ &= (\lambda_\pi^+ \alpha_{\pi i}^n + \lambda_\psi^+ \alpha_{\psi i}^n, -\lambda_\pi^+ \alpha_{\pi i}^n + \lambda_\psi^+ \alpha_{\psi i}^n)^T.\end{aligned}$$

Análogamente: $A^- \Delta \bar{\pi}_{i+1}^n = (\lambda_\pi^- \alpha_{\pi(i+1)}^n + \lambda_\psi^- \alpha_{\psi(i+1)}^n, -\lambda_\pi^- \alpha_{\pi(i+1)}^n + \lambda_\psi^- \alpha_{\psi(i+1)}^n)^T$,
entonces:

$$\begin{aligned}
A^+ \Delta \tilde{\pi}_i^{n(1)} &= \lambda_\pi^+ \alpha_{\pi i}^n + \lambda_\psi^+ \alpha_{\psi i}^n \\
A^- \Delta \tilde{\pi}_{i+1}^{n(1)} &= \lambda_\pi^- \alpha_{\pi(i+1)}^n + \lambda_\psi^- \alpha_{\psi(i+1)}^n \\
A^+ \Delta \tilde{\psi}_i^{n(2)} &= -\lambda_\pi^+ \alpha_{\pi i}^n + \lambda_\psi^+ \alpha_{\psi i}^n \\
A^- \Delta \tilde{\psi}_{i+1}^{n(2)} &= -\lambda_\pi^- \alpha_{\pi(i+1)}^n + \lambda_\psi^- \alpha_{\psi(i+1)}^n ,
\end{aligned}$$

donde:

$$\begin{aligned}
\lambda_\pi^+ &= \max(\lambda_\pi, 0) = \max((-b + a), 0) \\
\lambda_\psi^+ &= \max((-b - a), 0) .
\end{aligned}$$

Además:

$$\begin{aligned}
\Delta(\pi_i^n, \psi_i^n)^T &= (\pi_i^n, \psi_i^n)^T - (\pi_{i-1}^n, \psi_{i-1}^n)^T = (\pi_i^n - \pi_{i-1}^n, \psi_i^n - \psi_{i-1}^n)^T \\
&= \sum_k \alpha_{ki}^n \mathbf{r}_k \\
&= \alpha_{\pi i}^n \mathbf{r}_\pi + \alpha_{\psi i}^n \mathbf{r}_\psi \\
&= \alpha_{\pi i}^n (1, -1)^T + \alpha_{\psi i}^n (1, 1)^T ,
\end{aligned}$$

como los eigenvectores son ortogonales (pues A es simétrica), se puede usar (5.17), y se obtiene:

$$\begin{aligned}
\alpha_{\pi i}^n &= \frac{1}{2} [\Delta \pi_i^n - \Delta \psi_i^n] = \frac{1}{2} [(\pi_i^n - \pi_{i-1}^n) - (\psi_i^n - \psi_{i-1}^n)] \\
\alpha_{\psi i}^n &= \frac{1}{2} [\Delta \pi_i^n + \Delta \psi_i^n] = \frac{1}{2} [(\pi_i^n - \pi_{i-1}^n) + (\psi_i^n - \psi_{i-1}^n)] .
\end{aligned}$$

El último ingrediente para la iteración es:

$$\begin{aligned}
\tilde{\mathbf{F}}_i^n &= \frac{1}{2} \sum_k |\lambda_k| \left(1 - |\lambda_k| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{ki}^n) \alpha_{ki}^n \mathbf{r}_k \\
&= \frac{1}{2} \left[|\lambda_\pi| \left(1 - |\lambda_\pi| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{\pi i}^n) \alpha_{\pi i}^n (1, -1)^T + |\lambda_\psi| \left(1 - |\lambda_\psi| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{\psi i}^n) \alpha_{\psi i}^n (1, 1)^T \right] ,
\end{aligned}$$

de ahí:

$$\begin{aligned}\tilde{\mathbf{F}}_i^{n(1)} &= \frac{1}{2} \left[|\lambda_\pi| \left(1 - |\lambda_\pi| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{\pi i}^n) \alpha_{\pi i}^n + |\lambda_\psi| \left(1 - |\lambda_\psi| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{\psi i}^n) \alpha_{\psi i}^n \right] \\ \tilde{\mathbf{F}}_i^{n(2)} &= \frac{1}{2} \left[-|\lambda_\pi| \left(1 - |\lambda_\pi| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{\pi i}^n) \alpha_{\pi i}^n + |\lambda_\psi| \left(1 - |\lambda_\psi| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{\psi i}^n) \alpha_{\psi i}^n \right].\end{aligned}$$

Usando el limitador *minmod*, es decir: $\phi(\theta_{ki}^n) = \minmod(1, \theta_{ki}^n)$, donde:

$$\begin{aligned}\theta_{\pi i}^n &= \frac{\alpha_{\pi I}^n}{\alpha_{\pi i}^n} \quad I = \begin{cases} i-1 & \text{si } \lambda_\pi \geq 0, \\ i+1 & \text{si } \lambda_\pi < 0. \end{cases} \\ \theta_{\psi i}^n &= \frac{\alpha_{\psi J}^n}{\alpha_{\psi i}^n} \quad J = \begin{cases} i-1 & \text{si } \lambda_\psi \geq 0, \\ i+1 & \text{si } \lambda_\psi < 0. \end{cases}\end{aligned}$$

Se hicieron los cálculos para $a = 0.7$, $b = 0.2$, $\Delta x = 0.01$, $\Delta t = 0.005$, tomando como dato inicial un perfil gaussiano centrado en $x = 0.0$, de amplitud 1.0 y 0.1 de desviación. La Figura 5.2 muestra los resultados: notar que como era de esperarse, los resultados son los mismos que cuando se desacopla el sistema.

5.4 Sistemas lineales de coeficientes variables

La diferencia entre un sistema de coeficientes variables y uno de coeficientes constantes es que ahora se tiene $A = A(x)$ en (5.1), y por lo tanto, el k -ésimo eigenvalor es $\lambda_k = \lambda_k(x)$, así como $r_k = r_k(x)$. El sistema sigue llamándose hiperbólico siempre que, para x dada, $\lambda_k(x) \neq \lambda_j(x)$ cuando $k \neq j$.

Antes de estudiar el sistema y las modificaciones pertinentes al método, es necesario entender antes cómo se resuelve una sola ecuación de advección cuando la velocidad u depende de la posición, es decir, $u = u(x)$. Considérese la ecuación escrita en forma conservativa:

$$\partial_t q + \partial_x(u(x)q) = 0,$$

con una velocidad u_i para cada celda C_i . Se supone que $u_i \neq u_j$ para $i \neq j$. Considérese el problema de Riemman en la interface de la i -ésima celda:

$$\begin{aligned}\partial_t q + u_i \partial_x q &= 0 \quad \text{para } x \in C_i, \quad u_i > 0 \\ \partial_t q + u_{i-1} \partial_x q &= 0 \quad \text{para } x \in C_{i-1}.\end{aligned}$$

Haciendo un balance de flujos promediados en la celda C_i , se tiene que $u_i \Delta Q_i^n = \Delta u_i Q_i^n + \bar{F}_i^n$, donde Δu_i es la diferencia de velocidades entre las celdas y $\bar{F}_i^n$ es el flujo promedio que

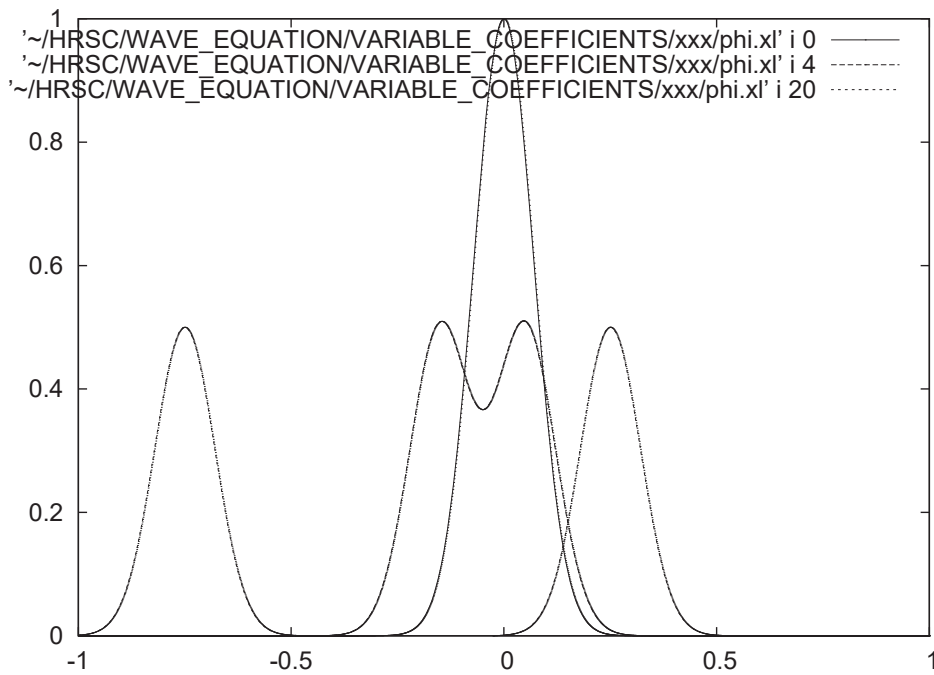


Figura 5.2: Solución de la Ecuación de Onda aplicando directamente la generalización del método de Alta Resolución con el limitador de flujos MC. $\Delta x = 0.01$, $\Delta t = 0.005$. La línea continua son los datos iniciales; la de puntos largos es la solución a $t = 0.25$ y la de puntos cortos a $t = 0.6$.

entra a la celda C_i proveniente de la celda vecina C_{i-1} , es decir, $\bar{F}_i^n = u_{i-1}\Delta Q_i^n$. Entonces:

$$\begin{aligned} u_i\Delta Q_i^n &= \Delta u_i Q_i^n + u_{i-1}\Delta Q_i^n \\ &= (u_i - u_{i-1})Q_i^n + u_{i-1}(Q_i^n - Q_{i-1}^n) \\ &= u_i Q_i^n - u_{i-1} Q_{i-1}^n, \end{aligned}$$

de donde se obtiene que $\Delta Q_i^n = Q_i^n - \left(\frac{u_{i-1}}{u_i}\right) Q_{i-1}^n$, que constituye la diferencia fundamental del caso de velocidades variables. Es bonito comprobar que cuando las velocidades son iguales (coeficientes constantes) se recupera: $\Delta Q_i^n = Q_i^n - Q_{i-1}^n$. El tedió viene al recordar que sólo se ha considerado $u_i > 0$. Si ahora $u_i \leq 0$, el balance de flujos promediados en la celda C_i es análogo: $u_i\Delta Q_i^{n*} = \Delta u_i Q_i^n + \bar{F}_i^n$, sin embargo el hecho de que el tránsito de información sea en el sentido contrario, implica calcular las diferencias en ese mismo sentido, es decir: $\Delta Q_i^{n*} = Q_i^n - Q_{i+1}^n$, $\Delta u_i = u_i - u_{i+1}$ y $\bar{F}_i^n = u_{i+1}\Delta Q_i^{n*}$, porque el flujo que entra ahora proviene de la celda C_{i+1} . Entonces:

$$u_i(Q_i^n - Q_{i+1}^n) = (u_i - u_{i+1})Q_i^n + u_{i+1}(Q_i^n - Q_{i+1}^n).$$

Para usar los mismos índices que en el caso anterior, se trasladan las cantidades a la celda vecina C_{i-1} , es decir:

$$\begin{aligned} u_{i-1}(Q_{i-1}^n - Q_i^n) &= (u_{i-1} - u_i)Q_{i-1}^n + u_i(Q_{i-1}^n - Q_i^n) \\ -u_{i-1}\Delta Q_i^n &= u_{i-1}Q_{i-1}^n - u_iQ_{i-1}^n + u_iQ_{i-1}^n - u_iQ_i^n \\ u_{i-1}\Delta Q_i^n &= u_iQ_i^n - u_{i-1}Q_{i-1}^n, \end{aligned}$$

de donde $\Delta Q_i^n = \left(\frac{u_i}{u_{i-1}}\right) Q_i^n - Q_{i-1}^n$. Notar que cuando las velocidades son iguales, se recupera $\Delta Q_i^n = Q_i^n - Q_{i-1}^n$.

El plan es usar (5.12), tomando los saltos ΔQ_i^n recién calculados, según el signo de la velocidad en cada celda.

Al tratarse de sistemas de ecuaciones de coeficientes variables se pretende usar una versión más general todavía de (5.16), donde habrá que introducir otra notación para el salto en Q_i^n y considerar tres casos en cada celda: $\lambda_{ki} > 0$, $\lambda_{ki} \leq 0$ y $\lambda_{ki} = 0$, donde λ_{ki} es el valor que toma k -ésimo eigenvalor de A en la i -ésima celda de la malla.

Si $\lambda_{ki} > 0$, defínase la velocidad de advección como $s_{ki} \equiv \lambda_{ki}$. La k -ésima componente del salto en este caso se define como $\tilde{\Delta} \mathbf{Q}_i^{n(k)} = \mathbf{Q}_i^{n(k)} - \left(\frac{s_{k(i-1)}}{s_{ki}}\right) \mathbf{Q}_{i-1}^{n(k)}$, pero

$$\mathbf{Q}_i^{n(k)} - \left(\frac{s_{k(i-1)}}{s_{ki}} \right) \mathbf{Q}_{i-1}^{n(k)} = (\mathbf{Q}_i^{n(k)} - \mathbf{Q}_{i-1}^{n(k)}) + \mathbf{Q}_{i-1}^{n(k)} - \left(\frac{s_{k(i-1)}}{s_{ki}} \right) \mathbf{Q}_{i-1}^{n(k)} = \Delta \mathbf{Q}_i^{n(k)} + \left(1 - \frac{s_{k(i-1)}}{s_{ki}} \right) \mathbf{Q}_{i-1}^{n(k)},$$

así que $\tilde{\Delta} \mathbf{Q}_i^{n(k)} = \Delta \mathbf{Q}_i^{n(k)} + \left(1 - \frac{s_{k(i-1)}}{s_{ki}} \right) \mathbf{Q}_{i-1}^{n(k)}$, e igual que antes, se descompone el salto en una suma vectorial de las direcciones características de propagación: $\tilde{\Delta} \mathbf{Q}_i^n = \sum_{k=1}^m \alpha_{ki}^n \mathbf{r}_{ki}$, donde $\mathbf{r}_{ki}$ es el valor que toma el k -ésimo eigenvector de A en la i -ésima celda. Entonces:

$$\alpha_{ki}^n = \frac{1}{|\mathbf{r}_{ki}|^2} (\tilde{\Delta} \mathbf{Q}_i^n \cdot \mathbf{r}_{ki}) = \frac{1}{|\mathbf{r}_{ki}|^2} \left[\Delta \mathbf{Q}_i^n \cdot \mathbf{r}_{ki} + \left(1 - \frac{s_{k(i-1)}}{s_{ki}} \right) \mathbf{Q}_{i-1}^n \cdot \mathbf{r}_{ki} \right], \quad \lambda_{ki} > 0. \quad (5.25)$$

En caso de tener un sistema cuya matriz no sea simétrica, puede suceder que los eigenvectores no sean ortogonales; en tal caso no funciona la sencillísima fórmula anterior y no queda más que trabajar un poco más haciendo cálculos explícitos para encontrar α_{ki}^n de $\tilde{\Delta} \mathbf{Q}_i^n = \sum_{k=1}^m \alpha_{ki}^n \mathbf{r}_{ki}$.

Sea el caso en que $\lambda_{ki} < 0$, entonces $s_{ki} \equiv \lambda_{k(i-1)}$. Análogamente al salto del caso anterior, ahora la k -ésima componente del salto en $\mathbf{Q}_i^n$ es: $\tilde{\Delta} \mathbf{Q}_i^{n(k)} = \Delta \mathbf{Q}_i^{n(k)} + \left(\frac{s_{ki}}{s_{k(i-1)}} - 1 \right) \mathbf{Q}_i^{n(k)}$, y haciendo $\tilde{\Delta} \mathbf{Q}_i^n = \sum_{k=1}^m \alpha_{ki}^n \mathbf{r}_{ki}$, se concluye:

$$\alpha_{ki}^n = \frac{1}{|\mathbf{r}_{ki}|^2} \left[\Delta \mathbf{Q}_i^n \cdot \mathbf{r}_{ki} + \left(\frac{s_{ki}}{s_{k(i-1)}} - 1 \right) \mathbf{Q}_i^n \cdot \mathbf{r}_{ki} \right], \quad \lambda_{ki} < 0. \quad (5.26)$$

Igual que cuando $\lambda_{ki} > 0$, esta elegante fórmula sólo es válida cuando los eigenvectores $\mathbf{r}_{ki}$ son ortogonales. Si no hay tal suerte, se tendrá que calcular α_{ki}^n de $\tilde{\Delta} \mathbf{Q}_i^n = \sum_{k=1}^m \alpha_{ki}^n \mathbf{r}_{ki}$, por algún método algebraico menos directo.

Si en cambio $\lambda_{ki} = 0$, es razonable tomar $s_{ki} = 0$, entonces todas las α 's y los saltos se reducen teóricamente a los valores que tendrían si los eigenvalores fueran constantes sobre toda la malla, el inconveniente es más bien práctico, pues el lenguaje de programación elegido pudiera interpretar las razones del tipo $\frac{s_{ki}}{s_{kj}}$ como valores indeterminados. Por ello se debe poner cuidado al especificar que si en alguna celda se tiene $\lambda_{ki} = 0$, los saltos y los coeficientes de la combinación lineal de eigenvectores se calculan como si en esa celda se tuviera un problema de coeficientes constantes.

Solamente falta un ingrediente para aplicar la fórmula (5.16), se trata del término relacionado con el limitador de flujos. En particular hay un problema con la manera de calcular θ_{ki}^n , la cual es una cantidad que mide la *suavidad* de la de la k -ésima componente característica de la solución; así cuando los eigenvectores eran constantes, lo único que quedaba por comparar eran los coeficientes α_{ki}^n (con $I = i \pm 1$, dependiendo del signo de la velocidad de advección) y α_{ki}^n , mediante la razón: $\theta_{ki}^n = \frac{\alpha_{ki}^n}{\alpha_{ki}^n}$. La magnitud de las α 's de-

pende de la magnitud del salto $\Delta \mathbf{Q}_i^n$, entonces cuando la solución es suave se tiene $\theta_{ki} \sim 1$, pero cerca de una discontinuidad, el valor de θ crece considerablemente y debe ser limitado mediante la función ϕ . Ahora que la magnitud y dirección de los eigenvectores cambia con la posición, la comparación de las componentes características involucra no solo los coeficientes α , sino también los $\mathbf{r}$. Para tener un parámetro que mida el grado de suavidad de la solución es razonable tomar la proyección de $\alpha_{kl}^n \mathbf{r}_{kl}$ sobre $\alpha_{ki}^n \mathbf{r}_{ki}$ y calcular el cociente con la proyección de $\alpha_{ki}^n \mathbf{r}_{ki}$ consigo mismo, es decir:

$$\theta_{ki}^n = \frac{\alpha_{kl}^n \mathbf{r}_{kl}^n \cdot \alpha_{ki}^n \mathbf{r}_{ki}^n}{\alpha_{ki}^n \mathbf{r}_{ki}^n \cdot \alpha_{ki}^n \mathbf{r}_{ki}^n} \quad \text{donde} \quad I = \begin{cases} i-1 & \text{si } s_{ki} \geq 0, \\ i+1 & \text{si } s_{ki} < 0, \end{cases}$$

En una regiones suaves, por ejemplo, los eigenvectores de celdas vecinas tienen direcciones y coeficientes equiparables, lo que devuelve $\theta \sim 1$. En la práctica se tiene:

$$\theta_{ki}^n = \frac{\alpha_{kl}^n (\mathbf{r}_{kl}^n \cdot \mathbf{r}_{ki}^n)}{\alpha_{ki}^n |\mathbf{r}_{ki}^n|^2} \quad \text{donde} \quad I = \begin{cases} i-1 & \text{si } s_{ki} \geq 0, \\ i+1 & \text{si } s_{ki} < 0. \end{cases}$$

Para implementar el método basta referirse a la sección (5.2.3) y aplicar prácticamente las mismas fórmulas dadas ahí con el propósito de usar (5.16), las diferencias están solo en la manera de calcular las α 's, y el parámetro θ .

5.4.1 Pruebas con la Ecuación de Onda

De (5.20), se tiene que la ecuación de onda en $1 + 1D$ puede escribirse como:

$$\partial_t \mathbf{u} + \partial_x (A \mathbf{u}) = \mathbf{0},$$

donde $\mathbf{u} = (\pi, \psi)^T$ y

$$A = - \begin{pmatrix} b & a \\ a & b \end{pmatrix}.$$

Como se mencionó en esa misma sección, las cantidades a y b son parámetros de la métrica: la función de lapso a es el coeficiente de dt^2 , y por lo tanto mide la separación entre *rebanadas* espaciales en el tiempo. Si por ejemplo el dominio $[-1, 1]$ se divide hipotéticamente en dos, y en cada mitad se pone un valor diferente para a , el resultado debiera ser que los pulsos que se propagan a cada lado, lo hagan a velocidades distintas. Ese efecto se logra tomando a como una función del tipo $\tanh(x)$, lo que garantiza una transición suave en el origen entre dos valores muy diferentes para a , a la derecha e izquierda de 0. Por otro lado, el vector de corrimiento b (con una sola componente en este caso, por tratarse de una dimensión espacial) está relacionado con la velocidad de las coordenadas en relación al observador. Un audaz experimento sería tomar $b = b(x) = x$, lo que implica

que en las fronteras del dominio (± 1), las coordenadas se alejan con la misma velocidad que los pulsos, es decir, éstos nunca alcanzan las fronteras.

Pueden hacerse elecciones arbitrarias de los parámetros a y b , según sea conveniente para la simulación de algún fenómeno. Lo relevante para este trabajo es que se ha considerado un sistema cuya matriz admite coeficientes dependientes de la posición. Incluso los eigenvalores $(-b \pm a)$ son ahora funciones de x . Esto permite aplicar la teoría desarrollada en el apartado anterior.

Como se ha dicho antes, las diferencias en la implementación del método no son muchas, de hecho, las cantidades calculadas en la subsección (5.3.2) son las mismas que se requieren en este caso, execepto las α 's y θ 's, ahora hay que hacer la siguiente distinción en cada una de las $N_x - 1$ celdas:

si $\lambda_{ki} > 0$, entonces la velocidad de advección en la i -ésima celda será $s_{ki} = \lambda_{ki}$ y la k -ésima componente del salto en $\mathbf{Q}_i^n$: $\tilde{\Delta}\mathbf{Q}_i^{n(k)} = \Delta\mathbf{Q}_i^{n(k)} + \left(1 - \frac{s_{k(i-1)}}{s_{ki}}\right)\mathbf{Q}_{ik}^{n(k)}$. Recordando que los eigenvalores de la matriz son $\lambda_{\pi i} = (-b_i + a_i)$ y $\lambda_{\psi i} = (-b_i - a_i)$ y los eigenvectores: $r_{\pi} = (1, -1)^T$ y $r_{\psi} = (1, 1)^T$ (notar que para la ecuación de onda los eigenvectores siguen siendo constantes), y aplicando (5.25):

$$\begin{aligned}\alpha_{\pi i}^n &= \frac{1}{2} \left[(\Delta\pi_i^n - \Delta\psi_i^n) + \left(1 - \frac{s_{\pi(i-1)}}{s_{\pi i}}\right) (\pi_{i-1}^n - \psi_{i-1}^n) \right], \quad \lambda_{\pi i} > 0 \\ \alpha_{\psi i}^n &= \frac{1}{2} \left[(\Delta\pi_i^n + \Delta\psi_i^n) + \left(1 - \frac{s_{\psi(i-1)}}{s_{\psi i}}\right) (\pi_{i-1}^n - \psi_{i-1}^n) \right], \quad \lambda_{\psi i} > 0.\end{aligned}$$

En cambio si $\lambda_{ki} < 0$, la velocidad de advección en la i -ésima celda será $s_{ki} = \lambda_{k(i-1)}$ y la k -ésima componente del salto en $\mathbf{Q}_i^n$: $\tilde{\Delta}\mathbf{Q}_i^{n(k)} = \Delta\mathbf{Q}_i^{n(k)} + \left(\frac{s_{ki}}{s_{k(i-1)}} - 1\right)\mathbf{Q}_{ik}^{n(k)}$, y aplica (5.26):

$$\begin{aligned}\alpha_{\pi i}^n &= \frac{1}{2} \left[(\Delta\pi_i^n - \Delta\psi_i^n) + \left(\frac{s_{\pi i}}{s_{\pi(i-1)}} - 1\right) (\pi_i^n - \psi_i^n) \right], \quad \lambda_{\pi i} < 0 \\ \alpha_{\psi i}^n &= \frac{1}{2} \left[(\Delta\pi_i^n + \Delta\psi_i^n) + \left(\frac{s_{\psi i}}{s_{\psi(i-1)}} - 1\right) (\pi_i^n - \psi_i^n) \right], \quad \lambda_{\psi i} < 0.\end{aligned}$$

Si sucediera que $\lambda_{ki} = 0$ para alguna i , la velocidad será $s_{ki} = 0$ y entonces:

$$\begin{aligned}\alpha_{\pi i}^n &= \frac{1}{2} [(\Delta\pi_i^n - \Delta\psi_i^n) + (\pi_i^n - \psi_i^n)], \quad \lambda_{\pi i} = 0 \\ \alpha_{\psi i}^n &= \frac{1}{2} [(\Delta\pi_i^n + \Delta\psi_i^n) + (\pi_i^n - \psi_i^n)], \quad \lambda_{\psi i} = 0.\end{aligned}$$

Sólo faltan:

$$\theta_{\pi i}^n = \frac{\alpha_{\pi I}^n (\mathbf{r}_{\pi}^n \cdot \mathbf{r}_{\pi}^n)}{\alpha_{\pi i}^n |\mathbf{r}_{\pi}^n|^2} = \frac{\alpha_{\pi I}^n}{\alpha_{\pi i}^n}$$

$$\theta_{\psi i}^n = \frac{\alpha_{\psi I}^n (\mathbf{r}_{\psi}^n \cdot \mathbf{r}_{\psi}^n)}{\alpha_{\psi i}^n |\mathbf{r}_{\psi}^n|^2} = \frac{\alpha_{\psi I}^n}{\alpha_{\psi i}^n}, \quad \text{donde } I = \begin{cases} i - 1 & \text{si } s_{ki} \geq 0 \\ i + 1 & \text{si } s_{ki} < 0. \end{cases}$$

Capítulo 6

Sistemas hiperbólicos no-lineales de leyes de conservación

En el capítulo anterior se asumió que las entradas de la matriz del sistema (5.1) dependen de la posición. Resolver ese problema requiere de una reformulación cuidadosa de algunos puntos del método de Alta Resolución que se ha usado. El problema se complica aún más cuando las entradas de la matriz son funciones de las cantidades conservadas (cantidades contenidas en las componentes del vector conservado $\mathbf{q}$), y precisamente en ese hecho radica la no-linealidad del sistema. En general, en cada tiempo y en cada celda de la malla, se tendrán eigenvalores y eigenvectores dependientes de las variables conservadas. Sin embargo la estrategia para resolver este problema no difiere mucho de lo que se ha hecho hasta ahora: reformular los puntos necesarios para aplicar (5.16), el resultado más importante y general que hemos obtenido.

Las similitudes son muchas con el caso de los sistemas de coeficientes variables, lo que simplifica considerablemente las modificaciones al método. Basta ver que para el caso de sistemas no-lineales en el fondo se tiene $A = A(x, t)$. Los eigenvalores serán ahora funciones del espacio y del tiempo, así el k -ésimo, en la i -ésima celda, al tiempo $t = t^n$ se denotará por: λ_{ki}^n , análogamente el k -ésimo eigenvector en la i -ésima celda, al tiempo $t = t^n$: $\mathbf{r}_{ki}^n$.

La buena noticia es que en cada paso temporal, mientras $t = t^n$ es fijo, las cantidades únicamente dependen de la posición, y entonces es razonable aplicar los mismos razonamientos expuestos en la sección (5.4). ¡El trabajo pesado está hecho! Sin embargo hay que poner mucho cuidado en el cálculo de las velocidades de advección s_{ki}^n . Para ver esto, primero es necesario entender los tipos de discontinuidades que pueden desarrollar las soluciones de los sistemas no-lineales.

6.1 Discontinuidades en la solución

Como se ha dicho antes, las soluciones de los sistemas no-lineales presentan varios fenómenos interesantes: los más relevantes, por las complicaciones numéricas que implican, son las discontinuidades. Toda discontinuidad suele llamarse *frente de choque*, el cual es un término heredado de la hidrodinámica, pues en su argot se llama así al frente de onda que delimita dos regiones espacio-temporales con valores muy distintos de las cantidades físicas del fluido; se trata pues de un frente de onda con altos gradientes para tales cantidades.

Dentro de los frentes de choque se distingue un tipo particular de discontinuidad: las llamadas *discontinuidades de contacto*. Se trata de frentes de choque cuya velocidad de propagación es la misma que las velocidades características de los frentes de onda a cada lado de la discontinuidad. En el contexto de la hidrodinámica se entiende mejor: una discontinuidad es de contacto si la discontinuidad de cada una de las cantidades que describen al fluido (presión, densidad, etcétera) se propaga a la misma velocidad que las demás, y esa velocidad coincide a su vez con la de propagación de la onda; así pueden distinguirse claramente dos regiones muy diferentes. Esto pudiera no ser así, pues el fluido pudiera *enrarecerse* en las zonas cercanas a las discontinuidades, si tal sucediera, se dice que se tiene ya sea una *onda de rarefacción* (cuando la densidad decrece cerca de la discontinuidad) o una *onda de compresión* (cuando la densidad crece en esa región), que son fenómenos frecuentemente presentados por las ecuaciones no-lineales. A diferencia de los flujos enrarecidos, las discontinuidades de contacto marcan claramente, en todo tiempo, el lugar donde ocurren los cambios bruscos en los parámetros del fluido.

6.2 Cálculo de las velocidades de propagación

En el caso más sencillo, cuando se tenían sistemas lineales y de coeficientes constantes, la velocidad de propagación s_k de la k -ésima *eigen-componente* de la información era exactamente $s_k = \lambda_k$ en todas las celdas del dominio. En el caso de sistemas de coeficientes variables, como la velocidad depende de la posición, se obtuvo:

$$s_{ki} = \begin{cases} \lambda_{ki} & \text{si } \lambda_{ki} > 0 \\ \lambda_{k(i-1)} & \text{si } \lambda_{ki} < 0 . \end{cases}$$

Si el sistema es no-lineal, se asume además una dependencia en el tiempo: s_{ki}^n . Ahora, la enseñanza de la sección anterior es que la naturaleza de las discontinuidades y rarefacciones cerca de las discontinuidades no es necesariamente la misma, de hecho se definió una discontinuidad de contacto como aquella para la que la velocidad de propagación de la información coincide con la velocidad de propagación a cada lado de la discontinuidad.

En cambio, en las regiones donde se presentan ondas de rarefacción o de compresión, la velocidad varía de manera continua y esa variación pudiera ser complicada. Una idea razonable es tomar un promedio de los eigenvalores en celdas vecinas: $\frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n)$ y considerar su signo, luego hacer:

$$s_{ki} = \begin{cases} \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) & \text{si } \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) > 0 \\ \frac{1}{2}(\lambda_{k(i-1)}^n + \lambda_{k(i-2)}^n) & \text{si } \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) < 0, \end{cases}$$

Con esta última corrección se termina con la construcción de un método capaz de resolver en general sistemas de ecuaciones no-lineales. El siguiente capítulo muestra los resultados para un sistema particular, que modela el flujo de un gas con ciertas particularidades y ejemplifica la utilidad del método aquí discutido.

Capítulo 7

Las Ecuaciones de Euler

Las ecuaciones de Euler describen en forma general la dinámica de un gas compresible y libre de viscosidad. Se trata de un sistema no-lineal de tres ecuaciones en derivadas parciales para tres cantidades conservadas a lo largo de la evolución: la masa, el momento lineal y la energía total.

7.1 Deducción de las leyes de conservación

Considérese un fluido en general en tres dimensiones que se mueve con una velocidad $\mathbf{u} = \mathbf{u}(\mathbf{r}, t)$ ($\mathbf{r}$ es un vector de posición tres-dimensional), descrito desde el enfoque lagrangiano, es decir, los cambios de las propiedades del fluido se describen desde un referencial que se mueve con éste. Sea V_t un volumen de control fijo que se mueve con el fluido, y sean $f : V_t \times \mathbb{R} \rightarrow \mathbb{R}$ un campo escalar y $\mathbf{F} : V_t \times \mathbb{R} \rightarrow \mathbb{R}^3$ un campo vectorial. Un resultado matemático útil para la deducción de las ecuaciones de Euler es el teorema del transporte [12], en su versión para campos escalares y vectoriales por separado:

$$\begin{aligned}\frac{d}{dt} \int_{V_t} f dV &= \int_{V_t} [\partial_t f + \nabla \cdot (f\mathbf{u})] dV \\ \frac{d}{dt} \int_{V_t} \mathbf{F} dV &= \int_{V_t} [\partial_t \mathbf{F} + (\mathbf{u} \cdot \nabla) \mathbf{F} + \mathbf{F} \cdot (\nabla \cdot \mathbf{u})] dV.\end{aligned}$$

Conservación de la masa

Si el campo escalar $\rho(\mathbf{r}, t)$ corresponde a la densidad volumétrica de masa del fluido, entonces dentro de V_t se tiene:

$$m(V_t) = \int_{V_t} \rho(\mathbf{r}, t) dV,$$

en ausencia de fuentes o sumideros de masa, el principio de conservación de esta misma establece que:

$$\frac{d}{dt}m(V_t) = \frac{d}{dt} \int_{V_t} \rho(\mathbf{r}, t) dV = 0 .$$

Aplicando el teorema del transporte:

$$0 = \frac{d}{dt} \int_{V_t} \left[\frac{D\rho}{Dt} + \rho(\nabla \cdot \mathbf{u}) \right] dV = 0 ,$$

donde se define el operador diferencial $\frac{D}{Dt} := \partial_t + \mathbf{u} \cdot \nabla$. La sustitución de $\frac{D}{Dt}$ y el hecho de que V_t es arbitrario llevan a la conclusión de que:

$$\partial_t \rho + \nabla \cdot (\rho \mathbf{u}) = 0 ,$$

que en una dimensión espacial ($\mathbf{u}(\mathbf{r}, t) = u(x, t)$) se escribe:

$$\partial_t \rho + \partial_x(\rho u) = 0 . \quad (7.1)$$

Conservación del momento

El momento lineal total dentro de V_t está dado por la integral:

$$\int_{V_t} \rho u(\mathbf{r}, t) dV .$$

La versión de la Segunda Ley de Newton en hidrodinámica establece que la variación temporal del momento en el volumen V_t es igual a la fuerza neta ejercida sobre dicho volumen. Esta fuerza neta incluye el campo tensorial de esfuerzos y las fuerzas externas. Al tratarse de un gas libre de viscosidad, es decir, que todos los esfuerzos de corte son nulos, solo queda por tomar en cuenta la presión P , que es una fuerza por unidad de área, así que la fuerza neta sobre V_t debida a P es:

$$- \int_{\partial V_t} P d\mathbf{A} ,$$

donde ∂V_t es la frontera de V_t .

Si denotamos con $\mathbf{G}$ al campo vectorial correspondiente a las fuerzas externas por unidad de volumen, su contribución sobre V_t es:

$$\int_{V_t} \mathbf{G} dV ,$$

luego, de la segunda ley de Newton:

$$\frac{d}{dt} \int_{V_t} \rho \mathbf{u}(\mathbf{r}, t) dV = - \int_{\partial V_t} P d\mathbf{A} + \int_{V_t} \mathbf{G} dV,$$

No es difícil ver que:

$$- \int_{\partial V_t} P d\mathbf{A} + \int_{V_t} \mathbf{G} dV = \int_{V_t} (\mathbf{G} - \nabla P) dV,$$

Aplicando el teorema del transporte:

$$\int_{V_t} [\partial_t(\rho \mathbf{u}) + (\mathbf{u} \cdot \nabla)(\rho \mathbf{u}) + \rho \mathbf{u}(\nabla \cdot \mathbf{u})] dV = \int_{V_t} (\mathbf{G} - \nabla P) dV,$$

lo cual vale para cualquier V_t , por lo tanto:

$$\partial_t(\rho \mathbf{u}) + (\mathbf{u} \cdot \nabla)(\rho \mathbf{u}) + \rho \mathbf{u}(\nabla \cdot \mathbf{u}) = \mathbf{G} - \nabla P,$$

En ausencia de fuerzas externas, haciendo un poco de álgebra y reduciendo todo a una dimensión espacial se llega a:

$$\partial_t(\rho u) + \partial_x(\rho u^2 + P) = 0. \quad (7.2)$$

Conservación de la energía

Sea E la energía total por unidad de volumen del fluido en V_t . La ley de la conservación de la energía establece que la variación temporal de la energía total debe ser igual al trabajo realizado sobre el sistema por la fuerza neta, por unidad de tiempo. En ausencia de fuerzas externas y de esfuerzos de corte, matemáticamente se tiene:

$$\frac{d}{dt} \int_{V_t} \rho E dV = - \int_{\partial V_t} P \mathbf{u} \cdot d\mathbf{A},$$

aplicando el teorema de la divergencia al lado derecho y el teorema del transporte al izquierdo:

$$\int_{V_t} (\partial_t E + \nabla \cdot (E \mathbf{u})) = - \int_{V_t} \nabla \cdot (P \mathbf{u}) dV,$$

como V_t es arbitrario:

$$\partial_t E + \nabla \cdot (E \mathbf{u}) + \nabla(P \mathbf{u}) = \partial_t E + \nabla \cdot [(E + P) \mathbf{u}] = 0,$$

y en una dimensión:

$$\partial_t E + \partial_x[(E + P)u] = 0. \quad (7.3)$$

Las ecuaciones (7.1), (7.2) y (7.3) forman el sistema llamado Ecuaciones de Euler:

$$\partial_t \begin{pmatrix} \rho \\ \rho u \\ E \end{pmatrix} + \partial_x \begin{pmatrix} \rho u \\ \rho u^2 + P \\ (E + P)u \end{pmatrix} = \mathbf{0}. \quad (7.4)$$

Este sistema es una simplificación de uno más general: el de las ecuaciones de Navier-Stokes, que incluyen todos los términos del tensor de esfuerzos (que se traduce en la consideración de los términos de viscosidad) y fuerzas externas. Los términos eliminados son proporcionales a derivadas de segundo orden, que de ser incluidas, el sistema sería del tipo parabólico y por lo tanto las soluciones serían suaves para todo tiempo.

En el sistema (7.4), las cantidades conservadas son ρ , ρu y E , sin embargo también están involucradas las cantidades u y P , de modo que la consistencia del sistema sólo se logra al encontrar una relación entre la energía por unidad de masa, la presión y la densidad (y posiblemente otras cantidades). Enseguida la discusión al respecto.

La energía total y la ecuación de estado

En general, la energía total por unidad de volumen es una función $E = E(\rho, P)$. Ésta puede descomponerse en una suma de la energía cinética del gas por unidad de volumen $E_K = \frac{1}{2}\rho u^2$ y la energía interna por unidad de volumen $E_{int} = \rho e$, donde e es la energía interna por unidad de masa, llamada comunmente *energía interna específica*. En general, ésta incluye energía traslacional, rotacional, vibracional y cualquier otro tipo de energía, sin embargo para un gas que localmente se encuentra en equilibrio químico y termodinámico, se asume que e es una función conocida $e = e(\rho, P)$. Esa función se llama *ecuación de estado* del gas (depende de la naturaleza de éste).

Considérese un gas ideal con presión P , densidad ρ y temperatura T . Su ecuación de estado es bien conocida [13] y una de sus versiones es:

$$P = R\rho T \quad (7.5)$$

(incluso funciona muy bien para modelar gases altamente diluidos), donde R es la constante universal de los gases por unidad de masa molecular. El objetivo es obtener una relación $e(\rho, P)$.

Si además se asume que el gas es tal que su energía interna específica es directamente proporcional a la temperatura cuando su volumen es constante, es decir: $e = c_v T$, donde c_v es una constante llamada *capacidad calorífica a volumen constante*, el gas se llama *ideal politrópico*. Entonces tomando T de (7.5) se tiene un resultado parcial:

$$e = \left(\frac{c_v}{R}\right) \frac{P}{\rho}.$$

Por otro lado, para un gas politrópico se define la entalpía como $h = e + P/\rho$, y se sabe [13] que a presión constante, h es directamente proporcional a la temperatura: $h = c_p T$, donde c_p

es una constante llamada *capacidad calorífica a presión constante*, entonces $e + P/\rho = c_p T$, pero $e = c_v T$, así que $c_v T + P/\rho = c_p T$, luego, $P = (c_p - c_v)\rho T$, y comparando con (7.5) se deduce que para un gas politrópico: $R = (c_p - c_v)$, lo que se sustituye en la expresión para e y se obtiene:

$$e = \left(\frac{c_v}{c_p - c_v} \right) \frac{P}{\rho} = \left(\frac{1}{\frac{c_p}{c_v} - 1} \right) \frac{P}{\rho}.$$

Sea $\gamma := c_p/c_v$ el llamado *exponente adiabático* del gas. Así se obtiene la forma común de la ecuación de estado para el gas ideal politrópico:

$$e = \frac{P}{(\gamma - 1)\rho}. \quad (7.6)$$

Volviendo a la descomposición de la energía total por unidad de volumen en la suma de la energía interna y la cinética (ambas por unidad de volumen): $E = E_{int} + E_K$, se tiene:

$$E = \rho e + \frac{1}{2}\rho u^2,$$

sustituyendo (7.6):

$$E = \frac{P}{\gamma - 1} + \frac{1}{2}\rho u^2.$$

Un estudio más detallado [13] demuestra que el parámetro γ tiene una bonita y sencilla relación con los n grados de libertad de las moléculas del gas:

$$\gamma = \frac{n + 2}{n},$$

de modo que si se considera un gas monoatómico, es decir, que sus grados de libertad son solo los traslacionales, entonces $n = 3$ y $\gamma = 5/3$. Para un gas diatómico también se tienen dos grados de libertad para rotaciones, entonces $n = 5$ y $\gamma = 7/5$. Valga anotar que bajo circunstancias ordinarias, el aire está compuesto principalmente por N_2 y O_2 , así que $\gamma \approx 7/5$.

Para probar el método numérico se ha tomado un gas ideal, politrópico y monoatómico, es decir:

$$E = \frac{P}{\frac{5}{3} - 1} + \frac{1}{2}\rho u^2 = \frac{3}{2}P + \frac{1}{2}\rho u^2. \quad (7.7)$$

7.2 La forma conservativa del sistema

El sistema de las ecuaciones de Euler (7.4) puede escribirse en forma conservativa como:

$$\partial_t \begin{pmatrix} \rho \\ \rho u \\ E \end{pmatrix} + \partial_x \left[\begin{pmatrix} u & 0 & 0 \\ 0 & u & P/E \\ 0 & P/\rho & u \end{pmatrix} \begin{pmatrix} \rho \\ \rho u \\ E \end{pmatrix} \right] = \mathbf{0}, \quad (7.8)$$

que para un gas ideal, politrópico y monoatómico se cierra con (7.7).

Denótese la matriz de este sistema con A . Notar que ésta no es simétrica y que algunas de sus entradas están en términos de las cantidades conservadas: en ese hecho radica la no-linealidad del sistema.

Los eigenvalores de la matriz resultan ser:

$$\lambda_\rho = u, \quad \lambda_{\rho u} = u - \frac{P}{\sqrt{\rho E}}, \quad \lambda_E = u + \frac{P}{\sqrt{\rho E}}. \quad (7.9)$$

Como se ha dicho antes, el sistema es hiperbólico si para todo tiempo estos eigenvalores son reales. Por otro lado los eigenvectores son:

$$\mathbf{r}_\rho = (1, 0, 0)^T, \quad \mathbf{r}_{\rho u} = (0, \sqrt{\rho/E}, -1)^T, \quad \mathbf{r}_E = (0, \sqrt{\rho/E}, 1)^T. \quad (7.10)$$

7.3 El Tubo de Choques de Sod

Un primer problema para probar el método de Alta Resolución presentado en esta Tesis es el llamado *Tubo de Choques de Sod*, propuesto por primera vez en [8]. El problema consiste en tomar un gas contenido en un tubo unidimensional de extensión arbitraria y con condiciones iniciales discontinuas en la misma posición para la densidad ρ y la presión P . Suponiendo que se conservan la masa, el momento y la energía, y que el gas es ideal, politrópico y monoatómico, se calcula la solución numérica de las ecuaciones de Euler, que precisamente rigen el comportamiento del fluido. El hecho de que el sistema sea no-lineal y que se tomen datos iniciales discontinuos hace necesaria la implementación de un método capaz de lidiar consistentemente con las discontinuidades.

Por supuesto que este experimento numérico no admite periodicidad, por lo que se eligen condiciones de frontera absorbentes.

Se toma el dominio $[-1, 1]$ y los datos iniciales:

$$\begin{aligned} u &= 0 \\ \rho &= \begin{cases} 1.0 & , \quad x \leq 0 \\ 0.125 & , \quad x > 0 \end{cases} \\ P &= \begin{cases} 1.0 & , \quad x \leq 0 \\ 0.1 & , \quad x > 0. \end{cases} \end{aligned}$$

Los detalles técnicos sobre la solución del problema se presentan enseguida.

7.4 Implementación del método de Alta Resolución y resultados

Dados los datos iniciales para u , ρ y P , se calcula fácilmente el momento ρu y la energía E usando la ecuación de estado: $E = \frac{P}{\gamma-1} + \frac{1}{2}\rho u^2$. Enseguida se calculan los promedios de las cantidades conservadas (ρ , ρu , E) en cada celda, usando (2.1). Denótese a las cantidades promediadas con $\bar{\rho}$, $\overline{\rho u}$ y $\bar{E}$. También hay que decidir cuál será la velocidad de advección en la i -ésima celda:

$$s_{ki}^n = \begin{cases} \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) & \text{si } \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) > 0 \\ \frac{1}{2}(\lambda_{k(i-1)}^n + \lambda_{k(i-2)}^n) & \text{si } \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) < 0, \end{cases}$$

Ahora, aprovechando la naturaleza conservativa del método de Alta Resolución que se ha presentado y haciendo caso al énfasis que se ha hecho sobre la importancia de escribir el sistema en forma conservativa, tomamos (7.8) como el punto de partida junto con sus eigenvalores y eigenvectores.

Para la implementación es necesario primero calcular los coeficientes α , lo cual se hace en cada celda, calculando en cada paso temporal el valor numérico de los eigenvalores λ_{ki}^n de cada una de las cantidades en evolución ($k = \rho, \rho u, E$).

Dado que la matriz del sistema no es simétrica, era de esperarse que los eigenvectores no fueran ortogonales, así que desafortunadamente no se puede echar mano de las fórmulas (5.25) y (5.26), sin embargo no es difícil resolver $\tilde{\Delta}\mathbf{Q}_i^n = \sum_{k=1}^m \alpha_{ki}^n \mathbf{r}_{ki}^n$ con $\tilde{\Delta}\mathbf{Q}_i^{n(k)} = \Delta\mathbf{Q}_i^{n(k)} + \left(1 - \frac{s_{ki}^n}{s_{ki}^n}\right)\mathbf{Q}_{i-1}^{n(k)}$ si $\frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) > 0$ y $\tilde{\Delta}\mathbf{Q}_i^{n(k)} = \Delta\mathbf{Q}_i^{n(k)} + \left(\frac{s_{ki}^n}{s_{k(i-1)}^n} - 1\right)\mathbf{Q}_i^{n(k)}$ si $\frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) < 0$, para llegar a que:

$$\begin{aligned} \alpha_{\rho i}^n &= \Delta\bar{\rho}_i^n + \left(1 - \frac{s_{\rho(i-1)}^n}{s_{\rho i}^n}\right)\bar{\rho}_i^n, & \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) > 0 \\ \alpha_{\rho u i}^n &= \frac{1}{2} \left[\sqrt{\frac{E_i^n}{\rho_i^n}} \left[\Delta(\overline{\rho u})_i^n + \left(1 - \frac{s_{\rho u(i-1)}^n}{s_{\rho u i}^n}\right)\overline{\rho u}_i^n \right] - \left(1 - \frac{s_{E(i-1)}^n}{s_{E i}^n}\right)\bar{E}_i^n - \Delta\bar{E}_i^n \right], & \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) > 0 \\ \alpha_{E i}^n &= \frac{1}{2} \left[\sqrt{\frac{E_i^n}{\rho_i^n}} \left[\Delta(\overline{\rho u})_i^n + \left(1 - \frac{s_{\rho u(i-1)}^n}{s_{\rho u i}^n}\right)\overline{\rho u}_i^n \right] + \left(1 - \frac{s_{E(i-1)}^n}{s_{E i}^n}\right)\bar{E}_i^n + \Delta\bar{E}_i^n \right], & \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) < 0. \end{aligned}$$

además:

$$\begin{aligned}
\alpha_{\rho i}^n &= \Delta \bar{\rho}_i^n + \left(\frac{s_{\rho i}^n}{s_{\rho(i-1)}^n} - 1 \right) \bar{\rho}_i^n, \quad \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) < 0 \\
\alpha_{\rho u i}^n &= \frac{1}{2} \left[\sqrt{\frac{E_i^n}{\rho_i^n}} \left[\Delta(\bar{\rho}u)_i^n + \left(\frac{s_{\rho u i}^n}{s_{\rho u(i-1)}^n} - 1 \right) \bar{\rho}u_i^n \right] - \left(\frac{s_{E i}^n}{s_{E(i-1)}^n} - 1 \right) \bar{E}_i^n - \Delta \bar{E}_i^n \right], \quad \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) < 0 \\
\alpha_{E i}^n &= \frac{1}{2} \left[\sqrt{\frac{E_i^n}{\rho_i^n}} \left[\Delta(\bar{\rho}u)_i^n + \left(\frac{s_{\rho u i}^n}{s_{\rho u(i-1)}^n} - 1 \right) \bar{\rho}u_i^n \right] + \left(\frac{s_{E i}^n}{s_{E(i-1)}^n} - 1 \right) \bar{E}_i^n + \Delta \bar{E}_i^n \right], \quad \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) < 0.
\end{aligned}$$

Si en alguna celda sucediera que $\frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) = 0$, entonces:

$$\begin{aligned}
\alpha_{\rho i}^n &= \Delta \bar{\rho}_i^n, \quad \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) = 0 \\
\alpha_{\rho u i}^n &= \frac{1}{2} \left[\sqrt{\frac{E_i^n}{\rho_i^n}} \Delta(\bar{\rho}u)_i^n - \Delta \bar{E}_i^n \right], \quad \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) = 0 \\
\alpha_{E i}^n &= \frac{1}{2} \left[\sqrt{\frac{E_i^n}{\rho_i^n}} \Delta(\bar{\rho}u)_i^n + \Delta \bar{E}_i^n \right], \quad \frac{1}{2}(\lambda_{ki}^n + \lambda_{k(i-1)}^n) = 0.
\end{aligned}$$

Como se indica en la subsección (5.2.3), lo anterior se usa para calcular $A^+ \Delta \mathbf{Q}_i^n = \sum_{k=1}^m \lambda_{ki}^{n+} \alpha_{ki}^n \mathbf{r}_k^n$ y $A^- \Delta \mathbf{Q}_{i+1}^n = \sum_{k=1}^m \lambda_{k(i+1)}^{n-} \alpha_{k(i+1)}^n \mathbf{r}_{k(i+1)}^n$, al igual que:

$$\tilde{\mathbf{F}}_{ki}^n = \frac{1}{2} \sum_{k=1}^m |s_{ki}^n| \left(1 - |s_{ki}^n| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{ki}^n) \alpha_{ki}^n \mathbf{r}_{ki}^n,$$

donde ϕ es alguna función limitadora de flujo y θ_{ki}^n esta dada por:

$$\theta_{ki}^n = \frac{\alpha_{kI}^n (\mathbf{r}_{kI}^n \cdot \mathbf{r}_{ki}^n)}{\alpha_{ki}^n |\mathbf{r}_{ki}^n|^2} \quad \text{donde} \quad I = \begin{cases} i-1 & \text{si } s_{ki}^n \geq 0, \\ i+1 & \text{si } s_{ki}^n < 0. \end{cases}$$

Todo se hace para poder aplicar finalmente (5.16), que explícitamente se ve:

$$\begin{aligned}
(\Delta \bar{\rho}_i^{n+1}, \Delta \bar{\rho u}_i^{n+1}, \Delta \bar{E}_i^{n+1})^T &= (\Delta \bar{\rho}_i^n, \Delta \bar{\rho u}_i^n, \Delta \bar{E}_i^n)^T \\
&- \frac{\Delta t}{\Delta x} \left[\sum_k \lambda_{ki}^{n+} \alpha_{ki}^n \mathbf{r}_{ki}^n + \sum_k \lambda_{k(i+1)}^{n-} \alpha_{k(i+1)}^n \mathbf{r}_{k(i+1)}^n \right] \\
&- \frac{\Delta t}{\Delta x} \left[\frac{1}{2} \sum_k |s_{k(i+1)}^n| \left(1 - |s_{k(i+1)}^n| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{k(i+1)}^n) \alpha_{k(i+1)}^n \mathbf{r}_{k(i+1)}^n \right] \\
&- \frac{\Delta t}{\Delta x} \left[\frac{1}{2} \sum_k |s_{ki}^n| \left(1 - |s_{ki}^n| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{ki}^n) \alpha_{ki}^n \mathbf{r}_{ki}^n \right]. \tag{7.11}
\end{aligned}$$

y más explícitamente, cada componente de la ecuación vectorial es la fórmula de evolución para cada variable conservada:

$$\begin{aligned}
\Delta \bar{\rho}_i^{n+1} &= \Delta \bar{\rho}_i^n - \frac{\Delta t}{\Delta x} \left(\lambda_{\rho i}^{n+} \alpha_{\rho i}^n + \lambda_{\rho(i+1)}^{n-} \alpha_{\rho(i+1)}^n \right) \\
&- \frac{\Delta t}{\Delta x} \frac{1}{2} \left[|s_{\rho(i+1)}^n| \left(1 - |s_{\rho(i+1)}^n| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{\rho(i+1)}^n) \alpha_{\rho(i+1)}^n - |s_{\rho i}^n| \left(1 - |s_{\rho i}^n| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{\rho i}^n) \alpha_{\rho i}^n \right].
\end{aligned}$$

$$\begin{aligned}
\Delta \bar{\rho u}_i^{n+1} &= \\
&= \Delta \bar{\rho u}_i^n - \frac{\Delta t}{\Delta x} \sqrt{\frac{\rho_i^n}{E_i^n}} \left(\lambda_{\rho u i}^{n+} \alpha_{\rho u i}^n + \lambda_{E i}^{n+} \alpha_{E i}^n + \lambda_{\rho u(i+1)}^{n-} \alpha_{\rho u(i+1)}^n + \lambda_{E(i+1)}^{n-} \alpha_{E(i+1)}^n \right) \\
&- \frac{\Delta t}{\Delta x} \frac{1}{2} \sqrt{\frac{\rho_i^n}{E_i^n}} \left[|s_{\rho u(i+1)}^n| \left(1 - |s_{\rho u(i+1)}^n| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{\rho u(i+1)}^n) \alpha_{\rho u(i+1)}^n + |s_{E(i+1)}^n| \left(1 - |s_{E(i+1)}^n| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{E(i+1)}^n) \alpha_{E(i+1)}^n \right] \\
&- \frac{\Delta t}{\Delta x} \frac{1}{2} \sqrt{\frac{\rho_i^n}{E_i^n}} \left[|s_{\rho u i}^n| \left(1 - |s_{\rho u i}^n| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{\rho u i}^n) \alpha_{\rho u i}^n + |s_{E i}^n| \left(1 - |s_{E i}^n| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{E i}^n) \alpha_{E i}^n \right].
\end{aligned}$$

$$\begin{aligned}
\Delta \bar{E}_i^{n+1} &= \\
&= \Delta \bar{E}_i^n - \frac{\Delta t}{\Delta x} \left(-\lambda_{\rho u i}^{n+} \alpha_{\rho u i}^n + \lambda_{E i}^{n+} \alpha_{E i}^n + (-\lambda_{\rho u(i+1)}^{n-} \alpha_{\rho u(i+1)}^n) + \lambda_{E(i+1)}^{n-} \alpha_{E(i+1)}^n \right) \\
&- \frac{\Delta t}{\Delta x} \frac{1}{2} \left[-|s_{\rho u(i+1)}^n| \left(1 - |s_{\rho u(i+1)}^n| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{\rho u(i+1)}^n) \alpha_{\rho u(i+1)}^n + |s_{E(i+1)}^n| \left(1 - |s_{E(i+1)}^n| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{E(i+1)}^n) \alpha_{E(i+1)}^n \right] \\
&- \frac{\Delta t}{\Delta x} \frac{1}{2} \left[-|s_{\rho u i}^n| \left(1 - |s_{\rho u i}^n| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{\rho u i}^n) \alpha_{\rho u i}^n + |s_{E i}^n| \left(1 - |s_{E i}^n| \frac{\Delta t}{\Delta x} \right) \phi(\theta_{E i}^n) \alpha_{E i}^n \right].
\end{aligned}$$

Antes de comenzar con la siguiente iteración hay que recuperar las cantidades *no promediadas*, interpolando para obtener su valor en los puntos originales de la malla. Una interpolación lineal es suficiente mientras se tenga una resolución alta. Enseguida deben calcularse u y P , la primera a partir del momento y la segunda a partir de la ecuación de estado:

$$\begin{aligned}u &= \frac{\rho u}{\rho} \\P &= (\gamma - 1) \left(E - \frac{1}{2} \rho u^2 \right) .\end{aligned}$$

Capítulo 8

Conclusiones

La Tesis que con este párrafo termino ha resultado ser una inagotable y prolífica fuente de investigación y documentación para su autor. Las expectativas iniciales, las pretenciones y los métodos a seguir se reformulan continuamente a causa de las dificultades encontradas y los cambios de prioridades e intereses.

Sucede que este trabajo se ha convertido en un tratado amplio sobre métodos numéricos de Alta Resolución que pueden ser aplicados no sólo en hidrodinámica, sino incluso en las más diversas áreas de la ciencia donde sea preciso resolver un sistema hiperbólico de ecuaciones no-lineales escrito en forma conservativa. Se ha comenzado con ecuaciones escalares *simples*, relativamente, como la ecuación de advección, la cual resultó ser el problema paradigmático y central. Las personas interesadas en el estudio de este texto han de notar que la ecuación de advección y la ecuación de Burgers han sido los dos únicos casos para los cuales se ha comenzado desde cero en la tarea de construir el esquema numérico que calcule su solución numérica, ilustrando el funcionamiento de una malla discretizada en volúmenes finitos e ilustrando las ideas elementales y nociones de Alta Resolución mediante reconstrucciones lineales continuas a trozos, discretizando las ecuaciones y calculando los flujos explícitamente en las fronteras de cada celda de la malla.

La razón por la cual la ecuación de advección se ha convertido en el eje principal de la teoría es que los resultados obtenidos al resolver este problema han sido reformulados, reutilizados y generalizados cuando se han abordado problemas más complicados, a saber: sistemas de ecuaciones. El soporte de muchos de los argumentos es la posibilidad de encontrar una forma en la que los sistemas puedan escribirse como sistemas de ecuaciones de advección. A partir de ese momento se cambió el enfoque de tal manera que la potencia del método se multiplicara y fuera capaz de atacar problemas más generales: sistemas de ecuaciones de coeficientes constantes y variables, culminando con el esquema más complejo: aquel que resuelve sistemas no-lineales. El punto importante fue escribir el método en términos de limitadores de flujo y no de pendientes, como se hizo para los esquemas primitivos e intuitivos que se formularon para la ecuación de advección.

La construcción del método es progresiva y pretende ser didáctica, lo que le hace comple-

tamente transparente, abriendo la posibilidad de enriquecerse con nuevas contribuciones y análisis más profundos.

En cuanto a los resultados obtenidos en cada problema abordado, pueden hacerse algunas generalizaciones. Por experiencia vívida: los problemas no-lineales son más difíciles de tratar y su convergencia a segundo orden requiere de condiciones más variadas que los sistemas lineales. Lo anterior era de esperarse, pues las no-linealidades dan lugar a la aparición de discontinuidades en la solución, para las cuales sólo se aspira a obtener convergencia a primer orden. Al respecto de la convergencia hay que mencionar que hay factores que la afectan directamente y que representan potenciales fuentes de investigación posteriores. En particular, si toda extrapolación hacia celdas fantasma (ya sea para calcular valores en la frontera para pendientes de reconstrucciones o de la misma cantidad conservada), es lineal, se cometerán errores numéricos de primer orden que se propagarán hacia el centro del dominio y opacarán los errores a segundo orden cometidos al evolucionar el perfil, el cual representa el interés principal. Otra fuente muy interesante de pruebas e investigaciones es el algoritmo para calcular los valores de la función evolucionada en los puntos de la malla donde originalmente fueron definidos los datos iniciales, es decir, después de evolucionar los promedios de la función de interés, es necesario encontrar los valores de la función *no promediada* en las interfaces de cada celda: para ello se han implementado interpolaciones lineales: importantes candidatas a ser las responsables de no obtener una convergencia *limpia* a segundo orden.

Para finalizar, valga mencionar que las pruebas realizadas en el séptimo capítulo para las ecuaciones de Euler marcan el comienzo de una etapa posterior de investigación de vanguardia en el área de la hidrodinámica relativista. Una vez resuelto el problema hidrodinámico de la propagación de ondas de choque en simetrías varias, en presencia de fuerzas gravitacionales internas y externas, y fusionando esto con los métodos existentes para la evolución de la geometría del espacio-tiempo cerca de fuentes gravitacionales con cierta simetría, se abren amplísimas posibilidades para explorar diversas situaciones astrofísica, como la evolución del espacio-tiempo en presencia de materia autogravitante. Ejemplos particulares son: la simulación del colapso de una estrella de neutrones, supernovas o explosiones de rayos gamma.

Es la sincera esperanza del autor, el haber contribuido a la consolidación y entendimiento de un área importante en la investigación de fenómenos de relevancia actual y servir como referencia a estudiantes interesados en el tema. Si al menos esta Tesis ayuda a alguien a la comprensión de algún concepto o a la implementación de algún algoritmo, se habrá cumplido el objetivo.

Referencias

- [1] Randall J. LeVeque. *Finite Volume Methods for Hyperbolic Problems*, Cambridge University Press, 2002.
- [2] Heinz-Otto Kreiss et al. *Time Dependent Problems and Difference Methods* Wiley-Interscience, 1995.
- [3] F. S. Guzmán, *Introduction to numerical relativity through examples*. Proceedings of the VI workshop of the DGFM-SMF. Rev. Mex. Fís. in press.
- [4] S. K. Godunov, Mat. Sb., 47:271, 1959.
- [5] W. H. Press, S. A. Teukolsky, W. T. Vetterling and B. P. Flannery, *Numerical Recipes in Fortran*. Cambridge University Press, 1992.
- [6] Kenneth Hoffman, Ray Kunze *Linear Algebra* Prentice-Hall, 1973.
- [7] E. F. Toro. *Riemann solvers and numerical methods for fluid dynamics* Springer-Verlag, Berlin, Heidelberg, 1997.
- [8] G. Sod. *Numerical Methods for Fluid Dynamics* Cambridge University Press, 1985.
- [9] J. A. Font. *Numerical hydrodynamics in general relativity* Living Reviews in Relativity (www.livingreviews.org), 2003.
- [10] B. van Leer. *Towards the ultimate conservative difference scheme IV. A new approach to numerical convection* J. Comput. Phys., 23:276-299, 1977.
- [11] Miguel Alcubierre Moya. *Investigations in Numerical Relativity* Tesis Doctoral, University of Wales, 1994.
- [12] A. J. Chorin y J. E. Marsden. *A Mathematical Introduction to Fluid Mechanics* Springer-Verlag New York, Inc. 1990, Tercera edición.
- [13] Russell, Lynn y George Adebisi. *Termodinámica Clásica* Ed. Addison Wesley Iberoamericana, S. A. México, 1997.