

Santiago Esteban Contreras Aranda
Bertha Ulloa Rubio

**EL MUNDO DE LA
TOMA DE
DECISIONES Y
LOS OPERADORES
MATEMÁTICOS**

Matemáticas

Santiago Esteban Contreras Aranda,
Bertha Ulloa Rubio

El mundo de la toma de decisiones y los operadores matemáticos

Atik Editorial



E15D N49-59 y Olivos, San Isidro. Código postal 170515.
Quito, Ecuador

Atik Editorial, es una iniciativa del Centro de Investigaciones
CICSHAL y está a cargo del departamento de Comunicación y
Difusión Científica.

www.atikeditorial.com

Primera Edición: 2026

Derechos de autor/Copyright:

Editorial: Atik Editorial

Materia Dewey: 650 - Gerencia y servicios auxiliares

Clasificación Thema: KJMD - Toma de decisiones en la gestión de empresas | PBW - Matemáticas
aplicadas | KJT - Investigación operativa

Público objetivo: Profesional/Académico

BISAC: BUS019000

Colección: Matemáticas

Soporte: Digital

Formato: Epub (.epub)/PDF (.pdf)

Publicado: 2026-07-06

ISBN: 978-9907-820-07-2

Disponible para su descarga gratuita en <http://atikeditorial.com>

El presente libro tienen el aval del Centro de Investigaciones en Ciencias y Humanidades
desde América Latina - CICSHAL.



Título: El mundo de la toma de decisiones y los operadores matemáticos

The world of decision-making and mathematical operators

O mundo da tomada de decisão e dos operadores matemáticos

(APA 7)

Contreras Aranda, S. E., & Ulloa Rubio, B. (2026). *El mundo de la toma de decisiones y los operadores matemáticos*. Atik Editorial. <https://doi.org/10.46652/atikbook27>



Este título se publica bajo una licencia de Atribución 4.0 Internacional (CC BY 4.0) la cual está disponible en: <https://creativecommons.org/licenses/by/4.0/deed.es>

Se debe dar crédito de manera adecuada, brindar un enlace a la licencia, e indicar si se han realizado cambios. Puede hacerlo en cualquier forma razonable, pero no de forma tal que sugiera que usted o su uso tienen el apoyo de la licenciante.

Las consultas relativas a la reproducción fuera del ámbito de esta licencia deberán enviarse al Departamento de Comunicación y Difusión Científica de CICS HAL a la siguiente casilla de correo: info@atikeditorial.com

Los enlaces a sitios web de terceros son facilitados por **Atik** Editorial de buena fe y a título meramente informativo. **Atik** Editorial declina toda responsabilidad por el material contenido en cualquier sitio web de terceros al que se haga referencia en esta obra.

AVAL DE REVISIÓN POR PARES

El presente libro constituye el resultado de un riguroso proceso de investigación académica, cuya calidad metodológica y solidez argumental han sido validadas mediante un sistema de revisión por pares externos implementado bajo el protocolo de doble ciego, bajo la supervisión del Centro de Investigaciones en Ciencias y Humanidades desde América Latina (CICSHAL). Como garantía de transparencia y rigor científico, los informes de evaluación realizados por los especialistas designados se conservan en el archivo institucional de la editorial, a disposición de las instancias que así lo requieran.

This book is the result of a rigorous academic research process, whose methodological quality and argumentative solidity have been validated through an external peer-review system implemented under a double-blind protocol, under the supervision of the Center for Research in Sciences and Humanities from Latin America (CICSHAL). As a guarantee of transparency and scientific rigor, the evaluation reports prepared by the designated specialists are preserved in the publisher's institutional archives, available to any party that may require them.

info@atikeditorial.com



AUTORES

AUTHORS

- Santiago Esteban Contreras Aranda

Universidad Nacional Mayor de San Marcos | Lima | Perú

<https://orcid.org/0000-0002-0698-0893>

scontrerasa@unmsm.edu.pe

Doctor en Ingeniería de Sistemas UNFV, Maestro en Ingeniería de Transportes Universidad Federal de Rio de Janeiro Brasil, Licenciado en Matemática Pura en Universidad Nacional Mayor de San Marcos, experiencia en modelaje matemático con enfoque holístico y pensamiento Sistémico. Con muchas publicaciones y es RENACYT en Perú.

- Bertha Ulloa Rubio

Universidad César Vallejo | Trujillo | Perú

<https://orcid.org/0000-0002-3438-9371>

bulloa@ucvvirtual.edu.pe

Doctora en Administración de la Educación, Maestra en Ingeniería de Sistemas y Computación UFRJ - Brasil Doctorado(e) en Ing. de Producción UENF- Brasil Docente Universitario Exdocente e investigadora Universidad Estadual do Norte Fluminense.

RESUMEN

Este libro ofrece una base matemática rigurosa para el cálculo fraccionario y su aplicación en la toma de decisiones. Explora los operadores diferenciales e integrales de Riemann-Liouville y Caputo, fundamentales para modelar sistemas con memoria en fenómenos del mundo real: viscoelasticidad, neurociencia, finanzas y dinámicas sociales. A través de teoremas de existencia, unicidad y estabilidad, se analizan ecuaciones diferenciales fraccionarias de uno y varios términos. Incluye funciones de Mittag-Leffler y métodos numéricos Adams-Bashforth-Moulton. Ideal para matemáticos, ingenieros y científicos aplicados que buscan herramientas avanzadas más allá del cálculo tradicional.

Palabras clave:

Cálculo fraccionario; Operadores de Caputo; Ecuaciones diferenciales fraccionarias; Funciones de Mittag-Leffler; Métodos numéricos de Adams.

ABSTRACT

This book provides a rigorous mathematical foundation for fractional calculus and its application in decision-making. It explores the Riemann-Liouville and Caputo fractional differential and integral operators, which are essential for modeling systems with memory in real-world phenomena such as viscoelasticity, neuroscience, finance, and social dynamics. Through existence, uniqueness, and stability theorems, fractional differential equations of single and multi-term are analyzed. The book includes Mittag-Leffler functions and Adams-Bashforth-Moulton numerical methods. Ideal for mathematicians, engineers, and applied scientists seeking advanced tools beyond traditional calculus.

Keywords:

Fractional calculus; Caputo operators; Fractional differential equations; Mittag-Leffler functions; Adams numerical methods.

RESUMO

Este livro oferece uma base matemática rigorosa para o cálculo fracionário e sua aplicação na tomada de decisões. Explora os operadores diferenciais e integrais fracionários de Riemann-Liouville e Caputo, fundamentais para modelar sistemas com memória em fenômenos do mundo real: viscoelasticidade, neurociência, finanças e dinâmicas sociais. Por meio de teoremas de existência, unicidade e estabilidade, são analisadas equações diferenciais fracionárias de um e múltiplos termos. Inclui funções de Mittag-Leffler e métodos numéricos de Adams-Bashforth-Moulton. Ideal para matemáticos, engenheiros e cientistas aplicados que buscam ferramentas avançadas além do cálculo tradicional.

Palavras-chave:

Cálculo fracionário; Operadores de Caputo; Equações diferenciais fracionárias; Funções de Mittag-Leffler; Métodos numéricos de Adams.

CONTENIDO

AVAL DE REVISIÓN POR PARES 6

Prefacio 19

01 Capítulo

ESTRUCTURA MATEMÁTICA DE LOS OPERADORES DIFERENCIALES E INTEGRALES

FRACCIONARIOS

Introducción 22

Estructura matemática del cálculo fraccionario 25

Incitación 25

La idea básica 32

Teorema A. Fundamental del Cálculo "TFC" 33

Definición 1.1 33

Lema 1.1 35

Lema 1.2 36

Teorema 1.3 37

Definición 1.2 37

Definición 1.3 37

Teorema B. Fundamental en Espacios de Lebesgue "TFEL" 38

Definición 1.4 38

02 Capítulo

OPERADORES INTEGRAL Y DIFERENCIAL DE RIEMANN-LIOUVILLE

44

Operador integral de Riemann-Liouville 45

Definición 2.1 45

Teorema 2.1 46

Teorema 2.2 47

Teorema 2.3 47

Teorema 2.4 47

Teorema 2.5 49

Teorema 2.6 54

Teorema 2.7 57

Teorema 2.8 58

Teorema 2.9 60

Teorema 2.10 62

Derivadas de Riemann-Liouville 66

Definición 2.2 67

Lema 2.11 67

Lema 2.12 68

Teorema 2.13 70

Teorema 2.14 72

Teorema 2.15 73

Teorema 2.16 74

Teorema 2.17 76

Teorema 2.A. Fórmula de Leibniz 76

Teorema 2.18. Fórmula de Leibniz para operadores de Riemann-Liouville 77

Teorema 2.B. Fórmula de Faà di Bruno 79

Teorema 2.19 Fórmula de Faà di Bruno para operadores de Riemann-Liouville.	81
Teorema 2.20	82
Teorema 2.21	84
Relaciones entre los operadores integrales y derivadas de Riemann-Liouville	86
Teorema 2.22	87
Teorema 2.23	87
Teorema 2.C. Expansión de Taylor.	89
Teorema 2.24. Expansión fraccionaria de Taylor	89
Operadores de Grunwald-Letnikov	90
Definición 2.3	92
Teorema 2.25	92
Teorema 2.26	95
Definición 2.4	96

03 Capítulo

OPERADORES DIFERENCIALES E INTEGRALES FRACCIONARIOS DE CAPUTO

99

Definición y propiedades básicas	100
Definición 3.1.	100
Teorema 3.1	102
Definición 3.2	103
Demostración (Teorema 3.1)	104
Lema 3.2	106
Lema 3.3	106
Lema 3.4	107
Lema 3.5	107
Lema 3.6	108
Teorema 3.7	108
Teorema 3.8	109
Teorema 3.9. Expansión de Taylor para derivados de Caputo	109
Definición 3.3	111
Teorema 3.10	112
Lema 3.11	113
Lema 3.12	113
Lema 3.13	114
Teorema 3.14	116
Teorema 3.15	117
Teorema 3.16	118
Observación 3.5	119
Representaciones no clásicas de operadores de Caputo	120
Teorema 3.18	121
Teorema 3.19	123
Teorema 3.20	124
Teorema 3.21	126

04 Capítulo

FUNCIONES DE MITTAG-LEFFLER

130

Definición 4.1	131
Teorema 4.1	133

Teorema 4.2	135
Teorema 4.3	136
Teorema 4.4	138
Teorema 4.5	139
Teorema 4.6	140

05 Capítulo

ECUACIONES DIFERENCIALES FRACCIONARIAS 143

Existencia y singularidad para ecuaciones diferenciales fraccionarias de Riemann-Liouville
144

Teorema 5.1	144
Lema 5.2	146
Lema 5.3	149

06 Capítulo

ECUACIONES DIFERENCIALES FRACCIONARIAS DESDE LA ÓPTICA DE CAPUTO 155

Teorema 6.1	157
Lema 6.2	158
Corolario 6.3	167
Observación 6.6.	168
Corolario 6.4	168
Teorema 6.5	169
Corolario 6.6	173
Corolario 6.7	173
Teorema 6.8	174
Corolario 6.9	175
Lema 6.10	175
Teorema 6.11	182
Teorema 6.A	184
Teorema 6.12.	186
Teorema 6.13	191
Teorema 6.14	192
Teorema 6.15	192
Corolario 6.16	194
Teorema 6.17	194
Teorema 6.18	195
Influencia de los datos perturbados	199
Lema 6.19	201
Teorema 6.20.	203
Teorema 6.21	204
Teorema 6.22	206
Corolario 6.23	208
Corolario 6.24	208
Suavidad de las soluciones	209
Teorema 6.B	209
Teorema 6.25	210
Teorema 6.26	211
Teorema 6.27	212

Teorema 6.28	213
Teorema 6.29	214
Corolario 6.30	215
Corolario 6.31	216
Teorema 6.32	217
Corolario 6.33	217
Corolario 6.34	218
Teorema 6.35.	223
Corolario 6.36	224
Teorema 6.37.	226
Teorema 6.38	227
Problemas de valores en la frontera	228
Lema 6.39	229
Teorema 6.40	231
Teorema 6.41	232
Lema 6.42	234
Teorema 6.43	235
Teorema 6.44	236

07 Capítulo

ECUACIONES DIFERENCIALES FRACCIONARIOS DE CAPUTO DE UN SOLO TÉRMINO: RESULTADOS AVANZADOS PARA CASOS ESPECIALES

238

Problemas de valores iniciales para ecuaciones lineales 239

Teorema 7.1 240

Teorema 7.2 242

Teorema 7.3 248

Teorema 7.4 250

Teorema 7.5 251

Teorema 7.6 253

Teorema 7.7 254

Teorema 7.8 255

Teorema 7.10 258

Definición 7.1 258

Lema 7.11 258

Teorema 7.12 262

Teorema 7.13 264

Teorema 7.14 268

Teorema 7.15 271

Teorema 7.16 274

Teorema 7.17 276

Teorema 7.18 277

Teorema 7.19 278

Estabilidad de ecuaciones diferenciales fraccionarias 278

Definición 7.2 280

Teorema 7.20 282

Teorema 7.21 285

Teorema 7.22 287

Ecuaciones singulares 291

Teorema 7.23	294
Teorema 7.24	295
Teorema 7.25	296

08 Capítulo

ECUACIONES DIFERENCIALES FRACCIONARIA DE CAPUTO DE VARIAS ETAPAS 299

Definición 8.1	301
Teorema 8.1	304
Teorema 8.2	307
Teorema 8.3	308
Teorema 8.4	309
Lema 8.5. Primera desigualdad de Gronwall para ecuaciones de dos términos.	310
Lema 8.6. Segunda desigualdad de Gronwall para ecuaciones de dos términos.	312
Teorema 8.7	313
Teorema 8.8	315
Teorema 8.9	318
Teorema 8.10	322
Teorema 8.11	325
Lema 8.12	328

Referencias 335

Apéndices 352

Apéndice A	353
Apéndice B	374
Apéndice C	376
Apéndice D	423

LISTADO DE FIGURAS

Figura 1. Gráficas de las soluciones de los cinco problemas de valor inicial del Ejemplo 191

Figura 2. Gráficas de soluciones aproximadas vecinas para el problema de valor terminal de Ejemplo 6.3 198

Figura 3. Gráficas de $E_n(-x^n)$ 249

Figura 4. Esquema de las gráficas de $u_0(x)$ 250

Figura 5. Esquemas de las gráficas de $u_1(x)$ 251

Figura 6. Figura de número de ceros de u_0 253

LISTADO DE TABLAS

Tabla 1. Resultados de la búsqueda de bisección para el Ejemplo 6.3 197

Prefacio

“Existe un universo infinito de modelos mentales y micro mundos conceptuales con sus complejidades donde el todo es mucho más que la suma sus partes.”

“Examina, analiza tu niño interno, ahí se encuentra la belleza y fuente del bien que debe emanar continuamente y que debes profundizar siempre, con honestidad, modestia, liberalidad, no desdeñar la enseñanza pública, la tolerancia en el trabajo y despreciar ciertas artes inútiles y vanas. Finalmente corregir y aprender de tus costumbres y cuidar de ellas.”

Burseca

Mis amigos, encontrar el camino conductor es fácil, es tumultuoso, complejo, explorarlo más aún, si por infortunio no logras todas tus metas, recuerda que las cosas materiales se desvanecen, pero las obras que estructuran el pensamiento permanecen, y serás el pionero de lo que sabemos, existen muchos hombres, y pocos timoneles. Estas reflexiones están dedicado a algunas cuestiones del pensamiento complejo de la toma de decisiones, es decir, del ¿Qué? ¿Para qué? y el ¿Cómo? Es el conocimiento complejo del mundo que nos rodea, es infinito, destacando que los primeros pasos de esta teoría en sí se remontan a 380 a 350 a.C., pero en realidad el tema sólo surgió realmente en las últimas décadas del siglo anterior. Por una característica particular de ingenieros, científicos y algunos humano inquieto, curiosos que han desarrollado nuevos modelos que involucran al Pensamiento Sistémico, sacando de su confort al Pensamiento Científico Mecanicista, Reduccionista. Nuevos modelos que se aplican con éxito, en mecánica, en bio-química, en ingeniería eléctrica, en medici-

na, en sociología y en cualquier área donde participa directamente el hombre como actor interno o externo; en Neurociencia, con la emulación del cerebro humano, etc.

“**El mundo de la toma de decisiones y los operadores matemáticos**” es la tercera de las reflexiones, en esta dirección, quiero dar rienda suelta a su imaginación con la siguiente pregunta: ¿Por qué estás en este mundo?, considero que están pensando en una infinidad de respuestas, unas muy simples y otras complejas, sin embargo, en realidad todas tienen su propio sesgo que estarán en función de sus modelos mentales o micro mundos conceptuales, estarán en una discusión con su yo en tomar decisiones, pero estarán de acuerdo, que para dar respuestas a esas inquietudes deben usar sus mejores instrumentos y se preguntarán: ¿estos instrumentos en verdad son los mejores?, entonces dirán: ¿Cómo han sido contruidos, esos instrumentos?, ¿Han sido contruidos con las mejores herramientas?, y entrarán en pánico, es allí, donde identificarán que en la estructura de los instrumentos se han utilizado una diversidad de conocimientos y tecnología de punta, sin embargo, se volverán a preguntar ¿en realidad son los mejores?. Por ejemplo: en matemáticas de repente no lo son, pues lo último que conozco son ecuaciones diferenciales con derivadas parciales y su diversidad de matices como la teoría de semigrupos y otros, no utilizamos tales como: **operadores diferenciales e integrales fraccionarios** que son muy buenos como herramientas para la estructura de modelos que mantienen un atributo: “la memoria”. Y es en esta dirección que apuntan las reflexiones de este libro, y seguirá la interrogante, cómo si cerráramos los ojos y observamos el mundo que nos rodea, pues no consideramos la incertidumbre al estructurar el modelo, entonces no contamos con los mejores instrumentos para tomar decisiones, puesto que

no consideramos las brillantes herramientas matemáticas, sólo son usadas en las matemáticas y ¿Qué pasó en las otras áreas del conocimiento?, sería necesario echar un vistazo en esa dirección.

01 Capítulo

***ESTRUCTURA MATEMÁTICA DE LOS
OPERADORES DIFERENCIALES E
INTEGRALES FRACCIONARIOS***

“Escucha a tu niño interno que te enseñara a obrar con libertad de espíritu desembrado de vanos respetos, fijarte en tus resoluciones sin perplejidad y no gobernarse por otros principios; sino por la buena razón, aún en las cosas mínimas”. BURSECA

Introducción

En estas reflexiones hablaremos sobre algunos puntos de cálculo fraccionario, la teoría de operadores diferenciales e integrales de orden no entero, serán las que nortearán nuestro pensamiento hacia la toma de decisiones en el Mundo Real Consensual Organizacional Dinámico Sensitivo e Inteligible “MRCODSI”, y en particular a operadores de ecuaciones diferenciales (OED). Destacando que, los primeros pasos de esta teoría en sí se remontan a la primera mitad del siglo XIX, aunque en realidad el tema surgió realmente en las últimas décadas del siglo XX. Como consecuencia del desarrollo de nuevos modelos generados por ingenieros y científicos involucrando a las Ecuaciones Diferenciales Fraccionarias (EDFs), las cuales han sido aplicados en modelamiento de fenómenos como: viscoelasticidad, visco plasticidad, polímeros, proteínas, transmisión de ondas de ultrasonido, tejido humanos con cargas mecánicas, emulación del cerebro humano, identificación de rostros, etc. Incluso en una diversidad de fenómenos simples y complejos del Mundo Real Consensual Organizacional Dinámico Sensitivo e Inteligible, buscando la competitividad productiva, que le permitan sobrevivir en este presente siglo de globalización, centralización, agobiante por la diversidad de incertidumbre, azaroso, caótico; en donde, fenómenos como: el hambre, la desocupación, el hacinamiento, el analfabetismo, la prostitución, la corrupción, la contaminación ambiental, la desequida

económica, social, educativa, y de salud, son fenómenos muy comunes que caracterizan al siglo XXI. Por consiguiente, se deben de considerar herramientas muy representativas para construir los modelos en la toma de las decisiones. Y aquellas que más se aproximan para representar estos fenómenos son los operadores fraccionarios, es más, son los Operadores Fraccionarios Difusos (OFD), los cuales serán más adelante, materia de análisis en una cuarta oportunidad.

En estas circunstancias la teoría matemática, parece estar rezagada detrás de las necesidades de estas aplicaciones. Aunque, existen algunos libros tratando los aspectos que se pueden resumir como el lado “matemático puro” de los fenómenos citados, sin tener en consideración, el punto de vista del ingeniero sin una rigurosa justificación matemática. Estas reflexiones intentan llenar el vacío entre estos dos enfoques, estableciendo una teoría matemática sólida de las ecuaciones diferenciales que han demostrado ser relevantes en la práctica, y proporcionan un minucioso análisis matemático. Para ser autónomos, repetimos los fundamentos del cálculo fraccionario antes de pasar al tema principal. El objetivo de estas reflexiones es proporcionar una base sólida matemática, que luego pueda usarse para la construcción de métodos numéricos eficientes y confiables para ecuaciones diferenciales fraccionarias “EDF”. Nosotros creemos firmemente en un desarrollo exitoso, y una minuciosa comprensión de tales esquemas numéricos, es posible con un estable antecedente analíticos.

Es en ese sentido, el lector estará familiarizado con el cálculo clásico diferencial, cálculo integral y la teoría elemental de las ecuaciones diferenciales. El conocimiento de la teoría de la integración de Lebesgue es útil, sin embargo, es necesario incluir algunas otras indicadas a continuación.

Estructura matemática del cálculo fraccionario

Incitación

Estas reflexiones tratan sobre problemas que surgen en el área del cálculo fraccionario, una rama de las matemáticas, tan antigua como el cálculo clásico. Su origen se remonta (Ross, 1977), a finales del siglo XVII, cuando Newton y Leibniz desarrollaron los fundamentos de diferencial y cálculo integral. En particular, Gottfried Wilhelm Leibniz introdujo el símbolo:

$\frac{d^n}{dx^n}f(x)$ para denotar la n -ésima derivada de una función f . Cuando informó esto en una carta al Marqués de L'Hôpital considerando aparentemente la suposición implícita que $n \in \mathbb{N}$, pero el Marqués de L'Hospital respondió: ¿Qué significa, $\frac{d^n}{dx^n}f(x)$ si $n = 1/2$? Esta carta del Marqués de L'Hospital, escrita en 1695, es considerada hoy en día como la primera ocurrencia llamada derivada fraccionaria, y el hecho de que el Marqués de L'Hospital preguntó específicamente para $n = 1/2$, es decir, una fracción, número racional, en realidad dio origen al nombre de esta parte de las matemáticas. Este nombre ha permanecido en uso desde entonces, aunque es bien conocido por ahora, que no hay razón para restringir n al conjunto de números racionales.

De hecho, como veremos en estas reflexiones, cualquier número real, racional o irracional, es el escenario de los operadores fraccionarios, al menos por las consideraciones analíticas que concentraremos en ciertos puntos, pero no en todos. Los métodos numéricos para la solución de algunos tipos de las ecuaciones diferenciales pueden encontrar problemas cuando se admiten números reales arbitrarios (Diethelm & Ford, s. f.). De hecho, incluso

los números complejos pueden ser permitidos, sin embargo, estos acápites están más allá del alcance de estas reflexiones.

Entonces, ¿Cuál es el alcance de estas reflexiones? La escritura de estas reflexiones tiene esencialmente: la motivación de la enorme cantidad de aplicaciones muy interesantes y novedosas de Ecuaciones Diferenciales Fraccionarias (EDFs) en física, química, ingeniería, finanzas, ciencias sociales, ciencias biológicas y otras ciencias que se han desarrollado en las últimas décadas. Algunos ejemplos tempranos se dan en el libro de procesos de difusión de Oldham con Spanier (1974); los artículos clásicos de Bagley con Torvik (2000); (Torvik & Bagley, 1984), Caputo (1967) y Caputo con Mainardi (1971); (Caputo & Mainardi, 2008), estos artículos que tratan sobre el modelado de la mecánica, propiedades de los materiales; procesamiento de señales de Olmstead con Handelsman (1976), que también se ocupan de problemas la difusión; se describen resultados más recientes, por ejemplo, en el trabajo de Benson advección y dispersión de solutos en medios naturales porosos o fracturados (1998); Bai con Feng (2007) y Cuesta con Finat Codes (2003), también en procesamiento de imágenes; Chern (1993), Diethelm (2002a) y Freed (2002), Diethelm con Luchko en el modelado del comportamiento de materiales viscoelásticos y viscoplásticos bajo influencias condiciones externas (2004); Dokoumetzidis et al. (2010); Popovic et al., y Verotta en la farmacocinética (2010); Freed y Diethelm (1999a) con Magin en la bioingeniería (2006); Gaul et al. (1991), sobre la descripción de sistemas mecánicos sujetos a amortiguamiento, Glöckle y Nonnenmacher (1995), en relación a la relajación y cinética de reacción de polímeros; Gorenflo y Rutman (1995), con los llamados procesos ultra lentos; Gorenflo y Mainardi asociado a las conexiones con la teoría de paseos aleatorios (Gorenflo & Vivoli, 2003) (Gorenflo

& Mainardi, 2008); estos artículos especialmente con respecto a aplicaciones a modelos matemáticos en finanzas (Mainardi et al., 2000) (Scalas et al., 2000); Joulin y Roquejoffre modelado de combustión (Joulin, 1985) (Audounet & Roquejoffre, 1998); Metzler con la relajación en redes de polímeros rellenos (Metzler et al., 1995); Podlubny y Caponetto asociada a la teoría de control (Podlubny, 1995) (Caponetto et al., 2010); la última publicación también con detalles sobre el hardware implementación de controladores de orden fraccional (Podlubny, 1994); Podlubny, propagación de calor (Podlubny, 1995) (Podlubny, 1999) y Shaw, Warby y Whiteman con el modelado de materiales viscoelásticos (Shaw et al., 1997).

Un campo de aplicación completamente diferente y muy novedoso es el área de las matemáticas en la psicología, donde los sistemas de orden fraccional pueden usarse para modelar el comportamiento de seres humanos, tales como los siguientes artículos: “Modelos dinámicos de felicidad con orden fraccional” (Song et al., 2010) y “Modelos dinámicos de amor de orden fraccional” (Ahmad & El-Khazali, 2007).

Específicamente, la forma en que una persona reacciona a las influencias depende de la experiencia que él o ella haya hecho en el pasado. En otras palabras, los humanos tienen recuerdos, y lo veremos más adelante en estas reflexiones que los operadores fraccionarios son una herramienta muy natural para modelar fenómenos dependientes de la memoria.

Las encuestas o colecciones de tales aplicaciones también se pueden encontrar en Baleanu (2009), Gorenflo y Mainardi (1996); (Gorenflo & Mainardi, 2008); (Mainardi & Gorenflo, 2000); Hilfer (2000); Klages et al. (2008); Mainardi (2012); Matignon (1998); (Matignon, 1994) y Montseny et al. (1995); (Metzler et al., 1995);

Podlubny (2001); Sabatier et al. (2007); Tas et al. (2007). Además, hay algunas aplicaciones del cálculo fraccionario dentro de varios campos de las propias matemáticas, por ejemplo: en la investigación analítica de varios tipos de funciones especiales. Finalmente, nos referimos al trabajo de Woon que esencialmente menciona aplicaciones matemáticas, a su vez, tienen importantes implicaciones en otras ciencias como: la física, en la cual se resaltó que muchas de estas aplicaciones dieron lugar a un tipo de ecuaciones no cubiertas en la literatura matemática estándar; esto está conectado al hecho de que, en cierto sentido, la respuesta a la pregunta al Marqués de L'Hospital, que Leibniz no pudo encontrar, excepto por el caso especial $f(x) = x$ no es único. Hay muchas generalizaciones posibles de $\frac{d^n}{dx^n} f(x)$ el caso . Lo haremos y sólo discutiremos dos de ellos, la derivada de Riemann-Liouville y la de Caputo. El primer concepto es históricamente, desarrollado en obras de Abel, Riemann y Liouville en la primera mitad del siglo XIX, y aquel donde la teoría matemática ha sido establecida bastante bien por ahora, sin embargo, tiene ciertas características que generan dificultades al aplicarlo al “mundo real consensual”. Problemáticas que, como consecuencia, se desarrolló este último concepto es de cerca relacionado con la idea de Riemann-Liouville, pero se introdujeron ciertas modificaciones para evitar las dificultades antes mencionadas. Las implicaciones matemáticas de estas modificaciones no han sido investigadas completamente hasta ahora. En estas reflexiones pretendemos dar tanta información sobre este tema como sea posible actualmente.

La estructura de estas reflexiones se organiza de la siguiente manera: **Primero**, comenzamos en el capítulo 1, recordando algunos hechos clásicos del cálculo que forman la base de nuestra generalización prevista, y la ejemplificación de la aplicación del

cálculo fraccionario en mecánica; **Segundo**, introduciremos en el capítulo 2 y 3, los conceptos y definiciones fundamentales de cálculo fraccionario, la cual incluye, algunos resultados básicos relativos a operadores diferencial e integral de Riemann-Liouville. Tener en cuenta que, el objetivo principal de este libro es presentar una descripción completa de las propiedades de los operadores del tipo de Caputo y sobre la teoría de las ecuaciones diferenciales que involucran tales operadores, adicionalmente, por ello hemos decidido incluir los operadores de Riemann-Liouville por las siguientes razones:

Primero, esta decisión nos permite proporcionar una comparación de los dos enfoques. Por lo cual, inducimos que debería ser útil tanto para el lector que está familiarizado con la teoría clásica de los operadores de Riemann-Liouville; así como, para el novato en cálculo fraccionario, a quien le permitirá obtener una comprensión más profunda de las derivadas e integrales de Riemann-Liouville.

Segundo, por la aplicación de ambas teorías, aunque la versión de Caputo del cálculo fraccionario requiere una modificación de los operadores diferenciales de tipo Riemann-Liouville de orden fraccionario, los operadores integrales de Riemann-Liouville de orden fraccionario no necesitan someterse a algún cambio. Se pueden usar en el cálculo fraccionario de Caputo en su forma original.

Tercero, debido a los resultados de varias pruebas asociados a los operadores de Caputo. Los operadores se pueden dar de una manera relativamente simple, mediante el uso de propiedades relacionadas de operadores de Riemann-Liouville y el conocimiento preciso de cómo funcionan los dos tipos de operadores que están interrelacionados.

Continuando con la organización del texto: **Tercero**, introducimos en el capítulo 4, una clase de funciones con fundamental importancia en la teoría de ecuaciones diferenciales fraccionarias, las funciones de Mittag-Leffler, estas funciones nuevamente se verán en varios capítulos en adelante de nuestro texto, brindando algunos conocimientos básicos sobre ello; **Cuarto**, el núcleo de las reflexiones en los siguientes capítulos estarán dedicadas al estudio analítico de ecuaciones diferenciales fraccionarias EDFs, explicadas su conceptualización en el capítulo 5, antes de dirigir nuestra atención a las ecuaciones con operadores de Caputo entre los capítulos 6 al 8, en donde, abordamos cuestiones de existencia y unicidad de soluciones, e investigamos las propiedades de estas, propiciando un adecuado conocimiento de estas propiedades, no sólo es de interés por derecho propio, sino que también juega un papel importante en la construcción exitosa de métodos numéricos; **Quinto**, finalizaremos con un conjunto de apéndices en donde se describen cada uno de los métodos numéricos específicos presentados y utilizados en este texto, lo cual permitirá ayudar a superar los problemas causados por la escasez de métodos analíticos para el cálculo de soluciones a ecuaciones diferenciales fraccionarias.

Por supuesto, algunos enfoques analíticos propiamente entendidos existen, y los presentaremos en este texto; por ejemplo: el esquema de iteración de Picard y las fórmulas para ecuaciones lineales mencionadas anteriormente, sin embargo, como se puede esperar de la observación sobre la teoría clásica de ecuaciones diferenciales de primer orden: “No existe un método generalmente aplicable para encontrar una solución analítica a una ecuación diferencial fraccionaria dada arbitrariamente”, por tanto, necesitamos mencionar, varios otros métodos numéricos y analíticos para resolver ecuaciones diferenciales que se han desarrollado

en el orden fraccionario. La cual incluye: el método de descomposición generalmente atribuido a Adomian (1994); (Rèpaci, 1990), aunque, en realidad se remonta a una serie de documentos mucho más antiguos de Perrón, y es de conocimiento que tienen una convergencia bastante pobre con respecto a las propiedades en general, tal como: el método de iteración variacional (Dehghan & Tatari, 2007), el método de análisis de homotopía (Liang & Jeffrey, 2009), el método de perturbación de homotopía (Momani & Odibat, 2007), el método de transformada diferencial generalizada (Odibat et al., 2008), y algunos otros para todos los cuales las pruebas de convergencia están disponibles, sólo bajo condiciones bastante restrictivas, además muchos de los cuales se sabe que no convergen de manera satisfactoria de todos modos.

Por lo tanto, nos abstendremos de tratar estos enfoques a detalle en estas reflexiones. Además, considerando completar nuestro tratamiento de la teoría de las Ecuaciones Diferenciales Fraccionarias (EDFs) al proporcionar apéndices adicionales, dedicados a la recopilación de varios otros tipos de información útil: Una lista de todos los símbolos utilizados en las reflexiones y una breve tabla de las derivadas de Caputo de ciertas funciones importantes, sus correspondientes tablas para Riemann-Liouville y otras derivadas fraccionarias ya existentes en la literatura. Finalmente haremos uso de algunos resultados clásicos del análisis sobre temas como la función Gamma y la transformada de Laplace (Rasof, 1962); en aras de la exhaustividad, son incluidos dichos resultados en un apéndice.

Por supuesto, es posible establecer una teoría, y en base a esta teoría, discutir métodos numéricos para ecuaciones diferenciales fraccionarias parciales, es decir, diferenciales ecuaciones en más de una variable, donde al menos una de las derivadas par-

ciales involucrada no es un operador de orden entero. Muchos de los resultados presentados a continuación se pueden demostrar que las unidades son bloques de construcción útiles, en tal contexto, el tratamiento completo de estos problemas no es lo que pretendemos en estas reflexiones, sino dar un tratamiento integral de los, en otras palabras, ecuaciones diferenciales ordinarias. Se debe tener en cuenta, que algunos autores prefieren reservar el término ordinario para ecuaciones de orden entero, y usar la expresión extraordinario para ecuaciones diferenciales fraccionarias en una variable. No seguiremos su nomenclatura porque es la opinión del autor, que estas ecuaciones pueden y deben usarse como una herramienta matemática absolutamente normal, y de ninguna manera extraordinaria, la cual es útil para muchas aplicaciones.

Durante el transcurso de las reflexiones en el presente libro, ocasionalmente enunciamos teoremas del clásico análisis con fines de comparación con sus contrapartes fraccionarias o con el fin de ilustrar las ideas detrás de las generalizaciones. Estos teoremas clásicos suelen ser bien conocidos y no se proporcionarán pruebas de ellos. Para dejar clara la distinción, no seguirán el esquema de numeración estándar de los otros teoremas; en cambio se les asignará una etiqueta consistente en el número del capítulo donde debe ser encontrado, seguido de una letra mayúscula como, por ejemplo, el Teorema A.

La idea básica

La idea básica detrás del cálculo fraccionario está íntimamente relacionada con un estándar clásico de los resultados del

cálculo diferencial e integral clásico, el teorema fundamental del cálculo.

Teorema A. Fundamental del Cálculo “TFC”

Sea $f: [a, b] \rightarrow \mathbb{R}$ una función continua, y sea $F: [a, b] \rightarrow \mathbb{R}$ definida por,

$$F(x) := \int_a^x f(t) dt$$

Entonces, F es diferenciable y,

$$F' = f.$$

Obsérvese: que tenemos una relación muy estrecha entre los operadores diferenciales e integrales.

Entonces uno de los objetivos del cálculo fraccionario es mantener esta relación y una sistematización adecuada para su generalización.

Existe la necesidad de tratar con operadores fracciones, es decir resulta útil primero discutir propiedades antes de iniciar con los Operadores Diferenciales Fraccionarios ODF, todo esto proporciona una respuesta a la pregunta realizada por el Marqués de L'Hospital.

Iniciaremos nuestras reflexiones con las definiciones singulares.

Definición 1.1

Primero. Por $\mathcal{D}$, denotamos el operador diferencial que mapea una función diferenciable sobre su derivada, es decir: $\mathcal{D}f(x) := f'(x)$

Segundo. Por J_a , denotamos el operador integral que mapea una función f , que se supone que es Riemann integrable en el intervalo compacto $[a, b]$, con una primitiva centrada en a , es decir:

$$J_a f(x) := \int_a^x f(t) dt, \text{ para } a \leq x \leq b$$

Tercero. Para $n \in \mathbb{N}$ usamos los símbolos D^n y J_a^n para denotar, las n iteraciones de D y J_a , respectivamente, es decir, establecemos

$$\begin{aligned} D^1 &:= D \text{ y } J_a^1 := J_a \\ D^n &:= D D^{n-1} \\ J_a^n &:= J_a J_a^{n-1} \text{ para } n \geq 2 \end{aligned}$$

La pregunta clave ahora es: ¿Cómo podemos extender estos conceptos de la definición anterior, principalmente Tercera para $n \notin \mathbb{N}$?

Una vez que hayamos proporcionado dicha respuesta, debemos solicitar las propiedades de mapeo de los operadores resultantes, y en particular, estas deben de incluir respuestas a las preguntas sobre su dominio y rango.

Obsérvese: que el Teorema A dice, según nuestra denotación,

$$D J_a f = f$$

lo que implica que $D^n J_a^n f = f \dots (1.1)$

Para $n \in \mathbb{N}$, es decir, D^n es el inverso izquierdo de J_a^n en un espacio adecuado de funciones.

Deseamos conservar esta propiedad. Sin embargo, como veremos, no es de modo alguno sencillo generalizar las condiciones del Teorema A, al caso fraccionario, de tal manera que todo se puede mantener intacto fácilmente. Es un error clásico que se comete, muy a menudo, que las propiedades conocidas del cálculo clásico se generalizan a la configuración fraccionaria directamente y sin la suficiente precaución necesaria.

Para continuar con la parte Tercera: para $n \notin \mathbb{N}$. Como ya se mencionó, hay varias generalizaciones posibles. Sólo discutimos aquellos que tienen mayor importancia para aplicaciones prácticas.

Una investigación exhaustiva de muchas otras posibilidades, sin embargo, excluye el caso prácticamente muy importante del operador Caputo que se va a presentar posteriormente, la cual está contenida en la monografía enciclopédica de Samko, Kilbas y Marichev, y en el libro más reciente de Kilbas, Srivastava y Trujillo.

Siguiendo el esquema dado anteriormente, comenzamos con el operador integral . En el caso $n \in \mathbb{N}$, es bien conocido y fácilmente probado por inducción, que podemos reemplazar la definición recursiva de la Definición 1.1 Tercero por la siguiente fórmula explícita.

Lema 1.1

Sea f una función Riemann integrable en $[a, b]$. Entonces, para $a \leq x \leq b$ y $n \in \mathbb{N}$, tenemos:

$$J_a^n f(x) = \frac{1}{(n-1)!} \int_a^x (x-t)^{n-1} f(t) dt$$

Además, es una consecuencia inmediata de (1.1) y por lo tanto una consecuencia del teorema fundamental del cálculo, cumpliéndose la siguiente relación para los operadores D y J_a .

Lema 1.2

Sea $m, n \in \mathbb{N}$ tal que $m > n$, y sea f una función que tiene una n -ésima derivada continua en el intervalo $[a, b]$. Entonces,

$$D^n f = D^m J_a^{m-n} f$$

Prueba

Por (1.1), tenemos $f = D^{m-n} J_a^{m-n} f$. Aplicando el operador D^n a ambos lados de esta relación $D^n f = D^n D^{m-n} J_a^{m-n} f$ y usando el hecho de que $D^n D^{m-n} = D^m$. Estos dos lemas son fundamentales para las generalizaciones que surgen en adelante. Mirando al Lema 1.1 sería importante y útil generalizar los argumentos de la factorial a los no enteros. Tal generalización existe y es bien conocida llamada función Gamma de Euler.

Función Gamma de Euler

La definición es la siguiente

La función $\Gamma: (0, \infty) \rightarrow \mathbb{R}$, definida por

$$\Gamma(x) := \int_0^\infty t^{x-1} e^{-t} dt$$

se llama función Gamma de Euler o integral de Euler de segunda clase.

Para dar un tratamiento razonablemente autónomo del tema, indicamos algunas de las propiedades claves de la función Gamma en el Apéndice D.1. El más importante, para nuestros propósitos, es el siguiente, cuya prueba también se da en Apéndice D.1

Teorema 1.3

Para $n \in \mathbb{N}$, tenemos $\Gamma(n) = (n-1)!$.

Antes de comenzar el trabajo principal, introduciremos algunos espacios funcionales en donde vamos a discutir asuntos. Dado que estos son espacios clásicos, podemos ser más bien breves.

Definición 1.2

Sean

$$0 < \mu \leq 1, k \in \mathbb{N}_0 \text{ y } 1 \leq p. \mathbb{N}_0 = \mathbb{N} \cup \{0\},$$

$$L_p[a, b] := \{f: [a, b] \rightarrow \mathbb{R}; f \text{ es medible en } [a, b] \text{ y la } \int_0^\infty |f(x)|^p dx < \infty\},$$

$$L_\infty[a, b] := \{f: [a, b] \rightarrow \mathbb{R}; f \text{ es medible y esencialmente acotada en } [a, b]\}$$

$$H_\mu[a, b] := \{f: [a, b] \rightarrow \mathbb{R}; \exists c > 0 \forall x, y \in [a, b] : |f(x) - f(y)| \leq c|x - y|^\mu\},$$

$$C^k[a, b] := \{f: [a, b] \rightarrow \mathbb{R}; f \text{ tiene una } k.\text{ésima derivada continua}\},$$

$$C[a, b] := C^0[a, b],$$

$$H_0[a, b] := C[a, b],$$

En otras palabras, L_p es para $1 \leq p \leq \infty$ el espacio de Lebesgue habitual, mientras que $H_\mu[a, b]$ es un espacio de Holder o espacio de Lipschitz de orden μ .

Cuando la pregunta es clara por el contexto, el intervalo $[a, b]$, a menudo optaremos por no mencionarla explícitamente y usaremos la notación más simple L_p en lugar de $L_p[a, b]$, etc. De vez en cuando también usaremos un espacio de funciones un poco menos estándar.

Definición 1.3

Por $H^* \circ H^*[a, b]$ denotamos el conjunto de funciones $f: [a, b] \rightarrow \mathbb{R}$ con la propiedad de que existe alguna constante $L > 0$ tal que

$$|f(x+h) - f(x)| \leq L|h| \ln|h|^{-1}$$

siempre que ocurra $|h| < 1/2$ y $x, x+h \in [a,b]$,

Obviamente, este conjunto es un poco más grande que H_1 . Cuando se trabaja en un espacio de Lebesgue en lugar de un espacio continuo de funciones, todavía podemos retener la parte principal de la declaración del teorema fundamental.

Teorema B. Fundamental en Espacios de Lebesgue “TFEL”

Sea $f \in L_p[a,b]$. Entonces, es derivable casi en todas partes en $[a, b]$ y $DJ_a f = f$ también se cumple casi en todas partes en $[a, b]$.

Ocasionalmente también usaremos el siguiente conjunto de funciones.

Definición 1.4

Por A^n o $A^n[a,b]$ denotamos el conjunto de funciones con $(n-1)$ derivadas absolutamente continuas, es decir, las funciones f , para las que existe casi en todas partes una función $g \in L_1[a,b]$ tal que

$$f^{n-1}(x) = f^{n-1}(a) + \int_a^x g(t)dt$$

En este caso, llamamos a g la n -ésima derivada, generalizada de f , y simplemente escribimos $g = f^{(n)}$.

Una nota de precaución está en orden aquí. Es claro a partir de esta definición que una función posee, casi en todas partes, una derivada $f \in A^1$. Sin embargo, esta implicación no se puede revertir en general. Existen, por ejemplo, funciones no constantes

f que son diferenciables en casi todas partes, con $f' \equiv 0$, o que seguramente es una función L^1 . Obviamente, en tal caso f no puede representarse como la primitiva de f' como se requiere en la Definición, y por lo tanto tal función f no está en A^1 .

Ejemplo de aplicación del cálculo fraccionario

Antes de llegar a un estudio detallado de las propiedades matemáticas de operadores fraccionarios diferenciales y ecuaciones diferenciales fraccionarias, echemos un breve vistazo a un ejemplo simple, pero no poco realista de un modelo que surge en mecánica donde las derivadas fraccionarias se pueden utilizar con éxito.

El modelo ha sido originalmente propuesto como una teoría de Nutting (1943); los trabajos de Scott Blair et al. estaban entre los primeros en confirmar su valor en la práctica.

Específicamente, queremos describir el comportamiento de ciertos materiales bajo la influencia de fuerzas externas. La forma tradicional de abordar estas cuestiones en mecánica utiliza las leyes de Hooke y Newton.

La relación que nos interesa es la relación entre la tensión $\sigma(t)$ y la deformación $\varepsilon(t)$, las cuales se toman como funciones del tiempo t . Si estamos tratando con líquidos viscosos, entonces la ley de Newton

$$\sigma(t) = \eta D^1 \varepsilon(t) \dots (1.2)$$

Como observamos tenemos:

$\sigma(t)$ = la tensión,

$\varepsilon(t)$ = la deformación.

η = es la viscoelasticidad, constante.

es la herramienta, modelo de nuestra elección. Aquí la constante del material η es la llamada viscosidad del material. Ley de Hooke

$$\sigma(t) = E D^0 \varepsilon(t) \dots (1.3)$$

por otro lado, es la forma correcta de modelar la relación tensión-deformación para sólidos elásticos. La constante E se conoce como constante de elasticidad del material.

Por supuesto, es una práctica común no mencionar el operador D^0 es decir, el operador identidad en (1.3) explícitamente, pero nos hemos desviado de este camino para enfatizar la forma similitud entre las dos leyes.

Ahora considere un experimento en el que la deformación se manipula de manera controlada tal que, digamos, $\varepsilon(t) = t$ para $t \in [0, T]$ con algo de $T > 0$. Entonces se sigue que la tensión se comporta como

$$\sigma(t) = E t$$

en el caso de un sólido elástico y

$$\sigma(t) = \eta = \text{constante}$$

para un líquido viscoso. Podemos resumir estas ecuaciones en la forma

$$\Psi_k = \frac{\sigma(t)}{\varepsilon(t)} t^k \dots (1.4)$$

Donde $\Psi_0 = E$ y $\Psi_1 = \eta$. Evidentemente el caso $k = 0$ corresponde a la ley de Hooke para sólidos y $k = 1$ se refiere a la ley de Newton para líquidos.

En la práctica, no es raro encontrar los llamados materiales viscoelásticos que exhiben un comportamiento en algún lugar entre el líquido viscoso puro y el sólido elástico puro, es decir, donde se observaría una relación de la forma (1.4) con $0 < k < 1$. En este caso es apropiado interpretar k como una segunda constante material además de Ψ_k .

Los ejemplos clásicos son los polímeros, además algunos tipos de tejido biológico también pueden compartir esta propiedad con una serie de metales como, por ejemplo: el aluminio, bajo ciertas condiciones de temperatura y presión. Cabe señalar que para el caso de una deformación constante ε , la tensión en dicho material se desarrollaría de acuerdo con la formula

$$\sigma(t) = \text{constante } t^{-k}$$

y por lo tanto converge a cero para tiempos de observación muy largos. A este respecto, una vez nuevamente se encuentra entre un líquido viscoso para el cual σ se desvanece de manera idéntica y un elástico sólido cuya tensión σ es una constante distinta de cero.

En vista de todas estas propiedades de “interpolación”, es natural suponer que es también posible modelar la relación entre la tensión y la deformación para tal material viscoelástico a través de una ecuación de la forma

$$\sigma(t) = \nu D^k \varepsilon(t) \dots (1.5)$$

donde ν es una constante del material y $k \in (0,1)$ es el parámetro introducido anteriormente.

Esta ecuación “interpola” entre (1.2) y (1.3) con un espíritu similar. En vista de los fundamentos teóricos antes mencionados de (1.5) establecidos por Nutting, esta relación es frecuentemente llamada ley de Nutting.

Veremos en los capítulos siguientes que, en efecto, está justificado argumentar de esta manera si definimos correctamente el operador diferencial . En particular resultará que todas las relaciones mencionadas anteriormente se pueden mantener intactas.

02

Capítulo

***OPERADORES INTEGRAL Y
DIFERENCIAL DE RIEMANN-
LIOUVILLE***

Escucha a tu niño interno que te enseñara a ser siempre el mismo en los dolores agudos, en la pérdida de tus hijos, en las largas enfermedades y en él mismo; como un vivo ejemplo.

Burseca

Ahora estamos en posición de dar una primera definición para Operador integral fraccionaria y para el operador Fraccionaria diferencial, J_a^n y D^n , $n \notin \mathbb{N}$. Como se indicó anteriormente, comenzamos con el Operador integral fraccionaria.

Operador integral de Riemann-Liouville $J_{a,0}^n {}^{RL}J_x^k$

En vista de las consideraciones anteriores, el siguiente concepto parece de manera natural.

Definición 2.1

Sea $n \in \mathbb{R}_+$. El operador $J_{a,0}^n$ definido en $L_1[a, b]$ por

$$J_a^n f(x) := \frac{1}{\Gamma(n)} \int_a^x (x-t)^{n-1} f(t) dt$$

para $a \leq x \leq b$, se llama operador integral fraccionario de Riemann-Liouville de orden n .

Primero: para $n = 0$, establecemos $J_a^n := J$, es el operador de identidad. La definición de $n = 0$ es bastante conveniente para futuras manipulaciones. Es evidente que la integral fraccionaria de Riemann-Liouville coincide con la definición clásica de J_a^n en el caso $n \in \mathbb{N}$, excepto por el hecho de que hemos extendido el dominio desde Funciones integrables de Riemann a funciones inte-

grables de Lebesgue que nos conducirán a cualquier problema en nuestro desarrollo.

Segundo: además, en el caso de $n \geq 1$ es obvio que la integral $J_a^n f(x)$ existe para todo $x \in [a, b]$ porque el integrando es el producto de una función integrable f y la función continua $(x-t)^{n-1}$.

Tercero: en el caso $0 < n < 1$ sin embargo, la situación es menos clara a primera vista. Pero el siguiente resultado afirma que esta definición está justificada.

Teorema 2.1

Sea $f \in L_1[a, b]$ y $n > 0$. Entonces, la integral $J_a^n f(x)$ existe para casi todo $x \in [a, b]$. Además, la función $J_a^n f(x)$ en sí, también es un elemento de $L_1[a, b]$.

Prueba

Escribimos la integral en cuestión como:

$$\frac{1}{\Gamma(n)} \int_a^x (x-t)^{n-1} f(t) dt = \int_{-\infty}^{\infty} \phi_1(x-t) \phi_2(t) dt$$

$$\phi_1(u) = \begin{cases} u^{n-1}, & \text{para } 0 < u \leq b-a \\ 0 & \text{en otro caso} \end{cases}$$

$$\phi_2(u) = \begin{cases} f(u), & \text{para } a \leq u \leq b \\ 0 & \text{en otro caso} \end{cases}$$

Por construcción, $\phi_j \in L_1(\mathbb{R})$ para $j \in \{1, 2\}$, y por lo tanto por un resultado clásico en integración de Lebesgue.

Nuestra propiedad conserva una propiedad importante de los operadores integrales de orden entero para su generalización.

Teorema 2.2

Sean $m, n \geq 0$ y $\phi \in L_1[a, b]$ Entonces, $J_a^m J_a^n \phi = J_a^{m+n} \phi$

se mantiene en casi todas partes de . Si además $\phi \in C[a, b]$ o $m + n \geq 1$, entonces la identidad se mantiene en todas partes en $[a, b]$.

Teorema 2.3

Bajo los supuestos del Teorema 2.2,

$$J_a^m J_a^n \phi = J_a^n J_a^m \phi$$

Hay una manera algebraica de enunciar este resultado.

Teorema 2.4

Los operadores $\{J_a^n: L_1[a, b] \rightarrow L_1[a, b]; n \geq 0\}$ forman un semigrupo conmutativo con respecto a las operaciones consideradas. El operador de identidad es el elemento neutro de este semigrupo.

Demostración (del Teorema 2.2).

Tenemos

$$J_a^m J_a^n \phi(x) = \frac{1}{\Gamma(m)} \int_a^x (x-t)^{m-1} \frac{1}{\Gamma(n)} \int_a^t \phi(\tau) d\tau dt$$

En vista del Teorema 2.1, las integrales existen, y por el teorema de Fubini podemos intercambiar el orden de integración, obteniendo,

La sustitución $t = \tau + s(x - \tau)$ da como resultado

$$\begin{aligned}
J_a^m J_a^n \phi(x) &= \frac{1}{\Gamma(m)} \frac{1}{\Gamma(n)} \int_a^x \phi(\tau) \int_0^1 [(x-\tau)(1-s)]^{m-1} \times [s(x-\tau)]^{n-1} ds d\tau \\
&= \frac{1}{\Gamma(m)} \frac{1}{\Gamma(n)} \int_a^x \phi(\tau) (x-\tau)^{m+n-1} \int_0^1 (1-s)^{m-1} s^{n-1} ds d\tau
\end{aligned}$$

En vista del Teorema D.6, que dice,

Sea $\alpha, \beta \in R_+$. Entonces

$$\int_0^1 t^{\alpha-1} (1-t)^{\beta-1} dt = \frac{\Gamma(\alpha) \Gamma(\beta)}{\Gamma(\alpha + \beta)}$$

$$\int_0^x t^{\alpha-1} (x-t)^{\beta-1} dt = x^{\alpha+\beta-1} \frac{\Gamma(\alpha) \Gamma(\beta)}{\Gamma(\alpha + \beta)}$$

La integral de la primera ecuación del teorema D.6 se conoce como integral de Euler del primer tipo o función Beta de Euler $B(\alpha, \beta)$

Se puede encontrar más información sobre la función Gamma, por ejemplo, en la obra clásica de Artin (Hewitt et al., 1964), o en las obras de referencia habituales sobre funciones especiales.

Considerando nuestro caso,

$$\int_0^1 (1-s)^{m-1} s^{n-1} ds = \frac{\Gamma(m) \Gamma(n)}{\Gamma(n+m)} \text{ y por lo tanto}$$

$$J_a^m J_a^n \phi(x) = \frac{1}{\Gamma(n+m)} \int_a^x \phi(\tau) (x-\tau)^{m+n-1} d\tau = J_a^{m+n} \phi(x)$$

casi en todas partes en $[a, b]$.

Además, por los teoremas clásicos sobre integrales de parámetros, si $\phi \in C[a, b]$ entonces también $J_a^n \phi \in C[a, b]$ y por lo tanto $J_a^m J_a^n \phi \in C[a, b]$ y $J_a^{m+n} \phi \in C[a, b]$ también. Por lo tanto, dado que estas dos funciones continuas coinciden en casi todas partes, deben coincidir en todas partes.

Finalmente, si $\phi \in L_1[a, b]$ y $m+n \geq 1$ tenemos, por el resultado anterior

$$J_a^m J_a^n \phi(x) = J_a^{m+n} \phi(x) = J_a^{m+n-1} J_a^1 \phi(x)$$

Casi en cualquier parte. Desde $J_a^1 \phi$ es continua, también tenemos que $J_a^{m+n} \phi = J_a^{m+n-1} J_a^1 \phi$ es continua, y una vez más podemos concluir que las dos funciones en cualquiera lado de la igualdad casi en todas partes son continuos; por lo tanto, deben ser idénticos en todas partes.

Ahora consideramos algunas propiedades de mapeo del operador J_a^n . Veremos que la integración fraccionaria mejora las propiedades de suavidad de las funciones.

Para ser un poco más precisos, podemos decir que es la suma de dos expresiones uno de los cuales, denotado por Φ en el teorema a continuación, generalmente se comporta mejor que si misma, mientras que el otro puede no ser liso en el punto a , con un comportamiento conocido ahí y es una función en otros lugares.

Teorema 2.5

Sea para algún $\mu \in [0, 1]$, y sea $0 < n < 1$. Entonces

$$J_a^n \phi(x) = \frac{\phi(a)}{\Gamma(n+1)} (x-a)^n + \Phi(x)$$

con alguna función Φ . Esta función Φ satisface

$$\Phi(x) = O((x-a)^{\mu+n})$$

como $x \rightarrow a$. Es más

$$\Phi \in \begin{cases} H_{\mu+n}[a, b], & \text{si } \mu + n < 1 \\ H^*[a, b], & \text{si } \mu + n = 1 \\ H_1[a, b], & \text{si } \mu + n > 1 \end{cases}$$

Observación 2.1.

Es posible mostrar una declaración aún más fuerte en el caso $\mu + n > 1$:

Bajo este supuesto tenemos que $\phi \in C^1[a, b]$ y $\phi' \in H_{\mu+n-1}[a, b]$, cf. [167, Teorema 3.1].

Demostración (del Teorema 2.5).

Tenemos

$$J_a^n \phi(x) = \frac{\phi(a)}{\Gamma(n)} \int_a^x (x-t)^{n-1} + \frac{1}{\Gamma(n)} \int_a^x \frac{\phi(t) - \phi(a)}{(x-t)^{1-n}} dt$$

Esto produce la representación deseada con

$$\phi(x) = \frac{1}{\Gamma(n)} \int_a^x \frac{\phi(t) - \phi(a)}{(x-t)^{1-n}} dt$$

En vista de $\phi \in H_\mu$

$$\begin{aligned} |\phi(x)| &\leq \frac{1}{\Gamma(n)} \int_a^x \frac{L|t-a|^\mu}{(x-t)^{1-n}} dt = \frac{L}{\Gamma(n)} \int_a^x (t-a)^\mu (x-t)^{n-1} dt \\ &= \frac{L}{\Gamma(n)} (x-a)^{\mu+n} \int_a^x s^\mu (1-s)^{n-1} dt = O((x-a)^{\mu+n}) \end{aligned}$$

Ahora establecemos $g(x) := \frac{\phi(x) - \phi(a)}{\Gamma(n)}$

Además, sea $h > 0$, y $x; x+h \in [a, b]$.

Después,

$$\begin{aligned}
\Phi(x+h) - \Phi(x) &= \int_a^{x+h} g(t)(x+h-t)^{n-1} dt - \int_a^{x+h} g(t)(x-t)^{n-1} dt \\
&= \int_a^x g(t)[(x+h-t)^{n-1} - (x-t)^{n-1}] dt + \int_a^{x+h} g(t)(x+h-t)^{n-1} dt \\
&= \int_a^x (g(t) - g(x))[(x+h-t)^{n-1} - (x-t)^{n-1}] dt \\
&\quad + \int_x^{x+h} (g(t) - g(x))(x+h-t)^{n-1} dt + k_1 + k_2 + K_3
\end{aligned}$$

Donde contiene los términos restantes. Un cálculo explícito muestra que

$$K_3 = g(x) \left(\int_a^x [(x+h-t)^{n-1} - (x-t)^{n-1}] dt + \int_a^{x+h} (x+h-t)^{n-1} dt \right)$$

Estimamos los términos K_1, K_2 y K_3 por separado. En vista de nuestra suposición, es claro que también, por lo tanto

$$\begin{aligned}
|K_1| &= \left| \int_0^{x-a} (g(x-\mu) - g(x))[(\mu+h)^{n-1} - \mu^{n-1}] d\mu \right| \\
&\leq L \int_0^{x-a} \mu^\mu [\mu^{n-1} - (\mu+h)^{n-1}] d\mu \\
&= Lh \int_0^{\frac{(x-a)}{h}} (ht)^\mu [(ht)^{n-1} - (ht+h)^{n-1}] dt \\
&= Lh^{\mu+n} \int_0^{\frac{(x-a)}{h}} t^\mu [t^{n-1} - (t+1)^{n-1}] dt
\end{aligned}$$

En $t \rightarrow 0$, no hay problema con la convergencia de la integral ya que el integrando se comporta como $t^{\mu+n-1}$ ahí (el exponente es estrictamente mayor que -1). En el caso $x-a < h$, la integral está acotada por $\int_0^1 t^{\mu+n-1} dt = \frac{1}{(\mu+n)}$ y por lo tanto $K_1 = O(h^{\mu+n})$ en este caso. Si $x-a \geq h$, encontramos,

$$\begin{aligned}
& \int_0^{\frac{(x-a)}{h}} t^\mu [t^{n-1} - (t+1)^{n-1}] dt \\
&= \int_0^1 t^\mu [t^{n-1} - (t+1)^{n-1}] dt + \int_1^{\frac{(x-a)}{h}} t^\mu [t^{n-1} - (t+1)^{n-1}] dt \\
&< \frac{1}{\mu+n} + (1-n) \int_0^{\frac{(x-a)}{h}} t^{\mu+n-2} dt
\end{aligned}$$

en vista del teorema del valor medio del cálculo diferencial. La integral restante se puede calcular fácilmente de forma explícita. Encontramos para $\mu+n < 1$ que

$$\begin{aligned}
|K_1| &\leq O(h^{\mu+n}) \left(\frac{1}{\mu+n} + (1-n) \int_0^{\frac{(x-a)}{h}} t^{\mu+n-2} dt \right) \\
&\leq O(h^{\mu+n}) \left(\frac{1}{\mu+n} + (1-n) \int_0^\infty t^{\mu+n-2} dt \right) \leq O(h^{\mu+n})
\end{aligned}$$

debido a que $\mu+n < 1$ y $x-a \geq h$. Para $\mu+n = 1$ un cálculo similar tenemos,

$$|K_1| \leq O(h^{\mu+n}) \left(1 + \int_0^{\frac{(x-a)}{h}} t^{-1} dt \right) \leq O(h \ln h^{-1})$$

Finalmente, para $\mu+n > 1$ tenemos, ya que $x \leq b$,

$$|K_1| \leq O(h^{\mu+n}) \left(\int_0^{\frac{(b-a)}{h}} t^{\mu+n-2} dt \right) = O(h^{\mu+n}) \left(\frac{b-a}{h} \right)^{\mu+n-1} = O(h)$$

Por lo tanto,

$$K_1 = \begin{cases} O(h^{\mu+n}), & \text{para } \mu+n < 1 \\ O(h \ln h^{-1}), & \text{para } \mu+n = 1 \\ O(h), & \text{para } \mu+n > 1 \end{cases}$$

A continuación, estimamos K_2 . Nuevamente usamos el hecho de que $g \in H_\mu$ para derivar (usando la sustitución $s = \frac{(t-x)}{h}$).

$$|K_1| \leq L \int_x^{x+h} (t-x)^\mu (x+h-t)^{n-1} dt = L h^{\mu+n} \int_1^1 s^\mu (1-s)^{n-1} ds = O(h^{\mu+n})$$

independientemente de μ y n . Obviamente, este límite es más fuerte que el límite anterior para K_1 .

Finalmente, para K_3 usamos la suposición de Hölder sobre que implica que $|g(x)| \leq L(x-a)^\mu$ para alguna constante L . Evaluando las integrales analíticamente, encontramos

$$|K_3| \leq \frac{L}{n} (x-a)^\mu [(x-a+h)^n - (x-a)^n]$$

En el caso $x-a \leq h$, esto está acotado superiormente por $O(h^{\mu+n})$.

Si $x-a > h$, tenemos que estimar el término entre paréntesis por el teorema del valor medio del cálculo diferencial y encontrar (teniendo en cuenta que $n < 1$)

$$K_3 = O(1)(x-a)^\mu h(x-a)^{n-1} = O(h)(x-a)^{\mu+n-1}$$

Una vez más observamos los tres casos por separado y encontramos

$$K_3 = O(h)$$

$$\text{para } \mu + n = 1, |K_3| \leq O(h)h^{\mu+n-1} = O(h^{\mu+n})$$

$$\text{para } \mu + n < 1, \text{ y } |K_3| \leq O(h)(b-a)^{\mu+n-1} = O(h)$$

para $\mu + n > 1$. Nuevamente, estas estimaciones son más sólidas que las que obtuvimos para K_1 .

Combinando todas las estimaciones, obtenemos

$$\Phi(x+h) - \Phi(x) = \begin{cases} O(h^{\mu+n}), & \text{para } \mu + n < 1 \\ O(h \ln h^{-1}), & \text{para } \mu + n = 1 \\ O(h), & \text{para } \mu + n > 1 \end{cases}$$

Se puede obtener un resultado similar cuando asumimos que el integrando está en una clase de Lebesgue adecuada.

Teorema 2.6

Sean $n > 0, p > \max\{1, \frac{1}{n}\}$ y $\Phi \in L_p[a, b]$. Entonces,

$$J_a^n \Phi(x) = O((x-a)^{n-1/p})$$

Como $x \rightarrow a_+$.

Si adicionalmente se tiene que $n - 1/p \notin \mathbb{N}$, entonces $J_a^n \Phi(x) \in C^{\lfloor n - \frac{1}{p} \rfloor}[a, b]$ y

$$D^{\lfloor n - \frac{1}{p} \rfloor} J_a^n \Phi \in H_{n - \frac{1}{p} - \lfloor n - 1 \rfloor p}[a, b]$$

Observación 2.2

En el caso de $n - 1/p \in \mathbb{N}$, se puede mostrar un enunciado ligeramente modificado.

La condición de Hölder sobre la derivada de $J_a^n \Phi$, pero se puede ser reemplazado por una condición similar a la que define el conjunto H^* .

Demostración (del Teorema 2.6)

Para un p dada, introducimos el exponente conjugado $q \in [1, \infty)$ por la relación $p^{-1} + q^{-1} = 1$. Entonces, por definición del operador integración Riemann-Liouville y desigualdad de Hölder,

$$\begin{aligned} |J_a^n \Phi(x)| &\leq \frac{1}{\Gamma(n)} \left(\int_a^x |\Phi(t)|^p dt \right)^{1/p} \left(\int_a^x |x-t|^{(n-1)q} dt \right)^{1/q} \\ &\leq \frac{(x-a)^{n-\frac{1}{p}}}{\Gamma(n)((n-1)q+1)^{\frac{1}{q}}} \left(\int_a^x |\Phi(t)|^p dt \right)^{\frac{1}{p}} = O((x-a)^{n-1/p}) \end{aligned}$$

Para la prueba del resultado de la continuidad, discutimos primero el caso $n-1/p < 1$. Aquí, encontramos eso

$$\begin{aligned}
 J_a^n \Phi(x+h) - J_a^n \Phi(x) &= \frac{1}{\Gamma(n)} \underbrace{\int_x^{x+h} (x+h-t)^{n-1} \Phi(t) dt}_{=: K_1} \\
 &+ \frac{1}{\Gamma(n)} \underbrace{\int_a^x [(x+h-t)^{n-1} - (x-t)^{n-1}] \Phi(t) dt}_{=: K_2}
 \end{aligned}$$

Como arriba, usamos la desigualdad de Hölder para derivar que

$$|K_1| \leq \left(\int_x^{x+h} |\Phi(t)|^p dt \right)^{1/p} \left(\int_x^{x+h} (x+h-t)^{(n-1)q} dt \right)^{1/q} \leq ch^{n-1/p}$$

con alguna constante c . Además, también por la desigualdad de Hölder,

$$\begin{aligned}
 |K_1| &\leq \|\Phi\|_{L_p} \left(\int_a^x |(x+h-t)^{(n-1)} - (x-t)^{n-1}|^q dt \right)^{1/q} \\
 &\leq \|\Phi\|_{L_p} h^{n-1/p} \left(\int_0^{(x-a)/h} |s^{n-1} - (s+1)^{n-1}|^q ds \right)^{1/q}
 \end{aligned}$$

En el caso $x-a \leq h$ la última integral está acotada por una constante a . En el caso complementario $x-a > h$, usamos (como en la demostración del teorema anterior) el teorema valor medio del cálculo diferencial y hallamos.

$$\begin{aligned}
 &\int_0^{(x-a)/h} |s^{n-1} - (s+1)^{n-1}|^q ds \\
 &= \int_0^1 |s^{n-1} - (s+1)^{n-1}|^q ds + \int_1^{(x-a)/h} |s^{n-1} - (s+1)^{n-1}|^q ds \\
 &\leq C + |n-1| \int_1^{(x-a)/h} s^{(n-1)q} ds \\
 &= C + \frac{|n-1|}{q(n-2)+1} \{s^{q(n-2)+1}\}_1^{(x-a)/h}
 \end{aligned}$$

con alguna constante C . Miramos el denominador en la última expresión y encontramos que

$$q(n-2)+1 = q\left(n-1-1+\frac{1}{q}\right) = q\left(n-1-\frac{1}{q}\right) < 0$$

en vista de $q > 0$ y nuestra suposición de que $n-1/p < 1$. Así podemos continuar con la estimación para la integral por

$$\begin{aligned} & \int_0^{(x-a)/h} |S^{n-1} - (S+1)^{n-1}|^q ds \leq \\ & \leq c + \left| \frac{n-1}{q(n-2)+1} \right| \left(1 - \left(\frac{x-a}{h} \right)^{q(n-2)+1} \right) = o(1) \end{aligned}$$

Por eso:

$$K_2 = o(h^{n-1/p})$$

también, y por lo tanto en este caso $J_a^n \phi \in H_{n-1/p}$.

Ahora discutimos el caso restante $n-1/p > 1$. Entonces, en particular, $n > 1$, y en vista de la propiedad de semigrupo de integración fraccionaria podemos escribir

$$J_a^n \phi = J_a^{\lfloor n-\frac{1}{p} \rfloor} J_a^{n-\lfloor n-\frac{1}{p} \rfloor} \phi$$

Por lo tanto, haciendo uso del teorema fundamental del cálculo tenemos,

$$D^{\lfloor n-\frac{1}{p} \rfloor} J_a^n \phi = J_a^{n-\lfloor n-\frac{1}{p} \rfloor} \phi = J_a^{\hat{n}} \phi$$

Aquí tenemos que $\hat{n} = n - \lfloor n - \frac{1}{p} \rfloor \geq n - |n| \geq 0$ y $\hat{n} - \frac{1}{p} = n - \frac{1}{p} - \lfloor n - \frac{1}{p} \rfloor < 1$ porque, por suposición, $n - \frac{1}{p} \notin \mathbb{N}$. Por lo tanto, podemos aplicar el resultado que acabamos de probar, con n reemplazada por $\hat{n}$, y encontrar que

$$D^{\lfloor n-\frac{1}{p} \rfloor} J_a^n \phi = J_a^{\hat{n}} \phi \in H_{\hat{n}-\frac{1}{p}-\lfloor \hat{n}-\frac{1}{p} \rfloor} = H_{n-\frac{1}{p}-\lfloor n-\frac{1}{p} \rfloor}$$

Los índices de los espacios de Hölder son idénticos aquí porque la diferencia es un número entero.

Ejemplo

Sea $f(x) = (x - a)^\beta$ para algún $\beta > -1$ y $n > 0$, entonces,

$$J_a^n f(x) = \frac{\Gamma(\beta + 1)}{\Gamma(n + \beta + 1)} (x - a)^{n+\beta}$$

En vista del conocido resultado correspondiente en el caso $n \in \mathbb{N}$, este resultado es precisamente lo que uno esperaría de una generalización sensata del operador integral En vista del Teorema D.6, la derivación es directa:

$$\begin{aligned} J_a^n f(x) &= \frac{1}{\Gamma(n)} \int_a^x (t - a)^\beta (x - t)^{n-1} dt \\ &= \frac{1}{\Gamma(n)} (x - a)^{n+\beta} \int_a^x s^\beta (1 - s)^{n-1} ds = \frac{\Gamma(\beta + 1)}{\Gamma(n + \beta + 1)} (x - a)^{n+\beta} \end{aligned}$$

A continuación, discutimos el intercambio de operación límite e integración fraccionaria.

Teorema 2.7

Sea $n > 0$. Suponga que $(f_k)_{k=1}^\infty$ es una secuencia uniformemente convergente de funciones continuas en $[a, b]$. Entonces podemos intercambiar el operador integral fraccionario y el proceso límite, es decir,

$$\left(J_a^n \lim_{k \rightarrow \infty} f_k \right)(x) = \left(\lim_{k \rightarrow \infty} J_a^n f_k \right)(x)$$

En particular, la secuencia de funciones $(J_a^n f_k)_{k=1}^\infty$ es uniformemente convergente.

Prueba

Denotamos el límite de la sucesión (f_k) por f . Es bien sabido que f es continuo. Luego encontramos

$$\begin{aligned}
|J_a^n f_k(x) - J_a^n f(x)| &\leq \frac{1}{\Gamma(n)} \int_a^x |f_k(t) - f(t)| (x-t)^{n-1} dt \\
&\leq \frac{1}{\Gamma(n)} \|f_k - f\|_\infty \int_a^x (x-t)^{n-1} dt \\
&= \frac{1}{\Gamma(n+1)} \|f_k - f\|_\infty (x-a)^n \\
&\leq \frac{1}{\Gamma(n+1)} \|f_k - f\|_\infty (b-a)^n
\end{aligned}$$

que converge a cero cuando $k \rightarrow \infty$ uniformemente para todo $x \in [a, b]$.

Teorema 2.8

Sea f analítica en $(a-h, a+h)$ para algún $h > 0$, y sea $n > 0$.

Entonces

$$J_a^n f(x) = \sum_{k=0}^{\infty} \frac{(-1)^k (x-a)^{k+n}}{k! (n+k) \Gamma(n)} D^k f(x)$$

para $a \leq x < a+h/2$, y

$$J_a^n f(x) = \sum_{k=0}^{\infty} \frac{(x-a)^{k+n}}{\Gamma(k+1+n)} D^k f(x)$$

para $a \leq x < a+h$. En particular, $J_a^n f$ es analítica en $(a, a+h)$.

Prueba

Para la primera declaración, usamos la definición del operador integral de Riemann-Liouville J_a^n , a saber

$$J_a^n f(x) = \frac{1}{\Gamma(n)} \int_a^x f(t) (x-t)^{n-1} dt$$

y expandir $f(t)$ en una serie de potencias alrededor de x . Como $x \in [a, a+h/2]$, la serie de potencias converge en todo el intervalo de integración. Así, por el Teorema 2.7, es legal la suma e integración de cambios. Entonces usamos la representación explícita para la integral fraccionaria de la función de potencia que habíamos derivado en el Ejemplo 2.1 para encontrar el resultado final.

Para la segunda afirmación, procedemos de manera similar, pero ahora ampliamos la serie de potencias en a y no en x . Esto nos permite nuevamente concluir la convergencia de la serie en el intervalo requerido.

La analiticidad de $J_a^n f$ sigue inmediatamente de la segunda declaración.

Los enunciados de estos dos últimos resultados nos permiten mirar otro instructivo ejemplo.

Ejemplo

Sea $f(x) = \exp(\lambda x)$ con algún $\lambda > 0$. Calcule $J_0^n f(x)$ para $n > 0$.

En el caso $n \in \mathbb{N}$ obviamente tenemos $J_0^n f(x) = \lambda^n \exp(\lambda x)$. Sin embargo, este resultado no generaliza de manera directa al caso $n \notin \mathbb{N}$. Más bien, en vista del conocido desarrollo en serie de la función exponencial, Teorema 2.7 y Ejemplo 2.1, encontramos

$$\begin{aligned} J_0^n f(x) &= J_0^n \left[\sum_{k=0}^{\infty} \frac{(\lambda \cdot)^k}{k!} \right] (x) = \sum_{k=0}^{\infty} \frac{\lambda^k}{k!} J_0^n [(\cdot)^k] (x) \\ &= \sum_{k=0}^{\infty} \frac{\lambda^k}{\Gamma(k+n+1)} x^{k+n} = \lambda^{-n} \sum_{k=0}^{\infty} \frac{(\lambda x)^{k+n}}{\Gamma(k+n+1)} \end{aligned}$$

y aquí la serie del lado derecho no es $\exp(\lambda x)$.

Más adelante encontraremos una clase de funciones que podemos usar convenientemente para expresar la serie, pero por el momento solo notamos que los operadores integrales fracciona-

rios de tipo de Riemann-Liouville no reproduce funciones exponenciales de la misma manera que las integrales de orden enteros.

Por cierto, surge el mismo problema cuando calculamos integrales fraccionarias de otras funciones no polinómicas que tienen integrales de orden entero muy simples como como la función seno o coseno. Animamos al lector a trabajar en los detalles como un ejercicio.

En los dos últimos teoremas de esta sección discutimos otra propiedad importante de operadores integrales fraccionarios, a saber, la continuidad con respecto al orden del operador. Primero nos fijamos en el caso de que trabajamos en los espacios de Lebesgue. Bajo esto suponiendo que la situación es muy simple.

Teorema 2.9

Sea $1 \leq p < \infty$ y sea $(m_k)_{k=1}^\infty$ sea una secuencia convergente de números no negativos con límite m . Entonces, para todo $f \in L_p[a, b]$,

$$\lim_{k \rightarrow \infty} J_a^{m_k} f = J_a^m f$$

donde la convergencia es en el sentido de la norma $L_p[a, b]$.

Una prueba de este resultado puede encontrarse en [Teorema 2.6]

Sin embargo, cuando tratamos con el espacio $C[a, b]$ equipado con la norma de Chebyshev (que es mucho más interesante e importante para las aplicaciones que tenemos en mente), entonces la situación es más complicada. Para ilustrar esto, proporcionamos un ejemplo muy simple antes de llegar al teorema mismo que describirá el caso general.

Ejemplo

Sea $f(x)=1$ y considere un sistema estrictamente monótono y convergente secuencia $(m_k)_{k=1}^{\infty}$ de números no negativos con límite $m \geq 0$. Por el Ejemplo 2.1 tenemos encuentra eso

$$J_a^{m_k} f(x) = \frac{1}{\Gamma(m_k + 1)} (x - a)^{m_k}$$

Entonces tenemos que introducir una distinción de dos casos que exhiben tipos muy diferentes de comportamiento:

Primero. Si $m > 0$ entonces

$$\|J_a^{m_k} f - J_a^m f\| = \sup_{x \in [a, b]} \left| \frac{(x-a)^{m_k}}{\Gamma(m_k+1)} - \frac{(x-a)^m}{\Gamma(m+1)} \right| \rightarrow 0, \quad (2.1)$$

como $\rightarrow m$, que se puede mostrar con una estimación larga pero simple. Dejamos los detalles de este caso especial para el lector (el resultado, por supuesto, seguirá del Teorema 2.10 a continuación) y solo tenga en cuenta que tenemos convergencia en la norma de Chebyshev en este caso.

Segundo. Si $m = 0$ entonces la sucesión debe ser decreciente. Además, tenemos eso $J_a^{m_k} f(a) = 0$ para todo k , mientras que $J_a^{m_k} f(a) = J_a^0 f(a) = f(a) = 1$, es decir, ni siquiera tienen convergencia puntual, y mucho menos convergencia en la norma de Chebyshev (convergencia uniforme).

Una descripción completa de la situación es la siguiente.

Teorema 2.10

Sea $f \in C[a, b]$ y $m \geq 0$. Además, suponga que (m_k) es una sucesión de números positivos tales que $\lim_{k \rightarrow \infty} m_k = m$. Entonces, para todo $\varepsilon > 0$

$$\lim_{k \rightarrow \infty} \sup_{x \in [a+\varepsilon, b]} |J_a^{m_k} f(x) - J_a^m f(x)| = 0$$

Si además $m > 0$ $f(x) = O((x-a)^\delta)$ cuando $x \rightarrow a$ para algún $\delta > 0$ entonces

$$\lim_{k \rightarrow \infty} \|J_a^{m_k} f(x) - J_a^m f(x)\|_\infty = 0$$

Prueba

Comenzamos con el caso $m > 0$. Sin pérdida de generalidad podemos suponer que $m_k > \frac{m}{2}$ para todo k . Entonces se sigue para arbitrario un $x \in [a, b]$ que

$$\begin{aligned} |J_a^{m_k} f(x) - J_a^m f(x)| &\leq \int_a^x |f(t)| \left| \frac{(x-a)^{m_k}}{\Gamma(m_k+1)} - \frac{(x-a)^m}{\Gamma(m+1)} \right| dt \\ &\leq \|f\|_\infty \int_a^x (x-t)^{m/4} \left| \frac{(x-a)^{m_k-1-\frac{m}{4}}}{\Gamma(m_k)} - \frac{(x-a)^{-1+3m/4}}{\Gamma(m)} \right| dt \\ &= \|f\|_\infty \int_0^{x-a} u^{m/4} \left| \frac{u^{m_k-1-\frac{m}{4}}}{\Gamma(m_k)} - \frac{u^{-1+3m/4}}{\Gamma(m)} \right| du \\ &\leq \|f\|_\infty \int_0^{b-a} u^{m/4} \left| \frac{u^{m_k-1-\frac{m}{4}}}{\Gamma(m_k)} - \frac{u^{-1+3m/4}}{\Gamma(m)} \right| du \\ &\leq \|f\|_\infty (b-a)^{m/4} \int_0^{b-a} \left| \frac{u^{m_k-1-\frac{m}{4}}}{\Gamma(m_k)} - \frac{u^{-1+3m/4}}{\Gamma(m)} \right| du \end{aligned}$$

Es fácil ver que el integrando solo cambia su signo por

$$u = c_k := \left(\frac{\Gamma(m_k)}{\Gamma(m)} \right)^{\frac{1}{(m_k-m)}}$$

y un cálculo explícito usando la regla de L'Hospital revela que $\lim_{k \rightarrow \infty} c_k = \exp(-\Psi(m))$ donde $\Psi = \frac{\Gamma'}{\Gamma}$ denota la función Digamma. Por lo tanto, podemos continuar nuestra estimación anterior de acuerdo con,

$$\begin{aligned}
 |J_a^{m_k} f(x) - J_a^m f(x)| &\leq \|f\|_\infty (b-a)^{m/4} \int_0^{b-a} \left| \frac{u^{m_k-1-\frac{m}{4}}}{\Gamma(m_k)} - \frac{u^{-1+3m/4}}{\Gamma(m)} \right| du \\
 &\leq \|f\|_\infty (b-a)^{m/4} \left| \int_0^{c_k} \left| \frac{u^{m_k-1-\frac{m}{4}}}{\Gamma(m_k)} - \frac{u^{-1+3m/4}}{\Gamma(m)} \right| du \right. \\
 &\quad \left. - \int_{c_k}^{b-a} \left| \frac{u^{m_k-1-\frac{m}{4}}}{\Gamma(m_k)} - \frac{u^{-1+3m/4}}{\Gamma(m)} \right| du \right| \\
 &= \|f\|_\infty (b-a)^{m/4} \left| \frac{2 c_k^{\frac{m_k-m}{4}}}{\Gamma(m_k(m_k-\frac{m}{4}))} - \frac{8 c_k^{\frac{3m}{4}}}{3m \Gamma(m)} - \frac{2 (b-a)^{m_k-\frac{m}{4}}}{\Gamma(m_k(m_k-\frac{m}{4}))} \right. \\
 &\quad \left. + \frac{4(b-a)^{\frac{3m}{4}}}{3m \Gamma(m)} \right|
 \end{aligned}$$

y evidentemente el término dentro de la operación de valor absoluto converge a 0 cuando $m_k \rightarrow m$, lo que completa la demostración en el caso $m > 0$.

En el caso $m = 0$ este argumento no es aplicable. En cambio, procedemos de la siguiente manera.

Para algún $\varepsilon_*, \varepsilon_* \in [0, x-a]$ que se especificará con más precisión más adelante, escribimos,

$$\begin{aligned}
 |J_a^{m_k} f(x) - f(x)| &= \left| \int_a^x f(t) \frac{(x-t)^{m_k-1}}{\Gamma(m_k)} dt - f(x) \right| \\
 &\leq \left| \int_a^{x-\varepsilon_*} f(t) \frac{(x-t)^{m_k-1}}{\Gamma(m_k)} dt \right| + \left| \int_{x-\varepsilon_*}^x f(t) \frac{(x-t)^{m_k-1}}{\Gamma(m_k)} dt - f(x) \right| \\
 &\leq \frac{(\varepsilon_*)^{m_k-1} \|f\|_\infty (x-\varepsilon_*-1)}{\Gamma(m_k)} + \left| \int_0^{\varepsilon_*} f(x-u) \frac{u^{m_k-1}}{\Gamma(m_k)} du - f(x) \right|
 \end{aligned}$$

Usando el teorema generalizado del valor medio, encontramos que la integral se puede representar en la forma.

$$\int_0^{\varepsilon_*} f(x-u) \frac{u^{m_k-1}}{\Gamma(m_k)} du = f(\xi) \int_0^{\varepsilon_*} f(x-u) \frac{u^{m_k-1}}{\Gamma(m_k)} du = f(\xi) \frac{\varepsilon_*^{m_k}}{\Gamma(m_k+1)}$$

con algún $\xi \in [x-\varepsilon, x]$. De este modo,

$$|J_a^{m_k} f(x) - f(x)| \leq \frac{(\varepsilon_*)^{m_k-1} \|f\|_{\infty} (x-\varepsilon_*-1)}{\Gamma(m_k)} + \left| \frac{f(\xi) \varepsilon_*^{m_k}}{\Gamma(m_k+1)} - f(x) \right|, \dots (2.2)$$

para todo $x \in [a + \varepsilon, b]$, donde $\xi \in [x - \varepsilon, x]$. Ahora establecemos $\varepsilon_* := \frac{m_k}{2}$ por el valor ε mencionado en la reclamación (que se ha fijado de antemano), sabemos que (ya que) existe algún $k_0 \in \mathbb{N}$ tal que para todo $k \geq k_0$ tenemos $\varepsilon_* = m_k^{\frac{1}{2}} < \varepsilon$.

Así, para estos k la desigualdad (3.2) es válida para todo $x \in [a + \varepsilon, b]$. Con nuestra especial elección de ε toma la forma

$$\begin{aligned} |J_a^{m_k} f(x) - f(x)| &\leq \frac{m_k^{(m_k-1)/2}}{\Gamma(m_k+1)} (b-a) \|f\|_{\infty} + \left| \frac{f(\xi) m_k^{m_k/2}}{\Gamma(m_k+1)} - f(\xi) + f(\xi) - f(x) \right| \\ &\leq \frac{m_k^{\frac{m_k-1}{2}}}{\Gamma(m_k)} (b-a) \|f\|_{\infty} + \|f\|_{\infty} \left(\frac{m_k^{\frac{m_k}{2}}}{\Gamma(m_k+1)} - 1 \right) + \omega(f; x-\xi), \dots (2.3) \end{aligned}$$

Donde,

$$\omega(g; h) = \sup\{|g(y_1) - g(y_2)|, y_1, y_2 \in [a, b], |y_1 - y_2| \leq h\}$$

denota el módulo clásico de continuidad de la función $g: [a, b] \rightarrow \mathbb{R}$. Un explícito cálculo, utilizando la regla de L'Hospital, reproducimos

$$\frac{m_k^{\frac{(m_k-1)}{2}}}{\Gamma(m_k)} \rightarrow 0 \text{ y } \frac{m_k^{\frac{m_k}{2}}}{\Gamma(m_k+1)} - 1 \rightarrow 0$$

como $k \rightarrow \infty$. Además, nuestros resultados anteriores implican que $\xi \rightarrow x$ cuando $k \rightarrow \infty$, y dado que f es continua esto implica que $\omega(f; x-\xi) \rightarrow 0$ también. Así tenemos la convergencia uniforme a 0 en $[a+\varepsilon, b]$.

Finalmente tenemos que probar la convergencia uniforme en todo $[a, b]$ en el caso $m = 0$ si $f(x) = O((x-a)^\delta)$. Para ello se procede de forma muy similar al caso anterior y en particular todavía

encontramos el límite (3.3) para $x \geq a + \varepsilon_* = a + m_k^{\frac{1}{2}}$. En el caso $x = a$ tenemos trivialmente

$$|J_a^{m_k} f(x) - f(x)| = 0 - 0 = 0$$

y por lo tanto queda probar la convergencia uniforme para

$$x \in \left[a, a + m_k^{\frac{1}{2}} \right].$$

Desde

$$|f(x)| \leq C(x-a)^\delta \leq m_k^{\frac{\delta}{2}}$$

para este rango de x (donde C es una constante independiente de k) y

$$\begin{aligned} |J_a^{m_k} f(x)| &= \frac{1}{\Gamma(m_k)} \left| \int_a^x (x-t)^{m_k-1} f(t) dt \right| \\ &\leq \frac{C}{\Gamma(m_k)} \int_a^x (x-t)^{m_k-1} (t-a)^\delta dt \\ &= \frac{C}{\Gamma(m_k)} (x-t)^{m_k+\delta} \int_0^1 (1-s)^{m_k-1} s^\delta ds \\ &= C \frac{\Gamma(\delta+1)}{\Gamma(m_k+\delta+1)} (x-t)^{m_k+\delta} \\ &\leq C \frac{\Gamma(\delta+1)}{\Gamma(m_k+\delta+1)} m_k^{(m_k+\delta)/2} \end{aligned}$$

por la sustitución $s = (t-a)/(x-a)$ y las propiedades fundamentales de la función Beta, encontramos para $x \in [a, a+m_n^{\frac{1}{2}}]$

$$\begin{aligned} |J_a^{m_k} f(x) - f(x)| &\leq |J_a^{m_k} f(x) + f(x)| \\ &\leq C \frac{\Gamma(\delta+1)}{\Gamma(m_k+\delta+1)} m_k^{(m_k+\delta)/2} + C m_k^{\delta/2} \end{aligned}$$

Recordando que, por la regla de L'Hospital, $m_k^{\delta/2} \rightarrow 1$ y por lo tanto $m_k^{(m_k+\delta)/2} \rightarrow 0$ como $k \rightarrow \infty$, de hecho, obtenemos el resultado de convergencia uniforme requerido.

Una demostración alternativa de la segunda parte de este teorema se dará en Observación 6.11.

Derivadas de Riemann-Liouville

Habiendo establecido las propiedades fundamentales para los operadores fraccionarios de integración de Riemann-Liouville, ahora llegamos a los operadores diferenciales fraccionarios.

La motivación viene de la definición dada en el Teorema 2.2 que bajo ciertas condiciones establecemos la identidad

$$D^n f = D^m J_a^{m-n} f$$

donde m y n eran números enteros tales que $m > n$. Ahora suponga que n no es un número entero.

Entonces aún podemos elegir un entero m tal que $m > n$. En vista de la teoría desarrollada en la sección anterior, el lado derecho de la identidad sigue siendo importante.

Sin embargo, hay una gran diferencia con el caso clásico en el que tanto m como n son enteros: ahora encontramos que el operador obtenido de esta manera depende de la elección del punto a . Por tanto, llegamos a la siguiente definición si elegimos el valor del entero m sea lo más pequeño posible:

Definición 2.2

Sea $n \in \mathbb{R}_+$ y $m = [n]$, función llamado techo $\min\{z \in \mathbb{Z}, z \geq x\}$. El operador definido por, $D_a^n f := D^m J_a^{m-n} f$

Es llamado Operador Diferencial Fraccionario de Riemann Liouville de orden n . Para $n=0$ tenemos $D_a^0 = I$, operador identidad.

Una vez más vemos que, como consecuencia del Lema 2.2, el recién definido con el nombre del operador coincide con el operador diferencial clásico D^n siempre que $n \in \mathbb{N}$.

En el Lema 2.2 no habíamos requerido que m fuera lo más pequeño posible; de hecho, arbitrario los números naturales para m estaban permitidos siempre que se cumpliera la desigualdad $m > n$. afirmación similar se cumple aquí.

Lema 2.11

Sea $n \in \mathbb{R}_+$ y sea $m \in \mathbb{N}$ tal que $m > n$. Entonces

$$D_a^n = D^m J_a^{m-n}$$

Prueba

Nuestras suposiciones sobre m implican que $m \geq [n]$, de este modo,

$$\begin{aligned} D^m J_a^{m-n} &= D^{[n]} D^{m-[n]} J_a^{m-[n]} J_a^{[n]-n} \\ &= D^{[n]} J_a^{m-[n]} D^{m-[n]} J_a^{[n]-n} \\ &= D^{[n]} J_a^{m-[n]} D^{m-[n]} J_a^{[n]-n} = D^{[n]} J_a^{[n]-n} = D_a^n \\ &D^m J_a^{m-n} = D_a^n \end{aligned}$$

en vista de la propiedad de semigrupo de integración fraccionaria y (1.1).

El siguiente resultado contiene una condición suficiente muy simple para la existencia de $D_a^n f$.

Lema 2.12

Sea $f \in A^1[a, b]$ y $0 < n < 1$. Entonces $D_a^n f$ existe en casi todas partes en $[a, b]$. Además $D_a^n f \in L^p[a, b]$ para $1 \leq p < 1/n$ y

$$D_a^n f(x) = \frac{1}{\Gamma(1-n)} \left(\frac{f(a)}{(x-a)^n} + \int_a^x f'(t) (x-a)^{-n} dt \right)$$

Prueba

Usamos la definición del operador diferencial de Riemann-Liouville y el hecho de que $f \in A^1[a, b]$. Esto produce

$$\begin{aligned} D_a^n f(x) &= \frac{1}{\Gamma(1-n)} \frac{d}{dx} \int_a^x f(t) (x-a)^{-n} dt \\ &= \frac{1}{\Gamma(1-n)} \frac{d}{dx} \int_a^x \left(f(a) + \int_a^t f'(u) du \right) (x-a)^{-n} dt \\ &= \frac{1}{\Gamma(1-n)} \frac{d}{dx} \left(f(a) \int_a^x \frac{dt}{(x-a)^n} + \int_a^x \int_a^t f'(u) (x-a)^{-n} du dt \right) \\ D_a^n f(x) &= \frac{1}{\Gamma(1-n)} \left(\frac{f(a)}{(x-a)^n} + \frac{d}{dx} \int_a^x \int_a^t f'(u) (x-a)^{-n} du dt \right) \end{aligned}$$

Por el Teorema de Fubini podemos intercambiar el orden de integración en la doble integral. Esto produce

$$D_a^n f(x) = \frac{1}{\Gamma(1-n)} \left(\frac{f(a)}{(x-a)^n} + \frac{d}{dx} \int_a^x f'(u) \frac{(x-u)^{1-n}}{1-n} dt \right)$$

Las reglas estándar sobre la diferenciación de integrales de parámetros dan entonces el valor deseado representación. El enunciado de integrabilidad es una consecuencia inmediata de esta representación utilizando los resultados clásicos de la teoría de integración de Lebesgue.

Como consecuencia inmediata de esta definición, enunciaremos las derivadas fraccionarias de algunas funciones elementales.

Ejemplo

Sea $f(x) = (x - a)^\beta$ para algún $\beta > -1$ y $n > 0$. Entonces, en vista del Ejemplo 2.1,

$$D_a^n f(x) = D_a^{[n]} J_a^{[n]-n} f(x) = \frac{\Gamma(\beta + 1)}{\Gamma([n] - n + \beta + 1)} D^{[n]}[(x - a)^{[n]-n+\beta}](x)$$

Específicamente, si $n - \beta \in \mathbb{N}$, el lado derecho es el n -ésima derivada de un polinomio clásico de grado $[n] - n + \beta \in \{0, 1, 2, \dots, [n] - 1\}$, por lo que la expresión se desvanece, es decir

$$D^{[n]}[(x - a)^{[n]-n+\beta}](x) = 0 \text{ para todo } n > 0, m \in \{0, 1, 2, \dots, [n]\}.$$

Por otro lado, si $n - \beta \notin \mathbb{N}$, encontramos,

$$D^{[n]}[(x - a)^\beta](x) = \frac{\Gamma(\beta + 1)}{\Gamma(\beta + 1 - n)} (x - a)^{\beta - n}$$

Ambas relaciones son generalizaciones directas de lo que sabemos por derivadas de orden entero. Nótese que en la última expresión el argumento de la función gamma puede ser negativa. Discutimos la interpretación de esto en Apéndice D.1.

El siguiente ejemplo muestra, sin embargo, que no todas las relaciones pueden llevarse de manera directa del marco clásico a las derivadas fraccionarias

Sea $f(x) = \exp(\lambda x)$ para algún $\lambda > 0$, y sea $n > 0$, $n \notin \mathbb{N}$. Entonces

$$D_a^n f(x) = \frac{\exp(\lambda a)}{\Gamma(1 - n)} (x - a)^{-n} {}_1F_1(1; 1 - n; \lambda(x - a))$$

Aquí ${}_1F_1$ denota la función hipergeométrica confluyente de Kummer. Esta expresión se ve algo diferente de lo familiar.

$$D^n f(x) = \lambda^n \exp(\lambda x), \dots (2.4)$$

que se cumple para $n \in \mathbb{N}$.

Este extraño comportamiento está en cierto sentido relacionado con la no unicidad de operadores fraccionarios integrales (y diferenciales). Si hubiésemos elegido $a = -\infty$, entonces podríamos haber obtenido (2.4). Los operadores correspondientes que son conocidos como integrales fraccionarias de Liouville y han sido investigadas. Para nuestro caso, sin embargo, no tienen un papel importante que desempeñar porque conducirían naturalmente a un análisis sobre intervalos ilimitados, y ese no sería un escenario natural para las ecuaciones diferenciales que se considerarán más adelante. Así que no los usaremos más.

Hemos visto en el Teorema 2.2 que los operadores integrales de Riemann-Liouville forman un semigrupo. Es una consecuencia inmediata de su definición que lo clásico de los operadores diferenciales $\{D^n; n \in \mathbb{N}_0\}$ también tienen una propiedad de semigrupo. Por tanto, es natural preguntar cuándo los operadores diferenciales de Riemann-Liouville tienen tal estructura. Comenzamos nuestras investigaciones en esta dirección con un resultado positivo.

Teorema 2.13

Suponga que $n_1, n_2 \geq 0$. Además, sea $\phi \in L_1[a, b]$ y $f = J_a^{n_1+n_2} \phi$

Entonces,

$$D_a^{n_1} D_a^{n_2} f = D_a^{n_1+n_2} f$$

Nótese que: para aplicar esta identidad no necesitamos conocer la función explícitamente; es suficiente saber que tal función exista. En vista del Teorema 2.6, la condición sobre f implica no sólo un cierto grado de suavidad, sino también el hecho de que, como $x \rightarrow a$, $f(x) \rightarrow 0$ es suficientemente.

Prueba.

Por nuestra suposición sobre f y la definición de Riemann-Liouville difieren en

$$D_a^{n_1} D_a^{n_2} f = D_a^{n_1} D_a^{n_2} J_a^{n_1+n_2} \phi = D_a^{[n_1]} J_a^{[n_1]-n_1} D_a^{[n_2]} J_a^{[n_2]-n_2} J_a^{n_1+n_2} \phi$$

La propiedad del semigrupo de los operadores integrales nos permite reescribir esta expresión como

$$\begin{aligned} D_a^{n_1} D_a^{n_2} f &= D_a^{[n_1]} J_a^{[n_1]-n_1} D_a^{[n_2]} J_a^{[n_2]+n_1} \phi = \\ &= D_a^{[n_1]} J_a^{[n_1]-n_1} D_a^{[n_2]} J_a^{[n_2]} J_a^{n_1} \phi \end{aligned}$$

Debido a (1.1) y el hecho de que las órdenes de los operadores integral y diferenciales involucrados son números naturales, encontramos que esto es equivalente a

$$D_a^{n_1} D_a^{n_2} f = D_a^{[n_1]} J_a^{[n_1]-n_1} J_a^{n_1} \phi = D_a^{[n_1]} J_a^{[n_1]} \phi$$

donde hemos usado una vez más la propiedad de semigrupo de integración fraccionaria. Nosotros ahora podemos usar (1.1) una vez más y encontrar que

$$D_a^{n_1} D_a^{n_2} f = \phi$$

La prueba de que $D_a^{n_1+n_2} f = \phi$ sigue líneas similares.

La condición de suavidad y cero en este teorema no es solo un tecnicismo. Los siguientes ejemplos muestran algunos casos en los que no se cumple la condición. Ellos prueban que una propie-

dad semigrupo incondicional de diferenciación fraccionaria en el sentido de Riemann-Liouville no se sostiene.

Ejemplo

Sea $f(x) = x^{-\frac{1}{2}}$ y $n_1 = n_2 = \frac{1}{2}$. Entonces, como se muestra en el Ejemplo 1.4, $D_0^{n_1}f(x) = D_0^{n_2}f(x) = 0$, y por lo tanto también $D_0^{n_1}D_0^{n_2}f(x) = 0$, pero $D_0^{n_1+n_2}f(x) = D^1f(x) = -(2x^2)^{-1}$.

Ejemplo

Sea $f(x) = x^{\frac{1}{2}}$ y $n_1 = \frac{1}{2}$; $n_2 = \frac{3}{2}$. Luego, nuevamente usando Ejemplo 1.4,

$$D_0^{n_1}f(x) = \frac{\sqrt{\pi}}{2} \text{ y } D_0^{n_2}f(x) = 0. \text{ Esto implica } D_0^{n_1}D_0^{n_2}f(x) = 0 \text{ pero } D_0^{n_2}D_0^{n_1}f(x) = -\frac{x^{-3/2}}{4} = D^2f(x) = D_0^{n_1+n_2}f(x)$$

En otras palabras, el primero de estos dos ejemplos muestra que es posible tener

$$D_a^{n_1}D_a^{n_2}f(x) = D_a^{n_2}D_a^{n_1}f(x) \neq D_a^{n_1+n_2}f(x)$$

mientras que el segundo ejemplifica el caso donde

$$D_a^{n_1}D_a^{n_2}f(x) \neq D_a^{n_2}D_a^{n_1}f(x) = D_a^{n_1+n_2}f(x)$$

Sostiene.

Recuerde que una de las características clave que queríamos obtener era (2.1). Da vueltas que las definiciones de Riemann-Liouville tienen esta propiedad

Teorema 2.14

Sea $n \geq 0$. Entonces, para todo $f \in L_1[a, b]$,

$$D_a^n J_a^n f = f$$

Casi en cualquier parte

Prueba.

El caso $n = 0$ es trivial para entonces D_a^n y J_a^n son ambos el operador de identidad.

Para $n > 0$ se procede como en la demostración del Teorema 2.13: Sea $m = [n]$.

Entonces, por la definición de D_a^n , la propiedad del semigrupo de integración fraccionaria y (2.1) (que puede aplicarse aquí ya que $m \in \mathbb{N}$),

$$D_a^n J_a^n f(x) = D_a^m J_a^{m-n} J_a^n f(x) = D_a^m J_a^m f(x) = f(x)$$

Esencialmente, este resultado y su prueba ya han sido conocidos por Abel Adomian (1994), incluso aunque no ha denotado los operadores involucrados como integrales y derivadas de orden fraccionario, respectivamente.

Ahora llegamos a un análogo del Teorema 2.7

Teorema 2.15

Sea $n > 0$. Suponga que $(f_k)_{k=1}^\infty$ es una secuencia uniformemente convergente de funciones continuas en $[a, b]$, y que $D_a^n f_k$ para cada k . Además, asumir que $(D_a^n f_k)_{k=1}^\infty$ converge uniformemente en $[a + \varepsilon, b]$ para todo $\varepsilon > 0$. Entonces, para todo $x \in (a, b]$, tenemos

$$\left(\lim_{k \rightarrow \infty} D_a^n f_k \right)(x) = \left(D_a^n \lim_{k \rightarrow \infty} f_k \right)(x)$$

Prueba

Recordamos que $D_a^n = D^{[n]} J_a^{[n]-n}$. Por el Teorema 1.7, la sucesión $(J_a^{[n]-n} f_k)_k$ es uniformemente convergente, y podemos intercambiar la operación límite y la operación fraccionaria integral.

Por suposición, la $[n]$ -ésima derivada de esta serie converge uniformemente en cada subintervalo compacto de (a,b) . Por lo tanto, por un teorema estándar del análisis, también podemos intercambiar el operador límite y el operador diferencial siempre que $a < x \leq b$.

Podemos deducir inmediatamente un análogo del Corolario 2.8.

Teorema 2.16

Sea f analítica en $(a - h, a + h)$ para algún $h > 0$, y sea $n > 0$, $n \notin \mathbb{N}$, Entonces

$$D_a^n f(x) = \sum_{k=0}^{\infty} \binom{n}{k} \frac{(x-a)^{k-n}}{\Gamma(k+1-n)} D^k f(x)$$

$$D_a^n f(x) = \sum_{k=0}^{\infty} \frac{(x-a)^{k-n}}{\Gamma(k+1-n)} D^k f(a)$$

para $a < x < a+h$. En particular, $D_a^n f$ es analítica en $(a, a+h)$.

En este resultado, los coeficientes binomiales $\binom{n}{k}$ para $n \in \mathbb{R}$ y $k \in \mathbb{N}_0$ están definidos por,

$$\binom{n}{k} = \frac{n(n-1)(n-2)\dots(n-k+1)}{k!}, \dots \quad (2.5)$$

Prueba

Usamos el Corolario 2.8 y la definición del operador D_a^n

$$D_a^n f(x) = D^{[n]} J_a^{[n]-n} f(x)$$

Esto produce inmediatamente las dos últimas declaraciones. Para la primera reivindicación, procedemos en una mane-

ra similar, usando el hecho de que $k! \Gamma(n)(n+k) \binom{n}{k} = (-1)^k \Gamma(k+1+n)$. Ejercicio 2.2). Esto nos permite reescribir el primer enunciado del Corolario 2.8 como,

$$J_a^{[n]-n} f(x) = \sum_{k=0}^{\infty} \binom{n-[n]}{k} \frac{(x-a)^{k+[n]-n}}{\Gamma(k+1+[n]-n)} D^k f(x)$$

Diferenciando $[n]$ veces con respecto a x , encontramos

$$D_a^n f(x) = \sum_{k=0}^{\infty} \binom{n-[n]}{k} \frac{1}{\Gamma(k+1+[n]-n)} D^{[n]} [(x-a)^{k+[n]-n} D^k f](x)$$

La versión clásica de la fórmula de Leibniz (del Teorema 2.A a continuación) produce entonces,

$$\begin{aligned} D_a^n f(x) &= \sum_{k=0}^{\infty} \binom{n-[n]}{k} \frac{1}{\Gamma(k+1+[n]-n)} \times \sum_{j=0}^{[n]} \binom{[n]}{j} D^{[n]-j} [(x-a)^{k+[n]-n}] (x) D^{k+j} f(x) \\ &= \sum_{k=0}^{\infty} \binom{n-[n]}{k} \sum_{j=0}^{[n]} \binom{[n]}{j} \frac{(x-a)^{k+j-n}}{\Gamma(k+1+[n]-n)} D^{k+j} f(x) \end{aligned}$$

Por definición, $\binom{\mu}{j} = 0$ si $\mu \in \mathbb{N}$ y $\mu < j$. Por lo tanto, podemos reemplazar el límite superior en la suma interna por ∞ sin cambiar la expresión. La sustitución $j = \ell - k$ da

$$\begin{aligned} D_a^n f(x) &= \sum_{k=0}^{\infty} \sum_{\ell=k}^{\infty} \binom{n-[n]}{k} \binom{[n]}{\ell-k} \frac{(x-a)^{\ell-n}}{\Gamma(\ell+1-n)} D^{\ell} f(x) \\ &= \sum_{\ell=0}^{\infty} \sum_{k=0}^{\ell} \binom{n-[n]}{k} \binom{[n]}{\ell-k} \frac{(x-a)^{\ell-n}}{\Gamma(\ell+1-n)} D^{\ell} f(x) \end{aligned}$$

Un cálculo explícito produce,

$$= \sum_{k=0}^{\ell} \binom{n-[n]}{k} \binom{[n]}{\ell-k} = \binom{[n]}{\ell-k}$$

(ver también el Ejercicio 1.2) y así sigue la primera afirmación.

Otra peculiaridad interesante de la diferenciación fraccionaria surge cuando observe las generalizaciones de las reglas clásicas para diferenciar funciones que son compuesto de otras funciones de una manera determinada. El primer resultado en este sentido es trivial.

Teorema 2.17

Sean dos funciones definidas en $[a, b]$ tales que $D_a^n f_1$ y $D_a^n f_2$ existe en casi todas partes. Además, sea $c_1, c_2 \in \mathbb{R}$ Entonces, $D_a^n(c_1 f_1 + c_2 f_2)$, existe casi en todas partes, y

$$D_a^n(c_1 f_1 + c_2 f_2) = c_1 D_a^n f_1 + c_2 D_a^n f_2.$$

Prueba

Esta propiedad de linealidad del operador diferencial fraccionario es una consecuencia de la definición de D_a^n .

Cuando se trata de productos de funciones, la situación es completamente diferente. En el caso clásico, tenemos el siguiente resultado bien conocido (que ya hemos usado en la demostración del Corolario 3.16):

Teorema 2.A. Fórmula de Leibniz

Sea $n \in \mathbb{N}$, y sea $f, g \in C^n[a, b]$. Entonces,

$$D^n[f \cdot g] = \sum_{k=0}^n \binom{n}{k} (D^k f) (D^{n-k} g)$$

Señalamos dos propiedades especiales de este resultado: La fórmula es simétrica, es decir podemos intercambiar f y g en am-

bos lados de la ecuación sin alterar la expresión, y para evaluar la n -ésima derivada del producto $f g$, solo necesitamos derivadas hasta el orden n de ambos factores. En particular, ninguno de los factores necesita tener una $(n + 1)$ st derivada. El siguiente teorema transfiere la fórmula de Leibniz a la configuración fraccionaria, y es inmediatamente evidente que estas dos propiedades se confunden.

Teorema 2.18. Fórmula de Leibniz para operadores de Riemann-Liouville

Sea $n > 0$, y supongamos que f y g son analíticas en $(a-h, a+h)$ para algún $h > 0$. Entonces

$$D_a^n[f \cdot g](x) = \sum_{k=0}^{[n]} \binom{n}{k} (D^k f)(x) (D^{n-k} \cdot g)(x) + \sum_{k=[n]+1}^{\infty} \binom{n}{k} (D^k f)(x) (J_a^{k-n} \cdot g)(x)$$

Para $a < x < a + \frac{h}{2}$.

Nótese en el teorema que k pasa por los enteros no negativos; por lo tanto, nosotros podríamos haber escrito $D^k f$ en lugar de $D_a^n f$ en el lado derecho. Además, k corre a través de todos los enteros no negativos, y por lo tanto necesitamos tener $f \in C^\infty[a, b]$ para tener que el lado derecho tenga sentido. Los requisitos de suavidad de g parecen ser mucho menos restrictivo (las derivadas de g solo se requieren hasta el orden n), pero para nuestra prueba necesitamos analiticidad del producto $f g$, y esto generalmente es solo seguro si g es analítico también. Esto muestra que las dos propiedades principales mencionadas anteriormente son en efecto similares. Finalmente mencionamos que los operadores integrales de Riemann-Liouville surgen en el lado derecho. Tales expresiones no estaban presentes en la formulación clásica.

A pesar de todas estas diferencias recuperamos el resultado clásico de la fraccionaria resultado usando un valor entero para n porque entonces los coeficientes binomiales $\binom{n}{k}$ son cero para $k > n$, de modo que la segunda suma (la que causa todas las diferencias) desaparece.

Prueba

$$D_a^n[f.g](x) = \sum_{k=0}^{\infty} \binom{n}{k} \frac{(x-a)^{k-n}}{\Gamma(k+1-n)} D^k[f.g](x)$$

Ahora aplicamos la fórmula estándar de Leibniz a e intercambiamos el orden de suma. Esto produce

$$\begin{aligned} D_a^n[f.g](x) &= \sum_{k=0}^{\infty} \binom{n}{k} \frac{(x-a)^{k-n}}{\Gamma(k+1-n)} \sum_{j=0}^k \binom{k}{j} D^k f(x) D^{k-j} g(x) \\ &= \sum_{j=0}^{\infty} \sum_{k=j}^{\infty} \binom{n}{k} \frac{(x-a)^{k-n}}{\Gamma(k+1-n)} \binom{k}{j} D^k f(x) D^{k-j} g(x) \\ &= \sum_{j=0}^{\infty} D^j f(x) \sum_{\ell=0}^{\infty} \binom{n}{\ell+j} \frac{(x-a)^{\ell+j-n}}{\Gamma(\ell+j+1-n)} \binom{\ell+j}{j} D^{\ell} g(x) \end{aligned}$$

La observación,

$$\binom{n}{\ell+j} \binom{\ell+j}{j} = \binom{n}{j} \binom{n-j}{\ell}$$

Nos proporciona,

$$\begin{aligned} D_a^n[f.g](x) &= \sum_{j=0}^{\infty} D^j f(x) \binom{n}{j} \sum_{\ell=0}^{\infty} \binom{n-j}{\ell} \frac{(x-a)^{\ell+j-n}}{\Gamma(\ell+j+1-n)} D^{\ell} g(x) \\ &= \sum_{j=0}^{\infty} \binom{n}{j} D^j f(x) \sum_{\ell=0}^{\infty} \binom{n-j}{\ell} \frac{(x-a)^{\ell+j-n}}{\Gamma(\ell+j+1-n)} D^{\ell} g(x) \\ &\quad + \sum_{j=[n]+1}^{\infty} \binom{n}{j} D^j f(x) \sum_{\ell=0}^{\infty} \binom{n-j}{\ell} \frac{(x-a)^{\ell+j-n}}{\Gamma(\ell+j+1-n)} D^{\ell} g(x) \end{aligned}$$

Por las primeras partes de los Corolarios 2.16 y 2.8, respectivamente, podemos reemplazar el interior sumas, y el resultado deseado sigue.

La otra regla importante para la evaluación es la regla de la cadena,

$$D[g(f(\cdot))](x) = (Dg)(f(x))Df(x)$$

En el mismo sentido que la fórmula de Leibniz es la generalización de la regla del producto para primeras derivadas a derivadas de orden arbitrario $n \in \mathbb{N}$, la regla de la cadena también se puede generalizado al caso de n -ésimas derivadas con $n \in \mathbb{N}$. El resultado se conoce como Faà di Bruno's fórmula de Di Bruno. Se puede escribir de varias formas, la más conocida es probablemente la llamada versión de partición establecida que ahora recordamos.

Teorema 2.B. Fórmula de Faà di Bruno

Si g y f son funciones con suficiente número de derivadas y $n \in \mathbb{N}$, entonces

$$D^n[g(f(\cdot))](x) = \sum (D^k g)(f(x)) \prod_{\mu \in \mathcal{P}_n} (D^\mu f(x))^{b_\mu}$$

donde la suma es sobre todas las particiones de $\{1, 2, \dots, n\}$ y para cada partición, k es su número de bloques y b_j es el número de bloques con exactamente j elementos.

Un relato agradablemente legible de la historia de esta fórmula, que incluye una prueba y una discusión de muchos aspectos relacionados, se da en (Johnson, 2002). Para nuestros propósitos, solo ilustrar la fórmula con un ejemplo:

Ejemplo

$$\frac{d^4}{dx^4} g(f(x)) = g'(f(x))f^4(x) + 4g''(f(x))f'(x)f'''(x) + 6g'''(f(x))[f'(x)]^2f''(x) + g^{(4)}(f(x))[f'(x)]^4$$

Para confirmar esta afirmación a través de la fórmula de Fa`a di Bruno, necesitamos escribir levanta todas las particiones de $\{1,2,3,4\}$ y cuente sus bloques y los elementos en cada bloque.

Las posibles particiones son

$\{1,2,3,4\};$
 $\{1\}, \{2,3,4\}; \{2\}, \{1,3,4\}; \{3\}, \{1,2,4\}; \{4\}, \{1,2,3\};$
 $\{1,2\}, \{3,4\}; \{1,3\}, \{2,4\}; \{1,4\}, \{2,3\};$
 $\{1\}, \{2\}, \{3,4\}; \{1\}, \{3\}, \{2,4\}; \{1\}, \{4\}, \{2,3\};$
 $\{2\}, \{3\}, \{1,4\}; \{2\}, \{4\}, \{1,3\}; \{3\}, \{4\}, \{1,2\};$
 $\{1\}, \{2\}, \{3\}, \{4\}.$

Por lo tanto, tenemos una partición que consta de un solo bloque con cuatro elementos (que corresponde a $b_4 = 1, b_1 = b_2 = b_3 = 0$), cuatro particiones que consisten en dos bloques con uno y tres elementos, respectivamente (es decir, con $b_1 = b_3 = 1, b_2 = b_4 = 0$), tres tabiques formados por tres bloques de 1, 1 y 2 elementos respectivamente ($b_1 = 2, b_2 = 1, b_3 = b_4 = 0$), seis particiones formadas por dos bloques de dos elementos cada uno ($b_2 = 2, b_1 = b_3 = b_4 = 0$), y finalmente una partición que consta de cuatro bloques con un elemento cada uno ($b_1 = 4, b_2 = b_3 = b_4 = 0$). Tomando la suma de todas las particiones así obtener la fórmula requerida.

Para esta fórmula, también se conoce un reemplazo en el ajuste fraccionario (Johnson, 2002). Lo declaramos aquí sin pruebas en aras de la exhaustividad.

Teorema 2.19 Fórmula de Faá di Bruno para operadores de Riemann-Liouville.

Bajo suposiciones adecuadas sobre las funciones f y g tenemos

$$D_a^n [f(g(\cdot))](x) = \sum_{k=1}^{\infty} \binom{n}{k} \frac{K!(x-a)^{k-n}}{\Gamma(1-n)} \sum_{\ell=1}^k (D^\ell f(x))(g(x)) \sum_{a_1, a_2, \dots, a_k \in A_{k,\ell}} \prod_{r=1}^k \frac{1}{a_k} \left(\frac{D^r g(x)}{r!} \right)^{a_k} + \frac{(x-a)^{-n}}{\Gamma(k-n+1)} f(g(x)), \dots \quad (2.6)$$

Donde $(a_1, \dots, a_k) \in A_k$, significa que

$$a_1, a_2, \dots, a_k \in \mathbb{N}_0, \sum_{r=1}^k r a_r = k \quad \sum_{r=1}^k r a_r = \ell$$

La estructura del lado derecho de (3.6) es tan compleja que difícilmente parece ser de alguna utilidad práctica. Por lo tanto, no daremos más detalles sobre esta regla, y en su lugar recurran a un tipo completamente diferente de problemas relacionados con operadores diferenciales fraccionario

Específicamente, una pregunta natural para hacer es: ¿Qué se puede decir sobre $D_a^n f$ como $n \rightarrow m \in \mathbb{N}$? Comenzamos mirando un caso especial

Ejemplo 2.9

Sea $f(x) = (x-a)^k$ para algún $k > 0$. Entonces tenemos, debido al Ejemplo 2.4

$$D^m f(x) = \frac{\Gamma(k+1)}{\Gamma(k+1-m)} (x-a)^{k-m} \quad \text{y} \quad D_a^n f(x) = \frac{\Gamma(k+1)}{\Gamma(k+1-n)} (x-a)^{k-n}$$

Comparando estas dos expresiones, encontramos en el límite $n \rightarrow m$: $D_a^n f(x) \rightarrow D_a^m f(x)$ uniformemente en $[a, b]$ si $m < k$.

Para $m = k$ obtenemos que $D^m f(x) = \Gamma(m+1)$, que es una constante distinta de cero, mientras que $D_a^n f(x) = \begin{cases} 0, & \text{si } n < m \\ \infty & \text{si } n > m \end{cases}$

Entonces tenemos convergencia puntual en $(a, b]$, pero no en el punto a , y por lo tanto la convergencia uniforme en el intervalo completo $[a, b]$ no es posible. en el restante caso $m > k$ solo obtenemos convergencia puntual en $(a, b]$ porque la diferencia de las dos expresiones siempre es ilimitada en a .

En el caso de que f sea de una forma muy general y no tan simple como en este ejemplo, podemos enunciar una observación similar, aunque un poco más débil

Teorema 2.20

Sea $f \in C^m[a, b]$ para algún $m \in \mathbb{N}$. Entonces,

$$\lim_{n \rightarrow m} D_a^n f = D^m f$$

en un sentido puntual en (a, b) . La convergencia es uniforme en $[a, b]$ si adicionalmente $f(x) = o((x - a)^{m+\delta})$ para algún $\delta > 0$ cuando $x \rightarrow a^+$.

Prueba

En vista de la suposición de suavidad en f , podemos realizar una expansión de

Taylor de f centrada en a y encontrar $f(x) = T_{m-1}[f; a](x) + R_{m-1}[f; a](x)$ donde

$$T_{m-1}[f; a](x) = \sum_{k=0}^{m-1} \frac{f^{(k)}(a)}{k!} (x - a)^k$$

y R_{m-1} denota el término restante. tenemos eso

$$\begin{aligned} D_a^n f(x) - D^m f(x) &= D_a^n T_{m-1}[f; a](x) - D^m T_{m-1}[f; a](x) + D_a^n R_{m-1}[f; a](x) \\ &\quad - D^m R_{m-1}[f; a](x) \end{aligned}$$

Como T_{m-1} es un polinomio de grado $m-1$, encontramos que $T_{m-1}[f; a] \equiv 0$. Además, por E

$$D_a^n T_{m-1}[f; a](x) = \sum_{k=0}^{m-1} \frac{f^{(k)}(a)}{\Gamma(k+1-n)} (x-a)^{k-n}$$

De este modo,

$$\begin{aligned} D_a^n f(x) - D^m f &= D^m J_a^{m-n} R_{m-1}[f; a](x) + \sum_{k=0}^{m-1} \frac{f^{(k)}(a)}{\Gamma(k+1-n)} (x-a)^{k-n} \\ &- D^m R_{m-1}[f; a](x) \end{aligned}$$

La suma es idénticamente cero bajo el supuesto de que $f(x) = O((x-a)^m)$ porque las derivadas relevantes de f desaparecen en a . Sin esa suposición adicional, se anula en el límite cuando $n \rightarrow m$ en un sentido puntual sobre $(a, b]$ porque los argumentos de las funciones Gamma convergen a un entero no positivo. Dado que la función Gamma tiene polos en los enteros no positivos, el límite es cero.

Queda por demostrar que la diferencia $D_a^n R_{m-1}[f; a] - D^m R_{m-1}[f; a]$ tiene las propiedades de convergencia requeridas. Para ello, utilizamos la representación integral de R_{m-1} que, en nuestra notación, se puede escribir en la forma

$$R_{m-1}[f; a](x) = \frac{1}{\Gamma(m)} \int_a^x f^{(m)}(u) (x-a)^{m-1} du$$

De este modo,

$$D^m R_{m-1}[f; a](x) = D^m J_a^m D^m f = D^m f \square$$

Para el otro término podemos escribir

$$\begin{aligned} D^m R_{m-1}[f; a] &= D^m J_a^{m,n} D^m R_{m-1}[f; a] \\ &= D^m J_a^{m-n} J_a^n D^m f = D^m J_a^m J_a^{m-n} D^m f = J_a^{m-n} D^m f \square \square \square \end{aligned}$$

en vista de la propiedad del semigrupo de integración fraccionaria y el Teorema 2.14. Por suposición, $D^m f$ es conti-

nua en $[a, b]$. Por lo tanto, la convergencia puntual de $\text{contra } J_a^{m-n} D^m f$ $\text{contra } J_a^n D^m f = D^m f$ en $(a, b]$ se sigue del Teorema 3.10.

Además, si $f(x) = O((x-a)^{m+\delta})$ entonces $D^m f(x) = O((x-a)^\delta)$ cuando $x \rightarrow a$, y por lo tanto, tenemos una convergencia uniforme en el intervalo completo $[a, b]$ por la última declaración de Teorema 3.10. Esto completa la prueba.

Observación 2.3

Hay otra diferencia fundamental entre los operadores diferenciales de orden entero y las derivadas fraccionarias de Riemann-Liouville: el primero son operadores locales, estos últimos no lo son. El significado de la palabra local aquí es el siguiente. Para calcular $D^m f(x)$ para $n \in \mathbb{N}$, es suficiente conocer f de forma arbitraria pequeño alrededor de x . Esto se sigue de la representación clásica de $D^m f(x)$ como límite de un cociente de diferencias. Sin embargo, para calcular $D_a^m f(x)$ para $n \notin \mathbb{N}$, la definición nos dice que necesitamos saber f en todo el intervalo $[a, x]$.

Esta propiedad se exhibe de manera aún más evidente en el siguiente Lema que proporciona una representación alternativa de la derivada de Riemann-Liouville. Esta nueva representación ha demostrado ser bastante útil en el desarrollo y presentación de ciertos algoritmos numéricos.

Teorema 2.21

Sea $n > 0$, $n \notin \mathbb{N}$, y $m = [n]$ Suponga que $f \in C^m[a, b]$ y $x \in [a, b]$. Entonces

$$D_a^n f(x) = \frac{1}{\Gamma(-n)} \int_a^x (x-t)^{-n-1} f(t) dt$$

En esta declaración, la integral necesita alguna explicación adicional. El integrando exhibe una singularidad de orden $n + 1$ que es estrictamente mayor que uno, y por lo tanto la voluntad integral en general no existe ni en el sentido propio ni en el impropio. Por lo tanto, definimos dicha integral de acuerdo con el concepto de integral de partes finitas de Hadamard como se explica en el Apéndice D.4 (1952).

Se puede dar una prueba muy corta y simple del Lema usando métodos de la teoría de funciones generalizadas (distribuciones). Sin embargo, no queremos presentar esta maquinaria aquí. Por lo tanto, damos una demostración más elemental de que sigue esencialmente las ideas de Elliott. Tenga en cuenta que, de acuerdo con esa referencia, ciertas partes de la demostración ya se pueden encontrar en un artículo de Marchaud que data volver a 1927.

Prueba

Por definición, $D_a^m f(x) = D^m J_a^{m-n} f(x)$ y

$$J_a^{m-n} f(x) = \frac{1}{\Gamma(m-n+1)} \int_a^x (x-t)^{m-n-1} f(t) dt$$

En vista de los supuestos de suavidad en f , podemos integrar parcialmente en este integral y hallar que

$$\begin{aligned} J_a^{m-n} f(x) &= \frac{1}{\Gamma(m-n+1)} \int_a^x (x-t)^{m-n} f'(t) dt \\ &= \frac{1}{\Gamma(m-n+1)} (x-t)^{m-n} f(t) \Big|_{t=a}^{t=x} \\ &= \frac{1}{\Gamma(m-n+1)} (x-t)^{m-n} f(a) + J_a^{m-n+1} f'(x) \end{aligned}$$

Los supuestos de suavidad nos permiten repetir este procedimiento un total de m veces; nosotros fin

$$J_a^{m-n} f(x) = \sum_{k=0}^{m-1} \frac{(x-t)^{k+m-n}}{\Gamma(k+m-n+1)} f^{(k)}(a) + J_a^{2m-n} f^{(m)}(x)$$

De esta manera,

$$D_a^n f(x) = D^m J_a^{m-n} f(x) = \sum_{k=0}^{m-1} \frac{(x-t)^{k-n}}{\Gamma(k+m-n+1)} f^{(k)}(a) + J_a^{2m-n} f^{(m)}(x)$$

Pero de acuerdo con el Teorema D.14 la expresión del lado derecho de la ecuación es solo $\frac{\int_a^x (x-t)^{-n-1} f(t) dt}{\Gamma(-n)}$

Observación 2.4

Otra característica interesante de esta representación aparece si asumimos formalmente que n es negativa, $n = -$ con alguna > 0 . Entonces la identidad de El lema 2.21 toma la forma de

$$D_a^{n^*} f(x) = \frac{1}{\Gamma(n^*)} \int_a^x (x-t)^{n^*-1} f(t) dt$$

y encontramos que la expresión del lado derecho es simplemente $J_a^{n^*} f(x)$.

Relaciones entre los operadores integrales y derivadas de Riemann-Liouville

Habiendo establecido una teoría de los operadores diferencial e integral de Riemann-Liouville por separado, ahora investigamos cómo interactúan. Un primer resultado muy importante en este contexto ya se ha mostrado en el Teorema 2.14 anterior: D_a^n es el inverso izquierdo de J_a^n . Por supuesto, no podemos afirmar que es el inverso derecho. Más precisamente, tenemos la siguiente situación.

Teorema 2.22

Sean $\alpha > 0$. Si existe algún $\phi \in L_1[a, b]$ tal que $f = J_a^m f$ entonces

$$J_a^m D_a^n f = f$$

Casi en cualquier parte

Prueba

Esto es una consecuencia inmediata del resultado anterior: Tenemos, por definición de f y el Teorema 2.14, que

$$J_a^m D_a^n f = J_a^m [D_a^n J_a^n] \phi = J_a^m \phi = f$$

Si f no es como se requiere en los supuestos del Teorema 2.22, entonces obtenemos una representación diferente para $J_a^m D_a^n f$.

Teorema 2.23

Sean $\alpha > 0$ y $m = n + 1$. Supongamos que f es tal que $J_a^{m-n} f \in A^m[a, b]$. Entonces,

$$J_a^m D_a^n f(x) = f(x) - \sum_{k=0}^{m-1} \frac{(x-a)^{n-k-1}}{\Gamma(n-k)} \lim_{z \rightarrow a^+} D_a^{m-k-1} J_a^{m-n} f(z)$$

Específicamente, para $0 < n < 1$ tenemos

$$J_a^m D_a^n f(x) = f(x) - \frac{(x-a)^{n-1}}{\Gamma(n)} \lim_{z \rightarrow a^+} J_a^{1-n} f(z)$$

Prueba

Primero observamos que los límites en el lado derecho existen debido a nuestra suposición sobre f que implica la continuidad de $D_a^{m-1} J_a^n f(x)$. Además, debido a esto supuesto, existe algún $\phi \in L_1[a, b]$ tal que

$$D_a^{m-1} J_a^n f(x) = D_a^{m-1} J_a^n f(a) + J_a^1 \phi.$$

Esta es una ecuación diferencial clásica de orden $m-1$ para $J_a^{m-n} f(x)$; su solución es fácil visto como de la forma

$$D_a^{m-n} J_a^n f(x) = \sum_{k=0}^{m-1} \frac{(x-a)^k}{k!} \lim_{z \rightarrow a^+} D_a^k J_a^{m-n} f(z) + J_a^m \phi(x) \dots (2.7)$$

Así, por definición de D_a^n ,

$$\begin{aligned}
 J_a^n D_a^n f(x) &= J_a^n D_a^m J_a^{m-n} f(x) \\
 &= J_a^n D_a^m \left[\sum_{k=0}^{m-1} \frac{(-a)^k}{k!} \lim_{z \rightarrow a^+} D^k J_a^{m-n} f(z) + J_a^m \phi(x) \right] (x) \\
 &= J_a^n D_a^m J_a^m \phi(x) + \sum_{k=0}^{m-1} \frac{J_a^n D_a^m [(-a)^k] (x)}{k!} \lim_{z \rightarrow a^+} D^k J_a^{m-n} f(z) \\
 &= J_a^n \phi(x), \dots (2.8)
 \end{aligned}$$

debido al Teorema 2.14 (nótese que D^m aniquila todo sumando en la suma). Próximo aplicamos el operador

D^{m-n} a ambos lados de (2.7) y encontramos, en vista del Teorema 2.14, eso

$$\begin{aligned}
 f(x) &= \sum_{k=0}^{m-1} \frac{D_a^{m-n} [(-a)^k] (x)}{k!} \lim_{z \rightarrow a^+} D^k J_a^{m-n} f(z) + D_a^{m-n} J_a^m \phi(x) \\
 &= \sum_{k=0}^{m-1} \frac{D_a^{m-n} [(-a)^k] (x)}{k!} \lim_{z \rightarrow a^+} D^k J_a^{m-n} f(z) + D_a^1 J_a^{1-m+n} J_a^m \phi(x)
 \end{aligned}$$

Ahora invocamos el ejemplo 2.4 para evaluar los términos en la suma y el semigrupo propiedad de integración fraccionaria y el Teorema 2.14 para manipular el resto de los términos. Esto produce

$$f(x) = \sum_{k=0}^{m-1} \frac{(x-a)^{k+n-m} (x)}{\Gamma(k+n-m+1)} \lim_{z \rightarrow a^+} D^k J_a^{m-n} f(z) + J_a^n \phi(x), \dots (2.9)$$

Finalmente sustituimos k en la suma por $m - k - 1$, resolvemos para $J_a^n \phi(x)$ y combinamos los resultados con (2.8) para obtenemos cómo se desea.

$$J_a^n D_a^n f(x) = J_a^m \phi(x) = f(x) - \sum_{k=0}^{m-1} \frac{(x-a)^{n-k-1}}{\Gamma(n-k)} \lim_{z \rightarrow a^+} D^{m-k-1} J_a^{m-n} f(z) +$$

Otro resultado básico importante en el análisis clásico es el teorema de Taylor. Tenemos ya utilizó su versión clásica en la demostración del Teorema 2.20. Se puede afirmar en la siguiente

manera que es un poco más instructiva que la formulación estándar dada en la mayoría de los libros de texto en el sentido de que da una idea adicional de la estructura del conjunto A^m . La forma clásica de enunciar este teorema se puede obtener considerando sólo la implicación $(a) \Rightarrow (b)$ y eligiendo $y = a$ allí.

Para una prueba de la versión general presentada aquí nos referimos a la monografía de Sard que también trata con problemas relacionados en un marco más general.

Teorema 2.C. Expansión de Taylor.

Las siguientes declaraciones son equivalentes

$$f \in A^m[a, b]$$

Para todo $x, y \in [a, b]$

$$f(x) = \sum_{k=0}^{m-1} \frac{(x-y)^k}{k!} D^k f(y) + J_y^m D^m f(x)$$

Con base en los resultados obtenidos hasta ahora, podemos encontrar una generalización fraccionada de esta declaración.

Teorema 2.24. Expansión fraccionaria de Taylor

Bajo los supuestos del Teorema 2.23, tenemos,

$$f(x) = \frac{(x-a)^{n-m}}{\Gamma(n-m+1)} \lim_{z \rightarrow a^+} D^k J_a^{m-n} f(z) \\ + \sum_{k=1}^{m-1} \frac{(x-a)^{k+n-m}}{\Gamma(k+n-m+1)} \lim_{z \rightarrow a^+} D^{k+n-m} f(z) + J_a^n D_a^n f(x)$$

Tenga en cuenta que, en el caso de $n \in$ tenemos $m = n + 1$ y, por lo tanto, el límite fuera de la suma se desvanece. Por lo tanto, podemos recuperar el resultado clásico (con m reemplazado por $m-1$).

Prueba

De (2.8) y (2.9) encontramos

$$f(x) = \sum_{k=1}^{m-1} \frac{(x-a)^{k+n-m}}{\Gamma(k+n-m+1)} \lim_{z \rightarrow a^+} D^k J_a^{m-n} f(z) + J_a^n D_a^n f(x)$$

Ahora movemos el sumando para $k = 0$ fuera de la suma; para los términos restantes y aplicando el Lema 2.11. Esto da el resultado

Operadores de Grunwald-Letnikov

En el cálculo clásico es bien sabido que las derivadas se pueden expresar como cocientes diferenciales, es decir, como límites de cocientes diferenciales. Por ejemplo, podemos usar diferencias hacia atrás de orden n con tamaño de paso h , denotadas y definidas por

$$\Delta_a^n f(x) := \sum_{k=0}^n (-1)^k \binom{n}{k} f(x - kh), \dots (2.10)$$

para concluir el siguiente resultado clásico:

Sean $n \in \mathbb{N}$, $f \in C^n[a, b]$ y $a < x \leq b$. Entonces

$$D^n f(x) = \lim_{h \rightarrow 0} \frac{\Delta_a^n f(x)}{h^n}$$

Este resultado en realidad no solo es útil para investigaciones analíticas; usando un valor positivo finito para h en lugar de realizar la operación de límite $h \rightarrow 0$ también nos da una aproximación numérica directa para la derivada.

En vista de estas ventajas de esta representación es evidentemente deseable tener un análogo también para el caso fraccionario. Tal construcción es posible; se remonta a la obra de Grunwald y Letnikov (1987). De hecho, todo lo que tenemos que hacer es dar

un significado a la diferencia finita en (2.10) para $n \notin \mathbb{N}$. Para ello recordemos que el binomio los coeficientes con coeficiente superior no entero se han introducido en (2.5) anterior. Ya habíamos anotado y usado la propiedad de que $\binom{n}{k} = 0$ si $n \in \mathbb{N}$ y $n < k$. De este modo, para $n \in \mathbb{N}$ (2.10) es equivalente a

$$\Delta_a^n f(x) := \sum_{k=0}^n (-1)^k \binom{n}{k} f(x - kh), \dots (2.11)$$

Ahora recuerda que queremos tener una expresión para nuestra clase de función, que normalmente es un subconjunto de $C[a,b]$, es decir, una clase de funciones definidas en el intervalo finito $[a,b]$. En este contexto observamos que la representación (2.11) introduce dos problemas en el caso $n \notin \mathbb{N}$ donde ninguno de los coeficientes binomiales se anula, por lo que esta expresión realmente representa una serie infinita:

Para evaluar la expresión en (2.11) para todo $x \in (a,b]$, la función f necesita a definir en $(-\infty, b]$

La función f debe ser tal que la serie converja.

Estos dos problemas se pueden resolver simultáneamente mediante un concepto simple: dada una función $f: [a,b] \rightarrow \mathbb{R}$, define una nueva función

$$f^*: (-\infty, b] \rightarrow \mathbb{R}$$

$$x \rightarrow \begin{cases} f(x), & \text{si } x \in [a, b] \\ 0, & \text{si } x \in (-\infty, a] \end{cases}$$

y use esta función en lugar de la f original. En vista del hecho de que f y f^* coinciden en el intervalo donde ambas funciones están definidas, f^* como continuación de f y, abusando un poco de la notación, escribiremos de ahora en adelante f en lugar de f^* . Esto nos lleva a la necesaria generalización del concepto de

cociente diferencial. En aras de la simplicidad, imponemos una restricción en la forma en que $h \rightarrow 0$; específicamente para el valor de x bajo consideración asumimos que h toma sólo los valores $h_N = \frac{(x-a)}{N}$, $N = 1, 2, \dots$, Mediante un tedioso análisis es posible demostrar que esta restricción no es necesaria, pero no entraremos en detalles aquí.

Definición 2.3

Sea $n > 0$, $f \in C^{[n]}[z, b]$ y $a < x \leq b$. Entonces

$$D^n f(x) = \lim_{N \rightarrow \infty} \frac{\Delta_{h_N}^n f(x)}{h_N^n} = \lim_{N \rightarrow \infty} \frac{1}{h_N^n} \sum_{k=0}^n (-1)^k \binom{n}{k} f(x - kh_N)$$

con $h_N = \frac{(x-a)}{N}$ se llama derivada fraccionaria de Grunwald-Letnikov de orden n de la función f .

El siguiente resultado explica la relación entre esta nueva noción de la derivada fraccionaria y la que ya conocemos.

Teorema 2.25

Sea $n > 0$, $m = [n]$ y $f \in C^{[n]}[z, b]$ Entonces, para $x \in (a, b]$,

$$\tilde{D}_a^n f(x) = D_a^n f(x)$$

Donde $[n]$ denota la función suelo $[.] = \max\{z \in \mathbb{Z}, z \leq x\}$

Prueba

Si $n \in \mathbb{N}$ entonces esto es evidente porque, como se indicó anteriormente, el diferencial los cocientes se reducen a la versión clásica que cubre el Teorema 3.D. Así nosotros ahora nos concentramos en el caso $n \notin \mathbb{N}$.

Seguiremos a Elliott, primer lugar, notamos que no hay pérdida de generalidad al suponga que $a = 0$ y $x = 1$ porque cualquier otro intervalo $[a, x]$ puede asignarse a $[0, 1]$ por una transformación afín, y todo el análisis de convergencia a continuación permanecerá sin cambios por esta transformación. Así tenemos que demostrar que

$$\lim_{N \rightarrow \infty} N^n \sum_{k=0}^N (-1)^k \binom{n}{k} f\left(1 - \frac{k}{N}\right) = D_0^n f(1), \dots (3.12)$$

En este contexto, escribamos

$$\begin{aligned} D_0^n f(x) - N^n \sum_{k=0}^N (-1)^k \binom{n}{k} f\left(1 - \frac{k}{N}\right) \\ = \left[1 - \frac{N^n \Gamma(N+1-n)}{\Gamma(N+1)} \right] D_0^n f(1) + N^n \left[1 - \frac{\Gamma(N+1-n)}{\Gamma(N+1)} D_0^n f(1) - \sum_{k=0}^N (-1)^k \binom{n}{k} f\left(1 - \frac{k}{N}\right) \right] \end{aligned}$$

Como $N \rightarrow \infty$, la expresión entre corchetes en el lado derecho de este converge a cero porque, por la fórmula de Stirling (Teorema D.5) y de la regla de L'Hospital,

$$\begin{aligned} \frac{N^n \Gamma(N+1-n)}{\Gamma(N+1)} &= N^n \left(\frac{N-n}{e}\right)^{N-n} \left(\frac{e}{N}\right)^N \frac{\sqrt{2\pi(N-n)}}{\sqrt{2\pi N}} (1 + \sigma(1)) \\ &= e^n \left(\frac{N-n}{e}\right)^{N-n} (1 + \sigma(1)) = 1 + \sigma(1) \end{aligned}$$

Por lo tanto, para probar nuestro resultado deseado (2.12), basta con demostrar que

$$N^n \left(\frac{\Gamma(N+1-n)}{\Gamma(N+1)} D_0^n f(1) - \sum_{k=0}^N (-1)^k \binom{n}{k} f\left(1 - \frac{k}{N}\right) \right) N \xrightarrow{N \rightarrow \infty} \infty, \dots (2.13)$$

Para ello introducimos un concepto auxiliar de la teoría de la aproximación, el N -ésimo polinomio de Bernstein de la función f , denotado y definido por

$$B_N[f](t) = \sum_{k=0}^N \binom{n}{k} (t)^k (1-t)^{N-k} f\left(\frac{k}{N}\right)$$

La conexión con nuestro enfoque se establece de la siguiente manera. Para

$N \in \mathbb{N}_0$ y $k \in \{0, 1, \dots, N\}$ definimos $b_{k,N}(t) := t^k(1-t)^{N-k}$. Tenga en cuenta que estas funciones son relacionadas con el polinomio de Bernstein vía $B_N[f](t) = \sum_{k=0}^N \binom{N}{k} f\left(\frac{k}{N}\right) b_{k,N}$ un explícito el cálculo entonces produce, por el Lema 2.21,

$$D_0^n b_{k,N}(1) = \frac{1}{\Gamma(-n)} \int_0^1 (1-t)^{N-k-n-1} t^k dt = \frac{\Gamma(k+1)\Gamma(N-k-n)}{\Gamma(-n)\Gamma(N-n+1)}, \dots (2.14)$$

(Esto se puede demostrar usando la integral Beta y su continuación analítica.) Como consecuencia de esta relación podemos expresar la suma en el lado izquierdo de (3.13), de acuerdo con:

$$\begin{aligned} \sum_{k=0}^N (-1)^k \binom{N}{k} f\left(1 - \frac{k}{N}\right) &= \sum_{k=0}^N \frac{\Gamma(k-n)}{\Gamma(k+1)\Gamma(-n)} f\left(1 - \frac{k}{N}\right) \\ &= \sum_{k=0}^N \frac{\Gamma(N-k-n)}{\Gamma(N-k+1)\Gamma(-n)} f\left(\frac{k}{N}\right) \\ &= \frac{1}{\Gamma(N+1)} \sum_{k=0}^N \binom{N}{k} \frac{\Gamma(k+1)\Gamma(N-k-n)}{\Gamma(-n)} f\left(\frac{k}{N}\right) \\ &= \frac{\Gamma(N+1-n)}{\Gamma(N+1)} \sum_{k=0}^N \binom{N}{k} D_0^n b_{k,N}(1) f\left(\frac{k}{N}\right) \\ &= \frac{\Gamma(N+1-n)}{\Gamma(N+1)} D_0^n \left[\sum_{k=0}^N \binom{N}{k} f\left(\frac{k}{N}\right) b_{k,N} \right](1) \\ &= \frac{\Gamma(N+1-n)}{\Gamma(N+1)} D_0^n B_N[f](1) \end{aligned}$$

Por tanto, nuestra sustentada (2.13) se reduce a

$$\begin{aligned} &N^n \frac{\Gamma(N+1-n)}{\Gamma(N+1)} D_0^n f(1) - D_0^n B_N[f](1) \\ &= N^n \frac{\Gamma(N+1-n)}{\Gamma(N+1)} D_0^n (f - B_N[f])(1) \rightarrow 0 \end{aligned}$$

Ya lo habíamos visto arriba

$$N^n \frac{\Gamma(N+1-n)}{\Gamma(N+1)} \rightarrow 1$$

como $N \rightarrow \infty$; por lo que solo queda probar que cuando $N \rightarrow \infty$. Para que este último pase usamos la representación para del

Lema 3.21 y la fundamental propiedad de la integral de partes finitas descrita en el teorema D.14 que produce

$$D_0^n (f - B_N[f])(1) = \frac{1}{\Gamma(-n)} \int_0^1 (1-t)^{-n-1} (f(t) - B_N[f](t)) dt$$

$$\sum_{k=0}^{[n]-1} \frac{1}{\Gamma(k-n+1)} D^k [f - B_N[f]](0) + \int_0^{[n]-n} D^{[n]-n} D^{[n]} [f - B_N[f]](1)$$

Ahora podemos invocar el Teorema D.16 que nos dice que, bajo nuestras suposiciones, $D^k B_N[f]$ converge a $D^k f$ uniformemente en $[0,1]$ para $k = 0, 1, \dots, n$, y por lo tanto en toda la expresión del lado derecho converge a cero según sea necesario.

La definición 1.2 plantea inmediatamente la pregunta de qué sucede si reemplazamos n por $-n$. Resulta que esta pregunta tiene una respuesta simple.

Teorema 2.26

Sean $n > 0$, $f \in C[a, b]$ y $a \leq x \leq b$. Entonces, con $b_N = \frac{(x-a)}{N}$ tenemos

$$J_a^n f(x) = \lim_{N \rightarrow \infty} h_N^n \sum_{k=0}^N (-1)^k \binom{-n}{k} f(x - k h_N)$$

Prueba

Para $x = a$, ambos lados de la ecuación se anulan. Para $x > a$, se procede formalmente como en la demostración del Teorema 2.25, sustituyendo únicamente $-n$ por n a lo largo de todo el argumento y reemplazando (2.14) por

$$J_0^n b_{k,N}(1) = \frac{1}{\Gamma(n)} \int_0^1 (1-t)^{N-k+n-1} t^k dt = \frac{\Gamma(k+1)\Gamma(N-k+n)}{\Gamma(n)\Gamma(N+n+1)}, \dots (2.14)$$

que también se puede deducir con la ayuda de la integral

Beta. Luego llegamos a la conclusión de que nuestra afirmación es equivalente a la afirmación

$$J_0^n(f - B_N[f])(1) \rightarrow 0 \text{ como } N \rightarrow \infty$$

que se sigue del Teorema 2.7 ya que $B_N[f] \rightarrow f$ uniformemente por el Teorema D.16.

En vista de este teorema y de la relación

$$\begin{aligned} (-1)^k \binom{-n}{k} &= (-1)^k \frac{(-n)(-n-1), \dots, (-n-k+1)}{k!} \\ &= \frac{(-n)(-n-1), \dots, (-n-k+1)}{k!} \\ &= \frac{(n+k-1)(n+k-2), \dots, n}{k!} \\ &= \binom{n+k-1}{k} = \frac{\Gamma(n+k)}{\Gamma(n)\Gamma(k+1)}. \end{aligned}$$

se justifica la siguiente definición formal.

Definición 2.4

Sean $n > 0$, $f \in C[a, b]$ y $a < x \leq b$. Entonces

$$\widetilde{J}_a^n f(x) := \frac{1}{\Gamma(n)} \lim_{N \rightarrow \infty} h_N^n \sum_{k=0}^N \frac{\Gamma(n+k)}{\Gamma(k+1)} f(x - k h_N)$$

Con $h_N = \frac{(x-a)}{N}$ se llama integral fraccionaria de Grunwald-Letnikov de orden n de la función f .

De hecho, esta es la definición históricamente correcta introducida en los artículos originales de Grunwald (1987) y Letnikov (en Shukla & Prajapati, 2007), respectivamente. Por lo tanto, la representación de Grunwald-Letnikov permite una unificación formal de los conceptos de derivadas fraccionarias e integrales fraccionarias al admitir valores positivos y negativos para n .

Algunos autores se han sentido motivados por esta potencial unificación introducir un concepto de notación unificada para integrales fraccionarias y derivadas, es decir escribirían, por ejemplo, D_a^{-n} en lugar de I_a^n para $n > 0$. Hemos optado por ceñirnos a símbolos separados para operadores diferenciales e integrales porque entonces es inmediatamente claro de la notación si el operador bajo consideración es de tipo diferencial o tipo integral.

Para una discusión más detallada de Grunwald-Letnikov operadores diferencial e integral nos remitimos a la monografía de (Samko et al., s. f.).

03

Capítulo

***OPERADORES DIFERENCIALES E
INTEGRALES FRACCIONARIOS DE
CAPUTO***

Escucha a tu niño interno que te enseñara ver que un maestro en sus lecciones no debe mostrarse desabrido e impaciente, en no hacer alarde de su destreza, en saber proponer y acomodar las instrucciones a la capacidad de los oyentes, saber cómo recibir de un amigo los beneficios.

Burseca

Resulta que las derivadas de Riemann-Liouville tienen ciertas desventajas cuando tratamos de modelar fenómenos del Mundo Real Consensual Organizacional Dinámico Sensitivo e Inteligible “MRCODSI” usando Ecuaciones Diferenciales Fraccionarias “EDFs”. Es en ese sentido que discutiremos ahora un concepto modificado de derivada fraccionaria, las que permitirá con mayor éxito el modelamiento matemático de sistemas dinámicos y considerar un control óptimo para la toma de decisiones aplicando los conceptos, leyes, en las diferentes áreas del conocimiento como la física, química, biología, neurociencia, la ciencias sociales, ciencias políticas, sistemas de mercado, técnica e interrelacionarlo con ciencias de la computación, haciendo uso de una diversidad de Softwares, y algoritmos que se encuentran relacionados directamente con el cálculo numérico; es decir los diferentes métodos numéricos donde entran a tallar la importancia de las Ecuaciones Diferenciales Fraccionarias.

Definición y propiedades básicas

Comenzamos con una definición preliminar.

Definición 3.1.

Sea $n \geq 0$ y $m = [n]$, Entonces, definimos el operador $\hat{D}_a^n$ por

$$\hat{D}_a^n f := J_a^{m-n} D^m f$$

Siempre que $D^m f \in L^1[a, b]$

Comencemos mirando el caso $n \in \mathbb{N}$. Aquí tenemos $m = n$ y por lo tanto nuestra definición implica $\widehat{D}_a^n f := J_a^0 D^n f = D^n f$ es decir, recuperamos la definición estándar en el caso clásico.

Comenzamos el análisis de este operador en el caso estrictamente fraccionario $n \notin \mathbb{N}$ con un ejemplo sencillo.

Ejemplo

Sea $f(x) = (x - a)^\beta$ para algún $\beta \geq 0$. Entonces,

$$\widehat{D}_a^n f(x) = \begin{cases} 0, & \text{si } \beta \in \{0, 1, 2, \dots, m-1\} \\ \frac{\Gamma(\beta+1)}{\Gamma(\beta+1-n)} (x-a)^{\beta-n}, & \text{si } \beta \in \mathbb{N} \text{ y } \beta \geq m \\ 0, & \beta \notin \mathbb{N} \text{ y } \beta > m-1 \end{cases}$$

Se anima al lector a comparar esta afirmación con la correspondiente para operadores de Riemann-Liouville (Ejemplo). Nótese en particular que los dos operadores tienen kernels diferentes, y que los dominios de los dos operadores mostrado aquí en términos del rango permitido del parámetro β también son diferentes.

Algunos ejemplos adicionales de derivados de tipo Caputo de ciertas funciones importantes se recopilan para conveniencia del lector en el Apéndice B. es decir;

Observación 3.1

Hemos requerido $m = n$ en la Definición 3.1. La misma condición es impuesta en la definición $D_a^m := D^m J_a^{m-n}$ de la derivada de Riemann-Liouville en la Definición 2.2. Sin embargo, en este último caso habíamos visto en el Lema 2.11 que esta restricción en realidad no es necesaria; uno puede usar cualquier $m \in \mathbb{R}$ con $m \geq n$

en el Caso Riemann-Liouville. Para el operador recién introducido $\widehat{D}_a^n f := J_a^{m-n} D^m f$ de la Definición 2.1, la situación es diferente: aquí no podemos reemplazar $m = [n]$ por algún $m \in \mathbb{N}$ con $m = [n]$. Esto es evidente al observar el ejemplo simple $f(x) = (x-a)^\beta$. Para tal función tenemos, según el ejemplo

$$\widehat{D}_a^n f(x) = \frac{\Gamma([n] + 1)}{\Gamma([n] + 1 - n)} (x-a)^{[n]-n}$$

pero, para $m \in \mathbb{N}$ con $m < [n]$, obtenemos y por lo tanto también. La clave para la construcción del operador diferencial alternativo que estamos buscando es la siguiente identidad que implica derivados de Riemann-Liouville en por un lado y el operador recién definido por otro lado.

Teorema 3.1

Sea $n \geq 0$ y $m = [n]$. Además suponga que $f \in A^m[a, b]$. Entonces,

$$\widehat{D}_a^n f(x) = D_a^n [f - T_{m-1}[f; a]]$$

Casi en cualquier parte. Aquí, como en la demostración del Teorema 2.20, $T_{m-1}[f; a]$ denota el Polinomio de Taylor de grado $m-1$ para la función f , con centro en a ; en el caso $m = 0$ definimos $T_{m-1}[f; a] = 0$.

Tenga en cuenta que la expresión en el lado derecho de la ecuación existe si D_a^n existe una f y f posee $m-1$ derivadas en a , esta última condición asegurando que El polinomio de Taylor existe. Esta condición es más débil que la condición anterior que $f \in A^m$. Esto se deduce que $f \in A^m$ implica: (a) $f \in C^{m-1}$ y por lo tanto la existencia del polinomio de Taylor requerido y su derivada de Riemann-Liouville, y (b) la existencia de $D_a^m f$ casi en todas partes,

como se puede ver por una aplicación repetida de las ideas usadas en la prueba del Lema 2.12. Por lo tanto, de ahora en adelante, usaremos la última expresión. Se da una formalización de la siguiente manera:

Definición 3.2

Suponga que $n \geq 0$ y que f es tal que $D_a^m[f - T_{m-1}[f; a]]$ existe,

donde $m = [n]$ Entonces definimos la función $D_{*a}^n f$ por

$$D_{*a}^n f := D_a^n[f - T_{m-1}[f; a]]$$

El operador $D_{*a}^n f$ se denomina operador diferencial de Caputo de orden n . En realidad, este concepto ha sido introducido de forma independiente por muchos autores, incluyendo Caputo y Rabotnov que han basado sus desarrollos en el enfoque dado en la Definición 3.1 y por Dzherbashyan y Nersesian quienes han utilizado la Definición 3.2 como punto de partida; otros colaboradores que han tratado con tales operadores desde varios puntos de vista incluyen Gross y Gerasimov, e incluso se puede encontrar en un artículo muy antiguo de Liouville.

Sin embargo, parece que Liouville no vio la diferencia entre este operador y el operador de Riemann-Liouville ya que estaba principalmente interesado en aquellos casos donde los dos operadores coinciden.

Seguimos la convención más común de nombrarlo solo después de Caputo. El lector que esté interesado en un relato histórico detallado debe consultar el reciente artículo de Rossikhin y las referencias citadas en él (2010). La notación que hemos presentado aquí sigue la sugerencia generalmente aceptada de Gorenflo y me gustaría Mainardi.

Una vez más notamos para $n \in \mathbb{N}$ que $m = n$, por lo tanto

$$D_a^n f = D_a^n [f - T_{n-1}[f; a]] = D^n f - D^n (T_{n-1}[f; a]) = D^n f$$

porque $T_{n-1}[f; a]$ es un polinomio de grado $n-1$ que es aniquilado por la clásica operador. Entonces, en este caso, también recuperamos el operador diferencial habitual. En particular, D_a^n es una vez más el operador de identidad. Existen varios documentos y libros donde se mencionan algunas de las propiedades clave del Caputo.

Se han descrito operadores, por ejemplo, Haubold et al. (2009); Jacob & Krageloh (2002); Jin (2021). Por lo general, sin embargo, eran, como se indica explícitamente en el resumen, escrito “de una manera accesible para los científicos aplicados” y “evitar generalidades improductivas y excesivo rigor matemático”.

Parece que las pruebas matemáticamente rigurosas de muchas propiedades importantes no son disponibles en la literatura. Por eso tratamos de darlos aquí.

Demostración (Teorema 3.1)

En el caso $n \in \mathbb{N}$ el enunciado es trivial porque, como visto anteriormente, ambos lados de la ecuación se reducen a $D^n f$. Por lo tanto, solo tenemos considerar el caso $n \notin \mathbb{N}$, lo que implica que $m > n$.

$$\begin{aligned} D_a^n [f - T_{n-1}[f; a]](X) &= D^m J_a^{m-n} [f - T_{m-1}[f; a]](X) = \\ &= \frac{d^m}{dx^m} \int_a^x \frac{(x-t)^{m-n-1}}{\Gamma(m-n)} (f(t) - T_{m-1}[f; a](t)) dt, \dots \quad (3.1) \end{aligned}$$

Se permite una integración parcial de la integral y se obtiene

$$\begin{aligned} & \frac{d^m}{dx^m} \int_a^x \frac{1}{\Gamma(m-n)} (f(t) - T_{m-1}[f; a](t))(x-t)^{m-n-1} dt \\ &= \frac{1}{\Gamma(m-n+1)} [(f(t) - T_{m-1}[f; a](t))(x-t)^{m-n}] \Big|_{t=a}^{t=x} + \\ &+ \frac{1}{\Gamma(m-n+1)} \int_a^x (Df(t) - DT_{m-1}[f; a](t))(x-t)^{m-n} dt \end{aligned}$$

El término fuera de la integral es cero, el primer factor desaparece en el límite inferior, el segundo desaparece en el límite superior. De este modo,

$$J_a^{m-n} [f - T_{m-1}[f; a]] = J_a^{m-n+1} D[f - T_{m-1}[f; a]]$$

Bajo nuestras suposiciones, podemos repetir este proceso un número total de m veces, y esto resulta en

$$J_a^{m-n} [f - T_{m-1}[f; a]] = J_a^{2m-n} D^m [f - T_{m-1}[f; a]] = J_a^m J_a^{m-n} D^m [f - T_{m-1}[f; a]]$$

Observamos que $D^m [f - T_{m-1}[f; a]] \equiv 0$ porque $T_{m-1}[f; a]$ es un polinomio de grado $m-1$.

Por lo tanto, la última identidad se puede simplificar a

$$J_a^{m-n} [f - T_{m-1}[f; a]] = J_a^m J_a^{m-n} D^m f$$

Esto se puede combinar con (3.1) para obtener

$$D_a^n [f - T_{m-1}[f; a]](x) = D^m J_a^m J_a^{m-n} D^m f = J_a^{m-n} D^m f = \widehat{D}_a^n$$

en vista de (1.1).

Teniendo en cuenta la definición del operador Caputo y el Lema 2.21, tenemos obtener una consecuencia directa.

Lema 3.2

Bajo los supuestos del Lema 2.21, tenemos

$$D_{*a}^n f = \frac{1}{\Gamma(-n)} \int_a^x (x-t)^{-n-1} (f(t) - T_{m-1}[f; a](t)) dt$$

Observación 3.2

Como en el caso de los operadores de Riemann-Liouville, vemos que las derivadas de Caputo tampoco son locales.

Se puede obtener otra representación del operador Caputo combinando su definición con el Teorema 2.25:

Lema 3.3

Sea $n > 0$, $m = [n]$ y $f \in C^m[a, b]$. Entonces, para $x \in [a, b]$,

$$D_{*a}^n f(x) = \lim_{N \rightarrow \infty} \frac{1}{h_N^n} \sum_{k=0}^N (-1)^k \binom{n}{k} [f(x - kh_N) - T_{m-1}[f; a](x - kh_N)]$$

$$\text{Con } h_N = \frac{(x-a)}{N}$$

Las representaciones de estos dos Lemas han demostrado ser útiles para trabajar. También se conocen otras representaciones; presentaremos algunos de ellos en la sección. 3.2. Sin embargo, primero continuamos las investigaciones de los aspectos analíticos de los operadores de Caputo. En este contexto, nuestro próximo objetivo es expresar la relación entre el operador de Riemann-Liouville y el operador de Caputo de manera diferente.

Lema 3.4

Sea $n \geq 0$ y $m = [n]$. Supongamos que f es tal que tanto $D_{*a}^n f$ y $D_a^n f$ existen. Entonces,

$$D_{*a}^n f = D_a^n f - \sum_{k=0}^{m-1} \frac{D^k f(a)}{\Gamma(k-n+1)} (x-a)^{k-n}$$

Prueba

En vista de la definición de la derivada de Caputo y el Ejemplo

$$\begin{aligned} D_{*a}^n f &= D_a^n f - \sum_{k=0}^{m-1} \frac{D^k f(a)}{k!} D_a^n [(-a)^k](x) \\ &= D_a^n f(x) - \sum_{k=0}^{m-1} \frac{D^k f(a)}{k!} (x-a)^{k-n} \end{aligned}$$

Una consecuencia inmediata de este Lema es

Lema 3.5

Asumir las hipótesis del Lema 3.4. Entonces

$$D_a^n f(x) = D_{*a}^n f(x)$$

se cumple si y solo si f tiene un m -vez cero en a , es decir, si y solo si $D^k f(a) = 0$ para $k = 0, 1, \dots, m-1$.

También podemos combinar el Lema 4.4 con el Teorema 3.25 para deducir

Lema 3.6

Sea $n > 0$, $m = n$ y $f \in C^m[a, b]$. Entonces, para $x \in (a, b]$,

$$D_{*a}^n f(x) = \lim_{N \rightarrow \infty} \frac{1}{h_N^n} \sum_{k=0}^N (-1)^k \binom{n}{k} f(x - kh_N) - \sum_{k=0}^{m-1} \frac{D^k f(a)}{\Gamma(k-n+1)} (x-a)^{k-n}$$

Con

$$h_N = \frac{(x-a)}{N}.$$

Cuando se trata de la composición de las integrales de Riemann-Liouville y operadores diferenciales de Caputo, encontramos que la derivada de Caputo es también una inversa izquierda de la Integral de Riemann-Liouville:

Teorema 3.7

Si f es continua y $n \geq 0$, entonces

$$D_{*a}^n J_a^n f(x) = f(x)$$

Prueba

Sea $\phi = J_a^n f$. Por el Teorema 2.5, tenemos $D^k \phi(a) = 0$ para $k = 0, 1, \dots, m-1$, y así (en vista del Lema 2.5 y el Teorema 2.14)

$$D_{*a}^n J_a^n f(x) = D_{*a}^n \phi = D_a^n \phi = D_a^n J_a^n f = f$$

Una vez más, encontramos que la derivada de Caputo no es el inverso derecho de la Integral de Riemann-Liouville:

Teorema 3.8

Suponga que $n \geq 0$, $m = [n]$ y $f \in A^m[a, b]$. Entonces

$$J_a^n D_{*a}^n f(x) = f(x) - \sum_{k=0}^{m-1} \frac{D^k f(a)}{k!} (x-a)^k$$

Prueba

Por el Teorema 3.1 y la Definición 3.1, tenemos

$$D_{*a}^n f = \widehat{D}_a^n f = J_a^{m-n} D^m f$$

Así, aplicando el operador J_a^n a ambos lados de esta ecuación y usando el semigrupo propiedad de la integración fraccionaria, obtenemos

$$J_a^n D_{*a}^n f = J_a^n J_a^{m-n} D^m f = J_a^n D^m f$$

Por la versión clásica del teorema de Taylor (cf. Teorema 2.C), tenemos que

$$f(x) = \sum_{k=0}^{m-1} \frac{D^k f(a)}{k!} (x-a)^k + J_a^n D^m f$$

Combinando estas dos ecuaciones obtenemos la afirmación. Un análogo del teorema fraccionario de Taylor sigue inmediatamente.

Teorema 3.9. Expansión de Taylor para derivados de Caputo

Bajo los supuestos del Teorema 3.8,

$$f(x) = \sum_{k=0}^{m-1} \frac{D^k f(a)}{k!} (x-a)^k + J_a^n D_{*a}^n f(x)$$

Las relaciones mostradas en el Teorema 4.8 y el Corolario 4.9 tienen implicaciones importantes cuando se trata de la solución de ecuaciones diferenciales que involucren los dos tipos de operadores diferenciales. Específicamente, suponga que h es una función dada con la propiedad que existe alguna función g tal que $\bar{h} = D_a^n g$. Entonces, la solución de la Ecuación diferencial de Riemann-Liouville

$$D_a^n f = h$$

$$f(x) = g(x) + \sum_{j=1}^{[n]} c_j (x-a)^{n-j}$$

con constantes arbitrarias c_j . Esto sigue por las mismas técnicas que uno emplearía para ecuaciones diferenciales de orden entero porque la ecuación es lineal y no homogénea. Por construcción, g es una solución de la ecuación no homogénea, y por el ejemplo 2.4 cada uno de los términos de la suma resuelve la homogeneidad correspondiente ecuación.

De manera similar, si es una función dada con la propiedad de que $h_* = D_{*a}^n g$ y si quiero resolver

$$D_{*a}^n f_* = h_*$$

$$f_*(x) = g_*(x) + \sum_{j=1}^{[n]} c_j^* (x-a)^{[n]-j}$$

de nuevo con constantes arbitrarias c_j^* . Por lo tanto, para obtener una solución única, es más natural prescribir los valores $f_*(a), Df_*(a), \dots, D^{[n]-1}f_*(a)$ en el ajuste de Caputo, mientras que en el caso de Riemann-Liouville, uno preferiría prescribir derivadas fraccionarias de f en a . Esto será explorado de manera más detallada en los capítulos siguientes. Por el momento señalamos que

se suele preferir la versión de Caputo, cuando se describen modelos físicos porque la interpretación física de los datos prescritos es clara y, por lo tanto, en general es posible proporcionar estos datos, p.ej. por medidas adecuadas. Esto no es cierto para las condiciones iniciales de orden fraccionario, necesarios para el entorno de Riemann-Liouville.

Por ejemplo, en aplicaciones como el modelado de materiales viscoelásticos en mecánica, $f(x)$ es típicamente un desplazamiento en el tiempo x , por lo que $f'(x)$ y $f''(x)$ serían la velocidad correspondiente y aceleración, respectivamente—cantidades que son bien entendidas y fáciles de medir. Por otro lado, a pesar de las explicaciones intentadas recientemente, una derivada fraccionaria de un desplazamiento sigue siendo un objeto cuya naturaleza física es clara, por lo que no hay métodos de medición disponibles para tal cantidad.

Aparte de esta razón (motivada principalmente por argumentos en relación con aplicaciones) en realidad hay otras razones que vienen del lado “puro” de matemáticas por preferir la derivada de Caputo sobre el operador de Riemann-Liouville, pero no nos detendremos aquí en este tema. Más bien, continuaremos por establecer otra representación para la derivada de Caputo de una función bajo un supuesto natural. Con este fin requerimos la siguiente definición que es un especial caso de un concepto establecido.

Definición 3.3

Sea $n > 0$ y sea v una función completa con el desarrollo en serie de potencias $v(x) = \sum_{k=0}^{\infty} c_k x^k$. Entonces, el operador $\mathcal{D}^n$ que mapea esta función v a la función con

$$\mathcal{D}^n v := \sum_{k=1}^{\infty} c_k \frac{\Gamma(kn+1)}{\Gamma(kn+1-n)} x^{k-1}$$

se llama el operador Gel'fond-Leont'ev de orden n .

Las derivadas de Caputo de ciertas funciones se pueden expresar con la ayuda de estos operadores de una manera muy conveniente (Mainardi & Gorenflo, 2000, p. 139)

Teorema 3.10

Sea $n > 0$ y sea v una función completa con la serie de potencias expansión $v(x) = \sum_{k=0}^{\infty} c_k x^k$. Además, sea $f(x) := v(x^n)$ para $x \geq 0$. Entonces,

$$D_{*0}^n f(x) = \mathcal{D}^n v(x^n)$$

Prueba

Este resultado sigue usando un cálculo directo usando las definiciones de la función f y el operador diferencial de Caputo, la expansión en serie de potencias de v y el hecho de que, debido a que v es entero, podemos intercambiar el operador de serie infinita y el operador diferencial de Caputo (es decir, podemos aplicar el operador de Caputo a la serie de forma de término a término)

El siguiente resultado establece otra diferencia significativa entre las Derivadas de Riemann-Liouville y Caputo. Una comparación con, por ejemplo, el Ejemplo 3.4 para $f(x) = 1$ y $n > 0, n \notin \mathbb{N}$, revela que no se nos permite reemplazar D_{*a}^n por D_a^n

Lema 3.11

Sea $n > 0$, $n \notin N$ y $m = [n]$. Además, suponga $f \in C^m[a, b]$. Entonces $D_{*a}^n f \in C[a, b]$ y $D_{*a}^n f(a) = 0$.

Prueba

Por definición y el Teorema 4.1, $D_{*a}^n f = J_a^{m-n} D^m f$. El resultado se sigue del Teorema 2.5 porque se supone que $D^m f$ es continua. Podemos relajar ligeramente las condiciones en f .

Lema 3.12

Sea $n > 0$, $n \notin N$ y $m = [n]$. Además, sea $f \in A^m[a, b]$ y suponga que $D_{*a}^{\hat{n}} f \in [a, b]$ para algún $\hat{n} \in (n, m)$. Entonces, $D_{*a}^n f \in C[a, b]$ y $D_{*a}^n f(a) = 0$.

Prueba.

Por definición y los Teoremas 2.1 y 2.2,

$$D_{*a}^n f = J_a^{m-n} D^m f = J_a^{\hat{n}-n} J_a^{m-\hat{n}} D^m f = J_a^{\hat{n}-n} D_{*a}^{\hat{n}} f$$

Por lo tanto, la afirmación se sigue en virtud del Teorema 2.5.

Observación 3.3

Para evaluar las consecuencias de estos dos Lemas, señalemos considerando el siguiente hecho. Supongamos, por ejemplo, las hipótesis del Lema 3.11 y además que $n > 2$. Entonces el Lema afirma que todas las derivadas $D_{*a}^n f$, $0 < n < 3$, son continuos, por lo tanto, en particular, $D_{*a}^\ell f$ y $D_{*a}^{\ell+1} f$ son continuas siempre que $0 < \ell < 2$.

Sin embargo, esto no significa que $D_{*a}^{\ell}f \in C^1[a, b]$ para estos ℓ . Para contraejemplo, nos referimos al ejercicio 3.1. Como consecuencia de esta observación, obtenemos que no podemos deducir la identidad $DD_{*a}^{\ell}f = D_{*a}^{\ell+1}f$ para ser verdad bajo los supuestos de nuestro Lema porque la función en el lado derecho es continua mientras que la del lado izquierdo no necesita tener esta propiedad. Por eso encontramos que los operadores diferenciales de Caputo no forman un semigrupo en general.

En este contexto es importante la siguiente observación.

Lema 3.13

Sea $f \in C^k[a, b]$ para algún $a < b$ y algún $k \in \mathbb{N}$. Además, sea $n, \varepsilon > 0$ ser tal que existe algún $\ell \in \mathbb{N}$ con $\ell \leq k$ y $n, n + \varepsilon \in [\ell - 1, \ell]$. Entonces

$$D_{*a}^{\ell} D_{*a}^n f = D_{*a}^{n+\varepsilon} f$$

Observación 3.4

Es necesario hacer dos comentarios con respecto a este Lema:

No se puede esperar que tal resultado se cumpla en general si las derivadas de Riemann-Liouville se utilizaron en lugar de las derivadas de Caputo. Como ejemplo, considere la función f con $f(x) = 1$ y sea $a = 0$, $n = 1$ y $\varepsilon = 1/2$. si fuéramos a usar las derivadas de Riemann-Liouville, el lado izquierdo sería $(D_0^{1/2} f')(x) = D_0^{1/2} 0 = 0$, mientras que el lado derecho es $D_0^{3/2} f(x) = D^3 J_a^{1/2} f(x) = \frac{1}{\Gamma(-1/2)} x^{-3/2}$

La condición que requiere la existencia del número con las propiedades mencionadas en el Lema es esencial. Para ver lo que puede suceder sin él, considere el ejemplo $n = \varepsilon = 7/10$ (es decir, $7/10 = n < n + \varepsilon = 7/5$), $a = 0$ y $f(x) = x$. Entonces, el lado derecho es $D_{*0}^{7/5} f(x) = \left(J_0^{3/5} f'' \right)(x) = J_0^{3/5} 0 = 0$, pero como

$$(D_{*0}^{7/10} f)(x) = \frac{1}{\Gamma(13/10)} x^{13/10}$$

el lado izquierdo toma el valor

$$D_{*0}^{7/10} (D_{*0}^{7/10} f)(x) = \frac{1}{\Gamma(3/5)} x^{-2/5}$$

Demostración del Lema 3.13

El enunciado es trivial en el caso $n = l - 1$ y $n + \varepsilon = \ell$, así que solo tratamos las otras situaciones explícitamente. Luego observamos primero que nuestras suposiciones implican $0 < \varepsilon < 1$. Así, por el Lema 4.5 encontramos que

$$D_{*a}^{\varepsilon} z = D_a^{\varepsilon} z$$

siempre que $z(a) = 0$. Consideramos tres casos:

$n + \varepsilon \in \mathbb{N}$: En este caso tenemos que $n = n + \varepsilon$ y por lo tanto $n - n = \varepsilon$. Además,

por Lema 3.11, $D_{*a}^n f(a) = 0$. Así

$$\begin{aligned} D_{*a}^{\varepsilon} D_a^n f &= D_a^{\varepsilon} D_{*a}^n f = D_a^{\varepsilon} J_a^{[n]-n} D^{[n]} f \\ &= D_a^{\varepsilon} J_a^{\varepsilon-n} D^{[n]} f = D^{[n]} f = D^{n+\varepsilon} f = D_{*a}^{n+\varepsilon} f \end{aligned}$$

$n \in \mathbb{N}$: Aquí tenemos, usando el Teorema 3.1,

$$D_{*a}^{\varepsilon} D_{*a}^n f = D_a^{\varepsilon} D^n f = J_a^{1+\varepsilon} D^{n+1} f = D_{*a}^{n+\varepsilon} f$$

De lo contrario, tenemos $n = n + \varepsilon$, y por lo tanto encontramos por argumentos similares que

$$\begin{aligned}
D_{*a}^\varepsilon D_{*a}^n f &= D_a^\varepsilon D_{*a}^n f = D_a^\varepsilon J_a^{[n]-n} D^{[n]} f \\
&= D^1 J_a^{1-\varepsilon} J_a^{[n]-n} D^{[n]} f = D^1 J_a^1 J_a^{[n+\varepsilon]-(n+\varepsilon)} D^{[n+\varepsilon]} f \\
&= D_{*a}^{n+\varepsilon} f
\end{aligned}$$

El resultado anterior se ha ocupado de la concatenación de dos diferenciales Caputo operadores. En algunos casos, sin embargo, también puede ser útil para concatenar un Caputo operador con un operador diferencial de tipo Riemann-Liouville:

Teorema 3.14

Sea $f \in C^\mu[a, b]$ para algún $\mu \in \mathbb{N}$. Además, sea $n \in [0, \mu]$. Entonces,

$$D_a^{\mu-n} D_{*a}^n f = D^\mu f$$

Observe que el operador D^μ que aparece en el lado derecho de la afirmación es un operador diferencial clásico (de orden entero).

Prueba

Si n es un número entero, ambos operadores diferenciales del lado izquierdo se reducen a operadores de orden entero y por lo tanto obtenemos el resultado deseado mediante una aplicación de la definición de los operadores iterados, a saber. Definición 2.1 (c).

Si n no es un número entero, podemos invocar el Teorema 4.1 para concluir que

$$D_{*a}^n f = \widehat{D}_{*a}^n f = J_a^{[n]-n} D^{[n]} f$$

Combinando esto con la definición de la derivada de Riemann-Liouville y usando

la propiedad del semigrupo de la integración fraccionaria y la ec. (2.1) encontramos

$$\begin{aligned} D_a^{\mu-n} D_{*a}^n f &= D^{\mu-[n]+1} J_a^{n+1-[n]} J_a^{[n]-n} D^{[n]} f = D^{\mu-[n]+1} J_a^1 D^{[n]} f \\ &= D^{\mu-[n]} D^{[n]} f = D^n f \end{aligned}$$

En realidad, es posible explorar las propiedades de suavidad de bajo el supuesto de suavidad en f con más detalle que en el Lema 3.11:

Teorema 3.15

Si $f \in C^\mu[a, b]$, para algún $\mu \in \mathbb{N}$ y $0 < n < \mu$ entonces

$$D_a^n f(x) = \sum_{\ell=1}^{u-[n]-1} \frac{f^{(\ell+[n])}(a)}{\Gamma([n]-n+\ell+1)} (x-a)^{[n]-n+\ell} + g(x)$$

con alguna función $g \in C^{\mu-[n]}[a, b]$. Además, la $(\mu-[n])$ ésima derivada de g satisface una condición de Lipschitz de orden $[n]-n$.

Prueba

Esta es una consecuencia directa de la definición del operador diferencial de Caputo y los teoremas 3.1 y 2.5.

Las principales reglas de cálculo para la derivada de Caputo son similares, pero no idénticas, a los de la derivada de Riemann-Liouville.

Teorema 3.16

Sean $f_1, f_2: [a, b] \rightarrow \mathbb{R}$ tales que $D_{*a}^n f_1$ y $D_{*a}^n f_2$ existen casi en todas partes y sea $c_1, c_2 \in \mathbb{R}$. Entonces, $D_{*a}^n(c_1 f_1 + c_2 f_2)$ existe en casi todas partes, y

$$D_{*a}^n(c_1 f_1 + c_2 f_2) = c_1 D_{*a}^n f_1 + c_2 D_{*a}^n f_2$$

Prueba

Esta propiedad de linealidad del operador diferencial fraccionario es una consecuencia de la definición de D_{*a}^n . Para la fórmula de Leibniz, sólo enunciamos explícitamente el caso $0 < n < 1$.

Teorema 3.17. Fórmula de Leibniz para los operadores de Caputo

Sea $0 < n < 1$, y suponga que f y g son analíticas en $(a-h, a+h)$. Entonces,

$$\begin{aligned} D_{*a}^n[f g](x) &= \frac{(x-a)^{-n}}{\Gamma(1-n)} g(a)(f(x) - f(a))(D_{*a}^n g(x))f(x) \\ &\quad + \sum_{k=1}^{\infty} \binom{n}{k} (J_a^{k-n} g(x)) D_{*a}^k f(x) \end{aligned}$$

Prueba

Aplicamos la definición de la derivada de Caputo y encontramos

$$D_{*a}^n[f g] = D_{*a}^n[f g - f(a)g(a)] = D_{*a}^n[f g] - f(a)g(a)D_{*a}^n[1]$$

Luego usamos la fórmula de Leibniz para las derivadas de Riemann-Liouville y encontramos

$$D_{*a}^n[f g] = f(D_{*a}^n g) + \sum_{k=1}^{\infty} \binom{n}{k} (D_{*a}^k f) (J_a^{k-n} g) - f(a)g(a)D_{*a}^n[1]$$

Ahora sumamos y restamos $f \cdot g(a) (D_{*a}^n[1])$ y reorganizar para obtener

$$\begin{aligned}
D_a^n[f g] &= f(D_a^n[g - g(a)] + \sum_{k=1}^{\infty} \binom{n}{k} (D_a^k f)(J_a^{k-n} g) + g(a)(f - f(a))D_a^n[1] \\
&= f \times (D_a^n g) + \sum_{k=1}^{\infty} \binom{n}{k} (D_a^k f)(J_a^{k-n} g) + g(a)(f - f(a)) \times D_a^n[1]
\end{aligned}$$

donde hemos utilizado el hecho de que, para $k \in \mathbb{N}$, $D_a^k = D^k = D_{*a}^k$. Para finalmente completar la prueba solo resta usar la expresión explícita para $D_a^n[1]$ del Ejemplo 2.4.

Observación 3.5

Para la fórmula de Fa`a di Bruno (la regla de la cadena) para los operadores de Caputo tenemos puede combinar la ec. (2.6), es decir, la regla correspondiente para los operadores de Riemann-Liouville, y Lema 3.4. Esto produce que, una vez más bajo suposiciones adecuadas sobre las funciones f y g que no especificaremos explícitamente, tiene la forma

$$\begin{aligned}
&D_a^n[f(g(\cdot))](x) \\
&= \sum_{k=1}^{\infty} \binom{n}{k} \frac{k! (x-a)^{k-n}}{\Gamma(k-n+1)} \sum_{\ell=1}^{\infty} (D^{\ell} f)(g(x)) \sum_{(a_1, a_2, \dots, a_k) \in A_{K, \ell}} \prod_{r=1}^k \frac{1}{a_r!} \left(\frac{D^r g(x)}{r!} \right)^{a_r} \\
&\quad + \frac{(x-a)^{-n}}{\Gamma(1-n)} f(g(x)) + \sum_{k=0}^{|n|-1} \frac{D^k[f(g(\cdot))](a)}{\Gamma(k-n+1)} (x-a)^{k-n}
\end{aligned}$$

donde, como en el caso de Riemann-Liouville discutido en el Teorema 2.19,

$(a_1, a_2, \dots, a_k) \in A_{K, \ell}$ significa que

$$a_1, a_2, \dots, a_k \in \mathbb{N}_0 \quad \sum_{r=1}^k r a_r = k \text{ y } \sum_{r=1}^k a_r = \ell.$$

Representaciones no clásicas de operadores de Caputo

En la sección anterior, hemos desarrollado una serie de representaciones diferentes para Operadores diferenciales de Caputo bajo varios supuestos sobre la función a diferenciar; véase, por ejemplo, la Definición 3.2, la Definición 3.1 en combinación con el Teorema 3.1, Lema 3.2, Lema 3.3 o Lema 3.4. Esencialmente, todas estas representaciones consisten en una combinación de una integral de convolución y algún tipo de operador diferencial o el límite de una versión discreta de este.

Este enfoque revela inmediatamente la no localidad del operador de Caputo y proporciona un enfoque natural para manejar esta propiedad. Sin embargo, en algunas aplicaciones ha resultado que otras formas de expresar los que no son de localidad son más útiles. Dos de estas representaciones alternativas vienen desarrollado recientemente de forma independiente por Yuan y Agrawal (2002) y Singh y Chatterjee (2006).

Nuestro tratamiento de estos dos métodos estrechamente relacionados se basa en las generalizaciones proporcionadas y analizadas.

Primero recordamos los detalles del método propuesto por Yuan y Agrawal (2002). Solo han discutido el caso $0 < n < 1$. Se ha hecho una extensión a $1 < n < 2$ proporcionada por Trinks y Ruge (2002). No impondremos ninguna restricción de este tipo al tamaño de n . Sin embargo, nuestras técnicas requieren que a lo largo de esta sección.

Dado que esto solo excluye los casos que no son verdaderamente fraccionarios de todos modos, esto no es una limitación sustancial. Se puede ver fácilmente que nuestro enfoque se redu-

ce al original esquema de Yuan y Agrawal (2002), si $0 < n < 1$.

El enfoque se adapta a las funciones $f \in C^n[a, b]$. En vista del Teorema 3.1 esto:

La característica nos permite usar la representación de la Definición 3.1 para la derivada de Caputo.

Luego definimos una función bivariada auxiliar $\phi: (0, \infty) \times [a, b] \rightarrow \mathbb{R}$ por

$$\phi(w, x) = (-1)^{[a]} \frac{2 \sin \pi n}{\pi} w^{2n-2[n]+1} \int_a^x f^{([n])}(\tau) e^{-(x-\tau)w^2} d\tau, \dots, (3.2)$$

Con esta notación obtenemos la siguiente generalización de un resultado presentado.

Teorema 3.18

Bajo los supuestos anteriores,

$$D_a^n f(x) = \int_0^\infty \phi(w, x) dw, \dots, (3.3)$$

además, para $w > 0$ fijo, la función $\phi(w, \cdot)$ satisface la ecuación diferencial

$$\frac{\partial}{\partial x} \phi(w, x) = -w^2 \phi(w, x) + (-1)^{[a]} \frac{2 \sin \pi n}{\pi} w^{2n-2[n]+1} f^{([n])}(x), \dots (3.4)$$

sujeto a la condición inicial $\phi(w, a) = 0$.

Observe que la ecuación diferencial (3.4) es efectivamente una ecuación diferencial ordinaria, ecuación que ya asumimos que w es un parámetro fijo. Además, es una ecuación diferencial de orden 1 y por tanto de tipo clásico (no fraccionario en sentido

estricto). En además notamos que la ecuación diferencial es lineal y no homogénea y que tiene coeficientes constantes. Por lo tanto, es una cuestión simple de calcular la solución del problema de valor inicial explícitamente y, por supuesto, este cálculo reproduce la representación (3.2).

Prueba

Teniendo en cuenta la definición de la función Gamma (Definición 1.2), Teorema D.3 y la identidad obvia

$$\sin \pi(n - [n] + 1) = (-1)^{|a|} \sin \pi n, n \notin \mathbb{N}$$

obtenemos que

$$\begin{aligned} D_{*a}^n f(x) &= \frac{1}{\Gamma([n] - n)} \int_a^x (x - \tau)^{[n] - n - 1} f^{([n])}(\tau) d\tau, \\ &= \frac{1}{\Gamma(n - [n] + 1) \Gamma([n] - n)} \times \int_a^x \int_0^\infty e^{-z} z^{n - [n]} dz (x - \tau)^{[n] - n - 1} f^{([n])}(\tau) d\tau, \\ &= \frac{\sin \pi(n - [n] + 1)}{\pi} \int_a^x \int_0^\infty e^{-z} \left(\frac{z}{x - \tau}\right)^{n - [n] + 1} \frac{1}{z} f^{([n])}(\tau) dz d\tau \\ D_{*a}^n f(x) &= (-1)^{[n]} \frac{\sin \pi n}{\pi} \int_a^x \int_0^\infty e^{-z} \left(\frac{z}{x - \tau}\right)^{n - [n] + 1} \frac{1}{z} f^{([n])}(\tau) dz d\tau \end{aligned}$$

Ahora podemos aplicar la sustitución $z = (x - \tau)w^2$ en la integral interna y notar que El Teorema de Fubini nos permite intercambiar el orden de las integraciones ya que $f(n)$ se supone que es continuo. Esto produce

$$\begin{aligned} D_{*a}^n f(x) &= (-1)^{[n]} \frac{2 \sin \pi n}{\pi} \int_a^x \int_0^\infty e^{-(x-\tau)w^2} w^{2n-2[n]+1} \frac{1}{z} f^{([n])}(\tau) dw d\tau \\ D_{*a}^n f(x) &= \int_0^\infty (-1)^{[n]} \frac{2 \sin \pi n}{\pi} x \int_a^\infty e^{-(x-\tau)w^2} f^{([n])}(\tau) dw d\tau \end{aligned}$$

y, recordando la definición (3.2) de ϕ , deducimos (3.3).

A continuación, diferenciamos la definición (3.2) de con respecto a x . las clásicas reglas para la diferenciación de integrales de parámetros con respecto al parámetro inmediatamente dado (3.4). Finalmente, el hecho de que $\phi(w, a) = 0$ para $w > 0$ también es

una consecuencia de (3.2) porque el integrando de la integral del lado derecho de (3.2) es continua.

El enfoque propuesto por Chatterjee e investigado más a fondo es también basado en expresar la derivada fraccionaria de la función dada f en la forma de una integral sobre $(0, \infty)$ cuyo integrando se puede calcular como la solución de un problema de valor inicial de primer orden. Específicamente, se basa en el siguiente análogo del Teorema 3.18. El resultado esencialmente establece que podemos reemplazar el integrando por una función ϕ^* que se puede caracterizar como la solución de un primer orden diferente del Problema de valor inicial (Chatterjee, 2005).

Teorema 3.19

Sea $f \in C^n[a, b]$. Además, para $w > 0$ fijo, sea $\phi^*(w, \cdot)$ la solución de la ecuación diferencial

$$\frac{\partial}{\partial x} \phi^*(w, x) = -w^{1/(n-[n]-1)} \phi^*(w, x) + \frac{(-1)^{[a]} \sin \pi n}{\pi(n-[n]-1)} f^{([n])}(x), \dots, (3.5)$$

Sujeto a la condición inicial $\phi^*(w, a) = 0$. Entonces, tenemos

$$\phi^*(w, x) = \frac{(-1)^{[a]} \sin \pi n}{\pi(n-[n]-1)} \int_0^x f^{([n])}(\tau) \exp(-(x-\tau)w^{1/(n-[n]-1)}) d\tau, \dots (3.6)$$

y

$$D_{*a}^n f(x) = \int_0^\infty \phi^*(w, x) dw, \dots, (3.7)$$

Prueba

La demostración de este resultado es casi idéntica a la demostración del Teorema 3.18; uno solo necesita reemplazar la sustitución $z = (x - \tau)w^2$ por $z = (x - \tau)w^{1/(n-[n]-1)}$ y usamos la ecuación funcional de la función Gamma, $\Gamma(u) = \Gamma(u+1)$. Dejamos los detalles al lector.

En muchas aplicaciones de estas representaciones es importante tener algún conocimiento adicional sobre el comportamiento de la función ϕ en la ec. (3.2) o la función ϕ^* en la ec. (3.6), respectivamente. El más importante de estas propiedades se resumen en los siguientes teoremas. Aquí, el símbolo $\alpha(v) \sim \beta(v)$ significa que existen dos constantes estrictamente positivas A y B tales que $|\alpha(v)/\beta(v)| \in [A, B]$ cuando v tiende al límite indicado. Empezamos con la función que surge en la representación original de Yuan-Agrawal que habíamos dado en nuestro Teorema 3.18.

Teorema 3.20

Sea $x \in (a, b)$ fijo y $0 < n \in \mathbb{N}$, y supongamos que existe algún $C > 0$ tal que $|f^{(n)}(\bar{x})| > C$ para todo $\bar{x} \in [a, b]$.

(a) La función $\phi(\cdot, x)$ definida en (3.2) se comporta como

$$\phi(w, x) \sim w^{2n-2[n]+1} \text{ como } w \rightarrow 0, \dots, (3.8)$$

(b) Además

$$\phi(w, x) \sim w^{2n-2[n]+1} \text{ como } w \rightarrow \infty, \dots, (3.9)$$

(c) Tenemos

$$\phi(\cdot, x) \in C^\infty(0, \infty)$$

Observación 3.6

La condición de que $f(n)$ esté acotada en cero es un requisito técnico, condición requerida para mantener la prueba simple y mantener el resultado válido para todo $x \in (a, b)$. Usando técnicas más complicadas, uno podría mostrar que el mismo el comporta-

miento asintótico está presente para casi todos los $x \in (a, b)$ bajo condiciones sustancialmente más débiles condiciones. Por lo tanto, está justificado decir que el comportamiento asintótico descrito en El teorema 3.20 es el comportamiento que uno puede esperar razonablemente de la función a menos que la función dada f sea de una naturaleza muy excepcional (Hoog & Weiss, 1973).

$$\begin{aligned} & \int_0^x f^{(n)}(\tau) e^{-(x-\tau)w^2} d\tau \\ &= f^{(n)}(\tau) e^{-(x-\tau)w^2} \Big|_{\tau=a}^x - w^2 \int_0^x f^{(n)}(\tau) e^{-(x-\tau)w^2} d\tau \\ &= f^{(n)}(x) - f^{(n)}(a) e^{-xw^2} - w^2 \int_0^x f^{(n)}(\tau) e^{-(x-\tau)w^2} d\tau \end{aligned}$$

Dado que x es fijo, la integral más a la derecha obviamente permanece acotada como $w \rightarrow 0$, y por lo tanto concluimos

$$\lim_{w \rightarrow 0} \int_0^x f^{(n)}(\tau) e^{-(x-\tau)w^2} d\tau = f^{(n)}(x) - f^{(n)}(0), \dots, (3.10)$$

Insertando esta relación en la definición (3.2) de ϕ obtenemos la primera afirmación.

$$\begin{aligned} & w^2 \int_0^x f^{(n)}(\tau) e^{-(x-\tau)w^2} d\tau \\ &= w^2 \int_0^{x-w^{-1}} f^{(n)}(\tau) e^{-(x-\tau)w^2} d\tau \\ &+ w^2 \int_{x-w^{-1}}^x f^{(n)}(\tau) e^{-(x-\tau)w^2} d\tau \\ &= f^{(n)}(\xi_1) w^2 \int_0^{x-w^{-1}} e^{-(x-\tau)w^2} d\tau + f^{(n)}(\xi_2) w^2 \int_{x-w^{-1}}^x e^{-(x-\tau)w^2} d\tau \\ &= f^{(n)}(\xi_1) (e^{-w} - e^{-w^2}) + f^{(n)}(\xi_2) (1 - e^{-w}) \end{aligned}$$

con $\xi_1 \in [a, x - w - 1]$ y $\xi_2 \in [x - w - 1, x]$ debido al Teorema del valor medio. Ahora, como $w \rightarrow \infty$, $f^{(n)}(\xi_1)$ permanece acotada mientras que $e^{-w} - e^{-w^2} \rightarrow 0$.

Así desaparece el primer sumando del lado derecho. Para el segundo sumando

Tenemos $1 - e^{-w} \rightarrow 1$ y $f^{([n])}(\xi_2) \rightarrow f^{([n])}(\xi_1)$ porque $\xi_2 \in [x - w - 1, x]$.
Por lo tanto, nosotros concluir

$$\lim_{w \rightarrow 0} w^2 \int_0^x f^{([n])}(\tau) e^{-(x-\tau)w^2} d\tau = f^{([n])}(x)$$

Insertando esta relación en la definición de ϕ , llegamos a

$$\phi(w, x) = (-1)^{[a]} \frac{2 \sin \pi n}{\pi} w^{2n2-[n]-1} [f^{([n])}(x) + \sigma(1)], \dots (3.11)$$

lo que completa la prueba de (b).

Finalmente, la parte (c) se sigue directamente de la definición de que habíamos dado en ec. (3.2).

También podemos proporcionar un resultado correspondiente para la función ϕ^* utilizada en Representación de Chatterjee (Teorema 3.19). El comportamiento de esta función ϕ^* , que se define en la ec. (4.6), se puede describir de la siguiente manera.

Teorema 3.21

Sea $x \in (a, b)$ fijo y $0 < n \notin N$, y supongamos que existe algún $C > 0$ tal que $|f^{([n])}(\bar{x})| > C$ para todo $\bar{x} \in [0, X]$.

(a) La función $\phi^*(\cdot, x)$ descrita en la ec. (3.6) se comporta como

$$\phi^*(\cdot, x) \sim 1 \text{ como } w \rightarrow 0, \dots, (3.12)$$

(b) Además

$$\phi^*(\cdot, x) \sim w^{1/(n-[n]-1)} \text{ como } w \rightarrow \infty, \dots, (3.13)$$

(c) Tenemos

$$\phi^*(\cdot, x) \in C^\infty(0, \infty)$$

La prueba procede de la misma manera que la prueba del teorema 3.20.

Los teoremas 3.20 y 3.21 nos permiten comparar las propiedades analíticas de la función utilizada por Yuan y Agrawal (2002) y la función ϕ^* propuesta por Chatterjee.

En primer lugar, notamos que ambas funciones poseen infinitas derivadas (en el sentido clásico) con respecto a la primera variable. Sin embargo, hay diferencias significativas en el comportamiento asintótico de las funciones ya que la primera variable tiende a cualquier extremo del intervalo $(0, \infty)$ sobre el cual las funciones deben integrarse en para calcular la derivada de Caputo $D_{*a}^n f$.

En concreto, para $w \rightarrow 0$ la función $\phi(w, x)$ presenta un comportamiento asintótico de la forma $w^{2n-2[n]-1}$ según el Teorema 3.20 (a). El exponente de donde es siempre estrictamente entre -1 y $+1$. Esto afirma la integralidad de $\phi(w, x)$ con respecto a w cerca de $w = 0$ al menos en el sentido impropio.

Sin embargo, podemos esperar un comportamiento suave cerca de este punto final del intervalo de integración sólo si el exponente es un número entero, y este es el caso sí y solo si, $n = k+1/2$ con algún $k \in \mathbb{N}_0$. Para todos los demás valores de n el comportamiento es menos regular.

Esta irregularidad debe tomarse en cuenta cuidadosamente cuando se intenta utilizar este enfoque en un algoritmo numérico. El teorema 3.21 (a) demuestra que la función ϕ^* es más fácil de manejar en este respecto ya que aquí siempre tenemos que $\phi^*(w, x)$ permanece acotado para distinto de cero constantes de arriba y abajo.

El comportamiento de los integrandos $\phi(w, x)$ y $\phi^*(w, x)$ para $w \rightarrow \infty$ también exhibe diferencias sustanciales. Como se muestra en el Teorema 3.20 (b), el integrando Yuan-Agrawal $\phi(w, x)$ se comporta como $w^{2n-2[n]-1}$. El exponente de donde siempre está contenido en el intervalo $(-3, -1)$. Esto es lo suficientemente rápido para asegurarse de que el incorrecto integral existe. Por otro lado, de acuerdo con el Teorema 3.21 (b), el Chatterjee integrando $\phi^*(w, x)$ se comporta de una manera que depende de n en un algo más complicada moda: si $n = k + \varepsilon$ con algún $k \in \mathbb{N}_0$ y $0 < \varepsilon < 1$ entonces el exponente en la pregunta es $-1/\varepsilon$. Siempre es menor que -1 y, por lo tanto, la integral impropia converge (Singh & Chatterjee, 2006).

A este respecto, no tenemos ninguna diferencia con el método Yuan-Agrawal. Sin embargo, si ε está cerca de 0, entonces el exponente, por supuesto, sigue siendo negativo, pero puede ser arbitrariamente grande en módulo, lo que lleva a una descomposición mucho más rápida (pero aún algebraica) de la integrando. En particular, en contraste con el método Yuan-Agrawal no hay menor ligado al exponente cuando n pasa por todos los números admisibles. para numérico trabajo, es preferible un integrando que decaiga rápidamente, por lo que al menos para $n = k + \varepsilon$ con $k \in \mathbb{N}$ y ε cerca de 0, el enfoque a través del Teorema 3.19 tiene algunas ventajas sobre el camino vía el Teorema 3.18.

Sin embargo, cabe señalar que, desde el punto de vista de teoría de aproximación, un integrando ideal (es decir, un integrando que puede ser manejado muy bien por un algoritmo numérico) decaería exponencialmente como $w \rightarrow \infty$, es decir, mucho más rápido de lo que podemos esperar incluso para ϕ^* .

04 Capítulo

FUNCIONES DE MITTAG-LEFFLER

Escucha a tu niño interno que te enseñara un afecto cordial para todos, como gobernar a una familia de casa con amor de padre antes que, con la severidad de un amo, la idea de una vida conforme a la razón natural y de una gravedad sin afectación, de tener cuidado con acertar con el gusto de un amigo.

Burseca

Antes de que podamos llegar al núcleo de estas reflexiones, es decir, a la discusión de ecuaciones diferenciales fraccionarias “EDFs”, necesitamos introducir dos clases de funciones, una de las cuales puede considerarse como un caso especial del otro e investigar sus propiedades básicas.

Estas funciones resultarán de fundamental importancia en nuestras reflexiones, y se utilizarán en muchos lugares a lo largo de estas reflexiones. Nosotros comenzaremos con el más restrictivo de los dos conceptos.

Definición 4.1

Sea $n > 0$. La función E_n definida por

$$E_n(z) = \sum_{j=0}^{\infty} \frac{z^j}{\Gamma(jn + 1)}$$

siempre que la serie converge se llama función Mittag-Leffler de orden n . Esta función ha sido introducida por Mittag-Leffler. Inmediatamente hagamos cuenta de

$$E_1(z) = \sum_{j=0}^{\infty} \frac{z^j}{\Gamma(j+1)} = E_n(z) = \sum_{j=0}^{\infty} \frac{z^j}{j!} = \exp(z), \dots (4.1)$$

es simplemente la conocida función exponencial.

La clase más general de funciones se define como sigue.

Definición 4.2

Sea $n_1, n_2 > 0$. La función $E_{n_1, n_2}(z)$ es definida por

$$E_{n_1, n_2}(z) := \sum_{j=0}^{\infty} \frac{z^j}{\Gamma(jn_1, n_2 + 1)}$$

Formalmente cada vez que la serie converge se llama la función de dos parámetros Mittag-Leffler con parámetros $n_1, y n_2$.

Observación 4.1

Es evidente que las funciones de Mittag-Leffler de un parámetro pueden ser definido en términos de sus contrapartes de dos parámetros a través de la relación $E_n(z) = E_{n,1}(z)$.

El nombre de estas últimas funciones en honor a Mittag-Leffler se debe al hecho de que son una generalización muy simple y obvia de las funciones introducidas originalmente, por él y recordado aquí en la Definición 4.1. En aras de la corrección histórica, sin embargo, se debe mencionar que las funciones de dos parámetros fueron en realidad primero discutido por Wiman poco después de la publicación del trabajo original de Mittag-Leffler.

Observación 4.2

La mayoría de los resultados de la función Mittag-Leffler de un parámetro dada siguiente seguirá siendo válida si la restricción $n > 0$ se reemplaza por $n \in \mathbb{C}$ con $\text{Re } n > 0$.

De manera similar, las condiciones $n_1, n_2 > 0$, para la función Mittag-Leffler de dos parámetros puede relajarse a $n_1, n_2 \in \mathbb{C}$ con

Re $n_1 > 0$ y Re $n_2 > 0$. Sin embargo, para nuestros propósitos será suficiente trabajar con parámetros reales.

En primer lugar, necesitamos discutir el radio de convergencia de la serie de potencias dada anterior, es decir, los dominios de definición de las funciones de Mittag-Leffler. Resulta que podemos dar un resultado completo para la versión de dos parámetros. En vista de la Observación 4.1 entonces, por supuesto, también tiene un a posteriori para las funciones de Mittag-Leffler de un parámetro.

Teorema 4.1

Considere la función de dos parámetros Mittag-Leffler E_{n_1, n_2} para algunos $n_1, n_2 > 0$. La serie de potencias que define $E_{n_1, n_2}(z)$, es convergente para todo $z \in \mathbb{C}$. En otras palabras, E_{n_1, n_2} es una función completa.

Prueba

Por definición, la función Mittag-Leffler de dos parámetros tiene la serie de potencias representación

$$\sum_{j=0}^{\infty} a_j z^j, \text{ con } a_j = \frac{1}{\Gamma(jn_1, n_2)}$$

En vista de la fórmula de Stirling (ver Teorema D.5) encontramos que,

$$a_j^{1/j} = \left(\frac{e}{jn_1, n_2} \right)^{n_1, n_2/j} (2\pi(jn_1, n_2))^{-\frac{1}{(2j)}} (1 + o(1)) \rightarrow 0$$

como $j \rightarrow \infty$ ya que $n_1 > 0$. Así, por el criterio de la raíz, el radio de convergencia de la serie de potencias es infinita.

Observación 4.3

Es posible investigar la función de Mittag-Leffler E_{n_1, n_2} también para $n_1 = 0$. En este caso, la serie de potencias tiene un radio de convergencia finito. Si, por ejemplo, $n_2 = 1$, entonces el radio de convergencia es 1. De hecho, una inspección detallada de la serie de potencias rendimientos de representación.

$$E_{0,1}(z) = \sum_{j=0}^{\infty} \frac{1}{\Gamma(1)} z^j = \sum_{j=0}^{\infty} z^j = \frac{1}{1+z}$$

Sin embargo, en el contexto de las ecuaciones diferenciales fraccionarias, las funciones de Mittag-Leffler $E_{0, n_2}(z)$ no son muy importantes, por lo que no los discutiremos más.

Observación 4.4

Las llamadas funciones de índice múltiple de Mittag-Leffler

$$E_{(n_{11}, n_{12}, \dots, n_{1,k}), (n_{21}, n_{22}, \dots, n_{2,k})}(z) = \sum_{j=0}^{\infty} z^j \prod_{\mu=1}^k \frac{1}{\Gamma(j n_{1\mu}, n_{2\mu} + 1)}$$

forman una clase aún más general de funciones. Resulta que los métodos que usan conceptos del cálculo fraccionario se pueden utilizar convenientemente para analizar estas funciones.

Sin embargo, dado que estas funciones no juegan un papel significativo en el contexto de este proyecto, no entraremos en detalles con respecto a estas funciones aquí y solo nos referiremos a las referencias citadas en él.

Ejemplo 4.1

Para algunas elecciones especiales de los parámetros n_1 y n_2 , podemos recuperar ciertas funciones bien conocidas:

- a) Para $x \in \mathbb{C}$, $E_2(-x^2) = E_{2,1}(-x^2) = \cos x$.
- b) Para $x \in \mathbb{C}$, $E_2(x^2) = E_{2,1}(x^2) = \cosh x$.
- c) Para $x > 0$, $E_{1/2}(-x^{1/2}) = E_{1/2,1}(x^{1/2}) = (1 + \operatorname{erf}(x)) \exp(x^2)$. (Esta relación puede extenderse a $x \in \mathbb{C}$ si $x^{1/2}$ se interpreta como la rama principal del complejo función de raíz cuadrada.)
- d) Para $x \in \mathbb{C}$ y $r \in \mathbb{N}$,

$$E_{1,r}(x) = \frac{1}{x^{r-1}} \left(\exp(x) - \sum_{k=0}^{r-2} \frac{x^k}{k!} \right)$$

En el caso de $x = 0$, se deben tomar los límites apropiados en el lado derecho.

En (c), $\operatorname{erf}(x)$ denota la función de error definida por

$$\operatorname{erf}(x) = \frac{2}{\sqrt{\pi}} \int_0^x \exp(-t^2) dt$$

Dejamos la demostración de estas identidades como ejercicio para el lector.

Una última propiedad de las funciones de Mittag-Leffler que mencionamos antes de construir el puente al cálculo fraccionario es una relación entre dos funciones de Mittag-Leffler con diferentes parámetros.

Teorema 4.2

Sean $n_1, n_2 > 0$ y $x \in \mathbb{C}$. Entonces,

$$E_{n_1, n_2}(x) = x E_{n_1, n_1 + n_2}(x) + \frac{1}{\Gamma(n_2)}$$

Prueba

Esto se puede mostrar escribiendo explícitamente la serie de potencias en cada lado de la identidad reclamada y comparando los coeficientes. Omitimos los detalles.

El resultado clave que indica por qué funciona Mittag-Leffler (en particular aquellos con un parámetro) son tan importantes en el cálculo fraccionario es el siguiente teorema. Esencialmente establece que las funciones propias de los operadores diferenciales de Caputo pueden ser expresado en términos de funciones de Mittag-Leffler.

Teorema 4.3

Sea $n > 0$ y $\lambda \in \mathbb{R}$. Además, defina

$$y(x) := E_n(\lambda x^n), x \geq 0$$

Entonces

$$D_{*0}^n y(x) = \lambda y(x)$$

Prueba

Primero observamos el caso $\lambda = 0$ y observamos que en este caso $y(x) E_n(0) = 1$.

Por lo tanto, $D_{*0}^n y(x) = 0 = \lambda y(x)$ según se requiera. Si, por el contrario, $\lambda \neq 0$, entonces (usando la notación) $p_k(x) = x^k$

$$\begin{aligned} D_{*0}^n y(x) &= D_{*0}^n \left[\sum_{j=0}^{\infty} \frac{(\lambda p_n)^j}{\Gamma(1+jn)} \right] (x) = J_0^{m-n} D^m \left[\sum_{j=0}^{\infty} \frac{\lambda^j p_{nj}}{\Gamma(1+jn)} \right] (x) \\ &= J_0^{m-n} \left[\sum_{j=0}^{\infty} \frac{\lambda^j D^m p_{nj}}{\Gamma(1+jn)} \right] (x) = J_0^{m-n} \left[\sum_{j=1}^{\infty} \frac{\lambda^j D^m p_{nj}}{\Gamma(1+jn)} \right] (x) \\ &= J_0^{m-n} \left[\sum_{j=1}^{\infty} \frac{\lambda^j p_{nj-m}}{\Gamma(1+jn-m)} \right] (x) = \sum_{j=1}^{\infty} \frac{\lambda^j J_0^{m-n} p_{nj-m}(x)}{\Gamma(1+jn-m)} \\ &= \sum_{j=1}^{\infty} \frac{\lambda^j p_{nj-n}}{\Gamma(1+jn-n)} = \sum_{j=1}^{\infty} \frac{\lambda^j x^{nj-n}}{\Gamma(1+jn-m)} \\ &= \sum_{j=1}^{\infty} \frac{\lambda^{j+1} x^{nj}}{\Gamma(1+jn)} = \lambda \sum_{j=1}^{\infty} \frac{(j x^n)^j}{\Gamma(1+jn)} = \lambda y(x) \end{aligned}$$

Aquí hemos utilizado el hecho de que, en vista de las propiedades de convergencia de la serie Al definir la función de Mittag-Leffler, podemos intercambiar la primera suma y la diferenciación y posterior suma e integración.

Observación 4.5

Es evidente a partir de (4.1) que la función de Mittag-Leffler E_1 satisface la ecuación funcional

$$E_1(x - y) = \frac{E_1(x)}{E_1(y)}, \dots, (4.2)$$

Una generalización de este resultado a las funciones de Mittag-Leffler E_n con $n \notin \mathbb{N}$ no se conoce y probablemente tal relación no existe. La ecuación funcional (4.2) juega un papel muy importante en el análisis de ecuaciones diferenciales de primer orden (en particular en la teoría de ecuaciones lineales) porque permite escribir un núcleo de convolución que surge, por ejemplo, en el enfoque de variación de constantes, en forma de un producto de una solución fundamental en un punto y la inversa de la solución fundamental en algún otro punto. El hecho de que una generalización fraccionada de esta característica no sea disponible es un obstáculo importante en el desarrollo de una teoría integral para ecuaciones diferenciales lineales fraccionarias.

Con frecuencia es de interés, tener algún conocimiento sobre el comportamiento asintótico de estas funciones. En este contexto, recordamos el siguiente resultado sobre las funciones de Mittag-Leffler de un parámetro que se verá que son importantes más adelante en el capítulo 5 (Syedlyetskii, 1994).

Teorema 4.4

Sea $n > 0$. La función de Mittag-Leffler E_n se comporta de la siguiente manera:

- (a) $E_n(re^{i\phi}) \rightarrow 0$ para $r \rightarrow \infty$ si $|\phi| > n\pi/2$,
- (b) $E_n(re^{i\phi})$ permanece acotado para $r \rightarrow \infty$ si $|\phi| = n\pi/2$,
- (c) $|E_n(re^{i\phi})| \rightarrow \infty$ para $r \rightarrow \infty$ si $|\phi| < n\pi/2$.

Obviamente, en el caso clásico $n = 1$ esto se reduce al hecho bien conocido de que, como $|z| \rightarrow \infty$, $\exp(z)$ (a) tiende a cero si $\arg z > \pi/2$, (b) permanece acotado si $\arg z = \pi/2$ y (c) crece sin límite si $\arg z < \pi/2$.

Prueba

Esbozaremos la prueba dada por Wiman.

En el primer paso, consideremos el caso de que $n = 1/k$ con algún $k \in \mathbb{N}$. En este caso podemos ver que $E_n = E_{1/k}$ satisface la ecuación diferencial lineal de primer orden

$$E'_{\frac{1}{k}}(x) = k x^{k-1} E_{\frac{1}{k}}(x) + k \sum_{\mu=1}^{k-1} \frac{z^{\mu-1}}{\Gamma(\mu/k)}$$

Esto se puede mostrar fácilmente usando la representación en serie de potencias de $E_{\frac{1}{k}}$. Desde que nosotros también sabemos que $E_{\frac{1}{k}}$ satisface la condición inicial que corresponde a esta ecuación diferencial, podemos usar los métodos estándar para la solución de tales ecuaciones diferenciales y encontrar la representación

$$E'_{\frac{1}{k}}(x) = \exp(x^k) + \exp(x^k) \int_0^x k \exp(-z^k) \sum_{\mu=1}^{k-1} \frac{z^{\mu-1}}{\Gamma(\mu/k)}$$

Para cada uno de los k sumandos del lado derecho de esta ecuación, se puede establecer formular expresiones asintóticas que permitan concluir el resultado deseado.

Si ahora n es un número racional positivo, digamos $n = l/k$ con ℓ relativo primo y k , entonces podemos invocar la identidad

$$E_{\ell/k}(x) = \frac{1}{\ell} \sum_{\mu=1}^{\ell-1} E_{1/k} \left(x^{1/\ell} \exp \frac{2\mu\pi i}{\ell} \right), \dots, (4.3)$$

lo cual, usando el resultado para $E_{1/k}$ que ya hemos mostrado, implica la afirmación también para $n = \ell/k$.

Finalmente, para valores irracionales de n podemos elegir una sucesión $(n_j)_{j=0}^{\infty}$ de valores racionales números que convergen a n . Para cada uno de los E_{n_j} el resultado ya está en su lugar, y tomando los límites apropiados, derivamos el resultado también para E_n .

En los resultados finales de este breve capítulo, describiremos la interconexión entre una función de Mittag-Leffler de un parámetro y la operación de transformada de Laplace y una importante consecuencia de este teorema. Más información sobre Laplace la transformada se da en el Apéndice D.3; en este punto solo notamos que es una herramienta muy útil herramienta para la solución de ciertas clases de ecuaciones diferenciales fraccionarias. Deberíamos tratar con tal enfoque en el capítulo 7.

Teorema 4.5

Sea $n > 0$ y $\lambda \in \mathbb{C}$ y defina $y(x) = E_n(-\lambda x^n)$. Entonces, transformada Laplace de y está dada por

$$\mathcal{L}y(s) = \frac{s^{n-s}}{s^n \lambda}, \dots, (4.4)$$

Prueba

Esto se puede demostrar escribiendo explícitamente la expansión en serie de $y(x)$ en potencias de x^n y aplicando la transfor-

mada de Laplace de manera terminológica. Dejamos los detalles al lector.

En conjunto con el Teorema del valor final para la transformada de Laplace (Teorema D.13) obtenemos un enunciado sobre el comportamiento asintótico de la función y mencionada en el Teorema 5.5 ya que su argumento tiende a infinito:

Teorema 4.6

Sea $n > 0$, $r > 0$, $\phi \in [-\pi, \pi]$ y $\lambda = r \exp(i\phi)$. Denote $y(x) := e_n(-\lambda x^n)$. Entonces,

(a) $\lim_{x \rightarrow \infty} y(x) = 0$ si $|\phi| < n\pi/2$,

(b) $y(x)$ no está acotado cuando $x \rightarrow \infty$ si $|\phi| > n\pi/2$.

Prueba

Esta es una consecuencia inmediata del Teorema 5.5, Teorema D.13 y Observación D.1. Alternativamente, también podemos deducir este resultado del Teorema 4.4. Vale la pena señalar que la evaluación numérica de la función Mittag-Leffler $E_{n_1, n_2}(x)$ puede, dependiendo de los valores precisos de los parámetros y el argumento x , ser una tarea sumamente difícil. Un algoritmo útil para la solución de este problema ha sido proporcionado por Gorenflo et al. Más recientemente, una alternativa El método numérico ha sido desarrollado por Seybold y Hilfer.

Usaremos con frecuencia las funciones de Mittag-Leffler en los siguientes capítulos. Para resultados generales adicionales sobre las funciones de Mittag-Leffler nos referimos al Apéndice A, y el estudio detallado y más información se describe su conexión con el cálculo fraccionario. Algo, más resultados sobre el com-

portamiento a largo plazo de ciertas funciones especiales de Mittag-Leffler y sobre el número de sus ceros también se dará más adelante en este proyecto; ver Teoremas 7.3–8.8.

05

Capítulo

*ECUACIONES DIFERENCIALES
FRACCIONARIAS*

Escucha a tu niño interno que él te enseñara que uno debe ser dueño de sí mismo, sin dejarse jamás arrastrar de las ocasiones; te enseñara todas las virtudes, misericordia y obediencia, de constancia inalterable en las resoluciones tomadas con madurez, de indiferencia respecto a la gloria popular.

Burseca

Existencia y singularidad para ecuaciones diferenciales fraccionarias de Riemann-Liouville

En esta parte discutimos ahora las cuestiones clásicas relativas a las ecuaciones diferenciales que involucran derivadas fraccionarias, es decir, las cuestiones de existencia y unicidad de las soluciones. Nos interesarán principalmente los problemas de valor inicial del problema de Cauchy, buscando una solución que sea derivable y continua en particular en los resultados globales.

Para una discusión de otros tipos de condiciones nos referimos a Samko, Kilbas y Marichev, Kilbas y Trujillo o Agarwal, Benchohra y Hamani, en el presente capítulo se centrará en las ecuaciones diferenciales con operadores Riemann-Liouville; las derivadas de Caputo son el tema del siguiente capítulo.

Teorema 5.1

Sea $n > 0$, $n \notin \mathbb{N}$ y $m = [n]$. Además, sea $K > 0$, $h^* > 0$, y $b_1, b_2, \dots, b_m \in \mathbb{R}$. Definir

$$G := \left\{ (x, y) \in \mathbb{R}^2 : 0 \leq x \leq h^*, y \in \mathbb{R}, \text{ para } x = 0, \left| y \left| x^{m-n} y - \sum_{k=1}^m b_k x^{m-k} / \Gamma(n-k+1) \right| < K \right\}$$

Además, supongamos que la función $f: G \rightarrow \mathbb{R}$ es continua y acotada en G y que cumple una condición de Lipschitz con respecto a la segunda variable, es decir, existe una constante $L > 0$ tal que, para todo $(x, y_1), (x, y_2) \in G$, tenemos

$$|f(x, y_1) - f(x, y_2)| < L|y_1 - y_2|$$

Entonces la ecuación diferencial

$$D_0^n y(x) = f(x, y(x))$$

Con las condiciones iniciales siguientes,

$$D_0^{n-k} y(0) = b_k; k = 1, 2, \dots, m-1, \lim_{z \rightarrow 0^+} J_0^{m,n} y(z) = b_m$$

tiene una solución continua, definida de forma única y $\in C(0, h]$ donde

$$h := \min \left\{ h^*, \tilde{h}, \left(\frac{\Gamma(n+1)k}{M} \right)^{1/m} \right\}$$

con $M = \sup_{(x,y) \in G} |f(x, z)|$ y siendo $\tilde{h}$ un número positivo arbitrario que satisface la restricción

$$\tilde{h} < \left(\frac{\Gamma(2n-m+1)k}{\Gamma(n-m+1)L} \right)^{1/n}$$

El resultado es muy similar a los resultados clásicos conocidos para ecuaciones de primer orden.

Por lo tanto, probablemente no sea sorprendente encontrar que la prueba también es análoga.

Específicamente, primero transformaremos el problema de valor inicial en un problema equivalente. Ecuación integral de Volterra (Lema 5.2), y luego vamos a probar la existencia y unicidad de la solución de esta ecuación integral por un tipo de Picard,

proceso de iteración; es decir, mediante el uso de una variante del teorema del punto fijo de Banach en un espacio métrico completo elegido, cf. Lema 5.3. El teorema 5.1 es, por tanto, una consecuencia de estos dos lemas (Diethelm & Ford, 2012).

Lema 5.2

Suponga las hipótesis del Teorema 5.1 y sea $h > 0$. La función $y \in C(0, h]$ es una solución de la ecuación diferencial

$$D_0^n y(x) = f(x, y(x))$$

equipado con las condiciones iniciales

$$D_0^{n-k} y(0) = b_k, k = 1, 2, \dots, m-1, \lim_{z \rightarrow 0_+} J_0^{m-n} y(z) = b_m$$

sí y solo si es una solución de la ecuación integral de Volterra

$$y(x) = \sum_{k=1}^m \frac{b_k x^{m-k}}{\Gamma(n-k+1)} + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-2} f(t, y(t)) dt$$

Debemos decir que la Existencia y Unicidad de E.D.Fs. determinan las propiedades que deben tener las funciones f para poder asegurar que la E.D.Fs.

$$D_0^n y(x) = f(x, y(x))$$

tenga solución, y adicionalmente unicidad de soluciones. La idea de estos teoremas, que vamos a seguir, es la siguiente. Dado el problema de Cauchy denotado por.

$$\begin{cases} D_0^n y(x) = f(x, y(x)) \\ D_0^{n-k} y(0) = b_k, k = 1, 2, \dots, m-1, \lim_{z \rightarrow 0_+} J_0^{m-n} y(z) = b_m \end{cases}$$

Observación 5.1

Una mirada a la ecuación integral revela, por qué sólo hemos asumido y sea continua en el intervalo semiabierto $(0, h]$ y no en el intervalo cerrado $[0, h]$, cómo podríamos haber hecho para ecuaciones de orden entero: Si y fuera continua a lo largo de $[0, h]$ entonces el lado izquierdo de la ecuación integral sería continuo en este intervalo, y así sería la integral en el lado derecho, debido a la continuidad de f . Por lo tanto, la suma también debe ser continua en $[0, h]$.

En vista de definición de m , vemos fácilmente que los sumandos son de hecho continuos en $[0, h]$ para $k = 1, 2, \dots, m-1$, pero el restante ($k = m$) es ilimitado cuando $x \rightarrow 0$ porque $m > n$ si $n \notin \mathbb{N}$, a menos que $b_m = 0$.

Prueba

Suponga primero que y es una solución de la ecuación integral. Podemos reescribir esta ecuación en la forma más corta

$$y(x) = \sum_{k=1}^m \frac{b_k x^{m-k}}{\Gamma(n-k+1)} + J_0^n f(\cdot, y(\cdot))(x)$$

Ahora aplicamos el operador diferencial D_0^n a ambos lados de esta relación e inmediatamente obtener, en vista del Ejemplo 2.4 y el Teorema 2.14, que y también resuelve la ecuación diferencial Con respecto a las condiciones iniciales, nos fijamos en el caso $1 \leq k \leq m-1$ primero y encuentra, mediante una aplicación de D_0^{n-k} a la ecuación de Volterra, que

$$D_0^{n-k} y(x) = \sum_{j=1}^m \frac{b_j D_0^{n-k} (\cdot)^{m-k}(x)}{\Gamma(n-k+1)} + D_0^{n-k} J_0^n J_0^k f(\cdot, y(\cdot))(x)$$

en vista de la propiedad de semigrupo de integración fraccionaria. Por el ejemplo 2.4 encontramos que los sumandos se anulan idénticamente para $j > k$. Además, por el mismo ejemplo, los sumandos para $j < k$ desaparecen si $x = 0$. Así, de acuerdo con el teorema 2.14,

$$D_0^{n-k}y(x) = \frac{b_j D_0^{n-k}(\cdot)^{m-k}(0)}{\Gamma(n-k+1)} + J_0^k f(\cdot y(\cdot))(0)$$

Como $k \geq 1$, la integral se anula y, una vez más, aplicando el ejemplo 2.4 encontramos que $D_0^{n-k}(\cdot)^{m-k}(\cdot) = \Gamma(n-k+1)$. Así $D_0^{n-k}y(x) = b_k$ como requiere la condición inicial.

Finalmente, para $k = m$ aplicamos el operador J_0^{n-m} a ambos lados de la ecuación integral y encontramos que, en el límite $z \rightarrow 0$, todos los sumandos de la suma se anulan excepto para la m ésima.

La integral $J_0^{m-n} J_0^n f(\cdot y(\cdot))(z) = J_0^m f(\cdot y(\cdot))(z)$ también desaparece cuando $z \rightarrow 0$. Así encontramos

$$\lim_{z \rightarrow 0_+} J_0^{m-n}y(z) = \lim_{z \rightarrow 0_+} J_0^{m-n} \frac{b_m J_0^{m-n}(\cdot)^{m-n}(z)}{\Gamma(n-m+1)} = b_m$$

debido al Ejemplo 2.1.

Por lo tanto y resuelve el problema de valor inicial dado.

Si y es una solución continua del problema de valor inicial, entonces definimos $z(x) := f(x, y(x))$.

Por suposición, z es una función continua y $z(x) = f(x, y(x)) = D_0^n y(x) = D^m J_0^{m-n} y(x)$, Por lo tanto, $D^m J_0^{m-n}$ y también es continuo, es decir, $J_0^{m-n} y \in C^m(0, h)$.

Por lo tanto, podemos aplicar el Teorema 2.23 para derivar

$$y(x) = J_0^n D^n y(x) + \sum_{k=1}^m C_k x^{n-k} = J_0^n f(\cdot y(\cdot))(x) + \sum_{k=1}^m C_k x^{n-k},$$

con ciertas constantes $c_1, c_2, \dots, c_m$. Introduciendo las condiciones iniciales como se indica arriba, podemos determinar estas constantes C_k como $C_k = b_k/\Gamma(n - m + 1)$.

Lema 5.3

Bajo los supuestos del Teorema 5.1, la ecuación de Volterra

$$y(x) = \sum_{k=1}^m \frac{b_k x^{m-k}}{\Gamma(n-k+1)} + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt$$

posee una solución unívocamente determinada $y \in C(0, h)$.

Prueba

$$B := \left\{ y \in (0, h]: \sup_{0 < x \leq h} \left| x^{m-n} y(x) - \sum_{k=1}^m \frac{b_k x^{m-k}}{\Gamma(n-k+1)} \right| \leq K \right\}$$

$$Ay(x) := \sum_{k=1}^m \frac{b_k x^{m-k}}{\Gamma(n-k+1)} + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt$$

Entonces notamos que, para $y \in B$, Ay es también una función continua en $(0, h)$. Además,

$$\begin{aligned} \left| x^{m-n} Ay(x) - \sum_{k=1}^m \frac{b_k x^{m-k}}{\Gamma(n-k+1)} \right| &= \left| \frac{x^{m-n}}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt \right| \\ &\leq \frac{x^{m-n}}{\Gamma(n)} M \int_0^x (x-t)^{n-1} dt \\ &\leq \frac{x^{m-n}}{\Gamma(n)} M \frac{x^n}{n} = \frac{x^m M}{\Gamma(n+1)} \leq K \end{aligned}$$

para $x \in (0, h]$, donde la última desigualdad se deriva de la definición de h . Esto muestra que $Ay \in B$ si $y \in B$, es decir, el operador

A mapea el conjunto B en sí mismo.

A continuación, presentamos un nuevo conjunto

$$\hat{B} := \left\{ y \in (0, h]: \sup_{0 < x \leq h} |x^{m-n} y(x)| < \infty \right\}$$

y sobre este conjunto definimos una norma $\|\cdot\|_{\hat{B}}$ por

$$\|y\|_{\hat{B}} := \sup_{0 < x \leq h} |x^{m-n} y(x)|$$

Se ve fácilmente que $\hat{B}$, equipado con esta norma, es un espacio lineal normado, y que B es un subconjunto completo de este espacio.

Usamos la definición de A para reescribir la ecuación de Volterra de manera más compacta como

$$y = Ay$$

Por lo tanto, para probar el resultado deseado, es suficiente demostrar que el operador A tiene un único punto fijo.

Para este propósito, emplearemos el punto fijo de Weissinger teorema (Teorema D.7). En este contexto demostramos, para $y, \tilde{y} \in B$

$$\|A^j y - A^j \tilde{y}\|_{\hat{B}} \leq \left(\frac{L h^n \Gamma(n-m+1)}{\Gamma(2n-m+1)} \right)^j \|y - \tilde{y}\|_{\hat{B}}, \dots, (5.1)$$

Esto se puede demostrar por inducción: en el caso de $j = 0$, el enunciado es trivialmente verdadero. Para el paso de inducción $j-1 \rightarrow j$, procedemos de la siguiente manera. Escribimos

$$\begin{aligned}
\|A^j y - A^j \tilde{y}\|_{\tilde{B}} &\leq \frac{L}{\Gamma(n)} \sup_{0 < x \leq h} x^{m-n} \int_0^x (x-t)^{n-1} |(A^{j-1}y(t) - (A^{j-1}\tilde{y}(x))| dt \\
&\leq \frac{L}{\Gamma(n)} \sup_{0 < x \leq h} x^{m-n} \int_0^x (x-t)^{n-1} t^{m-n} t^{n-m} |(A^{j-1}y(t) - (A^{j-1}\tilde{y}(x))| dt \\
&\leq \frac{L}{\Gamma(n)} \|A^j y - A^j \tilde{y}\|_{\tilde{B}} \sup_{0 < x \leq h} x^{m-n} \int_0^x (x-t)^{n-1} t^{m-n} dt \\
&= \frac{L}{\Gamma(n)} \|A^j y - A^j \tilde{y}\|_{\tilde{B}} \sup_{0 < x \leq h} \frac{\Gamma(n) \Gamma(n-m+1)}{\Gamma(2n-m+1)} x^n \\
\|A^j y - A^j \tilde{y}\|_{\tilde{B}} &= \frac{L h^n \Gamma(n-m+1)}{\Gamma(2n-m+1)} \|A^j y - A^j \tilde{y}\|_{\tilde{B}}
\end{aligned}$$

por definición del operador A, y la condición de Lipschitz en f. En el siguiente paso nos estima más para encontrar

$$\begin{aligned}
\|A^j y - A^j \tilde{y}\|_{\tilde{B}} &= \sup_{0 < x \leq h} |x^{m-n} (A^j y(x) - A^j \tilde{y}(x))| \\
&= \sup_{0 < x \leq h} |x^{m-n} (AA^{j-1}y(x) - AA^{j-1}\tilde{y}(x))| \\
&= \sup_{0 < x \leq h} \frac{x^{m-n}}{\Gamma(n)} \left| \int_0^x (x-t)^{n-1} [f(t, A^{j-1}y(t)) - f(t, A^{j-1}\tilde{y}(x))] dt \right| \\
&\leq \sup_{0 < x \leq h} \frac{x^{m-n}}{\Gamma(n)} \int_0^x (x-t)^{n-1} |f(t, A^{j-1}y(t)) - f(t, A^{j-1}\tilde{y}(x))| dt \\
&\leq \frac{L}{\Gamma(n)} \sup_{0 < x \leq h} x^{m-n} \int_0^x (x-t)^{n-1} |(A^{j-1}y(t) - (A^{j-1}\tilde{y}(x))| dt
\end{aligned}$$

Ahora usamos la hipótesis de inducción, demostrando (5.1).

Por lo tanto, podemos aplicar Teorema D.7 con $\alpha_j = \gamma^j$ donde $\gamma = \frac{L h^n \Gamma(n-m+1)}{\Gamma(2n-m+1)}$. Permanece para probar que la serie $\sum_{j=0}^{\infty} \alpha_j$ es convergente. Esto, sin embargo, es trivial en vista del hecho de que $h \leq \tilde{h}$ y la definición de $\tilde{h}$ que implica $\gamma < 1$. Así, una aplicación del teorema del punto fijo produce la existencia y la unicidad de la solución de nuestra ecuación integral.

Observación 5.2

La demostración del Lema 5.3 también nos da, al menos en teoría, una explicación constructiva método para encontrar la so-

lución del problema de valor inicial mediante el cálculo de la sucesión $(A^j y_0)_{j=0}^{\infty}$, donde y_0 es un elemento arbitrario de B . El límite de esta secuencia es la solución deseada. Que normalmente uno elige

$$y_0 = \sum_{k=1}^m \frac{b_k x^{n-k}}{\Gamma(n-k+1)}$$

En este caso llamamos a la secuencia $(A^j y_0)_{j=0}^{\infty} = 1$ la secuencia de iteración de Picard correspondiente al problema de valor inicial dado.

Observación 5.3

El teorema 5.1 puede interpretarse como un análogo del Picard-Lineal para ecuaciones diferenciales de primer orden. Podemos preguntarnos si las condiciones son demasiado nítidas en la configuración fraccionaria.

Resulta que es posible probar un resultado más débil bajo suposiciones más débiles: si eliminamos la condición de Lipschitz en f , todavía se puede demostrar la existencia de la solución. Esto corresponde al teorema de Peano de la existencia en la teoría clásica.

La prueba es esencialmente similar, simplemente reemplazando el teorema de Weissinger. por el teorema del punto fijo de Schauder. para el correspondiente problema que involucra operadores de Caputo en lugar de derivados de Riemann-Liouville, se investigará esto explícitamente en el capítulo 7.

Para una ilustración adicional de este comentario, discutimos un ejemplo muy simple de una ecuación diferencial fraccio-

naria con un lado derecho que no cumple la condición de Lipschitz.

Ejemplo 5.1

Considere la ecuación diferencial

$$D_0^n y(x) = [y(x)]^\mu$$

donde $0 < \mu < 1$. En este caso, el lado derecho, la continuidad de la ecuación viola la condición de Lipschitz. Si seleccionamos la condición inicial correspondiente a esta ecuación diferencial como

$$\lim_{z \rightarrow 0_+} J_0^{1-n} y(z) = 0 \text{ y } D_0^{n-k} y(0) = 0, (k = 1, 2, \dots, [n] - 1)$$

vemos fácilmente que una solución es $y \equiv 0$. Sin embargo, un cálculo explícito revela que la función y dada por

$$y(x) = \sqrt[\mu-1]{\frac{\Gamma(j+1)}{\Gamma(j+1-n)}} x^j$$

con $j = n/(1-\mu)$ también resuelve el problema del valor inicial. Así, en efecto, vemos que, en general, la unicidad de la solución no se puede esperar sin el Lipschitz condición.

06

Capítulo

*ECUACIONES DIFERENCIALES
FRACCIONARIAS DESDE LA ÓPTICA
DE CAPUTO*

Escucha a tu niño interno que te enseñara, el método claro y el camino seguro de inventar y ordenar las máximas necesidades para una vida ajustada digna y feliz, así mismo que no trasluzca señal de ira u otra pasión, sino por el contrario libre de estos afectos.

Burseca

Habiendo establecido los fundamentos de una teoría para ecuaciones diferenciales fraccionarias con derivadas de Riemann-Liouville, llegamos ahora al problema correspondiente para Operadores según Caputo. Dado que estos últimos parecen ser mucho más importantes que el primero en lo que se refiere a aplicaciones fuera de las matemáticas, discutiremos este problema de una manera más detallada.

El énfasis principal será en problemas de valor inicial. En particular, las dos primeras secciones de este capítulo, se dedica a las cuestiones de existencia y unicidad, respectivamente, para un análisis más general clase de ecuaciones mientras que en la tercera sección nos ocuparemos de la estabilidad estructural de las soluciones: ¿Cómo dependen de los datos? Las propiedades de suavidad de las soluciones serán luego discutidas en la Sección especial. Finalmente, lo haremos abandonar el área de los problemas de valor inicial y proporcionar algunos resultados fundamentales sobre problemas de valores en la frontera.

Los resultados de este capítulo se utilizarán en el siguiente para establecer a fondo más teoremas específicos que describen las propiedades de ciertos casos especiales prácticamente muy importantes.

Posteriormente, en el siguiente capítulo, las consideraciones se extenderán a un ámbito más general clase de ecuaciones, a saber, ecuaciones que contienen más de un operador diferencial.

Existencia de soluciones

Comenzamos una vez más con ecuaciones de la forma

$$D_0^n y(x) = f(x, y(x)), \dots (6.1A)$$

combinado con las condiciones iniciales apropiadas. Como se indica en el capítulo 3 estas condiciones tienen la forma

$$D^k y(0) = y_0^{(k)}, k = 1, 2, \dots, m-1, \dots (6.1B)$$

donde como siempre hemos establecido $m = [n]$. En el capítulo 6, un problema más general será considerado. Los resultados presentados en esta sección y en la primera parte de la Sección 6.2 son principalmente versiones ligeramente modificadas de los hallazgos de Diethelm y Ford.

El primer resultado es un resultado de existencia que corresponde a la existencia clásica del teorema Peano para ecuaciones de primer orden.

Teorema 6.1

Sea $0 < n$ y $m = [n]$. Además, sea $y_0^{(0)}, y_0^{(1)}, \dots, y_0^{(m-1)} \in R$, $K > 0$ y $h^* > 0$. Defina

$G := \left\{ (x, y) : x \in [0, h^*], \left| y - \sum_{k=0}^{m-1} \frac{x^k y_0^{(k)}}{k!} \right| \leq K \right\}$, y sea la función $f : G \rightarrow R$ continua.

Además, definimos $M := \sup_{(x,z) \in G} |f(x, z)|$ y

$$h = \left[\min \left\{ h^*, \left(\frac{k \Gamma(n+1)}{M} \right)^{1/n}, \text{ caso contrario} \right\} \right]$$

Entonces, existe una función $y \in C[0, h]$ que resuelve el problema de valor inicial (6.1).

$$\begin{cases} D_0^n y(x) = f(x, y(x)), \\ D^k y(0) = y_0^{(k)}, k = 1, 2, \dots, m-1, \end{cases} (6.1)$$

Observación 6.1

Teniendo en consideración la simplicidad de la presentación, sólo tratamos en esta oportunidad el caso escalar explícitamente. Sin embargo, todos los resultados en este capítulo y en el siguiente pueden extenderse a funciones vectoriales y es decir, sistemas de ecuaciones, sin ningún problema.

Observación 6.2

En muchas aplicaciones en ciencia e ingeniería, tenemos $0 < n \leq 1$. En este caso, el conjunto G definido en el Teorema 6.1 es simplemente el rectángulo simple $G = [0, h^*] \times [y_0^{(0)} - K, y_0^{(0)} + K]$.

Para las demostraciones de la mayoría de los teoremas en este capítulo y en la próxima sección, usaremos el siguiente lema que adapta el enunciado del Lema 5.2 a la situación actual.

Lema 6.2

Suponga las hipótesis del Teorema 6.1. la función $y \in C[0, h]$ es una solución del problema de valor inicial (6.1) si y solo sí, es una solución de la ecuación integral no lineal de Volterra de segunda clase:

$$y(x) = \sum_{k=0}^{m-1} \frac{x^k}{k!} y_0^{(k)} + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt, \dots (6.2)$$

Con $m = [n]$

Antes de llegar a la demostración de este lema, nos gustaría hacer algunos comentarios sobre su relación con los resulta-

dos correspondientes para ecuaciones con Operadores de Riemann-Liouville y ecuaciones de orden entero, respectivamente.

Observación 6.3

Recordando el comentario 5.1 y usando los lemas 5.2 y 6.2, podemos comparar el comportamiento de las soluciones de ecuaciones diferenciales fraccionarias tipo Caputo con los del tipo Riemann-Liouville. Al hacerlo, encontramos que el problema de continuidad en el origen que encontramos para el escenario de Riemann-Liouville no surge en el Entorno Caputo. Más bien, resulta que la continuidad de la función f implica continuidad de la solución y a lo largo del intervalo cerrado $[0, h]$

Observación 6.4

Veamos (6.2) para algún $n \in (0, 1]$ y para dos valores diferentes de x , digamos x_1 , y x_2 con $x_1 < x_2$, y reste la segunda de estas ecuaciones de la primera. Esto produce

$$\begin{aligned} y(x_1) - y(x_2) &= \frac{1}{\Gamma(n)} \int_0^{x_2} (x_2 - t)^{n-1} f(t, y(t)) dt \\ &\quad - \frac{1}{\Gamma(n)} \int_0^{x_1} (x_1 - t)^{n-1} f(t, y(t)) dt \\ &= \frac{1}{\Gamma(n)} \int_0^{x_2} [(x_2 - t)^{n-1} - (x_1 - t)^{n-1}] f(t, y(t)) dt + \frac{1}{\Gamma(n)} \int_{x_1}^{x_2} (x_3 - \\ &\quad t)^{n-1} f(t, y(t)) dt, \dots (6.3) \end{aligned}$$

Ahora consideremos primero el caso clásico (no fraccionario) $n = 1$. Aquí el término en corchetes en el lado derecho de (6.3) es cero, y por lo tanto toda la primera integral desaparece.

Esta ecuación implica entonces el hecho bien conocido de

que, si ya sabemos la solución $y(x_1)$ de nuestro problema de valor inicial dado (6.1) en el punto $x_1 > 0$ entonces puede calcular la solución en el punto $x_1 < x_2$ exclusivamente sobre la base de $y(x_1)$ y la función f . No necesitamos usar ninguna información sobre $y(x)$ para $x \in [0, x_1]$.

Esta observación, que es solo otra forma de expresar la localidad del operador diferencial de orden entero, es la base de casi todos los métodos clásicos para la solución de ecuaciones diferenciales de primer orden, y también es de importancia fundamental en el modelado matemático de muchos sistemas en física, ingeniería y otras ciencias porque establece que es suficiente observar el estado de un primer orden del sistema en un punto arbitrario en el tiempo para calcular su comportamiento en el futuro.

Sin embargo, cuando observamos el caso fraccionario $0 < n < 1$, la situación es fundamentalmente diferente. Aquí la primera integral del lado derecho de (6.3) no se anula en general. Por lo tanto, siempre que queramos calcular la solución $y(x_2)$ en algún punto es necesario tener en cuenta toda la historia de y desde el punto de partida 0 hasta el punto de interés x_2 .

Esto refleja la no localidad del operador diferencial fraccionario de Caputo que ya habíamos observado en la observación 3.2. Obviamente, esta observación tiene una influencia sustancial en la construcción de métodos numéricos para tales ecuaciones. Además, para un sistema de orden fraccionario que modela algunos fenómenos del mundo real uno puede llegar a la conclusión de que aquí uno se vería forzado para medir el estado del sistema en el punto inicial y no se le permitiría medir en un punto arbitrario. Esta última conclusión, sin embargo, no es correcta como lo veremos el Teorema 6.17.

Por lo tanto, se deduce que las ecuaciones de orden entero son herramientas apropiadas para el modelado de sistemas sin memoria, mientras que las ecuaciones de orden fraccionario son el método de elección para la descripción de sistemas con memoria.

Por supuesto, dado que habíamos notado que los derivados de Riemann-Liouville no son locales o bien, esta observación se aplica a las ecuaciones diferenciales fraccionarias de Riemann-Liouville también.

Observación 6.5

Como se indicó en el comentario anterior, (6.3) tiene un gran impacto en la base teórica de métodos numéricos para ecuaciones diferenciales fraccionarias. En los métodos estándar suelen tener esto en cuenta correctamente, pero a menudo no lo hacen. hacer uso de toda la potencia de esta ecuación. Específicamente, como señaló Deng (2007), aunque el término entre paréntesis en el lado derecho de (6.3) no sea cero, será ser bastante pequeña en magnitud en ciertas situaciones, por ejemplo, si x_1 y x_2 son bastante grandes en comparación con $x_2 - x_1$, una situación que es bastante común para los métodos numéricos donde $x_2 - x_1$ puede ser un tamaño de paso pequeño.

Por lo tanto, uno puede ser capaz de aproximar la primera integral en el lado derecho de (6.3) por un método menos exacto pero barato, así reduciendo la complejidad total sin perder significativamente la precisión.

Demostración (del Lema 6.2)

La prueba de que cada solución continua de la ecuación de Volterra también resuelve el problema de valor inicial está muy cerca de la prueba de la correspondiente parte del Lema 5.2; por lo tanto, dejamos los detalles al lector.

Para la otra dirección, definimos $z(x) := f(x, y(x))$ y una vez más tengamos en cuenta que $z \in C[0, h]$ por nuestras suposiciones sobre y , y f . Entonces, usando la definición del operador diferencial de Caputo, la ecuación diferencial se puede reescribir como

$$z(x) = f(x, y(x)) = D_{*0}^n y(x) = D_0^n (y - T_{m-1}[y; 0])(x)$$

Como estamos tratando con funciones continuas, podemos aplicar el operador a , ambos lados de la ecuación y hallar

$$J_0^m z(x) = J_0^{m-n} (y - T_{m-1}[y; 0])(x) + q(x)$$

con algún polinomio q de grado no superior a $m-1$. Como z es continua, la función $J_0^m z$ en el lado izquierdo de esta ecuación tiene un cero de orden (al menos) m en el origen. Además, la diferencia $y - T_{m-1}[y; 0]$ tiene la misma propiedad por construcción, y por lo tanto la función $J_0^{m-n} (y - T_{m-1}[y; 0])$ en el lado derecho de nuestra ecuación debe tener tal m -ésimo orden cero también. Así, el polinomio q tiene la misma propiedad, e inmediatamente deducimos ya que su grado no es mayor que $m-1$ que $q = 0$. En consecuencia,

$$J_0^m z(x) = J_0^{m-n} (y - T_{m-1}[y; 0])(x)$$

y aplicando el operador diferencial de Riemann-Liouville D_0^{m-n} a esta ecuación encontramos

$$y(x) - J_0^{m-n}(y - T_{m-1}[y; 0])(x) = D_0^{m-n} J_0^m z(x) = D^1 J_0^{1+n-m} J_0^m z(x) \\ = D J_0^{1+n} z(x) = J_0^n z(x)$$

Recordando las definiciones de z y el polinomio de Taylor $T_{m-1}[y, 0]$, este es solo la ecuación requerida de Volterra.

Demostración (del Teorema 5.1)

Si $M = 0$ entonces $f(x, y) = 0$ para todo $(x, y) \in G$. En este caso es evidente que la función $y : [0, h] \rightarrow \mathbb{R}$ con $y(x) = \sum_{k=0}^{m-1} \frac{y_0^{(k)} x^k}{k!}$ es una solución del problema de valor inicial (6.1). Por lo tanto, concluimos, como se requiere, que una solución existe en este caso.

De lo contrario, aplicamos el Lema 6.2 y vemos que nuestro problema de valor inicial (6.1) es equivalente a la ecuación de Volterra (6.2). Introducimos así el polinomio T que satisface las condiciones iniciales, a saber

$$T(x) := \sum_{k=0}^{m-1} \frac{x^k}{k!} y_0^{(k)}, \dots (6.4)$$

y el conjunto $U := \{y \in \mathcal{C}[0, h] : \|y - T\|_\infty \leq K\}$. Es evidente que U es un cerrado y subconjunto convexo del espacio de Banach de todas las funciones continuas en $[0, h]$, equipado con la norma de Chebyshev. Por lo tanto, U también es un espacio de Banach. Dado que el polinomio T es un elemento de U , también vemos que U no está vacío. Sobre este conjunto U definimos el operador A por

$$(Ay)(x) := T(x) + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt, \dots (6.5)$$

Usando este operador, la ecuación cuya solución necesitamos probar, a saber. En la ecuación de Volterra (6.2), se puede reescribir como $y = Ay$ y así, para probar nuestro resultado deseado de existencia, tenemos que mostrar que A tiene un punto fijo. Por lo tanto, procedemos investigando las propiedades del operador

A más cerca.

Nuestro primer objetivo en este contexto es mostrar que $Ay \in U$ para $y \in U$. Con este fin, Comencemos observando que, para $0 \leq x_1 \leq x_2 \leq h$,

$$\begin{aligned} & |(Ay)(x_1) - (Ay)(x_2)| \\ &= \frac{1}{\Gamma(n)} \left| \int_0^{x_1} (x_1 - t)^{n-1} f(t, y(t)) dt - \int_0^{x_2} (x_2 - t)^{n-1} f(t, y(t)) dt \right| \\ &= \frac{1}{\Gamma(n)} \left| \int_0^{x_1} ((x_1 - t)^{n-1} - (x_2 - t)^{n-1}) f(t, y(t)) dt - \int_{x_1}^{x_2} (x_2 - t)^{n-1} f(t, y(t)) dt \right| \\ &\leq \frac{M}{\Gamma(n)} \left(\int_0^{x_1} |(x_1 - t)^{n-1} - (x_2 - t)^{n-1}| dt - \int_{x_1}^{x_2} (x_2 - t)^{n-1} dt \right), \dots (7.6) \end{aligned}$$

La segunda integral en el lado derecho de (6.6) tiene el valor $(x_2 - x_1)^n / n$. Para la primera integral, miramos los tres casos $n = 1$, $n < 1$ y $n > 1$ por separado. En el primero

caso $n = 1$, el integrando se anula de forma idéntica y, por lo tanto, la integral tiene el valor cero.

En segundo lugar, para $n < 1$, tenemos $n-1 < 0$, y por lo tanto $(x_1 - t)^{n-1} \geq (x_2 - t)^{n-1}$. De este modo,

$$\begin{aligned} \int_0^{x_1} |(x_1 - t)^{n-1} - (x_2 - t)^{n-1}| dt &= \int_0^{x_1} ((x_1 - t)^{n-1} - (x_2 - t)^{n-1}) dt \\ &= \frac{1}{n} (x_1^n - x_2^n + (x_2 - x_1)^n) \leq \frac{1}{n} (x_2 - x_1)^n \end{aligned}$$

Finalmente, si $n > 1$ entonces $(x_1 - t)^{n-1} \leq (x_2 - t)^{n-1}$, y por lo tanto

$$\begin{aligned} \int_0^{x_1} |(x_1 - t)^{n-1} - (x_2 - t)^{n-1}| dt &= \int_0^{x_1} ((x_2 - t)^{n-1} - (x_1 - t)^{n-1}) dt \\ &= \frac{1}{n} (-x_1^n + x_2^n - (x_2 - x_1)^n) \leq \frac{1}{n} (x_2 - x_1)^n \end{aligned}$$

Una combinación de estos resultados produce

$$|(Ay)(x_1) - (Ay)(x_2)| = \begin{cases} \frac{2M}{\Gamma(n+1)} (x_2 - x_1)^n, & \text{si } n \leq 1 \\ \frac{2M}{\Gamma(n+1)} ((x_2 - x_1)^n + x_2^n - x_1^n), & \text{si } n > 1 \end{cases}, \dots (6.7)$$

En cualquier caso, la expresión del lado derecho de (6.7) converge a 0 cuando $x_2 \rightarrow x_1$ lo que prueba que Ay es una función continua. Además, para $y \in U$ y $x \in [0, h]$, encontramos

$$\begin{aligned} |(Ay)(x) - T(x)| &= \frac{1}{\Gamma(n)} \left| \int_0^x (x-t)^{n-1} f(t, y(t)) dt \right| \leq \frac{1}{\Gamma(n+1)} M x^n \\ &\leq \frac{1}{\Gamma(n+1)} M h^n \leq \frac{1}{\Gamma(n+1)} M \frac{K \Gamma(n+1)}{M} = K \end{aligned}$$

Así, hemos demostrado que $Ay \in U$ si $y \in U$, es decir, A mapea el conjunto U a sí mismo. Como queremos aplicar el Teorema del Punto Fijo de Schauder (Teorema D.9), todo eso, lo que queda ahora es mostrar que $A(U) := \{Au : u \in U\}$ es un conjunto relativamente compacto. Este se puede hacer mediante el Teorema de Arzelà-Ascoli (Teorema D.10). Para $z \in A(U)$ encontramos que, para todo $x \in [0, h]$,

$$\begin{aligned} |z(x)| &= |(Ay)(x)| \leq \|T\|_\infty + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt \\ &\leq \|T\|_\infty + \frac{1}{\Gamma(n+1)} M h^n \leq \|T\|_\infty + K \end{aligned}$$

que es la propiedad de acotación requerida. Además, la propiedad de equicontinuidad se puede derivar de (6.7) anterior. Específicamente, para $0 \leq x_1 \leq x_2 \leq h$, hemos encontrado en el caso $n \leq 1$ que

$$|(Ay)(x_1) - (Ay)(x_2)| \leq \frac{2M}{\Gamma(n+1)} (x_2 - x_1)^n$$

Por lo tanto, si $|x_2 - x_1| < \delta$, entonces

$$|(Ay)(x_1) - (Ay)(x_2)| \leq 2 \frac{M}{\Gamma(n+1)} \delta^n$$

Observando que la expresión del lado derecho es independiente de y , x_1 y x_2 , tenemos ver que el conjunto $A(U)$ es equicontinuo. De manera similar, en el caso de $n > 1$ podemos usar el Teorema del valor medio para concluir que

$$\begin{aligned} |(Ay)(x_1) - (Ay)(x_2)| &\leq \frac{M}{\Gamma(n+1)} ((x_2 - x_1)^n + x_2^n - x_1^n) \\ &= \frac{M}{\Gamma(n+1)} ((x_2 - x_1)^n + n(x_2 - x_1)\xi^{n-1}) \\ &\leq \frac{M}{\Gamma(n+1)} ((x_2 - x_1)^n + n(x_2 - x_1)h^{n-1}) \end{aligned}$$

con algún $\xi \in [x_1, x_2] \subseteq [0, h]$. Por lo tanto, si una vez más $|x_2 - x_1| < \delta$, entonces

$| (Ay)(x_1) - (Ay)(x_2) | \leq \frac{M}{\Gamma(n+1)} ((\delta)^n + n \delta h^{n-1})$ y el lado derecho es una vez más independiente de y , x_1 y x_2 , demostrando equicontinuidad. En cualquier caso, el teorema de Arzelà-Ascoli arroja que $A(U)$ es relativamente compacto, y por lo tanto el teorema del punto fijo de Schauder afirma que A tiene un punto fijo. Por construcción, un punto fijo de A es una solución de nuestro problema de valor inicial.

Notamos dos casos especiales importantes del Teorema 6.1. El primero de estos estados que, bajo ciertas suposiciones, la solución existe en todo el intervalo $[0, h^*]$ (y entonces para todo x para

el cual $f(x,y)$ está definido) y no solo para un subintervalo $[0,h]$ con alguna $h \leq h^*$.

Corolario 6.3

Suponga las hipótesis del Teorema 6.1, excepto que el conjunto G , es decir, el dominio de definición de la función f en el lado derecho de la diferencial ecuación (6.1A), ahora se toma como $G := [0, h^*] \times \mathbb{R}$. Además, suponemos que f es continua y que existen constantes $c_1 \geq 0, c_2 \geq 0$ y $0 \leq \mu < 1$, tales que

$$|f(x,y)| \leq c_1 + c_2 |y|^\mu \text{ para todo } (x,y) \in G, \dots (6.8)$$

Entonces, existe una función $y \in C[0, h^*]$ que resuelve el problema de valor inicial (6.1)

Prueba

Usamos el polinomio T definido en (6.4) en la prueba anterior. Dado que $\mu < 1$ podemos encontrar algún $K > 0$ tal que

$$c_1 + c_2 \left(K + \max_{x \in [0, h^*]} |T(x)| \right)^\mu \leq \frac{K\Gamma(n+1)}{h^*}$$

Usando este valor de K , restringimos nuestra función f al conjunto $G_K := \{(x,y) : x \in [0, h^*], |y - T(x)| \leq K\}$ (este es el conjunto que fue denotado por G en el Teorema 6.1). Entonces vemos eso

$$\begin{aligned} M &:= \sup_{(x,y) \in G_K} |f(x,y)| \leq c_1 + c_2 \sup_{(x,y) \in G_K} |y|^\mu \\ &\leq c_1 + c_2 \left(K + \max_{x \in [0, h^*]} |T(x)| \right)^\mu \leq \frac{K\Gamma(n+1)}{h^{*n}} \end{aligned}$$

Así podemos aplicar el Teorema 6.1 con este valor de K y el h^* dado y ver que $\left(\frac{K\Gamma(n+1)}{M} \right)^{1/n} \geq h^*$ lo que implica que

$$h^* = \min \left\{ h^*, \left(\frac{K\Gamma(n+1)}{M} \right)^{1/n} \right\} = h$$

Observación 6.6.

Son necesarios tres comentarios con respecto a la condición (6.8):

(a) Una condición suficiente sobre f para que se cumpla (6.8) es que f sea continua y acotado en G .

(b) En el teorema 6.1, se requería que f fuera una función continua en un conjunto compacto, y por lo tanto fue acotado automáticamente. En este corolario, f sigue siendo continua pero ahora se define en un conjunto no compacto, y por lo tanto tenemos que exigir un adecuado enlazado explícitamente.

(c) La condición (6.8) puede considerarse una restricción bastante severa ya que se viola incluso por algunos tipos de ecuaciones muy elementales y prácticamente importantes como, por ejemplo, ecuaciones lineales. Por lo tanto, las investigaciones adicionales son necesarias.

El segundo corolario del Teorema 6.1 afirma la existencia de una solución en todo el semieje $[0, \infty)$ en condiciones apropiadas.

Corolario 6.4

Suponga las hipótesis del Corolario 6.3, excepto que el conjunto G , es decir, el dominio de definición de la función f en el lado derecho de la diferencial ecuación (6.1A), ahora se toma como $G := \mathbb{R}^2$. Entonces, existe una función $y \in C[0, \infty)$ resolver el problema de valor inicial (6.1).

Por supuesto, la observación 5.6 también se aplica a este corolario.

Prueba

Sea $h^* > 0$. Bajo nuestras suposiciones, podemos aplicar el Corolario 6.3 para esta y concluya que existe una solución continua en $[0, h^*]$. Como h^* se puede elegir arbitrariamente grande, encontramos que existe una solución continua en $[0, \infty)$.

Unicidad de las soluciones

A continuación, llegamos a un teorema de unicidad que corresponde al conocido resultado del teorema de Picard- Lindel, puede verse como un análogo a la afirmación que se muestra para Riemann: Operadores de Liouville del capítulo anterior (Teorema 5.1).

Teorema 6.5

Sea $0 < n$ y $m = [n]$. Además, sea $y_0^{(0)}, y_0^{(1)}, \dots, y_0^{(m-1)} \in R, K > 0$

y . Defina el conjunto G como en el Teorema 6.1 y sea la función $f : G \rightarrow R$

continua y cumple una condición de Lipschitz con respecto a la segunda variable, es decir

$$|f(x, y_1) - f(x, y_2)| \leq L|y_1 - y_2|$$

con alguna constante $L > 0$ independiente de x , . Entonces, denotando h como en el teorema 6.1, existe una función definida unívocamente $y \in C[0, h]$ que resuelve el valor inicial problema (6.1).

Observación 6.7

En la Observación 6.3 que habíamos dicho anteriormente en relación con la existencia de soluciones también se aplica aquí en la discusión de cuestiones de unicidad.

Demostración (del teorema 6.5)

Primero observamos que el teorema 6.1 afirma que el valor inicial problema tiene solución. Para probar la unicidad de esta solución, comenzamos con argumentos similares a los de la demostración del Teorema 6.1. En particular, utilizamos el mismo polinomio T (definido en (6.4)) y el mismo operador A (definido en (6.5)) y recuerda que mapea el conjunto $U = \{y \in C[0, h] \text{ no vacío, convexo y cerrado: } \|y - T\|_\infty \leq K\}$ a sí mismo. Ahora tenemos que demostrar que A tiene un único punto fijo. Con el fin de hacer esto, probaremos primero que, para todo $J \in N_0$, todo $x \in [0, h]$ y todo $y, \tilde{y} \in U$, nosotros tenemos

$$\|A^J y - A^J \tilde{y}\|_{L_\infty[0, x]} \leq \frac{(Lx^n)^J}{\Gamma(1+nJ)} \|y - \tilde{y}\|_{L_\infty[0, x]}, \dots (6.9)$$

Esto se puede ver por inducción. En el caso $j = 0$, el enunciado es trivialmente verdadero. Para el paso de inducción $j-1 \rightarrow j$, escribimos

$$\begin{aligned} \|A^J y - A^J \tilde{y}\|_{L_\infty[0, x]} &= \|A(A^{J-1}y) - A(A^{J-1}\tilde{y})\|_{L_\infty[0, x]} \\ &= \frac{1}{\Gamma(n)} \sup_{0 \leq w \leq x} \left| \int_0^w (w-t)^{n-1} [f(t, A^{J-1}y(t)) - f(t, A^{J-1}\tilde{y}(t))] dt \right| \end{aligned}$$

Procedemos en el paso de inducción usando la suposición de Lipschitz sobre f y la hipótesis de inducción. Esto nos permite estimar las cantidades bajo consideración de la siguiente manera:

$$\begin{aligned}
\|A^j y - A^j \tilde{y}\|_{L_\infty[0,x]} &\leq \frac{L}{\Gamma(n)} \sup_{0 \leq w \leq x} \int_0^w (w-t)^{n-1} |A^{j-1}y(t) - A^{j-1}\tilde{y}(t)| dt \\
&\leq \frac{L^j}{\Gamma(n)} \int_0^x (x-t)^{n-1} \sup_{0 \leq w \leq t} |y(w) - \tilde{y}(w)| dt \\
&\leq \frac{L^j}{\Gamma(n)\Gamma(1+n(j-1))} \int_0^x (x-t)^{n-1} (t)^{n(j-1)} \sup_{0 \leq w \leq t} |y(w) - \tilde{y}(w)| dt \\
&\leq \frac{L^j}{\Gamma(n)\Gamma(1+n(j-1))} \sup_{0 \leq w \leq t} |y(w) - \tilde{y}(w)| \int_0^x (x-t)^{n-1} t^{n(j-1)} dt \\
&= \frac{L^j}{\Gamma(n)\Gamma(1+n(j-1))} \|y - \tilde{y}\|_{L_\infty[0,x]} \frac{\Gamma(n)\Gamma(1+n(j-1))}{\Gamma(1+n)} x^{nj}
\end{aligned}$$

Este es nuestro resultado deseado (6.9). Como consecuencia, encontramos, tomando las normas de Chebyshev sobre nuestro intervalo fundamental $[0, h]$, $\|A^j y - A^j \tilde{y}\|_\infty = \frac{(Lh^n)^j}{\Gamma(1+nj)} \|y - \tilde{y}\|_\infty$

Ahora hemos demostrado que el operador A cumple los supuestos del Teorema D.7 con $\alpha_j = \frac{(Lh^n)^j}{\Gamma(1+nj)}$. Para aplicar ese teorema, solo necesitamos verificar que la serie converge. Para ello notamos que $\sum_{j=0}^\infty \alpha_j$ con α_j como arriba es simplemente la representación en serie de potencias de la función Mittag-Leffler $E_n(Lh^n)$, y por lo tanto la convergencia requerida de la serie se sigue inmediatamente de Teorema 5.1. Por lo tanto, podemos aplicar el teorema del punto fijo de Weissinger y deducir la unicidad de la solución de nuestra ecuación diferencial.

Observación 6.8

Una observación similar a la realizada para el Riemann-Liouville El caso de la observación 5.2 también se cumple aquí: Las pruebas del teorema de unicidad 6.5 una vez de nuevo dar un método constructivo para encontrar una secuencia de iteraciones de Picard que converge contra la solución exacta del problema

de valor inicial. En particular, los elementos individuales de esta secuencia pueden ser considerados como aproximaciones para la solución exacta, al menos en teoría. Dado que no es necesariamente posible evaluar las integrales requeridas en forma cerrada, es poco probable que estas aproximaciones sean útiles numéricamente.

Observación 6.9.

Con respecto a la conexión entre la declaración de unicidad y la Condición de Lipschitz nos fijamos en la ecuación diferencial

$$D_{*0}^n y = y^\mu$$

Con $0 < \mu < 1$. Habíamos discutido el análogo de Riemann-Liouville de esta ecuación en el ejemplo. En la versión Caputo volvemos a utilizar condiciones iniciales homogéneas

$$y(0) = 0, y D_0^{n-k} y(0) = 0; k = 1, 2, 3, \dots, [n] - 1$$

y encuentre las mismas dos soluciones que en el caso de Riemann-Liouville, a saber.

$$y(x) = 0, y y(x) = \sqrt{\mu-1} \frac{\Gamma(j+1)}{\Gamma(j+1-n)} x^j$$

$$\text{Con } j = \frac{n}{1-\mu}$$

Una debilidad aparente en el enunciado del teorema 6.5 es que produce la existencia y unicidad de la solución del problema de valor inicial no en el intervalo $[0, h^*]$ de donde se permitió que viniera el primer argumento de la función dada f , pero solo en el intervalo posiblemente más pequeño $[0, h]$ con $h = \min\{h^*, (K\Gamma(n+1)/M)^{1/n}\}$. En este sentido, el Teorema 6.5 es completamente análogo a la

existencia de tipo Peano Teorema 6.1. Para este último, habíamos derivado los Corolarios 6.3 y 6.4 que dieron suficientes condiciones para que la solución exista en el intervalo completo $[0, h^*]$ o incluso en $[0, \infty)$. Los resultados correspondientes también se pueden mostrar para la pregunta de unicidad:

Corolario 6.6

Suponga las hipótesis del Teorema 6.5, excepto que el conjunto G , es decir, el dominio de definición de la función f en el lado derecho de la diferencial ecuación (6.1a), ahora se toma como $G := [0, h^*] \times R$. Además, suponemos que f es continua y que existen constantes $c_1 \geq 0; c_2 \geq 0$ y $0 \leq \mu < 1$ tales que $|f(x, y)| \leq c_1 + c_2|y|^\mu$ para todo $(x, y) \in G$.

Entonces, existe una función $y \in C[0, h^*]$ que resuelve el problema de valor inicial (6.1).

Corolario 6.7

Suponga las hipótesis del Corolario 6.6, excepto que el conjunto G , es decir, el dominio de definición de la función f en el lado derecho de la diferencial ecuación (6.1a), ahora se toma como $G := R^2$. Entonces, existe una función $y \in C[0, \infty)$ que resuelve el problema de valor inicial (6.1).

Las demostraciones de estos dos corolarios siguen exactamente las líneas de las demostraciones de los Corolarios.6.3 y 6.4, respectivamente. Omitimos los detalles.

Observación 6.10

Como en la observación 6.6, queremos comentar la condición (6.10):

(a) Dado que tenemos que suponer la continuidad de f de todos modos, una condición suficiente en f para que se satisfaga (6.10) es que f está acotada en G . Como G es ahora ilimitada, la acotación de f no es una consecuencia automática de su continuidad.

(b) La condición (6.10) puede considerarse una restricción bastante severa ya que se viola incluso por algunos tipos de ecuaciones muy elementales y prácticamente importantes como, por ejemplo, ecuaciones lineales.

La última observación aquí nos motiva a incluir otro teorema de existencia y unicidad. Utilizando un conjunto ligeramente diferente de supuestos que en realidad nos permiten derivar un resultado de existencia global y unicidad para una gran clase de ecuaciones que ahora incluye las ecuaciones lineales.

Teorema 6.8

Sea $0 < n$ y $m = [n]$. Además, sea $y_0^{(0)}, y_0^{(1)}, \dots, y_0^{(m-1)} \in \mathbb{R}$ y $h^* > 0$. Defina el conjunto $G := [0, h^*] \times \mathbb{R}^G$ y sea continua la función $f: G \rightarrow \mathbb{R}^n$ y cumpla una condición de Lipschitz con respecto a la segunda variable con una constante de Lipschitz $L > 0$ que es independiente de $x, y_1, \dots, y_m$. Entonces existe una función definida unívocamente $y \in C[0, h^*]$ resolviendo el problema de valor inicial (6.1).

En particular, este teorema es aplicable a ecuaciones lineales, es decir, ecuaciones de la forma

$D_{*0}^n y(x) = f(x)y(x) + g(x)$ con ciertas funciones $f, g \in C[0, h^*]$, porque aquí podemos elegir $L = \|f\|_\infty < \infty$

Obtenemos una consecuencia inmediata.

Corolario 6.9

Sea $0 < n$ y $m = [n]$. Además, sea $y_0^{(0)}, y_0^{(1)}, \dots, y_0^{(m-1)} \in R$ y $h^* > 0$.

Defina el conjunto $G := [0, h^*] \times R$ y sea continua la función $f: G \rightarrow \mathbb{R}$ y cumpla una condición de Lipschitz con respecto a la segunda variable con una constante de Lipschitz $L > 0$ que es independiente de x, y_1, y_2 . Entonces existe una función definida unívocamente $y \in C[0, \infty)$ resolviendo el problema de valor inicial (6.1).

Prueba

Sea $h^* > 0$. Bajo los supuestos del corolario, podemos aplicar el Teorema 6.8 en el intervalo $[0,]$ y concluir la existencia de una única solución continua en este intervalo. Dado que h^* puede elegirse arbitrariamente grande, derivamos la existencia y unicidad en todo el semieje $[0, \infty)$.

Para la demostración del Teorema 6.8 recogemos algunos resultados auxiliares

Lema 6.10

Suponga las hipótesis del Teorema 6.8. Además, denotamos $p(x, t) := (x - t)^{n-1} / \Gamma(n)$. Esta función p tiene las siguientes propiedades:

- (a) Para todo $x \geq 0$, $p(x, \cdot)$ es absolutamente integrable en $[0, x]$.
 (b) Para toda función $k \in C[0, h^*]$ y todo $\xi_1, \xi_2 \in [0, h^*]$, las expresiones

$$\int_{\xi_1}^{\xi_2} p(x, t) f(t, k(t)) dt \text{ y } \int_0^x p(x, t) f(t, k(t)) dt$$

Son continuas con respecto a x .

- (c) Existen números $0 = h_0 < h_1 < h_2 < \dots, < h_N = h^*$, tales que para todo $i \in$

$\{0, 1, \dots, N-1\}$ y todo $x \in [h_i, h_{i+1}^*]$ tenemos

$$L \int_{h_i}^{\min\{x, h_{i+1}^*\}} |p(x, t)| dt \leq \frac{1}{2}$$

- (d) Para cada $x \geq 0$,

$$\lim_{\delta \rightarrow 0^+} \int_x^{x+\delta} |p(x+\delta, t)| dt$$

Prueba

La parte (a) es una consecuencia inmediata de la definición de p y del hecho de que $n > 0$.

Para (b) consideraremos los dos casos $n \geq 1$ y $n < 1$. El primero es muy simple, ya que en este caso todas las funciones bajo la operación integral son continuas en sus respectivos dominios de definición, y por lo tanto la continuidad de las integrales es trivial.

En el otro caso podemos proceder como en la parte final de la demostración del Teorema 2.2 y usar los hechos de que p es integrable y $f(\cdot, k(\cdot))$ es continua para concluir el resultado.

Para la parte (c) distinguimos los mismos dos casos. En el caso de $n \geq 1$, p es continua y por lo tanto acotado en el conjunto compacto $\{(x, t): 0 \leq t \leq x \leq \cdot\}$. Así, eligiendo $N := [2h^*, L\|p\|_\infty]$ donde la norma de Chebyshev de p se toma sobre el conjunto mencionado, podemos definir $h_i := ih^*/N$. Esto implica $0 = h_0 < h_1 < h_2 < \dots, < h_N = h^*$ y

$$\begin{aligned}
L \int_{h_i}^{\min\{x, h_{i+1}\}} |p(x, t)| dt &\leq L \int_{h_i}^{h_{i+1}} |p(x, t)| dt \leq L \|p\|_{\infty} (h_{i+1} - h_i) \\
&= \frac{L \|p\|_{\infty} h^*}{N} \leq \frac{L \|p\|_{\infty} h^*}{2L \|p\|_{\infty} h^*} = \frac{1}{2}
\end{aligned}$$

según sea necesario.

En el caso de que $n < 1$ se procede de manera ligeramente diferente. Aquí nosotros usamos $N := \left[h^* \left(\frac{2L}{\Gamma(n+1)} \right)^{1/n} \right]$ y $h_i := ih^*/N$. Esto implica una vez más implica $0 = h_0 < h_1 < h_2 < \dots, < h_N = h^*$. Además, tenemos, para $x \leq h_{i+1}$

$$\begin{aligned}
L \int_{h_i}^{\min\{x, h_{i+1}\}} |p(x, t)| dt &= \frac{L}{\Gamma(n)} \int_{h_i}^{h_{i+1}} (x - h_i)^n dt = \frac{L}{\Gamma(n+1)} (x - h_i)^n \\
&\leq \frac{L}{\Gamma(n+1)} (h_{i+1} - h_i)^n = \frac{L}{\Gamma(n+1)} (h^*/N)^n \\
&\leq \frac{L}{\Gamma(n+1)} = \frac{\Gamma(n+1)}{2L} = \frac{1}{2}
\end{aligned}$$

Además, para $x \geq h_{i+1}$, encontramos

$$\begin{aligned}
\Psi(x) &:= L \int_{h_i}^{\min\{x, h_{i+1}\}} |p(x, t)| dt = \frac{L}{\Gamma(n)} \int_{h_i}^{h_{i+1}} (x - h_i)^{n-1} dt \\
&= \frac{L}{\Gamma(n+1)} [(x - h_i)^n - (x - h_{i+1})^n]
\end{aligned}$$

De este modo,

$$\Psi'(x) = \frac{L}{\Gamma(n+1)} [(x - h_i)^{n-1} - (x - h_{i+1})^{n-1}] < 0$$

ya que $n < 0$. Se sigue que, para $x \geq h_{i+1}$; $\Psi(x) \leq \Psi(h_{i+1}) \leq 1/2$ en vista de nuestro resultado arriba. Esto completa la prueba de (c) para $n < 1$ también.

Finalmente, para (d) vemos que

$$\int_x^{x+\delta} |p(x+\delta, t)| dt = \frac{1}{\Gamma(n+1)} \delta^n$$

lo cual, dado que $n > 0$, implica el resultado deseado.

Demostración (del Teorema 6.8)

Nuestro enfoque está inspirado en la demostración esbozada de Nkamnang (1999), Teorema 5.8. Recordemos los puntos del Lema 6.10 (c). Primero nos concentramos en el intervalo $[h_0, h_1]$. Comenzamos definiendo las funciones

$$y_0(x) := T(x) = \sum_{k=0}^{m-1} y_0^{(k)} \frac{x^k}{k!}$$

$$y_1(x) := T(x) + \frac{L}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y_{j-1}(t)) dt ; j = 1, 2, 3, \dots$$

y observamos que el Lema 6.10 (b) implica la continuidad de estas funciones en $[h_0, h_1]$. Además, definimos

$$\phi_j(x) := y_j(x) - y_{j-1}(x) ; j = 1, 2, 3, \dots$$

Y

$$\phi_0(x) := T(x) = y_0(x)$$

Evidentemente estas funciones también son continuas en $[h_0, h_1]$, y para $j = 0, 1, \dots$ es obvio que

$$y_j(x) = \sum_{\mu=0}^j \phi_\mu(x)$$

Además, para $j = 2, 3, \dots$ tenemos

$$\phi_j(x) = \frac{L}{\Gamma(n)} \int_0^x (x-t)^{n-1} (f(t, y_{j-1}(t)) - f(t, y_{j-2}(t))) dt$$

Para $x \in [h_0, h_1]$ entonces encontramos, usando (6.11), la condición de Lipschitz en f y el enunciado del Lema 6.10 (c)

$$\begin{aligned}
 |\phi_j(x)| &\leq \frac{L}{\Gamma(n)} \int_0^x (x-t)^{n-1} |f(t, y_{j-1}(t)) - f(t, y_{j-2}(t))| dt \\
 &\leq \frac{L}{\Gamma(n)} \int_0^x (x-t)^{n-1} |y_{j-1}(t) - y_{j-2}(t)| dt \\
 &\leq \frac{L}{\Gamma(n)} \max_{t \in [h_0, h_1]} |y_{j-1}(t) - y_{j-2}(t)| \int_0^x (x-t)^{n-1} dt \\
 &\leq \frac{1}{2} \max_{t \in [h_0, h_1]} |y_{j-1}(t) - y_{j-2}(t)| = \frac{1}{2} \max_{t \in [h_0, h_1]} |\phi_{j-1}(t)|
 \end{aligned}$$

Para $j=1, 2, 3, \dots$ lo que implica que

$$\max_{x \in [h_0, h_1]} |\phi_{j-1}(x)| \leq \frac{1}{2^{j-1}} \max_{t \in [h_0, h_1]} |\phi_1(t)|$$

Por lo tanto, hemos encontrado un mayorante convergente para la serie $\sum_{\mu=0}^{\infty} \phi_{\mu}$ en el intervalo $[h_0, h_1]$, y por lo tanto la serie es uniformemente convergente allí. Como es el j -ésima parcial suma de esta serie, se sigue que la sucesión también es uniformemente convergente en $[h_0, h_1]$, y los límites coinciden. Denotemos el límite de esta secuencia por y . Hemos visto anteriormente que y_j es continua en $[0, h^*]$, y por lo tanto en $[h_0, h_1]$, para todo j . Por lo tanto, en vista de la convergencia uniforme.

Nuestro próximo objetivo es mostrar que esta función resuelve la ecuación de Volterra (6.2), y por lo tanto el problema de valor inicial (6.1). Esto entonces, en particular, implicará la existencia de una solución continua. Para ello observamos que la convergencia uniforme de la sucesión $(y_j)_{j=1}^{\infty}$ y la propiedad de Lipschitz de f implica $|f(x, y(x)) - f(x, y_j(x))| \leq L|y(x) - y_j(x)| \rightarrow 0$ uniformemente para $x \in [h_0, h_1]$.

En otras palabras, la secuencia $(f(\cdot, y_j(\cdot)))_{j=0}^{\infty}$ converge uniformemente frente a $f(\cdot, y(\cdot))$. Así podemos intercambiar la operación límite y la integración fraccionaria. Esto produce

$$\begin{aligned} y(x) &= \lim_{j \rightarrow \infty} y_j(x) = \lim_{j \rightarrow \infty} (T(x) + J_0^n f(\cdot, y_{j-1}(\cdot))(x)) \\ &= T(x) + J_0^n (\lim_{j \rightarrow \infty} f(\cdot, y_{j-1}(\cdot))(x)) \\ &= \sum_{k=0}^{m-1} y_0^{(k)} \frac{x^k}{k!} + J_0^n f(\cdot, y_{j-1}(\cdot)) \end{aligned}$$

que es la relación requerida (6.2).

Para el intervalo $[h_0, h_1]$ queda probar que esta solución es única. Por lo tanto, suponemos que $\tilde{y}$ también es una solución de (6.2). Luego sigue

$$\begin{aligned} |y(x) - \tilde{y}(x)| &\leq \frac{L}{\Gamma(n)} \int_0^x (x-t)^{n-1} |(f(t, y_{j-1}(t)) - f(t, \tilde{y}_{j-2}(t)))| dt, \\ &\leq \frac{L}{\Gamma(n)} \int_0^x (x-t)^{n-1} |y(t) - \tilde{y}(t)| dt, \\ &\leq \frac{1}{2} \max_{t \in [h_0, h_1]} |y(t) - \tilde{y}(t)|. \end{aligned}$$

por el Lema 6.10 (c). Como esta desigualdad es uniforme para todo $x \in [h_0, h_1]$, tenemos deducir

$$\max_{x \in [h_0, h_1]} |y(x) - \tilde{y}(x)| \leq \frac{1}{2} \max_{t \in [h_0, h_1]} |y(t) - \tilde{y}(t)|.$$

lo que implica la declaración de unicidad requerida $y \equiv \tilde{y}$

Ahora necesitamos extender este resultado de existencia y unicidad de la inicial

Subintervalo $[h_0, h_1]$ a los subintervalos restantes $[h_{i-1}, h_i]$ ($i = 1, 2, 3, \dots, N$). Lo haremos esto de manera inductiva, usando nuestro resultado para el intervalo $[h_0, h_1]$ como base.

Por lo tanto, asumimos que la afirmación se cumple en $[h_{i-1}, h_i]$, para algún i y mostraremos que, si $i < N$, también se cum-

ple en $[h_i, h_{i+1}]$. Necesitamos demostrar que, para x en este intervalo, (7.2) tiene una única solución continua. En este contexto es conveniente reescribir la Ecuación de Volterra en cuestión en la forma

$$y(x) = T_i(x) + \frac{1}{\Gamma(n)} \int_{h_i}^x (x-t)^{n-1} f(t, y_{j-1}(t)) dt, \dots (6.12a)$$

$$T_i(x) = \sum_{k=0}^{m-1} y_0^{(k)} \frac{x^k}{k!} + \frac{1}{\Gamma(n)} \int_{h_i}^x (x-t)^{n-1} f(t, y_{j-1}(t)) dt, \dots (6.12b)$$

La observación esencial aquí es que la función es una función conocida porque su representación contiene solo datos dados y los valores de la solución y en el intervalo $[0, h_i]$ que ya ha sido calculado. Además, el Lema 6.10 (b) implica que T_i es continua en $[h_i, h_{i+1}]$. Recordando las definiciones de las funciones y entonces podemos proceder de una manera muy similar a nuestro enfoque anterior.

Específicamente definimos

$$y_0^{(i)}(x) := T_i(x)$$

Y

$$y_j^{(i)}(x) := T_i(x) + \frac{1}{\Gamma(n)} \int_{h_i}^x (x-t)^{n-1} f(t, y_{j-1}^{(i)}(t)) dt ; j = 1, 2, \dots$$

$$\phi_j^{(0)}(x) = y_j^{(i)}(x) - y_{j-1}^{(i)}(x), j=1, 2, \dots$$

y

$$\phi_0^{(i)}(x) := T_i(x) = y_0^{(i)}(x)$$

Evidentemente, por el Lema 6.10 (b), todas estas funciones son continuas en $[h_i, h_{i+1}]$ también, y para $j = 0, 1, \dots$ es obvio que

$$y_j^{(i)}(x) = \sum_{\mu=0}^j \phi_{\mu}^{(i)}(x)$$

Un mayorante convergente para $\sum_{\mu=0}^{\infty} \phi_{\mu}^{(i)}(x)$ se puede en-

contrar como arriba, lo que proporciona la existencia del límite uniforme $y := \lim_{j \rightarrow \infty} y_j^{(t)}$ en el intervalo . Lema 6.10 (d) implica que la función y , que hasta ahora ha sido definida por partes en los intervalos $[0, h_i]$ $[0, h_i]$ y $[h_i, h_{i+1}]$, respectivamente, es realmente continua en el punto y por lo tanto a lo largo de todo el intervalo $[0, h_{i+1}]$.

La prueba de que esta función resuelve la ecuación de Volterra y que cualquier otra solución continua de la ecuación de Volterra debe ser idéntica y sigue entonces lo mismo como líneas arriba.

De hecho, podemos encontrar la solución de una clase simple de fraccionario tipo Caputo ecuaciones diferenciales explícitamente, a saber, para la clase de ecuaciones lineales homogéneas con coeficientes constantes donde, como máximo, la condición inicial de orden más bajo es no homogéneo.

Teorema 6.11

Sea $n > 0, m = [n]$ y $\lambda \in \mathbb{R}$. La solución del problema de valor inicial

$$D_{+0}^n y(x) = \lambda y(x), y(0) = y_0, y^{(k)}(0) = 0, (k = 1, 2, \dots, m - 1)$$

es dado por

$$y(x) = y_0 E_n(\lambda y(x^n)), x \geq 0$$

La ecuación diferencial bajo consideración en el Teorema 6.11 es muy simple ejemplo de una ecuación diferencial fraccionaria lineal. Daremos una información más detallada investigación de ecuaciones lineales más generales en la Secc. 7.1.

Prueba

Es evidente a partir de nuestra existencia y resultado único en el Corolario 6.9 anterior que el problema de valor inicial tiene solución única. Por lo tanto, solo tenemos que verificar que la función y indicada anteriormente es una solución. Para la condición inicial, de hecho, vemos que $y(0) = y_0$ $E_n(0) = y_0$, ya que

$$E_n(z) = 1 + \frac{z}{\Gamma(1+n)} + \frac{z^2}{\Gamma(1+2n)} + \dots$$

además, en el caso $m \geq 2$ (es decir, $n > 1$) tenemos $y^{(k)} = 0$ para $k = 1, 2, \dots, m-1$ desde

$$y(x) = 1 + \frac{\lambda x^n}{\Gamma(1+n)} + \frac{\lambda^2 x^{2n}}{\Gamma(1+2n)} + \dots$$

Lo que implica,

$$y^{(k)}(x) = 1 + \frac{\lambda x^{n-k}}{\Gamma(1+n-k)} + \frac{\lambda^2 x^{2n-k}}{\Gamma(1+2n-k)} + \dots$$

para $k = 1, 2, \dots, m-1 < n$.

Finalmente, el hecho de que nuestra función y realmente resuelva la ecuación diferencial es una consecuencia inmediata del teorema 4.3.

Para concluir esta sección, recordamos otro resultado clásico de la teoría de ecuaciones diferenciales de primer orden y observe las posibles generalizaciones de este resultado al ajuste fraccionario. Es un teorema bien conocido estrechamente relacionado con la pregunta de unicidad.

Teorema 6.A

Sea continua y satisfaga una condición de Lipschitz con respecto a la segunda variable. Considere dos soluciones y de la ecuación diferencial

$$D^1 y_j = f(x, y_j(x)), j = 1, 2$$

sujeto a las condiciones iniciales $y_j(x_j) = y_j(0)$, respectivamente. Entonces las funciones y coinciden en todas partes o en ninguna.

Prueba

La demostración es muy simple: suponga que las funciones y coinciden en algún punto, es decir, $y_1(x^*) = y_2(x^*) =: y^*$, digamos. Entonces, ambas funciones resuelven el problema de valor inicial $D^1 y_j = f(x, y_j(x)), y_j(x^*) = y^*$. Dado que los supuestos afirman que este problema tiene solución única, y deben ser idénticas, es decir, coinciden en todos lados.

Interpretemos el enunciado del Teorema 6.A. Para ello comenzamos con una solución y_1 de la ecuación diferencial (6.13) con alguna condición inicial $y_1(x_1) = y_1(0)$. Entonces podemos elegir una abscisa arbitraria y prescribir una condición inicial $y_2(x_2) = y_2(0)$ en esta abscisa. Puede ocurrir que el punto se encuentre en la gráfica de . En este caso, las funciones y_1 y y_2 son idénticas. De lo contrario, los gráficos de y nunca se encontrarán ni se cruzarán. Específicamente llamamos la atención del lector sobre el hecho de que en el Teorema 6.A no nos importa si los valores x_1 y x_2 es decir, las abscisas donde las condiciones iniciales se especifican coincidan o no entre sí.

Se ha planteado la cuestión de una generalización fraccionaria de este teorema y respondido parcialmente (Diethelm, 2009). Se ha hecho un intento de dar una respuesta completa (Lubich, 2006), el Teorema 3.27. Sin embargo, una inspección minuciosa de la teoría desarrollada allí revela en una laguna la prueba. Por lo tanto, se ha desarrollado un enfoque alternativo de forma ligeramente configuración más general (Diethelm et al., 2002). Ahora describiremos la especialización de los resultados de ese documento a nuestra clase de problemas. Resulta que esto nos permite proporcionar una respuesta bastante satisfactoria para nuestra pregunta.

Como primer paso hacia la respuesta indicada a nuestra pregunta notamos que en la prueba del Teorema 6.A depende fuertemente de la localidad del operador diferencial D_I . Ya habíamos señalado en la observación 6.4 que tal propiedad no está disponible en el caso fraccionario, y por lo tanto la prueba no se puede llevar a cabo directamente. Además, como consecuencia de esta no-localidad, encontramos que ahora sí importa si las abscisas de las condiciones iniciales coinciden o no. De hecho, el siguiente ejemplo, tomado de Diethelm (2009), muestra que no podemos esperar un resultado comparable al Teorema 5.A mantener si $x_1 = x_2$:

Ejemplo

Sea $0 < n < 1$ y considere las ecuaciones diferenciales fraccionarias

$$D_{*0}^n y_1(x) = \Gamma(1+n), y_1(0) = 0, \dots \quad (6.14)$$

Y

$$D_{*0}^n y_2(x) = \Gamma(1+n), y_2(0) = 0, \dots \quad (6.15)$$

Aquí tenemos dos ecuaciones diferenciales con lados derechos idénticos, pero con condiciones inicial especificadas en diferentes puntos. Las soluciones se ven fácilmente como $y_1(x) = x^n$ y $y_2(x) = 1 + (x - 1)^n$. Es obvio que estas dos funciones coinciden en $x = 1$ pero en ningún otro lugar.

El teorema 6.A trata con ecuaciones diferenciales de orden 1, es decir, con ecuaciones sujeto a exactamente una condición inicial. Es bien sabido que una afirmación similar no nos vale para ecuaciones de orden superior (por ejemplo, la ecuación $D^2y(x) = -y(x)$ tiene soluciones $y(x) = \cos x$; $y(x) = \sin x$ y $y(x) = 0$ las dos primeras de las cuales oscilan y cuyas gráficas se cruzan entre sí). Efectos similares surgen en el contexto de ecuaciones diferenciales fraccionarias con más de una condición inicial (es decir, ecuaciones de orden $n > 1$); ver Secc. 7.1 a continuación. Por lo tanto, no podemos esperar más que el siguiente resultado para mantener:

Teorema 6.12.

Sea $0 < n < 1$ y suponga que $f : [0, b] \times [c, d] \rightarrow \mathbb{R}$ es continua y satisfacer una condición de Lipschitz con respecto a la segunda variable. Considere dos soluciones y_j de la ecuación diferencial

$$D_{*0}^n y_j(x) = f(x, y_j(x)), j = 1, 2$$

sujeto a las condiciones iniciales $y_j(0) = y_j 0$, respectivamente, donde $y_1(0) \neq y_2(0)$. Entonces para todo x donde existen tanto $y_1(x)$ como $y_2(x)$, tenemos $y_1(x) \neq y_2(x)$.

Prueba

Supongamos que existe algún x^* tal que $y_1(x^*) = y_2(x^*)$. Nosotros entonces tenemos que demostrar que $y_1 0 = y_2 0$. Con este fin, sea L la constante de Lipschitz de f con respecto a la segunda variable.

Entonces podemos elegir un número $\tau > 0$ que satisfaga las condiciones

$$\gamma := \frac{2L \tau^n}{\Gamma(n+1)} < 1$$

y $N := \frac{x^*}{\tau} \in \mathbb{N}$ y dividimos el intervalo fundamental $[0, x^*] = [0, N\tau]$ en los subintervalos $[(j-1)\tau, j\tau]$, $j = 1, 2, \dots, N$

Comenzamos observando el primero de estos subintervalos, es decir, en el intervalo $[0, \tau]$.

Nuestro objetivo aquí es mostrar que el problema que consiste en la ecuación diferencial (7.16) y la condición adicional $y(\tau) = y^*$ no puede tener más de una solución.

Para probar esto, recordemos que cualquier solución de la ecuación diferencial debe satisfacer la Ecuación de Volterra

$$y(x) = y_0 + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt, \dots (6.18)$$

con un número adecuado y_0 . La condición adicional en el punto τ implica entonces que

$$y^*(x) = y(\tau) + y_0 + \frac{1}{\Gamma(n)} \int_0^\tau (\tau-t)^{n-1} f(t, y(t)) dt$$

Es decir

$$y_0 = y^*(x) - \frac{1}{\Gamma(n)} \int_0^\tau (\tau-t)^{n-1} f(t, y(t)) dt$$

Insertando esto en (7.18) obtenemos

$$\begin{aligned} y(x) &= Ay(x): \\ &= y^*(x) \\ &+ \frac{1}{\Gamma(n)} \left(\int_0^x (x-t)^{n-1} f(t, y(t)) dt - \int_0^\tau (\tau-t)^{n-1} f(t, y(t)) dt \right) \end{aligned}$$

Ahora debemos demostrar que el operador A definido aquí no tiene más que un punto fijo. Es claro que A mapea $C[0, \tau]$ a sí mismo, y además encontramos eso

$$\begin{aligned} |A\tilde{y}(x) - A\tilde{y}(x)| &= \frac{1}{\Gamma(n)} \left(\int_0^x (x-t)^{n-1} |f(t, \tilde{y}(t)) - f(t, \tilde{y}(t))| dt \right. \\ &\quad \left. - \int_0^\tau (\tau-t)^{n-1} |f(t, \tilde{y}(t)) - f(t, \tilde{y}(t))| dt \right) \\ &\leq \frac{L}{\Gamma(n)} \|\tilde{y} - \tilde{y}\|_{L_\infty[0, \tau]} \left(\int_0^x (x-t)^{n-1} dt + \int_0^\tau (\tau-t)^{n-1} dt \right) \\ &\leq \frac{2L\tau^n}{\Gamma(n+1)} \|\tilde{y} - \tilde{y}\|_{L_\infty[0, \tau]} \\ &= \gamma \|\tilde{y} - \tilde{y}\|_{L_\infty[0, \tau]} \end{aligned}$$

y de (6.17) vemos que A es una contracción.

Así, por el teorema de punto fijo de Banach, obtenemos la unicidad deseada de la solución de la ecuación diferencial (6.16) sujeto a una condición de la forma $y(\tau) = y^*$.

Como consecuencia de esta observación, concluimos que dos soluciones de (6.16) que coinciden en el punto τ deben también coincidir en todo el intervalo $[0, \tau]$ y, por lo tanto, en particular, en el punto 0.

Por lo tanto, si $N = 1$, la prueba está completa. En el caso $N > 1$ queda por demostrar para $j = 2, 3, \dots, N$ que dos soluciones de (6.16) que coinciden en el punto deben en realidad ser idéntico en

el intervalo anterior $[(j-1)\tau, j\tau]$ y, por lo tanto, en la cuadrícula anterior punto $(j-1)\tau$.

Una vez que hemos demostrado que esto es cierto, podemos decir que nuestra suposición

$$y_1(N\tau) = y_1(x^*) = y_2(x^*) = y_2(N\tau)$$

implica que las soluciones y coinciden en $[(N-1)\tau, N\tau], [(N-2)\tau, (N-1)\tau], \dots [0, \tau]$, que es nuestro resultado deseado. Para esta parte restante de la prueba, nos fijamos en la ecuación diferencial dada (6.16) en el intervalo $[0, j\tau]$, sujeto a la condición $y(j\tau) = y^*$.

Como arriba, encontramos que esto es equivalente a escribir

$$y(x) = y^*(x) + \frac{1}{\Gamma(n)} \left(\int_0^x (x-t)^{n-1} f(t, y(t)) dt - \int_0^{j\tau} (j\tau-t)^{n-1} f(t, y(t)) dt \right)$$

Como se indicó anteriormente, estamos interesados en este problema en el intervalo $[(j-1)\tau, j\tau]$. Entonces, para x en este intervalo, reescribimos la identidad como

$$\begin{aligned} y(x) &= B_j(x) : \\ &= g_j(x) \\ &+ \frac{1}{\Gamma(n)} \left(\int_{(j-1)\tau}^x (x-t)^{n-1} f(t, y(t)) dt \right. \\ &\quad \left. - \int_{(j-1)\tau}^{j\tau} (j\tau-t)^{n-1} f(t, y(t)) dt \right) \end{aligned}$$

Donde

$$g_j(x) = y^*(x) + \frac{1}{\Gamma(n)} \int_0^{(j-1)\tau} ((x-t)^{n-1} - (j\tau-t)^{n-1}) f(t, y(t)) dt$$

(observe que esta cantidad está bien definida porque solo usa valores de y en el intervalo $[0, (j-1)\tau]$ y ya sabemos qué y está

determinado de forma única allí), entonces podemos nuevamente proceder como en el primer paso de nuestra argumentación y concluya que $B_f(x)$ es una contracción que conduce a la propiedad de unicidad requerida.

Ejemplo 6.2

Verificamos el enunciado del teorema 6.12 observando la ecuación diferencial

$$D_{*0}^{0.28} y_f(x) = (0.5 - x) \sin y_f(x) + 0.8x^3$$

Con condiciones iniciales

$$y_1(0) = 1.6 ; y_2(0) = 1.5 ; y_3(0) = 1.4 ; y_4(0) = 1.3$$

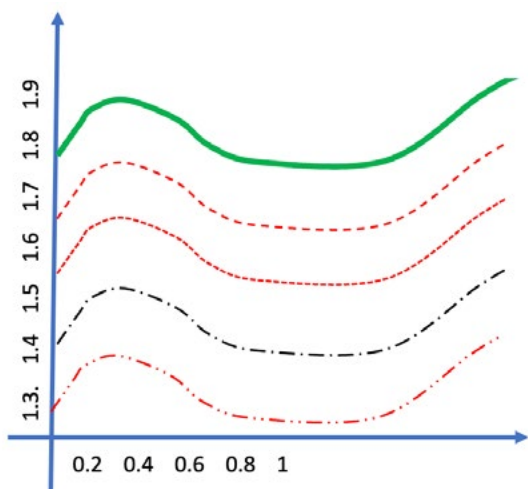
Dado que un método de aplicación general para determinar las soluciones analíticas de nuestros problemas de valor inicial no está disponible, tenemos que volver a algunas soluciones numéricamente aproximadas calculadas.

Con este fin, podemos utilizar el Adams-Bashforth- Método de Moulton descrito en el Apéndice C.1.

Las soluciones resultantes para los cinco problemas de valor inicial, obtenidos utilizando un tamaño de paso de $1/200$, se muestran en la figura 1.

La propiedad predicha por el Teorema 6.12 es evidente; en vista de la convergencia.

Figura 1. Gráficas de las soluciones de los cinco problemas de valor inicial del Ejemplo



Fuente: elaboración propia.

Nota. Las soluciones resultantes para los cinco problemas de valor inicial, obtenidos utilizando un tamaño de paso de 1/200 análisis para el método numérico proporcionado en el Apéndice C.1 podemos estar seguros de que esto muestra una imagen cualitativamente correcta de las soluciones exactas también.

Una formulación alternativa obvia del Teorema 6.12 dice lo siguiente.

Teorema 6.13

Sea $0 < n < 1$ y suponga que $f : [0, b] \times [c, d] \rightarrow \mathbb{R}$ es continua y satisface una condición de Lipschitz con respecto a la segunda variable. Considere dos soluciones y_j de la ecuación diferencial

$$D_*^n y_j(x) = f(x, y_j(x)), (j = 1, 2)$$

sujeto a las condiciones iniciales $y_j(0) = y_{j0}$ y, respectivamente. Sea $b^* \in (0, b]$ tal que ambas soluciones y_1 y y_2 existen en $[0, b^*]$. Entonces, ya sea $y_1(x) = y_2(x)$ para todo $x \in [0, b^*]$ o $y_1(x) = y_2(x)$ para todo $x \in [0, b^*]$.

Como resultado, también podemos reformular este resultado de una manera ligeramente diferente.

Presentamos aquí esta reformulación porque también es de interés en un contexto diferente. más tarde.

Teorema 6.14

Sea $0 < n < 1$ y suponga que $g : [0, b] \times [y^* - K, y^* + K] \rightarrow \mathbb{R}$ continua y satisface una condición de Lipschitz con respecto a la segunda variable.

Además, sea $g(x, 0) = 0$ para todo $x \in [0, b]$, y sea $y^* \neq 0$. Entonces, la solución y de la Ecuación de Volterra

$$y(x) = y^* = \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} g(t, y(t)) dt, \dots (6-19)$$

satisface $y(x) \neq 0$ para todo x .

La relación entre los teoremas 6.12 y 6.14 se describe fácilmente:

Teorema 6.15

El teorema 6.12 se cumple si y solo si se cumple el teorema 5.14.

Prueba.

En esta demostración seguimos las líneas. Primero demostramos que el teorema 6.12 implica el Teorema 6.14. Para ello recordemos que (6.19) es equivalente al problema de valor inicial fraccionario

$$D_{*0}^n y(x) = g(x, y(x)), y(0) = y^* \neq 0$$

Nuestra suposición de que $g(x, 0) = 0$ para todo x implica que la función $z(x) = 0$ resuelve el problema del valor inicial

$$D_{*0}^n z(x) = g(x, z(x)), z(0) = 0$$

Por lo tanto, si el Teorema 6.12 es verdadero, podemos concluir que $y(x) = z(x) = 0$ para todo x . Teorema, Por lo tanto, 6.14 también es cierto.

Para la otra dirección, sean y las soluciones de los problemas de valor inicial mencionado en el Teorema 6.12, y define

$$g(x, z) = f(x, z + y_1(x))$$

Además, establecemos $y(x) := y_2(x) - y_1(x)$ y $y^* := y_{20} - y_{10} = 0$. Ahora reescribimos los dos problemas de valor inicial del Teorema 6.12 en su correspondiente forma de Volterra, a saber.

$$y_j(x) = y_{j0} + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y_j(t)) dt, (j=1,2)$$

Restando estas dos ecuaciones, encontramos

$$\begin{aligned} y(x) &= y_2(x) - y_1(x) \\ &= y_{20} - y_{10} + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} [f(t, y_2(t)) - f(t, y_1(t))] dt \\ &= y^* + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} g(t, y_2(t) - y_1(t)) dt \\ &= y^* + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} g(t, y(t)) dt \end{aligned}$$

es decir, y resuelve la ecuación de Volterra (6.19). Dado que, para todo x , $g(x,0) = 0$, podemos aplicar Teorema 6.14. para concluir que $0 \neq y(x) := y_2(x) - y_1(x)$ para todo x . Una consecuencia trivial de nuestros resultados anteriores es

Corolario 6.16

Suponga las hipótesis del Teorema 6.12. Además, sea $y_{10} < y_{20}$. Entonces, $y_1(x) < y_2(x)$ para todo x para el cual existen ambas soluciones.

Terminamos esta sección echando un vistazo a este problema desde un punto de vista diferente. Específicamente del Teorema 7.12 podemos deducir el siguiente resultado.

Teorema 6.17

Sea $0 < n < 1$ y suponga que $f : [0,b] \times [c,d] \rightarrow \mathbb{R}$ es continua y satisfacer una condición de Lipschitz con respecto a la segunda variable. Entonces, para cada y y cada x , la ecuación diferencial

$$D_{*0}^n y(x) = f(x, y(x)), \dots (7.20)$$

Sujeta a la condición

$$y(x^*) = y^*, \dots (6.21)$$

tiene como máximo una solución.

Así tenemos, un teorema de unicidad para las soluciones de un diferencial fraccionario ecuación de la forma usual combinada con un valor prescrito de la solución desconocida en un punto que puede diferir del punto de partida del diferencial fracciona-

rio operador. En el caso de $x^* = 0$, esta es solo la condición inicial estándar que habíamos discutido a fondo al principio de esta sección, pero si $x^* > 0$ entonces tenemos un problema significativamente diferente que a veces se denomina problema de valor terminal porque normalmente uno está interesado en la solución en el intervalo $[0, x^*]$, es decir, uno proporciona una condición sobre la solución desconocida en el punto terminal del intervalo de interés. El siguiente teorema de equivalencia simple aclara la diferencia en carácter entre los problemas de valor inicial y terminal.

Teorema 6.18

Suponga las hipótesis del Teorema 6.17 y sea $0 < x^* \leq b$. Entonces, la ecuación diferencial fraccionaria (6.20) combinada con la condición (6.21) es equivalente a la ecuación integral débilmente singular

$$y(x) = y^* + \frac{1}{\Gamma(n)} \int_0^{x^*} G(x, t) f(t, y(t)) dt, \dots (6.22)$$

Donde

$$G(x, t) = \begin{cases} -(x^* - t)^{n-1}, & \text{para } t > x \\ (x - t)^{n-1} - (x^* - t)^{n-1}, & \text{para } t \leq x \end{cases}, \dots (6.23)$$

La diferencia sustancial entre la ecuación integral (6.22) y la ecuación integral (6.2) derivada del Lema 6.2 en el caso $x^* = 0$ es que ahora tenemos una ecuación integral Fredholm de tipo Hammerstein mientras que (6.2) era una ecuación de Volterra. Por lo tanto, en analogía con los resultados correspondientes para ecuaciones de orden entero, sería ser natural interpretar la condición (6.21) como una condición de contorno y no como una inicial condición. Por lo tanto, en contraste con la situación observada para

las ecuaciones diferenciales de primer orden, el problema de valor terminal que consiste en ecuaciones (6.20) y (6.21) es mucho más estrechamente relacionado con un problema de valores en la frontera que con un problema de valores iniciales.

Esto también es bastante natural, ya que necesitamos encontrar una solución en el intervalo $[0, x^*]$ cuyos puntos finales juegan un papel importante en la definición del problema: a la izquierda el punto final se utiliza para definir el operador diferencial y el punto final derecho proporciona la condición adicional que afirma la unicidad de la solución.

Proporcionaremos algunos resultados adicionales sobre problemas de valores en la frontera para Ecuaciones diferenciales fraccionarias tipo Caputo en la Secc. 6.5.

Prueba

Supongamos primero que la ecuación diferencial fraccionaria (6.20) combinada con la condición (6.21) tiene solución $y \in C[0, x^*]$. Entonces, podemos aplicar el operador integral a ambos lados de (6.20), lo que da

$$y(x) = y(0) + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt, \dots (6.24)$$

para todo $x \in [0, x^*]$ (Observe que la cantidad $y(0)$ es desconocida). Específicamente estableciendo $x = x^*$ y usando la condición (7.21), obtenemos

$$y^* = y(x^*) = y(0) + \frac{1}{\Gamma(n)} \int_0^{x^*} (x^* - t)^{n-1} f(t, y(t)) dt$$

Tabla 1. Resultados de la búsqueda de bisección para el Ejemplo 6.3

1	2	1.5	1.75	1.625	1.6875	1.71875
2.0556	2.6385	2.37738	2.51106	2.44567	2.47871	2.49496

Fuente: elaboración propia.

Nota. Da continuidad al Ejemplo 6.3 lo que implica.

$$y(0) = y^* - \frac{1}{\Gamma(n)} \int_0^{x^*} (x^* - t)^{n-1} f(t, y(t)) dt$$

Insertando esta relación en (6.24), obtenemos (6.22) y (6.23).

Por otra parte, si (6.22) con el núcleo G definido como en (6.23) tiene una solución continua entonces podemos aplicar el operador de Caputo $\bar{D}_{*0}^n$ a ambos lados de (6.22) que produce la ecuación diferencial fraccionaria (6.20). Finalmente, inmediatamente obtenga (6.21) haciendo $x = x^*$ en (6.22).

Para concluir esta sección, veamos un ejemplo simple de un problema de valor terminal

Ejemplo

Encuentre una solución al problema del valor terminal

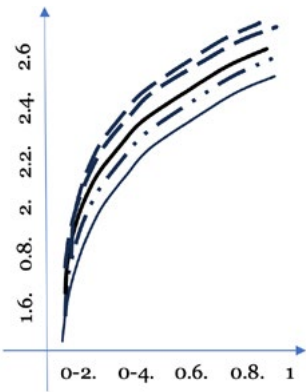
$$D_{*0}^{\frac{1}{2}} y(x) = \sin y(x). y(1) = 2.5$$

Una vez más, no podemos ofrecer un método analítico para encontrar la solución exacta; por lo tanto, volveremos a una técnica de aproximación. Empezamos eligiendo dos valores iniciales arbitrarios, digamos y , para la ecuación diferencial dada y calcular las soluciones correspondientes numéricamente en el intervalo $[0,1]$ por medio del método de Adams del Apéndice C.1 con un tamaño de paso de $1/200$. Resulta (ver Tabla 1) que una de estas opciones para los valores iniciales conduce a un valor de $y(1)$ que

es menor de lo deseado mientras que el otro produce un valor demasiado grande. Por lo tanto, empleamos una técnica de bisección simple para encontrar un nuevo valor inicial, y_0 , y calcular $y(l)$ de nuevo. Procediendo de manera iterativa como se indica en la Tabla 1 encontramos que podemos acercarnos tanto como queramos a la solución exacta deseada.

La Figura 2 muestra una visualización de este procedimiento. Hemos mostrado cinco de las soluciones vecinas de nuestra solución exacta, es decir, las obtenidas para $y_0 = 2$ (línea punteada; superior), $y_0 = 1.5$ (línea punteada y discontinua; segunda desde arriba), $y_0 = 1.625$ (línea punteada y de puntos dobles; segunda desde abajo) y $y_0 = 1.71875$ (línea continua; centro).

Figura 2. Gráficas de soluciones aproximadas vecinas para el problema de valor terminal de Ejemplo 6.3



Fuente: elaboración propia.

Nota. Es una muestra de la visualización del procedimiento de soluciones aproximadas vecinas para el problema de valor terminal de Ejemplo 6.3

Influencia de los datos perturbados

Habiendo establecido criterios sobre la existencia y unicidad de las soluciones a los problemas de valor inicial para ecuaciones diferenciales fraccionarias tipo Caputo, en este y los siguientes ítems llegamos a las cuestiones clásicas relativas a las más importantes propiedades de las soluciones. Estos incluyen, en particular, preguntas sobre la suavidad, continuidad y diferenciabilidad de las soluciones bajo varios supuestos sobre los datos proporcionados, así como investigaciones relativas a la buena postura de un problema de valor. Comenzamos con este último y recordamos los resultados clave en esta dirección.

Tradicionalmente, un problema se llama bien planteado si tiene las siguientes tres propiedades:

1. Existe una solución
2. La unicidad de la solución
3. La solución depende de los datos proporcionados de forma continua

Los dos primeros aspectos ya han sido tratados en los apartados anteriores; El tercero requiere más atención. En particular, notamos una diferencia importante entre el ajuste fraccionario y clásico, y ese es el significado preciso de la expresión “los datos proporcionados”.

En la teoría clásica de las ecuaciones diferenciales de orden entero, por lo general se asumen los valores iniciales y la función f del lado derecho de la ecuación diferencial

$$D^k y(x) = f(x, y(x), Dy(x), \dots, D^{k-1}y(x))$$

Dada y, luego el comportamiento de la solución bajo perturbaciones de estos. Se discuten las expresiones. Por supuesto, tenemos los mismos datos dados en la configuración fraccionaria, pero aquí se debe tener en cuenta un problema adicional: en una típica aplicación (Rudolf Gorenflo et al., 2002), una ecuación diferencial de la forma

$$D_{*0}^n y(x) = f(x, y(x))$$

surge aquí los parámetros principales de la ecuación, a saber, el valor o valores iniciales, posiblemente también el lado derecho f , y el orden del operador diferencial, dependen en constantes materiales que solo se conocen hasta cierta precisión, generalmente moderada. Por ejemplo, en los problemas que surgen en la viscoelasticidad considerados brevemente en Diethelm & Freed (1999b); Diethelm & Luchko (2004) y de forma más detallada en Rudolf Gorenflo et al. (2002), el conocimiento de los valores de n es solo impreciso y típicamente restringido a alrededor de dos dígitos decimales. Por tanto, es importante investigar cómo la solución depende también de este parámetro.

A lo largo de esta sección, asumimos que y es la solución exacta del problema de valor inicial

$$D_{*0}^n y(x) = f(x, y(x)), \dots (7.25a).$$

$$D^k y(0) = y_0^k, K = 0, 1, 2, \dots, m-1, \dots (7.25b).$$

donde como siempre hemos establecido $m = [n]$. Además, supongamos que f es tal que se satisfacen las hipótesis del teorema 7.5, de modo que podemos estar seguros de que solución existe en algún intervalo $[0, h]$, y que esta solución es única. Nosotros entonces tenemos que comparar esta solución y con la solución de otro problema de valor inicial involucrando de datos dados perturbados.

Resulta que de hecho seremos capaces de probar un resultado positivo en todos los casos: Pequeñas perturbaciones en cualquiera de los datos dados producen pequeñas perturbaciones en la solución.

Como requisito previo que encontraremos útil en la mayoría de las siguientes demostraciones, presentamos una desigualdad tipo Gronwall (Momani, 2000). Para el caso $0 < n < 1$ es (Dzherbashian & Nersesian, 2021) (Dzherbashian & Nersesian, 2021), el Teorema 5.1, pero aquí veremos el caso completamente general.

Lema 6.19

Sean $n, T, \varepsilon_1, \varepsilon_2 \in \mathbb{R}_+$. Además, suponga que $\delta : [0, T] \rightarrow \mathbb{R}$ es una función continua que satisface la desigualdad

$$|\delta(x)| \leq \varepsilon_1 + \frac{\varepsilon_2}{\Gamma(n)} \int_0^x (x-t)^{n-1} |\delta(t)| dt$$

para todo $x \in [0, T]$. Entonces

$$|\delta(x)| \leq \varepsilon_1 E_n(\varepsilon_2 x^n)$$

Para $x \in [0, T]$

Aquí, se denota una vez más la función Mittag-Leffler de orden n .

Prueba

Sea $\varepsilon_3 > 0$ e introduzca la función Φ con $\Phi(x) = (\varepsilon_1 + \varepsilon_3)E_n(\varepsilon_2 x^n)$. Aplicando el Teorema 7.11 y usando la linealidad del problema de valor inicial considerado ahí, esta función Φ se ve como la solución del problema de valor inicial

$D_{*0}^n \Phi(x) = \varepsilon_2 \Phi(x)$ con $\Phi(0) = \varepsilon_1 + \varepsilon_2$ y $D^k \Phi(0) = 0$, para $k = 1, 2, \dots, [n] - 1$

En vista del Lema 7.2 deducimos inmediatamente que Φ satisface la ecuación integral

$$\Phi(x) = \varepsilon_1 + \varepsilon_3 + \frac{\varepsilon_2}{\Gamma(n)} \int_0^x (x-t)^{n-1} \Phi(t) dt$$

Por nuestra suposición sobre δ , encontramos que

$$|\delta(0)| \leq \varepsilon_1 < \varepsilon_1 + \varepsilon_3 = \Phi(0)$$

Los argumentos de continuidad estándar producen que $|\delta(x)| < \Phi(x)$ para todo $x \in [0, \eta]$ con algún $\eta > 0$. Para probar que esta desigualdad se cumple en $[0, T]$, primero asumimos lo contrario y denotamos por el número positivo más pequeño con la propiedad de que $|\delta(x_0)| = \Phi(x_0)$. Entonces, para $0 \leq x \leq x_0$ tenemos $|\delta(x)| \leq \Phi(x)$ y por lo tanto

$$\begin{aligned} |\delta(x_0)| &\leq \varepsilon_1 + \frac{\varepsilon_2}{\Gamma(n)} \int_0^{x_0} (x_0 - t)^{n-1} |\delta(t)| dt \\ &\leq \varepsilon_1 + \frac{\varepsilon_2}{\Gamma(n)} \int_0^{x_0} (x_0 - t)^{n-1} \Phi(t) dt \\ &< \varepsilon_1 + \varepsilon_3 + \frac{\varepsilon_2}{\Gamma(n)} \int_0^{x_0} (x_0 - t)^{n-1} \Phi(t) dt = \Phi(x_0) \end{aligned}$$

lo cual no puede ser cierto en vista de nuestra elección de x_0 . Por lo tanto, la suposición debe ser falsa, y encontramos que efectivamente

$$|\delta(x)| < \Phi(x) = (\varepsilon_1 + \varepsilon_3) E_n(\varepsilon_2 x^n)$$

para todo $x \in [0, T]$. Dado que esto es válido para cada $\varepsilon > 0$, obtenemos el resultado deseado.

En nuestro primer resultado principal, investigamos la dependencia de la solución de una ecuación diferencial fraccionaria sobre los valores iniciales.

Teorema 6.20.

Sea y la solución del problema de valor inicial (6.25), y sea z la solución del problema de valor inicial

$$D_{*0}^n y(x) = f(x, y(x)), \dots (6.26a).$$

$$D^K z(0) = z_0^{(k)}, K = 0, 1, 2, \dots, m-1, \dots (6.26b).$$

Además, sea $\varepsilon := \max_{k=0,1,2,\dots,m-1} |y_0^k - z_0^k|$. Entonces, si ε es lo suficientemente pequeño, entonces existe algún $h > 0$ tal que las funciones y y z están definidas en $[0, h]$, y

$$\sup_{0 \leq k \leq h} |y(x) - z(x)| = O\left(\max_{k=0,1,2,\dots,m-1} |y_0^k - z_0^k|\right)$$

Prueba

Por construcción, existe una solución única de (6.25) en algún intervalo no vacío. Es claro del Teorema 6.5 que el conjunto G correspondiente al problema (6.26) no está vacío y, por lo tanto, este problema también tiene una solución determinada de manera única en algún intervalo. Ahora podemos elegir $[0, h]$ para que sea el menor de estos dos intervalos. Definiendo $\delta(x) := y(x) - z(x)$, encontramos que δ es una solución del problema de valor inicial

$$D_{*0}^n \delta(x) = f(x, \delta(x)) - f(x, z(x))$$

$$D^K \delta(0) = y_0^{(k)} - z_0^{(k)}, K = 0, 1, 2, \dots, m-1$$

En vista del Lema 6.2, este problema de valor inicial es equivalente a la ecuación integral

$$\delta(x) = \sum_{k=0}^{m-1} \frac{x^k}{k!} (y_0^{(k)} - z_0^{(k)}) + \frac{1}{\Gamma(n)} \int_0^{x_0} (x_0 - t)^{n-1} (f(t, y(t)) - f(t, z(t))) dt$$

Tomando valores absolutos y usando la desigualdad de Holder para la suma y la condición de Lipschitz en f para la integral, encontramos,

$$|\delta(x)| \leq \varepsilon \sum_{k=0}^{m-1} \frac{h^k}{k!} + \frac{L}{\Gamma(n)} \int_0^x (x-t)^{n-1} |\delta(t)| dt$$

donde L es la constante de Lipschitz de f . Así, por el Lema 6.19

$$|\delta(x)| \leq O(\varepsilon) E_n(L h^n) = O(\varepsilon)$$

cómo se desee.

A continuación, observamos la influencia de los cambios en la función dada en el lado derecho lado de la ecuación diferencial.

Teorema 6.21

Sea y la solución del problema de valor inicial (6.25), y sea z , la solución del problema de valor inicial

$$D_{+0}^n z(x) = \tilde{f}(x, z(x)), \dots (6.27a)$$

$$D^n z(0) = y_0^{(k)}, k=0,1,2,\dots,m-1,\dots(6.27b)$$

donde se supone que $\tilde{f}$ satisface las mismas hipótesis que f . Además, sea $\varepsilon := \max_{(x_1, x_2) \in G} |f((x_1, x_2)) - \tilde{f}(x_1, x_2)|$. Entonces, si ε es lo suficientemente pequeño, existe algún $h > 0$ tal que las funciones y y z están definidas en $[0, h]$, y tenemos que

$$\sup_{0 \leq x \leq h} |y(x) - z(x)| = O\left(\max_{(x_1, x_2) \in G} |f((x_1, x_2) - \tilde{f}(x_1, x_2))|\right)$$

Prueba

Sigue la existencia y unicidad de las soluciones de ambos problemas de valor inicial como en el teorema anterior. Para probar la desigualdad requerida, también procedemos en una manera similar. Una vez más definimos $\delta(x) := y(x) - z(x)$ y encontramos que δ resuelve el Problema de valor inicial

$$\begin{aligned} D_{+0}^n \delta(x) &= f(x, y(x)) - \tilde{f}(x, z(x)); \\ D^n \delta(0) &= 0, k = 0, 1, 2, \dots, m-1 \end{aligned}$$

que es equivalente a la ecuación integral

$$|\delta(x)| \leq \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} (f(t, y(t)) - \tilde{f}(t, z(t))) dt$$

Tomando valores absolutos y utilizando los supuestos de Lipschitz sobre f y $\tilde{f}$, deducimos

$$\begin{aligned} |\delta(x)| &= \int_0^x \frac{(x-t)^{n-1}}{\Gamma(n)} (|f(t, y(t)) - f(t, z(t))| + |f(t, z(t)) - \tilde{f}(t, z(t))|) dt \\ &\leq \frac{L}{\Gamma(n)} \int_0^x (x-t)^{n-1} |\delta(x)| dt + \varepsilon \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} dt \\ &\leq \frac{L}{\Gamma(n)} \int_0^x (x-t)^{n-1} |\delta(x)| dt + \varepsilon \frac{h^n}{\Gamma(n)} \end{aligned}$$

Ahora podemos aplicar el Lema 7.19 una vez más y encontrar según sea necesario.

$$|\delta(x)| \leq O(\varepsilon) E_n(L h^n) = O(\varepsilon)$$

Finalmente discutimos las consecuencias de una modificación del orden del diferencial ecuación. Aquí debemos tener especial cuidado porque un cambio en el orden de la ecuación di-

ferencial puede conducir a un cambio en el número de prescritos valores iniciales.

Teorema 6.22

Sea y la solución del problema de valor inicial (6.25), y sea z la solución del problema de valor inicial

$$D_{*0}^n z(x) = f(x, z(x)), \dots (6.28a)$$

$$D^n z(0) = y_0^{(k)}, k=0,1,2,\dots,\tilde{m}-1, \dots (6.28b)$$

Donde $\tilde{n} > n$ and $\tilde{m} := [\tilde{n}]$, además se $\varepsilon = \tilde{n} - n$ and

$$\varepsilon^* := \begin{cases} 0, m = \tilde{m} \\ \max \left\{ \left| y_0^{(k)} \right|, k = 0, 1, 2, \dots, \tilde{m} - 1 \right\}, \text{casi cintrario} \end{cases}$$

Entonces, si ε y ε^* son lo suficientemente pequeños, existe algún $h > 0$ tal que tanto las funciones ' y ' y ' z ' están definidas en $[0, h]$, y tenemos que

$$\sup_{0 \leq x \leq h} |y(x) - z(x)| = O(\tilde{n} - n) = O\left(\max \left\{ 0, \max \left\{ \left| y_0^{(k)} \right|, m \leq k \leq \tilde{m} - 1 \right\} \right\}\right)$$

Prueba.

La existencia y la unicidad de las soluciones se pueden deducir como se indicó anteriormente. Para el limitado por la diferencia, reescribimos los problemas de valor inicial (6.25) y (6.28) como ecuaciones integrales equivalentes según el Lema 6.2 y restar las dos resultantes ecuaciones entre sí. De este modo,

$$\begin{aligned}
\delta(x) &:= y(x) - z(x) \\
&= - \sum_{k=m}^{\tilde{m}-1} \frac{x^k}{k!} y_0^{(k)} + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt \\
&\quad - \frac{1}{\Gamma(\tilde{n})} \int_0^x (x-t)^{\tilde{n}-1} f(t, z(t)) dt \\
&= - \sum_{k=m}^{\tilde{m}-1} \frac{x^k}{k!} y_0^{(k)} + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} (f(t, y(t)) - f(t, z(t))) dt \\
&\quad + \int_0^x \left(\frac{(x-t)^{n-t}}{\Gamma(n)} - \frac{(x-t)^{\tilde{n}-1}}{\Gamma(\tilde{n})} \right) f(t, z(t)) dt
\end{aligned}$$

En vista de la propiedad de Lipschitz de f , esto implica

$$\begin{aligned}
|\delta(x)| &\leq \sum_{k=m}^{\tilde{m}-1} \frac{h^k}{k!} |y_0^{(k)}| + \frac{L}{\Gamma(n)} \int_0^x (x-t)^{n-1} |\delta(x)| dt \\
&\quad + \max_{(x_1, x_2) \in G} |f((x_1, x_2))| \int_0^x \left(\frac{(x-t)^{n-t}}{\Gamma(n)} - \frac{(x-t)^{\tilde{n}-1}}{\Gamma(\tilde{n})} \right) dt
\end{aligned}$$

Obviamente la suma es $O(\varepsilon^*)$. Además, podemos acotar la segunda integral, de acuerdo con:

$$\begin{aligned}
\int_0^x \left(\frac{(x-t)^{n-t}}{\Gamma(n)} - \frac{(x-t)^{\tilde{n}-1}}{\Gamma(\tilde{n})} \right) dt &= \int_0^x \left(\frac{u^{n-t}}{\Gamma(n)} - \frac{u^{\tilde{n}-1}}{\Gamma(\tilde{n})} \right) dt \leq \int_0^h \left(\frac{u^{n-t}}{\Gamma(n)} - \frac{u^{\tilde{n}-1}}{\Gamma(\tilde{n})} \right) dt \\
&= O(\varepsilon)
\end{aligned}$$

Para ver esto, uno puede calcular la integral explícitamente: primero tiene que encontrar el cero, del integrando, que se encuentra en $\left(\frac{\Gamma(\tilde{n})}{\Gamma(n)} \right)^{\frac{1}{(\tilde{n}-n)}}$ (una cantidad que converge a $\exp(\psi(n))$ como $\tilde{n} \rightarrow n$, donde $\psi = \frac{\Gamma'}{\Gamma}$ denota la función Digamma). si h es menor que este valor, entonces el integrando no tiene cambio de signo, y podemos movernos la operación de valor absoluto fuera de la integral. De lo contrario, debemos dividir el intervalo de integración en este punto, y cada integral se puede manejar de esta manera. En cualquiera caso, podemos usar el Teorema del Valor Medio del Cálculo Diferencial para ver que la expresión resultante

está limitada por $O(\tilde{n} - n) = O(\varepsilon)$.

$$|\delta(x)| \leq O(\varepsilon) + O(\varepsilon^*) + \frac{L}{\Gamma(n)} \int_0^x (x-t)^{n-1} |\delta(x)| dt$$

y así Lema 6.19 da el resultado deseado

Corolario 6.23

Suponga las hipótesis del Teorema 6.22. Además, sea $\tilde{m} = m$. Entonces

$$\sup_{0 \leq x \leq h} |y(x) - z(x)| = O(\tilde{n} - n)$$

Corolario 6.24

Suponga las hipótesis del Teorema 6.22. Además, sea $\tilde{m} > m$ y $y_0^{(k)} = 0$ para $k = m, m+1, \dots, \tilde{m}-1$. Entonces,

$$\sup_{0 \leq x \leq h} |y(x) - z(x)| = O(\tilde{n} - n)$$

Prueba (de los Corolarios 6.23 y 6.24). Bajo los supuestos de cualquiera de los corolarios, la cantidad ε^* en el Teorema 6.22 se anula y así obtenemos inmediatamente el límite más simple.

Observación 6.11

Sea z una función continua dada. Ya que la función $y := j_0^n z$ es la solución única de la ecuación diferencial $D_{*0}^n y = z$ con condiciones iniciales homogénea, el Corolario 6.24 da una prueba alternativa para la segunda declaración de Teorema 2.10.

Suavidad de las soluciones

Una pregunta interesante que surge con frecuencia es la cuestión de la suavidad de la solución de una ecuación diferencial bajo ciertas suposiciones sobre los datos proporcionados, principalmente en la función dada f en el lado derecho de la ecuación. Un típico resultado en el escenario clásico es el siguiente teorema bien conocido.

Teorema 6.B

Sean $k \in \mathbb{N}$ y $f \in C^{k-1}(G)$, donde $G = [y_0 - k, y_0 + k] \times \mathbb{R}$. Entonces, la solución y del problema de valor inicial

$$Dy(x) = f(x, y(x)), y(0) = y_0$$

es k veces continuamente diferenciable.

Un ejemplo simple muestra que no podemos esperar que este resultado sea cierto en general, para ecuaciones de orden fraccionario: Incluso para $f \in C^\infty$ puede ocurrir que $y \notin C^1$.

Ejemplo 6.4

Del Ejemplo 3.1 sabemos que $D_{*0}^{\frac{1}{2}}(\cdot) = \Gamma(3/2)$. Por lo tanto, la función no diferenciable y dada por $y(x) = x^{1/2}$ es la única solución del problema de valor inicial $D_{*0}^{\frac{1}{2}}(\cdot) = \Gamma(3/2)$, $y(0)=0$, cuya función dada f , al lado derecha de la ecuación diferencial, es analítico.

Ahora tratamos de encontrar algunas declaraciones positivas que se pueden decir acerca de las ecuaciones de orden fraccionario. Al igual que en la sección anterior, en la presente supondremos que y es la solución exacta del problema de valor inicial.

$$D_{*0}^n z(x) = f(x, z(x)), \dots (6.29a)$$

$$D^n z(0) = y_0^{(k)}, k=0,1,2,\dots,m-1, \dots (6.29b)$$

donde una vez más hemos puesto $m = [n]$. Además, supongamos que f es tal que se satisfacen las hipótesis del teorema 6.5, por lo que podemos estar seguros de que existe una solución continua en algún intervalo $[0, h]$, y que esta solución es única. Bajo estas supuestas comenzamos recordando algunos resultados de Diethelm (2008a). Un primer resultado en este la conexión es la siguiente.

Teorema 6.25

Bajo las hipótesis anteriores tenemos que $y \in C^{m-1} [0, h]$.

Prueba

En vista del Lema 6.2, la función, satisface la ecuación de Volterra $y(x) = p(x) + J_0^n [., y(.)](x)$

siendo p algún polinomio cuya forma precisa no es de interés en este momento.

Sea $k \in \{0, 1, 2, \dots, m-1\}$ (esto implica $k < n$) y derivar esta ecuación k veces:

$$\begin{aligned} D^k y(x) &= D^k p(x) + D^k J_0^n [., y(.)](x) \\ &= D^k p(x) + D^k J_0^k J_0^{n-k} [., y(.)](x) \\ &= D^k p(x) + J_0^{n-k} [., y(.)](x) \end{aligned}$$

en vista de la propiedad de semigrupo de integración fraccionaria y (1.1). ahora recuerda que y es continua; por tanto, el argumento del operador integral J_0^{n-k} es una función continua. Por lo

tanto, en vista de la estructura polinomial de p y el Teorema 2.5, la función en el lado derecho de la ecuación es continua, por lo que la función en la izquierda, a saber. $D^k y$, debe ser continuo también.

Teorema 6.26

Suponga las hipótesis del Teorema 6.25. Además, sea $n > 1$, $n \notin \mathbb{N}$ y $f \in C^1(G)$. Entonces $y \in C^m(0, h]$. Además, $y \in C^m[0, h]$, si y solo si $f(0, y(0)) = 0$.

Observación 6.12

Dado que la función f y el valor inicial $y_0^{(0)}$ se dan, es fácil de comprobar si la condición $f(0, y_0^{(0)}) = 0$ se cumple o no.

Observación 6.13

Una comparación con el Teorema 6.B revela una diferencia significativa entre el encuadre clásico y el fraccionario: en el primero, la suavidad de dada la función f implica suavidad de la solución y ; en este último esto sólo se cumple bajo ciertas condiciones adicionales. Esto no es inesperado debido a lo que encontramos en el ejemplo.

Prueba

Introducimos la abreviatura $z(x) := f(t, y(t))$ y escribimos la identidad establecida en la prueba anterior para $k = m-1$:

$$\begin{aligned} D^{m-1}y(x) &= D^{m-1}p(x) + J_0^{n-m+1}z(x) \\ &= D^{m-1}p(x) + \frac{1}{\Gamma(n-m+1)} \int_0^x (x-t)^{n-m} z(t) dt \end{aligned}$$

Derivamos una vez más, recordemos que p es un polinomio de grado $m-1$ y hallamos

$$\begin{aligned}
D^m y(x) &= D^m p(x) + \frac{1}{\Gamma(n-m+1)} \frac{d}{dx} \int_0^x (x-t)^{n-m} z(t) dt \\
&= \frac{1}{\Gamma(n-m+1)} \frac{d}{dx} \int_0^x u^{n-m} z(x-u) du \\
&= \frac{1}{\Gamma(n-m+1)} \left(x^{n-m} z(0) + \int_0^x u^{n-m} z'(x-u) du \right) \\
&= \frac{1}{\Gamma(n-m+1)} x^{n-m} f(0, y_0^{(0)}) + \int_0^{n-m+1} z'(x), \dots \quad (6.30)
\end{aligned}$$

Como $n > 1$ deducimos que $m \geq 2$, y así el teorema 6.25 afirma que $y \in C^1[0, h]$. Un cálculo explícito da que

$$z'(t) = \frac{\partial}{\partial t} f(t, y(t)) + \frac{\partial}{\partial t} f(t, y(t)) y'(t)$$

En consecuencia, por nuestra suposición de diferenciabilidad en f , la función z es continua, y entonces $J_0^{n-m+1} z \in \bar{C}[0, h]$ también del Teorema 2.5. El hecho de que $m > n$ finalmente da que el lado derecho de (6.30), y por lo tanto también el lado izquierdo de esta ecuación, es decir, la función $D^m y$, es siempre continua en el intervalo semiabierto $(0, h]$ mientras que es continua en el intervalo cerrado $[0, h]$ si y solo si $f(0, y_0^{(0)}) = 0$.

Es posible generalizar esta idea y mantener válidas las Observaciones 6.12 y 6.13:

Teorema 6.27

Suponga las hipótesis del Teorema 6.25. Además, sea $k \in \mathbb{N}$, $n > k$, $n \notin \mathbb{N}$ y $f \in C^k(G)$. Sea $z(t) := f(t, y(t))$. Entonces $y \in C^{m+k-1}([0, h])$ si y Además, $y \in C^{m+k-1}[0, h]$ si y solo si z tiene un k -pliegue cero en el origen.

La demostración se basa en una diferenciación repetida de (6.30); dejamos los detalles como un ejercicio para el lector.

Una característica común de los resultados anteriores es que siempre requieren un relativamente alto orden n del operador diferencial para probar que la solución y posee un gran número de derivadas en el intervalo semiabierto $(0, h]$ o incluso en el cerrado intervalo $[0, h]$. Sin embargo, también es posible obtener resultados similares si n es pequeño si imponer condiciones de suavidad más fuertes en f . Esto se deduce de nuestra próxima declaración.

Teorema 6.28

Suponga las hipótesis del Teorema 6.25. Además, sea $f \in C^k(G)$

Entonces $y \in C^k(0, h] \cap C^{m-1}[0, h]$ y para $\ell = m, m+1, \dots, k$ tenemos $y^{(\ell)}(x) = O(x^{n-\ell})$ como $x \rightarrow 0$, solo si z tiene un k -pliegue cero en el origen.

Este teorema es un caso especial de Caponetto et al. (2010), Teorema 3.1. La prueba dada allí es basada en métodos de la teoría de las ecuaciones integrales de Fredholm, que están más allá del alcance de este texto. Por lo tanto, no damos ningún detalle aquí y remitimos al lector al artículo original de Brunner et al. (1986); (Caponetto et al., 2010) en su lugar. Aquí debemos señalar un punto específico. De nuestros teoremas anteriores (y los que seguirá en el resto de esta sección) el lector puede ser llevado a la creencia que una función suave (diferenciable o incluso analítica) f en el lado derecho de la ecuación diferencial conducirá necesariamente a un comportamiento no uniforme de la solución, al menos en la vecindad del punto de partida. De hecho, en el clásico literatura uno a menudo encuentra afirmaciones como “se ve fácilmente que

y no puede ser suave [si f es analítica]” (Oldham, 1974, p. 89). El hecho de que la función analítica $y(x) = 1$ resuelva el problema de valor inicial $D_{*0}^n y(x) = y(x) - 1, y(0) = 1$, que tiene una función analítica $f(x, y) = y - 1$ en su lado derecho, da un contraejemplo fácil a estas declaraciones. Afortunadamente podemos dar una extensión del Teorema 6.27 que caracteriza completamente las situaciones en las que podemos tener suavidad para la función dada f y la solución y simultáneamente. Observamos que el siguiente enunciado es un caso especial de un resultado más general para una clase más grande de ecuaciones (Diethelm, 2008a).

Teorema 6.29

Considere el problema de valor inicial (6.29) y suponga que f es analítica en un conjunto G adecuado. Definir $T(x) := \sum_{j=0}^{m-1} \frac{y_0^{(j)} x_j}{j!}$. Entonces, y es analítico si y sólo si $f(x, T(x)) = 0$ para todo x .

Observación 6.14

Una vez más notamos que f es la función dada de la derecha lado de la ecuación diferencial bajo consideración, y T también se conoce porque es un polinomio cuyos coeficientes se pueden calcular a partir de la inicial dada condiciones. Por lo tanto, en la práctica siempre es posible comprobar si la condición $f(x, T(x)) = 0$ para todo x se cumple o no.

Prueba

La dirección “ $\Leftarrow$ ” es una simple consecuencia del hecho de que una solución (y por lo tanto, por unicidad, la solución) del problema de valor inicial es $y = T$ porque $D_{*0}^0 y = D_{*0}^n T = 0 = f(., T(.))$. En otras palabras, la solución es un polinomio clásico y por lo tanto analítica.

Para la otra dirección, asumimos que y es analítica. Entonces, como f es analítica, el lado derecho de la ecuación diferencial también es analítico. Por lo tanto, dado que se supone y ser la solución de la ecuación, el lado izquierdo $D_{*0}^n y$ debe ser analítico como Bueno. Pero el Teorema 3.15 afirma que y $D_{*0}^n y$ solo puede ser analítico simultáneamente si $D_{*0}^n y = 0$. Esto implica que y debe ser un polinomio de grado $m-1$, y como y satisface nuestras condiciones iniciales, encontramos que $y = T$. Por lo tanto, concluimos

$$0 = D_{*0}^n y(x) = f(x, y(x)) = f(x, T(x))$$

para todo x .

La demostración del teorema 6.29 en realidad nos brinda información adicional.

Corolario 6.30

Considere el problema de valor inicial (6.29), y suponga que f es analítica en un conjunto adecuado G . Si la solución y de este problema de valor inicial es analítica, entonces

$$y(x) = \sum_{j=0}^{m-1} \frac{y_0^{(j)}}{j!} x_j$$

En otras palabras, la única clase de funciones analíticas que puede surgir como solución de una ecuación diferencial fraccionaria tipo Caputo de orden n con una analítica dada La función del lado derecho consta de los polinomios de orden $m-1$. Podemos expresar este hecho de una manera diferente pero también bastante instructiva.

Corolario 6.31

Suponga que la solución del problema de valor inicial (6.29) es analítica, pero no un polinomio de grado $m-1$. Entonces, la función f no es analítica.

Así resulta que, en general (es decir, si no tratamos con el caso especial del Teorema 6.29), debemos esperar que alguna derivada de la solución tenga una discontinuidad al origen.

La naturaleza de esta discontinuidad se puede analizar para obtener alguna información útil sobre el comportamiento preciso de la solución cerca del origen. En este contexto tenemos los siguientes resultados que son muy similares a los de Lubich (Oldham, 1974, p. 2). Debe señalarse, sin embargo, que el artículo original de Lubich contiene una serie de pequeños errores. Daremos aquí las formulaciones correctas, con los elementos básicos de las pruebas influenciadas por las ideas de Hennecke. De hecho, es posible demostrar que estos resultados en realidad mantener para una gran clase de ecuaciones de Volterra que contiene nuestro diferencial fraccionario ecuaciones como un caso especial (Brunner et al., 1999). Sin embargo, no entraremos en detalles sobre este aspecto aquí y remita al lector interesado a Diethelm (2008a), en su lugar.

La naturaleza precisa de los resultados que podemos presentar depende de si o no el orden n de la ecuación diferencial es racional. Comenzamos analizando el caso. de valores racionales de n .

Teorema 6.32

Sea $n = p/q$ donde $p \geq 1$ y $q \geq 2$ son dos números primos relativos. Considere el problema de valor inicial (6.29) y suponga que f puede escribirse en la forma $f(x, y) = \tilde{f}(\frac{x_1}{q}, y)$ donde $\tilde{f}$ es analítica en un entorno de $(0, y_0^{(0)})$. Entonces, existe una función analítica unívocamente determinada $\tilde{f}: (-r, r) \rightarrow R$ con alguna $r > 0$ tal que $y(x) = \tilde{y}(\frac{x_1}{q})$ para $x \in [0, r)$.

Observe que nuestras suposiciones básicas afirman la existencia de una solución en algunos intervalos $[0, h)$, mientras que el Teorema 6.32 nos da una representación válida en $[0, r)$ con algunas r . Este parámetro r será el radio de convergencia de la serie de potencias representación de y . Nuestro método de prueba nos permitirá concluir que $r > 0$ pero no nos dirá nada sobre la relación entre r y h . la misma observación se aplica a los teoremas y corolarios siguientes a continuación.

Antes de llegar a la demostración del teorema 6.32, notamos dos consecuencias inmediatas de este resultado.

Corolario 6.33

Sea $n = p/q$ donde $p \geq 1$ y $q \geq 2$ son dos números primos relativos. Considere el problema de valor inicial (6.29) y suponga que f es analítica en una vecindad de $(0, y_0^{(0)})$. Entonces, existe una función analítica determinada unívocamente $\tilde{y}: (-r, r) \rightarrow R$ con algún $r > 0$ tal que $y(x) = \tilde{y}(x^{\frac{1}{q}})$ para $x \in [0, r)$.

Prueba

Dado que f en sí misma analítica, se puede escribir en forma de serie potencia convergente

$$f(x, y) = \sum_{j,k=0}^{\infty} f_{j,k} x^j y^k$$

$$\tilde{f}(x, y) = \sum_{j,k=0}^{\infty} f_{j,k} x^j y^k \text{ que la función } \tilde{f} \text{ con}$$

es analítico también y satisface la relación $f(x, y) = \tilde{f}(x^{\frac{1}{q}}, y)$. Así, la afirmación sigue directamente del teorema 6.32.

Nuestro segundo corolario es casi inmediatamente evidente:

Corolario 6.34

Bajo los supuestos del Teorema 6.32 o el Corolario 6.33, la solución y del problema de valor inicial (6.29) se puede escribir en una vecindad del origen en la forma de la serie convergente

$$y(x) = \sum_{i=1}^{\infty} \tilde{y}_i x^{\frac{i}{q}}$$

con ciertos coeficientes $\tilde{y}_i$. Además, $\tilde{y}_i = 0$ si $i < p$ y $i/q \notin \mathbb{N}$.

Prueba

El hecho de que y pueda representarse de la forma indicada es una consecuencia directa del teorema 6.32. El hecho de que $\tilde{y}_i = 0$ si $i < p$ y $i/q \notin \mathbb{N}$, se mostrará en la prueba de ese teorema; véase (6.35) a continuación.

Demostración (del Teorema 6.32)

Como se señaló anteriormente, nuestra prueba utiliza métodos similares a los utilizado por Lubich (Oldham, 1974) (Agarwal et al., 2008). Procedemos construyendo una solución formal en la

forma de una llamada serie P si o serie Puiseux (Linz, 1985, cap. 7). La estructura de esta serie, coincidirá con nuestros requisitos, es decir, será una serie de potencias en . los coeficientes de la serie se determinan recursivamente de tal manera que se puede ver que la serie es una solución formal del problema de valor inicial en cuestión. Finalmente mostraremos que esta serie formal es en realidad convergente. Esta parte de la prueba se basará en una modificación adecuada del método de la mayorante de Lindelöf (Liao & Sherif, 2004), para este método aplicado a ecuaciones diferenciales ordinarias de orden entero). Entonces se sigue que la función definida por esta serie convergente es una solución del problema de valor inicial y que puede representarse en la forma que hemos reclamado.

Consideremos entonces una función

$$y(x) = \sum_{i=1}^{\infty} \tilde{y}_i x^{\frac{i}{q}}, \dots (6.31)$$

Donde $y(0) = y_0^{(0)}$, necesitamos demostrar que podemos elegir los coeficientes tales que la serie converge y que la función representada por la serie resuelve nuestro problema de valor inicial.

Sustituyendo (6.31) en (6.2) (que sabemos que es equivalente al problema de valor inicial) y usando el desarrollo en serie de alrededor del punto $(0, y_0^{(0)})$ a saber.

$$f(x, y(x)) = \tilde{f}(x, y(x)) = \sum_{\rho, \sigma} f_{\rho, \sigma} x^{\frac{\rho}{q}} (y - y_0^{(0)})^{\sigma} = \sum_{\rho, \sigma} f_{\rho, \sigma} x^{\frac{\rho}{q}} \left(\sum_{i=1}^{\infty} \tilde{y}_i x^{\frac{i}{q}} \right)^{\sigma}$$

Nosotros encontramos

$$\sum_{i=1}^{\infty} \tilde{y}_i x^{\frac{i}{q}} = \sum_{k=0}^{\infty} \frac{x^k}{k!} + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} \sum_{\rho, \sigma} f_{\rho, \sigma} t^{\frac{\rho}{q}} \left[\sum_{i=1}^{\infty} \tilde{y}_i t^{\frac{i}{q}} \right]^{\sigma} dt, \dots (6.32)$$

Reorganizamos el término entre paréntesis como,

$$\left[\sum_{i=1}^{\infty} \tilde{y}_i t^{\frac{i}{q}} \right]^{\sigma} = \sum_{j=0}^{\infty} \left(\sum_{i_1+i_2+\dots+i_{\sigma}=j} \tilde{y}_{i_1} \dots \tilde{y}_{i_{\sigma}} \right) t^{\frac{j}{q}}, \dots (6.33)$$

Como $i \geq 1$, el caso $j = 0$ en la primera suma del lado derecho solo ocurre para $\sigma = 0$. En este caso fijamos el coeficiente al valor correcto 1, de acuerdo con la convención habitual para productos vacíos. En todos los demás casos, la segunda suma de la mano derecha se extiende sobre todos los productos $\tilde{y}_i \cdots \tilde{y}_\sigma$ con σ factores, donde la suma de los índices i_k es igual a j , porque exactamente estos términos contribuyen al coeficiente de $\frac{j}{t^q}$ en la serie reordenada. Los índices están etiquetados, por lo que no surgen multinomios, pero productos formalmente diferentes pueden ser iguales. Esta reorganización está permitida si podemos demostrar las propiedades de convergencia requeridas. Suponiendo una convergencia uniforme, podemos también intercambiar el orden de sumatoria e integración e integrar término por término. Si lo hacemos, obtenemos

$$\sum_{i=0}^{\infty} \tilde{y}_i t^{\frac{i}{q}} = \sum_{k=0}^{\infty} \frac{x^k}{k!} y_0^{(k)} + \sum_{\rho, \sigma} f_{\rho, \sigma} \left(\frac{\Gamma(\frac{\rho+j}{q}+1)}{\Gamma(\frac{\rho+j+p}{q}+1)} \sum_{i_1+i_2+\dots+i_\sigma=j} \tilde{y}_{i_1} \cdots \tilde{y}_{i_\sigma} \right) x^{\frac{\rho+j+p}{q}}, \quad (6.34)$$

Ahora comparamos los coeficientes de $x^{\frac{i}{q}}$ en ambos lados de esta ecuación. El resultado de esta comparación depende de i . Para $i < p$ obtenemos

$$\tilde{y}_i = \begin{cases} \frac{y_0^{(\frac{i}{q})}}{\frac{i}{q}!} & , \text{ si } i = q, 2q, \dots (m-1)q, \dots (6.35) \\ 0, & \text{ Caso contrario} \end{cases}$$

mientras que para $i \geq p$ encontramos,

$$\tilde{y}_i = \sum_{\rho+j+p=i} \sum_{\sigma=0}^{\infty} f_{\rho, \sigma} \left(\frac{\Gamma(\frac{\rho+j}{q}+1)}{\Gamma(\frac{\rho+j+p}{q}+1)} \sum_{i_1+i_2+\dots+i_\sigma=j} \tilde{y}_{i_1} \cdots \tilde{y}_{i_\sigma} \right), \dots (6.36)$$

Por lo tanto, los coeficientes $\tilde{y}_i$ se determinan únicamente para $0 \leq i < p$ debido a (6.35).

Para $i \geq p$ vemos que (6.36) es una relación de recurrencia: establece que su lado izquierdo, verbigracia. $\tilde{y}_i$, solo depende de los coeficientes donde

$$i_\ell \leq i_1 + i_2 + \dots + i_\sigma = j = i - p - \rho < i$$

porque $p > 0$. En otras palabras, $\tilde{y}_i$ se puede expresar en términos de coeficientes con índices más bajos. Esto significa que tenemos una solución formal única, y necesitamos mostrar que la serie así obtenida converge localmente de manera absoluta y uniforme.

Queda por demostrar que esta serie de potencias generalizada formal es absoluta y uniformemente convergente en una vecindad del origen. Para ello utilizaremos una Generalización del método de las mayores de Lindelof's El mayor de nuestra solución formal se basa en tomar valores absolutos de los coeficientes de f y de los valores iniciales.

$$F\left(x^{\frac{1}{q}}, y(x)\right) := \sum_{\rho, \sigma=0}^{\infty} |f_{\rho\sigma}| x^{\frac{p}{q}} \left(y - |y_0^{(0)}|\right)^{\sigma}$$

Se sabe que esta serie converge porque $\tilde{f}$ es analítica y, por lo tanto, tiene una Expansión en serie convergente. Luego observamos la ecuación de Volterra.

$$Y(x) = \sum_{k=0}^{\infty} \frac{x^k}{k!} \left| y_0^{(k)} \right| + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} F(t, y(t)) dt, \dots (6.37)$$

La solución formal Y de esta ecuación se puede calcular exactamente de la misma manera que arriba. Inmediatamente queda claro que, Y es mayorante para y , y que todos los coeficientes de Y son positivos. Por lo tanto, ahora necesitamos probar que la expansión en serie de $Y(r)$ converge para algún $r > 0$. La positividad

de los coeficientes de Y implica entonces que la serie la expansión de $Y(x)$ converge uniformemente para todo $x \in [0, r]$, y el criterio mayorista nos dice que la expansión de $y(x)$ converge absoluta y uniformemente para estos x también.

La idea en la prueba de convergencia de Lindelof's es que la suma parcial finita

$$P_{\ell+1}(x) = \sum_{i=0}^{\ell+1} Y_i x^{\frac{i}{q}}, \dots (6.38)$$

de la solución formal de (6.37) se puede acotar en términos de P , F y la inicial valores. La observación clave es la desigualdad.

$$P_{\ell+1}(x) = \sum_{k=0}^{m-1} \frac{x^k}{k!} \left| y_0^{(k)} \right| + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} F(t, P_\ell(t)) dt, \dots (6.39)$$

que se desprende del cálculo recursivo de los coeficientes. Si ampliamos En el lado derecho de una serie como la anterior, los términos de orden inferior hasta $x^{\frac{\ell+1}{q}}$ se cancelan exactamente el lado izquierdo por construcción, mientras que en general habrá términos adicionales de orden superior. Todos estos son no negativos porque todos los coeficientes involucrados no son negativos.

Elijamos algún $b > 0$ y definamos $C_1 := \sum_{k=0}^{m-1} \frac{b^k |y_0^{(k)}|}{k!}$. Además, sea

$$C_2 = \max_{(x,w) \in [0, 2C_1]} \frac{|F(x,w)|}{\Gamma(n+1)}. \text{ Finalmente, defina } r := \min_{\square} \left\{ b, \left(\frac{C_1}{C_2} \right) n^{\frac{1}{\square}} \right\}.$$

Nuestro objetivo ahora es probar la desigualdad

$$|P_\ell(x)| \leq 2C_1 \text{ para todo } \ell = 0, 1, 2, \dots \text{ t todo } x \in [0, r], \dots (6.40)$$

La prueba se basa en la inducción matemática sobre. La base de inducción es evidente ya que $P_0(x) = Y_0 = |y_0^{(0)}| \leq C_1$ por definición de . Para el paso de inducción de a recordamos que $|P_{\ell+1}(x)| = P_{\ell+1}(x)$ y escribe, usando (6.39),

$$\begin{aligned}
|P_{\ell+1}(x)| &\leq \sum_{k=0}^{m-1} \frac{x^k}{k!} |y_0^{(k)}| + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} F(t, P_{\ell}(t)) dt \\
&\leq \sum_{k=0}^{m-1} \frac{r^k}{k!} |y_0^{(k)}| + \frac{1}{\Gamma(n)} \max_{t \in [0, x]} |F(t, P_{\ell}(t))| \int_0^x (x-t)^{n-1} dt \\
&\leq C_1 + x^n \frac{1}{\Gamma(n+1)} \max_{t \in [0, x]} |F(t, P_{\ell}(t))| \\
&\leq C_1 + r^n \frac{1}{\Gamma(n+1)} \max_{(x, w) \in [0, r] \times [0, 2C_1]} |F(t, P_{\ell}(t))| \\
&\leq C_1 + r^n C_2 \leq 2C_1
\end{aligned}$$

Así, la secuencia de sumas parciales de la mayorante está uniformemente acotada en $[0, r]$.

En vista de la positividad de todos sus coeficientes es también monótono. Por lo tanto, el la mayorante es absolutamente convergente, y dado que tiene la estructura de una serie de potencias, también es uniformemente convergente en subconjuntos compactos de $[0, r]$. Discutiendo con lo de siempre criterio mayoritario, concluimos las mismas propiedades para el desarrollo de la serie (6.31) lo que finalmente nos dice que nuestro intercambio anterior de suma e integración fue legal. Así la prueba está completa.

Si n es irracional, encontramos un resultado ligeramente diferente.

Teorema 6.35.

Sea n un número irracional positivo. Considere el problema valor inicial (6.29) y suponga que f se puede escribir en la forma $f(x, y) = \bar{f}(x, x^n, y)$ donde $\bar{f}$ es analítica en un entorno de $(0, 0, y_0^{(0)})$. Entonces, existe una única función analítica

determinada $\bar{y} : (-r, r) \times (-r^n, r^n) \rightarrow R$ con algún $r > 0$ tal que $y(x) = \bar{y}(x, x^n)$ para $x \in [0, r]$.

Una vez más observamos dos corolarios antes de pasar a la demostración del teorema 6.35.

Corolario 6.36

Bajo los supuestos del teorema 6.35, y tiene la forma

$$y(x) = \sum_{\mu, v=0}^{\infty} \bar{y}_{\mu, v} x^{\mu+nv}$$

Demostración (del teorema 6.35)

Como en la demostración del Teorema 6.32, los supuestos implican la existencia local de una única solución continua y de (6.2) donde $y(0) = y_0^{(0)}$ en este tiempo sustituimos $\sum_{\mu, v=0}^{\infty} \bar{y}_{i_1, i_2} x^{i_1+i_2n}$ en (6.2), y por cálculos como en la prueba de Del Teorema 6.32 obtenemos relaciones similares a (6.35) y (6.36) para los coeficientes, a saber

$$\bar{y}_{i_1, 0} = \frac{y_0^{(i_1)}}{(i_1)!}, \text{ para } 0 \leq i_1 \leq m-1, \dots (6.41)$$

Y

$$\bar{y}_{i_1, i_2} = \sum_{\substack{m_1+j_1=i_1 \\ m_2+j_2+1=i_2}} \gamma_{m_1 m_2 j_1 j_2} f_{m_1 m_2 \sigma} \sum_{\substack{\mu_1+\mu_2+\dots+\mu_\sigma=j_1 \\ v_1+v_2+\dots+v_\sigma=j_2}} \bar{y}_{\mu_1, v_1} \dots \bar{y}_{\mu_\sigma, v_\sigma}, \dots (6.42)$$

para los casos restantes donde los coeficientes $\gamma_{m_1 m_2 j_1 j_2}$ vienen dados por

$$\gamma_{m_1 m_2 j_1 j_2} = \frac{\Gamma(m_1+j_1+(m_2+j_2)n+1)}{\Gamma(m_1+j_1+(m_2+j_2+1)n+1)}, \dots (6.43)$$

El coeficiente $\bar{y}_{i_1, i_2}$ se determina únicamente en términos de los coeficientes correspondientes a exponentes más pequeños. Por lo tanto, podemos ordenar los exponentes con respecto a magnitud creciente y calcular cada coeficiente únicamente a partir de un (a veces vacío) subconjunto de los coeficientes de los exponentes más pequeños que se han calculado de antemano.

El resto de la prueba es análoga al caso racional; por lo tanto, omitimos los detalles.

El teorema 6.35 nos da una representación de la solución en términos de una función analítica de dos variables, mientras que la función analítica del Teorema 6.32 (que trata con n racional) era una función de una sola variable. Ejemplos sencillos muestran que no puede esperarse que esto último se cumpla en general en el caso irracional. Podemos, sin embargo, construir una representación como la del Teorema 6.35 en el caso racional. El problema aquí es que esta representación ya no es única porque muchos de los exponentes $\mu + nv$ que aparecen en la forma de la expansión que se muestra en el Corolario 6.36 puede representarse como una combinación lineal de l y n en más de una forma (por ejemplo, si $n = p/q$ en términos mínimos, entonces $\mu + nv = (\mu - p) + n(v + q)$). Para ilustrar este fenómeno, nos fijamos en un ejemplo sencillo.

Ejemplo 6.5

Considere el problema del valor inicial

$$D_+^{\frac{1}{2}} y(x) = y(x); y(0) = 0, \dots (6.44)$$

La única solución continua es $y(x) \equiv 0$. Así, una posible representación de y de la forma indicada en el Corolario 6.36 se obtiene estableciendo $\bar{y}_{\mu, v} = 0$ para todo μ y v . Alternativamente podemos elegir $\bar{y}_{0,0} = 0$, $\bar{y}_{\mu,0} = s_\mu$ y $\bar{y}_{0,2\mu} = -s_\mu$ para $\mu \geq 1$ donde $(s_\mu)_{\mu=1}^\infty$ es

una secuencia arbitraria que decae lo suficientemente rápido. Entonces nosotros tenemos

$$\bar{y} \left(x, x^{\frac{1}{2}} \right) = \sum_{\mu=0}^{\infty} s_{\mu} x^{\mu} - \sum_{\mu=0}^{\infty} s_{\mu} (x^{\frac{1}{2}})^{2\mu} = 0$$

Por tanto, podemos representar la solución de infinitas maneras. Este tipo de la no unicidad obviamente se aplica a todas las ecuaciones de la forma (7.29) que obedecen a las condiciones impuestas en el Teorema 6.32.

Una inspección minuciosa de las pruebas revela la posibilidad de obtener resultados comparables. resulta bajo la suposición más débil de que f solo es una función C^k para algún $k \in \mathbb{N}$. Concretamente, en el caso racional tenemos:

Teorema 6.37.

Sea $n = p/q$ donde $p \geq 1$ y $q \geq 2$ son dos números primos relativos. Considere el problema de valor inicial (6.29) y suponga que f puede escribirse en la forma $f(x, y) = \tilde{f}(x^{\frac{1}{q}}, y)$, donde $\tilde{f} \in C^j([0, h^*] \times [y_0^{(0)} - K, y_0^{(0)} + K])$ - $f \in C^j([0, h^*] \times [y(0) - K, y(0) + K])$ con algunos $K > 0$, $h^* > 0$ y $j \in \mathbb{N}$. Entonces, la solución y de (6.29) tiene una expansión asintótica en potencias de $x^{\frac{1}{q}}$ cuando $x \rightarrow 0$. En particular, el exponente no entero más pequeño en esta expansión es n . Si n es irracional, podemos probar la siguiente afirmación.

Teorema 6.38

Sea n un número irracional positivo. Considere el valor inicial problema (6.29) y suponga que f se puede escribir en la forma $f(x, y) = \tilde{f}(x, x^n, y)$ donde $\tilde{f} \in C^j([0, h^*] \times [0, (h^*)^n] \times [y_0^{(0)} - K, y_0^{(0)} + K])$ con algunos $K > 0$, $h^* > 0$ y $j \in \mathbb{N}$. Entonces, la solución y de (6.29) tiene una expansión asintótica en potencias mixtas de x y x^n como $x \rightarrow 0$.

Demostración (de los Teoremas 6.38 y 6.39)

La demostración se basa en una combinación de las demostraciones de los teoremas 6.32 y 6.35, especialmente, y las ideas de Lubich (1983) y Oldham (1974), Corolario 3. En particular, en lugar de construir una expansión en serie infinita para $\tilde{y}$ como en (6.31) o su contraparte en el caso irracional, solo construimos una serie de potencias truncadas $\tilde{y}_N(x)$, digamos, más el término restante $R(x)$, donde el grado N de la serie truncada se calculará más adelante. De manera similar desarrollamos la función $\tilde{f}$ según Teorema de Taylor en forma de polinomio en $\frac{x_1}{q}$ e y alrededor del punto $(0, y_0^{(0)})$ (en el caso racional) o en x, x^n e y alrededor de $(0, 0, y_0^{(0)})$ (en el caso irracional), más el término restante correspondiente. Luego podemos proceder como se indicó anteriormente y llegar a una relación similar a (6.34), sólo que las series infinitas ahora están truncadas y que los términos restantes entran en la ecuación. Así podemos encontrar la definición directa para los primeros coeficientes de la expansión truncada como en (6.35) y la recurrencia relación para los siguientes coeficientes como en (6.36) donde este último ahora no se cumple para todo $i > p$ ya pero sólo hasta el punto en que el truncamiento de la expansión para permisos $\tilde{f}$. Esto determina el valor N donde termina nuestra serie truncada

$\tilde{y}_N$. El Entonces es evidente que (x) tiene una expansión asintótica de la forma requerida, y una estimación sencilla del término restante $R(x)$ basada en la generalización de (6.34) muestra que el resto sólo contiene términos de orden superior que no destruyen las asintóticas de $\tilde{y} = \tilde{y}_N + R$.

Los teoremas anteriores nos han proporcionado una gran cantidad de información sobre las propiedades de suavidad de las soluciones de ecuaciones diferenciales fraccionarias, y en particular sobre el comportamiento exacto de la solución como $x \rightarrow 0$, más notablemente la formal expansión asintótica. Sin embargo, en una aplicación concreta puede ser posible decir aún más. Un aspecto de especial trascendencia, por ejemplo, en vista del desarrollo de métodos numéricos, es la cuestión de los valores precisos de las constantes en esta expansión, y lo más importante, la cuestión de si ciertos coeficientes desaparecer. Una generalización adecuada de la técnica de expansión de Taylor para ordinario Las ecuaciones diferenciales descritas en Kiryakova (2010a; 2010b) (véase en capítulo I.8), podrían ser útiles en este contexto.

Sin embargo, los resultados precisos a este respecto parecen ser es conocidos en este momento.

Problemas de valores en la frontera

En esta sección final de este capítulo alejamos de los problemas de valor inicial PVI discutido hasta ahora y dirigir nuestra atención hacia los problemas de valores en la frontera PVE. Muchas y diferentes tipos de condiciones de contorno son concebibles; nos restringimos a esos tipos que parecen ser los más significativos. Para leer más, en particular con respecto a otras clases

de problemas, nos referimos a la encuesta de Agarwal et al. (2008). Como siempre en este capítulo, la ecuación diferencial bajo consideración es $D_{+0}^n y(x) = f(x, y(x))$ $x \in [0, T]$, ... (6.45) con algún $n > 0$.

Sabemos por la teoría de los problemas de valores iniciales PVF que el número de condiciones que debemos imponer para obtener resultados razonables de existencia y unicidad es n . Por lo tanto, primero miramos el caso $0 < n < 1$ donde es apropiado imponer exactamente una condición de contorno. La forma más natural para tal condición es

$$a y(0) + b y(T) = c, \dots (6.46)$$

con ciertas constantes. En este caso podemos reformular los valores límites del problema como una ecuación integral de Fredholm de la siguiente manera.

Lema 6.39

Sean $0 < n < 1$; $a, b, c \in \mathbb{R}$ y $a + b \neq 0$. Supongamos además que $f: [0, T] \times \mathbb{R} \rightarrow \mathbb{R}$ es continua. Entonces, la función $y \in C[0, T]$ es una solución del problema de valores en la frontera (6.45), (6.46) si y sólo si es una solución de la ecuación integral Fredholm EIF

$$y(x) = \frac{c}{a+b} + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt - \frac{1}{\Gamma(n)} \frac{b}{a+b} \int_0^T (T-t)^{n-1} f(t, y(t)) dt, \dots (6.47) \text{ EIF}$$

Es necesario mencionar dos casos especiales de este resultado: si $b=0$ entonces (6.46) reduce a una condición inicial, y así recuperamos el Lema 7.2. Y si $a = 0$ entonces (5.46) se convierte en una condición terminal y estamos en la situación del teorema 6.18.

Prueba.

La ecuación (6.47) implica

$$y(0) = \frac{c}{a+b} - \frac{1}{\Gamma(n)} \frac{b}{a+b} \int_0^T (T-t)^{n-1} f(t, y(t)) dt,$$

Y

$$y(T) = \frac{c}{a+b} + \frac{1}{\Gamma(n)} \frac{b}{a+b} \int_0^T (T-t)^{n-1} f(t, y(t)) dt,$$

de donde se sigue (6.46). Además, una aplicación del operador diferencial D_{*0}^n a ambos lados de (6.47) se obtiene (6.45). Por lo tanto, resuelve el problema del valor límite si resuelve la ecuación integral.

Por otro lado, si y resuelve la ecuación diferencial (6.45), entonces sabemos por Lema 6.2 que también satisface la ecuación de Volterra

$$y(x) = y(0) + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt, \dots, (6.48)$$

con alguna cantidad (actualmente desconocida) . Tomando en $x = T$ esto lee

$$y(T) = y(0) + \frac{1}{\Gamma(n)} \int_0^T (T-t)^{n-1} f(t, y(t)) dt,$$

Podemos conectar esto a la condición de frontera (6.46) y derivar

$$(a+b)y(0) + \frac{1}{\Gamma(n)} b \int_0^T (T-t)^{n-1} f(t, y(t)) dt = c, (6.49)$$

Ahora resolvemos (6.48) para e insertamos el resultado en (6.49) lo que produce

$$\begin{aligned} (a+b)y(x) + \frac{a+b}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt + \frac{1}{\Gamma(n)} \int_0^T (T-t)^{n-1} f(t, y(t)) dt \\ = c, \end{aligned}$$

que es equivalente a la ecuación integral de Fredholm (6.47) EIF.

Observación 6.15

En el caso $a + b = 0$ sólo podemos obtener el resultado más débil que cualquier. La solución del problema de valores en la frontera PVF (6.45), (6.46) satisface la ecuación integral.

$$c = \frac{b}{\Gamma(n)} \int_0^T (T-t)^{n-1} f(t, y(t)) dt$$

Esto sigue procediendo como en la segunda parte de la prueba del Lema 6.39 hasta a (6.49), que es justo lo que hemos afirmado.

Sin embargo, no podemos deducir que alguna, Solución de esta ecuación integral resuelve el problema del valor inicial PVI.

La diferencia sustancial entre (6.47) y la ecuación integral de Observación 6.15 es que la primera es una ecuación integral del segundo tipo mientras que la segunda es una ecuación del primer tipo.

Es bien sabido que las ecuaciones integrales de Fredholm EIF del segundo tipo se porta mucho mejor que los del primer tipo con respecto a en cuestiones de existencia y unicidad de las soluciones y la continuidad de las soluciones con respecto a los datos proporcionados (Matignon, 1998).

El lema 6.39 nos permite deducir inmediatamente una existencia simple y unicidad. teorema (Benchohra et al., 2008).

Teorema 6.40

Suponga las hipótesis del Lema 6.39. Si además f satisface una condición de Lipschitz con constante de Lipschitz L con respecto a su segunda variable, y si

$$LT^n \left(1 + \frac{|b|}{|a+b|}\right) \leq \Gamma(n+1), \dots \quad (6.50)$$

entonces el problema de valores en la frontera (6.45), (6.46) tiene una solución única $y \in C[0, T]$.

Prueba

Del Lema 6.39 podemos ver que necesitamos demostrar que el operador A , definido por

$$\begin{aligned} Ay(x) &= \frac{c}{a+b} \\ &+ \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt \\ &+ \frac{1}{\Gamma(n)} \frac{b}{a+b} \int_0^T (T-t)^{n-1} f(t, y(t)) dt \end{aligned}$$

tiene un único punto fijo. Está claro que el operador asigna $C[0, T]$ a sí mismo y que

$$\begin{aligned} &|Ay(x) - A\tilde{y}(x)| \\ &= \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} |f(t, y(t)) - f(t, \tilde{y}(t))| dt \\ &+ \frac{1}{\Gamma(n)} \frac{|b|}{|a+b|} \int_0^T (T-t)^{n-1} |f(t, y(t)) - f(t, \tilde{y}(t))| dt \\ &\leq \frac{L}{\Gamma(n)} \|y - \tilde{y}\|_\infty \left(\int_0^x (x-t)^{n-1} dt + \frac{|b|}{|a+b|} \int_0^T (T-t)^{n-1} dt \right) \\ &\leq \frac{L}{\Gamma(n)} T^n \left(1 + \frac{|b|}{|a+b|}\right) \|y - \tilde{y}\|_\infty \end{aligned}$$

lo que implica, bajo nuestro supuesto, que A es una contracción. Así, según teorema del punto fijo de Banach, obtenemos que A efectivamente tiene un punto fijo único. En condiciones ligeramente diferentes también podemos derivar un teorema de existencia.

Teorema 6.41

Suponga las hipótesis del Lema 6.39. Si además f es uniformemente acotada por una constante absoluta, entonces el problema del valor límite (6.45), (6.46) tiene solución $y \in C[0, T]$.

Prueba

Usamos el mismo operador A que en la prueba anterior. Nuestro objetivo ahora es mostrar que tenga al menos un punto fijo. Para este fin queremos invocar el teorema del punto fijo de Schauder.

Por lo tanto, todo lo que necesitamos demostrar es que $X := \{Ay : y \in C[0, T]\}$ es relativamente conjunto compacto. Una condición necesaria y suficiente para que esto se cumpla está contenida en el teorema de Arzelà-Ascoli: Necesitamos demostrar que X está uniformemente acotado y equicontinuo. Pero la acotación uniforme de X es una consecuencia trivial de la definición de A y la acotación de f , y la equicontinuidad se puede mostrar simplemente como en la demostración del teorema 6.1.

En la teoría de ecuaciones diferenciales de orden entero, la clase más importante de los problemas de valores en la frontera es el que tiene ecuaciones diferenciales de segundo orden, es decir ecuaciones donde se imponen dos condiciones de contorno y no solo una como hemos hecho hasta ahora. Transferido a ecuaciones diferenciales fraccionarias tipo Caputo, esto corresponde a ecuaciones de orden $n \in (1, 2)$.

Por lo tanto, ahora echaremos un vistazo a este tipo. de ecuaciones. Una vez más comenzamos convirtiendo el problema de valores en la frontera dado en una ecuación equivalente de Fredholm. Impondremos dos condiciones de contorno similares a los de (7.46); sin embargo, de acuerdo con la teoría clásica, ahora permitimos primero derivadas de la solución en los puntos finales 0 y T aparezcan además de los valores de función en sí mismos.

Por lo tanto, nuestro problema de valores en la frontera ahora tiene la forma

$$D_{\bullet 0}^n y(x) = f(x, y(x)), \dots (7.51a)$$

$$a_{j0}y(0) + a_{j1}y'(0) + b_{j0}T(0) + b_{j1}T'(0) = c_j, j = 1, 2, \dots (7.51b)$$

con algún $n \in (1, 2)$. El resultado de equivalencia queda entonces como sigue

Lema 6.42

Sea $1 < n < 2$, $a_{jk}, b_{jk}, c_{jk} \in \mathbb{R}$ ($j = 1, 2, k = 0, 1$) y supongamos que $f: [0, T] \times \mathbb{R} \rightarrow \mathbb{R}$ es continua.

Si la matriz

$$M := \begin{pmatrix} a_{10} + b_{10} & a_{10} + Tb_{10} + b_{11} \\ a_{20} + b_{20} & a_{20} + Tb_{20} + b_{21} \end{pmatrix}$$

es regular, entonces tenemos que $y \in C^1[0, T]$ es una solución del problema de valores en la frontera (6.51) si y sólo si es solución de la ecuación integral de Fredholm.

$$y(x) = \alpha_1 x + \alpha_2 x + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt, \dots (6.52)$$

Donde α_1, α_2 están determinados por el sistema lineal

$$M \begin{pmatrix} \alpha_1 \\ \alpha_2 \end{pmatrix} = \begin{pmatrix} c_1 - \int_0^x \left[b_{10} \frac{(T-1)^{n-1}}{\Gamma(n)} + b_{11} \frac{(T-1)^{n-2}}{\Gamma(n+1)} \right] f(t, y(t)) dt \\ c_2 - \int_0^x \left[b_{20} \frac{(T-1)^{n-1}}{\Gamma(n)} + b_{21} \frac{(T-1)^{n-2}}{\Gamma(n+1)} \right] f(t, y(t)) dt \end{pmatrix}$$

A primera vista (6.52) parece una ecuación de Volterra. Sin embargo, este no es el caso como es evidente a partir de la definición de las cantidades α_1 y α_2 : Los operadores de Fredholm en el lado derecho de (6.53) ingrese (6.52) de esta manera.

Prueba

La demostración es muy similar a la demostración del Lema 6.39. aplicamos el Lema 6.2 encontrar la equivalencia de la ecua-

ción diferencial con la ecuación integral (6.52) con ciertas constantes α_1 y α_2 . La diferenciabilidad de la solución se deriva de Teorema 6.25. Las dos condiciones de frontera proporcionan dos ecuaciones lineales que permiten determinar las constantes α_1 y α_2 ; el sistema de ecuaciones resultante es simplemente (6.53).

Dejamos los detalles al lector. Con base en este resultado, ahora podemos formular un análogo de los teoremas 6.41 y 6.42. Empezamos por este último

Teorema 6.43

Suponga las hipótesis del Lema 6.42. Además, sea f uniformemente acotado por una constante absoluta. Entonces, el problema de valores en la frontera (6.51) tiene una solución $y \in C^1[0, T]$.

Prueba

Procedemos como en la demostración del Teorema 6.41: Usamos el Lema 6.42 para reescriba el problema de valores en la frontera en la forma de la ecuación integral (6.52), observe que la existencia de una solución a esta ecuación puede interpretarse como la existencia de un punto fijo de un operador A apropiadamente elegido con $Ay(x)$ definido como el lado derecho de (6.52), y use la acotación de f para deducir la existencia de tal punto fijo a través del teorema de Schauder. Finalmente, el teorema de unicidad dice lo siguiente.

Teorema 6.44

Asumir las hipótesis del Lema 6.42. Además, sea f que satisfaga a Condición de Lipschitz con respecto a la segunda variable con un Lipschitz suficientemente pequeño constante. Entonces, el problema de valores en la frontera (6.51) tiene una solución única $y \in C^1[0, T]$.

Observación 6.16.

En el caso de problemas de valores en la frontera de orden $n \in (0, 1)$, también tienen existencia y unicidad solo si la constante de Lipschitz de f es suficientemente pequeño. Es evidente a partir de (7.50) que el valor máximo permitido L del Lipschitz constante en esta conexión depende de los parámetros de la condición de contorno, es decir, en los valores a , b y T . De manera similar, en el caso $1 < n < 2$ bajo consideración, ahora bien, el valor permitido de la constante de Lipschitz dependerá de a_{jk} , b_{jk} y T .

Prueba.

La demostración sigue las mismas líneas que la demostración del teorema 6.40. omitimos los detalles.

Concluimos el tratamiento de los problemas de valores en la frontera para el diferencial de las ecuaciones Caputo en este punto y sólo tenga en cuenta, para mayor referencia, que otros tipos de las condiciones de contorno también se utilizan ocasionalmente; ver, por ejemplo, la discusión.

Además, en principio, los problemas de valores en la frontera de orden superior se pueden tratar en una manera análoga. Sin embargo, en vista de su importancia menor en la práctica, No discutiré este tema aquí explícitamente.

07

Capítulo

***ECUACIONES DIFERENCIALES
FRACCIONARIOS DE CAPUTO DE
UN SOLO TÉRMINO: RESULTADOS
AVANZADOS PARA CASOS
ESPECIALES***

Escucha a tu niño interno que te enseñará, a comprender perfectamente cuál suele ser la envidia, la astucia, y la hipocresía propia y observar en general, quienes entre nosotros llevan el nombre de amigos que son en cierto modo insensibles a aquel amor que la naturaleza inspira para con las personas allegadas.

Burseca

Con los fundamentos de una teoría para ecuaciones diferenciales fraccionarias con Caputo EDFC. Ahora que las derivadas están en su lugar, a continuación, intentaremos dar alguna información adicional sobre ciertos casos especiales de ecuaciones particularmente importantes.

Específicamente esto incluye un análisis más preciso de las propiedades de las soluciones de ecuaciones lineales en el capítulo anterior, donde vimos problemas con valores iniciales PVI y, que se enfoca en problemas con valores en la frontera PVF, la investigación del comportamiento a largo plazo de soluciones definidas en intervalos suficientemente grandes y, en particular, ilimitados.

Problemas de valores iniciales para ecuaciones lineales

En las dos primeras secciones de este capítulo restringimos nuestra atención a una clase especial de ecuaciones que sin embargo es muy importante en muchas aplicaciones: ecuaciones diferenciales fraccionarias lineales EDFL.

Es una observación común en muchas áreas de las matemáticas, que la supuesta linealidad permite derivar enunciados más precisos. Lo mismo es cierto para las ecuaciones diferenciales fraccionarias EDFs.

En particular, las expresiones explícitas porque las soluciones de tales ecuaciones a menudo se pueden obtener. ya hemos encontrado un caso especial de esta situación en el Teorema 6.11.

En general, resulta que “la transformado de Laplace” es una herramienta extremadamente útil para el análisis de fracciones lineales ecuaciones diferenciales y lo usaremos aquí.

Las definiciones fundamentales y Las propiedades de esta transformada se repiten en el Apéndice D.3. Comenzamos las investigaciones en esta sección proporcionando generalizaciones del teorema de integración y el teorema de diferenciación para transformadas de Laplace, partes (c) y (d) del teorema D.11, a las integrales fraccionarias de Riemann-Liouville y derivadas fraccionarias de Caputo.

Es bastante obvio que para $n \in \mathbb{N}$ recuperamos los enunciados clásicos. Aquí y en todo el resto del texto denotamos el operador transformado de Laplace por $\mathcal{L}$.

Teorema 7.1

Supongamos que $f : [0, \infty) \rightarrow \mathbb{R}$ es tal que $\mathcal{L}f$ existe en $[s_0, \infty)$ con Algunos $s_0 \in \mathbb{R}$. Sea $n > 0$ y $m := [n]$. Entonces, para $s > \max\{0, s_0\}$ tenemos

$$\mathcal{L} J_0^n f(s) = \frac{1}{s^n} \mathcal{L} f(s)$$

y

$$\mathcal{L} D_{*0}^n f(s) = s^n \mathcal{L} f(s) - \sum_{k=1}^m s^{n-k} f^{(k-1)}(0)$$

Prueba

Sea $g(x) = \frac{x^{n-1}}{\Gamma(n)}$. Luego, por el ejemplo D.1 (b) y la linealidad de la transformada de Laplace, para $s > 0$. Además, por definición de la Operador integral Riemann-Liouville, $J_0^n f$ es la convolución de f y g , y por lo tanto el teorema de convolución (Teorema D.11 (b)) implica que

$$\mathcal{L} J_0^n f(s) = \mathcal{L} g(s) \cdot \mathcal{L} f(s) = \frac{1}{s^n} \mathcal{L} f(s)$$

Para $s > 0; \max\{0, s_0\}$

Además, recordamos que $D_{*0}^n f = J_0^{m-n} D^m f$, por lo tanto, en vista de lo que Acabo de mostrar y el teorema de diferenciación.

$$\begin{aligned} \mathcal{L} D_{*0}^n f(s) &= \mathcal{L} J_0^{m-n} D^m f = \frac{1}{s^{m-n}} \mathcal{L} D^m f(s) \\ &= s^{n-m} \left(s^m \mathcal{L} f(s) - \sum_{k=1}^m s^{m-k} f^{(k-1)}(0) \right) \\ &= s^n \mathcal{L} f(s) \sum_{k=1}^m s^{n-k} f^{(k-1)}(0) \end{aligned}$$

Usando estos resultados estamos en condiciones de resolver otra ecuación del ejemplo que es ligeramente más avanzada que las ecuaciones consideradas en el teorema 6.11.

Ejemplo 7.1

Buscamos la solución del problema del valor inicial.

$$\begin{cases} D_{*0}^n y(x) = -y(x) - q(x) \\ y(0) = 2 \end{cases}$$

con algún $n \in (0,1)$ y alguna función q .

El método se basa en el método de la transformada de Laplace. Aplicando la Laplace transformamos al problema de valor inicial, derivamos $s^n \mathcal{L} y(s) - 2s^{n-1} = -\mathcal{L} y(s) - \mathcal{L} q(s)$

en vista del Teorema 7.1 y el Ejemplo D.1. Resolvemos esta ecuación para $\mathcal{L} y(s)$ y obtenemos

$$\mathcal{L} y(s) = \frac{2s^{n-1}}{s^n + 1} - \frac{\mathcal{L} q(s)}{s^n + 1}$$

Ahora necesitamos encontrar las transformadas inversas de Laplace de las funciones en el lado derecho. Del teorema 4.5 sabemos que

$$z(x) = E_n(-x^n) \Rightarrow \mathcal{L} z(s) = \frac{s^{n-1}}{s^n + 1}$$

Además, tenemos para esta función z que

$$\mathcal{L} z'(s) = s\mathcal{L} z(s) - z(0) = \frac{s^{n-1}}{s^n + 1} - 1 = -\frac{1}{s^n + 1}$$

por el teorema de diferenciación de la transformada de Laplace. Combinando esto con el teorema de convolución para la transformada de Laplace, encontramos

$$y(x) = 2E_n(-x^n) + \int_0^x q(x-t) \frac{d}{dt} E_n(-x^n) dt$$

De hecho, podemos generalizar las observaciones de este ejemplo y desarrollar una explicación explícita, fórmula para la solución de una clase simple de ecuaciones, las ecuaciones lineales con coeficientes constantes. Descubriremos una vez más el significado de la Funciones de Mittag-Leffler E_n y E_{n_1, n_2}

Teorema 7.2

Sea $n > 0, m = [n]$ y $\lambda \in \mathbb{R}$. La solución del problema valor inicial $D_0^n y(x) = \lambda y(x) + q(x); y^{(k)}(0) = y_0^{(k)}, (k = 0, 1, 2, \dots, m-1), \dots (7.1)$ donde $q \in C[0, h]$ es una función dada, se puede expresar en la forma $y(x) = \sum_{k=0}^m y_0^{(k)} u_k(x) + \tilde{y}(x), \dots (7.2a)$

Con

$$\tilde{y}(x) = \begin{cases} J_0^n q(x) & \text{si } \lambda = 0 \\ \frac{1}{\lambda} \int_0^x q(x-t) u'_0(t) dt, & \text{si } \lambda \neq 0 \end{cases}, \dots, (7.2b)$$

Donde $u_k(x) := J_0^n e_n(x)$, $k = 0, 1, 2, \dots, m-1$ y $e_n(x) := E_n(\lambda x^n)$

Observación 7.1

En el caso $0 < n < 1$ podemos usar las definiciones mencionadas anteriormente de u_0 y e_n para reescribir la representación (7.2) de la solución del problema valor inicial en la forma

$$y(x) = y_0^{(0)} E_n(\lambda x^n) + n \int_0^x q(x-t) t^{n-1} E'_n(\lambda t^n) dt, \dots (7.3)$$

que es válido tanto para $\lambda = 0$ como para $\lambda \neq 0$. Usando la expansión en serie de potencias de la Funciones de Mittag-Leffler, se puede ver fácilmente que esto es equivalente a

$$y(x) = y_0^{(0)} E_n(\lambda x^n) + \int_0^x q(x-t) t^{n-1} E_{n,n}(\lambda t^n) dt$$

Alternativamente podemos reescribir (7.3) como

$$y(x) = y_0^{(0)} E_n(\lambda x^n) + n \int_0^x (x-t)^{n-1} E'_n(\lambda(x-t)^n) q(t) dt, \dots (7.4)$$

En el caso límite $n \rightarrow 1$ esto se reduce a la conocida fórmula

$$y(x) = y_0^{(0)} e^{\lambda x} + \int_0^x e^{\lambda(x-t)} q(t) dt$$

que generalmente se obtiene mediante el método de variación de constantes. Para referencia futura nosotros tenemos en cuenta que (7.4) sigue siendo válida en un entorno con valores vectoriales, es decir, si y , y y q son funciones mapeados en $\mathbb{R}^N$ para algunos N y $y_0^{(0)} \in \mathbb{R}^N$ y λ es una matriz $N \times N$. Más información So-

bre estos problemas multidimensionales se dará al final de este capítulo.

Demostración (del teorema 7.2)

En el caso $\lambda = 0$ tenemos que $e_n(x) = E_n(0) = 1$ y por tanto $E_n(x) = x^k / k!$. Para cada k . Por tanto, se puede verificarse mediante una solicitud directa, de los operadores diferenciales relevantes a la representación dada de y . Para $\lambda \neq 0$ demostraremos los siguientes hechos:

Las funciones u_n satisfacen la ecuación diferencial homogénea $D_{00}^n u_n = \lambda u_k$ ($k = 0, 1, \dots, m-1$), y cumplen las condiciones iniciales $u_k^{(j)}(0) = \delta_{jk}$ (Delta de Kronecker) para $j, k = 0, 1, \dots, m-1$

La función $\tilde{y}$ es una solución de la ecuación diferencial no homogénea con condiciones iniciales homogéneas.

Entonces, el principio de superposición da la afirmación.

Con respecto a (a), sabemos que

$$e_n(x) = \sum_{j=0}^{\infty} \frac{\lambda^j x^{nj+k}}{\Gamma(1+nj)}$$

$$u_k(x) = j_n^k e_n(x) = \sum_{j=0}^{\infty} \frac{\lambda^j x^{nj}}{\Gamma(1+nj+k)}$$

Utilizando esta representación, podemos ver que resuelve la ecuación diferencial homogénea, de la demostración del teorema 4.3. Además, vemos que $u_k^{(k)}(0) = D^k j_n^k e_n(0) = e_n(0) = 1$.

Para $j < k$ tenemos

$$u_k^{(k)}(0) = D^j j_0^k e_n(0) = j_0^{k-1} e_n(0) = 0,$$

debido a la continuidad de e_n y para $j > k$ encontramos

$$u_k^{(j)}(0) = D^j j_0^k e_n(0) = D^{j-k} e_n(0) = 0$$

Porque

$$e_n(x) = 1 + \frac{\lambda x^n}{\Gamma(1+n)} + \frac{\lambda^2 x^{2n}}{\Gamma(1+2n)} + \dots$$

Lo que implica,

$$D^{j-k} e_n(x) = \frac{\lambda x^{n+k-j}}{\Gamma(1+n+k-j)} + \frac{\lambda^2 x^{2n+k-j}}{\Gamma(1+2n+k-j)} + \dots \rightarrow 0$$

para $x \rightarrow 0$, recordemos que $1 \leq j-k \leq m-1 < n$. Esto completa la prueba de (a).

Con respecto al enunciado (b), tenemos.

$$\tilde{y}(x) = \frac{1}{\lambda} \int_0^x q(x-t) u'_0(t) dt + \frac{1}{\lambda} \int_0^x q(x-t) e'_n(t) dt = \frac{1}{\lambda} \int_0^x q(t) e'_n(x-t) dt$$

El primer factor en la integral es continuo por supuesto y el segundo factor es, al menos incorrectamente, integrable.

Así, la integral existe en todas partes y es una función continua de x , y $\tilde{y}(0) = 0$.

Además, para $n > 1$, es decir, $m \geq 2$ tenemos, por las reglas estándar para la diferenciación de integrales de parámetros con respecto al parámetro,

$$D \tilde{y}(x) = \frac{1}{\lambda} \int_0^x q(t) e''_n(x-t) dt + \frac{1}{\lambda} q(x) e'_n(0) \text{ donde } e'_n(0) = 0$$

Por el mismo argumento anterior, y continuidad de q y, en el peor de los casos, singularidad débil de e''_n , vemos que $\tilde{y}'(0) = 0$.

Procediendo de esta manera, encontramos

$$D^k \tilde{y}(x) = \frac{1}{\lambda} \int_0^x q(t) e_n^{(k+1)}(x-t) dt$$

para $k = 0, 1, \dots, m-1$, y en todos estos casos la argumentación da $D^k \tilde{y}(0) = 0$. Por lo tanto, satisface las condiciones iniciales homogéneas requeridas y queda por demostrar que resuelve la ecuación diferencial no homogénea.

Para ello escribimos

$$\begin{aligned} e'_n(u) &= \frac{d}{du} e_n(u) = \frac{d}{du} E_n(\lambda u^n) = \lambda n u^{n-1} E_n(\lambda u^n) \\ &= \lambda n u^{n-1} \sum_{j=1}^{\infty} \frac{j (\lambda u^n)^{j-1}}{\Gamma(1+nj)} = \lambda u^{n-1} \sum_{j=1}^{\infty} \frac{(\lambda u^n)^{j-1}}{\Gamma(jn)} = \sum_{j=1}^{\infty} \frac{\lambda^j u^{nj-1}}{\Gamma(jn)} \end{aligned}$$

Y por lo tanto

$$\begin{aligned} \tilde{y}(x) &= \frac{1}{\lambda} \int_0^x q(t) e'_n(x-t) dt = \frac{1}{\lambda} \int_0^x q(t) \sum_{j=1}^{\infty} \frac{\lambda^j (x-t)^{jn-1}}{\Gamma(jn)} \\ &= \sum_{j=1}^{\infty} \lambda^{j-1} \frac{1}{\Gamma(jn)} \int_0^x q(t) (x-t)^{jn-1} dt = \sum_{j=1}^{\infty} \lambda^{j-1} J_0^{jn} q(x) \end{aligned}$$

Por eso,

$$\begin{aligned} D_{*0}^n \tilde{y}(x) &= \sum_{j=1}^{\infty} \lambda^{j-1} D_{*0}^n J_0^{jn} q(x) + \sum_{j=1}^{\infty} \lambda^{j-1} J_0^{(j-1)n} q(x) \\ &= \sum_{j=1}^{\infty} \lambda^j J_0^{jn} q(x) = q(x) + \sum_{j=1}^{\infty} \lambda^j J_0^{jn} q(x) = q(x) + \lambda \tilde{y}(x) \end{aligned}$$

El intercambio de suma e integración o, respectivamente, suma y diferenciación es posible aquí debido a las propiedades de convergencia de la serie expansión de e'_n y la continuidad de q .

En la práctica, dos casos especiales de ecuaciones diferenciales fraccionarias lineales son de particular importancia porque pueden interpretarse como generalizaciones fraccionarias de dos ecuaciones diferenciales fundamentales de la física de orden entero. Estos dos clásicos las ecuaciones son la ecuación de relajación

$$\begin{cases} D y(x) = -u y(x) + q(x) \\ y(0) = y_0 \end{cases}, \dots, (7.5)$$

$$\text{y la ecuación de oscilación } \begin{cases} D^2 y(x) = -u y(x) + q(x); \\ y(0) = y_0^{(0)}, D y(0) = y_0^{(1)}, \dots, \end{cases} (7.6)$$

donde en ambos casos asumimos $\mu > 0$ por razones físicas.

Comenzamos nuestras consideraciones que conducirán a la generalización fraccionaria deseada, con la ecuación de relajación clásica (7.5). Como es bien sabido, su solución es

$$y(x) = e^{-ux} \left(y_0 + \int_0^x q(t) e^{ut} dt \right) = y_0 e^{-ux} + \int_0^x q(t) e^{u(t-x)} dt$$

que es exactamente el caso $n = 1$ y $\lambda = -\mu$ de (7.2). La solución decae exponencialmente, como $x \rightarrow \infty$.

Por cierto, este comportamiento de desintegración es la razón por la que elegimos $\mu > 0$; La elección $\mu < 0$ conduciría a una explosión de la solución y, por lo tanto, a inestabilidades que son físicamente imposibles en los sistemas generalmente descritos por tal ecuación.

Una generalización natural a nuestro entorno fraccionario consistiría en reemplazar el operador diferencial de primer orden por D_{*0}^n con alguna n adecuada manteniendo la condición inicial sin cambios.

Por lo tanto, dado que nuestra condición inicial sólo se relaciona con y , tenemos se puede elegir un $n \in (0,1)$ arbitrario.

De esta manera derivamos el llamado (simple) ecuación de relajación fraccionaria

$$\begin{cases} D_{*0}^n y(x) = -u y(x) + q(x); \\ y(0) = y_0 \end{cases}, \dots, (7.7)$$

con $0 < n < 1$ y $\mu > 0$. La solución de la ecuación homogénea viene dada por

$$y_{hom}(x) = y_0 E_n(-u x^n)$$

según el teorema 7.2, y resulta que esta solución se comporta de una manera que difiere de alguna manera del comportamiento mencionado anteriormente en el caso $n = 1$.

Teorema 7.3

Sea $\mu > 0$, $0 < n < 1$ y $u_0(x) = E_n(-u x^n)$

La función es completamente monótona en $(0, \infty)$, es decir $(-1)^k D^k u_0(x) \geq 0$ para cada $x > 0$ y cada $k \in \mathbb{N}_0$.

Para $x \rightarrow \infty$ tenemos

$$u_0(x) = \frac{x^{-n}}{u \Gamma(1-n)} (1 - o(1))$$

En otras palabras, la solución decae algebraicamente, es decir, mucho más lentamente que la decadencia exponencial conocida por la ecuación de relajación clásica de primer orden.

La declaración:

(a) Implica que la solución es monótonamente decreciente ya que Du_0 no es positiva; mientras, el enunciado

(b) Muestra la velocidad de la desintegración. Un proceso que muestra esta desintegración muy lenta

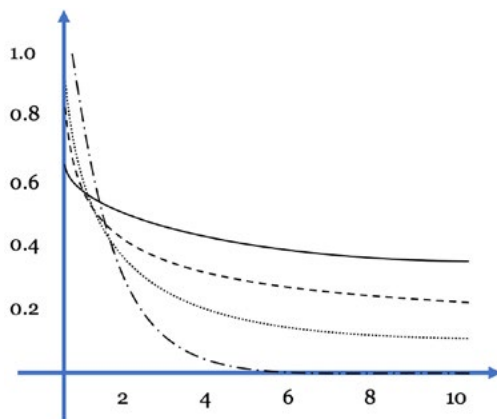
El comportamiento a veces se denomina ultralento (Hille, 1969). La demostración clásica de este teorema es basada en un análisis de las propiedades de la transformada de Laplace de en el plano complejo. Requiere métodos que van más allá del alcance de este texto y, por lo tanto, no incluirlo aquí. Los detalles de esta prueba se pueden encontrar en las obras originales de Gorenflo et al. (1999) y Mainardi (2012). Sin embargo, es posible proporcionar una prueba alternativa que evite el uso de técnicas de transformada de Laplace.

Proporcionaremos la maquinaria necesaria en el ítem 7.3 y volver allí a esta cuestión.

Por el momento sólo presentamos los gráficos de la Figura 3 que muestran el caso $\mu = 1$ para el propósito de la ilustración.

Observe que la línea de puntos y guiones correspondiente a $n = 1$ muestra la solución clásica e^{-x} ; es evidente que esta función decae mucho más rápido que los otros tres.

Figura 3. Gráficas de $E_n(-x^n)$



Fuente: elaboración propia

Nota. Referencia Kai Diethelm, Gráficas de $E_n(-x^n)$ para $n = 1/4$ (línea continua), $n = 1/2$ (línea discontinua), $n = 3/4$ (línea punteada) línea), y $n = 1$ (línea de guiones y puntos)

De manera similar, podemos ver que la simple ecuación de oscilación fraccionaria

$$\begin{cases} D_{+0}^n y(x) = -u y(x) + q(x); \\ y(0) = y_0^{(0)}, Dy(0) = y_0^{(1)}, \dots, \end{cases} \quad (7.8)$$

con $1 < n < 2$ y $\mu > 0$ generaliza la ecuación de oscilación clásica (7.6).. El problema homogéneo correspondiente tiene dos soluciones linealmente independientes $u_0(x) = E_n(-u x^n)$ y $u_1(x) = J_0 u_0(x)$ según el teorema 7.2. Podemos decir algo sobre el comportamiento de estos dos. funciones también.

Teorema 7.4

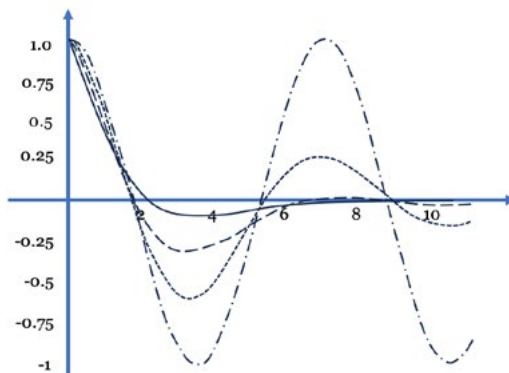
Sean $\mu > 0$, $1 < n < 2$ y $u_0(x) = E_n(-u x^n)$ y $u_1(x) = J_0 u_0(x)$

Entonces, para $x \rightarrow \infty$ tenemos

$$u_0(x) = \frac{x^{-n}}{u \Gamma(1-n)} (1 - \sigma(1)), \text{ y } u_1(x) = \frac{x^{1-n}}{u \Gamma(2-n)} (1 - \sigma(1))$$

Así, la desintegración es una vez más algebraica, pero, como se observa en las figuras 7.2 y 7.3, mostrando nuevamente el caso $\mu = 1$, no monótono. En particular, redescubrimos las soluciones clásicas $u_0(x) = \cos x$ y $u_1(x) = \sin x$ en las gráficas para $n = 2$; estos las funciones, por supuesto, no decaen en absoluto. En Hilfer (2000), se da una demostración del teorema. Los teoremas 7.3 y 7.4 nos dan información valiosa sobre la velocidad del decaimiento de la función $u_0(x) = E_n(-u x^n)$ y su primera primitiva u_1 una cuestión importante relacionada es lo relativa a los cambios de signo de u_0 . Algunas propiedades simples ya se pueden leer en las figuras. 3 y 4. Una declaración más precisa está contenida en el siguiente resultado

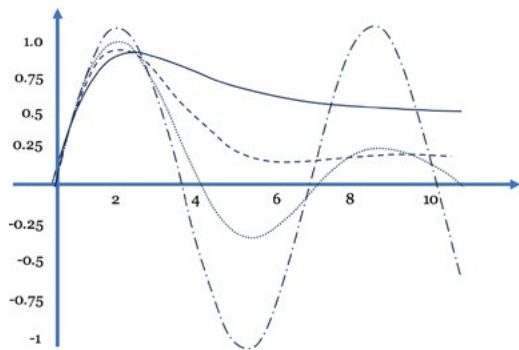
Figura 4. Esquema de las gráficas de $u_0(x)$



Fuente: elaboración propia.

Nota. Referencia Kai Diethelm. Esquemas de las gráficas de $u_0(x)$ para $n = 5/4$ (línea continua), $n = 3/2$ (línea discontinua), $n = 7/4$ (línea punteada), y $n = 2$ (línea de puntos y guiones). Se eligió $\mu = 1$ en todos los casos.

Figura 5. Esquemas de las gráficas de $u_1(x)$



Fuente: elaboración propia.

Nota. Referencia Kai Diethelm Esquemas de las gráficas de $u_1(x)$ para $n = 5/4$ (línea continua), $n = 3/2$ (línea discontinua), $n = 7/4$ (línea punteada), y $n = 2$ (línea de puntos y guiones). $\mu = 1$ fue elegido en todos los casos.

Teorema 7.5

Considere la función $u_0(x) = E_n(-u x^n)$ con algún $\mu > 0$.

En el caso $0 < n < 1$, u_0 no tiene ningún cero en $[0, \infty)$

En el caso $1 < n < 2$, el número de ceros de u_0 en $[0, \infty)$ (contando multiplicidades) es finito y extraño,

Recuerde que en el caso $n = 2$ teníamos que tiene infinitos ceros en $[0, \infty)$.

Prueba

El teorema 7.3 nos dice que, para $0 < n < 1$, u_0 decae monóto-

namente en $[0, \infty)$ y que $\lim_{x \rightarrow \infty} u_0(x) = 0$. Por lo tanto, no puede tener ceros, lo que prueba (a).

Para el enunciado (b), primero observamos que la propiedad $u_0(0) = 1$ es una consecuencia directa de la definición de u_0 y el desarrollo en serie de potencias de la función de Mittag-Leffler E_n .

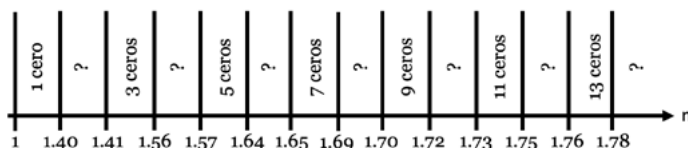
Además, observamos que el resultado asintótico del teorema 7.4 nos dice que, para x grande, $u_0(x)$ se comporta como $\frac{x^{-n}}{\Gamma(1-n)}$. Como ahora tenemos $1 < n < 2$ encontramos que $\Gamma(1-n) < 0$ y por tanto vemos que tiende a 0 desde abajo. Por lo tanto para x suficientemente grande. De ello se deduce que el número de ceros de debe ser finito y extraño.

Sin duda sería interesante saber la ubicación precisa de los ceros de u_0 en el caso $1 < n < 2$ o los valores de n para los cuales cambia el número de ceros. El Actualmente el autor sólo conoce algunos primeros pasos hacia una respuesta concluyente a estas preguntas se pueden encontrar en el trabajo de Gorenflo y Mainardi (Hille, 1976). Sus resultados tratan del caso normalizado $\mu = 1$. Se pueden utilizar sustituciones obvias para transferirlos al caso general. De hecho, han calculado el Mittag-Leffler funciones numéricamente y encontró las relaciones entre n y el número de ceros de u_0 indicado en la figura 7.4.

La transición de k a $k + 2$ ceros, digamos, toma lugar en los espacios marcados con un signo de interrogación. Desafortunadamente, análisis más precisos Los resultados relativos a la ubicación de los puntos de transición en el caso general parecen no estar disponible en este momento. Además de estos datos relacionados con los puntos de transición, Gorenflo y Mainardi también proporciona alguna información sobre la ubicación de los ceros de u_0 .

En particular prueban los siguientes resultados.

Figura 6. Figura de número de ceros de u_0



Fuente: elaboración propia

Notas. Número de ceros de u_0 para $\mu = 1$ y varios valores de n (eje horizontal no a escala) en los espacios marcados con un signo de interrogación.

Desafortunadamente, análisis más precisos

Los resultados relativos a la ubicación de los puntos de transición en el caso general parecen no estar disponible en este momento.

Además de estos datos relacionados con los puntos de transición, Gorenflo y Mainardi (1996), también proporcione alguna información sobre la ubicación de los ceros de u_0 . En particular prueban los siguientes resultados.

Teorema 7.6

Si $\varepsilon > 0$ es suficientemente pequeño entonces la función $u_0(x) = E_n(-u x^n)$ con $n = 1 + \varepsilon$ tiene exactamente un cero $x^* > 0$. Este cero satisface la relación $x^* = (1 + \sigma(1)) \ln(2/\varepsilon)$ como $\varepsilon \rightarrow 0$

Por tanto, resulta x^* crece hacia ∞ a un ritmo muy lento cuando $\varepsilon \rightarrow 0$. La prueba se basa nuevamente en la estimación de la transformada de Laplace de u_0 en el dominio complejo; nos referimos a Haubold et al. (2009), para más detalles. Observe sin embargo que la aproximación $x^* = (1 + \sigma(1)) \ln(2/\varepsilon)$ en realidad ya es bastante preciso para valores moderadamente grandes de ε .

Tomemos, por ejemplo, $\varepsilon = 1/4$, es decir, $n = 5/4$. Luego, el valor estimado para la ubicación. del único cero de u_0 es $\ln(2/\varepsilon) = \ln 8 \approx 2,08$, que coincide con el valor exacto muy muy bien, como puede verse al observar la gráfica de u_0 en la figura 6.2 (la línea continua).

De manera similar se puede investigar el comportamiento cuando $n \rightarrow 2$

Teorema 7.7

La función $u_0(x) = E_n(-u x^n)$ con $n = 2 - \varepsilon$ tiene $Z(\varepsilon)$ ceros en $[0, \infty)$, donde $Z(\varepsilon)$ es un número impar para todo $\varepsilon \in (0, 1)$ que satisface la

$$Z(\varepsilon) = 12\pi^{-2}\varepsilon^{-1}(1 + \sigma(1))\ln\varepsilon^{-1} \text{ cuando } \varepsilon \rightarrow 0.$$

Denotando el mayor de estos ceros por x^* , tenemos que

$$x^*(\varepsilon) = 12\pi^{-2}\varepsilon^{-1}(1 + \sigma(1))\ln\varepsilon^{-1} \text{ como } \varepsilon \rightarrow 0.$$

que la aproximación asintótica es menos precisa que la del teorema 7.6 en el sentido de que ahora ε debe estar mucho más cerca de 0 para que la aproximación sea buena en términos de error absoluto.

Una fórmula asintótica para la ubicación de todos los ceros de las funciones de Mittag-Leffler, es decir, no sólo los ceros situados en el eje real negativo, ha sido proporcionado por Sedlets-kij (1994); Seybold & Hilfer (2008). Su resultado se puede aplicar a funciones de Mittag-Leffler de dos parámetros con parámetros arbitrarios.

En el caso especial que hemos tratado en el dos teoremas anteriores, dice lo siguiente.

Teorema 7.8

Sea $0 < n_1 < 2$ y $n_2 > 2$ Entonces, para cada $m \in \mathbb{Z}$ tenemos $E_{n_1 n_2}(Z_m) = 0$ donde

$$Z_m(x) = \left(2\pi i m - n_1 \tau_{n_2} \left(\ln 2\pi |m| + i \frac{\pi}{2} \operatorname{sgn} m \right) + \gamma_{n_2} + O(|m|^{-n_1}) + O(|m|^{-1} \ln |m|) \right)^{n_1}$$

Con algunas τ_{n_2} y γ_{n_2} que solo dependen de n_2

Nos remitimos al artículo original de Sedletskij (1994); Seybold & Hilfer (2008), para una prueba de este teorema y para obtener información adicional, en particular para obtener más resultados sobre las funciones de Mittag Leffler con parámetros elegidos de manera diferente.

En este contexto también podemos señalar que una discusión sobre la ubicación de lo no real Es posible que ya se encuentren ceros de funciones de Mittag-Leffler de un parámetro en los primeros años. artículo de Wiman (Diethelm, 2010). Además, el trabajo reciente de Seybold y Hilfer (2005) (Miller & Ross, 1993), contiene algunos datos numéricos sobre la ubicación de los ceros de los dos parámetros Función Mittag-Leffler $E_{0,8,0,9}$. Sus métodos también podrían ser aplicables para obtener datos correspondientes a otras funciones de Mittag-Leffler. Finalmente mencionamos el trabajo de Gorenflo et al. (1999), quienes han proporcionado un algoritmo numérico para el cálculo de valores de funciones de Mittag-Leffler y sus derivadas. Por supuesto, una La combina-

ción de este método con un algoritmo estándar de búsqueda de raíces como el esquema de Newton Raphson puede producir valores numéricos para los ceros de las funciones de Mittag-Leffler. con parámetros dados.

Las consideraciones en esta sección hasta ahora sólo se han referido a ecuaciones lineales. con coeficientes constantes. Si permitimos ecuaciones con coeficientes no constantes, entonces la teoría se vuelve mucho más engorrosa. Sin embargo, todavía hay algunos resultados que podemos probar. Por tanto, el objeto de interés es ahora la ecuación diferencial fraccionario.

$$\begin{cases} D_{+0}^n y(x) = f(x)y(x) + q(x), (7.9a) \\ y^{(k)}(0) = y_0^{(k)}, (K = 0, 1, 2, 3, \dots, [n] - 1), \dots (7.9b) \end{cases}$$

Sujeto a las condiciones iniciales,

$$y^{(k)}(0) = y_0^{(k)}, (K = 0, 1, 2, 3, \dots, [n] - 1), \dots (7.9b)$$

comenzamos con un teorema muy simple.

Teorema 7.9

Sea $n > 0$, $m = n$ y $f, g \in C[0, h]$ con algún $h > 0$. Entonces, el problema con valores iniciales (7.9) tiene una solución única $y \in C^{m-1}[0, h]$.

Prueba

El hecho de que el problema de valor inicial tenga una solución continua única en el intervalo completo $[0, h]$ se deriva de la existencia básica y unicidad resultado dado en el teorema 6.8. La propiedad de diferenciabilidad es una consecuencia de Teorema 6.25.

Un método teórico que en ocasiones proporciona información útil sobre la solución del problema del valor inicial (7.9) se

basa en un concepto cuyos fundamentos son similares a un concepto utilizado en la demostración del teorema 7.8.

Específicamente reescribimos el inicial, problema de valor en la forma de Volterra y reorganizar algunos de los términos, obteniendo así

$$\begin{aligned} y(x) &= \sum_{k=0}^{m-1} y_0^{(k)} \frac{x^k}{k!} + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} (f(t)y(t) + g(t)) dt \\ &= \sum_{k=0}^{m-1} y_0^{(k)} \frac{x^k}{k!} J_0^n g(x) + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t)y(t) dt, \dots (8.10) \end{aligned}$$

Aquí ahora definimos

$$T^*(x) := \sum_{k=0}^{m-1} y_0^{(k)} \frac{x^k}{k!} J_0^n g(x), \dots (7.11)$$

y vea que T^* es continua en $[0, h]$ si g es continua en el mismo intervalo. Introducción a la nota

$$k(x, t) := \frac{1}{\Gamma(n)} (x-t)^{n-1} f(t)$$

luego escribimos,

$$\phi_0(x) := T^*(x)$$

Y

$$\phi_j(x) := J_0^n (f \cdot \phi_{j-1})(x) = \int_0^x k(x, t) \phi_{j-1}(t) dt, (j = 1, 2, \dots)$$

Finalmente definimos el j -ésimo núcleo iterado k_j para $j = 1, 2, \dots$ mediante la relación de recurrencia

$$k_1(x, t) = k(x, t) \text{ y } k_j(x, t) := \int_0^x k(x, \tau) k_{j-1}(\tau, t) d\tau, (j = 1, 2, \dots) \text{ y}$$

Armados con estas herramientas, podemos demostrar una representación explícita de la solución de (7.9)

Teorema 7.10

Sean $f, g \in C[0, h]$ y $n > 0$. Entonces, la solución única de la ecuación inicial 7.9 y del problema de valores en el intervalo $[0, h]$ se puede expresar en la forma

$$y(x) = T^*(x) + \int_0^x R(x, t) T^*(t) dt$$

donde se define en (7.11) y la función R está dada por

$$R(x, t) = \sum_{j=1}^{\infty} k_j(x, t)$$

En particular, esta serie es uniformemente convergente.

Definición 7.1

La función R introducida en el teorema 7.10 se llama núcleo resolutivo de la ecuación de Volterra (7.10).

Para la prueba necesitamos algunas declaraciones auxiliares.

Lema 7.11

Sean $f, g \in C[0, h]$ y $n > 0$. Para $j = 1, 2, \dots$, tenemos:

a) k_j es continua en $\{(x, t): 0 \leq t < x \leq h\}$

b) $|k_j(x, t)| \leq \frac{1}{\Gamma(jn)} \|f\|_0^j (x - t)^{jn-1}$

Y k_j es continua en el triángulo cerrado $\{(x, t); 0 \leq t \leq x \leq h\}$ para $j \geq 1/n$.

$$\phi_j(x) = \int_0^x R(x, t) T^*(t) dt$$

Prueba

La parte (a) se deriva directamente de la definición de k_j .

La prueba de (b) se basa en el principio de inducción.

La base de inducción ($j = 1$) es una consecuencia inmediata de la definición de k_1 .

Para el paso de inducción ($j-1 \rightarrow j$) tenemos, en vista de la definición de k y la hipótesis de inducción,

$$\begin{aligned} |k_j(x, t)| &= \left| \int_0^x k(x, \tau) k_{j-1}(\tau, t) d\tau \right| \\ &\leq \|f\|_\infty \frac{1}{\Gamma(n)} \frac{\|f\|_\infty^{j-1}}{\Gamma((j-1)n)} \int_0^x (x-\tau)^{n-1} (\tau-t)^{(j-1)n} dt \\ &\leq \frac{1}{\Gamma(n)} \frac{\|f\|_\infty^{j-1}}{\Gamma((j-1)n)} \int_0^x (x-\tau)^{n-1} (\tau-t)^{(j-1)n} dt \\ &= \frac{1}{\Gamma(n)} \frac{\|f\|_\infty^{j-1}}{\Gamma((j-1)n)} (x-t)^{jn} \frac{\Gamma(n)\Gamma((j-1)n)}{\Gamma(jn)} \end{aligned}$$

lo que demuestra la desigualdad deseada. La continuidad es una consecuencia directa de esto.

Para el enunciado (c), también procedemos por inducción matemática. la inducción base ($j = 1$) es una consecuencia inmediata de las definiciones de T^* , K_1 y ϕ_1 .

Para el paso de inducción ($j-1 \rightarrow j$), usamos la definición de ϕ_j y la hipótesis de inducción y encontrar

$$\phi_j(x) = \int_0^x k(x, t) \phi_{j-1}(t) dt = \int_0^x k(x, t) \int_0^t k_{j-1}(t, \tau) T^*(\tau) d\tau dt$$

En vista de las declaraciones (a) y (b), la integrabilidad de todas las funciones involucradas aquí, no es un problema y, en particular, podemos intercambiar el orden de integración. Este rendimiento

$$\begin{aligned}
\phi_j(x) &= \int_0^x \int_{\tau}^x k(x, t) k_{j-1}(t, \tau) T^*(\tau) dt d\tau \\
&= \int_0^x T^*(\tau) \int_{\tau}^x k(x, t) k_{j-1}(t, \tau) dt d\tau \\
&= \int_0^x T^*(\tau) k_j(x, \tau) d\tau
\end{aligned}$$

Como es requerido.

Prueba (del teorema 7.10)

Primero mostramos la convergencia uniforme de la serie que define el núcleo resolutivo.

En este contexto, los primeros sumandos $1/n$ no juegan un papel importante, rol; basta con mirar la serie $\sum_{j=[1/n]+1}^{\infty} k_j(x, t)$. Para esta serie, Lema 8.11 (b) proporciona al mayor

$$\sum_{j=[1/n]+1}^{\infty} \frac{\|f\|_{\infty}^{j-1}}{\Gamma(jn)} (x-t)^{jn-1} \leq \sum_{j=[1/n]+1}^{\infty} \frac{\|f\|_{\infty}^{j-1}}{\Gamma(jn)} h^{jn-1} = \frac{1}{h} \sum_{j=[1/n]+1}^{\infty} \frac{(\|f\|_{\infty} h^n)^j}{\Gamma(jn)}$$

que, utilizando la prueba de la raíz, se ve fácilmente que converge.

Esto implica la convergencia uniforme, destacamos en este punto que este resultado de convergencia uniforme no implica la continuidad de R desde los primeros sumandos de la serie puede ser ilimitada y por tanto discontinuo.

Sin embargo, en vista del Lema 7.11 (a), R es al menos continua en el conjunto

$$\{(x, t): 0 \leq t < x \leq h\}$$

Además, podemos combinar este resultado de convergencia uniforme con el enunciado del Lema 7.11 (c) para concluir que la serie $\sum_{j=1}^{\infty} \phi_j$ también es uniformemente convergente.

Para demostrar que la solución y se puede expresar de la manera indicada, observamos que (7.10) y (7.11) revelan que

$$y(x) = T^*(x) + \int_0^x k(x, t)y(t)dt$$

Por otro lado tenemos,

$$\begin{aligned} \int_0^x R(x, t) T^*(t) dt &= \sum_{j=1}^{\infty} \int_0^x k_j(x, t) T^*(t) dt = \sum_{j=1}^{\infty} \phi_j(x) \\ &= \sum_{j=1}^{\infty} \int_0^x k(x, t) \phi_j(t) dt = \int_0^x k(x, t) \sum_{j=1}^{\infty} \phi_j(t) dt \end{aligned}$$

en vista de la definición de R , Lema 7.11 (c) y la definición de ϕ_j . Así, Para completar la prueba basta demostrar que

$$\sum_{j=0}^{\infty} \phi_j(x) = y$$

Ya habíamos notado que la serie en el lado izquierdo de esta ecuación es uniformemente convergente.

Como todos sus sumandos son continuos en $[0, h]$, el límite la función también es continua. Nuestra prueba de identidad se basa en nuestro conocimiento de que y es la solución única de la ecuación integral

$$y(x) = \sum_{k=0}^{m-1} y_0^{(k)} \frac{x^k}{k!} + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} (f(t)y(t) + g(t)) dt$$

Si podemos demostrar que la serie $\sum_{j=1}^{\infty} \phi_j$ también resuelve esta ecuación integral, habríamos logrado nuestro objetivo.

Para ello utilizamos las definiciones de T^*, k, ϕ_0 y ϕ_j , $j=1, 2, 3, \dots$ y nótese que la convergencia uniforme nos permite intercambiar suma e integración para escribir

$$\begin{aligned}
\sum_{j=0}^{\infty} \phi_j(x) &= T^*(x) + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t) \sum_{j=1}^{\infty} \phi_{j-1}(t) dt \\
&= \sum_{k=0}^{m-1} y_0^{(k)} \frac{x^k}{k!} + J_0^n g(x) + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t) \sum_{j=1}^{\infty} \phi_j(t) dt \\
&= \sum_{k=0}^{m-1} y_0^{(k)} \frac{x^k}{k!} + J_0^n g(x) + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} (f(t) \sum_{j=1}^{\infty} \phi_j(t) + g(t)) dt
\end{aligned}$$

Ocasionalmente resulta interesante aumentar la representación en serie explícita del núcleo resolutorio anterior por el siguiente resultado que indica que el núcleo resolutorio mismo es la solución de una ecuación integral que está estrechamente relacionada con la original.

Teorema 7.12

Bajo los supuestos del Teorema 6.10, el núcleo resolutorio R satisface la ecuación resolutoria

$$R(x, t) = k(x, t) + \int_t^x k(x, \tau) R(\tau, t) d\tau, \text{ para } 0 \leq t < x \leq h$$

Prueba

De la representación en serie del teorema 7.10 y del hecho de que es uniforme La convergencia nos permite intercambiar suma e integración, deducimos

$$\begin{aligned}
\int_t^x k(x, \tau) R(\tau, t) d\tau &= \int_t^x k(x, \tau) \sum_{j=1}^{\infty} k_j(\tau, t) d\tau = \sum_{j=1}^{\infty} \int_t^x k(x, \tau) k_j(\tau, t) d\tau = \\
&= \sum_{j=1}^{\infty} k_{j+1}(x, t) = R(x, t) - k(x, t)
\end{aligned}$$

En el teorema 7.2 y la observación 7.1 habíamos desarrollado una serie de representaciones para la solución de problemas con valores iniciales para ecuaciones lineales con coeficientes constantes.

En el caso multidimensional, es decir, en el caso en que la solución desconocida y es un mapeo a $\mathbb{R}^N$ para algún N y λ es una matriz $N \times N$, de ahora en adelante denotada por Λ , estas representaciones requieren la evaluación de funciones de Mittag-Leffler de Argumentos valorados en matrices. Si bien esto no supone ningún problema en teoría, ya que se sabe que las series relevantes son convergentes, en la práctica es un método sumamente engorroso. Nosotros, Por lo tanto, presentamos ahora un enfoque alternativo que suele ser mucho más sencillo de manejar. Nuestro camino sigue las líneas trazadas por Odibat (2010); (Nonnenmacher & Metzler, 1995) y Bonilla et al. (2007); (Basset, 1888). Con el fin de conservar una estrecha relación con la técnica clásica para ecuaciones de primer orden (Chen et al., 2007), restringimos nuestra atención a ecuaciones diferenciales de orden $n \in (0,1)$.

Consideremos así la ecuación diferencial fraccionaria

$$D_{+0}^n y(x) = \Lambda y(x) + q(x), \dots (7.12)$$

con $0 < n < 1$, una matriz Λ , $N \times N$, una función dada $q: [0, h] \rightarrow \mathbb{C}^N$ y una incógnita

solución $y: [0, h] \rightarrow \mathbb{C}^N$. Primero discutiremos la construcción de la solución general. Para esta ecuación; la condición inicial que luego conduce a una solución especial será añadido más tarde.

Como siempre comenzamos con el problema homogéneo correspondiente a (7.12), es decir con el caso $q \equiv 0$.

En la situación clásica $n = 1$ sabemos que podemos usar una aproximación de la forma $y(x) = u \exp(\Lambda x)$ con un vector adecuado u . Desde que tenemos Ya hemos encontrado que la función de Mittag-Leffler $E_n(\lambda x^n)$ toma el papel de $\exp(\lambda x)$. En el caso unidimensional, es natural buscar una solución que sea una combinación lineal. de expresiones de la forma

$$y(x) = u E_n (\lambda x^n) ,(7.13)$$

con vectores adecuados $u \in \mathbb{C}^N$ y escalares $\lambda \in \mathbb{C}$ que deben determinarse. Insertar (7.13) en la ecuación diferencial homogénea dada

$$D_{\alpha}^n y(x) = \Lambda y(x) ,(7.14)$$

obtenemos, en vista del teorema 4.3,

$$u \lambda E_n (\lambda x^n) = \Lambda u E_n (\lambda x^n)$$

Dado que $E_n (\lambda x^n) \neq 0$, esto se deduce del Teorema 7.5 para $\lambda < 0$ y directamente de la representación en serie de potencias de E_n para $\lambda \geq 0$; para el complejo λ también se puede demostrar, esto implica

$$\lambda u = \Lambda u$$

es decir, λ debe ser un valor propio de la matriz Λ , y u debe ser un vector propio correspondiente.

Ahora, si todos los valores propios k veces mayores de Λ tienen k vectores propios, entonces sabemos que el conjunto de todos estos vectores propios es linealmente independiente y por lo tanto forma una base de $\mathbb{C}^N$. Así hemos mostrado:

Teorema 7.13

Sean $\lambda_1, \lambda_2, \dots, \lambda_N$ los valores propios de Λ y $u(1), \dots, u(N)$ sea los vectores propios correspondientes. Entonces, la solución general de la ecuación diferencial (7.14) tiene la forma

$$y(x) = \sum_{\ell=1}^N c_{\ell} u^{(\ell)} E_n (\lambda_{\ell} x^n)$$

con ciertas constantes $c_\ell \in \mathbb{C}$. La solución única de esta ecuación diferencial sujeta a la condición inicial $y(0) = y_0$ se caracteriza por el sistema lineal.

$$y_0 = (u^{(1)}, u^{(2)}, \dots, u^{(N)}) (c_{(1)}, c_{(2)}, \dots, c_{(N)})^T$$

El sistema lineal para el cálculo de los coeficientes c de la solución de Si el problema con valor inicial se obtiene estableciendo $x = 0$ en la solución general. el hecho de que el sistema lineal tenga una solución única se debe a la independencia lineal mencionada anteriormente de los vectores propios que forman la matriz de coeficientes de este sistema.

Ejemplo 7.2

Los valores propios de

$$A = \begin{pmatrix} 2 & -1 \\ 4 & -3 \end{pmatrix}$$

son 1 (con vector propio $(1,1)^T$ y -2 (con vector propio $(1,4)^T$; así la solución general de la ecuación diferencial

$$D_{x0}^n y(x) = \Lambda y(x), \text{ con } 0 < n < 1 \text{ es}$$

$$y(x) = c_1 \begin{pmatrix} 1 \\ 1 \end{pmatrix} E_n(1x^n) + c_2 \begin{pmatrix} 1 \\ 4 \end{pmatrix} E_n(-2x^n)$$

con constantes arbitrarias c_1 y c_2 .

La única solución que satisface las necesidades iniciales.

La condición $y(0) = (4, -3)^T$ y se obtiene eligiendo $c_1 = -\frac{7}{3}$ y $c_2 = 19/3$.

Por supuesto, es bien sabido que existen matrices que no tienen un conjunto completo de vectores propios, es decir, matrices que tienen un valor propio k veces mayor con algún $k > 1$ al que Hay menos de k vectores propios disponibles. En este caso, la teoría del teorema 7.13 no es aplicable, y debemos recurrir a una representación diferente de la solución general.

Específicamente sabemos por Álgebra lineal elemental que, dado cualquier matriz Λ cuadrado (real o compleja), existe una matriz B no singular tal que $\Phi = B^{-1} \Lambda B$ tiene la llamada forma Jordan

$$\Phi = \begin{pmatrix} \Phi_1 & 1 & \dots & 0 \\ 0 & \Phi_2 & \ddots & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & \Phi_k \end{pmatrix}, \dots (7.15)$$

donde las cajas de Jordan son matrices cuadradas de la forma

$$\Phi = \begin{pmatrix} \lambda_u & 1 & 0 & \dots & 0 \\ 0 & \lambda_u & 1 & \dots & 0 \\ 0 & 0 & \lambda_u & \ddots & 0 \\ \vdots & \ddots & \ddots & \ddots & \vdots \\ 0 & 0 & 0 & \dots & 0 \\ 0 & 0 & 0 & \lambda_u & 1 \\ 0 & 0 & 0 & 0 & \lambda_u \end{pmatrix}, \dots (7.16)$$

con r_u filas y columnas.

En particular tenemos $r_1 + r_2 + \dots + r_k = N$, y los valores propios de Λ se pueden encontrar en la diagonal principal de Φ en sus respectivas multiplicidades.

Observamos que todas las entradas diagonales de cada bloque de Jordan Φ_k coinciden entre sí.

Sin embargo, diferentes bloques de Jordan pueden tener las mismas entradas diagonales.

Por ejemplo, en la representación de Jordan de la matriz unitaria N -dimensional tenemos $k = N$, $r_u = 1$, para todos μ y $\lambda_\mu = 1$ para todos μ .

En vista de la estructura simple de un bloque de Jordan, el sistema

$$D_{*0}^n z(x) = \Phi_k z(x), \text{ es decir } \begin{cases} D_{*0}^n z_1(x) = \lambda_k z_1(x) + z_2(x) \\ \vdots \\ D_{*0}^n z_{r_k-1}(x) = \lambda_k z_{r_k-1}(x) + z_{r_k}(x) \\ D_{*0}^n z_{r_k}(x) = \lambda_k z_{r_k}(x) \end{cases}$$

se resuelve fácilmente mediante sustitución hacia atrás y una aplicación del teorema 7.2 para cada componente de z . Al evaluar explícitamente las integrales que surgen en el teorema 7.2, tenemos que encontrar que las columnas de la matriz $Z = (z_{i,j})_{i,j=1}^{rk}$ con

$$z_{i,j} = \begin{cases} 0, & \text{si } i > j \\ x^{(j-1)n} D^{j-i} E_n(\lambda_k x^n), & \text{caso contrario} \end{cases}$$

formar soluciones linealmente independientes de $D_{*0}^n z(x) = \Phi_k z(x)$. Observando eso

$$E_n(\Phi x^n) = \begin{pmatrix} E_n(\Phi_1 x^n) & 1 & \cdots & 0 \\ 0 & E_n(\Phi_2 x^n) & \ddots & \vdots \\ \vdots & \ddots & & \\ 0 & \cdots & 0 & E_n(\Phi_k x^n) \end{pmatrix}$$

De ello se deduce que: podemos combinar los resultados de los bloques individuales de Jordan en uno solo sistema de soluciones linealmente independientes (y por tanto una base para el espacio de soluciones) del sistema completo $D_{*0}^n z(x) = \Phi z(x)$ de forma sencilla. Ilustremos esto procedimiento por el ejemplo

$$\Phi = \begin{pmatrix} \lambda_1 & 1 & 0 & 0 & 0 & 0 \\ 0 & \lambda_1 & 1 & 0 & 0 & 0 \\ 0 & 0 & \lambda_1 & 0 & 0 & 0 \\ 0 & 0 & 0 & \lambda_2 & 0 & 0 \\ 0 & 0 & 0 & 0 & \lambda_3 & 1 \\ 0 & 0 & 0 & 0 & 0 & \lambda_3 \end{pmatrix}$$

Aquí tenemos $k=3$, $r_1=3$, $r_2=1$, $r_3=2$.

Esto da lugar a que las soluciones sean las columnas de la matriz

$$\Phi = \begin{pmatrix} E_n(\lambda_1 x^n) & x^n E'_n(\lambda_1 x^n) & x^n E''_n(\lambda_1 x^n) & 0 & 0 & 0 \\ 0 & E_n(\lambda_1 x^n) & x^n E'_n(\lambda_1 x^n) & 0 & 0 & 0 \\ 0 & 0 & E_n(\lambda_1 x^n) & 0 & 0 & 0 \\ 0 & 0 & 0 & E_n(\lambda_2 x^n) & 0 & 0 \\ 0 & 0 & 0 & 0 & E_n(\lambda_3 x^n) & x^n E'_n(\lambda_3 x^n) \\ 0 & 0 & 0 & 0 & 0 & E_n(\lambda_3 x^n) \end{pmatrix}$$

Pero ahora, en vista de la relación $\Phi = B^{-1} \Lambda B$, la solución y del sistema (7.14) es relacionado con la solución z creada anteriormente de una manera muy simple:

Tenemos $\Lambda = B\Phi B^{-1}$ y por tanto (7.14) se convierte en $D_{*0}^n y = B\Phi B^{-1}y$ es decir, $D_{*0}^n B^{-1}y = \Phi B^{-1}y$.

Presentando la sustitución $z = B^{-1}y$ o equivalentemente, $y = Bz$, llegamos a la ecuación $D_{*0}^n z(x) = \Phi z(x)$, que acabamos de resolver.

Así, obtenemos la solución deseada y de (7.14) en la forma

$$y(x) = Bz(x)$$

donde $z(x)$ tiene la forma indicada. Así hemos demostrado el siguiente resultado.

Teorema 7.14

Para cada valor propio k veces mayor λ de la matriz Λ tenemos: k linealmente soluciones independientes de la ecuación diferencial lineal homogénea (7.14) que se puede representar en la forma

$$y_\ell(x) = \pi^{(\ell)}(x), \text{ con } \ell = 1, 2, 3, \dots, k$$

donde el $\pi^{(\ell)}(x)$ son vectores N -dimensionales cuyas funciones componentes son $\pi_j^{(\ell)}, j = 1, 2, 3, \dots, N$, son de la forma

$$\pi_j^{(\ell)}(x) = \sum_{u=1}^{\ell-1} c_j^{(\ell,u)} x^{un} D^n E_n(\lambda x^n), \dots, (7.17)$$

La combinación de estas soluciones para todos los valores propios conduce a N soluciones linealmente independientes del

sistema (7.14), es decir, a una base del espacio de todas las soluciones de este sistema.

Tenga en cuenta que $\pi^{(1)}(x) = c^{(1,0)} E_n(\lambda x^n)$, por lo que $c^{(1,0)}$ debe ser un vector propio de Λ asociado al valor propio λ .

Observación 7.2

En el caso límite $n \rightarrow 1$, el teorema 7.14 se reduce al bien conocido observación clásica de que las k soluciones de $Dy = \Lambda y$ que corresponden al k -pliegue El valor propio λ de la matriz Λ tiene la forma:

$$y(x) = \pi_{k-1} \exp(\lambda x)$$

Donde π_{k-1} es un polinomio vectorial de grado como máximo $k-1$ con coeficientes adecuadamente elegidos.

Ejemplo 7.3

$$\Lambda = \begin{pmatrix} 1 & 1 & 1 \\ 2 & 1 & 1 \\ 0 & 1 & 1 \end{pmatrix}$$

tiene el valor propio único -1 con vector propio $(-3, 4, 2)^T$ y el valor propio doble 2 con vector propio $(0, 1, -1)^T$ pero sin segundo vector propio.

Así, la solución general para la ecuación diferencial $D_{*0}^n y(x) = \Lambda y(x)$ con $0 < n < 1$ toma la forma:

$$y(x) = c_1 \begin{pmatrix} -3 \\ 4 \\ 2 \end{pmatrix} E_n(-x^n) + c_2 \begin{pmatrix} 0 \\ 1 \\ -1 \end{pmatrix} E_n(2x^n) \\ + c_3 [c^{(2,0)} E_n(2x^n) + c^{(2,1)} x^n E'_n(2x^n)]$$

Al insertar esta representación en la ecuación diferencial, reescribir las funciones de Mittag Leffler en forma de series de potencias y comparar los coeficientes de las resultantes series de

potencias en ambos lados de la ecuación resultante, encontramos que

$$c^{(2,0)} = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}^T \text{ y } c^{(2,1)} = \begin{pmatrix} 0 \\ 1 \\ -1 \end{pmatrix}^T$$

Observación 7.3

Ya sabemos que los coeficientes vectoriales de las funciones $E_n(\lambda x^n)$ en la representación de la solución son solo los vectores propios de Λ con respecto a el valor propio λ . Obviamente, estos vectores son independientes de n .

El ejemplo anterior indica que los coeficientes $c^{(\ell,u)}$ que surgen en la representación del Teorema 7.14 tampoco depende de n .

Por lo tanto, podemos calcular estos coeficientes observando el caso especial $n = 1$ del problema y utilizar los valores obtenidos de esta manera para $n \in (0,1)$ también.

La ventaja de este enfoque es que el engorroso cálculo, se puede evitar por completo la representación en series de potencias de la función de Mittag-Leffler y la correspondiente comparación de coeficientes. Más bien, basta con invoquemos la teoría clásica (Chen et al., 2007, pp. 4-6), que nos dice que los vectores $c^{(\ell,u)}$ pueden obtenerse de forma sencilla como vectores propios y múltiplos adecuados de la generalización de vectores propios de Λ .

El teorema 7.14 se cumple para cualquier matriz Λ con coeficientes complejos. Si dado el problema es real, es decir, si los coeficientes de Λ son reales, entonces a uno normalmente le interesa en una solución real, es decir, en una solución con vectores reales

$y^{(\ell)}$ y funciones básicas con valores reales en lugar de funciones con valores complejos $E_n(\lambda_\ell)$ y sus derivadas.

El enfoque a través del Teorema 7.2 y la Observación 7.1 nunca nos lleva fuera de los números reales, en tal caso y por lo tanto proporciona automáticamente una solución real, pero implica el engorroso cálculo de funciones de Mittag-Leffler de argumentos matriciales.

El método de los teoremas 7.13 y 7.14 evita esta complicación, pero no siempre produce una solución real directamente ya que los valores y vectores propios de una matriz real no necesitan ser reales.

Sin embargo, podemos construir una solución real a partir de la solución proporcionado por el Teorema 7.13 de la forma habitual que se emplea comúnmente para las ecuaciones de primer orden (Chen et al., 2007, pp. 80-81), es decir, considerando las partes real e imaginaria de las soluciones complejas, respectivamente.

Teorema 7.15

Si A es una matriz real y el conjunto de soluciones del sistema (7.14) construido en el teorema 7.14 contiene funciones complejas, entonces una base puramente real del conjunto de soluciones se puede obtener mediante la siguiente manipulación:

Para cada valor propio no real k veces $\lambda = a + ib$, elimine los vectores de solución $\pi^{(\ell)} \ell = 1, 2, 3, \dots, k$, asociados a este valor propio y los correspondientes al valor propio del conjunto de soluciones construidas en el teorema 7.13, e inserte los vectores $\tilde{\pi}^{(\ell)}$ y $\tilde{\pi}^{(\ell)} \ell = 1, 2, 3, \dots, k$ en su lugar, donde

$$\begin{aligned}\hat{\pi}_j^{(\ell)}(x) &= \sum_{u=0}^{\ell-1} \frac{1}{2} \operatorname{Re} c_j^{(\ell,u)} x^{un} D^u E_n((a+ib)x^n) + D^u E_n((a-ib)x^n) \\ &\quad + \sum_{u=0}^{\ell-1} \frac{1}{2i} \operatorname{Im} c_j^{(\ell,u)} x^{un} D^u E_n((a+ib)x^n) - D^u E_n((a-ib)x^n)\end{aligned}$$

Y

$$\begin{aligned}\tilde{\pi}_j^{(\ell)}(x) &= \sum_{u=0}^{\ell-1} \frac{1}{2} \operatorname{Im} c_j^{(\ell,u)} x^{un} D^u E_n((a+ib)x^n) + D^u E_n((a-ib)x^n) \\ &\quad + \sum_{u=0}^{\ell-1} \frac{1}{2i} \operatorname{Re} c_j^{(\ell,u)} x^{un} D^u E_n((a+ib)x^n) - D^u E_n((a-ib)x^n)\end{aligned}$$

Sólo ilustramos el Teorema 7.15 mediante un ejemplo y dejamos la demostración como ejercicio.

Aquí simplemente observamos que la representación en serie de potencias de las funciones de Mittag-Leffler implican que expresiones de la forma $D^u E_n(z) + D^u E_n(\bar{z})$ y $(D^u E_n(z) - D^u E_n(\bar{z}))/i$ que surgen en este teorema son de hecho reales.

Ejemplo 7.4.

$$\Lambda = \begin{pmatrix} -1 & 0 & 0 \\ 2 & 1 & -9 \\ 3 & 6 & 1 \end{pmatrix} \text{ es de la matriz.}$$

son -1 y $1 \mp \sqrt{54}i$ con vectores propios $(58, -31, 6)^T$ y $(2, \pm\sqrt{6}i, 2)^T$ respectivamente. Por tanto, el teorema 7.13 es aplicable y deducimos que el complejo general solución de la ecuación diferencial $D_{*0}^n y(x) = \Lambda y(x)$ con $0 < n < 1$ es

$$\begin{aligned}y(x) &= c_1 \begin{pmatrix} 58 \\ -31 \\ 6 \end{pmatrix} E_n(-x^n) + c_2 \begin{pmatrix} 0 \\ \sqrt{6}i \\ 2 \end{pmatrix} E_n((1 + \sqrt{54}i)x^n) \\ &\quad + c_3 \left[\begin{pmatrix} 0 \\ -\sqrt{6}i \\ 2 \end{pmatrix} E_n((1 - \sqrt{54}i)x^n) \right]\end{aligned}$$

Con $c_1, c_2, c_3 \in \mathbb{C}$ arbitrario. Una solución que utiliza funciones de valores reales se puede dar como

$$\begin{aligned}
 y(x) = c_1 \begin{pmatrix} 58 \\ -31 \\ 6 \end{pmatrix} E_n(-x^n) \\
 + \frac{1}{2} \left[c_2 \begin{pmatrix} 0 \\ 0 \\ 2 \end{pmatrix} + c_3 \begin{pmatrix} 0 \\ -\sqrt{6} \\ 0 \end{pmatrix} \right] \left[E_n((1 + \sqrt{54}i)x^n) \right. \\
 \left. - E_n((1 - \sqrt{54}i)x^n) \right] \\
 + \frac{1}{2i} \left[c_2 \begin{pmatrix} 0 \\ 0 \\ 2 \end{pmatrix} + c_3 \begin{pmatrix} 0 \\ \sqrt{6} \\ 0 \end{pmatrix} \right] \left[E_n((1 + \sqrt{54}i)x^n) \right. \\
 \left. - E_n((1 - \sqrt{54}i)x^n) \right]
 \end{aligned}$$

Con $c_1, c_2, c_3 \in \mathbb{R}$

Hemos completado así la teoría fundamental de los sistemas lineales homogéneos. de ecuaciones diferenciales fraccionarias con coeficientes constantes. En particular tenemos encontrados métodos para calcular todas sus soluciones. Por lo tanto, estamos en condiciones de manejar el correspondiente problema no homogéneo (7.12) de una manera muy sencilla invocando el enfoque de variación por constantes descrito en el teorema 7.2 y la observación 7.1.

Problemas de valores en la frontera para ecuaciones lineales

Dejemos brevemente el área de los problemas de valor inicial y dirijamos nuestra atención hacia problemas de valores en la frontera por un tiempo corto. Sin embargo, mantenemos nuestra restricción sobre la clase de ecuaciones diferenciales en consideración a las ecuaciones lineales.

Específicamente, consideramos el problema

$$\{D_{+0}^n y(x) = f(x)y(x) + g(x), \dots, (7.18a)$$

$$\{U_j[y] = c_j, (j = 1, 2, 3, \dots, \sigma), \dots, (7.18b)$$

Donde

$$U_j[y] := a_{j0}y(0) + a_{j1}y'(0) + b_{j0}y(T) + b_{j1}y'(T), \dots, (7.18c)$$

con algunos $\sigma \in \mathbb{N}$, $n > 0$, funciones dadas $f, g \in C[0, T]$ y parámetros dados a_{jk}, b_{jk} , y c_j (ver también el teorema 6.51).

Evidentemente buscaremos una solución en el intervalo $[0, T]$. Está claro que el conjunto de soluciones del correspondiente problema homogéneo, es decir, el caso especial de (7.18) donde $g \equiv 0$ y $c_j = 0$ para todo j , es un espacio lineal que lo denotaremos por D_{hom} .

Nuestro primer resultado fundamental entonces es el siguiente. Aquí y a continuación escribimos $m = [n]$ como de costumbre.

Teorema 7.16

Sean $y_1, \dots, y_m$ soluciones linealmente independientes de la solución ecuación diferencial homogénea asociada a (7.18^a), y sea r el rango de la matriz

$$M_{hom} := \begin{pmatrix} U_1[y] & \cdots & U_1[y_m] \\ \vdots & \cdots & \vdots \\ U_\sigma[y] & \cdots & U_\sigma[y_m] \end{pmatrix}$$

Entonces, el espacio de solución D_{hom} del problema de valores en la frontera homogéneo tiene la propiedad

$$\dim D_{hom} = m - r$$

Prueba. La solución general y de la ecuación diferencial homogénea tiene la forma $y = \sum_{k=1}^m \alpha_k y_k$, ($\alpha_k \in \mathbb{R}$). Las condiciones de contorno homogéneas conducen al sistema.

$$M_{hom} \begin{pmatrix} \alpha_1 \\ \vdots \\ \alpha_m \end{pmatrix}, \dots (7.19)$$

Por suposición, el rango de M_{hom} es r y, por tanto, el sistema tiene $m - r$ linealmente vectores de solución independientes

$$\alpha^{(j)} = (\alpha_1^{(j)}, \dots, \alpha_m^{(j)})^T, j = 1, 2, \dots, m - r$$

Ahora considere las funciones $\tilde{y}_j = \sum_{k=1}^m \alpha_k^{(j)} y_k, j = 1, 2, \dots, m - r$. Estas son $m-r$ soluciones del problema del valor límite homogéneo. Para ver que son linealmente independientes, procedemos de la siguiente manera. Asumir que

$$0 = \sum_{j=1}^m \beta_j \tilde{y}_j = \sum_{k=1}^m \left(\sum_{j=1}^m \alpha_k^{(j)} \beta_j \right) y_k$$

La independencia lineal de y_k implica que todos los coeficientes entre paréntesis deben desaparecer.

Podemos combinar estas m condiciones escalares en una condición en vector notación.

$$\sum_{j=1}^m \alpha_k^{(j)} \beta_j = 0$$

ya sabemos que los $\alpha^{(j)}$ son linealmente independientes, por lo que concluimos $\beta_j = 0$ para todo j . Esto completa la prueba de la independencia lineal de las funciones $\tilde{y}_j, j=1, 2, \dots, m-r$.

Queda por demostrar que cada solución y del valor límite homogéneo. El problema se puede expresar como una combinación lineal de las funciones $\tilde{y}_1, \tilde{y}_2, \dots, \tilde{y}_{m-r}$

Para este fin, sea y una solución del problema de valores en la frontera homogéneos. Dado que debe ser en particular una solución de la ecuación diferencial homogénea, debe tener una representación en la forma $y = \sum_{k=1}^m \alpha_k y_k$, y dado que

también satisface las condiciones de frontera, los coeficientes tienen la propiedad (7.19), y habíamos visto arriba que el vector $\alpha = (\alpha_1, \dots, \alpha_m)^T$ se puede escribir como una combinación lineal del $\alpha^{(i)}$, $j = 1, 2, \dots, m - r$, con coeficientes β_j , digamos. Pero esto implica

$$y(x) = \sum_{k=1}^m \alpha_k y_k = \sum_{k=1}^m \sum_{j=1}^m \alpha_k^{(j)} \beta_j y_k = \sum_{j=1}^m \beta_j \sum_{k=1}^m \alpha_k^{(j)} y_k = \sum_{j=1}^m \beta_j \tilde{y}_j$$

La cual es la propiedad requerida.

Para el problema de valores límite no homogéneos podemos enunciar lo siguiente resultado.

Teorema 7.17

La solución general del problema de valores en la frontera (7.18) tiene la siguiente forma $y = y_{hom} + y_{inhom}$ donde y_{hom} es la solución general del problema homogéneo asociado y es una solución particular del problema no homogéneo.

Prueba

En vista de la linealidad del problema, esta afirmación se puede probar mediante los métodos elementales habituales del álgebra lineal.

Tenga en cuenta que, no afirmamos que exista una solución del problema no homogéneo. Sólo describe la estructura del conjunto de soluciones si existe una solución particular. El siguiente teorema proporciona un criterio que nos permite determinar si este es el caso o no.

Teorema 7.18

Sean $y_1, \dots, y_m$ soluciones linealmente independientes de la solución homogénea ecuación diferencial asociada a (7.18a), y sea $\tilde{y}$ una solución particular de la propia ecuación diferencial no homogénea (7.18a). Además, denota

$$M_{inhom} := \begin{pmatrix} U_1[y] & \cdots & U_1[y_m] & c_1 - U_1[\tilde{y}] \\ \vdots & \cdots & \vdots & \vdots \\ U_\sigma[y] & \cdots & U_\sigma[y_m] & c_\sigma - U_\sigma[\tilde{y}] \end{pmatrix}$$

Entonces, el problema de valores límite no homogéneos (7.18) tiene solución si y sólo si la matriz M_{inhom} tiene el mismo rango que la matriz M_{hom} introducida en el teorema 7.16

Prueba

Cualquier solución y de la ecuación diferencial no homogénea (7.18a) puede ser expresado en la forma $y = \tilde{y} + \sum_{k=1}^m \alpha_k y_k$. Una solución al límite no homogéneo. El problema de valor (7.18) existe entonces si y sólo si podemos encontrar constantes $\alpha_1, \dots, \alpha_m$ tales que

$$c_j = U_j[y] = U_j[\tilde{y}] + \sum_{k=1}^m \alpha_k U_j[y_k], \quad (j = 1, 2, \dots, \sigma)$$

Claramente, este es el caso si y sólo si:

$$M_{hom} = \begin{pmatrix} \alpha_1 \\ \vdots \\ \alpha_m \end{pmatrix} = \begin{pmatrix} c_1 - U_1[\tilde{y}] \\ \vdots \\ c_\sigma - U_\sigma[\tilde{y}] \end{pmatrix}$$

Pero, es bien sabido por el álgebra lineal que esto equivale al requisito que la matriz de coeficientes M_{hom} de este sistema tiene el mismo rango que la matriz extendida matriz obtenida sumando el vector columna en el lado derecho de la ecuación a la matriz, es decir, la matriz M_{inhom} .

Obviamente, el caso de que la matriz M_{hom} sea una matriz cuadrada, es decir, el caso $\sigma = m$, es particularmente importante. En esta situación podemos dar la siguiente información.

Teorema 7.19

Considere el problema de valores en la frontera (7.18) sujeto a la condición $\sigma = m$. Definamos la M_{hom} como en el teorema 7.16.

Si $\det M_{hom} \neq 0$ entonces el problema de valores límite homogéneos asociado a (7.18) sólo tiene la solución trivial $y_{hom} \equiv 0$ y el problema no homogéneo (7.18) en sí tiene una solución única

Si $\det M_{hom} = 0$ entonces el problema de valores límite homogéneos asociado a (7.18) tiene soluciones no triviales

Prueba

Usando álgebra lineal elemental, esto se sigue inmediatamente del teorema anterior

Estabilidad de ecuaciones diferenciales fraccionarias

Ahora dejaremos la clase de problemas de valores en la frontera PVF y volveremos a centrarnos en nuestra atención a los problemas de valor inicial PVI. Habiendo obtenido una idea del comportamiento de ecuaciones lineales, ahora consideramos el caso general y dirigimos nuestra atención hacia la cuestión de la estabilidad de una ecuación diferencial fraccionaria EDF dada o, más frecuentemente, para un sistema de ecuaciones diferenciales fraccionarias SEDFs. En el caso clásico de ecuaciones de orden entero, esta es un área de investigación bien conocida e importante.

Comúnmente se discuten varias nociones de estabilidad; nos referimos a (Chen et al., 2007, cap. 5) y las referencias allí para una introducción agradablemente legible. Como se indica en este caso referencia, los problemas de estabilidad generalmente se investigan para ecuaciones de primer orden, es decir, para ecuaciones que requieren sólo una condición inicial para garantizar la unicidad de la solución.

Restringimos nuestra atención aquí a una clase de problemas que es lo más cercano a este caso como sea posible. Así, nuestro objeto de estudio en esta sección es la ecuación diferencial:

$$D_{+0}^n y(x) = f(x, y(x)) \text{ con } n \in (0,1), \dots (7.20)$$

Aquí y puede ser una función que mapea a R^N para un $N \in \mathbb{N}$ arbitrario; por supuesto f debe luego definirse en un subconjunto adecuado de $\mathbb{R}^{N+1}$.

En el caso clásico $n = 1$ uno sería permitir puntos iniciales arbitrarios (Chen et al., 2007, cap. 5); El teorema 6.17 nos permite hacer lo mismo en el caso fraccionario.

Cuando se habla de estabilidad, uno está interesado en el comportamiento de las soluciones, de (7.20) para $x \rightarrow \infty$.

Por lo tanto, solo consideraremos problemas cuyas soluciones y existen en $[0, \infty)$. Además, se requieren algunas suposiciones adicionales que vamos a imponer a lo largo de esta sección.

La primera de estas suposiciones es que f está definida en un conjunto

$$G := [0, \infty) \times \{w \in R^N : \|w\| < W\}$$

con algún $0 < W \leq \infty$. La norma en esta definición de G puede ser una norma arbitraria en R^N .

Nuestra segunda suposición es que f es continua en su dominio de definición y que satisface una condición de Lipschitz allá.

Esto afirma que el problema de valor inicial PVI que consta de (7.20) y el valor inicial la condición $f(0) = y_0$ tiene solución única en el intervalo $[0, b)$ con algún $b \leq \infty$ si $\|y_0\| < W$.

Y finalmente suponemos que:

$$f(x, 0) = 0, \text{ para todo } x \geq 0$$

Esta condición implica que la función $y(x) = 0$ es una solución de (7.20). Bajo estas hipótesis podemos formular nuestros conceptos principales.

Definición 7.2

La solución $y(x) = 0$ de la ecuación diferencial (7.20), sujeta a los supuestos mencionados anteriormente, se llama estable si, para cualquier $\varepsilon > 0$ existe algún $\delta > 0$ tal que la solución del problema de valor inicial PVI que consiste en la ecuación diferencial (7.20) y la condición inicial $y(0) = y_0$ satisface $y(x) < \varepsilon$ para todo $x \geq 0$ siempre que $\|y_0\| < \delta$.

La solución $y(x) = 0$ de la ecuación diferencial (7.20), sujeta a los supuestos mencionados anteriormente, se dice asintóticamente estable si es estable y existe algún $\gamma > 0$ tal que $\lim_{x \rightarrow \infty} y(x) = 0$ siempre que $\|y_0\| < \gamma$.

En el ítem 6.3 habíamos visto que, bajo la continuidad habitual y Lipschitz supuestos sobre f , la solución de una ecuación diferencial fraccionaria no cambia mucho en algún intervalo finito si perturbamos los valores iniciales en una pequeña magnitud.

La esencia de la noción de estabilidad es la extensión de esta idea a límites ilimitados de intervalos:

La solución trivial es estable si un pequeño cambio en el valor inicial conduce a un pequeño cambio de la solución sobre la media línea positiva completa. Obviamente esto es mucho más fuerte que la dependencia continua de los datos proporcionados discutidos en el ítem 6.3.

La estabilidad asintótica es aún más fuerte ya que requiere la solución del problema perturbado no sólo para permanecer cerca de la solución original sino también para converger hacia este último.

Observación 7.4

En la Definición 7.2 sólo hemos analizado las propiedades de elementos idénticamente solución evanescente de una ecuación diferencial fraccionaria que satisface la condición de (7.21).

Podemos transferir estos conceptos y los resultados siguientes al comportamiento de soluciones arbitrarias de ecuaciones que pueden satisfacer o no (7.21) mediante un procedimiento similar a la utilizada en la demostración del teorema 7.15:

Una solución y de la ecuación diferencial $D_{0+}^{\alpha}y(x) = g(x, y(x))$ se dice que es asintóticamente estable si y sólo si la solución cero de $D_{0+}^{\alpha}z(x) = f(x, z(x))$ con $f(x, z) := g(x, z + y(x)) - g(x, y(x))$ es asintóticamente estable.

Comenzamos nuestro análisis observando un caso especial muy simple, la ecuación diferencial homogéneo lineal con coeficientes constantes (Mainardi, 2012).

Teorema 7.20

Considere el sistema de ecuaciones diferenciales fraccionarias N -dimensionales $D_{+0}^n y(x) = \Lambda y(x)$ donde Λ es una matriz $N \times N$ constante arbitraria.

La solución $y(x) = 0$ del sistema es asintóticamente estable si y sólo si todos los valores propios $\lambda_j, (j = 1, 2, 3, \dots, N)$ de Λ satisfacen $|\arg \lambda_j| > \frac{n\pi}{2}$.

La solución $y(x) = 0$ del sistema es estable si y sólo si los valores propios satisfacen $|\arg \lambda_j| \geq \frac{n\pi}{2}$ y todos los valores propios con $|\arg \lambda_j| = \frac{n\pi}{2}$ tienen una multiplicidad geométrica que coincide con su multiplicidad algebraica (es decir, un valor propio que es un ℓ – *dobles* ceros del polinomio característico tiene linealmente independientes vectores propios).

Observe que en el caso límite $n \rightarrow 1$ recuperamos el conocido resultado clásico (Chen et al., 2007, p. 94), que los valores propios deben tener partes reales negativas en el caso (a) y no positivas partes reales y un conjunto completo de vectores propios si las partes reales son cero en el caso (b).

Prueba

Los teoremas 7.13 y 7.14 nos dan información sobre la forma precisa de todas las soluciones de la ecuación diferencial considerada.

El hecho de que estas soluciones se comportan de la manera requerida, como se desprende de los teoremas 4.4 y 4.6.

Este resultado nos permite investigar las propiedades de estabilidad de los problemas discutidos en los Ejemplos 7.2, 7.3 y 7.4.

Ejemplo

La solución $y \equiv 0$ del ejemplo 7.2 es inestable porque la matriz de coeficientes tiene el valor propio $\lambda_1 = 1$ que es real y estrictamente positivo, es decir, tiene $\arg \lambda_1 = 0$. Esto corresponde al hecho de que el componente asociado de la solución general, a saber, la función $(1,1)^T E_n(x^n)$, crece sin límite como $x \rightarrow \infty$.

De manera similar, la matriz de coeficientes del ejemplo 7.3 tiene un valor propio real y positivo 2, por lo que también observamos inestabilidad en este ejemplo.

La matriz de coeficientes del ejemplo 7.2 tiene tres valores propios:

$$\begin{aligned}\lambda_1 &= -1 \text{ (con } \arg \lambda_1 = \pi, \lambda_2 = 1 + \sqrt{54} i \text{ (con } \arg \lambda_2 = \arccos(\frac{1}{\sqrt{55}}) \approx 1.43 \text{) y} \\ \lambda_2 &= \bar{\lambda}_2 = 1 - \sqrt{54} i \text{ (con } \arg \lambda_3 = -\arg \lambda_2 = -\arccos(\frac{1}{\sqrt{55}})\end{aligned}$$

Por lo tanto, tenemos estabilidad asintótica si y sólo si $n < \frac{2|\arg \lambda_2|}{\pi} = 2 \arccos(\frac{1}{\sqrt{55}})/\pi \approx 0.9139$ y estabilidad, ya que todos los valores propios son simples si y sólo si $n \leq \frac{2|\arg \lambda_2|}{\pi}$.

Observación 7.5

Del caso (c) concluimos que, las propiedades de estabilidad de la solución cero de una ecuación diferencial fraccionaria EDF puede depender del orden n de la ecuación.

Específicamente, como se desprende claramente del teorema 7.20, en el caso de un sistema lineal homogéneo con coeficientes constantes podemos decir que existe un valor umbral n^* , digamos, tal que el sistema es asintóticamente estable si $n < n^*$ y inestable si $n > n^*$.

En otras palabras, las propiedades de estabilidad se pueden mejorar reduciendo el orden n del operador diferencial.

Una afirmación análoga se aplica a otros Tipos no lineales de ecuaciones diferenciales.

Esta es una observación común en la teoría de los sistemas dinámicos fraccionarios (Linz, 1985), donde los sistemas tienden a exhibir caóticos, comportamiento si el orden de los operadores diferenciales es mayor que el valor umbral y permanecen estables si el orden es menor que n^* (Lu, 2006).

En el caso no fraccionario, es decir, para $n = 1$, se conocen resultados bastante profundos en el caso de ecuaciones lineales homogéneas con coeficientes no constantes.

La mayoría de las pruebas de estos resultados se basan en el hecho de que las expresiones explícitas para las soluciones de la que se conocen las ecuaciones diferenciales correspondientes.

En el caso fraccionario, tal explícita las expresiones no parecen estar disponibles.

Por lo tanto, queda por hacer una transferencia de estos resultados, un problema abierto.

Sin embargo, existen otros métodos que se pueden utilizar para obtener resultados en el comportamiento a largo plazo de soluciones de ecuaciones diferenciales fraccionarias.

Un posible enfoque se indica en el Teorema 4.6 de Kiryakova (2010b). Transferido a nuestro entorno, el resultado lee lo siguiente.

Teorema 7.21

Considere el problema del valor inicial.

$$\begin{cases} D_{*0}^n y(x) = f(x, y(x)), \\ y(0) = y_0 > 0 \end{cases}$$

donde $0 < n < 1$ y $f: [0, \infty) \times [0, y_0] \rightarrow (-\infty, 0]$ es continua y satisface una condición de Lipschitz con respecto a la segunda variable. Además, supongamos que $f(x, 0) = 0$ para todo x . Entonces, existe la solución única y de este problema de valor inicial en $[0, \infty)$ y satisface

$$0 < y(x) \leq y_0, \text{ para todo } x \geq 0$$

Prueba

Reescribimos el problema de valor inicial en la forma de Volterra correspondiente,

$$y(x) = y_0 + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt$$

Es decir:

$$f(t, z) = \begin{cases} f(t, y_0), & \text{para } z > y_0 \\ 0, & \text{para } z < 0 \end{cases}$$

extendiendo así el dominio de definición de f a $[0, \infty) \times \mathbb{R}$. Obviamente, esto se extendió. La función f sigue siendo continua y cumple una condición de Lipschitz con respecto a su segunda variable.

Así, según el Corolario 6.9, el problema de valor inicial PVI con este método extendido la función f tiene una solución única y.

Ahora necesitamos demostrar que esta solución realmente satisface la desigualdad $0 \leq y(x) \leq y_0$ para todo x .

Esto implica entonces que $(x, y(x))$ es siempre en el dominio original no extendido, de definición de f del cual concluya que y también es la solución única del problema con valores iniciales originales.

Para demostrar las desigualdades requeridas procedemos de la siguiente manera:

Nuestra primera observación es que, dado que $y(0) = y_0 > 0$ y es continua, existe algún $x_0 > 0$ tal que $y(x) > 0$ para $x \in [0, x_0]$.

Ahora supongamos que y tiene un cambio de signo. Entonces, podemos encontrar algunos x_1 y x_2 tales que $0 < x_1 < x_2$ y.

$$y(x) \begin{cases} > 0, \text{ para } 0 \leq x < x_1 \\ = 0, \text{ para } x = x_1 \\ < 0, \text{ para } x_1 < x \leq x_2 \end{cases}$$

Así, por definición de f y su extensión

$$f(x, y(x)) = \begin{cases} \leq 0, \text{ para } 0 \leq x < x_1, \\ = 0, \text{ para } x_1 \leq x \leq x_2 \end{cases}$$

Al introducir esta observación en la ecuación de Volterra anterior, encontramos que

$$\begin{aligned} 0 > y(x_2) &= y_0 + \frac{1}{\Gamma(n)} \int_0^{x_2} (x_2 - t)^{n-1} f(t, y(t)) dt \\ &= y_0 + \frac{1}{\Gamma(n)} \int_0^{x_1} (x_2 - t)^{n-1} f(t, y(t)) dt + \frac{1}{\Gamma(n)} \int_{x_1}^{x_2} (x_2 - t)^{n-1} f(t, y(t)) dt \\ &= y_0 + \frac{1}{\Gamma(n)} \int_0^{x_1} (x_2 - t)^{n-1} f(t, y(t)) dt \\ &\geq y_0 + \frac{1}{\Gamma(n)} \int_0^{x_1} (x_1 - t)^{n-1} f(t, y(t)) dt = y(x_1) = 0 \end{aligned}$$

que es la contradicción requerida que muestra que y no puede tener un cambio de signo. Por tanto, $y(x) \geq 0$ para todo $x \geq 0$.

Usando esta observación, entonces concluimos de la forma Volterra de nuestro problema de valor inicial, utilizando el hecho de que $f(t,y) \leq 0$ para todo t e y y por definición, que $y(x) \leq y_0$.

Finalmente, la desigualdad estricta $y(x) > 0$ se desprende del Corolario 7.16

Bajo supuestos adicionales también podemos decir algo sobre el límite de $y(x)$ como $x \rightarrow \infty$.

Teorema 7.22

Suponga las hipótesis del teorema 7.21. Además, supongamos que para

todo $a > 0$ y todas las funciones continuas $Y : [0, \infty) \rightarrow [a, y_0]$ se cumple

$$\lim_{x \rightarrow \infty} J_0^n[f(., Y(.))] (x) = -\infty$$

Entonces

$$\lim_{x \rightarrow \infty} y(x) = 0$$

Si el límite existe.

Prueba

En vista del teorema 7.21, está claro que $0 \leq y(x) \leq y_0$ para todo x , y por tanto $\lim_{x \rightarrow \infty} y(x) \geq 0$ si el límite existe.

Probaremos indirectamente que el límite debe ser cero. Para este fin supongamos que $\lim_{x \rightarrow \infty} y(x) = \lambda > 0$. Entonces podemos encontrar un x_0 tal que $y(x) \geq \lambda/2$ para todo $x \geq x_0$. Ahora definimos

$$Y(x) = \begin{cases} y(x), & \text{para } x > x_0, \\ y(x_0), & \text{para } x \leq x_0 \end{cases}$$

Obviamente, Y es una función continua que satisface $Y(x) \geq \lambda/2$ para todo $x \geq 0$, y de En la forma de Volterra del problema de valor inicial bajo consideración vemos que

$$y(x) = y_0 + J_0^n [f(., y(.))](x) = J_0^n [f(., y(.)) - f(., Y(.))](x) + J_0^n [f(., Y(.))](x)$$

En el lado derecho de esta ecuación, el primer término es una constante y el último El término tiende a $-\infty$ cuando $x \rightarrow \infty$ por supuesto.

A continuación, mostraremos que el segundo El término permanece acotado a medida que x crece, de esta observación concluimos entonces que $y(x) \rightarrow -\infty$ como $x \rightarrow \infty$ lo que contradice el hecho de que $y(x) \geq 0$ que se sigue de Teorema 7.21.

Por tanto, nuestra suposición $\lim_{x \rightarrow \infty} y(x) > 0$ debe ser falsa. Para demostrar la acotación del segundo término en la ecuación anterior para $x \rightarrow \infty$, nosotros escribimos

$$J_0^n [f(., y(.)) - f(., Y(.))](x) = \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} [f(t, y(t)) - f(t, Y(t))] dt$$

y observe que el término entre paréntesis desaparece para por definición de Y . Por lo tanto,

$$\begin{aligned} J_0^n [f(., y(.)) - f(., Y(.))](x) &= \frac{1}{\Gamma(n)} \left| \int_0^x (x-t)^{n-1} [f(t, y(t)) - f(t, Y(t))] dt \right| \\ &\leq \frac{2}{\Gamma(n)} \sup_{t \in [0, x_0], z \in [0, y_0]} |f(t, z)| \int_0^{x_0} (x-t)^{n-1} dt \\ &= \frac{2}{\Gamma(n)} \sup_{t \in [0, x_0], z \in [0, y_0]} |f(t, z)| (x^n - (x_0 - t)^{n-1}) \end{aligned}$$

Se ve fácilmente que la expresión del lado derecho es una función positiva y monótonamente decreciente de x y, por tanto:

Observación 7.6

Una situación particular donde los supuestos del teorema 7.21 son satisfecho es el caso de una ecuación diferencial fraccionaria lineal homogénea, es decir, el caso $f(x,z) = -\mu(x)z$, con alguna función μ continua y acotada no negativa.

Si, además, $\mu(x) \geq \mu_0 > 0$ para todo $x \geq 0$ con una constante adecuada μ_0 , entonces Teorema 7.22 también es aplicable. Una condición suficiente para que esto se cumpla es que sea un constante positiva.

Estas observaciones nos sitúan ahora, como se indica en la Sección. 7.1, en condiciones de demostrar Teorema 7.3 sin utilizar técnicas de transformada de Laplace.

Prueba (del teorema 7.3)

La función considerada en el teorema 7.3 es la solución única de la ecuación de relajación fraccionaria homogénea $D_{+0}^{\alpha} u_0(x) = -\mu u_0(x)$ con algún $\mu > 0$ sujeto a la condición inicial $u_0(x) = 1$.

Esto corresponde al caso $f(t,z) = -\mu z$ de los teoremas 7.21 y 7.22. Se ve fácilmente que Se cumplen los supuestos de ambos teoremas. Obtenemos así $0 < u_0(x) \leq 1$ para todo $x \in [0, \infty)$ por el Teorema 7.21 y $\lim_{x \rightarrow \infty} u_0(x) = 0$ por el Teorema 7.22.

De la representación explícita de en términos de la función Mittag-Leffler y la expansión en serie de potencias de este último encontramos que u_0 tiene infinitos derivadas continuas en el intervalo abierto $(0, \infty)$ y que $\lim_{x \rightarrow 0_+} u'_0(x) = \infty$. Nosotros Luego podemos diferenciar la forma de Volterra del problema de valor inicial para $x > 0$, cuyos rendimientos

$$u'_0(x) = -\frac{u}{\Gamma(n)} x^{n-1} - \frac{u}{\Gamma(n)} \int_0^x (x-t)^{n-1} u'_0(t) dt$$

Esta es una ecuación de Volterra para u'_0 . Hemos visto arriba que $u'_0(x) < 0$ para $x \in (0, \varepsilon)$ con algún $\varepsilon > 0$. Con este conocimiento podemos entonces manejar la ecuación de Volterra para u'_0 de la misma manera que lo habíamos hecho con la ecuación para u_0 y en la demostración del teorema 7.21 para mostrar que u'_0 no tiene cambio de signo. Una aplicación repetida de este paso (diferenciación de la ecuación de Volterra y demostración de que la solución no cambiar su signo) entonces da las propiedades de las derivadas de u_0 establecidas en el inciso (a) de Teorema 6.3.

Una consecuencia particular de los resultados que acabamos de demostrar es el hecho de que es una función monótonamente decreciente que converge a cero a medida que su argumento crece a ∞ . Por lo tanto, vale la pena intentar modelar el comportamiento asintótico de $u_0(x)$ en este caso límite a través del enfoque.

$$u_0(x) = c x^{-\beta} + \sigma(x^{-\beta}), \text{ como } x \rightarrow \infty$$

con ciertas constantes $\beta > 0$ y $c \in \mathbb{R}$. Podemos insertar esta relación en el modelo de forma de Volterra del problema de valor inicial para obtener

$$\begin{aligned} c x^{-\beta} + \sigma(x^{-\beta}) &= u_0(x) = 1 - \frac{u}{\Gamma(n)} \int_0^x (x-t)^{n-1} u_0(t) dt \\ &= 1 - \frac{cu}{\Gamma(n)} \int_0^x (x-t)^{n-1} t^{-\beta} dt + \int_0^x (x-t)^{n-1} \sigma(t^{-\beta}) dt \\ &= 1 - \frac{cu}{\Gamma(n)} x^{n-\beta} \frac{\Gamma(n)\Gamma(1-\beta)}{\Gamma(n+1-\beta)} + \sigma(x^{n-\beta}) \end{aligned}$$

Si no tenemos un término de orden en el lado izquierdo, no debemos tener tal término en el lado derecho tampoco. Por lo tanto, la constante 1 y el siguiente término más alto en el lado dere-

cho, es decir los términos que involucran $x^{n-\beta}$, deben cancelarse entre sí. Esto produce

$$\beta = n \text{ y } c = \frac{1}{\mu \Gamma(1-n)}$$

Lo que demuestra parte (o) del teorema.

Ecuaciones singulares

Todos los problemas considerados hasta ahora han sido regulares en el sentido de que la función f en el lado derecho de la ecuación diferencial (7.1a) ha sido al menos continua y, en la mayoría de los casos, incluso diferenciable un cierto número de veces.

En algunas aplicaciones, sin embargo, uno encuentra ecuaciones donde este no es el caso. Por lo tanto, nosotros. Concluiremos ahora este capítulo con una sección dedicada a algunos resultados sobre problemas tan singulares.

No proporcionaremos un análisis completamente general; bastante restringiremos nuestra atención a algunos casos especiales particularmente importantes que, sin embargo, demuestran una rica variedad de fenómenos que pueden encontrarse en la investigación de ecuaciones diferenciales fraccionarias singulares.

Un resultado importante y de uso frecuente en la teoría de ecuaciones sin singularidades fue la equivalencia entre el problema de valor inicial PVI (7.1) y la correspondiente formulación de la ecuación integral de Volterra (7.2) que habíamos establecido en el Lema 6.2. Nuestra primera observación en el contexto de esta sección es que en la presencia de singularidades podemos perder esta equivalencia. Esto se puede ver desde el siguiente ejemplo.

Ejemplo

Para $1 < n < 2$, considere la ecuación integral

$$y(x) = 1 + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} \frac{1}{y(t)-1} dt$$

que corresponde formalmente al problema de valor inicial

$$\text{PVI} \quad \begin{cases} D_0^n y(x) = \frac{1}{y(x)-1} \\ y(0) = 1 \\ y'(0) = 0 \end{cases}, \dots, (7.23)$$

O también,

$$D_0^n y(x) = \frac{1}{y(x)-1}, y(0) = 1, y'(0) = 0, \dots (7.23)$$

Un cálculo explícito muestra que (7.22) se resuelve mediante las funciones

$$y(x) = 1 \pm \frac{\sqrt{\Gamma(1-n/2)}}{\sqrt{\Gamma(1+n/2)}} x^{n/2}$$

Así, observamos que la formulación de la ecuación integral tiene más de una solución continua.

Sin embargo, también vemos que ninguna de estas soluciones es en $x = 0$, y por lo tanto estas soluciones no satisfacen la segunda condición inicial.

Así, hemos encontrado funciones que resuelven la ecuación integral (7.22) pero no el correspondiente problema de valor inicial PVI (7.23).

Un ejemplo prototípico diferente de una ecuación diferencial fraccionaria singular que tiene importantes aplicaciones en la práctica ha sido introducido por Joulin (Hilfer, 2000), quien lo utiliza modelar la propagación de una llama en el contexto de un modelo termodifusivo con altas energías de activación utilizando una mezcla gaseosa con química simple $A \rightarrow B$.

En estas circunstancias, muestra que el radio de la llama en el tiempo x está dado por $y(x)$, donde la función y es la solución del problema de valor inicial

$$y(x)D_{*0}^{\frac{1}{2}}y(x) = y(x)\ln y(x) + Eq(x), y(0) = 0, \dots, (7.24)$$

O

$$\begin{cases} y(x)D_{*0}^{1/2} y(x) = y(x)\ln y(x) + Eq(x) \\ y(0) = 0 \end{cases}, \dots, (7.24)$$

Aquí la función q describe una fuente de energía puntual dependiente del tiempo y , por lo tanto, se supone que es no negativo, continuo e integrable en $\mathbb{R}_+$ y E representa la intensidad de esta fuente de calor tal que $E \cdot \|q\|_{L_1(0,\infty)}$ es la cantidad total de energía introducido en el sistema.

A veces se supone que q está normalizado de modo que $\|q\|_{L_1(0,\infty)} = 1$; consideramos conveniente no imponer este requisito.

Para nuestros propósitos será útil enunciar la ecuación diferencial en forma explícita, es decir, resolverlo para $D_{*0}^{1/2} y$ y obtener el valor inicial del problema en la representación

$$D_{*0}^{\frac{1}{2}}y(x) = f(x, y(x)), \text{ con } f(x, w) = \ln w + E \frac{q(x)}{w}, y(0) = 0, \dots (7.25)$$

que es equivalente a la forma original (7.24).

Es inmediatamente evidente que la función f no es continua en ninguna vecindad del punto inicial $(0,0)$.

El modelo descrito por (7.24), que puede justificarse de manera matemáticamente rigurosa (Kilbas & Trujillo, 2002), tiene asociadas algunas cuestiones importantes bastante naturales.

Aparte desde el más obvio para una solución, exacta o aproximada, para un problema específico elección de los parámetros E y q , a menudo uno está muy interesado en la calidad cualitativa comportamiento de la solución. Analíticamente es posible demostrar el siguiente resultado, eso indica que tenemos que lidiar con un fenómeno de bifurcación (Audounet et al., 1998), Teorema 0.2:

Teorema 7.23

Supongamos que existe algún $x_0 > 0$ tal que $q(x) > 0$ para $x \in (0, x_0)$ y $q(x) = 0$, además, Entonces, el problema con valores iniciales PVI (7.24) tiene una solución continua única y . Además, existe un valor crítico $E_{crit}q(x)$ tal que:

- si $E > E_{crit}q(x)$ entonces y se define en $[0, \infty)$ y $\lim_{x \rightarrow \infty} y(x) = \infty$
- si $E = E_{crit}q(x)$ entonces y se define en $[0, \infty)$ y $\lim_{x \rightarrow \infty} y(x) = 1$
- si $E < E_{crit}q(x)$ entonces existe algún $x_{max} > x_0$ finito tal que y esté definido en $[0, x_{max}]$ y 6

La demostración de este teorema se puede encontrar en Audounet et al. (1998). Se basa en la observación que la solución del problema con valores iniciales fraccionarios puede escribirse como la solución de una ecuación diferencial parcial parabólica que tiene la forma clásica para tales ecuaciones y que, en particular, no contenga derivadas fraccionarias. En las propiedades requeridas se derivan entonces de la teoría estándar de ecuaciones diferenciales parciales parabólicas.

Además de las propiedades descritas en el teorema 7.23 anterior, existen otros resultados sobre los aspectos analíticos de (7.25) son de interés.

La primera observación trata del comportamiento asintótico de la solución y de nuestro problema (7.25) cuando $x \rightarrow 0$. Está tomado de proposición 1.2.

Teorema 7.24

Supongamos que $q(x) \geq 0$ para $x \geq 0$ y que $q(x) = q_0 x^\beta (1 + \sigma(1))$ como $x \rightarrow 0$ con algún $\beta \in [0, 1/2)$. Entonces, cuando $x \rightarrow 0$, la solución y del valor inicial del problema (7.24) se comporta como,

$$y(x) = \rho_\beta x^{\frac{1}{4+\beta/2}} (1 + \sigma(1))$$

Donde,

$$\rho_\beta = \left(E q_0 \frac{\Gamma(\frac{3}{4} - \frac{\beta}{2})}{\Gamma(\frac{5}{4} - \frac{\beta}{2})} \right)^{1/2}$$

El caso especial más simple de este resultado, $\beta = 0$, ya revela que la ecuación asintótica.

El comportamiento de la solución es significativamente diferente del comportamiento que esperamos de la solución de una ecuación no singular como se describe en los teoremas 6.38 y 6.39. Observe que este resultado con un factor omitido accidentalmente.

La importancia de la condición $\beta < 1/2$ en el teorema 7.24 es explicado por el siguiente resultado tomado de Ahmad & El-Khazali (2007).

Teorema 7.25

Supongamos que $q(x) = q_0 x^\beta (1 + \sigma(1))$ cuando $x \rightarrow 0$ con algún $\beta \geq 1/2$. Entonces, cuando $x \rightarrow 0$, la solución y del problema de valores iniciales (8.24) se comporta como:

$$y(x) = \frac{1}{\sqrt{\pi}} x^{1/2} |\ln x|. (1 + \sigma(1))$$

Entonces, el comportamiento asintótico de la solución exacta cerca del origen cambia a medida que el parámetro β cruza el valor $1/2$.

Se pueden derivar resultados similares en el caso de que el soporte de q sea ilimitado (Ahmad & El-Khazali, 2007). Puede encontrar más información sobre preguntas relacionadas en Mainardi (2010). De desde el punto de vista de las aplicaciones, sin embargo, la situación analizada en el teorema 7.23 es con diferencia el más relevante.

Dicho explícitamente, dice que la llama se apagará en un tiempo finito, tiempo si la energía agregada al sistema es menor que el nivel crítico E_{crit} , y arderá persistentemente si la energía está por encima de E_{crit} .

Por consideraciones de seguridad, es muy importante averiguar el valor de o al menos sus límites inferiores, si uno está interesado en mantener el incendio bajo control.

Por otra parte, a veces uno está interesado en construir una llama que arda permanentemente, y entonces uno necesita saber E_{crit} , o al menos límites superiores para él, para encontrar un proceso eficiente que utilice como la menor cantidad de energía posible.

Por lo tanto, ciertamente está justificado investigar el problema valor inicial PVI (7.24) a fondo, utilizando enfoques tanto

analíticos como numéricos. Esto es aún más enfatizado (Yuan & Agrawal, 2002, p. 570), que las ideas de Joulin pueden ser trasladado a una clase mucho más amplia de experimentos, por lo que uno debería esperar que un procedimiento exitoso para (7.24) también podrá manejar modelos de esta gran clase también.

Se ha descrito un estudio de este tipo de métodos numéricos para este problema, por ejemplo, en Diethelm (2001).

Sin embargo, desde el punto de vista analítico, el conocimiento actual parece ser muy limitado.

08

Capítulo

*ECUACIONES DIFERENCIALES
FRACCIONARIA DE CAPUTO DE VARIAS
ETAPAS*

Escucha a tu niño interno que te enseñará, a despreciarse las quejas de los amigos, aun cuando sucediese que se quejen sin razón, sino que, al contrario, se debe satisfacerlos y tratar de reducirlos a una buena armonía acostumbrada.

Burseca

Hasta este punto, solo hemos considerado las llamadas ecuaciones de un solo término, es decir ecuaciones en las que sólo interviene un operador diferencial. Aunque en ciertos casos necesitamos resolver ecuaciones que contengan más de un operador diferencial. Un ejemplo clásico es la llamada ecuación de Bagley-Torvik “EBT” (Torvik & Bagley, 1984).

$$\begin{cases} AD_{*0}^2 y(x) + BD_{*0}^{\frac{3}{2}} y(x) + C y(x) = f(x), \\ y(0) = y_0 \end{cases}$$

o

$$AD_{*0}^2 y(x) + BD_{*0}^{\frac{3}{2}} y(x) + C y(x) = f(x), y(0) = y_0$$

donde A, B y C son ciertas constantes y f es una función dada. esta ecuación surge, por ejemplo, en el modelado del movimiento de una placa rígida sumergida en un fluido newtoniano.

Originalmente fue propuesto en y se discute a fondo, por ejemplo, en , desde un punto de vista numérico, en. Otro ejemplo para una aplicación de ecuaciones con más de una derivada fraccionaria es la ecuación Basset

$$D^1 y(x) + b D_*^n y(x) + cy(x) = f(x), y(0) = y_0$$

donde $0 < n < 1$.

Esta ecuación se usa con mayor frecuencia, pero no exclusivamente, con $n = 1/2$. Describe las fuerzas que ocurren cuando un objeto esférico se hunde en un fluido viscoso incompresible, relativamente menos denso (Bagley & Torvik, 2000); (Podlubny, 1994).

Un tipo de ecuaciones más general que incluye los ejemplos mencionados anteriormente, sería

$$g(x, y(x), D_{*0}^{n_1} y(x), D_{*0}^{n_2} y(x), D_{*0}^{n_3} y(x), \dots, D_{*0}^{n_k} y(x)) = 0$$

con $0 < n_1 < n_2 < \dots < n_k$ y una determinada función g .

Una ecuación de este tipo denominase ecuación diferencial fraccionaria de términos múltiples o, más precisamente, de términos k .

Dado que para una ecuación tan general casi no se conocen resultados, restringimos nuestra atención a ciertos casos especiales, en su mayor parte veremos ecuaciones explícitas, es decir, ecuaciones que pueden resolverse para la derivada de orden más alto $D_{*0}^{n_k} y(x)$ y nosotros Investigamos las propiedades más fundamentales.

Primero presentamos un concepto que será útil a lo largo de la investigación de ecuaciones multitérminos y que es aplicable a la clase completamente general de ecuaciones multitérminos, ecuaciones mencionadas anteriormente, por lo que no es necesario especializarse en un subconjunto de estas ecuaciones todavía.

Definición 8.1

La ecuación diferencial fraccionaria

$$g(x, y(x), D_{*0}^{n_1} y(x), D_{*0}^{n_2} y(x), D_{*0}^{n_3} y(x), \dots, D_{*0}^{n_k} y(x)) = 0$$

con $0 < n_1 < n_2 < \dots < n_k$ y una cierta función g se llama proporcional si los números $n_1, n_2, \dots, n_k$ son conmensurables, es decir, si los cocientes n_μ / n_ν son números racionales para todo $\mu, \nu \in \{1, 2, 3, \dots, k\}$

Observación 8.1

Algunos autores ver, por ejemplo, Nonnenmacher & Metzler (1995), usan esta terminología de una manera diferente. sentido; lo aplican a sistemas de ecuaciones diferenciales fraccionarias de la forma

$$D_{*0}^{n_k} y(x) = f_k(x, y_1(x), y_2(x), \dots, y_N(x)), k = 1, 2, 3, \dots, N$$

Y es decir que tal sistema es proporcional si $n_1 = n_2 = \dots = n_N$. Nos abstenemos de la noción de este uso y nos atenemos a nuestra Definición 8.1 porque esta última está más en manteniendo el uso tradicional del concepto que es común en la teoría de números.

La ecuación de Bagley-Torvik es un ejemplo de ecuación diferencial fraccionaria conmensurada en nuestro sentido; un contraejemplo es $D_{*0}^1 y(x) + D_{*0}^{1/\pi} y(x) = 0$

Si la ecuación de Basset es proporcional o no depende del valor preciso de n : La ecuación es proporcional si y sólo si n es un número racional.

La conmensuración es importante porque nos permite seguir un enfoque que es bien conocido por la teoría de las ecuaciones diferenciales de orden entero: podemos transformar la ecuación dada en un sistema de ecuaciones diferenciales que involucren sólo un operador diferencial.

De esta manera podemos invocar la teoría de ecuaciones de un término como descrito en el capítulo 6 con el fin de investigar cuestiones de existencia y unicidad o derivar otras propiedades de las soluciones.

Además, podemos utilizar métodos numéricos. para ecuaciones de un término como se describe, por ejemplo, en el artículo de la encuesta (Diethelm & Ford, 2004) (Diethelm & Ford, 2001), para aproximar las soluciones.

Por supuesto, es cierto que el supuesto de conmensurabilidad restringe significativamente la clase de ecuaciones diferenciales que se pueden considerar, pero podemos decir que muchas de las ecuaciones derivadas en aplicaciones físicas o de ingeniería tienen esta propiedad.

Además de la ecuación de Bagley-Torvik, nos referimos al modelo de materiales viscoelásticos de Koeller (1986), por ejemplo, copolímeros, descritos en Kilbas & Trujillo (2001), ec. (3.9) y (3.10). Específicamente Llamamos la atención del lector sobre la motivación física de la conmensuración supuesto dado por Bagley y Calico en la segunda sección de su artículo (1991); (Ahmad & El-Khazali, 2007). Por cierto, Rossikhin y Shitikova (2007); (Samko et al., s. f.), han señalado recientemente que la teoría de Koeller (1986).

La ecuación es solo un caso especial de un modelo más general para el comportamiento viscoelástico. Rabotnov lo introdujo de manera formalmente diferente pero equivalente.

Para simplificar un poco las consideraciones, de ahora en adelante introducimos la hipótesis mencionada anteriormente sobre la ecuación diferencial dada: Suponemos que puede resolverse explícitamente para la derivada de orden más alto, es decir, que la ecuación diferencial puede reescribirse en la forma:

$$D_{*0}^{n_k} y(x) = f(x, y(x), D_{*0}^{n_1} y(x), D_{*0}^{n_2} y(x), D_{*0}^{n_3} y(x), \dots, D_{*0}^{n_{k-1}} y(x))$$

con una función adecuada f .

Nuestro enfoque se basa esencialmente en lo siguiente teorema de equivalencia tomado de Diethelm & Ford (2004), que ampliaremos más adelante.

La declaración clave de este teorema y el teorema 7.2 siguiente, es que la situación con medidas proporcional a las ecuaciones diferenciales fraccionarias EDFs de varios términos sujetas a valores iniciales elegidos apropiadamente.

Los valores son esencialmente los mismos que para las ecuaciones diferenciales de orden entero superior: Estos problemas de valor inicial PVI pueden reescribirse en la forma de un sistema de orden único ecuaciones.

Teorema 8.1

Considere la ecuación

$$D_{*0}^{n_k} y(x) = f(x, y(x), D_{*0}^{n_1} y(x), D_{*0}^{n_2} y(x), D_{*0}^{n_3} y(x), \dots, D_{*0}^{n_{k-1}} y(x)), \dots (8.1a)$$

Sujeta a las condiciones iniciales

$$y^{(j)}(0) = y_0^{(j)}, j = 1, 2, \dots, [n_k], \dots, (8.1.b)$$

donde

$$n_k > n_{k-1} > \dots > n_1, n_j - n_{j-1} \leq 1 \text{ para todo } j = 2, 3, \dots, k \text{ y } 0 < n_1 \leq 1$$

Supongamos que para todo $j = 1, 2, \dots, k$, defina M como el mínimo común múltiplo de los denominadores de y conjunto y Luego El problema de valor inicial PVI es equivalente al sistema de ecuaciones.

$$\begin{aligned}
D_{*0}^Y y_0(x) &= y_1(x) \\
D_{*0}^Y y_1(x) &= y_2(x) \\
&\vdots \\
D_{*0}^Y y_{N-2}(x) &= y_{N-1}(x), (8.2a) \\
D_{*0}^Y y_{N-1}(x) &= f(x, y_0(x), \underset{Y}{y_{n_1}}(x), \dots, \underset{Y}{y_{n_{k-1}}}(x))
\end{aligned}$$

junto con las condiciones iniciales,

$$y_j(0) = \begin{cases} y_0^{(j/M)}, & \text{si } j/M \in \mathbb{N} \\ 0, & \text{en otros casos} \end{cases}$$

O

$$\left\{ \begin{array}{l} D_{*0}^Y y_0(x) = y_1(x) \\ D_{*0}^Y y_1(x) = y_2(x) \\ \vdots \\ D_{*0}^Y y_{N-2}(x) = y_{N-1}(x) \\ D_{*0}^Y y_{N-1}(x) = f(x, y_0(x), \underset{Y}{y_{n_1}}(x), \dots, \underset{Y}{y_{n_{k-1}}}(x)) \end{array} \right., \dots, (8.2a)$$

junto con las condiciones iniciales

$$y_j(0) = \begin{cases} y_0^{(j/M)}, & \text{si } j/M \in \mathbb{N} \\ 0, & \text{en otros casos} \end{cases}$$

en el siguiente sentido.

1°. Siempre que $Y := (y_0, \dots, y_{N-1})^T$ con $y_0 \in C^{[n_k]}[0, b]$ para algún $b > 0$ es la solución del sistema (8.2), la función $y :=$ resuelve la ecuación multitérmino problema de valor inicial PVI (8.1).

2°. Siempre que $y \in C^{[n_k]}[0, b]$ sea una solución del problema de valor inicial PVI de varios términos (8.1), la función vectorial

$Y := (y_0, \dots, y_{N-1})^T := (y, D_{*0}^Y y, D_{*0}^{2Y} y, \dots, D_{*0}^{(N-1)Y} y)^T$ resuelve el problema de valor inicial PVI multidimensional (8.2)

Prueba

En el caso $n_j \in \mathbb{N}$ para todo j tenemos $M = 1$ y $\gamma = 1$, y por lo tanto recuperamos un resultado estándar de la teoría de ecuaciones diferenciales ordinarias, de orden números enteros.

Así, de ahora en adelante suponemos que al menos uno de los no es un número entero. Como consecuencia tenemos $M \geq 2$ y por tanto $0 < \gamma \leq 1/2$.

Para probar la primera afirmación, debemos suponer que $(y_0, \dots, y_{N-1})^T$ es una solución del sistema, y definimos $y := y_0$. Luego, mediante una aplicación repetida de Lema 3.13, que es legal ya que $\gamma < 1$ en combinación con el sistema (8.2a) tenemos

$$\begin{aligned}
 D_{*0}^{\gamma} y(x) &= D_{*0}^{\gamma} y_0(x) = y_1(x) \\
 D_{*0}^{2\gamma} y(x) &= D_{*0}^{\gamma} D_{*0}^{\gamma} y(x) = D_{*0}^{\gamma} y_1(x) = y_2(x) \\
 D_{*0}^{3\gamma} y(x) &= D_{*0}^{\gamma} D_{*0}^{2\gamma} y(x) = D_{*0}^{\gamma} y_2(x) = y_3(x) \\
 &\vdots \\
 D_{*0}^{(N-1)\gamma} y(x) &= D_{*0}^{\gamma} D_{*0}^{(N-2)\gamma} y(x) = D_{*0}^{\gamma} y_{N-2}(x), \dots, (8.3) \\
 D_{*0}^{N\gamma} y(x) &= D_{*0}^{\gamma} D_{*0}^{(N-1)\gamma} y(x) = D_{*0}^{\gamma} y_{N-1}(x) \\
 &= f(x, y_0(x), y_{n_1/\gamma}(x), \dots, y_{n_k/\gamma}(x)) \\
 &= f(x, y_0(x), D_{*0}^{n_1} y(x), \dots, D_{*0}^{n_k-1} y(x))
 \end{aligned}$$

Por definición, $N_{\gamma} = \frac{M n_k}{M} = n_k$ y, por tanto, el lado izquierdo de la última ecuación es simplemente $D_{*0}^{n_k} y(x)$. Por tanto, la función y satisfacen la ecuación diferencial (8.1a). Además, del sistema de ecuaciones (8.3) se desprende claramente que, para $\ell = 0, 1, 2, \dots, [n_k] - 1$ tenemos $y^{(\ell)}(0) = D_{*0}^{\ell} y(0) = y_{\ell/\gamma}(0) = y_{\ell/M}(0) = y_0^{(\ell)}$ y entonces la función y también satisface las condiciones iniciales (8.1b).

Para la segunda afirmación tenemos que suponer que y satisface la ecuación de varios términos (8.1a) y las condiciones iniciales (8.1b). El sistema de ecuaciones (8.3) es válido en este caso. también, y por lo tanto se sigue que la función vectorial Y de hecho satisface (8.2a) y las condiciones iniciales $y_j(0) = y_0^{(j/M)}$

siempre que $j/M \in \mathbb{N}_0$.

Finalmente, una aplicación del Lema 3.11 revela que $y_j(0) = 0$ en los otros casos.

En muchas aplicaciones prácticas se tiene $n_k < 1$.

Una inspección minuciosa de la prueba revela que podemos relajar la condición de racionalidad del Teorema 8.1 en tal caso. Nosotros entonces podemos establecer la siguiente modificación.

Teorema 8.2

Considere la ecuación (8.1a) sujeta a las condiciones iniciales (8.1b) como en el teorema 8.1. Sea ahora $1 \geq n_k > n_{k-1} > \dots > n_1 > 0$, y supongamos que la ecuación es conmensurable. Definir $\tilde{n}_j := \frac{n_j}{n_1}$ para $j = 1, \dots, k$, sea $\tilde{M}$ el mínimo múltiplo común de los denominadores de los valores $\tilde{n}_1, \tilde{n}_2, \dots, \tilde{n}_k$ y conjunto $\gamma := \frac{n_1}{\tilde{M}}$ y $N := \frac{\tilde{M}n_k}{n_k}$.

Entonces este problema de valor inicial PVI es equivalente al sistema de ecuaciones (8.2a) junto con las condiciones iniciales (8.2b) en el mismo sentido que en Teorema 8.1.

Observación 8.2

En el caso del Teorema 8.2, debido a la restricción $n_k \leq 1$, el valor inicial de las condiciones (8.2b) del sistema se pueden simplificar a

$$y_j(0) = \begin{cases} y_0^{(0)}, & \text{si } j = 0 \\ 0, & \text{en otros casos} \end{cases}$$

Prueba

La demostración es idéntica a la demostración del teorema anterior. En ese resulta que nosotros tuvimos que imponer el su-

puesto de racionalidad sólo porque necesitábamos asegurarnos de que es una fracción unitaria, es decir, $1/\gamma$ es un número entero.

La razón de esto fue que tuvimos que combinar los valores iniciales dados, correspondientes derivadas de orden entero, con los sistemas (8.2a) de tal manera que teníamos ecuaciones correspondientes al orden de los números enteros derivados también.

En el presente caso esto no es necesario porque la única información inicial dada La condición está relacionada con la función y misma.

Como consecuencia inmediata de este resultado, podemos deducir una existencia teorema y un teorema de unicidad

Teorema 8.3

Suponga las hipótesis del Teorema 8.1 o del Teorema 8.2. Además, deja $K > 0$, $h^* > 0$ y $G := [0, h^*] \times [y_0^{(0)} - K, y_0^{(0)} + K] \times \prod_{j=1}^k T_j$ donde $T_j = [-K, K]$ si $n_j \notin \mathbb{N}_0$ y $T_j = [y_0^{(n_j)} - K, y_0^{(n_j)} + K]$ más. Si $f : G \rightarrow \mathbb{R}$ es continua, entonces El problema de valor inicial PVI de varios términos (8.1) tiene una solución en el intervalo $[0, h]$ con algunos $h > 0$.

Prueba

Por el enunciado de equivalencia del teorema 8.1 o 8.2, respectivamente, podemos reducir la pregunta de existencia para la ecuación de términos múltiples a la pregunta de existencia para la ecuación vectorial (8.2a) con condiciones iniciales (8.2b).

Entonces es fácil ver que para esta ecuación podemos usar el resultado de existencia del Teorema 6.1, ver también Observación 6.1, para deducir que existe una solución en el intervalo $[0, h]$ con una solución adecuada $h > 0$.

Teorema 8.4

Suponga las hipótesis del Teorema 8.1 o del Teorema 8.2. Además, defina el conjunto G como en el teorema 8.3.

Si $f : G \rightarrow \mathbb{R}$ es continua y satisface una condición de Lipschitz con respecto a todas las variables excepto la primera, entonces existe algún $h > 0$ tal que el problema de valor inicial PVI de varios términos (8.1) tenga una solución única en el intervalo $[0, h]$.

Prueba

Como en la demostración del teorema 8.3, utilizamos el enunciado de equivalencia de Teoremas 8.1 u 8.2, respectivamente, para reducir la pregunta de unicidad para la ecuación de múltiples términos a la misma pregunta para la ecuación con valores vectoriales (8.2a) sujeta a las condiciones iniciales (8.2b). Esta vez invocamos el resultado de unicidad del Teorema 5.5 de Picard-Lindelof también el teorema 6.5. Observación 6.1, para deducir que el problema (8.2) efectivamente tiene una solución única.

Nuestro próximo objetivo es obtener algunos resultados tipo Gronwall (Momani, 2000). Específicamente demostramos que, bajo pequeñas variaciones en los órdenes en la ecuación diferencial multitérminos (8.1a) pero sujeto a la suposición de que todos los demás datos proporcionados permanecen sin cambios, podemos dar un límite uniforme en el cambio en la solución en cualquier intervalo acotado cerrado $[0, h]$.

Enunciamos y demostramos el resultado de una ecuación con dos términos, pero La generalización a ecuaciones de varios términos es sencilla en Benchohra et al. (2008).

Tenga en cuenta que se puede encontrar una discusión estrechamente relacionada de resultados similares que se aplican a ecuaciones integrales, por ejemplo, en Brunner & Houwen (1986).

Lema 8.5. Primera desigualdad de Gronwall para ecuaciones de dos términos.

Sea $n_2 > 0$ y $n_1, \tilde{n}_1 \in (0, n_2)$ se elige de manera que las ecuaciones

$$D_{*0}^{n_2} y(x) = f(x, y(x), D_{*0}^{n_1} y(x)), \dots (8.4a)$$

Sujeto a las condiciones iniciales,

$$y(x) = y_0, y'(0) = y'_0, \dots, y^{([n_2]-1)}(0) = y_0^{([n_2]-1)}, \dots (8.4b)$$

Y,

$$D_{*0}^{n_2} z(x) = f(x, z(x), D_{*0}^{n_1} z(x)), \dots (8.5a)$$

Sujeta a las mismas condiciones iniciales

$$z(x) = y_0, z'(0) = y'_0, \dots, z^{([n_2]-1)}(0) = y_0^{([n_2]-1)}, \dots (8.5b)$$

(donde f satisface una condición de Lipschitz en su segundo y tercer argumento en un dominio adecuado) tienen soluciones continuas únicas $y, z : [0, h] \rightarrow \mathbb{R}$. Suponemos además ese $n_1 = \tilde{n}_1$. Entonces existen constantes k_1 y k_2 tales que

$$|y(x) - z(x)| \leq k_1 |n_1 - \tilde{n}_1| E_{n_2}(k_2 h^{n_2}), \dots (8.6)$$

Para todo $x \in [0, h]$

Observación 8.3

Tenga en cuenta que si $n_1, \tilde{n}_1, n_2$ son todos racionales, entonces las ecuaciones (8.4) y (8.5) pueden reescribirse como sistemas de ecuaciones como se describe en el teorema 8.1, y ambos las ecuaciones tienen soluciones continuas únicas en vista del teorema 8.4.

Prueba

Los pasos esenciales de la prueba son muy similares a los que se encuentran en el teorema en la sección. 6.3: Escribimos las soluciones y y z en forma de equivalente Ecuaciones integrales de Volterra

$$y(x) = \sum_{j=0}^{|n_2|-1} \frac{y_j}{j!} x^j + \frac{1}{\Gamma(n_n)} \int_0^{n_2} (x-t)^{n_2-1} f(t, y(t), D_{*0}^{n_1} y(t)) dt, \dots (8.7)$$

Y,

$$z(x) = \sum_{j=0}^{|n_2|-1} \frac{y_j}{j!} x^j + \frac{1}{\Gamma(n_n)} \int_0^{n_2} (x-t)^{n_2-1} f(t, y(t), D_{*0}^{\tilde{n}_1} z(t)) dt, \dots (8.8)$$

y fijamos $h > 0$. Restando obtenemos la relación

$$\begin{aligned} y(x) - z(x) &= \frac{1}{\Gamma(n_n)} \int_0^{n_2} (x-t)^{n_2-1} f(t, y(t), D_{*0}^{\tilde{n}_1} y(t)) dt \\ &\quad + \frac{1}{\Gamma(n_n)} \int_0^{n_2} (x-t)^{n_2-1} f(t, y(t), D_{*0}^{\tilde{n}_1} z(t)) dt \end{aligned}$$

Ahora, con $m \in \mathbb{N}$ elegido de modo que $m-1 < n_1, \tilde{n}_1 < m$, y teniendo en cuenta que y es la solución única a (8.4) en $[0, h]$ podemos estimar, usando la condición de Lipschitz en f y la definición de la derivada de Caputo, el primer término del lado derecho:

$$\left| \frac{1}{\Gamma(n_n)} \int_0^{n_2} (x-t)^{n_2-1} \left(f(t, y(t), D_{*0}^{n_1} y(t)) - f(t, y(t), D_{*0}^{\tilde{n}_1} z(t)) \right) dt \right| \leq k_1 |n_1 - \tilde{n}_1|$$

uniformemente para $x \in [0, h]$. Además, podemos usar las condiciones de Lipschitz en f en el segundo término en el lado derecho para dar

$$\left| \frac{1}{\Gamma(n_n)} \int_0^{n_2} (x-t)^{n_2-1} \left(f(t, y(t), D_{*0}^{n_1} y(t)) - f(t, y(t), D_{*0}^{\tilde{n}_1} z(t)) \right) dt \right| \leq k_2 J_0^{n_2} |y - z|(x)$$

evaluando las representaciones integrales de las derivadas fraccionarias de y y z . Si ponemos $\delta(x) = y(x) - z(x)$ se deduce que

$$|\delta(x)| \leq k_1 |n_1 - \tilde{n}_1| k_2 J_0^{n_2} |\delta|(x), \dots (8.9)$$

La ecuación (8.9) ahora nos permite concluir mediante el Lema 6.19 que

$$|\delta(x)| \leq k_1 |n_1 - \tilde{n}_1| E_{n_2}(k_2 h^{n_2})$$

uniformemente para $x \in [0, h]$ y la prueba está completa.

Una afirmación análoga es cierta si variamos el orden de los otros operadores

diferenciales.

Lema 8.6. Segunda desigualdad de Gronwall para ecuaciones de dos términos.

Sea $n_1 > 0$, y $\tilde{n}_2, n_2 \in (n_1, \infty)$ se eligen de modo que $\tilde{n}_2 = n_2$ y de modo que las ecuaciones (8.4) y

$$D_{*0}^{n_2} z(x) = f(x, z(x), D_{*0}^{n_1} z(x)), \dots (8.10a)$$

Sujeta a las mismas condiciones iniciales,

$$z(x) = y_0, z'(0) = y'_0, \dots, z^{(n_2-1)}(0) = y_0^{(n_2-1)}, \dots (8.10b)$$

(donde f satisface una condición de Lipschitz en su segundo y tercer argumento en un dominio adecuado) tienen soluciones continuas únicas $y, z : [0, h] \rightarrow \mathbb{R}$. Entonces

$$\sup_{x \in [0, h]} |y(x) - z(x)| = O(n_1 - \tilde{n}_1), \dots (8.11)$$

Prueba

La demostración es esencialmente la misma que la del Teorema 6.22; omitimos los detalles.

Usamos las conclusiones de los Lemas 8.5 y 8.6 varias veces. Primero derivamos un Teorema de existencia y unicidad para ecuaciones no lineales de términos múltiples más generales. (con derivadas múltiples no conmensuradas). Como antes, damos una prueba para la ecuación de dos términos; el resultado se puede generalizar fácilmente para ecuaciones con más términos.

Teorema 8.7

Sea la función acotada $f: [0, h] \times \mathbb{R}^2 \rightarrow \mathbb{R}$ satisfacer una función uniforme condición de Lipschitz en su segundo y tercer argumento y ser continua en su primero argumento. De ello se deduce que la ecuación

$$D_{*0}^{n_2} z(x) = f(x, z(x), D_{*0}^{n_1} z(x)), \dots (8.12)$$

Donde $n_2 > n_1 > 0$, sujeto a las condiciones iniciales

$$D^k y(0) = y_0^{(k)}, k = 1, 2, \dots, [n_2] - 1$$

tiene una solución continua única en $[0, h]$.

Prueba.

Caso 1: $n_1, n_2 \in \mathbb{Q}$ observamos que el resultado ya está establecido cuando son racionales debido al teorema 8.4.

Caso 2: . Construimos una sucesión $(n_1, j)_{j=1}^{\infty}$ de números racionales cuyo límite es n_1 . Sin pérdida de generalidad, podemos suponer que todos los miembros de la secuencia están contenidos en el intervalo $(\bar{n}_1, n_1 + 1)$. Claramente en el caso 1 la ecuación (8.12) con n_1 reemplazado a su vez por cada j , sujeto a la inicial dada condiciones, tiene una solución continua única $y[j]$ en $[0, h]$ cuya derivada de Caputo de orden n_2 también es continua. Ahora usamos la ecuación (8.9) junto con el hecho que $n_{1,j} - n_1 \rightarrow 0$ cuando $j \rightarrow \infty$ para concluir que la secuencia de soluciones converge uniformemente en $[0, h]$ a una función continua y con $D_{*0}^{n_2} y(x)$ siendo también continuo. Él queda por demostrar que esta función y es la solución de (8.12).

Para ello definamos $r_j(z) := \|D_{*0}^{n_2} z - f(\cdot, z(\cdot), D_{*0}^{n_2, j} z(\cdot))\|_{L_{\infty}[0, h]}$ para cualquier z tal que $D_{*0}^{n_2} z$ es continuo. Por definición, obtenemos inmediatamente $r_j(y^{[j]}) = 0$ para todo j . Dado que $y^{[j]} \rightarrow y$ uniformemente cuando $j \rightarrow \infty$ y r es continua en el espacio que consideramos, tenemos que encontrar eso

$$r_j(y^{[j]}) \rightarrow r_j(y) \rightarrow 0, \text{ como } j \rightarrow \infty$$

Por lo tanto y es una solución del problema de valor inicial dado.

Debido a la condición de Lipschitz en f , podemos probar la unicidad de la solución mediante las técnicas habituales de iteración de Picard, es decir, procedemos como en la prueba de Teorema 6.5. Caso 3: $n_2 \notin \mathbb{Q}$. Aquí usamos una secuencia (n_2, j) de números racionales que satisfacen $n_2 < n_{2,j} < [n_2]$ para todo j y

$\lim_{j \rightarrow \infty} n_2, j = n_2$. Para cada j podemos usar cualquier caso 1 o el caso 2 (dependiendo de si $n_1 \in \mathbb{Q}$ o no) para proporcionar una solución única a la ecuación perturbada obtenida reemplazando n_2 por n_2, j sujeta a la ecuación sin modificar condiciones iniciales. Luego procedemos como en el caso 2 (pero aplicando el Lema 8.6 en lugar de Lema 8.5) para mostrar que la secuencia correspondiente de soluciones converge a la solución única del problema original.

Más adelante en esta sección daremos una prueba alternativa de este resultado, cf. Observación 8.5. Ahora podemos utilizar los lemas de Gronwall para demostrar la estabilidad estructural del problema de valor inicial (7.1) incluso bajo pequeñas perturbaciones en los órdenes de los derivados.

Teorema 8.8

Sea y la solución de

$$D_{*0}^{n_k} y(x) = f(x, y(x), D_{*0}^{n_1} y(x), D_{*0}^{n_2} y(x), D_{*0}^{n_3} y(x), \dots, D_{*0}^{n_{k-1}} y(x))$$

Con condiciones iniciales,

$$y^{(j)}(0) = y_0^{(j)}, j = 0, 1, 2, \dots, [n_k] - 1$$

Y sea z la solución de

Con condiciones iniciales,

$$z^{(j)}(0) = z_0^{(j)}, j = 0, 1, 2, \dots, [n_k] - 1$$

donde $|n_j - \tilde{n}_j| < \varepsilon$ para todo $j = 1, 2, \dots, k$. Entonces existe algún $h > 0$ tal que ambas ecuaciones tienen una solución continua única en el intervalo $[0, h]$, y

$$\|y - z\|_{L_\infty[0, h]} = O(\varepsilon), \varepsilon \rightarrow 0$$

Prueba

El teorema se deriva de la observación de que la diferencia $y - z$ es una Función de Lipschitz de $n_j - \tilde{n}_j, j, j = 1, 2, \dots, k$, debido a los Lemas tipo Gronwall 8.5 y 8.6.

Observación 8.4

De la aplicación del teorema 8.8 se deduce que la solución de cualquier una ecuación diferencial fraccionaria de orden múltiple no conmensurada puede ser arbitrariamente estrechamente aproximado en un intervalo finito $[0, h]$ mediante soluciones de ecuaciones de racional orden (que a su vez puede resolverse mediante la conversión a un sistema de ecuaciones de bajo orden).

Hasta este punto, nuestro enfoque principal para manejar ecuaciones de términos múltiples se basó sobre reescribirlos en forma de una ecuación de un solo término para una función con valores vectoriales.

Esta última ecuación tiene una estructura formalmente muy simple y atractiva, y Todo el proceso es bastante natural, pero el enfoque tiene dos desventajas. En primer lugar, sólo funciona exactamente en el caso de ecuaciones proporcionales (si $n_k < 1$), o incluso más restrictivamente, en el caso de ecuaciones de orden racional (si $n_k \geq 1$). En todos otros casos no podemos reemplazar la ecuación de múltiples términos dada por una correspondiente sistema de orden único exactamente; debemos contentarnos con una aproximación.

El segundo problema potencial no es un obstáculo importante desde el punto de vista analítico, pero puede tener efectos muy desagradables cuando se intenta emplearlo para la solución nu-

mérica de ecuaciones de términos múltiples como se analiza, por ejemplo, en Diethelm (2003) y Diethelm (s. f.). Dependiendo de los valores exactos de los órdenes de los operadores diferenciales, el parámetro γ en el Teorema 8.1 o 8.2 puede ser muy pequeño.

Esto significa que la dimensión del sistema resultante, que es proporcional a $1/\gamma$, puede ser extremadamente grande.

En vista de estas posibles dificultades Ahora veremos un enfoque diferente para manejar ecuaciones de términos múltiples que también nos proporcionará una prueba alternativa de la existencia general y única resultado indicado en el teorema 8.7 anterior.

Este segundo enfoque que se describe, por ejemplo, en Diethelm (2008b), se basa en una alternativa forma de establecer un modelo matemático que involucra más de una derivada fraccionaria. Específicamente, podemos usar un sistema de ecuaciones diferenciales fraccionarias donde cada ecuación tiene un orden que puede coincidir o no con el orden de las otras ecuaciones. Para decirlo más formalmente, esto conduce a un modelo del tipo

$$D_{*0}^{n_1} y_1(x) = f_1(x, y_1(x), \dots, y_k(x))$$

$$\vdots \dots, \dots (8.13a)$$

$$D_{*0}^{n_k} y_k(x) = f_k(x, y_1(x), \dots, y_k(x))$$

Como veremos, es suficiente para nuestros propósitos suponer que $0 < n_j \leq 1$ para todo j . Esto implica que las condiciones iniciales para el sistema de ecuaciones diferenciales 8.13a) tiene la forma

$$y_j(0) = y_{j,0}, (j = 1, 2, \dots, k), \dots (8.13b)$$

Un sistema de esta clase se denominará sistema diferencial fraccionario de varios órdenes. Los sistemas diferenciales fraccionarios de orden múltiple parecen investigarse con menos frecuencia que las ecuaciones multitérmino, pero ahora describiremos algunas conexiones estrechas entre los dos conceptos y, por lo tanto, los primeros merecen cierta atención, al menos a la vista. de que pueden ser herramientas muy útiles para el tratamiento analítico y numérico de estos últimos.

Una vez más, queremos reescribir la ecuación diferencial fraccionaria de varios términos dada en forma de un sistema de ecuaciones de un solo término. En contraste con el método Como se describió anteriormente, este sistema ahora será un sistema de órdenes múltiples. Para ello será útil suponer que todos los números enteros que están contenidos en el intervalo $(0, n_k)$ también son miembros de la secuencia finita $(n_j)_{j=1}^k$. En otras palabras, es imposible que dos elementos consecutivos de la secuencia finita (n_j) se encuentren en lados opuestos. lados de un número entero. Es obvio que tal suposición no conduce a cualquier pérdida de generalidad. Entonces, dada la ecuación (8.1a), podemos escribir $\beta_1 := n_1, \beta_j := n_j - n_j - 1, (j=2,3,\dots,k)$, $y_1 := y, y_j := D_*^{n_j-1} y, j = 2,3,\dots,k$. Tenga en cuenta que Según nuestros supuestos sobre β_j , está claro que $0 < \beta_j \leq 1$ para todo j . Entonces podemos concluir el siguiente resultado de equivalencia.

Teorema 8.9

Sujeto a las condiciones anteriores, la ecuación de términos múltiples (9.1a) con condiciones iniciales (8.1b) es equivalente al sistema

$$\begin{aligned}
D_{*0}^{\beta_1} y_1(x) &= y_2(x) \\
D_{*0}^{\beta_2} y_2(x) &= y_3(x) \\
D_{*0}^{\beta_3} y_3(x) &= y_4(x) \\
&\vdots \\
D_{*0}^{\beta_{k-1}} y_{k-1}(x) &= y_k(x) \quad , \dots, (8.14a) \\
D_{*0}^{\beta_k} y_k(x) &= f(x, y_1(x), y_2(x), \dots, y_k(x))
\end{aligned}$$

Con las condiciones iniciales,

$$y_j(0) = \begin{cases} y_0^{(0)}, si \ j = 1 \\ y_0^{(\ell)}, si \ n_{j-1} = \ell \in \mathbb{N}, \dots \\ 0, EN \ OTROS \ CASOS \end{cases} \quad (8.14b)$$

en el siguiente sentido:

1°. Siempre que la función $y \in C^{[n_k]}[0, X]$ sea una solución de la ecuación multitérmino (8.1a) con las condiciones iniciales (8.1b), la función vectorial $Y := (y_1, \dots, y_k)^T$ con

$$y_j(0) = \begin{cases} y(x), si \ j = 1 \\ D_{*0}^{n_{k-1}} y(x), si \ j \geq 2, \dots \end{cases} \quad (8.15)$$

es una solución del sistema diferencial fraccionario de varios órdenes (8.14a) con inicial condiciones (8.14b).

2°. Siempre que la función vectorial $Y := (y_1, \dots, y_k)^T$ sea una solución del sistema diferencial fraccionario de orden múltiple (8.14a) con condiciones iniciales (8.14b), la función $y := y_1$ es una solución de la ecuación de múltiples términos (8.1a) con condiciones iniciales (8.1b).

Prueba

La prueba sigue exactamente los mismos pasos que la prueba de Teorema 8.1.

$$\begin{aligned}
& D_{*0}^{3.3} y(x) \\
& = f \left(x, y(x), D_{*0}^{0.1} y(x), D_{*0}^1 y(x), D_{*0}^{1.2} y(x), \right. \\
& \quad \left. D_{*0}^{1.5} y(x), D_{*0}^{1.7} y(x), D_{*0}^2 y(x), D_{*0}^{2.2} y(x), D_{*0}^{2.6} y(x), D_{*0}^3 y(x) y(x) \right), \dots \quad (8.16) \\
& \quad y^{(j)}(0) = y_0^{(j)}, j = 1, 2, 3
\end{aligned}$$

en la forma indicada en el teorema 8.9.

En este caso, el enfoque del Teorema 8.9 nos da el sistema

$$\begin{aligned}
D_{*0}^{0.1} y_1(x) &= y_2(x), D_{*0}^{0.9} y_2(x) = y_3(x) \\
D_{*0}^{0.2} y_3(x) &= y_4(x), D_{*0}^{0.3} y_4(x) = y_5(x) \\
D_{*0}^{0.2} y_5(x) &= y_6(x), D_{*0}^{0.3} y_6(x) = y_7(x) \\
D_{*0}^{0.2} y_7(x) &= y_8(x), D_{*0}^{0.4} y_8(x) = y_9(x) \\
D_{*0}^{0.4} y_9(x) &= y_{10}(x), D_{*0}^{0.3} y_{10}(x) = f(x, y_1(x), y_2(x), \dots, y_{10}(x))
\end{aligned}$$

Con las condiciones iniciales,

$$\begin{aligned}
y_1(0) &= y_0^{(0)}, y_3(0) = y_0^{(1)} \\
y_7(0) &= y_0^{(2)}, y_{10}(0) = y_0^{(3)} \\
y_j(0) &= 0, \text{ para } j \in \{2, 4, 5, 6, 8, 9\}
\end{aligned}$$

Además, observamos la correspondencia exacta entre las funciones componentes de la solución del sistema de varios términos, por un lado, y las derivadas fraccionarias de la solución y de la ecuación de varios términos, por un lado

$$\begin{aligned}
y_1 &= y, y_2 = D_{*0}^{0.1} y, y_3 = D_{*0}^1 y \\
y_4 &= D_{*0}^{1.2} y, y_5 = D_{*0}^{1.5} y, y_6 = D_{*0}^{1.7} y \\
y_7 &= D_{*0}^2 y, y_8 = D_{*0}^{1.2} y, y_9 = D_{*0}^{2.6} y, \dots \quad (8.17) \\
y_{10} &= D_{*0}^3 y
\end{aligned}$$

Edwards et al. (2002); Ford & Simpson (2000), han desarrollado un enfoque alternativo para la conversión de ecuaciones de términos múltiples a sistemas de órdenes múltiples. Para descri-

bir este método asumimos, en aras de la simplicidad, que el operador diferencial de orden más alto no es una derivada de orden entero. De lo contrario, son necesarias algunas pequeñas modificaciones formales en la notación, pero el concepto básico y los resultados principales permanecen sin cambios.

La idea fundamental se explica mejor observando el problema del ejemplo 8.1 nuevamente. Tenemos dos objetivos en mente. El primero es conservar la estructura del sistema (8.17) (en particular, la dimensión del sistema y la estructura de las condiciones iniciales). Por otro lado, queremos combinar los componentes en una manera diferente de modo que se minimice el número de derivadas de orden no entero. Esto nos lleva al sistema.

$$\begin{aligned}
 D_{*0}^{0.1} y_1(x) &= y_2(x), D_{*0}^1 y_1(x) = y_3(x) \\
 D_{*0}^{0.2} y_3(x) &= y_4(x), D_{*0}^{0.5} y_3(x) = y_5(x) \\
 D_{*0}^{0.7} y_3(x) &= y_6(x), D_{*0}^1 y_3(x) = y_7(x) \\
 D_{*0}^{0.2} y_7(x) &= y_8(x), D_{*0}^{0.6} y_7(x) = y_9(x) \\
 D_{*0}^1 y_7(x) &= y_{10}(x), D_{*0}^{0.3} y_{10}(x) = f(x, y_1(x), y_2(x), \dots, y_{10}(x))
 \end{aligned}$$

El enfoque mediante el teorema 8.9 creó diez ecuaciones diferenciales de órdenes estrictamente fraccionales, mientras que ahora tenemos siete estrictamente fraccionarias y tres de primer orden. Estas ecuaciones su característica parece ser atractiva desde el punto de vista del cálculo costo cuando tales sistemas se resuelven numéricamente: una ecuación de primer orden es más barata resolver que una ecuación de orden fraccionario porque los operadores involucrados en esta última son no locales. Sin embargo, como indican Ford y Connolly (Dokoumetzidis et al., 2010), la estructura del sistema puede ser tan inconveniente que la ventaja obtenida por esta localidad puede perderse completamente.

ara una descripción formal general de este método es ventajoso expresar la ecuación de varios términos en la forma

$$\begin{aligned} D_{*0}^{k+\delta_{k,\ell_k}} y(x) &= f(x, D^0 y(x), D_{*0}^{\delta_{0,1}} y(x), \dots, D_{*0}^{\delta_{0,\ell_0}} y(x), \\ &D^1 y(x), D_{*0}^{1+\delta_{1,1}} y(x), \dots, D_{*0}^{1+\delta_{1,\ell_1}} y(x), \dots, (8.18a) \\ &D^k y(x), \dots, D_{*0}^{1+\delta_{k,\ell_k}} y(x) \end{aligned}$$

donde $0 < \delta_{j,1} < \delta_{j,2} < \dots < \delta_{j,\ell_j} < 1$ para todo j . Las condiciones iniciales correspondientes. son entonces

$$y_j(0) = y_0^{(j)}, \text{ para } j = 0, 1, 2, 3, \dots, k, \dots (8.18b)$$

Para lograr nuestro objetivo, definimos,

$$s(\mu, \sigma) := \sigma + \mu + 1 + \sum_{j=0}^{\mu-1} \ell_j, \text{ y } N := s(k, \ell_k) - 1 = k + \sum_{j=0}^k \ell_j$$

Observe que N es el número total de operadores diferenciales de orden estrictamente positivo en (9.18a). Por tanto, en nuestra terminología, (9.18a) es una ecuación de N términos. Usando esta notación, llegamos a la siguiente afirmación.

Teorema 8.10

El problema de valor inicial de varios términos (8.18) es equivalente al sistema n -dimensional

$$\begin{aligned} D_{*0}^{\delta_{\mu,\sigma}} y_{s(\mu,0)}(x) &= y_{s(\mu,\sigma)}(x), \mu = 0, 1, 2, \dots, k, \sigma = 1, 2, \dots, \hat{\sigma}_\mu \\ D^1 y_{s(\mu,0)}(x) &= y_{s(\mu+1,0)}(x), \mu = 0, 1, 2, \dots, k-1, \dots, (8.19a) \\ D_{*0}^{\delta_{k,n_k}} y_{s(k,0)}(x) &= f(x, y_1(x), y_2(x), \dots, y_N(x)) \end{aligned}$$

Donde, $\hat{\sigma}_\mu := \ell_\mu$, sí $0 \leq \mu < k$, y $\hat{\sigma}_k := \ell_k - 1$, con las condiciones iniciales

$$y_j(0) = \begin{cases} y_0^{(k)}, & \text{si existe } k \text{ tal que } j = s(k, 0), \dots \\ 0, & \text{si caso contrario} \end{cases} \quad (8.19 \text{ b})$$

en el siguiente sentido:

Siempre que la función $y \in C^{k+1}[0, X]$ es una solución de la ecuación multitérmino (8.18a) con condiciones iniciales (8.18b), la función vectorial $Y = (y_1, \dots, y_N)^T$ con

$$y_{s(\mu, \sigma)}(x) := \begin{cases} D^\mu y(x), & \text{para } \sigma = 0 \\ D_{\sigma_0}^{\mu + \delta_{\mu, \sigma}} y(x), & \text{para } \sigma = 1, 2, \dots, \hat{\sigma}_\mu, \end{cases} \quad \mu = 0, 1, 2, \dots, k, \dots \quad (8.20)$$

es una solución del sistema multiorden (8.19a) con condiciones iniciales (7.19b).

Siempre que la función vectorial $Y = (y_1, \dots, y_N)^T$ es una solución del sistema multiorden (8.19a) con condiciones iniciales (8.19b), la función $y := y_1$ es una solución de la ecuación de múltiples términos (8.18a) con condiciones iniciales (8.18b)

Este resultado se ha expresado sin prueba para un subconjunto de la clase de ecuaciones descrito (8.18a) (Doetsch, 1971). Es evidente que el método utilizado en la prueba del teorema 8.1 anterior, es decir, una aplicación repetida de nuestro Lema 3.13, puede una vez más utilizarse para dar una demostración formal del teorema 8.10. Dejamos los detalles al lector. Se puede encontrar una aplicación útil de los conceptos desarrollados anteriormente, como lo señala Ford et al. (2006); (Mainardi, 2012), en el contexto de ecuaciones diferenciales fraccionarias de orden único cuyo el orden es mayor que 1. De hecho, dado un problema de valor inicial de la forma

$$D_{\sigma_0}^n y(x) = f(x, y(x)), y^{(j)}(0) = y_0^{(k)}, (k = 0, 1, 2, \dots, [n] - 1), \dots \quad (8.21)$$

con algún número no entero $n > 1$, podemos interpretar el lado derecho del diferencial ecuación formalmente en función de x y $D^k y(x), k = 0, 1, 2, \dots, [n] - 1$, que en realidad no depende de $D^k y(x), k = 0, 1, 2, \dots, [n] - 1$. Luego, usando nuestras técnicas desarrollado en el teorema 8.9 podemos reescribir el problema dado en la forma equivalente,

$$\begin{aligned} D^1 y_1(x) &= y_2(x) \\ D^1 y_2(x) &= y_3(x) \\ &\vdots, \dots (8.22a) \\ D^1 y_{[n]-1}(x) &= y_{[n]}(x) \\ D_{*0}^{n-[n]} y_{[n]}(x) &= f(x, y_1(x)) \end{aligned}$$

Con las condiciones iniciales,

$$y_k(0) = y_0^{(k-1)}, (k = 1, 2, \dots, [n]), \dots (8.22b)$$

es decir, como sistema de múltiples pedidos con pedidos menores o iguales a 1. Por cierto, podríamos también he aplicado el Teorema 8.10 en lugar del Teorema 8.9; en este caso especial Ambos enfoques conducen al mismo sistema.

El sistema (8.22) es algo más fácil. objeto de trabajo numérico ya que algunos algoritmos tienden a comportarse mucho peor cuando aplicado a ecuaciones de orden superior. Además, este concepto puede considerarse una extensión de la conocida técnica clásica para la solución numérica de problemas iniciales problemas de valores de orden entero superior que consiste en reescribir el problema en la forma de un sistema de primer orden y resolver este sistema numéricamente con la ayuda de un algoritmo para problemas de valores iniciales de primer orden.

No hemos abordado las cuestiones de existencia y unicidad de las soluciones para sistemas generales de orden múltiple

(8.13a) con las condiciones iniciales correspondientes (8.13b) todavía. Para sistemas de la forma especial (8.14a) o (8.19a) podemos usar una reducción a ecuaciones de términos múltiples mediante el Teorema 8.9 o el Teorema 8.10, respectivamente. Sin embargo, ninguno de estos enfoques cubre el caso general (8.13a). Para demostrar un resultado de existencia y unicidad para este último, proponemos utilizar un método de prueba completamente diferente. Al hacerlo, es posible mostrar el siguiente resultado que en realidad cubre una clase de ecuaciones que es más general que el dado en (8.13).

Teorema 8.11

Sea $n_j > 0$ para $j = 1, 2, \dots, k$ y considere el problema del valor inicial dado por el sistema diferencial fraccionario de varios órdenes

$$D_*^{n_j} y_j(x) = f_j(x, y_1(x), \dots, y_k(x)), j = 1, 2, \dots, k, (8.23a)$$

Con condiciones iniciales,

$$y_j^{(\ell)} = y_{j,0}^{(\ell)}, \ell = 0, 1, 2, \dots, [n] - 1, j = 1, 2, \dots, k, \dots (8.23b)$$

Supongamos que las funciones $f_j: [0, X] \times \mathbb{R}^k \rightarrow \mathbb{R}$ $j = 1, 2, \dots, k$, son continuas y satisfacen las condiciones de Lipschitz con respecto a todos sus argumentos excepto el primero. Entonces el problema con valor inicial tiene una solución continua determinada de forma única.

Prueba

Una aplicación por componentes del Lema 6.2 muestra que el valor inicial El problema (8.23) es equivalente al sistema de ecuaciones de Volterra.

$$y_j(x) = \sum_{\ell=0}^{[n_j]-1} y_{j,0}^{(\ell)} \frac{x^\ell}{\ell!} + \frac{1}{\Gamma(n_j)} + \int_0^x (x-t)^{n_j-1} f_j(t, y_1(t), \dots, y_k(t)) dt, \dots \quad (8.24)$$

$j = 1, 2, \dots, k$. Luego definimos y reescribimos (8.24) en la forma

$$y_j(x) = \sum_{\ell=0}^{[n_j]-1} y_{j,0}^{(\ell)} \frac{x^\ell}{\ell!} + \frac{1}{\Gamma(n_j)} + \int_0^x (x-t)^{n_j-1} \tilde{f}_j(t, y_1(t), \dots, y_k(t)) dt$$

Para $j=1, 2, \dots, k$, donde

$$\tilde{f}_j(t, y_1(t), \dots, y_k(t)) := \frac{\Gamma(n)}{\Gamma(n_j)} (x-t)^{n_j-1} f_j(t, y_1(t), \dots, y_k(t))$$

Introduciendo la notación vectorial $Y := (y_1(t), \dots, y_k(t))^T$, $\tilde{F} := (\tilde{f}_1, \tilde{f}_2, \dots, \tilde{f}_k)^T$, y una expresión correspondiente para los valores iniciales, podemos combinar estas k ecuaciones escalares en una ecuación vectorial, a saber

$$Y(x) = \sum_{\ell=0}^{\max j [n_j]-1} Y_0^{(\ell)} \frac{x^\ell}{\ell!} + \frac{1}{\Gamma(n)} + \int_0^x (x-t)^{n_j-1} \tilde{F}(t, Y(t)) dt$$

En vista de nuestras suposiciones sobre f_j y el hecho de que $n_j \geq n$ para todo j , concluimos que todas las funciones f_j son continuos y satisfacen las condiciones de Lipschitz con respecto a $y_1, \dots, y_k$. Por tanto, F es continua y satisface una condición de Lipschitz con respecto a Y , y en vista de algunos resultados estándar de la teoría de las ecuaciones integrales de Volterra (Nkamnang, 1999), Teorema 5.8; y en Baleanu & Tenreiro Machado (2009) (Sección 4.5) podemos concluir la existencia y unicidad de una solución continua $Y := (y_1, \dots, y_k)^T$

Observación 8.5

Evidentemente, la combinación de los teoremas 8.11 y 8.9 nos da un segundo método para probar la existencia general y el resultado de unicidad para ecuaciones de términos múltiples, a saber, Teorema 8.7. Por lo tanto, hemos presentado tres formas diferentes de construir un sistema de ecuaciones diferenciales fraccionarias de un solo término a partir de una ecuación de varios términos dada.

El primer enfoque, descrito en el teorema 8.1, conduce a un sistema cuyas ecuaciones constituyentes todos tienen el mismo orden. Si bien esta característica puede parecer atractiva, sus desventajas son que la dimensión del sistema puede ser muy grande y que es aplicable sólo si los órdenes de los operadores diferenciales satisfacen ciertas condiciones de la teoría de números.

El segundo enfoque está contenido en el teorema 8.9. En general, conducirá a un sistema de ecuaciones diferenciales de varios órdenes, es decir, un sistema cuyas ecuaciones pueden contener operadores diferenciales de diferentes órdenes, pero normalmente ninguno de estos pedidos serán números enteros. Este concepto siempre es aplicable y tiende a producir muchos sistemas más pequeños que el primer enfoque.

Finalmente, el tercer enfoque (ver Teorema 8.10) es muy similar al segundo, pero reemplaza el diferencial de orden fraccionario ecuaciones mediante ecuaciones de primer orden siempre que sea posible. Siempre es aplicable también, y produce sistemas cuyas dimensiones son las mismas que las creadas por el segundo abordaje.

Una comparación detallada de estos tres métodos ha sido proporcionada por Ford y Connolly (Diethelm et al., 2004), cuyo objetivo principal era descubrir cuál de estos esquemas era más adecuado para combinarse con algoritmos numéricos para (orden único o múltiple) sistemas diferenciales fraccionarios en un intento de crear una estrategia eficiente para la solución numérica de ecuaciones de varios términos. Su principal resultado fue que la reformulación de la ecuación de múltiples términos dada en la forma indicada en el Teorema 8.10 por lo general resultó dar el desempeño más débil, mientras que los otros dos enfoques fueron mucho más eficientes.

Si la aplicación del teorema 8.1 debe ser preferir sobre el teorema 8.9 o viceversa depende de la naturaleza precisa de la situación dada ecuación de términos múltiples, particularmente en la distribución de los órdenes de la ecuación diferencial operadores involucrados.

Nos referimos al Apéndice C.2 para una discusión más detallada de este tema Para ecuaciones lineales, es posible derivar expresiones explícitas para la solución.

Las funciones de Mittag-Leffler de dos parámetros resultan ser herramientas muy útiles en este contexto. Requeriremos el siguiente resultado fundamental con respecto a la transformación de Laplace de ciertas funciones estrechamente relacionadas con estas funciones de Mittag-Leffler:

Lema 8.12

Considere la función de Mittag-Leffler de dos parámetros E_{n_1, n_2} para algunos $n_1, n_2 > 0$. Sean $j \in \mathbb{N}_0, a \in \mathbb{R}$ y

$$z_{\pm}(x) := x^{j n_1 + n_2 - 1} E_{n_1, n_2}^{(j)}(\pm x^{n_1})$$

Entonces, para $j > |a|$

$$\mathcal{L} z_{\pm}(x) = \frac{j! s^{n_1 - n_2}}{(s^{n_1} \pm a)^{j+1}}$$

Dejamos la prueba como ejercicio al lector y procedemos utilizando estos resultados. para examinar explícitamente un caso particularmente importante.

Ejemplo. Resuelve la ecuación de Bagley-Torvik

$$D_{*0}^2 y(x) + 2D_{*0}^{3/2} y(x) + 2y(x) = \sin x$$

Con condiciones iniciales,

$$y(0) = y'(0) = 0$$

Como primera aproximación, recordamos que, según el teorema 8.1, la ecuación puede ser transformado en un sistema de ecuaciones de cuatro dimensiones de orden 1/2. el preciso la forma del sistema es

$$D_{*0}^{1/2} y_0(x) = y_1(x)$$

$$D_{*0}^{1/2} y_1(x) = y_2(x)$$

$$D_{*0}^{1/2} y_2(x) = y_3(x), \dots \dots \dots (8.25)$$

$$D_{*0}^{1/2} y_3(x) = -2 y_0(x) - 2y_3(x) + \sin x$$

combinado con las condiciones iniciales $y_j(0) = 0$ para $j=0,1,2,3$, Entonces podremos adaptarnos el enunciado del Teorema 7.2 a este entorno multidimensional, es decir, tenemos que tomar teniendo en cuenta que el parámetro λ que aparece allí ahora es una matriz,

$$\lambda = \begin{pmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ -2 & 0 & 0 & -2 \end{pmatrix}, \text{ siendo } \lambda^{-1} = \begin{pmatrix} 0 & 0 & -1 & -1/2 \\ 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \end{pmatrix}$$

Como tenemos condiciones iniciales homogéneas, deducimos que la solución del sistema es

$$Y(x) = \begin{pmatrix} y_0(x) \\ y_1(x) \\ y_2(x) \\ y_3(x) \end{pmatrix} = \lambda^{-1} \int_0^x u(t) \begin{pmatrix} 0 \\ 0 \\ 0 \\ \sin(x-t) \end{pmatrix} dt$$

Donde

$$u(t) = \frac{d}{dx} E_{\frac{1}{2}} \left(\lambda x^{\frac{1}{2}} \right) = \frac{d}{dx} \sum_{j=0}^{\infty} \frac{\lambda^j x^{\frac{j}{2}}}{\Gamma(1 + j/2)}$$

es una función valorada en matriz. Por construcción, λ^{-1} conmuta con $u(t)$ (para cualquier t), y por lo tanto

$$Y(x) = \int_0^x u(t) \lambda^{-1} \begin{pmatrix} 0 \\ 0 \\ 0 \\ \sin(x-t) \end{pmatrix} dt = -\frac{1}{2} \int_0^x u(t) \begin{pmatrix} \sin(x-t) \\ 0 \\ 0 \\ 0 \end{pmatrix} dt$$

Porque originalmente estábamos interesados en la solución de la ecuación de Bagley-Torvik dada. ecuación, sólo necesitamos mirar el primer componente (con índice 0) del vector Y ; encontramos

$$Y(0) = -\frac{1}{2} \int_0^x u_{11}(t) \sin(x-t) dt$$

donde $u_{11}(t)$ es el componente superior izquierdo de la matriz $u(t)$.

Dado que la ecuación es una ecuación lineal con coeficientes constantes, alternativamente podemos usar el método de la

transformada de Laplace. Aplicando transformadas de Laplace a ambos lados de la ecuación encontramos

$$s^2 \mathcal{L}y(s) = 2s^{3/2} \mathcal{L}y(s) + \mathcal{L}y(s) = \mathcal{L} \sin(s) = \frac{1}{s^2 + 1}$$

De este modo,

$$\mathcal{L}y(s) = \frac{1}{(s^2 + 1)(s^2 + 2s^{3/2} + 2)}$$

Y nosotros encontramos

$$y(x) = \int_0^x z(x-t) \sin t dt$$

Donde,

$$\mathcal{L}z(s) = \frac{1}{(s^2 + 2s^{3/2} + 2)} = \frac{1}{2} \frac{2s^{-3/2}}{(s^{1/2} + 2)} = \frac{1}{1 + \frac{2s^{-3/2}}{s^{1/2} + 2}} = \frac{1}{2} \gamma \frac{1}{1 + \gamma}$$

Con $\gamma := 2s^{-3/2}/(s^{1/2} + 2)$. Para s suficientemente grandes (y para el propósito de Laplace teoría de transformadas es suficiente considerar solo estos valores) tenemos $|\gamma| < 1$, y por lo tanto, podemos expandir $\mathcal{L}z(s)$ usando la conocida serie geométrica. Esto produce

$$\begin{aligned} \mathcal{L}z(s) &= \frac{1}{2} \gamma \sum_{k=0}^{\infty} (-1)^k \gamma^k = \frac{1}{2} \gamma \sum_{k=0}^{\infty} (-1)^k 2^{k+1} \frac{s^{-3(k+1)/2}}{(s^{1/2} + 2)^{k+1}} \\ &= \sum_{k=0}^{\infty} (-1)^k 2^k \frac{s^{-3(k+1)/2}}{(s^{1/2} + 2)^{k+1}} \end{aligned}$$

Aquí es posible una transformada de Laplace inversa término por término y da, en vista de Lema 8.12.

$$z(x) = \sum_{k=0}^{\infty} (-1)^k \frac{2^k}{k!} x^{2k+1} E_{\frac{1}{2}, 2+3k/2}^{(k)} (-2 x^{1/2})$$

Combinando estas ecuaciones encontramos el resultado final,

$$z(x) = \sum_{k=0}^{\infty} (-1)^k \frac{2^k}{k!} \int_0^x (x-t)^{2k+1} E_{\frac{1}{2}, 2+3k/2}^{(k)} (-2 (x-t)^{\frac{1}{2}}) \sin t dt$$

Dejamos la prueba como ejercicio al lector y procedemos utilizando estos resultados para examinar explícitamente un caso particularmente importante.

Se anima al lector a dar una prueba explícita de que esta solución es idéntica a la solución obtenida por el primer enfoque.

Otra posibilidad más para encontrar la solución de la ecuación de Bagley-Torvik es proporcionado por una aplicación del Teorema 7.13 a la ecuación homogénea asociada a 8.25, seguida por el uso del método de variación de constantes como se indica en la Observación 7.1 para encontrar una solución particular del problema no homogéneo. Los detalles como ejercicio.

Casos especiales adicionales de ecuaciones lineales de términos múltiples se analizan en Nkamnang (Shukla & Prajapati, 2007), las fórmulas resultantes suelen ser muy complicadas y no los repetimos aquí. Se pueden encontrar más resultados, por ejemplo, en el artículo de Luchko y Gorenflo (1999); Lubich (2006), y en el libro de Podlubny (Olmstead & Handelsman, 1976) (capítulos 5 y 6).

Cerremos este capítulo con una breve perspectiva: Observando una ecuación lineal general de términos múltiples.

$$\sum_{k=0}^{\infty} \alpha_j(x) D_{*0}^{n_j} y(x) = g(x)$$

con algunas funciones dadas $\alpha_0, \alpha_1, \dots, \alpha_k$ y g y $0 = n_0 < n_1 < \dots < n_k$ vemos que podemos reescribir una ecuación como

$$\int_0^{n_k} \alpha_j(x) D_{*0}^{n_j} y(x) du(j) = g(x)$$

donde μ es una función escalonada con saltos de altura unitaria en los puntos si permitimos μ es una medida más general, entonces esto nos lleva inmediatamente a la llamada ecuaciones diferenciales de orden distribuido. Primeros pasos en este tema, tanto en relación con Métodos de solución analítica y modelado matemático de problemas físicos. con esas ecuaciones, han sido descritas por Bagley y Torvik (2000); (Torvik & Bagley, 1984). En los artículos (Diethelm, s. f.); (Atanackovic et al., 2009); (Cuesta-Montero & Finat, 2003), se puede encontrar un marco para un enfoque numérico. sin embargo, la discusión de este tema está más allá del alcance de este libro.

Referencias

- Adomian, G. (1994). *Solving frontier problems of physics: The decomposition method*. Springer. <https://doi.org/10.1007/978-94-015-8289-6>
- Agarwal, R. P., Benchohra, M., & Hamani, S. (2008). A survey on existence results for boundary value problems of nonlinear fractional differential equations and inclusions. *Acta Applicandae Mathematicae*, 109(3), 973-1033. <https://doi.org/10.1007/s10440-008-9356-6>
- Ahmad, W. M., & El-Khazali, R. (2007). Fractional-order dynamical models of love. *Chaos, Solitons and Fractals*, 33(4), 1367-1375. <https://doi.org/10.1016/j.chaos.2006.01.098>
- Atanackovic, T. M., Budincevic, M., & Pilipovic, S. (2009). Numerical analysis for distributed-order differential equations. *Journal of Computational and Applied Mathematics*, 225(1), 96-104. <https://doi.org/10.1088/0305-4470/38/30/006>
- Audounet, J., Giovangigli, V., & Roquejoffre, J. M. (1998). A threshold phenomenon in the propagation of a point source initiated flame. *Physica D: Nonlinear Phenomena*, 121(3-4), 295-316. [https://doi.org/10.1016/S0167-2789\(98\)00153-5](https://doi.org/10.1016/S0167-2789(98)00153-5)
- Audounet, J., & Roquejoffre, J.-M. (1998). An asymptotic fractional differential model of spherical flame. *ESAIM: Proceedings*, 5, 15-27. <https://doi.org/10.1051/proc:1998009>
- Bagley, R. L., & Calico, R. A. (1991). Fractional order state equations for the control of viscoelasticallydamped structures. *Journal of Guidance, Control, and Dynamics*, 14(2), 304-311. <https://doi.org/10.2514/3.20641>
- Bai, J., & Feng, X. C. (2007). Fractional-order anisotropic diffusion for image denoising. *IEEE Transactions on Image Processing*, 16(10), 2492-2502. <https://doi.org/10.1109/TIP.2007.904971>

- Baleanu, D., & Tenreiro Machado, J. A. (2009). Fractional differentiation and its applications (FDA08). *Physica Scripta*, (136). <https://doi.org/10.1088/0031-8949/2008/T136/011001>
- Basset, A. B. (1888). On the motion of a sphere in a viscous liquid. *Proceedings of the Royal Society of London*, 43(258-265), 174-175. <https://doi.org/10.1098/rspl.1887.0120>
- Benchohra, M., Hamani, S., & Ntouyas, S. K. (2008). Boundary value problems for differential equations with fractional order. *Surveys in Mathematics and its Applications*, 3, 1-12. <http://www.utgjiu.ro/math/sma>
- Benson, D. (1998). *The fractional advection-dispersion equation: Development and application* [Doctoral dissertation, University of Nevada, Reno].
- Bonilla, B., Rivero, M., & Trujillo, J. J. (2007). On systems of linear fractional differential equations with constant coefficients. *Applied Mathematics and Computation*, 187(1), 68-78. <https://doi.org/10.1016/j.amc.2006.08.104>
- Brunner, H., & van der Houwen, P. J. (1986). *The numerical solution of Volterra equations*. North-Holland.
- Brunner, H., Pedaş, A., & Vainikko, G. (1999). The piecewise polynomial collocation method for nonlinear weakly singular Volterra equations. *Mathematics of Computation*, 68(227), 1079-1095. <https://doi.org/10.1090/s0025-5718-99-01073-x>
- Caponetto, R., Dongola, G., Fortuna, L., & Petráš, I. (2010). *Fractional order systems: Modeling and control applications*. World Scientific. <https://doi.org/10.1142/7709>
- Caputo, M. (1967). Linear models of dissipation whose Q is almost frequency independent—II. *Geophysical Journal International*, 13(5), 529-539. <https://doi.org/10.1111/j.1365-246X.1967.tb02303.x>

- Caputo, M., & Mainardi, F. (1971). A new dissipation model based on memory mechanism. *Pure and Applied Geophysics*, 91(1), 134-147. <https://doi.org/10.1007/BF00879562>
- Caputo, M., & Mainardi, F. (2008). Linear models of dissipation in anelastic solids. *La Rivista del Nuovo Cimento (1971-1977)*, 1(2), 161-198. <https://doi.org/10.1007/BF02820620>
- Chatterjee, A. (2005). Statistical origins of fractional derivatives in viscoelasticity. *Journal of Sound and Vibration*, 284(3-5), 1239-1245. <https://doi.org/10.1016/j.jsv.2004.09.019>
- Chen, J., Lundberg, K. H., Davison, D. E., & Bernstein, D. S. (2007). The final value theorem revisited: Infinite limits and irrational functions. *IEEE Control Systems*, 27(3), 97-99. <https://doi.org/10.1109/MCS.2007.365008>
- Chern, J.-T. (1993). *Finite element modeling of viscoelastic materials on the theory of fractional calculus* [Doctoral dissertation, University of Akron].
- Collatz, L. (2013). *Funktionalanalysis und numerische Mathematik*. Springer-Verlag.
- Corduneanu, C. (1994). *Principles of differential and integral equations*. Chelsea Publishing.
- Cuesta-Montero, E., & Finat, J. (2003). *Image processing by means of a linear integro-differential equation. Proceedings of the IASTED International Conference on Signal and Image Processing*.
- de Hoog, F., & Weiss, R. (1973). Asymptotic expansions for product integration. *Mathematics of Computation*, 27(122), 295-306. <https://doi.org/10.2307/2005616>
- Dehghan, M., & Tatari, M. (2007). On the convergence of He's variational iteration method. *Journal of Computational and Applied Mathematics*, 207(1), 121-128. <https://doi.org/10.1088/0031-8949/74/3/003>

- Deng, W. (2007). Short memory principle and a predictor–corrector approach for fractional differential equations. *Journal of Computational and Applied Mathematics*, 206(1), 174-188. <https://doi.org/10.1016/j.cam.2006.06.008>
- Diethelm, K. (s.f.). Smoothness properties of solutions of Caputo-type fractional differential equations. *Fractional Calculus and Applied Analysis*, 10(2), 151-160.
- Diethelm, K. (1997). An algorithm for the numerical solution of differential equations of fractional order. *Electronic Transactions on Numerical Analysis*, 5, 1-6.
- Diethelm, K. (2001). *Predictor-corrector strategies for single and multi-term fractional differential equations. Proceedings of the Hellenic-European Conference on Computer Mathematics and its Applications.*
- Diethelm, K. (2003). Efficient solution of multi-term fractional differential equations using P(EC)mE methods. *Computing*, 71(4), 305-319. <https://doi.org/10.1007/s00607-003-0033-3>
- Diethelm, K. (2008a). An investigation of some nonclassical methods for the numerical approximation of Caputo-type fractional derivatives. *Numerical Algorithms*, 47(4), 361-390. <https://doi.org/10.1007/s11075-008-9193-8>
- Diethelm, K. (2008b). On the separation of solutions of fractional differential equations. *Fractional Calculus and Applied Analysis*, 11(3), 259-268.
- Diethelm, K. (2009). An improvement of a nonclassical numerical method for the computation of fractional derivatives. *Journal of Vibration and Acoustics*, 131(1). <https://doi.org/10.1115/1.2981167>
- Diethelm, K. (2010). Existence and uniqueness results for Riemann-Liouville fractional differential equations. *Lecture Notes in Mathematics*, 2004, 77-83. https://doi.org/10.1007/978-3-642-14574-2_5

- Diethelm, K., & Ford, N. J. (2001). Numerical solution methods for distributed order differential equations. *Fractional Calculus and Applied Analysis*, 4(4), 531-542.
- Diethelm, K., & Ford, N. J. (2002a). Analysis of fractional differential equations. *Journal of Mathematical Analysis and Applications*, 265(2), 229-248. <https://doi.org/10.1006/jmaa.2000.7194>
- Diethelm, K., & Ford, N. J. (2002b). Numerical solution of the Bagley-Torvik equation. *BIT Numerical Mathematics*, 42(3), 490-507. <https://doi.org/10.1023/a:1021973025166>
- Diethelm, K., & Ford, N. J. (2004). Multi-order fractional differential equations and their numerical solution. *Applied Mathematics and Computation*, 154(3), 621-640. [https://doi.org/10.1016/S0096-3003\(03\)00739-2](https://doi.org/10.1016/S0096-3003(03)00739-2)
- Diethelm, K., & Ford, N. J. (2012). Volterra integral equations and fractional calculus: Do neighboring solutions intersect? *Journal of Integral Equations and Applications*, 24(1), 25-37. <https://doi.org/10.1216/JIE-2012-24-1-25>
- Diethelm, K., Ford, N. J., & Freed, A. D. (2002). A predictor-corrector approach for the numerical solution of fractional differential equations. *Nonlinear Dynamics*, 29(1-4), 3-22. <https://doi.org/10.1023/A:1016592219341>
- Diethelm, K., Ford, N. J., & Freed, A. D. (2004). Detailed error analysis for a fractional Adams method. *Numerical Algorithms*, 36(1), 31-52. <https://doi.org/10.1023/B:NUMA.0000027736.85078.be>
- Diethelm, K., & Freed, A. D. (1999a). On the solution of nonlinear fractional-order differential equations used in the modeling of viscoplasticity. In F. Keil, W. Mackens, H. Voß, & J. Werther, (eds.). *Scientific computing in chemical engineering II* (pp. 217-224). Springer. https://doi.org/10.1007/978-3-642-60185-9_24

- Diethelm, K., & Freed, A. D. (1999b). *The FracPECE subroutine for the numerical solution of differential equations of fractional order*. Software y documentación.
- Diethelm, K., & Luchko, Y. (2004). Numerical solution of linear multi-term initial value problems of fractional order. *Journal of Computational Analysis and Applications*, 6(3), 243-263.
- Diethelm, K., & Walz, G. (1997). Numerical solution of fractional order differential equations by extrapolation. *Numerical Algorithms*, 16(3), 231-253. <https://doi.org/10.1023/a:1019147432240>
- Doetsch, G. (1971). *Anleitung zum praktischen Gebrauch der Laplace-Transformation und der Z-Transformation*. Springer-Verlag.
- Dokoumetzidis, A., Magin, R., & Macheras, P. (2010). A commentary on fractionalization of multi-compartmental models. *Journal of Pharmacokinetics and Pharmacodynamics*, 37(2), 203-207. <https://doi.org/10.1007/s10928-010-9153-5>
- Dzherbashian, M. M., & Nersesian, A. B. (2021). Fractional derivatives and Cauchy problem for differential equations of fractional order. *Fractional Calculus and Applied Analysis*, 23(6), 1810-1836. <https://doi.org/10.1515/fca-2020-0090>
- Edwards, J. T., Ford, N. J., & Simpson, C. A. (2002). The numerical solution of linear multi-term fractional differential equations: Systems of equations. *Journal of Computational and Applied Mathematics*, 148(2), 401-418. [https://doi.org/10.1016/S0377-0427\(02\)00558-7](https://doi.org/10.1016/S0377-0427(02)00558-7)
- Ford, N. J., & Connolly, J. A. (2006). Comparison of numerical methods for fractional differential equations. *Communications on Pure and Applied Analysis*, 5(2), 289-307. <https://doi.org/10.3934/cpaa.2006.5.289>

- Ford, N. J., & Connolly, J. A. (2009). Systems-based decomposition schemes for the approximate solution of multi-term fractional differential equations. *Journal of Computational and Applied Mathematics*, 229(2), 382-391. <https://doi.org/10.1016/j.cam.2008.04.003>
- Ford, N. J., & Simpson, C. (2000). *The approximate solution of fractional differential equations of order greater than 1. Proceedings of the 16th IMACS World Congress on Scientific Computation.*
- Freed, A. D., & Diethelm, K. (2006). Fractional calculus in biomechanics: A 3D viscoelastic model using regularized fractional derivative kernels with application to the human calcaneal fat pad. *Biomechanics and Modeling in Mechanobiology*, 5(4), 203-215. <https://doi.org/10.1007/s10237-005-0011-0>
- Freed, A., & Diethelm, K. (2002). *Fractional-order viscoelasticity (FOV): Constitutive development using the fractional calculus: First annual report.* (NASA/TM-2002-211914). NASA Glenn Research Center.
- Galeone, L., & Garrappa, R. (2008). Fractional Adams–Moulton methods. *Mathematics and Computers in Simulation*, 79(4), 1358-1367. <https://doi.org/10.1016/j.matcom.2008.03.008>
- Gaul, L., Klein, P., & Kemple, S. (1991). Damping description involving fractional operators. *Mechanical Systems and Signal Processing*, 5(2), 81-88. [https://doi.org/10.1016/0888-3270\(91\)90016-X](https://doi.org/10.1016/0888-3270(91)90016-X)
- Glöckle, W. G., & Nonnenmacher, T. F. (1995). A fractional calculus approach to self-similar protein dynamics. *Biophysical Journal*, 68(1), 46-53. [https://doi.org/10.1016/S0006-3495\(95\)80157-8](https://doi.org/10.1016/S0006-3495(95)80157-8)
- Gluskin, E. (2003). Let us teach this generalization of the final-value theorem. *European Journal of Physics*, 24(6), 591-597. <https://doi.org/10.1088/0143-0807/24/6/005>

- Gorenflo, R., De Fabritiis, G., & Mainardi, F. (1999). Discrete random walk models for symmetric Levy-Feller diffusion processes. *Physica A: Statistical Mechanics and its Applications*, 269(1), 79-89. [https://doi.org/10.1016/S0378-4371\(99\)00082-5](https://doi.org/10.1016/S0378-4371(99)00082-5)
- Gorenflo, R., & Mainardi, F. (1996). Fractional oscillations and Mittag-Leffler functions. *Proceedings of the International Workshop on the Recent Advances in Applied Mathematics*, 193-208.
- Gorenflo, R., & Mainardi, F. (2008). *Fractional calculus: Integral and differential equations of fractional order*. arXiv. <http://arxiv.org/abs/0805.3823>
- Gorenflo, R., Loutchko, J., & Luchko, Y. (2002). Computation of the Mittag-Leffler function and its derivatives. *Fractional Calculus and Applied Analysis*, 5(4), 491-518.
- Gorenflo, R., & Rutman, R. (1995). On ultraslow and intermediate processes. In P. Rusev, I. Dimovski, & V. Kiryakova, (eds.). *Transform methods and special functions* (pp. 61-81). Science Culture Technology.
- Gorenflo, R., & Vivoli, A. (2003). Fully discrete random walks for space-time fractional diffusion equations. *Signal Processing*, 83(11), 2411-2420. [https://doi.org/10.1016/S0165-1684\(03\)00193-2](https://doi.org/10.1016/S0165-1684(03)00193-2)
- Grunwald, A. K. (1867). Ueber "begrenzte" Derivationen und deren Anwendung. *Zeitschrift für Mathematik und Physik*, 12, 441-480.
- Hadamard, J. (1952). *Lectures on Cauchy's problem in linear partial differential equations*. Dover Publications.
- Hadid, S. B., & Momani, S. M. (1996). Some existence theorems on differential equations of generalized order through a fixed-point theorem. *Journal of the Faculty of Science, United Arab Emirates University*, 9(1), 1-12.

- Haubold, H. J., Mathai, A. M., & Saxena, R. K. (2009). Mittag-Leffler functions and their applications. *Journal of Applied Mathematics*, 2011. <http://arxiv.org/abs/0909.0230>
- Hewitt, E. (1964). *The gamma function*. In *Athena selected topics in mathematics*. Holt, Rinehart and Winston.
- Hilfer, R. (Ed.). (2000). *Applications of fractional calculus in physics*. World Scientific. <https://doi.org/10.1142/3779>
- Hille, E. (1969). *Lectures on ordinary differential equations*. Addison-Wesley.
- Hille, E. (1976). *Ordinary differential equations in the complex domain*. John Wiley & Sons.
- Jachymski, J., Jóźwik, I., & Terepeta, M. (2024). The Banach fixed point theorem: Selected topics from its hundred-year history. *Revista de la Real Academia de Ciencias Exactas, Físicas y Naturales. Serie A. Matemáticas*, 118(4). <https://doi.org/10.1007/s13398-024-01636-6>
- Jacob, N., & Kr̃ageloh, A. M. (2002). The Caputo derivative, Feller semigroups, and the fractional power of the first order derivative on $C_\infty(\mathbb{R}^+)$. *Fractional Calculus and Applied Analysis*, 5(4), 395-410.
- Jin, B. (2021). *Fractional differential equations: An approach via fractional derivatives*. Springer. <https://doi.org/10.1007/978-3-030-76034-9>
- Johnson, W. P. (2002). The curious history of Faà di Bruno's formula. *American Mathematical Monthly*, 109(3), 217-234. <https://doi.org/10.2307/2695352>
- Joulin, G. (1985). Point-source initiation of lean spherical flames of light reactants: An asymptotic theory. *Combustion Science and Technology*, 43(1-2), 99-113. <https://doi.org/10.1080/00102208508947000>

- Kilbas, A. A., & Trujillo, J. J. (2001). Differential equations of fractional order: Methods, results and problem—I. *Applicable Analysis*, 78(1-2), 153-192. <https://doi.org/10.1080/00036810108840931>
- Kilbas, A. A., & Trujillo, J. J. (2002). Differential equations of fractional order: Methods, results and problems. II. *Applicable Analysis*, 81(2), 435-493. <https://doi.org/10.1080/0003681021000022032>
- Kilbas, A. A., Srivastava, H. M., & Trujillo, J. J. (2006). *Theory and applications of fractional differential equations*. Elsevier.
- Kiryakova, V. (2010a). The multi-index Mittag-Leffler functions as an important class of special functions of fractional calculus. *Computers & Mathematics with Applications*, 59(5), 1885-1895. <https://doi.org/10.1016/j.camwa.2009.08.025>
- Kiryakova, V. (2010b). The special functions of fractional calculus as generalized fractional calculus operators of some basic functions. *Computers & Mathematics with Applications*, 59(3), 1128-1141. <https://doi.org/10.1016/j.camwa.2009.05.014>
- Klages, R., Radons, G., & Sokolov, I. M. (Eds.). (2008). *Anomalous transport: Foundations and applications*. Wiley-VCH.
- Kochubei, A. N. (2008). Fractional differential equations: entire solutions, regular and irregular singularities. *Fractional Calculus and Applied Analysis*, 11(3), 259-268.
- Koeller, R. C. (1986). Polynomial operators, Stieltjes convolution, and fractional calculus in hereditary mechanics. *Acta Mechanica*, 58(3-4), 251-264. <https://doi.org/10.1007/BF01176603>
- Liang, S., & Jeffrey, D. J. (2009). Comparison of homotopy analysis method and homotopy perturbation method through an evolution equation. *Communications in Nonlinear Science and Numerical Simulation*, 14(12), 4057-4064. <https://doi.org/10.1016/j.cnsns.2009.02.016>

- Liao, S. (2004). *Beyond perturbation: Introduction to the homotopy analysis method*. Chapman & Hall/CRC.
- Linz, P. (1985). *Analytical and numerical methods for Volterra equations*. SIAM. <https://doi.org/10.1137/1.9781611970852>
- Lu, J. G. (2006). Chaotic dynamics of the fractional-order Lü system and its synchronization. *Physics Letters A*, 354(4), 305-311. <https://doi.org/10.1016/j.physleta.2006.01.068>
- Lubich, C. (1983). Runge-Kutta theory for Volterra and Abel integral equations of the second kind. *Mathematics of Computation*, 41(163), 87-102.
- Lubich, C. (1986). Discretized fractional calculus. *SIAM Journal on Mathematical Analysis*, 17(3), 704-719. <https://doi.org/10.1137/0517050>
- Luchko, Y., & Gorenflo, R. (1999). An operational method for solving fractional differential equations with the Caputo derivatives. *Acta Mathematica Vietnamica*, 24(2), 207-233.
- Magin, R. (2006). *Fractional calculus in bioengineering*. Begell House.
- Mainardi, F. (2010). *Fractional calculus and waves in linear viscoelasticity: An introduction to mathematical models*. Imperial College Press. <https://doi.org/10.1142/P614>
- Mainardi, F. (2012). *Fractional calculus: Some basic problems in continuum and statistical mechanics*. arXiv. <http://arxiv.org/abs/1201.0863>
- Mainardi, F., & Gorenflo, R. (2000). On Mittag-Leffler-type functions in fractional evolution processes. *Journal of Computational and Applied Mathematics*, 118(1-2), 283-299. [https://doi.org/10.1016/S0377-0427\(00\)00294-6](https://doi.org/10.1016/S0377-0427(00)00294-6)

- Mainardi, F., Raberto, M., Gorenflo, R., & Scalas, E. (2000). Fractional calculus and continuous-time finance II: The waiting-time distribution. *Physica A: Statistical Mechanics and its Applications*, 287(3-4), 468-481. [https://doi.org/10.1016/S0378-4371\(00\)00386-1](https://doi.org/10.1016/S0378-4371(00)00386-1)
- Matignon, D. (1994). *Représentations en variables d'état de modèles de guides d'ondes avec dérivation fractionnaire* [Doctoral dissertation, Université Paris XI].
- Matignon, D. (1998). Stability properties for generalized fractional differential systems. *ESAIM: Proceedings*, 5, 145-158. <https://doi.org/10.1051/proc:1998004>
- Metzler, R., Schick, W., Kilian, H. G., & Nonnenmacher, T. F. (1995). Relaxation in filled polymers: A fractional calculus approach. *The Journal of Chemical Physics*, 103(16), 7180-7186. <https://doi.org/10.1063/1.470346>
- Miller, K. S., & Ross, B. (1993). *An introduction to the fractional calculus and fractional differential equations*. John Wiley & Sons.
- Momani, S. M. (2000). Local and global uniqueness theorems on differential equations of non-integer order via Bihari's and Gronwall's inequalities. *Revista Técnica de la Facultad de Ingeniería. Universidad del Zulia*, 23(1), 66-69.
- Momani, S., & Odibat, Z. (2007). Homotopy perturbation method for nonlinear partial differential equations of fractional order. *Physics Letters A*, 365(5-6), 345-350. <https://doi.org/10.1016/j.physleta.2007.01.046>
- Nkamnang, A. R. (1999). *Diskretisierung von mehrgliedrigen Abelschen Integralgleichungen und gewöhnlichen Differentialgleichungen gebrochener Ordnung* [Doctoral dissertation, Freie Universität Berlin]. <https://doi.org/10.17169/REFUBIUM-7701>

- Nonnenmacher, T. F., & Metzler, R. (1995). On the Riemann-Liouville fractional calculus and some recent applications. *Fractals*, 3(3), 557-566. <https://doi.org/10.1142/s0218348x95000497>
- Nutting, P. G. (1943). A general stress-strain-time formula. *Journal of the Franklin Institute*, 235(5), 513-524. [https://doi.org/10.1016/s0016-0032\(43\)91483-8](https://doi.org/10.1016/s0016-0032(43)91483-8)
- Odibat, Z. M. (2010). Analytic study on linear systems of fractional differential equations. *Computers & Mathematics with Applications*, 59(3), 1171-1183. <https://doi.org/10.1016/j.camwa.2009.06.035>
- Odibat, Z., Momani, S., & Erturk, V. S. (2008). Generalized differential transform method: Application to differential equations of fractional order. *Applied Mathematics and Computation*, 197(2), 467-477. <https://doi.org/10.1016/j.amc.2007.07.068>
- Oldham, K. B., & Spanier, J. (1974). *The fractional calculus: Theory and applications of differentiation and integration to arbitrary order*. Academic Press.
- Olmstead, W. E., & Handelsman, R. A. (1976). Diffusion in a semi-infinite region with nonlinear surface dissipation. *SIAM Review*, 18(2), 275-291. <https://www.jstor.org/stable/2028788>
- Podlubny, I. (1994). *Fractional-order systems and fractional-order controllers* (UEF-03-94). Slovak Academy of Sciences.
- Podlubny, I. (1995). Application of fractional-order derivatives to calculation of heat load intensity change in blast furnace walls. *Transactions of Technical University of Košice*, 5(2), 137-144.
- Podlubny, I. (1999). *Fractional differential equations: An introduction to fractional derivatives, fractional differential equations, to methods of their solution and some of their applications*. Academic Press.

- Podlubny, I. (2001). *Geometric and physical interpretation of fractional integration and fractional differentiation*. arXiv.
- Rasof, B. (1962). The initial- and final-value theorems in Laplace transform theory. *Journal of the Franklin Institute*, 274(3), 165-177. [https://doi.org/10.1016/0016-0032\(62\)90939-0](https://doi.org/10.1016/0016-0032(62)90939-0)
- Rèpaci, A. (1990). Nonlinear dynamical systems: On the accuracy of Adomian's decomposition method. *Applied Mathematics Letters*, 3(4), 35-39. [https://doi.org/10.1016/0893-9659\(90\)90042-A](https://doi.org/10.1016/0893-9659(90)90042-A)
- Ross, B. (1977). The development of fractional calculus 1695–1900. *Historia Mathematica*, 4(1), 75-89. [https://doi.org/10.1016/0315-0860\(77\)90039-8](https://doi.org/10.1016/0315-0860(77)90039-8)
- Rossikhin, Y. A. (2010). Reflections on two parallel ways in the progress of fractional calculus in mechanics of solids. *Applied Mechanics Reviews*, 63(1). <https://doi.org/10.1115/1.4000246>
- Rossikhin, Yu. A., & Shitikova, M. V. (2007). Comparative analysis of viscoelastic models involving fractional derivatives of different orders. *Fractional Calculus and Applied Analysis*, 10(2), 111-121. <https://doi.org/10.31857/s2308114724020025>
- Sabatier, J., Agrawal, O. P., & Tenreiro Machado, J. A. (Eds.). (2007). *Advances in fractional calculus: Theoretical developments and applications in physics and engineering*. Springer. <https://doi.org/10.1007/978-1-4020-6042-7>
- Samko, S. G., Kilbas, A. A., & Marichev, O. I. (1993). *Fractional integrals and derivatives: Theory and applications*. Gordon and Breach Science Publishers.
- Scalas, E., Gorenflo, R., & Mainardi, F. (2000). Fractional calculus and continuous-time finance. *Physica A: Statistical Mechanics and its Applications*, 284(1-4), 376-384. [https://doi.org/10.1016/S0378-4371\(00\)00255-7](https://doi.org/10.1016/S0378-4371(00)00255-7)

- Seybold, H. J., & Hilfer, R. (2005). *Numerical results for the generalized Mittag-Leffler function*. Preprint.
- Seybold, H. J., & Hilfer, R. (2008). Numerical algorithm for calculating the generalized Mittag-Leffler function. *SIAM Journal on Numerical Analysis*, 47(1), 69-88. <https://doi.org/10.1137/070700280>
- Shaw, S., Warby, M., & Whiteman, J. (1997). *A comparison of hereditary integral and internal variable approaches to numerical linear solid viscoelasticity*. BICOM, Brunel University.
- Shukla, A. K., & Prajapati, J. C. (2007). On a generalization of Mittag-Leffler function and its properties. *Journal of Mathematical Analysis and Applications*, 336(2), 797-811. <https://doi.org/10.1016/j.jmaa.2007.03.018>
- Singh, S. J., & Chatterjee, A. (2006). Galerkin projections and finite elements for fractional order derivatives. *Nonlinear Dynamics*, 45(1), 183-206. <https://doi.org/10.1007/s11071-005-9002-z>
- Song, L., Xu, S., & Yang, J. (2010). Dynamical models of happiness with fractional order. *Communications in Nonlinear Science and Numerical Simulation*, 15(3), 616-628. <https://doi.org/10.1016/j.cnsns.2009.04.029>
- Syedlyetskii, A. M. (1994). Asymptotic formulae for zeros of a function of Mittag-Leffler's type. *Analysis Mathematica*, 20(2), 117-132. <https://doi.org/10.1007/BF01908643>
- Taş, K., Tenreiro Machado, J. A., & Baleanu, D. (Eds.). (2007). *Mathematical methods in engineering*. Springer.
- Torvik, P. J., & Bagley, R. L. (1984). On the appearance of the fractional derivative in the behavior of real materials. *Journal of Applied Mechanics*, 51(2), 294-298. <https://doi.org/10.1115/1.3167615>

- Trinks, C., & Ruge, P. (2002). Treatment of dynamic systems with fractional derivatives without evaluating memory-integrals. *Computational Mechanics*, 29(6), 471-476. <https://doi.org/10.1007/s00466-002-0356-5>
- Verotta, D. (2010). Fractional compartmental models and multi-term Mittag-Leffler response functions. *Journal of Pharmacokinetics and Pharmacodynamics*, 37(2), 209-215. <https://doi.org/10.1007/s10928-010-9155-3>
- Yuan, L., & Agrawal, O. P. (2002). A numerical scheme for dynamic systems containing fractional derivatives. *Journal of Vibration and Acoustics*, 124(2), 321-324. <https://doi.org/10.1115/1.1448322>

Apendice A

Lista de Símbolos

En este apéndice recopilamos una lista de todos los símbolos que se han utilizado. Si es necesario, también se dan referencias cruzadas.

Funciones

$\| \cdot \|_{\infty}$ norma de Chebyshev, $\|f\|_{\infty} = \sup_{a \leq x \leq b} |f(x)|$

$\| \cdot \|_p$ norma L_p , $(1 \leq p < \infty)$; $\|f\|_p = \left(\int_a^b |f(x)|^p dx \right)^{1/p}$

$[\cdot]$ Función suelo, $[x] = \max\{z \in \mathbb{Z}, z \leq x\}$

$[\cdot]$ Función techo $[\cdot] = \min\{z \in \mathbb{Z}, z \geq x\}$

$\binom{n}{k}$ coeficiente binomial $\binom{n}{k} = \frac{n(n-1)(n-2)\dots(n-k+1)}{k!}$, para $n \in \mathbb{R}$ y para $n \in \mathbb{N}_0$

$B_N[f]$ Enésimo polinomio de Bernstein para la función f ,

$B_N[f] = \sum_{k=0}^N \binom{N}{k} t^k (1-t)^{N-k} f(k/N)$ ver apéndice D5

B Función Beta de Euler $B(x, y) = \frac{\Gamma(x)\Gamma(y)}{\Gamma(x+y)}$ ver apéndice D1.

Γ función gamma de Euler, $\Gamma(x) = \int_0^{\infty} t^{x-1} e^{-t} dt$ ver definición 2.2

Ψ función Digama, $\Psi(x) = \frac{\Gamma'(x)}{\Gamma(x)}$

E_n Función de Mittag-Leffler de orden n , $E_n(x) = \sum_{j=0}^{\infty} \frac{x^j}{\Gamma(jn+1)}$ ver definición 4.1

E_{n_1, n_2} función Mittag-Leffler de dos parámetros, $E_{n_1, n_2}(x) = \sum_{j=0}^{\infty} \frac{x^j}{\Gamma(jn_1, n_2+1)}$ ver definición 4.2

${}_2F_1$ Función hipergeométrica confluyente de Kummer

$${}_2F_1(a, b, c) = \frac{\Gamma(c)}{\Gamma(a)\Gamma(b)} \sum_{k=0}^{\infty} \frac{\Gamma(a+k)\Gamma(b+k)}{\Gamma(c+k)k!} z^k, a \in \mathbb{R}, -b \notin \mathbb{N}_0$$

$\mathcal{O}$, Símbolos de Landau

$T_j[f; a]$, Polinomio de Taylor de grado j para la función f centrado en el punto a

Conjuntos

$A^n, A^n[a, b]$, Conjunto de funciones con derivada absolutamente continua de

orden $n-1$ (ver Definición 1.4).

$C; C[a, b]$, Conjunto de funciones continuas (de Definición 1.2)

$C^k; C^k[a, b]$ Conjunto de funciones con k -ésima derivada continua (ver Definición 1.2)

$H^*; H^*[a, b]$ Cfr. Definición 1.3

$H_\mu; H_\mu[a, b]$ Espacio de Holder (ver Definición 1.2),

$L_p; L_p[a, b]$ Espacio de Lebesgue (cf. Definición 1.2)

$\mathbb{N} = \{1, 2, 3, \dots\}$ el conjunto de los números naturales

$\mathbb{N}_0 = \mathbb{N} \cup \{0\}$

$\mathbb{R}$, El conjunto de los números reales.

$\mathbb{R}_0 = \{x \in \mathbb{R}, x > 0\}$, el conjunto de números reales estrictamente positivos

$\mathbb{Z} = \{0, \pm 1, \pm 2, \dots\}$, el conjunto de números enteros

Operadores

Δ_h^n Diferencia finita de orden n ; cf. (2.11)

D Operador diferencial, $Df(x) = f'(x)$ (ver Definición 1.1)

$D^n = n \in \mathbb{N}$: iteración n veces del operador diferencial D (de Definición 1.1)

$D_a^n, n \in \mathbb{R}_+$: operador diferencial fraccionario de Riemann-Liouville (de Definición 2.2)

$\tilde{D}_a^n, n \in \mathbb{R}_+$: operador diferencial fraccionario de Gruunwald-Letnikov (cf. Definición 2.3)

$\widehat{D}_a^n$: cf. Definición 3.1

$D_{*a}^n, n \in \mathbb{R}_+$: Operador diferencial fraccionario de Caputo (cf. Definición 3.2)

$\mathcal{D}_n^n, n \in \mathbb{R}_+$: Operador Gel'fond-Leont'ev (de Definición 3.3)

I operador de identidad

J_a : Operador integral, $J_a f(x) = \int_a^x f(t)dt$: (cf. Definición 1.

$J_a^n = n \in \mathbb{N}$, : iteración n veces del operador integral J_a (cf. Definición 1.1)

$n \in \mathbb{R}_+ \setminus \mathbb{N}$: Operador integral fraccionario de Riemann-Liouville (cf. Definición 2.1)

$n = 0$: operador de identidad (cf. Definición 2.1)

$\tilde{J}_a^n, n \in \mathbb{R}_+$: operador integral fraccionario de Grünwald-Letnikov (ver Definición 2.4)

$\mathcal{L}$ Operador de transformada de Laplace (cf. Apéndice D.3)

ω Módulo de continuidad de la función $g : [a,b] \rightarrow \mathbb{R}$,

$\omega(g, h); = \sup\{|g(y_1) - g(y_2)|, y_1, y_2 \in [a, b], |y_1 - y_2| \leq h\} ..$

Otros símbolos

$\sim, \sim a_j \sim a_j \Leftrightarrow \exists A, B > 0, \exists j_0 \in \mathbb{N} \forall j \geq j_0 : A \leq |a_j / b_j| \leq B;;$

Observaciones

1. Las series de potencias para ambos tipos de funciones de Mittag-Leffler convergen en todo el plano complejo (cf. Teorema 3.1).
2. La serie de potencias de la función hipergeométrica confluyente de Kummer converge en todo el plano complejo.
3. Para la función hipergeométrica de Gauss, la serie de potencias converge para todo complejo z con $|z| < 1$ y puede ampliarse analíticamente a todo el plano complejo con una rama cortada a lo largo del eje real positivo de $+1$ a $+\infty$. (En las fórmulas de Apéndice B, necesitamos evaluar esta función para $z < 0$, por lo que el corte de rama para $z \geq 1$ no da problemas.

A. Preliminares Matemáticos

En este apéndice, recordamos varios espacios de funciones, por ejemplo, espacios AC, espacios Hölder y espacios de Sobolev, dos transformadas integrales y varios teoremas de punto fijo.

A.1 Espacios AC y Espacios Hölder

A.1.1 Espacios AC

El espacio de funciones absolutamente continuas (AC) se usa ampliamente en el estudio de cálculo fraccionario. En todo momento, $D = (a, b)$, con $a < b$, es un intervalo acotado abierto.

Definición A.1

Una función $v: D \rightarrow \mathbb{R}$ se llama absolutamente continua en D , si para cualquier $\varepsilon > 0$, existe un $\delta > 0$ tal que para cualquier conjunto finito de intervalos disjuntos por pares $[a_k, b_k] \subset \bar{D}$; $k = 1, 2, 3, \dots, n$, tal que $\sum_{k=1}^n (b_k - a_k) < \delta$, ahí se cumple la desigualdad

$$\sum_{k=1}^n |v(b_k) - v(a_k)| < \varepsilon$$

El espacio de estas funciones se denota por $AC(\bar{D})$.

Se sabe que el espacio $AC(\bar{D})$ coincide con el espacio de primitivas de Funciones integrables de Lebesgue, es decir,

$$v(x) \in AC(\bar{D}) \Leftrightarrow v(x) = c + \int_a^x \phi(s) ds \quad \forall x \in \bar{D} \text{ con } \int_a^x \phi(s) ds < \infty$$

para alguna constante c . Por tanto, las funciones AC tienen una derivada sumable $v'(x)$ casi en todas partes (a.e.). Claramente, $C^1(\bar{D}) \rightarrow AC(\bar{D})$, pero lo contrario no es cierto. Para ejemplo, para $\alpha \in (0, 1)$, $v(x) = (a - x)^\alpha \in AC(\bar{D})$, $a \in AC(\bar{D})$, pero $v(x) \notin C^1(\bar{D})$. Sin embargo, si v es continua y $v'(x) \in L^1(D)$ existe, es decir, v no es necesaria $AC(\bar{D})$. un contraejemplo es la función singular v de Lebesgue (también conocida como función de Cantor-Vitali o función de la escalera del diablo) que es continua en $[0, 1]$ y tiene una derivada cero a.e., pero no es $AC(\bar{D})$, de hecho, $v(0) = 0$, $v(1) = 1$, por lo tanto $v(1) - v(0) \neq \int_0^1 v'(s) ds$. Los siguientes hechos son bien conocidos en un intervalo acotado: $C^1(\bar{D}) \subset C^{0,1}(\bar{D}) \subset$ diferencial a.e., y $AC(\bar{D}) \subset$ uniformemente continua $\subset C(\bar{D})$. También se sabe que, un intervalo de forma acotada, la suma y el producto puntual de funciones en $AC(\bar{D})$ pertenece a $AC(\bar{D})$, y si $v \in AC(\bar{D})$ y $f \in C^{0,1}(\bar{D})$, entonces $f \circ v \in AC(\bar{D})$. Sin embargo, la composición de dos funciones $AC(\bar{D})$ no necesita es de $AC(\bar{D})$.

Definición A.2

$AC^n(\bar{D}), n \in N$, denota el espacio de funciones $v(x)$ con continuas derivadas hasta el orden $n - 1$ en $\bar{D}$ y la $(n - 1)$ ésima derivada $v^{n-1}(x) \in AC(\bar{D})$

El espacio $AC^n(\bar{D})$ se puede caracterizar de la siguiente manera.

Teorema A.1

El espacio $AC^n(\bar{D})$, $n \in \mathbb{N}$ consta de aquellas y sólo aquellas funciones $v(x)$ que se representan en la forma $v(x) = \frac{1}{(n-1)!} \int_a^x (x-s)^{n-1} \varphi(s) ds + \sum_{j=0}^{n-1} c_j (x-a)^j$, donde $\varphi(t) \in L^1(D)$ y c_j , $j=0, 1, \dots, n-1$ son constantes arbitrarias.

A.1.2 Espacios Hölder

Sea $\Omega \subset \mathbb{R}^d$ un dominio abierto acotado, $\partial\Omega$ su límite y $\bar{\Omega}$ su clausura.

Entonces para cualquier $\gamma \in (0, 1]$, definimos el espacio de Hölder $C^{0,\gamma}(\bar{\Omega})$ por

$$C^{0,\gamma}(\bar{\Omega}) := \{v \in C(\bar{\Omega}) : |v(x) - v(y)| \leq L|x - y|^\gamma, \forall x, y \in \bar{\Omega}\}$$

donde L es una constante no negativa, equipada con la norma

$$\|v\|_{C^{0,\gamma}(\bar{\Omega})} = \|v\|_{C(\bar{\Omega})} + [v]_{C^\gamma(\bar{\Omega})}$$

Siendo $[\cdot]_{C^\gamma(\bar{\Omega})}$ la seminorma $C^\gamma(\bar{\Omega})$, definida por

$$[v]_{C^\gamma(\bar{\Omega})} = \sup_{x, y \in \bar{\Omega}, x \neq y} \frac{|v(x) - v(y)|}{|x - y|^\gamma}$$

El caso $\gamma = 1$ se llama Lipschitz, y la seminorma correspondiente es el Lipschitz

constante. Para cualquier índice múltiple

$$\ell = (\ell_1, \ell_2, \dots, \ell_d) \in \mathbb{N}_0^d, |\ell| = \sum_{i=1}^d \ell_i, y D^\ell = \frac{\partial^{|\ell|}}{\partial x_1^{\ell_1} \partial x_2^{\ell_2} \dots \partial x_d^{\ell_d}}$$

Entonces denotamos por $C^{k,\gamma}(\Omega)$, $k \geq 1$, el espacio

$$C^{k,\gamma}(\bar{\Omega}) := \{v \in C^k(\bar{\Omega}) : |D^\ell v(x) - D^\ell v(y)| \leq L|x - y|^\gamma, \forall x, y \in \bar{\Omega}, \forall \ell \in \mathbb{N}_0^d \text{ con } |\ell| = k\}$$

con la norma $\|\cdot\|_{C^{k,\gamma}(\bar{\Omega})}$ definido de manera similar por

$$\|v\|_{C^{k,\gamma}(\bar{\Omega})} = \|v\|_{C^{k,\gamma}(\bar{\Omega})} + \sum_{\ell \in \mathbb{N}_0^d, |\ell|=k} [D^\ell v]_{C^{0,\gamma}(\bar{\Omega})}$$

A menudo identificamos $C^\gamma(\bar{\Omega})$, con $\gamma > 0$, con el espacio $C^{k,\gamma}(\bar{\Omega})$, donde k es un número entero, $\gamma' \in (0,1]$ y $\gamma = k + \gamma'$. Los espacios de Hölder son espacios de Banach. Además, para cualquier $0 < \alpha < \beta \leq 1$ y $\gamma \in \mathbb{N}_0$, $C^{k,\beta}(\bar{\Omega}) \rightarrow C^{k,\beta}(\bar{\Omega})$ (es decir, incrustación compacta).

Los espacios de Hölder $C^{k,\gamma}(\bar{\Omega})$ pueden definirse de manera análoga, si las funciones y sus derivadas no son continuas hasta el límite $\partial\Omega$. Nótese que si $v \in C^{0,\gamma}(\Omega)$, $0 < \gamma < 1$, entonces v es uniformemente continua en Ω , y por lo tanto el criterio de Cauchy implica que v puede extenderse únicamente a $\partial\Omega$ para dar una función continua en Ω .

Esto nos permite hablar de valores en la frontera $\partial\Omega$ de $v \in C^{0,\gamma}(\Omega)$, y también se escribe $C^{0,\gamma}(\Omega) = C^{0,\gamma}(\bar{\Omega})$. $C^{0,\gamma}(\Omega) = C^{0,\gamma}(\bar{\Omega})$.

A.2 Espacios de Sobolev

Ahora revisamos varios resultados útiles sobre los espacios de Sobolev. Las referencias estándar sobre el tema incluyen [AF03, Gri85]. Sea $\Omega \subset \mathbb{R}^d$ un dominio abierto acotado con un límite $\partial\Omega$, y $\Omega^c = \mathbb{R}^d \setminus \Omega$. Se dice que el dominio Ω es $C^{k,\gamma}$ (respectivamente, Lipschitz) si $\partial\Omega$ es compacto y localmente cerca de cada punto límite, Ω puede ser representado como el conjunto sobre el gráfico de una función $C^{k,\gamma}$ (respectivamente, Lipschitz).

A.2.1 Espacios de Lebesgue

El espacio de Lebesgue $L^p(\Omega)$; $1 \leq p \leq \infty$, consta de funciones $v: \Omega \rightarrow \mathbb{R}$ con un finito $\|v\|_{L^p(\Omega)}$, donde

$$\|v\|_{L^p(\Omega)} = \begin{cases} \left(\int_{\Omega} |v|^p dx \right)^{\frac{1}{p}}, & \text{si } 1 \leq p < \infty \\ \sup_{\Omega} |v|, & \text{si } p = \infty \end{cases}$$

Se dice que una función $v: \Omega \rightarrow \mathbb{R}$ está en $L^p_{loc}(\Omega)$, si para cualquier subconjunto compacto $\omega \subset \subset \Omega$, se cumple $v \in L^p(\omega)$.

En espacios $L^p(\Omega)$, la desigualdad de Hölder es una herramienta indispensable: si $p, q \in [1, \infty]$ con $p^{-1} + q^{-1} = 1$ (es decir, p, q son exponentes conjugados y, a menudo, se escriben como $q = p'$), entonces para todo $v \in L^p(\Omega)$, y $w \in L^q(\Omega)$, se cumple $vw \in L^1(\Omega)$, y

$$\|vw\|_{L^1(\Omega)} \leq \|v\|_{L^p(\Omega)} \|w\|_{L^q(\Omega)}$$

La igualdad se cumple si y sólo si v y w son linealmente dependientes. Una generalización útil es la desigualdad de Young para la convolución $v * w$ definida en $\mathbb{R}^d$, es decir,

$$(v * w)(x) = \int_{\mathbb{R}^d} v(x-y)w(y)dy$$

Se puede obtener una versión de la desigualdad en dominio acotado por extensión cero.

Teorema A.2

Sean $v \in L^p(\mathbb{R}^d)$ y $w \in L^q(\mathbb{R}^d)$, con $1 \leq p, q \leq \infty$. Para cualquier

$1 \leq r \leq \infty$ con $p^{-1} + q^{-1} = r^{-1} + 1$,

$$\|v * w\|_{L^r(\mathbb{R}^d)} \leq \|v\|_{L^p(\mathbb{R}^d)} \|w\|_{L^q(\mathbb{R}^d)} \dots (A1)$$

Ocasionalmente usamos el débil $L^p(\Omega)$, $1 \leq p < \infty$, denotado por $L^{p,\infty}(\Omega)$. Para $v: \Omega \rightarrow \mathbb{R}$, la distribución λ_v se define para $t > 0$ por

$$\lambda_v(t) = |\{x \in \Omega: |v(x)| > t\}|$$

donde $|\cdot|$ denota la medida de Lebesgue de un conjunto en $\mathbb{R}^d$. se dice que v está en el espacio $L^{p,\infty}(\Omega)$, si existe $c > 0$ tal que,

para todo $t > 0$, $\lambda_v(t) \leq c^p t^{-p}$. la mejor constante c es la norma $L^{p,\infty}(\Omega)$, de v , y se denota por

$$\|v\|_{L^{p,\infty}(\Omega)} = \sup_{t>0} t \lambda_v^p(t)$$

La norma $L^{p,\infty}(\Omega)$ no es una norma verdadera, ya que la desigualdad del triángulo no se cumple. Para

$$v \in L^p(\Omega), \|v\|_{L^{p,\infty}(\Omega)} \leq \|v\|_{L^p(\Omega)},$$

$L^p(\Omega) \subset L^{p,\infty}(\Omega)$ Con la convención de que dos funciones son iguales si son iguales a.e., entonces los espacios $L^{p,\infty}(\Omega)$ son completos, y además, para $p > 1$, $L^{p,\infty}(\Omega)$ son espacios de Banach.

Ejemplo A.1

Sea $\omega_\alpha(t) = \frac{t^{\alpha-1}}{\Gamma(\alpha)}$, $\alpha \in (0, 1)$ y $T > 0$ sean fijos. Entonces, $\omega_\alpha \in L^p(0, T)$, para cualquier $p \in [0, (1 - \alpha)^{-1})$, pero no pertenece a $L^{\frac{1}{1-\alpha}}(0, T)$. Sin embargo, $\omega_\alpha \in L^{\frac{1}{1-\alpha},\infty}(0, T)$. En efecto, la distribución λ_{ω_α} viene dada por $\lambda_{\omega_\alpha}(t) = |\{s \in (0, T): |\omega_\alpha(s)| > t\}| = (\Gamma(\alpha)t)^{-\frac{1}{1-\alpha}}$. Por eso,

$$\|\omega_\alpha\|_{L^{\frac{1}{1-\alpha},\infty}(0,T)} = \sup_{t \in (0,T)} t \lambda_{\omega_\alpha}^{1-\alpha}(t) = \Gamma(\alpha)^{-1} < \infty$$

Esto muestra la afirmación deseada.

A.2.2 Espacios de Sobolev

Sea $C_0^\infty(\Omega)$ denota el espacio de funciones infinitamente diferenciables con compacto apoyo en Ω . Si $v, w \in L_{loc}^1(\Omega)$ y $\ell = (\ell_1, \ell_2, \dots, \ell_d) \in \mathbb{N}_0^d$, a índice múltiple, entonces w es llamó ℓ ésima derivada parcial débil de v , denotada por $w = D^\ell v$, si se cumple

$$\int_{\Omega} v D^\ell \phi dx = (-1)^{|\ell|} \int_{\Omega} w \phi dx, \forall \phi \in C_0^\infty(\Omega)$$

Con $|\ell| = \sum_{i=1}^d \ell_i$. Tenga en cuenta que una derivada parcial débil, si existe, es única hasta un conjunto de Lebesgue medida

cero. Definimos el espacio de Sobolev $W^{k,p}(\Omega)$, para cualquier $1 \leq p \leq \infty$ y $k \in \mathbb{N}$. Consta de todo $v \in L^1_{loc}(\Omega)$ tal que para cada $\ell \in \mathbb{N}_0^d$ con $|\ell| \leq k$, $D^\ell v$ existe en sentido débil y pertenece a $D^p(\Omega)$, equipado con la norma

$$\|v\|_{W^{k,p}(\Omega)} = \begin{cases} \left(\sum_{|\ell| \leq k} \int_{\Omega} |D^\ell v|^p dx \right)^{\frac{1}{p}}, & \text{si } 1 \leq p < \infty \\ \sum_{|\ell| \leq k} \text{eso} \sup_{\Omega} |D^\ell v|, & p = \infty \end{cases}$$

El espacio $W^{k,p}(\Omega)$, para cualquier $k \in \mathbb{N}$ y $1 \leq p \leq \infty$, es un espacio de Banach. el subespacio $W_0^{k,p}(\Omega)$ denota la clausura de $C_0^\infty(\Omega)$ con respecto a la norma $\|\cdot\|_{W_0^{k,p}(\Omega)}$, es decir, $W_0^{k,p}(\Omega) = \overline{C_0^\infty(\Omega)}^{W^{k,p}(\Omega)}$. El caso $p = 2$, es decir, $W^{k,2}(\Omega)$, y $W_0^{k,2}(\Omega)$, es Hilbert espacio, y a menudo denotado por $H^k(\Omega)$ y $H_0^k(\Omega)$, respectivamente.

A continuación, recordamos algunos resultados útiles. La primera es la desigualdad de Poincaré. Cuando $p = 2$, la constante óptima $c(\Omega, d)$ en la desigualdad coincide con el menor valor propio del Dirichlet Laplaciano negativo $-\Delta$ en el dominio Ω .

Proposición A.1 Si el dominio Ω está acotado, entonces para cualquier $1 \leq p < \infty$

$$\|v\|_{L^p(\Omega)} \leq c(\Omega, d) \|\nabla v\|_{L^p(\Omega)}, \forall v \in W_0^{1,p}(\Omega), \dots (A,2)$$

Las desigualdades incrustadas de Sobolev representan una de las herramientas más poderosas en el análisis. La suposición de regularidad C^2 en el dominio se puede relajar.

Teorema A.3

Sea Ω un dominio C^1 acotado. Entonces se cumplen las siguientes afirmaciones.

(i) Si $k < \frac{d}{p}$, entonces la incrustación $W^{k,p}(\Omega) \rightarrow L^q(\Omega)$ es continua, con $\frac{1}{q} = \frac{1}{p} - \frac{k}{d}$

(ii) Si $k > \frac{d}{p}$, entonces la incrustación $W^{k,p}(\Omega) \rightarrow C^{k-\lceil \frac{d}{p} \rceil - 1, \gamma}(\Omega)$ es continua, con $\gamma = \lceil \frac{d}{p} \rceil - \frac{d}{p} + 1$, si $\frac{d}{p} \notin \mathbb{N}$, y para cualquier $\gamma \in (0, 1)$ en caso contrario.

(iii) Si $k > l$ y $\frac{1}{p} - \frac{k}{d} < \frac{1}{q} - \frac{\ell}{d}$, entonces la incrustación $W^{k,p}(\Omega) \rightarrow W^{\ell,p}(\Omega)$ es compacta.

El tercer resultado es la incrustación multiplicativa de Gagliardo-Nirenberg

Teorema A.4

Sea Ω un dominio de Lipschitz acotado y $p \geq 1$. Entonces, para todo $s \geq 1$, se cumple

$$\|v\|_{L^q(\Omega)} \leq c(d, p, s) \|Dv\|_{L^p(\Omega)}^\alpha \|v\|_{L^s(\Omega)}^{1-\alpha}, \forall v \in W_0^{l,p}(\Omega)$$

donde $\alpha \in [0, 1]$, $p, q \geq 1$ están unidos por $\alpha = \frac{s^{-1}-q^{-1}}{s^{-1}-p^{-1}+d^{-1}}$, con el rango admisible

$$\begin{cases} q \in [s, \infty], \alpha \in \left[0, \frac{p}{p+s(p-1)}\right], \text{ si } d = 1 \\ q \in \left[\min\left(s, \frac{dp}{d-p}\right), \max\left(s, \frac{dp}{d-p}\right)\right], \alpha \in [0, 1], \text{ si } 1 \leq p < d, \\ q \in [s, \infty], \alpha \in \left[0, \frac{p}{dp+s(p-1)}\right], \text{ si } p \geq d > 1. \end{cases}$$

El siguiente resultado es un corolario del Teorema A.4, con $\alpha = 1$ y $s = 1$.

Corolario A.1

Sea $1 \leq p < d$, se cumple

$$\|v\|_{L^{\frac{dp}{d-p}}(\Omega)} \leq c(d, p) \|v\|_{L^p(\Omega)}, \forall v \in W_0^{l,p}(\Omega)$$

Si el límite $\partial\Omega$ es suave por partes, las funciones v en $W^{k,p}(\Omega)$ son en realidad definidas hasta $\partial\Omega$ a través de sus trazas, denotado por γv . El espacio $W_0^{l,p}(\Omega)$ se puede definir equivalentemente como el conjunto de funciones en $W^{k,p}(\Omega)$ cuya traza es cero.

Teorema A.5

Si $\partial\Omega$ es suave por partes, entonces para $q \in \left[1, \frac{(d-1)p}{d-p}\right]$, si $1 < p < d$ y $q \in [0, \infty)$, si $p = d$, se cumple

$$\|\gamma v\|_{L^q(\partial\Omega)} \leq c(d, p, \Omega) \|v\|_{W^{1,p}(\Omega)}$$

A.2.3 Espacios fraccionarios de Sobolev

Los espacios fraccionarios de Sobolev juegan un papel central en el análisis de fdes para obtener una descripción general de los espacios fraccionarios de Sobolev. Para cualquier $s \in (0, 1)$, el espacio de Sobolev $W^{s,p}(\Omega)$ está definido por

$$W^{s,p}(\Omega) = \{v \in L^p(\Omega) : |v|_{W^{s,p}(\Omega)} < \infty\}, \dots (A, 3)$$

donde el seminorm Sobolev-Slobodeckij $|v|_{W^{s,p}(\Omega)}$ está definida por

$$|v|_{W^{s,p}(\Omega)} = \left(\iint_{\Omega \times \Omega} \frac{(v(x) - v(y))^p}{|x - y|^{d+ps}} dx dy \right)^{\frac{1}{p}}$$

El espacio $W^{s,p}(\Omega)$ es un espacio de Banach, cuando está equipado con la norma

$$\|v\|_{W^{s,p}(\Omega)} = \left(\|v\|_{L^p(\Omega)}^p + |v|_{W^{s,p}(\Omega)}^p \right)^{\frac{1}{p}}$$

De manera equivalente, el espacio $W^{s,p}(\Omega)$ puede considerarse como la restricción de funciones en $W^{s,p}(\mathbb{R}^d)$ a Ω . El subespacio $W_0^{s,p}(\Omega) \subset W^{s,p}(\Omega)$ se define como la clausura de $C_0^\infty(\Omega)$ con respecto a la norma $W^{s,p}(\Omega)$, es decir,

$$W_0^{s,p}(\Omega) = \overline{C_0^\infty(\Omega)}^{W^{s,p}(\Omega)}, \dots (A, 4)$$

Denotamos $W^{s,2}(\Omega)$ por $H^s(\Omega)$, y de manera similar $W_0^{s,2}(\Omega)$ por $H_0^s(\Omega)$. Estos son Hilbert espacios y de uso frecuente. Si el límite $\partial\Omega$ es suave, pueden ser equivalentes definido a través de la interpolación real de espacios por el método K (Baleanu & Tenreiro Machado, 2009). Específicamente, establezca $H^0(\Omega) = L^2(\Omega)$, entonces los espacios de Sobolev con índice real $0 \leq s \leq 1$ pueden definirse como espacios de interpolación de índices para el par $[L^2(\Omega), H^1(\Omega)]$, es decir,

$$H^s(\Omega) = [L^2(\Omega), H^1(\Omega)], \dots (A, 5)$$

De manera similar, para $s \in [0,1] \setminus \{\frac{1}{2}\}$, los espacios $H_0^s(\Omega)$ se definen como espacios de interpolación de índice s para el par $[L^2(\Omega), H_0^1(\Omega)]$, es decir, $H_0^s(\Omega) = [L^2(\Omega), H_0^1(\Omega)]_s, s \in [0,1] \setminus \{\frac{1}{2}\}, \dots (A, 6)$

El espacio $[L^2(\Omega), H_0^1(\Omega)]_{\frac{1}{2}}$, es el llamado espacio de Lions-Magenes, $H_{0,0}^{\frac{1}{2}}(\Omega) = [L^2(\Omega), H_0^1(\Omega)]_{\frac{1}{2}}$, que se puede caracterizar como (Baleanu & Tenreiro Machado, 2009)

$$H_{0,0}^{\frac{1}{2}}(\Omega) = \left\{ w \in H^{\frac{1}{2}}(\Omega) : \int_{\Omega} \frac{w^2(x')}{\text{dist}(x', \partial\Omega)} dx' < \infty, \dots (A, 7) \right.$$

Si $\partial\Omega$ es Lipschitz, entonces (i) la caracterización (A.7) es equivalente a la definición mediante interpolación, y las definiciones (A.5) y (A.6) también son equivalentes a (A.3) y (A.4), respectivamente; (ii) el espacio, $C_0^\infty(\Omega)$ es denso en $H^s(\Omega)$ si y solo si $s \leq \frac{1}{2}$, es decir, $H_0^s(\Omega) = H^s(\Omega)$. Si $s > \frac{1}{2}$, $H_0^s(\Omega)$ está estrictamente contenido en $H^s(\Omega)$ [LM72, Teorema 11.1], y se cumplen las siguientes inclusiones:

$$H_{00}^{\frac{1}{2}}(\Omega) \subsetneq H_0^{\frac{1}{2}}(\Omega) = H^{\frac{1}{2}}(\Omega)$$

ya que $1 \in H_0^{\frac{1}{2}}(\Omega)$ pero $1 \notin H_{00}^{\frac{1}{2}}(\Omega)$.

Los espacios fraccionarios de Sobolev de orden mayor que uno se definen de manera similar. Dado $k \in \mathbb{N}_0$ y $0 < s < 1$, el espacio $H^{k+s}(\Omega)$ está definido por

$$H^{k+s}(\Omega) = \{u \in H^k(\Omega): |D^\ell u| \in H^s(\Omega), \forall \ell \in \mathbb{N}_0^d \text{ s. t. } |\ell| = k\}$$

Equipada con la norma

$$\|v\|_{H^{s,p}(\Omega)} = \|v\|_{H^k(\Omega)} + \sum_{|\ell|=k} |D^\ell v|_{H^s(\Omega)}$$

Cuando $d = 1$, el dominio Ω es un intervalo abierto acotado $D = (a, b)$. Entonces el espacio $H_0^s(D)$ consta de funciones en $H^s(\Omega)$ cuya extensión por cero a $\mathbb{R}$ está en $H^s(\mathbb{R})$. Definimos $H_{0,L}^s(D)$, (respectivamente, $H_{0,\mathbb{R}}^s(D)$) como el conjunto de funciones cuya la extensión por cero está en $H^s(-\infty, b)$ (respectivamente, $H^s(a, \infty)$).

A.2.4 $\dot{H}^s(\Omega)$ espacios

Usamos los espacios $\dot{H}^s(\Omega)$, $s \geq 0$, frecuentemente en el capítulo 6. Estos son espacios con tipo especial de condiciones de contorno, asociado con operadores elípticos adecuados. Dejar $A: H^s(\Omega) \cap H_0^1(\Omega) \rightarrow L^2(\Omega)$, sea un segundo orden simétrico uniformemente coercitivo operador elíptico con coeficientes suaves. El espectro de A consiste enteramente en valores propios $\{\lambda_j\}_{j=1}^\infty$ (ordenado no decreciente, contando las multiplicidades). Por $\varphi_j \in H^2(\Omega) \cap H_0^1(\Omega)$, denotamos las funciones propias ortonormales $L^2(\Omega)$ correspondientes a λ_j . Para cualquier $s \geq 0$, definimos el espacio $\dot{H}^s(\Omega)$ por

$$\dot{H}^s(\Omega) = \left\{ v \in L^2(\Omega) : \sum_{j=1}^\infty \lambda_j^s |(v, \varphi_j)|^2 < \infty \right\}$$

y que $\dot{H}^s(\Omega)$ es un espacio de Hilbert con la norma

$$\|v\|_{\dot{H}^s(\Omega)}^2 = \sum_{j=1}^\infty \lambda_j^s |(v, \varphi_j)|^2 < \infty$$

Por definición, tenemos la siguiente forma equivalente:

A.3 Transformadas Integrales

Ahora recordamos las transformadas de Laplace y Fourier.

A.3.1 Transformada de Laplace

La transformada de Laplace de una función $f: R_+ \rightarrow R$, denotada por $\hat{f}$ o $\mathcal{L}[f]$, se define por

$$\hat{f}(z) = \mathcal{L}[f](z) = \int_0^{\infty} e^{-zt} f(t) dt$$

Supongamos que $f \in \mathcal{L}_{loc}^1(R_+)$ y existe algún $\lambda \in R_+$ tal que $|f(t)| \leq c e^{\lambda t}$ para t grande, entonces $\hat{f}(z)$ existe y es analítica para $\Re(z) > \lambda$. La transformada de Laplace de toda la función de tipo exponencial se puede obtener transformando término a término la expansión de Taylor de la función original alrededor del origen (Doetsch, 1971). En este caso la transformada de Laplace resultante es analítica y se desvanece en el infinito

Ejemplo A.2

Calculemos la transformada de Laplace del kernel de Riemann-Liouville

$w_\alpha(t) = \Gamma(\alpha)^{-1} t^{\alpha-1}$ para $t > 0$. Luego, para $\alpha > 0$, la sustitución $s = tz$ da

$$\begin{aligned} \hat{f}(z) = \mathcal{L}[w_\alpha](z) &= \frac{1}{\Gamma(\alpha)} \int_0^{\infty} e^{-zt} t^{\alpha-1} dt = \frac{z^{-\alpha}}{\Gamma(\alpha)} \int_0^{\infty} e^{-s} s^{\alpha-1} ds \\ &= z^{-\alpha} \end{aligned}$$

donde la última identidad se deriva de la definición de $\Gamma(z)$.

La transformada de Laplace de la convolución en R_+ , es decir

$$(f * g)(t) = \int_0^t f(t-s)g(s)ds = \int_0^t f(s)g(t-s)ds$$

de dos funciones f y g que se anulan para $t < 0$ satisface la regla de convolución:

$$\widehat{f * g}(z) = \hat{f}(z)\hat{g}(z), \dots (A.8)$$

Siempre que existan tanto $\hat{f}(z)$ y $\hat{g}(z)$. Esta regla es útil para evaluar la transformada de la Laplace de integrales fraccionarias y derivadas.

Ejemplo A.3

La integral fraccionaria de Riemann-Liouville ${}_0I_t^\alpha f(t)$ es una convolución con w_α , es decir, ${}_0I_t^\alpha f(t) = (w_\alpha * g)(t)$. Por lo tanto, por (A.8), la transformada de Laplace de ${}_0I_t^\alpha f(t)$ es dada por,

$$\mathcal{L}[{}_0I_t^\alpha f](z) = e^{-\alpha} \hat{f}(z)$$

La transformada de Laplace de la derivada de n -ésimo orden $f^{(n)}$ viene dada por

$$\mathcal{L}[f^{(n)}](z) = z^n \hat{f}(z) - \sum_{k=0}^{n-1} z^{n-k-1} f^{(k)}(0^+)$$

Se sigue integración por partes, si todas las integrales involucradas tienen sentido.

También tenemos una fórmula de inversión para la transformada de Laplace

$$f(t) = \frac{1}{2\pi i} \int_{a-i\infty}^{a+i\infty} e^{zt} \hat{f}(z) dz$$

donde la integral está a lo largo de la línea vertical $\mathcal{R}(z) = a$ en $\mathbb{C}$ (también conocido como contorno de Bromwich), tal que $a > \lambda$, es decir, es mayor que la parte real de todas las singularidades de $\hat{f}(z)$ y $\hat{f}(z)$ está acotado en la línea. La evaluación directa de la fórmula de inversión es a menudo inconveniente. El contorno a menudo se puede deformar para facilitar la evaluación analítica y computación numérica. Por ejemplo, se puede deformar a un contorno de Hankel, es decir, un camino en el plano complejo $\mathbb{C}$ que se extiende desde $(-\infty, a)$, dando vueltas alrededor del origen en sentido antihorario y de regreso a $(-\infty, a)$.

La transformada de Laplace es una herramienta importante para el análisis matemático. A continuación, damos dos

ejemplos, es decir, monotonidad completa y asintótica. Recuerda que una función

$f: R_+ \rightarrow \overline{R_+}$ se dice que es completamente monótono si $(-1)^k f^{(k)}(t) \geq 0$, para cualquier $t \in R_+$,

$k = 0, 1, 2, \dots$. El teorema de Bernstein establece que una función f es completamente monótona si y solo si la transformada de Laplace de una medida es no negativa (Chen et al., 2007).

Teorema A.8

Una función $f: R_+ \rightarrow \overline{R_+}$ es completamente monótona si y sólo si tiene la representación $f(t) = \int_0^\infty e^{-ts} d\mu(s)$, donde μ es una medida no negativa en R_+ tal que la integral converge para todo $t > 0$.

El comportamiento asintótico de una función $f(t)$ cuando $t \rightarrow \infty$ se puede determinar, bajo condiciones adecuadas, observando el comportamiento de su transformada de Laplace $\hat{f}(z)$ como $z \rightarrow 0$, y viceversa. Esto se describe mediante el teorema tauberiano de Karamata-Feller. (Gorenflo et al., 1999)(Joulin, 1985)

Teorema A.9 Sea $L: R_+ \rightarrow R_+$ una función que varía lentamente en ∞ , es decir, para todo x fijo > 0 , $\lim_{t \rightarrow \infty} L(xt)L(t)^{-1} = 1$. Sea $\beta > 0$ y $f: R_+ \rightarrow R$ una función monótona cuya transformada de Laplace $\hat{f}(z)$ existe para todo $z \in C_+$. Entonces

$\hat{f}(z) \sim z^{-\beta} L(z^{-1})$ como $z \rightarrow 0$ si y solo si $f(t) \sim \frac{t^{\beta-1}}{\Gamma(\beta)} L(t)$ como $t \rightarrow \infty$

donde los enfoques están en R_+ y $f(t) \sim g(t)$ como $t \rightarrow t^*$ denota $\lim_{t \rightarrow t^*} \frac{f(t)}{g(t)} = 1$

A.3.2 Transformada de Fourier

La transformada de Fourier de una función $f \in L^1(\mathbb{R})$, denotada por $\tilde{f}$ o $\mathcal{F}[f]$, está definida por

$$\tilde{f}(\xi) = \mathcal{F}[f](\xi) = \int_{-\infty}^{\infty} e^{i\xi x} f(x) dx$$

Tenga en cuenta que esta elección no involucra el factor 2π que a menudo aparece en la literatura.

Se toma para ser consistente con varios libros de texto populares sobre fraccionario cálculo (Samko et al., s.f.) (A. Kilbas et al., 2006). Si $f \in L^1(\mathbb{R})$, entonces $\tilde{f}$ es continua en $\mathbb{R}$, y $\tilde{f}(\xi) \rightarrow 0$ como $|\xi| \rightarrow \infty$. El teorema de Plancherel establece que la transformada de Fourier se extiende únicamente a un operador lineal acotado $\mathcal{F}: L^2(\mathbb{R}) \rightarrow L^2(\mathbb{R})$ que satisface.

$$\frac{1}{2\pi} \int_{-\infty}^{\infty} \tilde{f}(\xi) \overline{\tilde{g}(\xi)} d\xi = \int_{-\infty}^{\infty} f(x) g(x) dx \dots (A.9)$$

donde la línea superior denota conjugado complejo. La fórmula de inversión de la transformación Fourier está dada por

$$f(x) = \frac{1}{2\pi} \lim_{R \rightarrow \infty} \int_{-R}^R e^{-i\xi x} \tilde{f}(\xi) d\xi, x \in \mathbb{R}$$

y la igualdad se cumple en todo punto continuo de f .

Ejemplo A.4

La transformada de Fourier de $f(x) = e^{-\lambda x^2}$, $\lambda > 0$, viene dada por

$$\tilde{f}(\xi) = \int_{-\infty}^{\infty} e^{i\xi x} e^{-\lambda x^2} dx = \int_{-\infty}^{\infty} \cos(\xi x) e^{-\lambda x^2} dx - i \int_{-\infty}^{\infty} \sin(\xi x) e^{-\lambda x^2} dx$$

La segunda integral se anula ya que el integrando es una función impar. Utilizando el

identidad para cualquier $\lambda > 0$ (Véase en (Gorenflo & Rutman, 1995), pág. 391):

$$\int_{-\infty}^{\infty} e^{-\lambda x^2} \cos(bx) dx = \frac{1}{2} \sqrt{\frac{\pi}{\lambda}} e^{-\frac{b^2}{4\lambda}}$$

Deducimos:

$$\tilde{f}(\xi) = \int_{-\infty}^{\infty} \cos(\xi x) e^{-\lambda x^2} dx = \frac{1}{2} \sqrt{\frac{\pi}{\lambda}} e^{-\frac{b^2}{4\lambda}}$$

De manera similar, la transformada inversa de Fourier $\tilde{g}(\xi) = e^{-\lambda \xi^2}$ viene dada por,

$$\mathcal{F}[\tilde{g}](x) = \frac{1}{\sqrt{4\lambda\pi}} e^{-\frac{x^2}{4\lambda}}.$$

Tenga en cuenta que las transformadas de Fourier y Fourier inversa difieren en un factor de 2π

A.4 Teoremas del punto fijo

Ahora describimos varios teoremas de punto fijo (Hadid & Momani, 1996), para obtener información completa demostraciones y amplias aplicaciones. Estos teoremas son centrales para mostrar varios resultados de existencia y unicidad (S. M. Momani, 2000) para ecuaciones de operadores en espacios de Banach, en los que fdes se puede reformular. La herramienta más poderosa es el teorema del punto fijo de Banach.

Hay varias versiones, con condiciones ligeramente diferentes. El más básico se ocupa de los mapeos de contracción. Sea (X, ρ) un espacio métrico. un mapeo

$T: X \rightarrow X$ se dice que es una contracción si existe una constante $\gamma \in (0, 1)$ tal que Ver detalles

$$\rho(Tx_1, Tx_2) \leq \gamma \rho(x_1, x_2), x_1, x_2 \in X$$

Teorema A.10

Sea (X, ρ) un espacio métrico completo y T una contracción. Entonces T tiene un único punto fijo, dado por $\lim_{n \rightarrow \infty} T^n x_0$, para cualquier $x_0 \in X$, y además

$$\rho(x_n, x) \leq (1 - \gamma)^{-1} \gamma^n \rho(x_0, Tx_0)$$

El siguiente resultado debido a Schauder da la existencia de un punto fijo, sin unicidad.

Teorema A.11

Sea K un subconjunto no vacío, cerrado, acotado y convexo de un Banach espacio X , y supongamos que $T : K \rightarrow K$ es un operador compacto. Entonces T tiene un punto fijo.

El siguiente resultado es una versión alternativa del teorema del punto fijo de Schauder.

Teorema A.12

Sea K un subconjunto no vacío, compacto y convexo de un espacio de Banach,

y sea $T : K \rightarrow K$ una aplicación continua. Entonces T tiene un punto fijo.

El siguiente resultado de Krasnoselskii maneja operadores completamente continuos.

Teorema A.13

Sean X un espacio de Banach, $C \subset X$ un cono y B_1, B_2 dos espacios acotados.

bolas abiertas de X centradas en el origen, con $\overline{B_1} \subset B_2$. Supongamos que $T : C \cap (\overline{B_2} \setminus B_1) \rightarrow C$

es un operador completamente continuo tal que,

$$\|Tx\| \leq \|x\| \quad x \in C \cap \partial B_1 \quad \text{y} \quad \|Tx\| \geq \|x\| \quad x \in C \cap \partial B_2 \quad \text{o}$$

$$\|Tx\| \geq \|x\| \quad x \in C \cap \partial B_1 \quad \text{y} \quad \|Tx\| \leq \|x\| \quad x \in C \cap \partial B_2 \quad \text{o}$$

se cumple, entonces T tiene un punto fijo en $C \cap (\overline{B_2} \setminus B_1)$.

El siguiente resultado es útil para probar la existencia en un cono.

Teorema A.14

Sea X un espacio de Banach, $C \subset X$ un cono. Sea B_R una bola de radio $R > 0$ centrado en el origen. Supongamos que $T: C \cap B_R \rightarrow C$ es completamente continua operador tal que $Tx \neq \lambda x$ para todo $x \in \partial(B_R \cap C)$ y $\lambda \geq 1$. Entonces T tiene un punto en $C \cap B_R$.

La siguiente alternativa no lineal de Leray y Schauder es poderosa.

Teorema A.15

Sea $K \subset X$ un conjunto convexo de un espacio lineal normado X , y Z sea un subconjunto abierto en K tal que $0 \in Z$. Entonces cada aplicación compacta continua $T: \bar{Z} \rightarrow K$ tiene al menos una de las siguientes propiedades:

- (i) T tiene un punto fijo en $\bar{Z}$
- (ii) Existen $u \in \partial Z$ y $\mu \in (0, 1)$ tales que $u = \mu Tu$.

Referencias (Agarwal et al., 2008). Algunos teoremas de existencia en diferentes.

Apendice B

Una tabla de derivadas de Caputo

Para comodidad del lector, proporcionamos este apéndice donde damos algunas derivadas de tipo Caputo de ciertas funciones importantes. No nos esforzamos por la integridad, en cualquier sentido, pero queremos dar al menos las derivadas clásica y ejemplos

A lo largo de este apéndice, n siempre denotará el orden del tipo Caputo operador diferencial bajo consideración. Sólo consideraremos el caso $n > 0$ y $n \notin \mathbb{N}$, y usamos la notación $m := [n]$ para denotar el entero más pequeño mayor que (o igual a) n . Recuerde que para $n \in \mathbb{N}$, el operador diferencial de Caputo coincide con el operador diferencial habitual de orden entero, y para $n < 0$, el diferencial de Caputo

El operador de orden negativo se puede interpretar como el operador diferencial de Riemann-Liouville de la misma orden. Las tablas de este último se dan en varios lugares en la literatura (cf., por ejemplo, Podlubny o Samko et al.); no vamos a repetir esos resultados aquí.

Varias funciones especiales surgirán en este sentido; para las definiciones precisas nos referimos al Apéndice A. Por $i = \sqrt{-1}$ denotamos la unidad imaginaria.

1. Sea $f(x) = x^j$. Aquí tenemos que distinguir algunos casos:

$$(D_{*0}^n f)(x) = \begin{cases} 0, & \text{si } j \in \mathbb{N}_0 \text{ y } j < m \\ \frac{\Gamma(j+1)}{\Gamma(j+1-n)} x^{j-n}, & \text{si } j \in \mathbb{N}_0 \text{ y } j \geq m. \text{ o si } j \notin \mathbb{N}_0 \text{ y } j > m \end{cases}$$

2. Sea $f(x) = (x+c)^j$. para $c > 0$ arbitrario y $j \in \mathbb{R}$. Entonces

$$(D_{*0}^n f)(x) = \frac{\Gamma(j+1)}{\Gamma(j+1-m)} \frac{c^{j-m-1} x^{m-n}}{\Gamma(m-n+1)} {}_2F_1(1, m-j; m-n+1; -x/c),$$

3. Sea $f(x) = \exp(jx)$ para alguna $j \in \mathbb{R}$. Entonces

$$(D_{*0}^n f)(x) = j^m x^{m-n} E_{1, m-n+1}(jx)$$

4. Sea $f(x) = x^j \ln x$ para algún $j > m-1$. Entonces

$$\begin{aligned} (D_{*0}^n f)(x) &= x^{j-n} \sum_{k=0}^{m-1} (-1)^{m-k+1} \binom{j}{k} \frac{m!}{m-k} \frac{\Gamma(j-m+1)}{\Gamma(j-n+1)} \\ &\quad + \frac{\Gamma(j+1)}{\Gamma(j-n+1)} x^{j-n} (\Psi(j-m+1) \\ &\quad - \Psi(j-n+1) + \ln x) \end{aligned}$$

5. Sea $f(x) = \sin jx$ para algún $j \in \mathbb{R}$. Aquí nuevamente tenemos dos casos:

$$\begin{aligned} &(D_{*0}^n f)(x) \\ &= \begin{cases} \frac{j^m i (-1)^{m/2} x^{m-n}}{2\Gamma(m-n+1)} [{}_1F_1(1; m-n+1; jx) + {}_1F_1(1; m-n+1; -jx)] & \text{incluso } m \\ \frac{j^m i (-1)^{(m-1)/2} x^{m-n}}{2\Gamma(m-n+1)} [{}_1F_1(1; m-n+1; jx) + {}_1F_1(1; m-n+1; -jx)] & m \text{ impar} \end{cases} \end{aligned}$$

6. Finalmente consideramos $f(x) = \cos jx$ con algún $j \in \mathbb{R}$. Como en el ejemplo anterior, obtenemos dos casos:

$$\begin{aligned} &(D_{*0}^n f)(x) \\ &= \begin{cases} \frac{j^m (-1)^{m/2} x^{m-n}}{2\Gamma(m-n+1)} [{}_1F_1(1; m-n+1; jx) + {}_1F_1(1; m-n+1; -jx)] & \text{incluso } m \\ \frac{j^m i (-1)^{(m-1)/2} x^{m-n}}{2\Gamma(m-n+1)} [{}_1F_1(1; m-n+1; jx) + {}_1F_1(1; m-n+1; -jx)] & m \text{ impar} \end{cases} \end{aligned}$$

Apendice C

C. Solución numérica de Ecuaciones diferenciales fraccionarias

Para la mayoría de las ecuaciones diferenciales fraccionarias no podemos proporcionar métodos para calcular las soluciones exactas analíticamente. Por tanto, es necesario volver a los métodos numéricos.

Para brindar al lector una herramienta que pueda aplicarse a una clase muy amplia de ecuaciones. Ahora presentamos un método que se comprende bien y que ha demostrado ser eficiente en muchas aplicaciones prácticas (Diethelm & Freed, 1999b), (Diethelm & Luchko, 2004), (Glöckle & Nonnenmacher, 1995). Comenzamos discutiendo este problema para ecuaciones de un solo término y luego extender nuestra idea a problemas de múltiples términos

C.1 Un algoritmo para ecuaciones de un solo término

El método puede llamarse indirecto porque, en lugar de discretizar la Ecuación diferencial

$$D_{*0}^n y(x) = f(x, y(x))$$

con condiciones iniciales apropiadas

$$D^n y(0) = y_0^{(k)}, k = 0, 1, 2, \dots, [n] - 1$$

directamente, requiere alguna manipulación analítica preliminar, es decir, una aplicación del Lema 7.2 para convertir el problema del valor inicial para la ecuación diferencial en una ecuación integral de Volterra equivalente

$$y(x) = \sum_{k=0}^{m-1} \frac{x^k}{k!} D^k y(0) + \frac{1}{\Gamma(n)} \int_0^x (x-t)^{n-1} f(t, y(t)) dt, \dots \quad (C.1)$$

dónde $m = [n]$. Por lo tanto, ahora veremos un método para la solución numérica de (C.1).

El algoritmo que consideraremos puede interpretarse como una variante fraccionaria del método clásico de Adams-Bashforth-Moulton de segundo orden. ha sido introducido y

discutido brevemente en (Diethelm & Freed, 1999b); se proporciona más información en (Diethelm & Luchko, 2004). Algunos resultados adicionales para un problema de valor inicial específico están contenidos en ((Diethelm & Ford, 2004), pág. 490–507), una descripción detallada El análisis matemático se proporciona en (Diethelm & Freed, 1999a), y se pueden hacer comentarios prácticos adicionales, encontrado en (Diethelm et al., 2004). Se realizan experimentos numéricos y comparaciones con otros métodos. informado en (Edwards et al., 2002), (Glöckle & Nonnenmacher, 1995). Aquí haremos un análisis aún más detallado. bajo supuestos bastante generales.

Observación C.I.

Antes de iniciar las investigaciones es necesario dar una nota de cautela. Es común construir métodos para ecuaciones diferenciales fraccionarias tomando métodos para ecuaciones clásicas (típicamente de primer orden) y luego generalizar los conceptos de forma adecuada. A las fórmulas resultantes generalmente se les da el mismo nombre que el algoritmo clásico subyacente, posiblemente extendido por el adjetivo "fraccionario". Sin embargo, muchos esquemas numéricos clásicos se pueden ampliar en más de una manera. Esto puede conducir al problema de que, en dos artículos diferentes de la literatura, dos algoritmos diferentes se denotan de manera idéntica. Por supuesto, este es una potencial fuente de confusión y el lector debe tener mucho cuidado a este respecto. Por ejemplo, las reglas fraccionarias de Adams-Moulton de Galeone y Garrappa (Gorenflo & Mainardi, 1996) no coinciden con los métodos del mismo nombre que desarrollaremos a continuación.

Formulación clásica

Para motivar la construcción del método, primero recordaremos brevemente la idea detrás del algoritmo clásico de Adams-Bashforth-Moulton para ecuaciones de primer orden. Entonces, para empezar, centraremos nuestra atención en el conocido problema de valor inicial para la ecuación diferencial de primer orden:

$$Dy(x) = f(x, y(x)), \dots (C.2a)$$

$$y(x) = 0, \dots (C.2b)$$

Suponemos que la función f es tal que existe una solución única en algún intervalo $[0, T]$, digamos. Siguiendo (Kiryakova, 2010b), sugerimos utilizar la técnica predictor-corrector de Adams donde, en aras de la simplicidad, asumimos que estamos trabajando en una cuadrícula uniforme $\{t_j = jh; j = 0, 1, 2, \dots, N\}$ con algún número entero N y $h = T/N$. En algunas aplicaciones puede ser más eficiente utilizar una cuadrícula no uniforme, y lo haremos Desarrollando las fórmulas de aproximación numérica en este sentido generalizado. Sin embargo, para el análisis posterior de las propiedades del esquema restringiremos Pasando al caso equiespaciado.

La idea básica es, suponiendo que ya hemos calculado las aproximaciones $y(t_j)$ ($j = 1, 2, \dots, k$), que intentamos obtener la aproximación y_{k+1} mediante la ecuación

$$y(t_{k+1}) = y(t_k) + \int_{t_k}^{t_{k+1}} f(z, y(z)) dz, \dots (C.3)$$

Esta ecuación surge de la integración de (C.2a) en el intervalo $[t_k, t_{k+1}]$. Por supuesto, No conocemos exactamente ninguna de las expresiones del lado derecho de (C.3), pero Tenemos una aproximación para $y(t_k)$, concretamente y_k , que podemos usar en su lugar. la integral luego se reemplaza por la fórmula de cuadratura trapezoidal de dos puntos.

$$\int_a^b g(z) dz \approx \frac{b-a}{2} (g(a) + g(b)), \dots (C.4)$$

dando así una ecuación para la aproximación desconocida y_{k+1} , siendo

$$y_{k+1} = y_k + \frac{t_{k+1} - t_k}{2} (f(t_k, y(t_k)) + f(t_{k+1}, y(t_{k+1}))), \dots (C.5)$$

donde nuevamente tenemos que reemplazar $y(t_k)$ y $y(t_{k+1})$ por sus aproximaciones y_k y y_{k+1} , respectivamente. Esto produce la ecuación para el Adams-Moulton implícito de un paso método, que es

$$y_{k+1} = y_k + \frac{t_{k+1} - t_k}{2} (f(t_k, y_k) + f(t_{k+1}, y_{k+1})), \dots (C.6)$$

El problema con esta ecuación es que la cantidad desconocida y_{k+1} aparece en ambos lados, y debido a la naturaleza no lineal de la función f , no podemos resolver para y_{k+1} directamente en general. Por lo tanto, podemos usar (C.6) en un proceso iterativo, insertando una aproximación preliminar para y_{k+1} en el lado derecho para determinar una mejor aproximación que luego podemos utilizar.

La aproximación preliminar y_{k+1}^P , el llamado predictor, se obtiene de una manera muy similar, solo reemplazando la fórmula de la cuadratura trapezoidal por la regla del rectángulo,

$$\int_a^b g(z)dz \approx (b-a)g(a), \dots (C.7)$$

dando el método explícito (Euler directo o Adams-Bashforth de un paso)

$$y_{k+1}^P = y_k + hf(t_k, y_k), \dots (C.8)$$

Es bien sabido (véase en (Kiryakova, 2010b), pág. 372), que el proceso definido por (C.8) y

$$y_{k+1} = y_k + \frac{h}{2} (f(t_k, y_k) + f(t_{k+1}, y_{k+1}^P)), \dots (C.9)$$

conocida como técnica de Adams-Bashforth-Moulton de un solo paso, es convergente de orden 2, es decir

$$\max_{j=1,2,\dots,N} |y(t_j) - y_j| = O(h^2), \dots (C.10)$$

Además, este método se comporta satisfactoriamente desde el punto de vista de su valor numérico. estabilidad en el cap. IV del (Klages et al., 2008). Se dice que es del PECE (Predecir, Evaluar, Corregir, Evaluar) tipo porque, en una implementación concreta, comenzaríamos calculando el predictor en (C.8), entonces evaluamos $f(t_{k+1}, y_{k+1}^P)$, usa esto para calcular el corrector en (C.9), y finalmente evaluar $f(t_{k+1}, y_{k+1})$. Este resultado se almacena para uso futuro. en el siguiente paso de integración.

Formulación fraccionaria

Habiendo introducido este concepto, ahora intentamos trasladar las ideas esenciales al Problema de orden fraccionario con algunas modificaciones inevitables. La clave es derivar una ecuación similar a (C.3). Afortunadamente, existe una ecuación de este tipo, a saber (C.1). Esta ecuación se ve algo diferente de (C.3), porque el rango de integración ahora comienza en 0 en lugar de t_k . Esto es una consecuencia de la estructura no local de los operadores diferenciales de orden fraccionario. Sin embargo, esto no causa grandes problemas en nuestros intentos de generalizar el método Adams. Lo que hacemos es simplemente usar la fórmula del producto de cuadratura trapezoidal para reemplazar la integral, es decir, usamos los nodos $t_j, j = 0, 1, 2, \dots, k+1$ e interpreta la función $(t_{k+1} - z)^{n-1}$ como un peso función para la integral. En otras palabras, aplicamos la aproximación

$$\int_0^{t_{k+1}} (t_{k+1} - z)^{n-1} g(z) dz \approx \int_0^{t_{k+1}} (t_{k+1} - z)^{n-1} \tilde{g}_{k+1}(z) dz, \dots (C.11)$$

donde $\tilde{g}_{k+1}$ es la interpolación lineal por partes para g con nodos y nudos elegidos en el $t_j, j = 0, 1, 2, \dots, k+1$.

Está claro por construcción que la fórmula de cuadratura trapezoidal ponderada requerida se puede representar como una suma ponderada de valores de función del integrando g , tomado en los puntos t_j . Específicamente, encontramos que podemos escribir la integral en el lado derecho de (C.11) como

$$\int_0^{t_{k+1}} (t_{k+1} - z)^{n-1} \tilde{g}(z) dz = \sum_{j=0}^{k+1} a_{j,k+1} g(t_j), \dots (C.12a)$$

Donde,

$$a_{j,k+1} = \int_0^{t_{k+1}} (t_{k+1} - z)^{n-1} \phi_{j,k+1}(z) dz, \dots (C.12b)$$

Y

$$\phi_{j,k+1}(z) = \begin{cases} \frac{(z-t_j)}{(t_j-t_{j-1})}, & \text{si } t_{j-1} < z \leq t_j \\ \frac{(t_{j+1}-z)}{(t_{j+1}-t_j)}, & \text{si } t_j < z \leq t_{j+1}, \dots \text{(C.12c)} \\ 0, & \text{ptros casos} \end{cases}$$

Esto es claro porque las funciones $\phi_{j,k+1}$ satisfacen

$$\phi_{j,k+1}(t_u) = \begin{cases} 0, & \text{si } j \neq u \\ 1, & \text{si } j = u \end{cases}$$

y que son continuos y lineales por partes con puntos de interrupción en los nodos t_u , por lo que deben integrarse exactamente por nuestra fórmula.

Un cálculo explícito sencillo produce que, para una elección arbitraria de t_j , (C.12b) y (C.12c) producir

$$a_{0,k+1} = \frac{(t_{k+1}-t_1)^{n+1} + t_{k+1}^n [nt_1 + t_1 - t_{k+1}]}{t_1 n(n+1)}, \dots \text{(C.13a)}$$

$$a_{j,k+1} = \frac{(t_{k+1}-t_{j-1})^{n+1} + (t_{k+1}-t_j)^n + [n(t_{j-1}-t_j) + t_{j-1} - t_{k+1}]}{(t_j - t_{j-1})n(n+1)} + \frac{(t_{k+1}-t_{j+1})^{n+1} + (t_{k+1}-t_j)^n + [n(t_j - t_{j+1}) - t_{j+1} + t_{k+1}]}{(t_{j+1} - t_j)n(n+1)}, \dots \text{(C.13b)}$$

$$\text{si } 1 \leq j \leq k, \text{ y } a_{k+1,k+1} = \frac{(t_{k+1}-t_k)^n}{n(n+1)}, \dots \text{(C.13c)}$$

En el caso de nodos equiespaciados ($t_j = jh$ con alguna h fija), estas relaciones se reducen a

$$a_{j,k+1} = \begin{cases} \frac{h^n}{n(n+1)} (k^{n+1} - (k-n)(k+1)^n), & \text{si } j = 0 \\ \frac{h^n}{n(n+1)} ((k-j+2)^{n+1} - (k-j)^{n+1} - 2(k-j+1)^{n+1}), & \text{si } 1 \leq j \leq k \\ \frac{h^n}{n(n+1)}, & \text{si } j = k+1 \end{cases} \dots \text{(C.14)}$$

Esto nos da nuestra fórmula correctora (es decir, la variante fraccionaria de la fórmula de un solo paso). método Adams-Moulton), que es

$$y_{k+1}(x) = \sum_{j=0}^{m-1} \frac{t_{k+1}^j}{j!} y_0^{(j)} + \frac{1}{\Gamma(n)} \left(\sum_{j=0}^k a_{j,k+1} f(t_j, y_j) + a_{k+1,k+1} f(t_{k+1}, y_{k+1}^P) \right), \dots (C.15)$$

El problema restante es la determinación de la fórmula predictiva que requiere calcular el valor y_{k+1}^P . La idea que utilizamos para generalizar el paso único es el método Adams-Bashforth, el mismo que el descrito anteriormente para Adams-Técnica de Moulton: Sustituimos la integral del lado derecho de (C.1) por la regla del rectángulo del producto

$$\int_0^{t_{k+1}} (t_{k+1} - z)^{n-1} g(z) dz \approx \sum_{j=0}^k b_{j,k+1} g(t_j), \dots (C.16)$$

donde ahora,

$$b_{j,k+1} = \int_0^{t_{k+1}} (t_{k+1} - z)^{n-1} g(z) dz = \frac{(t_{k+1} - t_j)^n - (t_{k+1} - t_{j+1})^n}{n}, \dots (C.17)$$

Esta expresión para pesos se puede derivar de una manera similar al método utilizado en la derivación de (C.13). Sin embargo, aquí estamos tratando con una constante por partes aproximación y no lineal por partes, por lo que tenemos que reemplazar la Funciones “en forma de sombrero” $b_{k,j}$ mediante funciones de valor constante 1 en $[t_j, t_{j+1}]$ y 0 en las partes restantes del intervalo $[0, t_{j+1}]$. Nuevamente, en el caso equiespaciado, tenemos tener la expresión más simple

$$b_{j,k+1} = \frac{h^n}{n} ((k+1-j)^n - (k-j)^n), \dots (C.18)$$

Por tanto, el predictor y_{k+1}^P se determina mediante el método fraccionario de Adams-Bashforth

$$y_{k+1}^P = \sum_{j=0}^{m-1} \frac{t_{k+1}^j}{j!} y_0^{(j)} + \frac{1}{\Gamma(n)} \sum_{j=0}^k b_{j,k+1} f(t_j, y_j), \dots (C.19)$$

Por lo tanto, nuestro algoritmo básico, el método fraccionario de Adams-Bashforth-Moulton, es descrito completamente ahora por (C.19) y (C.15) con los pesos $a_{j,k+1}$ y $b_{j,k+1}$ definiéndose según (C.13) y (C.17), respectivamente.

Análisis de errores

Para el análisis de errores de este algoritmo, restringimos nuestra atención al caso de una cuadrícula equiespaciada, es decir, de ahora en adelante asumimos que $t_j = jh = jT/N$ con algo de $N \in \mathbb{N}$. Esencialmente seguimos la estructura de (Diethelm & Freed, 1999a) y comenzamos indicando algunos elementos auxiliares. resultados.

Lo que necesitamos para nuestros propósitos es alguna información sobre los errores de cuadratura fórmulas que hemos utilizado en la derivación del predictor y el corrector, respectivamente. Primero damos una afirmación sobre la regla del rectángulo del producto que tenemos utilizado para el predictor.

Teorema C.1.

a) Sea $z \in C^1[0, T]$. Entonces

$$\left| \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} z(t) dt - \sum_{j=0}^k b_{j,k+1} z(t_j) \right| \leq \frac{1}{n} \|z'\|_{\infty} t_{k+1}^n h$$

b) Sea $z(t) = t^p$ para algún $p \in (0, 1)$. Entonces

$$\left| \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} z(t) dt - \sum_{j=0}^k b_{j,k+1} z(t_j) \right| \leq C_{n,p}^{Re} t_{k+1}^{n+p-1} h$$

Donde $C_{n,p}^{Re}$ es una constante que depende sólo de n y p .

Prueba.

Al construir la fórmula del producto rectángulo, encontramos en ambos casos que el error de cuadratura tiene la representación

$$\int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} z(t) dt - \sum_{j=0}^k b_{j,k+1} z(t_j) = \sum_{j=0}^k \int_{jh}^{(j+1)h} (t_{k+1} - t)^{n-1} (z(t) - z(t_j)) dt, \dots (C.20)$$

Para probar el enunciado (a), aplicamos el teorema del valor medio del cálculo diferencial al segundo factor del integrando en el lado derecho de (C.20) y derivar

$$\begin{aligned}
& \left| \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} z(t) dt - \sum_{j=0}^k b_{j,k+1} z(t_j) \right| \\
& \leq \|z'\|_\infty \sum_{j=0}^k \int_{jh}^{(j+1)h} (t_{k+1} - t)^{n-1} (t - jh) dt \\
& = \|z'\|_\infty \frac{h^{n+1}}{n} \sum_{j=0}^k \left(\frac{1}{1+n} [(k+1-j)^{1+n} - (k-j)^n] \right) \\
& = \|z'\|_\infty \frac{h^{n+1}}{n} \left(\frac{(k+1)^{1+n}}{1+n} - \sum_{j=0}^k (j)^n \right) \\
& = \|z'\|_\infty \frac{h^{n+1}}{n} \left(\int_0^{k+1} t^n dt - \sum_{j=0}^k (j)^n \right)
\end{aligned}$$

Aquí el término entre paréntesis es simplemente el resto del rectángulo estándar.

fórmula de cuadratura, aplicada a la función t^n , y tomada en el intervalo $[0, k+1]$.

Dado que el integrando es monótono, podemos aplicar algunos resultados estándar de cuadratura teoría (véase el Thm97 en (Brunner & Houwen, 1986)), para encontrar que este término está acotado por la variación total del integrando, a saber la cantidad $(k+1)^n$. De este modo,

$$\begin{aligned}
& \left| \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} z(t) dt - \sum_{j=0}^k b_{j,k+1} z(t_j) \right| \\
& \leq \|z'\|_\infty \frac{h^{n+1}}{n} (k+1)^n
\end{aligned}$$

De manera similar, para probar (b), usamos la monotonidad de z en (C.20) y derivamos.

$$\begin{aligned}
& \left| \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} z(t) dt - \sum_{j=0}^k b_{j,k+1} z(t_j) \right| \\
& \leq \sum_{j=0}^k |z(t_{j+1}) - z(t_j)| \int_{jh}^{(j+1)h} (t_{k+1} - t)^{n-1} dt
\end{aligned}$$

$$\begin{aligned}
&= \frac{h^{n+p}}{n} \sum_{j=0}^k ((j+1)^p - j^p) ((k+1-j)^n - (k-j)^n) \\
&\leq \frac{h^{n+p}}{n} \left((k+1)^n - k^n + (k+1)^p - k^p \right. \\
&\quad \left. + pn \sum_{j=1}^{k-1} j^{p-1} (k-j+q)^{n-1} \right) \\
&\leq \frac{h^{n+p}}{n} \left(n(k+q)^{n-1} + pk^{p-1} + pn \sum_{j=1}^{k-1} j^{p-1} (k-j+q)^{n-1} \right)
\end{aligned}$$

mediante aplicaciones adicionales del teorema del valor medio. Aquí $q = 0$ si $n \leq 1$, y $q = 1$ en caso contrario. En cualquier caso, un breve análisis asintótico utilizando el método de Euler-Fórmula de MacLaurin produce que el término entre paréntesis está acotado desde arriba por $C_{n,p}^{Re} (k+1)^{p+n-1}$ donde $C_{n,p}^{Re}$ es una constante que depende de n y p pero no de k .

A continuación, llegamos al resultado correspondiente para la fórmula trapezoidal del producto que hemos utilizado para el corrector. La demostración de este teorema es muy similar a la demostración del Teorema C.1; por lo tanto, omitimos los detalles.

Teorema C.2.

- a) Si $z \in C^2[0, T]$ entonces hay una constante C_n^{Tr} dependiendo sólo de n tal que

$$\begin{aligned}
&\left| \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} z(t) dt - \sum_{j=0}^k a_{j,k+1} z(t_j) \right| \\
&\leq C_n^{Tr} \|z''\|_{\infty} t_{k+1}^n h^2
\end{aligned}$$

- b) Sea $z \in C^1[0, T]$ y supongamos que z cumple una condición de Lipschitz de orden μ para algunos $\mu \in (0, 1)$. Entonces existen constantes positivas $C_{n,\mu}^{Tr}$ (dependiendo sólo de n y μ) y $M(z, \mu)$ (dependiendo sólo de z y μ) tales que

$$\begin{aligned}
&\left| \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} z(t) dt - \sum_{j=0}^k a_{j,k+1} z(t_j) \right| \\
&\leq B_{n,\mu}^{Tr} M(z, \mu) t_{k+1}^n h^{1+\mu}
\end{aligned}$$

c) Sea $z(t) = t^p$ para algún $p \in (0,2)$ y $\rho := \min(2, p+1)$. Entonces

$$\left| \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} z(t) dt - \sum_{j=0}^k a_{j,k+1} z(t_j) \right| \leq C_{n,p}^{Tr} t_{k+1}^{n+p} h^p$$

Donde $C_{n,p}^{Tr}$ es una constante que depende sólo de n y p .

Observación C.2.

Observe que en el inciso (c) del teorema C.2 puede suceder que $n < 1$ y $p < 1$. Esto implica $\rho = p+1$. Por tanto, el exponente de t_{k+1} en el lado derecho de la desigualdad es igual a $n-1$ que es negativo. A primera vista esto puede parecer Contraintuitivo porque significa que el error de integración general se vuelve mayor. si el tamaño del intervalo de integración se vuelve más pequeño. La explicación para esto El fenómeno es que al hacer t_{k+1} más pequeño no sólo acortamos la longitud de el intervalo de integración (lo que debería conducir a un error menor) pero también cambiamos la función de peso de una manera que hace que la integral sea más difícil, y esta segunda característica conduce a un aumento en el error.

Se puede hacer una observación similar en el teorema C.1 (b).

Presentamos ahora los principales resultados relacionados con el error de nuestro esquema de Adams. Él Es útil distinguir varios casos. Específicamente, veremos que la precisión El comportamiento del error difiere dependiendo de si $n < 1$ o $n > 1$. Además, las propiedades de suavidad de la función dada f y la solución desconocida y play, un rol importante. A la vista de los resultados del art. 6.4, encontramos que la suavidad de una de estas funciones implicará la falta de suavidad de la otra a menos que se requiera alguna se cumplen las condiciones. Por lo tanto, también investigaremos el error bajo esos dos diferentes supuestos de suavidad.

Con base en las estimaciones de error anteriores, primero presentaremos una convergencia general. resultado del método Adams-Bashforth-Moulton. En los teoremas siguientes veremos especializar este resultado en casos especiales particularmente importantes.

Lema C.3.

Suponga que la solución y del problema de valor inicial es tal que

$$\left| \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} D_{*0}^n y(t) dt - \sum_{j=0}^k b_{j,k+1} D_{*0}^n y(t_j) \right| \leq C_1 t_{k+1}^{\gamma_1} h^{\delta_1}$$

Y

$$\left| \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} D_{*0}^n y(t) dt - \sum_{j=0}^k a_{j,k+1} D_{*0}^n y(t_j) \right| \leq C_2 t_{k+1}^{\gamma_2} h^{\delta_2}$$

con algunos $\gamma_1, \gamma_2 \geq 0$ y $\delta_1, \delta_2 \geq 0$. Entonces, para algunos $T > 0$ adecuadamente elegidos, tenemos

$$\max_{0 \leq j \leq N} |y(t_j) - y_j| = O(h^q)$$

donde $q = \min\{\delta_1 + n, \delta_2\}$ y $N = T/h$.

Prueba.

Demostraremos que, para h suficientemente pequeño,

$$|y(t_j) - y_j| \leq C h^q, \dots (C,21)$$

para todo $j \in \{0, 1, \dots, N\}$, donde C es una constante adecuada. La prueba se basará en inducción matemática. En vista de la condición inicial dada, la base de inducción ($j = 0$) se presupone. Ahora supongamos que (C.21) es verdadera para $j = 0, 1, \dots, k$ para algunos $k \leq N - 1$. Entonces debemos demostrar que la desigualdad también se cumple para $j = k+1$. Hacer Para esto, primero miramos el error del predictor Y_{k+1}^p . Por construcción del predictor. encontramos eso

$$|y(t_{j+1}) - y_{k+1}^p| = \frac{1}{\Gamma(n)} \left| \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} f(t, y(t)) dt - \sum_{j=0}^k b_{j,k+1} f(t_j, y_j) \right|$$

$$\begin{aligned}
&\leq \frac{1}{\Gamma(n)} \left| \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} D_{*0}^n y(t) dt - \sum_{j=0}^k b_{j,k+1} D_{*0}^n y(t_j) \right| \\
&\quad + \frac{1}{\Gamma(n)} \sum_{j=0}^k b_{j,k+1} |f(t, y(t)) - f(t_j, y_j)| \\
&\leq \frac{C_1 t_{k+1}^{\gamma_1} h^{\delta_1}}{\Gamma(n)} + \frac{1}{\Gamma(n)} \sum_{j=0}^k b_{j,k+1} L C h^q \\
&\leq \frac{C_1 T^{\gamma_1}}{\Gamma(n)} h^{\delta_1} + \frac{CL T^{\gamma_1}}{\Gamma(n+1)} h^q, \dots (C.22)
\end{aligned}$$

Aquí hemos utilizado la propiedad de Lipschitz de f , el supuesto sobre el error de la fórmula del rectángulo y los hechos que, mediante la construcción de la fórmula de la cuadratura subyacente al predictor, $b_{j,k+1} > 0$ para todos j y k

$$\sum_{j=0}^k b_{j,k+1} = \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} dt = \frac{1}{n} t_{k+1}^n \leq \frac{1}{n} T^n$$

Sobre la base del límite (C.22) para el error del predictor comenzamos el análisis del error del corrector. Recordamos la relación (C.14) que utilizaremos en particular para $j = k+1$ y encuentre, argumentando de manera similar a lo anterior, que

$$\begin{aligned}
|y(t_{k+1}) - y_{k+1}| &= \left| \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} f(t, y(t)) dt - \sum_{j=0}^k a_{j,k+1} f(t_j, y_j) - a_{k,k+1} f(t_{k+1}, y_{k+1}^p) \right| \\
&\leq \frac{1}{\Gamma(n)} \left| \int_0^{t_{k+1}} (t_{k+1} - t)^{n-1} D_{*0}^n y(t) dt - \sum_{j=0}^k a_{j,k+1} D_{*0}^n y(t_j) \right| \\
&\quad + \frac{1}{\Gamma(n)} \sum_{j=0}^k a_{j,k+1} |f(t_j, y(t_j)) - f(t_j, y_j)| \\
&\quad + \frac{1}{\Gamma(n)} a_{k+1,k+1} |f(t_{k+1}, y(t_{k+1})) - f(t_{k+1}, y_{k+1}^p)| \\
&\leq \frac{C_2 t_{k+1}^{\gamma_2} h^{\delta_2}}{\Gamma(n)} + \frac{CL}{\Gamma(n)} h^q \sum_{j=0}^k a_{j,k+1} \\
&\quad + a_{k+1,k+1} \frac{L}{\Gamma(n)} \left(\frac{C_1 T^{\gamma_1}}{\Gamma(n)} h^{\delta_1} + \frac{CL T^n}{\Gamma(n+1)} h^q \right)
\end{aligned}$$

$$\leq \left(\frac{C_2 T^{\gamma_2}}{\Gamma(n)} + \frac{CL T^n}{\Gamma(n+1)} + \frac{C_1 T^{\gamma_1}}{\Gamma(n)\Gamma(n+2)} + \frac{CL^2 T^n}{\Gamma(n+1)\Gamma(n+2)} h^n \right) h^q$$

en vista de la no negatividad de γ_1 y γ_2 y las relaciones $\delta_2 \leq q$ y $\delta_1 + n \leq q$. Al elegir T lo suficientemente pequeño, podemos asegurarnos de que el segundo sumando en el paréntesis está delimitado por $C/2$. Una vez fijado este valor para T , podemos hacer la suma de las expresiones restantes entre paréntesis menores que $C/2$ también (por suficientemente pequeño h) simplemente eligiendo C suficientemente grande. Es entonces obvio que todo el límite superior no excede $C h^q$.

Como primera aplicación de este Lema asumimos que los datos dados son tales que la solución y en sí misma es suficientemente diferenciable. Como se mencionó anteriormente, el resultado depende sobre si $n > 1$ o $n < 1$.

Teorema C.4.

Sea $0 < n$ y supongamos $D_{*0}^n y \in C^2[0, T]$ para alguna T adecuada. Entonces

$$\max_{0 \leq j \leq N} |y(t_j) - y_j| = \begin{cases} O(h^2), & \text{si } n \geq 1 \\ O(h^{1+n}), & \text{si } n < 1 \end{cases}$$

Antes de pasar a la prueba, observamos un punto en particular: el orden de convergencia depende de n , y es una función no decreciente de n . Esto se debe al hecho que discretizamos el operador integral en (C.1) que se comporta más suavemente (y por lo tanto, se puede aproximar con mayor precisión) a medida que n aumenta. En cambio, los llamados Los métodos directos como el método de diferenciación hacia atrás de Diethelm, K utilizan un método diferente. acercarse; como su nombre indica, discretizan directamente el operador diferencial en el problema de valor inicial dado. Las propiedades de suavidad de tales operadores (y (de ahí la facilidad con la que pueden aproximarse) se deterioran a medida que n aumenta, y entonces encontramos que el orden de convergencia del método de Diethelm, K es no creciente función de n ; en

particular, no se logra convergencia allí para $n \geq 2$. Es un distintivo La ventaja del esquema de Adams presentado aquí es que converge para todo $n > 0$.

Observación C.3.

Recuperamos formalmente el límite de error (C.10) si establecemos $n = 1$.

Demostración (del teorema C.4). En vista de los teoremas C.1 y C.2, podemos aplicar Lema C.3 con $\gamma_1 = \gamma_2 = n > 0$, $\delta_1 = 1$ y $\delta_2 = 2$. Así, definiendo

$$q = \min\{1 + n, 2\} = \begin{cases} 2, & \text{si } n \geq 1 \\ 1 + n, & \text{si } n < 1 \end{cases}$$

encontramos un límite de error $O(h^q)$.

Tenga en cuenta que, en cierto sentido, el teorema anterior trata de la situación "óptima": La función que aproximamos en nuestro proceso es $f(\cdot, y(\cdot)) = D_{*0}^n$. Con el fin de obtener límites de error muy buenos, debemos asegurarnos de que los errores de cuadratura para esta función son (asintóticamente) lo más pequeños posible. Una condición suficiente para esto es, como es bien sabido por la teoría de la cuadratura, que esta función está en C2 sobre el intervalo de integración. Este es precisamente el escenario discutido en el teorema C.4. Entonces, este teorema nos muestra qué tipo de desempeño puede ofrecer el método Adams. dar en circunstancias óptimas, y también establece condiciones suficientes para tal resultado para mantener.

Una objeción contra el uso de este esquema Adams-Bashforth-Moulton puede ser la muy lenta tasa de convergencia si n es cercano a 0. Sin embargo, una inspección cuidadosa de la prueba del límite del error revela un hecho que es bien conocido en el análisis del error para métodos de esta estructura para ecuaciones de primer orden. La aplicación de la fórmula correctora mejora la precisión de su entrada (el predictor) en un factor de h^n hasta que se alcanza un orden de $O(h^2)$ (es decir, una saturación). Así podremos reemplazar la estructura PECE simple mediante un método $P(EC)^{\mu}E$, es decir, introduciendo un corrector adicional iteraciones.

Observación C.4.

Una observación interesante aquí es que, al elegir un mayor número de iteraciones del corrector, esencialmente dejamos la complejidad computacional sin cambios.

Una iteración del corrector tiene la forma

$$y_{j+1}^{|\ell|} = \sum_{r=1}^{[n]-1} \frac{t_{j+1}^n}{r!} y_0^{(r)} + \frac{h^n}{\Gamma(n+2)} f(t_{j+1}, y_{j+1}^{|\ell|}) \\ + \frac{h^n}{\Gamma(n+2)} \sum_{r=0}^j a_{j,k+1} f(t_r, y_r)$$

cf. (C.15). Aquí $y_{j+1}^{|\ell|}$ denota la aproximación después de los pasos del corrector, $y_{j+1}^{|\ell|} = y_{j+1}^p$ es el predictor, y, $y_{j+1}^{|\mu|}$ es la aproximación final después de los pasos del corrector μ que realmente utilizamos. Podemos reescribir esto como

$$y_{j+1}^{|\ell|} = \beta_{j+1} + \frac{h^n}{\Gamma(n+2)} f(t_{j+1}, y_{j+1}^{|\ell-1|})$$

Donde

$$\beta_{j+1} = \sum_{r=1}^{[n]-1} \frac{t_{j+1}^n}{r!} y_0^{(r)} + \frac{h^n}{\Gamma(n+2)} \sum_{r=0}^j a_{r,j+1} f(t_r, y_r)$$

es independiente. Así, la complejidad aritmética total de la parte correctora del $(j+1)$ el primer paso (que nos lleva de t_j a t_{j+1} es $O(j)$ para el cálculo de β_{j+1} más $O(\mu)$ para los pasos del corrector μ , que (dado que μ es constante) es asintóticamente lo mismo que la complejidad en el caso $\mu = 1$.

Para el error del esquema descrito en la Observación C.4 encontramos, como se indicó anteriormente, por una aplicación repetida de las consideraciones de la demostración del Teorema C.4 (ver (Diethelm, s. f.) para más detalles):

Teorema C.5.

Bajo los supuestos del Teorema C.4, la aproximación obtenida por el método $P(EC)^\mu$ E descrito anteriormente satisface

$$\max_{0 \leq j \leq N} |y(t_j) - y_j| = O(h^q),$$

Donde $q = \min\{2, 1 + \mu\}$.

Por el momento dejamos el tema de los métodos generales $P(EC)^\mu E$ en favor de una investigación más detallada de su caso especial $\mu = 1$, es decir, el Adams original– Método Bashforth–Moulton introducido en (C.19) y (C.15). En este contexto nosotros observamos una aparente desventaja en la formulación de las hipótesis de los teoremas arriba: se expresan en términos de la solución y (o, más precisamente, su derivada de Caputo de orden n), que se desconoce en general. Aunque a veces es posible determinar las propiedades de suavidad de $D_{*0}^n y$ de los datos dados, hay Todavía existe cierta necesidad de una teoría del error correspondiente para el método de Adams supuestos formulados directamente en términos de los datos dados, es decir, en términos de función f . Dichos resultados se derivarán más adelante en esta sección.

Sin embargo, antes de llegar a esos resultados, queremos brindar más información. bajo supuestos similares a los del Teorema C.4. Específicamente queremos afirmar la conjetura de que el error de nuestro esquema, tomado en una abscisa fija, posee una expansión asintótica en potencias del tamaño de paso h bajo suavidad adicional condiciones en $D_{*0}^n y$. Si esto fuera cierto (como, por ejemplo, los resultados numéricos del Ejemplo C.1 a continuación se indica), podríamos construir un algoritmo de extrapolación de Richardson basado en el método Adams de una manera similar a la construcción en (Diethelm & Walz, 1997) que se basa en el esquema de (Diethelm, 1997).

El uso de este procedimiento de extrapolación nos permitiría entonces obtener aproximaciones numéricas más precisas para la solución deseada.

Conjetura C.I.

Sea $n > 0$ y supongamos que $D_{*0}^n y \in C^k[0T]$ para algunos $k \geq 3$ y alguna T adecuada. Entonces,

$$y(T) - y_T = \sum_{j=1}^{k_1} c_j h^{2j} + \sum_{j=2}^{k_2} d_j h^{j+n} + O(h^{k_3})$$

donde k_1, k_2 y k_3 son ciertas constantes que dependen sólo de k y satisfacen $k_3 > \max(2k_1, k_2 + n)$.

Observe que la expansión asintótica comienza con un término h^2 y continúa con h^{1+n} para $1 < n < 3$, mientras que comienza con h^{1+n} , seguido de h^2 , para $0 < n < 1$.

Ejemplo C.1.

Considere la ecuación

$$D_{*0}^n y(x) = \frac{40320}{\Gamma(9-n)} x^{8-n} - 3 \frac{\Gamma(5+n/2)}{\Gamma(5-\frac{n}{2})} x^{4-\frac{n}{2}} + \frac{9}{4} \Gamma(n+1) \\ + \left(\frac{3}{2} x^{\frac{n}{2}} - x^4 \right)^3 - [y(x)]^{3/2}$$

para $x \in [0,1]$ con condiciones iniciales homogéneas ($y(0) = 0$, $y'(0) = 0$; este último sólo en el caso $n > 1$).

La solución exacta de este problema de valor inicial es

$$y(x) = x^8 - 3x^{4+\frac{n}{2}} + \frac{9}{4} x^n$$

y por lo tanto

$$D_{*0}^n y(x) = \frac{40320}{\Gamma(9-n)} x^{8-n} - 3 \frac{\Gamma(5+n/2)}{\Gamma(5-\frac{n}{2})} x^{4-\frac{n}{2}} + \frac{9}{4} \Gamma(n+1)$$

es decir, $D_{*0}^n y \in C^2[0,1]$ si $n \leq 4$, y por tanto se cumplen las condiciones del Teorema C.4. Además, suponiendo que se cumpla la conjetura C.1, la aplicación de la extrapolación de Richardson también está justificado.

Mostramos algunos de los resultados en las Tablas C.1 y C.2. donde, por ejemplo, la notación $-5,53(-3)$ representa -5.53×10^{-3} . En cada caso, el extremo izquierdo la columna muestra el tamaño de paso utilizado, la siguiente columna muestra el error de muestreo:

Tabla C.1. Errores para el ejemplo C.1 con $n = 1,25$, tomado en $x = 1$

Numero de paso	Error del Esquema de Adams del tamaño del paso	Valores extrapolados			
1/10	-5.53(-3)				
1/20	-1.59(-3)	-2.80(-4)			
1/40	-4.33(-4)	-2.460(-5)	1.64(-5)		
1/80	-1.14(-4)	-8.17(-6)	1.90(-6)	2.13(-7)	
1/160	-2.98(-5)	-1.54(-6)	2.24(-7)	2.71(-8)	1.47(-8)
1/320	-7.66(-6)	-3.04(-7)	2.56(-8)	2.28(-9)	6.24(-10)
1/640	-1.96(-6)	-6.16 (-8).	2.85(-9)	1.73(-10)	3.25(-11)
COE	1.97	2.30	3.17	3.72	4.26

Nota: Basada por el ejemplo C.1 con $n = 1,25$, tomado en $x = 1$

Tabla C.2. Errores para el ejemplo C.1 con $n = 0.25$, tomado en $x = 1$

Numero de paso	Error del Esquema de Adams del tamaño del paso	Valores extrapolados			
1/10	2.5(-1)				
1/20	.181(-2)	-1.50 (-1)			
1/40	3.61(-3)	-6.91 (-2)	4.04(-2)		
1/80	1.45(-3)	-1.10 (-3)	2.16 (-3)	-8.15 (-3)	
1/160	6.58(-4)	-8.19 54(-4)	1.46 (-4)	- 3.89 (-3)	1.28 (-4)
1/320	2.97(-4)	-3.49 (-5)	1.92 (-5)	- 1.45,(-5)	1.05 (-5)
1/640	1.31(-4)	-1.12 (-5).	3.37 (-6)	- 850(-7)	6.01 (-8)
COE	1.18	1.63	2.51	4.09	7.44

Nota: Basada por el ejemplo C.1 con $n = 0.25$, tomado en $x = 1$ esquema en $x = 1$, y las columnas posteriores dan los valores extrapolados. El fondo La línea (marcada “EOC”) indica el orden de convergencia determinado experimentalmente para cada una de las columnas a la derecha de la tabla. Según nuestras consideraciones teóricas, estos valores deberían ser $1+n$, 2 , $2+n$, $3+n$, 4 , $4+n$, $\dots$ en el caso $0 < n < 1$ y 2 , $1+n$, $2+n$, 4 , $3+n$, $4+n$, $\dots$ para $1 < n < 2$. Los datos numéricos en las siguientes tablas muestran que estos valores se reproducen aproximadamente al menos para $n > 1$ (ver Tabla C.1). En el caso $0 < n < 1$, mostrado en la Tabla C.2, la situación parece ser menos obvio. Aparentemente, necesitamos usar valores mucho más pequeños para h que en el caso $n > 1$ antes de que podamos ver que el comportamiento asintótico realmente establece in. Esto normalmente correspondería a la situación en la que los coeficientes de los términos principales son de magnitud pequeña en comparación con los coeficientes de los términos de orden superior.

Nuestra creencia en la verdad de la conjetura C.1 no sólo está respaldada por los resultados numéricos sino también por los resultados de Hoog y Weiss quienes muestran que las expansiones asintóticas de esta forma son válidas si utilizamos la fórmula fraccionaria de Adams-Moulton (Galeone & Garrappa, 2008). método (es decir, si resolvemos exactamente

la ecuación del corrector) y que se puede derivar una expansión similar para el método fraccionario de Adams-Bashforth (usando el predictor como la aproximación final en lugar de corregir una vez con Adams-Moulton fórmula (Galeone & Garrappa, 2008)). Sin embargo, por el momento dejamos la cuestión de la influencia del paso corrector (que combina los dos enfoques) en esta expansión abierta.

Más bien, dirigimos nuestra atención a otro problema relacionado. En los teoremas anteriores Habíamos formulado nuestras hipótesis en forma de supuestos de suavidad en $D_{*0}^n y$. Ahora queremos reemplazar esto por suposiciones similares sobre la propia y . En vista del teorema 4.15 debemos ser conscientes del hecho de que la suavidad de y en general implica la no suavidad de $D_{*0}^n y$ (la función que tenemos que aproximar), por lo que algunas dificultades

es probable. Afortunadamente, el teorema 4.15 también nos informa sobre la naturaleza precisa de la singularidad en las derivadas de $D_{*0}^n y$. Podemos explotar esta información para obtener los siguientes resultados.

Teorema C.6.

Sea $n > 1$ y supongamos que $y \in C^{1+[n]}[0, T]$ para algún T adecuado. Entonces,

$$\max_{0 \leq j \leq N} |y(t_j) - y_j| = O(h^{1+[n]-n}),$$

Prueba.

Por el teorema 4.15 encontramos que $D_{*0}^n y(x) = cx^{[n]-n} + g(x)$ donde $g \in C^1[0, T]$ y G cumple una condición de Lipschitz de orden $[n] - n$. Así, según Teoremas C.1 y C.2 podemos aplicar el Lema C.3 con $\gamma_1 = 0, \gamma_2 = n - 1 > 0, \delta_1 = 1, y \delta_2 = 1 + [n] - n$. Como $n > 1$ encontramos que $\delta_1 + n = 1 + n > 2 > \delta_2$, y por tanto $\min\{\delta_1 + n, \delta_2\} = \delta_2$. Entonces el error general está limitado por $O(h^{\delta_2})$.

Observe que una reformulación del Teorema C.6 produce que, si $1 < n = k_1 + k_2$ con $k_1 \in \mathbb{N}$ y $0 < k_2 < 1$, entonces el error es

$O(h^{2-k_2})$. Así, la parte fraccionaria de n juega el papel decisivo para el orden del error. En particular, encontramos una lenta convergencia. si la parte fraccionaria de n es grande. En consecuencia, bajo estos supuestos no podemos esperar que el orden de convergencia sea una función monótona de n . Sin embargo, podemos demostrar que el método converge para todo $n > 0$:

Teorema C.7.

Sea $0 < n < 1$ y suponga que $y \in C^2[0, T]$ para algún T adecuado.

Entonces, para $1 \leq j \leq N$ tenemos

$$|y(t_j) - y_j| \leq C t_j^{n-1} \times \begin{cases} h^{1+n}, & \text{si } 0 < n < 1/2 \\ h^{2-n}, & \text{si } 1/2 \leq n < 1 \end{cases}$$

donde C es una constante independiente de j y h .

Obtenemos dos consecuencias inmediatas.

Corolario C.8.

Bajo los supuestos del Teorema C.7, tenemos

$$\max_{0 \leq j \leq N} |y(t_j) - y_j| = \begin{cases} O(h^{2n}), & \text{si } 0 < n < 1/2 \\ O(h), & \text{si } 1/2 \leq n < 1 \end{cases}$$

Además, para cada $\varepsilon \in (0, T)$ tenemos

$$\max_{t_j \in [\varepsilon, T]} |y(t_j) - y_j| = \begin{cases} O(h^{1+n}), & \text{si } 0 < n < 1/2 \\ O(h^{2-n}), & \text{si } 1/2 \leq n < 1 \end{cases}$$

Demostración (del teorema C.7).

Los primeros pasos de la prueba son como en la prueba de Teorema C.6. La diferencia clave es que ahora $\gamma_2 < 0$ (tenga en cuenta que todavía tenemos $\gamma_2 = n - 1$, pero ahora $n < 1$). Por tanto, no podemos aplicar el Lema C.3. En lugar de eso modificamos su prueba para que se ajuste a nuestros requisitos: mantenemos la estructura inductiva y recordamos que nuestra afirmación ahora es (C.23) en lugar de (C.21). Con este cambio en la inducción hipótesis procedemos de manera muy similar a la prueba del Lema C.3. Sin embargo, debido a esta nueva

hipótesis, ahora tenemos que estimar términos de la forma $\sum_{j=1}^{k-1} b_{j,k+1} t_j^{\gamma_2}$ y $\sum_{j=1}^{k-1} a_{j,k+1} t_j^{\gamma_2}$. Por el teorema del valor medio tenemos $0 \leq b_{j,k+1} \leq h^n (k-j)^{n-1}$ y $0 \leq a_{j,k+1} \leq ch^n (k-j)^{n-1}$ para $1 \leq j \leq k-1$ (donde la constante c es independiente de j y k), respectivamente, de modo que el problema se reduce a encontrar un límite para $S_k := \sum_{j=1}^{k-1} j^{\gamma_2} (k-j)^{n-1}$. Bajo nuestras suposiciones, tanto los exponentes γ_2 como $n-1$ están en el intervalo $(-1, 0)$, y luego se ve fácilmente que $S_k = (k^{\gamma_2+n})$. Usando esta relación podemos completar la demostración del Teorema C.7 siguiendo las líneas del resto de la prueba del Lema C.3.

Concluimos la discusión sobre los límites de error con un resultado en el que formulamos la hipótesis en términos de los datos dados y no en términos de la solución desconocida. Nosotros damos un resultado en el caso $n > 1$ y luego discutimos las propiedades del esquema numérico cuando $n < 1$.

Teorema C.9.

Sea $n > 1$. Entonces, si $f \in C^2(G)$,

$$\max_{0 \leq j \leq N} |y(t_j) - y_j| = O(h^2),$$

Prueba (del teorema C.9).

Comenzamos discutiendo el caso $n \geq 2$. Luego, de acuerdo con Teorema 7.25, encontramos que $y \in C^2[0, T]$. Por lo tanto, en vista del supuesto de suavidad sobre f y la regla de la cadena, $D_{*0}^n y: f(\cdot, y(\cdot)) \in C^2[0, T]$ también, y la afirmación sigue por virtud del Teorema C.4.

Para el caso $1 < n < 2$, queremos aplicar el Lema C.3 y por lo tanto tenemos que determinar las constantes $\gamma_1, \gamma_2, \delta_2$, y δ_1 en sus hipótesis. Para ello necesitamos información más precisa sobre el comportamiento de y . Esta información se puede encontrar en el Teorema 7.38 que afirma que $y(x) = cx^n + \Psi(x)$ con algún $c \in \mathbb{R}$ y algún $\Psi \in C^2[0, T]$. Esto implica, en particular, que $y \in C^1[0, T]$. Como en el caso $n > 2$ anterior entonces podemos deducir $D_{*0}^n y \in C^1[0, T]$ también, y por el Teorema C.1(a), encontramos que podemos elegir $\gamma_1 = n$ y $\delta_1 = 1$. Además, usando

nuevamente el hecho de que $y(x) = cx^n + \Psi(x)$ con algunos $c \in \mathbb{R}$ y algunos $\Psi \in C^2[0, T]$ y aplicando el Teorema C.2(a) y (c), tenemos determinado los valores correctos para las cantidades restantes como $\gamma_2 = \min\{n, 2n - 2\} = 2n - 2 \geq 0$, $y\delta_2 = 2$. La afirmación se deriva entonces del Lema C.3.

En el caso $n < 1$ la situación parece menos clara. Según los teoremas presentado al final de la Sección. 7.4, las condiciones de suavidad en f implican que la precisión la solución es de la forma

$$y(t) = \Psi(t) + \sum_{v=1}^{\hat{v}} c_v t^{vn} + \sum_{v=1}^{\tilde{v}} d_v t^{1+vn}$$

donde ψ es dos veces diferenciable. La primera suma consta de términos que no son diferenciables, y la segunda suma está formada por términos que son diferenciables una vez, pero no dos. Como señaló Lubich (Oldham, 1974), parece poco probable que se utilicen esquemas numéricos. rápidamente convergente en cualquier intervalo que contenga el origen.

De hecho, podemos probar que el error $y(t_1) - y_1$ de la aproximación después de un solo paso se comporta como $O(h^{2n})$ si $f \in C^2[0, T]$. Experimentos numéricos simples indican que este resultado no puede mejorarse. Sin embargo, este error introducido en la fase inicial es transitorio y desde resultados numéricos informados en (Ahmad & El-Khazali, 2007) creemos que la siguiente conjetura, a decir verdad.

Conjetura C.2.

Sea $0 < n < 1$. Entonces, si $f \in C^3[0, T]$, para cada $\varepsilon > 0$ tenemos

$$\max_{t_j \in [\varepsilon, T]} |y(t_j) - y_j| = O(h^{1+n})$$

C.2 Esquemas numéricos para ecuaciones de términos múltiples

Llegamos ahora a la extensión de los métodos numéricos discutidos en el punto anterior. sección de ecuaciones de varios

términos. Las propiedades teóricas más importantes de estas ecuaciones multitérminos se han analizado en el capítulo 8. Como en ese capítulo restringimos nuestra atención a las ecuaciones de la forma

$$D_{*0}^n y(x) = f(x, y(x), D_{*0}^{n_1} y(x), D_{*0}^{n_2} y(x), \dots, D_{*0}^{n_{k-1}} y(x), \dots) \quad (C.24.a)$$

donde $0 < n_1 < n_2 < \dots < n_k$ con una función adecuada f y condiciones iniciales

$$y^{(j)}(0) = y_0^{(j)}, j = 0, 1, 2, \dots, [n_k] - 1, \dots (C.24.b)$$

Para problemas de valor inicial de este tipo discutiremos los diversos enfoques presentados en el capítulo 8 y descubres sus respectivas ventajas y desventajas.

Conversión a sistemas de orden único

Nuestro primer intento consiste en una aplicación directa del resultado del Teorema 8.1 o 8.2. (dependiendo de si $n_k > 1$ o no) al problema de valor inicial dado. De este modo

transformamos el problema de valor inicial dado en un sistema de ecuaciones de la forma

$$\begin{aligned} D_{*0}^\gamma y_0(x) &= y_1(x), \\ D_{*0}^\gamma y_1(x) &= y_2(x), \\ &\vdots \\ D_{*0}^\gamma y_{N-2}(x) &= y_{N-1}(x) \\ D_{*0}^\gamma y_{N-1}(x) &= f(x, y_0(x), y_{\frac{n_1}{\gamma}}(x), \dots, y_{\frac{n_{k-1}}{\gamma}}(x)), \dots (C.25.a) \end{aligned}$$

junto con las condiciones iniciales

$$y_j(0) = \begin{cases} y_0^{(j\gamma)} & \text{si } j\gamma \in \mathbb{N}_0 \\ 0, & \text{en otros caso} \end{cases}, \dots (C.25b)$$

siendo la elección precisa de los nuevos parámetros γ , N de acuerdo con los teoremas 9.1 u 9.2, según corresponda. De esta manera hemos obtenido formalmente una ecuación del tipo

$$D_{*0}^{\gamma} Y(x) = F(x, Y(x)), Y(0) = Y_0, \dots \text{(C.26)}$$

con ciertas funciones vectoriales F (conocida) e Y (desconocida) y una inicial vector de condición Y_0 , es decir, una ecuación de un solo término de orden γ con datos valorados por vectores. Así, podemos aplicar cualquier método numérico para tales ecuaciones de un solo término y calcular una solución aproximada para este sistema; En aras de la simplicidad, Nos limitaremos al esquema de Adams-Bashforth-Moulton desarrollado anteriormente. El primer componente del vector de solución (numérico) (con índice 0) es entonces la requerida solución aproximada de la ecuación original. Ilustramos el procedimiento con un ejemplo sencillo tomado de (Diethelm & Ford, 2002b).

Ejemplo C.2.

Resuelve la ecuación de Bagley-Torvik

$$Ay''(x) + BD_{*0}^{\frac{3}{2}}y(x) + Cy(x) = C(x + 1)$$

(donde $A = 0$ y $B, C \in \mathbb{R}$) con condiciones iniciales

$$y(0) = y'(0) = 1$$

con el enfoque descrito anteriormente.

Se puede comprobar fácilmente que la solución exacta de este problema es

$$y(x) = x + 1$$

independiente de la elección de los coeficientes A, B y C . Por lo tanto, el resultado el sistema es

$$D_{*0}^{\frac{1}{2}} \begin{pmatrix} y_0(x) \\ y_1(x) \\ y_2(x) \\ y_3(x) \end{pmatrix} = \frac{1}{A} \begin{pmatrix} y_1(x) \\ y_2(x) \\ y_3(x) \\ -By_3(x) - Cy_0(x) + C(x + 1) \end{pmatrix}$$

$$\begin{pmatrix} y_0(0) \\ y_1(0) \\ y_2(0) \\ y_3(0) \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \\ 1 \\ 0 \end{pmatrix}$$

Hemos resuelto este problema en el intervalo $[0,5]$ con el método Adams-Bashforth- Esquema Moulton. Los resultados numéricos en $x = 5$ fueron los que se muestran en la Tabla C.3.

Estos datos indican que el comportamiento de convergencia es $O(h^{3/2})$. Comprende cómo se relaciona esto con las estimaciones de error de la Sección C.1, debemos recordar que ahora construyamos una solución aproximada para todo el sistema y no sólo para su primer componente.

Así vemos que, como subproducto del método, no sólo obtenemos información sobre “y”, también sobre sus derivadas fraccionarias de orden $\gamma, 2\gamma, \dots, (N-1)\gamma$. Dependiendo de la tarea en cuestión, esta información puede ser muy útil o absolutamente innecesaria. En cualquier caso, la estimación del error está dominada por la peor estimación del error de los cuatro componentes.

Dado que la solución exacta para el sistema es $(x + 1, \frac{x^{\frac{1}{2}}}{\Gamma(\frac{3}{2})}, 1, 0)^T$ vemos que la segunda, tercera y cuarta componentes de la solución exacta tiene derivadas suaves de orden $1/2$; por tanto, pueden aproximarse con orden $O(h^{3/2})$ según el teorema C.4.

El primer componente es liso. sí mismo; permite una aplicación del Corolario C.8 que da una estimación del error $O(h)$ en el intervalo completo $[0,T]$ y una estimación $O(h^{3/2})$ en cada intervalo de la forma $[\varepsilon, T]$ con $\varepsilon > 0$.

Dado que el último caso cubre el problema considerado en la Tabla C.3, tenemos tener concordancia entre los resultados teóricos y numéricos.

Para explicar las debilidades de este concepto, veamos un segundo ejemplo, ya considerado en (Diethelm & Ford, s. f.).

Ejemplo C.3.

Resuelve la ecuación no lineal de tres términos.

$$D_{*0}^{1.455}y(x) = -x^{0.1} \frac{E_{1.545}(-x)}{E_{1.455}(-x)} e^x y(x) D_{*0}^{0.555}y(x) + e^{-2x} - [D_{*0}^1y(x)]^2,...(C.27)$$

para $0 \leq x \leq 1$, equipado con las condiciones iniciales $y(0) = 1$ y $y'(0) = -1$ con el mismo algoritmo.

La solución exacta de este problema es $y(x) = \exp(-x)$. Al aplicar nuestra idea a esta ecuación, primero necesitamos calcular el orden γ del nuevo sistema como se describe en el teorema 8.1. En nuestro caso el resultado es $\gamma = 1/200$, y por tanto la dimensión del sistema resultante es $N = \frac{1.455}{\gamma} = 291$, un número bastante grande. En un primer intento Hemos intentado resolver el sistema con el esquema Adams-Bashforth-Moulton como

Tabla C.3. Ecuación de Bagley-Torvik resuelta con el método de Adams

Tamaño del paso r	Solución numérica	Error	Orden estimada de convergencia
0.5	6.15131473519232 2	-0.1513147351923	
0.25	6.04684102179946	-0.04684102179946	1.69
0.125	6.01602947553912	-0.01602947553912	1.55
0.0625	6.00562770408881	-0.00562770408881	1.51

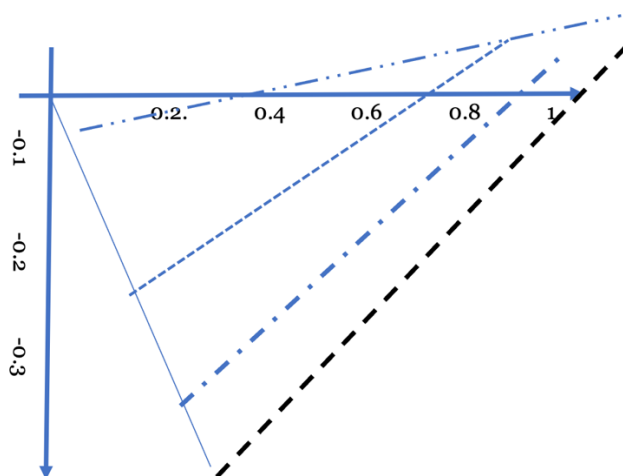
Nota: Utiliza el método de Adams

Tabla C.4. Resultados numéricos para el Ejemplo C.3 (sistema con $N = 291$, $\gamma = 0,005$) utilizando el tipo PECE estándar algoritmo de adams

Tamaño del paso	Error máximo	Tiempo de ejecución
1/200	0.3904	101.2 s
1/400	0.2193	368.4 s
1/800	0.1164	1358.0 s
1/1,600	0.0600	5017.4 s

Nota: Basada por el ejemplo C.3 (sistema con $N = 291$, $\gamma = 0,005$) utilizando el tipo PECE estándar algoritmo de Adams

Figura C.1. Errores de aproximación para el ejemplo C.3



Nota: Errores de aproximación para el ejemplo C.3 con $h = 1/100$ (línea continua), $h = 1/200$ (línea discontinua), $h = 1/400$ (línea de puntos y discontinuas), $h = 1/800$ (línea de puntos), $h = 1/1,600$ (línea de puntos y tres puntos) como anteriormente.

Los errores resultantes se presentan en la Tabla C.4 y en la Fig.C.1. A fin de En comparación con métodos posteriores, también hemos incluido información sobre el tiempo de ejecución. del algoritmo en una PC Pentium estándar de 500 MHz en aritmética de doble precisión.

Se ve claramente que los resultados para $h = 1/100$ y $h = 1/200$ son totalmente inaceptables. Hay una explicación sencilla para este fenómeno que se hace evidente cuando uno mira la solución numérica del problema de valor inicial y no en el error de aproximación: en cada caso, los primeros 99 pasos del algoritmo no cambia el primer componente de la solución aproximada.

En otras palabras, nos quedamos estancados en el valor inicial en lugar de seguir la solución exacta. La razón de este

comportamiento se puede encontrar en la estructura de la función F en el lado derecho del sistema (C.26) y de su condición inicial: El primer componente (índice 0) es y_0 , el componente con índice $200(=1/\gamma)$ es $y'_0 (\neq 0)$, y todos los componentes intermedios desaparecer.

La interacción del método Adams-Bashforth-Moulton con la función F ahora implica que el componente distinto de cero se propaga por una fila en cada paso del predictor y otra fila en cada paso del corrector, por lo que solo en los primeros 99 pasos se añaden un total de 198 ceros al valor inicial del componente principal de la solución numérica.

El último cero (199.^o) se utiliza luego en el predictor del número 100 pasos, y el corrector del paso 100 es en realidad la primera operación donde un valor distinto de cero, y la entrada llega al primer componente de la solución. Así siempre tenemos una inicial C.2 esquemas numéricos para ecuaciones de términos múltiples intervalo de $2/\gamma - 1$ pasos donde la solución numérica es constante antes de que pueda comenzar para avanzar hacia la solución exacta.

Observe que en nuestro ejemplo de Bagley-Torvik anterior teníamos $\gamma = 1/2$, por lo que (dado que aquí $2/\gamma - 1 = 3$) el efecto es insignificante. Esto significa que, si se quiere mantener la estructura del algoritmo y la uniformidad tamaño del paso, entonces la única forma de reducir este efecto es una reducción drástica del paso tamaño h , esencialmente por la regla $h = O(\gamma)$. Una mirada a la figura C.1 confirma esta observación.

Resumamos brevemente lo que hemos logrado hasta ahora: El enfoque descrito anterior tiene la ventaja de producir un sistema de ecuaciones con una muy simple estructura.

Como consecuencia de esta estructura, los esquemas numéricos para este sistema se pueden implementar en una arquitectura de computadora paralela de una manera bastante eficiente. Sin embargo, también hemos visto algunas desventajas:

- a) El método sólo funciona en el caso de ecuaciones proporcionales de términos múltiples (si $n_k \leq 1$) o bajo el

supuesto aún más restrictivo de que todos los n_j son racional (si $n_k > 1$)

- b) La dimensión del sistema puede ser muy grande (dependiendo de la precisión valores de los parámetros de la ecuación original); esto puede llevar a muy largo tiempos de ejecución en máquinas secuenciales
- c) La estructura de las condiciones iniciales del nuevo sistema puede ser problemática para algunos tipos de algoritmos numéricos; en particular, podemos vernos obligados a utilizar tamaños de paso excesivamente pequeños

Para superar estos problemas (al menos parcialmente), proponemos dos posibles estrategias.

El primero, tomado de (Diethelm & Ford, s. f.), es un ligero refinamiento de la idea utilizada hasta ahora. Hasta ahora eso funciona en dos etapas. Recordamos que el problema se debe principalmente a la gran dimensión del sistema, y esto a su vez es una consecuencia del tamaño del máximo común divisor de los órdenes de los operadores diferenciales en el sistema dado. Por lo tanto, introducimos una etapa de manipulaciones preliminares antes de comenzar. El algoritmo numérico.

Esta primera etapa consiste en sustituir el problema de valor inicial dado (C.24) por una nueva ecuación diferencial

$$D_{*0}^{\tilde{n}_k} \tilde{y}(x) = f(x, \tilde{y}(x), D_{*0}^{\tilde{n}_1} \tilde{y}(x), D_{*0}^{\tilde{n}_2} \tilde{y}(x), \dots, D_{*0}^{\tilde{n}_{k-1}} \tilde{y}(x)), \dots \quad (\text{C.28})$$

con condiciones iniciales idénticas (C.24b). De este modo perturbamos los órdenes de los operadores diferenciales, pero todos los demás parámetros del problema dado (la función f en el lado derecho y las condiciones iniciales) permanecen sin cambios.

La esencia de esta idea es que, según el teorema 8.8, la solución exacta $\tilde{y}$ de este nuevo problema de valor inicial y la solución exacta del problema original difieren solo por:

$$\|y - \tilde{y}\|_{\infty} = o\left(\max_{j=1,2,\dots,k} |n_j - \tilde{n}_j|\right)$$

Aquí por $\| \cdot \|_\infty$ denotamos la norma de Chebyshev tomada en un intervalo finito adecuado $[0, T]$, digamos, donde ambos problemas tienen solución.

Para explotar las capacidades de este enfoque, debemos elegir el nuevo parámetro $\tilde{n}_1, \tilde{n}_2, \dots, \tilde{n}_k$ de tal manera que tengan las siguientes tres propiedades:

- a) $\tilde{n}_1, \tilde{n}_2, \dots, \tilde{n}_k \in Q$
- b) El mínimo común múltiplo de los denominadores de $\tilde{n}_1, \tilde{n}_2, \dots, \tilde{n}_k$ es pequeño
- c) $\max_i |n_j - \tilde{n}_j|$ es pequeño

Aquí la condición (a) afirma que una conversión de (C.28) a un sistema de un solo término (como descrito en el capítulo 8) es posible. Específicamente, dado que sólo entran los nuevos valores $\tilde{n}_j$

En las últimas etapas del esquema, dicha conversión siempre es posible, sin ninguna restricción sobre los valores originales n_j .

El propósito de la condición (b) es mantener la dimensión N de este sistema pequeña (recordemos que en el Teorema 8.1 había que esencialmente $N = M_{n_k}$ donde M es el mínimo común múltiplo mencionado en la condición (b)), que – según nuestra idea inicial – era el punto principal del concepto.

La condición (c) finalmente asegura que el error introducido por esta perturbación permanezca pequeño, cf. (C.29). Por supuesto, hay que señalar que existe un conflicto entre las condiciones (b) y (c): En muchos casos será posible mejorar la aproximación requerida en (c) al precio de aumentar el mínimo común múltiplo mencionado en (b).

Una adecuada en este caso es necesario llegar a un compromiso. Sin embargo, parece imposible afirmar una estrategia generalmente válida para la solución de este conflicto; un buen compromiso probablemente dependa de los parámetros específicos de la ecuación bajo consideración.

Esto completa la primera etapa del algoritmo. Al final de esta etapa tenemos un nuevo problema de valor inicial que consiste

en la ecuación diferencial perturbada. (C.28) junto con las condiciones iniciales originales (no perturbadas) (C.24b).

La segunda etapa del algoritmo es entonces la etapa donde el problema de valor inicial que se construyó en la etapa 1 se resolverá numéricamente.

En la práctica lo haremos:

Primero utilice el enfoque descrito al principio de esta sección: convertimos el nuevo problema de valor inicial en un sistema de un solo término, y luego resolveremos este sistema numéricamente (por ejemplo, mediante el método Adams).

Ejemplo C.4.

Construya una solución aproximada para el problema a partir del ejemplo. C.3 mediante la estrategia de dos etapas descrita anteriormente.

Como primer intento de resolver el problema con nuestro método refinado, aproximamos (C.27) por

$$D_{*0}^{1.5} \tilde{y}(x) = -x^{0.1} \frac{E_{1.455}(-x)}{E_{1.445}(-x)} e^x \tilde{y}(x) D_{*0}^{0.5} \tilde{y}(x) + e^{-2x} - [D_{*0}^1 \tilde{y}(x)]^2, \dots (C.30)$$

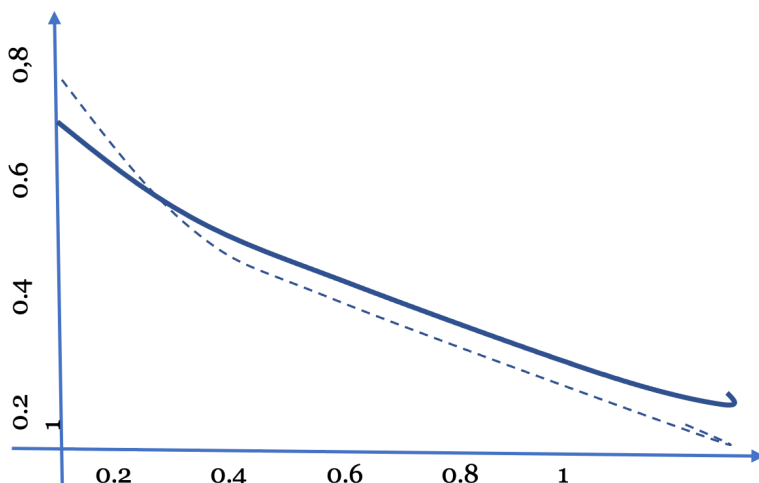
convierta (C.30) a un sistema tridimensional de orden $\gamma = 0.5$ y resuelva este sistema numéricamente con el método Adams en su forma estándar usando varios tamaños de paso. Los resultados se describen en la Tabla C.5.

Tabla C.5. Resultados numéricos de primera aproximación ($N = 3$, $\gamma = 0,5$)

Tamaño del paso	Error máximo	Tiempo de ejecución
1/10	0.136	0.07 s
1/20	0.124	0.18 s
1/40	0.118	0.56 s

Nota: Basada en $N = 3$, $\gamma = 0,5$

Figura C.2 Solución exacta y primera aproximación



Nota: Basada en: $N = 3$, $\gamma = 0,5$, tamaño de paso $h = 0,1$

Podemos ver que casi no hay mejora cuando cambiamos el tamaño del paso. de $1/20$ a $1/40$. Esto indica que el error del esquema Adams (es decir, el error introducido en la segunda etapa) ya es muy pequeño en comparación con el error de la primera etapa (es decir, el error introducido al perturbar la ecuación diferencial). Por lo tanto, no hay necesidad de buscar un esquema mejorado para la solución de este simple sistema. Nótese en particular (ver Figura C.2) que incluso la más cruda de estas tres aproximaciones (la línea discontinua) da una imagen cualitativamente correcta de la solución exacta. (la línea continua). Ciertamente no se puede decir que esto sea cierto para nuestra simple y llana primera enfoque discutido anteriormente.

Para obtener una mejor aproximación con nuestro método ahora debemos reducir el error de la etapa 1, es decir, necesitamos introducir perturbaciones más pequeñas en el orden de los operadores diferenciales. Por tanto, intentamos aproximar la ecuación dada (C.27) no por (C.30) sino por

$$D_{*0}^{14.5} \tilde{y}(x) = -x^{0.1} \frac{E_{1.545}(-x)}{E_{1.445}(-x)} e^x \tilde{y}(x) D_{*0}^{0.5} \tilde{y}(x) + e^{-2x} - [D_{*0}^1 \tilde{y}(x)]^2, \dots (C.31)$$

y proceda como arriba. En consecuencia, encontramos que tenemos que resolver un problema de 29 dimensiones. sistema de orden 0,05 numéricamente. Esta tarea es (en particular debido a la naturaleza de las condiciones iniciales) mucho más difícil que el anterior, y por lo tanto Es necesario poner más esfuerzo en el esquema numérico. Por el momento interpretamos esto requisito como demanda de un tamaño de paso más pequeño; se considerará una alternativa más tarde. Los resultados se dan en la Tabla C.6.

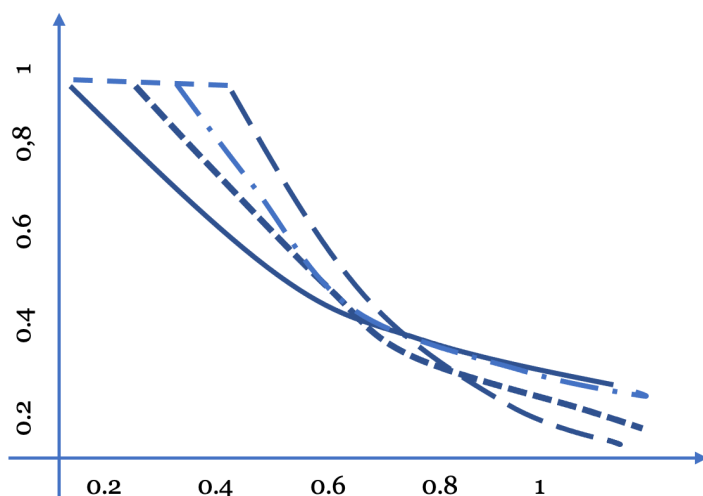
Para fines de comparación con el ejemplo anterior, hemos incluido el caso de un tamaño de paso de 1/40. Como puede verse al comparar las Tablas C.5 y C.6, el error es mucho mayor ahora que antes. La razón es el problema que nosotros Como se mencionó anteriormente: dado que se ha aumentado la dimensión del sistema, la solución numérica necesita más tiempo para alejarse del valor inicial. Un aún más Una imagen obvia de la situación aparece cuando miramos los datos gráficos proporcionados. en la figura C.3. Aquí nuevamente la línea continua es la solución exacta, las otras líneas corresponden a las soluciones numéricas (línea discontinua: $h = 1/40$; línea discontinua y punteada: $h = 1/100$; línea de puntos: $h = 1/200$). Por lo tanto, tenemos que decir que la gráfica para $h = 1/40$ no dar una imagen cualitativamente correcta de la verdadera solución.

Tabla C.6 Resultados numéricos de segunda aproximación

Tamaño del paso	Error máximo	Tiempo de ejecución
1/40	0.2015	0.9 s
1/100	0.0861	5.5 s
1/200	0.0440	21.5 s
1/400	0.0222	82.4 s

Nota: Basada en $N = 29$, $\gamma = 0,05$

Figura C.3 Solución exacta y segunda aproximación



Nota: Basada en $N = 29$, $\gamma = 0,05$, varios tamaños de paso

Como se señaló anteriormente, a veces nos veremos obligados a elegir los parámetros en de tal manera que la dimensión del sistema sea mayor de lo deseable. En este contexto el actual sistema de 29 dimensiones puede considerarse como tal caso. Eso significa que también el número de ceros en la condición inicial del sistema resultante es mayor de lo que a uno le gustaría que fuera, lo que nos obliga a utilizar un tamaño de paso muy pequeño.

Aquí es, donde nuestra segunda estrategia mencionada anteriormente entra en juego como posible alternativa. Específicamente, puede ser útil reemplazar la estructura PECE simple por una Método $P(EC)^{\mu}E$ (es decir, introduciendo iteraciones correctoras adicionales) como se describe en el teorema C.5. Esto permite una propagación más rápida de los elementos distintos de cero, y puede ser posible evitar el uso de tamaños de paso excesivamente pequeños. Deberíamos proporcionar un ejemplo numérico ahora. Esta flexibilidad en el

número de pasos del corrector es en realidad una de las razones principales por las que sugerimos el esquema Adams y no, por ejemplo, el método de (Diethelm, 2001). Recuerde que, como se deriva de la Observación C.4, por (por ejemplo) Al duplicar el número de iteraciones del corrector, esencialmente dejamos el cálculo complejidad sin cambios.

Si usáramos la otra opción y redujéramos el tamaño del paso en un factor de dos, entonces el tiempo de ejecución aumentaría en un factor de cuatro porque la complejidad del algoritmo es $O(h^{-2})$. Ambos enfoques reducirían el tamaño del intervalo inicial, donde la solución numérica se queda estancada en el valor inicial por un factor de $1/2$.

Los datos obtenidos mediante nuestro enfoque $P(EC)^{\mu}E$ se dan en la Tabla C.7. Tenga en cuenta que los datos de la Tabla C.6 corresponden a este método con $\mu = 1$.

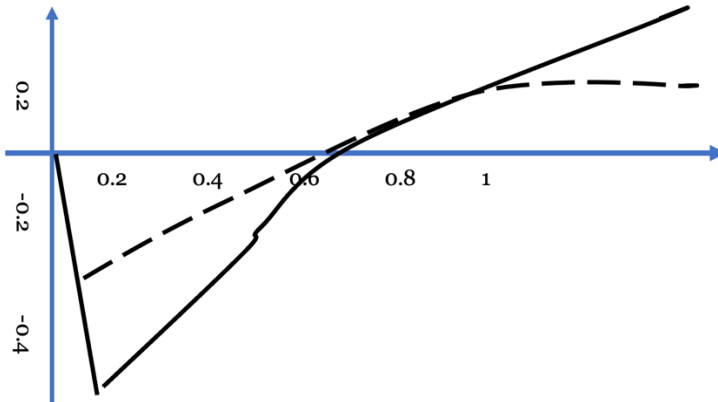
Se ve claramente que este enfoque tiene una ventaja significativa: eligiendo $\mu = 10$ y $h = 1/40$, por ejemplo, obtenemos un error absoluto de alrededor del 25%. menor que en el caso $\mu = 1$ y $h = 1/200$, y al mismo tiempo el tiempo de ejecución.

Tabla C.7. Resultados numéricos de segunda aproximación ($N = 29$, $\gamma = 0,05$) como en (C.31) con el algoritmo $P(EC)^{\mu}E$

Número μ de iteraciones del corrector	Tamaño del paso	Error máximo	Tiempo de ejecución
10	1/40	0.03175	4.9 s
10	1/100	0.01174	28.6 s
20	1/40	0.00989	9.3 s
20	1/100	0.00379	55.0 s

Nota: Basada en $N = 29$, $\gamma = 0,05$ con el algoritmo $P(EC)^{\mu}E$

Figura C.4. Errores para la segunda aproximación ($N = 29$, $\gamma = 0,05$)



Nota: Errores para la segunda aproximación ($N = 29$, $\gamma = 0,05$) como en (C.31) con arias combinaciones de tamaño de paso y número de correctores 75% más corto.

La razón es la siguiente: En el caso $\mu = 1$ la solución numérica se queda atascado en el valor inicial durante un intervalo bastante largo. Al final de este intervalo, la verdadera solución se ha alejado significativamente del valor inicial, y aquí el error alcanza su máximo.

Durante el resto del intervalo $[0,1]$, la solución numérica, tiene que avanzar lentamente hacia la solución exacta y el error se hace más pequeño. Si elegimos un valor mayor para μ , hacemos que el intervalo inicial problemático sea más pequeño y, por lo tanto, también disminuimos el error alcanzado durante este intervalo. Esto se desprende de la figura C.4. donde hemos comparado los errores absolutos para $\mu = 1$, $h = 1/200$ (línea continua) y $\mu = 10$, $h = 1/40$ (línea discontinua).

En este ejemplo, por supuesto, ahora se puede aplicar la idea de utilizar muchos correctores pasos también a la propia ecuación dada (C.27). Esto equivale a saltarse una etapa 1 de nuestro proceso de dos etapas. Está claro que los tiempos de ejecución

del esquema PECE simple (ver Tabla C.4) no son competitivos. Por lo tanto, volvemos una vez más al $P(EC)^\mu E$ estructura con valores más grandes para μ y tamaños de paso más grandes como antes. Algunos resultados son indicados en la Tabla C.8.

Comparando las Tablas C.4 y C.8 encontramos una vez más una ventaja significativa en el tiempo de ejecución en el método $P(EC)^\mu E$ en comparación con el método PECE sin perder precisión, pero incluso las aproximaciones obtenidas por el enfoque más rápido $P(EC)^\mu E$ son menos preciso y requiere más tiempo que los resultados presentados en la Tabla C.7 donde habíamos usado un sistema de ecuaciones diferenciales más simple.

Tabla C.8 Resultados numéricos para la ecuación no perturbada ($N = 291$, $\gamma = 0,005$), con algoritmo $P(EC)\mu E$

Número μ de iteraciones del corrector	Tamaño del paso	Error máximo	Tiempo de ejecución
10	1/100	0.16473	117.3 s
10	1/200	0.08607	446.8 s
20	1/100	0.08607	222.8 s
20	1/200	0.04400	811.0 s

Nota: Basado en $N = 291$, $\gamma = 0,005$), con algoritmo $P(EC)\mu E$

Basado en nuestras consideraciones teóricas aquí y en la Sección C.1 y en heurística argumentos provenientes de los resultados numéricos, ahora damos una descripción completa de un posible algoritmo para la solución aproximada del problema de valor inicial (C.24). El algoritmo seguirá las ideas básicas descritas anteriormente. El fundamental concepto es que asumimos un límite en la complejidad que se dará (expresada en términos del mínimo común múltiplo de los denominadores de los órdenes $\tilde{n}_j$) y que Intentamos lograr una alta precisión en la solución sin exceder la complejidad. límite.

Específicamente, asumimos que el usuario especifica un parámetro $M \in \mathbb{N}$ que interpretar como un límite superior para el mínimo común múltiplo de los denominadores de $\tilde{n}_1, \dots, \tilde{n}_k$. Dado que la dimensión N del sistema que construiremos en la

etapa 2 de el algoritmo viene dado por $N = M_{\tilde{n}_k} \approx M_{n_k}$. Nestos datos nos dan un límite superior en la dimensión, por tanto, un límite superior de la complejidad aritmética.

Comenzamos construyendo las perturbaciones requeridas para la primera etapa. Esto es muy simple; para $j = 1, 2, \dots, k$ solo tenemos que poner $\tilde{n}_j := \frac{\alpha_j}{M}$ donde $\alpha_j \in \mathbb{N}$ es elegido para ser el número natural más cercano a M_{n_j} (es decir, $\alpha_j = M n_j + 0.5$). De este modo nos aseguramos de que, para cada j , la cantidad $|\tilde{n}_j - n_j|$ se minimiza bajo la condición de que el mínimo común múltiplo de los denominadores de $\tilde{n}_1, \dots, \tilde{n}_k$ sea delimitado por M . Esto esencialmente completa la primera etapa.

La segunda etapa comienza reescribiendo las ecuaciones perturbadas como un sistema de orden $\gamma = 1/M$ y dimensión N como se describe en el teorema 8.1. este sistema esta solucionado por el esquema $P(EC)^\mu E$ indicado anteriormente. Para evitar los problemas causados por los grandes números de ceros en la nueva condición inicial, elegimos el parámetro μ de una manera eso depende del número de ceros (es decir, de M); específicamente establecemos $\mu := M$. Tenga en cuenta que de nuestras consideraciones se desprende que no es necesario ni útil introducir flexibilidad adicional eligiendo diferentes valores para el parámetro μ en cada paso. La elección que proponemos aquí es suficiente para evitar los problemas causados por el (posiblemente) gran número de ceros en la inicial condición. Elegir μ mayor que esto no daría un mejor orden de precisión, por lo que no tiene sentido hacer eso (cf. las consideraciones sobre (C.32) más adelante). Elegir μ más pequeño (permanente o temporalmente) significaría que el problema no puede resolverse. evitarse por completo, por lo que habría que suponer un deterioro de la calidad de la aproximación, pero por otro lado no conduciría a una velocidad significativamente más rápida. algoritmo porque la complejidad aritmética de todo el esquema es (asintóticamente) independiente de μ .

Otra ventaja del esquema $P(EC)^\mu E$ con la elección de μ indicada anteriormente puede explicarse echando un vistazo al Teorema C.5: El algoritmo converge a la verdadera solución de la ecuación perturbada con un error de

$$\max_{j=0,1,2,\dots,N} |y(t_j) - y_h(t_j)| = O(h^p); \text{ donde } p = \min(2, 1 + \gamma\mu)$$

Dado que aquí tenemos $\mu = M = 1/\gamma$ por construcción, esto significa que en la propuesta esquema tenemos en realidad $p = 2$ en todos los casos, por lo que encontramos una convergencia ligeramente mejor comportamiento que en el enfoque PECE simple; de hecho, este es el pedido máximo. de lo que se puede obtener mediante un algoritmo que utilice el método de aproximación subyacente a nuestro esquema.

Este enfoque es particularmente útil cuando se busca una solución computacional. aproximación económica, pero aun razonablemente precisa. En muchas aplicaciones esto será lo que se desea porque a menudo uno necesita resolver un gran número de tales problemas con valores iniciales cuyas soluciones luego se requieren como datos de entrada para otros problemas. Además, con frecuencia es imposible obtener una alta precisión porque los datos dados (en particular las órdenes n_j de los operadores diferenciales) son algo como constantes materiales conocidas sólo hasta una cierta precisión (generalmente moderada).

Conversión a sistemas multiorden

En (N. J. Ford & Connolly, 2009) se ha sugerido un enfoque alternativo. Son posibles dos variantes; nosotros comenzamos con el que sea más sencillo de describir y analice el otro más adelante.

La idea básica es utilizar una transformación completamente diferente de la inicial dada. El problema de valor en la forma de un sistema de ecuaciones diferenciales fraccionarias. Esencialmente esto equivale a reemplazar el camino que usa el Teorema 8.1 o 8.2 que teníamos utilizado anteriormente por la transformación a un sistema de múltiples órdenes según el Teorema 8.9 o Teorema 8.10. Comenzamos con el primero y, como arriba, asumimos que el original El problema de valor inicial se da en la forma.

$$D_{*0}^{n_k} y(x) = f(x, y(x), D_{*0}^{n_1} y(x), D_{*0}^{n_2} y(x), \dots, D_{*0}^{n_{k-1}} y(x)), \dots \text{ (C.33a)}$$

donde $0 < n_1 < n_2 < \dots < n_k$ con una función adecuada f y condiciones iniciales

$$y^{(j)}(0) = y_0^{(j)}, j = 0, 1, 2, \dots, [n_k - 1], \dots (C.33b)$$

Sin embargo, ahora asumimos (sin pérdida de generalidad) que además tenemos

$$\{1, 2, 3, \dots, [n_k] - 1\} \subset \{n_{1, \dots}, n_{2, \dots}, \dots, n_{k, \dots}\}, \dots (C.34)$$

Esto implica $n_j, n_{j-1} \leq 1$ para todo j . Consideremos ahora las diferencias $d_j = n_j - n_{j-1}$ para $j = 1, 2, 3, \dots, k$, donde hemos definido $n_0 := 0$ Luego introducimos las nuevas funciones

$$y_0 := y, y_j := D_{*0}^{d_j} y_{j-1}, (j = 1, 2, \dots, k)$$

tal que,

$$y_1 = D_{*0}^{d_1} y_0 = D_{*0}^{n_1 - n_0} y_0 = D_{*0}^{n_1} y$$

$$y_2 = D_{*0}^{d_2} y_1 = D_{*0}^{n_2 - n_1} D_{*0}^{n_1} y = D_{*0}^{n_2} y,$$

$$\vdots,$$

$$y_k = D_{*0}^{d_k} y_0 = D_{*0}^{n_k - n_{k-1}} D_{*0}^{n_{k-1}} y = D_{*0}^{n_k} y.$$

Para la derivación de estas identidades se puede proceder como en la prueba del Lema 4.13; El punto clave es que, debido a nuestro supuesto (C.34), nunca saltamos a través de un número entero al pasar de n_{j-1} , a n_j . Así podemos reescribir (C.33a) en la forma

$$D_{*0}^{d_1} y_0 = y_1$$

$$D_{*0}^{d_2} y_1 = y_2$$

$$\vdots$$

$$D_{*0}^{d_{k-1}} y_{k-2} = y_{k-1}$$

$$D_{*0}^{d_k} y_{k-1} = f(x, y_0(x), y_1(x), y_2(x), \dots, y_{k-1}(x)), \dots (C.35a)$$

Hemos encontrado el sistema requerido de ecuaciones diferenciales; los valores iniciales correspondientes obviamente tienen

$$y_j(0) = \begin{cases} y_0^{(n_j)}, & \text{si } n_j \in \mathbb{N} \\ 0, & \text{en otros casos} \end{cases}$$

en vista de que $y_j = D_{*0}^{n_j} y$, Lema 3.11 y los valores iniciales dados (C.33b). Una comparación de este problema de valor inicial multidimensional con su contraparte (C.25) construido anteriormente revela una serie de diferencias sustanciales, aunque formalmente son equivalentes en el sentido de que los primeros componentes de las soluciones de los dos problemas coinciden:

- a) La dimensión del nuevo sistema es k (un número pequeño en aplicaciones típicas), independiente de los valores de n_j ; habíamos visto anteriormente (ver, por ejemplo, el Ejemplo C.3) que el otro enfoque podría dar lugar a sistemas de dimensiones muy grandes incluso si k fuera pequeño.
- b) El número de ceros en la condición inicial (C.33b) se relaciona con el número de ceros en su homólogo (C.25b) del mismo modo que las dimensiones.
- c) La estructura del lado izquierdo del nuevo sistema (C.35a) es mucho más complicado que en el antiguo sistema (C.25a).
- d) Las estructuras formales de los lados derechos de los dos sistemas no difieren unos de otros en absoluto.

Resulta que, en vista de las consideraciones con respecto al enfoque que utiliza sistemas de orden único, los dos primeros puntos mencionados anteriormente indican la capacidad del nuevo enfoque para evitar los problemas potencialmente graves encontrados en la primera acercarse. Por otro lado, el tercer ítem revela que tenemos que pagar un precio determinado para esta mejora: en lugar de requerir una aproximación para un solo diferencial operador ahora necesitamos trabajar con operadores de orden $d_1, d_2, \dots, d_k$. Estas órdenes pueden coincidir o no entre sí.

Ejemplo C.5.

Reescribimos las ecuaciones de los ejemplos C.2 y C.3 como sistemas de ecuaciones de acuerdo con las ideas descritas anteriormente.

Para la ecuación Bagley-Torvik

$$Ay''(x) = -BD_{*0}^{\frac{3}{2}}y(x) - Cy(x) + C(x+1), y(0) = y'(0) = 1$$

del ejemplo C.2, tenemos

$$n_1 = 1; n_2 = \frac{3}{2}; n_3 = 2$$

$$d_1 = 1; d_2 = \frac{1}{2}; d_3 = 1/2$$

El sistema resultante es tridimensional y se lee

$$D_{*0}^1 y_0(x) = y_1(x)$$

$$D_{*0}^{1/2} y_1(x) = y_2(x)$$

$$D_{*0}^{1/2} y_2(x) = -\frac{B}{A}y_2(x) - \frac{C}{A}y_0(x) + \frac{C}{A}(x+1)$$

con condiciones iniciales

$$y_0(0) = y_1(0) = 1; y_2(0) = 0$$

Una comparación con el ejemplo C.2 muestra que las diferencias para este ejemplo simple son pequeños: la dimensión se reduce en uno y dos derivadas fraccionarias diferentes aparecen en el lado izquierdo del sistema.

En el otro ejemplo, la ecuación fue,

$$D_{*0}^{1.455} y(x) = -x^{0.1} \frac{E_{1.545}(-x)}{E_{1.445}(-x)} \exp(x) y(x) D_{*0}^{0.555} y(x) + \exp(-2x) - [D_{*0}^1 y(x)]^2$$

con condiciones iniciales $y(0) = 1$ y $y'(0) = -1$. Aquí el nuevo enfoque utiliza los parámetros

$$n_1 = 0.555, n_2 = 1 \text{ y } n_3 = 1.455$$

tal que

$$d_1 = 0.555, d_2 = 0.445 \text{ y } d_3 = 0.455$$

Nuevamente obtenemos un sistema tridimensional; esta vez tiene el par

$$D_{*0}^{0.555} y_0(x) = y_1(x)$$

$$D_{*0}^{0.445} y_1(x) = y_2(x)$$

$$D_{*0}^{0.455} y_2(x) = -x^{0.1} \frac{E_{1.545}(-x)}{E_{1.445}(-x)} \exp(x) y_0(x) y_1(x) + \exp(-2x) - [y_2(x)]^2$$

Con condiciones iniciales,

$$y_0(0) = 1; y_1(0) = 0; y_2(0) = -1$$

Ahora bien, la diferencia con el enfoque del sistema de orden único es enorme: la dimensión del sistema se reduce de 291 a 3, por supuesto ahora tenemos que trabajar con tres diferentes operadores diferenciales.

Antes de abordar la cuestión de un esquema numérico adecuado para sistemas de esta estructura, introduzcamos brevemente una pequeña modificación de la idea presentada hasta ahora que puede conducir a un esquema ligeramente más eficiente. Esto es motivado por la observación de los dos ejemplos anteriores de que el diferencial de orden entero de los operadores que son locales por naturaleza se descompone en dos (o más) operadores diferenciales fraccionarios no locales: en el ejemplo de Bagley-Torvik tenemos

$$y'' = D_{*0}^{1/2} y_2 = D_{*0}^{1/2} D_{*0}^{1/2} y_1 = D_{*0}^{1/2} D_{*0}^{1/2} D_{*0}^{1/2} y_0,$$

por lo que en cualquier método de aproximación se pierde la posibilidad de ahorrar tiempo haciendo uso de la localidad. Una descomposición similar

$$y_2 = y = D_{*0}^{0.445} y_1 = D_{*0}^{0.445} D_{*0}^{0.555} y_0,$$

se utiliza en el otro ejemplo. Así sería o será preferible utilizar sistemas alternativos.

$$D_{*0}^1 y_0(x) = y_1(x)$$

$$D_{*0}^{1/2} y_1(x) = y_2(x)$$

$$D_{*0}^1 y_1(x) = -\frac{B}{A} y_2(x) - \frac{C}{A} y_0(x) + \frac{C}{A} (x + 1)$$

Con condiciones iniciales

$$y_0(0) = y_1(0) = 1 ; y_2(0) = 0$$

para el problema Bagley-Torvik y

$$D_{*0}^{0.555} y_0(x) = y_1(x)$$

$$D_{*0}^1 y_1(x) = y_2(x)$$

$$D_{*0}^{0.455} y_2(x) = -x^{0.1} \frac{E_{1.545}(-x)}{E_{1.445}(-x)} \exp(x) y_0(x) y_1(x) + \exp(-2x) - [y_2(x)]^2$$

Con condiciones iniciales

$$y_0(0) = 1; y_1(0) = 0; y_2(0) = -1$$

para el otro ejemplo. De esta manera simplificamos un poco la estructura porque algunos de los operadores diferenciales fraccionarios del lado izquierdo se pueden reemplazar por operadores de orden entero. Esta modificación equivale a utilizar el Teorema 8.10. en lugar del Teorema 8.9 en nuestro enfoque de sistema de orden múltiple.

Para la solución numérica de estos sistemas de ecuaciones se puede utilizar un esquema numérico para ecuaciones diferenciales fraccionarias escalares para cada componente por separado. Por supuesto, hay que tener en cuenta que las ecuaciones individuales ahora normalmente no serán del mismo orden, por lo que los métodos numéricos para diferentes, se deben utilizar órdenes. En general, no parece haber ninguna ventaja en el uso métodos construidos de forma diferente para las ecuaciones individuales; más bien, normalmente se preferiría utilizar sólo una clase de esquemas numéricos y simplemente cambiar los órdenes. de los algoritmos según lo prescrito por las órdenes de los operadores diferenciales en el lado izquierdo del sistema.

En cualquiera de los dos enfoques de sistemas de múltiples, Edwards et al. órdenes (Edwards et al., 2002) investigó el uso de la fórmula de Diethelm, K (Diethelm, 1997) para la solución numérica del resultado sistema de ecuaciones; resultados posteriores Ford, N.J., Connolly J.A (N. J. Ford & Connolly, 2006) indican que nuestro Adams-Bashforth-Moulton es probable que el método sea más eficiente, en lo que respecta al error, resulta que el comportamiento de todo el esquema está dominado por el componente con peor comportamiento.

Una comparación con el enfoque del sistema de orden único muestra que el sistema de orden múltiple es el enfoque del sistema, que siempre es aplicable (no existen restricciones de la teoría de números sobre los órdenes $n_1, \dots, n_k$), y en muchos casos conducirá a un sistema con una dimensión considerablemente menor. Sin embargo, la estructura del lado izquierdo se vuelve más complicado.

Los experimentos numéricos de Ford y Connolly (N. J. Ford & Connolly, 2009) indican que la conversión a un sistema de orden múltiple mediante el teorema 8.10 tiende a ser computacionalmente el menos eficiente de los enfoques presentados aquí. Encontraron la conversión a un sistema de orden múltiple a través del Teorema 8.9 y el enfoque que utiliza la transformación a sistemas de orden único mediante el teorema 8.1 o 8.2 suele ser preferible. ¿Cuál de estas ideas funciona mejor? Parece depender del problema particular bajo análisis consideración, por lo que no se puede dar un consejo generalmente válido.

Apéndice D

D. Resultados útiles en el análisis

En este capítulo recopilamos información sobre algunos conceptos de análisis que son útiles en el resto del texto.

D.1 Función gamma de Euler

Comenzamos con la función Gamma. Recordamos la definición:

$$\Gamma(x) = \int_0^{\infty} t^{x-1} e^{-t} dt$$

para $x > 0$. Consideraciones elementales de la teoría de integrales impropias revelan que la integral existe. Además, al establecer $x = 1$ vemos fácilmente

$$\begin{aligned} \Gamma(1) &= \int_0^{\infty} e^{-t} dt = \lim_{z \rightarrow \infty} \int_0^z e^{-t} dt = \lim_{z \rightarrow \infty} [-e^{-t}]_0^z \\ &= \lim_{z \rightarrow \infty} (1 - e^{-z}) = 1, \dots (D.1) \end{aligned}$$

Además, podemos, para $x > 0$ arbitrario, manipular la integral en la definición de la función Gamma mediante una integración parcial. Esto produce

$$\begin{aligned} \Gamma(x+1) &= \int_0^{\infty} t^x e^{-t} dt = \lim_{z \rightarrow \infty, y \rightarrow 0+} \int_0^z t^x e^{-t} dt \\ &= \lim_{z \rightarrow \infty, y \rightarrow 0+} \left([-e^{-t} t^x]_{t=y}^{t=z} + x \int_y^z t^{x-1} e^{-t} dt \right) \\ &= x \int_y^z t^{x-1} e^{-t} dt = x \Gamma(x) \end{aligned}$$

Así hemos demostrado.

Teorema D.1. Ecuación funcional para Γ

Si $x > 0$ entonces $x\Gamma(x) = \Gamma(x+1)$. Ahora podemos probar la importante relación entre la función Gamma y la factorial

Prueba (del teorema 2.3)

La prueba utiliza la inducción matemática. La base de la inducción. ($n = 1$) lee $\Gamma(1) = 0! = 1$, lo cual es cierto teniendo en cuenta (D.1). Para el paso de inducción, usamos la ecuación funcional y la hipótesis de inducción.

$$\Gamma(n+1) = n\Gamma(n) = n(n-1)! = n!$$

cómo se desee.

Hay otra aplicación importante de la ecuación funcional de Gamma. Lo resolvemos para $\Gamma(x)$; luego lee

$$\Gamma(x) = \frac{\Gamma(x+1)}{x}, \dots (D.2)$$

si $x > 0$. Ahora bien, la expresión del lado derecho tiene significado no sólo si $x > 0$ pero también en el caso $-1 < x < 0$. Por lo tanto, podemos usarlo como definición para el lado izquierdo, es decir, para $\Gamma(x)$, en ese caso (que no está cubierto por la definición original porque la integral definitoria es divergente para $x < 0$). Habiendo hecho esta extensión, la función gamma también está definida para $-1 < x < 0$, y podemos volver a (D.2) con x en ese rango. Esto nos permite ampliar la definición a $-2 < x < -1$. Procediendo de esta manera, encontramos una definición para la función Gamma que se puede aplicar para todos $x \in \mathbb{R}$ con excepción de aquellos para los cuales $-x \in \mathbb{N}_0$

Como consecuencia de estas consideraciones, encontramos otra identidad importante involucrando la función Gamma:

Teorema D.2.

Sean $n \notin \mathbb{Z}$ y $k \in \mathbb{N}_0$ Entonces

$$(-1)^{k+1} \Gamma(n-k) \Gamma(k+1-n) = \Gamma(-n) \Gamma(n+1)$$

Otra identidad útil en este contexto es.

Teorema D.3. Fórmula de reflexión para Γ

Sea $0 < x < 1$. Entonces,

$$\Gamma(x) \Gamma(1-x) = \frac{\pi}{\operatorname{sen} \pi x}$$

También es posible encontrar una representación alternativa, debida a Gauss, para la Función gamma. Esta representación es válida para la extensión indicada. arriba. Sin embargo, en los cálculos prácticos se observa frecuentemente que la integral la representación es más fácil de manejar

Teorema D.4. Fórmula del producto de Gauss para Γ

Sea $x \in \mathbb{R}$, $-x \in \mathbb{N}_0$.. Entonces

$$\Gamma(x) = \lim_{n \rightarrow \infty} \frac{n! \, n^x}{x(x+1)(x+2) \dots (x+n)}$$

El comportamiento asintótico de $\Gamma(x)$ cuando $x \rightarrow \infty$ es a veces importante; puede ser descrito por el siguiente resultado (Agarwal et al., 2008), Capítulo 7.

Teorema D.5. Fórmula de Stirling

Para $x \rightarrow \infty$ tenemos

$$\Gamma(x+1) = \left(\frac{x}{e}\right)^x \sqrt{2\pi x} (1 + o(1))$$

Un último resultado que mencionaremos explícitamente es la siguiente identidad integral. Dejamos la demostración como ejercicio.

Teorema D.6.

Sea $\alpha, \beta \in \mathbb{R}_+$. Entonces

$$\int_0^1 t^{\alpha-1} (1-t)^{\beta-1} dt = \frac{\Gamma(\alpha) \Gamma(\beta)}{\Gamma(\alpha+\beta)}$$

y por lo tanto

$$\int_0^x t^{\alpha-1} (x-t)^{\beta-1} dt = x^{\alpha+\beta-1} \frac{\Gamma(\alpha) \Gamma(\beta)}{\Gamma(\alpha+\beta)}$$

La integral de la primera ecuación del teorema D.6 se conoce como integral de Euler del primer tipo o función Beta de Euler $B(\alpha, \beta)$

Se puede encontrar más información sobre la función Gamma, por ejemplo, en la obra clásica de Artin (Hewitt et al., 1964) o en las obras de referencia habituales sobre funciones especiales.

D.2 Teoremas del punto fijo

Las pruebas de varios teoremas de existencia y unicidad a lo largo de este texto tienen como base en teoremas clásicos que afirman la existencia o unicidad de puntos fijos de determinados operadores.

El primero de estos teoremas es la siguiente generalización del punto fijo de Banach teorema que tomamos de Weissinger, J. (Jachymski et al., 2024).

Teorema D.7. Teorema del punto fijo de Weissinger)

Supongamos que (U, d) es un espacio métrico completo no vacío, y sea $\alpha_j \geq 0$ para todo $j \in \mathbb{N}_0$ y tal que $\sum_{j=0}^{\infty} \alpha_j$ converge. Además, supongamos que la aplicación $A: U \rightarrow U$ satisfaga la desigualdad:

$$d(A^j u, A^j v) \leq \alpha_j d(u, v), \dots, (D.3)$$

para cada $j \in \mathbb{N}$ y cada $u, v \in U$. Entonces, A tiene un valor fijo unívocamente determinado punto u^* . Además, para cualquier $u_0 \in U$, la secuencia $(A^j u_0)_{j=1}^{\infty}$ converge a este fijo punto u^* .

Una consecuencia inmediata es:

Corolario D.8. Teorema del punto fijo de Banach

Supongamos que (U, d) es un espacio métrico completo no vacío, sea $0 \leq \alpha < 1$, y deje que el mapeo $A: U \rightarrow U$ satisfaga la desigualdad

$$d(Au, Av) \leq \alpha d(u, v), \dots (D.4)$$

para cada $u, v \in U$. Entonces, A tiene un punto fijo determinado unívocamente u^* además, para cualquier $u_0 \in U$ la secuencia $(A^j u_0)_{j=1}^\infty$ converge a este punto fijo u^* .

Además, utilizamos un resultado ligeramente diferente que afirma sólo la existencia, pero no la unicidad de un punto fijo. Aquí podemos trabajar con supuestos más débiles. del operador en cuestión.

Se puede encontrar una prueba, por ejemplo, en (Collatz, 2013) (Corduneanu, 1994).

Teorema D.9. Teorema del punto fijo de Schauder

Sea (E, d) un espacio métrico completo, sea U un subconjunto convexo cerrado de E , y sea $A : U \rightarrow U$ un mapeo tal que el conjunto $\{Au : u \in U\}$ es relativamente compacto en E . Entonces A tiene al menos un punto fijo.

En este contexto recordamos una definición:

Definición D.1.

Sea (E, d) un espacio métrico y $F \subseteq E$. El conjunto F se llama relativamente subconjunto compacto en E si,

Un resultado clásico útil del análisis en relación con tales conjuntos es el siguiente. La prueba se puede encontrar en muchos libros de texto estándar (Véase en (Corduneanu, 1994), pág. 30).

Teorema D.10. Arzela-Ascoli

Sea $F \subseteq C[a, b]$ para algunos $a < b$, y supongamos que Los conjuntos estarán equipados con la norma Chebyshev. Entonces, F es relativamente compacto en $C[a, b]$ si F es equicontinuo (es decir, para cada $\varepsilon > 0$ existe algún $\delta > 0$ tal que para todo $f \in F$ y todo $x, x^* \in [a, b]$ con $|x - x^*| < \delta$ tenemos $|f(x) - f(x^*)| < \varepsilon$) y uniformemente acotado (es decir, existe una constante $C > 0$ tal que $\|f\|_\infty \leq C$ para cada $f \in F$).

D.3 La transformada de Laplace

El método de la transformada de Laplace es una herramienta extremadamente útil para el análisis de Problemas de valor inicial (fraccionales o clásicos). En particular, nos permite reemplazar una ecuación diferencial por una ecuación algebraica. Tomamos la definición fundamental, del libro clásico de Doetsch (Doetsch, 1971), donde el lector interesado puede encontrar una Tratamiento integral de la transformada de Laplace.

Definición D.2.

Sea $f: [0, \infty) \rightarrow \mathbb{R}$. dada la función F definida por

$$F(s) := Lf(s) := \int_0^{\infty} f(x)e^{-sx} dx$$

se llama transformada de Laplace de f siempre que existe la integral. Es bastante sencillo calcular la transformada de Laplace de algunas funciones elementales.

Ejemplo D.1.

- a) Para $f(x) = \exp(ax)$ con $a \in \mathbb{R}$ tenemos $\mathcal{L}f(s) = \frac{1}{(s-a)}$ siempre que $s > a$
- b) Para $f(x) = x^k$ con $k > -1$ encontramos $\mathcal{L}f(s) = \frac{\Gamma(k+1)}{s^{k+1}}$ siempre que $s > 0$.
- c) Para $f(x) = \sin \omega x$ con $\omega > 0$ tenemos $\mathcal{L}f(s) = \frac{\omega}{(s^2 + \omega^2)}$, nuevamente para $s > 0$.

Citamos las reglas más importantes para las transformadas de Laplace.

Teorema D.II.

Supongamos que las funciones f_1, f_2 y f_3 están dadas en $[0, \infty)$ y son tal que sus transformadas de Laplace existen para todo $s \geq$

s_0 con algún $s_0 \in \mathbb{R}$ adecuado. Entonces tenemos las siguientes reglas.

- a) Si $f_3 = a_1 f_1 + a_2 f_2$ con constantes reales arbitrarias a_1 y a_2 entonces

$$\mathcal{L}f_3(s) = a_1 \mathcal{L}f_1(s) + a_2 \mathcal{L}f_2(s)$$

linealidad de la transformada de Laplace

- b) Si f_3 es la convolución de f_1 y f_2 es decir, si

$$f_3(x) = \int_0^x f_1(x-t) f_2(t) dt$$

Entonces

$$\mathcal{L}f_3(s) = \mathcal{L}f_1(s) \cdot \mathcal{L}f_2(s)$$

el teorema de la convolución. En otras palabras: la convolución en el original. El dominio corresponde al producto habitual en el dominio de Laplace.

- c) Si $f_3(x) = \int_0^x f_1(t) dt$ entonces tenemos para $s > \max\{0, s_0\}$

$$\mathcal{L}f_3(s) = \frac{1}{s} \mathcal{L}f_1(s).$$

el teorema de integración.

- d) Sea $m \in \mathbb{N}$ Si $f_3 = D^m f_1$ es la m -ésima derivada de f_1 entonces

$$\mathcal{L}f_3(s) = s^m \mathcal{L}f_1(s) - \sum_{k=1}^m s^{m-k} f_1^{(k-1)}(0)$$

el teorema de diferenciación.

- e) Sea $a > 0$ y $f_3(x) = f_1(ax)$ Entonces

$$\mathcal{L}f_3(s) = \frac{1}{a} \mathcal{L}f_1(s/a)$$

- f) Sea $a \in \mathbb{R}$ y $f_3(x) = e^{-ax} f_1(x)$ $f_3(x) = e^{-ax} f_1(x)$. Entonces

$$\mathcal{L}f_3(s) = \mathcal{L}f_1(s+a)$$

g) Sea $m \in \mathbb{N}$ Si $f_3(x) = x^m f_1(x)$. Entonces

$$\mathcal{L}f_3(s) = (-1)^m \frac{d^m}{ds^m} \mathcal{L}f_1(s)$$

h) Sea $f_3(x) = f_1(x)/x$. Entonces

$$\mathcal{L}f_3(s) = \int_s^\infty \mathcal{L}f_1(\sigma) d\sigma$$

i) Mar $a \in \mathbb{R}$ y

$$f_3(x) = \begin{cases} 0 & \text{para } x < a \\ f_1(x-a) & \text{para } x \geq a \end{cases}$$

Entonces,

$$\mathcal{L}f_3(s) = e^{-as} \mathcal{L}f_1(s)$$

La parte (d), el teorema de diferenciación es de particular interés para nosotros. Específicamente este resultado necesitaba generalizarse a $m \notin \mathbb{N}$ con una definición adecuada del operador diferencial. Hemos tratado esta cuestión en el Teorema 8.1

Por supuesto, no basta con tener la transformada de Laplace; para trabajos prácticos También se requiere la transformación inversa. Hay varias formas de expresar este inverso; una posibilidad está contenida en el siguiente resultado. Nos referimos a los libros estándar, en las transformadas de Laplace para obtener detalles sobre los "supuestos adecuados".

Teorema D.12.

Bajo supuestos adecuados sobre f tenemos

$$f(x) = \frac{1}{2\pi i} \int_{c-i\infty}^{c+i\infty} \exp(sx) \mathcal{L}f(s) ds$$

Bajo ciertas condiciones, el comportamiento a largo plazo de las funciones también puede expresarse con la ayuda de las transformadas de Laplace, ver en Chen J. (Chen et al., 2007), Gluskin (Gluskin, 2003); Enseñemos esta generalización del

teorema del valor final y las referencias. citado en el mismo (Gluskin, 2003)(Chen et al., 2007)(Kochubei, 2008).

Teorema D.13. Teorema del valor final

Supongamos que $\mathcal{L} f$ no tiene ninguna singularidad en el semiplano derecho cerrado $\{s \in \mathbb{C} : \operatorname{Re} s \geq 0\}$, excepto posiblemente un polo simple en el origen. Entonces,

$$\lim_{x \rightarrow \infty} f(x) = \lim_{s \rightarrow 0_+} s \cdot \mathcal{L} f(s)$$

Observación D.1. La condición sobre las singularidades de $\mathcal{L} f$ es esencial aquí: Si $\mathcal{L} f$ tiene un polo con parte real positiva, entonces $f(x)$ es ilimitado cuando $x \rightarrow \infty$, y si $\mathcal{L} f$ tiene un polo en el eje imaginario (pero no en el origen),

D.4 Integral de partes finitas de Hadamard

La integral $\int_a^b (x-a)^{-\mu} f(x) dx$ es divergente para $\mu \geq 1$ siempre que $f(a) \neq 0$. Sin embargo, a veces es útil asignar un valor finito a tales integrales. Esto tiene Hadamard J (A. A. Kilbas & Trujillo, 2002) ha observado en relación con métodos de solución para ciertas ecuaciones diferenciales parciales, e introdujo la siguiente idea para la solución de este problema, conocida como integral de partes finitas. Requeriremos este concepto principalmente para $\mu \notin \mathbb{N}$ y, por lo tanto, restringiremos nuestra atención a estos valores de μ . La consideración de valores enteros requiere algunas pequeñas modificaciones.

La parte finita de Hadamard de la integral (que, en aras de la simplicidad, denotamos por el mismo símbolo que la integral estándar) (Hadamard, 1952) se define, en términos generales, por una expansión de Taylor de f en $x = a$ donde las integrales singulares resultantes son definidos por

$$\int_a^b (x-a)^{-\mu} dx = \frac{1}{1-\mu} (b-a)^{1-\mu}, (\mu > 1), \dots (D.5)$$

Básicamente, esto significa que primero reemplazamos la integral $\int_a^b (x-a)^{-\mu} dx$ por la expresión $\int_{a+\varepsilon}^b (x-a)^{-\mu} dx$ para $\varepsilon > 0$. Esta es una integral convergente; su valor es simplemente

$(1 - \mu)^{-1} [(b - a)^{1-\mu} - \varepsilon^{1-\mu}]$. Entonces consideramos $\varepsilon \rightarrow 0$. Por supuesto que el límite no existe para $\mu \geq 1$, por lo que Hadamard sugirió simplemente ignorar la contribución ilimitada $\lim_{\varepsilon \rightarrow 0} \frac{\varepsilon^{1-\mu}}{(1-\mu)}$ y asignar el valor de la expresión restante (finita) $(1 - \mu)^{-1} (b - a)^{1-\mu}$ de ahí el nombre “integral de partes finitas”.

Una forma precisa de definir la integral de partes finitas es (para $\mu \notin \mathbb{N}$)

$$\begin{aligned} \int_a^b (x - a)^{-\mu} f(x) dx: \\ = \sum_{k=0}^{[\mu]-1} \frac{f^{(k)}(a)(b - a)^{k+1-\mu}}{(k + 1 - \mu)k!} \\ + \int_a^b (x - a)^{-\mu} R_{[\mu]-1}(x, a) dx \dots (D.6a) \end{aligned}$$

Donde

$$R_p(x, a) := \frac{1}{p!} \int_a^x (x - y)^p f^{(p+1)}(y) dy, \dots (D.6b)$$

es el resto del polinomio de Taylor de pésimo grado de f con punto de expansión a . Él es bien sabido que una condición suficiente para la existencia de la integral (D.6a) es que $f \in C^s[a, b]$ con $\mu - 1 < s \in \mathbb{N}$. Esto se debe a que entonces El término restante de la expansión de Taylor tiene un cero en un orden tan alto que la singularidad en el otro factor en la última integral de (D.6a) casi se está cancelando; la singularidad restante es débil e integrable en lo impropio sentido.

Una representación alternativa que nos resulta útil puede tomarse de véase en (A. D. Freed & Diethelm, 2006), eq. (A17):

Teorema D.14.

Sea $\mu > 1$ pero $\mu \notin \mathbb{N}$ y $m := \mu - 1$. Para $f \in C^m[a, b]$ tenemos

$$\begin{aligned} & \frac{1}{\Gamma(1-\mu)} \int_a^b (b-x)^{-\mu} f(x) dx \\ &= \sum_{k=0}^{m-1} \frac{(b-a)^{k-\mu+1}}{\Gamma(k-\mu+2)} f^{(k)}(a) + j_a^{k-\mu+1} f^{(m)}(b) \end{aligned}$$

Mencionamos aquí las propiedades más importantes de la integral de partes finitas:

- A diferencia de la integral clásica de Riemann o Lebesgue, la integral de partes finitas no es un funcional positivo, es decir, la desigualdad

$$\left| \int_a^b (x-a)^{-\mu} f(x) dx \right| \leq \int_a^b (x-a)^{-\mu} |f(x)| dx$$

no es cierto en general.

- La integral de partes finitas es una extensión consistente del concepto de integrales regulares, es decir, siempre que la integral $\int_a^b (x-a)^{-\mu} f(x) dx$ existe en el sentido clásico, entonces también existe en el sentido de partes finitas, y las dos integrales tienen el mismo valor.
- La integral de partes finitas es aditiva con respecto a la unión de intervalos de integración e invariante con respecto a la traslación.
- La integral de partes finitas es lineal.
- La regla habitual de cambio de variables sigue siendo válida si $\mu \notin \mathbb{N}$.

Un resultado muy útil sobre estas integrales es el siguiente.

Teorema D.15.

Sea $f \in C^k[a, b]$ para algún $k \in \mathbb{N}_0$, y sea $p < k$. Entonces, para $a < x < b$,

$$\frac{d}{dx} \int_a^x f(t)(x-t)^{-p} dt = -p \int_a^x f(t)(x-t)^{-p-1} dt$$

Dejamos la prueba como ejercicio a lector.

D.5 Teoría de la aproximación

Un concepto bien conocido de la teoría de la aproximación que tuvimos que usar en la prueba del teorema 3.25 fue el polinomio de Bernstein. Una referencia clásica es el libro de Lorentz. La definición es:

$$B_N[f](t) := \sum_{k=0}^N \binom{N}{k} t^k (1-t)^{N-k} f\left(\frac{k}{N}\right)$$

donde $f: [0,1] \rightarrow \mathbb{R}$. El resultado fundamental que requerimos es un teorema convergencia.

Teorema D.16.

Sea $f \in C^\ell[0,1]$ para algún $\ell \in \mathbb{N}_0$. Entonces, para todo $\mu \in \{0,1,2,\dots\}$, la secuencia $(D^\mu B_N)_{N=1}^\infty$ converge uniformemente hacia $D^\mu f$. Se puede encontrar una prueba en (Chen et al., 2007).

Atik Editorial



 **ATIK**
editorial

ISBN: 978-9907-820-07-2



9 789907 820072