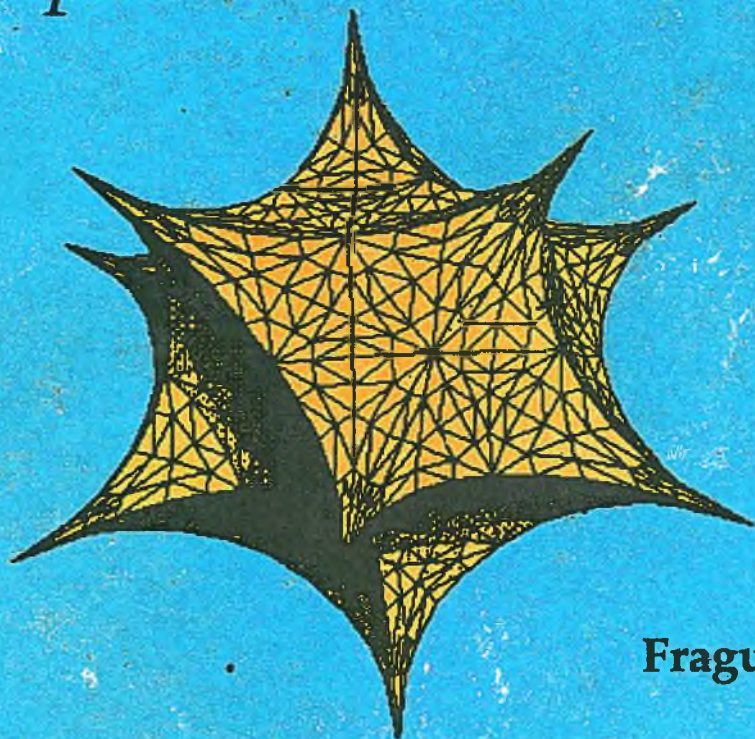


*VII Coloquio del Departamento
de Matemáticas*

29 de Julio al 16 de Agosto de 1991
La Trinidad Tlaxcala



*Análisis Funcional
Aplicado*



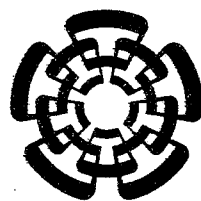
**Andrés
Fraguela Collar**



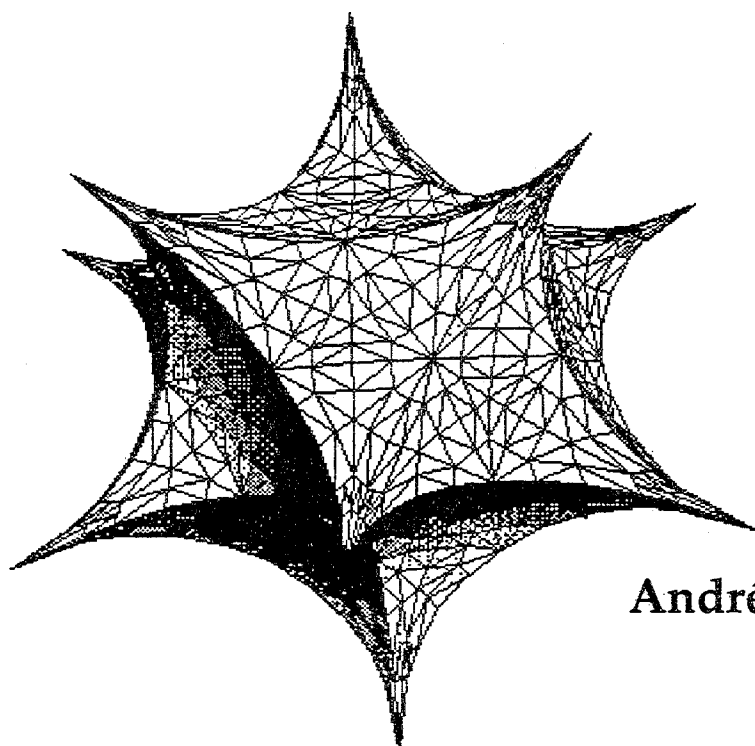
Centro de Investigación y de Estudios Avanzados del I.P.N.

*VII Coloquio del Departamento
de Matemáticas*

29 de Julio al 16 de Agosto de 1991
E. S. F. M. del I. P. N.



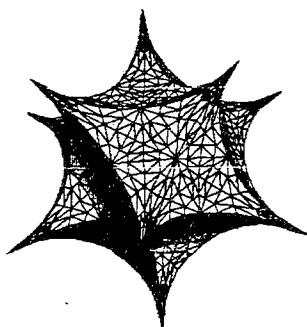
Teoría Espectral de Operadores Diferenciales



Andrés Fraguela
Collar



Centro de Investigación y de Estudios Avanzados del I.P.N.



Estas notas fueron expresamente preparadas para el VII Coloquio del Departamento de Matemáticas del CINVESTAV. La edición consta de 500 ejemplares.

Esperamos que estas notas cumplan los objetivos del Coloquio, los cuales son:

- Promover el desarrollo de la Matemática dentro de las Instituciones con necesidades e intereses en este campo.
- Promover la producción de textos de Matemáticas de nivel avanzado escritos en México.
- Promover la actualización y superación académica de los profesores de nivel superior, particularmente de las escuelas de ingeniería y ciencias.
- Promover la discusión interdisciplinaria entre ingeniería, ciencia, matemáticas y áreas afines.

Agradecemos el apoyo financiero del Consejo Nacional de Ciencia y Tecnología (CONACyT) y de la Secretaría de Educación Pública (SEP), sin el cual este Coloquio no hubiera podido realizarse en las condiciones en que se llevó a cabo.

El Comité Organizador
Julio de 1991

Índice

Introducción	v
1 Espacios métricos normados y euclidianos	1
§ 1.1 Espacios métricos. Conceptos topológicos en los espacios métricos.	2
§ 1.2 Completitud, compacidad y conexidad en los espacios métricos.	27
§ 1.3 Normas y espacios normados	49
§ 1.4 Producto escalar y espacios euclidianos. Ortogonalidad.	67
§ 1.5 Series de Fourier y bases en espacios euclidianos.	80
§ 1.6 Espacios de Hilbert	88
§ 1.7 Derivadas en sentido de Soboliev y espacios de Soboliev.	106
2 Operadores y funcionales lineales continuos	143
§ 2.1 Algunos resultados generales de la teoría de operadores y funcionales lineales continuos en espacios normados	144

§ 2.2	Extensión de operadores y funcionales lineales continuos. Principio de prolongación de Hahn–Banach	157
§ 2.3	Algunos teoremas de representación de funcionales lineales y bilineales continuos.	170
§ 2.4	Diferentes tipos de convergencia para operadores y funcionales lineales continuos	183
§ 2.5	Teoremas de la aplicación abierta y del grafo cerrado . .	204
3	Conceptos y resultados generales de la teoría de operadores en los espacios de Hilbert	211
§ 3.1	Algunos conceptos generales de la teoría de operadores lineales	212
§ 3.2	Operadores cerrados y cerrables	215
§ 3.3	Conjugado de un operador lineal	220
§ 3.4	Operadores autoadjuntos, simétricos y esencialmente autoadjuntos	226
§ 3.5	Resolvente y espectro	242
§ 3.6	Espectro de operadores simétricos y autoadjuntos	260
§ 3.7	Teoría de extensión de operadores simétricos y representación de formas sesquilineales	274
§ 3.8	La integral de Dunford	312
§ 3.9	Raíz cuadrada de operadores autoadjuntos positivos y representación polar de operadores cerrados.	319
§ 3.10	Separación del espectro, subespacios invariantes y reducción de operadores cerrados	328

§ 3.11 Valores propios aislados en el espectro	337
§ 3.12 Operadores J -autoadjuntos y espacios con métrica in- definida	345
Bibliografía	355

Introducción

El presente material está constituido por las notas del curso Teoría Espectral de Operadores Diferenciales dirigido a los participantes del VII Coloquio de Matemáticas del Centro de Investigación y de Estudios Avanzados del IPN.

Este volumen corresponde a la parte básica del curso dedicada a introducir al lector en la temática y abarca desde los conceptos primarios de espacios métricos y euclidianos hasta resultados generales de la teoría de operadores en espacios de Hilbert pasando por los principales tópicos relativos a operadores y funcionales lineales continuos.

En una segunda parte, de próxima aparición, se desarrollan la Teoría Espectral y sus aplicaciones a problemas de la Mecánica.

En el primer capítulo se estudian diferentes aspectos de la topología y de la estructura uniforme en los espacios métricos, en particular se hace especial énfasis en el caso de espacios normados y euclidianos. En el último epígrafe de este capítulo se introducen los espacios de Soboliev, una clase de espacios normados que juega un papel importante en el área de Ecuaciones Diferenciales y Física Matemática.

El siguiente capítulo contiene un estudio de los operadores y funcionales lineales y continuos definidos en espacios de Banach que incluye los teoremas principales del Análisis Funcional lineal: Hahn-Banach-Steinhaus, del grafo cerrado y de la aplicación abierta.

Para finalizar esta primera parte el tercer capítulo presenta de una manera sistemática y accesible los principales conceptos y resultados de la teoría de operadores en espacios de Hilbert. Se presta especial atención a los operadores simétricos y autoconjugados así como al estudio de su espectro.

Por último deseo agradecer la invitación formulada por el Departamento de Matemáticas del CINVESTAV-IPN a dictar este curso en el marco del VII Coloquio de Matemáticas, asimismo testimonio mi reconocimiento a la Srita. Norma Acosta y al Dr. Ismael Muñoz por el esfuerzo realizado en la edición de estas notas.

Andrés Fraguela

1 Espacios métricos normados y euclideanos

Este capítulo es de carácter introductorio. En él se estudian diferentes aspectos de la topología y la estructura uniforme en los espacios métricos y muy especialmente en el caso particular de los espacios normados y euclideanos. En el caso de los espacios normados se hace énfasis en las consecuencias de la compatibilidad entre su estructura topológica y la estructura vectorial subyacente. Para los espacios euclideanos se incluye en el análisis además la compatibilidad con la estructura geométrica.

Una parte importante del capítulo está dedicada al estudio de la teoría geométrica de los espacios euclideanos y de Hilbert (espacios euclideanos completos), es decir, al estudio de las propiedades derivadas de la noción de ortogonalidad.

A pesar de que en el resto del libro desarrollaremos la teoría de operadores en espacios de Hilbert y no en espacios normados cualesquiera, partiremos aquí de la introducción de los espacios normados por dos motivos:

- 1.- En primer lugar, porque muchos de los teoremas fundamentales del Análisis Funcional relativos a funcionales lineales continuos y operadores lineales continuos admiten un planteamiento más general en el contexto de los espacios normados, sin que esto implique complicaciones auxiliares sustanciales durante su desarrollo.

- 2.- En segundo lugar, en el estudio de los operadores diferenciales aparecen de manera natural ciertos espacios funcionales normados llamados espacios de Soboliev, los cuales serán estudiados al final de este capítulo, y resulta que al analizar la regularidad de la solución de los problemas de contorno, se hace necesaria la introducción de espacios de este tipo que no son euclidianos.

Muchos de los resultados expuestos en este capítulo serán dados sin demostración con el objetivo de abreviar la exposición. El lector interesado en ampliar sobre estos temas puede consultar el libro “Análisis matemático en espacios métricos” de A. Fraguera [1].

§ 1.1 Espacios métricos. Conceptos topológicos en los espacios métricos.

La introducción de la estructura métrica en conjuntos arbitrarios persigue como objetivo crear una base axiomática que permita la generalización más intuitiva posible, basándose en la noción de proximidad entre puntos, de conceptos fundamentales del Análisis Matemático tales como convergencia y continuidad, a conjuntos no numéricos y muy en particular a los espacios de funciones.

El concepto de convergencia es fundamental en el Análisis Matemático. Lo empleamos al hablar de sucesiones convergentes de números reales y complejos, de sucesiones de puntos en un espacio euclideo de dimensión finita, de sucesiones de funciones, etc. En este último caso incluso utilizamos diferentes tipos de convergencia, tales como: convergencia puntual, convergencia uniforme, convergencia en media y en media cuadrática, etc. El concepto de convergencia no solo es fundamental para los razonamientos teóricos del Análisis y el Análisis Funcional sino también en la Matemática Aplicada. Basta mencionar los problemas de aproximación en el Análisis Numérico y en el Cálculo Operacional. En la teoría de Optimización, por ejemplo, muchas veces no es posible determinar explícitamente el valor óptimo buscado, en cuyo caso el problema

consiste en hallar un algoritmo para construir una sucesión de valores que “converja” al valor óptimo desconocido. Para funciones en lugar de valores reales y vectoriales existen problemas análogos: aproximación de controles optimales en la teoría del control, aproximación de soluciones de ecuaciones diferenciales en Análisis Numérico, etc.

En todos los casos anteriormente mencionados se utiliza el concepto de “convergencia” basándose en una cierta noción de distancia: distancia entre dos números reales, distancia euclídeana entre dos vectores de $\mathbb{R}^n$, distancia uniforme entre dos funciones continuas sobre $[a, b]$, distancia en media o en media cuadrática entre dos funciones integrables sobre $[a, b]$, etc. En otros casos, como por ejemplo la convergencia puntual para funciones reales definidas sobre un intervalo $[a, b]$, no tenemos una sola medida de proximidad entre dos funciones, sino toda una familia de ellas que miden la proximidad entre dos funciones en cada punto $x \in [a, b]$ por separado.

Si excluimos conceptos como la convergencia puntual en una primera fase de nuestra consideración, entonces con respecto a las distancias de que hablábamos más arriba los conceptos de convergencia a ellas asociadas admiten una formulación general: una sucesión de elementos (x_i) (resp. (f_i)) converge a un elemento x (resp. f), si las x_i (resp. f_i) están tan cerca de x (resp. f) como se quiera, cuando los subíndices i son suficientemente grandes.

Seleccionando las propiedades comunes más elementales de las distancias antes mencionadas podemos enunciar a manera de definición la caracterización axiomática de una distancia o métrica sobre un conjunto arbitrario E .

D 1. (métrica, espacio métrico). Sea E un conjunto. Una aplicación $d : E \times E \rightarrow \mathbb{R}_+$ se llama métrica o distancia sobre E , si verifica los siguientes axiomas:

$$(D1) \quad d(x, y) = d(y, x) \quad (x, y \in E)$$

$$(D2) \quad d(x, y) = 0 \iff x = y \quad (x, y \in E)$$

$$(D3) \quad d(x, z) \leq d(x, y) + d(y, z) \quad (x, y, z \in E) \text{ (desigualdad triangular).}$$

Para $x, y \in E$, el número real positivo $d(x, y)$ se llama *distancia* entre x y y . Al par (E, d) se le llama *espacio métrico*.

Si A es un subconjunto de E , entonces a cada par x, y de elementos de A , le podemos asociar su distancia $d(x, y)$. Estas distancias $d(x, y) (x, y \in A)$ verifican trivialmente los axiomas (D1)–(D3), luego inducen una métrica sobre A :

$$d_A : A \times A \rightarrow \mathbb{R}_+ \quad d_A(x, y) = d(x, y) \quad (x, y \in A)$$

esta función d_A se le llama *métrica inducida* por d sobre el subconjunto A de E y (A, d_A) se llama *subespacio métrico* de (E, d) .

Ahora está claro cómo generalizar las nociones de convergencia y continuidad en espacios métricos cualesquiera.

D 2. (sucesión convergente). Sea (E, d) un espacio métrico y (x_n) una sucesión de elementos de E . Se dice que la sucesión (x_n) converge al punto $x \in E$, si para todo $\varepsilon > 0$ existe un número natural $N = N(\varepsilon)$ tal que $d(x, x_n) < \varepsilon$ para todo $n \geq N$. Al punto x se le llama *límite* de la sucesión (x_n) , utilizándose la notación $(x_n) \xrightarrow{d} x$ para indicar que (x_n) converge a x con respecto a la métrica d .

Notemos que mediante la métrica d en E , la *convergencia de sucesiones en (E, d) se reduce a la convergencia a cero de las sucesiones de números reales positivos $d(x_n, x)$* .

Una propiedad importante de la convergencia de sucesiones en los espacios métricos es que *el límite, cuando existe, es único*, lo cual se deduce obviamente de la desigualdad triangular.

No es difícil cerciorarse de que el concepto de espacio métrico no puede ser la estructura básica para la construcción axiomática de una

teoría general de la convergencia, ya que la noción de convergencia en un espacio métrico (E, d) no está biunívocamente determinada por la métrica d . En efecto, si d es una métrica en E también lo serán $\frac{d}{1+d}$ y $\inf\{1, d\}$. El lector puede verificar sin dificultad que para las sucesiones de puntos de E se cumple:

$$(x_n) \xrightarrow{d} x \iff (x_n) \xrightarrow{\frac{d}{1+d}} x \iff x_n \xrightarrow{\inf\{1, d\}} x$$

Pasemos ahora a estudiar, el segundo concepto fundamental del Análisis Matemático: la continuidad de funciones.

D 3. (continuidad puntual de funciones). Se dice que la función $f : (E, d) \rightarrow (F, \rho)$ es continua en $x_0 \in E$ si para cualquier $\varepsilon > 0$ es posible hallar un número real $\delta = \delta(\varepsilon, x_0) > 0$ tal que $d(x, x_0) < \delta$ implique $\rho(f(x), f(x_0)) < \varepsilon$.

La continuidad puntual de funciones está estrechamente vinculada a la convergencia de sucesiones en un espacio métrico, lo cual nos muestra el teorema siguiente:

T 1. La función $f : (E, d) \rightarrow (F, \rho)$ es continua en el punto $x_0 \in E$ si y solo si, f transforma toda sucesión (x_n) de E , convergente a x_0 , en una sucesión $(f(x_n))$ convergente en (F, ρ) a $f(x_0)$.

Demostración. Probemos primeramente la necesidad de esta condición. Sea f continua en x_0 y (x_n) una sucesión en (E, d) convergente a x_0 . Dado $\varepsilon > 0$, existe $\delta = \delta(\varepsilon, x_0) > 0$ tal que si $d(x, x_0) < \delta$ entonces $\rho(f(x), f(x_0)) < \varepsilon$. Pero como $(x_n) \xrightarrow{d} x_0$, se puede encontrar un número natural $N = N(\delta)$ tal que para $n \geq N$ se tiene que $d(x_n, x_0) < \delta$ y, por lo tanto, $\rho(f(x_n), f(x_0)) < \varepsilon$. Por la arbitrariedad de $\varepsilon > 0$ se concluye la convergencia de $f(x_n)$ a $f(x_0)$ en (F, ρ) .

Para demostrar que la condición del teorema es suficiente basta ver que si f no es continua en x_0 , entonces se puede hallar una sucesión (x_n)

en (E, d) que converge a x_0 y tal que $(f(x_n))$ no converge a $f(x_0)$ en (F, ρ) . En efecto, por ser f discontinua en x_0 , existe $\varepsilon > 0$ tal que para cada

$$\delta = \frac{1}{2}, \text{ es posible hallar un elemento } x_n \text{ en } E \text{ con } d(x_n, x_0) < \frac{1}{n}$$

y tal que $\rho(f(x_n), f(x_0)) > \varepsilon$. Obviamente $(x_n) \xrightarrow{d} x_0$ y sin embargo $(f(x_n))$ no converge a $f(x_0)$ en (F, ρ) . $\square$

Utilizando este resultado se obtienen inmediatamente las siguientes propiedades de las funciones continuas.

C1. (continuidad de la compuesta). Sea $(E, d), (F, \rho)$ y (G, θ) espacios métricos, $f : (E, d) \rightarrow (F, \rho)$ continua en el punto $x_0 \in E$, y $g : (F, \rho) \rightarrow (G, \theta)$ continua en el punto $y_0 = f(x_0)$ de F . Entonces la función compuesta $g \circ f(x) = g(f(x))$ es continua en x_0 . $\square$

C 2. (suma, producto y cociente de funciones reales continuas). Sean $f, g : (E, d) \rightarrow \mathbb{R}$ funciones continuas en el punto $x_0 \in E$ (considerando sobre $\mathbb{R}$ la métrica usual definida por el módulo). Entonces la suma $f + g$ y el producto $f \cdot g$ son funciones continuas en x_0 . Si además $g(x_0) \neq 0$, entonces el cociente f/g también es una función continua en x_0 . $\square$

D 4. (continuidad global). Se dice que la función $f : (E, d) \rightarrow (F, \rho)$ es continua en E si lo es en cada punto $x \in E$. No es difícil verificar que si $f : (E, d) \rightarrow (F, \rho)$ es continua y $A \subset E$, entonces también serán continuas

$$\begin{aligned} f|_A : (A, d_A) &\longrightarrow (F, \rho) \\ f : (E, d) &\longrightarrow (f[E], \rho_{f[E]}) \end{aligned}$$

donde $f|_A$ denota la restricción de la función f al subespacio (A, d_A) y $f[E] \subset F$ es la imagen de E por f .

Tanto en la definición de convergencia como en la de continuidad puntual en los espacios métricos hemos utilizado implícitamente ciertos conjuntos que permiten dar una expresión más abreviada de esos conceptos. Tales conjuntos son las llamadas *bolas abiertas* y *bolas cerradas*:

$$B_d(a, \delta) := \{x \in E | d(a, x) < \delta\}$$

(bola abierta de centro a y radio δ en (E, d))

$$B'_d(a, \delta) := \{x \in E | d(a, x) \leq \delta\}$$

(bola cerrada de centro a y radio δ en (E, d)).

Con estas definiciones, la definición de convergencia de sucesiones en (E, d) puede escribirse de la manera equivalente:

La sucesión (x_n) converge al punto x_0 en (E, d) si para todo $\varepsilon > 0$ existe un número natural $N = N(\varepsilon)$ tal que

$$x_n \in B_d(x_0, \varepsilon) \quad \text{para todo } n \geq N.$$

De forma análoga se tiene la siguiente expresión equivalente para la continuidad puntual:

La función $f : (E, d) \rightarrow (F, \rho)$ es continua en el punto $x_0 \in E$ si para todo $\varepsilon > 0$ existe un número real $\delta = \delta(\varepsilon, x_0) > 0$ tal que

$$f[B_d(x_0, \delta)] \subset B_\rho(f(x_0), \varepsilon)$$

(o también $f^{-1}[B_\rho(f(x_0), \varepsilon)] \supset B_d(x_0, \delta)$).

Notemos que diferentes métricas sobre un mismo conjunto pueden dar lugar a diferentes tipos de bolas y sin embargo, definir la misma convergencia.

Tal es el caso del plano $\mathbb{R}^2$, donde las métricas

$$\begin{aligned} d_1(x, y) &= |x_1 - y_1| + |x_2 - y_2|, \\ d_2(x, y) &= (|x_1 - y_1|^2 + |x_2 - y_2|^2)^{\frac{1}{2}}, \\ d_\infty(x, y) &= \text{máx } \{|x_1 - y_1|, |x_2 - y_2|\}, \end{aligned}$$

que satisfacen las desigualdades

$$d_\infty(x, y) \leq d_2(x, y) \leq d_1(x, y) \leq 2d_\infty(x, y),$$

generan bolas que son rombos, círculos y cuadrados respectivamente.

De la representación geométrica de esta situación y, utilizando la definición de sucesiones a través del concepto de bola, se puede concluir que las sucesiones convergentes al punto a en $\mathbb{R}^2$ con cada una de las métricas d_1, d_2 y d_∞ , son las mismas. Luego, esto nos indica que *en la definición de convergencia de sucesiones pueden sustituirse las bolas por conjuntos más generales siempre que estos contengan determinadas bolas*.

D 5. (vecindad) Sea (E, d) un espacio métrico y $x_0 \in E$. Se dice que $V \subset E$ es una *vecindad* de x_0 o que x_0 es un *punto interior* de V , si existe un número real $\delta > 0$, tal que $V \supset B_d(x_0, \delta)$.

La colección de todas las vecindades del punto x_0 en (E, d) se denota por $\mathcal{V}_d(x_0)$.

Se dice que V es vecindad del subconjunto A de E si $V \in \mathcal{V}_d(x)$ para todo $x \in A$. La colección de todas las vecindades del subconjunto A en (E, d) se denota por $\mathcal{V}_d(A)$.

Observación. Obviamente toda bola abierta en un espacio métrico es vecindad de cada uno de sus puntos.

Utilizando la noción de vecindad pueden expresarse las definiciones de convergencia y continuidad puntual de la manera siguiente:

La sucesión (x_n) en (E, d) converge al punto $x_0 \in E$, si para toda $V \in \mathcal{V}_d(x_0)$ existe un número natural $N = N(V)$ tal que

$$x_n \in V \quad \text{para todo } n \geq N.$$

La función $f : (E, d) \rightarrow (F, \rho)$ es continua en el punto $x_0 \in E$, si para toda $V \in \mathcal{V}_\rho(f(x_0))$ existe $W \in \mathcal{V}_d(x_0)$ tal que

$$f[W] \subset V$$

(es decir, $f^{-1}[V] \in \mathcal{V}_d(x_0)$)

Para caracterizar la continuidad global de una función el concepto de vecindad no es suficiente, porque se refiere a un punto particular.

D 6. (interior de un conjunto, conjuntos abiertos). Para un subconjunto A del espacio métrico (E, d) el conjunto de los puntos interiores de A , es decir

$$\{a \in A \mid A \in \mathcal{V}_d(a)\}$$

se denota por $\overset{\circ}{A}$ ó $\text{int } A$ y se llama *interior de A* .

El subconjunto A se llama *abierto* en (E, d) si es vecindad de cada uno de sus puntos, o sea si $A = \overset{\circ}{A}$. A la familia de todos los subconjuntos abiertos del espacio métrico (E, d) la denotaremos por θ_d . La familia θ_d tiene las siguientes propiedades fundamentales y que el lector puede verificar sin dificultad:

A1 $\emptyset, E$ pertenecen a θ_d .

A2 si $A_1, A_2, \dots, A_n$ es una subfamilia finita de elementos de θ_d , entonces $\bigcap_{i=1}^n A_i \in \theta_d$.

A3 si $\{A_i\}_{i \in I}$ es una subfamilia arbitraria de θ_d , entonces $\bigcup_{i \in I} A_i \in \theta_d$.

De la definición de punto interior de un conjunto en un espacio métrico se deduce que *los conjuntos abiertos son precisamente aquellos que pueden expresarse como unión de una colección cualquiera de bolas abiertas*. No resulta difícil demostrar la siguiente caracterización de las funciones globalmente continuas.

T 2. (caracterización de la continuidad global) La función $f : (E, d) \rightarrow (F, \rho)$ es continua en E , si y sólo si, la imagen recíproca de cada abierto en (F, ρ) es un abierto en (E, d) .

Demostración. La suficiencia de la condición es evidente. En efecto, para cada $a \in E$ la bola $B_\rho(f(a), \varepsilon)$ es un abierto en F y si su imagen recíproca $f^{-1}[B_\rho(f(a), \varepsilon)]$ es un conjunto abierto en E esto implica en particular que existe un número real $\delta > 0$, tal que

$$B_d(a, \delta) \subset f^{-1}[B_\rho(f(a), \varepsilon)],$$

de donde se concluye la continuidad de f .

Recíprocamente, sea f continua y A abierto en (F, ρ) . Para cada punto $a \in A$ existe un número $\delta_a > 0$ tal que

$$A = \bigcup_{a \in A} B_\rho(a, \varepsilon_a);$$

y, por lo tanto

$$f^{-1}[A] = \bigcup_{a \in A} f^{-1}B_\rho(a, \varepsilon_a) = \bigcup \{f^{-1}[B_\rho(a, \varepsilon_a)] | a \in f[f^{-1}[A]]\}.$$

Por ser f continua, para cada $a \in f[f^{-1}[A]]$ existe $b \in f^{-1}[A]$ y $\delta_a > 0$ tales que

$$f^{-1}[B_\rho(a, \varepsilon_a)] \supset B_d(b, \delta_a), \quad f(b) = a.$$

Pero entonces

$$f^{-1}[A] \supset \bigcup \{B_d(b, \delta_a) | b \in f^{-1}[A]\}$$

y como la inclusión en sentido opuesto es evidente, se tendrá que

$$f^{-1}[A] = \bigcup \{B_d(b, \delta_a) | b \in f^{-1}[A]\},$$

de donde se concluye que $f^{-1}(A)$ es abierto en (E, d) . $\square$

Hasta ahora hemos desarrollado las nociones de convergencia y continuidad con el basamento operativo de los conceptos de vecindad y conjunto abierto. Existen también otros conceptos de gran importancia en el análisis matemático y que pasaremos a definir a continuación.

D 7. (punto adherente, clausura) Sea A un subconjunto del espacio métrico (E, d) . El punto $a \in E$ se llama *punto adherente* del conjunto A si toda vecindad $V \in \mathcal{V}(a)$ tiene una intersección no vacía con A . El conjunto de todos los puntos adherentes de A se llama *adherencia* o *clausura* de A y se denota por $\overline{A}$.

Los puntos adherentes de un conjunto A en el espacio métrico (E, d) se clasifican de manera natural en *puntos de acumulación* y *puntos aislados*. Los puntos de acumulación de A son aquellos $a \in E$ para los cuales toda $V \in \mathcal{V}(a)$ contiene algún punto (y de hecho una cantidad infinita) de A diferente de a . En caso contrario, es decir cuando para alguna $V \in \mathcal{V}(a)$ se tiene que $V \cap A = \{a\}$, entonces se dice que a es un punto aislado de A .

A partir de un subconjunto cualquiera A de (E, d) , se pueden construir los siguientes:

$$\text{int } A, \quad \text{ext } A \quad \text{y} \quad \text{fr } A$$

donde $\text{ext } A$ es el llamado *exterior* de A y se define como

$$\text{ext } A = \text{int } (E \setminus A),$$

y $\text{fr } A$ es la llamada *frontera* de A y se define mediante

$$\text{fr } A = \overline{A} \cap \overline{E \setminus A}.$$

Obviamente los conjuntos $\text{int } A$, $\text{ext } A$ y $\text{fr } A$ forman una partición de (E, d) .

La operación clausura, definida sobre los subconjuntos de un espacio métrico cualquiera, cumple la siguiente propiedad evidente:

$$A \subset B \Rightarrow \overline{A} \subset \overline{B}.$$

T 3. (caracterizaciones de la clausura). En un espacio métrico (E, d) son equivalentes:

- i) $a \in \overline{A}$,
- ii) $d(a, A) = 0$,
- iii) Existe una sucesión (x_n) de elementos de A tal que $x_n \xrightarrow{d} a$.

Demostración. Se deja de ejercicio al lector. $\square$

El conjunto A se llama *denso* en (E, d) si $\overline{A} = E$.

Ejemplo 1. (espacios de funciones continuas y de funciones de cuadrado integrable). Denotamos por $C[a, b]$ al conjunto de todas las funciones reales continuas definidas sobre el intervalo $[a, b]$, y por $L_2[a, b]$ al conjunto de las funciones reales definidas sobre $[a, b]$ cuyo cuadrado es integrable en sentido de Lebesgue. Es conocido el *teorema de aproximación de Weierstrass* que dice que toda función real continua definida sobre el intervalo $[a, b]$ puede ser uniformemente aproximada por polinomios algebraicos (véase el libro Elementos de la teoría de funciones y el Análisis funcional de A.N. Kolmogorov [12]).

No es difícil comprobar que $d_\infty(f, g)$, definida como

$$d_\infty(f, g) = \sup_{a \leq x \leq b} |f(x) - g(x)|, \quad (1)$$

es una métrica sobre $C[a, b]$. Utilizando las definiciones introducidas anteriormente, el teorema de Weierstrass se expresa diciendo que el conjunto $\mathcal{P}[a, b]$ de los polinomios algebraicos sobre $[a, b]$ constituyen un subconjunto denso en el espacio métrico $(C[a, b], d_\infty)$.

Si para $f, g \in L_2[a, b]$ definimos

$$d_2(f, g) = \left(\int_a^b |f(x) - g(x)|^2 dx \right)^{\frac{1}{2}}, \quad (2)$$

donde la integral se interpreta en el sentido de Lebesgue, entonces se prueba en cualquier curso de *teoría de la medida*, que $\mathcal{P}[a, b]$ es también denso en $(L_2[a, b], d_2)$.

No resulta difícil concluir que tanto $(C[a, b], d_\infty)$ como $(L_2[a, b], d_2)$ tienen subconjuntos densos que son numerables; por ejemplo el subconjunto de $\mathcal{P}[a, b]$ constituido por los polinomios algebraicos cuyos coeficientes son números racionales.

En general, sustituyendo el intervalo $[a, b]$ por cualquier subconjunto abierto Ω de $\mathbb{R}^n$, puede probarse que $C_0^\infty(\Omega)$ es denso en $L_2(\Omega)$, donde $L_2(\Omega)$ se define análogamente a $L_2[a, b]$ y $C_0^\infty(\Omega)$ está constituido por las funciones reales infinitamente diferenciables definidas sobre Ω y con soporte compacto (véase el libro *Real and Complex Analysis* de Walter Rudin [3]). $\square$

P 1. (coincidencia de funciones continuas). Sean f, g funciones continuas definidas sobre el espacio métrico (E, d) y con valores en (F, ρ) y sea $A \subset E$ denso en (E, d) . Si f y g coinciden sobre A , entonces también coinciden sobre E .

Demostración. Se deja al lector. $\square$

Resulta natural pensar que para problemas de convergencia y continuidad, los conjuntos que contienen a todos sus puntos límites deben tener una importancia fundamental.

D 8. (conjunto cerrado). Un subconjunto A del espacio métrico (E, d) se llama *cerrado* si $A = \overline{A}$.

Obviamente el conjunto A es cerrado en (E, d) cuando contiene a todos sus puntos de acumulación o cuando contiene a todos los límites en E de sucesiones convergentes de elementos de A

T 4. (Caracterización de los cerrados). Un subconjunto A de (E, d) es cerrado si y sólo si $E \setminus A$ es abierto.

Demostración. Se deduce inmediatamente de la definición $\square$

El resultado anterior establece una dualidad entre las familias de los conjuntos abiertos y la de los conjuntos cerrados en un espacio métrico. En particular se concluyen las siguientes propiedades duales a (A1) – (A3), inherentes a la familia $\mathcal{T}_d$ de los subconjuntos cerrados en el espacio métrico (E, d) :

(C 1) $\phi, E \in \mathcal{T}_d$.

(C 2) Si $C_1, C_2, \dots, C_n$ es una subfamilia finita de elementos de $\mathcal{T}_d$ entonces $\bigcup_{i=1}^n C_i \in \mathcal{T}_d$.

(C 3) Si $(C_i)_{i \in I}$ es una subfamilia arbitraria de $\mathcal{T}_d$, entonces $\bigcap_{i \in I} C_i \in \mathcal{T}_d$.

Para cualquier subconjunto A de (E, d) , se tiene que $\overline{\overline{A}} = \overline{A}$ y, por lo tanto, $\overline{A}$ es cerrado. Más aún, $\overline{A}$ es el menor conjunto cerrado en (E, d) que contiene a A .

Debido a la dualidad entre las propiedades (A1)–(A3) de los abiertos y (C1)–(C3) de los cerrados, tenemos la siguiente caracterización de las funciones continuas (compárese con el Teorema 2).

T 5. (caracterización de funciones continuas). Para una función $f : (E, d) \rightarrow (F, \rho)$ son equivalentes:

- i) f es continua,
- ii) $f^{-1}[C] \in \mathcal{T}_d$ para todo $C \in \mathcal{T}_\rho$,
- iii) $f[\overline{A}] \subset \overline{f[A]}$ para todo $A \subset E$.

Demostración. Se deja al lector. $\square$

Un criterio muy utilizado en la práctica para probar que un conjunto es cerrado lo brinda la proposición siguiente:

P 2. Sean $f, g : (E, d) \rightarrow (F, \rho)$ funciones continuas. El conjunto de los puntos de coincidencia de f y g es cerrado en E , es decir,

$$\{x \in E : f(x) = g(x)\} \in \tau_d.$$

Demostración. La propiedad de unicidad del límite en los espacios métricos es una consecuencia inmediata del hecho que dos puntos cualesquiera diferentes pueden ser separados por vecindades abiertas disjuntas. Es fácil verificar que esta propiedad de separación es equivalente a que la diagonal D_f del espacio producto $F \times F$, provisto de la métrica

$$d \times d((x_1, y_1), (x_2, y_2)) = (d^2(x_1, x_2) + d^2(y_1, y_2))^{\frac{1}{2}},$$

es un subconjunto cerrado en $(F \times F, d \times d)$.

Si definimos

$$\tilde{f} : E \rightsquigarrow F \times F$$

mediante

$$\tilde{f}(x) = (f(x), g(x)),$$

no es difícil verificar que $\tilde{f}$ es continua por serlo f y g , y que además

$$\{x \in E : f(x) = g(x)\} = \tilde{f}^{-1}[D_F]$$

luego, aplicando el resultado del teorema 2, se concluye que el conjunto de los puntos de coincidencia de f y g es cerrado en (E, d) . $\square$

Como hemos visto en la demostración del teorema anterior, la propiedad de separación por abiertos es muy importante en el estudio de las funciones continuas. Resulta además que la propiedad que tienen ciertos subconjuntos en los espacios métricos, de poder ser separados por abiertos, está íntimamente relacionada con la posibilidad de prolongar una función continua a un dominio más amplio conservando la continuidad.

No es difícil demostrar que en un espacio métrico, cualquier conjunto cerrado puede ser separado de todo punto que no pertenece a él por medio de vecindades abiertas disjuntas. Esta propiedad es llamada *regularidad de los espacios métricos* y de ella se concluye el teorema siguiente

T 6. (prolongación de funciones continuas a la clausura) Sean A un subconjunto denso del espacio métrico (E, d) y $f : (A, d_A) \rightarrow (F, \rho)$ una aplicación continua. Entonces existe una única prolongación continua

$$\bar{f} : (E, d) \rightarrow (F, \rho) \quad (\bar{f}(x) = f(x), \text{ para todo } x \in A),$$

si y sólo si, para cada $z \in E$ existe $\lim_{\substack{x \rightarrow z \\ x \in A}} f(x)$ en (F, ρ) .

Demostración. Sea $f : (A, d_A) \rightarrow (F, \rho)$ continua y supongamos

que existe una prolongación continua $\bar{f}$ de f a todo E .

Consideremos el valor de $\bar{f}$ en un punto arbitrario $z \in E$. Puesto que A es denso en (E, d) , existe una sucesión (x_n) en A tal que $(x_n) \xrightarrow{d} z$ (ver Teorema 3). Para la prolongación continua $\bar{f}$, tendremos necesariamente

$$\bar{f}(z) = \lim f(x_n)$$

y como este límite no depende de la elección de la sucesión (x_n) en A siempre que $(x_n) \xrightarrow{d} z$, pondremos

$$\bar{f}(z) = \lim_{\substack{x \rightarrow z \\ x \in A}} f(x) \quad (3)$$

Luego el valor de la prolongación $\bar{f}$ en $z \in E$ puede calcularse por medio de los valores de f en los puntos de A que son cercanos a z .

Veamos ahora que, inversamente, mediante (3) se define una prolongación continua de f a E si existe $\lim_{\substack{x \rightarrow z \\ x \in A}} f(x)$ para todo $z \in E$.

Sea W una vecindad de $\bar{f}(z)$ en (F, ρ) . Por la propiedad de regularidad de (F, ρ) existe una vecindad cerrada V de $\bar{f}(z)$ tal que $V \subset W$ (¿Por qué?). Según la suposición de continuidad de f , hay una vecindad abierta $U \in \mathcal{U}_d(z)$ tal que $f[U \cap A] \subset V$. Es fácil ver que $\overline{U \cap A} = \bar{U}$ (demuéstrello). Por iii) del Teorema 5 se tiene que

$$\bar{f}[U] \subset \bar{f}[\bar{U}] \subset f[U \cap A] \subset V,$$

con lo cual se prueba la continuidad de $\bar{f}$ en z . $\square$

Ejemplo 2. (integral de Cauchy).

Consideremos sobre un intervalo $[a, b] \subset \mathbb{R}$ el espacio vectorial E de todas las funciones *escalonadas*:

$$h : [a, b] \rightarrow \mathbb{R}$$

tales que $h[[a, b]] = \{h_1, \dots, h_n\}$ es un conjunto finito y $h^{-1}[\{h_i\}]$ es para cada $i = 1, \dots, n$ una unión finita de subintervalos de $[a, b]$.

A estas funciones se les asigna de manera natural una *integral*

$$\int_a^b h(x)dx = \sum_{i=1}^n h_i(x_i - x_{i-1}) \text{ con } \{h_i\} = h[x_{i-1}, x_i].$$

El problema de la construcción de las integrales de *Cauchy*, *Riemann* o *Lebesgue*, consiste en generalizar esta integral natural a una clase más amplia de funciones. En el caso particular de la integral de Lebesgue, que generaliza a las dos anteriores, los conjuntos $h^{-1}[\{h_i\}]$ pueden tener una estructura más compleja (conjuntos medibles) y la posibilidad de establecer un concepto de *longitud* para tales conjuntos se estudia en la llamada *teoría de medida*.

Consideremos aquí solamente la integral de Cauchy, cuya construcción es elemental. Sea R el espacio vectorial de todas las funciones *regladas* sobre $[a, b]$, o sea, el espacio de todas las funciones que pueden aproximarse uniformemente sobre $[a, b]$ por funciones escalonadas. El conjunto provisto de la métrica uniforme (1), es un espacio métrico que contiene a E como subespacio métrico. No es difícil verificar que $C[a, b]$ es también un subespacio vectorial de R .

La aplicación real $I : h \rightsquigarrow I(h) = \int_a^b h(x)dx$ definida sobre E satisface la desigualdad

$$|I(h) - I(g)| \leq (b - a)d_\infty(h, g) \quad (4)$$

para todos $h, g \in E$, de donde se deduce la continuidad de

$$I : (E, d_\infty) \rightarrow \mathbb{R}.$$

Para probar que $\lim_{\substack{h \rightarrow f \\ h \in E}} I(h)$ existe para todo $f \in R$, es suficiente demostrar la existencia de $\lim I(h_n)$ para toda sucesión (h_n) de funciones escalonadas que converge a f en (R, d_∞) . Pero, tomando $h = h_n$ y $g = h_m$ en (4), se comprueba que $I(h_n)$ es una sucesión de Cauchy, y por tanto convergente, en $\mathbb{R}$.

Del Teorema 6 se concluye que I posee una prolongación continua $\bar{I} : (R, d_\infty) \rightarrow \mathbb{R}$. Se prueba sin dificultad que $\bar{I}$ es lineal. Para cada $f \in R$, se utiliza la notación $\bar{I}(f) = \int_a^b f(x)dx$ y se le llama *integral de Cauchy* de la función f . $\square$

Dos subconjuntos A y B de un conjunto E se dice que son *separados por la función* $f : E \rightarrow \mathbb{R}$, si existen $\alpha, \beta \in \mathbb{R}$ tales que $\alpha < \beta$ y

$$f(x) < \alpha \quad \text{para todo } x \in A, \quad (5)$$

$$f(x) > \beta \quad \text{para todo } x \in B.$$

Esto significa intuitivamente que A y B son separados por dos líneas (o superficies) de nivel de f , a saber $\{f = \alpha\}$ y $\{f = \beta\}$.

Si E es un espacio métrico provisto de la métrica d y f es continua, entonces la separación de los conjuntos A y B por f implica la separación de A y B mediante vecindades abiertas disjuntas. En efecto basta notar que la condición (5) es equivalente a que

$$A \subset \{f < \alpha\} \text{ y } B \subset \{f > \beta\}.$$

Pero

$$\{f < \alpha\} = f^{-1}(-\infty, \alpha) \quad \{f > \beta\} = f^{-1}(\beta, +\infty)$$

y por el Teorema 2 estos conjuntos son abiertos en E ya que f es continua.

Llamaremos *función de Urysohn* para el par de subconjuntos A y B del espacio métrico (E, d) a toda función real continua $f : E \rightarrow [0, 1]$ tal que

$$f[A] = \{0\} \text{ y } f[B] = \{1\}.$$

Obviamente una función de Urysohn para el par (A, B) separa a los conjuntos A y B en (E, d) . Recíprocamente, si existe alguna función

real continua sobre (E, d) que separe a los conjuntos A y B , entonces también hay una función de Urysohn para el par (A, B) .

En efecto si $f : (E, d) \rightarrow \mathbb{R}$ real continua es tal que

$$A \subset \{f < \alpha\}, \quad B \subset \{f > \beta\}$$

donde $\alpha, \beta \in \mathbb{R}, \alpha < \beta$, entonces la función $\tilde{f}(x) = \max(\alpha, \min(\beta, f(x)))$ es una función real continua sobre (E, d) que toma el valor α sobre A y el valor β sobre B . Luego la función $\hat{f} = (\tilde{f} - \alpha)/(\beta - \alpha)$ es una función de Urysohn para el par (A, B) en (E, d) .

Hasta el momento hemos visto que dos subconjuntos A y B del espacio métrico (E, d) pueden ser separados por una función real continua si y sólo si existe una función de Urysohn para el par (A, B) . Por otra parte vimos que la separación por funciones continuas implica la separación por vecindades abiertas disjuntas. Utilizando este hecho veamos que *los subconjuntos cerrados disjuntos de un espacio métrico pueden ser separados por vecindades abiertas disjuntas*. Esta propiedad de los espacios métricos es llamada *normalidad*. Este hecho es evidente cuando la distancia entre ambos conjuntos

$$d(A, B) = \inf_{\substack{x \in A \\ y \in B}} d(x, y)$$

es estrictamente positiva. Sin embargo, no resulta tan evidente cuando $d(A, B) = 0$. Para cerciorarse de ello basta considerar en $\mathbb{R}^2$, los conjuntos cerrados $A = \{(x, \frac{1}{x}) : x > 0\}$ y

$$B = \{(x, 0) : x \geq 0\}.$$

Se tiene el teorema siguiente:

T 7. (existencia de funciones de Urysohn) Para cada par de conjuntos cerrados disjuntos en un espacio métrico existe una función

de Urysohn.

Demostración. Sean A y B dos conjuntos cerrados disjuntos en (E, d) . Entonces

$$f(x) = \frac{d(x, A)}{d(x, A) + d(x, B)}$$

es una función de Urysohn para el par (A, B) . $\square$

No resulta difícil interpretar este teorema como *un teorema de prolongación continua para ciertas funciones reales continuas*.

En efecto, si A y B son subconjuntos cerrados disjuntos en (E, d) , la función

$$f : A \cup B \rightarrow [0, 1],$$

definida por

$$f(x) = \begin{cases} 0 & \text{si } x \in A, \\ 1 & \text{si } x \in B \end{cases} \quad (6)$$

es continua, considerando $A \cup B$ provisto de la métrica inducida por d .

Las funciones de Urysohn para el par (A, B) son precisamente las prolongaciones reales continuas de f a todo E que conservan a $[0, 1]$ como espacio de llegada. El Teorema 7 nos asegura que existen tales prolongaciones.

Como hemos visto $A \cup B$ es cerrado en (E, d) y f , definida en (6), es una función real continua muy especial sobre el cerrado $A \cup B$. Luego el Teorema 7 es un teorema de prolongación para una clase muy especial de funciones sobre un cerrado en un espacio métrico. Sin embargo se obtiene el resultado general siguiente.

T 8. (de Tietze) Sea C cerrado en el espacio métrico (E, d) y $f : (C, d_C) \rightarrow [a, b]$. Entonces existe una prolongación continua $\bar{f} : (E, d) \rightarrow [a, b]$, de f a todo E . $\square$

Existe una clase muy importante de espacios métricos que tienen localmente la estructura de un espacio euclideo $\mathbb{R}^n$, o sea en dichos espacios cada punto tiene una vecindad que está en correspondencia biyectiva y bicontinua con $\mathbb{R}^n$ (las llamadas variedades topológicas). Sobre tales espacios no es difícil definir lo que son funciones diferenciables, puesto que la diferenciabilidad es, al igual que la continuidad, una *noción local*. En cambio, la integral de una función depende del comportamiento de la función sobre todo el espacio, la integral es una *noción global*. Para introducir conceptos globales sobre *espacios localmente euclideos* se necesitan técnicas de vincular nociones y propiedades globales con nociones y propiedades locales. Uno de estos métodos está estrechamente vinculado con la *propiedad de normalidad* de los espacios métricos, es decir con la posibilidad de separar subconjuntos cerrados disjuntos por vecindades abiertas disjuntas.

Comencemos por analizar los cubrimientos abiertos de un espacio métrico. No resulta difícil verificar que la propiedad de normalidad en los espacios métricos es equivalente a que para cada cerrado C en (E, d) y V vecindad abierta de C , existe otra vecindad abierta W de C , tal que

$$C \subset W \subset \overline{W} \subset V \quad (7)$$

De (7) se concluye que la propiedad de normalidad puede ser interpretada como una propiedad que asegura que en todo cubrimiento (V_1, V_2) de (E, d) , uno de sus elementos, por ejemplo V_1 puede sustituirse por un abierto más pequeño W_1 tal que

$$\overline{W_1} \subset V_1$$

de manera que el par (W_1, V_2) continúa siendo un cubrimiento abierto de E .

En efecto si $V_1 \cup V_2 = E$ y $V_1, V_2 \in \theta_d$, entonces $C = E \setminus V_2 \subset V_1$ es cerrado, luego existe, según (7) un abierto W_1 tal que $C \subset W_1 \subset \overline{W_1} \subset V_1$. Obviamente este W_1 cumple las propiedades requeridas.

Aplicando este método sucesivamente a todos los abiertos de un cubrimiento abierto finito $(V_1, \dots, V_n)$ se obtiene otro cubrimiento abierto *más fino* $(W_1, \dots, W_n)$ llamado *encogimiento* de (V_i) , tal que

$\overline{W}_i \subset V_i, i = 1, \dots, n$. Esta propiedad de los espacios métricos es mucho más general y se cumple para todos los *cubrimientos abiertos localmente finitos*.

D 9. (cubrimiento puntualmente y localmente finito). Sea $(U_\alpha)_{\alpha \in I}$ un cubrimiento del espacio métrico (E, d) . Se dice que $(U_\alpha)_{\alpha \in I}$ es *puntualmente finito* si cada punto $x \in E$ está contenido solamente en un número finito de los conjuntos U_α . El cubrimiento $(U_\alpha)_{\alpha \in I}$ es *localmente finito* si para cada $x \in E$ existe una vecindad $V \in \mathcal{V}(x)$ que intersecta solamente a un número finito de los conjuntos U_α .

Como mencionamos anteriormente puede probarse que *cada cubrimiento abierto localmente finito de un espacio métrico admite un encogimiento*. Es decir si $(U_\alpha)_{\alpha \in I}$ es un cubrimiento de (E, d) , puntualmente finito, donde cada U_α es abierto, entonces puede hallarse otro cubrimiento abierto $(W_\alpha)_{\alpha \in I}$ de (E, d) , tal que

$$W_\alpha \subset \overline{W}_\alpha \subset U_\alpha, \text{ para todo } \alpha \in I. \quad (8)$$

Sean ahora $(U_\alpha)_{\alpha \in I}$ un cubrimiento abierto localmente finito del espacio métrico (E, d) y $(W_\alpha)_{\alpha \in I}$ un encogimiento de dicho cubrimiento. Según el Teorema 7 existe para todo $\alpha \in I$ una función continua $f_\alpha : E \rightarrow [0, 1]$ que es igual a 1 sobre $\overline{W}_\alpha$ y se anula fuera de U_α . Estas funciones nos servirán para descomponer toda función real continua sobre E en una familia de funciones que solo localmente son diferentes de cero. Comencemos por ver algunas definiciones.

D 10. (soporte de una función). Sea (E, d) un espacio métrico y $f : (E, d) \rightarrow \mathbb{R}$ una aplicación. Entonces el conjunto cerrado

$$\text{sop } (f) = \overline{\{x \in E : f(x) \neq 0\}}$$

se llama *soporte de f* . En otras palabras $\text{sop } f$ es el cerrado más pequeño fuera del cual f se anula.

No resulta difícil demostrar que si $(f_\alpha)_{\alpha \in I}$ es una familia de funciones reales continuas sobre (E, d) tal que la familia de sus soportes es localmente finita, entonces la función suma

$$f = \sum_{\alpha \in I} f_\alpha : (E, d) \rightarrow \mathbb{R} \text{ es continua.}$$

D 11. (partición de unidad, partición de unidad subordinada a un cubrimiento) Una familia $(\varphi_\alpha)_{\alpha \in I}$ de funciones continuas $\varphi_\alpha : (E, d) \rightarrow [0, 1]$ se llama *partición de unidad* sobre (E, d) si cumple:

- i) La familia $(\text{sop } \varphi_\alpha)_{\alpha \in I}$ es localmente finita
- ii) $\sum_{\alpha \in I} \varphi_\alpha(x) = 1$ para todo $x \in E$.

sea además $(U_\alpha)_{\alpha \in I}$ un cubrimiento de (E, d) . Entonces la partición de unidad $(\varphi_\alpha)_{\alpha \in I}$ se llama *subordinada al cubrimiento* $(U_\alpha)_{\alpha \in I}$ si para todo $\alpha \in I$ se tiene $\text{sop } \varphi_\alpha \subset U_\alpha$. Ahora podemos enunciar el teorema fundamental siguiente:

T 8. (existencia de particiones de unidad subordinadas) En un espacio métrico todo cubrimiento abierto, localmente finito, posee una partición de unidad subordinada.

Demostración. Sea $(U_\alpha)_{\alpha \in I}$ un cubrimiento abierto localmente finito de (E, d) . Existe un encogimiento $(W_\alpha)_{\alpha \in I}$, tal que cumple (8). Sabemos además por el Teorema 7 que existen funciones continuas $f_\alpha : (E, d) \rightarrow [0, 1]$ tales que $f_\alpha[\overline{W}] = 1$ y f_α se anula fuera de U_α . Utilizando la propiedad de normalidad de los espacios métricos no es difícil comprobar que f_α puede ser elegida de manera tal que $\text{sop } (f_\alpha) \subset U_\alpha$.

Como $(U_\alpha)_{\alpha \in I}$ es localmente finita, la familia $(\text{sop}(f_\alpha))_{\alpha \in I}$ también lo será y por esta razón la función suma $f = \sum_{\alpha \in I} f_\alpha : (E, d) \rightarrow \mathbb{R}$ es

continua. Para cada $x \in E$ se tiene que $f(x) \geq 1$ ya que x tiene que estar en algún $\overline{W}_\alpha$.

Definamos ahora $\varphi_\alpha = f_\alpha/f$. Obviamente $\varphi_\alpha : (E, d) \rightarrow [0, 1]$ es continua y cumple $\text{sop } (\varphi_\alpha = \text{sop } (f_\alpha) \subset U_\alpha (\alpha \in I)$. Por definición, tenemos $\sum_{\alpha \in I} \varphi_\alpha(x) = (1/f(x)) \sum_{\alpha \in I} f_\alpha(x) = 1$ para todo $x \in E$. Luego $(\varphi_\alpha)_{\alpha \in I}$ es una partición de unidad subordinada a $(U_\alpha)_{\alpha \in I}$. $\square$

Si (E, d) es un espacio métrico y A, B dos cerrados disjuntos en E , entonces $(E \setminus A, E \setminus B)$ es un cubrimiento abierto localmente finito de (E, d) . Supongamos que (φ_1, φ_2) es una partición de unidad subordinada al cubrimiento abierto $(E \setminus A, E \setminus B)$. Entonces $\text{sop } (\varphi_1) \subset E \setminus A$ y $\text{sop } (\varphi_2) \subset E \setminus B$, o sea $\varphi_1[A] = \varphi_2[B] = 0$. Como $\varphi_1(x) + \varphi_2(x) = 1$ para todo $x \in E$, concluimos que $\varphi_1[B] = \varphi_2[A] = 1$. Veamos que φ_1 es una función de Urysohn para el par (A, B) y φ_2 una función de Urysohn para el par (B, A) . Esto nos muestra que *el enunciado del Teorema 8 para todo cubrimiento de dos abiertos es una propiedad equivalente a la normalidad en los espacios métricos*.

En el transcurso de este epígrafe hemos podido constatar el papel central que juegan los conjuntos abiertos de un espacio métrico en el estudio de los problemas relacionados con convergencia y continuidad. La teoría matemática dedicada al estudio de la convergencia y continuidad en toda su amplitud se llama *Topología*. Se dice que sobre un conjunto E hay dada una estructura topológica si de alguna manera está determinada una familia de subconjuntos de E que reúnan las propiedades (A1) — (A3) que satisfacen los abiertos de un espacio métrico. O sea, la métrica no es más que una herramienta que permite definir una estructura topológica sobre un conjunto de manera que los conceptos de convergencia y continuidad a ella asociados tengan una expresión más cercana a los ya conocidos en los conjuntos numéricos.

Dos métricas se llaman *topológicamente equivalentes* si definen las mismas clases de conjuntos abiertos. No es difícil probar (¡ demuéstrelolo!)

que *dos métricas son topológicamente equivalentes si y sólo si definen las mismas sucesiones convergentes*. Luego, *continuidad y convergencia dependen esencialmente de la colección de los abiertos pero no de la métrica específica que la define*.

En toda estructura matemática; ya sea la *estructura conjuntista*, la *estructura topológica*, la *estructura uniforme*, la *estructura geométrica*, la *estructura diferencial*, la *estructura algebraica*, la *estructura bornológica*, la *estructura medible*, etc; existe siempre un concepto de *isomorfismo* que conserva las propiedades que son intrínsecas a una estructura dada. De esta forma, dentro de una misma estructura, dos entes isomorfos son indistinguibles.

Por ejemplo, desde el punto de vista de la teoría de conjuntos, dos conjuntos son isomorfos si existe una biyección entre ellos. La estructura métrica, por ser una estructura adicional sobre los conjuntos, debe tener un concepto de isomorfismo tal que conserve el ya aceptado para los conjuntos. Se dice entonces que dos espacios métricos (E, d) y (F, ρ) son isomorfos si existe una biyección $f : E \rightarrow F$ tal que

$$\rho(f(x), f(y)) = d(x, y), \text{ para todos } x, y \in E \quad (9)$$

Para la estructura topológica dos espacios métricos son isomorfos, si existe un *homeomorfismo* entre ellos. Es decir, si existe una aplicación biyectiva entre dichos espacios y tal que tanto ella como su inversa sean continuas.

Del Teorema 2 se concluye que si los espacios métricos (E, d) y (F, ρ) son homeomorfos, se tiene una correspondencia biunívoca no solamente entre los elementos de E y F , sino también entre los elementos de θ_d y θ_ρ .

Precisamente llamaremos *propiedades y nociones topológicas* en los espacios métricos a aquellas que se mantienen invariantes ante homeomorfismos. En el próximo epígrafe (1.2) estudiaremos otras nociones topológicas importantes en los espacios métricos: *compacidad* y *conexión*. También veremos en § 1.2 otras nociones no topológicas en los

espacios métricos: sucesión de *Cauchy* y *completitud*. Estas dos últimas nociones son propias de la llamada *estructura uniforme*. En este libro no se estudiarán las particularidades de la estructura uniforme. El lector interesado al respecto puede consultar el libro *Topología General* de D. Hinrichsen y otros [4].

En los epígrafes de este capítulo, posteriores al § 1.2, nos dedicaremos a estudiar cómo se enriquecen las propiedades métricas y topológicas de un espacio métrico cuando su estructura es compatible con otra cierta estructura algebraica (espacios normados) o geométrica (espacios euclidianos).

§ 1.2 Completitud, compacidad y conexidad en los espacios métricos.

La veracidad de muchos teoremas importantes del Análisis matemático (encaje de intervalos, Bolzano, Weierstrass, Rolle, etc.) depende en gran medida de propiedades de completitud, compacidad y conexidad, inherentes a ciertos subconjuntos de los espacios métricos $\mathbb{R}^n$. En este epígrafe veremos cómo se generalizan estos conceptos a los espacios métricos. Además estudiaremos sus propiedades fundamentales y cómo se relacionan entre sí.

Algunos de los resultados serán expuestos sin demostración con el objetivo de abreviar la exposición. El lector interesado puede hallar las demostraciones en el libro *Análisis matemático en espacios métricos* de A. Fraguera [1].

Los dos conceptos fundamentales asociados a la estructura uniforme de los espacios métricos son *sucesión de Cauchy* y *función uniformemente continua*.

D 1. (sucesión de Cauchy) La sucesión (x_n) del espacio métrico

(E, d) se llama de Cauchy si para cada $\varepsilon > 0$ es posible hallar un número natural $n_0 = n_0(\varepsilon)$, tal que

$$m, n \geq n_0 \Rightarrow d(x_m, x_n) < \varepsilon.$$

D 2. (función uniformemente continua) Una función

$$f : (E, d) \rightarrow (F, \rho)$$

se llama uniformemente continua si para cada $\varepsilon > 0$ es posible hallar un número real $\delta = \delta(\varepsilon, f) > 0$, tal que para todos $x, y \in E$:

$$d(x, y) < \delta \Rightarrow \rho(f(x), f(y)) < \varepsilon.$$

Notemos que, a diferencia de la definición de continuidad, el número δ en la Definición 2 depende de ε pero no de $x \in E$.

P 1. (propiedades de las sucesiones de Cauchy) Sea (E, d) un espacio métrico. Entonces:

- i) Toda sucesión convergente en (E, d) es de Cauchy;
- ii) Si la sucesión de Cauchy (x_n) tiene una subsucesión convergente en (E, d) ($x_{\varphi(n)} \rightarrow x$), entonces la propia sucesión (x_n) es convergente y además $(x_n) \rightarrow x$.

Demostración. Se deja al lector. $\square$

Ejemplo 1. (sucesión de Cauchy no convergente) La proposición anterior nos sugiere la existencia de sucesiones de Cauchy en un espacio métrico que no son convergentes. En efecto, consideremos el espacio $C[0, 1]$ provisto de la métrica

$$d_1(f, g) = \int_0^1 |f(x) - g(x)| dx, \quad (10)$$

de la convergencia en media.

Para cada $n \geq 2$, definamos

$$f_n(x) = \begin{cases} 1, & \text{si } x \in [0, \frac{1}{2} - \frac{1}{n}] \\ -\frac{n}{2}x + \frac{n}{4} + \frac{1}{2} & \text{si } x \in [\frac{1}{2} - \frac{1}{n}, \frac{1}{2} + \frac{1}{n}] \\ 0, & \text{si } x \in [\frac{1}{2} + \frac{1}{n}, 1]. \end{cases}$$

Un simple cálculo nos muestra que

$$d_1(f_n, f_m) = \frac{1}{2} \left(\frac{1}{m} - \frac{1}{n} \right) \xrightarrow{n, m \rightarrow \infty} 0$$

luego la sucesión (f_n) es de Cauchy en $(C[0, 1], d_1)$. Sin embargo, la sucesión (f_n) no es convergente en $(C[0, 1], d_1)$. En efecto, si $f \in C[0, 1]$ fuera el límite de (f_n) ($f_n \xrightarrow{d_1} f$), entonces para las restricciones

$$f_n^{(1)} \text{ y } f_n^{(2)} \text{ de } f_n \text{ a } [0, 1/2] \text{ y } [1/2, 1]$$

respectivamente, se tendría que $(f_n^{(1)})$ converge a $f^{(1)}$ en $(C[0, \frac{1}{2}], d_1)$ y $(f_n^{(2)})$ converge a $f^{(2)}$ en $(C[\frac{1}{2}, 1], d_1)$. No es difícil verificar que $(f_n^{(1)})$ converge a la función constante $\langle 1 \rangle$ en $(C[0, \frac{1}{2}], d_1)$, mientras que $(f_n^{(2)})$ converge a la función constante $\langle 0 \rangle$ en $(C[\frac{1}{2}, 1], d_1)$.

Por la unicidad del límite se tendría que:

$$f(x) = \begin{cases} 1, & \text{si } x \in [0, \frac{1}{2}] \\ 0 & \text{si } x \in]\frac{1}{2}, 1] \end{cases}$$

lo cual es una contradicción con la suposición de que f es continua sobre $[0, 1]$. $\square$

D 3. (espacio métrico completo). El espacio métrico (E, d) se llama completo, si toda sucesión de Cauchy en (E, d) es convergente en (E, d) .

La mayoría de los espacios métricos completos utilizados en la práctica poseen generalmente una estructura vectorial.

Ejemplo 2. (espacios de funciones acotadas). Dado un conjunto A y un espacio métrico (E, d) , se dice que la función $f : A \rightarrow E$ está acotada si su imagen $f[A]$ está contenida en alguna bola de E . El conjunto $B(A, E) = \{f : A \rightarrow E \text{ acotada}\}$ provisto de la métrica de la convergencia uniforme

$$d_\infty(f, g) = \sup_{x \in A} d(f(x), g(x)) \quad (11)$$

es completo si (E, d) es completo.

En efecto, si (f_n) es una sucesión de Cauchy en $(B(A, E), d_\infty)$, entonces para cada $x \in A$, la sucesión $(f_n(x))$ será de Cauchy en (E, d) . Por ser (E, d) completo, para cada $x \in A$, existe $\lim f_n(x)$, al cual denotaremos mediante $f(x)$.

De esta forma hemos definido una función $f : A \rightarrow E$. Demostremos que (f_n) converge a f uniformemente. Como (f_n) es de Cauchy en $(B(A, E), d_\infty)$, dado $\varepsilon > 0$ es posible hallar un número natural $n_0 = n_0(\varepsilon)$, tal que si $m, n \geq n_0$ entonces

$$d_\infty(f_m, f_n) < \varepsilon,$$

es decir

$$d(f_m(x), f_n(x)) < \varepsilon \text{ para todo } x \in A \quad (12)$$

Dejando m fijo y haciendo tender n a infinito en (12), se obtiene para $m \geq n_0$, que $d(f_m(x), f(x)) \leq \varepsilon$ para todo $x \in A$. Concluimos que (f_n) converge a f en $(B(A, E), d_\infty)$. $\square$

Consideremos ahora al conjunto A provisto de una métrica y definamos

$$C_a(A, E) = \{f : (A, \rho) \rightarrow (E, d) \text{ continua y acotada}\} \quad (13)$$

y continuemos con la suposición de que (E, d) es completo.

Si (f_n) es una sucesión de Cauchy en $(C_a(A, E), d_\infty)$ también lo será en $(B(A, E), d_\infty)$. Luego, por lo demostrado en la primera parte del ejemplo, se tendrá que (f_n) converge a una función f en $(B(A, E), d_\infty)$. Pero no es difícil probar que *el límite uniforme de una sucesión de funciones continuas es también continua* (¡ demuéstrela!) de lo cual se concluye que $(C_a(A, E), d_\infty)$ es completo si (E, d) lo es. $\square$

De los cursos de Análisis Matemático es conocido que $\mathbb{R}$ es completo (provisto de la métrica usual; módulo) y que toda función real continua definida sobre un intervalo cerrado y acotado $[a, b]$, es acotada. El resultado anterior nos muestra que $(C[a, b], d_\infty)$ es completo.

El ejemplo 2 nos sugiere la proposición siguiente:

P 2. (subespacios completos) Sea A subconjunto del espacio métrico completo (E, d) . El subespacio métrico (A, d_A) es completo, si y sólo si A es cerrado en (E, d) .

Demostración. Se deja al lector. $\square$

Hemos visto en el ejemplo 1 que el espacio métrico $(C[a, b], d_1)$ no es completo. Sin embargo, todo espacio métrico admite un *completamiento*. Este proceso de completamiento para espacios métricos cualesquiera es análogo al proceso mediante el cual se obtienen los números reales a partir de los números racionales.

T 1. (prolongación de funciones uniformemente continuas) Sean A un subespacio del espacio métrico (E, d) y (F, ρ) un espacio métrico completo. Entonces toda aplicación uniformemente continua $f : (A, d_A) \rightarrow (F, \rho)$ tiene una única prolongación $\bar{f} : (A, d_A) \rightarrow (F, \rho)$ que es además uniformemente continua.

Demostración. La primera parte del teorema es una consecuencia

inmediata de la continuidad uniforme de f y del Teorema 1.1.19. La demostración de la continuidad uniforme de $\bar{f}$ se deja al lector. $\square$

T 2. (del completamiento) Todo espacio métrico (E, d) admite un completamiento, es decir existe un par $(h, (F, \rho))$ con las siguientes propiedades:

- i) (F, ρ) es un espacio métrico completo,
- ii) $h : (E, d) \rightarrow (F, \rho)$ cumple que $\rho(h(x), h(y)) = d(x, y)$, $(x, y \in E)$, o sea $h : (E, d) \rightarrow (h[E], \rho_h[E])$ es una isometría,
- iii) $h[E]$ es denso en (F, ρ) .

Además, el completamiento de un espacio métrico es único salvo isometrías.

Demostración. Demostraremos primeramente la unicidad del completamiento. Sean $(h_1, (F_1, \rho_1))$ y $(h_2, (F_2, \rho_2))$ dos completamientos de (E, d) . Evidentemente $h = h_2 \circ h_1^{-1} : h_1[E] \rightarrow h_2[E]$ es una isometría y por tanto un isomorfismo uniforme. Según el Teorema 1, h puede extenderse a un isomorfismo uniforme $\bar{h} : (F_1, \rho_1) \rightarrow (F_2, \rho_2)$ el cual por continuidad es claramente una isometría.

Pasemos ahora a probar la existencia del completamiento. Sea $x_0 \in E$ fijo y definamos la aplicación

$$h : E \rightarrow B(E, \mathbb{R})$$

que asocia a todo $x \in E$ la función g_x definida por:

$$g_x(y) = d(y, x) - d(y, x_0).$$

Es claro que $g_x \in B(E, \mathbb{R})$ pues $\sup_{y \in E} |g_x(y)| = d(x, x_0)$.

Además h es una isometría de E sobre $h[E]$ ya que

$$d_\infty(g_x, g_y) = d(x, y).$$

Por tanto el par $(h, \overline{h[E]})$ es un completamiento de (E, d) . $\square$

Pasemos ahora a enunciar algunos teoremas importantes en los espacios métricos completos. De ellos el resultado fundamental lo constituye el teorema de categoría de Baire cuyas aplicaciones se ven fundamentalmente en el estudio de los operadores en espacios de Banach y en la demostración de teoremas de existencia. Primeramente necesitamos la caracterización siguiente de los espacios métricos completos.

T 3. (de Cantor) Un espacio métrico (E, d) es completo si y sólo si toda sucesión decreciente (F_n) de conjuntos cerrados no vacíos con diámetro $(\delta(F_n))$ que tiende a cero tiene intersección no vacía, donde

$$\delta(F_n) = \sup_{x, y \in F_n} d(x, y) \quad (14)$$

Más exactamente, $\bigcap_{n \in \mathbb{N}} F_n$ se reduce a un punto.

Demostración. Véase el libro Análisis matemático en espacios normados de A. Fraguera. $\square$

T 4. (de categoría de Baire) En un espacio métrico completo la intersección de cualquier familia numerable de conjuntos abiertos densos es densa.

Demostración. Sea (E, d) un espacio métrico completo y (D_n) una sucesión de conjuntos densos y abiertos en (E, d) .

Probaremos que $U \cap (\bigcap_{n \in \mathbb{N}} D_n) \neq \emptyset$ para todo abierto no vacío U en (E, d) . En efecto, por ser $U \cap D_1 \neq \emptyset$ existe una bola abierta B_1 tal

que $\overline{B}_1 \subset U \cap D_1$ y $\delta(\overline{B}_1) \leq 1$. Procediendo por inducción se puede construir una sucesión (B_n) de bolas abiertas tales que:

$$\overline{B}_1 \subset B_{n-1} \cap D_n \text{ y } \delta(\overline{B}_n) \leq \frac{1}{n}$$

entonces:

$$\bigcap_{n \in \mathbb{N}} \overline{B}_n \subset (U \cap \bigcap_{n \in \mathbb{N}} D_n))$$

pero las $\overline{B}_n$ satisfacen las condiciones del teorema de Cantor y por tanto

$$\bigcap_{n \in \mathbb{N}} \overline{B}_n \neq \emptyset$$

con lo cual queda probado el teorema. $\square$

D 4. (conjuntos de primera y segunda categoría) Sea (E, d) un espacio métrico. El subconjunto A de E se llama *nunca denso* si su clausura tiene interior vacío $\overline{A} = \emptyset$). Un conjunto de *primera categoría* en (E, d) es aquel que puede representarse mediante la unión de una familia numerable de conjuntos nunca densos. Todo subconjunto de E que no sea de primera categoría se llama de *segunda categoría*.

El Teorema 4 admite una formulación equivalente.

C 1. (Todo espacio métrico completo es de segunda categoría.)

Demostración. Se deja al lector. $\square$

Finalmente, enunciamos un resultado, de gran importancia en el estudio de la existencia y los métodos numéricos de solución de ecuaciones en espacios métricos completos.

Una aplicación f del espacio métrico (E, d) en sí mismo se dice *contractante de razón k* , si $k < 1$ y

$$d(f(x), f(y)) < kd(x, y) \quad \forall x, y \in E.$$

Claramente una tal aplicación es uniformemente continua.

Dada una aplicación $g : (E, d) \rightarrow (E, d)$, a cada punto $x \in E$ tal que $g(x) = x$ se le llama *punto fijo* de g .

T 5. (del punto fijo) Cada aplicación contractante de razón k definida en un espacio métrico completo (E, d) , tiene un punto fijo único x_0 , caracterizado por la propiedad siguiente:

Para cada punto $x \in E$ la sucesión de valores iterados

$$f^n(x) = \underbrace{f \circ f \circ \cdots \circ f}_{n\text{-veces}}(x)$$

converge al punto x_0 con el error de aproximación:

$$d(f^n(x), x_0) \leq \frac{k^n}{1-k} d(x_0, f(x_0)).$$

Demostración. Véase el libro Análisis matemático en espacios métricos de A. Fraguera. $\square$

Los subconjuntos cerrados y acotados C de $\mathbb{R}^N$ tienen dos propiedades fundamentales en las cuales se basan algunos de los teoremas más importantes del Análisis:

Propiedad I: Todo cubrimiento abierto $(U_\alpha)_{\alpha \in I}$ de C contiene un subcubrimiento finito $(U_{\alpha_1}, U_{\alpha_2}, \dots, U_{\alpha_2})$ de C .

Propiedad II: Toda sucesión (x_n) de puntos de C tiene un punto adherente en C (es decir, contiene una subsucesión $(x_{\varphi(n)})$ convergente a un punto de C).

Los pioneros del Análisis funcional como Ascoli y Volterra reconocieron la gran importancia que tienen estas propiedades para la investigación de conjuntos de funciones. En efecto, el descubrimiento de que toda una serie de teoremas clásicos sobre subconjuntos de $\mathbf{R}^N$ pueden trasladarse a conjuntos de funciones que verifican estas propiedades fué uno de los motivos centrales para el desarrollo del Análisis funcional.

Teniendo en cuenta la gran utilidad de estas propiedades I y II en el Análisis matemático y el Análisis funcional es natural que nos dediquemos a estudiar la clase de todos los espacios métricos que verifican estas mismas condiciones. Resulta que para los espacios métricos las propiedades I y II son equivalentes.

T 6. (criterios de compacidad) En cualquier espacio métrico (E, d) son equivalentes las propiedades siguientes:

- i) Todo cubrimiento abierto de E contiene un subcubrimiento finito,
- ii) Toda familia de cerrados de E con intersección vacía contiene una subfamilia finita, también de intersección vacía,
- iii) Toda familia de cerrados de E en la que cualquier subfamilia finita tiene intersección no vacía, es a su vez de intersección no vacía.
- iv) Toda sucesión en E posee una subsucesión convergente.

(Notemos que los tres primeros enunciados que hablan de familias de subconjuntos constituyen un grupo de propiedades correspondientes a I).

Demostración. i) $\Leftrightarrow$ ii), ya que i) y ii) son enunciados duales. En efecto, $(U_\alpha)_{\alpha \in I}$ es un cubrimiento abierto, si y sólo si $(E \setminus U_\alpha)_{\alpha \in I}$ es una familia de cerrados de intersección vacía.

ii) $\Leftrightarrow$ iii) ya que iii) no es más que otra versión de ii).

Queda por probar la equivalencia entre el grupo de propiedades i), ii), iii) y iv).

iii) $\Rightarrow$ iv) Sea (x_n) una sucesión en (E, d) y consideremos para cada n la sección final

$$S_n = \{x_m : m \geq n\}.$$

Es claro que $(\overline{S_n})_{n \in \mathbb{N}}$ es una familia decreciente de cerrados y por lo tanto cualquier subfamilia finita tiene intersección no vacía. Luego por iii) se tendrá que existe

$$x \in \bigcap_{n \in \mathbb{N}} \overline{S_n}.$$

Para entonces el punto x es un punto adherente de la sucesión (x_n) y, por lo tanto, existe una subsucesión $(x_{\varphi(n)})$ de la sucesión (x_n) tal que $(x_{\varphi(n)}) \rightarrow x$.

Por la propiedad iv) puede demostrarse que:

iv-a) Para todo cubrimiento abierto $U = \{U_\alpha\}_{\alpha \in I}$, de E existe un número $\delta > 0$ (llamado número de Lebesgue) tal que el cubrimiento $B = \{B(x, \delta) : x \in E\}$ formado por todas las bolas de radio δ es *más fino* que U (se dice que el cubrimiento B es *más fino* que U si para todo $V \in B$ existe $W \in U$ tal que $V \subset W$)

iv-b) Para cada $\varepsilon > 0$ existe un conjunto finito $\{x_1, \dots, x_n\}$ tal que las bolas $\{B(x_i, \varepsilon) : i = 1, \dots, n\}$ constituyen un cubrimiento de E . (El conjunto finito $x_1, \dots, x_n$ se llama ε -red)

(Demuestre los enunciados iv-a) y iv-b)).

Demostremos ahora la implicación iv) $\Rightarrow$ i) del teorema. Sea $U = (U_\alpha)_{\alpha \in I}$ un cubrimiento abierto de E . Según iv-a) existe un número de Lebesgue δ para U_0 . Según iv-b), existe un conjunto finito $\{x_1, \dots, x_n\}$ tal que $E = \bigcup_{i=1}^n B(x_i, \delta)$.

Para $i = 1, \dots, n$ existen abiertos $U_{\alpha_i} \in \mathcal{U}$ que contienen a las respectivas bolas $B(x_i, \delta)$. Concluimos que $\{U_{\alpha_1}, \dots, U_{\alpha_n}\}$ es un subcobrimiento finito de E , con lo cual se prueba i). $\square$

D 5. (espacios y subconjuntos compactos y relativamente compactos). El espacio métrico (E, d) se llama compacto si cumple alguna de las propiedades equivalentes i) – iv) del Teorema 6. El subconjunto C de (E, d) es compacto si lo es el subespacio métrico (C, d_C) . El subconjunto C de (E, d) se llama relativamente compacto si su clausura $\overline{C}$ es compacto.

Notemos que si d y ρ son métricas topológicamente equivalentes sobre el conjunto E entonces (E, d) es compacto si y sólo si (E, ρ) es compacto.

P 3. (propiedades de los espacios y subconjuntos compactos)

- i) En un espacio métrico compacto un subconjunto es compacto si y sólo si es cerrado.
- ii) Todo subconjunto compacto de un espacio métrico es cerrado y acotado.

Demostración. Se deja al lector. $\square$

Ejemplo 3. (caracterización de los subconjuntos compactos de $\mathbb{R}$)

No es difícil verificar que todo intervalo cerrado y acotado $[a, b]$ de $\mathbb{R}$ satisface la propiedad iv) del Teorema 6 luego es compacto. Por la propiedad i) de la Proposición 3 se tiene que todo subconjunto cerrado y acotado C en $\mathbb{R}$ es compacto, puesto que es un subconjunto cerrado del intervalo compacto $[\inf C, \sup C]$. Luego por P 3 ii) concluimos que *los compactos en $\mathbb{R}$ son los subconjuntos cerrados y acotados*. Este importante resultado puede ser generalizado para los espacios $\mathbb{R}^n$. En efecto,

si $C \subset \mathbb{R}^n$ es cerrado y acotado, entonces para un cierto par de números reales a, b se tiene que C es cerrado en el cubo n -dimensional $[a, b]^n$. Utilizando el criterio iv) del Teorema 6 y el hecho de que la convergencia en $\mathbb{R}^n$ es la convergencia por coordenadas, concluimos por P 3 i) que C es compacto en $\mathbb{R}^n$. Luego hemos probado que *un subconjunto de $\mathbb{R}^n$ es compacto si y sólo si es cerrado y acotado*.

Ejemplo 4. (conjunto no compacto en $(C[0, 1], d_\infty)$). La caracterización obtenida en el ejemplo anterior para los subconjuntos compactos de $\mathbb{R}^n$ no puede ser trasladada a espacios de dimensión infinita. Consideremos la bola cerrada B de centro cero y radio unidad en el espacio métrico $(C[0, 1], d_\infty)$.

El conjunto B no es compacto ya que la sucesión $f_n(x) = x^n, n = 1, 2, \dots$, no tiene ninguna subsucesión convergente en $(C[0, 1], d_\infty)$. En efecto, para cualquier sucesión n_k de números naturales, la sucesión de funciones $f_{n_k}(x) = x^{n_k}$ converge puntualmente a cero si $x \in [0, 1[$ y a 1 en el punto $x = 1$. Si (f_{n_k}) convergiera uniformemente a una cierta función f también convergería puntualmente a f , luego necesariamente

$$f(x) = \begin{cases} 0 & \text{si } x \in [0, 1[, \\ 1 & \text{si } x = 1, \end{cases}$$

y, por lo tanto, f no es continua. $\square$

El ejemplo anterior nos muestra que la función constante $\langle 0 \rangle$ no tiene ninguna vecindad compacta en $(C[0, 1], d_\infty)$. Puede concluirse que ninguna $f \in C[0, 1]$ tiene una vecindad compacta en $(C[0, 1], d_\infty)$ (¡demuéstrelo!). Este hecho motiva la definición siguiente.

D 6. (espacios y subespacios localmente compactos). El espacio métrico (E, d) se llama localmente compacto si cada punto $x \in E$ tiene una vecindad que es compacto como subconjunto de E . El subconjunto A de E es localmente compacto si lo es el subespacio métrico (A, d_A) .

Veremos, sin demostración, un importante resultado, que caracteriza los subconjuntos relativamente compactos en espacios de funciones continuas provistos de la métrica de la convergencia uniforme. Para ello comencemos por estudiar las propiedades de las funciones continuas definidas sobre conjuntos compactos.

T 7. (imagen de compactos por aplicaciones continuas) Sean $f : (E, d) \rightarrow (F, \rho)$ continua y K un subconjunto compacto de E . Entonces $f[K]$ es compacto en (F, ρ) .

Demostración. Sea $(V_\alpha)_{\alpha \in I}$ un cubrimiento abierto de $f[K]$. Por 1.1.2, la colección $(f^{-1}[V_\alpha])_{\alpha \in I}$ constituye un cubrimiento abierto de K . Por el criterio i) de T 6, se puede extraer un subcubrimiento finito $(f^{-1}[V_{\alpha_k}])_{k=1}^n$ de K . Entonces $(V_{\alpha_k})_{k=1}^n$ es un subcubrimiento finito de $f[K]$, luego $f[K]$ es compacto. $\square$

Del teorema anterior se concluye el corolario siguiente:

C 2. (principio del máximo y el mínimo para funciones reales continuas) Toda función real continua definida sobre un espacio métrico compacto alcanza sus valores máximo y mínimo. Es decir si (E, d) es compacto y $f : (E, d) \rightarrow \mathbb{R}$ es continua, entonces existen $x_0, y_0 \in E$, tales que:

$$f(x_0) = \min_{x \in E} f(x), \quad f(y_0) = \max_{x \in E} f(x).$$

Demostración. Se deja al lector. $\square$

Del Teorema 7 y la Proposición 3 ii) se concluye que si (E, d) es compacto, entonces toda función continua f definida sobre (E, d) y con valores en otro espacio métrico (F, ρ) es acotada (véase ejemplo 2), luego $C_a(E, F) = C(E, F)$.

Antes de enunciar el teorema sobre la caracterización de los subconjuntos relativamente compactos en $(C(E, F), d_\infty)$ con (E, d) compacto y (F, ρ) arbitrario, veamos las definiciones siguientes:

D 7. (conjuntos equicontinuos y equiacotados en $\mathcal{F}(E, F)$) Sean (E, d) y (F, ρ) espacios métricos, $\mathcal{F}(E, F)$ el conjunto de todas las funciones definidas sobre E con valores en F y H un subconjunto de $\mathcal{F}(E, F)$.

- i) El conjunto H se llama *equicontinuo* en $x \in E$ si para cada $\varepsilon > 0$ puede hallarse un número real $\delta > 0$ tal que para todo $f \in H$ se cumpla

$$d(x, y) < \delta \Rightarrow \rho(f(x), f(y)) < \varepsilon \quad (15)$$

El conjunto H se llama *equicontinuo* si es equicontinuo en todo $x \in E$.

- ii) H se llama *equiacotado* en el punto $x \in E$ si

$$H(x) = \{f(x) : f \in H\} \quad (16)$$

es relativamente compacto en (F, ρ) .

H se llama *equiacotado* si es equiacotado en todo punto $x \in E$.

Notemos que los subconjuntos equicontinuos de $\mathcal{F}(E, F)$ están formados por funciones que son uniformemente continuas.

T 8. (de Arzelá–Ascoli sobre la caracterización de los subconjuntos relativamente compactos en $(C(E, F), d_\infty)$) Sean (E, d)

compacto y (F, ρ) métrico. El subconjunto $H \subset C(E, F)$ es relativamente compacto en $(C(E, F), d_\infty)$ si y sólo si es equicontinuo y equicotado.

Demostración. Véase el libro Análisis matemático en espacios métricos de A. Fraguera. $\square$

El resultado siguiente nos muestra la relación existente entre continuidad uniforme para funciones definidas sobre un espacio métrico compacto.

T 9. (relación entre continuidad y continuidad uniforme) Sean (E, d) compacto y (F, ρ) métrico. Entonces toda función continua $f : (E, d) \rightarrow (F, \rho)$ es uniformemente continua.

Demostración. Supongamos que f no es uniformemente continua. Entonces, para cierto ε positivo y cualquier número natural n existirán elementos x_n, y_n en E , tales que

$$d(x_n, y_n) < \frac{1}{2} \quad \text{mientras que} \quad \rho(f(x_n), f(y_n)) \geq \varepsilon.$$

De la sucesión (x_n) se puede extraer, debido a la compacidad de E , una subsucesión $(x_{\rho(n)})$ convergente a un punto $x \in E$. En este caso la subsucesión $(x_{\rho(n)})$ también converge al punto x ; pero para cada n debe cumplirse al menos una de las desigualdades

$$\rho(f(x), f(x_{\rho(n)})) \geq \varepsilon/2, \quad \rho(f(x), f(y_{\rho(n)})) \geq \varepsilon/2$$

lo cual contradice la continuidad de f en el punto $x \in E$. $\square$

El teorema anterior nos sugiere la existencia de una relación entre las propiedades topológicas y uniformes de un espacio métrico compacto. En efecto, en un espacio métrico compacto toda sucesión de Cauchy tiene una subsucesión convergente y por P 1 ii) concluimos que

la propia sucesión es convergente. Luego, *todo espacio métrico compacto es completo*. Sin embargo, la afirmación inversa no se cumple como nos muestra el ejemplo 4.

Para estudiar más exactamente la relación entre compacidad y completitud se necesita la definición siguiente:

D 8. (espacios y subconjuntos precompactos o totalmente acotados) El espacio métrico (E, d) se llama *totalmente acotado* o *precompacto* si para todo $\varepsilon > 0$ hay una ε -red en E . Un subconjunto finito $\{x_1, \dots, x_n\}$ se llama ε -red en E si las bolas $B(x_i, \varepsilon); i = 1, \dots, n$ constituyen un cubrimiento de E . El subconjunto A de E se llama *totalmente acotado* si lo es el subespacio métrico (A, d_A) .

T 10. (relación entre compacidad y completitud) El espacio métrico (E, d) es compacto si y sólo si es completo y precompacto.

Demostración. Ya hemos visto que todo espacio métrico compacto es completo. Si (E, d) es compacto basta tomar el cubrimiento abierto $\{B(x, \varepsilon) : x \in E\}$ y utilizar el criterio i) del Teorema 6 para probar la precompacidad de (E, d) .

Demostremos ahora que todo espacio métrico completo y precompacto es también compacto. Para ello es suficiente demostrar que toda sucesión (x_n) de puntos de E tiene al menos un punto adherente. Sean F_1 una 1-red en (E, d) . Las bolas cerradas de radio 1 centradas en los puntos de F_1 constituyen un cubrimiento finito de E . Luego, al menos una de estas bolas, llamémosla B_1 , contiene una subsucesión infinita $x_1^{(1)}, \dots, x_n^{(1)}, \dots$ de la sucesión (x_n) . Escojamos ahora en $\overline{B_1}$ una $\frac{1}{2}$ -red F_2 y construyamos alrededor de todo punto de esta red una bola cerrada de radio $\frac{1}{2}$. Al menos una de estas bolas, llamémoslo B_2 , contiene una subsucesión infinita $x_1^{(2)}, \dots, x_n^{(2)}, \dots$ de la sucesión $(x_n^{(1)})$. Así sucesivamente, por inducción podemos encontrar para cada $k \in \mathbb{N}$ una bola

cerrada de radio $(\frac{1}{2})^{k-1}$ que contiene una subsucesión infinita $(x_n^{(k)})$ de la sucesión $(x_n^{(k-1)})$. Es fácil ver que la familia $\overline{B_k}$ es una sucesión decreciente de conjuntos cerrados con diámetro que tiende a cero y, por el teorema de Cantor (T 3) su intersección $\bigcap_{k \in \mathbb{N}} B_k$ consta de un solo punto x_0 . Este punto es un punto de adherencia de la sucesión inicial (x_n) , ya que toda vecindad suya contiene a alguna bola B_k y, por consiguiente, una subsucesión infinita $(x_n^{(k)})$ de la sucesión (x_n) . $\square$

No es difícil concluir del teorema anterior que un *subconjunto* A de un *espacio métrico completo* es *relativamente compacto* si y sólo si es *precompacto*. (¡Demuéstrelo!)

Como hemos podido apreciar, el trabajo sobre los espacios métricos compactos es muy cómodo debido al gran número de propiedades a las cuales puede recurrirse. En el Teorema 2 del completamiento vimos como todo espacio métrico puede ser *sumergido* mediante una isometría en un subespacio denso de un espacio métrico. Sin embargo, no es posible sumergir un espacio métrico cualquiera mediante un homeomorfismo en un subespacio denso de un espacio métrico compacto.

D 9. (compactificación) Sean (E, d) y (F, ρ) espacios métricos, (F, ρ) compacto y $f : (E, d) \rightarrow (F, \rho)$. El par $(f, (F, \rho))$ se llama una *compactificación* de (E, d) si se cumple:

- i) $f : (E, d) \rightarrow (f[E], \rho_{f[E]})$ es un homeomorfismo, es decir f es una inmersión de (E, d) en el espacio métrico (F, ρ) ,
- ii) $f[E]$ es denso en (F, ρ) .

Ejemplo 5. (Proyección estereográfica) Denotemos por E^1 la esfera de radio 1 y centro en $(0, 1)$ de $\mathbb{R}^2$ y sea S_1 la proyección estereográfica de dicha esfera sobre $\mathbb{R}$. No es difícil verificar que el par (S_1^{-1}, E^1) es una

compactificación de $\mathbb{R}$. Notemos que E^1 se obtiene *encogiendo* a $\mathbb{R}$ y *uniendo* sus extremos. $\square$

Tenemos el teorema siguiente sobre compactificación:

T 11. (de Urysohn) Todo espacio métrico que tenga un subconjunto denso numerable tiene una inmersión en el espacio producto $I^{\mathbb{N}}$ donde $I = [0, 1]$.

Demostración. Véase el libro Análisis matemático en espacios métricos de A. Fraguera. $\square$

Aclaremos que $I^{\mathbb{N}}$ se considera provisto de la *métrica producto* con respecto a la métrica usual sobre $[0, 1]$:

$$\begin{aligned} x &= (x_n) \in I^{\mathbb{N}}, & y &= (y_n) \in I^{\mathbb{N}}, \\ 0 &\leq x_n, & y_n &\leq 1, \\ D(x, y) &= \sum_{n=0}^{\infty} \frac{1}{2^n} \frac{|x_n - y_n|}{1 + |x_n - y_n|}. \end{aligned}$$

No es difícil verificar que *el espacio métrico $(I^{\mathbb{N}}, D)$ es compacto* (¡ demuéstrela!). Luego, si $f : (E, d) \rightarrow (I^{\mathbb{N}}, D)$ es una inmersión, el par $(f, \overline{f[E]})$ es una compactificación de (E, d) .

Pasemos finalmente al estudio de la *propiedad de conexión* en los espacios métricos. Precisamente es la propiedad de conexión de los intervalos en $\mathbb{R}$ la que permite demostrar un teorema tan importante para el Análisis clásico como en el *Teorema de Rolle* o *Teorema del valor medio*. Algunos espacios métricos pueden *dividirse de manera natural* mediante conjuntos abiertos disjuntos de manera que el análisis de la estructura topológica del espacio se reduzca al análisis de cada una de sus *componentes*. Los espacios métricos que no admiten una tal división

se llaman *conexos*.

D 10. (espacios y subconjuntos conexos) El espacio métrico (E, d) se llama conexo si no existe una partición de E formada por dos abiertos no vacíos. El subconjunto A de (E, d) se llama conexo si el subespacio (A, d_A) es conexo.

P 4. (caracterizaciones de los espacios métricos conexos) El espacio métrico (E, d) es conexo si y sólo si cumple una de las condiciones equivalentes siguientes:

- i) No existe una partición de E formada por subconjuntos no vacíos A y B de E tales que $\overline{A} \cap B = \emptyset = A \cap \overline{B}$,
- ii) No existe una partición $(\mathcal{U}_\alpha)_{\alpha \in I}$ de E formada por más de un abierto no vacío,
- iii) No existe una partición de E formada por un número finito ≥ 2 de cerrados no vacíos,
- iv) Los únicos subconjuntos que son a la vez abiertos y cerrados en E son el propio E y el $\emptyset$,
- v) No existe una sobreyección continua de (E, d) sobre un *espacio discreto* (F, ρ) con dos elementos (el espacio (F, ρ) es discreto si la métrica cumple: $\rho(x, y) = 0$ para $x = y$, $\rho(x, y) = 1$ para $x \neq y$),
- vi) Toda aplicación continua de (E, d) en un espacio discreto es constante.

Demostración. Se deja al lector. $\square$

La caracterización iv) resulta muy útil para la demostración de propiedades sobre los espacios métricos conexos. En efecto, si se desea probar que todos los puntos de un espacio métrico conexo (E, d) satisfacen

una cierta propiedad P , hasta demostrar que el conjunto

$$\Pi = \{x \in E : x \text{ satisface } P\}$$

es no vacío y además abierto y cerrado en (E, d) . Por la caracterización iv) se concluye entonces que

$$\Pi = E.$$

P 5. (propiedades de los conexos)

- i) Sea A un subconjunto conexo del espacio métrico (E, d) . Entonces todo subconjunto B de E , tal que $A \subset B \subset \overline{A}$, es conexo. En particular $\overline{A}$ es conexo.
- ii) Si $(C_\alpha)_{\alpha \in I}$ es una familia de subconjuntos conexos de (E, d) tal que para algún $\beta \in I$ fijo se cumple:

$$\overline{C_\alpha} \cap C_\beta \neq \emptyset \quad \text{ó} \quad C_\alpha \cap \overline{C_\beta} \neq \emptyset, (\alpha \in I)$$

entonces $C = \bigcap_{\alpha \in I} C_\alpha$ es conexo.

Demostración. i) Utilicemos el criterio v) de la proposición anterior. Sea f una aplicación continua de B en el espacio discreto $\{1, 2\}$. Sabemos que $f[A]$ es unitario, luego $f[B] \subset f[\overline{A}] \subset \overline{f[A]}$ (T 1.1.5-iii) también es unitario. Concluimos que f no es sobreyectiva y, por lo tanto, B es conexo.

ii) Se deja de ejercicio al lector. $\square$

Notemos que *la intersección de dos conjuntos conexos no siempre es conexo*. Veamos que *la conexión es una propiedad topológica*. Más

exactamente, se tiene:

T 12. (continuidad y conexión) Sea $f : (E, d) \rightarrow (F, \rho)$ continua. Si (E, d) es conexo, entonces $f[E]$ es conexo en (F, ρ) .

Demostración. Sea g una aplicación continua del subespacio $f[E]$ en el espacio discreto $\{1, 2\}$. Entonces $g \circ f$ es una aplicación continua de E en $\{1, 2\}$, luego g no puede ser sobreyectiva. $\square$

Este teorema nos permite obtener una generalización del conocido teorema del valor medio en el Análisis matemático. Pero antes necesitaremos el resultado siguiente:

T 13. (caracterización de los subconjuntos conexos de $\mathbb{R}$) Un subconjunto C de $\mathbb{R}$ es conexo si y sólo si es un intervalo.

Demostración. Recordemos que los intervalos C de $\mathbb{R}$ están caracterizados por la propiedad:

$$x, y, z \in \mathbb{R}; x, y \in C; x < z < y \Rightarrow z \in C.$$

Supongamos que el subconjunto C de $\mathbb{R}$ no es un intervalo. Entonces existen $x, y \in C$ y $z \in \mathbb{R} \setminus C$ tales que $x < z < y$. Los subconjuntos no vacíos

$$A = \{c \in C : c < z\}, \quad B = \{c \in C : c > z\}$$

constituyen una partición de C , luego C no es conexo.

Recíprocamente, consideremos un intervalo C y demostremos que es un subconjunto conexo de $\mathbb{R}$. Para ello supongamos por el absurdo que existen dos subconjuntos abiertos no vacíos A y B del espacio C tales que $A \cup B = C$ y $A \cap B = \emptyset$. Intercambiando A y B si es necesario, existen $a \in A$ y $b \in B$ tales que $a < b$. Sea $z = \sup\{A \cap [a, b]\}$; $z \in [a, b] \subset C$ es punto adherente de A en C . Puesto que A es cerrado en C (por ser el

complemento del abierto B) tenemos que $z \in A$ y, por lo tanto, $z < b$. Por otra parte, A es también abierto en C . Por ello existe $\varepsilon > 0$ tal que $[z, z + \varepsilon] \subset A$ y $z + \varepsilon < b$. Pero esto contradice la definición de z . Concluimos que C es conexo. $\square$

T 14. (del valor medio generalizado) Sean (E, d) un espacio métrico conexo y f una función real continua definida sobre E . Si $x, y \in E$ y $\alpha \in \mathbb{R}$ son tales que

$$f(x) \leq \alpha \leq f(y),$$

entonces existe $z \in E$ con $f(z) = \alpha$.

Demostración. Por T 12, $f[E]$ es conexo en $\mathbb{R}$, luego un intervalo (T 13), de donde se concluye la tesis del teorema. $\square$

§ 1.3 Normas y espacios normados

En este epígrafe nos dedicaremos al estudio de las propiedades métricas y topológicas de algunos espacios métricos cuya estructura métrica es compatible, en un sentido que definiremos más adelante, con alguna estructura vectorial definida sobre dicho espacio.

Comenzaremos por introducir algunas nociones algebraicas. Consideremos un conjunto arbitrario E sobre el que hay definidas dos operaciones: *suma* (ley interna) y *producto por escalares* (ley externa). Denotaremos la operación suma por la letra S y la de producto por escalares, por la letra P . Ambas operaciones se definen como aplicaciones

$$\begin{aligned} S : E \times E &\longrightarrow E, \\ P : K \times E &\longrightarrow E, \end{aligned} \tag{17}$$

donde K será eventualmente uno de los conjuntos numéricos $\mathbb{R}$ ó $\mathbb{C}$. En otras palabras, mediante las operaciones de suma y producto por

escalares, hacemos corresponder a cada par de elementos $x, y \in E$ y cada número $\lambda \in K$, los elementos unívocamente determinados $S(x, y)$ y $P(\lambda, x)$ de E . En lo que sigue denotaremos al elemento $S(x, y)$ por $x + y$:

$$S(x, y) = x + y \quad (18)$$

y le llamaremos *suma de los elementos* $x \in y$ de E . De la misma forma denotaremos al elemento $P(\lambda, x)$ por λx :

$$P(\lambda, x) = \lambda x \quad (19)$$

y le llamaremos *producto del elemento* x *por el número* λ .

A las operaciones (18) y (19) les exigiremos determinadas relaciones naturales de compatibilidad. Para toda triada de elementos $x, y, z \in E$ y todo par de números $\lambda, \mu \in K$ se debe cumplir:

(EV1) $x + y = y + x$ (conmutatividad de la suma)

(EV2) $(x + y) + z = x + (y + z)$ (asociatividad de la suma)

(EV3) $\lambda(x + y) = \lambda x + \lambda y$, $(\lambda + \mu)x = \lambda x + \mu x$ (distributividad del producto con respecto a la suma).

(EV4) $\lambda(\mu x) = (\lambda\mu)x$ (asociatividad del producto)

(EV5) existe un elemento $\Theta \in E$ llamado elemento nulo o neutro, tal que $0x = \Theta$ para cualquier $x \in E$, donde 0 es el cero de K .

(EV6) $1x = x$ para cualquier $x \in E$.

De esta forma hemos incluido en las propiedades exigidas a la suma y el producto por escalares en E , un grupo fundamental de leyes algebraicas que son válidas para las correspondientes operaciones en los espacios euclidianos de dimensión finita $\mathbb{R}^n$.

D 1. (espacio vectorial y subespacio vectorial) Un espacio vectorial sobre K es un conjunto E provisto de dos operaciones, una interna

llamada suma (18) y otra externa llamada producto por escalares (19), las cuales satisfacen las propiedades (EV1)–(EV6). El espacio vectorial se llama real si $K = \mathbb{R}$ y complejo si $K = \mathbb{C}$. Los elementos del espacio vectorial E se denominan vectores.

Un subconjunto F del espacio vectorial E es un subespacio vectorial de E si es un espacio vectorial con respecto a las operaciones suma y producto por escalares de E , restringidas a sus elementos.

No resulta difícil verificar que un subconjunto F del espacio vectorial E es un subespacio vectorial de E , si para todos $\lambda, \mu \in K$ y $x, y \in F$ se cumple

$$\lambda x + \mu y \in F. \quad (20)$$

Salvo en casos particulares, donde será indicado, en lo que sigue estudiaremos espacios vectoriales complejos. Luego, cuando hablemos simplemente de espacios vectoriales nos estaremos refiriendo a espacios vectoriales complejos.

Comencemos por enunciar algunas consecuencias de los axiomas (EV1)–(EV6) que pueden ser rápidamente verificadas por el lector:

- a) En cada espacio vectorial existe un único elemento neutro,
- b) El elemento neutro de E es el cero para la suma, es decir

$$\Theta + x = x, \quad \text{para todo } x \in E.$$

- c) Para todo $x \in E$ puede definirse el elemento $-x$ mediante $-x = (-1)x$ y se llama *elemento opuesto* a x . Entonces

$$x + (-x) = \Theta, \quad (21)$$

$$-(-x) = x, \quad (22)$$

Luego para dos elementos $x, y \in E$ se puede definir su *diferencia*

$$x - y = x + (-y). \quad (23)$$

De la definición (23) se deduce que

$$(x - y) + y = x.$$

Esta notación nos permite hablar de una operación de *sustracción* o *resta*, inversa a la suma o adición. Para cualesquiera $x, y \in E$, la ecuación $y + z = x$, de la incógnita z , tiene una única solución $z = x - y$.

- d) La propiedad (EV3) de distributividad del producto con respecto a la suma se cumple también para la sustracción.
- e) $\lambda\Theta = \Theta$ para todo $\lambda \in \mathbb{C}$.
- f) De $\lambda x = \Theta$ se deduce que $\lambda = 0$ ó $x = \Theta$. Luego si $\lambda x = \lambda y$ entonces $\lambda = 0$ ó $x = y$.

De la misma forma, si $\lambda x = \mu x$, entonces

$$x = \Theta \quad \text{ó} \quad \lambda = \mu.$$

Veamos de manera algo informal, algunas definiciones y resultados relacionadas con la estructura vectorial de los espacios vectoriales. Dados los vectores $x_1, \dots, x_n$ del espacio vectorial E y los números complejos $\lambda_1, \dots, \lambda_n$, al vector $\lambda_1 x_1 + \dots + \lambda_n x_n$ se le llama *combinación lineal* de los vectores $x_i : i = 1, \dots, n$ con *coeficientes* $\{\lambda_i : i = 1, \dots, n\}$.

La familia $\mathcal{F}$ de vectores de E se llama *linealmente independiente* si no es posible obtener una combinación lineal nula de elementos de $\mathcal{F}$ con coeficientes diferentes de cero:

$$\begin{aligned} x_1, \dots, x_n \in \mathcal{F}; \quad \lambda_1, \dots, \lambda_n \in \mathbb{C}, \\ \lambda_1 x_1 + \dots + \lambda_n x_n = 0 \Rightarrow \lambda_1 = \dots = \lambda_n = 0. \end{aligned}$$

En caso contrario se dice que la familia $\mathcal{F}$ es *linealmente dependiente*. Se puede demostrar que toda familia $\mathcal{F}$ de vectores de E contiene una *subfamilia maximal, no necesariamente única, de vectores linealmente independientes* y que, además, todas estas subfamilias maximales tienen

la misma potencia. Si en lugar de $\mathcal{F}$ consideramos al propio espacio vectorial E , estas subfamilias maximales son llamadas *bases algebraicas de E* y la potencia de cada una de ellas se llama *dimensión algebraica de E* .

No es difícil probar que *la intersección de cualquier familia de subespacios vectoriales de un espacio vectorial E , es también un subespacio vectorial de E* . Luego, para toda familia $\mathcal{F}$ de vectores de E existe un *menor subespacio vectorial* en E que contiene a $\mathcal{F}$ y que coincide con la intersección de todos los subespacios vectoriales de E que contienen a $\mathcal{F}$. A este menor subespacio vectorial se le llama *subespacio vectorial generado por $\mathcal{F}$* y se denota por $[\mathcal{F}]$. El lector puede comprobar que $[\mathcal{F}]$ coincide con el conjunto de todas las combinaciones lineales de elementos de $\mathcal{F}$. Luego, toda subfamilia maximal de vectores linealmente independientes de $\mathcal{F}$ constituye una base algebraica de $[\mathcal{F}]$.

Se dice que el subconjunto F del espacio vectorial E es *suma* de los subconjuntos F_1 y F_2 y se denota $F = F_1 + F_2$, si todo elemento $x \in F$ puede expresarse de la forma $x = x_1 + x_2$ con $x_i \in F_i; i = 1, 2$. La suma $F = F_1 + F_2$ se llama *directa* y se denota $F = F_1 \oplus F_2$, cuando la expresión de cada $x \in F$, como suma de elementos de F_1 y F_2 , es única.

Si F_1 y F_2 son subespacios vectoriales de E , su suma F es un subespacio vectorial de E . Además, en este caso, la suma es directa cuando, y solamente cuando, $F_1 \cap F_2 = \{\Theta\}$.

En lo que sigue utilizaremos las notaciones λA y $\wedge A$ con $\lambda \in \mathbb{C}, \wedge \subset \mathbb{C}$ y $A \subset E$ para denotar los subconjuntos de E

$$\lambda A = \{\lambda x : x \in A\} \quad (24)$$

$$\wedge A = \{\mu x : \mu \in \wedge, x \in A\}. \quad (25)$$

No es difícil verificar que, para las operaciones usuales de suma de números, vectores y funciones, así como el producto usual por escalares, los espacios euclidianos $\mathbb{R}^n$ y los espacios $C(E, \mathbb{R})$ de funciones reales continuas definidas sobre un espacio métrico (E, d) son espacios vectoriales reales. De la misma forma puede comprobarse que son espacios

vectoriales complejos los espacios euclidianos $\mathbb{C}^n$ y los espacios de funciones complejas continuos $C(E, \mathbb{C})$.

El lector puede fácilmente verificar que para las métricas d_1, d_2 y d definidas en los espacios vectoriales $\mathbb{R}^n, \mathbb{C}^n, C([a, b], \mathbb{R})$ y $C([a, b], \mathbb{C})$, las operaciones suma y producto por escalares (17) son continuas. Sin embargo, un espacio vectorial puede ser provisto de una estructura métrica para la cual las aplicaciones de suma y producto por escalares no sean continuas. Al estudiar estructuras métricas sobre espacios vectoriales resulta natural excluir estos casos singulares de nuestra consideración.

D 2. (métrica topológicamente compatible con la estructura vectorial) Sea E un espacio vectorial provisto de una métrica d . Se dice que la métrica d es *topológicamente compatible* con la estructura vectorial de E si las aplicaciones de suma y producto por escalares (17), son continuas con respecto a d .

En un espacio vectorial métrico (E, d) donde la métrica sea compatible con la estructura vectorial, se cumple que si

$$(x_n) \xrightarrow{d} x, (y_n) \xrightarrow{d} y, \lambda_n \longrightarrow \lambda \text{ (convergencia en } \mathbb{C}),$$

entonces

$$x_n + y_n \xrightarrow{d} x + y, \lambda_n x_n \xrightarrow{d} \lambda x. \quad (26)$$

La estructura topológica de un espacio vectorial E provisto de una métrica topológicamente compatible d queda totalmente determinada si se conoce el comportamiento de dicha métrica en un entorno del elemento nulo Θ . En efecto, de (26) se deduce, que en el espacio métrico (E, d) se cumple:

$$x_n \xrightarrow{d} x \iff x_n - x \xrightarrow{d} \Theta. \quad (27)$$

Luego, para el estudio de la convergencia en (E, d) , no es necesario conocer la distancia entre dos elementos cualesquiera, sino basta con

saber la distancia de cada elemento $x \in E$ al elemento nulo Θ de E . Es decir, es suficiente conocer el comportamiento de la función

$$\begin{aligned} N : E &\longrightarrow \mathbb{R}_+ \\ x &\rightsquigarrow N(x) = d(x, \Theta), \end{aligned} \quad (28)$$

ya que

$$x_n \xrightarrow{d} x \iff N(x_n - x) \longrightarrow 0 \quad (29)$$

(convergencia usual en $\mathbb{R}$)

Más aún, se tiene el resultado siguiente:

P 1. (invarianza de los abiertos y cerrados por traslaciones y homotecias) Sea E un espacio vectorial provisto de una métrica d topológicamente compatible con su estructura algebraica. Entonces las siguientes condiciones son equivalentes:

- i) A es abierto (resp. cerrado) en E ;
- ii) $\forall a \in E, a + A$ es abierto (resp. cerrado) en E ;
- iii) $\forall \lambda \in \mathbb{C}, \lambda \neq 0, \lambda A$ es abierto (resp. cerrado) en E .

Demostración. Se concluye inmediatamente del hecho que para cada $a \in A$ y $\lambda \in \mathbb{C}, \lambda \neq 0$; las aplicaciones:

$$S_a : E \longrightarrow E, \quad x \rightsquigarrow a + x, \quad P_\lambda : E \longrightarrow E, \quad x \rightsquigarrow \lambda x,$$

son homeomorfismos de (E, d) en sí mismo. $\square$

De la Proposición anterior se concluye que todo abierto en (E, d) se puede expresar como la unión de conjuntos de la forma $a + V$, donde $a \in E$ y V es una vecindad abierta del elemento neutro Θ .

Hasta el momento hemos visto como la función N , definida en (28), nos permite obtener información sobre la estructura topológica de todo el espacio métrico (E, d) a partir de su *estructura topológica local*, cuando la métrica d es topológicamente compatible con la estructura vectorial de E . Sin embargo, no es posible describir las *propiedades métricas* de d en todo E a partir de sus propiedades tomadas con respecto al elemento neutro Θ . En efecto, si se conoce la distancia de cualquier elemento $x \in E$ al elemento nulo $\Theta \in E$, esto en general no es suficiente para conocer la distancia de x a otro punto $y \in E$. Para poder obtener esta importante y útil propiedad es necesario imponer condiciones adicionales a la métrica d .

D 3. (distancias métricamente compatibles con la estructura algebraica) Sea E un espacio vectorial provisto de una métrica d . Se dice que la métrica d es *métricamente compatible* con la estructura vectorial de E si para todos $x, y, z \in E$ y $\lambda \in \mathbb{C}$, se cumple:

$$d(x + z, y + z) = d(x, y), \quad (30)$$

$$d(\lambda x, \lambda y) = |\lambda|d(x, y), \quad (31)$$

donde $|\cdot|$ denota el módulo de un número complejo.

Obviamente para toda distancia métricamente compatible con la estructura algebraica de E se tiene que la función N definida en (28) cumple

$$N(x - y) = d(x, y) \quad (32)$$

De (30), (31), (32) y las propiedades de una métrica se obtienen inmediatamente las propiedades siguientes para la correspondiente función N :

$$\begin{aligned} (N1) \quad & N(x) \geq 0, \text{ y } & N(x) = 0 &\iff x = \Theta, \\ (N2) \quad & N(\lambda x) = |\lambda|N(x), \\ (N3) \quad & N(x + y) \leq N(x) + N(y). \end{aligned}$$

Por la analogía entre las propiedades (N1) – (N3) con las propiedades características del módulo para los números complejos, se conviene en utilizar la notación

$$N(x) = \|x\| \quad (33)$$

que utilizaremos a partir de ahora para denotar la función N correspondiente a una distancia métricamente compatible con la estructura algebraica de E .

D 4. (norma, espacio normado y subespacio normado) Sea E un espacio vectorial. Toda función definida sobre E con valores en $\mathbb{R}_+$ y que satisfaga las propiedades (N1) – (N3) se llama una *norma* sobre E . Se dice que el espacio métrico (E, d) es *normado* si la distancia d es métricamente compatible con la estructura algebraica de E (en este caso la métrica d se expresa mediante (32) utilizando la norma N definida en (28): $d(x, y) = \|x - y\|$).

Se llama subespacio normado del espacio normado (E, d) a cualquier subespacio vectorial de E provisto de la métrica inducida por d . En lo sucesivo convendremos en denotar a los espacios normados por un par del tipo $(E, \|\cdot\|)$ y cuando hablemos de sus propiedades métricas o topológicas nos estaremos refiriendo a las propiedades correspondientes a la métrica $d(x, y) = \|x - y\|$.

No es difícil verificar que las aplicaciones

$$\|x\|_1 = \sum_{i=1}^n |x_i|, \quad (34)$$

$$\|x\|_2 = \left(\sum_{i=1}^n |x_i|^2 \right)^{\frac{1}{2}}, \quad (35)$$

$$\|x\|_\infty = \max_{1 \leq i \leq n} |x_i|, \quad (36)$$

son normas en $\mathbb{C}^n$ que corresponden a las métricas d_1, d_2 y d_∞ . Además

$$\|f\|_1 = \int_a^b |f(t)| dt, \quad (37)$$

$$\|f\|_2 = \left(\int_a^b |f(t)|^2 dt \right)^{\frac{1}{2}}, \quad (38)$$

$$\|f\|_\infty = \max_{a \leq t \leq b} |f(t)|, \quad (39)$$

son normas en $C([a, b], \mathbb{C})$ que corresponden a las métricas d_1 , d_2 y d_∞ . Luego, todas las métricas importantes vistas hasta el momento provienen de una norma.

Del hecho que se cumple la desigualdad

$$(N3') \quad |||x|| - ||y||| \leq \|x - y\|$$

(demuéstre esta desigualdad) se deduce que la aplicación

$$(E, \|\cdot\|) \longrightarrow \mathbb{R}_+, \quad x \mapsto \|x\|,$$

es continua.

Para un espacio normado $(E, \|\cdot\|)$ no es difícil verificar que *si V es una vecindad abierta y acotada del elemento neutro Θ (en calidad de V puede tomarse la bola abierta $B(\Theta, 1)$) y λ_n es una sucesión de números reales que tiende a cero, entonces todo abierto en E puede expresarse como la unión de conjuntos del tipo $a + \lambda_n V$ donde $a \in E$ (¡demuéstrelo!).*

Veamos, sin demostración, algunos resultados concernientes a la cerrabilidad de subespacios vectoriales de un espacio normado $(E, \|\cdot\|)$. El lector interesado puede hallar las demostraciones en el libro *Análisis matemático en espacios normados* de A. Fraguola:

- a) Si F es un subespacio vectorial de E , entonces $\overline{F}$ también es un subespacio vectorial de E ,
- b) Todo subespacio vectorial en $(E, \|\cdot\|)$ de dimensión finita, es cerrado.
- c) Si F y M son subespacios vectoriales cerrados de $(E, \|\cdot\|)$ y F es además de dimensión finita, entonces el subespacio vectorial suma $F + M$ también es cerrado.

D 5. (envoltura lineal cerrada de una familia de vectores, familias totales) Sea $\mathcal{F}$ una familia de vectores en el espacio normado $(E, \|\cdot\|)$. Se llama *envoltura lineal cerrada* de $\mathcal{F}$ a la clausura del espacio vectorial $[\mathcal{F}]$ generado por $\mathcal{F}$. La familia $\mathcal{F}$ se llama *total* en E , si su envoltura lineal cerrada coincide con E :

$$[\mathcal{F}] = E.$$

Ejemplo 1. El teorema de aproximación de Weierstrass expuesto en el ejemplo 1.1.1 puede ser expresado de la manera equivalente siguiente:

La familia de funciones $\{1, x, x^2, \dots, x^n, \dots\}$ es total en $(C[a, b], \|\cdot\|_\infty)$. $\square$

Veamos ahora cómo se expresa en términos de las normas, la equivalencia topológica entre sus métricas asociadas. Ya hemos visto que dos métricas d_1 y d_2 definen los mismos abiertos sobre un conjunto E si y sólo si ambos definen la misma convergencia. Si las métricas d_1 y d_2 son generadas por normas entonces se tiene el teorema siguiente:

T 1. (normas topológicamente equivalentes) Dos normas $\|\cdot\|_1$ y $\|\cdot\|_2$ definidas sobre el espacio vectorial E , dan lugar a la misma convergencia, si y sólo si existen constantes estrictamente positivas C_1 y C_2 , tales que

$$C_2\|x\|_1 \leq \|x\|_2 \leq C_1\|x\|_1, \quad (40)$$

para todo $x \in E$.

Demostración. Evidentemente la condición (40) es suficientemente para que ambas normas definan la misma convergencia sobre E . Demostremos que la condición (40) es también necesaria. En efecto, si las

normas $\|\cdot\|_1$ y $\|\cdot\|_2$ definen la misma convergencia en E , entonces la identidad

$$1_E : (E, \|\cdot\|_1) \rightarrow (E, \|\cdot\|_2)$$

es un homeomorfismo. Analicemos cómo se manifiesta la continuidad de 1_E en el elemento nulo $\Theta \in E$, y para ello denotemos por d_1 y d_2 las respectivas métricas asociadas a $\|\cdot\|_1$ y $\|\cdot\|_2$.

Por ser 1_E continua en Θ , para cada número real $\varepsilon > 0$ existe un número real $\delta > 0$, tal que:

$$x \in B_{d_1}(\Theta, \delta) \Rightarrow 1_E(x) \in B_{d_2}(1_E(\Theta), \varepsilon) \quad (41)$$

Pero teniendo en cuenta que

$$\begin{aligned} 1_E(x) &= x, \quad \text{para todo } x \in E, \\ B_d(\Theta, r) &= \{x \in E : \|x\| < r\}, \end{aligned}$$

la expresión (41) puede escribirse en la forma equivalente

$$\|x\|_1 < \delta \Rightarrow \|x\|_2 < \varepsilon \quad (42)$$

Luego, como para cada $x \in E$ el elemento $y = \frac{\delta}{2\|x\|_1}x$, cumple que $\|y\|_1 < \delta$, entonces, por (41)

$$\|y\|_2 < \varepsilon \quad (43)$$

Pero de (43) se tiene que

$$\|x\|_2 < \frac{2\varepsilon}{\delta} \|x\|_1$$

y queda probada una de las desigualdades en (40), poniendo $C_1 = 2\varepsilon/\delta$. La desigualdad en sentido opuesto se obtiene intercambiando los papeles de $\|\cdot\|_1$ y $\|\cdot\|_2$. $\square$

El resultado anterior motiva la definición siguiente:

D 6. (normas equivalentes) Dos normas $\|\cdot\|_1$ y $\|\cdot\|_2$ sobre un

espacio vectorial E son equivalentes si satisfacen las desigualdades (40) para algún par de constantes estrictamente positivas C_1 y C_2 .

Pasemos a ver una importante propiedad de los espacios normados de dimensión finita.

T 2. (equivalencias de las normas en los espacios normados de dimensión finita) Todas las normas definidas sobre un espacio vectorial de dimensión finita son equivalentes.

Demostración. Sea E un espacio vectorial de dimensión n y sea $\|\cdot\|$ una norma sobre E . Basta demostrar que $\|\cdot\|$ es equivalente a la norma $\|\cdot\|_0$ definida de la manera siguiente: si $\{e_1, \dots, e_n\}$ es una suma de E y $x \in E, x = \lambda_1 e_1 + \dots + \lambda_n e_n$ entonces

$$\|x\|_0 = \max_{1 \leq i \leq n} |\lambda_i|. \quad (44)$$

Se verifica fácilmente que la aplicación $f : (\mathbb{C}^n, \|\cdot\|_\infty) \longrightarrow (E, \|\cdot\|_0)$ definida mediante

$$f(\lambda_1, \dots, \lambda_n) = \lambda_1 e_1 + \dots + \lambda_n e_n$$

es un homeomorfismo. Por otra parte, si consideramos sobre E la norma $\|\cdot\|$, entonces para

$$\lambda = (\lambda_1, \dots, \lambda_n), \mu = (\mu_1, \dots, \mu_n),$$

se tiene que

$$\|f(\lambda) - f(\mu)\|_0 \leq \sum_{i=1}^n \|e_i\| \|\lambda - \mu\|_\infty$$

de donde se concluye la continuidad de f . La continuidad de $f^{-1} : (E, \|\cdot\|) \longrightarrow (\mathbb{C}^n, \|\cdot\|_\infty)$ se demuestra sin dificultad. Luego f es también un homeomorfismo entre $(E, \|\cdot\|)$ y $(\mathbb{C}^n, \|\cdot\|_\infty)$. Pero entonces $1_E :$

$(E, \|\cdot\|) \rightarrow (E, \|\cdot\|_0)$ es un homeomorfismo y por lo tanto, $\|\cdot\|$ y $\|\cdot\|_0$ son normas equivalentes. $\square$

El ejemplo siguiente nos muestra que este no es el caso cuando se trata de espacios vectoriales de dimensión infinita.

Ejemplo 2. En el espacio vectorial $C[0, 1]$ las normas $\|\cdot\|_1$ y $\|\cdot\|_\infty$, definidas en (37) y (39), no son equivalentes ya que no definen la misma convergencia. En efecto, la sucesión (x^n) converge a $\langle 0 \rangle$ con respecto a $\|\cdot\|_1$ pero no es convergente con respecto a la norma $\|\cdot\|_\infty$. $\square$

El Teorema 2 nos muestra la relación intrínseca entre las propiedades topológicas y algebraicas de un espacio normado. Esta relación se observa también si analizamos la compacidad en los espacios normados. Obviamente ningún espacio normado puede ser compacto ya que no es acotado.

L 1. (de Riesz) Sean $(E, \|\cdot\|)$ y F un subespacio vectorial cerrado propio de E . Entonces para cada número real $0 < \alpha < 1$, existe un vector x_α tal que $\|x_\alpha\| = 1$ y $\|x - x_\alpha\| \geq \alpha$ ($x \in F$).

Demostración. Seleccionemos $x_1 \in E \setminus F$ y sea

$$\rho = \inf_{x \in F} \|x - x_1\|.$$

Como F es cerrado, tenemos que $\rho > 0$ y además existe $x_0 \in F$ tal que $\|x_0 - x_1\| \leq \rho/\alpha$. Pongamos $x_\alpha = h(x_1 - x_0)$ donde $h = \|x_0 - x_1\|^{-1}$. Es claro que $\|x_\alpha\| = 1$. Si $x \in F$, entonces $h^{-1}x + x_0 \in F$ y, por lo tanto

$$\begin{aligned} \|x - x_\alpha\| &= \|x - hx_1 + hx_0\| = h\|(h^{-1}x + x_0) - x_1\| \\ &\geq h\rho = \|x_1 - x_0\|^{-1}\rho \geq \alpha \end{aligned}$$

con lo cual se completa la demostración. $\square$

T 3. En un espacio vectorial normado $(E, \|\cdot\|)$ son equivalentes:

- i) $(E, \|\cdot\|)$ es de dimensión finita,
- ii) La bola unidad $B_1 = \{x : \|x\| < 1\}$ es compacta,
- iii) $(E, \|\cdot\|)$ es localmente compacto.

Demostración. Las implicaciones $i) \Rightarrow ii) \Rightarrow iii) \Rightarrow ii)$ son evidentes y se dejan de ejercicio al lector.

$ii) \Rightarrow i)$ Supongamos por el absurdo que E no es de dimensión finita. Seleccionemos $x_1 \in S_1 = \{x \in E : \|x\| = 1\}$ y sea F_1 el subespacio generado por x_1 . Obviamente F_1 es un subespacio cerrado propio de E . Luego, por el lema de Riesz (L 1) existe $x_2 \in S_1$ tal que $\|x_2 - x_1\| \geq 1/2$. Sea F_2 el subespacio cerrado propio de E generado por x_1 y x_2 , entonces existe (L 1) $x_3 \in S_1$ tal que $\|x_3 - x\| \geq 1/2$ si $x \in F_2$. Procediente por inducción encontramos una sucesión infinita de elementos de S_1 tal que $\|x_n - x_m\| \geq 1/2$ si $n \neq m$. Evidentemente esta sucesión no puede contener ninguna subsucesión convergente, lo cual contradice la compacidad de S_1 . Luego E debe ser de dimensión finita. $\square$

Para los espacios normados de dimensión finita se tiene también el teorema siguiente:

T 4. (completitud en los espacios normados de dimensión finita) Todo espacio normado de dimensión finita es completo.

Demostración. La norma $\|\cdot\|$ es equivalente a $\|\cdot\|_0$ definida en (44), luego basta probar que $(E, \|\cdot\|_0)$ es completo. Pero esto es evidente

ya que la aplicación $f : (\mathbb{C}^n, \|\cdot\|_\infty) \rightarrow (E, \|\cdot\|_0)$ definida en el Teorema 2 es una isometría, es decir $\|\lambda\|_\infty = \|f(\lambda)\|_0$, y $(\mathbb{C}^n, \|\cdot\|_\infty)$ es completo. $\square$

La importancia de la noción de completitud en los espacios métricos justifica la definición siguiente:

D 7. (espacio de Banach) Se llama *espacio de Banach* a todo espacio normado completo.

Por el Teorema 1.2.2, todo espacio normado, considerado como espacio métrico, admite un completamiento. Puede probarse que *el completamiento de un espacio normado es un espacio de Banach*. En un espacio normado, todo par de puntos x, y pueden ser unidos por un segmento $\{\lambda x + (1 - \lambda)y : \lambda \in [0, 1]\}$. En particular todo punto puede ser unido al elemento nulo por un segmento, que es un subconjunto conexo. Luego, *todo espacio normado es conexo*.

La estructura vectorial de los espacios normados permite introducir el importante concepto de *serie convergente* ya conocido en los espacios numéricos.

Dada una sucesión (x_k) de elementos del espacio normado $(E, \|\cdot\|)$, se puede construir la sucesión de *sumas parciales*

$$S_n = x_1 + \cdots + x_n = \sum_{k=1}^n x_k. \quad (45)$$

Se dice que la sucesión (x_k) determina una *serie convergente* en $(E, \|\cdot\|)$ si la sucesión de sumas parciales S_n converge. Al límite S de la sucesión (S_n) , cuando éste existe, se le llama *suma de la serie* correspondiente a la sucesión (x_k) , y se denota por

$$S = \lim(S_n) = \sum_{k=1}^{\infty} x_k. \quad (46)$$

Diremos también que la serie $\sum_{k=1}^{\infty} x_k$ es *convergente* si el límite de la sucesión (S_n) existe. En caso contrario, se dice que la serie $\sum_{k=1}^{\infty} x_k$ es *divergente*. La serie $\sum_{k=1}^{\infty} x_k$ se llama *absolutamente convergente* si la serie de números reales positivos $\sum_{k=1}^{\infty} \|x_k\|$ converge.

Si la serie $\sum_{k=1}^{\infty} x_k$ es convergente, al elemento $S - S_n$ se le llama *resto de orden n de la serie* y se denota por

$$r_n = S - S_n = \sum_{k=n+1}^{\infty} x_k \quad (47)$$

Obviamente, *el resto de una serie convergente tiende a cero*. Utilizando la equivalencia entre la condición de Cauchy y la convergencia de sucesiones en un espacio de Banach, aplicada a la sucesión (S_n) , no es difícil ver que:

La serie $\sum_{k=1}^{\infty} x_k$ converge en el espacio de Banach $(E, \|\cdot\|)$ si para todo $\varepsilon > 0$ existe $N = N(\varepsilon) \in \mathbb{N}$ tal que $\|\sum_{k=n}^m x_k\| < \varepsilon$ cuando $m \geq n \geq N$.

Obviamente, de este criterio de convergencia de series, se concluye que *en un espacio de Banach toda serie absolutamente convergente es también convergente*.

Anteriormente hemos introducido el concepto de base algebraica de un espacio vectorial. Teniendo en cuenta la estructura topológica de los espacios normados podemos dar la definición siguiente:

D 8. (base de un espacio normado) Una familia de elementos $(x_n)_{n=1}^N$ ($N \leq \infty$) del espacio normado $(E, \|\cdot\|)$ es llamada una *base* de

$(E, \|\cdot\|)$, si cada $x \in E$ puede ser representado en la forma única:

$$x = \sum_{n=1}^N \lambda_n x_n, \quad \lambda_n \in \mathbb{C}. \quad (48)$$

Si $N = \infty$ la sumatoria se interpreta en el sentido de las series convergentes.

No es difícil verificar las siguientes *propiedades de las bases* que se dejan de ejercicio al lector:

- 1) Los elementos de una base cualquiera son vectores linealmente independientes;
- 2) Toda base de un espacio normado es también una familia total;
- 3) Toda base de un espacio normado tiene el mismo número de elementos,
- 4) En un espacio normado de dimensión algebraica finita toda base algebraica es también una base en el sentido de la Definición 8.

Sin embargo, pueden encontrarse familias totales constituidas por vectores linealmente independientes en espacios normados y que no son bases en el sentido de la Definición 8.

Ejemplo 3. (familias totales que no son bases) Hemos visto que la familia de funciones $f_k(t) = t^k, 0 \leq k < \infty$, es total en el espacio normado $(C[0, 1], \|\cdot\|_\infty)$. Sin embargo (t^k) no es una base en este espacio, ya que de la representación

$$f(t) = \sum_{k=0}^{\infty} C_k t^k \quad (49)$$

en la norma $\|\cdot\|_\infty$ se deduce la analiticidad de f para $|t| < 1$. La misma conclusión se obtiene si suponemos que la serie (49) converge en media cuadrática; luego (t^k) tampoco es una base en $(C[0,1], \|\cdot\|_2)$. $\square$

El problema sobre la existencia de bases en espacios de Banach es aún un problema abierto en la actualidad. El concepto de base de un espacio normado fue introducido por J. Schauder en 1927, quien encontró el primer ejemplo de base en el espacio $(C[a,b], \|\cdot\|_\infty)$: las funciones continuas y lineales a trozos.

En los próximos epígrafes estudiaremos la relación entre las familias totales y las bases en un tipo particular de espacios normados, llamados *espacios euclidianos*, a través de las llamadas *series de Fourier*.

§ 1.4 Producto escalar y espacios euclidianos. Ortogonalidad.

Pasaremos ahora al estudio de ciertos espacios normados con una *estructura geométrica* adicional que es compatible, en un sentido que precisaremos posteriormente, con la estructura vectorial y la estructura topológica. Esta estructura geométrica adicional se introduce mediante la noción de *ortogonalidad*. Es precisamente a partir de este nuevo concepto que se puede generalizar el conocido resultado de los cursos elementales de geometría que nos dice que la menor distancia de un punto a una recta o a un plano es la longitud del segmento que parte de dicho punto y corta perpendicularmente a la recta o el plano.

D 1. (producto escalar, espacio euclidiano). Un producto escalar en el espacio vectorial real (complejo) H es una aplicación de $H \times H$ en $\mathbb{R}(\mathbb{C})$ que a cada par de puntos $x, y \in H$ les hace corresponder el elemento $\langle x, y \rangle \in \mathbb{R}(\mathbb{C})$ y satisface las propiedades siguientes:

$$(PE\ 1) \quad \langle x, y \rangle = \overline{\langle y, x \rangle}$$

$$(PE\ 2) \quad \langle x + y, z \rangle = \langle x, z \rangle + \langle y, z \rangle$$

$$(PE\ 3) \quad \langle \lambda x, y \rangle = \lambda \langle x, y \rangle$$

$$(PE\ 4) \quad \langle x, x \rangle \geq 0 \text{ y además } \langle x, x \rangle = 0 \text{ si y sólo si } x = \Theta.$$

Al par $(H, \langle \cdot, \cdot \rangle)$ se le llama espacio euclideo. El espacio euclideo es real (complejo) si H es un espacio vectorial real (complejo). La barra sobre el miembro derecho del axioma (PE 1) denota el paso al complejo conjugado. Para productos escalares reales este axioma se expresa en la forma

$$(PE1) \quad \langle x, y \rangle = \langle y, x \rangle$$

ya que $\langle x, y \rangle$ es siempre un número real.

Mientras no se diga lo contrario nosotros nos referiremos en el texto a espacios euclideos complejos. Sin embargo la mayoría de los resultados que obtendremos son también válidos para espacios euclideos reales.

De (PE 2) y (PE 3) se deduce la linealidad del producto escalar con respecto al primer argumento. Si conjugamos estos dos axiomas con (PE 1), tendremos para el segundo argumento:

$$\begin{aligned} \langle x, \lambda y + \mu z \rangle &= \overline{\lambda} \langle x, y \rangle + \overline{\mu} \langle x, z \rangle, \\ \lambda, \mu &\in \mathbb{C}; x, y, z \in H. \end{aligned}$$

A partir de ahora, al número $\sqrt{\langle x, x \rangle}$ lo denotaremos por $\|x\|$. Nuestro primer objetivo es verificar que $\|x\|$ define efectivamente una norma sobre H . Obviamente, los axiomas (PE 1) y (PE 3) nos dan (N 2) mientras que de (PE 4) se deduce (N1).

La demostración de la desigualdad triangular (N3) es más compleja y requiere la demostración de un importante lema auxiliar.

L 1. (desigualdad de Cauchy–Buniakovsky) Para todo par de vectores $\langle x, y \rangle$ en un espacio euclideo se cumple

$$|\langle x, y \rangle| \leq \|x\| \|y\|. \quad (50)$$

La igualdad en (50) ocurre si y sólo si $x = \lambda y$ ó $y = \Theta$.

Demostración. Si $y = \Theta$ la desigualdad ocurre trivialmente. Supongamos entonces que $y \neq \Theta$. Para todo escalar $\lambda \in \mathbb{C}$, se tiene

$$\begin{aligned} 0 \leq \langle x - \lambda y, x - \lambda y \rangle &= \langle x, x \rangle - \lambda \langle y, x \rangle - \bar{\lambda} \langle x, y \rangle \\ &\quad + |\lambda|^2 \langle y, y \rangle. \end{aligned}$$

En particular, para $\lambda = \langle x, y \rangle / \langle y, y \rangle$, tendremos:

$$0 \leq \langle x, x \rangle - \frac{|\langle x, y \rangle|^2}{\langle y, y \rangle},$$

de donde se concluye (50).

Se deja de ejercicio al lector comprobar la segunda parte del lema. $\square$

P 1. En un espacio euclideo $(H, \langle, \rangle)$ la función $\|x\| = \sqrt{\langle x, x \rangle}$ es una norma.

Demostración. Queda solamente por establecer la desigualdad triangular (N3). Para cualesquiera $x, y \in H$, tenemos

$$\begin{aligned} \|x + y\|^2 &= \langle x + y, x + y \rangle = \langle x, x \rangle + \langle x, y \rangle + \langle y, x \rangle + \langle y, y \rangle \\ &= \|x\|^2 + 2\operatorname{Re}\langle x, y \rangle + \|y\|^2 \\ &\leq \|x\|^2 + 2|\langle x, y \rangle| + \|y\|^2. \end{aligned}$$

Pero de (50) se deduce que

$$\|x + y\|^2 \leq \|x\|^2 + 2\|x\|\|y\| + \|y\|^2 = (\|x\| + \|y\|)^2.$$

Extrayendo raíz cuadrada a ambos lados de esta desigualdad se obtiene el resultado deseado. $\square$

De ahora en adelante llamaremos también espacio euclideo a todo espacio normado $(E, \|\cdot\|)$ cuya norma pueda ser obtenida a partir de un producto escalar $\langle, \rangle$ mediante $\|x\| = \sqrt{\langle x, x \rangle}$. Introduciendo un producto escalar conveniente se puede mostrar que algunos de los espacios normados considerados hasta el momento son euclideos.

Ejemplo 1. (espacios euclideos) El espacio normado $\mathbb{C}^n$ provisto de la norma (35) es euclideo. En efecto, (35) es la norma correspondiente al producto escalar

$$\langle x, y \rangle = \sum_{i=1}^n x_i \bar{y}_i. \quad (51)$$

Similarmente, puede verse que la norma (38) en $C([a, b], \mathbb{C})$ proviene del producto escalar

$$\langle f, g \rangle = \int_a^b f(t) \overline{g(t)} dt. \quad (52)$$

El siguiente teorema nos brinda una caracterización de los espacios normados que son euclideos.

T 1. (identidad del paralelogramo) El espacio normado $(E, \|\cdot\|)$ es euclideo si y sólo si, para cada $x, y \in E$ se cumple la identidad del paralelogramo

$$\|x + y\|^2 + \|x - y\|^2 = 2\|x\|^2 + 2\|y\|^2 \quad (53)$$

Demostración. Si $(E, \|\cdot\|)$ es euclideo, (53) se demuestra desarrollando las normas a partir de las propiedades del producto escalar.

Inversamente, si se cumple (53), basta poner

$$4\langle x, y \rangle = \|x + y\|^2 - \|x - y\|^2 + i\|x + iy\|^2 - i\|x - iy\|^2$$

en el caso complejo, y

$$4\langle x, y \rangle = \|x + y\|^2 - \|x - y\|^2$$

en el caso real, y comprobar que $\langle, \rangle$ así definido es un producto escalar en E tal que $\|x\| = \sqrt{\langle x, x \rangle}$. $\square$

Ejemplo 2. (espacio normado no euclideo) El espacio normado $C[0, 2\pi]$ provisto de la norma $\|\cdot\|_\infty$ (39) no es euclideo. En efecto, tomando $f(x) = \sin^2 x, g(x) = \cos^2 x$, se tiene que

$$\|f + g\|_\infty = \|f - g\|_\infty = \|f\|_\infty = \|g\|_\infty = 1,$$

luego, no se satisface la identidad del paralelogramo. $\square$

Veamos la siguiente propiedad del producto escalar.

P 2. (continuidad del producto escalar) Supongamos que $x_n \rightarrow x, y_n \rightarrow y$ en el espacio euclideo $(H, \|\cdot\|)$, entonces

$$\langle x_n, y_n \rangle \longrightarrow \langle x, y \rangle$$

Demostración. Por las propiedades del producto escalar y utilizando la desigualdad de Cauchy-Buniakovsky, se concluye que:

$$|\langle x_n, y_n \rangle - \langle x, y \rangle| \leq \|x_n\| \|y_n - y\| + \|x_n - x\| \|y\|,$$

de donde se obtiene el resultado deseado, ya que la sucesión de las normas $\|x_n\|$ está acotada debido a que (x_n) es una sucesión convergente. $\square$

D 2. (convergencia débil en espacios euclidianos) La sucesión (x_n) converge débilmente al elemento x en el espacio euclidiano $(H, \langle, \rangle)$ si

$$\forall y \in H, \langle x_n, y \rangle \longrightarrow \langle x, y \rangle$$

Para la convergencia débil se utiliza la notación $(x_n) \rightharpoonup x$.

De acuerdo a esta definición a veces se conviene en llamar *convergencia fuerte* en un espacio euclidiano a la convergencia con respecto a la norma asociada al producto escalar. Obviamente, si (x_n) converge fuertemente al elemento x , también convergerá débilmente, debido a la continuidad del producto escalar.

Inversamente se tiene:

P 3. Sea (x_n) una sucesión débilmente convergente al elemento x en el espacio euclidiano $(H, \langle, \rangle)$ entonces

$$\|x_n\| \rightarrow \|x\| \Rightarrow (x_n) \rightarrow x \quad (\text{fuertemente}).$$

Demostración. Basta notar que

$$\|x_n - x\|^2 = \|x_n\|^2 - \langle x_n, x \rangle - \langle x, x_n \rangle + \|x\|^2$$

y utilizar el hecho que $\langle x_n, x \rangle \rightarrow \|x\|^2$. $\square$

Definiremos a continuación la noción de ortogonalidad y veremos algunas de sus propiedades.

D 3. (vectores ortogonales, complemento ortogonal) Dos vectores x, y del espacio euclidiano $(H, \langle, \rangle)$ se llaman ortogonales si $\langle x, y \rangle =$

0 y se denota $x \perp y$. Dado un subconjunto S de H , el conjunto de todos los vectores en H que son ortogonales a cada uno de los vectores de S es llamado complemento ortogonal de S y se denota por $S^\perp$

$$S^\perp = \{x \in H : \forall y \in S, \langle x, y \rangle = 0\} \quad (54)$$

Obviamente para dos vectores ortogonales $x \perp y$ se cumple que

$$\|x + y\|^2 = \|x\|^2 + \|y\|^2. \quad (55)$$

P 4. (propiedades del complemento ortogonal) Sean S y T subconjuntos de un espacio euclideo $(H, \langle, \rangle)$:

- i) $S = \Theta \iff S^\perp = H$,
- ii) $S^\perp$ es un subespacio vectorial cerrado de H ,
- iii) $S^\perp = \overline{S}^\perp$,
- iv) $S^\perp = [S]^\perp$,
- v) $S \cap S^\perp$ es vacío o igual al elemento nulo de H ,
- vi) $S \subset S^{\perp\perp}$,
- vii) Si $S \subset T$, entonces $T^\perp \subset S^\perp$,
- viii) $S^\perp = S^{\perp\perp\perp}$
- ix) Si S es una familia total en H , entonces $S^\perp = \{\Theta\}$

Demostración. Se deja de ejercicio al lector. $\square$

Veamos cómo se relaciona la ortogonalidad con la convergencia de

series en espacios euclidianos.

P 5. Si $x = \sum_{n=1}^{\infty} x_n$ es una serie convergente de vectores ortogonales dos a dos en el espacio euclideo $(H, \langle, \rangle)$, entonces

$$\|x\|^2 = \sum_{n=1}^{\infty} \|x_n\|^2 \quad (56)$$

Luego, una condición necesaria para la convergencia de una serie de vectores ortogonales es la convergencia de la serie numérica (56).

Demostración. Por la distributividad y continuidad del producto escalar, tenemos

$$\|x\|^2 = \langle x, \sum_{n=1}^{\infty} x_n \rangle = \sum_{n=1}^{\infty} \langle x, x_n \rangle = \sum_{n,m} \langle x_m, x_n \rangle.$$

Pero $\langle x_m, x_n \rangle = 0$ para $m \neq n$, y por lo tanto

$$\|x\|^2 = \sum_{n=1}^{\infty} \|x_n\|^2. \quad \square$$

D 4. (familias ortogonales, ortonormales y biortogonales) Una familia de vectores $S = \{x_\alpha\}_{\alpha \in I}$ en un espacio euclideo se llama *ortogonal* si cada par de vectores diferentes en S son ortogonales:

$$\alpha \neq \beta \Rightarrow \langle x_\alpha, x_\beta \rangle = 0 \quad (57)$$

La familia S se llama *ortonormal* si es ortogonal y además

$$\forall \alpha \in I, \langle x_\alpha, x_\alpha \rangle = \|x_\alpha\|^2 = 1; \quad (58)$$

es decir $\langle x_\alpha, x_\beta \rangle = \delta_{\alpha\beta}$, donde $\delta_{\alpha\beta}$ es la *delta de Kronecker*:

$$\delta_{\alpha\beta} = \begin{cases} 0 & \text{si } \alpha \neq \beta, \\ 1 & \text{si } \alpha = \beta. \end{cases}$$

La familia S se dice que es *biortogonal* a la familia $T = (y_\alpha)_{\alpha \in I}$, si

$$\forall \alpha, \beta \in I, \quad \langle x_\alpha, y_\beta \rangle = \delta_{\alpha\beta} \quad (59)$$

La existencia de un sistema biortogonal permite determinar los coeficientes del desarrollo de un elemento con respecto a una base. En efecto, si $f = \sum_{k=1}^{\infty} c_k f_k$ y (g_k) es un sistema biortogonal a (f_k) , entonces $\langle f, g_k \rangle = c_k$.

Sin embargo el siguiente ejemplo nos muestra que una sucesión arbitraria de vectores linealmente independientes de un espacio euclideo no tiene necesariamente alguna sucesión biortogonal.

Ejemplo 3. Consideremos el espacio euclideo $C[-1, 1]$ con el producto escalar

$$\langle f, g \rangle_2 = \int_{-1}^1 f(t) \overline{g(t)} dt.$$

La sucesión de funciones $1, t, t^2, \dots, t^n, \dots$ no posee ninguna sucesión biortogonal. En efecto, si una tal sucesión (g_n) existiera, entonces se tendría

$$\int_{-1}^1 g_m(t) t^m dt = 1, \quad (60)$$

$$\int_{-1}^1 g_m(t) t^k dt = 0 \quad (k \neq m) \quad (61)$$

Pero de (61), en particular, se deduce que

$$\int_{-1}^1 \{g_m(t) t^{m+1}\} t^j dt = 0 \quad (j = 0, 1, 2, \dots)$$

y como la sucesión $(t^k)_{k=0}^{\infty}$ es una familia total en $C[-1, 1]$, entonces por P 4 (ix) se tendría que $g_m(t)t^{m+1} = 0$, lo cual es una contradicción con (60). $\square$

En el § 1.6 veremos una condición necesaria y suficiente para que una sucesión de vectores en un espacio euclideo completo tenga una sucesión biortogonal.

Analicemos las siguientes nociones de numerabilidad en un espacio métrico:

I) (E, d) tiene *base topológica numerable*, es decir existe una familia numerable $\mathcal{U} = (\mathcal{U}_n)$ de abiertos en (E, d) tal que cualquier otro abierto $V \in \mathcal{O}_d$ se puede expresar como unión de abiertos de la familia.

II) (E, d) es *separable*, es decir tiene un subconjunto denso numerable.

Puede demostrarse (véase el libro Análisis Matemático en espacios métricos de A. Fraguera) que las propiedades I y II son equivalentes en cualquier espacio métrico. Si la métrica d proviene de una norma, entonces de la propiedad

III) $(E, \|\cdot\|)$ tiene *base numerable* (en el sentido de la definición 1.3.8);

se concluye cualquiera de las propiedades I y II. Para los espacios euclideos, de las propiedades equivalentes I, II se deduce la propiedad siguiente:

IV) Todo sistema ortogonal de puntos en $(H, \langle, \rangle)$ es a lo sumo numerable.

Demostración. Supongamos que $(H, \langle, \rangle)$ es separable y consideremos, sin perder generalidad, que el sistema $\{x_\alpha\}_{\alpha \in I}$ es ortonormal. En

caso contrario, podemos tomar $\{x_\alpha/\|x_\alpha\|\}$ en lugar de $\{x_\alpha\}$. En este caso $\|x_\alpha - x_\beta\| = \sqrt{2}$ para $\alpha \neq \beta$. Consideremos la familia de las bolas $B(x_\alpha, 1/2)$. Evidentemente estas bolas no se intersectan. Por ser H separable existe un subconjunto denso numerable $\{y_n\}$ en H . Luego, cada bola $B(x_\alpha, 1/2)$ debe contener al menos un elemento y_n . Por lo tanto, el número de elementos $\{x_\alpha\}$ es a lo sumo numerable. $\square$

Más adelante veremos que en los espacios euclidianos completos, también llamados *espacios de Hilbert*, estas cuatro propiedades (I–IV) son equivalentes.

A continuación mostraremos cómo pueden ser construidas familias ortonormales numerables en espacios euclidianos separables a partir de sistemas numerables de vectores linealmente independientes.

T 2. (proceso de ortonormalización de Gram–Schmidt) Sea $(H, \langle, \rangle)$ un espacio euclidiano complejo separable y $\{x_n\}$ un sistema de vectores linealmente independientes y a lo sumo numerable, de vectores de H . Entonces existe en H un sistema de vectores $\{\varphi_n\}$ que satisface las siguientes condiciones:

- i) Los sistemas $\{\varphi_n\}$ y $\{x_n\}$ tienen la misma cantidad de elementos;
- ii) El sistema φ_n es ortonormal;
- iii) Todo elemento φ_n es una combinación lineal de los elementos $\{x_n\}$ de la forma:

$$\varphi_n = a_{n1}x_1 + \cdots + a_{nn}x_n,$$

con $a_{nn} \neq 0$;

- iv) Todo elemento x_n se representa en la forma

$$x_n = b_{n1}\varphi_1 + \cdots + b_{nn}\varphi_n;$$

donde $b_{nn} \neq 0$.

Además todo elemento del sistema $\{\varphi_n\}$ queda determinado unívocamente por las condiciones ii), iii) y iv), salvo un factor de módulo 1.

Demostración. Obviamente cualquiera de las condiciones iii) o iv) implica que

$$[\{x_n\}] = [\{\varphi_n\}]$$

y además la propiedad i) se deduce de iii) y iv). Representemos el elemento φ_1 en la forma

$$\varphi_1 = |a_{11}|^2 x_1;$$

entonces, a_{11} se determina por la condición

$$\langle \varphi_1, \varphi_1 \rangle = a_{11} \langle x_1, x_1 \rangle = 1,$$

de donde

$$a_{11} = \frac{1}{b_{11}} = \frac{\pm 1}{\sqrt{\langle x_1, x_1 \rangle}}.$$

Es claro que φ_1 está unívocamente determinado (salvo un factor de módulo 1). Supongamos que ya han sido construidos los elementos φ_k con $k < n$ que verifican las condiciones ii), iii) y iv). En este caso, se puede representar x_n en la forma

$$x_n = b_{n1}\varphi_1 + \cdots + b_{n,n-1}\varphi_{n-1} + h_n,$$

donde $\langle h_n, \varphi_k \rangle = 0$ para $k < n$.

En efecto, los coeficientes correspondientes b_{nk} y, por consiguiente, el elemento h_n quedan unívocamente determinados por las condiciones

$$\begin{aligned} \langle h_n, \varphi_k \rangle &= \langle x_n - b_{n1}\varphi_1 - \cdots - b_{n,n-1}\varphi_{n-1}, \varphi_k \rangle \\ &= \langle x_n, \varphi_k \rangle - b_{nk} \langle \varphi_k, \varphi_k \rangle = 0. \end{aligned}$$

Es evidente que $\langle h_n, h_n \rangle \neq 0$, ya que la suposición $\langle h_n, h_n \rangle = 0$ estará en contradicción con la independencia lineal del sistema $\{x_n\}$. Pongamos

$$\varphi_n = \frac{h_n}{\sqrt{\langle h_n, h_n \rangle}}.$$

De esta construcción inductiva se desprende que h_n y, por consiguiente, también φ_n se expresan mediante $x_1, \dots, x_n$, es decir,

$$\varphi_n = a_{n1}x_1 + \dots + a_{nn}x_n,$$

donde

$$a_{nn} = \frac{1}{\sqrt{\langle h_n, h_n \rangle}} \neq 0.$$

Además

$$\langle \varphi_n, \varphi_n \rangle = 1, \quad \langle \varphi_n, \varphi_k \rangle = 0 \quad (k < n),$$

y $x_n = b_{n1}\varphi_1 + \dots + b_{nn}\varphi_n$, siendo

$$b_{nn} = \sqrt{\langle h_n, h_n \rangle} \neq 0,$$

es decir $\{\varphi_n\}$ verifica las condiciones del teorema. $\square$

Ejemplo 4. (polinomios ortogonales) Consideremos de nuevo el espacio euclideo $C[-1, 1]$ con el producto escalar

$$\langle f, g \rangle_2 = \int_{-1}^1 f(t) \overline{g(t)} dt.$$

Las funciones linealmente independientes $1, t, t^2, \dots, t^n, \dots$, generan el subespacio de los polinomios. El procedimiento de Gram-Schmidt puede ser aplicado a la sucesión $1, t, t^2, \dots$, y obtendremos una sucesión ortonormal $\{e_n\}_{n=0}^\infty$. Obviamente los e_n son polinomios por ser combinaciones lineales de polinomios. Puede demostrarse que

$$e_n(t) = \sqrt{\frac{2n+1}{2}} P_n(t), \quad n = 0, 1, 2, \dots,$$

donde los $P_n(t)$ son los bien conocidos polinomios de Legendre

$$P_n(t) = \frac{(-1)^n}{2^n n!} \frac{d^n}{dt^n} \{(1-t^2)^n\}.$$

Al ortonormalizar el sistema $\{1, t, t^2, \dots, t^n, \dots\}$ con respecto a otros productos escalares definidos sobre $C[-1, 1]$; por ejemplo:

$$\langle f, g \rangle = \int_{-1}^1 f(t)g(t)(1-t)^\alpha(1+t)^\beta dt,$$

con $\alpha, \beta \geq -1$, se obtienen nuevos sistemas de polinomios ortonormales clásicos conocidos con el nombre de polinomios de Jacobi de tipo α, β . En este caso para definir el producto escalar se ha utilizado una *función de peso* de tipo $(1-t)^\alpha(1+t)^\beta$. Utilizando otras funciones de peso, como por ejemplo e^{-x} en el intervalo $[0, \infty]$ en lugar de $[-1, 1]$, se obtienen los llamados polinomios de Legendre-Hermite. Los polinomios ortogonales han sido objeto de investigación de la teoría clásica de aproximación. $\square$

§ 1.5 Series de Fourier y bases en espacios euclidianos.

En este epígrafe estudiaremos la relación entre las bases y las familias totales en un espacio euclidiano, compuestas por vectores ortonormales. Para ello se hace necesario introducir la definición siguiente:

D 1. (coeficiente de Fourier, serie de Fourier con respecto a un sistema ortonormal) Sea $(H, \langle, \rangle)$ un espacio euclidiano separable y $\{e_n\}$ un sistema ortonormal de vectores en H . Para cada $x \in H$, se llama n -ésimo coeficiente de Fourier de x con respecto al sistema $\{e_n\}$, al número

$$a_n = \langle x, e_n \rangle, \quad n = 0, 1, \dots \quad (62)$$

y a la serie (por ahora formal)

$$\sum_n a_n e_n$$

se le llama serie de Fourier del elemento x con respecto al sistema ortonormal $\{e_n\}$.

Ejemplo 1. (sistema trigonométrico) En el espacio euclideo $C[-\pi, \pi]$ con el producto escalar integral

$$\langle f, g \rangle_2 = \int_{-\pi}^{\pi} f(t) \overline{g(t)} dt,$$

las funciones

$$e_k(t) = \begin{cases} \frac{1}{\sqrt{2\pi}}, & k = 0 \\ \frac{1}{\sqrt{\pi}} \operatorname{sen} kt, & k = 1, 2, \dots \\ \frac{1}{\sqrt{\pi}} \cos kt, & k = -1, -2, \dots \end{cases} \quad (63)$$

constituyen un sistema ortonormal en $(C[-\pi, \pi], \langle, \rangle_2)$, lo cual puede verificar el lector mediante un cálculo directo. Del teorema de aproximación de Weierstrass puede concluirse que los polinomios trigonométricos $\{e_k(t)\}$ constituyen un sistema ortonormal total en $(C[-\pi, \pi], \langle, \rangle_2)$.

El sistema (63) aparece en la teoría clásica de las series de Fourier. Como consecuencia de los resultados que veremos a continuación se tiene que cada $f \in C[-\pi, \pi]$ admite un desarrollo en serie de Fourier del tipo

$$f(t) = \frac{a_0}{2} + \sum_{n=1}^{\infty} (a_n \cos nt + b_n \operatorname{sen} nt) \quad (64)$$

donde

$$a_0 = \frac{1}{\pi} \int_{-\pi}^{\pi} f(t) dt \quad (65)$$

$$a_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(t) \cos ntdt \quad (66)$$

$$b_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(t) \operatorname{sen} ntdt, n = 1, 2, \dots \quad (67)$$

donde la serie (64) converge en el sentido de la norma $\|\cdot\|_2$ definida en (38) con $a = -\pi, b = \pi$. El estudio de las propiedades de la serie de Fourier y de su convergencia a f en diferentes sentidos (puntualmente, uniformemente, etc.) constituye uno de los objetivos de la teoría clásica de las series de Fourier.

Puede demostrarse que si denotamos por $S_n(t)$ a la suma parcial de orden n de la serie (64) correspondiente a una función periódica f de período 2π y definimos las *sumas de Cesaro* como

$$\sigma_n(t) = \frac{S_0(t) + S_1(t) + \cdots + S_n(t)}{n+1}, \quad (68)$$

entonces $\sigma_n(t)$ converge uniformemente a $f(t)$ en $[-\pi, \pi]$ (véase el libro de A. Fraguera *Análisis Matemático en espacios normados*). $\square$

Lógicamente, después de la Definición 1.5.1, nos surgen las preguntas siguientes: ¿Es convergente la serie (64)?, es decir ¿converge a algún límite la sucesión de sus sumas parciales en el sentido de la métrica del espacio H ?, y si converge ¿coincide su suma con el elemento inicial x ?

Para responder a estas preguntas veamos las dos proposiciones siguientes y el Teorema 1.

P 1. Sea $(H, \langle, \rangle)$ un espacio euclideo separable, $\{e_k\}$ un sistema ortonormal en H y $x \in H$. Entonces para cada $n \in \mathbb{N}$, el elemento de la forma

$$\tilde{S}_n = \sum_{k=1}^n \alpha_k e_k \quad (\alpha_k \in \mathbb{R})$$

que se encuentra a menor distancia de x , coincide con la suma parcial de orden n de la serie de Fourier de x (es decir, la distancia mínima de x a $\{e_k\}_1^n$ se alcanza tomando $\alpha_k = \langle x, e_k \rangle$).

De la proposición anterior se deduce que si ponemos

$$S_n = \sum_{k=1}^n \langle f, e_k \rangle e_k = \sum_{k=1}^n a_k e_k$$

entonces

$$\|f - S_n\|^2 = \|f\|^2 - \sum_{k=1}^n a_k^2 \quad (69)$$

P 2. (desigualdad de Bessel) Sean $(H, \langle, \rangle)$ y $\{e_k\}$ como en la proposición anterior. Entonces para toda $x \in H$ se cumple

$$\sum_k a_k^2 \leq \|x\|^2 \quad (70)$$

donde $a_k = \langle x, e_k \rangle$, $k = 1, 2, \dots$, son los coeficientes de Fourier de f con respecto al sistema $\{e_k\}$.

Demostración. De (69) se deduce que para todo $n \in \mathbb{N}$

$$\sum_{k=1}^n a_k^2 \leq \|f\|^2$$

luego (70) se satisface inmediatamente si $\{e_k\}$ es una familia finita y por paso al límite si es numerable. $\square$

Podemos ahora enunciar un teorema que nos explica la relación existente entre las bases y las familias totales compuestas por vectores ortonormales.

T 1. Sean $(H, \langle, \rangle), \{e_k\}$ como en la Proposición 1. Entonces son equivalentes:

i) Para todo $x \in H$ se cumple la identidad de Parseval

$$\sum_k a_k^2 = \|x\|^2 \quad (71)$$

donde a_k son los coeficientes de Fourier de x con respecto al sistema $\{e_k\}$;

- ii) Toda $x \in H$ se puede representar mediante su serie de Fourier. Es decir, para cada $x \in H$, la serie de Fourier correspondiente a x converge precisamente a x ;
- iii) El sistema $\{e_k\}$ es una base en H ;
- iv) El sistema $\{e_k\}$ es total en H .

Demostración. i) $\Rightarrow$ ii) se deduce inmediatamente de (69);

ii) $\Rightarrow$ iii) evidente;

iii) $\Rightarrow$ iv) pues toda base es siempre una familia total;

iv) $\Rightarrow$ i) supongamos que el sistema $\{e_n\}$ es total, es decir todo elemento $x \in H$ se puede aproximar con precisión arbitraria mediante combinaciones lineales del tipo

$$\sum_{k=1}^n \alpha_k e_k.$$

De P 1 sabemos que la suma parcial S_n de la serie de Fourier de x nos da una aproximación mejor a f .

Por consiguiente, la serie de Fourier de x converge a x y de (69) se deduce que tiene lugar la identidad de Parseval. $\square$

Como consecuencia de T 1.4.2 y T 1, obtenemos el corolario siguiente.

C 1. (existencia de bases ortonormales) En todo espacio euclideo

separable existe siempre una base ortonormal.

Demostración. En efecto, si $\{\varphi_n\}$ es un conjunto numerable denso en H , de él se puede extraer un sistema total constituido por vectores linealmente independientes. Para ello basta con omitir de la sucesión $\{\varphi_n\}$ todos aquellos elementos φ_k que pueden representarse como combinación lineal de φ_i con $i < k$.

Aplicando el proceso de ortonormalización de Gram-Schmidt (T 1.4.2) encontramos un sistema ortonormal total en H , que según T1 será una base ortonormal de H . $\square$

Observación. La existencia de bases ortonormales constituye un factor importante para la aplicación de los métodos numéricos a la solución de ecuaciones en espacios normados. Para tener una idea de la importancia del corolario anterior, aunque no veremos aquí la demostración por ser un resultado bastante más profundo, debemos señalar que hay espacios de Banach en los cuales no se puede demostrar la existencia de bases.

Veamos otra propiedad importante de las bases ortonormales.

P 3. En un espacio euclideo separable toda base ortonormal es un conjunto ortonormal maximal.

Demostración. Sea $\{e_k\}$ una base ortonormal en H . Si suponemos que dicha base no constituye un conjunto ortonormal maximal, entonces existirá $x \in H$ no nulo, ortogonal a todos los e_k . Por el Teorema 1, x puede expresarse mediante su serie de Fourier con respecto al sistema $\{e_k\}$

$$x = \sum_k \langle x, e_k \rangle e_k$$

pero como $\langle x, e_k \rangle = 0$ para todo k , entonces $x = 0$; contradicción. $\square$

Observación. En el próximo epígrafe veremos que en los espacios de Hilbert también se cumple el inverso de la proposición anterior y por lo tanto la maximalidad es también otra propiedad en los espacios de Hilbert separables, que junto a las expuestas en el Teorema 1, caracteriza a los conjuntos ortonormales que son bases. Notemos que en la proposición anterior nos referimos al concepto usual de maximalidad: un sistema ortonormal se llama maximal en H si deja de ser ortonormal al añadirle cualquier otro elemento de H no perteneciente a dicho sistema.

Pasemos finalmente a estudiar el segundo problema relacionado con el concepto de ortogonalidad en los espacios euclidianos, es decir el problema de la mejor aproximación, que hemos ya tratado en la Proposición 1. Notemos que el resultado obtenido en P1 puede ser interpretado geoméricamente del siguiente modo: el elemento $x - \sum_{k=1}^n \alpha_k e_k$ es ortogonal a todas las combinaciones lineales del tipo $\sum_{k=1}^n \beta_k e_k$ es decir, es ortogonal al subespacio generado por los elementos $\{e_1, e_2, \dots, e_n\}$ cuando, y solamente cuando, se cumple la condición $\alpha_k = \langle x, e_k \rangle$ (compruébese esta afirmación). Por consiguiente, el resultado obtenido en P1 representa una generalización del conocido teorema de la Geometría Elemental: la longitud de la perpendicular bajada de un punto dado a una recta o a un plano es menor que la longitud de cualquier oblicua trazada por el mismo punto. Precisamente, la expresión abstracta de este resultado es la primera versión del teorema general de optimización que expondremos a continuación en los espacios euclidianos. En el próximo epígrafe veremos otras versiones, con conclusiones más fuertes para los espacios de Hilbert.

T 2. Sea $(H, \langle, \rangle)$ un espacio euclidiano, M un subespacio vectorial de H y x un vector arbitrario de H . Si existe un vector $m_0 \in M$ tal que $\|x - m_0\| \leq \|x - m\|$ para todo $m \in M$, entonces m_0 es único.

Una condición necesaria y suficiente para que $m_0 \in M$ sea el vector minimizante (único) en M es que el vector error $x - m_0$ sea ortogonal a M ($x - m_0 \in M^\perp$).

Demostración. Probaremos primero que si m_0 es un vector minimizante, entonces $x - m_0$ es ortogonal a M . Supongamos por el contrario que existe un $m \in M$ que no es ortogonal a $x - m_0$. Sin pérdida de generalidad, podemos asumir que $\|m\| = 1$ y que $\langle x - m_0, m \rangle = \delta > 0$. Definamos el vector m_1 en M como $m_1 = m_0 + \delta m$.

Entonces

$$\begin{aligned} \|x - m_1\|^2 &= \|x - m_0 - \delta m\|^2 = \|x - m_0\|^2 - \langle x - m_0, \delta m \rangle \\ &\quad - \langle \delta m, x - m_0 \rangle + |\delta|^2 \\ &= \|x - m_0\|^2 - |\delta|^2 < \|x - m_0\|^2 \end{aligned}$$

Luego, si $x - m_0$ no es ortogonal a M , m_0 no es un vector minimizante.

Probemos ahora que si $x - m_0 \in M^\perp$, entonces m_0 es el único vector minimizante. Para cualquier $m \in M$, por el teorema de Pitágoras generalizado (P 1.4.5), tenemos

$$\|x - m\|^2 = \|x - m_0 + m_0 - m\|^2 = \|x - m_0\|^2 + \|m_0 - m\|^2$$

y por lo tanto

$$\|x - m\| > \|x - m_0\| \quad \text{para todo } m \neq m_0. \quad \square$$

Observación. Notemos que en el teorema anterior (T2) no hemos establecido la existencia del vector minimizante, lo cual solo podemos asegurar hasta el momento si M es de dimensión finita (P1). Hemos probado que si dicho vector m_0 existe, entonces el mismo es único y $x - m_0$ es ortogonal al subespacio M . En el próximo epígrafe introduciremos una clase de espacios euclidianos, llamados espacios de Hilbert, donde se garantiza la existencia del vector minimizante.

§ 1.6 Espacios de Hilbert

Introduzcamos ahora una de las definiciones más importantes en el desarrollo de la matemática actual.

D 1. (espacio de Hilbert) Se llama espacio de Hilbert a todo espacio euclideo completo con respecto a la norma asociada al producto escalar.

Ejemplo 1. Sabemos que el espacio euclideo $\mathbb{R}^n$ (resp. $\mathbb{C}^n$) con el producto escalar

$$\begin{aligned}\langle x, y \rangle &= \sum_{k=1}^n x_k y_k \\ \left(\langle x, y \rangle &= \sum_{k=1}^n x_k \bar{y}_k \right)\end{aligned}\tag{72}$$

es un espacio de Hilbert, por ser de dimensión finita. Consideremos el espacio $l_2(\mathbb{N})$ de las sucesiones $x = (x_n)$ de números complejos con cuadrado sumable

$$\sum_{n=0}^{\infty} |x_n|^2 < \infty\tag{73}$$

Este espacio es un espacio euclideo con el producto escalar

$$\langle x, y \rangle = \sum_{n=0}^{\infty} x_n \bar{y}_n\tag{74}$$

Demostremos que $l_2(\mathbb{N})$ es también completo. En efecto, sea $(x^{(m)}) = (x_n^{(m)})$ una sucesión de Cauchy en $l_2(\mathbb{N})$. Esto significa que para cada $\varepsilon > 0$ existe un m_0 tal que

$$p, q \geq m_0 \Rightarrow \sum_{n=0}^{\infty} |x_n^{(p)} - x_n^{(q)}|^2 \leq \varepsilon.$$

Para cada n fijo, esto implica en primer lugar que la sucesión $(x_n^{(m)})$ es de Cauchy en $\mathbb{R}$ y, por tanto, converge hacia un límite x_n .

En segundo lugar, dejando $p \geq m_0$ fijo y haciendo tender q a ∞ en $\sum_{n=0}^N |x_n^{(p)} - x_n^{(q)}|^2$ ($N \in \mathbb{N}$), obtenemos por la continuidad del módulo, que

$$\sum_{n=0}^N |x_n^{(p)} - x_n|^2 \leq \varepsilon \quad \forall N \in \mathbb{N},$$

luego

$$\sum_{n=0}^{\infty} |x_n^{(p)} - x_n|^2 \leq \varepsilon$$

para $p \geq m_0$. Esto demuestra que la sucesión $(x_n^{(p)} - x_n)$ pertenece a $l_2(\mathbb{N})$; por lo tanto $x = (x_n)$ pertenece también a $l_2(\mathbb{N})$ y se tiene que $x^{(m)} \rightarrow x$ en $l_2(\mathbb{N})$. $\square$

Más adelante veremos que, en cierto modo, los espacios $\mathbb{R}^n$ y $l_2(\mathbb{N})$ son modelos de espacios de Hilbert separables. Hemos visto que hay espacios euclidianos que no son de Hilbert. Sin embargo, a todo espacio euclidiano se le asocia una norma y, por tanto, una métrica. El espacio métrico asociado a un espacio euclidiano admite siempre un *completamiento* (T 1.2.2). Resulta que puede probarse que el *completamiento de un espacio euclidiano es un espacio de Hilbert* (¡ demuéstrello!). Por este motivo en ocasiones se conviene en llamar espacios *prehilbertianos* a los espacios euclidianos.

Ejemplo 2. Consideremos de nuevo el espacio $C[a, b]$ provisto del producto escalar

$$\langle f, g \rangle_2 = \int_a^b f(t) \overline{g(t)} dt.$$

Sabemos que $(C[a, b], \langle, \rangle_2)$ es un espacio euclidiano. Desarrollando los cálculos para la métrica d_2 asociada al producto escalar $\langle, \rangle_2$, en

el Ejemplo 1.2.1, puede comprobarse que este espacio euclideo no es completo. Sin embargo $(C[a, b], \langle, \rangle_2)$ admite un completamiento que denotaremos por $L_2[a, b]$, el cual es un espacio de Hilbert.

Un resultado importante es que $L_2[a, b]$ puede realizarse en un espacio de clases de funciones complejas definidas sobre $[a, b]$ y que el producto escalar en $L_2[a, b]$ puede espresarse mediante una noción de integral más general llamada *integral de Lebesgue*, la cual coincide para las funciones continuas con la *integral de Cauchy*, a partir de la cual se define $\langle, \rangle_2$.

La teoría de la integral de Lebesgue se desarrolla también para funciones definidas sobre conjuntos abiertos acotados Ω de $\mathbb{R}^n$, pudiéndose así introducir los espacios de Hilbert $L_2(\Omega)$ compuestos por funciones de cuadrado integrable en sentido de Lebesgue y que de tanta utilidad resultan en el estudio de los operadores diferenciales.

Los resultados fundamentales sobre la teoría de integración según Lebesgue pueden ser hallados por el lector en el libro de A.N. Kolmogorov y S.V. Fomin, *Elementos de la teoría de funciones y del análisis funcional* [2], y en el libro de W. Rudin, *Real and complex analysis* [3].

□

Veamos ahora algunas propiedades de los espacios de Hilbert que mejoran las correspondientes en los espacios euclideos. La obtención de este primer grupo de propiedades en los espacios de Hilbert se debe fundamentalmente a que hemos añadido la propiedad de completitud.

P 1. Es un espacio de Hilbert $(H, \langle, \rangle)$ para una sucesión (x_n) de elementos dos a dos ortogonales, son equivalentes:

- i) $\sum_{n=1}^{\infty} x_n$ es convergente;
- ii) $\sum_{n=1}^{\infty} \|x_n\|^2$ es convergente.

Demostración. La implicación i) $\Rightarrow$ ii) se cumple en cualquier espacio euclideo por la Proposición 1.4.5. Inversamente, supongamos que la serie $\sum_{n=1}^{\infty} \|x_n\|^2$ converge y pongamos

$$S_m = \sum_{n=1}^m x_n$$

para $m > p$ se tiene

$$S_m - S_p = \sum_{n=p+1}^m x_n$$

y por 1.4.5

$$\|S_m - S_p\|^2 = \sum_{n=p+1}^m \|x_n\|^2.$$

De la convergencia de la serie $\sum_{n=1}^{\infty} \|x_n\|^2$ se deduce que la sucesión (S_n) es de Cauchy en H , luego es convergente con lo cual probamos que la serie $\sum_{n=1}^{\infty} x_n$ converge. $\square$

De la desigualdad de Bessel (70) se deduce que para que los números $a_1, a_2, \dots, a_n, \dots$ representen los coeficientes de Fourier de un elemento en un espacio euclideo H es necesario que la serie

$$\sum_n a_n^2$$

converja. Resulta, que en un espacio de Hilbert esta condición no sólo es necesaria, sino también suficiente. Tiene lugar el teorema siguiente.

T 1. (Riesz–Fischer) Sea $\{e_n\}$ un sistema ortonormal arbitrario en un espacio de Hilbert H y sean los números

$$a_1, a_2, \dots, a_n, \dots$$

tales que la serie

$$\sum_{k=1}^{\infty} a_k^2$$

converge. Entonces existe un elemento $x \in H$ tal que $a_k = \langle x, e_k \rangle$ y $\sum_{k=1}^{\infty} a_k^2 = \langle x, x \rangle = \|x\|^2$

Demostración. Pongamos

$$x_n = \sum_{k=1}^n a_k e_k$$

Entonces por la Proposición 1.4.5

$$\|x_{n+p} - x_n\|^2 = \sum_{k=n+1}^{n+p} a_k^2.$$

Como la serie $\sum_{k=1}^{\infty} a_k^2$ converge, de aquí se deduce, en virtud de la completitud de H , la convergencia de la sucesión $\{x_n\}$ a un elemento $x \in H$. Además

$$\langle x, e_i \rangle = \langle x_n, e_i \rangle + \langle x - x_n, e_i \rangle \quad (75)$$

donde el primer sumando del miembro derecho es igual a a_i para $n \geq i$, mientras que el segundo sumando tiende a cero cuando $n \rightarrow \infty$ debido a la continuidad del producto escalar (P 1.4.2). El miembro izquierdo de la igualdad (75) no depende de n ; por eso, pasando al límite cuando $n \rightarrow \infty$ a ambos lados, obtenemos $\langle x, e_i \rangle = a_i$.

De acuerdo con la definición de x , $\|x - x_n\| \rightarrow 0$, cuando $n \rightarrow \infty$ y por eso

$$\sum_{k=1}^{\infty} a_k^2 = \langle x, x \rangle.$$

En efecto,

$$\langle x - \sum_{k=1}^n a_k e_k, x - \sum_{k=1}^n a_k e_k \rangle = \langle x, x \rangle - \sum_{k=1}^n a_k^2 \xrightarrow{n \rightarrow \infty} 0$$

con lo cual el teorema queda demostrado. $\square$

Del teorema de Riesz–Fischer se deduce el importante resultado siguiente:

T 2. Para un sistema ortonormal $\{e_n\}$ en un espacio de Hilbert separable $(H, \langle, \rangle)$, son equivalentes:

- i) $\{e_n\}$ es un sistema ortonormal maximal en H .
- ii) $\{e_n\}$ es una familia total en H ;
- iii) No existe ningún elemento diferente de Θ en H que sea ortogonal a todos los elementos del sistema $\{e_n\}$.

(compárese con P 1.4.4 ix) y con P 1.5.3).

Demostración. i) $\Rightarrow$ ii) Supongamos que la familia $\{e_n\}$ no es total. Entonces por T 1.5.1 (ii) existe $x \in H$ que no puede ser representado por su serie de Fourier con respecto a $\{e_n\}$:

$$x \neq \sum_{n=1}^{\infty} a_n e_n$$

donde $a_n = \langle x, e_n \rangle$. Por el Teorema de Riesz–Fischer (T1) sabemos que la serie de Fourier de x converge a un cierto elemento de H que en nuestro caso es diferente de x . Pero entonces el elemento $x - \sum_{n=1}^{\infty} a_n e_n$ es un vector no nulo de H , ortogonal a todos los e_n ; luego el sistema $\{e_n\}$ no es maximal.

ii) $\Rightarrow$ iii) es consecuencia inmediata de P 1.4.4. ix) y se cumple en todo espacio euclideo. Esto puede verse también de la siguiente forma: si $\{e_n\}$ es total, por la equivalencia i) $\Rightarrow$ iv) de T 1.5.1, se cumple la identidad de Parseval para todo $x \in H$. Luego, si x es ortogonal a

todos los elementos del sistema $\{e_n\}$, entonces todos sus coeficientes de Fourier a_k se anulan, y por la identidad de Parseval obtenemos

$$\langle x, x \rangle = \sum_{k=1}^{\infty} a_k^2 = 0, \quad \text{de donde } x = \Theta.$$

iii) $\Rightarrow$ i) se deduce de la definición de maximalidad. $\square$

En los espacios de Hilbert se cumple el teorema siguiente sobre la estabilidad de las bases ortonormales.

T 3. Si (e_n) es una base ortonormal en el espacio de Hilbert H , y (f_n) es un sistema ortonormal de vectores *cuadráticamente cerca* a la base (e_n) , es decir

$$\sum_{k=1}^{\infty} \|f_k - e_k\|^2 < \infty,$$

entonces (f_n) es también una base ortonormal de H .

Demostración. Probemos que si $f_0 \perp f_k$, para cualquier $k = 1, 2, \dots$, entonces $f_0 = \Theta$. Entonces por el Teorema 2 tendremos que el sistema (f_n) es total en H y de T 1.5.1 podremos concluir que (f_n) es una base ortonormal.

Supongamos que de la relación $f_0 \perp f_k$ para todo $k = 1, 2, \dots$, no se tiene que $f_0 = \Theta$, es decir $f_0 \neq \Theta$. Entonces $f_0, f_1, f_2, \dots$, es una familia ortogonal y, por lo tanto, está constituida por vectores linealmente independientes. Mostremos que esto no puede ocurrir. Tomemos $M \in \mathbb{N}$, tal que $\sum_{k>M} \|e_k - f_k\|^2 < 1$, y mostremos que los vectores $f_0, f_1, \dots, f_M$ son linealmente dependientes. Para ello pongamos

$$g_k = \sum_{j=1}^M \langle f_k, e_j \rangle e_j, \quad k = 0, 1, \dots, M.$$

Los $M+1$ vectores $g_k, k = 0, 1, \dots, M$, pertenecen al espacio vectorial generado por los M vectores $e_j, j = 1, \dots, M$ y por ello los vectores $g_k, k = 0, 1, \dots, M$ son linealmente dependientes. Consideremos una combinación lineal nula

$$\sum_{k=0}^M C_k g_k = \Theta.$$

Sustituyendo en esta igualdad por la expresión de g_k , llegamos a

$$\sum_{k=0}^M C_k \sum_{j=1}^M \langle f_k, e_j \rangle e_j = \sum_{j=1}^M \langle \sum_{k=0}^M C_k f_k, e_j \rangle e_j = \Theta.$$

Por la independencia lineal de los vectores $e_j, j = 1, \dots, M$, tenemos que

$$\langle \sum_{k=0}^M C_k f_k, e_j \rangle = 0, \quad j = 1, \dots, M.$$

El vector $h = \sum_{k=0}^M C_k f_k$ es ortogonal a $e_1, \dots, e_M$. Calculemos la norma del vector h . Se tiene

$$\|h\|^2 = \sum_{k=1}^{\infty} |\langle h, e_k \rangle|^2 = \sum_{k=M+1}^{\infty} |\langle h, e_k \rangle|^2.$$

Por otra parte, h pertenece al espacio vectorial generado por $f_0, \dots, f_M$ y, por lo tanto, es ortogonal a $f_{M+1}, f_{M+2}, \dots$. Luego

$$\begin{aligned} \|h\|^2 &= \sum_{k=M+1}^{\infty} |\langle h, e_k \rangle - \langle h, f_k \rangle|^2 \\ &\leq \sum_{k=M+1}^{\infty} \|h\|^2 \|e_k - f_k\|^2 < \|h\|^2 \end{aligned}$$

De aquí se deduce que $h = 0$, es decir, los vectores $f_0, f_1, \dots, f_M$ son linealmente dependientes. $\square$

Veamos otra consecuencia importante del Teorema de Riesz–Fischer que nos dice que desde el punto de vista de la estructura introducida por el producto escalar todos los espacios de Hilbert son equivalentes. El concepto de equivalencia entre espacios de Hilbert debe conservar:

- la estructura conjuntista,
- la estructura vectorial,
- la estructura métrica,
- la estructura geométrica asociada al producto escalar

y además debe mantener la compatibilidad entre estas estructuras.

A este efecto se introduce la definición siguiente:

D 2. (espacios de Hilbert isomorfos) Dos espacios de Hilbert $(H_1, \langle \cdot, \cdot \rangle_1)$ y $(H_2, \langle \cdot, \cdot \rangle_2)$ son isomorfos si existe una aplicación biyectiva

$$i : H_1 \longrightarrow H_2$$

tal que

$$\langle x, y \rangle_1 = \langle i(x), i(y) \rangle_2 \quad (76)$$

para todos $x, y \in H_1$.

Observación. Notemos que de la propiedad (76) se deduce que para todos $x_1, x_2, x \in H_1, \lambda \in \mathbb{R}(\mathbb{C})$ se tiene

$$i(x_1 + x_2) = i(x_1) + i(x_2); \quad (77)$$

$$i(\lambda x) = \lambda i(x). \quad (78)$$

En efecto, para cada $y \in H_1$

$$\begin{aligned}\langle i(x_1 + x_2), i(y) \rangle_2 &= \langle x_1 + x_2, y \rangle_1 = \langle x_1, y \rangle_1 + \langle x_2, y \rangle_1 \\ &= \langle i(x_1), i(y) \rangle_2 + \langle i(x_2), i(y) \rangle_2,\end{aligned}$$

de donde, agrupando convenientemente,

$$\langle i(x_1 + x_2) - i(x_1) - i(x_2), i(y) \rangle = 0.$$

Por la sobreyectividad de la aplicación i , la arbitrariedad de y , y utilizando el hecho que $H_2^\perp = \{\Theta\}$, concluimos que

$$i(x_1 + x_2) - i(x_1) - i(x_2) = 0,$$

o sea se cumple (77). No resulta difícil verificar (78).

Por otra parte de (76) se deduce que para las normas asociadas a los respectivos productos escalares, se tiene

$$\|x\|_1 = \|i(x)\|_2,$$

es decir i determina una isometría con respecto a las estructuras métricas subyacentes.

T 4. (isomorfismos entre espacios de Hilbert) Todos los espacios de Hilbert separables son isomorfos entre sí.

Demostración. En el caso de espacios de dimensión finita la demostración es evidente y por ello se deja al lector. Pasemos a analizar el caso de dimensión infinita. El lector puede verificar sin dificultad que la relación de isomorfía es una relación de equivalencia entre espacios de Hilbert. Luego basta demostrar que todo espacio de Hilbert separable H es isomorfo a un espacio modelo.

Como espacio modelo tomaremos el espacio $l_2(\mathbb{N})$ introducido en el Ejemplo 1. Por ser H un espacio de Hilbert separable, de C 1.5.1

concluimos que existe una base ortonormal numerable $\{e_n\}$ en H . Definamos la aplicación

$$i : H \longrightarrow l_2(\mathbb{N})$$

de la siguiente manera:

$$i(x) = (a_1, a_2, \dots, a_n, \dots) \quad (79)$$

donde $a_i = \langle x, e_i \rangle$ es el i -ésimo coeficiente de Fourier de x con respecto a la base ortonormal $\{e_n\}$. La aplicación i está bien definida ya que por la desigualdad de Bessel

$$\sum_{n=1}^{\infty} a_n^2 < \infty.$$

Mostremos que i , definida mediante (79) establece una correspondencia biunívoca entre H y $l_2(\mathbb{N})$. La inyectividad de i se deduce del Teorema 2 y la sobreyectividad del Teorema de Riesz–Fischer (T1). En efecto, si $x, y \in H$ con $x \neq y$, entonces $x - y \neq \Theta$. Por T2 y por ser $\{e_n\}$ una base de H , $x - y$ no puede ser ortogonal a todos los elementos e_n . Luego, para algún $n \in \mathbb{N}$ se tendrá

$$\langle x - y, e_n \rangle \neq 0,$$

es decir

$$\langle x, e_n \rangle \neq \langle y, e_n \rangle,$$

y por lo tanto $i(x) \neq i(y)$, con lo cual queda probada la inyectividad.

El lector puede verificar que la sobreyectividad de i es precisamente la tesis del Teorema de Riesz–Fischer.

Nos queda por probar que

$$\langle x, y \rangle = \langle i(x), i(y) \rangle_2 = \sum_{n=1}^{\infty} a_n \overline{b_n} \quad (80)$$

donde $\{a_n\}$ y $\{b_n\}$ son los coeficientes de Fourier de x y y respectivamente y además hemos denotado por $\langle, \rangle_2$ al producto escalar definido

en $l_2(\mathbb{N})$. Notemos que la identidad de Parseval (la cual se cumple en este caso según T 1.5.1) puede ser escrita también en la forma

$$\langle x, x \rangle = \langle i(x), i(x) \rangle_2 \quad (81)$$

Sean ahora $x, y \in H$. De (81) tenemos que se cumple

$$\langle x + \lambda y, x + \lambda y \rangle = \langle i(x) + \lambda i(y), i(x) + \lambda i(y) \rangle_2$$

para todo escalar λ . De aquí, por las propiedades del producto escalar

$$\overline{\lambda} \langle x, y \rangle + \lambda \langle y, x \rangle = \overline{\lambda} \langle i(x), i(y) \rangle_2 + \lambda \langle i(y), i(x) \rangle_2 \quad (82)$$

Tomemos $\lambda = 1$ y $\lambda = i$ en (82). Entonces de (82) se deduce que $\langle x, y \rangle$ y $\langle i(x), i(y) \rangle_2$ tienen la misma parte real e imaginaria y por tanto son iguales, con lo cual queda probado (80) y de hecho el teorema. $\square$

Pasemos ahora a ver otras propiedades de los espacios de Hilbert más vinculadas al concepto de ortogonalidad y los problemas de distancia mínima. Comencemos por el siguiente teorema que mejora al T 1.5.2 en el sentido de que garantiza la existencia del vector minimizante.

T 5. (teorema clásico sobre la proyección) Sea H un espacio de Hilbert y M un subespacio vectorial de H . Entonces para cada vector $x \in H$, existe un único vector $m_0 \in M$ tal que

$$\|x - m_0\| \leq \|x - m\| \quad \text{para todo } m \in M.$$

Una condición necesaria y suficiente para que $m_0 \in M$ sea el vector minimizante es que $x - m_0$ sea ortogonal a M .

Demostración. La unicidad y la ortogonalidad han sido ya establecidas en el Teorema 1.5.2. Luego, se necesita solamente establecer la existencia del vector minimizante.

Si $x \in M$, entonces $m_0 = x$ y se obtiene el resultado deseado. Supongamos que $x \notin M$ y sea $\delta = \inf_{m \in M} \|x - m\|$. Debemos hallar un elemento $m_0 \in M$ tal que $\|x - m_0\| = \delta$. Con este propósito, tomemos una sucesión (m_k) de vectores de M tal que $\|x - m_k\| \xrightarrow{k \rightarrow \infty} \delta$.

Por la ley del paralelogramo

$$\begin{aligned} \|(m_j - x) + (x - m_i)\|^2 + \|(m_j - x) - (x - m_i)\|^2 &= \\ &= 2\|m_j - x\|^2 + 2\|x - m_i\|^2 \end{aligned}$$

Reordenando convenientemente, obtenemos

$$\|m_j - m_i\|^2 = 2\|m_j - x\|^2 + 2\|x - m_i\|^2 - 4\left\|x - \frac{m_i + m_j}{2}\right\|^2.$$

Para todo par i, j de índices, el vector $(m_i + m_j)/2$ pertenece a M ya que M es un *subespacio vectorial*. Entonces, por la definición de δ , $\|x - (m_i + m_j)/2\| \geq \delta$ y obtenemos

$$\|m_j - m_i\|^2 \leq 2\|m_j - x\|^2 + 2\|x - m_i\|^2 - 4\delta^2.$$

Como $\|m_i - x\|^2 \rightarrow \delta^2$ cuando $i \rightarrow \infty$, concluimos que

$$\|m_j - m_i\|^2 \rightarrow 0 \quad \text{cuando } i, j \rightarrow \infty.$$

Luego, la sucesión (m_k) es de Cauchy en M y como M es un subespacio cerrado del espacio completo H , de aquí se obtendrá la existencia de un límite m_0 en M . Por continuidad de la norma se prueba que $\|x - m_0\| = \delta$. $\square$

Observación. Notemos que durante la demostración de la existencia del vector minimizante en el teorema anterior no hemos hecho referencia explícita al producto escalar, sino solamente se utiliza la norma asociada. Sin embargo, implícitamente hemos utilizado el producto escalar al hacer uso de la ley del paralelogramo.

Existe una amplia clase de espacios de Banach, donde también tiene sentido el teorema anterior, aunque el mismo no puede ser extendido a cualquier espacio de Banach.

Como consecuencia del teorema anterior veremos ahora un resultado que justifica el nombre de *complemento ortogonal* para denotar al conjunto de los vectores ortogonales a un conjunto dado. Si el conjunto dado es un subespacio vectorial cerrado de un espacio de Hilbert, entonces su complemento ortogonal contiene suficientes vectores adicionales como para generar a todo el espacio.

T 6. Si M es un subespacio vectorial cerrado de un espacio de Hilbert H , entonces $H = M \oplus M^\perp$ y además $M = M^{\perp\perp}$.

Demostración. Sea $x \in H$. Por el teorema de la proyección (T5), existe un único vector $m_0 \in M$ tal que $\|x - m_0\| \leq \|x - m\|$ para todo $m \in M$ y $n_0 = x - m_0 \in M^\perp$. Luego $x = m_0 + n_0$ con $m_0 \in M$ y $n_0 \in M^\perp$.

Para probar que esta representación es única basta probar que $M \cap M^\perp = \{\Theta\}$ lo cual fue probado en P 1.4.4 v).

Para probar que $M = M^{\perp\perp}$ es solamente necesario probar que $M^{\perp\perp} \subset M$ ya que por P 1.4.4 vi) se tiene la inclusión en sentido contrario. Sea $x \in M^{\perp\perp}$. Por la primera parte de este teorema $x = m + n$ donde $m \in M$ y $n \in M^\perp$. Como ambos x y m pertenecen a $M^{\perp\perp}$, tendremos que $x - m \in M^{\perp\perp}$ y por lo tanto $n \in M^{\perp\perp}$. Pero $n \in M^\perp$ y por ello $n \perp n$ lo cual implica que $n = \Theta$. Luego $x = m \in M$ y $M^{\perp\perp} \subset M$. $\square$

Teniendo en cuenta el resultado anterior se introduce la definición siguiente:

D 3. (proyección ortogonal) Sea M un subespacio cerrado de un

espacio de Hilbert y $x \in M$. El vector único $m_0 \in M$ tal que $x - m_0 \in M^\perp$ se llama proyección ortogonal de x sobre M .

Del Teorema 6 se desprende el siguiente corolario que completa la Proposición 1.4.4 para los espacios de Hilbert.

C 1. Sea S un subconjunto arbitrario de un espacio de Hilbert. Entonces se cumple:

- i) $S^{\perp\perp} = \overline{[S]}$ (donde por $\overline{[S]}$ se denota la clausura del subespacio vectorial generado por S);
- ii) $S^\perp = \{\Theta\} \Leftrightarrow S$ es una familia total en H .

Demostración. i) Por P 1.4.4 iii) y iv) se tiene que

$$S^{\perp\perp} = [S]^{\perp\perp} = \overline{[S]}^{\perp\perp}$$

y por T 6:

$$\overline{[S]}^{\perp\perp} = \overline{[S]}$$

de donde

$$S^{\perp\perp} = \overline{[S]}.$$

- ii) La implicación $\Leftarrow$ fué ya obtenida en P 1.4.4 ix) para un espacio euclideo cualquiera. Para probar la implicación en sentido contrario notemos que si $S^\perp = \{\Theta\}$, entonces

$$S^{\perp\perp} = \{\Theta\}^\perp = H.$$

Pero por T 6, $S^{\perp\perp} = \overline{[S]}$, luego $[S] = H$, y por lo tanto S es una familia total en H . $\square$

En el Ejemplo 1.4.3 vimos una familia total de vectores en un espacio de Hilbert separable que no admite una sucesión biortogonal. Como una

consecuencia de T6 veremos una caracterización de las sucesiones que admiten sucesiones biortogonales.

T 7. La sucesión (f_n) de vectores del espacio de Hilbert H , tiene una sucesión biortogonal si y sólo si

$$f_j \notin M_j, \quad \text{para cada } j \in \mathbb{N},$$

donde M_j es el espacio vectorial cerrado generado por los vectores

$$f_1, f_2, \dots, f_{j-1}, f_{j+1}, \dots$$

Demostración. Denotemos por M al espacio vectorial cerrado generado por la sucesión (f_n) . Si para cualquier número natural j , se tiene que

$$M_j \neq M,$$

entonces, tomando en el complemento ortogonal de M_j en M , un vector g_j normado por la condición

$$\langle f_j, g_j \rangle = 1,$$

llegamos a que

$$\langle f_j, g_k \rangle = \delta_{kj} \quad \text{delta de Kronecker} .$$

De esta forma obtenemos una sucesión (g_n) biortogonal a la sucesión (f_n) .

Inversamente, si existe una sucesión (g_n) biortogonal a (f_n) , entonces f_j no puede pertenecer a M_j para ningún j , ya que en este caso él sería ortogonal al vector g_j el cual es ortogonal a todos los vectores f_k con $k \neq j$. $\square$

Veamos finalmente cómo los resultados sobre distancia mínima pueden ser también obtenidos sustituyendo las variedades lineales por conjuntos cerrados convexos. Se dice que un subconjunto K de un espacio

vectorial es *convexo* si para cada par de elementos x, y en K , el segmento

$$[x, y] = \{\lambda x + (1 - \lambda)y; 0 \leq \lambda \leq 1\}.$$

que une a x con y , está totalmente contenido en K . El teorema que presentaremos a continuación es una generalización del teorema de la proyección.

T 8. (distancia mínima a un cerrado convexo) Sea x un vector en un espacio de Hilbert H y sea K un subconjunto cerrado convexo de H . Entonces existe un único vector $k_0 \in K$ tal que

$$\|x - k_0\| \leq \|x - k\|$$

para todo $k \in K$. Más aún, una condición necesaria y suficiente para que k_0 sea el único vector minimizante es que $\langle x - k_0, k - k_0 \rangle \leq 0$ para todo $k \in K$.

Demostración. Para probar la existencia, sea $\{k_i\}$ una sucesión en K tal que

$$\|x - k_i\| \rightarrow \delta = \inf_{k \in K} \|x - k\|.$$

Por la ley del paralelogramo.

$$\|k_i - k_j\|^2 = 2\|k_i - x\|^2 + 2\|k_j - x\|^2 - 4\|x - \frac{k_i + k_j}{2}\|^2.$$

Por la convexidad de K , $(k_i + k_j)/2$ está en K ; de aquí que

$$\|x - \frac{k_i + k_j}{2}\| \geq \delta$$

y por lo tanto

$$\|k_i - k_j\|^2 \leq 2\|k_i - x\|^2 + 2\|k_j - x\|^2 - 4\delta^2 \rightarrow 0.$$

Como la sucesión (k_i) es de Cauchy y K es cerrado, convergerá a un elemento $k_0 \in K$. Por continuidad de la norma, $\|x - k_0\| = \delta$.

Para probar la unicidad, supongamos que existe otro $k_1 \in K$ tal que $\|k_1 - x\| = \delta$. La sucesión

$$k_n = \begin{cases} k_0, & n \text{ par}, \\ k_1, & n \text{ impar}, \end{cases}$$

cumple que $\|x - k_n\| \rightarrow \delta$, luego por el argumento anterior $\{k_n\}$ es de Cauchy y por lo tanto converge. Obviamente esto puede ocurrir en este caso tan sólo si $k_1 = k_0$.

Probaremos ahora que si k_0 es el único vector minimizante, entonces

$$\langle x - k_0, k - k_0 \rangle \leq 0$$

para todo $k \in K$. Supongamos por el contrario que existe un vector $k_1 \in K$ tal que $\langle x - k_0, k_1 - k_0 \rangle = \varepsilon > 0$. Consideremos los vectores $k_\alpha = (1 - \alpha)k_0 + \alpha k_1$; con $0 \leq \alpha \leq 1$. Como K es convexo, cada $k_\alpha \in K$. Además

$$\begin{aligned} \|x - k_\alpha\|^2 &= \|(1 - \alpha)(x - k_0) + \alpha(x - k_1)\|^2 = \\ &= (1 - \alpha)^2 \|x - k_0\|^2 + 2\alpha(1 - \alpha) \langle x - k_0, x - k_1 \rangle \\ &\quad + \alpha^2 \|x - k_1\|^2 \end{aligned}$$

La expresión $\|x - k_\alpha\|^2$ es una función diferenciable de α con derivada en $\alpha = 0$ igual a

$$\begin{aligned} \frac{d}{d\alpha} \|x - k_\alpha\|^2|_{\alpha=0} &= -2\|x - k_0\|^2 + 2\langle x - k_0, x - k_1 \rangle \\ &= -2\langle x - k_0, k_1 - k_0 \rangle = -2\varepsilon < 0. \end{aligned}$$

Luego para algún α positivo suficientemente pequeño, $\|x - k_\alpha\| < \|x - k_0\|$ lo cual contradice la propiedad minimizante de k_0 . De aquí concluimos que no puede existir un tal k_1 .

Inversamente, supongamos que $k_0 \in K$ es tal que $\langle x - k_0, k - k_0 \rangle \leq 0$ para todo $k \in K$. Entonces para cualquier $k \in K, k \neq k_0$, tendremos

$$\begin{aligned} \|x - k\|^2 &= \|x - k_0 + k_0 - k\|^2 = \|x - k_0\|^2 \\ &\quad + 2\langle x - k_0, k_0 - k \rangle + \|k_0 - k\|^2 > \|x - k_0\|^2 \end{aligned}$$

y por lo tanto, k_0 es el único vector minimizante. $\square$

§ 1.7 Derivadas en sentido de Soboliev y espacios de Soboliev.

Dedicaremos este epígrafe al estudio de los fundamentos de la teoría de las funciones generalizadas y de los espacios de Soboliev, que serán utilizadas en el resto del texto para el desarrollo de la teoría de los operadores diferenciales.

Comencemos por introducir algunas notaciones que serán utilizadas en lo sucesivo.

Para cada subconjunto Ω *abierto y acotado* del espacio euclideo $\mathbb{R}^n$, consideraremos las siguiente clases de funciones.

$C(\overline{\Omega})$ - funciones complejas y continuas definidas sobre $\overline{\Omega}$

$C^m(\overline{\Omega})$ - funciones complejas, con derivadas continuas en Ω hasta el orden m inclusive, todas las cuales pueden ser prolongadas continuamente a $\overline{\Omega}$ ($1 \leq m < \infty$).

Notemos que $C^\infty(\overline{\Omega}) = \bigcap_{m=1}^{\infty} C^m(\overline{\Omega})$.

Recordemos que el *soporte de una función* f definida sobre Ω y con valores en $\mathbb{C}$ es el complemento en Ω del mayor subconjunto abierto de

$\mathbb{C}$ donde f se anula. Denotaremos al soporte de la función f mediante $\text{supp } f$.

$$C_0^\infty(\Omega) = \{f \in C^\infty(\overline{\Omega}) : \text{supp } f \text{ es compacto y está contenido en } \Omega\}.$$

Mediante $L_p(\Omega)$ denotaremos la clase de todas las funciones complejas medibles $u(x)$ definidas sobre Ω y tales que

$$\int_{\Omega} |u(x)|^p dx < \infty.$$

Es conocido que la aplicación

$$\|u\|_p = \left(\int_{\Omega} |u(x)|^p dx \right)^{1/p} \quad (83)$$

es una norma en $L_p(\Omega)$ que provee a este conjunto de funciones de una estructura de espacio de Banach. Utilizaremos la notación $\|\cdot\|_p$ sin hacer referencia al conjunto Ω cuando no haya lugar a confusión. La desigualdad triangular para $\|\cdot\|_p$, es decir

$$\|u + v\|_p \leq \|u\|_p + \|v\|_p, \quad (84)$$

es la llamada *desigualdad de Minkowsky*.

Para $\|\cdot\|_p$ se cumple también otra desigualdad muy importante llamada *desigualdad de Hölder*.

Si $p, q > 1$ son tales que $\frac{1}{p} + \frac{1}{q} = 1$, y $u \in L_p(\Omega), v \in L_q(\Omega)$, entonces:

$$\int_{\Omega} |u(x)v(x)| dx \leq \|u\|_p^p \|v\|_q^p \quad (85)$$

≤ \|u\|_p \|v\|_q

Una función medible $u(x)$ definida sobre Ω se llama *localmente sumable* en Ω si para cualquier abierto Ω_1 tal que $\overline{\Omega}_1 \subset \Omega$, se cumple

que

$$\int_{\Omega_1} |u(x)| dx < \infty.$$

En general se dice que $u(x)$ está *localmente en* $L_p(\Omega)$ (en símbolos, $u(x) \in L_{p,loc}(\Omega)$), si para cualquier abierto Ω_1 tal que $\overline{\Omega_1} \subset \Omega$, se cumple que la restricción de u a Ω_1 pertenece a $L_p(\Omega_1)$.

Diremos que la región Ω pertenece a la clase C^k ($k \geq 1$), si para cualquier punto $x_0 \in \partial\Omega$ existe una vecindad $V(x_0)$ de x_0 y un homeomorfismo

$$f: \mathbb{R}^{n-1} \longrightarrow V(x_0) \cap \partial\Omega \quad (86)$$

tal que $f \in C^k(\mathbb{R}^{n-1})$.

Diremos que $\Omega \subset \mathbb{R}^n$ es un abierto con frontera *suave* o *regular*, si $\partial\Omega$ es una variedad de clase C^∞ y dimensión $n-1$, y además Ω localmente se encuentra siempre a un mismo lado de $\partial\Omega$.

Para los abiertos Ω de clase C^k , se cumple la *fórmula de Ostrogradski-Gauss*, que se demuestra en los cursos básicos de Análisis Matemático y que formularemos de la forma siguiente. Sean $u_j(x)$, $j = 1, 2, \dots, n$, funciones pertenecientes a la clase $C^1(\overline{\Omega})$, $\Omega \in C^k$ ($k \geq 1$) y $\nu = (\nu_1, \dots, \nu_n)$ el vector normal exterior unitario a $\partial\Omega$. Entonces $n\eta = (n\eta_1, \dots, n\eta_n)$

$$\int_{\Omega} \sum_{j=1}^n \frac{\partial u_j}{\partial x_j} dx = \int_{\partial\Omega} \sum_{j=1}^n u_j \nu_j ds \quad (87)$$

donde ds denota el elemento de área en $\partial\Omega$. Esta fórmula también se escribe en la forma reducida

$$\int_{\Omega} \operatorname{div} u dx = \int_{\partial\Omega} u \cdot \nu ds \quad (88)$$

donde $u = (u_1, \dots, u_n)$ y el símbolo $u \cdot \nu$ representa el producto escalar de vectores en $\mathbb{R}^n$.

$$u_j = \begin{cases} 0 & j \neq r \\ uv & j = r \end{cases}$$

Poniendo $u_j = 0$ para $j \neq r$ y $u_r = uv$ en la fórmula de Ostrogradski-Gauss, se obtiene la siguiente *fórmula de integración por partes*

$$\int_{\Omega} u \frac{\partial v}{\partial x_r} dx = - \int_{\Omega} v \frac{\partial u}{\partial x_r} dx + \int_{\partial\Omega} uv \nu_r ds \quad (89)$$

Introduzcamos las siguientes notaciones que serán utilizadas en lo sucesivo. Sea $\alpha = (\alpha_1, \dots, \alpha_n)$ — un multíndice, donde $\alpha_j, j = 1, \dots, n$, son números enteros no negativos, $|\alpha| = \alpha_1 + \dots + \alpha_n$. Denotaremos

$$Dx_j = \frac{\partial}{\partial x_j}, j = 1, \dots, n; D_x^\alpha = D_{x_1}^{\alpha_1} \dots D_{x_n}^{\alpha_n} = \frac{\partial^{|\alpha|}}{\partial x_1^{\alpha_1} \dots \partial x_n^{\alpha_n}}$$

donde $|\alpha|$ es el orden de la derivada D .

Para obtener algunos teoremas importantes sobre separación de conjuntos y aproximación de funciones, por funciones diferenciables en lugar de continuas, es necesario introducir la definición siguiente:

D 1. (núcleo de promediación, promedio de una función) Se llama núcleo de promediación a toda función $w_h(x)$, definida para $x \in \mathbb{R}^n, h \in \mathbb{R}, h > 0$ y que satisfaga las condiciones:

- i) $w_h(x) \in C^\infty(\mathbb{R}^n)$ para cualquier $h > 0$;
- ii) $w_h(x) = 0$ para $|x| > h$;
- iii) $w_h(x) \geq 0$ en $\mathbb{R}^n$, para todo $h > 0$;
- iv) $\int_{\mathbb{R}^n} w_h(x) dx = 1$, para todo $h > 0$.

Para toda función $u(x)$ localmente sumable en $\mathbb{R}^n$, la función

$$u_h(x) = \int_{\mathbb{R}^n} w_h(x-y) u(y) dy \quad (90)$$

se llama *promedio de la función* $u(x)$ con radio de promediación h .

Ejemplo 1. El ejemplo clásico de núcleo de promediación es la función

$$w_h(x) = \begin{cases} C \exp \left\{ \frac{h^2}{|x|^2 - h^2} \right\}, & |x| < h \\ 0, & |x| \geq h \end{cases}$$

donde

$$C = h^{-n} \left[\int_{|x| < 1} \exp \left\{ \frac{1}{|y|^2 - 1} \right\} dy \right]^{-1}.$$

El lector puede verificar sin dificultad que w_h satisface las propiedades i) – iv) de la Definición 1.

T 1. Sean Ω una región acotada de $\mathbb{R}^n$ y Ω_1 un abierto tal que $\overline{\Omega}_1 \subset \Omega$.

- i) Si $u(x) \in L_p(\Omega), p \geq 1$, se anula fuera de Ω_1 , entonces $u_h(x) \in C_0^\infty(\Omega)$ para todo $h < \delta$, donde δ es la distancia de Ω_1 a $\partial\Omega$.
- ii) Si $u(x) \in C(\overline{\Omega})$ y se anula sobre $\partial\Omega$, entonces $u_h(x)$ converge uniformemente a $u(x)$ sobre Ω , cuando $h \rightarrow 0$.
- iii) Si $u(x) \in L_p(\Omega), p \geq 1$, entonces

$$\|u_h\|_p \leq \|u\|_p \quad \text{y} \quad \|u_h - u\|_p \rightarrow 0 \quad \text{cuando } h \rightarrow 0.$$

Demostración. i) Por las propiedades de w_h se tiene que

$$u_h(x) = \int_{|x-y| < h} w_h(x-y) u(y) dy.$$

Pero si $x \in \Omega$ es tal que su distancia a $\partial\Omega$ es menor que $\delta - h$, entonces $u(y) = 0$ para $|x - y| < h$ y, por consiguiente, la expresión

subintegral se anula. Por ello para tales puntos x y $h < \delta$ se cumple que $u_h(x) = 0$. La diferenciabilidad de u_h se deduce inmediatamente de los teoremas sobre diferenciación bajo el signo integral y del hecho que $w_h \in C^\infty(\mathbb{R}^n)$.

ii) Sea $u(x) \in C(\overline{\Omega})$, $u = 0$ en $\mathbb{R}^n \setminus \Omega$ entonces

$$\begin{aligned} |u_h(x) \overset{\text{es nula}}{u(x)}| &= \left| \int_{\mathbb{R}^n} w_h(x-y)[u(y) - u(x)]dy \right| \\ &\leq \sup_{|x-y| \leq h} |u(y) - u(x)| \end{aligned}$$

y la parte derecha de esta última desigualdad tiende a cero cuando $h \rightarrow 0$, uniformemente en x , en virtud de la continuidad uniforme de la función u en $\overline{\Omega}$.

iii) Sea $u(x) \in L_p(\Omega)$. Pongamos $u(x) = 0$ fuera de Ω . Entonces para $p > 1$,

$$\begin{aligned} |u_h(x)| &\leq \int_{\mathbb{R}^n} |w_h(x-y)u(y)|dy \\ &\leq \int_{\mathbb{R}^n} |w_h(x-y)|^{\frac{1}{p} + \frac{1}{q}} |u(y)|dy, \quad \frac{1}{p} + \frac{1}{q} = 1 \end{aligned}$$

Aplicando la desigualdad de Hölder (85) a esta última integral, obtendremos

$$\begin{aligned} |u_h(x)| &\leq \left(\int_{\mathbb{R}^n} w_h(x-y)dy \right)^{1/q} \left(\int_{\mathbb{R}^n} w_h(x-y)|u(y)|^p dy \right)^{1/p} \\ &= \left(\int_{\mathbb{R}^n} w_h(x-y)|u(y)|^p dy \right)^{1/p}, \end{aligned}$$

ya que

$$\int_{\mathbb{R}^n} w_h(x-y)dy = \int_{\mathbb{R}^n} w(\theta)d\theta = 1.$$

Por eso

$$\begin{aligned} \|u_h\|_p &\leq \left[\int_{\mathbb{R}^n} \left(\int_{\mathbb{R}^n} w_h(x-y)|u(y)|^p dy \right) dx \right]^{1/p} \\ &= \left[\int_{\mathbb{R}^n} |u(y)|^p \int_{\mathbb{R}^n} w_h(x-y) dx dy \right]^{1/p} = \|u\|_p. \end{aligned}$$

Si $p = 1$, entonces

$$\begin{aligned} |u_h(x)| &\leq \int_{\mathbb{R}^n} w_h(x-y)|u(y)|dy \\ \|u_h\|_1 &= \int_{\mathbb{R}^n} |u_h(x)|dx \leq \int_{\mathbb{R}^n} |u(y)| \int_{\mathbb{R}^n} w_h(x-y) dx dy \\ &= \|u\|_1. \end{aligned}$$

Es conocido que para cada función $u(x) \in L_p(\Omega)$ y cualquier $\varepsilon > 0$, se puede hallar una función continua $v(x)$, en Ω tal que $\|u-v\|_p < \varepsilon$ y $v(x)$ igual a cero en cierta vecindad de $\partial\Omega$. Pero como $\|u-v\|_p < \varepsilon$ y, de acuerdo a lo demostrado más arriba $\|u_h-v_h\|_p = \|(u-v)_h\|_p \leq \|u-v\|_p$, entonces si tomamos h_0 suficientemente pequeño de manera que

$$\|v-v_h\|_p \leq (\text{mes } \Omega)^{1/p} \|v-v_h\|_\infty < \varepsilon$$

para $h < h_0$, tendremos que $\|u_h-u\|_p < 3\varepsilon$.

La existencia de h_0 está asegurada por la afirmación ii) del teorema. $\square$

Introduzcamos ahora la noción de transformada de Fourier para funciones de $L_1(\mathbb{R}^n)$ y $L_2(\mathbb{R}^n)$, la cual juega un papel fundamental en la

teoría de ecuaciones en derivadas parciales y es ampliamente utilizada en las aplicaciones. Para estudiar las propiedades de la transformada de Fourier, definiremos el siguiente subconjunto denso en $L_p(\mathbb{R}^n)$, $p \geq 1$:

$$S(\mathbb{R}^n) = \left\{ \varphi \in C^\infty(\mathbb{R}^n) : \sup_x \{(1 + \|x\|^p) |D^\alpha \varphi(x)|\} < \infty \right\}$$

para todo múltiplo α , y todo número entero $p > 0$, es decir, existe una constante $C_{\alpha,p}$ tal que $(1 + |x|^p) |D^\alpha \varphi(x)| \leq C_{\alpha,p}$

D 2. (transformada de Fourier en $L_1(\mathbb{R}^n)$) Se llama *transformada de Fourier* de la función $f \in L_1(\mathbb{R}^n)$ a la función

$$\mathcal{F}[f](\xi) = \hat{f}(\xi) = \int_{\mathbb{R}^n} f(x) e^{i\xi \cdot x} dx \quad (91)$$

donde $\xi = (\xi_1, \dots, \xi_n)$, $x = (x_1, \dots, x_n)$, y

$$\xi \cdot x = \xi_1 x_1 + \dots + \xi_n x_n.$$

T 2. Si $\varphi(x) \in S(\mathbb{R}^n)$, entonces $\hat{\varphi}(\xi) \in S(\mathbb{R}^n)$. Además para $\varphi \in S(\mathbb{R}^n)$ se tiene

$$\begin{aligned} \widehat{D^\alpha \varphi} &= (-i\xi)^\alpha \hat{\varphi}(\xi) \\ \widehat{x^\alpha \varphi} &= (-i)^{|\alpha|} D^\alpha \hat{\varphi}(\xi), \end{aligned} \quad (92)$$

donde $x^\alpha = x_1^{\alpha_1} \dots x_n^{\alpha_n}$.

Demostración: Diferenciando bajo el signo de integral en (91), obtenemos

$$D^\alpha \hat{\varphi}(\xi) = \int_{\mathbb{R}^n} (ix)^\alpha \varphi(x) e^{i\xi \cdot x} dx \quad (93)$$

Esta diferenciación es posible ya que la integral (93) converge uniformemente por ser $\varphi \in S(\mathbb{R}^n)$. De aquí, se tiene que $\hat{\varphi}(\xi) \in C^\infty(\mathbb{R}^n)$ y

además,

$$\widehat{x^\alpha \varphi}(x) = (-i)^{|\alpha|} D^\alpha \hat{\varphi}(\xi).$$

Integrando por partes tenemos

$$\begin{aligned} \xi^\beta D^\alpha \hat{\varphi}(\xi) &= \int_{\mathbb{R}^n} (ix)^\alpha \xi^\beta \varphi(x) e^{i\xi \cdot x} dx \\ &= \int_{\mathbb{R}^n} D_x^\beta ((ix)^\alpha \varphi) i^{|\beta|} e^{i\xi \cdot x} dx. \end{aligned} \quad (94)$$

Como $D^\beta((ix)^\alpha \varphi) \in L_1(\mathbb{R}^n)$, entonces $|\xi|^\beta |D^\alpha \hat{\varphi}(\xi)|$ está acotada en $\mathbb{R}^n$, cualesquiera sean los multíndices α y β .

Por consiguiente $\hat{\varphi} \in S(\mathbb{R}^n)$. Poniendo $\alpha = 0$ en (94), obtenemos

$$\widehat{D^\beta \varphi} = (-i)^{|\beta|} \xi^\beta \hat{\varphi}(\xi)$$

con lo cual culmina la demostración del teorema. $\square$

D 3. (transformada de Fourier inversa) Se llama *transformada de Fourier inversa* de la función $f(\xi) \in L_1(\mathbb{R}^n)$ a la función $\mathcal{F}^{-1}[f](x)$ definida por la igualdad

$$\mathcal{F}^{-1}[f](x) = (2\pi)^{-n} \int_{\mathbb{R}^n} f(\xi) e^{-i\xi \cdot x} d\xi \quad (95)$$

T 3. (fórmula de inversión en $S(\mathbb{R}^n)$) Si $\varphi \in S(\mathbb{R}^n)$, entonces tiene lugar la siguiente fórmula de inversión de la transformada de Fourier:

$$\varphi(x) = (2\pi)^{-n} \int_{\mathbb{R}^n} \hat{\varphi}(\xi) e^{-i\xi \cdot x} d\xi, \quad (96)$$

es decir,

$$\varphi(x) = \mathcal{F}^{-1}[\mathcal{F}[\varphi]].$$

Demostración. Calculemos la integral iterada

$$\int_{\mathbb{R}^n} e^{-i\xi \cdot x} \left\{ \int_{\mathbb{R}^n} e^{i\xi \cdot y} \varphi(y) dy \right\} d\xi = (2\pi)^n \mathcal{F}^{-1}[\mathcal{F}[\varphi]].$$

Para ello examinemos la integral

$$\begin{aligned} I &= \int_{\mathbb{R}^n} \Psi_\varepsilon(\xi) e^{-i\xi \cdot x} \left\{ \int_{\mathbb{R}^n} e^{i\xi \cdot y} \varphi(y) dy \right\} d\xi \\ &= \int_{\mathbb{R}^n} \Psi_\varepsilon(\xi) \hat{\varphi}(\xi) e^{-i\xi \cdot x} d\xi \end{aligned} \quad (97)$$

donde $\varepsilon = \text{const} > 0$, $\Psi_\varepsilon(\xi) = \Psi(\varepsilon \xi) \in S(\mathbb{R}^n)$. Del hecho que la integral doble, correspondiente a la integral iterada (97), converge absolutamente, se tiene por el Teorema de Fubini que en (97) se puede cambiar el orden de integración.

Entonces tenemos

$$\begin{aligned} I &= \int_{\mathbb{R}^n} \varphi(y) \left\{ \int_{\mathbb{R}^n} \Psi_\varepsilon(\xi) e^{i\xi \cdot (y-x)} d\xi \right\} dy \\ &= \int_{\mathbb{R}^n} \varphi(y) \hat{\Psi}_\varepsilon(y-x) dy = \int_{\mathbb{R}^n} \varphi(x+\theta) \hat{\Psi}_\varepsilon(\theta) d\theta. \end{aligned}$$

Es fácil ver, que $\hat{\Psi}_\varepsilon(x) = \varepsilon^{-n} \hat{\Psi}(\frac{x}{\varepsilon})$, pues

$$\hat{\Psi}_\varepsilon(\xi) = \int_{\mathbb{R}^n} \Psi_\varepsilon(x) e^{ix \cdot \xi} dx = \int_{\mathbb{R}^n} \Psi(\varepsilon x) e^{ix \cdot \xi} dx$$

$$= \int_{\mathbb{R}^n} \Psi(y) e^{iy \cdot \frac{\xi}{\varepsilon}} \varepsilon^{-n} dy = \varepsilon^{-n} \hat{\Psi}\left(\frac{\xi}{\varepsilon}\right).$$

Por eso

$$I = \int_{\mathbb{R}^n} \varphi(\theta + x) \varepsilon^{-n} \hat{\Psi}\left(\frac{\theta}{\varepsilon}\right) d\theta = \int_{\mathbb{R}^n} \hat{\Psi}(s) \varphi(\varepsilon s + x) ds.$$

De esta forma, tenemos

$$I = \int_{\mathbb{R}^n} \Psi_\varepsilon(\xi) \hat{\varphi}(\xi) e^{-ix \cdot \xi} d\xi = \int_{\mathbb{R}^n} \hat{\Psi}(s) \varphi(\varepsilon s + x) ds \quad (98)$$

Pasando al límite en la igualdad (98), cuando $\varepsilon \rightarrow 0$, obtendremos

$$\varphi(x) \int_{\mathbb{R}^n} \hat{\Psi}(s) ds = \Psi(0) \int_{\mathbb{R}^n} e^{-ix \cdot \xi} \hat{\varphi}(\xi) d\xi \quad (99)$$

Pero el lector puede encontrar en cualquier tabla de transformadas de Fourier, que para $\Psi(x) = e^{-|x|^2}$ se tiene

$$\hat{\Psi}(\xi) = \int_{\mathbb{R}^n} e^{-|x|^2} e^{ix \cdot \xi} dx = (\sqrt{\pi})^n e^{-\frac{|\xi|^2}{4}}$$

como

$$\begin{aligned} \int_{\mathbb{R}^n} \hat{\Psi}(\xi) d\xi &= (\sqrt{\pi})^n \int_{\mathbb{R}^n} e^{-\frac{|\xi|^2}{4}} d\xi = (\sqrt{\pi})^n \prod_{j=1}^n \int e^{-\frac{|\xi_j|^2}{4}} d\xi_j \\ &= (2\pi)^n; \end{aligned}$$

entonces de la relación (99) obtenemos la igualdad (96). $\square$

Enunciaremos a continuación un importante resultado sobre la transformada de Fourier en $L_1(\mathbb{R}^n)$ cuya demostración puede hallarse en el

ligro de A.N. Kolmogorov y S.V. Fomin, *Elementos de la teoría de funciones y el análisis funcional* [2].

T 4. La transformada de Fourier es una aplicación inyectiva y continua de $L_1(\mathbb{R}^n)$ en el espacio $C_0(\mathbb{R}^n)$, de las funciones complejas y continuas $f(x)$, definidas sobre $\mathbb{R}^n$ y que tienden a cero cuando $|x| \rightarrow \infty$, provisto de la norma de la convergencia uniforme. $\square$

D 4. (producto convolución en $L_1(\mathbb{R}^n)$) Se llama *convolución* de las funciones f y g de $L_1(\mathbb{R}^n)$ a la función $f * g$, definida por la fórmula:

$$f * g(x) = \int_{\mathbb{R}^n} f(x-y)g(y)dy. \quad (100)$$

Notemos que $f * g(x)$ es también una función de $L_1(\mathbb{R}^n)$, ya que por el Teorema de Fubini

$$\begin{aligned} \|f * g\|_1 &= \int_{\mathbb{R}^n} \left| \int_{\mathbb{R}^n} f(x-y)g(y)dy \right| dx \\ &\leq \int_{\mathbb{R}^n} \int_{\mathbb{R}^n} |f(x-y)| |g(y)| dy dx \leq \|f\|_1 \|g\|_1 \end{aligned}$$

Además si f y g pertenecen al espacio $S(\mathbb{R}^n)$, entonces $f * g \in S(\mathbb{R}^n)$. En efecto, en ese caso $f * g \in C^\infty(\mathbb{R}^n)$, ya que la expresión (100) puede ser diferenciada tanto como se quiera bajo el signo integral:

$$D^\alpha(f * g) = \int_{\mathbb{R}^n} D_x^\alpha f(x-y)g(y)dy.$$

Por otra parte

$$|x^\beta D^\alpha(f * g)| = \left| \int_{\mathbb{R}^n} g(y) [(x-y) + y]^\beta D_x^\alpha f(x-y) dy \right| \leq C_{\alpha, \beta},$$

ya que $f, g \in S(\mathbb{R}^n)$.

T 5. Sean $\varphi(x)$ y $\Psi(x)$ funciones en $S(\mathbb{R}^n)$. Entonces:

$$\int_{\mathbb{R}^n} \hat{\varphi} \Psi dx = \int_{\mathbb{R}^n} \varphi \hat{\Psi} dx; \quad (101)$$

$$\int_{\mathbb{R}^n} \varphi \overline{\Psi} dx = (2\pi)^{-n} \int_{\mathbb{R}^n} \hat{\varphi} \overline{\hat{\Psi}}; \quad (102)$$

(identidad de Parseval)

$$\widehat{\varphi * \Psi} = \hat{\varphi} \hat{\Psi}, \quad (103)$$

$$\widehat{\varphi \Psi} = (2\pi)^{-n} \hat{\varphi} * \hat{\Psi} \quad (104)$$

Demostración. Demostremos la igualdad (101). Por el Teorema de Fubini tenemos

$$\begin{aligned} \int_{\mathbb{R}^n} \hat{\varphi}(x) \Psi(x) dx &= \int_{\mathbb{R}^n} \left(\int_{\mathbb{R}^n} \varphi(y) e^{ix \cdot y} dy \right) \Psi(x) dx \\ &= \int_{\mathbb{R}^n} (\varphi(y) \int_{\mathbb{R}^n} \Psi(x) e^{ix \cdot y} dx) dy = \int_{\mathbb{R}^n} \varphi(y) \hat{\Psi}(y) dy. \end{aligned}$$

Para la demostración de (102) pongamos en la igualdad (101) la función $h = (2\pi)^{-n} \overline{\hat{\Psi}}$ en lugar de Ψ . Entonces tenemos

$$(2\pi)^{-n} \int_{\mathbb{R}^n} \hat{\varphi} \overline{\hat{\Psi}} dx = \int_{\mathbb{R}^n} \varphi \hat{h} dx.$$

Utilizando la fórmula de inversión de la transformada de Fourier, obtenemos

$$\begin{aligned}\hat{h}(x) &= (2\pi)^{-n} \int_{\mathbb{R}^n} \overline{\Psi}(\xi) e^{ix \cdot \xi} d\xi = \overline{(2\pi)^{-n} \int_{\mathbb{R}^n} \hat{\Psi}(\xi) e^{-ix \cdot \xi} d\xi} \\ &= \overline{\Psi}(x)\end{aligned}$$

Veamos la demostración de (103). Cambiando el orden de integración de acuerdo al Teorema de Fubini, se tiene

$$\begin{aligned}\widehat{\varphi * \Psi} &= \int_{\mathbb{R}^n} \left(\int_{\mathbb{R}^n} \varphi(y) \Psi(x-y) dy \right) e^{ix \cdot \xi} dx \\ &= \int_{\mathbb{R}^n} \varphi(y) e^{i\xi \cdot y} \left(\int_{\mathbb{R}^n} \Psi(x-y) e^{i(x-y) \cdot \xi} dx \right) dy = \hat{\varphi} \hat{\Psi}\end{aligned}$$

Para la demostración de (104) observemos que

$$\mathcal{F}[\mathcal{F}[\varphi]] = (2\pi)^n \varphi(-x) \quad (105)$$

Por eso

$$\mathcal{F}[\mathcal{F}[\varphi \Psi]] = (2\pi)^n \varphi(-x) \Psi(-x) \quad (106)$$

Pero por la fórmula (103)

$$\begin{aligned}\mathcal{F}[(2\pi)^{-n}(\hat{\varphi} * \hat{\Psi})] &= (2\pi)^{-n} \mathcal{F}[\mathcal{F}[\varphi]] \mathcal{F}[\mathcal{F}[\Psi]] \\ &= (2\pi)^n \varphi(-x) \Psi(-x)\end{aligned} \quad (107)$$

Las igualdades (106) y (107) muestran que

$$\mathcal{F}[\widehat{\varphi \Psi}] = \mathcal{F}[(2\pi)^{-n} \hat{\varphi} * \hat{\Psi}]$$

y, por consiguiente, se cumple la igualdad (104). $\square$

Observación. Sabemos que $S(\mathbb{R}^n)$ es un subconjunto denso en $L_2(\mathbb{R}^n)$. Si consideramos $S(\mathbb{R}^n)$ como subespacio normado de $L_2(\mathbb{R}^n)$, entonces

poniendo $\varphi = \Psi$ en (102) tendremos que $(2\pi)^{-n/2}\mathcal{F}$ es una isometría en $S(\mathbb{R}^n)$. Luego, la transformada de Fourier puede ser extendida de manera única a todo $L_2(\mathbb{R}^n)$, conservándose las fórmulas (101) y (102). Es decir, si $f \in L_2(\mathbb{R}^n)$ y (φ_k) es una sucesión de funciones de $S(\mathbb{R}^n)$ que converge a f en la norma de $L_2(\mathbb{R}^n)$, entonces la sucesión $(\hat{\varphi}_k)$ converge a una cierta función g que no depende de la elección de la sucesión (φ_k) y a la que llamaremos *transformada de Fourier en $L_2(\mathbb{R}^n)$* de la función f : $\hat{f} = g$. Puede probarse que g se obtiene también como el límite en $L_2(\mathbb{R}^n)$ de la sucesión de funciones

$$g_k(x) = \int_{B_k} f(y) e^{ix \cdot y} dy,$$

donde B_k es la bola en $\mathbb{R}^n$ con centro en cero y radio k . Debido a la continuidad de la transformada de Fourier en $L_2(\mathbb{R}^n)$, es obvio que se cumplen las fórmulas (101) y (102).

Introduciremos unos subespacios de $L_2(\Omega)$ que juegan un papel fundamental en la teoría de ecuaciones en derivadas parciales; los llamados *espacios de Soboliev*.

D 5. (espacios de Soboliev) Se llama espacio de Soboliev de orden k ($k \in \mathbb{N}^*$), y se denota por $H^k(\Omega)$, al completamiento del conjunto $C^{(k)}(\bar{\Omega})$, provisto de la norma

$$\|\varphi\|_{2,k} = \left(\sum_{|\alpha| \leq k} \|D^\alpha \varphi\|_2^2 \right)^{1/2} \quad (108)$$

Utilizaremos las notaciones

$$D^0 \varphi = \varphi, \quad H^0(\Omega) = L_2(\Omega).$$

Si una función $u(x)$ pertenece a la clase $C^1(\bar{\Omega})$, entonces por la fórmula de integración por partes (89), tenemos

$$\int_{\Omega} \varphi(x) \frac{\partial u(x)}{\partial x_j} dx = - \int_{\Omega} u(x) \frac{\partial \varphi(x)}{\partial x_j} dx \quad (109)$$

para cualquier función $\varphi \in C_0^\infty(\Omega)$. Si $u(x) \in C^k(\overline{\Omega})$, entonces aplicando la fórmula de integración por partes k veces obtendremos

$$\int_{\mathbb{R}^n} \varphi(x) D^\alpha u(x) dx = (-1)^{|\alpha|} \int_{\Omega} u(x) D^\alpha \varphi(x) dx$$

para cualquier función $\varphi \in C_0^\infty(\Omega)$, $|\alpha| \leq k$. Es natural utilizar esta última igualdad como base para la definición de la derivada en sentido de Soboliev $D^\alpha u$ de la función u , aún cuando la derivada $D^\alpha u$ pueda no existir en sentido usual.

D 6. (derivada en sentido de Soboliev de funciones localmente sumables) La función $v(x)$ localmente sumable en Ω se llama *derivada en sentido de Soboliev* de la función localmente sumable $u(x)$ en Ω , y se denota $v = D^\alpha u$, si para cualquier función $\varphi \in C_0^\infty(\Omega)$ se satisface la igualdad

$$\int_{\Omega} v(x) \varphi(x) dx = (-1)^{|\alpha|} \int_{\Omega} u(x) D^\alpha \varphi(x) dx. \quad (110)$$

Mostremos que la función $v(x)$ que satisface la relación (110) es *única*. En efecto, si suponemos que existen dos funciones $v_1(x)$ y $v_2(x)$, las cuales satisfacen (110) para cualquier $\varphi \in C_0^\infty(\Omega)$, entonces, restando las igualdades correspondientes a v_1 y v_2 , obtendremos para $V = v_1 - v_2$, la igualdad

$$\int_{\Omega} V(x) \varphi(x) dx = 0 \quad (111)$$

para cualquier $\varphi \in C_0^\infty(\Omega)$. Pero no es difícil concluir de los puntos i) y iii) del Teorema 1 que $C_0^\infty(\Omega)$ es denso en $L_2(\Omega)$.

Notemos que (111) significa que V pertenece al complemento ortogonal de $C_0^\infty(\Omega)$ en $L_2(\Omega)$ y por ser $L_2(\Omega)$ un espacio de Hilbert, concluimos que $V = 0$ en $L_2(\Omega)$.

De la definición dada vemos que si $u(x) \in C^k(\overline{\Omega})$, entonces $u(x)$ posee en Ω derivadas en sentido de Soboliev hasta el orden k que coinciden con sus derivadas en sentido usual. Sin embargo, hay funciones para las cuales existe la derivada en sentido de Soboliev sin ser derivables en sentido usual. Observemos además, que a diferencia de la definición clásica de derivada, la derivada en sentido de Soboliev $D^\alpha u$ de orden $|\alpha|$ se define por la igualdad (110) independientemente de las derivadas generalizadas de orden inferior. Veamos algunas otras propiedades elementales de la derivada en sentido de Soboliev.

P 1. Sea $u(x) \in L_{1,\text{loc}}(\Omega)$ y supongamos que existe la derivada en sentido de Soboliev $D^\alpha u \in L_{1,\text{loc}}(\Omega)$. Entonces para cualquier abierto Ω_1 tal que $\overline{\Omega}_1 \subset \Omega$ y cualquier $h < \delta$, donde δ es la distancia entre Ω_1 y $\partial\Omega$, se tiene la igualdad

$$D^\alpha u_h(x) = (D^\alpha u)_h(x), \quad x \in \Omega_1 \quad (112)$$

Demostración. Por la definición de u_h tenemos

$$D^\alpha u_h(x) = \int_{\mathbb{R}^n} D_x^\alpha w_h(x-y)u(y)dy.$$

Pero como $D_x^\alpha w_h(x-y)u(y) = (-1)^{|\alpha|} D_y^\alpha w_h(x-y)$ y $w_h(x-y) \in C_0^\infty(\Omega)$ para $x \in \overline{\Omega}_1$ y $h \leq \delta$, entonces de acuerdo a la definición de derivada en sentido de Soboliev obtenemos:

$$\begin{aligned} D^\alpha u_h(x) &= \int_{\mathbb{R}^n} (-1)^{|\alpha|} D_y^\alpha w_h(x-y)u(y)dy \\ &= \int_{\mathbb{R}^n} w_h(x-y)D_y^\alpha u(y)dy = (D^\alpha u)_h(x). \quad \square \end{aligned}$$

P 2. Sea $u(x) \in L_{1,\text{loc}}(\Omega)$ y supongamos que las derivadas en sentido de Soboliev $D_{x_j} u = 0$ en Ω , $j = 1, 2, \dots, n$. Entonces $u = \text{cte}$ en Ω .

Demostración. Por la Proposición 1 tenemos que en cualquier región Ω_1 , tal que $\overline{\Omega}_1 \subset \Omega$ se cumple $D_{x_j} u_h = (D_{x_j} u)_h$ para $h < \delta$, $j = 1, \dots, n$. Por eso para $h < \delta$, $u_h = C_h = \text{cte}$ en Ω_1 . Por el Teorema 1 iii) se tiene que $u_h \rightarrow u$ en $L_1(\Omega_1)$. Pero las constantes C_h deben converger a una constante C cuando $h \rightarrow 0$, luego $u = C$ en Ω_1 y como Ω_1 es una subregión arbitraria de Ω se tendrá que $u = \text{cte}$ en Ω . $\square$

A continuación veremos cómo pueden ser trasladados los resultados expuestos en el Teorema 2 a funciones de $L_2(\mathbb{R}^n)$.

T 6. Sea $u(x) \in L_2(\mathbb{R}^n)$.

- i) Si existe la derivada en sentido de Soboliev $D^\alpha u$ y pertenece a $L_2(\mathbb{R}^n)$, entonces

$$\widehat{D^\alpha u} = (-i\xi)^\alpha \hat{u}(\xi) \quad (113)$$

- ii) Si existe la derivada en sentido de Soboliev $D^\alpha \hat{u}$ y pertenece a $L_2(\mathbb{R}^n)$, entonces

$$\widehat{x^\alpha u} = (-i)^{|\alpha|} D^\alpha \hat{u}(\xi) \quad (114)$$

Demostración. Probaremos i) y dejaremos ii) de ejercicio al lector. Aplicando el Teorema de Fubini no resulta difícil verificar que

$$(\hat{f})_h = \hat{f}_h \quad (115)$$

para toda $f \in L_2(\mathbb{R}^n)$. Para probar i) supondremos primeramente que $u \in L_2(\mathbb{R}^n)$ tiene soporte compacto. Entonces por (115) y P 1 se tiene que

$$(\widehat{D^\alpha u})_h = \widehat{D^\alpha}(u_h) \quad (116)$$

Pero de T1 i) se tiene que $u_h \in C_0^\infty(\mathbb{R}^n)$ y aplicando la fórmula (92) en (116), obtendremos

$$(\widehat{D^\alpha u})_h = \widehat{D^\alpha} u_h = (-i\xi)^\alpha \hat{u}_h(\xi) = (-i\xi)^\alpha (\hat{u})_h.$$

De donde, la veracidad de (113) se concluye de la siguiente

$$(\widehat{D^\alpha u} - (-i\xi)^\alpha \hat{u})_h = 0.$$

Tomemos ahora $u \in L_2(\mathbb{R}^n)$, no necesariamente de soporte compacto para la cual existe la derivada $D^\alpha u$ en sentido de Soboliev. Obviamente también existirá la derivada en sentido de Soboliev $D^\alpha(u\varphi)$ para todo $\varphi \in C_0^\infty(\mathbb{R}^n)$. Consideremos una sucesión de funciones (φ_n) de $C_0^\infty(\mathbb{R}^n)$, tales que $0 \leq \varphi_N(x) \leq 1$ y $\varphi_N(x) = \begin{cases} 1 & \text{si } |x| \leq N \\ 0 & \text{si } |x| \geq N+1 \end{cases}$.

Si Ponemos $u_N(x) = u(x)\varphi_N(x)$ entonces no es difícil verificar que $D^\alpha u_N$ converge débilmente a $D^\alpha u$, es decir para toda $\varphi \in C_0^\infty(\mathbb{R}^n)$

$$\int_{\mathbb{R}^n} D^\alpha u_N \varphi dx \longrightarrow \int_{\mathbb{R}^n} D^\alpha u \varphi dx. \quad (117)$$

Pero de (117) se concluye que

$$\int_{\mathbb{R}^n} \widehat{D^\alpha} u_N \varphi d\xi \rightarrow \int_{\mathbb{R}^n} \widehat{D^\alpha} u \varphi d\xi, \quad (118)$$

$$\int_{\mathbb{R}^n} (-i\xi)^\alpha \hat{u}_N \varphi d\xi \rightarrow \int_{\mathbb{R}^n} (-i\xi)^\alpha \hat{u} \varphi d\xi \quad (119)$$

para toda $\varphi \in C_0^\infty(\mathbb{R}^n)$.

De la primera parte de la demostración se tiene que

$$\widehat{D^\alpha u_N} = (-i\xi)^\alpha \hat{u}_N \quad (120)$$

para todo número natural N . Luego multiplicando escalarmente por $\varphi \in C_0^\infty(\mathbb{R}^n)$ y haciendo tender $N \rightarrow \infty$, tendremos de (118) y (119) que

$$\int_{\mathbb{R}^n} (\hat{D^\alpha u} - (-i\xi)^\alpha \hat{u}) \varphi d\xi = 0$$

para toda $\varphi \in C_0^\infty(\mathbb{R}^n)$, de donde se concluye (113). $\square$

Observación. Notemos que si se cumple la hipótesis i) del Teorema 6, entonces de (113) se concluye que no sólo $\hat{u}(\xi)$ pertenece a $L_2(\mathbb{R}^n)$ sino también la función $(-i\xi)^\alpha \hat{u}(\xi)$. Recíprocamente, de (113) se deduce que si $u \in L_2(\mathbb{R}^n)$ es tal que $(-i\xi)^\alpha \hat{u}(\xi) \in L_2(\mathbb{R}^n)$, entonces existe la derivada en sentido de Soboliev de orden α de u y $D^\alpha u \in L_2(\mathbb{R}^n)$. Un análisis similar puede realizarse con la fórmula (114).

Puede ser probado el resultado siguiente que se deja de ejercicio al lector.

T 7. Sean $u, v \in L_1(\mathbb{R}^n)$. Si existe derivada en sentido de Soboliev $D^\alpha u$ y pertenece a $L_1(\mathbb{R}^n)$, entonces también existirá la derivada en sentido de Soboliev de orden α del producto convolución $u * v$, la cual pertenecerá a $L_1(\mathbb{R}^n)$. En este caso, se cumple la igualdad

$$D^\alpha(u * v) = (D^\alpha u) * v. \quad (121)$$

Demostración. Ejercicio. $\square$

El teorema siguiente nos muestra una útil caracterización de los espacios $H^k(\Omega)$ mediante la noción de derivada en sentido de Soboliev.

T 8. Para todo $k \in \mathbb{N}$ se tiene que

$$H^k(\Omega) = \left\{ \begin{array}{l} u \in L_2(\Omega) : \text{ existe la derivada en} \\ \text{ sentido de Soboliev} \\ D^\alpha u, |\alpha| \leq k \\ \text{ y pertenece a } L_2(\Omega). \end{array} \right\} \quad (122)$$

Demostración. Denotemos por $W_2^k(\Omega)$ al conjunto descrito entre corchetes. Veamos primeramente que $H^k(\Omega) \subset W_2^k(\Omega)$.

En efecto, si $f \in H^k(\Omega)$, entonces existe una sucesión (φ_n) en $C^k(\overline{\Omega})$ que es de Cauchy en la norma (108) y cuyo límite en la norma de $L_2(\Omega)$ coincide con f . Para cada $n \in \mathbb{N}$ y $\varphi \in C_0^\infty(\Omega)$, se tiene:

$$\int_{\Omega} (D^\alpha \varphi_n) \varphi dx = (-1)^\alpha \int_{\Omega} \varphi_n D^\alpha \varphi dx \quad (123)$$

Pero por ser (φ_n) una sucesión de Cauchy en la norma (107) tenemos que las sucesiones $(D^\alpha \varphi_n) (|\alpha| \leq k)$ son de Cauchy en $L_2(\Omega)$ y, por lo tanto, convergen a ciertas funciones $g_\alpha \in L_2(\Omega); (|\alpha| \leq k)$. Luego, pasando al límite en (123) cuando $n \rightarrow \infty$, se tiene que

$$\int_{\Omega} g_\alpha \varphi dx = (-1)^\alpha \int_{\Omega} f D^\alpha \varphi dx, \quad (124)$$

de donde se deduce que existe la derivada en sentido de Soboliev $D^\alpha f = g_\alpha \in L_2(\Omega)$.

Para demostrar la inclusión inversa $W_2^k(\Omega) \subset H^k(\Omega)$, basta notar que por (111), para cada $u \in W_2^k(\Omega)$ la sucesión $(D^\alpha u)_h = D^\alpha u_h \in C^\infty(\overline{\Omega})$ y, por T1 iii), $\|D^\alpha u_h - D^\alpha u\|_2 \rightarrow 0; (|\alpha| \leq k)$, cuando $h \rightarrow 0$.

De aquí se deduce que la sucesión (u_h) converge a u en la norma (107).
□

Cuando la región $\Omega \subset \mathbb{R}^n$ no es acotada, se puede definir el espacio $H^k(\Omega)$, como el completamiento de $C^k(\Omega) \cap L_2(\Omega)$ con respecto a la norma (108). El lector puede verificar que, en el caso $\Omega = \mathbb{R}^n$ se obtiene el mismo espacio $H^k(\mathbb{R}^n)$ haciendo el completamiento de $C_0^\infty(\mathbb{R}^n)$ o de $S(\mathbb{R}^n)$ con respecto a la norma (108).

Utilizando T 8 y T 6 i) se obtiene sin dificultad la caracterización siguiente de los espacios $H^k(\mathbb{R}^n)$.

T 9. Para todo $k \in \mathbb{N}$, se tiene

$$H^k(\mathbb{R}^n) = \{u \in L_2(\mathbb{R}^n) : (1 + |\xi|^2)^{k/2} \hat{u}(\xi) \in L_2(\mathbb{R}^n)\} \quad (125)$$

Además la norma $\|\cdot\|_{2,k}$ definida por (108) es topológicamente equivalente (40) a la norma

$$\|u\|_2^{(k)} = \|(1 + |\xi|^2)^{k/2} \hat{u}\|_2. \quad (126)$$

Demostración. Ejercicio. □

Este teorema nos sugiere cómo definir los espacios $H^k(\mathbb{R}^n)$, para todo $k > 0$. En efecto se puede definir $H^k(\mathbb{R}^n)$ ($k > 0$) como el subespacio de $L_2(\mathbb{R}^n)$ definido por (125) y provisto de la norma (126).

En el libro de J.L. Lions y E. Magenes, *Problemas de contorno no homogéneos y sus aplicaciones (vol. I)* [5] se puede hallar el resultado siguiente que expondremos sin demostración.

T 10. Si Ω es una región acotada con frontera regular, entonces para cada $k \in \mathbb{N}$, el espacio $H^k(\Omega)$ coincide con el espacio de las restricciones

a Ω de funciones de $H^k(\mathbb{R}^n)$. Además la norma en $H^k(\Omega)$ definida por la fórmula (108) es equivalente a la norma

$$\|u\|_{2,k} = \inf\{\|U\|_{H^k(\mathbb{R}^n)} : U(x) = u(x), x \in \Omega\} \quad (127)$$

y se puede definir una aplicación llamada *prolongación*

$$P : H^k(\Omega) \rightarrow H^k(\mathbb{R}^n) \quad (128)$$

que es continua y cumple:

$$(pu)(x) = u(x), \quad \text{para todo } x \in \Omega.$$

El Teorema 10 y la definición de los espacios $H^k(\mathbb{R}^n)$ para todo $k > 0$, nos permite definir $H^k(\Omega)$ para cualquier región acotada con frontera regular y cualquier $k > 0$. En efecto, se puede definir $H^k(\Omega)$ ($k > 0$) como el subespacio de $L_2(\Omega)$ compuesto por las restricciones a Ω de funciones de $H^k(\mathbb{R}^n)$ provisto de la norma (127).

Observación. Obviamente, con la definición que hemos dado para los espacios $H^k(\Omega)$ ($k > 0$) se sigue cumpliendo el enunciado del Teorema 10 para todo $k > 0$.

De la definición dada de los espacios $H^k(\Omega)$ ($k > 0$), se deduce que

$$H^{k_1}(\Omega) \subset H^{k_2}(\Omega) \quad (k_1 > k_2 \geq 0) \quad (129)$$

Además de la desigualdad evidente entre las normas

$$\|u\|_{2,k_2} \leq \|u\|_{2,k_1}, \quad u \in H^{k_1}(\Omega) \quad (130)$$

para $k_1 > k_2 \geq 0$ se tiene que la *inmersión* (128) es continua, es decir la aplicación $i : H^{k_1}(\Omega) \rightarrow H^{k_2}(\Omega); i(u) = u$, es continua. El resultado siguiente es muy utilizado en la teoría de operadores diferenciales.

T 11. (compacidad de la inmersión entre espacios de Soboliev)

Sea Ω una región acotada de $\mathbb{R}^n$ con frontera regular y sean $k_1 > k_2 \geq$

0. Entonces la inmersión $H^{k_1}(\Omega) \rightarrow H^{k_2}(\Omega)$, es compacta, es decir todo conjunto acotado en $H^{k_1}(\Omega)$ es relativamente compacto (D1.2.5) en $H^{k_2}(\Omega)$.

Demostración. No es difícil verificar que la compacidad de la inmersión es equivalente a probar que si la sucesión (u_n) de $H^{k_1}(\Omega)$ converge débilmente a cero en $H^{k_1}(\Omega)$ entonces ella también converge a cero en la norma de $H^{k_2}(\Omega)$ (¡ demuéstrela!). Pero la sucesión (u_n) converge débilmente a cero en $H^{k_1}(\Omega)$, si ella está acotada:

$$\|u_n\|_{2,k_1} \leq C$$

y además, para toda $\varphi \in C^{k_1}(\overline{\Omega})$, $\int_{\Omega} u_n \varphi dx \xrightarrow{n \rightarrow \infty} 0$

Por el Teorema 10, la sucesión $v_n = pu_n$, converge débilmente a cero en $H^{k_1}(\mathbb{R}^n)$ y además los soportes de las funciones v_n puede suponerse que están contenidos dentro de un compacto $K \supset \Omega$. Pero obviamente el operador restricción

$$r : H^k(\mathbb{R}^n) \rightarrow H^k(\Omega) \quad (131)$$

es continuo para todo $k > 0$ y, por ello, es suficiente mostrar que la convergencia débil de la sucesión (v_n) a cero en $H^{k_1}(\mathbb{R}^n)$, implica la convergencia de la sucesión (v_n) en la norma de $H^{k_2}(\mathbb{R}^n)$. Luego, debemos probar que para cualquier $\varepsilon > 0$ existe un número natural $n(\varepsilon)$, tal que para $n \geq n(\varepsilon)$

$$A_n = \int_{\mathbb{R}^n} (1 + |\xi|^2)^{k_2} |\hat{v}_n(\xi)|^2 d\xi \leq \varepsilon, \quad (132)$$

donde v_n es la transformada de Fourier de la función v_n . Pero

$$\begin{aligned} A_n &= \int_{|\xi| \geq M} (1 + |\xi|^2)^{k_2 - k_1} (1 + |\xi|^2)^{k_1} |\hat{v}_n(\xi)|^2 d\xi \\ &\quad + \int_{|\xi| \leq M} (1 + |\xi|^2)^{k_2} |\hat{v}_n(\xi)|^2 d\xi \end{aligned}$$

de donde

$$\begin{aligned} A_n \leq & (1 + M^2)^{k_2 - k_1} \int_{\mathbb{R}^n} (1 + |\xi|^2)^{k_1} |\hat{v}_n(\xi)|^2 d\xi \\ & + (1 + M^2)^{k_2} \int_{|\xi| \geq M} |\hat{v}_n(\xi)|^2 d\xi \end{aligned} \quad (133)$$

Por la acotación uniforme $\|v_n\|_2^{(k_1)} \leq C_1$, se tiene que

$$A_n \leq C_1(1 + M^2)^{k_2 - k_1} + (1 + M^2)^{k_2} \int_{|\xi| \geq M} |\hat{v}_n(\xi)|^2 d\xi \quad (134)$$

Tomemos M tal que $C_1(1 + M^2)^{k_2 - k_1} \leq \varepsilon/2$. Entonces tendremos la desigualdad (132), si mostramos que

$$\int_{|\xi| \geq M} |\hat{v}_n(\xi)|^2 d\xi \leq \frac{\varepsilon}{2(1 + M^2)^{k_2}}, \quad n \geq n(\varepsilon) \quad (135)$$

Pero como el soporte de todas las funciones v_n está incluido en un compacto K , entonces si tomamos una función $\Theta \in C_0^\infty(\mathbb{R}^n)$ que sea igual a 1 sobre K , podemos escribir

$$\hat{v}_n(\xi) = \int_{\mathbb{R}^n} v_n \{ \Theta e^{-ix \cdot \xi} \} dx. \quad (136)$$

Por la convergencia débil de (v_n) a cero en $H^{k_1}(\mathbb{R}^n)$ y por ser $|\xi| \leq M$, se tiene que (136) converge uniformemente a cero en $|\xi| \leq M$. Luego

$$\int_{|\xi| \geq M} |\hat{v}_n(\xi)|^2 d\xi \xrightarrow{n \rightarrow \infty} 0$$

lo cual prueba (135) y con ello, el teorema. $\square$

Demostremos un lema auxiliar que nos permitirá obtener una importante desigualdad de interpolación entre espacios de Sobolev.

L 1. Sean E, F y G , espacios de Banach, que satisfacen la condición

$$E \subset F \subset G$$

y además que la inmersión $E \rightarrow F$ es compacta. Entonces para cualquier número real $\theta > 0$, existe una constante $C(\theta) \geq 0$, tal que

$$\|u\|_F \leq \theta \|u\|_E + C(\theta) \|u\|_G, \quad \forall u \in E. \quad (137)$$

Demostración. Supongamos que (137) no es cierta. Entonces, para algún $\theta > 0$, existen vectores $u_n \in E$ y números $C_n \rightarrow \infty$, tales que

$$\|u_n\|_F > \theta \|u_n\|_E + C_n \|u_n\|_G.$$

Poniendo $v_n = u_n / \|u_n\|_E$, obtenemos

$$\|v_n\|_F > \theta + C_n \|v_n\|_G. \quad (138)$$

Pero como $\|v_n\|_F \leq C \|v_n\|_E = C$, una constante, entonces de (138) se deduce que

$$\|v_n\|_G \rightarrow 0. \quad (139)$$

Finalmente, del hecho que $\|v_n\|_E = 1$ y que la inmersión $E \rightarrow F$ es compacta, se puede elegir una subsucesión $(v_{\varphi(n)})$ de la sucesión (v_n) , que converge en la norma de F . Obviamente, de (139) se tendrá que $\|v_{\varphi(n)}\|_F \rightarrow 0$, lo cual está en contradicción con (138) y, por lo tanto, el lema queda demostrado. $\square$

T 12. (desigualdad de interpolación) Sea Ω una región acotada de $\mathbb{R}^n$ con frontera regular y sean $k_i, i = 1, 2, 3$, números reales tales que

$k_1 > k_2 > k_3 \geq 0$. Entonces para cualquier $\theta > 0$, existe una constante $C(\theta) > 0$, tal que para toda $u \in H^{k_1}(\Omega)$

$$\|u\|_{2,k_2} \leq \theta \|u\|_{2,k_1} + C(\theta) \|u\|_{2,k_3}. \quad (140)$$

Demostración. Es consecuencia del Teorema 11 y del Lema 1. $\square$

Pasemos ahora a estudiar algunas propiedades de regularidad de las funciones de $H^k(\Omega)$. Para ello compararemos los espacios $H^k(\Omega)$ con otros espacios frecuentemente utilizados, entre los cuales se encuentran los espacios de funciones continuamente diferenciables en sentido usual. Comencemos por un resultado simple.

T 13. (regularidad de las funciones de $H^k(\Omega)$) Supongamos que Ω es una región acotada de $\mathbb{R}^n$ con frontera regular. Entonces para $k > n/2$, tiene lugar la inclusión

$$H^k(\Omega) \subset C(\overline{\Omega}) \quad (141)$$

Además la inmersión $H^k(\Omega) \rightarrow C(\overline{\Omega})$ es compacta si se considera en $C(\overline{\Omega})$ la norma $\|\cdot\|_\infty$.

Demostración. Sea $u \in H^k(\Omega)$. Pondremos $v = pu \in H^k(\mathbb{R}^n)$, donde p es el operador de prolongación (128). Por definición $v = u$ casi dondequiera en Ω . Sea $\hat{v}$ la transformada de Fourier de la función v . Entonces

$$\hat{v} = (1 + |\xi|^2)^{-k/2} (1 + |\xi|^2)^{k/2} \hat{v} \in L_1(\mathbb{R}^n) \quad (142)$$

de donde

$$\|\hat{v}\|_1 \leq C \|v\|_2^{(k)}.$$

De aquí, se deduce que v como transformada inversa de Fourier de una función de $L_1(\mathbb{R}^n)$ es una función continua en $\mathbb{R}^n$, la cual tiende a cero cuando $|\xi| \rightarrow \infty$ (T4). Por consiguiente, tomando la restricción

a Ω obtenemos la inclusión (141). Para probar la compacidad de la inmersión utilizaremos el Teorema de Arzela – Ascoli (T1.2.8). Sea un subconjunto acotado en $H^k(\Omega)$. Probaremos que B es relativamente compacto en $C(\overline{\Omega})$ provisto de la norma $\|\cdot\|_\infty$. El conjunto $p[B] = \{pu : u \in B\}$, continuará siendo acotado en $H^k(\mathbb{R}^n)$. Por (142), para toda $v = pu, u \in B$ se tendrá:

$$v(x) - v(y) = (2\pi)^{-n} \int_{\mathbb{R}^n} (1 + |\xi|^2)^{-k/2} (1 + |\xi|^2)^{k/2} \hat{v}(\xi) (e^{i\xi \cdot x} - e^{i\xi \cdot y}) d\xi.$$

Luego

$$|v(x) - v(y)| \leq (2\pi)^{-n} \left(\int_{\mathbb{R}^n} 4(1 + |\xi|^2)^{-k} \sin^2 \frac{\xi(x-y)}{2} d\xi \right)^{1/2} \|v\|_2^{(k)} \quad (143)$$

Pero, por el teorema de la convergencia dominada de Lebesgue, la integral en el miembro derecho de (143) converge a cero cuando $|x-y| \rightarrow 0$, de donde se concluye la equicontinuidad de la familia B en Ω . La equiacotación de B es evidente. Luego, aplicando el Teorema de Arzela – Ascoli (T1.2.8) se tiene que B es una familia relativamente compacta en $C(\overline{\Omega})$. Así obtenemos la compacidad de la inmersión, con lo cual queda probado el teorema. $\square$

Del Teorema 13 se obtiene inmediatamente el corolario siguiente:

C 1. Supongamos que Ω satisface las condiciones del Teorema 13 y que $k > \frac{n}{2} + m, m$ entero, $m > 0$. Entonces tiene lugar la inclusión

$$H^k(\Omega) \subset C^m(\overline{\Omega}) \quad (144)$$

Además, la inmersión $H^k(\Omega) \rightarrow C^m(\overline{\Omega})$ es compacta, si se provee a

$C^m(\overline{\Omega})$ de la norma $\sum_{|\alpha| \leq m} \|D^\alpha v\|_\infty$.

Demostración. Ejercicio. $\square$

Observación. Notemos que las inclusiones (141) y (144) se interpretan como que toda función de $H^k(\Omega)$ coincide casi dondequiera en Ω con una función de $C(\overline{\Omega})$ (resp. de $C^m(\overline{\Omega})$).

C 2. Supongamos, que Ω satisface las condiciones del Teorema 13. Entonces

$$\bigcap_{k=0}^{\infty} H^k(\Omega) = C^\infty(\overline{\Omega}) \quad (145)$$

Demostración. Ejercicio. $\square$

Para obtener los llamados *teoremas de inmersión de espacios de Soboliev en diferentes dimensiones*, pasaremos a definir los espacios $H^k(\Gamma)$, para cualquier variedad compacta Γ de clase C^∞ y de dimensión p en $\mathbb{R}^n$. Para mayor comodidad supondremos que $p = n - 1$ y que Γ es la frontera de una región acotada Ω con frontera regular. Sea $U_j, j = 1, 2, \dots, \nu$, un cubrimiento abierto en $\mathbb{R}^n$, de la variedad Γ . Sea, para cada j , una aplicación infinitamente diferenciable

$$\varphi_j : U_j \rightarrow U = \{y \in \mathbb{R}^n : y = (y', y_n), |y'| < 1, -1 < y_n < 1\}$$

la cual posee un inverso infinitamente diferenciable

$$\varphi_j^{-1} : U \rightarrow U_j.$$

Supondremos que la aplicación φ_j aplica $U_j \cap \Omega$ sobre $U_+ = \{y : y \in U, y_n > 0\}$ y a la vez, aplica $U_j \cap (\mathbb{R}^n \setminus \Omega)$ sobre $U_- = \{y : y \in U, y_n < 0\}$. Por consiguiente, φ_j aplica $U_j \cap \Gamma$ sobre $U \cap \{y_n = 0\}$.

Además, deben cumplirse las siguientes condiciones de compatibilidad: si $U_j \cap U_i \neq 0$, entonces existe un homeomorfismo infinitamente diferenciable J_{ij} de $\varphi_i(U_i \cap U_j)$ en $\varphi_j(U_i \cap U_j)$, con jacobiano positivo y tal que

$$\varphi_j(x) = J_{ij}(\varphi_i(x)) \quad \forall x \in U_i \cap U_j.$$

Sea $\{\alpha_j\}$ una partición de unidad sobre Γ con las propiedades siguientes:

$$\alpha_j \in C^\infty(\Gamma) = \left\{ \begin{array}{l} \text{espacio de las funciones infinitamente} \\ \text{diferenciable sobre } \Gamma \end{array} \right\},$$

α_j tiene soporte compacto en $U_j \cap \Gamma$, y

$$\sum_{j=1}^{\nu} \alpha_j = 1 \quad \text{en } \Gamma.$$

Para cualquier función u , definida sobre Γ , se tiene el desarrollo

$$U = \sum_{j=1}^{\nu} \alpha_j U.$$

Definamos

$$\varphi_j^*(\alpha_j u)(y', 0) = (\alpha_j u)(\varphi_j^{-1}(y', 0)), \quad y' \in U \cap \{y_n = 0\}.$$

Como las funciones α_j tienen soporte compacto en $\Gamma \cap U_j$, las funciones $\varphi_j^*(\alpha_j u)$ tienen también soporte compacto en $U \cap \{y_n = 0\}$ y, por consiguiente $\varphi_j^*(\alpha_j u)$ puede ser prolongada a todo $\mathbb{R}_y^{n-1}$, haciéndola igual a cero fuera de $U \cap \{y_n = 0\}$.

Para cualquier número real $k > 0$, daremos la definición siguiente:

$$H^k(\Gamma) = \{u : \varphi_j^*(\alpha_j u) \in H^k(\mathbb{R}_y^{n-1}), j = 1, \dots, \nu\} \quad (146)$$

No es difícil demostrar que la definición algebraica (146) no depende de la elección del *sistema de cartas locales* $\{U_j, \varphi_j\}$ y de la partición de unidad $\{\alpha_j\}$. En $H^k(\Gamma)$, se puede introducir la norma siguiente:

$$\|U\|_{H^k(\Gamma)} = \left(\sum_{j=1}^{\nu} \|\varphi_j^*(\alpha_j u)\|_{H^k(\mathbb{R}^{n-1}_{y'})}^2 \right)^{1/2} \quad (147)$$

que depende, claramente, del sistema $\{U_j, \varphi_j, \alpha_j\}$. Es fácil probar, que $H^k(\Gamma)$ con la norma (147) *es un espacio de Hilbert*, y que *diferentes normas (147) son equivalentes entre sí*.

Notemos que $C^\infty(\Gamma)$ *es denso en* $H^k(\Gamma)$. En particular, para $k = 0$, hemos dado implícitamente la definición del espacio $L_2(\Gamma)$ con respecto a la medida $d\sigma$, generada sobre la superficie Γ por la medida de Lebesgue dx .

Con un ligero aumento de complejidad teórica *puede extenderse la definición de* $H^k(\Gamma)$ *a cualquier superficie compacta* Γ *en* $\mathbb{R}^n$, *de clase* C^∞ *y de dimensión* $p < n$ *arbitraria*.

T 4. (trazas de las funciones de $H^k(\Omega)$.) Sean Ω una región arbitraria de $\mathbb{R}^n$ con frontera regular Γ y una superficie compacta de clase C^∞ y de dimensión $p < n$ contenida en $\bar{\Omega}$. En particular puede ser $\Omega = \mathbb{R}^n$.

- i) Si $|\alpha| < k - \frac{n-p}{2}$, la aplicación $u \rightarrow D^\alpha u|_\Gamma$, definida de $C^\infty(\bar{\Omega}) \rightarrow C^\infty(\Gamma)$, se extiende continuamente a una aplicación lineal

$$T_\alpha : H^k(\Omega) \rightarrow H^{k-|\alpha|-\frac{n-p}{2}}(\Gamma) \quad (148)$$

de manera tal que la inmersión (148) es compacta.

- ii) Si Ω es además acotada $\Gamma = \partial\Omega$ y $j < k - \frac{1}{2}$, la aplicación $u \rightarrow \frac{\partial^j u}{\partial \nu^j}$, definida de $C^\infty(\bar{\Omega}) \rightarrow C^\infty(\partial\Omega)$, donde $\frac{\partial^j}{\partial \nu^j}$ denota la derivada de

orden j en la dirección normal exterior ν a $\partial\Omega$, puede extenderse continuamente a una aplicación lineal

$$N_j : H^k(\Omega) \rightarrow H^{k-j-1/2}(\Gamma) \quad (149)$$

de manera tal que la inmersión (149) es compacta y sobreyectiva. Existe una aplicación lineal y continua de levantamiento

$$P_j : H^{k-j-1/2}(\Gamma) \rightarrow H^k(\Omega) \quad (150)$$

que es inversa de N_j a la derecha.

Demostración. La demostración del inciso ii) puede encontrarse en el libro de J.L. Lions y E. Magenes: *Problemas de contorno no homogéneos y sus aplicaciones*. El inciso i) puede ser obtenido de manera similar, atendiendo a la definición de los espacios $H^s(\Gamma)$, y se deja de ejercicio al lector. $\square$

D 7. (trazas en sentido de Soboliev) Sean Ω y Γ como en el Teorema 14. Si $|\alpha| < k - \frac{n-p}{2}$, a la función $T_\alpha(u) \in H^{k-|\alpha|-\frac{n-p}{2}}(\Gamma)$ se le llama *traza en sentido de Soboliev* de la derivada de orden α de la función $u \in H^k(\Omega)$.

En particular, para $\alpha = 0$, la función $T_0(u) = u|_\Gamma$ cuando existe, se llama traza de la función $u \in H^k(\Omega)$ a Γ .

Observación. Consideremos la variedad compacta suave Γ de dimensión $p < n$ contenida en $\bar{\Omega} \subset \mathbb{R}^n$. Si $u \in H^k(\Omega)$, se puede encontrar una sucesión $u_n \in C^\infty(\bar{\Omega})$ tal que $D^\alpha u_n \rightarrow D^\alpha u$ en $L_2(\Omega)$ para todo α con $|\alpha| \leq k$. Tiene sentido considerar las restricciones $D^\alpha u_n|_\Gamma$. El Teorema 14 i), nos dice que si $|\alpha| < k - \frac{n-p}{2}$, entonces la sucesión de funciones $D^\alpha u_n|_\Gamma$ es una sucesión de Cauchy en $L_2(\Gamma)$, por lo tanto, converge a una función g_α en $L_2(\Gamma)$ a la cual se llama traza de $D^\alpha u$ a Γ : $T_\alpha u := D^\alpha u|_\Gamma = g_\alpha$. Esta definición tiene sentido, ya que g_α

depende solamente de u y no de la elección de la sucesión u_n en $C^\infty(\overline{\Omega})$. Resulta que si $|\alpha| \geq k - \frac{n-p}{2}$ entonces la sucesión de funciones $D^\alpha u_n|_\Gamma$ ó en caso de que converja, su límite dependerá no sólo de u sino de la sucesión u_n escogida.

Consideremos de nuevo la región $\Omega \subset \mathbb{R}^n$ y la superficie $\Gamma \subset \overline{\Omega}$, como en el Teorema 14. Los espacios de Banach

$$H_0^k(\Omega; \Gamma), \quad (151)$$

definidos como la clausura de

$$C_0^\infty(\overline{\Omega} \setminus \Gamma) = \left\{ u \in C^\infty(\overline{\Omega}) : \begin{array}{l} u \text{ se anula en alguna} \\ \text{vecindad de } \Gamma \text{ en } \overline{\Omega} \end{array} \right\} \quad (152)$$

en $H^k(\Omega)$, nos serán de gran utilidad en el estudio de los operadores diferenciales. En particular, si Ω es acotado y $\Gamma = \partial\Omega$, denotaremos

$$H_0^k(\Omega) = H_0^k(\Omega; \partial\Omega) \quad (153)$$

Para cualquier función $u \in C_0^\infty(\overline{\Omega} \setminus \Gamma)$ se tiene que $D^\alpha u|_\Gamma = 0$, luego por el Teorema 14, para las trazas de las derivadas de orden α de funciones de $H^k(\Omega)$, se tendrá

$$T_\alpha(u) = D^\alpha u|_\Gamma = 0; \quad |\alpha| \leq k - \frac{n-p}{2}.$$

Más exactamente se tiene el teorema siguiente:

T 15. Sean Ω y Γ como en el Teorema 14

i) Si $k > \frac{n-p}{2}$, entonces

$$\begin{aligned} H_0^k(\Omega; \Gamma) &= \{u \in H^k(\Omega) : T_\alpha(u) = D^\alpha u|_\Gamma = 0\} \\ &\quad 0 \leq |\alpha| < k - \frac{n-p}{2}, \end{aligned} \quad (154)$$

- ii) Si $k \leq \frac{n-p}{2}$, $C_0^k(\overline{\Omega} \setminus \Gamma)$ es denso en $H^k(\Omega)$.
 iii) Si Ω es acotado, $\Gamma = \partial\Omega$ y $k > \frac{1}{2}$, entonces

$$\begin{aligned} H_0^k(\Omega) &= \{u \in H^k(\Omega) : N_j(u) = \frac{\partial^j u}{\partial \nu^j} = 0\} \\ &\quad 0 \leq j < k - 1/2 \end{aligned} \quad (155)$$

Demostración. La inclusión

$$\begin{aligned} H_0^k(\Omega; \Gamma) &\subset \{u \in H^k(\Omega) : T_\alpha(u) = D^\alpha u|_\Gamma = 0\} \\ &\quad 0 \leq |\alpha| < k - \frac{n-p}{2} \end{aligned}$$

en (i) se obtiene, según el razonamiento previo al enunciado del Teorema 15, como una consecuencia inmediata del Teorema 14. La inclusión inversa se demuestra para iii) en el libro de J.L. Lions y E. Magenes: *Problemas de contorno no homogéneos y sus aplicaciones* [5]. La extensión de este resultado a i) se lleva a cabo sin dificultad haciendo uso de cartas locales, teniendo en cuenta la definición de T_α .

La demostración de ii) requiere del desarrollo de una técnica algo compleja y que nos aparta de la línea central del libro. Por ello no la citaremos aquí. $\square$

A continuación propondremos, en calidad de ejercicio, algunas propiedades de los espacios de Soboliev que serán útiles en el estudio de los operadores diferenciales.

E 1. Denotemos mediante $\|\cdot\|_{2,\Omega}$ la norma usual de una función en $L_2(\Omega)$.

- a) Demuestre que para cada $u \in C_0^\infty(a, b)$ se cumple la desigualdad

$$\|u\|_{2,(a,b)} \leq C \left\| \frac{du}{dx} \right\|_{2,(a,b)}$$

donde la constante C sólo depende de a y de b .

- b) Si K es el cubo n -dimensional $K = \prod_{i=1}^n (a_i, b_i)$, demuestre que para cada $u \in C_0^\infty(K)$ tiene lugar la desigualdad

$$\|u\|_{2,K} \leq C_K \|\nabla u\|_{2,K}$$

donde la constante C_K sólo depende de a_i y b_i , $i = 1, \dots, n$.

- c) Obtenga la desigualdad del inciso b) cuando en lugar de K tiene un subconjunto abierto arbitrario $\Omega \subset \mathbb{R}^n$ con frontera suave y que tiene lugar la llamada *desigualdad de Friedrichs*:

$$\|u\|_{2,\Omega} \leq C_\Omega \|\nabla u\|_{2,\Omega}, \quad \forall u \in H_0^1(\Omega).$$

- d) Concluya del inciso anterior que mediante $\|\nabla u\|_{2,\Omega}$ se puede definir en $H_0^1(\Omega)$ una norma equivalente a la norma usual de Sobolev.

El siguiente ejercicio permite obtener una caracterización importante de los espacios de Sobolev definidos sobre intervalos de $\mathbb{R}$:

E 2. a) Demuestre que si una función compleja u definida sobre el intervalo acotado $[a, b]$ satisface las propiedades: $u \in L_2(a, b)$; $u, u', \dots, u^{(k-1)}$ son absolutamente continuos y $u^{(k)} \in L_2(a, b)$; entonces $u^{(j)} \in L_2(a, b)$, $j = 1, \dots, k-1$ y además se tiene la desigualdad

$$\|u^{(j)}\|_2 \leq C \{\|u\|_2 + \|u^{(k)}\|_2\}, \quad j = 1, \dots, k-1.$$

- b) Demuestre la igualdad:

$$H^k(a, b) = \left\{ u \in L_2(a, b) \mid \begin{array}{l} \text{i) } u, u', \dots, u^{(k-1)} \text{ absolutamente continuas} \\ \text{ii) } u^{(k)} \in L_2(a, b) \end{array} \right\}$$

y utilizando el inciso a) pruebe que mediante $(\|u\|_2 + \|u^{(k)}\|_2)^{1/2}$ se define una norma en $H^k(a, b)$ equivalente a la norma usual $\|u\|_{2,k}$.

c) Concluir la igualdad

$$H_0^k(a, b) = \left\{ \begin{array}{l} u \in H^k(a, b) \mid \\ \quad \mid \end{array} \begin{array}{l} u(a) = u'(a) = \dots = u^{(k-1)}(a) = 0 \\ u(b) = u'(b) = \dots = u^{(k-1)}(b) = 0 \end{array} \right\}.$$

2 Operadores y funcionales lineales continuos

El estudio de los operadores y funcionales lineales es fundamental en el tratamiento de innumerables problemas dentro de los espacios vectoriales. Cuando se combinan la estructura vectorial de los espacios normados con la estructura métrica, aparece la necesidad de introducir un tipo particular de aplicaciones continuas, que tienen la propiedad de conservar la estructura vectorial; tales aplicaciones son los operadores lineales continuos y los funcionales lineales continuos. El objetivo fundamental del Análisis Funcional Lineal es precisamente el estudio de las propiedades de los operadores y funcionales lineales definidos en espacios de funciones.

En este capítulo, estudiaremos un grupo importante de propiedades que tienen los operadores funcionales y lineales continuos, actuando sobre espacios de Banach. El § 2.1 se dedica al estudio de algunas de sus propiedades generales y se hace énfasis en la relación entre continuidad y acotación de ellos.

En el § 2.2 se estudia el problema de la prolongación continua de operadores y funcionales lineales. En este epígrafe se estudia el Teorema de Hahn-Banach y se dan algunas de sus consecuencias.

En el § 2.3 se estudian algunos teoremas de representación para funcionales lineales continuos definidos sobre espacios de Hilbert y espacios

de funciones continuas, los cuales son de gran utilidad en las aplicaciones.

En el § 2.4 se examinan diferentes tipos de convergencia para sucesiones de operadores lineales continuos, obteniéndose algunos resultados importantes del Análisis Funcional, como son los teoremas de Banach–Steinhaus y Banach–Alaoglu. Al final de este epígrafe se introducen los espacios de funciones generalizadas. El capítulo termina con el § 2.5 donde se estudian otros teoremas importantes relacionados con los operadores lineales continuos en los espacios de Banach: el teorema de la aplicación abierta y el teorema del grafo cerrado.

§ 2.1 Algunos resultados generales de la teoría de operadores y funcionales lineales continuos en espacios normados

En todo este epígrafe consideraremos generalmente espacios vectoriales sobre $\mathbb{C}$. No obstante todos los resultados que obtendremos serán válidos para espacios vectoriales sobre $\mathbb{R}$.

D 1. (operador lineal y funcional lineal) Sean E y F espacios vectoriales. Llamaremos *operador lineal* de E en F a toda aplicación

$$A : E \longrightarrow F$$

que cumpla la propiedad siguiente:

$$A(\lambda x + \mu y) = \lambda A(x) + \mu A(y) \tag{1}$$

para todos $x, y \in E$; $\lambda, \mu \in \mathbb{C}$. Si $F = \mathbb{C}$ a A se le llama *funcional lineal*.

Observación. Para no complicar la notación, en la definición anterior hemos denotado de la misma forma las operaciones de suma y producto por escalares en E y en F . A partir de este momento adoptaremos esta

notación salvo si se señala lo contrario. También denotaremos por Θ al elemento neutro en todos los espacios vectoriales citados.

De la propiedad (1) de los operadores lineales se deduce que para cualesquiera elementos $x_1, x_2, \dots, x_n \in E$ y números complejos $\lambda_1, \lambda_2, \dots, \lambda_n$:

$$A\left(\sum_{k=1}^n \lambda_k x_k\right) = \sum_{k=1}^n \lambda_k A(x_k). \quad (2)$$

De (2) se tiene que para cualquier subespacio vectorial M de E , $A[M]$ es un subespacio vectorial de F .

No es difícil verificar que la compuesta $B \circ A$ de dos operadores lineales

$$A : E \rightarrow F, \quad B : F \rightarrow G$$

es también un operador lineal de E en G . Luego si $A : E \rightarrow E$ es un operador lineal, cualquier potencia de $A : A^k = A \circ A \circ \dots \circ A$ (k veces; $k \in \mathbb{N}$), es también un operador lineal de E en sí mismo.

Ejemplo 1. (matrices) Sea A una aplicación lineal de $\mathbb{R}^n$ en $\mathbb{R}^m$. Una base de $\mathbb{R}^n$ está constituida por los vectores coordenados

$$\varphi_k = (0, \dots, 0, 1, 0, \dots, 0) \in \mathbb{R}^n, \quad k = 1, 2, \dots, n$$

donde para cada k el número 1 ocupa el k ésimo lugar en el n -uplo φ_k . Ésta es la llamada base canónica de $\mathbb{R}^n$. De manera similar se construye la base canónica $(\Psi_j)_{j=1}^m$ de $\mathbb{R}^m$.

Por la propiedad 2, A está completamente determinado si se da la expresión de $A(\varphi_k)$ en términos de $(\Psi_j)_{j=1}^m$, para cada $k = 1, \dots, n$.

Si ponemos

$$A(\varphi_k) = \sum_{j=1}^m a_{jk} \Psi_j,$$

entonces la matriz

$$(A) = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix}$$

describe completamente al operador A .

En efecto, el producto de una matriz de orden $n \times m$ por un vector $1 \times n$ se define de manera tal que si para los elementos de $\mathbb{R}^n$, $x = x_1\varphi_1 + \cdots + x_n\varphi_n$, resp. $\mathbb{R}^m$, $y = y_1\Psi_1 + \cdots + y_m\Psi_m$ utilizamos la notación

$$x := \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}, y := \begin{pmatrix} y_1 \\ \vdots \\ y_m \end{pmatrix}$$

entonces se obtenga el mismo resultado al multiplicar (A) por el vector x que al evaluar el operador A en el vector x . La teoría de matrices, es pues, la teoría de operadores en espacios vectoriales de dimensión finita. $\square$

Ejemplo 2. (operadores de proyección ortogonal). A partir de la definición de proyección ortogonal de un vector sobre un subespacio cerrado de un espacio de Hilbert, se introducen los llamados operadores de proyección ortogonal. Si M es un subespacio vectorial cerrado del espacio de Hilbert H , se llama proyector ortogonal asociado a M , y se denota por p_M , al operador

$$p_M : H \longrightarrow H,$$

donde $p_M(x)$ es la proyección ortogonal de x sobre M , es decir $p_M(x)$ es la componente de x en M para la descomposición $H = M \oplus M^\perp$. No es difícil comprobar que p_M es lineal y por ello se deja de ejercicio al lector.

Además

$$p_M^2 = p_M \circ p_M = p_M, \quad (3)$$

es decir $p_M(x) = x$ para todo $x \in M$. $\square$

Ejemplo 3. (operadores simétricos en espacios de Hilbert) Un operador A de un espacio de Hilbert $(H, \langle \cdot \rangle)$ en sí mismo se llama *simétrico* si se cumple

$$\langle Ax, y \rangle = \langle x, Ay \rangle, \quad (4)$$

para todo par x, y de elementos de H .

No es difícil comprobar que todo operador simétrico es lineal (ejercicio). Un resultado mucho más profundo y que se deriva del llamado teorema del grafo cerrado, cuya demostración veremos en el § 2.5, dice que todo operador simétrico definido sobre un espacio de Hilbert es continuo. Este ejemplo nos muestra una vez más la relación intrínseca que existe sobre conceptos de orígenes diferentes como son la continuidad (de carácter topológico) y la simetría (de carácter geométrico).

Si se consideran en calidad de espacio de Hilbert, los espacios euclidianos de dimensión finita $\mathbb{C}^n$, no resulta difícil verificar que a los operadores simétricos de $\mathbb{C}^n$ en sí mismo, corresponden las matrices simétricas de orden $n \times n$, es decir, aquellas tales que $a_{ij} = \overline{a_{ji}}$, para todo par i, j con $1 \leq i, j \leq n$. $\square$

Ejemplo 4. (operadores unitarios) Son aquellos operadores definidos de un espacio normado $(E, \|\cdot\|)$ sobre sí mismo y tales que

$$\|Ax\| = \|x\|, \quad (5)$$

para todo $x \in E$. Si E es además euclidiano y A lineal, no es difícil comprobar que la propiedad (5) implica

$$\langle Ax, Ay \rangle = \langle x, y \rangle \quad (6)$$

para todos $x, y \in E$ (ejercicio),

El lector puede comprobar que en el caso en que $E = \mathbb{C}^n$, a los

operadores lineales unitarios corresponden las matrices de orden $n \times n$ que tienen determinante igual a 1. $\square$

Ejemplo 5. (operadores integrales) Sea $E = C([0, 1], \mathbb{C})$ y definamos el operador $A : E \rightarrow E$ mediante

$$Ax = \int_0^1 k(x, t)x(t)dt, \quad (7)$$

donde la función k es continua en el cuadrado unidad $0 \leq x \leq 1, 0 \leq t \leq 1$.

El lector puede verificar, a partir de las propiedades de la integral de Cauchy para funciones continuas, que el operador A es lineal. $\square$

Ejemplo 6. (funcional lineal de evaluación puntual) Sea E el espacio vectorial del Ejemplo 5. Un simple cálculo nos muestra que para cada $t_0 \in [0, 1]$ fijo, la aplicación

$$\delta_{t_0} : E \longrightarrow \mathbb{C},$$

definida por

$$\delta_{t_0}(x) = x(t_0), \quad (8)$$

es un funcional lineal, llamado evaluación puntual en el punto t_0 . $\square$

Ejemplo 7. (funcional lineal en un espacio de Hilbert) Sea H un espacio de Hilbert y $a \in H$ un elemento fijo. Utilizando las propiedades del producto escalar se puede probar que la aplicación

$$e_a : H \longrightarrow \mathbb{C}, \quad e_a(x) = \langle x, a \rangle \quad (9)$$

es un funcional lineal sobre H . $\square$

Para los operadores lineales en espacios vectoriales de dimensión

finita se tiene el importante resultado siguiente:

T 1. Todo operador lineal que actúa entre dos espacios normados de dimensión finita es continuo.

Demostración. Sabemos que en un espacio vectorial de dimensión finita todas las normas son equivalentes. Por ello, para analizar la continuidad del operador lineal $A : E \longrightarrow F$ $\dim E = m$, $\dim F = n$, podemos considerar la norma euclídeana en E . Sea $\{e_1, \dots, e_m\}$ una base ortonormal en E . Entonces, cada $x \in E$ tiene una expresión única del tipo

$$x = \sum_{k=1}^m \lambda_k e_k$$

y, por lo tanto

$$A(x) = \sum_{k=1}^m \lambda_k A(e_k).$$

Si ahora la sucesión $(x^{(p)}) \xrightarrow{p \rightarrow \infty} x$ en E , donde

$$x^{(p)} = \sum_{k=1}^m \lambda_k^{(p)} e_k$$

se tiene $\lambda_k^{(p)} = \langle x^{(p)}, e_k \rangle \xrightarrow{p \rightarrow \infty} \langle x, e_k \rangle = \lambda_k$.

Pero entonces

$$A(x^{(p)}) = \sum \lambda_k^{(p)} A(e_k) \xrightarrow{p \rightarrow \infty} \sum_{k=1}^m \lambda_k A(e_k) = A(x)$$

lo cual prueba la continuidad de A . $\square$

Cuando introdujimos la noción de continuidad en espacios métricos cualesquiera comenzamos por dar la definición de continuidad puntual

y posteriormente la de continuidad global. Resulta que para operadores lineales definidos sobre espacios normados ambas nociones son equivalentes. Más exactamente se tiene el teorema siguiente:

T 2. (relación entre continuidad global, continuidad puntual y acotación, para operadores lineales) Para un operador lineal $A : (E, \|\cdot\|_1) \rightarrow (F, \|\cdot\|_2)$ son equivalentes las propiedades siguientes:

- i) A es continuo en algún elemento x_0 de E ,
- ii) A es continuo,
- iii) A transforma conjuntos acotados de E en conjuntos acotados de F ,
- iv) $A[B_E]$ es un conjunto acotado de F , donde

$$B_E = \{x \in E : \|x\|_1 \leq 1\} \text{ es la bola unidad en } E,$$

- v) Existe una constante $M > 0$, tal que

$$\|Ax\|_2 \leq M\|x\|_1, \quad (10)$$

para todo $x \in E$.

Demostración. Probemos que i) $\Rightarrow$ ii). Sea x otro punto arbitrario de E y $(x^{(n)})$ una sucesión en E que converge a x , es decir $(x^{(n)}) \rightarrow x$. Por la continuidad de la suma, se tiene que $x_n - x + x_0 \rightarrow x_0$. Por otra parte, como hemos supuesto la continuidad de A en x_0 , tendremos que

$$A(x_n - x + x_0) = A(x_n) - A(x) + A(x_0) \rightarrow A(x_0) \quad (11)$$

Pero de nuevo basándonos en la continuidad de la suma, de (11) se deduce que

$$A(x_n) - A(x) \rightarrow 0,$$

es decir, $A(x_n) \rightarrow A(x)$, de donde concluimos la continuidad de A en un punto arbitrario $x \in C$.

Sabemos que un conjunto B de un espacio normado $(E, \|\cdot\|_1)$ está acotado si existe una constante $M > 0$ tal que $\|x\|_1 \leq M$ para todo $x \in B$.

Para probar ii) $\Rightarrow$ iii), supongamos que $A : E \rightarrow F$ es un operador lineal para el cual existe un conjunto acotado B en E tal que $A[B]$ no es un subconjunto acotado de F . Entonces para cada $n \in \mathbb{N}$, se puede encontrar un elemento $x_n \in B$ tal que

$$\|A(x_n)\|_2 \geq n \quad (12)$$

Como B es acotado, existe una constante $M > 0$ tal que

$$\|x_n\|_1 \leq M \quad \text{para todo } n \in \mathbb{N}. \quad (13)$$

De (13) se deduce obviamente que la sucesión $(\frac{x_n}{n})$ converge a cero en E . Pero de (12) se tiene que $\|A(\frac{x_n}{n})\| \geq 1$ y, por lo tanto, $A(\frac{x_n}{n})$ no converge a $A(\Theta) = \Theta$, luego A no sería continuo, con lo cual llegamos a una contradicción.

iii) $\Rightarrow$ iv) es evidente.

iv) $\Rightarrow$ v) Supongamos que no existe una constante $M > 0$, tal que se cumpla (10). Entonces, para cualquier entero positivo n , existe un elemento x_n en E , tal que

$$\|A(x_n)\|_2 > n\|x_n\|_1 \quad (14)$$

Si ponemos

$$x'_n = \frac{x_n}{\|x_n\|_1},$$

entonces $x'_n \in B_E$ y, por lo tanto (14) contradice iv).

Finalmente para probar que $v) \Rightarrow i)$ basta notar que de (10) se deduce inmediatamente la continuidad de A en $x_0 = \Theta$. $\square$

Observación. Un operador lineal $A : (E, \|\cdot\|_1) \rightarrow (F, \|\cdot\|_2)$ se llama *compacto* si transforma conjuntos acotados de E en conjuntos relativamente compactos de F . Como todo subconjunto relativamente compacto de un espacio normado es también acotado, del inciso iii) del Teorema 2, se deduce que todo operador compacto es continuo.

En el capítulo 4 nos dedicaremos al estudio de la teoría espectral de operadores compactos en espacios de Hilbert.

Ejemplo 8. Puede probarse sin dificultad que los operadores y funcionales lineales definidos en los Ejemplos 2 y 4 – 7, son todos continuos. Para demostrarlo aplicaremos el criterio $v)$ del Teorema 2. En efecto, como para todo $x \in H$, se tiene que $x = p_M(x) + (x - p_M(x))$, donde $p_M(x) \perp (x - p_M(x))$, entonces por el Teorema de Pitágoras obtenemos

$$\|x\|^2 = \|p_M(x)\|^2 + \|x - p_M(x)\|^2$$

y, por lo tanto

$$\|p_M(x)\| \leq \|x\|, \quad \text{para todo } x \in H. \quad (15)$$

Esta última desigualdad demuestra la continuidad del operador p_M del Ejemplo 2.

La propia definición (5) de los operadores unitarios nos muestra que todo operador unitario es continuo.

Para el operador integral (7), no es difícil comprobar la desigualdad

$$\|Ax\|_\infty \leq \max_{0 \leq s \leq 1} \left\{ \int_0^1 |K(s, t)| dt \right\} \|x\|_\infty, \quad (16)$$

para toda $x \in C([0, 1], \mathbb{C})$ lo cual prueba su continuidad.

Para el funcional de evaluación puntual δ_{t_0} definido en (8) se tiene

$$|\delta_{t_0}(x)| \leq \|x\|_\infty \quad (17)$$

para toda $x \in C([0, 1], \mathbb{C})$ lo cual nos muestra que es continuo.

De la desigualdad de Cauchy–Buniakovsky se concluye que

$$|e_a(x)| \leq \|a\| \|x\| \quad (18)$$

de donde se obtiene la continuidad de (9). $\square$

Notemos que del Teorema 2 se deduce que para todo operador lineal y continuo

$$A : (E, \|\cdot\|_1) \rightarrow (E, \|\cdot\|_2),$$

la cantidad

$$\|A\| = \sup_{x \in E, x \neq 0} \frac{\|Ax\|_2}{\|x\|_1}, \quad (19)$$

es un número positivo. Precisamente $\|A\|$ coincide con el ínfimo de todas las constantes $M > 0$ para las cuales tiene lugar el punto v) del Teorema 2.

El lector puede comprobar sin dificultad que

$$\|A\| = \sup_{\substack{x \in E \\ \|x\|_1 \leq 1}} \|A(x)\|_2 = \sup_{\substack{x \in E \\ \|x\|_1 = 1}} \|A(x)\|_2 \quad (20)$$

Consideremos ahora el espacio $\mathcal{L}(E, F)$ de todos los operadores lineales y continuos (acotados) de E en F . Este espacio puede proveerse de una estructura natural de espacio vectorial mediante la suma y el producto por escalares usual de operadores: Si $A, B \in \mathcal{L}(E, F)$; $\lambda \in \mathbb{C}$, entonces

$$\begin{aligned} (A + B)(x) &:= A(x) + B(x), \\ (\lambda A)(x) &:= \lambda A(x) \end{aligned} \quad (21)$$

En el caso particular en que $F = \mathbb{C}$, el espacio $\mathcal{L}(E, \mathbb{C})$ se llama *dual topológico* de E y se denota por E^* .

T 3. Para todo par de espacios normados $(E, \|\cdot\|_1)$ y $(F, \|\cdot\|_2)$ el espacio vectorial $\mathcal{L}(E, F)$ de los operadores lineales y continuos de E en F , provisto de la aplicación

$$\|\cdot\| : \mathcal{L}(E, F) \longrightarrow \mathbb{R}_+, \quad A \rightsquigarrow \|A\|$$

definida en (19), es un espacio normado.

Demostración. Se deja al lector verificar, en calidad de ejercicio, que la aplicación (19) satisface los requerimientos de una norma. $\square$

Ejemplo 9. Calculemos las normas de los operadores lineales continuos analizados en el Ejemplo 8.

De (15) se deduce que $\|p_M\| \leq 1$, pero como para $x \in M$, se tiene que $\|p_M\| = \|x\|$, entonces tendremos

$$\|p_M\| = 1. \quad (22)$$

La propia definición (5) muestra que la norma de un operador unitario es igual a 1.

La desigualdad (16) nos dice que la norma del operador integral (7) es

$$\|A\| \leq \max_{0 \leq s \leq 1} \int_0^1 |K(s, t)| dt \quad (23)$$

Podemos probar que la cantidad que se encuentra en la parte derecha de la desigualdad (23) es precisamente la norma de A . En efecto, sea s_0

el punto del intervalo $[0, 1]$ en el cual la función continua

$$\int_0^1 |K(s, t)| dt$$

alcanza su máximo. Dado $\varepsilon > 0$, sea p un polinomio que aproxima uniformemente a la función $K(s_0, \cdot)$ con un orden menor que ε .

$$\max_{0 \leq t \leq 1} |K(s_0, t) - p(t)| < \varepsilon$$

y sea x una función en $C([0, 1], \mathbb{R})$ con $\|x\|_\infty \leq 1$ y que aproxima a la función discontinua

$$\operatorname{sgn} p(t) = \begin{cases} 1 & \text{si } p(t) \geq 0, \\ -1 & \text{si } p(t) < 0 \end{cases}$$

en el sentido que

$$\left| \int_0^1 p(t)x(t)dt - \int_0^1 |p(t)|dt \right| < \varepsilon.$$

Esta última aproximación se construye fácilmente ya que p tiene solamente un número finito de cambios de signo (los detalles se dejan al lector).

Para este x tenemos

$$\begin{aligned} \left| \int_0^1 K(s_0, t)x(t)dt \right| &\geq \left| \int_0^1 p(t)x(t)dt \right| - \left| \int_0^1 [K(s_0, t) - p(t)]x(t)dt \right| \\ &\geq \left| \int_0^1 p(t)x(t)dt \right| - \varepsilon \geq \int_0^1 |p(t)|dt - 2\varepsilon \\ &\geq \int_0^1 |K(s_0, t)|dt - \left| \int_0^1 [K(s_0, t) - |p(t)|]dt \right| - 2\varepsilon \\ &\geq \int_0^1 |K(s_0, t)|dt - 3\varepsilon. \end{aligned}$$

De aquí, como $\|x\|_\infty \leq 1$, tenemos que

$$|||A||| \geq \int_0^1 |K(s_0, t)| dt - 3\varepsilon.$$

Pero como $\varepsilon > 0$ es arbitrario y la desigualdad en sentido inverso fue establecida más arriba, tenemos que

$$|||A||| = \max_{0 \leq s \leq 1} \int_0^1 |K(s, t)| dt. \quad (24)$$

De la desigualdad (17) se tiene que la norma del funcional lineal de evaluación puntual δ_{t_0} es

$$|||\delta_{t_0}||| \leq 1.$$

Si consideramos la función constante $x(t) = \langle 1 \rangle$, entonces

$$|\delta_{t_0}(x)| = 1 = \|x\|_\infty,$$

y, por lo tanto,

$$|||\delta_{t_0}||| = 1 \quad (25)$$

Finalmente, de (18) se deduce que $|||e_a||| \leq \|a\|$, pero tomando $x = a$, se tiene que

$$|e_a(x)| = \|a\| \|x\|,$$

de donde obtenemos

$$|||e_a||| = \|a\| \quad (26)$$

Se propone al lector probar que el funcional δ_{t_0} , definido sobre el espacio $C([0, 1], \mathbb{C})$ provisto de la norma (38) (Capítulo 1) de la convergencia en media cuadrática, no es continuo. $\square$

§ 2.2 Extensión de operadores y funcionales lineales continuos. Principio de prolongación de Hahn–Banach

Si en un espacio vectorial E está dado un operador lineal o un funcional lineal definido sobre un subespacio $M \subset E$, entonces, de manera natural podemos preguntarnos sobre la posibilidad de su prolongación a E conservando ciertas propiedades. En otras palabras, se trata de construir un nuevo operador o funcional, definido en todo el espacio E , el cual posee determinadas propiedades y que además coincide sobre M con el operador o funcional inicial.

Este problema de prolongación ha sido ya estudiado para funciones continuas en el Capítulo 1, (P1.1.1), (T1.1.6). En el Capítulo 1 vimos también la relación entre los problemas de prolongación de funciones continuas definidas sobre subconjuntos cerrados de un espacio métrico y la separación de cerrados por vecindades abiertas o por funciones continuas (T1.1.7).

En este epígrafe veremos una aplicación del Teorema 1.1.6 a la prolongación lineal continua de operadores lineales continuos y funcionales lineales continuos, definidos sobre un subespacio normado, a su clausura. En segundo lugar estudiaremos la relación entre los problemas de prolongación lineal continua y separación de conjuntos convexos en un espacio normado, por funcionales lineales continuos. El teorema fundamental que enunciaremos es el llamado Teorema de Hahn–Banach, una de las herramientas fundamentales del Análisis Funcional Lineal y que juega un papel muy importante en el estudio de los problemas de optimización en espacios vectoriales.

T 1. (prolongación lineal continua a la clausura) Sean E un espacio normado, M un subespacio vectorial denso en E , F un espacio de Banach, y $A : M \rightarrow F$ un operador lineal y continuo. Entonces A puede prolongarse de manera única, conservando su norma, a un operador

lineal y continuo $\tilde{A} : E \rightarrow F$.

Demostración. La existencia de la prolongación continua de A a todo E se concluye del Teorema 1.1.6. En efecto, para cada $x \in E$ y cada sucesión (x_n) en M tal que $(x_n) \rightarrow x$, se tiene que

$$\|A(x_n) - A(x_m)\|_F \leq \|A\| \|x_n - x_m\|_E,$$

de donde se concluye que la sucesión $(A(x_n))$ es de Cauchy y, por lo tanto, converge en F . Además, de la desigualdad

$$\|A(x_n) - A(y_n)\|_F \leq \|A\| \|x_n - y_n\|_E,$$

válida para todo par de sucesiones $(x_n), (y_n)$ en M que converjan al mismo elemento $x \in E$, se deduce que el límite de $(A(x_n))$ en F no depende de la elección de la sucesión (x_n) en M , siempre que $(x_n) \rightarrow x$. Del Teorema 1.1.6 concluimos que existe una prolongación continua única $\tilde{A} : E \rightarrow F$ de A a todo E .

La linealidad $\tilde{A}$ así como la conservación de la norma se demuestra fácilmente utilizando argumentos de continuidad y por ello se dejan de ejercicio al lector. $\square$

Pasemos ahora a estudiar la llamada *forma algebraica del Teorema de Hahn-Banach* del cual se deducen los teoremas de extensión de funcionales lineales continuos, pero antes veamos la siguiente definición.

D 1. (prolongación o extensión lineal) Sea f un funcional lineal definido sobre un subespacio vectorial M de un espacio vectorial E . Un funcional lineal g se llama una extensión lineal de f , si g está definido sobre un subespacio vectorial N de E que contiene a M y además la restricción de g a M ($g|_M$) coincide con f . En este caso se dice que g es una extensión de f desde M hasta N .

Dicho en forma simple, el Teorema de Hahn-Banach nos asegura que todo funcional lineal continuo f definido sobre un subespacio vectorial

M de un espacio normado E puede ser extendido a un funcional lineal continuo g definido sobre el espacio E y con la misma norma que la norma de f sobre M ; es decir,

$$\|g\| = \|f\|_M = \sup_{m \in M} \frac{|f(m)|}{\|m\|}.$$

Nosotros veremos una versión más general de este resultado en el cual se reemplazan las normas por funcionales sublineales. Esta generalización nos servirá más adelante para probar la versión geométrica del Teorema de Hahn-Banach.

D 2. (funcional sublineal) Una función real p definida sobre un espacio vectorial real E se llama funcional sublineal si satisface:

$$p(x+y) \leq p(x) + p(y) \quad \text{para todos } x, y \in E; \quad (27)$$

$$p(\lambda x) = \lambda p(x) \quad \text{para todo } \lambda \geq 0, x \in E. \quad (28)$$

Evidentemente, cualquier norma es un funcional sublineal.

T 2. (Hahn-Banach para espacios normados reales) Sea E un espacio vectorial normado real y p un funcional sublineal continuo sobre E . Sea f un funcional lineal definido sobre un subespacio vectorial M de E y que satisface

$$f(m) \leq p(m),$$

para todo $m \in M$. Entonces existe una extensión g de f a todo E tal que

$$g(x) \leq p(x)$$

sobre E .

Demostración. Aunque el teorema es válido en cualquier espacio normado, en nuestra demostración asumiremos que E es separable. El

resultado general es obtenido exactamente por el mismo método junto con una simple aplicación del Lema de Zorn. El lector familiarizado con el Lema de Zorn no debe tener mucha dificultad para generalizar la prueba. La idea básica es extender f en una dimensión cada vez y aplicar la hipótesis de inducción.

Supongamos que y es un vector de E que no está en M . Consideremos todos los elementos del subespacio vectorial generado por M y y , es decir $[M + y]$. Un elemento x de $[M + y]$ tiene una representación única de la forma $x = m + \lambda y$, donde $m \in M$ y λ es un número real. Toda extensión f_0 de f desde M hasta $[M + y]$ tiene la forma

$$f_0(x) = f(m) + \lambda f_0(y)$$

y, por lo tanto, f_0 está totalmente determinada por la constante $f_0(y)$. Nosotros debemos probar que esta constante puede ser seleccionada de manera tal que $f_0(x) \leq p(x)$ sobre $[M + y]$.

Para dos elementos cualesquiera m_1 y m_2 en M , tenemos

$$\begin{aligned} f(m_1) + f(m_2) &= f(m_1 + m_2) \leq p(m_1 + m_2) \\ &\leq p(m_1 - y) + p(m_2 + y), \end{aligned}$$

de donde

$$f(m_1) - p(m_1 - y) \leq p(m_2 + y) - f(m_2).$$

Luego

$$\sup_{m \in M} [f(m) - p(m - y)] \leq \inf_{m \in M} [p(m + y) - f(m)].$$

Por lo tanto, existe una constante C tal que

$$\sup_{m \in M} [f(m) - p(m - y)] \leq C \leq \inf_{m \in M} [p(m + y) - f(m)].$$

Entonces, para el vector $x = m + \lambda y \in [M + y]$, definiremos

$$f_0(x) = f(m) + \lambda C.$$

Ahora debemos probar que $f_0(m + \lambda y) \leq p(m + \lambda y)$. Si $\lambda > 0$, entonces

$$\begin{aligned}\lambda C + f(m) &= \lambda[C + f(\frac{m}{\lambda})] \\ &\leq \lambda[p(\frac{m}{\lambda} + y) - f(\frac{m}{\lambda}) + f(\frac{m}{\lambda})] \\ &= \lambda p(\frac{m}{\lambda} + y) \\ &= p(m + \lambda y)\end{aligned}$$

Si $\lambda = -\mu < 0$, entonces

$$\begin{aligned}-\mu C + f(m) &= \mu[-C + f(\frac{m}{\mu})] \leq \mu[p(\frac{m}{\mu} + y) - f(\frac{m}{\mu}) + f(\frac{m}{\mu})] \\ &= \mu p(\frac{m}{\mu} + y) \\ &= p(m - \mu y)\end{aligned}$$

Luego $f_0(m + \lambda y) \leq p(m + \lambda y)$ para todo $\lambda \in \mathbb{R}$ y f_0 es una extensión de f desde M hasta $[M + y]$.

Ahora sea $\{x_1, x_2, \dots, x_n, \dots\}$ un conjunto numerable denso en E . De este conjunto de vectores seleccionaremos un subconjunto $\{y_1, y_2, \dots, y_n, \dots\}$ formado por vectores linealmente independientes y que no pertenecen a M . El conjunto $\{y_1, y_2, \dots, y_n, \dots\}$ junto con el subespacio M generan un subespacio vectorial S denso en E .

El funcional f puede ser extendido a un funcional g_0 sobre el subespacio S de manera que

$$g_0(x) \leq p(x)$$

para todo $x \in S$. Para ello basta con extender f desde M hasta $[M + y_1]$, después hasta $[[M + y_1] + y_2]$ y así sucesivamente aplicando la hipótesis de inducción. Finalmente, el funcional lineal g_0 (el cual es continuo ya que p lo es) puede ser extendido por continuidad desde el subespacio denso S a todo el espacio E . Para cada $x \in E$, existe una sucesión (x_n)

en S que converge a x . La extensión g buscada viene dada por

$$g(x) = \lim g_0(x_n).$$

Obviamente g es lineal y por la continuidad de p se tendrá que

$$g(x) \leq p(x)$$

para todo $x \in E$, con lo cual queda demostrado el teorema. $\square$

T 3. (Hahn–Banach para espacios normados complejos) Sea M un subespacio vectorial del espacio normado complejo E , p un funcional sublineal continuo sobre E y f un funcional lineal definido sobre M , tal que

$$|f(m)| \leq p(m), \forall m \in M.$$

Entonces sobre E existe un funcional lineal que extiende a f , tal que

$$|g(x)| \leq p(x), \forall x \in M.$$

Demostración. Notemos que la afirmación de este teorema se deduce inmediatamente para espacios normados reales del Teorema 2 ya que $p(x) = p(-x)$ para toda $x \in E$. No es difícil comprobar que la parte real u del funcional lineal complejo f , es un funcional lineal real, y del hecho que $z = \operatorname{Re} z - i \operatorname{Re}(iz)$ para cada $z \in \mathbb{C}$, se cumple que

$$f(m) = u(m) - iu(im), \quad \forall m \in M. \quad (29)$$

Inversamente, si $u : M \rightarrow \mathbb{R}$ es un funcional lineal real, entonces un simple cálculo muestra que el funcional f , definido por (29) es un funcional lineal complejo cuya parte real coincide con u .

Sea pues $u = \operatorname{Re} f$. Por el Teorema 2, existe un funcional lineal U sobre E , tal que $U = u$ sobre M y $U \leq p$ sobre E . Sea g el funcional

lineal complejo cuya parte real coincide con U . Por el razonamiento hecho más arriba, se tiene que $g = f$ sobre M . Además, para cada $x \in E$, existe $\alpha \in \mathbb{C}$ con $|\alpha| = 1$ y, tal que $\alpha g(x) = |g(x)|$. Por eso,

$$|g(x)| = g(\alpha x) = U(\alpha x) \leq p(\alpha x) = p(x). \quad \square$$

C 1. (extensión de funcionales lineales continuos en espacios normados) Sea f un funcional lineal continuo definido sobre un subespacio vectorial M de un espacio normado E . Entonces existe un funcional lineal continuo g definido sobre todo E que es una extensión de f y cuya norma es igual a la norma de f sobre M .

Demostración. Basta tomar $p(x) = \|f\|_M \|x\|$ y aplicar el Teorema de Hahn-Banach (T3). $\square$

C 2. (existencia de elementos no nulos en E^*) Sea x un elemento de un espacio normado $(E, \|\cdot\|)$. Entonces existe un funcional lineal continuo g sobre E de norma 1 y tal que

$$g(x) = \|x\|.$$

Demostración. Supongamos que $x \neq \Theta$. Sobre el subespacio unidimensional generado por el vector x , definamos el funcional lineal f mediante

$$f(\lambda x) = \lambda \|x\|.$$

Entonces f es un funcional lineal acotado con norma igual a uno. Por el Corolario 1, f puede ser extendido a un funcional lineal continuo g sobre E con norma unidad. Este funcional satisface los requerimientos del corolario. Si $x = \Theta$, nos sirve cualquier funcional lineal continuo sobre E (cuya existencia acabamos de demostrar). $\square$

Observación. El corolario anterior nos muestra la riqueza en elemen-

tos del espacio dual correspondiente a cualquier espacio normado. Este resultado es importante ya que muchos problemas en los espacios normados se pueden plantear de manera equivalente en los correspondientes duales. Precisamente gracias a la riqueza del espacio dual a veces resulta más fácil resolver el problema dual y después de obtenida la solución pasar de nuevo al espacio original.

Para finalizar este epígrafe veremos el problema de separación de conjuntos convexos en espacios normados mediante funcionales lineales continuos. La base de este tipo de resultado se encuentra también en el Teorema de Hahn-Banach pero en su forma geométrica.

Existe una diferencia conceptual entre los funcionales lineales vistos como elementos de un espacio dual o vistos como hiperplanos generados en el espacio primario. Estos diferentes puntos de vista nos permiten relacionar los aspectos más relevantes de un espacio y su dual a través de una interpretación geométrica de los mismos.

D 3. (Hiperplanos) Un hiperplano H en un espacio vectorial E es un conjunto de la forma $H = x_0 + M$, donde M es un subespacio vectorial propio maximal en E , es decir M es un subespacio vectorial diferente de H y tal que si N es cualquier otro subespacio vectorial conteniendo a M , entonces $M = N$ o $N = E$.

Esta definición de hiperplano está enunciada sin hacer referencia explícita a los funcionales lineales ya que la misma refleja la interpretación geométrica de un hiperplano. Sin embargo los hiperplanos y los funcionales lineales están íntimamente relacionados, como lo muestran las dos proposiciones siguientes:

P 1. Sea H un hiperplano en el espacio vectorial E . Entonces existe un funcional lineal f sobre E y una constante c tales que

$$H = \{x \in E : f(x) = c\}.$$

Inversamente, si f es un funcional lineal no nulo sobre E , el conjunto $\{x \in E : f(x) = c\}$ es un hiperplano en E .

Demostración. Sea H un hiperplano en E . Entonces H es la traslación de un subespacio vectorial M de E :

$$H = x_0 + M.$$

Si $x_0 \notin M$, entonces $[x_0 + M] = E$, y para $x = \lambda x_0 + m$ con $m \in M$, definimos $f(x) = \lambda$. Entonces

$$H = \{x \in E : f(x) = 1\}.$$

Inversamente sea f un funcional no nulo sobre E y sea $M = \{x \in E : f(x) = 0\}$. Obviamente M es un subespacio vectorial de E . Sea $x_0 \in E$ tal que $f(x_0) = 1$. Entonces para cualquier $x \in E$, $f[x - f(x)x_0] = 0$ y, por lo tanto $x - f(x)x_0 \in M$. Luego, $E = [x_0 + M]$, de donde se deduce que M es un subespacio vectorial propio maximal en E . Para cualquier número real c , sea x_1 algún elemento para el cual $f(x_1) = c$. Entonces $\{x \in E : f(x) = c\} = \{x \in E : f(x - x_1) = 0\} = M + x_1$, que es un hiperplano. $\square$

La proposición anterior es de carácter general y es aplicable a hiperplanos en un espacio vectorial arbitrario. Un hiperplano H en un espacio normado E debe ser cerrado o denso en E ya que como H es una variedad lineal maximal y por ser $\overline{H}$ también una variedad lineal, entonces $\overline{H} = H$ o $\overline{H} = E$. Para nuestro propósito nos interesan fundamentalmente los hiperplanos cerrados en un espacio normado E . Estos hiperplanos corresponden a los funcionales lineales continuos sobre E .

P 2. Sea f un funcional lineal no nulo sobre el espacio normado E . Entonces el hiperplano $H = \{x \in E : f(x) = c\}$ es cerrado para cada constante c si y sólo si f es continuo.

Demostración. Supongamos primeramente que f es continuo. En-

tonces H es cerrado por ser la preimagen mediante f del subconjunto cerrado unitario $\{c\}$.

Inversamente, supongamos que

$$M = \{x \in E : f(x) = 0\}$$

es cerrado. Sea $E = [x_0 + M]$ y supongamos que $(x_n) \rightarrow x$ en E . Entonces $x_n = \lambda_n x_0 + m_n$, $x = \lambda x_0 + m$, y si denotamos por d la distancia de x_0 a M (la cual es positiva ya que M es cerrado), tendremos

$$|\lambda_n - \lambda|d \leq \|x_n - x\| \rightarrow 0,$$

y, por lo tanto, $\lambda_n \rightarrow \lambda$. Pero, por otra parte, $f(x_n) = \lambda_n f(x_0) + f(m_n) = \lambda_n f(x_0) \rightarrow \lambda f(x_0) = f(x)$. Luego, f es continuo sobre E . $\square$

Si f es un funcional lineal no nulo sobre un espacio vectorial E , se pueden asociar al hiperplano $H = \{x \in E : f(x) = c\}$ los siguientes cuatro conjuntos

$$\begin{aligned} \{x \in E : f(x) \leq c\}, & \quad \{x \in E : f(x) < c\}, \\ \{x \in E : f(x) \geq c\}, & \quad \{x \in E : f(x) > c\} \end{aligned}$$

llamados *semiespacios* determinados por H . Los dos primeros se llaman *semiespacios negativos* determinados por f y los otros dos *semiespacios positivos*. Si f es continuo, entonces los semiespacios $\{x \in E : f(x) < c\}$ y $\{x \in E : f(x) > c\}$ son abiertos mientras que $\{x \in E : f(x) \leq c\}$ y $\{x \in E : f(x) \geq c\}$ son cerrados.

Como hemos visto hasta el momento, existe una correspondencia entre los hiperplanos cerrados en un espacio normado E y los elementos de su espacio dual E^* . Luego es de esperar que aquellos resultados en que juegue un papel fundamental el espacio dual E^* puedan ser visualizados en términos de hiperplanos cerrados. Por este motivo se puede demostrar una versión geométrica del Teorema de Hahn–Banach la cual en su forma más simple nos dice que *dado un conjunto convexo c con interior no vacío y un punto x_0 que no pertenece a C^0 , existe un hiperplano cerrado que contiene a x_0 pero disjunto con C^0 .*

Antes de demostrar este importante teorema de separación veamos la siguiente definición auxiliar.

D 4. (funcional de Minkowsky) Sea C un conjunto convexo en un espacio vectorial normado E y supongamos que el elemento nulo Θ es un punto interior de C . Entonces el funcional de Minkowsky p de C es una función real definido sobre E mediante

$$p(x) = \inf \left\{ r \in \mathbb{R} : \frac{x}{r} \in C, r > 0 \right\} \quad (30)$$

Notemos que para C igual a la bola unidad de E , el correspondiente funcional de Minkowsky es $p(x) = \|x\|$. En el caso general, $p(x)$ define una especie de distancia desde el origen a x medida con respecto a C ; $p(x)$ es el factor por el cual debe ser multiplicado C hasta incluir a x .

L 1. Sea C un conjunto convexo que contiene al elemento nulo Θ en su interior. Entonces el funcional de Minkowsky p de C satisface:

- i) $\infty > p(x) \geq 0$ para todo $x \in E$,
- ii) $p(\lambda x) = \lambda p(x)$ para $\lambda > 0$,
- iii) $p(x + y) \leq p(x) + p(y)$
- iv) p es continuo
- v) $\overline{C} = \{x \in E : p(x) \leq 1\}$, $C^0 = \{x \in E : p(x) < 1\}$

Demostración. i) Como C contiene a una bola de centro en Θ , entonces dado $x \in E$ existe un $r > 0$ tal que $x/r \in C$. Luego $p(x)$ es finito para todo x . Obviamente $p(x) \geq 0$.

ii) Para $\lambda > 0$,

$$p(\lambda x) = \inf \left\{ r \in \mathbb{R} : \frac{\lambda x}{r} \in C, r > 0 \right\}$$

$$\begin{aligned}
&= \inf\{\lambda r' \in \mathbb{R} : \frac{x}{r'} \in C, r' > 0\} \\
&= \lambda \inf\{r' \in \mathbb{R} : \frac{x}{r'} \in C, r' > 0\} = \lambda p(x)
\end{aligned}$$

iii) Dados $x, y \in E$ y $\varepsilon > 0$, seleccionemos $t, s \in \mathbb{R}$ tales que

$$\begin{aligned}
p(x) &< t < p(x) + \varepsilon, \\
p(y) &< s < p(y) + \varepsilon
\end{aligned}$$

Por ii) $p(\frac{x}{t}) < 1$ y $p(\frac{y}{s}) < 1$ y por lo tanto $\frac{x}{t}, \frac{y}{s} \in C$. Sea $r = t + s$. Por la convexidad de C , $(\frac{t}{r})(\frac{x}{t}) + (\frac{s}{r})(\frac{y}{s}) = \frac{(x+y)}{r} \in C$. Entonces, $p(x+y)/r \leq 1$. Por ii) tendremos que $p(x+y) \leq r < p(x) + p(y) + 2\varepsilon$. De la arbitrariedad de ε se deduce la subaditividad de p .

iv) Sea ε el radio de una bola cerrada centrada en Θ y contenida en C . Entonces para cualquier $x \in E$, $\varepsilon x / \|x\| \in C$ y, por lo tanto, $p(\varepsilon x / \|x\|) \leq 1$. De aquí, por ii) se tiene que $p(x) \leq (1/\varepsilon)\|x\|$. Esto prueba que p es continuo en Θ . Sin embargo, por iii), tenemos que

$$p(x) = p(x - y + y) \leq p(x - y) + p(y)$$

y

$$p(y) = p(y - x + x) \leq p(y - x) + p(x).$$

De aquí concluimos que

$$-p(y - x) \leq p(x) - p(y) \leq p(x - y),$$

de donde se deduce la continuidad sobre E de la continuidad en Θ .

v) Se obtiene inmediatamente de iv). □

T 4. (Teorema de Mazur; forma geométrica del Teorema de Hahn–Banach) Sea E un espacio vectorial normado real y C un conjunto convexo en E con interior no vacío. Supongamos que V es una variedad lineal en E que no contiene puntos interiores de C . Entonces

existe un hiperplano cerrado H en E que contiene a V pero que no contiene puntos interiores de C ; es decir, existe un elemento $f \in E^*$ y una constante k tales que

$$\begin{aligned} f(v) &= k \text{ para todo } v \in V, \\ f(x) &< k \text{ para todo } x \in C. \end{aligned}$$

Demostración. Haciendo una traslación apropiada podemos suponer que $\Theta \in C^0$. Sea M el subespacio vectorial de E generado por V . Entonces V es un hiperplano en M que no contiene al Θ ; luego existe un funcional lineal continuo f_0 sobre M tal que

$$V = \{x \in M : f_0(x) = 1\}.$$

Sea p el funcional de Minkowsky de C . Como V no contiene puntos interiores de C , tendremos que

$$f_0(x) = 1 \leq p(x) \text{ para todo } x \in V.$$

Por homogeneidad, $f_0(\lambda x) = \lambda \leq p(\lambda x)$ para $x \in V$ y $\lambda > 0$. Además, para $\lambda < 0$, $f_0(\lambda x) \leq 0 \leq p(\lambda x)$. Entonces

$$f_0(x) \leq p(x) \text{ para todo } x \in M.$$

Por el Teorema de Hahn-Banach (T 2), existe una extensión f de f_0 desde M hasta E que cumple $f(x) \leq p(x)$ para todo $x \in E$. Sea $H = \{x \in E; f(x) = 1\}$.

Como $f(x) \leq p(x)$ sobre E y ya que p es continuo por el Lema 1, tendremos que f es continuo y además $f(x) < 1$ para todo $x \in C^0$. Luego H es el hiperplano cerrado buscado. $\square$

Existen varios corolarios y modificaciones de este importante teorema de los cuales nosotros veremos el siguiente.

T 5. (teorema de separación de Eidelheit) Sean C_1 y C_2 conjuntos

convexos en el espacio vectorial normado E , tales que C_1 tiene interior no vacío y C_2 no contiene puntos interiores de C_1 . Entonces existe un hiperplano cerrado H separando C_1 y C_2 , es decir, existe $f \in E^*$ tal que

$$\sup_{x \in C_1} f(x) \leq \inf_{x \in C_2} f(x).$$

En otras palabras, C_1 y C_2 están en semiespacios opuestos determinados por el hiperplano H .

Demostración. Sea $C = C_1 - C_2$; entonces C contiene un punto interior y $0 \notin C^0$. Del Teorema 4 no resulta difícil concluir que existe un $f \in E^*$ no nulo, tal que $f(x) \leq 0$ para todo $x \in C$ (¡verifíquelo!). Entonces para $x_1 \in C_1, x_2 \in C_2$,

$$f(x_1) \leq f(x_2).$$

Consecuentemente, existe un número real tal que

$$\sup_{x_1 \in C_1} f(x_1) \leq C \leq \inf_{x_2 \in C_2} f(x_2).$$

Luego, el hiperplano deseado será $H = \{x \in E : f(x) = C\}$. $\square$

§ 2.3 Algunos teoremas de representación de funcionales lineales y bilineales continuos.

En este epígrafe estudiaremos cuatro resultados que son de utilidad en las aplicaciones:

- I. La representación de los funcionales lineales continuos en un espacio de Hilbert.

- II. La relación entre los funcionales bilineales continuos y los operadores lineales continuos en un espacio de Hilbert.
- III. La representación de los funcionales lineales continuos en espacios de funciones continuas $(C([a, b], \mathbb{C}), \|\cdot\|_\infty)$.
- IV. La representación de los funcionales lineales continuos sobre $L_p(\Omega)$ ($p > 1$).

T 1. (F. Riesz) Sea $(H, \langle, \rangle)$ un espacio de Hilbert. Entonces para todo funcional lineal continuo $f \in H^*$ existe un único vector $a \in H$ tal que

$$f(x) = f_a(x) = \langle x, a \rangle, \quad (x \in H). \quad (31)$$

Más aún $\|f\| = \|a\|$ y todo $a \in H$ determina un único funcional lineal continuo mediante (31). En otras palabras, para todo espacio de Hilbert, la aplicación

$$H \rightsquigarrow H^*, \quad a \rightsquigarrow f_a \quad (32)$$

define una isometría de H sobre H^* .

Demostración. La segunda parte del teorema es una consecuencia de los Ejemplos 2.1. (18) y 2.9 (26). Pasemos entonces a demostrar (31). Sea $f \in H^*$ y denotemos por H_0 al conjunto de todos los $y \in H$ tales que $f(y) = 0$. Por la linealidad de f , H_0 será un subespacio vectorial de H . Además por ser f continuo, H_0 es un subespacio vectorial cerrado; en efecto $H_0 = f^{-1}[\{0\}]$. Si $H_0 = E$, entonces $f \equiv 0$ y el teorema estaría probado tomando $a = 0$.

Si $H \neq H_0$ entonces, de acuerdo al Teorema 1.6.6, podemos escribir $H = H_0 \oplus H_0^\perp$. Como $H \neq H_0$, existe un vector no nulo $z \in H_0^\perp$. Como z es no nulo y $z \notin H_0$, necesariamente $f(z) \neq 0$. Del hecho que H_0 es un subespacio vectorial de H , podemos elegir z de manera conveniente de forma tal que $f(z) = 1$. Probaremos que z es un múltiplo escalar del vector buscado a .

Dado cualquier $x \in H$, tenemos que $x - f(x)z \in H_0$, y como $z \in H_0^\perp$, tendremos que $\langle x - f(x)z, z \rangle = 0$, es decir $\langle x, z \rangle = f(x)\|z\|^2$, de donde $f(x) = \langle x, z/\|z\|^2 \rangle$. Luego, definiendo $a = z/\|z\|^2$ se tiene que $f(x) = \langle x, a \rangle$. Con esto demostramos la existencia de a que cumple (32)

El vector a es claramente único, ya que si b es otro vector tal que $f(x) = \langle x, b \rangle$ para todo $x \in H$, entonces $\langle x, b \rangle = f(x) = \langle x, a \rangle$, y por lo tanto $\langle x, b - a \rangle = 0$ para todo $x \in H$. De aquí que $b - a \in H^\perp$, de donde $b - a = \Theta$, lo cual implica que $b = a$. $\square$

Utilizando el teorema anterior se puede establecer una relación entre operadores y funcionales bilineales continuos en un espacio de Hilbert. Comencemos por la siguiente definición.

D 1. (funcional bilineal) Un funcional bilineal sobre el espacio vectorial E es una aplicación $h : E \times E \rightarrow \mathbb{C}$ que cumple:

$$h(x + y, z) = h(x, z) + h(y, z) \quad (33)$$

$$h(\lambda x, y) = \lambda h(x, y); \quad (34)$$

$$h(x, y) = \overline{h(y, x)} \quad (35)$$

para todos $x, y, z \in E$ y $\lambda \in \mathbb{C}$, donde la raya sobre h en (35) denota el paso al número complejo conjugado.

La siguiente proposición es análoga a 2.1.2.

P 1. Sea $(E, \|\cdot\|)$ un espacio vectorial normado y h un funcional bilineal sobre E . Entonces son equivalentes:

- i) h es continuo (globalmente);
- ii) h es continuo en un punto $(a, b) \in E \times E$;
- iii) para todo conjunto acotado C en $E \times E$ existe una constante $M > 0$ tal que si $(x, y) \in C$

$$|h(x, y)| \leq M;$$

iv) si B es la bola unidad en E ($B = \{x \in E / \|x\| \leq 1\}$), entonces h está acotada sobre $B \times B$;

v) existe una constante $M > 0$ tal que para todo $x, y \in E$

$$|h(x, y)| \leq M \|x\| \|y\|.$$

Demostración. Se deja de ejercicio al lector. Indicación: seguir la idea de las demostraciones dadas en el Teorema 2.1.2. $\square$

Del punto v) de la proposición anterior se deduce que se puede definir la norma de un funcional bilineal h mediante

$$\|h\| = \sup_{\substack{x, y \in E \\ x \neq 0 \\ y \neq 0}} \frac{|h(x, y)|}{\|x\| \|y\|} \quad (36)$$

P 2. Sea $(H, \langle, \rangle)$ un espacio euclideo y A un operador lineal de H en sí mismo. Entonces la aplicación

$$h : H \times H \longrightarrow \mathbb{C},$$

definida por

$$h(x, y) = \langle Ax, y \rangle, \quad (37)$$

es un funcional bilineal. Además h es continuo si y sólo si A es continuo y

$$\|h\| = \|A\|. \quad (38)$$

Demostración. La bilinealidad de h es una consecuencia directa de la linealidad de A y las propiedades del producto escalar. De (37) y la continuidad del producto escalar se prueba inmediatamente que h

es continuo si A lo es. Recíprocamente si h es continuo, por (P 1 - v) existe una constante $M > 0$ tal que para todos $x, y \in H$

$$|h(x, y)| = |\langle Ax, y \rangle| \leq M \|x\| \|y\|.$$

Tomando $y = Ax$, tendremos

$$\|Ax\|^2 \leq M \|x\| \|Ax\|,$$

de donde

$$\|Ax\| \leq M \|x\|,$$

y de aquí concluimos la continuidad de A . La demostración de (38) se deduce inmediatamente de la definición y se deja de ejercicio al lector. $\square$

El siguiente teorema de representación es muy importante para la solución de ecuaciones lineales en espacios de Hilbert.

T 2 (II) (Lax–Milgram) Para cada funcional bilineal continuo h definido sobre un espacio de Hilbert existe un único operador lineal y continuo $A : H \rightarrow H$ que *representa* a h en el sentido de (37) y además se cumple (38).

Demostración. Para cada $y \in H$ fijo, la aplicación

$$x \rightsquigarrow h(x, y)$$

es un funcional lineal continuo sobre H . Por el teorema de representación de F. Riesz (T 1), existe un vector $a \in H$ tal que para todo $x \in H$,

$$h(x, y) = \langle x, a \rangle.$$

Como el vector a está unívocamente determinado por y , podemos establecer la correspondencia

$$y \rightsquigarrow a.$$

Denotado al vector a correspondiente a y , por $A(y)$, demostremos que la aplicación así definida $A : H \longrightarrow H$ es un operador lineal. En efecto,

$$h(x, y) = \langle x, A(y) \rangle$$

para todos $x, y \in H$. Luego

$$h(x, \lambda y + \mu z) = \langle x, A(\lambda y + \mu z) \rangle \quad (39)$$

para todos $x, y, z \in H; \lambda, \mu \in K (K = \mathbb{R} \text{ ó } \mathbb{C})$. Por otra parte de la bilinealidad de h se tiene que

$$h(x, \lambda y + \mu z) = \bar{\lambda}h(x, y) + \bar{\mu}h(x, z) \quad (40)$$

$$\begin{aligned} &= \bar{\lambda}\langle x, A(y) \rangle + \bar{\mu}\langle x, A(z) \rangle \\ &= \langle x, \lambda A(y) + \mu A(z) \rangle. \end{aligned} \quad (41)$$

De (39) y (41) se deduce que

$$\langle x, A(\lambda y + \mu z) + \lambda A(y) + \mu A(z) \rangle = 0$$

para todo $x \in H$, de donde se concluye que

$$A(\lambda y + \mu z) = \lambda A(y) + \mu A(z),$$

o sea A es lineal. La tesis del teorema se obtiene ahora aplicando el resultado de la Proposición 2. $\square$

Veamos cómo se aplican los resultados sobre representación de funcionales lineales continuos a la solución de ecuaciones en espacios de Hilbert.

Ejemplo 1. (soluciones débiles de ecuaciones operacionales en espacios de Hilbert) Sea A un operador acotado en el espacio de Hilbert $(H, \langle, \rangle)$ y sea $\mathcal{H}$ un subespacio vectorial cerrado de H . Diremos que $x \in \mathcal{H}$ es una *solución débil* de la ecuación

$$x + A^2x = y, \quad (y \in H) \quad (42)$$

con respecto a $\mathcal{H}$, si para todo $z \in \mathcal{H}$ se cumple que

$$\langle x, z \rangle + \langle Ax, Az \rangle = \langle y, z \rangle. \quad (43)$$

Demostremos la existencia de soluciones débiles para la ecuación (42). Para ello consideremos $\mathcal{H}$ provisto del nuevo producto escalar

$$[x, y] = \langle x, y \rangle + \langle Ax, Ay \rangle \quad (44)$$

La comprobación de que (44) define un producto escalar en $\mathcal{H}$ se deja de ejercicio al lector. La norma asociada al producto escalar (44) es

$$\|x\|_{\mathcal{H}} = \sqrt{[x, x]} = \sqrt{\|x\|^2 + \|A(x)\|^2}$$

luego, para cada $y \in H$ fijo y todo $z \in \mathcal{H}$ se tiene:

$$|\langle y, z \rangle| \leq \|y\| \|z\| \leq \|y\| \|z\|_{\mathcal{H}}.$$

De esta desigualdad se tiene que para cada y fijo, la aplicación

$$z \rightsquigarrow \langle y, z \rangle$$

es un funcional lineal y continuo en el espacio euclideo $(\mathcal{H}, [\cdot, \cdot])$. El lector puede comprobar sin mucha dificultad por la continuidad de A y por ser $\mathcal{H}$ cerrado en H , que $(\mathcal{H}, [\cdot, \cdot])$ es también un espacio de Hilbert. Luego por el teorema de representación de Riesz (1), para cada $y \in H$ existe $x \in \mathcal{H}$ único tal que para todo $z \in \mathcal{H}$ $[x, z] = \langle y, z \rangle$. Esto último es precisamente lo que queríamos demostrar: x es la solución débil de la ecuación (42) con respecto a $\mathcal{H}$.

En la práctica, en ocasiones se trata de elegir el subespacio $\mathcal{H}$ convenientemente de manera que la solución débil (la cual siempre existe) coincida con la solución de (42) en sentido usual.

Esta técnica de encontrar soluciones de ecuaciones a partir de sus soluciones débiles es muy utilizada en la teoría de los problemas de contorno elípticos para ecuaciones diferenciales en derivadas parciales. Para

más información al respecto consulte el libro de V.P. Mijailov *Ecuaciones diferenciales en derivadas parciales* [6]. $\square$

Pasemos ahora a estudiar el problema (III) enunciado en la introducción al § 2.3. Comencemos por las definiciones siguientes.

D 2. (funciones de variación acotada, variación total) Una función compleja f definida sobre el intervalo $[a, b]$ se llama de *variación acotada*, si existe una constante C , tal que cualquiera que sea la partición

$$a = x_0 < x_1 < \dots < x_n = b,$$

del intervalo $[a, b]$, se cumple la desigualdad

$$\sum_{k=1}^n |f(x_k) - f(x_{k-1})| \leq C. \quad (45)$$

Se llama *variación total* en $[a, b]$ de la función de variación acotada f , al número

$$V_a^b[f] = \sup \sum_{k=1}^n |f(x_k) - f(x_{k-1})|, \quad (46)$$

donde el supremo se toma en todas las particiones finitas del intervalo $[a, b]$.

D 3. (integral de Riemann–Stieltjes) Sea ϕ una función continua a la izquierda y de variación acotada en $[a, b]$ y sea f una función compleja arbitraria, definida sobre el intervalo $[a, b]$.

Para cualquier partición finita $a = x_0 < x_1 < \dots < x_n = b$ del intervalo $[a, b]$ en subintervalos $[x_{i-1}, x_i]$, elijamos $\mathcal{E}_i \in [x_{i-1}, x_i]$ y confeccionemos la suma

$$\sum_{k=1}^n f(\mathcal{E}_i) [\phi(x_i) - \phi(x_{i-1})] \quad (47)$$

(por $\phi(x_n)$ se entiende el límite por la izquierda $\phi(b-0)$).

Si estas sumas (47) convergen a un cierto límite, cuando $\max(x_i - x_{i-1}) \rightarrow 0$, que no depende de la partición elegida ni del punto $\mathcal{E}_i$ seleccionado en $[x_{i-1}, x_i[$, se dice que la función f es *integrable en el sentido de Riemann-Stieltjes* en el intervalo $[a, b]$ con respecto a la función de variación acotada ϕ . Al valor del límite se le llama *integral de Riemann-Stieltjes* de la función f en el intervalo $[a, b]$ con respecto a la función de variación acotada ϕ y se denota

$$\int_a^b f(x) d\phi(x) \quad (48)$$

No es difícil comprobar que la integral (48) es un funcional lineal de f , definido sobre el espacio vectorial de todas las funciones integrables en sentido de Riemann-Stieltjes con respecto a ϕ en $[a, b]$.

Además, toda función continua sobre $[a, b]$ es obviamente integrable en sentido de Riemann-Stieltjes sobre $[a, b]$ con respecto a cualquier función de variación acotada ϕ . De la desigualdad evidente

$$\left| \int_a^b f(x) d\phi(x) \right| \leq V_a^b[\phi] \|f\|_\infty, \quad (49)$$

se concluye que *el funcional lineal (48) con ϕ fijo, es continuo en*

$$(C([a, b], \mathbb{C}), \|\cdot\|_\infty)$$

y su norma esta acotada por $V_a^b[\phi]$. Más exactamente se tiene el teorema siguiente:

T 3. (III) (F. Riesz) Todo funcional lineal continuo F definido en el espacio $(C([a, b], \mathbb{C}), \|\cdot\|_\infty)$ se representa de manera única en la forma

$$F(f) = \int_a^b f(x) d\phi(x), \quad (50)$$

donde $\phi(x)$ es una cierta función continua a la izquierda y de variación acotada sobre $[a, b]$. Además se cumple que

$$|||F||| = V_a^b[\phi]. \quad (51)$$

Demostración. Sea F un funcional lineal y continuo definido sobre $(C([a, b], \mathbb{C}), \|\cdot\|_\infty)$. Por el Teorema de Hahn-Banach (Corolario 2.2.1), F puede ser prolongado conservando su norma, al espacio $(K([a, b], \mathbb{C}), \|\cdot\|_\infty)$ de las funciones complejas definidas sobre $[a, b]$ y con un número finito de discontinuidades de primer orden. En particular, esa extensión de F (que también denotaremos por F) estará definida sobre todas las funciones del tipo

$$h_a(x) = 0, \quad h_\tau(x) = \begin{cases} 1 & \text{si } x \leq \tau, \\ 0 & \text{si } x > \tau, \text{ cuando } \tau > a. \end{cases} \quad (52)$$

Pongamos

$$\phi(\tau) = F(h_\tau), \quad (53)$$

y mostremos que ϕ es una función de variación acotada en el intervalo $[a, b]$. En efecto, tomemos una partición arbitraria

$$a = x_0 < x_1 < \cdots < x_n = b$$

del intervalo $[a, b]$ y designemos

$$\alpha_k = \operatorname{sgn}(\phi(x_k) - \phi(x_{k-1})), \quad k = 1, \dots, n.$$

Entonces

$$\begin{aligned} \sum_{k=1}^n |\phi(x_k) - \phi(x_{k-1})| &= \sum_{k=1}^n \alpha_k (\phi(x_k) - \phi(x_{k-1})) \\ &= \sum_{k=1}^n \alpha_k F(h_{x_k} - h_{x_{k-1}}) \end{aligned}$$

$$\begin{aligned}
&= F\left(\sum_{k=1}^n \alpha_k (h_{x_k} - h_{x_{k-1}})\right) \\
&\leq \|F\| \cdot \left\| \sum_{k=1}^n \alpha_k (h_{x_k} - h_{x_{k-1}}) \right\|.
\end{aligned}$$

Pero la función $\sum_{k=1}^n \alpha_k (h_{x_k} - h_{x_{k-1}})$, pertenece a $K([a, b], \mathbb{C})$ y toma solamente los valores 0 y 1. Luego:

$$\sum_{k=1}^n |\phi(x_k) - \phi(x_{k-1})| \leq \|F\|. \quad (54)$$

De la validez de (54) para cualquier partición del intervalo $[a, b]$, se tiene que

$$V_a^b[\phi] \leq \|F\|. \quad (55)$$

Así pues, hemos construido una función de variación acotada ϕ a partir del funcional F . Mostremos que precisamente con ayuda de ϕ , se puede expresar F en la forma de una integral de Riemann–Stieltjes (50). Sea f una función continua arbitraria definida sobre $[a, b]$. Sabemos que f es uniformemente continua y por lo tanto, dado $\varepsilon > 0$, se puede elegir $\delta > 0$, tal que $|f(x'') - f(x')| < \varepsilon$ cuando $|x' - x''| < \delta$.

Elijamos ahora la partición de $[a, b]$ de manera tal, que la longitud de cada subintervalo $[x_{i-1}, x_i]$ sea menor que δ , y examinemos la función escalonada f_ε :

$$\begin{aligned}
f_\varepsilon(x) &= f(x_k) \quad \text{para } x_{k-1} < x \leq x_k, k = 1, 2, \dots, n, \\
f_\varepsilon(a) &= f(a).
\end{aligned}$$

Esta función puede ser escrita en la forma

$$f_\varepsilon(x) = \sum_{k=1}^n f(x_k) [h_{x_k}(x) - h_{x_{k-1}}(x)],$$

donde h_τ es la función descrita en (52). Es claro que $\|f(x) - f_\varepsilon(x)\|_\infty < \varepsilon$. Calculemos el valor de $F(f_\varepsilon)$. En virtud de la linealidad del funcional F y la definición de h_τ , se tiene

$$\begin{aligned} F(f_\varepsilon) &= \sum_{k=1}^n f(x_k)[F(h_{x_k}) - F(h_{x_{k-1}})] \\ &= \sum_{k=1}^n f(x_k)[\phi(x_k) - \phi(x_{k-1})], \end{aligned}$$

es decir, $F(f_\varepsilon)$ es una suma integral para la integral de Riemann Stieltjes

$$\int_b^a f(x) d\phi(x).$$

Por ello, para una partición suficientemente pequeña del intervalo $[a, b]$ se tendrá que

$$|F(f_\varepsilon) - \int_a^b f(x) d\phi(x)| < \varepsilon.$$

Pero, a la vez

$$|F(f) - F(f_\varepsilon)| \leq \|F\| \|f - f_\varepsilon\|_\infty \leq \|F\| \varepsilon.$$

Por lo tanto

$$|F(f) - \int_a^b f(x) d\phi(x)| < \varepsilon(1 + \|F\|),$$

de donde en virtud de la arbitrariedad de ε , se tiene la igualdad

$$F(f) = \int_a^b f(x) d\phi(x)$$

que queríamos probar. La igualdad (51) se obtiene de (49) y (55), y la unicidad de la representación (50) se deja de ejercicio al lector. $\square$

Observación. Del teorema anterior se concluye inmediatamente que el espacio $V([a, b], \mathbb{C})$ de las funciones complejas de variación acotada y continuas a la izquierda es un espacio de Banach provisto de la norma $V_a^b[\phi]$.

La desigualdad de Hölder (84) (Capítulo 1), nos muestra que si $p > 1$, entonces toda función $g \in L_q(\Omega)$, donde $\frac{1}{p} + \frac{1}{q} = 1$, define un funcional lineal y continuo

$$F_g(f) = \int_{\Omega} fg dx$$

sobre $L_p(\Omega)$, a través de la integral de Lebesgue. Más aún, la desigualdad de Hölder nos muestra que

$$|||F_g||| = \|g\|_q.$$

Se tiene el teorema siguiente:

T 4. (IV) (dual de $L_p(\Omega)$) Para cada $p > 1$ y $F \in (L_p(\Omega))^*$ existe una única función $g \in L_q(\Omega)$ (con $\frac{1}{p} + \frac{1}{q} = 1$), tal que

$$F(f) = \int_{\Omega} fg dx, \text{ para toda } f \in L_p(\Omega).$$

Además $|||F||| = \|g\|_q$. En otras palabras, se tiene la identificación: $(L_p(\Omega))^* = L_q(\Omega)$ en donde $(\frac{1}{p} + \frac{1}{q} = 1)$.

Demostración. Véase el libro de Walter Rudin *Real and Complex Analysis* [3]. $\square$

§ 2.4 Diferentes tipos de convergencia para operadores y funcionales lineales continuos

Para sucesiones de operadores lineales continuos no resulta interesante estudiar la convergencia uniforme sobre todo el espacio de partida. En efecto, si la sucesión $A_n : E \rightarrow F$ de operadores lineales continuos, converge uniformemente sobre E al operador lineal continuo idénticamente nulo: $A(x) = \Theta$ para todo $x \in E$; entonces para n suficientemente grande los conjuntos $A_n[E]$ están acotados en F . Pero, por ser A_n lineal, $A_n[E]$ es un subespacio vectorial de F , por lo tanto:

$$A_n[E] = \{\Theta\}, \text{ para } n \geq n_0.$$

Luego, de la convergencia a $\langle \Theta \rangle$ uniformemente sobre todo E de la sucesión (A_n) , se concluye que $A_n = \langle \Theta \rangle$ a partir de cierto n_0 . Estudiemos entonces tres tipos de convergencia para sucesiones de operadores lineales continuos y funcionales lineales continuos: *la convergencia uniforme sobre la bola unidad del espacio de partida*, *la convergencia puntual* y *la convergencia débil*. Las dos primeras se definen de una manera bastante natural y la tercera es de gran importancia en las aplicaciones al Análisis Funcional. A diferencia de la convergencia débil y la puntual, veremos que la convergencia uniforme sobre la bola unidad puede ser definida por una norma.

T 1. (convergencia uniforme sobre la bola unidad) Sean E y F espacios normados y (A_n) una sucesión de operadores lineales continuos de E en F . Si la sucesión (A_n) converge uniformemente sobre la bola $B_E = \{x \in E : \|x\|_E \leq 1\}$ a una aplicación $A : B_E \rightarrow F$, entonces A puede ser prolongado de manera única, de B_E a todo E , a un operador lineal y continuo. Luego la convergencia uniforme sobre B_E coincide con la convergencia definida en $\mathcal{L}(E, F)$ por la norma $||| \cdot |||$ (19).

Demostración. Sean $A_n : E \rightarrow F$ como en el enunciado del teo-

rema, y supongamos que (A_n) converge a la aplicación $A : B_E \rightarrow F$.

Para cada $x \in E$, $x/\|x\|_E \in B_E$ y, por lo tanto,

$$A_n(x/\|x\|_E) \xrightarrow{F} A(x/\|x\|_E).$$

De la linealidad de A_n se tiene que

$$A_n(x) \xrightarrow{F} \|x\|_E A(x/\|x\|_E), \text{ para todo } x \in E.$$

Definamos

$$\tilde{A}(x) = \|x\|_E A(x/\|x\|_E), \text{ para todo } x \in E.$$

Entonces, para todos $\lambda, \mu \in \mathbb{C}$, $x, y \in E$, se tiene:

$$A_n(\lambda x + \mu y) \xrightarrow{F} \|\lambda x + \mu y\|_E A\left(\frac{\lambda x + \mu y}{\|\lambda x + \mu y\|_E}\right), \quad (56)$$

$$\begin{aligned} A_n(\lambda x + \mu y) &= \lambda A_n(x) + \mu A_n(y) \\ &\xrightarrow{F} \lambda \|x\|_E A\left(\frac{x}{\|x\|_E}\right) + \mu \|y\|_E A\left(\frac{y}{\|y\|_E}\right), \end{aligned} \quad (57)$$

De (56) y (57) se tiene que

$$\tilde{A}(\lambda x + \mu y) = \lambda \tilde{A}(x) + \mu \tilde{A}(y)$$

de donde se concluye la linealidad de $\tilde{A}$. Para probar la continuidad de A , notemos que, para todo $x \in B_E$:

$$\|\tilde{A}(x)\|_F = \|A(x)\|_F \leq \|A(x) - A_n(x)\|_F + \|A_n(x)\|_F.$$

Seleccionemos n_0 suficientemente grande, para que $\|A(x) - A_{n_0}(x)\|_F \leq 1$. Entonces, para todo $x \in B_E$:

$$\|\tilde{A}(x)\|_F \leq 1 + \|A_{n_0}\|_F,$$

de donde, por T2.1.2 iv) se concluye la continuidad de $\tilde{A}$. $\square$

Pasemos ahora a estudiar las propiedades de la convergencia puntual para sucesiones de operadores lineales continuos. El resultado fundamental que enunciaremos es el llamado Teorema de S. Banach y H. Steinhauss. Comencemos por el lema auxiliar siguiente:

L 1. Sean E y F espacios normados y $A_n : E \rightarrow F$ una sucesión de operadores lineales continuos. Si la sucesión (A_n) converge puntualmente en E a alguna aplicación $A : E \rightarrow F$, entonces A es lineal. Si además las normas $\|A_n\|$, están uniformemente acotadas, entonces el operador límite también es continuo.

Demostración. La linealidad de A se demuestra fácilmente y por ello se deja de ejercicio. Ahora, teniendo en cuenta que $\|A_n\| \leq M$ para todo n , donde M es una constante positiva, y pasando al límite en la desigualdad

$$\|A_n(x)\|_F \leq M\|x\|_E,$$

llegamos a que

$$\|A(x)\|_F \leq M\|x\|_E,$$

de donde, por T2.1.2 (v), se tiene la continuidad de A . $\square$

Veamos ahora bajo qué condiciones se puede asegurar la acotación uniforme de las normas de una sucesión de operadores lineales continuos.

T 2. (principio de acotación uniforme) Sean E un espacio de Banach, F un espacio normado y $A_\alpha : E \rightarrow F$ ($\alpha \in I$) una familia de operadores lineales y continuos. Si los conjuntos

$$\{A_\alpha(x) : \alpha \in I\} \tag{58}$$

son acotados en F para cada $x \in E$, entonces es también acotado el

conjunto $\{\|A_\alpha\| : \alpha \in I\}$.

Demostración. Sea $\varepsilon > 0$, y para cada entero positivo k pongamos

$$E_k = \{x \in E : \sup_{\alpha \in I} \|\frac{1}{k} A_\alpha(x)\|_F \leq \varepsilon\}.$$

Por la continuidad de A_α , cada conjunto E_k es cerrado. Del hecho que cada conjunto (58) es acotado, se deduce que

$$E = \bigcup_{k=1}^{\infty} E_k.$$

Por ser E completo, es de segunda categoría, por lo tanto, para algún natural k_0 , el conjunto E_{k_0} debe contener una vecindad $U = x_0 + \delta B$ de cierto punto $x_0 \in E$, donde B es la bola de radio 1 centrada en Θ y $\delta > 0$. Pero para $x \in B$ tendremos que

$$\sup_{\alpha \in I} \|\frac{1}{k_0} A_\alpha(x_0 + \delta x)\|_F \leq \varepsilon.$$

Luego

$$\begin{aligned} \|A_\alpha(\delta x/k_0)\|_F &= \|A_\alpha(\frac{1}{k_0}(x_0 + \delta x - x_0))\|_F \\ &\leq \|\frac{1}{k_0} A_\alpha(x_0 + \delta x)\|_F + \|\frac{1}{k_0} A_\alpha(x_0)\|_F \leq 2\varepsilon, \end{aligned}$$

para todo $x \in B, \alpha \in I$. De aquí se concluye finalmente que para todo $\alpha \in I$, $\|A_\alpha\| \leq 2\varepsilon k_0/\delta$, para todo $\alpha \in I$. $\square$

T 3. (Banach–Steinhaus) Sean (A_n) una sucesión de operadores lineales continuos definidos sobre un espacio de Banach E y con valores

en un espacio normado F . Si (A_n) converge puntualmente en E a una aplicación $A : E \rightarrow F$, entonces A es lineal y continuo.

Demostración. Por L1 basta probar la acotación uniforme de las normas $|||A_n|||$, la cual se tiene aplicando T2, ya que las sucesiones convergentes $\{A_n(x)\}(x \in E)$ son acotadas en F . $\square$

C 1. Para que una sucesión (A_n) de operadores lineales continuos definidos sobre un espacio de Banach E y con valores en un espacio normado F , converja puntualmente al operador lineal y continuo A , es necesario y suficiente que:

- i) La sucesión $\{|||A_n|||\}$ sea acotada;
- ii) $A_n(x) \rightarrow A(x)$ en F , para cualquier x en un subconjunto total en E .

Demostración. Ejercicio. $\square$

El ejemplo siguiente es una aplicación importante del principio de acotación uniforme.

Ejemplo 1. Sea f una función compleja medible en el abierto acotado $\Omega \subset \mathbb{R}^n$. Si la integral $\int_{\Omega} g \bar{f} dx$ existe para cada $g \in L_2(\Omega)$, entonces $f \in L_2(\Omega)$. En efecto para cada $n \in \mathbb{N}$, definamos

$$f_n(x) = \begin{cases} f(x) & \text{si } |f(x)| \leq |n| \\ 0 & \text{si } |f(x)| > |n|. \end{cases}$$

Las funciones $f_n \in L_2(\Omega)$, definen una sucesión de funcionales lineales continuos sobre $L_2(\Omega)$, mediante

$$L_n(g) = \int_{\Omega} g \bar{f}_n dx.$$

La sucesión (f_n) converge puntualmente casi dondequiera sobre Ω a la función medible f . Por la convergencia de la integral $\int_{\Omega} g \bar{f} dx$ y aplicando el teorema de la convergencia dominada de Lebesgue, se tiene que $L_n(g)$ converge a $\int_{\Omega} \bar{f} dx$ para cada $g \in L_2(\Omega)$. Aplicando el principio de acotación uniforme (T2) concluimos que $\|L_n\| = \|f_n\|_2$ están uniformemente acotados. Pero de aquí se concluye que $f \in L_2(\Omega)$ y $\|f_n\|_2 \rightarrow \|f\|_2$. $\square$

T 4. Sea E y F espacios normados. Si F es de Banach, entonces también lo es $(\mathcal{L}(E, F), \|\cdot\|)$.

Demostración. Sea (A_n) una sucesión de Cauchy en $\mathcal{L}(E, F)$. Como para cualquier $x \in E$ se tiene

$$\begin{aligned} \|A_n(x) - A_m(x)\|_F &= \|(A_n - A_m)(x)\|_F \\ &\leq \|A_n - A_m\| \|x\|_E, \end{aligned}$$

y como además $\|A_n - A_m\| \rightarrow 0$ cuando $n, m \rightarrow \infty$, entonces la sucesión $(A_n(x))$ es una sucesión de Cauchy en F . Por ser F un espacio de Banach, existe el límite $\lim A_n(x)$, para todo $x \in E$.

Como toda sucesión de Cauchy es acotada, entonces $\|A_n\|$ esta acotada uniformemente por una constante $M > 0$. De acuerdo al Lema 1 el operador A definido por la fórmula

$$A(x) = \lim A_n(x),$$

pertenece a $\mathcal{L}(E, F)$.

Probemos que (A_n) converge a A con respecto a la norma $\|\cdot\|$. Para un $\varepsilon > 0$ arbitrario, podemos seleccionar $N \in \mathbb{N}$ tal que $\|A_n - A_m\| < \varepsilon$ para $n, m \geq N$. Esta desigualdad significa que

$$\|A_n(x) - A_m(x)\|_F \leq \varepsilon \|x\|_E, \quad (59)$$

para todo $x \in E$. Pasando al límite en (59) cuando $m \rightarrow \infty$, tendremos que

$$\|A_n(x) - A(x)\|_F \leq \varepsilon \leq \|x\|_E$$

de donde se deduce que

$$|||A_n - A||| \leq \varepsilon,$$

o lo que es lo mismo: A_n converge a A en $(\mathcal{L}(E, F), ||| \cdot |||)$. Con esto queda probado que $(\mathcal{L}(E, F), ||| \cdot |||)$ es completo. $\square$

Utilizando el Teorema 4 se obtiene el siguiente resultado que es muy utilizado en la resolución de ecuaciones operacionales.

T 5. Sean E un espacio de Banach, I el operador identidad de E y A un operador lineal continuo que aplica E en sí mismo tal que $|||A||| < 1$. Entonces existe el operador inverso $(I - A)^{-1}$, el cual es lineal y continuo y se representa en la forma

$$(I - A)^{-1} = \sum_{k=0}^{\infty} A^k \quad (60)$$

donde por A^k se denota la composición de A consigo mismo k veces y la serie converge en la norma $||| \cdot |||$ del espacio de Banach $\mathcal{L}(E, F)$.

Demostración. Demostraremos la existencia y acotación del operador $(I - A)^{-1}$ con lo cual quedará probada su linealidad ya que el inverso de cualquier operador lineal es también un operador lineal (¡compruébelo!).

Como $|||A||| < 1$, tenemos que

$$\sum_{k=0}^{\infty} |||A^k||| \leq \sum_{k=0}^{\infty} |||A|||^k < \infty.$$

Por la completitud de E , aplicando el Teorema 4, tenemos que la serie $\sum_{k=0}^{\infty} A^k$ converge en $(\mathcal{L}(E, E), ||| \cdot |||)$ y por lo tanto su suma representa un operador lineal acotado. Para $n \in \mathbb{N}$ cualquiera tenemos

$$(I - A) \sum_{k=0}^n A^k = \sum_{k=0}^n A^k (I - A) = I - A^{n+1}.$$

Pasando al límite cuando $n \rightarrow \infty$ y tomando en consideración que $|||A^{n+1}||| \leq |||A|||^{n+1} \rightarrow 0$, encontramos

$$(I - A) \sum_{k=0}^{\infty} A^k = \sum_{k=0}^{\infty} A^k (I - A) = I,$$

de donde $(I - A)^{-1} = \sum_{k=0}^{\infty} A^k$, que es lo que queríamos demostrar. $\square$

Pasemos ahora a estudiar el tercer tipo de convergencia de que hablábamos en la introducción de este epígrafe: *la convergencia débil*. Sean E un espacio normado y E^* su dual. Definamos ciertos subconjuntos de E a los que llamaremos *vecindades débiles* de los elementos de E .

D 1. (vecindad débil y abierto débil en E) Sean $f_1, \dots, f_n$ una cantidad finita arbitraria de elementos de E^* , y ε un número real positivo. Cualquier subconjunto de E que contenga a un conjunto de la forma

$$x_0 + \{x : |f_i(x)| < \varepsilon, i = 1, \dots, n\}, \quad (61)$$

se llama una *vecindad débil* de $x_0 \in E$.

A todo subconjunto de E que sea vecindad débil de todos sus puntos en el sentido de la definición anterior se le llama un *abierto débil* en E . No resulta difícil demostrar (ejercicio) que *la familia de todos los conjuntos abiertos débiles de un espacio normado satisface las propiedades*

(A1)–(A3) de la familia Θ_d de los conjuntos abiertos en un espacio métrico (E, d) (ver Capítulo 1).

La familia de los abiertos débiles en el espacio normado E , constituye la llamada *topología débil* en E . Esta topología permite definir una noción de convergencia en E a la que llamaremos *convergencia débil*.

D 2. (convergencia débil en E) Se dice que la sucesión (x_n) del espacio normado E converge débilmente al elemento $x_0 \in E$, si para cada vecindad débil $V = V(x_0)$ de x_0 en E , se puede hallar un número natural $N = N(V)$, tal que $x_n \in V(x_0)$ para todo $n \geq N$. (Compárese con la observación a la Definición 1.1).

P 1. (caracterización de la convergencia débil) La sucesión (x_n) converge débilmente en el espacio normado E al elemento x_0 , si para cualquier $f \in E^*$, la sucesión numérica $\{f(x_n)\}$ converge a $f(x_0)$.

Demostración. En efecto, considerando por simplicidad $x_0 = \Theta$, supongamos que $f(x_n) \rightarrow 0$ en $\mathbb{C}$. Entonces para cada vecindad débil

$$U = \{x : |f_i(x)| < \varepsilon, i = 1, \dots, n\}$$

del punto Θ , se puede hallar $N \in \mathbb{N}$, tal que $x_n \in U$ si $n \geq N$ (para ello es suficientemente elegir N_i , tal que $|f_i(x_n)| < \varepsilon$ cuando $n \geq N_i$ y después poner $N = \max N_i$).

Inversamente, si para cada vecindad débil U de Θ , existe N , tal que $x_n \in U$ cuando $n \geq N$, entonces la condición $f(x_n) \rightarrow 0$ se satisface obviamente, para cualquier $f \in E^*$. $\square$

Observación. La proposición anterior nos muestra que toda sucesión convergente con respecto a la norma de E es también débilmente convergente al mismo límite.

Sin embargo, *solamente en los espacios normados de dimensión finita coinciden ambos tipos de convergencia* (ejercicio). En efecto el lector puede tratar de probar (ejercicio) que en los espacios normados de dimensión infinita hay sucesiones débilmente convergentes que no son convergentes en el sentido de la norma. Por este motivo se acostumbra a llamar *convergencia fuerte* a la convergencia con respecto a la norma en un espacio normado.

T 6. Si (x_n) es una sucesión débilmente convergente en el espacio normado E , entonces existe una constante C , tal que

$$\|x_n\| \leq C, n = 1, 2, \dots.$$

En otras palabras, toda sucesión débilmente convergente en un espacio normado es acotada.

Demostración. Por la Proposición 1, si (x_n) converge débilmente a $x_0 \in E$, entonces $f(x_n) \rightarrow f(x_0)$ para todo $f \in E^*$. Definamos para cada n , el funcional lineal L_n sobre E^* , mediante

$$L_n(f) = f(x_n) \tag{62}$$

De la desigualdad

$$|L_n(f)| \leq \|x_n\| \|f\|$$

se tiene que L_n es continuo y $\|L_n\| \leq \|x_n\|$.

Del Corolario 2.2.2, tenemos que existe $g \in E^*$ con $\|g\| = 1$, tal que $L_n(g) = \|x_n\|$, luego $\|L_n\| = \|x_n\|$. Pero por T4, E^* es un espacio de Banach y del principio de acotación uniforme (T2), concluimos que existe una constante $C > 0$, tal que:

$$\|x_n\| = \|L_n\| \leq C. \quad \square$$

El teorema siguiente resulta útil para comprobar la convergencia débil de una sucesión.

T 7. La sucesión (x_n) del espacio normado E converge débilmente a $x \in E$, si:

- i) $\|x_n\|$ está uniformemente acotada por una constante C ;
- ii) $f(x_n) \rightarrow f(x)$ para todo f perteneciente a un subconjunto total en E^* .

Demostración. Ejercicio. $\square$

De manera análoga puede definirse la *topología débil* en el espacio dual E^* . Para ello observemos primeramente, que *las vecindades de $f \in E^*$, con respecto a la norma $\|\cdot\|$ coinciden precisamente con aquellos subconjuntos de E^* que contienen a algún conjunto de la forma*

$$U_{\varepsilon, A}(f) = f + \{\varphi \in E^* : |\varphi(x)| < \varepsilon, x \in A\} \quad (63)$$

donde A es un subconjunto acotado arbitrario en E y ε un número positivo cualquiera (ejercicio).

Si en lugar de considerar en (63) todos los conjuntos acotados de E , consideramos solamente los subconjuntos finitos $A \subset E$, entonces obtenemos las llamadas *vecindades débiles* en el espacio dual de E^* .

Se llama *abierto débil* en E^* a todo subconjunto de E^* que sea vecindad débil de todos sus puntos. La familia de todos los abiertos débiles de E^* también satisface las propiedades (A1)–(A3) de la familia de los abiertos en un espacio métrico y constituye la llamada *topología débil en el espacio dual*. Al igual que en D2 se puede definir la *convergencia débil de funcionales* en el espacio dual E^* .

Siguiendo la idea de la demostración de T6 y T7 se pueden probar los teoremas siguientes:

T 8. Si (f_n) es una sucesión débilmente convergente de funcionales lineales continuos sobre un espacio de Banach, entonces existe una constante C , tal que

$$|||f_n||| \leq C, \quad n = 1, 2, \dots$$

En otras palabras, toda sucesión débilmente convergente de funcionales lineales continuos sobre un espacio de Banach, es acotada en norma.

Demostración. Ejercicio. $\square$

T 9. La sucesión de funcionales lineales (f_n) de E^* , converge débilmente a $f \in E^*$, si:

- i) la sucesión (f_n) está acotada en E^* , es decir $|||f_n||| \leq C, \quad n = 1, 2, \dots$
- ii) se cumple que $f_n(x) \rightarrow f(x)$, para todo $x \in E$ perteneciente a un subconjunto total de E .

Demostración. Ejercicio. $\square$

En las aplicaciones del concepto de convergencia débil de funcionales lineales continuos, juega un importante papel el teorema siguiente:

T 10. (Banach-Alaughlu) Si E es un espacio normado separable, entonces toda sucesión acotada de funcionales lineales continuas sobre

E , contiene una subsucesión débilmente convergente.

Demostración. Elijamos en E un subconjunto numerable denso $\{x_1, x_2, \dots, x_n, \dots\}$. Si la sucesión (f_n) está acotada en norma, entonces la sucesión numérica

$$f_1(x_1), f_2(x_1), \dots, f_n(x_1), \dots$$

está acotada. Por ello se puede elegir una subsucesión

$$f_1^{(1)}, f_2^{(1)}, \dots, f_n^{(1)}, \dots$$

de la sucesión (f_n) , de manera tal que la sucesión numérica

$$f_1^{(1)}(x_1), f_2^{(1)}(x_1), \dots, f_n^{(1)}(x_1), \dots$$

converja. Con el mismo razonamiento se puede elegir una subsucesión $f_1^{(2)}, f_2^{(2)}, \dots, f_n^{(2)}, \dots$ de la sucesión $(f_n^{(1)})$, tal que converja la sucesión numérica

$$f_1^{(2)}(x_2), f_2^{(2)}(x_2), \dots, f_n^{(2)}(x_2), \dots$$

Continuando este proceso, obtendremos un sistema de sucesiones

$$\begin{array}{ccccccc} f_1^{(1)} & f_2^{(1)} & \dots & f_n^{(1)} & \dots & & \\ f_1^{(2)} & f_2^{(2)} & \dots & f_n^{(2)} & \dots & & \\ & \dots & \dots & \dots & & & \end{array}$$

(cada una de las cuales es subsucesión de la anterior), tal que $(f_n^{(k)})$ converge en los puntos $x_1, x_2, \dots, x_k$. Entonces tomando la sucesión diagonal

$$f_1^{(1)}, f_2^{(2)}, \dots, f_n^{(n)}, \dots$$

obtendremos una subsucesión de la sucesión (f_n) tal que la sucesión numérica $(f_n^{(n)}(x_k))$ es convergente para todo $k = 1, 2, \dots$. Pero entonces,

en virtud del Teorema 9, la sucesión $(f_n^{(n)}(x))$ converge para cualquier $x \in E$. $\square$

Ejemplo 1. Probar que si E es un espacio normado separable y $B^* = \{f \in E^* : |||f||| \leq 1\}$ es la bola unidad en el espacio dual E^* , entonces la familia de los conjuntos abiertos débiles en B^* coincide con la familia de los abiertos definidos en B^* por la métrica

$$d(f, g) = \sum_{n=0}^{\infty} 2^{-n} |(f - g)(x_n)| \quad (64)$$

donde (x_n) es un subconjunto numerable denso en la bola unidad cerrada B de E ,

Obviamente este resultado puede extenderse a todo subconjunto acotado del dual de un espacio normado separable (ejercicio). Por lo tanto, la convergencia débil para sucesiones de funcionales lineales continuos definidos sobre un espacio normado separable corresponde a la convergencia con respecto a una métrica del tipo (64) (ejercicio).

De los comentarios anteriores se deduce el teorema siguiente:

T 11. Sea E un espacio normado separable. Todo subconjunto acotado M de E^* es relativamente compacto en el sentido de cualquier métrica que describa la convergencia débil sobre M .

Demostración. Ejercicio. $\square$

El teorema siguiente es una aplicación importante de los resultados vistos en este epígrafe.

T 12. (teorema de Banach sobre las bases) Toda sucesión biortogonal (f_n) a una base (e_n) de un espacio de Hilbert H es también una

base de H .

Demostración. Todo vector $f \in H$ admite un desarrollo en serie convergente del tipo

$$f = \sum_{k=1}^{\infty} \langle f, f_k \rangle e_k,$$

luego, la serie numérica

$$\langle f, h \rangle = \sum_{k=1}^{\infty} \langle f, f_k \rangle \langle e_k, h \rangle$$

es convergente para todo $h \in H$. Pero entonces

$$\langle f, h \rangle = \lim_{n \rightarrow \infty} \langle f, \sum_{k=1}^n \langle h, e_k \rangle f_k \rangle$$

para cualquier $h \in H$, lo cual significa que la sucesión

$$h_n = Q_n(h) = \sum_{k=1}^n \langle h, e_k \rangle f_k, \quad n = 1, 2, \dots; \quad h \in H,$$

converge débilmente al vector h (T 2.3.1). Pero por T6, tenemos que la sucesión (h_n) está acotada en H . Aplicando el principio de acotación uniforme (T2) a la sucesión de operadores Q_n , tendremos que

$$|||Q_n||| \leq M, \quad n = 1, 2, \dots$$

Por ser (f_n) una sucesión biortogonal a la base (e_n) se tiene que (f_n) es una familia total en H (ejercicio). Entonces, para cada $\varepsilon > 0$ y cualquier $h \in H$, se pueden encontrar números $c_j^\varepsilon \in \mathbb{C}$ ($j = 1, 2, \dots, N_\varepsilon$), tales que

$$\|h - \sum_{j=1}^{N_\varepsilon} c_j^\varepsilon f_j\| < \varepsilon. \quad (65)$$

Luego

$$Q_n(h - \sum_{j=1}^{N_\varepsilon} c_j^\varepsilon f_j) = Q_n(h) - \sum_{j=1}^{N_\varepsilon} c_j^\varepsilon f_j, \quad n > N_\varepsilon$$

y de aquí obtenemos

$$\|Q_n(h) - \sum_{j=1}^{N_\varepsilon} c_j^\varepsilon f_j\| < M\varepsilon, \quad n > N_\varepsilon. \quad (66)$$

De (65) y (66), se tiene que

$$\|Q_n(h) - h\| < (1 + M)\varepsilon,$$

de donde concluimos que cualquier $h \in H$ puede ser desarrollado en una serie convergente del tipo

$$h = \sum_{j=1}^{\infty} c_j f_j$$

cuyos coeficientes c_j se determinan unívocamente por las igualdades $c_j = \langle h, e_j \rangle$. $\square$

Para finalizar este epígrafe veremos una especie de generalización de la noción de convergencia débil para un tipo particular de funcionales lineales: *las funciones generalizadas*, también llamadas distribuciones. La teoría de distribuciones juega un papel fundamental en la teoría de las ecuaciones diferenciales y en el análisis moderno de muchos problemas de la física matemática.

No pretendemos dar aquí una aplicación detallada de esta teoría sino solamente enunciaremos algunos resultados relevantes de la misma. Para las demostraciones y un estudio más profundo de la teoría de distribuciones remitimos al lector al magnífico libro de V.S. Vladimirov, *Ecuaciones de la física matemática* [7].

Para cualquier abierto Ω de $\mathbb{R}^n$ introdujimos en el Capítulo 1, el espacio $C_0^\infty(\Omega)$ de las funciones complejas infinitamente diferenciables sobre

Ω y con soporte compacto contenido en Ω . En la teoría de funciones generalizadas con frecuencia se denota este espacio por $\mathcal{D}(\Omega)$ y se llama *espacio de funciones de prueba*. Definamos una noción de convergencia para las sucesiones de funciones de $\mathcal{D}(\Omega)$ de la siguiente forma. Diremos que la sucesión (φ_n) en $\mathcal{D}(\Omega)$, converge cuando $n \rightarrow \infty$ a la función $\varphi \in \mathcal{D}(\Omega)$ si:

- (I) Para cualquier multi-índice α , la sucesión $\mathcal{D}^\alpha \varphi_n$ converge uniformemente a $\mathcal{D}^\alpha \varphi$, cuando $n \rightarrow \infty$;
- (II) Existe un compacto $K \subset \Omega$, tal que los soportes de todas las funciones φ_n , salvo un número finito, están contenidos en K .

D 3. (funciones generalizadas o distribuciones) Se llama función generalizada (distribución) en Ω a todo funcional lineal y continuo sobre $\mathcal{D}(\Omega)$. El espacio de todas las funciones generalizadas en Ω se denota por $\mathcal{D}'(\Omega)$.

La continuidad de las funciones generalizadas $u \in \mathcal{D}'(\Omega)$ se interpreta en el sentido de que $u(\varphi_n) \rightarrow u(\varphi)$ en $\mathbb{C}$, para toda sucesión (φ_n) que converge a φ en $\mathcal{D}(\Omega)$. La notación (u, φ) se usa para denotar la evaluación de la distribución u sobre la función de prueba φ .

Se dice que dos funciones generalizadas $u_1, u_2 \in \mathcal{D}'(\Omega_1)$ coinciden sobre el abierto $\Omega_1 \subset \Omega$, si para cualquier $\varphi \in \mathcal{D}(\Omega)$, se cumple

$$(u_1, \varphi) = (u_2, \varphi). \quad (67)$$

Se llama *soporte de la función generalizada* $u \in \mathcal{D}'(\Omega)$, y se denota por $\text{supp } u$, al complemento en Ω del conjunto

$$\begin{aligned} & \{x \in \Omega : u \text{ se anula en una vecindad de } x\} = \\ & = \{x \in \Omega : \exists V_x \text{ vecindad abierta de } x, (u, \varphi) = 0 \forall \varphi \in \mathcal{D}(V_x)\} \end{aligned}$$

En el espacio $\mathcal{D}(\Omega)$ se define una *convergencia de tipo débil* para funcionales lineales: la sucesión (u_n) en $\mathcal{D}'(\Omega)$ converge a $u \in \mathcal{D}'(\Omega)$ si

$$u_n(\varphi) = (u_n, \varphi) \rightarrow u(\varphi) = (u, \varphi) \quad (68)$$

para toda $\varphi \in \mathcal{D}(\Omega)$.

Ejemplo 2. (delta de Dirac). Un ejemplo importante de función generalizada en $\mathcal{D}'(\mathbb{R}^n)$ lo constituye la llamada *función delta* $(\delta(x))$ de Dirac, la cual se define de la manera siguiente:

$$(\delta, \varphi) = \varphi(0), \quad \forall \varphi \in \mathcal{D}(\mathbb{R}^n).$$

Es fácil ver que si w_h es un núcleo de promediación (D1.7.1), entonces w_h converge a la δ -función en el sentido definido en (68). En efecto

$$(w_h, \varphi) = \int_{\mathbb{R}^N} w_h \varphi dx = \varphi_h(0) \rightarrow \varphi(0) = (\delta, \varphi),$$

donde la convergencia $\varphi_h(0) \rightarrow \varphi(0)$ se deduce de T1.7.1 ii).

Toda sucesión de funciones definidas en $\mathbb{R}^n$ y convergente a la δ -función en el sentido de la convergencia en $\mathcal{D}'(\mathbb{R}^n)$, se llama *sucesión de unidades aproximativas*. $\square$

Siguiendo el mismo razonamiento que nos condujo en el § 1.7 a la definición de la derivada en sentido de Sobolev de funciones localmente integrables, vemos que para toda distribución $u \in \mathcal{D}'(\Omega)$ se puede definir su *derivada generalizada* de cualquier orden α mediante la igualdad

$$(\mathcal{D}^\alpha u, \varphi) = (-1)^{|\alpha|} (u, \mathcal{D}^\alpha \varphi) \quad (69)$$

(ejercicio) No es difícil demostrar que $\mathcal{D}^\alpha u \in \mathcal{D}'(\Omega)$ para toda $u \in \mathcal{D}'(\Omega)$ y todo multíndice α .

Para $u \in \mathcal{D}'(\Omega)$ y $a(x) \in C^\infty(\Omega)$ se puede definir también la función generalizada au mediante

$$(au, \varphi) = (u, a\varphi). \quad (70)$$

Proponemos al lector probar que $au \in \mathcal{D}'(\Omega)$ para toda $u \in \mathcal{D}'(\Omega)$ y $a \in C^\infty(\Omega)$.

Toda función localmente integrable f en Ω , define una distribución T_f por la fórmula (ejercicio)

$$(T_f, \varphi) = \int_{\Omega} f\varphi dx \quad (71)$$

Hemos visto en el § 1.7 que *no toda función de $L_{1,loc}(\Omega)$ tiene una derivada en sentido de Soboliev, aunque siempre tiene una derivada generalizada, la cual no necesariamente coincide con una función en sentido usual*. Sin embargo, es obvio que *para las funciones derivables en sentido de Soboliev ambos conceptos coinciden*.

Uno de los aspectos más relevantes de la teoría de distribuciones lo constituye la posibilidad de definir el producto de convolución para cualquier par de funciones generalizadas siempre que una de ellas tenga soporte compacto. En efecto, no resulta muy difícil probar que si $v(y) \in \mathcal{D}'(\mathbb{R}_y^n)$ y $\varphi(x, y) \in \mathcal{D}(\mathbb{R}_{x,y}^{n+m})$, entonces $\Psi(x) = (v(y), \varphi(x, y))$ pertenece al espacio $\mathcal{D}(\mathbb{R}_x^n)$. Luego, si $u(x) \in \mathcal{D}'(\mathbb{R}_x^n)$ tiene soporte compacto, entonces puede definirse el producto convolución $u * v$ como la función generalizada en $\mathcal{D}'(\mathbb{R}^n)$:

$$(u * v, \varphi) = (u(x), (v(y), \theta(x)\varphi(x + y))) \quad (72)$$

donde $\theta(x) \in \mathcal{D}(\mathbb{R}_x^n)$ y $\theta(x) = 1$ en una vecindad de $\text{supp } u$. El lector puede verificar (ejercicio) que la definición (72) no depende de la elección de $\theta(x)$ siempre que satisfaga las propiedades anteriores mencionadas.

Pondremos, por definición,

$$v * u = u * v. \quad (73)$$

Se puede probar que *la convolución de dos funciones generalizadas, una de las cuales tiene soporte compacto, es continua con respecto a la convergencia débil de $\mathcal{D}'(\Omega)$* , en cada uno de sus factores. Es decir, si $v_k(y) \rightarrow v(y)$ en $\mathcal{D}'(\mathbb{R}_y^n)$ y $u(x) \in \mathcal{D}'(\mathbb{R}_x^n)$ tiene soporte compacto, entonces $u * v_k \rightarrow u * v$ en $\mathcal{D}'(\mathbb{R}_x^n)$.

Análogamente, si $u_k \rightarrow u$ en $\mathcal{D}'(\mathbb{R}_x^n)$ y existe un compacto $K \subset \mathbb{R}_x^n$, tal que $\text{supp } u_k \subset K$ para todo k , y $\text{supp } u \subset K$, entonces $u_k * v \rightarrow u * v$ en $\mathcal{D}'(\mathbb{R}_x^n)$ cuando $k \rightarrow \infty$.

Obviamente $\mathcal{D}(\mathbb{R}_x^n) \subset \mathcal{D}'(\mathbb{R}_x^n)$ y para cualquier función generalizada $v \in \mathcal{D}'(\mathbb{R}_y^n)$ se tiene que $\varphi * v = (v(y), \varphi(x-y))$ y además $\varphi * v \in C^\infty(\mathbb{R}_y^n)$.

Por este motivo, podemos definir el *promedio de orden h de una función generalizada $v \in \mathcal{D}'(\mathbb{R}_y^n)$* , como la función $u_h(x) \in C^\infty(\mathbb{R}_y^n)$ dada por la fórmula

$$u_h(x) = (u(y), w_h(x-y)), \quad (74)$$

donde w_h es un núcleo de promediación.

De la propiedad de continuidad del producto convolución de distribuciones y del hecho que $w_h \rightarrow \delta(x)$ en $\mathcal{D}'(\mathbb{R}_x^n)$, vemos que u_h converge a $u * \delta$ cuando $h \rightarrow 0$. Pero

$$(u * \delta, \varphi) = (\delta(x), (u(y), \theta(x)\varphi(x+y))) = (u(y), \varphi(y)),$$

y, por lo tanto

$$u * \delta = u. \quad (75)$$

Luego, toda distribución $u \in \mathcal{D}'(\mathbb{R}_x^n)$ puede ser aproximada, en el sentido de (68), por funciones $u_h \in C^\infty(\mathbb{R}_x^n)$. Obviamente de aquí se deduce que u también puede ser aproximada, en el sentido de (68) por funciones de $\mathcal{D}(\mathbb{R}_x^n)$. En efecto, basta tomar $a_h \in \mathcal{D}(\mathbb{R}_x^n)$, $|a_h(x)| \leq 1$, $a_h(x) = 1$ si $|x| \leq 1/h$, y notar que $v_h = a_h u_h \in \mathcal{D}(\mathbb{R}_x^n)$ y además $\lim_{h \rightarrow 0} v_h = \lim_{h \rightarrow 0} u_h$ en $\mathcal{D}'(\mathbb{R}_x^n)$ (ejercicio).

Otra propiedad importante del producto convolución es:

$$D^\alpha(u * v) = (D^\alpha u) * v = u * D^\alpha v. \quad (76)$$

Para generalizar a las distribuciones algunas propiedades como las enunciadas en T 1.7.5 iii)–iv) y T1.7.6, es necesario introducir la noción de *transformada de Fourier para distribuciones*. Resulta que no es posible dar una definición general de transformada de Fourier para toda $u \in \mathcal{D}'(\mathbb{R}^n)$. Por este motivo se define el espacio $S(\mathbb{R}^n)$ de las *funciones de prueba de decrecimiento lento* que ya vimos en el § 1.7. En $S(\mathbb{R}^n)$ se puede definir una noción de convergencia de la manera siguiente: la sucesión $\varphi_k \in S(\mathbb{R}^n)$ converge a $\varphi \in S(\mathbb{R}^n)$ cuando $k \rightarrow \infty$, si para cualquier múltiplo α y cualquier número $p \geq 0$, la sucesión

$$(1 + |x|^p)D^\alpha \varphi_k(x) \rightarrow (1 + |x|^p)D^\alpha \varphi(x) \quad (77)$$

uniformemente en $\mathbb{R}^n$, cuando $k \rightarrow \infty$.

Se llama *espacio de funciones generalizadas de decrecimiento lento* y se denota por $S'(\mathbb{R}^n)$ al conjunto de todos los funcionales lineales sobre $S(\mathbb{R}^n)$, que son continuos con respecto a la noción de límite definida en (77). No es difícil demostrar (ejercicio) que

$$S'(\mathbb{R}^n) \subset \mathcal{D}'(\mathbb{R}^n). \quad (78)$$

Para las distribuciones $u \in S'(\mathbb{R}^n)$, se puede definir la transformada de Fourier mediante la fórmula:

$$(\hat{u}, \varphi) = (u, \hat{\varphi}), \quad \forall \varphi \in S(\mathbb{R}^n). \quad (79)$$

Remitimos al lector a la bibliografía señalada anteriormente, donde se demuestra que *la transformada de Fourier en $S'(\mathbb{R}^n)$ es una aplicación lineal continua y biyectiva*. Más exactamente, se cumple la fórmula de inversión

$$u(-x) = (2\pi)^{-n} \hat{\hat{u}}, \quad \forall u \in S'(\mathbb{R}^n). \quad (80)$$

Finalmente, podemos decir que la transformada de Fourier en $S'(\mathbb{R}^n)$ satisface las propiedades enunciadas en 1.7.2 y T1.7.5 iii)–iv) para la transformada de Fourier en $S(\mathbb{R}^n)$.

§ 2.5 Teoremas de la aplicación abierta y del grafo cerrado

En este epígrafe veremos dos teoremas, que junto con los Teoremas de Hahn–Banach, Banach–Steinhaus y Banach–Alaughlou, constituyen una de las herramientas más potentes del Análisis Funcional. La demostración de estos teoremas depende fuertemente de la propiedad de completitud en los espacios de Banach.

Durante todo este epígrafe consideraremos espacios de Banach E y F y A un operador lineal de E en F . Mediante un desarrollo elemental se puede verificar que si el operador inverso $A^{-1} : A[E] \rightarrow E$ existe, entonces él también es lineal.

T 1. (aplicación abierta) Sea A un operador lineal continuo que aplica de manera sobreyectiva el espacio de Banach E sobre el espacio de Banach F . Entonces A es una aplicación abierta, es decir A transforma subconjuntos abiertos de E en subconjuntos abiertos de F .

Demostración. Desarrollemos la demostración en tres etapas:

i) Comenzaremos por probar que para cualquier vecindad V de Θ en F , existe una vecindad W de Θ en E , tal que $\overline{A[V]} \supset W$.

Escojamos una vecindad V_0 de Θ en E tal que $V_0 - V_0 \subset V$. En efecto, si la bola $B(\Theta, \varepsilon) \subset V$, basta tomar $V_0 = B(\Theta, \varepsilon/2)$. Para todo $x \in E$, es posible seleccionar $n \in \mathbb{N}$ suficientemente grande, tal que

$x/n \in V_0$, es decir $x \in n V_0$. En otras palabras

$$E = \bigcup_{n=1}^{\infty} nV_0,$$

y, por lo tanto

$$F = A[E] = \bigcup_{n=1}^{\infty} nA[V_0].$$

Pero también $F = \bigcup_{n=1}^{\infty} \overline{nA[V_0]}$. Por el teorema de categoría de Baire, alguno de los conjuntos $\overline{nA[V_0]}$ debe contener a un conjunto abierto. Luego $\overline{A[V_0]}$ también debe contener a algún subconjunto abierto W_0 . Entonces

$$\overline{A[V]} \supset \overline{A[V_0] - A[V_0]} \supset \overline{A[V_0] - A[V_0]} \supset W_0 - W_0.$$

El conjunto $W = W_0 - W_0$ es abierto ya que cada conjunto de la forma $\{a\} - W_0 = \{a - y : y \in W_0\}$ es abierto, y

$$W = \bigcup_{a \in W_0} \{a\} - W_0.$$

Además W es una vecindad de Θ en F puesto que $\Theta \in W_0 - W_0 = W$.

ii) Probemos ahora que la imagen por A de cualquier vecindad de Θ en E es una vecindad de Θ en F . Es fácil ver que es suficiente demostrar que la imagen por A de toda bola

$$V(\varepsilon) = \{x \in E : \|x\|_E < \varepsilon\},$$

contiene a alguna bola

$$W(\delta) = \{y \in F : \|y\|_F < \delta\}.$$

Sea $\varepsilon_0 > 0$ y sea $(\varepsilon_n)_{n=1}^{\infty}$ una sucesión arbitraria de números positivos, tal que $\sum_{n=1}^{\infty} \varepsilon_n < 2\varepsilon_0$. Por lo demostrado en la parte i), existe una sucesión decreciente de números positivos $(\delta_n), \delta_n \rightarrow 0$, tal que $\overline{A[V(\varepsilon_n)]} \supset W(\delta_n)$, para todo $n = 0, 1, 2, \dots$

Si $y \in W(\delta_0)$, entonces se puede hallar $x \in V(2\varepsilon_0)$, tal que $A(x) = y$. Como $\overline{A[V(\varepsilon_0)]} \supset W(\delta_0)$, existe $x_0 \in V(\varepsilon_0)$, tal que $\|y - A(x_0)\|_F < \delta_1$. De la misma forma, como $y - A(x_0) \in W(\delta_1)$, se puede hallar $x_1 \in V(\varepsilon_1)$, tal que $\|y - A(x_0) - A(x_1)\|_F < \delta_2$.

Continuando este proceso, se encuentra una sucesión $(x_n), x_n \in V(\varepsilon_n)$, tal que

$$\|y - A(\sum_{i=1}^n x_i)\|_F < \delta_{n+1}, n = 0, 1, 2, 3, \dots$$

Pongamos

$$z_k = \sum_{i=0}^k x_i,$$

entonces

$$\|z_{k+p} - z_k\|_E = \left\| \sum_{i=k+1}^{k+p} x_i \right\|_E \leq \sum_{i=k+1}^{k+p} \varepsilon_i, \quad p > 0.$$

Luego, la sucesión (z_k) es de Cauchy en E . Por la completitud de E concluimos que (z_k) converge a algún elemento $x \in E$. Además

$$\|x\|_E = \lim_{k \rightarrow \infty} \|z_k\|_E \leq \lim_{k \rightarrow \infty} \sum_{i=0}^k \varepsilon_i < 2\varepsilon_0.$$

Por ser A un operador continuo, tendremos que

$$\lim_{n \rightarrow \infty} \|y - A(\sum_{i=0}^n x_i)\|_F \leq \lim_{n \rightarrow \infty} \delta_{n+1} = 0,$$

es decir,

$$y = A\left(\lim_{n \rightarrow \infty} \sum_{i=0}^n x_i\right) = A(x).$$

Con esto hemos probado que la imagen por A de la bola $V(2\varepsilon_0)$ de E contiene a la bola $W(\delta_0)$ de F . Por la arbitrariedad de ε_0 se obtiene el resultado deseado.

iii) Finalmente notemos que si U es un subconjunto abierto no vacío de E y $x \in E$, entonces existe una bola $V(\varepsilon)$ de centro Θ en E , tal que $\{x\} + V(\varepsilon) \subset U$. Sea $W(\delta)$ una bola de centro Θ en F contenida en $A[V(\varepsilon)]$. La existencia de $W(\delta)$ está asegurada por la demostración de la parte ii). Entonces

$$A[U] \supset A[\{x\} + V(\varepsilon)] = \{A(x)\} + A[V(\varepsilon)] \supset \{A(x)\} + W(\delta).$$

Pero observemos que el conjunto $\{A(x)\} + W(\delta)$ es una vecindad de $A(x)$ en F y, por lo tanto, hemos demostrado que $A[U]$ es vecindad de cada uno de sus puntos $A(x)$; luego $A[U]$ es abierto. $\square$

El teorema siguiente es un corolario del teorema anterior.

T 2. (de Banach sobre el operador inverso) Sean E y F espacios de Banach y A un operador lineal y continuo que transforma E en F de manera biunívoca. Entonces el operador inverso $A^{-1} : F \rightarrow E$ es también lineal y continuo.

Demostración. La linealidad de A^{-1} es evidente. Para probar la continuidad de A^{-1} basta notar que si U es un abierto en F , la preimagen por A^{-1} de $U : [A^{-1}]^{-1}[U] = A[U]$, es abierto en E , como consecuencia del Teorema 1. $\square$

Ejemplo 1. (Bases de Riesz en un espacio de Hilbert) Si $A : H \rightarrow H$ es

un operador lineal continuo e inversible y (e_n) es una base ortonormal en H , entonces $(A(e_n))$ es también una base en H que se llama *base de Riesz*. En efecto, todo $y \in H$ se escribe de manera única $y = A(x)$, $x \in H$. Por ser (e_n) una base

$$x = \sum_{n=1}^{\infty} c_n e_n$$

y, por la continuidad de A ,

$$y = A(x) = \sum_{n=1}^{\infty} c_n A(e_n).$$

Si ponemos $A(e_n) = \varphi_n$, obviamente

$$\begin{aligned} \left\| \sum_{n=1}^{\infty} c_n \varphi_n \right\|^2 &= \|A(x)\|^2 \leq \|A\|^2 \|x\|^2 \\ &= \|A\|^2 \sum_{n=1}^{\infty} |c_n|^2 \end{aligned}$$

Pero, por T2, A^{-1} es también continuo y, por lo tanto

$$\begin{aligned} \sum_{n=1}^{\infty} |c_n|^2 &= \|x\|^2 = \|A^{-1}(A(x))\|^2 \\ &\leq \|A^{-1}\|^2 \|A(x)\|^2 \leq \|A^{-1}\|^2 \sum_{n=1}^{\infty} |c_n \varphi_n|^2. \end{aligned}$$

En fin, para toda base de Riesz (φ_n) , existen constantes positivas a_1 y a_2 , tales que

$$a_1 \sum_{n=1}^{\infty} |c_n|^2 \leq \left\| \sum_{n=1}^{\infty} c_n \varphi_n \right\|^2 \leq a_2 \sum_{n=1}^{\infty} |c_n|^2,$$

para toda sucesión (c_n) de números complejos tal que $\sum_{n=1}^{\infty} |c_n|^2 < \infty$. $\square$

Pasemos ahora a demostrar el segundo resultado importante de este epígrafe.

T 3. (del grafo cerrado) Sean E y F espacios de Banach y A un operador lineal de E en F . Entonces A es continuo si y sólo si, el grafo de A es cerrado en el espacio producto $E \times F$.

Demostración. Obviamente, el grafo de cualquier aplicación continua $f : E \rightarrow F$

$$G(f) = \{(x, f(x)) : x \in E\}$$

es un subconjunto cerrado en $E \times F$, ya que si la sucesión $(x_n, f(x_n))$ converge a un elemento (x, y) en $E \times F$, se tendrá que $x_n \rightarrow x$, $f(x_n) \rightarrow y$. Pero como además f es continua $f(x_n) \rightarrow f(x)$ y, por lo tanto, $(x, y) = (x, f(x)) \in G(f)$.

Supongamos inversamente que $A : E \rightarrow F$ es lineal y tiene grafo cerrado. Por ser A lineal, $G(A)$ es un subespacio vectorial cerrado en $E \times F$. Pero $E \times F$ es un espacio de Banach y, por lo tanto, $G(A)$ es también un espacio de Banach, provisto de la norma

$$\|(x, A(x))\| = \|x\|_E + \|A(x)\|_F.$$

Los operadores proyección:

$$\begin{aligned} \pi_E : G(A) &\longrightarrow E, & (x, A(x)) &\rightsquigarrow x, \\ \pi_F : G(A) &\longrightarrow F, & (x, A(x)) &\rightsquigarrow A(x), \end{aligned}$$

son lineales y continuas. Además π_E es biyectivo y por el Teorema de Banach sobre el operador inverso se concluye que existe π_E^{-1} el cual es también lineal y continuo. Pero $A = \pi_F \circ \pi_E^{-1}$, luego A es continuo. $\square$

La consecuencia siguiente del Teorema 3 fué ya anunciada en el Ejem-

plo 2.1.3.

C 1. Todo operador simétrico en un espacio de Hilbert es continuo.

Demostración. Según el Teorema 3, es suficiente probar que todo operador simétrico $A : H \rightarrow H$ (H espacio de Hilbert) tiene grafo cerrado. En efecto, supongamos que $x_n \rightarrow x$ y $A(x_n) \rightarrow y$ en H . Entonces, por la continuidad del producto escalar

$$\langle A(x_n), z \rangle \rightarrow \langle y, z \rangle,$$

para todo $z \in H$. Pero por ser A simétrico,

$$\langle A(x_n), z \rangle = \langle x_n, A(z) \rangle \rightarrow \langle x, A(z) \rangle = \langle A(x), z \rangle;$$

luego, por la unicidad del límite, tendremos que

$$\langle A(x), z \rangle = \langle y, z \rangle, \text{ para todo } z \in H.$$

De esto último, se deduce que $y = A(x)$ y, por lo tanto, el grafo de A es cerrado en $H \times H$. $\square$

3 Conceptos y resultados generales de la teoría de operadores en los espacios de Hilbert

En este capítulo nos dedicaremos a estudiar algunos conceptos y resultados generales de la teoría de operadores lineales en los espacios de Hilbert.

En los tres primeros epígrafes 3.1, 3.2 y 3.3, se introduce el concepto general de operador lineal; y se estudian las propiedades fundamentales de la clausura y el adjunto de un operador lineal. El § 3.4 está dedicado a estudiar un caso particular de mucha importancia para la Física: los operadores simétricos y los operadores autoadjuntos, también llamados autoconjugados.

En el § 3.5 se estudian las nociones de resolvente y espectro de un operador lineal, que resultan de gran utilidad en la Física-Matemática. Se hace especial énfasis en el estudio del espectro de los operadores simétricos y autoconjugados, a cuyo análisis se dedica el § 3.6. En el § 3.7 se introduce la integral de Dunford que juega un importante papel en el desarrollo de un cálculo operacional con operadores cerrados. En el epígrafe § 3.8 se ven dos importantes aplicaciones: la existencia de raíz cuadrada de operadores autoadjuntos positivos y la descomposición

polar de operadores cerrados. La relación entre los subespacios invariantes de operadores cerrados y la separación de diferentes partes de su espectro se estudia en el § 3.9, aplicándose en el § 3.10 al desarrollo de la resolvente en un entorno de los valores propios aislados. El Capítulo termina con el § 3.11 donde se introducen las primeras nociones y resultados para el estudio de los operadores J autoadjuntos en espacios con métrica indefinida.

A lo largo del capítulo hemos tratado de ejemplificar las nuevas nociones introducidas utilizando el operador de diferenciación $-i\frac{d}{dx}$ en un intervalo y el operador de multiplicación por una función en $L_2[a, b]$.

Utilizaremos las notaciones $G(A)$, $D(A)$, $\text{Ker } A$ y $R(A)$ para denotar el grafo, el dominio, el núcleo y el rango respectivamente del operador A .

§ 3.1 Algunos conceptos generales de la teoría de operadores lineales

En el capítulo anterior nos hemos dedicado al estudio de los operadores lineales continuos actuando sobre espacios normados. A partir de este momento comenzaremos a preparar las condiciones para estudiar los operadores diferenciales lineales, es decir operadores lineales definidos sobre ciertos espacios funcionales y que actúan mediante la aplicación de algún concepto de diferenciación ya sea clásico o generalizado.

Resulta que en los espacios funcionales donde se estudian los operadores diferenciales, ellos no son continuos.

Ejemplo 1. Consideremos el operador de diferenciación

$$d(f) = i\frac{df}{dx} \tag{1}$$

definido en el espacio normado $(C^\infty[0, 2\pi], \|\cdot\|_2)$ y con valores en $L_2[0, \pi]$,

donde $i = \sqrt{-1}$ y $\|\cdot\|_2$ es la norma en $L_2[0, 2\pi]$. Obviamente este operador lineal no es continuo puesto que transforma la sucesión acotada $(\sin nx)$ cuya norma es

$$\|\sin nx\|_2 = \frac{1}{\sqrt{2}}$$

en la sucesión no acotada $(n \cos nx)$.

$$\|n \cos nx\|_2 = \frac{n}{\sqrt{2}}. \quad \square$$

Muchas de las propiedades de los operadores lineales no acotados dependen, no solamente de la forma en que ellos actúan sobre los elementos de su dominio de definición, sino también de las propiedades de dicho dominio de definición en sí. En efecto, a la misma expresión diferencial (1) se le pueden asociar diferentes operadores variando el dominio de definición, conservando como espacio de llegada al propio espacio $L_2[0, 2\pi]$. Así podríamos definir operadores P_0, P_C, P_z y P con dominios de definición

$$\mathcal{D}(P_0) = C_0^\infty(0, 2\pi), \quad (2)$$

$$\mathcal{D}(P_C) = \{\varphi \in H^1[0, 2\pi] : \varphi(0) = \varphi(2\pi) = 0\}, \quad (3)$$

$$\mathcal{D}(P_z) = \{\varphi \in H^1[0, 2\pi] : \varphi(2\pi) = z\varphi(0)\} \quad (4)$$

$$\mathcal{D}(P) = H^1[0, 2\pi]. \quad (5)$$

todos ellos provistos de la norma $\|\cdot\|_2$. Aquí $H^1[0, 2\pi]$ es un espacio de Sobolev que ya ha sido definido en el Capítulo 1.

No es difícil comprobar (ver E1.7.2) que, en el caso unidimensional, $H^1[0, 2\pi]$ coincide con el conjunto de las funciones absolutamente continuas en $[0, 2\pi]$ y que pertenecen a $L_2[0, 2\pi]$ junto con su derivada, la cual existe casi dondequiera. Por el Teorema 1.7.13, tenemos que $H^1[0, 2\pi] \subset C[0, 2\pi]$, por lo cual tienen sentido las condiciones de contorno impuestas a $\varphi(0)$ y $\varphi(2\pi)$ en la expresión de los dominios de definición de D_C y D_z .

Conjugando T 1.7.15 con el teorema de inmersión 1.7.13 para el caso $\Omega = (0, 2\pi)$, se concluye además, que

$$\mathcal{D}(P_c) = H_0^1[0, 2\pi]. \quad (6)$$

Obviamente, para los operadores D_c , D_z y D , la operación de diferenciación corresponde a la derivación de una función absolutamente continua. Veremos en los próximos epígrafes algunas propiedades de los operadores (2)–(5). En particular veremos algunas ventajas que presentan los operadores D_z con respecto a los restantes operadores D_0 , D_c y D .

D 1. (operador lineal) Un *operador lineal* en un espacio de Hilbert H es un par $(A, D(A))$, donde $D(A)$ es un subespacio vectorial de H llamado *dominio de A* y $A : D(A) \rightarrow H$ es una aplicación lineal. Cuando no haya lugar a dudas acerca del dominio de definición denotaremos al operador simplemente por A .

En el conjunto de los operadores lineales sobre un espacio de Hilbert H se define una *relación de orden* mediante

$$A \subset B \iff \mathcal{D}(A) \subset \mathcal{D}(B) \text{ y } A(x) = B(x) \text{ para todo } x \in \mathcal{D}(A).$$

La demostración de que esta relación es de orden se deja de ejercicio al lector. Obviamente $A \subset B \iff G(A) \subset G(B)$. Si $A \subset B$, se dice que B es una *extensión* de A o que A es una *restricción* de B .

Obviamente para los operadores introducidos en (2)–(5) se tienen las inclusiones

$$D_0 \subset D_c \subset D_z \subset D. \quad (7)$$

La *suma* y la *composición* de operadores lineales, se define, al menos formalmente, de la misma forma que para operadores definidos sobre todo el espacio de partida:

$$(A + B)(x) = A(x) + B(x),$$

$$\begin{aligned}(\lambda A)(x) &= \lambda A(x), \\ (AB)(x) &= A(B(x)).\end{aligned}$$

La dificultad aquí aparece en relación con el hecho de que los dominios de definición de los operadores en general no coinciden con todo el espacio de partida.

Por este motivo se define

$$\begin{aligned}\mathcal{D}(A + B) &= \mathcal{D}(A) \cap \mathcal{D}(B), \\ \mathcal{D}(AB) &= \{x \in \mathcal{D}(B) : B(x) \in \mathcal{D}(A)\}.\end{aligned}$$

Puede ocurrir que $\mathcal{D}(A + B)$ o $\mathcal{D}(AB)$ se reduzca al espacio trivial $\{\Theta\}$. De manera evidente se generaliza el concepto de operador inverso para cualquier operador A con la propiedad de que tome valores diferentes en puntos diferentes de su dominio. Si esta condición se cumple, entonces, por definición

$$\begin{aligned}A^{-1}(g) &= f, \text{ si } A(f) = g, \\ \mathcal{D}(A^{-1}) &= R(A).\end{aligned}$$

Obviamente A^{-1} es lineal si A lo es y puede ocurrir que A^{-1} sea acotado sin serlo A .

§ 3.2 Operadores cerrados y cerrables

Sabemos que en un espacio de Banach los operadores lineales y continuos coinciden con los que tienen grafo cerrado (T 2.5.3). Para operadores lineales no continuos definidos en un espacio de Hilbert H y cuyo dominio de definición no coincide con todo H , parece natural estudiar, en primer lugar, aquellos que tienen grafo cerrado. Tales operadores se llaman

cerrados. Más exactamente el operador A en H se llama cerrado si del hecho que:

$$f \in \mathcal{D}(A), \quad \lim_{n \rightarrow \infty} f_n = f, \quad \lim_{n \rightarrow \infty} A(f_n) = g$$

se tiene

$$f \in \mathcal{D}(A), \quad g = A(f).$$

Ejemplo 1. El operador P_c definido en (3) es cerrado. En efecto, si $f_n \in \mathcal{D}(P_c)$ es tal que

$$f_n \rightarrow f \text{ y } i \frac{df_n}{dx} \rightarrow g, \text{ en } L_2[0, 2\pi] \quad (8)$$

entonces la sucesión (f_n) es de Cauchy en $H^1[0, 2\pi]$. Pero entonces, la sucesión (f_n) converge a una cierta función h en $H^1[0, 2\pi]$. Como la norma en $H^1[0, 2\pi]$ se define mediante

$$\|h\|_{2,1} = (\|h\|_2^2 + \|\frac{dh}{dx}\|_2^2)^{1/2},$$

entonces de (8) se tiene que

$$h = f, \quad g = i \frac{df}{dx}$$

donde la derivada a f se aplica en sentido de Soboliev. Para ver que $g = P_c(f)$ basta probar que $f(0) = f(2\pi) = 0$. Pero por el teorema de inmersión 1.7.13 tenemos la desigualdad

$$\|\varphi\|_\infty \leq K \|\varphi\|_{2,1} \quad (9)$$

donde la constante K es la misma para toda $\varphi \in H^1[0, 2\pi]$.

De (9) se concluye que $\|f_n - f\|_\infty \rightarrow 0$, y como $f_n(0) = f_n(2\pi) = 0$, entonces $f(0) = f(2\pi) = 0$. $\square$

P 1. El operador A es cerrado si y sólo si el grafo $G(A)$, provisto de la norma

$$\|(x, A(x))\|_G = (\|x\|^2 + \|A(x)\|^2)^{1/2} \quad (10)$$

es un espacio de Hilbert.

Demostración. Del hecho que $\|\cdot\|$ es la norma en el espacio de Hilbert donde está definido A se tiene que $\|\cdot\|_G$ proviene de un producto escalar.

La completitud del espacio $(G(A), \|\cdot\|_G)$ es obviamente equivalente al hecho que A es cerrado. $\square$

Para un operador lineal cerrado A en un espacio de Hilbert, a diferencia de los operadores lineales continuos, la existencia de $\lim_{n \rightarrow \infty} f_n$ ($f_n \in \mathcal{D}(A)$) no asegura la existencia de $\lim_{n \rightarrow \infty} A(f_n)$. De la definición de operador lineal cerrado, solamente podemos asegurar *que si dos sucesiones $(f_n^{(1)})$ y $(f_n^{(2)})$ en $\mathcal{D}(A)$ convergen al mismo elemento $f \in H$, y si existen $\lim_{n \rightarrow \infty} A(f_n^{(1)})$ y $\lim_{n \rightarrow \infty} A(f_n^{(2)})$, entonces ambos límites coinciden, ya que en este caso $f \in \mathcal{D}(A)$ y*

$$\lim_{n \rightarrow \infty} A(f_n^{(1)}) = \lim_{n \rightarrow \infty} A(f_n^{(2)}) = A(f).$$

Sin embargo, la propiedad subrayada no caracteriza a los operadores cerrados como nos muestra el ejemplo siguiente.

Ejemplo 2. El operador P_0 definido en (2) no es cerrado. En efecto, para toda

$$f \in H_0^1[0, 2\pi] \setminus C_0^\infty(0, 2\pi),$$

existe una sucesión (f_n) en $C_0^\infty(0, 2\pi)$ que converge a f en la norma de $H^1[0, 2\pi]$. (Véase (153) en el Capítulo 1). La sucesión $f_n \in \mathcal{D}(D_0)$, $\lim_{n \rightarrow \infty} f_n = f$ en $L_2[0, 2\pi]$ y $\lim_{n \rightarrow \infty} P_0(f_n)$ existe en $L_2[0, 2\pi]$ y coincide con $i \frac{df}{dx}$ (derivada en sentido de Soboliev). Pero hemos tomado f precisamente de manera tal que $f \notin \mathcal{D}(P_0)$, luego P_0 no es cerrado. Sin embargo, si $(f_n^{(1)})$ y $(f_n^{(2)})$ son sucesiones de $\mathcal{D}(P_0)$ que convergen en

$L_2[0, 2\pi]$ a la misma función f y tales que existen $\lim_{n \rightarrow \infty} D_0(f_n^{(1)})$ y $\lim_{n \rightarrow \infty} D_0(f_n^{(2)})$ en $L_2[0, 2\pi]$, entonces ambas sucesiones son de Cauchy en $H^1[0, 2\pi]$ y, por lo tanto

$$\lim_{n \rightarrow \infty} P_0(f_n^{(1)}) = \lim_{n \rightarrow \infty} P_0(f_n^{(2)}) = i \frac{df}{dx}. \quad \square$$

D 1. (operador lineal cerrable) Un operador lineal se llama *cerrable* si admite alguna extensión que sea cerrado.

Si el operador lineal A es cerrable entonces *existe una extensión lineal cerrada minimal* de A , la cual se denota por $\overline{A}$ y se llama *clausura* de A . En efecto, basta definir $\overline{A}$ mediante

$$\begin{aligned} \mathcal{D}(\overline{A}) &= \bigcap \{ \mathcal{D}(B) : B \text{ cerrado y } B \supset A \} \\ \overline{A}(f) &= B(f), \forall f \in \mathcal{D}(\overline{A}), \end{aligned} \quad (11)$$

donde B es cualquier extensión cerrada de A . El lector puede verificar en calidad de ejercicio que el operador lineal $\overline{A}$ es la extensión cerrada minimal del operador cerrable A .

P 2. Un operador lineal A es cerrable si y sólo si se cumple la condición siguiente: para toda sucesión (f_n) en $\mathcal{D}(A)$ tal que $\lim_{n \rightarrow \infty} f_n = \Theta$ y $\lim_{n \rightarrow \infty} A(f_n)$ existe, se tiene que $\lim_{n \rightarrow \infty} A(f_n) = \Theta$.

Cuando se cumple esta condición, $\mathcal{D}(\overline{A})$ se obtiene uniendo a $\mathcal{D}(A)$ todos los elementos $f \in \overline{\mathcal{D}(A)}$, para los cuales es posible hallar al menos una sucesión $(f_n) \in \mathcal{D}(A)$, tal que

$$\lim_{n \rightarrow \infty} f_n = f \text{ y } \lim_{n \rightarrow \infty} A(f_n) \text{ existe.}$$

Para tales f se pone

$$\overline{A}(f) = \lim_{n \rightarrow \infty} A(f_n).$$

Demostración. La condición enunciada es obviamente necesaria. Si ahora suponemos que se cumple dicha condición, entonces también se cumplirá la propiedad subrayada antes del enunciado del Ejemplo 1. De aquí se concluye que tiene sentido la definición del operador $\overline{A}$ dada en la formulación de la proposición. Se deja al lector probar que $\overline{A}$ así definido es una extensión cerrada de A y más aún, coincide con la clausura de A . $\square$

Para un operador arbitrario A , el conjunto $\overline{G(A)}$ no es necesariamente el grafo de algún operador lineal. Sin embargo se tiene la proposición siguiente:

P 3. Si el operador A es cerrable entonces $G(\overline{A}) = \overline{G(A)}$.

Demostración. En efecto, si A es cerrable su clausura viene dada por (11), de cuya expresión se concluye sin dificultad el resultado deseado. $\square$

Ejemplo 3. El operador A definido en $L_2[0, 1]$ por la fórmula

$$A(f)(x) = f(1)x$$

y con dominio

$$\mathcal{D}(A) = \mathcal{C}[0, 1],$$

no es cerrable. En efecto, se pueden construir dos sucesiones de funciones continuas (f_n) y (g_n) , que converjan en $L_2[0, 1]$ a un límite común, pero tales que $f_n(1) \rightarrow 1$ y $g_n(1) \rightarrow 0$. Entonces $\lim_{n \rightarrow \infty} A(f_n) = x$, mientras que $\lim_{n \rightarrow \infty} A(g_n) = 0$, lo cual muestra que ambos límites existen pero son diferentes. $\square$

Ejemplo 4. Tomando $f_n^{(1)} = f_n$ con $\lim_{n \rightarrow \infty} f_n = \Theta$ y $f_n^{(2)} \equiv \Theta$, en

el Ejemplo 2, podemos concluir utilizando P 2, que el operador P_0 es cerrable. Además, de la segunda parte de P2 se concluye que $\mathcal{D}(P_0)$ coincide con la clausura de $C(0, 2\pi)$ en $H^1[0, 2\pi]$; luego

$$\mathcal{D}(\overline{P}_0) = H_0^1[0, 2\pi] = \mathcal{D}(P_c). \quad (12)$$

De (12) se concluye que

$$\overline{P}_0 = P_c. \quad (13)$$

§ 3.3 Conjugado de un operador lineal

Para todo operador lineal continuo A definido sobre un espacio de Hilbert H , el funcional bilineal

$$Ly(x) = \langle A(x), y \rangle$$

con $y \in H$ fijo, es continuo. Por el teorema de representación de los funcionales lineales continuos sobre un espacio de Hilbert, 2.3.1), concluimos que existe un elemento en H , al que denotaremos mediante $A^*(y)$, tal que

$$\langle A(x), y \rangle = \langle x, A^*(y) \rangle, \quad \forall x, y \in H.$$

No es difícil verificar que la aplicación

$$A^* : H \longrightarrow H$$

define un operador sobre H que es también lineal y continuo y además

$$|||A||| = |||A^*|||.$$

El operador A^* se llama *conjugado* del operador lineal acotado A . Resulta que el concepto de operador conjugado puede extenderse también a muchos operadores no acotados.

D 1. (operador adjunto o conjugado) Sea A un operador lineal definido sobre un subespacio vectorial $\mathcal{D}(A)$ denso en el espacio de

Hilbert H . Sea $\mathcal{D}(A^*)$ el conjunto de los $y \in H$, para los cuales existen $z \in H$, que cumplen

$$\langle A(x), y \rangle = \langle x, z \rangle, \quad \forall x \in \mathcal{D}(A). \quad (14)$$

Para cada $y \in \mathcal{D}(A^*)$, pongamos

$$A^*(y) = z. \quad (15)$$

A^* se llama el *operador conjugado (adjunto)* del operador A .

Observación. La condición $\mathcal{D}(A)$ denso en H es fundamental, para que tenga sentido definir la aplicación (15). En efecto, de no cumplirse esa condición podrían existir más de un elemento z en H correspondiente al mismo $y \in \mathcal{D}(A^*)$, en cuyo caso no tendría sentido la aplicación (15). Es obvio que todo operador lineal A con dominio denso en H admite un operador conjugado cuyo dominio puede ser no denso. Por el Teorema de Riesz sobre representación de funcionales lineales continuos en un espacio de Hilbert, $\mathcal{D}(A^*)$ está compuesto por aquellos $y \in H$, tales que

$$|\langle Ax, y \rangle| \leq C \|x\|, \quad \forall x \in \mathcal{D}(A)$$

donde la constante C no depende de x .

Un operador A con dominio denso en H , se dice que esta *densamente definido* en H .

Ejemplo 1. Calculemos el conjugado del operador P_0 definido en (2). Según la Definición 1, $\mathcal{D}(P_0^*)$ consiste de las funciones $f \in L_2[0, 2\pi]$, para las cuales existe $g \in L_2[0, 2\pi]$, que cumple

$$\int_0^{2\pi} i \frac{d\varphi}{dx} \bar{f} dx = \int_0^{2\pi} \varphi \bar{g} dx, \quad \forall \varphi \in C_0^\infty(0, 2\pi).$$

Por la definición de derivada en sentido de Soboliev, tenemos que

$$P_0^*(f) = g = i \frac{df}{dx}$$

y, como $f, g \in L_2[0, 2\pi]$, entonces

$$f \in H^1[0, 2\pi].$$

En fin, hemos probado que

$$P_0^* = P, \tag{16}$$

donde P es el operador definido en (5). $\square$

El lector puede verificar en calidad de ejercicio que también

$$P_c^* = P. \tag{17}$$

Las proposiciones siguientes nos muestran algunas propiedades de la operación de conjugación.

P 1. Sean A y B operadores lineales con dominios densos en un espacio de Hilbert H . Entonces

- i) A^* es lineal,
- ii) $A \subset B \Rightarrow B^* \subset A^*$,
- iii) A^* es siempre cerrado, aunque A no lo sea,
- iv) A cerrable $\Rightarrow A^* = (\overline{A})^*$.

Demostración. Se deduce inmediatamente de la definición de conjugado y por ello se deja de ejercicio al lector. $\square$

Veamos como actúa la conjugación ante las operaciones elementales definidas con operadores.

P 2. Sean A y B como en la proposición y $\lambda \in \mathbb{C}$. Entonces

- i) $(\lambda A)^* = \bar{\lambda} A^*$ ($\lambda \neq 0$);
- ii) $(A + B)^* \supset A^* + B^*$;
- iii) $(AB)^* \supset B^* A^*$;
- iv) Si uno de los dos operadores es acotado entonces se cumple la inclusión inversa en ii) y iii).
- v) Si existe A^{-1} y además $\mathcal{D}(A^{-1}) = R(A)$ es denso en H , entonces existe $(A^*)^{-1}$ y

$$(A^*)^{-1} = (A^{-1})^*.$$

Demostración. i) – iii) Ejercicio.

- iv) Supongamos que A es acotado y apliquemos ii) a la suma $(A + B) + (-A)$, la cual coincide con B por el hecho que A está definido en todo H . Entonces obtendremos

$$B^* \supset (A + B)^* - A^*.$$

Si sumamos A^* a ambos lados de esta última inclusión, llegamos a la relación

$$A^* + B^* \supset (A + B)^* - A^* + A^* = (A + B)^*$$

para lo cual hemos utilizado el hecho que A^* también está definido en todo H . Ahora tomemos $f \in \mathcal{D}((AB)^*)$, entonces para cualquier $g \in \mathcal{D}(B)$ se tiene

$$\langle Bg, A^*f \rangle = \langle ABg, f \rangle = \langle g, (AB)^*f \rangle,$$

luego $A^*f \in \mathcal{D}(B^*)$ y $B^*A^*f = (AB)^*f$. Con esto se prueba la afirmación iv).

v) Sean $f \in \mathcal{D}(A)$ y $g \in \mathcal{D}((A^{-1})^*)$. En este caso

$$\langle f, g \rangle = \langle A^{-1}Af, g \rangle = \langle Af, (A^{-1})^*g \rangle.$$

Pero esta última igualdad nos muestra que

$$\begin{aligned} (A^{-1})^*g &\in \mathcal{D}(A^*) \\ A^*(A^{-1})^*g &= g. \end{aligned} \tag{18}$$

Por otra parte, si $f \in \mathcal{D}(A^{-1})$ y $h \in \mathcal{D}(A^*)$, entonces

$$\langle f, h \rangle = \langle AA^{-1}f, h \rangle = \langle A^{-1}f, A^*h \rangle,$$

de donde se deduce que

$$\begin{aligned} A^*h &\in \mathcal{D}((A^{-1})^*) \\ (A^{-1})^*A^*h &= h. \end{aligned} \tag{19}$$

Las relaciones (18) y (19) muestran que

$$(A^*)^{-1} = (A^{-1})^*. \quad \square$$

No es difícil probar (ejercicio), que si el operador $A^{**} = (A^*)^*$ existe entonces $A \subset A^{**}$. Este simple resultado nos muestra que una condición necesaria para la existencia de A^{**} es la cerrabilidad del operador A . Veamos que esta condición es también suficiente.

T 1. El operador lineal A , densamente definido en el espacio de Hilbert H , es cerrable si y sólo si $\mathcal{D}(A^*)$ es también denso en H . En este caso se tiene que

$$\overline{A} = A^{**}. \tag{20}$$

Demostración. Por el comentario previo al enunciado del teorema basta probar que la cerrabilidad de A es una condición suficiente para que $\mathcal{D}(A^*)$ sea denso en H . Consideremos el operador unitario V definido en $H \times H$ mediante

$$V(f, g) = (-g, f). \quad (21)$$

Obviamente, $V^2 = -I$, y en virtud de la unitariedad de V se tiene que, para cualquier subespacio E de $H \times H$,

$$V[E^\perp] = (V[E])^\perp. \quad (22)$$

Si A es un operador lineal en H y $(f, g) \in H \times H$, entonces $(f, g) \in (V[G(A)])^\perp$ si y sólo si

$$\ll (f, g), (-A\Psi, \Psi) \gg = 0, \quad \forall \Psi \in \mathcal{D}(A) \quad (23)$$

donde $\ll, \gg$ denota el producto escalar en el espacio producto $H \times H$. Pero (23) se cumple si y sólo si,

$$\langle f, A\Psi \rangle = \langle g, \Psi \rangle, \quad \forall \Psi \in \mathcal{D}(A),$$

es decir, solamente cuando $(f, g) \in G(A^*)$. En fin, hemos probado que

$$G(A^*) = (V[G(A)])^\perp. \quad (24)$$

Por ser $G(A)$ un subespacio vectorial del espacio de Hilbert $H \times H$, se tiene que

$$\begin{aligned} \overline{G(A)} &= (G(A))^{\perp\perp} (V^2[G(A)])^{\perp\perp} \\ &= (V[V[G(A)]])^{\perp\perp} = (V[G(A^*)])^\perp, \end{aligned} \quad (25)$$

donde la segunda igualdad se debe a (21) y la cuarta a (22) y (24).

Obviamente, si $\mathcal{D}(A^*)$ es denso en H , entonces existe A^{**} y aplicando (24) y (25) a A^* en lugar de A , llegamos a que

$$\overline{G(A)} = G(A^{**}). \quad (26)$$

Pero de (26) se desprende que A es cerrable y $\overline{A} = A^{**}$.

Recíprocamente, supongamos que $\mathcal{D}(A^*)$ no es denso en H . Entonces existe $0 \neq \Psi \in \mathcal{D}(A^*)^\perp$ y, por lo tanto $(\Psi, 0) \in (G(A^*))^\perp$. Pero de aquí se tiene que $(0, \Psi) \in V[G(A^*)]^\perp$, y de (25) concluimos que $\overline{G(A)}$ no es el grafo de ningún operador, o sea A no es cerrable. $\square$

P 3. Para todo operador A densamente definido en un espacio de Hilbert se cumple

$$R(A)^\perp = \text{Ker} A^*. \quad (27)$$

Demostración. La inclusión $g \in \text{Ker} A^*$ es equivalente a que

$$\langle f, A^* g \rangle = 0$$

para todo $f \in \mathcal{D}(A)$. Pero $g \in R(A)$ significa, que $\langle Af, g \rangle = 0$ para todo $f \in \mathcal{D}(A)$. Entonces (27) se desprende de la definición de operador conjugado. $\square$

§ 3.4 Operadores autoadjuntos, simétricos y esencialmente autoadjuntos

Para la física resultan de gran importancia los llamados operadores autoadjuntos. Ciertas magnitudes físicas como, por ejemplo, la energía, pueden ser representadas mediante operadores. A esos operadores se les asocian determinados conjuntos de escalares (espectro), que corresponden a aquellos valores de la magnitud física representada por el operador, para los cuales el sistema físico es observable en el experimento.

Las magnitudes físicas deben estar expresadas por valores que puedan ser medibles en el experimento y, por lo tanto, deben estar dadas en números reales. Resulta que precisamente son los operadores autoadjuntos los que tienen su espectro constituido por números reales.

D 1. (operador autoadjunto, extensión autoadjunta) Un operador lineal A densamente definido en un espacio de Hilbert H se llama *autoadjunto* si

$$A = A^*. \quad (28)$$

Se dice que el operador B , densamente definido en H , es una *extensión autoadjunta* de A , si B es autoadjunto y $A \subset B$.

Ejemplo 1. Calculemos todas las posibles extensiones autoadjuntas del operador P_0 definido en (2). Sea $\tilde{P}$ una tal extensión, entonces

$$P_0 \subset \tilde{P} = \tilde{P}^* \subset P$$

donde la segunda inclusión sale del hecho que

$$P_0 \subset \tilde{P} \Rightarrow \tilde{P} = \tilde{P}^* \subset \tilde{P}_0^* = P.$$

Además, $\tilde{P}$ es cerrado por ser autoadjunto y, por lo tanto

$$P_c = \tilde{P}_0 \subset \tilde{P}^*.$$

En fin, llegamos a la cadena de inclusiones

$$P_0 \subset P_c \subset \tilde{P} \subset P.$$

Por los teoremas de inmersión de espacios de Soboliev, tenemos que para cualquier par de funciones $\varphi, \Psi \in \mathcal{D}(P)$ se puede aplicar la fórmula

de integración por partes:

$$\begin{aligned}
 \langle \tilde{P}\varphi, \Psi \rangle &= \int_0^{2\pi} i\varphi'(t)\overline{\Psi(t)}dt \\
 &= i[\varphi(2\pi)\overline{\Psi(2\pi)} - \varphi(0)\overline{\Psi(0)}] + \int_0^{2\pi} \varphi(t)i\overline{\Psi'(0)}dt \\
 &= i[\varphi(2\pi)\overline{\Psi(2\pi)} - \varphi(0)\overline{\Psi(0)}] + \langle \varphi, \tilde{P}\Psi \rangle.
 \end{aligned}$$

Por ser $\tilde{P}$ autoadjunto, debe cumplirse la relación

$$\varphi(2\pi)\overline{\Psi(2\pi)} - \varphi(0)\overline{\Psi(0)} = 0. \quad (29)$$

El lector puede verificar sin dificultad que P_c no es autoadjunto (ejercicio) y, por lo tanto, existe $\Psi_0(t) \in \mathcal{D}(\tilde{P})$ tal que $\Psi_0(2\pi) \neq 0$. Poniendo $\Psi(t) = \Psi_0(t)$ en (29), obtenemos la relación

$$\varphi(2\pi) = z\varphi(0) \quad (30)$$

para toda $\varphi \in \mathcal{D}(\tilde{P})$, donde la constante

$$z = \frac{\overline{\Psi_0(0)}}{\overline{\Psi_0(2\pi)}}. \quad (31)$$

Como la condición (30) debe cumplirse también para $\Psi(t) = \Psi_0(t)$, concluimos que

$$|z| = 1. \quad (32)$$

Hemos probado que

$$\mathcal{D}(\tilde{P}) \subset \{\varphi \in H^1[0, 2\pi] : \varphi(2\pi) = z\varphi(0)\}; \quad |z| = 1.$$

Demostremos que es válida la inclusión en sentido contrario. Sean $\Psi \in H^1[0, 2\pi]$, $\Psi(2\pi) = z\Psi(0)$ y α una constante tal que

$$\Psi(0) - \alpha\Psi_0(0) = 0.$$

Pongamos $\varphi(t) = \Psi(t) - \alpha\Psi_0(t)$. Entonces es fácil ver que $\varphi(t) \in H^1[0, 2\pi]$ y que $\varphi(0) = \varphi(2\pi) = 0$, luego $\varphi \in \mathcal{D}(P_c) \subset \mathcal{D}(\tilde{P})$. De aquí se sigue que $\Psi(t) = \varphi(t) + \alpha\Psi_0(t) \in \mathcal{D}(\tilde{P})$. En fin, se tiene que

$$\mathcal{D}(\tilde{P}) = \{\varphi \in H^1[0, 2\pi] : \varphi(2\pi) = z\varphi(0)\}. \quad (33)$$

No es difícil verificar que, para cada z con $|z| = 1$, el operador $i\frac{d}{dx}$ con un dominio del tipo (33) es autoadjunto. Luego, *toda extensión autoadjunta de P_0 es del tipo P_z* (4). $\square$

Observación. No todo operador tiene necesariamente extensiones autoadjuntas. Por ejemplo, el operador P definido en (5) no las tiene (ejercicio). El problema de la existencia y caracterización de las extensiones autoconjugadas de cierta clase de operadores, llamados simétricos, será estudiado en el Capítulo 4.

D 2. (operadores simétricos y esencialmente autoadjuntos, dominio esencial) Sea A un operador lineal densamente definido en el espacio de Hilbert H . Entonces

i) El operador A se llama *simétrico* si

$$\langle Ax, y \rangle = \langle x, Ay \rangle, \quad \forall x, y \in \mathcal{D}(A) \quad (34)$$

ii) El operador A se llama *esencialmente autoadjunto*, si $\overline{A}$ es autoadjunto.

iii) Si A es cerrado, se llama *dominio esencial* del operador A , a cualquier subconjunto $\mathcal{D} \subset \mathcal{D}(A)$, tal que la clausura de la restricción de A a $\mathcal{D}$ coincide con A

$$\overline{A|_{\mathcal{D}}} = A. \quad (35)$$

Observación. No deben confundirse las definiciones de operador simétrico y autoadjunto. Ambas definiciones coinciden para los operadores lineales acotados, no siendo así para los no acotados. En efecto, no resulta difícil comprobar que el operador P_0 definido en (2) es simétrico pero no autoadjunto, ya que como mostramos en el Ejemplo 3.3.1, $P_0^* = P$. De ambas definiciones se concluye inmediatamente que *el operador A es simétrico si y sólo si $A \subset A^*$.*

Si el Operador A es esencialmente autoadjunto, entonces posee una y sólo una extensión autoadjunta.

En efecto, supongamos que B es una extensión autoadjunta de A . Entonces B es cerrado y, por consiguiente, de $B \supset A$ se tiene que $B \supset A^{**} = \overline{A}$. Por eso $B = B^* \subset (A^{**})^* = A^{**}$, de donde concluimos que

$$B = A^{**} = \overline{A}.$$

Puede probarse, aunque no lo veremos aquí, que se cumple también la afirmación inversa, o sea *si A posee una única extensión autoadjunta, entonces él es esencialmente autoadjunto.*

Para A esencialmente autoadjunto, se tiene

$$A^* = \overline{A^*} = A^{***} = A^{**} = \overline{A} \subset A^*.$$

Luego, podemos concluir que *A es esencialmente autoadjunta si y sólo si*

$$\overline{A} = A^*. \quad (36)$$

La autoadjunticidad esencial representa una gran ayuda cuando se trabaja con un operador simétrico no cerrado A . Si se logra demostrar que A es esencialmente autoadjunto, entonces a él le corresponde de manera única el operador autoadjunto $\overline{A} = A^{**}$. Pero un operador autoadjunto puede ser definido totalmente describiendo algún dominio

esencial, lo cual en general es mucho más simple que describir todo su dominio de definición. Es precisamente esto lo que se hace en la Física y en parte ello justifica la importancia de demostrar la autoadjunticidad esencial de los operadores que aparecen en muchos problemas de la Física.

Supongamos ahora que A es un operador autoadjunto y que existe $x \in \mathcal{D}(A^*) = \mathcal{D}(A)$, tal que $A^*(x) = ix$ ($i = \sqrt{-1}$). Entonces $A(x) = ix$, de donde

$$\begin{aligned} i\langle x, x \rangle &= \langle ix, x \rangle = \langle A(x), x \rangle = \langle x, A^*(x) \rangle \\ &= \langle x, A(x) \rangle = -i\langle x, x \rangle \end{aligned}$$

y, por lo tanto, $x = \Theta$. De manera análoga se prueba que la ecuación $A^*(x) = -ix$ no tiene soluciones diferentes de la trivial. La afirmación inversa: *si A es un operador cerrado simétrico tal que las ecuaciones $A^*(x) = \pm ix$ no tienen soluciones diferentes de la trivial, entonces A es autoadjunto*, es el criterio fundamental de autoadjunticidad.

T 1. (criterio fundamental de autoadjunticidad) Sea A un operador simétrico en el espacio de Hilbert H . Entonces las afirmaciones siguientes son equivalentes:

- i) A es autoadjunto,
- ii) A es cerrado y $\text{Ker}(A^* \pm ix) = \{\Theta\}$,
- iii) $R(A \pm iI) = H$

Demostración. Hemos acabado de ver que i) $\Rightarrow$ ii). Supongamos ahora que ocurre ii) y demostremos iii). En efecto de P 3.3.3 se tiene que

$$\overline{R(A \pm iI)} = [\text{Ker}(A^* \mp iI)]^\perp = H$$

luego $R(A \pm iI)$ es denso en H . Queda por probar que $R(A \pm iI)$ es cerrado en H . Pero para todo $x \in \mathcal{D}(A)$,

$$\|(A - iI)x\|^2 = \|Ax\|^2 + \|x\|^2.$$

Luego, si $(x_n) \in \mathcal{D}(A)$ y $(A - iI)(x_n) \rightarrow y_0$, entonces (x_n) converge a un cierto vector $x_0 \in H$ y $A(x_n)$ también converge. Por ser A cerrado, $x_0 \in \mathcal{D}(A)$ y $(A - iI)(x_0) = y_0$, con lo cual queda probado que $R(A - iI)$ es cerrado y con ello iii). Análogamente se tiene $R(A + iI) = H$. Mostremos finalmente que

iii) $\Rightarrow$ i). Sea $y \in \mathcal{D}(A^*)$. Por ser $R(A - iI) = H$, existe $x \in \mathcal{D}(A)$, tal que $(A - iI)(x) = (A^* - iI)(y)$.

Pero $\mathcal{D}(A) \subset \mathcal{D}(A^*)$, luego $y - x \in \mathcal{D}(A^*)$, y por tanto se sigue que $(A^* - iI)(y - x) = 0$. Por otra parte,

$$\{\Theta\} = H^\perp = R(A + iI)^\perp = \text{Ker}(A^* - iI),$$

de donde se concluye que $y = x \in \mathcal{D}(A)$. Esto muestra que $\mathcal{D}(A^*) = \mathcal{D}(A)$, o sea A es autoadjunto. $\square$

C 1. Sea A un operador simétrico en el espacio de Hilbert H . Entonces son equivalentes:

- i) A es esencialmente autoadjunto,
- ii) $\text{Ker}(A^* \pm iI) = \{\Theta\}$
- iii) $R(A^* \pm iI)$ son densos en H .

Demostración. Ejercicio. $\square$

Una condición suficiente para la autoadjunticidad nos la brinda la

proposición siguiente:

P 1. Un operador simétrico A definido en un espacio de Hilbert H , y tal que $R(A) = H$, es autoadjunto.

Demostración. Obviamente $A \subset A^*$. Probemos la inclusión inversa. Tomemos $y \in \mathcal{D}(A^*)$ y $A^*(y) = g$. Como $R(A) = H$, entonces existe $x \in \mathcal{D}(A)$, tal que $A(x) = g$. Entonces para cualquier $f \in \mathcal{D}(A)$,

$$\langle A(f), y \rangle = \langle f, g \rangle = \langle f, A(x) \rangle = \langle Af, x \rangle.$$

Por ser $R(A) = H$, se sigue que $y = x$, es decir, $y \in \mathcal{D}(A)$. $\square$

E 1. Demuestre que una variedad lineal $\mathcal{D}$, con $\mathcal{D}(A) \subset \mathcal{D} \subset \mathcal{D}(A^*)$ es el dominio de alguna extensión autoadjunta para el operador simétrico A si y solo si el conjunto de todos los elementos $f \in \mathcal{D}(A^*)$ tales que

$$\langle A^*f, g \rangle = \langle f, A^*g \rangle \quad \forall g \in \mathcal{D}, \quad \text{coincide con } \mathcal{D}.$$

Consideremos un operador simétrico cerrado A en un espacio de Hilbert H . Sea $x \in \mathcal{D}(A)$ y pongamos

$$(A + iI)(x) = f \tag{37}$$

$$(A - iI)(x) = g. \tag{38}$$

Repitiendo la demostración de la implicación ii) $\Rightarrow$ iii) del Teorema 1, se ve que $R(A + iI)$ y $R(A - iI)$ son subespacios vectoriales cerrados de H y, por lo tanto, son subespacios de Hilbert. Además, los operadores

$$\begin{aligned} A + iI : \mathcal{D}(A) &\longrightarrow R(A + iI) \\ A - iI : \mathcal{D}(A) &\longrightarrow R(A - iI) \end{aligned} \tag{39}$$

establecen correspondencias biyectivas ya que

$$\|(A \pm iI)(x)\|^2 = \|A(x)\|^2 + \|x\|^2$$

y si $(A \pm iI)(x) = \Theta$, entonces $x = \Theta$.

Por este motivo, para cada $f \in R(A + iI)$, se puede hallar un único elemento $x \in \mathcal{D}(A)$ tal que se cumple (37). Una vez hallado este elemento x , definamos $g \in H$ mediante la fórmula (38). De esta forma hemos definido un operador

$$V : R(A + iI) \longrightarrow R(A - iI) \quad (40)$$

con $V(f) = g$. Obviamente

$$V(f) = (A - iI)(A + iI)^{-1}(f). \quad (41)$$

El operador V definido de esta forma mediante (41) se llama *transformada de Cayley* del operador simétrico cerrado A .

Es fácil ver que V es un operador isométrico. En efecto, V es lineal y además

$$\|f\|^2 = \|A(x)\|^2 + \|x\|^2 = \|g\|^2.$$

Utilizando V se puede expresar el operador A por la fórmula

$$A(x) = i(I + V)(I - V)^{-1}(x). \quad (42)$$

En fin, a todo operador isométrico A se le puede asociar un operador isométrico V , tal que se cumple (42). Recíprocamente se cumple el teorema siguiente:

T 2. Sean F y G subespacios vectoriales cerrados del espacio de Hilbert H , y $V : F \rightarrow G$ un operador isométrico. Si $R(V - I)$ es denso en H , entonces el operador definido por la fórmula (42) es simétrico y V es su transformada de Cayley.

Demostración. Por ser $R(V - I)$ denso en H , existe el operador inverso $(V - I)^{-1}$. En efecto, si $V - I$ no fuera inversible, existiría $f \in F$,

tal que $V(f) = f$. Pero entonces, para cualquier $g \in F$, se tendría

$$\begin{aligned}\langle V(g) - g, f \rangle &= \langle V(g), f \rangle - \langle g, f \rangle \\ &= \langle g, V(f) \rangle - \langle g, f \rangle = 0\end{aligned}$$

es decir $g \perp R(V - I)$.

Por existir el operador $(V - I)^{-1}$, tiene sentido definir

$$A = i(I + V)(I - V)^{-1}$$

cuyo dominio de definición es denso en H . Demostremos que este operador es simétrico.

Sean $f, g \in \mathcal{D}(A) = R(V - I)$ y escribamos $-f = V(x) - x$, y $-g = V(y) - y$, con $x, y \in \mathcal{D}(V) = F$. Entonces,

$$\begin{aligned}A(f) &= i(I + V)(x) = ix + iV(x), \\ A(g) &= i(I + V)(y) = iy + iV(y).\end{aligned}$$

Por lo tanto

$$\langle A(f), g \rangle = i[\langle V(x), y \rangle - \langle x, V(y) \rangle] = \langle f, A(g) \rangle$$

lo cual prueba que A es simétrico.

La demostración de la relación (41) no presenta ninguna dificultad, luego el operador V es la transformada de Cayley de A . $\square$

T 3. Un operador simétrico y cerrado, definido en un espacio de Hilbert, es autoadjunto si y sólo si su transformada de Cayley es un operador unitario.

Demostración. Sean A un operador cerrado y simétrico definido en el espacio de Hilbert H , y V su transformada de Cayley. Por ser V isométrico, la unitariedad de V es equivalente a que

$$R(A + iI) = R(A - iI). \quad (43)$$

Si suponemos que A es autoadjunto, entonces (43) es una consecuencia de T 1 iii). Inversamente, supongamos que V es isométrico. Sabemos que $R(A \pm iI)$ son subespacios vectoriales cerrados en H . Pero de (27) y (43) se tiene que

$$Ker(A^* + i) = R(A - iI)^\perp = R(A + iI)^\perp = Ker(A^* - iI) \quad (44)$$

y como

$$Ker(A^* + iI) \cap Ker(A^* - iI) = \{\Theta\}$$

entonces de (44) obtenemos $R(A \pm iI)^\perp = \{\Theta\}$ y, por lo tanto $R(A \pm iI) = H$, de donde se concluye (T1) que A es autoadjunto. $\square$

Para todo operador simétrico A (en particular autoadjunto), se tiene que

$$\{\langle A(x), x \rangle : x \in \mathcal{D}(A)\} \subset \mathbb{R}. \quad (45)$$

Debido a (45), podemos comparar dos operadores simétricos poniendo:

$$A \leq B \iff \mathcal{D}(A) \subset \mathcal{D}(B), \langle A(x), x \rangle \leq \langle B(x), x \rangle, \quad \forall x \in \mathcal{D}(A). \quad (46)$$

Diremos que $A \geq B$ si $-A \leq -B$.

No es difícil probar que $\leq$ es una relación de orden en el conjunto de todos los operadores simétricos definidos en un espacio de Hilbert H . Para demostrar esto es suficiente ver que

$$\langle A(x), x \rangle = 0 \quad \forall x \in \mathcal{D}(A) \Rightarrow A \equiv \Theta.$$

Pero esto último se cumple puesto que

$$0 = \langle A(x + iy), x + iy \rangle = i\langle A(y), x \rangle - i\langle A(x), y \rangle$$

para todos $x, y \in \mathcal{D}(A)$. De aquí se tiene

$$\text{Im } \langle A(x), y \rangle = 0, \quad \forall x, y \in \mathcal{D}(A).$$

Sustituyendo en la última expresión iy en lugar de y , llegamos a que

$$\langle A(x), y \rangle = 0 \quad \forall x, y \in \mathcal{D}(A)$$

de donde $A(x) \equiv \Theta, \forall x \in \mathcal{D}(A)$.

D 3. Se dice que el operador simétrico A está acotado *inferiormente* si existe una constante $m \in \mathbb{R}$ tal que $A \geq mI$; es decir

$$m\langle x, x \rangle \leq \langle Ax, x \rangle, \quad \forall x \in \mathcal{D}(A). \quad (47)$$

La constante m es una *cota inferior* para el operador A . El operador A está *acotado superiormente* si $-A$ está acotado inferiormente.

El operador A se llama *positivo (estrictamente positivo)* si $A \geq 0$ (si $A > 0$) y se llama *negativo (estrictamente negativo)* si $-A \geq 0$ (si $-A > 0$).

Observación 1. Si A está acotado superiormente, entonces existe una constante M (*cota superior*), tal que

$$\langle A(x), x \rangle \leq M\langle x, x \rangle, \quad \forall x \in \mathcal{D}(A). \quad (48)$$

Observación 2. Para todo operador A simétrico y positivo, la aplicación $h(x, y) = \langle A(x), y \rangle$ define un producto escalar en $\mathcal{D}(A)$ que aunque no cumpla que $h(x, x) = 0 \Rightarrow x = \Theta$, sí satisface la desigualdad de Cauchy–Buniakovsky y, por lo tanto, se tiene que

$$|\langle A(x), y \rangle| \leq \sqrt{\langle A(x), x \rangle} \sqrt{\langle A(y), y \rangle}, \quad \forall x, y \in \mathcal{D}(A). \quad (49)$$

T 4. Si A es un operador lineal cerrado y densamente definido en el espacio de Hilbert H , entonces la composición A^*A , es un operador

autoadjunto positivo.

Demostración. Ante todo notemos, que para cualesquiera $x, y \in \mathcal{D}(A^*A)$, se tiene

$$\langle A^*A(x), y \rangle = \langle A(x), A(y) \rangle = \langle x, A^*A(y) \rangle$$

y, en particular

$$\langle A^*A(x), x \rangle = \langle A(x), A(x) \rangle \geq 0$$

lo cual prueba la positividad de A^*A .

Para ver que A^*A es simétrico, probemos que $\mathcal{D}(A^*A)$ es denso en H . Por ser A cerrado, de (25) tendremos

$$H \times H = G(A) \oplus G(A)^\perp = G(A) \oplus V[G(A^*)],$$

donde V es el operador unitario definido en (21). Luego, para todo $h \in H$, el elemento (h, Θ) de $H \times H$, se representa de manera única

$$\begin{aligned} (h, \Theta) &= (x, A(x)) + V(-y, A^*(-y)) \\ &= (x, A(x)) + (A^*(y), -y). \end{aligned}$$

Por lo tanto

$$\begin{aligned} h &= x + A^*(y) \\ \Theta &= A(x) - y, \end{aligned}$$

de donde

$$h = (I + A^*A)(x).$$

Luego hemos probado que, para cualquier $h \in H$, la ecuación

$$(I + A^*A)(f) = h, \tag{50}$$

tiene solución única y, por lo tanto

$$R(I + A^*A) = H. \tag{51}$$

Veamos que de este hecho se concluye que $\mathcal{D}(A^*A)$ es denso en H . En efecto, supongamos que existe un vector $h \neq \Theta$ en H , tal que $h \in \mathcal{D}(A^*A)^\perp$. Pero h se puede representar en la forma (50). Luego, para todo $g \in \mathcal{D}(A^*A)$

$$0 = \langle h, g \rangle = \langle (I + A^*A)(f), g \rangle = \langle f, (I + A^*A)(g) \rangle$$

y, tomando $g = f$, obtendremos

$$0 = \langle f, f \rangle + \langle f, A^*A(f) \rangle = \langle f, f \rangle + \langle A(f), A(f) \rangle$$

de donde concluimos, que $f = \Theta$, es decir $h = \Theta$, lo cual es una contradicción. Entonces, por P1, tendremos que $I + A^*A$ es autoadjunto, de donde se deduce la autoadjunticidad de A^*A ya que $A^*A = (I + A^*A) - I$. $\square$

T 5. Toda sucesión monótona (con respecto al orden $\leq$), acotada superior e inferiormente, de operadores simétricos y acotados, definidos en un espacio de Hilbert, converge puntualmente a algún operador simétrico y acotado.

Demostración. Sin pérdida de generalidad se puede suponer que

$$0 \leq A_1 \leq A_2 \leq \dots \leq I.$$

Sean $m < n$, entonces $A_{mn} = A_n - A_m \geq 0$. Por (49) tendremos que, para todo $x \in H$,

$$\begin{aligned} \|A_{mn}(x)\|^4 &= \langle A_{mn}(x), A_{mn}(x) \rangle^2 \\ &\leq \langle A_{mn}(x), x \rangle \langle A_{mn}^2(x), A_{mn}(x) \rangle. \end{aligned}$$

Del hecho que, $0 \leq A_{mn} \leq I$, se tiene que $\|A_{mn}\| \leq 1$ y

$$\|A_n(x) - A_m(x)\|^4 \leq [\langle A_n(x), x \rangle - \langle A_m(x), x \rangle] \|x\|^2.$$

La sucesión numérica $\langle A_n(x), x \rangle$ es creciente y acotada superiormente, luego converge. Pero entonces la sucesión $(A_n(x))$ es de Cauchy y, por lo tanto, convergente en H . Obviamente el operador A , definido mediante la igualdad $A(x) = \lim_{n \rightarrow \infty} A_n(x)$ es lineal y simétrico. La continuidad de A se desprende entonces del Corolario 2.5.1. $\square$

Es fácil ver que el cuadrado de todo operador simétrico es también un operador simétrico y además positivo, ya que

$$\langle A^2(x), x \rangle = \langle A(x), A(x) \rangle \geq 0.$$

Más adelante veremos un cierto recíproco de este resultado: *todo operador autoadjunto y positivo admite una raíz cuadrada*. Por ahora demostraremos este enunciado para operadores acotados.

T 6. (raíz cuadrada de operadores simétricos acotados) A cada operador simétrico, acotado y positivo A , definido en un espacio de Hilbert, le corresponde un único operador simétrico, acotado y positivo, llamado raíz cuadrada del operador A , que se denota mediante $A^{1/2}$ y que cumple

$$(A^{1/2})^2 = A. \quad (52)$$

El operador $A^{1/2}$ es el límite puntual de una sucesión de polinomios en A con coeficientes reales y por ello $A^{1/2}$ conmuta con todo operador que conmute con A .

Demostración. Sin pérdida de generalidad supondremos que $0 \leq A \leq I$. Pongamos $A = I - B$ ($0 \leq B \leq I$) y $X = A^{1/2} = I - Y$. Entonces la solución de la ecuación $X^2 = (A^{1/2})^2 = A$ se obtiene resolviendo la ecuación equivalente

$$Y = 1/2(B + Y^2). \quad (53)$$

Resolvamos esta última ecuación por el método de las aproxima-

ciones sucesivas:

$$Y_0 = 0, Y_1 = 1/2B, \dots, Y_{n+1} = 1/2(B + Y_n^2), \quad n \geq 0.$$

Mostremos que la sucesión (Y_n) converge y su límite será la solución de la ecuación (53). Veamos ante todo, que Y_n y $Y_n - Y_{n-1}$ son polinomios con respecto a B con coeficientes reales no negativos. En efecto, esta afirmación es cierta para $n = 1$. Supongamos que también es cierta para $n = m$ y demostremos su validez para $n = m + 1$. Para el operador Y_{m+1} la afirmación es obvia, pues se deduce de su expresión en función de Y_m .

Para $Y_{m+1} - Y_m$ se tiene

$$\begin{aligned} Y_{m+1} - Y_m &= 1/2(B + Y_m^2) - 1/2(B + Y_{m-1}^2) \\ &= 1/2(Y_m^2 - Y_{m-1}^2) = 1/2(Y_m + Y_{m-1})(Y_m - Y_{m-1}). \end{aligned}$$

Aquí hemos utilizado la conmutatividad de Y_m y Y_{m-1} , ya que ambos son polinomios de B . En la parte derecha de la igualdad tenemos el producto de dos polinomios de B con coeficientes reales no negativos. Por la hipótesis de inducción tenemos el resultado para todo n .

Por otra parte, si $B \geq 0$, entonces $B^n \geq 0, n = 2, 3, \dots$, ya que para $n = 2k, \langle B^{2k}(x), x \rangle = \|B^k(x)\|^2$, y para $n = 2k + 1, \langle B^{2k+1}(x), x \rangle = \langle BB^k(x), B^k(x) \rangle$.

Luego $Y_n \geq 0$ y $Y_n - Y_{n-1} \geq 0$. Probemos finalmente que $\|Y_n\| \leq 1$ para todo n . Para $n = 0$ esto es cierto, y para los restantes valores de n , el resultado se demuestra por inducción con ayuda de la igualdad

$$Y_{n+1} = 1/2(B + Y_n^2), \quad n \geq 0.$$

Si aplicamos el Teorema 5 a la sucesión (Y_n) , tendremos que (Y_n) converge puntualmente a un operador simétrico, acotado y positivo, que satisface la ecuación

$$Y = 1/2(B + Y^2).$$

Además

$$X = I - Y, \quad X^2 = A.$$

Demostremos la unicidad de la raíz cuadrada. Supongamos que X_1 es otro operador simétrico, acotado y positivo, tal que $X_1^2 = A$. Como $X_1 A = X_1 (X_1)^2 = (X_1)^2 X_1 = A X_1$, entonces X_1 conmuta con A y, por lo tanto, también conmuta con X por ser el límite puntual de una sucesión de polinomios del operador A . Sean Z y Z_1 , raíces cuadradas de X y X_1 , construidas de la misma forma que fué construida la raíz cuadrada X de A . Pongamos $g = (X - X_1)(f)$, para cualquier $f \in H$. Entonces

$$\begin{aligned} \|Z(g)\|^2 + \|Z_1(g)\|^2 &= \langle Z^2(g), g \rangle + \langle Z_1^2(g), g \rangle \\ &= \langle X(g), g \rangle + \langle X_1(g), g \rangle \\ &= \langle (X + X_1)(X - X_1)(f), g \rangle \\ &= \langle (X^2 - X_1^2)(f), g \rangle \\ &= \langle (A - A)(f), g \rangle = 0 \end{aligned}$$

de donde tenemos, $Z(g) = Z_1(g) = 0$ y, por consiguiente, $X(g) = ZZ(g) = 0, X_1(g) = Z_1 Z_1(g) = 0$. De aquí, concluimos que

$$\|(X - X_1)(f)\|^2 = \langle (X - X_1)^2(f), f \rangle = \langle (X - X_1)(g), f \rangle = 0$$

es decir, $(X - X_1)(f) = 0, \forall f \in H$, luego $X = X_1$. $\square$

§ 3.5 Resolvente y espectro

En la introducción al § 3.4. hemos ya hablado de la importancia para la Física del estudio del espectro de operadores lineales. Si A es un operador lineal en $\mathbb{C}^n$, entonces se llama *espectro de A* al conjunto de sus valores propios, es decir, aquellos números complejos λ , para los cuales la ecuación $Ax - \lambda x = 0$ tiene solución no trivial $x \in \mathbb{C}^n$, es decir el operador $A - \lambda I$ no es inyectivo.

Pero se sabe que para una matriz $A = (a_{ij})_{i,j=1}^n$, la *invertibilidad* de $A - \lambda I$ es equivalente a su *inyectividad*, lo cual a su vez es equivalente a que el determinante $\det(A - \lambda I)$ sea diferente de cero. Notemos que esto, que hemos expresado en unas pocas palabras, es gran parte de los contenidos de cualquier curso de Algebra Lineal. Es precisamente, este resultado el que nos permite afirmar que para una matriz A se cumple:

$$\sigma(A) = \{\lambda \in \mathbb{C} : \det(A - \lambda I) = 0\}.$$

Pero $\det(A - \lambda I) = 0$ es una ecuación polinomial de orden n , y de ahí la “sencillez” de la estructura del espectro de los operadores en dimensión finita.

Recordemos que a la noción de espectro de una matriz están vinculados resultados fundamentales del Algebra Lineal:

1° *El teorema sobre la posibilidad de diagonalizar cualquier matriz simétrica*, es decir cualquier matriz cuyos coeficientes satisfagan la relación $a_{ij} = \overline{a_{ji}}$, $j = 1, \dots, n$,

Este teorema nos dice que para tales matrices puede encontrarse una base de $\mathbb{C}^n$ en la que la matriz A se puede expresar en la forma diagonal:

$$A_0 = \begin{pmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{pmatrix},$$

donde $(\lambda_i)_1^n$ son las raíces de la ecuación $\det(A - \lambda I) = 0$. Dicho de otra forma se puede encontrar una matriz invertible B , también llamada *matriz de cambio de base* tal que se tiene la relación:

$$A = BA_0B^{-1}.$$

2° *El teorema de descomposición de Jordán* sobre la posibilidad de encontrar una base en la cual la matriz A (no necesariamente simétrica) adquiera una expresión triangular donde debajo de la diagonal sólo

aparecen ceros, mientras que encima pueden aparecer ceros y unos, y en la diagonal aparecen de nuevo las raíces de la ecuación $\det (A - \lambda I) = 0$.

Recomendamos al lector revisar la teoría de operadores en espacios de dimensión finita en el libro de I. M. Gelfand; Lecciones de álgebra lineal [8].

Precisamente la generalización de estos resultados a operadores en espacios de dimensión infinita va a ser objeto de estudio de los capítulos siguientes. Podremos cerciorarnos de la incidencia tan grande que tienen estas generalizaciones en los problemas de dinámica y estabilidad de sistemas mecánicos.

La teoría espectral de operadores en espacios normados de dimensión infinita es mucho más compleja e interesante, y a su vez de gran utilidad para el estudio de las propiedades fundamentales de los propios operadores.

En este epígrafe desarrollaremos la teoría espectral para operadores en espacios de Hilbert, aunque muchos de los resultados que obtendremos serán válidos en espacios de Banach cualesquiera.

D. 1. (resolvente y espectro) Sea A un operador lineal con dominio denso en el espacio de Hilbert H . Diremos que el número complejo λ pertenece al *conjunto resolvente* $\rho(A)$ del operador A , si $A - \lambda I$ es una biyección de $\mathcal{D}(A)$ sobre $H = R(A)$, con inverso $(A - \lambda I)^{-1} = R_\lambda(A)$ acotado. El operador $R_\lambda(A)$ se llama *resolvente de A* en el punto $\lambda \in \rho(A)$. Llamaremos *espectro de A* al complemento de $\rho(A)$ en $\mathbb{C}$

$$\sigma(A) = \mathbb{C} \setminus \rho(A). \quad (54)$$

Observación. Si el operador A es cerrado, entonces

$$\rho(A) = \{ \lambda \in \mathbb{C} \mid A - \lambda I : \mathcal{D}(A) \rightarrow H \text{ es una biyección} \}$$

ya que, si $A - \lambda I$ es una biyección, $(A - \lambda I)^{-1}$ es un operador cerrado definido en todo H y por el teorema del grafo cerrado es acotado.

Veremos más adelante que desde el punto de vista conjuntista la estructura del espectro de un operador no acotado puede ser muy compleja. Desgraciadamente esto no es todo el problema, pues para las aplicaciones de la teoría espectral a distintas ramas de la Física y de la propia Matemática se hace necesario distinguir, clasificar y hasta caracterizar ciertos subconjuntos del espectro cuya estructura puede ser aún más compleja.

Comenzaremos dando la siguiente clasificación para el espectro de un operador cerrado arbitrario.

D 2. (clasificación del espectro) El número complejo λ pertenece al *espectro puntual* de A ($\lambda \in \sigma_p(A)$) si $A - \lambda I$ no es inyectivo. Los elementos de $\sigma_p(A)$ se llaman valores propios de A . Si $\lambda \in \sigma_p(A)$, al subespacio $N_\lambda(A) = \overline{\text{Ker}(A - \lambda I)}$ se le llama *subespacio propio* correspondiente al valor propio λ y la dimensión algebraica del subespacio propio es la *multiplicidad* de λ . Los elementos de $N_\lambda(A)$ se llaman *vectores propios* correspondientes al valor propio λ .

Si el operador $A - \lambda I$ tiene un inverso no acotado con dominio denso en H , diremos que λ pertenece al *espectro continuo* de A ($\lambda \in \sigma_c(A)$).

Llamaremos *espectro residual* de A al conjunto de los números complejos λ para los cuales el operador $A - \lambda I$ tiene un inverso cuyo dominio no es denso en H y lo denotaremos mediante $\sigma_r(A)$.

Al conjunto de los puntos aislados en el espectro se conviene en llamar *espectro directo* y se denota mediante $\sigma_d(A)$.

Observación. Nótese que las diferentes partes del espectro introducidas en la definición anterior constituyen una partición del espectro de A :

$$\sigma(A) = \sigma_p(A) \cup \sigma_c(A) \cup \sigma_r(A)$$

donde la intersección entre dos cualesquiera de las componentes es vacía.
□

Una noción importante para nuestro objetivo y vinculada a la de vector propio es la de *vector radical*. Si λ es un valor propio de A y e_0 un vector propio correspondiente a λ , entonces toda solución e_1 de la ecuación

$$Ae_1 - \lambda e_1 = e_0, \quad e_1 \in \mathcal{D}(A), \quad \text{etc.}$$

es un vector radical asociado al vector propio e_0 . A partir de cada vector radical e_1 se pueden construir otros vectores radicales resolviendo la ecuación

$$Ae_2 - \lambda e_2 = e_1, \quad e_2 \in \mathcal{D}(A).$$

De esta forma se construyen *cadena de vectores radicales*, partiendo de los vectores propios. Estas cadenas pueden ser infinitas. El espacio vectorial generado por todos los vectores propios y radicales correspondientes a un valor propio λ se llama *subespacio radical* y se denota mediante $L_\lambda(A)$. A la dimensión algebraica de $L_\lambda(A)$ se le llama *rango del valor propio* λ .

E 2. Demuestre que: $L_\lambda(A) = \bigcup_{k=1}^{\infty} \text{Ker}(A - \lambda I)^k$.

Recuerde que el punto clave en la demostración del teorema de descomposición de Jordán para endomorfismos A en espacios de dimensión finita, es precisamente la demostración de que la sucesión de los núcleos $\text{Ker}(A - \lambda_i I)^k, k = 1, \dots, \infty$, se mantiene estacionaria a partir de un cierto k_i para cada λ_i solución de la ecuación $\det(A - \lambda I) = 0$. Si se denota mediante

$$L_i = \text{Ker}(A - \lambda_i I)^{k_i}$$

resulta que se puede probar que

$$H = \bigoplus_{i=1}^m L_i \quad (\text{se supone que } \dim H = n),$$

donde $m \leq n$ es el número de raíces diferentes de la ecuación $\det(A - \lambda I) = 0$.

Escogiendo en cada L_i una base de vectores propios y radicales conveniente se puede dar una expresión matricial del endomorfismo A que adquiere la forma señalada en el punto 2°. Después de esta digresión “finito dimensional” pasaremos a ver algunos ejemplos.

Ejemplo 1. Consideremos el operador P_z definido en (4) para $z = 1$ y calculemos su espectro. La igualdad

$$P_1 \varphi = \lambda \varphi$$

significa que

$$i\varphi'(t) = \lambda \varphi(t), \text{ con } \varphi(2\pi) = \varphi(0) \text{ y } \varphi \in H^1[0, 2\pi].$$

De aquí se tiene que

$$\lambda = \lambda_k = k, \quad \varphi(t) = \varphi_k(t) = e^{-ikt} \quad (\pm k = 0, 1, 2, \dots).$$

No resulta difícil verificar que para $\lambda \neq \lambda_k$, la ecuación $(P_1 - \lambda I)f = g$, $g \in L_2[0, 2\pi]$, o equivalentemente,

$$if'(t) - \lambda f(t) = g(t), \text{ con } g \in L_2[0, 2\pi], f(2\pi) = f(0), f \in H^2[0, 2\pi]$$

tiene solución única para cualquier $g \in L_2[0, 2\pi]$. Luego, todo $\lambda \neq \lambda_k$ pertenece al conjunto $\rho(P_1)$ y, por lo tanto,

$$\sigma(P_1) = \sigma_d(P_1) = \{\lambda_k : \pm k = 0, 1, \dots\}. \quad \square$$

Ejemplo 2. Definamos ahora un operador análogo a P_1 en los intervalos $[0, \infty[$ y $]-\infty, \infty[$ y analicemos su espectro. Notemos primeramente que

para toda función $\varphi \in H^1[0, \infty]$ se tiene que el producto $\varphi(t)\overline{\varphi'(t)}$ es integrable en el intervalo $[0, \infty[$. Integrando por partes, tendremos

$$\int_0^t \varphi(s)\overline{\varphi'(s)}ds = |\varphi(t)|^2 - |\varphi(0)|^2 - \int_0^t \varphi'(s)\overline{\varphi(s)}ds$$

lo cual nos muestra que $|\varphi(t)|$ tiene un límite cuando $t \rightarrow \infty$. Pero como $\varphi(t) \in L_2[0, \infty]$, entonces

$$\lim_{t \rightarrow \infty} \varphi(t) = 0.$$

Así vemos, que en infinito, se satisface automáticamente una condición de contorno, como condición de contorno en el otro extremo del intervalo $[0, \infty]$, tomaremos de manera natural

$$\varphi(0) = 0. \quad (55)$$

De esta forma hemos definido un operador P_c , análogo al definido en (3), mediante

$$P_c(\varphi) = i\varphi'(t), \quad (56)$$

$$\mathcal{D}(P_c) = \{\varphi \in H^1[0, \infty]; \varphi(0) = 0\}. \quad (57)$$

No es difícil verificar (ejercicio) que P_c es un operador cerrado y más aún, P_c es la clausura del operador P_0 definido por (56), con dominio de definición

$$\mathcal{D}(P_0) = C_0^\infty(0, \infty). \quad (58)$$

El operador P_c es además simétrico, ya que la relación

$$\begin{aligned} \langle P_c\varphi, \Psi \rangle &= i \int_0^\infty \varphi'(t)\overline{\Psi(t)}dt = \int_0^\infty \varphi(t)\overline{i\Psi(t)} \\ &= \langle \varphi, P_c\Psi \rangle \end{aligned}$$

se cumple para todo par $\varphi, \Psi \in \mathcal{D}(P_c)$.

Al igual que en el caso del intervalo finito $[0, 2\pi]$, no es difícil demostrar que $\mathcal{D}(P_c^*) = H^1[0, \infty]$ y que $P_c^*\Psi = i\Psi'$ (ejercicio). Luego P_c no es autoadjunto, ya que $P_c \neq P_c^*$. Pero resulta, que en el caso del intervalo $[0, \infty]$, el operador P_c no tiene extensiones autoadjuntas. En efecto, si existiera una tal extensión autoadjunta $\tilde{P}$ de P_c , entonces $\mathcal{D}(\tilde{P})$ debería contener alguna función $\varphi_0(t)$ diferente de cero en $t = 0$. En este caso tendríamos

$$\begin{aligned} \langle \tilde{P}\varphi_0, \varphi_0 \rangle &= i \int_0^\infty \varphi_0'(t) \overline{\varphi_0(t)} dt \\ &= i|\varphi_0(0)|^2 + \int_0^\infty \varphi_t(t) \overline{i\varphi_0'(t)} dt \neq \langle \varphi_0, \tilde{P}\varphi_0 \rangle \end{aligned}$$

lo cual no es posible.

Calculemos ahora el espectro del operador P_c . Notemos primeramente que P_c no tiene valores propios, ya que para cualquier $\lambda \in \mathbb{C}$, las únicas soluciones de la ecuación

$$i\varphi'(t) - \lambda\varphi(t) = 0$$

son de la forma

$$\varphi_\lambda(t) = \text{const } e^{-i\lambda t}$$

que, aunque pertenecen a $L_2[0, \infty]$ para $\text{Im } \lambda < 0$, no satisfacen la condición de contorno (55), salvo en el caso $\text{const} = 0$.

Por otra parte, para cualquier función $g(t)$, localmente integrable en $[0, \infty]$, la ecuación

$$i\varphi'(t) - \lambda\varphi(t) = g(t) \tag{59}$$

tiene la única solución

$$\varphi(x) = -ie^{-i\lambda x} \int_0^x g(t) e^{i\lambda t} dt, \tag{60}$$

que satisface la condición de contorno (55) y es además absolutamente continua.

Notemos que de (59) se deduce, que para cada $g \in L_2[0, \infty]$, basta demostrar la pertenencia de la solución φ dada por (60) a $L_2[0, \infty]$ para concluir que $\varphi \in H^1[0, \infty]$.

De lo anterior concluimos que $(P_c - \lambda I)$ es un operador inyectivo para cualquier $\lambda \in \mathbb{C}$ y, por lo tanto tiene un inverso

$$(P_c - \lambda I)^{-1} : R(P_c - \lambda I) \longrightarrow \mathcal{D}(P_c) = \mathcal{D}(P_c - \lambda I).$$

Es obvio que $R(P_c - \lambda I) \neq L_2[0, \infty]$ para cualquier $\lambda \in \mathbb{C}$, luego

$$\sigma(P_c) = \mathbb{C}.$$

Además

$$[R(P_c - \lambda I)]^\perp = \text{Ker}(P_c^* - \bar{\lambda}I) = \begin{cases} \{\Theta\} & \text{si } \text{Im}\lambda \leq 0 \\ \neq \{\Theta\} & \text{si } \text{Im}\lambda > 0. \end{cases}$$

Luego $\sigma(P_c) = \sigma_c(P_c) \cup \sigma_r(P_c)$, donde

$$\begin{aligned} \sigma_c(P_c) &= \{\lambda \in \mathbb{C} : \text{Im}\lambda \leq 0\} \\ \sigma_r(P_c) &= \{\lambda \in \mathbb{C} : \text{Im}\lambda > 0\}. \end{aligned}$$

Analícemos ahora un operador análogo a P_c en la recta real $(-\infty, \infty)$. Notemos primeramente que para toda función $\varphi \in H^1(\mathbb{R})$, se cumple automáticamente

$$\lim_{t \rightarrow \pm\infty} \varphi(t) = 0.$$

Por este motivo podemos definir el operador P_c en $\mathbb{R}$, mediante

$$P_c(\varphi) = i\varphi'(t) \tag{61}$$

$$\mathcal{D}(P_c) = H^1[\mathbb{R}]. \tag{62}$$

Obviamente, el operador P_c así definido es autoadjunto. Este operador no tiene valores propios, ya que la ecuación

$$i\varphi'(t) = \lambda\varphi(t)$$

no posee soluciones no triviales en $L_2(\mathbf{R})$. El lector puede verificar sin dificultad que

$$\sigma(P_c) = \sigma_c(P_c) = \mathbf{R}. \quad \square \quad (63)$$

Ejemplo 3. A continuación introduciremos el llamado operador de multiplicación por la variable independiente, al cual denotaremos mediante Q . Si $[a, b]$ es un intervalo finito, entonces el operador Q se define sobre todo $\varphi \in L_2[a, b]$, mediante la igualdad

$$Q(\varphi) = t\varphi(t). \quad (64)$$

En este caso Q es un operador lineal acotado y autoadjunto en $L_2[a, b]$ cuya norma es igual a $\max\{|a|, |b|\}$. En el caso de un intervalo (a, b) infinito consideraremos el operador Q actuando según la fórmula (64), y con dominio $\mathcal{D}(Q)$ constituido por las funciones $\varphi \in L_2[a, b]$, tales que $t\varphi(t) \in L_2[a, b]$.

Obviamente $\mathcal{D}(Q)$ es denso en $L_2[a, b]$ y Q es simétrico. El lector puede verificar en calidad de ejercicio que el operador Q es autoadjunto. El operador Q no posee valores propios, ya que la ecuación

$$Q(\varphi) = \lambda\varphi$$

significa que

$$\int_a^b |t - \lambda|^2 |\varphi(t)|^2 dt = 0$$

de donde se concluye que $\varphi(t) \equiv 0$ en $L_2[a, b]$.

Todos los puntos del intervalo $[a, b]$ pertenecen al espectro continuo del operador Q , ya que $R(Q - \lambda I)$, $a < \lambda < b$, está constituido por las

funciones $\varphi(t) \in L_2[a, b]$, tales que

$$\frac{\varphi(t)}{t - \lambda} \in L_2[a, b].$$

Obviamente $R(Q - \lambda I)$ es denso en $L_2[a, b]$, pero no coincide con $L_2[a, b]$ puesto que no contiene a las funciones que son constantes en un entorno de $t = \lambda$. $\square$

Ejemplo 4. Observemos que en lugar del operador Q en el espacio $L_2[a, b]$, podríamos haber introducido el operador Q_γ de multiplicación por la variable independiente en el espacio $L_{2,\gamma}[a, b]$ de las funciones de cuadrado integrable en el sentido de Lebesgue–Stieltjes con respecto a la función γ de variación acotada en $[a, b]$ (para la definición de la integral de Lebesgue–Stieltjes véase el libro de A.N. Kolmogorof y S.V. Fomin: *Fundamentos de la teoría de funciones y el análisis funcional* [2]).

E.3 Demostrar los resultados siguientes:

- 1.- Q_γ es autoadjunto,
- 2.- Los valores reales del conjunto resolvente $\rho(Q_\gamma)$ coinciden con los puntos de $[a, b]$, en un entorno de los cuales la función γ es constante
- 3.- Los valores propios del operador Q_γ son precisamente los puntos de discontinuidad de la función γ .
- 4.- El espectro continuo del operador Q_γ coincide con el conjunto de puntos no aislados en los cuales la función γ es continua pero no localmente constante. $\square$

En el Capítulo 6 veremos que los operadores del tipo Q_γ juegan un papel fundamental en la teoría de operadores autoadjuntos.

Decíamos anteriormente que la estructura del espectro de un operador arbitrario ya no resulta tan simple al encontrarnos en espacios de dimensión infinita. Para cerciorarnos de ello, proponemos al lector el siguiente ejercicio:

E 4. Sea $f(x)$ una función compleja medible definida en un intervalo $[a, b]$ con la propiedad de que el conjunto

$$\mathcal{D}_f = \{\varphi \in L_2[a, b] : f\varphi \in L_2[a, b]\}.$$

es denso en $L_2[a, b]$. Demostrar que el operador T_f de multiplicación por la función f en $L_2[a, b] : T_f\varphi = f\varphi$, con dominio $\mathcal{D}(T_f) = \mathcal{D}_f$, tiene como espectro

$$\sigma(T_f) = \overline{\{f(x) : x \in [a, b]\}}$$

donde la barra horizontal denota clausura en $\mathbb{C}$ del conjunto. $\square$

Pasemos ahora a estudiar algunas propiedades de la resolvente y el espectro de operadores cerrados.

P 1. Para que el operador lineal continuo T conmute con A es necesario que

$$R_\lambda T = T R_\lambda \tag{65}$$

para todo $\lambda \in \rho(A)$, y es suficiente que (65) ocurra para algún $\mu \in \rho(A)$.

Demostración. Supongamos que los operadores A y T conmutan, es decir que

$$TA(x) = AT(x)$$

para cualquier $x \in \mathcal{D}(A)$.

Si $\lambda \in \rho(A)$, $x = R_\lambda(y)$, entonces x recorre todo $\mathcal{D}(A)$ cuando y recorre H . Pero de la conmutatividad de A y T se tiene

$$(A - \lambda I)T(x) = T(A - \lambda I)(x), \quad \forall x \in \mathcal{D}(A)$$

es decir,

$$R_\lambda(A - \lambda I)T(x) = R_\lambda T(A - \lambda I)(x),$$

de donde

$$TR_\lambda(y) = R_\lambda T(y).$$

La demostración de la segunda parte del teorema es sencilla y se deja de ejercicio al lector. $\square$

Existe una relación muy simple entre la resolvente de un operador cerrado densamente definido y la de su adjunto. Notemos, que si T es un operador cerrado inversible con $R(T) = H$, entonces, por el teorema del grafo cerrado, T^{-1} es acotado. Luego, por P 3.3.2 v), en este caso existe $(T^*)^{-1}$ el cual es un operador acotado y

$$(T^*)^{-1} = (T^{-1})^*. \quad (66)$$

Recíprocamente, si $(T^*)^{-1}$ existe y es acotado, entonces también existe el operador acotado T^{-1} y se cumple (66). $\square$

De este razonamiento se concluye inmediatamente el resultado siguiente:

P 2. (resolvente del adjunto) Sea A un operador cerrado en H . Entonces los conjuntos $\rho(A^*)$ y $\sigma(A^*)$ se obtienen de $\rho(A)$ y $\sigma(A)$ mediante reflexión con respecto al eje real, y

$$R_\lambda(A^*) = R_{\bar{\lambda}}(A)^*, \quad \bar{\lambda} \in \rho(A) \quad (67)$$

Demostración. Ejercicio. $\square$

T 1. (analiticidad de la resolvente) Sea A un operador cerrado en H . Entonces $\rho(A)$ es un conjunto abierto del plano complejo y $R_\lambda(A)$

es analítica en $\lambda \in \rho(A)$.

Demostración. Sea $\lambda \in \rho(A)$. Entonces para todo $\mu \in \mathbb{C}$ con

$$|\mu - \lambda| < 1/\|R_\lambda\| \quad (68)$$

se tiene que

$$\begin{aligned} R_\mu(A) &= (A - \mu I)^{-1} = [(A - \lambda I) - (\mu - \lambda)I]^{-1} \\ &= R_\lambda(I - (\mu - \lambda)R_\lambda)^{-1} \\ &= \sum_{n=0}^{\infty} (\mu - \lambda)^n R_\lambda^{n+1}. \end{aligned} \quad (69)$$

La existencia de R_μ , para todo μ que satisface (68) nos muestra que $\rho(A)$ es un conjunto abierto. La expresión (69) para $R_\mu(A)$ muestra la analiticidad de R_μ en $\rho(A)$. $\square$

Observación. Del teorema anterior se concluye que el espectro de cualquier operador cerrado es un subconjunto cerrado de $\mathbb{C}$. Sin embargo $\sigma(A)$ puede ser vacío (trate de encontrar un ejemplo en que $\sigma(A) = \emptyset$).

T 2. (espectro de operadores acotados) El espectro de cualquier operador acotado A es un subconjunto compacto no vacío de $\mathbb{C}$, contenido en el círculo de radio

$$\text{spr } A = \overline{\lim_{n \rightarrow \infty}} \|A^n\|^{1/n} \quad (70)$$

donde $\text{spr } A = \sup_{\lambda \in \sigma(A)} |\lambda|$ es el llamado *radio espectral de A*. Además

$R_\lambda(A)$ es holomorfa en infinito y $R_\infty(A) = 0$.

Demostración. Por el criterio de Cauchy para convergencia de series en espacios de Banach, se tiene que si

$$|\lambda| > \lim_{n \rightarrow \infty} \|A^n\|^{1/n}$$

entonces existe $R_\lambda(A)$ y tiene un desarrollo en un entorno de $\lambda = \infty$, de la forma

$$R_\lambda(A) = -\frac{1}{\lambda} \left(I - \frac{A}{\lambda}\right)^{-1} = -\sum_{n=0}^{\infty} \frac{A^n}{\lambda^{n+1}}$$

luego $R_\infty(A) = 0$. Pero entonces no puede ocurrir, que $\sigma(A) = \phi$, ya que en este caso $R_\lambda(A)$ sería una función entera que se anula en infinito y por el Teorema de Liouville, tendríamos que

$$R_\lambda(A) \equiv 0, \quad \forall \lambda \in \mathbb{C}$$

lo cual es imposible. La demostración de la desigualdad

$$\sup_{\lambda \in \sigma(A)} \geq \lim_{n \rightarrow \infty} \|A^n\|^{1/n}$$

se deja de ejercicio al lector. $\square$

Hemos visto que todo operador A determina una partición del plano complejo en los conjuntos $\rho(A)$ y $\sigma(A)$. Para algunos propósitos es útil considerar el papel del punto infinito en esta partición.

T 3. Sea A un operador cerrado en H y supongamos que $\rho(A)$ contiene al exterior de un círculo. Entonces hay dos posibilidades:

- i) A es acotado; $R_\lambda(A)$ es holomorfa en infinito y $R_\infty(A) = 0$.
- ii) $R_\lambda(A)$ tiene una singularidad esencial en $\lambda = \infty$.

Demostración. El Teorema 2 nos muestra que en el caso en que A es acotado se cumple i). Supongamos que $\lambda = \infty$ no es una singularidad esencial de A y veremos que estamos en el caso i). En efecto, como $R_\lambda(A)$ no es idénticamente nulo, tendremos el desarrollo

$$R_\lambda(A) = \lambda^k A_0 + \lambda^{k-1} A_1 + \dots,$$

donde $A_0, A_1, \dots$ son operadores acotados y $A_0 \neq 0, |\lambda|$ es suficientemente grande y k es un entero. Entonces

$$AR_\lambda(A) = (A - \lambda I + \lambda I)R_\lambda(A) = I + \lambda R_\lambda(A) = I + \lambda^{k+1}A_0 + \lambda^k A_1 + \dots$$

Probemos primeramente que $k \leq -1$. Si fuera $k \geq 0$, entonces tendríamos que

$$\lambda^{-k-1}R_\lambda(A) \rightarrow 0, \quad A\lambda^{-k-1}R_\lambda(A) \rightarrow A_0$$

cuando $|\lambda| \rightarrow \infty$, lo cual implica que $A_0 = 0$ por ser A cerrado. Esto contradice la suposición $A_0 \neq 0$. Pero entonces $k \leq -1$ y por lo tanto

$$R_\lambda(A) \rightarrow 0, \quad AR_\lambda(A) \rightarrow I + \left(\lim_{\lambda \rightarrow \infty} \lambda^{k+1}\right)A_0$$

cuando $\lambda \rightarrow \infty$. De aquí concluimos que $I + (\lim_{\lambda \rightarrow \infty} \lambda^{k+1})A_0 = 0$, por ser A cerrado, lo cual es posible solamente si $k = -1$ y $A_0 = -I$. Luego $\lambda R_\lambda(A)(x) \rightarrow -x$ ($\lambda \rightarrow \infty$), y $A[\lambda R_\lambda(A)(x)] \rightarrow A_1(x)$ para todo $x \in H$. De nuevo, el ser A cerrado implica que $x \in \mathcal{D}(A)$ y $A(x) = -A_1(x)$. En fin, hemos probado que $A = -A_1$ y, por lo tanto, A es un operador acotado. $\square$

Observación. Como consecuencia del teorema anterior es natural incluir el punto $\lambda = \infty$ en el conjunto resolvente $\rho(A)$ si A es acotado y en el espectro $\sigma(A)$ si no lo es. De esta manera llegamos a una *noción ampliada de resolvente y espectro*, a la cual nos referiremos implícitamente en todo lo que sigue. De paso hemos demostrado que el espectro de un operador cerrado es un conjunto compacto de $\mathbb{C}$ si y sólo si el operador es acotado.

T 4. Sea A un operador cerrado e inversible en H . Entonces $\sigma(A)$ y $\sigma(A^{-1})$ se aplican uno en el otro mediante la aplicación $\lambda \rightarrow \lambda^{-1}$ definida sobre el plano complejo ampliado.

Demostración. Sea $\lambda \in \rho(A), \lambda \neq 0$. Entonces existe $R_\lambda(A)$. Consideremos el operador acotado $B(\lambda) = AR_\lambda(A) = I + \lambda R_\lambda(A)$. Para todo

$x \in H$, tenemos $B(\lambda)(x) = AR_\lambda(A)(x)$ y $A^{-1}B(\lambda)(x) = R_\lambda(A)(x) = \lambda^{-1}(B(\lambda) - I)(x)$. De aquí se tiene

$$-\lambda(A^{-1} - \lambda^{-1}I)B(\lambda)(x) = x. \quad (71)$$

Esto prueba que la imagen de $A^{-1} - \lambda^{-1}I$ coincide con H . Pero este operador también es inversible, ya que $(A^{-1} - \lambda^{-1}I)(y) = 0$ implica $y = \lambda A^{-1}y$, o sea $A(y) = \lambda y$, de donde $y = 0$.

De (71) se concluye que $(A^{-1} - \lambda^{-1}I)^{-1} = -\lambda B(\lambda)$ es un operador acotado y que $\lambda^{-1} \in \rho(A^{-1})$. Si $\lambda = 0 \in \rho(A)$, entonces A^{-1} es acotado y $0^{-1} = \infty \in \sigma(A^{-1})$ por definición. Si $\lambda = \infty \in \rho(A)$, entonces A es acotado y, por lo tanto $0 = \infty^{-1} \in \rho(A^{-1})$. Luego $\rho(A)$ se aplica mediante $\lambda \rightarrow \lambda^{-1}$ en $\rho(A^{-1})$. Lo mismo puede ser probado de manera análoga para los conjuntos complementarios $\sigma(A)$ y $\sigma(A^{-1})$. $\square$

Para operadores en espacios de Hilbert la noción de rango numérico es importante en diversas aplicaciones.

D 3. (rango numérico de un operador) Sea A un operador lineal en H . El rango numérico de A es el conjunto de todos los números complejos

$$r(A) = \{\langle A(x), x \rangle : x \in \mathcal{D}(A), \|x\| = 1\}. \quad (72)$$

En general, $r(A)$ no es ni abierto ni cerrado, aunque A sea un operador cerrado. El teorema siguiente es importante pero omitiremos su demostración.

T 5. (de Hausdorff) Para cualquier operador lineal A en un espacio de Hilbert, el conjunto $r(A)$ es convexo. $\square$

Del teorema anterior se concluye que la clausura de $r(A)$ es un subconjunto cerrado convexo de $\mathbb{C}$ y, por lo tanto, su complemento $\mathbb{C} \setminus \overline{r(A)}$ es un abierto conexo, salvo en el caso especial en que $\overline{r(A)}$ sea una banda

limitada por dos rectas paralelas, incluyendo el caso límite cuando ambas rectas coinciden. En este caso excepcional $\mathbb{C} \setminus \overline{r(A)}$ estará compuesto por dos semiplanos S_1 y S_2 .

T 6. Sea A un operador cerrado en H . Para cualquier $\lambda \in \mathbb{C} \setminus \overline{r(A)}$, el conjunto $R(A - \lambda I)$ es cerrado en H , $\text{Ker}(A - \lambda I) = \{\emptyset\}$ y $\dim [R(A - \lambda I)]^\perp$ es constante en todo $\mathbb{C} \setminus \overline{r(A)}$, salvo en la situación excepcional señalada más arriba, en cuyo caso $\dim [R(A - \lambda I)]^\perp$ es constante en cada uno de los semiplanos S_1 y S_2 . Además, si $\dim [R(A - \lambda I)]^\perp = 0$ para algún $\lambda \in \mathbb{C} \setminus \overline{r(A)}$ ($\lambda \in S_1$ o $\lambda \in S_2$), entonces $\mathbb{C} \setminus \overline{r(A)}(S_1 \text{ ó } S_2)$ es un subconjunto de $\rho(A)$ y

$$|||R_\lambda(A)||| \leq 1/\text{dist}(\lambda, \overline{r(A)}). \quad (73)$$

Demostración. Véase el teorema sobre invarianza del índice de defecto en el § 3.7 y T 3.6.1. $\square$

Observación. Del Teorema 6 concluimos que si para un operador cerrado existe

$$\lambda \in [\mathbb{C} \setminus \overline{r(A)}] \cap \rho(A) \quad (74)$$

entonces la componente conexa de $\mathbb{C} \setminus \overline{r(A)}$ que contiene a λ está contenida en $\rho(A)$ y, por lo tanto, si $\mathbb{C} \setminus \overline{r(A)}$ es conexo, tendremos que

$$\sigma(A) \subset \overline{r(A)}. \quad (75)$$

Evidentemente, (74) ocurre para cualquier operador acotado en cuyo caso además, $\mathbb{C} \setminus \overline{r(A)}$ es conexo. Luego, para los operadores acotados se cumple (75).

En el próximo epígrafe veremos que (75) también es válido para cualquier operador autoadjunto.

§ 3.6 Espectro de operadores simétricos y autoadjuntos

Dos operadores autoadjuntos A_1 y A_2 definidos en un espacio de Hilbert H de dimensión n , se llaman *unitariamente equivalentes*, si existe una matriz unitaria U en H , tal que

$$A_1 = UA_2U^{-1}.$$

Es bien conocido que dos operadores A_1 y A_2 en el caso de dimensión finita, son unitariamente equivalentes, si ambos tienen los mismos valores propios con la misma multiplicidad algebraica. Luego, el conjunto de los valores propios junto con sus respectivas multiplicidades, constituyen los *invariantes unitarios* de las matrices autoadjuntas.

D 1. (operadores autoadjuntos unitariamente equivalentes)
Sean A_1 y A_2 operadores autoconjugados en los espacios de Hilbert H_1 y H_2 respectivamente. Diremos, que A_1 y A_2 son *unitariamente equivalentes*, si existe un operador isométrico V que transforma H_1 sobre H_2 , (en particular puede ocurrir que $H_1 = H_2$), tal que

$$\begin{aligned}\mathcal{D}(A_2) &= V[\mathcal{D}(A_1)] \\ A_1 &= V^{-1}A_2V.\end{aligned}$$

Se llama *invariante unitario* de un operador autoadjunto A a cualquier propiedad de A que se conserva para todo operador unitariamente equivalente a él, es decir que se conserva para todo operador del tipo UAU^{-1} donde U es un operador unitario en H . Luego, los invariantes unitarios son aquellas propiedades de los operadores autoadjuntos que no dependen de su *representación*.

Según vimos anteriormente, para los operadores autoadjuntos en dimensión finita, toda la información sobre los invariantes unitarios se encuentra en el espectro. En el caso de dimensión infinita *el espectro es*

también un invariante unitario de cualquier operador autoadjunto. En otras palabras, si A es autoadjunto

$$\sigma(A) = \sigma(UAU^{-1}). \quad (76)$$

para cualquier operador U (ejercicio).

Ejemplo 1. El operador de diferenciación P_c , definido sobre toda la recta real mediante (61), (62), es unitariamente equivalente al operador Q de multiplicación por la variable independiente (64) en $L_2(\mathbb{R})$. En efecto, en el § 1.7, vimos que se cumple la identidad de Parseval en $L_2(\mathbb{R})$ para la transformada de Fourier; es decir $\frac{1}{\sqrt{2\pi}}\mathcal{F}$ es un operador unitario en $L_2(\mathbb{R})$, donde $\mathcal{F}$ denota la transformada de Fourier. Además, por el Teorema 1.7.6, tenemos que

$$P_c(u) = \mathcal{F}^{-1}(Q\mathcal{F}(u))$$

para todo $u \in H^1(\mathbb{R})$

En el Ejemplo 3.5.3 vimos que el espectro de Q es continuo y coincide con todo el intervalo $(-\infty, \infty)$. Luego, lo mismo puede decirse para $\sigma(P_c)$, debido a la equivalencia unitaria entre P_c y Q . Hemos obtenido una demostración de (63). $\square$

Ejemplo 2. En el caso de dimensión infinita, el espectro es un invariante unitario bastante pobre para los operadores autoadjuntos, ya que no brinda mucha información acerca de la naturaleza del propio operador. Por ejemplo, el operador de multiplicación por la variable independiente en $L_2[0, 1]$ tiene el mismo espectro que cualquier operador autoadjunto en $L_2[0, 1]$ que tenga como espectro puntual un subconjunto numerable denso en $[0, 1]$ y las funciones propias asociadas constituyan una familia total en $L_2[0, 1]$. En ambos casos el espectro coincide con el intervalo $[0, 1]$, a pesar de que los dos operadores pueden ser totalmente diferentes.

El lector interesado puede encontrar en el libro de M. Reed y B. Simon: *Métodos de la Física Matemática moderna* (T-1) [9], la explicación detallada de porqué $\sigma(A)$ es, en general, un mal invariante. En el

libro citado, el lector podrá ver que los invariantes unitarios dependen de las propiedades de cierta familia de *medidas espectrales* con soporte en $\sigma(A)$ y no precisamente de $\sigma(A)$. $\square$

En el Ejemplo 3.5.2 definimos el operador P_c mediante (56), (57) y vimos que es simétrico pero no autoadjunto y que su espectro coincide con todo $\mathbb{C}$. El teorema siguiente nos muestra que esta situación es característica de una amplia clase de operadores simétricos.

T 1. Sea A un operador simétrico y cerrado en el espacio de Hilbert H . Entonces el espectro de A coincide con uno de los conjuntos siguientes:

- i) el semiplano cerrado $\text{Im} z \geq 0$,
- ii) el semiplano cerrado $\text{Im} z \leq 0$,
- iii) todo el plano,
- iv) un subconjunto del eje real.

El operador A es autoadjunto, si y sólo si se tiene el caso iv).

Demostración. Sea $\lambda = a + ib, b \neq 0$. Como A es simétrico, entonces

$$\|(A - \lambda I)(x)\|^2 \geq b^2 \|x\|^2, \quad (77)$$

para todo $x \in D(A)$. De esta desigualdad y el hecho que A es cerrado, se concluye que $R(A - \lambda I)$ es un subespacio cerrado en H . Además,

$$\text{Ker}(A^* - \lambda I) = [R(A - \bar{\lambda} I)]^\perp. \quad (78)$$

Probemos ahora que la cantidad $\dim \text{Ker}(A^* - \lambda I)$ permanece constante cuando λ varía en cada semiplano abierto $\text{Im} \lambda > 0$, $\text{Im} \lambda < 0$. En efecto, sea $\mu \in \mathbb{C}$ con módulo suficientemente pequeño y tomemos $u \in \mathcal{D}(A^*)$, tal que $u \in \text{Ker}((A^* - (\lambda + \mu)I), \|u\| = 1$. Supongamos que

$\langle u, v \rangle = 0$ para todo $v \in \text{Ker}(A^* - \lambda I)$. Entonces por (78), $u \in R(A - \bar{\lambda}I)$ y, por lo tanto, existe $x \in D(A)$ con la propiedad $(\bar{\lambda}I - A)x = u$.

Tendremos que

$$\begin{aligned} 0 &= \langle ((\lambda + \mu)I - A^*)u, x \rangle = \langle u, (\bar{\lambda}I - A)x \rangle + \mu \langle u, x \rangle \\ &= \|u\|^2 + \mu \langle u, x \rangle, \end{aligned}$$

lo cual nos conduce a una contradicción para $|\mu| < |b|$, ya que por (77),

$$\|x\| \leq \|u\|/|b|.$$

Luego, si $|\mu| < |b|$

$$\text{Ker}(A^* - (\lambda + \mu)I) \cap [\text{Ker}(A^* - \lambda I)]^\perp = \emptyset.$$

Pero entonces un breve razonamiento, utilizando las propiedades de los proyectores ortogonales, (ejercicio) nos muestra que

$$\dim \text{Ker}(A^* - (\lambda + \mu)I) \leq \dim \text{Ker}(A^* - \lambda I). \quad (79)$$

Un razonamiento análogo nos muestra que si $|\mu| < |b|/2$, se cumple también la desigualdad en sentido contrario. En fin hemos probado que la función $\dim \text{Ker}(\lambda I - A^*)$ es localmente constante en los semiplanos superior e inferior de donde se concluye que es constante en cada uno de estos semiplanos abiertos.

De (77) tenemos que para $\text{Im}\lambda \neq 0$ el operador $A - \lambda I$ siempre posee un inverso, definido sobre $R(A - \lambda I)$ que es acotado. De (78) se deduce que el operador inverso está definido en todo H , si y sólo si $\dim \text{Ker}(A^* - \bar{\lambda}I) = 0$. Del hecho que $\dim \text{Ker}(A^* - \lambda I)$ es constante en cada semiplano $\text{Im}\lambda > 0$ y $\text{Im}\lambda < 0$, se tiene que cada uno de esos semiplanos pertenece completamente a $\sigma(A)$ o a $\rho(A)$. Entonces la primera afirmación del teorema se tiene por ser $\sigma(A)$ cerrado. La

segunda parte del teorema es una consecuencia inmediata de T 3.4.1. iii). $\square$

Para los operadores simétricos el rango numérico $r(A)$ (D3.5.3) es siempre un intervalo real que, en general, no brinda información sobre la localización del espectro. De la definición de $r(A)$, para un operador simétrico solamente podemos concluir que

$$\inf r(A)\langle u, u \rangle \leq \langle Au, u \rangle \leq \sup r(A)\langle u, u \rangle. \quad (80)$$

Sin embargo, si sabemos que A es autoadjunto, entonces $\sigma(A)$ es un subconjunto cerrado de $\mathbb{R}$ y no es difícil ver que en este caso se tiene siempre

$$\sigma(A) \subset \overline{r(A)}. \quad (81)$$

En efecto, (81) es obvia si $r(A) = \mathbb{R}$. Si $r(A) \neq \mathbb{R}$, entonces $\mathbb{C} \setminus \overline{r(A)}$ es un subconjunto conexo que tiene intersección no vacía con $\rho(A)$, y (79) se concluye de la observación al Teorema 3.5.6.

T 2. Sea A un operador autoadjunto en el espacio de Hilbert H . Se cumplen entonces las propiedades siguientes:

i) $r(A)$ es un intervalo real y se cumplen (78) y (79);

ii) Si las constantes $-\infty \leq m \leq M \leq +\infty$ son tales que para todo $u \in D(A)$,

$$m\langle u, u \rangle \leq \langle Au, u \rangle \leq M\langle u, u \rangle \quad (82)$$

entonces

$$\sigma(A) \subset [m, M]; \quad (83)$$

iii) El espectro de A incluye a los extremos de $r(A)$, es decir,

$$\inf r(A) \in \sigma(A), \quad (84)$$

$$\sup r(A) \in \sigma(A). \quad (85)$$

iv) Si A es además acotado, entonces

$$\sup_{\|u\|=\|v\|=1} |\langle Au, v \rangle| = \sup_{\|u\|=1} |\langle Au, u \rangle| = \sup\{|\lambda| : \lambda \in r(A)\} \quad (86)$$

En otras palabras,

$$|||A||| = \max\{|\inf r(A)|, |\sup r(A)|\}. \quad (87)$$

Demostración. i) ha sido demostrado previamente al enunciado del teorema. La propiedad ii) se deduce inmediatamente de i), ya que de (82) se tiene que $r(A) \subset [m, M]$.

Demostremos iii). Notemos que es suficiente probar (84), ya que (85) se deduce de (84) por la igualdad evidente

$$\sup r(A) = \inf r(-A). \quad (88)$$

El caso $\inf r(A) = -\infty$ puede darse solamente si A es no acotado, pero en este caso (84) es una consecuencia de la observación al Teorema 3.5.3. Supongamos que $\inf r(A) = m_0 \in \mathbb{R}$. Si $m_0 \notin \sigma(A)$, entonces $m_0 \in \rho(A)$ y por ser $\rho(A)$ abierto existirá un número real $m > m_0$, tal que $m \in \rho(A)$. Por ser $r(A)$ un intervalo, m podría también ser elegido de manera que $m \in r(A)$. Pero esto es una contradicción con la suposición $m \in \rho(A)$, pues en este caso, existiría $u \in D(A)$ con $\|u\| = 1$, tal que $\langle Au, u \rangle = m$ y, por lo tanto

$$\langle Au - mu, u \rangle = 0,$$

es decir $u \in [R(A - mI)]^\perp = \text{Ker}(A - mI)$.

La propiedad iv) se demuestra fácilmente y por ello se deja de ejercicio al lector. $\square$

Observación. En general sabemos de (70), que para un operador acotado A

$$\text{spr } A \leq |||A|||. \quad (89)$$

Según el Teorema 2, la igualdad en (89) se alcanza si A es además autoadjunto. El lector puede verificar que si A es autoadjunto y no necesariamente acotado, también se cumple

$$|||R_\lambda||| = \text{spr} R_\lambda(A) \quad (90)$$

para todo $\lambda \in \rho(A)$. Pero como además (ejercicio)

$$\text{spr} R_\lambda(A) = \frac{1}{\text{dist}(\lambda, \sigma(A))}, \quad (91)$$

entonces tenemos

$$|||R_\lambda||| \leq 1/|\text{Im}\lambda| \quad (92)$$

para todo A autoadjunto, $\lambda \in \rho(A)$. $\square$

Pasemos ahora a caracterizar las diferentes partes del espectro de un operador autoadjunto.

T 3. Sean A un operador autoadjunto en el espacio de Hilbert H y $\lambda \in \mathbb{R}$. Entonces

- i) $\lambda \in \sigma_p(A) \Leftrightarrow \overline{R(A - \lambda I)} \neq H$,
- ii) $\lambda \in \sigma_c(A) \Leftrightarrow R(A - \lambda I) \neq \overline{R(A - \lambda I)} = H$
- iii) $\lambda \in \rho(A) \Leftrightarrow R(A - \lambda I) = H$,
- iv) $\sigma_r(A) = \emptyset$

Demostración. i) Para A autoadjunto y λ real se tiene

$$[R(A - \lambda I)]^\perp = \text{Ker}(A - \lambda I), \quad (93)$$

y el resultado se obtiene de la siguiente cadena de equivalencias:

$$\overline{R(A - \lambda I)} \neq H \Leftrightarrow [R(A - \lambda I)]^\perp \neq \{\emptyset\} \Leftrightarrow \lambda \in \sigma_p(A),$$

donde la última equivalencia es una consecuencia de (93).

ii) Por definición se tiene

$$\begin{aligned} \lambda \in \sigma_c(A) &\Leftrightarrow Ker(A - \lambda I) = \{\Theta\} \text{ y} \\ R(A - \lambda I) &\neq \overline{R(A - \lambda I)} = H. \end{aligned}$$

Pero de (93) se desprende que

$$Ker(A - \lambda I) = \{\Theta\} \Leftrightarrow \overline{R(A - \lambda I)} = H$$

y, por lo tanto, de la definición de $\sigma_c(A)$ puede eliminarse el enunciado $Ker(A - \lambda I) = \{\Theta\}$.

iii) La implicación $\Rightarrow$ es evidente. Para probar la implicación en sentido opuesto, notemos que de (93) se deduce que $Ker(A - \lambda I) = \{\Theta\}$, luego $(A - \lambda I)^{-1}$ existe y es un operador cerrado definido sobre todo H . Por el teorema del grafo cerrado se tiene que $(A - \lambda I)^{-1}$ es continuo y, por lo tanto $\lambda \in \rho(A)$. iv) Los enunciados i) - iii) son excluyentes y agotan todos los casos posibles. Como $\sigma_r(A)$ es disjunto con $\sigma_c(A)$ y $\sigma_p(A)$, entonces $\sigma_r(A) = \phi$. $\square$

Observación. Nótese que el espectro residual está constituido por aquellos valores complejos λ , tales que $\overline{R(A - \lambda I)} \neq H$ pero λ no es un valor propio. El teorema anterior nos muestra que si A es autoadjunto no es posible la existencia de tales puntos. Sin embargo $\sigma_r(A)$ puede ser no vacío para un operador cerrado no autoadjunto (¡encuentre un ejemplo!). $\square$

En el Ejemplo 2 vimos que para operadores autoadjuntos definidos en espacios de dimensión infinita el espectro es un invariante unitario muy pobre. Esta pobreza se manifiesta también en el hecho de que sus componentes $\sigma_p(A)$ y $\sigma_c(A)$ son muy inestables ante perturbaciones. Una muestra de este hecho la brinda el teorema siguiente de Von Neumann.

T 4. A cualquier operador autoadjunto A actuando en un espacio de

Hilbert separable le puede ser añadido un operador compacto, autoadjunto K , con norma arbitrariamente pequeña, de manera que $\sigma_p(A + K)$ sea un subconjunto numerable denso en $\sigma(A + K)$ y además las funciones propias correspondientes a los valores propios $\lambda \in \sigma_p(A + K)$ constituyan una familia total en H (ver D 1.3.5).

Demostración. Vease el libro “Teoría de operadores lineales en espacios de Hilbert” de N.I. Akhiezer y I.M. Glazman. $\square$

Este resultado nos muestra cuan bruscamente puede variar la naturaleza del espectro de un operador autoadjunto ante una “buena perturbación”. Sería conveniente poder caracterizar ciertos subconjuntos del espectro de un operador autoadjunto que fueran “estables” ante tales perturbaciones. Este deseo motiva la siguiente definición.

D 2. (espectro esencial de un operador autoadjunto) El punto λ pertenece al *espectro esencial* del operador autoadjunto A ($\lambda \in \sigma_e(A)$), si se satisface una de las dos condiciones siguientes:

- i) λ es un valor propio de A de multiplicidad infinita.
- ii) $R(A - \lambda I) \neq \overline{R(A - \lambda I)}$.

Observación. En algunas monografías, como por ejemplo en [10] se llama espectro continuo del operador autoadjunto A a lo que nosotros hemos definido como espectro esencial. Nuestra definición de espectro esencial coincide con la dada en [13]. $\square$

De D 2 y T 3 se obtiene inmediatamente el teorema siguiente

T 5. Para el operador autoadjunto A definido en H se cumple:

- i) $\lambda \in \rho(A) \Leftrightarrow R(A - \lambda I) = H$,
- ii) $\lambda \in \sigma(A) \Leftrightarrow R(A - \lambda I) \neq H$,
- iii) $\sigma(A) = \sigma_p(A) \cup \sigma_e(A)$,
- iv) $\lambda \in \sigma_p(A) \Leftrightarrow \overline{R(A - \lambda I)} \neq H$,
- v) $\lambda \in \sigma_e(A) \Leftrightarrow R(A - \lambda I) \neq \overline{R(A - \lambda I)}$ o bien λ es un valor propio de A de multiplicidad infinita. $\square$

Observación. A diferencia de la clasificación del espectro establecida en la Definición 2, la descomposición de $\sigma(A)$ obtenida en T 5 iii) no constituye una partición. Luego, *pueden haber valores propios de multiplicidad finita de A contenidos en $\sigma_e(A)$ (¡busque un ejemplo de tal situación!)*.

Por otra parte veremos en el § 3.11 que para un operador autoadjunto se cumple que *todo punto aislado de su espectro es un valor propio para el cual $R[A - \lambda I]$ es cerrado* y una de estas dos afirmaciones ocurre necesariamente para un operador cerrado arbitrario. En este caso la codimensión de $R[A - \lambda I]$ corresponde a la multiplicidad del valor propio λ y por ello es posible encontrar valores propios aislados de multiplicidad infinita.

De la propiedad enunciada se concluye que, *todo valor propio de multiplicidad finita contenido en el espectro esencial, es necesariamente un punto no aislado del espectro y se caracteriza por las propiedades*

$$\begin{aligned} R(A - \lambda I) &\neq \overline{R(A - \lambda I)} \neq H \\ \text{codim } \overline{[R(A - \lambda I)]} &< +\infty. \end{aligned}$$

Luego, *si para un valor propio λ contenido en el espectro esencial se cumple que $R(A - \lambda I)$ es cerrado, entonces o bien λ es un elemento aislado del espectro o es de multiplicidad infinita.*

Del hecho que $\sigma(A)$ es cerrado y sus únicos posibles puntos aislados pueden ser valores propios se concluye que $\sigma_e(A)$ también es cerrado, no ocurriendo lo mismo en general para $\sigma_c(A)$.

Nótese que $\sigma_e(A)$ se obtiene agregando a $\sigma_c(A)$ todos los valores propios de A de multiplicidad infinita, así como todos los valores propios λ de multiplicidad finita para los cuales $R(A - \lambda I)$ no es cerrado, los cuales son puntos de acumulación en $\sigma_e(A)$. $\square$

El siguiente Teorema de H. Weyl nos muestra que el espectro esencial de cualquier operador autoadjunto es “estable” ante “perturbaciones” compactas autoadjuntas.

T 6. (H. Weyl) Sean A y K operadores autoadjuntos definidos en H y K compacto. Entonces

$$\sigma_e(A + K) = \sigma_e(A).$$

Demostración. Veremos en el Capítulo 4 que todo operador compacto se puede expresar como el límite (con respecto a la norma en el espacio de los operadores acotados en H) de una sucesión de operadores de rango finito. Demostremos entonces primeramente que *para cualquier operador K_n de rango finito $\leq n$ se cumple el enunciado del teorema.*

Haremos la demostración por inducción en n . El operador K_1 se escribe en la forma

$$K_1 = \langle \cdot, \Psi \rangle \phi$$

con $\Psi, \phi \in H$.

Sea $\lambda \in \sigma_e(A)$. Si λ es un valor propio de multiplicidad infinita de A , entonces en el subespacio propio $N_\lambda = \text{Ker}(A - \lambda I)$ se puede encontrar un subespacio de dimensión infinita N_λ^0 constituido por vectores que son ortogonales a Ψ . Obviamente

$$N_\lambda^0 \subset \text{Ker}(A + K_1 - \lambda I),$$

luego λ es también un valor propio de multiplicidad infinita para $A + K_1$.

Si suponemos que $R(A - \lambda I)$ no es cerrado, entonces del hecho que

$$R(A + K_1 - \lambda I) = R(A - \lambda I) + [\phi],$$

donde $[\phi]$ es el espacio vectorial generado por ϕ y la suma es directa si $\phi \notin R(A - \lambda I)$, se tendrá que $R(A + K_1 - \lambda I)$ tampoco es cerrado, luego $\lambda \in \sigma_e(A + K_1)$. En fin hemos probado que $\sigma_e(A) \subset \sigma_e(A + K_1)$. La inclusión en sentido contrario se obtiene intercambiando los papeles de A y $A + K_1$.

Ahora la hipótesis de inducción en n se aplica de manera trivial, y con esto hemos probado la tesis del teorema para operadores compactos de rango finito.

Demostremos ahora que si B es un operador acotado autoadjunto, entonces $\sigma_e(A + B)$ está contenido en una $|||B|||$ -vecindad de $\sigma_e(A)$.

Observemos primeramente que en este enunciado pueden ser intercambiados los operadores A y $A + B$. Entonces es suficiente probar que una $|||B|||$ -vecindad cerrada de cada punto del conjunto $\sigma_e(A)$ contiene al menos un punto del conjunto $\sigma_e(A + B)$. Consideremos lo contrario y llegaremos a una contradicción.

Supongamos que en una $|||B|||$ -vecindad cerrada de algún punto $\lambda \in \sigma_e(A)$ no hay puntos de $\sigma_e(A + B)$. Pero entonces por los comentarios hechos en la observación al Teorema 5, tendremos que en esa vecindad cerrada existen a lo sumo valores propios aislados de multiplicidad finita del operador $A + B$. Sean $(\lambda_i)_1^n$ sus valores propios y P_i los proyectores sobre los respectivos subespacios propios. Por ser λ_i valores propios aislados en $\sigma(A + B)$, se puede encontrar $\varepsilon > 0$, tal que todos los puntos del intervalo $|x - \lambda| < |||B||| + \varepsilon$ pertenecen al conjunto resolvente del operador $T = A + B - \sum_{i=1}^n \lambda_i P_i$.

Pero entonces obviamente

$$\|(T - \lambda I)^{-1}\| < \frac{1}{\|B\| + \varepsilon}.$$

Si utilizamos la notación $S = A - \sum_{i=1}^n \lambda_i P_i$, entonces de la acotación anterior obtendremos para cualquier $f \in \mathcal{D}(A)$

$$\|(S - \lambda I)f\| \geq \|(S + B - \lambda I)f\| - \|Bf\| > \varepsilon \|f\|$$

de lo cual concluimos que $\lambda \in \rho(S)$. Luego no puede ocurrir que $\lambda \in \sigma_e(A)$ ya que A y S se diferencian en un operador de rango finito y por lo demostrado en la primera parte tendríamos que $\lambda \in \sigma_e(S)$. Hemos llegado a una contradicción con lo cual se obtiene el resultado deseado.

Para concluir la demostración del teorema basta combinar los dos resultados parciales obtenidos con la afirmación hecha al inicio de la demostración sobre la caracterización de los operadores compactos. $\square$

Observación. Nótese que si en el teorema anterior hubiéramos definido el espectro esencial de un operador cerrado arbitrario A como el conjunto de puntos de $\sigma(A)$ que satisfacen las condiciones exigidas en la Definición 2, hubiéramos podido obtener el mismo resultado sin exigir la autoadjunticidad de la perturbación compacta K . $\square$

Observación. Más adelante estudiaremos una clasificación más detallada de diferentes partes del espectro de un operador autoadjunto y veremos un teorema muy importante sobre la estabilidad de la llamada *parte absolutamente continua* del espectro ante perturbaciones *nucleares* (véase el *Teorema de Kato-Rosenblum*).

El lector interesado en la teoría de perturbación del espectro para operadores cerrados, así como su relación con la teoría del índice para operadores de tipo Fredholm o semi-Fredholm puede consultar la monografía de T. Kato “Perturbation Theory for Linear Operators” [13].

Es importante destacar que las diferentes clasificaciones de partes del espectro de un operador A tienen un significado físico concreto. Por ejemplo si el operador A describe la dinámica de un sistema mecánico en un entorno de cierta posición de equilibrio de dicho sistema, entonces el espectro de A está asociado a las oscilaciones del sistema alrededor de esa posición de equilibrio. Generalmente el espectro esencial corresponde a “ondas interiores” en el sistema y el espectro discreto a “ondas superficiales”. Los valores propios contenidos en el espectro esencial corresponden a ondas producidas por “estratificación” o la presencia de “inhomogeneidades” en el medio. La parte absolutamente continua del espectro está asociada con la “dispersión” de las ondas interiores al “chocar” con una inhomogeneidad u obstáculo. $\square$

Veamos sin demostración el resultado siguiente, que en cierta medida es el recíproco del T 6 de H. Weyl:

T 7. Sean A y B operadores autoadjuntos y acotados definidos en el espacio de Hilbert separable H . Si $\sigma(A) = \sigma(B)$, entonces se puede hallar un operador unitario U , tal que el operador

$$K = B - UAU^{-1},$$

es compacto.

Demostración. Véase el libro N.I. Ajiezer y I.M. Glazman: *Teoría de operadores lineales en espacios de Hilbert* [10]. $\square$

Concluiremos el presente epígrafe con el siguiente teorema cuya demostración será dada en el Capítulo 6.

T 8. El punto $\lambda \in \mathbb{R}$ pertenece al espectro esencial del operador autoadjunto A , si y sólo si en $\mathcal{D}(A)$ existe una sucesión ortonormada e infinita

de elementos (x_n) para los cuales

$$\lim_{n \rightarrow \infty} (Ax_n - \lambda x_n) = 0.$$

Observación. Si para un operador cerrado cualquiera se define el espectro esencial como el conjunto de los $\lambda \in \mathbb{C}$ para los cuales tiene lugar el enunciado de T 8, entonces se puede probar que sigue siendo válido el teorema de perturbación de H. Weyl sin exigir la autoadjunticidad de los operadores A y K . $\square$

§ 3.7 Teoría de extensión de operadores simétricos y representación de formas sesquilineales

Las demostraciones no expuestas en el texto de este epígrafe pueden ser encontradas en [10], [11], [12] y [13]. A lo largo de todo el epígrafe consideraremos operadores A en H con dominio $D(A)$ denso en H . Comencemos con la siguiente definición:

D1. Se dice que $\lambda \in \mathbb{C}$ es un *punto regular* del operador A si existe una constante $M > 0$ tal que $\|Ax - \lambda x\| \geq M\|x\|$, para todo $x \in D(A)$. El conjunto de los puntos regulares de A se llama *dominio de regularidad* de A y lo denotaremos mediante $DR(A)$.

Obviamente, el dominio de regularidad del operador A está compuesto por aquellos puntos $\lambda \in \mathbb{C}$, tales que existe el operador $(A - \lambda I)^{-1}$ y es además acotado, pero a diferencia de la resolvente no tiene que estar definido en todo H . Por este motivo el conjunto resolvente de A es, en general, un subconjunto de su dominio de regularidad:

$$\rho(A) \subseteq DR(A). \quad (94)$$

E 1. Demostrar que para un operador autoadjunto A se cumple la igualdad en (94).

E 2. Demostrar que el dominio de regularidad de un operador es un subconjunto abierto de $\mathbb{C}$.

E 3. Demostrar que para un operador simétrico A y $z \in \mathbb{C}$ con $\text{Im} z = y \neq 0$ se cumple que

$$\|(A - zI)f\| \geq |y|\|f\|, \quad f \in D(A),$$

y concluir que para un operador simétrico $A : \mathbb{C} \setminus \mathbb{R} \subset DR(A)$

El siguiente teorema, no trivial, es fundamental en lo que sigue:

T 1. En cada componente conexa del dominio de regularidad de un operador A , la dimensión del subespacio $(R[A - \lambda I])^\perp$ permanece constante.

En fin, si $DR(A) = \bigcup_{k=1}^{\infty} U_k$, donde U_k es una componente conexa de $DR(A)$, entonces

$$\dim (R[A - \lambda I])^\perp = \theta_k, \quad \lambda \in U_k$$

donde el entero θ_k no depende de $\lambda \in U_k$.

De esta forma podemos asociar a cada operador A una sucesión de números (θ_k) llamados *índices de defecto de A* . Al subespacio $N_\lambda := (R[A - \lambda I])^\perp$ se le llama *subespacio de defecto de A con respecto a λ* .

Sabemos que

$$N_\lambda = (R[A - \lambda I])^\perp = \text{Ker}(A^* - \bar{\lambda}I) \quad (95)$$

y, por lo tanto, si $\theta_k \neq 0$, entonces todo número complejo $\bar{\lambda}$ tal que $\lambda \in U_k$, es un valor propio de A^* con multiplicidad θ_k .

Si combinamos el resultado del Ejercicio 3 con el Teorema 1, resulta que *un operador simétrico A tiene a lo sumo dos índices de defecto que denotaremos $m(A)$ y $n(A)$:*

$$\dim (R[A - zI])^\perp = \begin{cases} m(A) & \text{si } I_m z < 0 \\ n(A) & \text{si } I_m z > 0. \end{cases}$$

y al par $(m(A), n(A))$ le llamaremos *índices de defecto* del operador simétrico A . Se cumplen las siguientes propiedades que dejamos de ejercicio al lector:

- I. Si el operador simétrico A tiene algún punto regular en el eje real, entonces $m(A) = n(A)$.
- II. Para un operador simétrico A , cualquier $\lambda \in \mathbb{C}$ con $I_m \lambda \neq 0$ es un valor propio de A^* de multiplicidad $m(A)$ si $I_m \lambda > 0$ y $n(A)$ si $I_m \lambda < 0$.
- III. Un operador simétrico A es autoadjunto si y sólo si $m(A) = n(A) = 0$.
- IV. Un operador simétrico A tiene extensiones autoadjuntas, si y solo si $m(A) = n(A)$.

Más adelante veremos que un problema mixto (con condiciones iniciales y condiciones de contorno) se reduce al estudio de una ecuación operacional

$$\frac{du}{dt} = Au \tag{96}$$

sujeta a la condición inicial

$$u(0) = u_0 \tag{97}$$

donde u_0 pertenece a un cierto espacio de Hilbert H cuyos elementos son funciones y $u(t) \in H$ para todo $t > 0$. Como veremos próximamente al definir el operador A debemos dar su dominio de definición $D(A)$ y es precisamente en $D(A)$ donde se recogen las condiciones de contorno que

deben satisfacer las soluciones del problema mixto. En general dentro de $D(A)$ están aquellas funciones de H que satisfacen las condiciones de contorno y a las que “tiene sentido aplicar” la expresión diferencial L (ver Cap. 5), dando lugar a una nueva función de H .

Pero como la solución del problema (96)–(97) debe cumplir que $u(t) \in D(A)$ para todo $t > 0$, puede ocurrir que dicho problema no tenga solución y que haya necesidad de ampliar el dominio de definición del operador A , obteniéndose así una extensión de este operador.

Entonces, generalmente, en lugar de resolver el problema (96)–(97) se resuelve un problema modificado (98)–(97) donde:

$$\frac{du}{dt} = \tilde{A}u \quad (98)$$

y $\tilde{A}$ es una extensión conveniente del operador A .

La extensión del dominio de A está vinculada con la necesidad de añadir a $D(A)$ otras funciones a las que “tenga sentido aplicar” la operación diferencial L en algún sentido generalizado. En muchas ocasiones el operador A en (96) es simétrico y las extensiones $\tilde{A}$ necesarias para poder darle sentido matemático a la solución del problema de Cauchy (96)–(97), son *extensiones autoadjuntas* de A .

Esta exigencia matemática de ampliar el campo de posible soluciones a veces no es muy comprensible por el físico o el ingeniero y conduce siempre a un análisis de lo que representa la “solución matemática” $u(t) \in D(\tilde{A})$ del problema (98)–(97) con respecto a la “solución física” que corresponde al problema (96)–(97). Resulta que es posible justificar la existencia de buenas extensiones autoadjuntas. En efecto, consideremos un ejemplo que se sale un tanto de la temática de este texto pero que es bastante ilustrativo: si A_0 es el operador que representa la energía de una partícula libre en la mecánica cuántica, entonces $A = A_0 + B$ representa la energía de la partícula sometida a un potencial externo que se relaciona con el operador B . Este operador es tal, que A es un operador simétrico pero no necesariamente autoadjunto. Sin embargo, si la partícula no está sumergida en un paso infinito de potencial, entonces

el operador $A_0 + B$ está acotado inferiormente con respecto a la relación de orden $\leq$, introducida en el § 3.4 (ver (46)), es decir:

$$\langle (A_0 + B)x, x \rangle \geq \alpha \langle x, x \rangle, \alpha \in \mathbb{R}, \alpha = \text{const}, x \in D(A_0 + B).$$

Pero para un operador simétrico A acotado inferiormente por αI , se tiene que todo $\beta \in \mathbb{R}$ con $\beta < \alpha$ es un punto regular de A (¡ demuéstrello!), y, por lo tanto, para estos operadores sus índices de defecto coinciden (véase propiedad I) de donde se concluye que admiten extensiones autoadjuntas (véase propiedad IV).

El análisis realizado nos habla, no tan sólo de la importancia de las extensiones autoadjuntas de un operador simétrico, sino también de su existencia en situaciones concretas. Los siguientes teoremas de Von-Neumann permiten caracterizar todas las extensiones autoadjuntas de operadores simétricos, en caso de que éstas existan.

T 2. Sea A simétrico en H y $N_z, N_{\bar{z}}$ cualquier par de subespacios de defecto de A con $I_m z > 0$. Entonces

$$D(A^*) = D(A) \oplus N_{\bar{z}} \oplus N_z$$

(donde la suma directa no es necesariamente ortogonal).

Del Teorema 2 y de (95) se concluye que si $\varphi \in D(A^*)$ entonces φ se expresa de manera única en la forma

$$\varphi = \varphi_0 + g_z + g_{\bar{z}}$$

donde $\varphi_0 \in D(A)$, $g_z \in N_{\bar{z}}$, $g_{\bar{z}} \in N_z$, y

$$A^* \varphi = A \varphi_0 + z g_z + \bar{z} g_{\bar{z}}. \quad (99)$$

T 3. Sea A simétrico y supongamos que $m(A) = n(A)$. Para cada extensión autoadjunta $\tilde{A}$ del operador A , existe un operador isométrico $V : N_{\bar{z}} \rightarrow N_z$, tal que:

$$D(\tilde{A}) = D(A) \oplus (I + V)N_{\bar{z}}$$

y si $f \in D(\tilde{A})$, $f = f_0 + g_z + Vg_{\bar{z}}$, entonces

$$\tilde{A}f = Af_0 + zg_z + \bar{z}Vg_z.$$

La segunda parte de este teorema se obtiene de (99) y del hecho, que toda extensión autoadjunta $\tilde{A}$ del operador simétrico A cumple

$$A \subset \tilde{A} \subset A^*.$$

Este último teorema nos dice que hay una correspondencia biunívoca entre los operadores isométricos de $N_{\bar{z}}$ en N_z y las extensiones autoadjuntas de un operador simétrico A .

Como veremos más adelante, la diferencia fundamental entre la teoría de operadores diferenciales ordinarios y en derivadas parciales, radica en que, para los primeros $\dim N_{\bar{z}} < +\infty$, $\dim N_z < +\infty$, mientras que para los últimos, generalmente

$$\dim N_{\bar{z}} = \dim N_z = +\infty.$$

Luego, para los operadores diferenciales ordinarios simétricos A con $m(A) = n(A)$, el estudio de las extensiones autoadjuntas se reduce a la caracterización del conjunto de los operadores isométricos de $N_{\bar{z}}$ en N_z , el cual coincide con el conjunto de las matrices unitarias en $\mathbb{C}^{n(A)}$.

Ejemplo 1. Hemos visto en los ejemplos 3.3.1 y 3.5.2 que si definimos el operador P_0 en $L_2(a, b)$ con dominio de definición $C_0^\infty(a, b)$, mediante

$$P_0\varphi = i\frac{d\varphi}{dx}, \quad \varphi \in C_0^\infty(a, b),$$

entonces en los casos en que (a, b) es acotado, $(a, b) = (0, +\infty)$ y $(a, b) = (-\infty, +\infty)$ se tiene que $\bar{P}_0 = P_c$ donde $P_c\varphi = i\frac{d\varphi}{dx}$, $\varphi \in \mathcal{D}(P_c)$ donde

$$\mathcal{D}(P_c) = \{\varphi \in H^1(a, b) : \varphi(a) = \varphi(b) = 0\},$$

entendiéndose por $\varphi(a) = \lim_{x \rightarrow -\infty} \varphi(x)$ si $a = -\infty$ y $\varphi(b) = \lim_{x \rightarrow +\infty} \varphi(x)$ si $b = +\infty$.

Además, en los tres casos se cumple que $P_0^* = P_c^* = P$, donde $P\varphi = i\frac{d\varphi}{dx}$, $\varphi \in \mathcal{D}(P) = H^1(a, b)$.

De aquí podemos concluir que los índices de defecto de P_0 y P_c coinciden y son iguales a

$$m(P_0) = m(P_c) = \dim \text{Ker}(P - iI) = \begin{cases} 1 & \text{si } (a, b) \text{ es acotado} \\ 0 & \text{si } (a, b) = (0, +\infty) \\ 0 & \text{si } (a, b) = (-\infty, +\infty) \end{cases}$$

$$n(P_0) = n(P_c) = \dim \text{Ker}(P + iI) = \begin{cases} 1 & \text{si } (a, b) \text{ es acotado} \\ 1 & \text{si } (a, b) = (0, +\infty) \\ 0 & \text{si } (a, b) = (-\infty, +\infty) \end{cases}$$

ya que toda solución de la ecuación

$$i\frac{d\varphi}{dx} \mp i\varphi = 0$$

es proporcional a $\varphi_{\mp}(x) = e^{\pm x}$ y obviamente $\varphi_{-}(x)$ y $\varphi_{+}(x)$ pertenecen a $L_2(a, b)$ cuando el intervalo (a, b) es acotado, pero si $(a, b) = (0, +\infty)$ se tiene que solamente $\varphi_{+}(x) \in L_2(0, +\infty)$, en el caso $(a, b) = (-\infty, +\infty)$ ninguna de las dos funciones pertenece a $L_2(-\infty, +\infty)$.

Luego, de las propiedades III y IV de los índices de defecto obtenemos inmediatamente que en el caso $(a, b) = (-\infty, +\infty)$ el operador P_c es autoadjunto, en el caso $(a, b) = (0, +\infty)$, P_c no tiene extensiones autoadjuntas y en el caso (a, b) acotado sí las tiene.

Veremos ahora cómo se pueden obtener todas las extensiones autoadjuntas del operador P_0 en el caso $(a, b) = (0, 2\pi)$ acotado utilizando los teoremas de representación de Von-Neumann (compárese con Ejemplo 3.4.1). Para ello veremos quiénes son las isometrías V de N_{-i} en N_i , donde en este caso $N_{\mp} = \text{Ker}(P \mp iI)$ coincide con el espacio vectorial generado por la función $\varphi_{\mp}(t) = e^{\pm t}$. Una tal isometría debe actuar de

manera que $Ve^t = \gamma e^{-t}$, donde la constante γ se define de la relación:

$$\int_0^{2\pi} |e^t|^2 dt = |\gamma|^2 \int_0^{2\pi} |e^{-t}|^2 dt. \quad (100)$$

De (100) se deduce inmediatamente que

$$\gamma = \theta e^{2\pi} \quad \text{donde} \quad |\theta| = 1.$$

Luego, si $\tilde{P}_0$ es una extensión autoadjunta del operador P_0 , podemos encontrar un número complejo θ con $|\theta| = 1$, tal que toda $\varphi \in D(\tilde{P}_0)$ se expresa en forma única

$$\varphi = \varphi_0 + \alpha(e^t + \theta e^{2\pi-t}), \quad \alpha \in \mathbb{C}, \varphi_0 \in D(P_0),$$

y, además

$$\tilde{P}_0 \varphi = i\varphi'_0(t) + \alpha i(e^t - \theta e^{2\pi-t}) = i\varphi'(t).$$

Proponemos al lector, en calidad de ejercicio, probar que $D(\tilde{P}_0) = \{\varphi \in H^1(0, 2\pi) : \varphi(2\pi) = \theta_0 \varphi(0)\}$ donde

$$\theta_0 = \frac{\theta + e^{2\pi}}{1 + e^{2\pi}\theta}, \quad |\theta_0| = 1.$$

Además calcule el espectro y la resolvente del operador $\tilde{P}_0$. $\square$

Veamos ahora algunos resultados relativos a las extensiones autoadjuntas de operadores simétricos que son útiles en las aplicaciones a la teoría de operadores diferenciales.

Resulta que si A es un operador simétrico en H con índice de defecto finito (m, m) entonces se puede dar una fórmula llamada *fórmula de M.G. Krein para la resolvente de cualquier extensión autoadjunta $\tilde{A}$ de A* .

De esta fórmula se deduce que *la diferencia entre los resolventes de dos extensiones autoadjuntas $\tilde{A}_1$ y $\tilde{A}_2$ de A es un operador de rango finito $\leq m$* .

A partir de este resultado se prueba el teorema siguiente:

T 4. Todas las extensiones autoadjuntas de un operador simétrico A con dominio denso en H y con índice de defecto finito (m, m) tienen el mismo espectro continuo.

Además, si A es acotado inferiormente: $A \geq \alpha I$ ($\alpha \in \mathbb{R}$), entonces la parte del espectro de cualquier extensión autoadjunta $\tilde{A}$ de A que se encuentra en el semieje $(-\infty, \alpha)$ puede consistir a lo sumo de una cantidad finita de valores propios, la suma de cuyas multiplicidades no excede a m . En general cualquier valor propio de A continúa siendo un valor propio de $\tilde{A}$ cuya multiplicidad puede aumentar a lo sumo en una cantidad que no exceda a m .

Ya hemos visto que si el operador autoadjunto $\tilde{A}$ está acotado inferiormente: $\tilde{A} \geq \beta I$, entonces todo el espectro de $\tilde{A}$ se encuentra contenido en el intervalo $[\beta, +\infty]$. Recíprocamente en virtud de (83), (86) y (87), si para el operador autoadjunto $\tilde{A}$ se cumple

$$-\infty < \beta_- = \inf\{\lambda : \lambda \in \sigma(\tilde{A})\}$$

o bien

$$+\infty > \beta_+ = \sup\{\lambda : \lambda \in \sigma(\tilde{A})\},$$

entonces, en el primer caso se tendrá que $\tilde{A} \geq \beta_- I$, y en el segundo, $\tilde{A} \leq \beta_+ I$.

De este resultado y el teorema anterior se concluye que *toda extensión autoadjunta de un operador simétrico acotado inferiormente y con índice de defecto finito, es también acotada inferiormente.*

Utilizando los teoremas de representación de Von-Neumann se puede probar que si A es un operador simétrico acotado inferiormente por un número $\alpha \in \mathbb{R}$ (no necesariamente de índice finito), entonces *para cada $\mu < \alpha$ se puede encontrar una extensión autoadjunta $\tilde{A}$ de A acotada inferiormente por μ .*

Este resultado nos hace pensar en la existencia de extensiones autoadjuntas de A que preserven su cota inferior α . En efecto, se tiene el teorema siguiente que fué formulado como conjetura por J. Von Neumann:

T 5. Sea A un operador simétrico acotado inferiormente con dominio denso en H y $\alpha_0 = \sup\{\alpha : A \geq \alpha I\}$. Entonces existe al menos una extensión autoadjunta de A que conserva la misma cota inferior α_0 .

Pasemos a explicar el contenido de este importante teorema, para lo cual basta considerar el caso $\alpha_0 = 0$ ya que en caso contrario aplicamos nuestro razonamiento al operador $A - \alpha_0 I$ en lugar de A . Obviamente el supremo de las cotas inferiores de A es α_0 si y solo si el supremo de las cotas inferiores de $A - \alpha_0 I$ es cero.

Si suponemos que el operador simétrico es positivo, entonces el punto $\lambda = -1$ pertenece a su dominio de regularidad y, por lo tanto, tiene sentido definir el operador

$$S = (I - A)(I + A)^{-1}.$$

Obviamente el operador S cuyo dominio de definición es $D(S) = (I + A)D(A)$, satisface la condición de simetría a pesar de que su dominio no tiene que ser denso en H . Además es fácil probar que $S : D(S) \rightarrow H$ es acotado y su norma no excede a 1.

M.G. Krein desarrolló una teoría constructiva a partir de la cual se puede probar que *cualquier operador simétrico acotado S cuyo dominio $D(S) \subset H$ no es denso en H , puede ser extendido a un operador autoadjunto $\tilde{S}$ con $D(\tilde{S}) = H$ preservando la norma, es decir con $\|\tilde{S}\| = \|S\|$.*

Para describir todas las extensiones autoadjuntas del operador simétrico acotado S que preservan su norma se necesita un resultado de I.M.

Guelfand y M.G. Krein:

L 1. Sea B un operador acotado y positivo en H , sea G un subespacio de H . Denotaremos mediante $0(G)$ el conjunto de todos los operadores autoadjuntos C tales que

$$C \leq B \text{ y } Cg = 0 \quad (g \in G).$$

Entonces existe un operador maximal en $0(G)$, es decir un operador que es $\geq$ que cualquier otro operador en el conjunto $0(G)$ y este operador maximal $\tilde{C}$ puede ser expresado en la forma

$$\tilde{C} = B^{1/2} Q B^{1/2},$$

donde Q es el proyector ortogonal sobre $F^\perp$, donde

$$F = \{B^{1/2}g : g \in G\}.$$

Supongamos ahora, que $\tilde{S}_0$ es una extensión autoadjunta fija del operador S y $\tilde{S}$ otra extensión arbitraria, ambas preservando su norma ($\|S\| = 1$). Entonces el operador C definido mediante

$$\tilde{S} = \tilde{S}_0 + C$$

satisface las propiedades

$$-(I + \tilde{S}_0) \leq C \leq (I - \tilde{S}) \quad Cf = 0, \quad f \in D(S).$$

Inversamente, si un operador autoadjunto C tiene estas dos últimas propiedades, entonces el operador $\tilde{S}$ definido por $\tilde{S} = \tilde{S}_0 + C$ será una extensión autoadjunta de S preservando su norma. Luego, usando el Lema 1, obtenemos una descripción de todas las extensiones autoadjuntas del operador S que preservan su norma.

En efecto, si denotamos mediante Q_1 y Q_2 los proyectores ortogonales sobre $[(I + \tilde{S}_0)^{1/2} D(S)]^\perp$ y $[(I - \tilde{S}_0)^{1/2} D(S)]^\perp$ respectivamente, y

construimos los operadores

$$\begin{aligned}\tilde{C}_1 &= (I + \tilde{S}_0)^{1/2} Q_1 (I + \tilde{S}_0)^{1/2}, \\ \tilde{C}_2 &= (I - \tilde{S}_0)^{1/2} Q_2 (I - \tilde{S}_0)^{1/2},\end{aligned}$$

se tiene el teorema siguiente:

T 6. El operador autoadjunto T es una extensión del operador S que preserva la norma de S ($\|S\| = 1$) si y solo si satisface la doble desigualdad

$$\tilde{S}_{\min} \leq T \leq \tilde{S}_{\max},$$

donde

$$\tilde{S}_{\min} = \tilde{S}_0 - \tilde{C}_1, \quad \tilde{S}_{\max} = \tilde{S}_0 + \tilde{C}_2$$

siendo $\tilde{C}_1$ y $\tilde{C}_2$ los operadores definidos más arriba. Además T coincide con $\tilde{S}_{\min}$ si y solo si $D(S)$ es denso en H con respecto a la norma generada por el producto escalar

$$(f, g)_T = (f, g) + (Tf, g).$$

Volvamos de nuevo a considerar el operador acotado y simétrico S construido a partir del operador cerrado, simétrico y positivo A según la fórmula

$$S = (I - A)(I + A)^{-1}$$

con $D(S) = (I + A)D(A)$. Para cada $g \in D(S)$ de la forma $g = f + Af$, se tiene $Sg = f - Af$.

Las tres propiedades fundamentales de S son su *simetría*, el hecho que $\|S\| < 1$ y que $\lambda = -1$ *no es valor propio* de S . Esta última propiedad se deriva del hecho que

$$Sg + g = 2f \neq 0 \quad \text{si} \quad g \neq 0.$$

Notemos que para obtener estas tres propiedades de S no se necesita que $D(A)$ sea denso en H . Recíprocamente, a cada operador simétrico

en H teniendo las propiedades anteriores se le puede hacer corresponder un operador simétrico y positivo A , cuyo dominio no necesariamente será denso en H , mediante

$$D(A) = (I + S)D(S) \ni f = \frac{1}{2}(g + Sg) \rightsquigarrow Af = \frac{1}{2}(g - Sg).$$

En fin, hemos visto que: *la transformación definida por las fórmulas $S = (I - A)(I + A)^{-1}$, $A = (I - S)(I + S)^{-1}$, establece una correspondencia biunívoca entre el conjunto de todos los operadores simétricos acotados S con $\|S\| < 1$ y que no tienen a $\lambda = -1$ como valor propio y el conjunto de todos los operadores simétricos y positivos A .*

Veamos ahora cómo se obtiene la importante conclusión del Teorema 5. Consideremos que $\alpha_0 = 0$ en el Teorema 5 y sea S el operador definido por la fórmula $S = (I - A)(I + A)^{-1}$. Sabemos que S tiene una extensión autoadjunta $\tilde{S}$ definida en todo H con $\|\tilde{S}\| = \|S\| \leq 1$. El operador $\tilde{S}$ satisface la condición de que $\lambda = -1$ no es un valor propio.

En efecto, si para algún $g \neq 0$ ocurriera que $\tilde{S}g + g = 0$, entonces se tendría

$$(g, f + \tilde{S}f) = (g + \tilde{S}g, f) = 0, \quad f \in H,$$

de donde en particular $(g, f + Sf) = 0$, para cualquier $f \in D(S)$ y esto no podría ocurrir ya que $(I + S)D(S) = D(A)$ es denso en H . Luego, el operador $\tilde{S}$ cumple las tres condiciones que permiten definir el operador $\tilde{A} = (I - \tilde{S})(I + \tilde{S})^{-1}$ que es una extensión simétrica positiva de A . Como $\tilde{S}$ es autoadjunto entonces también $\tilde{A}$ es autoadjunto lo cual finalmente prueba el enunciado del Teorema 5. $\square$

Uniando el enunciado del Teorema 5 con el del Teorema 6 concluimos que *todas las extensiones autoadjuntas positivas $\tilde{A}$ de un operador simétrico positivo A con dominio denso en H se expresan en la forma*

$$\tilde{A} = (I - \tilde{S})(I + \tilde{S})^{-1},$$

donde $\tilde{S}$ es cualquier extensión autoadjunta de $S = (I - A)(I + A)^{-1}$

con $\|\tilde{S}\| = \|S\| \leq 1$, es decir $\tilde{S}$ satisface la doble desigualdad

$$\tilde{S}_{\min} \leq \tilde{S} \leq \tilde{S}_{\max} \quad (\text{ver Teorema 6}).$$

Denotemos las extensiones autoadjuntas $\tilde{A}$ correspondientes a los operadores extremos $\tilde{S}_{\min}$ y $\tilde{S}_{\max}$ mediante $\tilde{A}_{\min}$ y $\tilde{A}_{\max}$:

$$\begin{aligned}\tilde{A}_{\min} &= (I - \tilde{S}_{\min})(I + \tilde{S}_{\min})^{-1} \\ \tilde{A}_{\max} &= (I - \tilde{S}_{\max})(I + \tilde{S}_{\max})^{-1}.\end{aligned}$$

El operador simétrico y positivo A con dominio denso en H tendrá una única extensión autoadjunta positiva si y solo si $\tilde{A}_{\min} = \tilde{A}_{\max}$.

Sin embargo, en general $\tilde{A}_{\min} \neq \tilde{A}_{\max}$ y, en este caso, ambos operadores tienen propiedades extremales que lo caracterizan. De ellos dos la extensión más importante es $\tilde{A}_{\min}$ que es llamada *extensión de Friedrichs* del operador simétrico positivo A y se denota también mediante A_F ; $A_F := \tilde{A}_{\min}$. Esta fué obtenida por K. Friedrichs al probar la conjetura de J. Von Neumann enunciada en el Teorema 5. Utilizando la caracterización de $\tilde{S}_{\min}$ dada en el Teorema 6 se puede dar la siguiente caracterización de A_F :

T 7. Entre todas las posibles extensiones autoadjuntas acotadas inferiormente de un operador simétrico positivo A con dominio denso en H , existe una única extensión $\tilde{A}$ tal que $D(\tilde{A}) \subset \mathcal{D}[A]$, donde $\mathcal{D}[A]$ representa el completamiento de $D(A)$ con respecto a la A -norma $\|f\|_A^2 = \|f\|^2 + (Af, f)$. Esta extensión coincide con el operador A_F y para ella se cumple que

$$D[A_F] = \mathcal{D}[A].$$

E 4. Considere el operador A definido en $L_2(0,1)$ mediante

$$\begin{aligned}Au &= \frac{d^2u}{dx^2} = -u''(x) \\ D(A) &= \{u \in C^2[0,1] : u(0) = u'(0) = u(1) = u'(1) = 0\}.\end{aligned}$$

Demuestre que A es simétrico y positivo pero no esencialmente autoadjunto. ¿Cuál es la extensión de Friedrichs del operador A ?

Hemos visto que a cada operador simétrico y positivo A con dominio denso en H se le asocia un espacio normado $(\mathcal{D}(A), \|\cdot\|_A)$ cuyo completamiento denotamos mediante $\mathcal{D}[A]$.

Obviamente la norma en $\mathcal{D}[A]$ proviene de un producto escalar al que denotaremos mediante $\langle \cdot, \cdot \rangle_A$. Resulta que la extensión de Friedrichs A_F es el único operador autoadjunto con $\mathcal{D}(A_F) \subset \mathcal{D}[A]$ y tal que

$$\langle (I + A_F)u, v \rangle = \langle u, v \rangle_A, \quad \forall u \in \mathcal{D}(A_F), \quad v \in \mathcal{D}[A], \quad (101)$$

donde $\langle \cdot, \cdot \rangle$ es una forma sesquilineal, en general no acotada en H . El resultado anterior no es más que una generalización del Teorema de Lax–Milgram (T 2.3.2) sobre representación de formas lineales acotadas mediante operadores acotados, al caso no acotado.

Sin embargo, está lejos de ocurrir que cualquier forma sesquilineal no acotada en H pueda ser representada por algún operador con buenas propiedades. A continuación daremos una caracterización de aquellas formas sesquilineales que pueden ser representadas por operadores autoadjuntos en un sentido que será precisado. Este resultado será una generalización del teorema de extensión 7 y nos proveerá de una técnica para definir operadores autoadjuntos asociados a expresiones diferenciables. Veremos en el Capítulo 5, que este resultado es particularmente útil en aquellas cosas en que no es posible dar una caracterización completa de todas las extensiones autoadjuntas de un operador simétrico, tal y como ocurre generalmente con los operadores diferenciales en derivadas parciales.

En la monografía de T. Kato, *Perturbation theory for linear operations* [13] se hace un estudio más amplio de la teoría de representación de formas sesquilineales mediante operadores llamados sectoriales que son una importante generalización de los operadores autoadjuntos. Nosotros

seguiremos en lo fundamental al estilo de exposición de [13].

D 2. (Formas sesquilineales y cuadráticas) Una forma sesquilineal τ en H es una aplicación $\tau : \mathcal{D} \times \mathcal{D} \rightarrow \mathbb{C}$, donde $\mathcal{D}$ es un subespacio vectorial de H llamado dominio de $\tau : \mathcal{D} = \mathcal{D}(\tau)$. Además τ es lineal con respecto a $u \in \mathcal{D}$ para cada $v \in \mathcal{D}$ fijo y semilineal en $v \in \mathcal{D}$ para cada $u \in \mathcal{D}$ fijo, es decir

$$\tau[u, \lambda v_1 + \mu v_2] = \bar{\lambda}\tau[u, v_1] + \bar{\mu}\tau[u, v_2]$$

donde $\lambda, \mu \in \mathbb{C}$; $u, v_1, v_2 \in \mathcal{D}$ y $\bar{\lambda}, \bar{\mu}$ denotan los respectivos complejos conjugados.

Se dice que τ está densamente definida si $\mathcal{D}(\tau)$ es denso en H . Se llama forma cuadrática asociada con τ a la expresión

$$\tau[u, u] := \tau[u], u \in \mathcal{D}(\tau).$$

Obviamente τ está unívocamente determinada por su forma cuadrática asociada ya que se tiene la identidad:

$$\tau[u, v] = \frac{1}{4}(\tau[u + v] - \tau[u - v] + i\tau[u + iv] - i\tau[u - iv]). \quad (102)$$

Diremos que dos formas τ_1 y τ_2 son iguales si $\mathcal{D}(\tau_1) = \mathcal{D}(\tau_2) = \mathcal{D}$ y además

$$\tau_1[u, v] = \tau_2[u, v], \quad \forall u, v \in \mathcal{D}.$$

Notemos, que de (102) se concluye que para la igualdad de las formas τ_1 y τ_2 es necesario y suficiente la igualdad de las formas cuadráticas asociadas. El concepto de extensión o restricción de formas ($\tau_1 \subset \tau_2$ ó $\tau_2 \subset \tau_1$) se define de manera obvia como en el caso de los operadores.

El hecho de que el dominio de una forma no tiene que coincidir con todo H , dificulta la definición de operaciones algebraicas con formas tal y como ocurre para los operadores no acotados.

La *suma* $\tau = \tau_1 + \tau_2$ de dos formas se define como

$$\tau[u, v] = \tau_1[u, v] + \tau_2[u, v], \quad \mathcal{D}(\tau) = \mathcal{D}(\tau_1) \cap \mathcal{D}(\tau_2).$$

El *producto* $\alpha\tau$ de la forma τ por el escalar α , se define como

$$(\alpha\tau)[u, v] = \alpha\tau[u, v], \quad \mathcal{D}(\alpha\tau) = \mathcal{D}(\tau).$$

Por definición la *forma unidad* es igual al producto escalar en H :

$$1[u, v] = \langle u, v \rangle, \quad \mathcal{D}(1) = H,$$

y la forma 0 es la que toma el valor cero para todos $u, v \in H$:

$$0[u, v] = 0, \quad \mathcal{D}(0) = H.$$

De estas definiciones se concluye que para cualquier forma $\tau, 0\tau \subset 0$ y además tiene sentido definir $\tau + \alpha := \tau + \alpha I$, mediante

$$(\tau + \alpha)[u, v] = \tau[u, v] + \alpha\langle u, v \rangle, \quad \mathcal{D}(\tau + \alpha) = \mathcal{D}(\tau).$$

La forma τ se llama *simétrica*, si

$$\tau[u, v] = \overline{\tau[v, u]}, \quad u, v \in \mathcal{D}(\tau).$$

Obviamente τ es una forma simétrica si y solo si $\tau[u] \in \mathbb{R}, \forall u \in \mathcal{D}(\tau)$ (¡ demuéstrello!). A cada forma τ se le asocia la llamada forma *adjunta*, mediante

$$\tau^*[u, v] = \overline{\tau[v, u]}, \quad \mathcal{D}(\tau^*) = \mathcal{D}(\tau).$$

Luego τ es simétrica si y solo si $\tau = \tau^*$. Se cumple la identidad

$$(\alpha_1\tau_1 + \alpha_2\tau_2)^* = \overline{\alpha_1}\tau_1^* + \overline{\alpha_2}\tau_2^*.$$

A cada forma τ se le asocian dos formas

$$\operatorname{Re}\tau := \frac{1}{2}(\tau + \tau^*), \quad \operatorname{Im}\tau = \frac{1}{2i}(\tau - \tau^*)$$

llamada *parte real e imaginaria de τ* respectivamente, las cuales son simétricas y cumplen que

$$\tau = \operatorname{Re} \tau + i \operatorname{Im} \tau.$$

Nótese que esta expresión no significa que $\operatorname{Re} \tau[u, v]$ ó $\operatorname{Im} \tau[u, v]$ sean números reales para todos $u, v \in \mathcal{D}(\tau)$, sino únicamente cuando $u = v$.

La forma simétrica τ se llama *acotada inferiormente* si el conjunto numérico

$$\Theta(\tau) = \{\tau[u] : u \in \mathcal{D}(\tau), \|u\| = 1\} \quad (103)$$

está acotado inferiormente. Nótese que si γ es una cota inferior del conjunto $\Theta(\tau)$, entonces

$$\tau[u] \geq \gamma \|u\|^2, \quad u \in \mathcal{D}(\tau) \quad (104)$$

lo cual escribiremos simplemente $\tau \geq \gamma$. El ínfimo del conjunto $\Theta(\tau)$ coincide con el supremo de los $\gamma \in \mathbb{R}$ para los cuales tiene lugar la propiedad (104) y se llama *cota inferior* de τ : γ_τ . Si $\tau \geq 0$, diremos que τ es no negativa y de forma análoga definiremos la *acotación superior*, *cota superior*, *no positividad*, etc.

E 5. Demuestre que si τ es una forma simétrica acotada superior e inferiormente, entonces ella es acotada, es decir

$$|\tau[u, v]| \leq M \|u\| \|v\|, \quad u, v \in \mathcal{D}(\tau)$$

donde M es el valor máximo entre el módulo de la cota superior y el módulo de la cota inferior. Más generalmente, pruebe que si τ_1 y τ_2 son simétricas, τ_2 es no negativa y para todo $u \in \mathcal{D} = \mathcal{D}(\tau_1) = \mathcal{D}(\tau_2)$ se cumple que

$$|\tau_1[u]| \leq M \tau_2[u],$$

entonces

$$|\tau_1[u, v]| \leq M (\tau_2[u])^{1/2} (\tau_2[v])^{1/2}.$$

E 6. Pruebe que si τ es una forma simétrica no negativa, se cumplen las desigualdades

$$\begin{aligned} |\tau[u, v]| &\leq (\tau[u])^{1/2}(\tau[v])^{1/2} \leq \frac{1}{2}(\tau[u] + \tau[v]), \\ (\tau[u + v])^{1/2} &\leq (\tau[u])^{1/2} + (\tau[v])^{1/2}, \\ \tau[u + v] &\leq 2\tau[u] + 2\tau[v], \end{aligned}$$

para todos $u, v \in \mathcal{D}(\tau)$.

Para una forma τ no necesariamente simétrica se define el *rango numérico* de τ mediante la fórmula (103).

Al igual que en el caso de los operadores, $\Theta(\tau)$ es un subconjunto convexo en $\mathbb{C}$. Luego para una forma simétrica, la acotación inferior es equivalente a que $\Theta(\tau)$ sea un intervalo de $\mathbb{R}$ acotado inferiormente. Generalizando esta propiedad diremos que la forma τ es *dissipativa* si

$$\Theta(\tau) \subset \{\operatorname{Re} \lambda \geq 0\}$$

y más generalmente es *sectorial* si $\Theta(\tau)$ está contenido en un sector de la forma

$$|\arg(\lambda - \gamma)| \leq \vartheta, \quad 0 \leq \vartheta < \pi/2, \quad \gamma \in \mathbb{R},$$

lo cual significa que

$$\operatorname{Re} \tau \geq \gamma, \quad \text{y} \quad (105)$$

$$|(\operatorname{Im} \tau)[u]| \leq (\tan \vartheta)[(\operatorname{Re} \tau) - \gamma][u], \quad u \in \mathcal{D}(\tau). \quad (106)$$

E 7. Pruebe que si τ es una forma sectorial se cumplen las desigualdades:

$$\begin{aligned} |(\operatorname{Re} \tau - \gamma)[u, v]| &\leq \{(\operatorname{Re} \tau - \gamma)[u]\}^{1/2} \{(\operatorname{Re} \tau - \gamma)[v]\}^{1/2}, \\ |(\operatorname{Im} \tau)[u, v]| &\leq (\tan \vartheta) \{(\operatorname{Re} \tau - \gamma)[u]\}^{1/2} \{(\operatorname{Re} \tau - \gamma)[v]\}^{1/2}, \end{aligned}$$

$$\begin{aligned}
|(\tau - \gamma)[u, v]| &\leq (1 + \tan \vartheta) \{(\operatorname{Re} \tau - \gamma)[u]\}^{1/2} \{(\operatorname{Re} \tau - \gamma)[v]\}^{1/2}, \\
(\operatorname{Re} \tau - \gamma)[u] &\leq |(\tau - \gamma)[u]| \leq (\sec \vartheta)(\operatorname{Re} \tau - \gamma)[u], \\
|(\tau - \gamma)[u + v]|^{1/2} &\leq (\sec \vartheta)^{1/2} \{|(\tau - \gamma)[u]|^{1/2} + |(\tau - \gamma)[v]|^{1/2}\}, \\
|(\tau - \gamma)[u + v]| &\leq 2(\sec \vartheta) \{|(\tau - \gamma)[u]| + |(\tau - \gamma)[v]|\}.
\end{aligned}$$

Ejemplo 2. Sea T un operador en H y definamos

$$\tau[u, v] = \langle Tu, v \rangle, \quad \mathcal{D}(\tau) = \mathcal{D}(T).$$

Obviamente T y τ tienen el mismo rango numérico. En particular T es simétrico, acotado inferiormente o acotado superiormente, si y solo si τ lo es. $\square$

Ejemplo 3. Sea S un operador definido en H_1 con imagen en H_2 y pongamos

$$\tau[u, v] = \langle Su, Sv \rangle_2, \quad \mathcal{D}(\tau) = \mathcal{D}(S). \quad (107)$$

La forma τ es simétrica y no negativa en H_1 . En particular, si tomamos $H_1 = L_2(0, 1)$, $H_2 = \mathbb{C}$ con

$$S_0 u = u(0), \quad \mathcal{D}(S) = C[0, 1], \quad (108)$$

entonces

$$\tau_0[u, v] = u(0)\overline{v(0)}, \quad \mathcal{D}(\tau_0) = C[0, 1], \quad (109)$$

es la forma asociada al operador S_0 según la definición (107). Obviamente τ_0 es simétrica, no negativa y densamente definida. $\square$

Ejemplo 4. Sea $H = L_2(\mathbb{R}^3)$ y definamos

$$\tau[u, v] = \int_{\mathbb{R}^3} [\operatorname{grad} u(x) \operatorname{grad} \overline{v(x)} + q(x)u(x)\overline{v(x)}] dx,$$

donde q es una función compleja localmente integrable. Definiendo $\mathcal{D}(\tau) = C_0^1(\mathbb{R}^3)$ ó $\mathcal{D}(\tau) = C_0^\infty(\mathbb{R}^3)$, en cualquier caso tendremos una forma densamente definida que será simétrica si q es real y sectorial si la imagen de q se encuentra contenida en algún sector. Estas propiedades de τ pueden conservarse extendiendo su dominio a algún conjunto de funciones con soporte no compacto. $\square$

Consideremos ahora que τ es una forma sectorial. Entonces diremos que la sucesión de vectores (u_n) es τ -convergente al elemento $u \in H$, y lo denotaremos mediante el símbolo

$$u_n \xrightarrow[\tau]{}, (n \rightarrow \infty),$$

si $u_n \in \mathcal{D}(\tau)$, $u_n \rightarrow u$ en H , y $\tau[u_n - u_m] \rightarrow 0$ cuando $n, m \rightarrow \infty$. Nótese que no necesariamente $u \in \mathcal{D}(\tau)$ y compárese esta noción de convergencia con la noción de convergencia respecto a la A -norma introducida en el Teorema 7. En efecto, *la sucesión (u_n) es de Cauchy con respecto a la A -norma, si y solo si, es τ -convergente en H para la forma τ definida a partir del operador A como en el Ejemplo 2.* Luego $\mathcal{D}[A]$ en el Teorema 7 no es más que el conjunto de los límites de las respectivas sucesiones τ -convergentes.

Es inmediato comprobar que la τ -convergencia es equivalente a la $\tau + \alpha$ -convergencia para cualquier escalar α y a la $\text{Re}\tau$ -convergencia. Esta última afirmación se concluye de la cuarta desigualdad en el Ejercicio 7. De este resultado obtenemos que *la τ -convergencia para cualquier forma sectorial τ es equivalente a la τ_0 -convergencia para alguna forma simétrica y no negativa τ_0 .* Nótese que basta tomar

$$\tau_0 = \text{Re}\tau - \gamma. \quad (110)$$

Luego, según el resultado del Ejercicio 6, el conjunto en H de los límites de sucesiones τ -convergentes para una forma sectorial τ constituye un espacio vectorial.

En fin, hemos visto que la sucesión (u_n) es τ -convergente en H , si y

sólo si es una sucesión de Cauchy con respecto a la norma

$$\|u\|_0 = (\|u\|^2 + \tau_0[u])^{1/2} \quad (111)$$

donde $\|\cdot\|$ denota la norma asociada al producto escalar en H y τ_0 es la forma simétrica no negativa definida en (110).

Si denotamos mediante $\mathcal{D}[\tau]$ al conjunto de los límites de sucesiones τ -convergentes en H , entonces obviamente $\mathcal{D}[\tau]$ es el completamiento de $\mathcal{D}(\tau)$ con respecto a la norma (111).

D 3. Diremos que la forma sectorial τ es *cerrada* si $u_n \xrightarrow{\tau} u$ implica que

$$u \in \mathcal{D}(\tau) \text{ y } \tau[u_n - u] \rightarrow 0.$$

La forma sectorial τ se llama *cerrable* si tiene una extensión cerrada.

De los comentarios anteriores se concluye inmediatamente las propiedades siguientes:

P 1. Sea τ una forma sectorial en H . Entonces se cumple:

- i) τ es cerrada si y solo si $\operatorname{Re} \tau$ es cerrada.
- ii) Si τ es cerrada entonces $\tau + \alpha$ es cerrada para cualquier $\alpha \in \mathbb{C}$. Recíprocamente, si $\tau + \alpha$ es cerrada para algún $\alpha \in \mathbb{C}$, entonces τ es cerrada.
- iii) τ es cerrada si y solo si el espacio prehilbertiano $\mathcal{D}(\tau)$ provisto de la norma (111), es completo
- iv) τ es cerrada si y solo si $\mathcal{D}(\tau) = \mathcal{D}[\tau]$.
- v) Si τ es además acotada, entonces será cerrada, si y solo si $\mathcal{D}(\tau)$ es un subespacio vectorial cerrado en H .

- vi) La forma τ es cerrable si y solo si $u_n \xrightarrow{\tau} 0$ implica que $\tau[u_n] \rightarrow 0$. Cuando esta condición es satisfecha, τ tiene una *clausura* (menor extensión cerrada) $\tilde{\tau}$ cuyo dominio coincide con $\mathcal{D}[\tau]$: $\mathcal{D}(\tilde{\tau}) = \mathcal{D}[\tau]$. Más exactamente $\tilde{\tau}$ está definida de la manera siguiente: $\mathcal{D}(\tilde{\tau})$ es el conjunto de aquellos $u \in H$ tales que existe una sucesión (u_n) con $u_n \xrightarrow{\tau} u$, y se define, para cualesquiera $u_n \xrightarrow{\tau} u, v_n \xrightarrow{\tau} v$,

$$\tilde{\tau}[u, v] = \lim_{n \rightarrow \infty} \tau[u_n, v_n]. \quad (112)$$

Cualquier extensión cerrada de τ es también una extensión de $\tilde{\tau}$.

- vii) Sea τ una forma sectorial cerrada, $u_n \in \mathcal{D}(\tau), u_n \rightarrow u$ (donde $\rightarrow$ denota la convergencia débil en H) y $(\tau[u_n])$ acotada. Entonces $u \in \mathcal{D}(\tau)$ y si además $u_n \rightarrow u$ en H , entonces $\text{Re} \tau[u] < \liminf \text{Re} \tau[u_n]$.

Demostración. i) y ii) se dejan de ejercicio al lector. Para la demostración de iii) notemos primeramente que τ es cerrada si y sólo si es cerrada la forma simétrica no negativa τ_0 definida en (110). Pero obviamente τ_0 es cerrada si y sólo si $\mathcal{D}(\tau_0) = \mathcal{D}(\tau)$ es completo con respecto a la norma (111). El punto iv) es otra versión de iii).

Para probar v), notemos que si τ es acotada entonces se puede elegir γ en (110) de manera que la norma (111) sea equivalente a la norma usual en H . Luego, de iii) se concluye inmediatamente el resultado de v).

Pasemos a demostrar el punto vi). Obviamente la condición enunciada es necesaria ya que si τ_1 es una extensión cerrada de τ , entonces de $u_n \xrightarrow{\tau} 0$ se obtiene que

$$\tau[u_n] = \tau_1[u_n] = \tau_1[u_n - 0] \rightarrow 0.$$

Para demostrar la suficiencia consideremos $\mathcal{D}(\tilde{\tau})$ definido como en el enunciado del punto vi). El lector puede verificar en calidad de ejercicio,

que para cualquier forma sectorial τ se cumple, que si $u_n \xrightarrow{\tau} u, u_n \xrightarrow{\tau} v$ entonces $\lim \tau[u_n, v_n]$ existe y que además si τ es cerrada entonces este límite coincide con $\tau[u, v]$. Luego el límite en el miembro derecho de (112) existe para todo $u, v \in \mathcal{D}(\tilde{\tau})$. Demostremos que si se satisface la condición impuesta a τ en el punto vi) entonces el límite en (112) depende solamente de u y v y no de la elección de la sucesión (u_n) en $\mathcal{D}(\tau)$ tales que $u'_n \xrightarrow{\tau} u, v'_n \xrightarrow{\tau} v$. Entonces $u'_n - u_n \xrightarrow{\tau} 0, v'_n - v_n \xrightarrow{\tau} 0$, y por lo tanto, $\tau[u'_n - u_n] \rightarrow 0, \tau[v'_n - v_n] \rightarrow 0$ debido a la suposición hecha sobre τ .

Pero notemos que

$$\tau[u'_n, v'_n] - \tau[u_n, v_n] = \tau[u'_n - u_n, v'_n] + \tau[u_n, v'_n - v_n]$$

el miembro derecho en esta igualdad tiende a cero, de donde concluimos que el límite en (112) no depende de la elección de las sucesiones (u_n) y (v_n) . En fin, hemos podido definir una forma sectorial $\tilde{\tau}$ unívocamente determinada por τ y que obviamente es una extensión de τ .

Si ahora denotamos mediante H_τ (resp. $H_{\tilde{\tau}}$) al espacio normado definido sobre el espacio vectorial $\mathcal{D}(\tau)$ (resp. $\mathcal{D}(\tilde{\tau})$) provisto de la norma (111) correspondiente al operador τ (resp. $\tilde{\tau}$), entonces claramente se ve que $H_{\tilde{\tau}}$ es el completamiento de H_τ . Pero por el punto iii) concluimos que $\tilde{\tau}$ es cerrada.

Demostremos finalmente el punto vii). Notemos que por ser τ cerrada,

$$\mathcal{D}(\tau) = \mathcal{D}[\tau],$$

y sea H_τ el correspondiente espacio de Hilbert definido más arriba.

Por ser $(\tau[u_n])$ y $(\|u_n\|)$ sucesiones numéricas acotadas (ver T 2.4.8), tendremos que la sucesión (u_n) está acotada con respecto a la norma (111) del espacio H_τ . Pero H_τ es un espacio de Hilbert por ser τ cerrada, y del teorema de Banach-Alaoglu (T 2.4.10) concluimos la existencia de una subsucesión $(u_n^{(1)})$ de la sucesión (u_n) , que converge a un elemento u_0 en H_τ . Si aplicamos el resultado que obtendremos en el primer teorema

de representación (T 11), concluimos que existe un operador autoadjunto y no negativo T_0 en H con $\mathcal{D}(T_0) \subset \mathcal{D}(\tau_0)$ y que representa a la forma τ_0 definida en (110) en el sentido siguiente:

$$\langle T_0 u, v \rangle = \tau_0[u, v]$$

para todo $u \in \mathcal{D}(T_0)$, $v \in \mathcal{D}(\tau_0)$.

Como la sucesión $(u_n^{(1)})$ converge débilmente a u_0 en H_τ , tendremos que para cualquier $v \in \mathcal{D}(T_0)$ se cumple

$$\begin{aligned} \langle u_n^{(1)}, (T_0 + I)v \rangle &= (\tau_0 + 1)[u_n^{(1)}, v] \\ &= \langle u_n^{(1)}, v \rangle_0 \rightarrow \langle u_0, v \rangle_0 \\ &= (\tau_0 + 1)[u_0, v] \\ &= \langle u_0, (T_0 + I)v \rangle, \end{aligned} \quad (113)$$

donde $\langle \cdot, \cdot \rangle$ denota el producto escalar en H y $\langle \cdot, \cdot \rangle_0$ el producto escalar asociado a la norma definida en (111).

Por otra parte como la sucesión $(u_n^{(1)})$ converge débilmente a u en H , tendremos que

$$\langle u_n^{(1)}, (T_0 + I)v \rangle \rightarrow \langle u, (T_0 + I)v \rangle. \quad (114)$$

Comparando (113) con (114), obtendremos que

$$\langle u - u_0, (T_0 + I)v \rangle = 0 \text{ para todo } v \in \mathcal{D}(T_0). \quad (115)$$

Pero el punto -1 pertenece al conjunto resolvente de T_0 ya que $T_0 \geq 0$ y, por lo tanto, la imagen de $T_0 + I$ recorre todo el espacio H . Entonces de (115) tendremos que

$$u = u_0 \in H_\tau = \mathcal{D}(\tau).$$

Por otra parte tenemos que $\|u\|_0 = \|u_0\|_0 \leq \liminf \|u_n^{(1)}\|_0$, luego si $u_n \rightarrow u$, se concluye inmediatamente que $\text{Re} \tau[u] \leq \liminf \text{Re} \tau[u_n^{(1)}]$.

Como esta desigualdad es válida para cualquier subsucesión de (u_n) siempre que converge débilmente a u en H_τ , de esta desigualdad se deduce la desigualdad requerida en el punto vii), con lo cual se culmina la demostración. $\square$

Observación. El punto vii) de la proposición anterior resulta muy útil para probar que un elemento $u \in H$ pertenece al dominio de alguna forma sectorial cerrada. Sin embargo de este resultado no podemos concluir que la sucesión (u_n) (ni tan siquiera una subsucesión de ella) converja a u en H_τ . $\square$

Ejemplo 5. Mostremos a continuación una aplicación del punto vii) de la Proposición 1 a la caracterización de funciones del espacio de Soboliev $H^1(\Omega)$, donde Ω es un abierto acotado de $\mathbb{R}^n$. Consideremos la forma definida en $L_2(\Omega)$ teniendo $\mathcal{D}(\tau) = C^1(\overline{\Omega})$, mediante

$$\tau[u, v] = \int_{\Omega} \nabla u \nabla \bar{v} dx. \quad (116)$$

La forma τ es simétrica, no negativa y densamente definida, pero no es cerrada como fácilmente puede comprobar el lector.

Si la sucesión (u_n) en $C^1(\overline{\Omega})$ cumple que $u_n \xrightarrow{\tau} 0$, entonces (u_n) es una sucesión de Cauchy en $H^1(\Omega)$ y, por lo tanto, converge a la función idénticamente nula en $H^1(\Omega)$, de donde se concluye que $\tau[u_n] \rightarrow 0$. Pero entonces aplicando el punto vi) de la Proposición 1, tendremos que τ es cerrable y su clausura $\bar{\tau}$ tiene como dominio precisamente a $H^1(\Omega)$: $\mathcal{D}(\bar{\tau}) = \mathcal{D}[\tau] = H^1(\Omega)$.

Supongamos ahora que tenemos una sucesión (f_n) en $C^1(\overline{\Omega})$ que converge a la función $f \in L_2(\Omega)$ en sentido de funciones generalizadas,

es decir:

$$\int_{\Omega} f_n \varphi dx \xrightarrow{n \rightarrow \infty} \int_{\Omega} f \varphi dx, \quad (n \rightarrow \infty) \quad \forall \varphi \in C_0^\infty(\Omega), \quad (117)$$

y además, que se cumple

$$\|f_n\|_2 + \|\nabla f_n\|_2 \leq M, \quad n = 1, 2, \dots \quad (118)$$

donde la constante M no depende de n y $\|\cdot\|_2$ denota la norma usual en $L_2(\Omega)$. Si aplicamos el Teorema 2.4.9, entonces de (117), (118) y teniendo en cuenta que $C_0^\infty(\Omega)$ es un subconjunto denso en $L_2(\Omega)$, concluimos que la sucesión (f_n) converge débilmente a f en $L_2(\Omega)$. Por otra parte de (118) se tiene que la sucesión numérica $\tau[f_n] = \|\nabla f_n\|_2^2$ está acotada. Aplicando el punto vii) de la Proposición 1, obtenemos que $f \in H^1(\Omega)$.

En fin hemos obtenido el enunciado siguiente: *la función $f \in H^1(\Omega)$ si y solo si existe una sucesión (f_n) en $C^1(\overline{\Omega})$ con las propiedades (117) y (118).* $\square$

E 8. a) Demuestre una caracterización *similar* a la obtenida en la observación anterior, para las funciones de $(H^0)^1(\Omega)$, sustituyendo $C^1(\overline{\Omega})$ por $C_0^1(\overline{\Omega})$ (funciones continuamente derivables en Ω con soporte compacto).

b) ¿Como se enunciaría esta caracterización en el caso $\Omega = \mathbb{R}^n$?

Ejemplo 6. La forma simétrica no negativa τ definida en el Ejemplo 3 es cerrada (resp. cerrable) si y solo si el operador S es cerrado (resp. cerrable).

Para comprobar esto es suficiente notar que $u_n \xrightarrow{\tau} u$ es equivalente a $u_n \rightarrow u, (Su_n)$ es de Cauchy, y la propiedad $\tau[u_n - u] \rightarrow 0$ es equivalente a $Su_n \rightarrow Su$. La forma (109) asociada al operador $S_0 u = u(0)$ no es cerrable en $L_2(0,1)$ por no serlo el operador S_0 . En efecto, puede ocurrir que la sucesión de funciones continuas (u_n) converja a la función

idénticamente nula en $L_2(0, 1)$ sin que necesariamente $(u_n(0))$ converja a cero en $\mathbb{C}$.

Si definiéramos la forma

$$\tau_1[u, v] = \int_0^1 u'(x) \overline{v'(x)} dx + u(0) \overline{v(0)}, \quad (119)$$

en $L_2(0, 1)$ con $\mathcal{D}(\tau_1) = C^1[0, 1]$, resulta que τ_1 sí es cerrable. En efecto si $u_n \rightarrow 0$ en $L_2(0, 1)$ y $\tau_1[u_n - u_m] \rightarrow 0$ cuando $n, m \rightarrow \infty$, entonces la sucesión (u_n) será una sucesión de Cauchy en $H^1(0, 1)$, luego es convergente en $H^1(0, 1)$ y debe converger a la función idénticamente nula. Por otra parte, del teorema de inmersión de Soboliev (Corolario 1.7.1) se tiene la desigualdad

$$\|u\|_\infty \leq C \|u\|_{H^1(0,1)}, \quad \forall u \in H^1(0,1),$$

donde la constante C no depende de u y $\|\cdot\|_\infty$ denota la norma de la convergencia uniforme en $C[0, 1]$. De esta desigualdad se obtiene que $|u_n(0)| \rightarrow 0$ y, por lo tanto, $\tau_1[u_n] \rightarrow 0$ de donde concluimos que τ_1 es cerrable. Dejamos de ejercicio al lector comprobar que

$$\mathcal{D}[\tau_1] = \mathcal{D}(\bar{\tau}_1) = H^1(0, 1). \quad \square$$

E 9. Demuestre que la forma definida en el Ejemplo 4 con q real y acotado inferiormente es cerrable y calcule $\mathcal{D}[\tau]$.

E 10. Sea τ una forma sectorial cerrable. Demuestre que el rango numérico $\Theta(\tau)$ de τ es un subconjunto denso en el rango numérico $\Theta(\bar{\tau})$ de la clausura $\bar{\tau}$ de τ .

E 11. Demuestre que toda forma sectorial cerrable cuyo dominio coincide con todo el espacio H , es acotada.

Cuando τ es una forma sectorial cerrada diremos que una subvariedad $\mathcal{D}'$ de $\mathcal{D}(\tau)$ constituye un *núcleo* (en inglés, *core*) de τ si la restricción τ' de τ con dominio $\mathcal{D}'$ tiene como clausura a $\tau : \tilde{\tau}' = \tau$. Es claro que $\mathcal{D}'$ es un núcleo de τ si lo es de $\tau + \alpha$ para cualquier escalar α . Además si τ es acotada, $\mathcal{D}'$ es denso en $\mathcal{D}(\tau)$, que en este caso es una subvariedad cerrada en H .

E 12. Sea τ una forma sectorial cerrada. Demuestre que la subvariedad lineal $\mathcal{D}'$ de $\mathcal{D}(\tau)$ es un núcleo para τ , si y solo si $\mathcal{D}'$ es denso en el espacio de Hilbert H_τ asociado a τ .

Un operador T se llama *sectorial* si su rango numérico (ver D 3.5.3) es un subconjunto de algún sector de la forma

$$|\arg(\lambda - \gamma)| \leq \vartheta, \quad 0 \leq \vartheta < \pi/2, \quad \gamma \in \mathbb{R}.$$

Dado el operador sectorial T se puede definir la forma sesquilineal

$$\tau[u, v] = \langle Tu, v \rangle \text{ con } \mathcal{D}(\tau) = \mathcal{D}(T) \quad (120)$$

la cual obviamente es sectorial con vértice γ y semiángulo ϑ .

En el caso particular con que $\vartheta = 0$ sabemos que T es cerrable, lo cual no tiene que ocurrir para $\vartheta \neq 0$. Sin embargo, se tiene el siguiente resultado.

T 8. Todo operador sectorial es cerrable en sentido de formas sesquilineales (form-cerrable), es decir, la forma τ definida en (120) es cerrable.

Demostración. Sin pérdida de generalidad podemos asumir que T tiene su vértice en 0, de manera que $\operatorname{Re} \tau \geq 0$. Sea $(u_n) \in \mathcal{D}(\tau) = \mathcal{D}(T)$ una sucesión τ -convergente a 0 : $u_n \rightarrow 0, \tau[u_n - u_m] \rightarrow 0$; y debemos probar que $\tau[u_n] \rightarrow 0$. Del resultado de E7 se obtiene

$$|\tau[u_n]| \leq |\tau[u_n, u_n - u_m]| + |\tau[u_n, u_m]|$$

$$\leq (1 + \tan \vartheta)(\operatorname{Re} \tau)[u_n]^{1/2}(\operatorname{Re} \tau)[u_n - u_m]^{1/2} + |\langle Tu_n, u_m \rangle|.$$

Para cualquier $\varepsilon > 0$ existe un entero N tal que $(\operatorname{Re} \tau)[u_n - u_m] \leq \varepsilon^2$ para $n, m \geq N$. Como $(\operatorname{Re} \tau)[u_n]$ es una sucesión acotada entonces de la desigualdad anterior concluimos que

$$|\tau[u_n]| \leq M\varepsilon + |\langle Tu_n, u_m \rangle|, \quad n, m \geq N$$

donde M es una constante. Haciendo tender m a infinito, obtenemos $|\tau[u_n]| \leq M\varepsilon$ lo cual prueba que $\tau[u_n] \rightarrow 0$. $\square$

Ejemplo 7. Según el teorema anterior todo operador simétrico T acotado inferiormente es form-cerrable además de ser cerrable.

Sin embargo el dominio de su clausura $\mathcal{D}(\tilde{T})$ no tiene que coincidir con el dominio $\mathcal{D}[\tau]$ de la clausura de la forma τ definida en (120). Proponemos al lector demostrar esto para el operador diferencial $T = -\frac{d^2}{dx^2}$ definido en $L_2(0, 1)$ con $\mathcal{D}(T) = C_0^\infty(0, 1)$. Repitiendo el razonamiento del Ejemplo 6 se puede comprobar que $\mathcal{D}(\tilde{T}) = H_0^2(0, 1)$ mientras que $\mathcal{D}[\tau] = H_0^1(0, 1)$. $\square$

E 13. A partir de una familia de formas sesquilineales se define su suma mediante

$$\tau = \tau_1 + \tau_2 + \cdots + \tau_n$$

con $\mathcal{D}(\tau) = \mathcal{D}(\tau_1) \cap \cdots \cap \mathcal{D}(\tau_n)$.

- Demuestre que si τ_i es sectorial $i = 1, \dots, n$, entonces también lo es τ .
- Si todas las formas τ_i son sectoriales y cerradas entonces τ es cerrada.

- c) Si todas las formas τ_i son sectoriales y cerrables entonces τ también es cerrable y se cumple que

$$\tilde{\tau} \subset \tilde{\tau}_1 + \cdots + \tilde{\tau}_n$$

donde la inclusión no puede ser reemplazada por igualdad ni tan siquiera cuando todas las formas τ_i son simétricas. Busque un ejemplo que confirme esta afirmación. $\square$

En muchos problemas de la Física resulta que el estudio de la dinámica y estabilidad de una magnitud o sistema físico está asociado a las propiedades espectrales de algún operador T definido en un espacio de Hilbert (o generalmente de Banach) conveniente y no necesariamente acotado. Es importante conocer qué ocurre cuando en el sistema se introduce alguna *perturbación*. En ese caso el operador asociado ya no será el mismo T sino una cierta perturbación de T que en general se expresa en la forma $T + A$, donde A es otro operador que corresponde a la perturbación. Entonces, interesa saber cómo se comportan las propiedades espectrales de $T + A$ con respecto a las del operador T .

Es bastante frecuente encontrarse con la situación en que A es de la forma εV donde V es otro operador y ε un parámetro pequeño. El estudio del comportamiento del espectro de $T + \varepsilon V$ cuando ε es suficientemente pequeño es a lo que se conoce en la Física como Teoría de Perturbaciones. También puede considerarse el caso en que $A = V(\varepsilon)$ siendo más compleja la dependencia del parámetro pequeño ε . Sin embargo, en muchos problemas importantes de perturbaciones no aparece un parámetro pequeño. Con vistas a estudiar este tipo general de perturbaciones introduciremos la siguiente definición:

D 3. (acotación relativa de operadores y formas sesquilineales)

Sean T y A operadores con el mismo espacio dominio H (pero no necesariamente el mismo espacio imagen) tales que $\mathcal{D}(T) \subset \mathcal{D}(A)$ y

$$\|Au\|_1 \leq a\|u\|_0 + b\|Tu\|_2, \quad u \in \mathcal{D}(T) \quad (121)$$

donde a y b son constantes no negativas. Entonces diremos que A es relativamente acotado con respecto a T , o simplemente T -acotado. El ínfimo b_0 de todas las posibles constantes b en (121) es llamada la cota relativa de A con respecto a T o simplemente la T -cota de A .

Sea τ una forma sectorial en H y τ' otra forma sesquilineal en H no necesariamente sectorial. Diremos que τ' es relativamente acotada con respecto a τ , o simplemente τ -acotada, si $\mathcal{D}(\tau') \supset \mathcal{D}(\tau)$ y

$$|\tau'[u]| \leq a\|u\|^2 + b|\tau[u]|, \quad u \in \mathcal{D}(\tau) \quad (122)$$

donde a y b son constantes no negativas. El ínfimo b_0 de todos los posibles valores b en (122) se llama τ -cota de τ' .

Nótese que si b es seleccionado suficientemente cerca de b_0 en (121) o (122), entonces en general la constante a puede ser muy grande y por ello no siempre es posible tomar $b = b_0$.

Obviamente τ -acotación (T -acotación) es equivalente a $(\tau + \alpha)$ -acotación ($T + \alpha$ -acotación) para cualquier escalar α . También τ -acotación es equivalente a $\text{Re}\tau$ -acotación, ya que (122) es equivalente a

$$|\tau'[u]| \leq a'\|u\|^2 + b'\text{Re}\tau[u], \quad (123)$$

$u \in \mathcal{D}(\tau) = \mathcal{D}(\text{Re}\tau)$, con constantes a' y b' que no coinciden necesariamente con las constantes a y b en (122). Si ambas formas τ y τ' son sectoriales y cerrables, entonces (122) se extiende a todo $u \in \mathcal{D}(\tilde{\tau})$ con las mismas constantes a y b y con τ y τ' reemplazadas por sus respectivas clausuras $\tilde{\tau}$ y $\tilde{\tau}'$. (¡demuestre en calidad de ejercicio todas las afirmaciones arriba señaladas!)

La demostración de los teoremas siguientes puede ser hallada en la monografía [13] de T. Kato.

T 9. Sea τ una forma sectorial y supongamos que τ' es τ -acotada con $b' < 1$ en (123). Entonces $\tau + \tau'$ es sectorial. Además $\tau + \tau'$ es cerrada

(resp. cerrable) si y solo si τ es cerrada (resp. cerrable). En el caso en que $\tau + \tau'$ es cerrable se cumple que $\mathcal{D}(\tau + \tau') = \mathcal{D}(\bar{\tau})$.

T 10. Sea T un operador autoadjunto, acotado inferiormente y sea A simétrico y T -acotado con T -cota b . Entonces la forma $\langle Au, u \rangle$ es relativamente acotada con respecto a la forma $\langle Tu, u \rangle$ con cota relativa $\leq b$, y lo mismo es cierto para sus clausuras.

Nótese que hemos definido las nociones de acotación relativa para operadores y formas sesquilineales por separado. Sin embargo, ambas nociones pueden ser aplicadas a operadores sectoriales S y S' . Por ello puede ocurrir que S' sea acotado con respecto a S o que la forma sectorial $\langle S'u, u \rangle$ sea relativamente acotada con respecto a la forma $\langle Su, u \rangle$. El Teorema 9 nos dice que bajo determinadas condiciones de cerrabilidad y pequeñez de la cota relativa, la clausura de la forma $\langle (S + S')u, u \rangle$ tiene el mismo dominio que la clausura de la forma $\langle Su, u \rangle$.

En general no está clara la relación entre los dos tipos de cota relativa asociada a un operador. El Teorema 10 nos muestra que la acotación relativa en sentido de formas es más débil que la acotación relativa en sentido de operadores. Ahora nos encontramos preparados para enunciar los teoremas de representación de formas sesquilineales no acotados a que hicimos referencia después del enunciado del teorema de representación de Friedrichs (T7).

T 11. (primer teorema de representación) Sea $\tau[u, v]$ una forma sesquilineal sectorial, densamente definida y cerrada en el espacio de Hilbert H . Entonces existe un operador sectorial cerrado T en H con las propiedades siguientes:

- i) $\sigma(T) \subset \overline{\tau(T)}$;
- ii) $\mathcal{D}(T) \subset \mathcal{D}(\tau)$ y $\tau[u, v] = \langle Tu, v \rangle$, para todo $u \in \mathcal{D}(T)$, $v \in \mathcal{D}(\tau)$;
- iii) $\mathcal{D}(T)$ es un núcleo para τ ;

- iv) si $u \in \mathcal{D}(\tau)$, $w \in H$ y $\tau[u.v] = \langle w, v \rangle$ para todo v perteneciente a un núcleo de τ , entonces $u \in \mathcal{D}(T)$ y $Tu = w$. El operador sectorial T está unívocamente determinado por la condición i).

Demostración. Sin pérdida de generalidad podemos asumir que τ tiene vértice cero y por lo tanto, $\operatorname{Re} \tau \geq 0$. Consideremos el espacio de Hilbert H_τ definido sobre el conjunto $\mathcal{D}(\tau)$ provisto del producto escalar

$$\langle u, v \rangle_\tau = \langle u, v \rangle + (\operatorname{Re} \tau)[u, v].$$

Del resultado del Ejercicio 7 se obtiene

$$\begin{aligned} |\tau[u, v]| &\leq (1 + \tan \vartheta)(\operatorname{Re} \tau)[u]^{1/2} \operatorname{Re} \tau[v]^{1/2} \\ &\leq (1 + \tan \vartheta) \|u\|_\tau \|v\|_\tau, \end{aligned}$$

de lo cual concluimos que τ es una forma sesquilineal acotada en H_τ . Obviamente la forma sesquilineal $1[u, v] = \langle u, v \rangle$ es también acotada en H_τ . Si denotamos mediante $\tau_1 := \tau + 1$ y aplicamos el teorema de Lax-Milgram (T 2.3.2), tendremos que existe un operador B acotado en H_τ tal que

$$\tau_1[u, v] = \langle Bu, v \rangle_\tau, \quad u, v \in H_\tau = \mathcal{D}(\tau).$$

Como $\|u\|_\tau^2 = (\operatorname{Re} \tau + 1)[u] = (\operatorname{Re} \tau_1)[u] = \operatorname{Re} \langle Bu, u \rangle_\tau \leq \|Bu\|_\tau \|u\|_\tau$, entonces $\|u\|_\tau \leq \|Bu\|_\tau$. Luego B tiene una inversa acotada B^{-1} con dominio cerrado en H_τ . Veamos que $\mathcal{D}(B^{-1}) = H_\tau$. Para ver esto es suficiente probar que un $u \in H_\tau$ que sea ortogonal en H_τ a $\mathcal{D}(B^{-1}) = R(B)$ debe ser nulo, lo cual es obvio porque $\|u\|_\tau^2 = \operatorname{Re} \langle Bu, u \rangle_\tau = 0$. Pero entonces $B^{-1} \in \mathcal{B}(H_\tau)$ con $\|B^{-1}\|_\tau \leq 1$. Para cualquier $u \in H$ fijo consideremos el funcional semilineal $v \mapsto l_u(v) = \langle u, v \rangle$ definido para $v \in H_\tau$. Obviamente l_u es un funcional semilineal acotado en H_τ con norma $\leq \|u\|$, ya que $|l_u(v)| \leq \|u\| \|v\| \leq \|u\| \|v\|_\tau$. Por el teorema de representación de Riesz en espacios de Hilbert, existe un único $u' \in H_\tau$ tal que $\langle u, v \rangle = l_u(v) = \langle u', v \rangle_\tau$, $\|u'\|_\tau \leq \|u\|$.

Definamos ahora un operador A mediante $Au = B^{-1}u'$. El operador A es lineal con dominio H y rango en H_τ . Considerado como operador

en H , A pertenece a $\mathcal{B}(H)$ con $\|A\| \leq 1$, ya que $\|Au\| = \|B^{-1}u'\| \leq \|B^{-1}u'\|_\tau \leq \|u'\|_\tau \leq \|u\|$. De la definición de A se sigue que

$$\langle u, v \rangle = \langle u', v \rangle_\tau = \langle BAu, v \rangle_\tau = \tau_1[Au, v] = (\tau + 1)[Au, v]. \quad (124)$$

Luego

$$\tau[Au, v] = \langle u - Au, v \rangle, \quad u \in H, v \in H_\tau = \mathcal{D}(\tau). \quad (125)$$

El operador A es inversible, ya que por (124) la condición $Au = 0$ implica que $\langle u, v \rangle = 0$ para todo $v \in \mathcal{D}(\tau)$ y $\mathcal{D}(\tau)$ es denso en H . Escribiendo $w = Au$, $u = A^{-1}w$ en (125), llegamos a que

$$\tau[w, v] = \langle (A^{-1} - I)w, v \rangle = \langle Tw, v \rangle \quad (126)$$

donde $T := A^{-1} - I$, para todo $w \in \mathcal{D}(T) = R(A) \subset \mathcal{D}(\tau)$ y $v \in \mathcal{D}(\tau)$. Esto prueba el inciso ii) del teorema. El operador T es cerrado en H por ser $A \in \mathcal{B}(H)$. T es sectorial y $r(T) \subset \Theta(\tau)$ pues $\langle Tu, u \rangle = \tau[u]$ por (126). Además del hecho que $R(T + I) = R(A^{-1}) = \mathcal{D}(A) = H$ se concluye que $-1 \in \rho(T)$, y del Teorema 3.5.6 concluimos que $\sigma(T) \subset r(T)$, lo cual prueba i).

Para demostrar el inciso iii) del teorema es suficiente probar que $\mathcal{D}(T) = R(A)$ es denso en H_τ (ver E12). Pero como B aplica H_τ sobre sí mismo de manera bicontinua, basta con probar que $BR(A) = R(BA)$ es denso en H_τ . Sea $v \in H_\tau$ ortogonal en H_τ a $R(BA)$. Entonces de (126) se tiene que $\langle u, v \rangle = 0$ para todo $u \in H$ y por tanto $v = 0$, por lo cual $R(BA)$ es denso en H_τ .

Para probar el inciso iv) del teorema resulta conveniente considerar τ^* la forma adjunta a τ . Como τ^* es también densamente definida, sectorial, con vértice cero y cerrada, podemos construir un operador sectorial T' asociado a τ^* de la misma forma que construimos T para τ . Entonces para cualesquiera $u \in \mathcal{D}(\tau^*) = \mathcal{D}(\tau)$ y $v \in \mathcal{D}(T')$, tendremos

$$\tau^*[u, v] = \langle T'u, v \rangle,$$

o lo que es lo mismo

$$\tau[u, v] = \langle u, T'v \rangle. \quad (127)$$

Si en particular en (127) tomamos $u \in \mathcal{D}(T) \subset \mathcal{D}(\tau)$ y $v \in \mathcal{D}(T') \subset \mathcal{D}(\tau)$, entonces de ii) y (127) tendremos que $\langle Tu, v \rangle = \langle u, T'v \rangle$, lo cual implica que $T' \subset T^*$. Pero como T y T' son ambos sectoriales, cerrados y cumplen que $\sigma(T) \subset \Theta(T)$, $\sigma(T') \subset \Theta(T')$, entonces se cumple que $T' = T^*$ o también $(T')^* = T$ (¡ejercicio!)

Supongamos ahora que tiene lugar la relación $\tau[u, v] = \langle w, v \rangle$ para todo v perteneciente a un núcleo de τ , donde $u \in \mathcal{D}(\tau)$ y $w \in H$. Obviamente esta relación puede ser extendida por continuidad a todo $v \in \mathcal{D}(\tau)$. Si en particular tomamos $v \in \mathcal{D}(T') \subset \mathcal{D}(\tau^*) = \mathcal{D}(\tau)$, entonces tendremos que $\langle u, T^*v \rangle = \tau[u, v] = \langle w, v \rangle$. De aquí se concluye que $u \in \mathcal{D}(T'^*) = \mathcal{D}(T)$ y $w = T'^*u = Tu$ por la definición de T'^* . Con esto queda probada la primera parte de iv). La unicidad del operador T satisfaciendo ii) es trivial y se deja al lector. $\square$

En lo sucesivo denotaremos mediante $T = T_\tau$ al operador asociado a la forma sectorial τ según el teorema anterior.

Según T 3.5.6 la condición necesaria y suficiente para que se satisfaga la condición i) de este teorema es que $\rho(T) \cap \{\mathbb{C} \setminus r(T)\} \neq \emptyset$.

E 14. Definamos una forma τ_0 a partir del operador T obtenido en el Teorema 11 de representación, mediante

$$\tau_0[u, v] = \langle Tu, v \rangle, \quad \mathcal{D}(\tau_0) = \mathcal{D}(T).$$

Demostrar que $\tau = \tilde{\tau}_0$.

E 15. Demostrar que el rango numérico $r(T)$ de T en el Teorema 11 es un subconjunto denso del rango numérico $\Theta(\tau)$ de τ .

E 16. Demuestre que si $T = T_\tau$ entonces $T^* = T_{\tau^*}$. Note que si τ es una forma densamente definida, sectorial y cerrada también lo será la

forma adjunta τ^* .

E 17. Sea τ una forma cerrada, densamente definida, simétrica y acotada inferiormente. Demuestre que el operador $T = T_\tau$ asociado con τ según el Teorema 11 es un operador autoadjunto y acotado inferiormente. Del Teorema 11 y los Ejercicios E 14 y E17 se concluye el resultado siguiente:

T 12. La correspondencia $\tau \rightarrow T = T_\tau$ es una biyección entre el conjunto de las formas sectoriales cerradas densamente definidas y el conjunto de los operadores sectoriales cerrados T para los cuales $\sigma(T) \subset \overline{r(T)}$. Más exactamente, τ es acotada si y solo si T_τ es acotado y τ es simétrica si y solo si T_τ es autoadjunto.

Este resultado muestra que las formas sectoriales cerradas densamente definidas son un instrumento conveniente para construir operadores sectoriales cerrados tales que $\sigma(T) \subset \overline{r(T)}$ (en particular para construir operadores autoadjuntos acotados inferiormente). Una importante aplicación de este método la veremos al estudiar operadores diferenciales en derivadas parciales.

Ejemplo 8. Supongamos que τ_1 y τ_2 son formas sectoriales cerradas. Según el resultado del Ejercicio 13 tenemos que $\tau = \tau_1 + \tau_2$ también es sectorial y cerrada. Si suponemos que τ es densamente definida, también lo serán τ_1 y τ_2 y, por lo tanto, existen operadores sectoriales, cerrados T , T_1 y T_2 con $\sigma(T) \subset \overline{r(T)}$, $\sigma(T_i) \subset \overline{r(T_i)}$ $i = 1, 2$, que representan a las formas τ , τ_1 y τ_2 en el sentido del Teorema 11. El operador T puede ser considerado como una suma en sentido generalizado de T_1 y T_2 :

$$T = T_1 + T_2. \quad (128)$$

Notemos, sin embargo, que la suma en sentido usual $S = T_1 + T_2$ no tiene porqué ser densamente definida ni mucho menos autoadjunta

aún siendo densamente definida. El operador T definido en (128) es una extensión autoadjunta de S y además es su única extensión autoadjunta cuando $T_1 + T_2$ es esencialmente autoadjunto.

Luego, el proceso para calcular T a partir del teorema de representación nos brinda una forma de calcular la extensión autoadjunta de un operador esencialmente autoadjunto de la forma $T_1 + T_2$ con T_1 y T_2 autoadjuntos. Proponemos al lector en calidad de ejercicio probar que si T_1 y T_2 son autoadjuntos, acotados inferiormente y $T_1 + T_2$ está densamente definido, entonces la extensión de Friedrichs T_F de $T_1 + T_2$ satisface la inclusión $T_F \subset T$ y en general T_F no tiene que coincidir con el operador T definido en (128).

Supongamos recíprocamente que se tienen dos operadores autoadjuntos T_1 y T_2 acotados inferiormente y sean τ_{10}, τ_{20} las formas simétricas definidas mediante:

$$\tau_{i0}[u, v] = \langle T_i u, v \rangle, \quad \mathcal{D}(\tau_i) = \mathcal{D}(T_i), i = 1, 2.$$

Las formas τ_{i0} son cerrables. Denotemos por $\tau_i = \tilde{\tau}_{i0}$ sus clausuras. Según el resultado del Ejercicio 14, el operador que representa a la forma τ_i es precisamente T_i . Hemos visto que resulta común que el operador suma $T_1 + T_2$ no sea autoadjunto. Sin embargo, si ocurre que la forma simétrica y acotada inferiormente $\tau = \tau_1 + \tau_2$ es densamente definida (lo cual es un requerimiento más débil que la exigencia de que sea densamente definido el operador $T_1 + T_2$), entonces existe la suma en sentido generalizado de T_1 y T_2 : $T = T_1 + T_2$ y es una extensión autoadjunta de $T_1 + T_2$. $\square$

Consideremos ahora que tenemos una forma τ cerrada, densamente definida, simétrica y acotada inferiormente. Según el resultado del Ejercicio 17 el operador T que representa a la forma τ en el sentido del Teorema 11, es un operador autoadjunto. El siguiente teorema nos brinda una descripción más completa de T en este caso; para una demostración

véase la monografía de Kato [13].

T 13. (segundo teorema de representación) Sea τ una forma cerrada, densamente definida, simétrica y positiva ($\tau \geq 0$) y sea $T = T_\tau$ el operador autoadjunto asociado a τ . Entonces se tiene que $\mathcal{D}(T^{1/2}) = \mathcal{D}(\tau)$ y

$$\tau[u, v] = \langle T^{1/2}u, T^{1/2}v \rangle, \quad u, v \in \mathcal{D}(\tau). \quad (129)$$

Además, un subconjunto $\mathcal{D}'$ de $\mathcal{D}(\tau)$ es un núcleo para τ si y solo si es un núcleo para $T^{1/2}$, en el sentido de que

$$\overline{T^{1/2}}|_{\mathcal{D}'} = T^{1/2}. \quad \square$$

Ejemplo 9. Sea A un operador simétrico, positivo y densamente definido. Si definimos la forma $\tau_0[u, v] = \langle Au, v \rangle$, $\mathcal{D}(\tau_0) = \mathcal{D}(A)$, entonces τ_0 es cerrable y su clausura $\tau = \bar{\tau}_0$ es una forma que satisface los requerimientos del Teorema 13. El operador autoadjunto asociado a la forma τ es precisamente la extensión de Friedrichs del operador $A : T = A_F$. $\square$

§ 3.8 La integral de Dunford

Hasta ahora hemos visto cómo definir una importante función de un operador cerrado. En efecto, si $\lambda \in \rho(A)$, entonces la resolvente $(A - \lambda I)^{-1}$ puede interpretarse de manera informal, como la evaluación de la función $t \rightsquigarrow (t - \lambda)^{-1}$ en el operador A . Notemos que la relación de Hilbert (65) no es más que la evaluación en A de la relación

$$\frac{1}{t - \lambda} - \frac{1}{t - \mu} = (\lambda - \mu) \frac{1}{t - \lambda} \cdot \frac{1}{t - \mu}.$$

Es deseable contar con una definición general de función de un operador cerrado que permita desarrollar un *cálculo operacional* más amplio. En este epígrafe introduciremos la noción de *integral de Dunford* de un

operador cerrado A , mediante la cual podremos construir familias de operadores acotados que podrán interpretarse como funciones de A .

Con este propósito, comencemos por denotar mediante $\mathcal{H}_C[A]$ la familia de todas las funciones complejas que son holomorfas en alguna vecindad, no necesariamente conexa, de una parte aislada C del espectro $\sigma(A)$ del operador cerrado A . En alguna ocasión C puede coincidir con todo $\sigma(A)$ en cuyo caso denotaremos esta familia de funciones por $\mathcal{H}[A]$. La vecindad a que hacemos referencia puede depender de la función $f \in \mathcal{H}_C[A]$ considerada. Sea $f \in \mathcal{H}_C[A]$ y U un subconjunto abierto de $\mathbb{C}$, tal que $C \subset U$, f es holomorfa en $\bar{U}$ y $\bar{U} \cap [\sigma(A) \setminus C] = \emptyset$. Supongamos, para simplificar que la frontera ∂U consiste de un número finito de curvas de Jordán, orientadas en sentido positivo. Denotaremos mediante $F_C[A]$ la familia de todas las funciones $f \in \mathcal{H}_C[A]$, para las cuales puede ser hallado algún abierto U con las propiedades anteriormente expuestas y tal que

$$\int_{\partial U} |f(\lambda)| \| (A - \lambda I)^{-1} \| |d\lambda| < \infty. \quad (130)$$

En el caso $C = \sigma(A)$, pondremos

$$F[A] = F_{\sigma(A)}[A]. \quad (131)$$

D 1. (integral de Dunford) Sea C una parte aislada del espectro $\sigma(A)$ del operador cerrado A y $f \in F_C[A]$. Entonces el operador lineal acotado $f[A]$ definido mediante

$$f[A] = \frac{-1}{2\pi i} \int_{\partial U} f(\lambda) (A - \lambda I)^{-1} d\lambda, \quad (132)$$

donde U es un abierto del plano con las propiedades citadas más arriba, se llama *integral de Dunford* de A con respecto a la función f .

Observación 1. Por el teorema integral de Cauchy, la definición de

$f[A]$ depende solamente de la función f y del operador A , pero no de la selección del abierto U , siempre que se satisfaga (130).

Observación 2. Notemos que si C es una parte aislada y acotada de $\sigma(A)$, entonces

$$\mathcal{H}_C[A] = F_C[A]. \quad (133)$$

En este caso las funciones constantes pertenecen a $F_C[A]$. En particular, si A es acotado, siempre se cumple (133).

El teorema siguiente nos muestra cómo desarrollar un cálculo operacional con funciones de operadores cerrados.

T 1. (N. Dunford) Sean A un operador cerrado y C una parte aislada de $\sigma(A)$. Si $f, g \in F_C\{A\}$ y $\alpha, \beta \in \mathbb{C}$, entonces:

- i) $\alpha f + \beta g \in F_C[A]$ y $\alpha f[A] + \beta g[A] = (\alpha f + \beta g)[A]$;
- ii) $f \cdot g \in F_C[A]$ y $f[A] \circ g[A] = (f \cdot g)[A]$;
- iii) $f^*(z) \in F_{C^*}[A^*]$, donde $C^* = \{\bar{\lambda} : \lambda \in C\}$, $f^*(z) = \overline{f(\bar{z})}$ y

$$(f[A])^* = f^*[A^*];$$

- iv) Si además suponemos que las funciones constantes pertenecen a $F_C[A]$, entonces para toda sucesión $(f_n) \in F_C[A]$, holomorfas en una misma vecindad U de C y tal que (f_n) converge uniformemente a f en $\overline{U}$, se tiene que $f \in F_C[A]$ y

$$|||f_n[A] - f[A]||| \rightarrow 0, \quad (n \rightarrow \infty).$$

Demostración. i) es evidente.

- ii) Sean U_1 y U_2 vecindades abiertas de $C \subset \sigma(A)$, cuyas fronteras ∂U_1 , y ∂U_2 están formadas por un número finito de curvas de Jordán y

asumamos que $\partial U_1 \subset U_2$ y que $U_2 \cup \partial U_2$ está contenido en el dominio de analiticidad de f y g . Entonces, utilizando la relación de Hilbert para la resolvente y el teorema integral de Cauchy para funciones analíticas, tendremos

$$\begin{aligned}
 f[A] \cdot g[A] &= -\frac{1}{4\pi^2} \int_{\partial U_1} f(\lambda) R_\lambda(A) d\lambda \int_{\partial U_2} g(\mu) R_\mu(A) d\mu \\
 &= -\frac{1}{4\pi^2} \int_{\partial U_1} \int_{\partial U_2} f(\lambda) g(\mu) (\lambda - \mu)^{-1} (R_\lambda(A) - R_\mu(A)) d\lambda d\mu \\
 &= -\frac{1}{2\pi i} \int_{\partial U_1} f(\lambda) R_\lambda(A) \left\{ \frac{1}{\pi i} \int_{\partial U_2} (\mu - \lambda)^{-1} g(\mu) d\mu \right\} d\lambda \\
 &\quad + \frac{1}{2\pi i} \int_{\partial U_2} g(\mu) R_\mu(A) \left\{ \frac{1}{2\pi i} \int_{\partial U_1} (\mu - \lambda)^{-1} f(\lambda) d\lambda \right\} d\mu \\
 &= \frac{1}{2\pi i} \int_{\partial U_1} f(\lambda) g(\lambda) R_\lambda(A) d\lambda \\
 &= (f \cdot g)[A].
 \end{aligned}$$

iii) Se concluye de P3.5.2 y la definición de $f[A]$.

iv) Es consecuencia inmediata de (i) y de la desigualdad evidente

$$\|f[A]\| \leq \sup_{z \in \partial U} |f(z)| \int_{\partial U} \|R_\lambda(A)\| |d\lambda|. \quad \square$$

Observación. Se puede dar una definición más general de la integral de Dunford de un operador cerrado que permite ampliar la clase de funciones $F_C[A]$ a otro espacio vectorial $\tilde{F}_C[A]$. El teorema anterior nos asegura que la integral (132) induce un homomorfismo entre el álgebra $F_C[A]$ y el álgebra $\mathcal{B}(H)$ de los operadores acotados en H . Es posible ampliar la clase $F_C[A]$ de manera que la integral (132) defina un ope-

rador lineal del espacio vectorial $\tilde{F}_C[A]$ en el espacio vectorial $C(H)$ de los operadores cerrados sobre H . $\square$

Para los operadores acotados se tienen los importantes corolarios siguientes:

C 1. (teorema de la aplicación espectral) Sea A acotado y $f \in \mathcal{H}[A]$. Entonces

$$\sigma(f[A]) = f[\sigma(A)]. \quad (134)$$

Demostración. Sea $\lambda \in \sigma(A)$ y definamos la función g mediante $g(\mu) = (f(\lambda) - f(\mu))/(\lambda - \mu)$. Por el teorema anterior,

$$f(\lambda)I - f[A] = (\lambda I - A)g[A].$$

De aquí se deduce que, si $(f(\lambda)I - f[A])$ tuviera inverso acotado B , entonces $g[A]B$ sería el inverso acotado de $(\lambda I - A)$. Luego, $\lambda \in \sigma(A)$ implica que $f(\lambda) \in \sigma(f[A])$. Inversamente, sea $\lambda \in \sigma(f[A])$ y supongamos que $\lambda \notin f[\sigma(A)]$. Entonces la función $g(\mu) = (f(\mu) - \lambda)^{-1}$ debería pertenecer a $\mathcal{H}[A]$, y por el teorema anterior, $g[A](f[A] - \lambda I) = I$, lo cual contradice la suposición de que $\lambda \in \sigma(f[A])$. $\square$

Ejemplo 1.6 (función exponencial de un operador acotado) En la teoría de las ecuaciones diferenciales juega un papel muy importante la función operacional e^{At} (A acotado en H), la cual puede ser definida de cualquiera de las dos formas siguientes:

$$e^{At} = -\frac{1}{2\pi i} \int_{\Gamma_A} e^{\lambda t} R_{\lambda}(A) d\lambda \quad (135)$$

$$e^{At} = \sum_{k=0}^{\infty} \frac{A^k t^k}{k!}. \quad (136)$$

(Demuestre la equivalencia entre las definiciones (135) y (136)).

De la propiedad multiplicativa enunciada en T1 ii) se obtiene que los operadores e^{At} forman un grupo uniparamétrico:

$$e^{At}e^{A\tau} = e^{A(t+\tau)} \quad \forall t, \tau \in (-\infty, \infty) e^{At}|_{t=0} = I.$$

Observemos, sin embargo, que en general $e^{A+B} \neq e^A e^B$ salvo en el caso cuando A y B conmutan (¡ demuéstrela!).

Obviamente la serie (136) converge uniformemente y para ella se obtiene el estimado

$$\|e^{At}\| \leq \sum_{k=0}^{\infty} \frac{\|A^k\| t^k}{k!} \leq e^{\|A\|t}. \quad (137)$$

Además la serie (136) puede ser derivada término a término con respecto a t , de donde se obtiene la fórmula

$$\frac{de^{At}}{dt} = Ae^{At}. \quad (138)$$

Esta fórmula nos dice que $u(t) = e^{At}u_0$ es la solución del problema de Cauchy

$$\frac{du}{dt} = Au, \quad u(0) = u_0 \in H. \quad (139)$$

En el Capítulo 7 nos dedicaremos a estudiar las propiedades espectrales de A que implican la acotación respecto a t de las soluciones del problema de Cauchy (139). $\square$

C 2. Sean A acotado, $f \in \mathcal{H}[A]$ y $g \in \mathcal{H}[f[A]]$. Si ponemos $h(\lambda) = g(f(\lambda))$, entonces $h \in \mathcal{H}[A]$ y $h[A] = g[f[A]]$.

Demostración. El hecho que $h \in \mathcal{H}[A]$ se concluye del Corolario 1. Sea U_1 una vecindad abierta de $\sigma(f[A])$ cuya frontera está constituida por un número finito de curvas rectificables de Jordán tales que $U_1 \cup$

∂U_1 está contenido en el dominio de analiticidad de g . Análogamente consideremos una vecindad abierta U_2 de $\sigma(A)$ cuya frontera también está formada por un número finito de curvas rectificables de Jordán tales que $U_2 \cup \partial U_2$ está contenido en el dominio de analiticidad de f y

$$f[U_2 \cup \partial U_2] \subset U_1.$$

Entonces, para todo $\lambda \in \partial U_1$, tendremos

$$R_\lambda(f[A]) = \frac{1}{2\pi i} \int_{\partial U_2} (f(\mu) - \lambda)^{-1} R_\mu(A) d\mu \quad (140)$$

ya que según T 1 ii), el operador S en el miembro derecho de (140) satisface la ecuación operacional

$$(f[A] - \lambda I)S = S(f[A] - \lambda I) = I.$$

Luego, por el teorema integral de Cauchy, se tiene que

$$\begin{aligned} g(f[A]) &= \frac{1}{2\pi i} \int_{\partial U_1} g(\lambda) R_\lambda(f[A]) d\lambda \\ &= -\frac{1}{4\pi^2} \int_{\partial U_1} \int_{\partial U_2} g(\lambda) (f(\mu) - \lambda)^{-1} R_\mu(A) d\mu d\lambda \\ &= \frac{1}{2\pi i} \int_{\partial U_2} R_\mu(A) g(f(\mu)) d\mu = h[A]. \quad \square \end{aligned}$$

Veamos una propiedad importante de la integral (132) en el caso en que C es acotado.

C 3. Sea C una parte aislada y acotada del espectro de un operador cerrado A . Entonces el operador P correspondiente a la integral de Dunford de A con respecto a la función $f(z) \equiv 1$,

$$P = \frac{1}{2\pi i} \int_{\partial U} R_\lambda(A) d\lambda, \quad (141)$$

es un operador de proyección, es decir $P^2 = P$. Si A es autoadjunto, $P = P^*$, y por lo tanto, P es un proyector ortogonal.

Demostración. La primera parte del enunciado se verifica directamente de manera análoga a T 1 ii) y por ello se deja de ejercicio al lector. Para probar la segunda parte notemos primeramente que por ser A autoadjunto, C es un subconjunto compacto de $\mathbb{R}$. Tomando el abierto $U \supset C$, simétrico con respecto al eje real y aplicando T 1 iii) se obtiene que $P = P^*$. $\square$

En los dos epígrafes siguientes veremos algunas aplicaciones importantes de la integral de Dunford.

§ 3.9 Raíz cuadrada de operadores autoadjuntos positivos y representación polar de operadores cerrados.

En este epígrafe estudiaremos dos importantes teoremas. El primer resultado ha sido ya obtenido en T 3.4.6 para el caso de operadores acotados.

T 1. (raíz cuadrada de operadores autoadjuntos positivos) Sea A autoadjunto y positivo, $\langle A(x), x \rangle \geq 0, \forall x \in \mathcal{D}(A)$. Entonces existe una única raíz cuadrada $A^{1/2}$ de A , autoadjunta y positiva, tal que $(A^{1/2})^2 = A$ y además satisface las propiedades siguientes:

- i) $\sigma(A^{1/2}) = \sqrt{\sigma(A)} = \{\lambda^{1/2} : \lambda \in \sigma(A)\}$;
- ii) $A^{1/2}$ conmuta con cada operador acotado que conmute con A ;
- iii) $\mathcal{D}(A)$ es un dominio esencial de $A^{1/2}$;
- iv) $\text{Ker } A = \text{Ker } A^{1/2}$.

Demostración. Haremos primeramente la demostración cuando $0 \in \rho(A)$, que en el caso de un operador autoadjunto positivo es equivalente a decir que sea *estrictamente positivo*, es decir para algún $\delta > 0$

$$\langle A(x), x \rangle \geq \delta \langle x, x \rangle, \quad \forall x \in \mathcal{D}(A).$$

En este caso, el operador $T = A^{-1}$ es también autoadjunto y acotado. Además por T 3.5.4, $\sigma(A^{-1}) = \{1/\lambda : \lambda \in \sigma(A)\} \subset \mathbb{R}_+$ y, por lo tanto de T 3.6.2 se tiene que A^{-1} es positivo. Por T 3.4.6 concluimos entonces que existe el operador $T^{1/2} = B$. Pongamos por definición $A^{1/2} \equiv B^{-1}$ y $A^{-1/2} \equiv (A^{1/2})^{-1}$; entonces $A^{-1/2} = B$.

Obviamente se cumple la siguiente cadena de igualdades:

$$(A^{1/2} A^{1/2})^{-1} = (A^{1/2})^{-1} (A^{1/2})^{-1} = (A^{-1})^{1/2} (A^{-1})^{1/2} = B \cdot B = A^{-1}$$

y, por lo tanto

$$A = (A^{1/2})^2,$$

luego $A^{1/2}$ es una raíz cuadrada de A . La autoadjunticidad de $A^{1/2}$ se deduce inmediatamente de P 3.3.2 v), por ser el inverso del operador autoadjunto B . La unicidad de $A^{1/2}$ se desprende de la unicidad de B , demostrada en T 3.4.5.

De la unicidad de B aplicando el Corolario 3.8.2 del Teorema 3.8.1 a las funciones $g(\lambda) = \sqrt{\lambda}$ y $f(\lambda) = \lambda^2$, analíticas en $\operatorname{Re} \lambda > \delta$, y al operador T , se tiene que B es precisamente el operador definido por la integral de Dunford siguiente:

$$B = \frac{1}{2\pi i} \int_{\Gamma} z^{-1/2} R_z(A) dz, \quad (142)$$

donde Γ es cualquier arco de Jordan contenido en $\rho(A)$ y que encierra a $\sigma(A)$. Los valores de $z^{-1/2}$ deben ser seleccionados de forma tal que

$z^{-1/2} > 0$, en los puntos donde Γ corta al eje real positivo. Pero entonces, por el teorema de la aplicación espectral,

$$\sigma(B) \equiv \frac{1}{\sqrt{\sigma(A)}} = \left\{ \frac{1}{\lambda} : \lambda \in \sigma(A) \right\}. \quad (143)$$

De (143) y T3.5.4 concluimos la veracidad de i) y por ende, la positividad de $A^{1/2}$. Del hecho que B conmuta con cualquier operador acotado que conmute con A^{-1} se obtiene la afirmación ii).

Para probar iii) es necesario demostrar que para todo $x \in D(A^{1/2})$ existe una sucesión (x_n) en $D(A)$, tal que $x_n \rightarrow x$ y $A^{1/2}(x_n - x) \rightarrow \Theta$. Una sucesión (x_n) con estas propiedades puede ser construida poniendo

$$x_n = \left(I + \frac{1}{n}A\right)^{-1}(x) = n(A + nI)^{-1}(x).$$

En efecto, es obvio que $x_n \in D(A)$. Pongamos $y = A^{1/2}(x)$, entonces $x = B(y)$, luego $x_n = n(A + nI)^{-1}B(y) = nB(A + nI)^{-1}(y)$, donde la conmutación de B con la resolvente de A se deduce de la expresión (142). Pero entonces

$$A^{1/2}(x_n) = n(A + nI)^{-1}y \rightarrow y = A^{1/2}(x)$$

(ejercicio), de donde se concluye la demostración de iii).

Pasemos ahora a probar que en la demostración realizada anteriormente puede ser eliminada la suposición $0 \in \rho(A)$ y que la misma es válida siempre que:

$$\langle A(x), x \rangle \geq 0, \quad \forall x \in \mathcal{D}(A). \quad (144)$$

En efecto, si se cumple (144), entonces para cualquier $\varepsilon > 0$, $A_\varepsilon = A + \varepsilon I$ es estrictamente positivo, luego puede ser construido el operador $A_\varepsilon^{1/2}$ que satisface las propiedades i) – iii). Probaremos que $\lim_{\varepsilon \rightarrow 0} A_\varepsilon^{1/2}$ existe en un cierto sentido y coincide con un operador que satisface las

condiciones exigidas en el teorema. Para ello notemos primeramente que la integral (142) ofrece una expresión más conveniente para $B_\varepsilon = A_\varepsilon^{-1/2}$, si se toma en lugar de Γ el arco obtenido al recorrer el semieje real negativo desde 0 a $-\infty$ y desde $-\infty$ hasta 0. Poniendo $z = -\lambda$ en (148) y observando que $z^{-1/2} = \mp i\lambda^{-1/2}$, obtenemos la expresión:

$$A_\varepsilon^{-1/2} = \frac{1}{\pi} \int_0^\infty \lambda^{-1/2} (A_\varepsilon + \lambda I)^{-1} d\lambda. \quad (145)$$

Aplicando (145) al elemento $A_\varepsilon(x)$ en lugar de x , y notando que $A_\varepsilon^{-1/2} A_\varepsilon(x) = A_\varepsilon^{1/2}(x)$, se tiene:

$$A_\varepsilon^{1/2}(x) = \frac{1}{\pi} \int_0^\infty \lambda^{-1/2} (A_\varepsilon + \lambda I)^{-1} A_\varepsilon(x) d\lambda, \quad (146)$$

para todo $x \in \mathcal{D}(A_\varepsilon) = \mathcal{D}(A)$. Pero como

$$\begin{aligned} (A_\varepsilon + \lambda I)^{-1} A_\varepsilon(x) &= (A_\varepsilon + \lambda I)^{-1} (A_\varepsilon + \lambda I - \lambda I)(x) \\ &= x - \lambda (A_\varepsilon + \lambda I)^{-1}(x) \\ &= x - \lambda (A + \varepsilon I + \lambda I)^{-1}(x), \end{aligned}$$

tenemos entonces

$$\begin{aligned} \frac{d}{d\varepsilon} (A_\varepsilon + \lambda I)^{-1} A_\varepsilon(x) &= \frac{d}{d\varepsilon} [x - \lambda (A + \varepsilon I + \lambda I)^{-1}(x)] \\ &= \lambda (A + \varepsilon I + \lambda I)^{-2}(x) \\ &= \lambda (A_\varepsilon + \lambda I)^{-2}(x) \\ &= -\lambda \frac{d}{d\lambda} (A_\varepsilon + \lambda I)^{-1}(x). \end{aligned} \quad (147)$$

Por lo tanto,

$$\frac{d}{d\varepsilon} A_\varepsilon^{1/2}(x) = -\frac{1}{\pi} \int_0^\infty \lambda^{1/2} \frac{d}{d\lambda} (A_\varepsilon + \lambda I)^{-1}(x) d\lambda$$

$$\begin{aligned}
&= \frac{1}{2\pi} \int_0^{\infty} \lambda^{-1/2} (A_{\varepsilon} + \lambda I)^{-1}(x) d\lambda \\
&= \frac{1}{2} A_{\varepsilon}^{-1/2}(x)
\end{aligned} \tag{148}$$

donde para obtener (148) hemos integrado por partes y utilizado la expresión (145) para $A_{\varepsilon}^{-1/2}$. Integrando (148) entre θ y ε , con $0 < \theta < \varepsilon$, obtendremos

$$A_{\varepsilon}^{1/2}(x) - A_{\theta}^{1/2}(x) = \frac{1}{2} \int_{\theta}^{\varepsilon} A_{\varepsilon}^{-1/2}(x) d\varepsilon. \tag{149}$$

Pero notemos que de la expresión (145) y el hecho que

$$|||(A_{\varepsilon} + \lambda I)^{-1}||| \leq (\lambda + \varepsilon)^{-1}$$

se tiene

$$|||A_{\varepsilon}^{-1/2}||| \leq \varepsilon^{-1/2}. \tag{150}$$

De (149) y (150) obtenemos

$$\begin{aligned}
\|A_{\varepsilon}^{1/2}(x) - A_{\theta}^{1/2}(x)\| &\leq \frac{1}{2} \int_{\theta}^{\varepsilon} \|A_{\varepsilon}^{-1/2}(x)\| d\varepsilon \\
&\leq (\varepsilon^{1/2} - \theta^{1/2}) \|x\|.
\end{aligned} \tag{151}$$

Esto prueba que $\lim A_{\varepsilon}^{1/2} = A'$ existe y que el operador $A' - A_{\varepsilon}^{1/2}$ es acotado sobre $D(A)$ con norma $\leq \varepsilon^{1/2}$. Del hecho que $D(A) = D(A_{\varepsilon})$ es un dominio esencial de $A_{\varepsilon}^{1/2}$ según iii), y por ser $S_{\varepsilon} = A' - A_{\varepsilon}^{1/2}$ acotado, se puede probar (ejercicio) que $A' = S_{\varepsilon} + A_{\varepsilon}^{1/2}$ con dominio $D(A)$ es cerrable y su clausura $\overline{A'}$ tiene el mismo dominio que $A_{\varepsilon}^{1/2}$, lo cual de hecho implica que $D(A_{\varepsilon}^{1/2})$ no depende de ε . Pero entonces podemos poner

$$\overline{A'} = S_{\varepsilon} + A_{\varepsilon}^{1/2} \tag{152}$$

donde $A_\varepsilon^{1/2}$ es autoadjunto y S_ε simétrico y acotado con

$$|||S_\varepsilon||| \leq \varepsilon^{1/2}.$$

De (152) se deduce que $\overline{A'}$ es una raíz cuadrada autoadjunta y positiva de A que satisface las propiedades i) – iii):

$$\overline{A'} = A^{1/2}. \quad (153)$$

Finalmente probemos la unicidad de la raíz cuadrada autoadjunta y positiva en el caso general cuando A satisface la propiedad (144). Para ello denotemos mediante R cualquier raíz cuadrada autoadjunta y positiva de A . Entonces $R + \varepsilon I$ es estrictamente positivo para cualquier $\varepsilon > 0$, y

$$(R + \varepsilon I)^2 = R^2 + 2\varepsilon R + \varepsilon^2 I = A + 2\varepsilon R + \varepsilon^2 I = Q_\varepsilon$$

es también autoadjunto y estrictamente positivo con dominio $D(A)$ (ejercicio). Por la unicidad de la raíz cuadrada demostrada en la primera parte del teorema para operadores estrictamente positivos, se tiene que $R + \varepsilon I$ es la única raíz cuadrada autoadjunta y positiva de Q_ε , o sea

$$R + \varepsilon I = Q_\varepsilon^{1/2}. \quad (154)$$

Sea $x \in D(A)$. Entonces $x \in D(R)$ y $(R + \varepsilon I)(x) \rightarrow R(x)$ cuando $\varepsilon \rightarrow 0$. Si probamos que

$$Q_\varepsilon^{1/2}(x) \rightarrow A^{1/2}(x) \quad (155)$$

entonces de (154) y (155) tendremos que

$$R(x) = A^{1/2}(x), \quad \forall x \in D(A). \quad (156)$$

Pero como $D(A)$ es un dominio esencial para cualquier raíz cuadrada autoadjunta positiva de A , entonces de (156) se tendrá

$$R = A^{1/2}.$$

Pasemos a probar (155). Para ello haremos uso de una expresión similar a (146) para $Q_\varepsilon^{1/2}(x)$ y $A^{1/2}(x)$ con $x \in D(A)$.

Notemos que

$$[(Q_\varepsilon + \lambda I)^{-1}(x) - (A + \lambda I)^{-1}(x)] = \begin{cases} (Q_\varepsilon + \lambda I)^{-1}[x - (Q_\varepsilon + \lambda I)(A + \lambda I)^{-1}(x)] \\ (Q_\varepsilon + \lambda I)^{-1}[-(2\varepsilon R + \varepsilon^2 I)(A + \lambda I)^{-1}(x)], \end{cases}$$

luego

$$\begin{aligned} Q_\varepsilon^{1/2}(x) - A^{1/2}(x) &= -\frac{1}{\pi} \int_0^\infty \lambda^{1/2} [(Q_\varepsilon + \lambda I)^{-1}(x) - (A + \lambda I)^{-1}(x)] d\lambda \\ &= \frac{\varepsilon}{\pi} \int_0^\infty \lambda^{1/2} (Q_\varepsilon + \lambda I)^{-1} (2R + \varepsilon I) (A + \lambda I)^{-1}(x) d\lambda. \end{aligned} \quad (158)$$

Para hallar un estimado de esta última integral es conveniente separarla en dos partes $\int_0^\varepsilon$ y $\int_\varepsilon^\infty$ y usar el integrando en la forma del término de (158) para la primera parte y en la otra forma para la segunda parte. Entonces por la conmutatividad de R con $(A + \lambda I)^{-1}$, obtenemos

$$\begin{aligned} \pi \|Q_\varepsilon^{1/2}(x) - A^{1/2}(x)\| &\leq \int_0^\varepsilon \lambda^{1/2} [\|(Q_\varepsilon + \lambda I)^{-1}\| + \|(A + \lambda I)^{-1}\|] \|x\| d\lambda \\ &\quad + \varepsilon \int_\varepsilon^\infty \lambda^{1/2} \|(Q_\varepsilon + \lambda I)^{-1}\| \|(A + \lambda I)^{-1}\| \|(2R + \varepsilon I)(x)\| d\lambda \\ &\leq \left(\int_0^\varepsilon 2\lambda^{-1/2} d\lambda \right) \|x\| + \varepsilon \left(\int_\varepsilon^\infty \lambda^{-3/2} d\lambda \right) \|(2R + \varepsilon I)(x)\| \\ &= 4\varepsilon^{1/2} \|x\| + 2\varepsilon^{1/2} \|(2R + \varepsilon I)(x)\| \rightarrow 0, \end{aligned}$$

cuando $\varepsilon \rightarrow 0$, como deseábamos probar. $\square$

Observación. De manera general se pueden definir las potencias fraccionarias A^α , $0 < \alpha < 1$, de un operador autoadjunto positivo mediante

una construcción similar a la dada en el teorema anterior para $A^{1/2}$. Si A es estrictamente positivo, entonces es suficiente reemplazar $z^{1/2}$ por $z^{-\alpha}$ en (142) para obtener $A^{-\alpha}$. La fórmula correspondiente a (145) será entonces

$$A^{-\alpha} = \frac{\sin \pi \alpha}{\pi} \int_0^{\infty} \lambda^{-\alpha} (A + \lambda I)^{-1} d\lambda. \quad (159)$$

Además las potencias fraccionarias de operadores cerrados pueden ser definidas para una clase mucho más amplia que la de los operadores autoadjuntos positivos. El lector interesado puede consultar el libro de T. Kato: *Perturbation theory for linear operators* [13].

Antes de enunciar el próximo teorema, que es una consecuencia de T 1 y T 3.4.4, daremos la definición siguiente:

D 1. (operador parcialmente isométrico) Un operador lineal W definido sobre el espacio de Hilbert H_1 y con valores en el espacio de Hilbert H_2 , se llama *parcialmente isométrico*, si existe un subespacio vectorial cerrado M en H_1 , tal que

$$\begin{aligned} \|W(x)\|_2 &= \|x\|_1, \quad \forall x \in M, \\ W(x) &= \Theta, \quad \forall x \in M^{\perp}. \end{aligned}$$

T 2. (representación polar de operadores cerrados) Sea A un operador cerrado y densamente definido en H . Entonces A puede ser representado en la forma

$$A = W|A|, \quad (160)$$

donde mediante $|A|$ se representa al operador autoadjunto positivo

$$|A| = (A^* A)^{1/2}, \quad (161)$$

y $W : H \rightarrow H$ es un operador parcialmente isométrico.

Demostración. De T 3.4.4 y T 1 se deduce la existencia del operador autoadjunto positivo $|A|$. Para cada par de elementos x, y en $D(A^*A)$ se tiene

$$\langle A(x), A(y) \rangle = \langle A^*A(x), y \rangle = \langle |A|(x), |A|(y) \rangle. \quad (162)$$

Pero de la demostración del Teorema 3.4.4 se obtiene sin dificultad que $D(A^*A)$ es un dominio esencial de A y en T 1 vimos que $D(A^*A)$ es también un dominio esencial de $|A|$. Luego

$$D(A) = D(|A|), \quad (163)$$

y la igualdad (160) se cumple para todos $x, y \in D(A)$. De la igualdad

$$\|A(x)\| = \||A|(x)\|, \quad \forall x \in D(A), \quad (164)$$

se concluye que la aplicación

$$|A|(x) \longrightarrow A(x), \quad (165)$$

define una isometría W de $R(|A|) \subset H$ sobre $R(A) \subset H$ tal que

$$A(x) = W|A|(x), \quad x \in D(A). \quad (166)$$

Por continuidad W puede ser extendida a una isometría de $\overline{R(|A|)}$ sobre $\overline{R(A)}$, a la cual denotaremos también por W . Si ponemos $W(x) = 0$ para todo $x \in R(|A|)^\perp$, entonces obtendremos un operador parcialmente isométrico W para el cual se cumple la relación (166), llamada *representación polar del operador A* . $\square$

Observación 1. La representación polar de un operador cerrado es, en cierto sentido, análoga a la representación polar de números complejos en la forma $z = \rho e^{i\vartheta}$ con $\rho = |z|$, $\vartheta = \arg z$. $\square$

Observación 2. En la demostración del Teorema 3 se ve que la representación polar no es única salvo en el caso cuando $R(A)$ y $R(|A|)$ son densos en H , lo cual ocurre si $\text{Ker } A = \text{Ker } A^* = \{\Theta\}$ (ejercicio).

Se deja al lector, en calidad de ejercicio la demostración de las propiedades siguientes:

$$\text{Ker} A = \text{Ker} |A|, \quad (167)$$

$$R(A) = R(|A^*|), \quad (168)$$

$$A^* = W^* |A^*|, \quad (169)$$

donde W es el operador parcialmente isométrico que aparece en (166).

□

§ 3.10 Separación del espectro, subespacios invariantes y reducción de operadores cerrados

Los subespacios propios $N_\lambda = \text{Ker}(A - \lambda I)$ correspondientes a los valores propios λ de un operador cerrado A , tienen una importante propiedad: $A(x) \in N_\lambda$ para todo $x \in N_\lambda$. Una noción más general que la de subespacio propio es la de subespacios invariantes:

D 1. (subespacio invariante) El subespacio vectorial M del espacio de Hilbert H , es un subespacio invariante del operador A si

$$x \in D(A) \cap M \Rightarrow A(x) \in M.$$

Se puede decir que A genera un cierto operador A_1 en el subespacio invariante M , para el cual

$$D(A_1) = D(A) \cap M, \quad A_1 \subset A.$$

Notemos que todo subespacio invariante de dimensión finita de un operador A , contiene al menos un valor propio de A como es conocido del álgebra lineal.

Si M es un subespacio invariante del operador A , entonces el complemento ortogonal $M^\perp$ puede no ser un subespacio invariante para A . Supongamos, que ambos subespacios M y $M^\perp$ son invariantes para el operador A y sean A_1 y A_2 las restricciones de A a M y $M^\perp$ respectivamente. Nos preguntamos, si en este caso el estudio de A se reduce al estudio de los operadores A_1 y A_2 . La respuesta a esta pregunta es positiva si el operador A está definido en todo H . En efecto, en este caso si $x \in H$, podemos expresar x en la forma única $x = x_1 + x_2$ con $x_1 \in M, x_2 \in M^\perp$, y obtenemos

$$A(x) = A_1(x_1) + A_2(x_2).$$

Si $D(A) \neq H$, la afirmación anterior se mantiene, con una condición auxiliar, como vemos en la siguiente proposición, cuya fácil demostración queda como un ejercicio.

P 1. Si el subespacio cerrado M de H y su complemento ortogonal $M^\perp$, son subespacios invariantes del operador A , y si la proyección sobre M transforma elementos de $D(A)$ en elementos de $D(A)$, entonces para cualquier $x \in D(A)$

$$A(x) = A_1(x_1) + A_2(x_2)$$

donde A_1 y A_2 son las restricciones de A a M y $M^\perp$ respectivamente, y x_1, x_2 son las proyecciones de x sobre M y $M^\perp$. $\square$

D2. Si el subespacio cerrado M de H satisface las condiciones de la Proposición 1, diremos que M *reduce* al operador A .

Es fácil ver que si el subespacio M reduce al operador A , entonces $M^\perp$ también reduce a A . Los subespacios triviales que reducen a cualquier operador A son el subespacio nulo y el propio espacio H . Si el operador A no admite una reducción, entonces diremos que A es *irreducible*.

T 1. Sea P el operador de proyección ortogonal sobre el subespacio

cerrado M de H . Para que el subespacio M reduzca al operador A es necesario y suficiente que P conmute con A , es decir:

- i) $P(x) \in D(A), \quad \forall x \in D(A),$
- ii) $PA(x) = AP(x), \quad \forall x \in D(A).$

Demostración. Comencemos por probar la necesidad de las condiciones i) y ii). Si el subespacio M reduce a A , entonces de $x \in D(A)$ se deduce que $P(x) \in D(A)$, lo cual prueba la condición i). Para demostrar ii), pongamos

$$x = y + z$$

donde $y = P(x)$. Como M reduce a A , entonces $A(x) = A(y) + A(z)$, donde $A(y) \in M, A(z) \in M^\perp$. Por eso

$$PA(x) = PA(y) = A(y),$$

es decir, $PA(x) = AP(x)$.

La suficiencia de las condiciones i) y ii) se demuestra muy simplemente y se deja de ejercicio al lector. $\square$

Observación. Si P es un operador de proyección ($P^2 = P$) pero no ortogonal ($P^* \neq P$), entonces la demostración anterior puede ser ligeramente adaptada para obtener el resultado siguiente:

El espacio H se descompone en la suma directa $H = M_1 + M_2$, no necesariamente ortogonal, donde

$$M_1 = P[H], \quad M_2 = (I - P)[H].$$

Si P conmuta con el operador A , entonces ambos subespacios M_1 y M_2 son invariantes para el operador A y se cumple

$$A(x) = A_1(x_1) + A_2(x_2)$$

donde $x = x_1 + x_2$, $x_1 \in M_1$, $x_2 \in M_2$ y A_1, A_2 son las restricciones de A a M_1 y M_2 respectivamente. $\square$

Diremos que el *operador de proyección ortogonal* P *reduce al operador* A , si A es reducible por el subespacio sobre el cual P proyecta. El estudio de la estructura de cualquier operador A es equivalente al estudio de los subespacios que reducen a A y de las restricciones de A a dichos subespacios, según nos muestra el teorema siguiente.

T 2. Sean M_k ($k = 1, 2, \dots, n; n \leq \infty$) subespacios cerrados y mutuamente ortogonales en H , tales que

$$H = \bigoplus_{k=1}^n M_k.$$

Sea A un operador, que es reducido por cada uno de los subespacios M_k y que supondremos cerrado si $n = \infty$. Finalmente, denotemos mediante P_k el operador de proyección ortogonal sobre M_k , y A_k la restricción de A a M_k . Entonces se cumple:

$$\text{i) } x \in D(A) \iff P_k(x) \in D(A_k) \text{ y}$$

$$\sum_{k=1}^n \|A_k P_k(x)\|^2 < \infty,$$

$$\text{ii) } A(x) = \sum_{k=1}^n A_k P_k(x), \forall x \in D(A),$$

$$\text{iii) } \sigma(A) = \bigcup_{k=1}^n \sigma(A_k).$$

Demostración. i) Sea $x \in D(A)$. Como M_k reduce a A , entonces $P_k(x) \in D(A)$ y, por lo tanto $P_k(x) \in D(A) \cap M_k = D(A_k)$. Además,

$P_k A(x) = A P_k(x)$. Luego

$$A(x) = \sum_{k=1}^n P_k A(x) = \sum_{k=1}^n A P_k(x) = \sum_{k=1}^n A_k P_k(x).$$

De aquí, en el caso $n = \infty$, se deduce la convergencia de la serie

$$\sum_{k=1}^{\infty} \|A_k P_k(x)\|^2.$$

Probemos simultáneamente la implicación en sentido inverso y ii). Supongamos primeramente que $n < \infty$. En este caso la pertenencia de x a $D(A)$ y la veracidad de ii) son consecuencia de que $D(A)$ es un subespacio vectorial. Si $n = \infty$, todas las sumas parciales

$$\sum_{k=1}^r P_k(x), (r = 1, 2, \dots),$$

pertenecen a $D(A)$. Pero de la convergencia de la serie $\sum_{k=1}^{\infty} \|A_k P_k(x)\|^2$ se tiene la convergencia de la serie

$$\sum_{k=1}^{\infty} A_k P_k(x) = \lim_{r \rightarrow \infty} A \left(\sum_{k=1}^r P_k(x) \right),$$

de donde, en virtud de ser A cerrado se concluye que

$$x \in D(A) \text{ y } A(x) = \sum_{k=1}^{\infty} A_k P_k(x).$$

En lugar de iii) probaremos el enunciado equivalente

$$\rho(A) = \bigcap_{k=1}^n \rho(A_k)$$

donde $\rho(A_k)$ denota el conjunto resolvente de A_k , considerado como operador definido sobre el subespacio M_k . De i) y ii) se deduce que la

solución de la ecuación $A(x) - \lambda x = y$ (con $\lambda \in \mathbb{C}, x \in D(A), y \in H$) existe si y solo si para cada $k = 1, 2, \dots, n$, existe la solución de la ecuación $A_k(x_k) - \lambda x_k = y_k, (x_k \in D(A_k), y_k \in M_k)$. En este caso se tiene $x_k = P_k(x)$ y, por lo tanto,

$$x = \sum_{k=1}^n x_k.$$

Pero entonces A es inyectivo si y solo si cada A_k es inyectivo y tendremos

$$(A - \lambda I)^{-1} = \sum_{k=1}^n (A_k - \lambda I_k)^{-1} P_k, \quad (170)$$

donde I_k denota la identidad en M_k . De (170) se deduce

$$\|(A - \lambda I)^{-1}(x)\|^2 = \sum_{k=1}^n \|(A_k - \lambda I_k)^{-1} P_k(x)\|^2 \quad (171)$$

para cada $x \in R(A - \lambda I)$. Si $\lambda \in \rho(A)$, entonces de (171) se tiene

$$\sum_{k=1}^n \|(A_k - \lambda I_k)^{-1} P_k(x)\|^2 \leq M \|x\|^2 \quad (172)$$

para todo $x \in H$. Tomando $x = x_k \in M_k$ en (172), llegamos a que

$$\|(A_k - \lambda I_k)^{-1}(x_k)\|^2 \leq M \|x_k\|^2$$

y, por lo tanto, $\lambda \in \rho(A_k)$ para todo $k = 1, 2, \dots, n$.

Supongamos ahora que $\lambda \in \bigcap_{k=1}^n \rho(A_k)$ y consideremos primeramente que n es finito. Entonces, para todo $x_k \in M_k$

$$\|(A_k - \lambda I_k)^{-1}(x_k)\|^2 \leq M_k \|x_k\|^2. \quad (173)$$

Tomando $M = \max M_k$, de (171) se deduce que

$$\|(A - \lambda I)^{-1}(x)\|^2 \leq M \|x\|^2$$

luego $\lambda \in \rho(A)$.

Analicemos ahora el caso $n = \infty$. Por ser A cerrado, $\lambda \in \rho(A)$ si y solo si la aplicación $A - \lambda I$ es una biyección de $D(A)$ sobre H . Pero del comentario hecho al inicio de la demostración de iii) se concluye que la biyectividad de $A_k - \lambda I_k : D(A_k) \rightarrow M_k$ implica la biyectividad de $A - \lambda I$, de donde se obtiene la inclusión para el caso $n = \infty$

$$\bigcap_{k=1}^n \rho(A_k) \subset \rho(A). \quad \square$$

Observación. Notemos que en las condiciones del Teorema 2 puede ocurrir que

$$\sigma(A_k) \cap \sigma(A_j) \neq \emptyset, \quad j \neq k.$$

Para ello es suficiente con que los subespacios M_k y M_j sean isomorfos y las restricciones A_k y A_j de A , tengan una estructura análoga. $\square$

A continuación veremos un teorema que nos muestra cómo puede ser reducido un operador cerrado mediante subespacios ortogonales, de manera tal que los espectros de las restricciones a diferentes subespacios sean disjuntos.

T 3. Sea σ_1 una parte acotada del espectro del operador cerrado A . Supondremos que σ_1 está aislada del resto del espectro de A , de tal forma que puede ser hallada una curva simple cerrada y rectificable Γ , que encierra a σ_1 separándolo de $\sigma_2 = \sigma(A) \setminus \sigma_1$. Entonces se puede encontrar una descomposición en suma directa del espacio $H = M_1 + M_2$, donde M_1 y M_2 son subespacios vectoriales cerrados que reducen a A en el sentido de la observación al Teorema 1, y para las restricciones A_1 y A_2 de A a M_1 y M_2 respectivamente, se cumple

$$\sigma(A_1) = \sigma_1, \quad \sigma(A_2) = \sigma_2.$$

Demostración. Pongamos

$$P = -\frac{1}{2\pi i} \int_{\Gamma} R_{\lambda}(A) d\lambda.$$

Por el Corolario 3.7.3 al Teorema 3.7.1, tenemos que P es un operador de proyección. No es difícil probar que P conmuta con $R_{\lambda}(A)$ y según P 3.5.1, P conmuta con A . Entonces de la observación hecha al Teorema 1 se deduce la existencia de una descomposición del tipo $H = M_1 + M_2$ con las propiedades anteriormente enunciadas.

Queda por probar que $\sigma(A_1) = \sigma_1$ y $\sigma(A_2) = \sigma_2$. Para ello notemos primeramente que las restricciones de $R_{\lambda}(A)$ a M_1 y M_2 coinciden con $R_{\lambda}(A_1)$ y $R_{\lambda}(A_2)$ respectivamente. Esto prueba que $\rho(A_1)$ y $\rho(A_2)$ contienen a $\rho(A)$. Veremos además que $\rho(A_1)$ también contiene a σ_2 . En efecto, se tiene que

$$R_{\lambda}(A_1)(x) = R_{\lambda}(A)(x) = R_{\lambda}(A)P(x),$$

para todo $x \in M_1$ y $\lambda \in \rho(A)$. Pero para cualquier $\lambda \in \rho(A)$ tal que $\lambda \notin \Gamma$, tenemos

$$\begin{aligned} R_{\lambda}(A)P &= -\frac{1}{2\pi i} \int_{\Gamma} R_{\lambda}(A)R_z(A)dz \\ &= -\frac{1}{2\pi i} \int_{\Gamma} (R_{\lambda}(A) - R_z(A)) \frac{dz}{\lambda - z}, \end{aligned} \quad (174)$$

donde (174) se obtiene de la expresión de P y la relación de Hilbert para la resolvente. Si λ está fuera de la región encerrada por Γ , de (174) se obtiene

$$R_{\lambda}(A)P = \frac{1}{2\pi i} \int_{\Gamma} R_z(A) \frac{dz}{\lambda - z}. \quad (175)$$

Como el miembro derecho de (175) es una función analítica fuera de la región encerrada por Γ , entonces $R_{\lambda}(A)P$ y, por lo tanto, también

$R_\lambda(A_1)$, tienen una prolongación analítica a esa región. Del hecho que $\rho(A)$ es el dominio natural de analiticidad de $R_\lambda(A_1)$, se concluye que la prolongación analítica a que hacíamos referencia anteriormente es precisamente la resolvente de A_1 . Pero entonces $\rho(A_1)$ contiene al exterior de Γ y, por lo tanto $\sigma(A_1) \subset \sigma_1$.

De manera similar, de (174) se obtiene que

$$R_\lambda(A)P = R_\lambda(A) + \frac{1}{2\pi i} \int_{\Gamma} R_z(A) \frac{dz}{\lambda - z}, \quad (176)$$

si λ pertenece a la región encerrada por Γ . Esto prueba que $R_\lambda(A)(I-P)$ tiene una prolongación analítica al interior de Γ . Al igual que antes llegamos a la conclusión de que $\sigma(A_2) \subset \sigma_2$.

Por otra parte, un punto $\lambda \in \sigma(A)$ no puede pertenecer simultáneamente a $\rho(A_1)$ y $\rho(A_2)$. En efecto, si $\lambda \in \rho(A_1) \cap \rho(A_2)$, entonces $R_\lambda(A_1)P + R_\lambda(A_2)(I-P)$ sería el inverso de $A - \lambda I$, de donde se tendría que $\lambda \in \rho(A)$. En fin, concluimos que

$$\sigma(A_1) = \sigma_1 \text{ y } \sigma(A_2) = \sigma_2. \quad \square$$

Observación. De la observación hecha al Teorema 3.5.3 podemos concluir que si σ_1 es acotado, entonces el operador A_1 obtenido en el teorema anterior, es también acotado por ser

$$\sigma(A_1) = \sigma_1. \quad \square$$

Observación. Se tienen algunos resultados acerca de la “continuidad” de la descomposición de un operador cerrado con respecto a partes separadas de su espectro. Para las demostraciones puede consultarse el § del Capítulo 4 de la monografía [13].

La descomposición $H = M_1 + M_2$ del Teorema 3, donde $M_i = M_i(A)$, $i = 1, 2$, es continua con respecto al operador A en el sentido

siguiente: si (B_n) es una sucesión de operadores acotados con $\|B_n\| \rightarrow 0$, entonces $\|P(A + B_n) - P(A)\| \rightarrow 0$ donde $P(A + B_n)$ y $P(A)$ son los proyectores correspondientes a los subespacios $M_1(A + B_n)$ y $M_1(A)$, descritos en el Teorema 3.

Por otra parte, si Γ es un subconjunto compacto del conjunto resolvente $\rho(A)$ del operador cerrado A , entonces existe $\varepsilon > 0$ tal que $\Gamma \subset \rho(T)$ para todo $T = A + B$ con B acotado y $\|B\| < \varepsilon$. Este resultado nos indica que el espectro de un operador cerrado no puede “crecer bruscamente” ante “pequeñas perturbaciones”. Esta propiedad es conocida como semicontinuidad superior del espectro. Sin embargo, se pueden dar ejemplos en los que aún ante muy pequeñas perturbaciones el espectro puede “reducirse bruscamente”. $\square$

§ 3.11 Valores propios aislados en el espectro

Supongamos que el espectro $\sigma(A)$ de un operador cerrado A tiene un punto aislado λ . Obviamente, $\sigma(A)$ se divide en dos partes aisladas entre sí, si tomamos $\sigma_1 = \{\lambda\}$ y $\sigma_2 = \sigma(A) \setminus \{\lambda\}$.

Tomemos como Γ cualquier curva cerrada que encierre a λ y no encierre a ningún otro punto de $\sigma(A)$. Entonces para el operador A_1 , descrito en el Teorema 3.10.3, y correspondiente a $\sigma_1 = \{\lambda\}$, su espectro está constituido únicamente por el punto λ . Si denotamos por I_i , la identidad en el subespacio invariante M_i de A ; $i = 1, 2$, entonces (72):

$$\text{spr}(A_1 - \lambda I_1) = 0.$$

Por el T 3.5.2, se tiene que $R_z(A_1) = (A_1 - zI_1)^{-1}$ admite un desarrollo en serie del tipo

$$R_z(A_1) = - \sum_{n=0}^{\infty} (z - \lambda)^{-n-1} (A_1 - \lambda I_1)^n \quad (177)$$

que converge excepto para $z = \lambda$.

Notemos que (177) es equivalente a

$$R_z(A_1) = R_z(A)P_\lambda = -\frac{P_\lambda}{z-\lambda} - \sum_{n=1}^{\infty} \frac{D^n}{(z-\lambda)^{n+1}} \quad (178)$$

donde

$$P_\lambda = -\frac{1}{2\pi i} \int_{\Gamma} R_\mu(A) d\mu, \quad (179)$$

$$\begin{aligned} D^n &= [(A - \lambda I)P_\lambda]^n = (A - \lambda I)^n P_\lambda \\ &= -\frac{1}{2\pi i} \int_{\Gamma} (\mu - \lambda)^n R_\mu(A) d\mu. \end{aligned} \quad (180)$$

Por otra parte, $R_z(A_2)$ es holomorfa en $z = \lambda$ y admite un desarrollo de Taylor en un entorno de este punto. Entonces nos queda

$$R_z(A_2) = R_z(A)(I - P_\lambda) = \sum_{n=0}^{\infty} (z - \lambda)^n S^{n+1}, \quad (181)$$

donde

$$S = R_\lambda(A_2)(I - P_\lambda) = \lim_{z \rightarrow \lambda} R_z(A)(I - P_\lambda). \quad (182)$$

No es difícil verificar que

$$S^n = \frac{1}{2\pi i} \int_{\Gamma} R_\mu(A) \frac{d\mu}{(\mu - \lambda)^n}. \quad (183)$$

Notemos que $R_\lambda(A)$ no existe y por eso a $R_\lambda(A_2)$ se le llama *resolvente reducida* de A para el valor propio λ .

Uniendo (178) y (181) obtenemos el teorema siguiente:

T 1. Sean A un operador cerrado y λ un punto aislado en $\sigma(A)$. Entonces la resolvente $R_z(A)$ admite un desarrollo en serie en un entorno

del punto $z = \lambda$, del tipo

$$R_z(A) = -\frac{P_\lambda}{z - \lambda} - \sum_{n=1}^{\infty} \frac{D^n}{(z - \lambda)^{n+1}} + \sum_{n=0}^{\infty} (z - \lambda)^n S^{n+1}, \quad (184)$$

donde los operadores acotados D y S están definidos en (180) y (182) respectivamente. Además se cumplen las relaciones

$$D = DP_\lambda = P_\lambda D, \quad (185)$$

$$(A - \lambda I)S = I - P_\lambda, \quad (186)$$

$$SP_\lambda = P_\lambda S = 0, \quad (187)$$

$$SA \subset AS, \quad (188)$$

siendo AS un operador acotado.

Demostración. Se deja de ejercicio al lector. $\square$

Observación. El desarrollo de Laurent (184) es similar al que se obtiene para operadores en espacios de dimensión finita, con la única diferencia de que en el caso infinito dimensional, la parte principal de (184) (la correspondiente a las potencias negativas de $(z - \lambda)$) puede ser una serie infinita. $\square$

Observación. Notemos, que si $P_\lambda[H] = M_1$ es dimensión finita, entonces λ es un valor propio de A . En efecto, como $\lambda \in \sigma(A_1)$, entonces λ sería un valor propio de A_1 y, por lo tanto, un valor propio de A . Sin embargo, si $\dim P_\lambda[H] = \infty$, puede ocurrir que λ no sea un valor propio de A . Sugerimos al lector buscar un ejemplo de esta situación. $\square$

Introduzcamos la definición siguiente que será muy importante en lo sucesivo.

D 1. (multiplicidad algebraica de un valor propio, subespacio radical) Sean A un operador cerrado y λ un valor propio aislado de

A. El subespacio $P_\lambda[H]$ al cual denotaremos mediante $\mathcal{L}_\lambda$, se llama *subespacio radical* correspondiente al valor propio λ . Los elementos de $\mathcal{L}_\lambda$ se llaman *vectores radicales* correspondientes al valor propio λ . El número $\dim \mathcal{L}_\lambda$ se llama *multiplicidad algebraica* del valor propio λ y se denota por $a(\lambda)$. El valor propio aislado λ se llama *normal* si su multiplicidad algebraica es finita.

Pasemos a estudiar la relación entre los subespacios radicales y el desarrollo analítico (184) obtenido para la resolvente. Comencemos por el lema siguiente:

L 1. Sean A un operador cerrado y λ un valor propio aislado en $\sigma(A)$. Entonces,

$$\text{Ker}(A - \lambda I)^n = \text{Ker} D^n,$$

para todo $n \in \mathbb{N}$, donde D es el operador definido en (180).

Demostración. Consideremos la descomposición de H en la suma directa

$$H = \mathcal{L}_\lambda + (I - P_\lambda)[H] \quad (189)$$

donde los subespacios coordenados son invariantes para el operador A . Con respecto a esta descomposición, todo $x \in H$ puede ser escrito de manera única en la forma $x = x_1 + x_2$.

Por ser D la restricción de $(A - \lambda I)$ a $\mathcal{L}_\lambda$ es obvio que $\text{Ker} D^n \subset \text{Ker}(A - \lambda I)^n$. La inclusión en sentido contrario se deduce del hecho que $(A - \lambda I)^n(I - P_\lambda)$ es un operador inversible sobre el subespacio $(I - P_\lambda)[H]$. Luego

$$(A - \lambda I)^n(x_2) = \Theta \Rightarrow x_2 = \Theta. \quad \square$$

T 2. Sean A y λ como en el Lema 1. Se cumplen las propiedades siguientes:

i) Si λ es un valor propio normal, entonces

$$\begin{aligned}\mathcal{L}_\lambda &= \{x \in H : D^n(x) = 0, \text{ para algún } n \in \mathbb{N}\} \\ &= \{x \in H : (A - \lambda I)^n(x) = 0, \text{ para algún } n \in \mathbb{N}\} \\ &= \text{Ker}(A - \lambda I)^{a(\lambda)}.\end{aligned}\tag{190}$$

ii) El valor propio λ es normal si y solo si, la parte principal del desarrollo (184) es finita y además, $\dim \text{Ker}(A - \lambda I)^n < \infty$ para todo $n \in \mathbb{N}$,

iii) Si λ es un valor propio normal, entonces:

$$\begin{aligned}a(\lambda) &= \inf\{n \in \mathbb{N} : D^n \equiv 0\} = \\ &= \text{multiplicidad de } \lambda \text{ como polo de la resolvente } R_z(A).\end{aligned}$$

iv) Si λ es un valor propio normal, entonces

$$\text{Ker}(A - \lambda I) \subset \mathcal{L}_\lambda \tag{191}$$

y la igualdad en (191) ocurre si A es autoadjunto.

Demostración. i) Consideremos la descomposición (189) de H y sean A_1 y A_2 las restricciones de A a $\mathcal{L}_\lambda$ y $(I - P_\lambda)[H]$ respectivamente. Para A_1 se tiene que $\sigma(A_1) = \{\lambda\}$. Luego, si consideramos que $\mathcal{L}_\lambda$ es de dimensión finita, entonces de la teoría de operadores lineales en espacios vectoriales de dimensión finita, concluimos que

$$D^{a(\lambda)}[\mathcal{L}_\lambda] = (A_1 - \lambda I_1)^{a(\lambda)}[\mathcal{L}_\lambda] = \{\Theta\} \tag{192}$$

y además

$$D^r[\mathcal{L}_\lambda] = (A_1 - \lambda I_1)^r[\mathcal{L}_\lambda] \neq 0$$

si $r < a(\lambda)$. De (191) y el Lema 1 concluimos que

$$\mathcal{L}_\lambda = \text{Ker}(A - \lambda I)^{a(\lambda)}.$$

La primera igualdad en (190) se deduce inmediatamente del hecho que la sucesión $\text{Ker} D^k$ es una sucesión creciente de subespacios de $\mathcal{L}_\lambda$. La segunda igualdad en (190) es consecuencia inmediata del Lema 1.

Del hecho que $\text{Ker} D^k$ es una sucesión creciente de subespacios se obtiene la necesidad de la Proposición ii). Para probar la suficiencia basta notar que si tomamos n , tal que $D^n \equiv 0$, entonces $\text{Ker} D^n = \mathcal{L}_\lambda$ y como $\dim \text{Ker} D^n < \infty$, tendremos que $a(\lambda) < \infty$.

La primera igualdad en iii) se obtiene de ii) teniendo en cuenta (192) y el hecho que $\text{Ker} D^n$ es una sucesión creciente. La segunda igualdad de iii) es consecuencia directa de la primera y del desarrollo en serie de Laurent (184) para la resolvente.

La inclusión (191) es una consecuencia trivial de (190) lo cual prueba la primera parte de iv). Supongamos que A es autoadjunto y probemos que se tiene la igualdad en (191). En efecto, si A es autoadjunto, entonces $\lambda \in \mathbb{R}$ y P es autoadjunto. Por este motivo D será un operador autoadjunto en el espacio de dimensión finita $\mathcal{L}_\lambda$ y cuyo único valor propio es el cero. Del algebra lineal concluimos entonces que $D \equiv 0$. Pero por iii) tendremos que $a(\lambda) = 1$ y por i), $\mathcal{L}_\lambda = \text{Ker}(A - \lambda I)$. $\square$

Observación. A continuación pasaremos a demostrar la propiedad enunciada en la observación al Teorema 3.6.5: todo punto aislado λ del espectro de un operador autoadjunto es un valor propio para el cual $R(A - \lambda I)$ es cerrado.

En efecto, sean M_1 y M_2 los subespacios definidos en el Teorema 3.10.3 y correspondientes a las partes $\sigma_1 = \{\lambda\}$ y $\sigma_2 = \sigma(A) \setminus \{\lambda\}$ del espectro del operador A . Obviamente la suma directa $H = M_1 \oplus M_2$ es ortogonal ya que por ser A autoadjunto el proyector P definido en T 3.10.3 es un proyector ortogonal. Pero para un operador autoadjunto se tiene que $\mathcal{L}_\lambda = M_1 = \text{Ker}(A - \lambda I)$ y, por lo tanto $(A - \lambda I)P \equiv 0$. Por otra parte $(A - \lambda I)(I - P)$ tiene un inverso acotado en M_2 . Entonces

se cumple que

$$(A - \lambda I)[H] = (A - \lambda I)(I - P)[M_2] = M_2$$

y $AP = \lambda P$, de donde se concluye que $R(A - \lambda I) = M_2$ es cerrado y que λ es un valor propio cuyo subespacio propio coincide con M_1 . $\square$

Notemos que los resultados obtenidos en este epígrafe pueden ser extendidos al caso en que se consideren varios puntos aislados del espectro $\lambda_1, \dots, \lambda_k$. En este caso obtendremos el desarrollo

$$R_z(A) = - \sum_{j=1}^k \left[\frac{P_j}{z - \lambda_j} + \sum_{n=1}^{\infty} \frac{D_j^n}{(z - \lambda_j)^{n+1}} \right] + R_0(z) \quad (193)$$

donde los P_j son proyectores y los D_j son operadores con radio espectral nulo definidos mediante

$$(A - \lambda_j I)P_j = D_j \quad (194)$$

y tales que

$$P_j D_j = D_j P_j = D_j; \quad (195)$$

La función operacional $R_0(z)$ es holomorfa en λ_j ; $j = 1, \dots, k$ y

$$R_0(z) = P_0 R_z(A) = R_z(A) P_0, \quad (196)$$

$$P_0 = I - (P_1 + \dots + P_k). \quad (197)$$

De nuevo tenemos que λ_k es un valor propio de A si $\mathcal{L}_k = P_k[H]$ es de dimensión finita. Además se cumple una especie de *representación espectral* del operador A en un sentido restringido

$$AP = \sum_{j=1}^k (\lambda_j P_j + D_j), \quad P = P_1 + \dots + P_k. \quad (198)$$

La representación (198) no es tan completa como en el caso de dimensión finita, donde corresponde a la llamada *descomposición de Jordan* de una matriz, ya que en dimensión infinita el espectro no tiene que estar constituido únicamente por puntos aislados. Sin embargo, la representación (198) ofrece una descripción bastante completa de la restricción de A a $P[H]$, cuando $\lambda_1, \dots, \lambda_k$ son un número finito de valores propios aislados de A con multiplicidad algebraica finita.

No es difícil concluir de (198) que la correspondiente expresión para la resolvente del operador adjunto A^* en un entorno de los puntos aislados de su espectro $\bar{\lambda}_1, \dots, \bar{\lambda}_k$, es:

$$R_z(A^*) = - \sum_{j=1}^k \left[\frac{P_j^*}{z - \bar{\lambda}_j} + \sum_{n=1}^{\infty} \frac{D_j^{*n}}{(z - \bar{\lambda}_j)^{n+1}} \right] + R_0^*(z) \quad (199)$$

donde los operadores P_j^* siguen siendo proyectores y $R_0^*(z) = R_0(\bar{z})^*$ es holomorfa en $z = \bar{\lambda}_j, j = 1, \dots, k$. Si $\mathcal{L}_\lambda = P_j[H]$ es de dimensión finita, lo mismo será cierto para $\mathcal{L}_\lambda^* = P_j^*[H]$ y además

$$\dim \mathcal{L}_j = \dim \mathcal{L}_j^*. \quad (200)$$

Luego, si λ_j es un valor propio normal de A con multiplicidad algebraica m , entonces $\bar{\lambda}_j$ será un valor propio normal de A^* con multiplicidad algebraica m . Los valores propios normales de un operador cerrado tienen propiedades muy similares a los valores propios de operadores en espacios de dimensión finita. El lector puede verificar sin dificultad que para los valores propios normales también se cumple que *la multiplicidad geométrica de λ respecto a A coincide con la multiplicidad geométrica de $\bar{\lambda}$ respecto a A^** (ejercicio).

Un estudio detallado de la teoría de perturbación para valores propios normales puede ser encontrado en la monografía [14].

§ 3.12 Operadores J -autoadjuntos y espacios con métrica indefinida

En este epígrafe introduciremos las nociones primarias con respecto a una importante clase de operadores que aparecen con frecuencia en el estudio de la estabilidad de sistemas mecánicos. La mayor parte de los resultados aquí expuestos pueden ser hallados en el trabajo de M.G. Krein: *Introducción a la geometría de J -espacios indefinidos y a la teoría de operadores en estos espacios* [15].

D 1. (producto escalar indefinido y métrica indefinida) Un funcional bilineal $[\ , \]$ en el espacio de Hilbert $(H, \langle \ , \ \rangle)$ se llama *producto escalar indefinido*, si es continuo con respecto a la métrica determinada por $\langle \ , \ \rangle$ y además $[x, x] \in \mathbb{R}$ para todo $x \in H$, pudiendo alcanzar valores con signos opuestos. La aplicación

$$H \longrightarrow \mathbb{R} \quad x \rightsquigarrow [x, x] \quad (201)$$

se llama una *métrica indefinida* en H .

Por el Teorema de Lax-Milgram sobre la representación de funcionales bilineales continuos en espacios de Hilbert, se tiene que existe un operador acotado autoadjunto e indefinido (es decir, no necesariamente positivo o negativo) tal que para todos $x, y \in H$,

$$[x, y] = \langle G(x), y \rangle. \quad (202)$$

La teoría que expondremos a continuación adquiere una forma más simple cuando $G = J$ es un *operador de involución*, o sea, un operador con las propiedades siguientes:

$$J = J^*, \quad J^2 = I. \quad (203)$$

Notemos que el producto escalar $\langle \ , \ \rangle$ se da solamente para introducir en H una topología. Sin embargo, lo que en esencia nos interesa es

el producto escalar indefinido $[,]$, con ayuda del cual se definen los conceptos fundamentales de esta teoría, clasificándose los operadores de manera análoga a como se hace en los espacios de Hilbert con respecto al producto escalar.

D 2. (operadores J -simétricos y J -unitarios) El operador A , definido en H , se llama J -simétrico, si

$$[A(x), y] = [x, A(y)] \quad (204)$$

para todos $x, y \in D(A)$. Diremos que el operador U , definido en todo H , es J -unitario, si

$$[U(x), U(y)] = [x, y] \quad (205)$$

para todos $x, y \in H$.

Observación. Si el producto escalar indefinido $[,]$ puede ser representado mediante (202) por un operador de involución J (203), entonces para los operadores acotados tendremos que:

$$\begin{aligned} A \text{ es } J - \text{simétrico} &\Leftrightarrow JA \text{ es autoadjunto} . \\ U \text{ es } J - \text{unitario} &\Leftrightarrow JU \text{ es unitario} . \quad \square \end{aligned}$$

Podemos ahora plantearnos el problema de si puede ser cambiado el producto escalar $\langle , \rangle$ en H por otro que defina una métrica topológicamente equivalente a la definida por $\langle , \rangle$ y con relación al cual el producto escalar indefinido pueda ser representado mediante (202) con $G = J$ (203). La respuesta a esta pregunta viene dada en el teorema siguiente que veremos sin demostración.

T 1. Para que el producto escalar $\langle , \rangle$ pueda ser sustituido por otro producto escalar $\langle , \rangle_1$ topológicamente equivalente, tal que

$$[x, y] = \langle G(x), y \rangle = \langle J(x), y \rangle_1, \quad (206)$$

$(x, y \in H)$, es necesario y suficiente, que el operador G tenga inverso acotado.

Nosotros consideraremos precisamente el caso cuando G^{-1} existe y es un operador acotado definido en todo H . Por este motivo, podemos suponer que H es un espacio de Hilbert con una J -métrica, es decir

$$[x, y] = \langle J(x), y \rangle, \quad J = J^*, \quad J^2 = I.$$

No resulta difícil describir la estructura del operador J . Si ponemos

$$P_+ = \frac{I + J}{2} \quad (207)$$

entonces $P_+^* = P_+$ y $P_+^2 = P_+$, luego P_+ es un proyector ortogonal. Poniendo

$$P_- = I - P_+ = \frac{I - J}{2}, \quad (208)$$

tendremos que

$$J = P_+ - P_-, \quad I = P_+ + P_-. \quad (209)$$

La igualdad $J^2 = I$ nos muestra, que el espectro del operador J está formado por los dos puntos ± 1 . Denotemos mediante N_+ y N_- a los subespacios propios del operador indefinido J correspondientes a los valores propios 1 y -1 respectivamente. Obviamente P_+ y P_- son los operadores de proyección ortogonal sobre N_+ y N_- . Cada vector $x \in H$ se escribe en forma única

$$x = P_+(x) + P_-(x) = x_+ + x_-, \quad (210)$$

donde $P_+(x) = x_+$ y $P_-(x) = x_-$. Además

$$J(x) = P_+(x) - P_-(x) = x_+ - x_-, \quad (211)$$

y también

$$\langle x, x \rangle = \langle x_+, x_+ \rangle + \langle x_-, x_- \rangle, \quad [x, x] = \langle x_+, x_+ \rangle - \langle x_-, x_- \rangle. \quad (212)$$

Los primeros resultados importantes de la teoría de operadores en espacios de dimensión infinita con métrica indefinida fueron obtenidos en trabajos de matemáticos soviéticos. Fue precisamente L.S. Soboliev quien primero tuvo que ver con una J -métrica con $\dim N_- = 1$, en relación con algunos problemas de balística. Sin embargo sus resultados no fueron publicados. En 1944 apareció un trabajo de L.S. Pontriaguin dedicado a los operadores J -simétricos para el caso, cuando $\min \dim N_{\pm} < \infty$. En este trabajo fue demostrado un importante teorema que resultó ser la base para muchas investigaciones posteriores relacionadas con métricas indefinidas. Diferentes generalizaciones de estos resultados fueron obtenidos por M.G. Krein, utilizando métodos topológicos. El desarrollo ulterior de esta teoría se ha debido al trabajo de un gran número de matemáticos soviéticos y de otros países.

En lo que sigue necesitaremos algunos resultados de la geometría de los espacios con J -métrica y de la teoría de operadores en esos espacios. Comencemos por las definiciones siguientes:

D 3. Un elemento $x \in H$ se llama J -positivo, si $[x, x] = \|x_+\|^2 - \|x_-\|^2 > 0$, J -negativo, si $[x, x] < 0$ y *neutral*, si $[x, x] = 0$. Denotaremos mediante $H_+(H_-)$ al conjunto de todos los elementos J -no negativos (J -no positivos) de H :

$$H_+ = \{x \in H : [x, x] \geq 0\}, \quad H_- = \{x \in H : [x, x] \leq 0\}. \quad (213)$$

D 4. Un subespacio vectorial $\mathcal{L} \subset H$ se llama:

- i) J -no negativo, si $\mathcal{L} \subset H_+$;
- ii) J -positivo, si para cualquier $x \in \mathcal{L}, x \neq 0, [x, x] > 0$;
- iii) J -no negativo maximal, si es J -no negativo y no está estrictamente contenido en ningún otro subespacio J -no negativo.

Análogamente se introducen los conceptos *J-no positivo*, *J-negativo* y *J-no positivo maximal*. En lo sucesivo denotaremos mediante $\mathcal{L}_+$ al conjunto:

$$\mathcal{L}_+ = \{P_+(x)\}_{x \in \mathcal{L}}.$$

L 1. Sea el subespacio $\mathcal{L} \subset H_+$. Entonces la aplicación $P_+ : \mathcal{L} \rightarrow \mathcal{L}_+$ es lineal, biyectiva y bicontinua.

Demostración. Como $\mathcal{L} \subset H_+$, entonces para cualquier $x \in \mathcal{L}$, se tiene $\|x_+\|^2 \geq \|x_-\|^2$. Por otra parte $\|x\|^2 = \|x_+\|^2 + \|x_-\|^2 \leq 2\|x_+\|^2$.

De aquí se concluye que

$$\frac{1}{2}\|x\|^2 \leq \|x_+\|^2 \leq \|x\|^2$$

de donde se deduce la tesis del Lema. $\square$

T 1. Para que un subespacio $\mathcal{L} \subset H_+$ sea maximal, es necesario y suficiente que $P_+\mathcal{L} = N_+$.

Demostración. Suficiencia: Supongamos que $\mathcal{L} \subset \mathcal{L}_1 \subset H_+$. Entonces $N_+ \supset P_+\mathcal{L}_1 \supset P_+\mathcal{L} = N_+$. Luego $P_+\mathcal{L}_1 = P_+\mathcal{L}$, y por el Lema 1 concluimos que $\mathcal{L} = \mathcal{L}_1$.

Necesidad: Supongamos que $\mathcal{L}_+ = P_+\mathcal{L} \neq N_+$. Por ser $\mathcal{L}$ cerrado y del Lema 1 se tiene que $\mathcal{L}_+$ también es cerrado. Como $\mathcal{L}_+ \neq N_+$, entonces existe un vector $y \in N_+$ ortogonal a $\mathcal{L}_+$. Consideremos el subespacio $\mathcal{L}_1 = \mathcal{L} \oplus [y]$. Obviamente $\mathcal{L}_1 \supset \mathcal{L}$. Queda por probar que $\mathcal{L}_1$ es *J-no negativo*. Pero para un elemento arbitrario $z = \lambda y + x$ ($x \in \mathcal{L}, \lambda \in \mathbb{C}$) tenemos:

$$\begin{aligned} [\lambda y + x, \lambda y] &= \langle \lambda y + x_+, \lambda y + x_+ \rangle - \langle x_-, x_- \rangle \\ &= |\lambda|^2 \langle y, y \rangle + \langle x_+, x_+ \rangle - \langle x_-, x_- \rangle \geq 0, \end{aligned}$$

lo cual se exigía demostrar. $\square$

Del Lema 1 y el Teorema 1 se deduce inmediatamente la siguiente ley de inercia.

T 2 (ley de inercia) La dimensión de cualquier subespacio no negativo (no positivo) maximal de H coincide con la dimensión de $N_+(N_-)$. Luego la dimensión de cualquier subespacio no negativo (no positivo) de H no puede ser superior que la dimensión de $N_+(N_-)$. $\square$

Consideremos todos los subespacios $\mathcal{L} \subset H$, que se proyectan sobre N_+ de manera continua y biunívoca: $P_+\mathcal{L} = N_+$. Por el Teorema de Banach la restricción de P_+ a $\mathcal{L}$ tiene un inverso acotado, al cual denotaremos por $F : x = F(x_+), x \in \mathcal{L}, x_+ \in N_+$. Por eso $x_- = P_-x = P_-Fx_+ = Kx_+$, donde $K = P_-F$ es un operador continuo. El operador K se llama *operador angular* del subespacio $\mathcal{L}$ con respecto a N_+ . Es natural denotar al operador K mediante $K = \tan(\mathcal{L}, \widehat{N_+})$. Entonces se puede representar al subespacio $\mathcal{L}$ como el conjunto de los elementos del tipo $x_+ + Kx_+$, donde x_+ recorre todo N_+ . El operador K actúa, de manera continua, de N_+ en N_- . En fin, tenemos que

$$\mathcal{L} = \{x = x_+ + Kx_+ : x_+ \in N_+\}. \quad (214)$$

Inversamente, si un subespacio $\mathcal{L} \subset H$ se escribe en la forma (214) con un operador continuo K de N_+ en N_- , es fácil ver que puede ser aplicado de manera continua y biunívoca sobre N_+ . Por lo tanto, la fórmula (214) describe todos los subespacios $\mathcal{L}$ de H , que se proyectan ortogonalmente de manera continua y biunívoca sobre todo N_+ . Consideremos $\mathcal{L} \subset H_+$ maximal. Por el Teorema 1, $\mathcal{L}$ se proyecta de manera continua y biunívoca sobre N_+ . Luego $\mathcal{L}$ se puede expresar en la forma (214) con un operador K , el cual en virtud de la J -no negatividad de $\mathcal{L}$ cumple: $\|K\| \leq 1$. Inversamente, la fórmula (214) con $\|K\| \leq 1$, nos da

un subespacio no negativo maximal. Así llegamos al teorema siguiente:

T 3. Para que el subespacio $\mathcal{L}$ sea no negativo y maximal en H , es necesario y suficiente que $\mathcal{L}$ admita una representación del tipo

$$\mathcal{L} = \{x \in H : x = x_+ + Kx_+, \quad x_+ \in N_+\},$$

donde K es un operador contractante ($\|K\| \leq 1$) de N_+ en N_- . $\square$

En lo sucesivo denotaremos mediante $\mathcal{M}_+$ (resp. $\mathcal{M}_-$) a la familia de todos los subespacios maximales de H_+ (resp. H_-).

El Teorema 3 se puede generalizar para cualquier subespacio no negativo. En efecto: si $\mathcal{L} \subset H_+$ entonces $\mathcal{L} = \{y \in H : y = x + Kx, x \in \mathcal{L}_+\}$ donde $\mathcal{L}_+$ es un cierto subespacio de N_+ y K es un operador contractante de $\mathcal{L}_+$ en N_- . De aquí se concluye el importante teorema siguiente:

T 4. Cualquier subespacio $\mathcal{L}$ de H_+ admite una extensión no negativa maximal. Existe una correspondencia biunívoca entre las extensiones no negativas maximales de $\mathcal{L}$ y las extensiones contractantes de su operador angular K a todo N_+ .

Observación. Utilizando el concepto de operador angular, la positividad de un subespacio $\mathcal{L}$ significa que $\|Kx_+\| \leq \|x_+\|, x_+ \neq 0, x_+ \in \mathcal{L}_+$. $\square$

Finalmente enunciaremos en forma de teorema una serie de resultados cuya demostración puede ser hallada en [15].

T Sean H autoadjunto y U J -unitario. Entonces tienen lugar las propiedades siguientes:

- i) El espectro de H (resp. de U) es simétrico con respecto al eje real (resp. con respecto al círculo unidad).

- ii) Si λ es un valor propio normal de H (resp. de U) entonces $\bar{\lambda}$ (resp. $1/\bar{\lambda}$) es también un valor propio normal de H (resp. de U).
- iii) Si $\lambda \neq \bar{\mu}$ (resp. $\lambda\bar{\mu} \neq 1$) entonces $[\mathcal{L}_\lambda(H), \mathcal{L}_\mu(H)] = 0$ (resp. $[\mathcal{L}_\lambda(U), \mathcal{L}_\mu(U)] = 0$) y en particular $[\mathcal{L}_\lambda(H), \mathcal{L}_\lambda(H)] = 0$ si $\lambda \notin \mathbb{R}$ (resp. $[\mathcal{L}_\lambda(U), \mathcal{L}_\lambda(U)] = 0$ si $|\lambda| \neq 1$). Aquí hemos denotado mediante $\mathcal{L}_\lambda$ al subespacio radical correspondiente al valor propio λ .
- iv) Si P_+HP_- (resp. P_+UP_-) es un operador compacto (ver definición en Capítulo 4), entonces la parte no real del espectro de H (resp. la parte del espectro de U que está fuera del círculo unidad) es simétrica con respecto al eje real (resp. con respecto al círculo unidad) y consiste de valores propios normales.
- v) Si P_+HP_- (resp. P_+UP_-) es un operador compacto, entonces para cualquier descomposición $\Lambda \cup \bar{\Lambda}$ (resp. $\Lambda \cup 1/\bar{\Lambda}$) de la parte no real del espectro de H (resp. de la parte del espectro de U que está fuera del círculo unidad), tal que $\Lambda \cap \bar{\Lambda} = \emptyset$ (resp. $\Lambda \cap 1/\bar{\Lambda} = \emptyset$), se pueden encontrar subespacios $\mathcal{L}^\pm \in \mathcal{M}_\pm$ invariantes para H (resp. para U) y que satisfacen la propiedad siguiente:

$$\sigma(H|\mathcal{L}^+) \cap \{Im\lambda \neq 0\} = \Lambda, \quad (215)$$

$$\sigma(H|\mathcal{L}^-) \cap \{Im\lambda \neq 0\} = \bar{\Lambda} \quad (216)$$

$$(\text{resp. } \sigma(U|\mathcal{L}^+) \cap \{|\lambda| \neq 1\} = \Lambda, \quad (217)$$

$$\sigma(U|\mathcal{L}^-) \cap \{|\lambda| \neq 1\} = 1/\bar{\Lambda}). \quad (218)$$

- vi) Si $\min\{\dim N_+, \dim N_-\} = \mathcal{H} < +\infty$ y denotamos mediante $\rho(\lambda)$, la dimensión del subespacio radical $\mathcal{L}_\lambda$ con $\lambda \in \Lambda$, donde λ viene definido en el punto anterior, entonces se cumple que

$$\sum_{\lambda \in \Lambda} \rho(\lambda) \leq \mathcal{H}. \quad (219)$$

Observación. Si Λ satisface las propiedades exigidas en el punto v) del teorema anterior, entonces según iii) tendremos que $[\mathcal{L}_\lambda(H), \mathcal{L}_\mu(H)] = 0$

para cualesquiera $\lambda, \mu \in \Lambda$ con $\lambda \neq \mu$. Pero entonces el subespacio

$$\mathcal{L}_\Lambda = \sum_{\lambda \in \Lambda} \mathcal{L}_\lambda \subset H_+. \quad (220)$$

De la misma forma se obtiene que

$$\mathcal{L}_{\bar{\Lambda}} = \sum_{\lambda \in \bar{\Lambda}} \mathcal{L}_\lambda \subset H_-. \quad (221)$$

La suma en las expresiones (220) y (221) es directa pero no necesariamente ortogonal con respecto al producto escalar definido $\langle \cdot, \cdot \rangle$. Por el Teorema 4 se tiene que $\mathcal{L}_\Lambda$ y $\mathcal{L}_{\bar{\Lambda}}$ tienen extensiones maximales $\mathcal{L}^\pm \in \mathcal{M}_\pm$.

Lo importante del punto v) del Teorema 5 es que nos asegura que pueden ser halladas extensiones maximales de los subespacios $\mathcal{L}_\Lambda$ y $\mathcal{L}_{\bar{\Lambda}}$ que continúan siendo invariantes para el operador H y además satisfacen las importantes propiedades (215) y (216). $\square$

Observación. Del punto vi) del Teorema 5 se concluye una propiedad, que según veremos más adelante, es fundamental en el estudio de la estabilidad de sistemas mecánicos:

Si H es un operador J -autoadjunto (resp. U es J -unitario) y

$$\min\{\dim N_+, \dim N_-\} = \mathcal{H} < +\infty \quad (222)$$

entonces H (resp. U) tiene a lo sumo $2\mathcal{H}$ valores propios no reales (resp. con módulo diferente de 1) teniendo en cuenta la multiplicidad algebraica.

Si un espacio con métrica indefinida satisface la propiedad (220), se dice que es un *espacio de Pontryaguin de orden $\mathcal{H}$* , o también un espacio con $\mathcal{H}$ cuadrados negativos y se denota mediante $\pi_{\mathcal{H}}$. $\square$

Bibliografía

- [1] Fraguera Collar, A. *Análisis matemático en espacios métricos*. Editorial Pueblo y Educación, Habana Cuba 1985.
- [2] Kolmogorof, A. N. y Fomín, S. V. *Elementos de la teoría de funciones y del análisis funcional*. Editorial MIR. Moscú 1975.
- [3] Rudin W. *Real and complex analysis*. Second Ed., Mc. Graw Hill, 1983
- [4] Hinrichsen D., Fernández J. L., Fraguera A. *Topología general*. Editorial Grijalbo, España 1976
- [5] Lions J. L., Magenes E. *Problemes aux limites non homogenes et applications Vol I*. Dunford, Paris 1968. (English trans. Springer-Verlag 1972)
- [6] Mikhailov, V. P. *Partial differential equations*. MIR, Moscú 1978.
- [7] Vladimirov, V. S. *Equations of mathematical physics*. M. Dekker, New York 1971.
- [8] Guelfand, I. M. *Lectures on linear algebra (english translation)*. Interscience, New York 1961.
- [9] Reed, M., Simon, B. *Method of modern mathematical physics I. Functional Analysis*. Academic Press 1980.

- [10] Akhiezer, N. I., Glazman, I. M. *Theory of linear operators in Hilbert space*. Pitman advanced publishing program. Vol. I, II, 1981.
- [11] Sadovnichy, V. A. *Teoría de operadores (en ruso)*. Editorial Nauka, Moscú 1985.
- [12] Birman, M. S., Solomfak, M. E. *Teoría de operadores autoadjuntos (en ruso)*. Editora de la Universidad Estatal de Leningrado, 1980.
- [13] Kato, T. *Perturbation theory of linear operators*. Second edition. Springer-Verlag, 1976.
- [14] Baumgärtel, H. *Analytic perturbation theory for matrices and operators*. Birkhäuser, Basel 1985.
- [15] Krein, M. G. *Introduction to the geometri of indefinite J -spaces and the theory of operators on these spaces*. Second Math. summer School. (Katsiveli 1964) Part I, "Naukova Dumka", Kiev 1965 pp. 15-92. English translation in Amer. Math. Soc. Transl. (2) 93 (1970).

Esta obra se terminó de imprimir en el mes de julio de 1991, en la sección de Offset del Centro de Investigación y de Estudios Avanzados del I. P. N.

La edición consta de 500 ejemplares y sobrantes para reposición.

1. The following information is being furnished to you
 2. as a result of a request made by you on 1/31/81.
 3. The information is being furnished to you as a result
 4. of a request made by you on 1/31/81.
 5. The information is being furnished to you as a result
 6. of a request made by you on 1/31/81.

[illegible]

VII Coloquio de Matemáticas



Centro de Investigación y de Estudios Avanzados del I.P.N.
Departamento de Matemáticas.

