

UNIVERSIDAD NACIONAL DE LA PLATA
FACULTAD DE CIENCIAS EXACTAS
DEPARTAMENTO DE MATEMATICA

Tesis Doctoral

Teorema de Schur-Horn, mayorización conjunta, modelado de
operadores y aplicaciones

Pedro G. Massey

2006

Director: Dr. Demetrio Stojanoff

A Marina, Candelaria y Agustín
A mis padres

Agradecimientos

Quiero agradecer especialmente al Dr. Demetrio Stojanoff por haber compartido conmigo su profunda visión de la matemática y del quehacer matemático, por haber compartido sus conocimientos, por sus consejos acertados, y por la claridad de pensamiento con que ha sabido guiarme en estos inicios a la labor matemática.

También quiero agradecer al Dr. Gustavo Corach por haberme apoyado en varios proyectos académicos importantes para mí, así como por haber sabido prodigar con sus conocimientos el ambiente en que se gestó el grupo de teoría de operadores del IAM, del cual me siento parte y en el cual he aprendido mucha de la matemática que aparece en esta tesis.

Es muy grato para mí expresar mis agradecimientos a mis amigos y compañeros, Francisco Martínez Pería, Mariano Ruíz y Jorge Antezana con quienes he compartido las incertidumbres pero también los pequeños logros obtenidos en estos años de estudio.

Mucho quiero agradecer a Martín Argerami y Luis Silvestre quienes compartieron conmigo su conocimiento y entusiasmo por la matemática y la investigación, así como por la amistad que nos une.

También quiero agradecer a Alejandro Varela, Gabriel Larotonda, Celeste González, Laura Arias, Cristian Conde y Guillermina Fongi por compartir las tardes de seminario de los viernes en el IAM. Quiero agradecer a mis alumnos de la Licenciatura en Matemática, de quienes también he aprendido, por su copioso entusiasmo por el aprendizaje de la matemática. Un agradecimiento muy especial le corresponde al Departamento de Matemática como un todo, por haberme formado como Licenciado y por haberme apoyado desde recién graduado.

Finalmente, quiero agradecer con todo mi corazón por una infinidad de razones a mis padres Alicia y Horacio, a mis Hermanas Gabriela y Natalia, y a Marina quien con su ternura me ha apoyado incondicionalmente y me ha brindado los tesoros de mi vida, Candelaria y Agustín.

P.M.

Índice general

1. Introducción	1
2. Preliminares	27
2.1. Mayorización y submayorización	27
2.1.1. Mayorización de vectores en $\mathbb{R}^n$	27
2.1.2. Matrices no negativas. Teorema de Schur-Horn	30
2.1.3. Mayorización de matrices autoadjuntas	31
2.1.4. Mayorización en $\ell_{\mathbb{R}}^{\infty}(\mathbb{N})$ y el teorema de Schur-Horn	33
2.2. Marcos en espacios de Hilbert	36
2.2.1. Definiciones básicas	36
2.2.2. El exceso de un marco	38
2.2.3. El problema de la reconstrucción de marcos	39
2.3. Álgebras de von Neumann	40
2.3.1. Una introducción abstracta a las álgebras de von Neumann	40
2.3.2. Topologías en $L(\mathcal{H})$. Álgebras de von Neumann concretas	42
2.3.3. Comparación de proyecciones y Clasificación de factores	44
2.3.4. Funcionales y pesos	46
2.3.5. Esperanzas Condicionales	48
2.4. Hechos básicos sobre C^* -álgebras abelianas	49
2.4.1. Presentaciones de las álgebras abelianas	49
2.4.2. Medidas espectrales y representaciones	50
2.4.3. Cálculo funcional en varias variables	52
2.4.4. Valores singulares, preorden espectral y (sub)mayorización	54
2.4.5. Hechos básicos de la teoría de Choquet	56
3. Refinamientos de resoluciones espectrales	59
3.1. Refinamientos y modelado de operadores	59
3.1.1. Refinamientos de resoluciones espectral	59
3.1.2. Modelado de operadores	60
3.1.3. Algunas observaciones sobre el caso semifinito no finito	65
3.2. Refinamientos de resoluciones espectrales en factores II_1	66

4. Desigualdades de tipo Jensen	73
4.1. Funciones monótonas convexas-preorden espectral	73
4.2. Funciones convexas arbitrarias-submayorización	76
4.3. El caso finito dimensional	81
5. Teoremas de tipo Schur-Horn	85
5.1. Teoremas de tipo Schur-Horn en $L(\mathcal{H})$	86
5.2. Un teorema de tipo Schur-Horn para factores II_1	92
5.2.1. Reordenamientos y aproximaciones de funciones	93
5.2.2. Un teorema de tipo Schur-Horn en factores II_1	99
6. Marcos con operador de marco y normas prefijadas	103
6.1. Reformulación de la admisibilidad de marcos	103
6.2. El caso finito dimensional	105
6.3. La admisibilidad en el caso de dimensión infinita	107
6.4. Algunos ejemplos	111
6.4.1. El exceso de marcos en $F(S, \mathbf{c})$	113
7. Mayorización de matrices	115
7.1. Mayorización de matrices	116
7.1.1. Mayorización débil de matrices	117
7.1.2. Convexidad y mayorización de matrices	119
7.1.3. Cuando la mayorización débil implica la mayorización fuerte	122
7.1.4. Relaciones de equivalencias asociadas a las mayorizaciones de matrices	125
7.2. Mayorizaciones conjuntas	127
7.2.1. Mayorización conjunta entre familias abelianas en $\mathbf{H}(\mathbf{n})$	127
7.2.2. Caracterizaciones de las mayorizaciones conjuntas	128
7.2.3. Mayorizaciones conjuntas y funciones convexas	131
7.2.4. Relaciones de equivalencias asociadas a las mayorizaciones conjuntas	133
8. Mayorización conjunta en factores factores II_1	135
8.0.5. Medidas espectrales conjuntas y distribuciones conjuntas	135
8.1. Forma local de las transformaciones doble estocásticas	136
8.2. Núcleos doble estocásticos	140
8.2.1. Comparación con la mayorización conjunta de Alberti y Uhlmann	147
8.3. Algunas herramientas técnicas	148
8.3.1. Refinamientos de medidas espectrales	148
8.3.2. Aproximaciones discretas en álgebras abelianas, separables y difusas	151

Capítulo 1

Introducción

En este trabajo se desarrollan algunos de los muchos tópicos relacionados con la mayorización. Para poder poner en contexto al material contenido en este trabajo hacemos una breve descripción del desarrollo histórico de esta noción, así como de algunos otros conceptos relacionados.

Una breve introducción histórica de la mayorización

Un primer desarrollo sistemático de la mayorización puede hallarse en el libro *Inequalities* (1934) por Hardy, Littlewood y Polya. En este trabajo se define la mayorización entre funciones esencialmente acotadas definidas en un espacio de medida (I, μ) donde $I \subseteq \mathbb{R}$ es un intervalo acotado y μ es una medida finita en I . Se establece aquí la relación entre la mayorización y una familia de desigualdades integrales con respecto a una clase amplia de funciones convexas. Estos resultados fueron mejorados en trabajos de Ky Fan y G.G. Lorentz (1954), D. Brunk (1956), L. Mirsky (1966), W.A. Luxemburg (1967). En 1974, K.M. Chong introduce una interesante generalización de la mayorización, independiente de la noción de reordenada de una función con la que se definió la mayorización originalmente.

Paralelamente al desarrollo de la mayorización en espacios de medida abstractos, la sencilla noción de mayorización entre vectores adquiere más interés en tanto era un recurso importante al momento de probar desigualdades generales. De hecho, las investigaciones de von Neumann (1933) sobre las normas unitariamente invariantes muestran la utilidad de este concepto, que además sucede con relativa frecuencia dentro del análisis matricial. Por otro lado las “matrices doble estocásticas”, que es una clase de matrices relacionadas con la mayorización, atraen la atención de varios investigadores como I. Schur (1923), G. Birkhoff (1946), J.V. Ryff, Y. Sakai y T. Shimogaki (1972). En este clima, A. Horn desarrolla en 1954 un resultado que será clave para nuestro estudio posterior relacionando la mayorización con la comparación del espectro de una matriz autoadjunta con su diagonal principal.

En la década del 50, S. Sherman estudia fenómenos relacionados con extensiones naturales de la mayorización de vectores al caso de mayorización entre n -uplas de vectores en espacios de dimensión finita. En particular, considera problemas relacionados con las hoy llamadas mayorización multivariada y direccional, y problemas relacionados con la teoría de comparación de medidas de Choquet. Lamentablemente, algunos de sus trabajos contienen errores de impresión, lo que crea algún recelo entre los matemáticos de esa época. En particular, en 1954 Sherman publica un

contraejemplo de A. Horn que invalida un resultado suyo obtenido dos años antes, que establecía la equivalencia de la mayorización multivariada con la mayorización direccional. Posiblemente la buena labor de investigación de Sherman en otras áreas del análisis funcional hayan disuadido a otros investigadores de continuar con estos temas. Son relativamente pocos los trabajos de investigación relacionados con estos temas hasta la actualidad.

En 1982 T. Ando, un destacado investigador del análisis funcional y matricial, publica un trabajo en el que se desarrolla una extensión “no conmutativa” de la noción de mayorización entre vectores al caso de mayorización entre matrices autoadjuntas. Esta monografía ha sido el punto de partida de una gran cantidad de trabajos hasta el día de hoy, fundamentalmente relacionados con desigualdades traciales que son de importancia para las aplicaciones. Actualmente, esta noción está ganando espacio dentro de la investigación del fenómeno de “entrelace” de la computación cuántica. Cabe destacar que ya en 1955 A. Grothendieck había estudiado las reordenadas de funciones en relación con desigualdades traciales en las álgebras de von Neumann y esta técnica también había sido considerada por B. Simon (1979).

En 1982, poco tiempo después de que Ando formalizara la mayorización de matrices autoadjuntas, E. Kamei generaliza esta noción a las álgebras de von Neumann denominadas “factores finitos”, dentro de los cuales se encuentran las álgebras de matrices. En este trabajo inicial Kamei prueba, en el contexto de los factores finitos, varios de los resultados obtenidos por T. Ando en las álgebras de matrices. Pero el desarrollo sistemático de la mayorización en el contexto más general de los factores semifinitos es realizada en una serie de trabajos por F. Hiai (1986, 1988, 1990).

Es pertinente remarcar que el estudio de la mayorización, tanto en espacios de medida abstractos como en los factores semifinitos no sólo produjo una serie de resultados importantes, sino que además motivó a considerar una serie de conceptos útiles dentro de estas áreas de la matemática como lo son las reordenadas de funciones en espacios de medida abstracta, los núcleos estocásticos entre espacios funcionales como en las álgebras de operadores, los valores singulares de operadores en álgebras de von Neumann semifinitas y algunas técnicas de aproximación de operadores en factores semifinitos. Además las aplicaciones de este concepto varían desde la comparación de medidas (en términos de medidas de dispersión), desigualdades traciales de operadores (de importancia para las aplicaciones de la teoría de álgebras de operadores a la física), así como aplicaciones a la teoría de potenciales de marcos en espacios de Hilbert y a problemas de entrelaces relacionados con la computación cuántica.

Actualmente, la mayorización en espacios de dimensión infinita se ha considerado desde dos perspectivas distintas. Por un lado, A. Neumann (1999) desarrolla una noción general de mayorización para sucesiones en $\ell^\infty(\mathbb{R})$ y extiende esta noción a operadores autoadjuntos en un espacio de Hilbert separable. De esta forma, obtiene un teorema de tipo Schur-Horn general, en espacios de dimensión infinita. Vale remarcar que la motivación de Neumann proviene de la geometría, más precisamente del problema de extender la validez del teorema de Konstant. Por otro lado, W. Arveson y R. Kadison (2005), dos investigadores reconocidos dentro del análisis funcional y las álgebras de operadores, han considerado problemas de mayorización en las álgebras de operadores, desigualdades de tipo Jensen y teoremas de tipo Schur-Horn. Más aún, plantearon posibles reformulaciones del teorema de Schur-Horn dentro de los llamados factores finitos de tipo II_1 , o factores finitos continuos. Este trabajo es una continuación de trabajos de Kadison (2002).

Nuestro estudio de la mayorización nos llevó a considerar los marcos en espacios de Hilbert. Esta noción surge en un trabajo de Duffin y Schaeffer (1952) con el estudio de desarrollos de tipo L^2 con

respecto a sistemas exponenciales, denominados desarrollos no armónicos, similares a los desarrollos de Fourier. Sin embargo los marcos no fueron reconsiderados hasta el trabajo fundamental de Daubechies, Grossmann y Meyer (1986) sobre onditas, que son marcos en espacios de Hilbert con una estructura determinada. Junto con los desarrollos de marcos de Gabor, las onditas y sus generalizaciones son una fuente de investigación importante del análisis funcional. Actualmente, la estructura adicional de varias clases de marcos se ha dejado a un lado, y se consideran marcos generales (abstractos) en espacios de Hilbert. Algunos de los investigadores más importantes de esta reciente noción son P. Casazza, E. Christensen, H. Han, D. Larson. Recientemente se ha considerado el problema de hallar marcos con cierta estructura, más precisamente con operador de marco y normas de los elementos del marco pre-establecidas. Como veremos, este problema está íntimamente relacionado con el teorema de Schur-Horn de la teoría de mayorización y en particular con extensiones del teorema de A. Neumann.

Contexto general del trabajo

En esta sección hacemos una descripción sintética del contenido del trabajo por capítulos y los situamos en contexto con respecto a resultados relacionados desarrollados en distintos trabajos de investigación. Aclaramos que la organización del contenido de los capítulos no respeta el orden cronológico en que se obtuvieron los resultados allí descritos, sino más bien que está basada en una presentación conceptual de dichos resultados. Es por eso que en la descripción de las motivaciones de los distintos desarrollos tengamos que saltar de capítulos según el orden propuesto.

En el Capítulo 3 introducimos y desarrollamos los aspectos básicos de los *refinamientos de resoluciones espectrales a derecha y acotadas*. Brevemente, una resolución espectral refina a otra resolución si ésta induce una partición más fina de la identidad que aquella. En el caso en que ambas resoluciones se encuentren en un factor $\text{II}_1(\mathcal{M}, \tau)$ establecemos algunas comparaciones en términos de ciertas funciones de acumulación asociadas a estas resoluciones, que están inducidas por la traza del factor. Obtenemos además refinamientos continuos (maximales) de resoluciones arbitrarias en factores de tipo II_1 .

Las resoluciones continuas son el punto de partida para la técnica que denominamos *modelado de operadores*. Este método consiste en lo siguiente: dado un factor $\text{II}_1(\mathcal{M}, \tau)$ y un operador positivo $a \in \mathcal{M}^+$ cuya resolución espectral (a derecha y acotada) es continua entonces, dado cualquier otro operador positivo $b \in \mathcal{M}^+$ construimos el operador $c = h_b(a)$ como cálculo funcional de a de forma que c (el modelo de b) y b comparten los valores singulares o equivalentemente, c y b son aproximadamente unitariamente equivalentes (ver la Observación 2.1.19) en $\mathcal{M}$.

Si bien en los trabajos [41, 42] Hiai y Nakamura consideran algunas resoluciones espectrales continuas y modelos específicos de operadores en el sentido antes descrito, no se realiza un desarrollo sistemático de estas nociones. Más aún, la noción natural de refinamiento no parece haber sido considerada previamente. Además de hacer un desarrollo sistemático de algunos aspectos de estas nociones aplicamos los resultados de refinamientos y modelado obteniendo nuevas caracterizaciones del preorden espectral y la submayorización en factores de tipo II_1 , de características distintas a las obtenidas por Hiai en los trabajos [39, 40]; tanto el preorden espectral introducido por Fujii y Kasahara en [34]. Tanto como la submayorización introducida por Kamei en [49] son preórdenes que se definen entre los operadores positivos de un factor semifinito $(\mathcal{M}, \tau)$ y ocurren con frecuencia en este contexto; ejemplos recientes son los trabajos de Antezana, Massey y Stojanoff [7], Kadison y

Arveson [11], Aujla y Silva [12], Farenick y Manjegani [33] (ver también los trabajos [19, 32, 42]). En particular, nuestras caracterizaciones del preorden espectral dan una respuesta parcial afirmativa a una pregunta realizada por D. Farenick y M. Manjegani en [33] al respecto. Más aún, estas caracterizaciones permiten hallar nuevas reformulaciones de desigualdades del tipo Young [33] y Jensen [7, 11] obtenidas recientemente (ver el Capítulo 4).

Estos resultados fueron motivados por una pregunta de D. Farenick con respecto a una instancia particular del preorden espectral (en su estudio de las desigualdades de tipo Young en factores semifinitos). Al momento de considerar ese problema, resultó claro que había que desarrollar un método para eliminar los átomos de una medida espectral. Motivados por algunas ideas desarrolladas en el estudio de refinamientos de medidas espectrales conjuntas (ver Capítulo 8) es que desarrollamos los refinamientos de resoluciones espectrales. El contenido de este capítulo forma parte del trabajo “Refinements of spectral resolutions and modelling of operators in II_1 factors” que se encuentra en referato para su publicación.

En el Capítulo 4 obtenemos *desigualdades de tipo Jensen* para funciones convexas con respecto al preorden espectral y la submayorización, en el contexto de las álgebras de operadores. De esta forma, no solo mostramos dos instancias naturales en las que estos preórdenes ocurren, sino que además obtenemos un ámbito donde aplicar los resultados del Capítulo 3.

Esta clase de desigualdades ha generado interés por parte de varios investigadores. Ejemplos recientes del estudio de desigualdades de tipo Jensen son los trabajos Arveson y Kadison [11], Aujla y Silva [12], los trabajos de Hansen y Pedersen [36, 37] y el trabajo de Pedersen [59] (ver también los trabajos [16, 19]). En particular, los trabajos recientes [36, 37, 59] motivaron nuestro estudio de este fenómeno de convexidad. Recordemos que la desigualdad de Jensen es de tipo integral y se da en un espacio de probabilidad (espacio de medida de masa igual a 1). En el ámbito de las álgebras de operadores, la medida se reemplaza por una integral abstracta (transformación positiva y unital).

En la primera parte de este capítulo se desarrollan desigualdades del tipo Jensen para funciones convexas monótonas, y en este caso se obtiene una comparación entre los términos de esta desigualdad con respecto al preorden espectral. Estos resultados generalizan los obtenidos en [19]. Utilizando las caracterizaciones desarrolladas en el Capítulo 3, obtenemos reformulaciones de estas desigualdades con respecto al orden usual de operadores en ciertos casos.

En la segunda parte obtenemos una desigualdad de tipo Jensen para funciones convexas de varias variables, de características más bien técnicas. Este resultado generaliza algunas desigualdades obtenidas en los trabajos [11, 12, 16, 36, 37]. En el caso particular que el álgebra de operadores considerada sea un factor finito, entonces esta desigualdad implica una desigualdad de tipo Jensen con respecto a la submayorización. Este resultado generaliza resultados de los trabajos de Arveson y Kadison [11] y Aujla y Silva [12]. Algunos corolarios de este caso particular de desigualdad de Jensen serán importantes en nuestro desarrollo posterior de un teorema de tipo Schur-Horn en factores II_1 . El contenido de este capítulo forma parte del trabajo “Jesen inequality and majorization” que se encuentra actualmente en realización en co-autoría con J. Antezana y D. Stojanoff.

En el Capítulo 5 desarrollamos dos versiones diferentes de *extensiones del teorema de Schur-Horn al caso de espacios de dimensión infinita*. Comenzamos con el estudio del conjunto de posibles diagonales de operadores autoadjuntos en un espacio de Hilbert separable $\mathcal{H}$ de dimensión infinita. Los resultados de esta parte del capítulo están basados en los obtenidos por A. Neumann en [56] en relación a la mayorización de vectores en $\ell_{\mathbb{R}}^{\infty}(\mathbb{N})$, que generaliza la noción usual de mayorización en

$\mathbb{R}^n$, y sus extensiones del teorema de Schur-Horn a operadores autoadjuntos en $L(\mathcal{H})_h$. En este caso, obtenemos caracterizaciones del conjunto de posibles diagonales de un operador autoadjunto acotado arbitrario, con respecto a sus representaciones matriciales inducidas por bases ortonormales; estas descripciones están especialmente desarrolladas para ser aplicadas al problema de admisibilidad de marcos que estudiamos en el Capítulo 6, que es la motivación de nuestro estudio del teorema de Schur-Horn en este contexto. Algunos de nuestros resultados están también relacionados con los trabajos de Kadison [46, 47] y el trabajo reciente de Arveson y Kadison [11].

En la segunda parte de este capítulo obtenemos una versión del *teorema de Schur-Horn en factores de II_1* , que da una respuesta parcial afirmativa a una conjetura propuesta por Arveson y Kadison en [11]. En realidad obtenemos una versión un poco más débil de la conjetura allí propuesta (para más detalles ver la discusión del comienzo de la sección 5.2) pero representa una primera descripción del conjunto de operadores de una subálgebra abeliana maximal (masa) mayorizados por un operador positivo fijo en los términos deseados. A tal fin, estudiamos algunas propiedades sencillas de las reordenadas de una función monótona y desarrollamos algunos resultados de aproximación de funciones integrables en espacios de medida (I, ν) donde $I \subseteq \mathbb{R}$ es un intervalo compacto y ν es una medida difusa (es decir, los puntos tienen masa cero). Nuestro enfoque está fuertemente influenciado por el desarrollo de las técnicas de refinamientos y modelado de operadores desarrolladas en el Capítulo 3 y depende de algunas de las versiones de la desigualdad de Jensen desarrolladas en el Capítulo 4. El contenido de la primera parte de este capítulo forma parte del trabajo “The Schur-Horn theorem for operators and frames with prescribed norms and frame operator.”^{en} co-autoría con J. Antezana, M. Ruiz y D. Stojanoff, que está aceptado para su publicación en la revista Illinois Journal of Mathematics. La segunda parte del capítulo está basada en el artículo “A Schur-Horn theorem in II_1 factors.”^{en} co-autoría con M. Argerami, que se encuentra en realización.

En el Capítulo 6 utilizamos las versiones extendidas del teorema de Schur-Horn en $L(\mathcal{H})_h$ para hallar condiciones necesarias para la *existencia de marcos con ciertas condiciones prefijadas*. Adelantamos brevemente que los marcos son ciertas sucesiones de vectores en un espacio de Hilbert $\mathcal{H}$ que permiten la descomposición y reconstrucción de vectores de $\mathcal{H}$ (de forma similar a las bases ortonormales). Actualmente, estos objetos son estudiados por varios grupos de investigadores a nivel internacional debido a sus aplicaciones en problemas de transferencia de datos.

En trabajos recientes de investigación de Casazza y Leon [21, 22], Dykema, Freeman, Korleson, Larson, Ordower y Weber [30], Larson y Korleson [50] y Topp, Dhillon, Health y Strohmer [67] se considera el problema de la existencia y construcción (algorítmica) de marcos con ciertas condiciones adicionales, más explícitamente con operador de marco y normas de sus elementos prefijadas. Este problema es el denominado “problema de la admisibilidad”. Es conocido el hecho que (ver [21], [67]) , en el caso de espacios de Hilbert $\mathcal{H}$ de dimensión finita y marcos de una cantidad finita de elementos, hay una conexión entre el problema de la admisibilidad y la teoría de mayorización, en particular con el teorema de Schur-Horn. En este capítulo hacemos explícita esta conexión tanto en el caso de dimensión finita como en el contexto de dimensión infinita. Esta presentación del problema nos permite obtener condiciones necesarias para la admisibilidad. Más aún, fortaleciendo estas condiciones necesarias obtenemos condiciones suficientes para la admisibilidad que son menos restrictivas que las halladas en los trabajos [30, 50]. Hacia el final del capítulo, desarrollamos una serie de ejemplos que muestran que las condiciones suficientes halladas anteriormente no pueden relajarse más en general. También mostramos, haciendo un análisis del exceso de los marcos, la variedad de posibles soluciones cualitativamente distintas al problema de la admisibilidad. El con-

tenido de este capítulo forma parte del trabajo antes mencionado “The Schur-Horn theorem for operators and frames with prescribed norms and frame operator.”^{en} co-autoría con J. Antezana, M. Ruiz y D. Stojanoff.

Los Capítulos 7 y 8 corresponden al desarrollo de lo que llamamos mayorizaciones conjuntas, que constituyen una extensión de la mayorización entre operadores al caso de mayorización entre n -uplas de operadores que conmutan entre sí.

La mayorización de vectores en $\mathbb{R}^n$ admite dos extensiones naturales al caso de m -uplas de vectores en $\mathbb{R}^n$ denominadas mayorización fuerte y direccional que ya han sido consideradas en trabajos de Sherman [62] y Shreiber [63] y en los libros de Alberti y Uhlman [2] y Marshall y Olkin [54]. Más recientemente, se han estudiado en los trabajos de Beasley y Lee [14], Cheon y Lee [23], Dahl [26], Hwang y Pyo [45], Li y Poon [52] y Martínez, Massey y Silvestre [55]. En el Capítulo 7 desarrollamos una noción original, que llamamos mayorización débil, relacionada con las extensiones anteriores. Si bien la mayorización débil tiene interpretaciones geométricas sencillas resulta ser un concepto auxiliar útil en la comparación de la mayorización fuerte y direccional. Utilizando técnicas elementales de convexidad, establecemos algunas propiedades y relaciones que existen entre estas tres mayorizaciones. En particular, consideramos el problema de hallar condiciones bajo las cuales la mayorización débil (o direccional) implica la mayorización fuerte (este problema ha generado alguna controversia en los comienzos de esta teoría, ver la sección 1). Nuestros resultados generalizan los trabajos [45, 62, 63].

En la segunda parte de este capítulo consideramos posibles extensiones de la mayorización de matrices autoadjuntas al caso de m -uplas de matrices autoadjuntas que conmutan entre sí. Luego de desarrollar las definiciones, establecemos algunas propiedades básicas así como algunas relaciones entre éstas. Obtenemos también caracterizaciones de estas nociones en términos del cálculo funcional multivariado, que en este contexto tiene una descripción particularmente sencilla; estos resultados están relacionados parcialmente con el trabajo de Dahl [26]. Estos desarrollos pueden considerarse una preparación más bien elemental para abordar la mayorización conjunta de operadores en factores II_1 que se realiza en el capítulo siguiente, y que es la motivación por la cual estudiamos estos conceptos. El contenido de este capítulo forma parte del trabajo “Weak matrix majorization.”^{en} co-autoría con F.D. Martínez Pería y L.E. Silvestre, que ha sido publicado en la revista *Linear Algebra and its Applications* (2005).

Recientemente, la mayorización ha vuelto a conquistar el interés debido a sus relaciones con las clausuras en norma de las órbitas unitarias de operadores autoadjuntos y las esperanzas condicionales sobre subálgebra abelianas. Algunos ejemplos recientes de trabajos relacionados con esta noción son Antezana, Massey y Stojanoff [7], Antezana, Massey, Ruiz y Stojanoff [8], Arveson y Kadison [11], Farenick y Manjegani [33], los trabajos de Kadison [46, 47] y el trabajo de Neumann [56]. Uno de los objetivos del Capítulo 8 es obtener una extensión de la noción de mayorización entre operadores autoadjuntos al caso de una comparación entre n -uplas de operadores autoadjuntos que conmutan entre sí en un factor II_1 , que denominamos *mayorización conjunta* de familias abelianas (en el caso finito dimensional, esta noción coincide con la mayorización conjunta fuerte considerada en el Capítulo 7).

Para obtener caracterizaciones de esta noción extendida describimos la *forma local* de una transformación doble estocástica (DS), i.e. obtenemos una familia de transformaciones DS particularmente sencillas que aproximan la restricción de cualquier transformación DS a una C^* -subálgebra

abeliana y separable de un factor II_1 . Como herramienta para obtener esta forma local, construimos refinamientos abelianos, separables y *difusos* de una C^* -subálgebra abeliana y separable de un factor II_1 $\mathcal{M}$. Algunas de las técnicas desarrolladas en este caso son nuevas, incluso cuando son restringidas a la mayorización de operadores autoadjuntos en factores II_1 . Nuestros resultados generalizan, en el caso de factores II_1 , algunos de los obtenidos en los trabajos [2, 7, 11, 39, 40, 49, 61].

Nos hemos restringido al caso de factores II_1 porque, por un lado, aspectos técnicos del trabajo se hacen más simples y por otro, pues éste es el contexto donde la mayorización tiene su significado pleno: este fenómeno se aprecia incluso en la mayorización entre operadores autoadjuntos, donde la (doble) mayorización en factores II_1 caracteriza la equivalencia aproximadamente unitaria. Señalamos también que, como cada álgebra de von Neumann actuando en un espacio de Hilbert separable tiene una descomposición en integral directa de factores finitos, el estudio de factores finitos provee de información útil para el caso de álgebras más generales. El contenido de este capítulo forma parte del trabajo “The local form of doubly stochastic maps and joint majorization in II_1 factors.”^{en} co-autoría con M. Argerami, que se encuentra en referato para su publicación.

Resumen de los resultados originales

En esta sección resumimos los resultados originales que aparecen en este trabajo respetando los contenidos por capítulo. Vamos a introducir informalmente las definiciones y conceptos necesarios para la exposición de los siguientes resultados y referimos al lector a los Preliminares del Capítulo 2, para más detalles. En algunos casos, para clarificar la exposición, presentamos versiones restringidas de los resultados obtenidos y dejamos algunos detalles técnicos adicionales para ser desarrollados en los capítulos correspondientes. A fin de enfatizar el enunciado de los resultados, éstos aparecen en renglón aparte con el signo “•”.

Capítulo 3: Refinamientos de resoluciones espectrales

En este contexto, una resolución espectral a derecha y acotada de p en un álgebra de von Neumann $\mathcal{M}$ (que abreviamos REDA de p en $\mathcal{M}$) es una función decreciente $E_{(\cdot)} : I \rightarrow \mathcal{P}(\mathcal{M})$ de un intervalo compacto $I = [\alpha, \beta] \subseteq \mathbb{R}$ a valores en el reticulado de proyecciones de un álgebra de von Neumann, que es continua a derecha cuando dotamos a $\mathcal{P}(\mathcal{M})$ de la topología fuerte de operadores y es tal que $E_\alpha = p$ y $E_\beta = 0$ (ver el comienzo de la sección 3.1.1). Si $E_{(\cdot)} : I \rightarrow \mathcal{P}(\mathcal{M})$ es una REDA entonces por simplicidad la describimos como $\{E_\lambda\}_{\lambda \in I}$; y si el intervalo I está claro del contexto escribimos simplemente $\{E_\lambda\}$.

El ejemplo de REDA en el que estamos interesados (que de hecho es un modelo general para las REDAs) es el siguiente: dado un operador positivo $a \in \mathcal{M}^+$ en un álgebra de von Neumann $\mathcal{M}$ entonces si definimos $E_\lambda = P^a(\lambda, \infty)$, donde $P^a(\Delta)$ denota la proyección espectral del operador a correspondiente al conjunto (medible Borel) $\Delta \subseteq \mathbb{R}$, entonces $\{E_\lambda\}_{\lambda \in [0, \|a\|]}$ es una REDA de $P_{\overline{R(a)}}$, el proyector al rango de a .

En el capítulo 1 introducimos la siguiente noción de refinamiento entre REDAs

• Sean $\{E_\lambda\}_{\lambda \in I}$ y $\{E'_\lambda\}_{\lambda \in I'}$ REDAs con $I = [\alpha, \beta]$, $I' = [\alpha', \beta']$. Dada una función $f : I \rightarrow I'$, decimos que el par $(\{E'_\lambda\}, f)$ es un *refinamiento* de $\{E_\lambda\}$ si

- f es creciente, continua a derecha y $f(\beta) = \beta'$;

- $E_\lambda = E'_{f(\lambda)}$ para cada $\lambda \in I$.

□

En lo que sigue $(\mathcal{M}, \tau)$ denota un factor de tipo II_1 es decir un álgebra de von Neumann de centro trivial y dotada de una traza fiel, normal y finita tal que para todo $\alpha \in [0, 1]$ existe una proyección $p \in \mathcal{P}(\mathcal{M})$ tal que $\tau(p) = \alpha$. Consideramos a continuación los valores singulares de operadores positivos en $(\mathcal{M}, \tau)$ junto con los preórdenes espectral y la submayorización. Si $a \in \mathcal{M}^+$ un operador positivo entonces definimos los τ -valores singulares de a por

$$\mu_a(t) = \min\{s \in \mathbb{R}_0^+ : \tau(p^a(s, \infty)) \leq t\}, \quad t \geq 0.$$

Si $a, b \in \mathcal{M}^+$ son operadores positivos entonces decimos que

- b domina espectralmente a a si $\mu_a(t) \leq \mu_b(t)$ para $t \geq 0$ y notamos $a \preceq b$.
- b submayoriza a a si $\int_0^s \mu_a(t) dt \leq \int_0^s \mu_b(t) dt$ para $s \geq 0$ y notamos $a \prec_w b$.
- Si $a \prec_w b$ y además $\tau(a) = \tau(b)$ entonces decimos que b mayoriza a a y notamos $a \prec b$.

Hemos mencionado que dado $a \in \mathcal{M}^+$, éste induce una REDA $\{P^a(\lambda, \infty)\}_{\lambda \in [0, \|a\|]}$. En este caso, decimos que a tiene distribución continua si la REDA inducida por él resulta continua, o equivalentemente, si la función de $\tau(P^a(\lambda, \infty))$ es una función continua.

El siguiente resultado es el denominado modelado de operadores en factores de tipo II_1 . En su enunciado la expresión *álgebra abeliana maximal* es abreviada por *masa*. Las masas del factor juegan un papel importante en relación al teorema de Schur-Horn en factores de tipo II_1 , que consideramos en el Capítulo 5.

• Sea $(\mathcal{M}, \tau)$ un factor II_1 y sea $a \in \mathcal{M}^+$. Entonces, existe un operador $a' \in \mathcal{M}^+$ con distribución continua tal que

- * La resolución espectral de a' refina la resolución espectral de a .
- * Si $b \in \mathcal{M}^+$ entonces, existe una función h_b creciente y continua a izquierda tal que, si $\tilde{b} = h_b(a')$ entonces $\mu_b = \mu_{\tilde{b}}$. Más aún, la resolución espectral de a' refina la resolución espectral de b si y solo si $h_b(a') = b$.
- * Si $c^+ \in \mathcal{M}$ entonces $c \preceq b$ (resp $c \prec_w b$, $c \prec b$) si y solo si $\tilde{c} \leq \tilde{b}$ (resp $\tilde{c} \prec_w \tilde{b}$, $\tilde{c} \prec \tilde{b}$).

Más aún, si $\mathcal{A} \subseteq \mathcal{M}$ es una masa y $a \in \mathcal{A}^+$ entonces podemos elegir a $a' \in \mathcal{A}^+$.

□

Como consecuencia del modelado de operadores antes expuesto obtenemos las siguientes caracterizaciones del preorden espectral y la submayorización. En lo que sigue, $\overline{\mathcal{U}_{\mathcal{M}}(b)}$ denota la clausura en norma de la órbita unitaria de b i.e.

$$\overline{\mathcal{U}_{\mathcal{M}}(b)} = \overline{\{u^*bu : u \in \mathcal{M} \text{ unitario}\}}^{\|\cdot\|}$$

- Sea $(\mathcal{M}, \tau)$ un factor II_1 y sean $a, b \in \mathcal{M}^+$. Entonces

* b domina espectralmente a a ($a \lesssim b$) si y solo si existe

$$c \in \overline{\mathcal{U}_{\mathcal{M}}(b)} \quad \text{tal que} \quad a \leq c$$

o, equivalentemente, si existe

$$\overline{\mathcal{U}_{\mathcal{M}}(a)} \ni d \quad \text{tal que} \quad d \leq b.$$

Más aún , podemos asumir que a y c conmutan y, b y d conmutan.

* b submayoriza a a ($a \prec_w b$) si y solo si existe $c \in \mathcal{M}^+$ tal que

$$a \leq c \prec b.$$

Más aún , podemos asumir que a y c conmutan.

□

De esta forma se tiene la siguiente caracterización de la dominación espectral que complementa las anteriores

• Sea $(\mathcal{M}, \tau)$ un factor II_1 y sean $a, b \in \mathcal{M}^+$. Entonces b domina espectralmente a a ($a \lesssim b$) si y solo si existe una REDA $\{E_\lambda\}_{\lambda \in I}$, donde $I = [0, \|a\|]$ tal que $\tau(E_\lambda) = \tau(P^a(\lambda, \infty))$ para cada $\lambda \in I$ y

$$\lambda E_\lambda \leq E_\lambda b E_\lambda, \quad \forall \lambda \geq 0.$$

□

Si bien en el caso de un factor semifinito puede mostrarse que las condiciones anteriores no caracterizan la dominación espectral, se tiene el siguiente resultado en el caso especial del factor semifinito $L(\mathcal{H})$.

• Sea $\mathcal{H}$ un espacio de Hilbert. Dados a, b operadores positivos de $L(\mathcal{H})$ tales que a es compacto, b es diagonalizable y $a \lesssim b$ entonces, existe una isometría parcial $u \in L(\mathcal{H})$ con espacio inicial $\overline{R(a)}$ tal que

$$uau^* \leq b \quad \text{y} \quad (uau^*)b = b(uau^*).$$

Más aún, si $\mathcal{H}$ tiene dimensión finita, la misma desigualdad vale para operadores autoadjuntos $a, b \in L_{sa}(\mathcal{H})$ y la isometría puede cambiarse por un unitario.

□

El modelado de operadores está basado en ciertas construcciones en términos de refinamientos de resoluciones espectrales a derecha y acotadas (REDAs) en $(\mathcal{M}, \tau)$. En lo que sigue describimos los resultados correspondientes a estos desarrollos mas bien técnicos.

Sea $I = [\alpha, \beta]$ y sea $\{E_\lambda\}_{\lambda \in I}$ una REDA de una proyección $p \in \mathcal{P}(\mathcal{M})$. Decimos que $\lambda_0 \in (\alpha, \beta]$ es un átomo de $\{E_\lambda\}$, si la resolución no es continua a la izquierda en λ_0 i.e, si

$$\lim_{\lambda \rightarrow \lambda_0^-} E_\lambda = E_{\lambda_0} + p(\lambda_0), \quad p(\lambda_0) \neq 0.$$

En este caso $p(\lambda_0)$ es la *proyección salto* de $\{E_\lambda\}$ en λ_0 . Si $p \neq 1$ entonces α es considerado un átomo. El conjunto de átomos de $\{E_\lambda\}_{\lambda \in I}$ es denotado por $\text{At}(\{E_\lambda\})$; este conjunto es vacío si y solo si la resolución es continua.

Al número real

$$\mathcal{J}(\{E_\lambda\}) := \sum_{\lambda \in \text{At}(\{E_\lambda\})} \tau(p(\lambda)) = \tau \left(\sum_{\lambda \in \text{At}(\{E_\lambda\})} p(\lambda) \right) \leq 1$$

lo llamamos el *salto total* de la resolución; este número da una idea de las discontinuidades de la resolución. Hemos probado

- Sean $\{E_\lambda\}_{\lambda \in I}$, $\{E'_\lambda\}_{\lambda \in I'} \subseteq \mathcal{M}$ REDAs, donde $(\mathcal{M}, \tau)$ es un factor II_1 . Si existe una función f tal que $(\{E'_\lambda\}, f)$ refina a $\{E_\lambda\}$ entonces $\mathcal{J}(\{E_\lambda\}) \geq \mathcal{J}(\{E'_\lambda\})$. □

- Decimos que un refinamiento $(\{E'_\lambda\}, f)$ de $\{E_\lambda\}$ es un *refinamiento fuerte* si f satisface

- $f(\lambda) \geq \lambda$ para cada $\lambda \in I$, y
- $f(\lambda) - f(\mu) \geq \lambda - \mu$, para cada $\lambda > \mu \in I$.

Si $\{\alpha_k\}_{k \in \mathbb{N}} \in l^1(\mathbb{R}^+)$, decimos que la sucesión $(\{E_\lambda^k\}_{\lambda \in I_k})_{k \in \mathbb{N}} \subseteq \mathcal{M}$ de REDAs es $\{\alpha_k\}_{k \in \mathbb{N}}$ -*compatible* si valen las siguientes condiciones:

- $\exists 0 \leq \alpha < \beta \in \mathbb{R}$ tal que $I_k = [\alpha, \beta + \sum_{i=1}^k \alpha_i]$ para cada $k \in \mathbb{N}$.
- $(\{E_\lambda^{k+1}\}, f_k)$ es un refinamiento fuerte de $\{E_\lambda^k\}$ para cada $k \in \mathbb{N}$.
- $f_k(\lambda) - \lambda \leq \alpha_k$, para cada $\lambda \in I_k$ y cada $k \in \mathbb{N}$.

El siguiente resultado expresa que una sucesión compatible de resoluciones se amalgama en una resolución que refina simultáneamente a cada miembro de la sucesión. En lo que sigue consideramos además el caso especial en que las resoluciones están dentro de una masa (subálgebra abeliana maximal) $\mathcal{A} \subseteq \mathcal{M}$.

- Sea $(\mathcal{M}, \tau)$ un factor II_1 y sea $\{\alpha_k\}_{k \in \mathbb{N}} \in l^1(\mathbb{R}^+)$. Si $(\{E_\lambda^k\}_{\lambda \in I_k})_{k \in \mathbb{N}} \subseteq \mathcal{M}$ es $\{\alpha_k\}_{k \in \mathbb{N}}$ -compatible entonces existe una REDA $\{E_\lambda\}_{\lambda \in I}$ de p en $\mathcal{M}$ tal que $(\{E_\lambda\}, h_k)$ es un refinamiento fuerte de $\{E_\lambda^k\}$, para cada $k \in \mathbb{N}$. Más aún, si $\mathcal{A} \subseteq \mathcal{M}$ es una masa y cada miembro de la sucesión es una REDA en $\mathcal{A}$ entonces podemos elegir $\{E_\lambda\}$ en $\mathcal{A}$. □

El siguiente resultado muestra que, dado un salto de una resolución espectral, éste puede ser eliminado considerando un refinamiento conveniente.

- Sea $\{E_\lambda\}_{\lambda \in [\alpha, \beta]} \subseteq \mathcal{M}$ una REDA de $p \in \mathcal{P}(\mathcal{M})$, donde $(\mathcal{M}, \tau)$ es un factor II_1 , y sea $\lambda_0 \in \text{At}(\{E_\lambda\})$. Entonces, existe un refinamiento fuerte $(\{E'_\lambda\}_{\lambda \in I'}, f)$ de $\{E_\lambda\}$, donde $I' = [\alpha, \beta + \tau(p(\lambda_0))]$ tal que

$$\mathcal{J}(\{E'_\lambda\}) = \mathcal{J}(\{E_\lambda\}) - \tau(p(\lambda_0))$$

y para cada $\lambda \in I$ tenemos $f(\lambda) - \lambda \leq \tau(p(\lambda_0))$. Más aún, si $\mathcal{A} \subseteq \mathcal{M}$ es una masa y $\{E_\lambda\}$ es una REDA en $\mathcal{A}$ entonces podemos elegir $\{E'_\lambda\}$ en $\mathcal{A}$. □

Iterando el procedimiento anterior se puede construir una sucesión compatible de resoluciones. Considerando la resolución que surge de amalgamar esta sucesión compatible se prueba que pueden eliminarse todos los saltos de una resolución espectral mediante un refinamiento adecuado.

• Sea $(\mathcal{M}, \tau)$ un factor II_1 . Dada una REDA $\{E_\lambda\}_{\lambda \in I}$ de p en $\mathcal{M}$, existe una REDA continua $\{E'_\lambda\}_{\lambda \in I'}$ en $\mathcal{M}$ y una función $f : I \rightarrow I'$ tal que $(\{E'_\lambda\}, f)$ es un refinamiento fuerte de $\{E_\lambda\}$. Más aún, si $\mathcal{A} \subseteq \mathcal{M}$ es una masa y $\{E_\lambda\}$ es una REDA en $\mathcal{A}$ entonces podemos elegir a $\{E'_\lambda\}$ también en $\mathcal{A}$. $\square$

El resultado anterior es la base para poder desarrollar el resultado sobre modelado de operadores.

Capítulo 4: Desigualdades de tipo Jensen

En este capítulo estudiamos desigualdades de tipo Jensen en el ámbito de las álgebras de operadores. En este contexto, consideramos integrales “no conmutativas” de probabilidad representadas por transformaciones lineales $\phi : \mathcal{A} \rightarrow \mathcal{B}$ positivas ($\phi(a) \geq 0$ siempre que $a \geq 0$) y unitales ($\phi(1) = 1$) entre C^* -álgebras unitales $\mathcal{A}, \mathcal{B}$. En lo que sigue consideramos el preorden espectral que ha sido definido en la descripción del capítulo anterior (ver también la Definición 2.4.15 en los Preliminares).

• Sea $\mathcal{A}$ una C^* -álgebra unital, $\mathcal{B}$ un álgebra de von Neumann y $\phi : \mathcal{A} \rightarrow \mathcal{B}$ una transformación positiva y unital. Entonces, para cada función convexa y monótona f , definida en algún intervalo I , y para cada operador autoadjunto $a \in \mathcal{A}$ cuyo espectro está contenido en I , se tiene que $\phi(f(a))$ domina espectralmente a $f(\phi(a))$ y en símbolos

$$f(\phi(a)) \preceq \phi(f(a)).$$

$\square$

Como consecuencia de la desigualdad de arriba y de las caracterizaciones de la dominación espectral en el caso de $L(\mathcal{H})$ y los factores de tipo II_1 del capítulo anterior obtenemos

• Sea $\mathcal{A}$ una C^* -álgebra unital y $\phi : \mathcal{A} \rightarrow \mathcal{M}$ una transformación positiva y unital en un factor semifinito $(\mathcal{M}, \tau)$. Entonces, para cada función convexa y monótona f definida en $[0, +\infty)$ se tiene:

- * Si $\mathcal{M} = L(\mathcal{H})$ con $\mathcal{H}$ un espacio de Hilbert separable, f tal que $f(0) = 0$, y para cada operador positivo a de $\mathcal{A}$ tal que $\phi(a)$ y $\phi(f(a))$ son compactos, entonces existe una isometría parcial $u \in L(\mathcal{H})$ con espacio inicial $R(f(\phi(a)))$ tal que:

$$uf(\phi(a))u^* \leq \phi(f(a)) \quad \text{y} \quad \left(uf(\phi(a))u^*\right)\phi(f(a)) = \phi(f(a))\left(uf(\phi(a))u^*\right).$$

- * Si $(\mathcal{M}, \tau)$ es un factor II_1 entonces existen sucesiones $(u_n)_{n \in \mathbb{N}}, (v_n)_{n \in \mathbb{N}} \in \mathcal{M}$ de operadores unitarios tales que

$$f(\Phi(a)) \leq \lim_{n \rightarrow \infty} u_n^* \Phi(f(a)) u_n$$

y

$$\lim_{n \rightarrow \infty} v_n^* f(\Phi(a)) v_n \leq \Phi(f(a)).$$

$\square$

En el caso de funciones convexas arbitrarias de una variable los resultados anteriores no valen; sin embargo obtenemos desigualdades de tipo Jensen para las funciones convexas continuas de varias variables con respecto a la submayorización, que es un preorden más débil que el espectral.

Para establecer los siguientes enunciados referimos al lector a la Definición 2.3.27 de esperanza condicional y la sección 2.4.3 en donde se describe el espectro conjunto de una familia abeliana y el cálculo funcional en varias variables. Adelantamos aquí que una esperanza es una proyección de una álgebra sobre una subálgebra y constituye un caso particular de transformación positiva y unital. Por otro lado, el espectro conjunto de una familia de operadores autoadjuntos que conmutan entre sí (familia abeliana) es una generalización de la noción de espectro de un operador y en este contexto el cálculo funcional (continuo) usual de un operador se extiende al caso de funciones (continuas) de varias variables aplicadas a una familia abeliana.

• Sean $\mathcal{A}$ y $\mathcal{B}$ C^* -álgebras unitales, $\phi : \mathcal{A} \rightarrow \mathcal{B}$ una transformación positiva y unital, $U \subseteq \mathbb{R}^n$ un conjunto convexo y abierto y $f : U \rightarrow \mathbb{R}$ una función convexa. Sea $(a_1, \dots, a_n) \subseteq \mathcal{A}_{sa}$ una familia abeliana tal que $\prod_{i=1}^n \sigma(a_i) \subseteq U$ y tal que $(\phi(a_1), \dots, \phi(a_n), \phi(f(a_1, \dots, a_n))) \subseteq \mathcal{B}_{sa}$ resulta también una familia abeliana. Entonces tenemos

$$f(\phi(a_1), \dots, \phi(a_n)) \leq \phi(f(a_1, \dots, a_n)).$$

Más aún, si $\tilde{0} = (0, \dots, 0) \in U$ y $f(\tilde{0}) \leq 0$ entonces la ecuación anterior vale si ϕ es positiva y contractiva. □

Dada una esperanza condicional $\mathcal{E} : \mathcal{B} \rightarrow \mathcal{C}$, el centralizador de $\mathcal{E}$ es la C^* -subálgebra de $\mathcal{B}$ definida por:

$$\mathcal{B}^{\mathcal{E}} = \{b \in \mathcal{B} : \mathcal{E}(ba) = \mathcal{E}(ab), \forall a \in \mathcal{B}\}$$

Notemos que, para cada $b \in \mathcal{B}^{\mathcal{E}}$, $\mathcal{E}(b) \in \mathcal{Z}(\mathcal{C})$; donde $\mathcal{Z}(\mathcal{C}) = \{c \in \mathcal{C} : cb = bc, \forall b \in \mathcal{C}\}$ es el centro de $\mathcal{C}$. El siguiente resultado, de características técnicas tiene consecuencias interesantes

• Sean $\mathcal{A}$, $\mathcal{B}$ C^* -álgebras unitales y $\phi : \mathcal{A} \rightarrow \mathcal{B}$ una transformación positiva y unital. Supongamos que existe una esperanza condicional $\mathcal{E} : \mathcal{B} \rightarrow \mathcal{C}$, de $\mathcal{B}$ sobre la C^* -subálgebra $\mathcal{C}$. Entonces para cada función convexa $f : U \rightarrow \mathbb{R}$ definida en algún conjunto abierto y convexo $U \subseteq \mathbb{R}^n$ se verifica que

$$\mathcal{E}(g[f(\phi(a_1), \dots, \phi(a_n))]) \leq \mathcal{E}(g[\phi(f(a_1, \dots, a_n))]).$$

donde $(a_1, \dots, a_n) \subseteq \mathcal{A}_{sa}$ es una familia abeliana tal que $\prod_{i=1}^n \sigma(a_i) \subseteq U$, $\phi(a_1), \dots, \phi(a_n) \in \mathcal{B}^{\mathcal{E}}$ también forman una familia abeliana, y $g : I \rightarrow \mathbb{R}$ es una función convexa y creciente definida en algún intervalo abierto, con $\text{Im}(f) \subseteq I$. □

El siguiente resultado, que es una consecuencia del anterior, nos será indispensable en el estudio de teoremas de tipo Schur-Horn en factores II_1 . En lo que sigue consideramos el preorden de la submayorización que ha sido definido en la descripción de los contenidos del capítulo anterior (ver también la Definición 2.4.17 en los Preliminares).

• Sea $\mathcal{A}$ una C^* -álgebra unital, $(\mathcal{M}, \tau)$ un factor finito y $\phi : \mathcal{A} \rightarrow \mathcal{M}$ una transformación positiva y unital. Entonces para cada función convexa $f : U \rightarrow \mathbb{R}$ definida en algún conjunto abierto y convexo $U \subseteq \mathbb{R}^n$ se verifica que

$$f(\phi(a_1), \dots, \phi(a_n)) \prec_w \phi(f(a_1, \dots, a_n)).$$

donde $(a_1, \dots, a_n) \subseteq \mathcal{A}_{sa}$ es una familia abeliana tal que $\prod_{i=1}^n \sigma(a_i) \subseteq U$ y $\phi(a_1), \dots, \phi(a_n) \in \mathcal{M}$ también forman una familia abeliana.

□

De esta forma obtenemos el siguiente corolario. En este caso, la norma $\|\cdot\|_1$ está definida por $\|a\|_1 = \tau(|a|)$ que es la extensión natural de la norma $\|\cdot\|_1$ en el caso de $L(\mathcal{H})$.

• Sea $\mathcal{A} \subseteq \mathcal{M}$ una masa del factor $\text{II}_1(\mathcal{M}, \tau)$ y sea $E_{\mathcal{A}}$ la esperanza condicional que preserva la traza sobre $\mathcal{A}$. Sea $b \in \mathcal{M}$ un operador autoadjunto entonces

$$* \quad E_{\mathcal{A}}(b) \prec b.$$

$$* \quad \|E_{\mathcal{A}}(b)\|_1 \leq \|b\|_1.$$

□

Algunos casos particulares de las desigualdades de Jensen anteriores, obtenidas fijando transformaciones positivas y uniales son de interés e incluso generalizan algunos resultados de trabajos recientes.

Capítulo 5: Teoremas de tipo Schur-Horn

En la primera parte de este capítulo caracterizamos (la clausura en norma $\|\cdot\|_{\infty}$ de) el conjunto de posibles diagonales de un operador autoadjunto en $L(\mathcal{H})$; naturalmente, el caso de interés es cuando $\mathcal{H}$ es de dimensión infinita. Este problema está resuelto por el teorema de Schur-Horn en el caso de matrices. Es por esto que nuestro objetivo es hallar generalizaciones de este último. A tal fin introducimos las siguientes nociones.

- Dado $S \in L(\mathcal{H})$ un operador autoadjunto, definimos para cualquier $k \in \mathbb{N}$,

$$U_k(S) = \sup_{P \in \mathcal{P}_k} \text{tr}(SP) \quad \text{y} \quad L_k(S) = \inf_{P \in \mathcal{P}_k} \text{tr}(SP) = -U_k(-S),$$

donde $\mathcal{P}_k$ es el conjunto de proyecciones ortogonales de $L(\mathcal{H})$ tales que $\text{tr}(p) = k$

□

Para enunciar los siguientes resultados introducimos la siguiente notación (que también será usada en la descripción de los contenidos del siguiente capítulo): sea $\mathbb{M}$ un conjunto, sea $\mathcal{K}$ un espacio de Hilbert con $\dim \mathcal{K} = |\mathbb{M}|$ y sea $\mathcal{B} = \{e_n\}_{n \in \mathbb{M}}$ una base ortonormal de $\mathcal{K}$

- Para cualquier $\mathbf{a} = (a_n)_{n \in \mathbb{M}} \in \ell^{\infty}(\mathbb{M})$, denotemos por $M_{\mathcal{B}, \mathbf{a}} \in L(\mathcal{K})$ el operador diagonal dado por $M_{\mathcal{B}, \mathbf{a}} e_n = a_n e_n$, $n \in \mathbb{M}$. Cuando sea claro del contexto qué base se está usando abreviaremos $M_{\mathcal{B}, \mathbf{a}} = M_{\mathbf{a}}$.
- La compresión diagonal $\mathcal{C}_{\mathcal{B}} : L(\mathcal{K}) \rightarrow L(\mathcal{K})$ asociada a la base $\mathcal{B}$, está definida por $\mathcal{C}_{\mathcal{B}}(T) = M_{\mathcal{B}, \mathbf{a}}$, donde $\mathbf{a} = (\langle T e_n, e_n \rangle)_{n \in \mathbb{M}}$.

- Sea $\mathcal{B} = \{e_n\}_{n \in \mathbb{N}}$ una base ortonormal de un espacio de Hilbert $\mathcal{H}$. Si $\mathbf{a} \in \ell_{\mathbb{R}}^{\infty}(\mathbb{N})$ entonces, para cada $k \in \mathbb{N}$

$$U_k(M_{\mathcal{B}, \mathbf{a}}) = U_k(\mathbf{a}).$$

donde $U_k(\mathbf{a}) = \sup_{F \in \Phi_k} \sum_{i \in F} a_i$ y Φ_k es el conjunto formado por todos los subconjuntos de $\mathbb{N}$ de cardinal k .

□

Este último resultado indica que las definiciones de $U_k(S)$, $L_k(S)$ generalizan las sumas que aparecen en las ecuaciones (2.1) y (2.9) (ver la Observación 2.1.2 que establece la relación entre las nociones determinadas en las dos ecuaciones anteriores). Como veremos resultan ser conceptos adecuados para caracterizar el conjunto de posibles diagonales de un operador autoadjunto en el sentido de Neumann. Por otro lado, las siguientes propiedades están relacionadas con condiciones necesarias para el problema de la admisibilidad de marcos que estudiamos en el Capítulo 6.

Denotamos por $L_0(\mathcal{H})$ el $*$ -ideal cerrado de los operadores compactos y consideremos el cociente $\mathcal{A}(\mathcal{H}) = L(\mathcal{H})/L_0(\mathcal{H})$, que es una C^* -álgebra unital, conocida como el álgebra de Calkin. Dado $T \in L(\mathcal{H})$, el *espectro esencial* de T , notado $\sigma_e(T)$, es el espectro de la clase $T + L_0(\mathcal{H})$ en el álgebra $\mathcal{A}(\mathcal{H})$. La norma esencial $\|T\|_e = \inf\{\|T + K\| : K \in L_0(\mathcal{H})\}$ de T es la norma (cociente) de $T + L_0(\mathcal{H})$ en $\mathcal{A}(\mathcal{H})$. Notemos que $\sigma_e(T)$ es un subconjunto compacto de $\sigma(T)$. Si $S \in L(\mathcal{H})_h$ definimos

$$\alpha^+(S) = \max \sigma_e(S) = \|S\|_e \quad \text{y} \quad \alpha_-(S) = \min \sigma_e(S) ,$$

- Sea $S \in L(\mathcal{H})$ un operador autoadjunto. Entonces, para cada $k \in \mathbb{N}$,
- * $U_k(S) = U_k(S^+) + k \alpha^+(S)$
- * $L_k(S) = L_k(S_-) + k \alpha_-(S)$

En particular,

$$\lim_{k \rightarrow \infty} \frac{U_k(S)}{k} = \alpha^+(S) = \|S\|_e \quad \text{y} \quad \lim_{k \rightarrow \infty} \frac{L_k(S)}{k} = \alpha_-(S) .$$

□

A partir de los resultados anteriores obtenemos el siguiente teorema de tipo Schur-Horn en $L(\mathcal{H})$ para operadores $S \in L(\mathcal{H})$ autoadjuntos arbitrarios. Para enunciarlo definimos los siguientes conjuntos: sea $\mathcal{H}$ un espacio de Hilbert, $S \in L(\mathcal{H})$ y $\mathcal{B}$ una base ortonormal de $\mathcal{H}$. Entonces,

- $\mathcal{U}_{\mathcal{H}}(S) = \{U^* S U : U \in \mathcal{U}(\mathcal{H})\}$.
- $\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)] = \{\mathbf{c} \in \ell^\infty(\mathbb{N}) : M_{\mathcal{B}}, \mathbf{c} \in \mathcal{C}_{\mathcal{B}}(\mathcal{U}_{\mathcal{H}}(S))\}$.

En este contexto, el conjunto $\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]$ se interpreta como el conjunto de las posibles diagonales de las representaciones matriciales de un operador con respecto a todas las bases ortonormales de $\mathcal{H}$. El siguiente resultado caracteriza la clausura en norma infinito de este conjunto. Un problema interesante es hallar condiciones suficientes para que un vector en $\ell^\infty(\mathbb{N})$ sea una posible diagonal.

• Sea $S \in L(\mathcal{H})$ un operador autoadjunto y $\mathbf{c} \in \ell_{\mathbb{R}}^\infty(\mathbb{N})$. Luego, las siguientes condiciones son equivalentes:

- * $\mathbf{c} \in \overline{\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]}^{\|\cdot\|_\infty}$.
- * $U_k(S) \geq U_k(\mathbf{c})$ y $L_k(S) \leq L_k(\mathbf{c})$ para cada $k \in \mathbb{N}$.

En este caso,

$$\max \sigma_e(S) \geq \limsup \mathbf{c} \quad \text{y} \quad \min \sigma_e(S) \leq \liminf \mathbf{c} .$$

□

En el caso particular en que $S \in L(\mathcal{H})$ sea un operador autoadjunto de tipo traza (es decir $\text{tr}(|S|) < \infty$) entonces se tiene una caracterización de la clausura en norma $\|\cdot\|_1$ del conjunto de posibles diagonales de S .

• Sea $S \in L^1(\mathcal{H})$ un operador autoadjunto, y $\mathbf{b} \in \ell_{\mathbb{R}}^1(\mathbb{N})$. Entonces, los siguiente son equivalentes

* $\mathbf{b} \in \overline{\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]}^{\|\cdot\|_1}$.

* $U_k(S) \geq U_k(\mathbf{b})$, $L_k(S) \leq L_k(\mathbf{b})$ para cada $k \in \mathbb{N}$, y $\sum_{k=1}^{\infty} b_k = \text{tr } S$.

□

En la segunda parte de este capítulo consideramos una versión del teorema de Schur-Horn en el contexto de los factores II_1 . El desarrollo de este problema estuvo motivado por una conjetura de Kadison y Arveson al respecto. Nuestra versión difiere en algún sentido de la conjetura original.

Adelantamos una posible interpretación de nuestros resultados. En el caso del teorema de Schur-Horn anterior tenemos la siguiente situación: fijamos una base ortonormal $\mathcal{B} = \{e_k\}_{k \in \mathbb{N}}$ de $\mathcal{H}$, y consideramos la compresión a la diagonal $\mathcal{C}_{\mathcal{B}}$ determinada por $\mathcal{B}$, que es la esperanza condicional que conmuta con la traza usual de $L(\mathcal{H})$ (es decir $\text{tr}(A) = \text{tr}(\mathcal{C}_{\mathcal{B}}(A))$) y que proyecta sobre el “álgebra diagonal determinada por $\mathcal{B}$ ”. Esta álgebra diagonal tiene la siguiente propiedad: si $A \in L(\mathcal{H})$ es tal que $AD = DA$ para todo D en el álgebra diagonal determinada por $\mathcal{B}$ entonces A es un operador diagonal con respecto a la base $\mathcal{B}$. Esto implica que no hay álgebras abelianas en $L(\mathcal{H})$ que contengan estrictamente al álgebra diagonal determinada por $\mathcal{B}$. De aquí que esta álgebra diagonal es una subálgebra abeliana maximal de $L(\mathcal{H})$ o más brevemente una *masa* de $L(\mathcal{H})$.

En el caso de factores de tipo II_1 se conoce la existencia de masas en ellos con distintas propiedades: en particular, dada una masa $\mathcal{A} \subseteq \mathcal{M}$ de un factor II_1 $(\mathcal{M}, \tau)$ entonces existe una única esperanza condicional $E_{\mathcal{A}}$ que conmuta con la traza en el sentido anterior (que es un ejemplo especial de transformación doble estocástica en $(\mathcal{M}, \tau)$). Una cuestión natural es la de caracterizar el conjunto de posibles diagonales con respecto a esta álgebra de un operador autoadjunto a es decir el conjunto

$$E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(a)).$$

Este planteo del teorema de Schur-Horn es desarrollado en el trabajo de Arveson y Kadison [11] donde se propone estudiar (la clausura en norma de) el conjunto de posibles diagonales con respecto a la masa $\mathcal{A}$ en términos de mayorización (para una discusión más detallada ver el comienzo de la sección 5.2). Nuestro resultado al respecto de este problema es el siguiente

• Sea $(\mathcal{M}, \tau)$ un factor II_1 y sea $\mathcal{A} \subseteq \mathcal{M}$ una masa. Si $b \in \mathcal{M}^+$ y $a \in \mathcal{A}$ entonces $a \prec b$ si y solo si $a \in \overline{E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|_1}$, con lo cual tenemos

$$\overline{E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|_1} = \{a \in \mathcal{A} : a \prec b\}.$$

□

A partir de este resultado desarrollamos algunos corolarios relacionados con este problema. Para obtener este teorema de tipo Schur-Horn utilizamos el método de modelado de operadores en factores II_1 descrito previamente (en el resumen correspondiente al Capítulo 3). Además necesitamos algunos resultados de aproximación de funciones integrables por funciones simples en ciertos

espacios de medida. En lo que sigue describimos los resultados obtenidos con respecto a estos problemas.

Consideramos espacios de medida difusa (es decir, sin átomos) del tipo $([\alpha, \beta], \nu)$ donde $I = [\alpha, \beta] \subseteq \mathbb{R}$ es un intervalo compacto y ν es una medida de probabilidad. Además consideramos una función $h : I \rightarrow \mathbb{R}_{\geq 0}$ creciente y continua a izquierda. En lo que sigue consideramos valores singulares μ_f de funciones a valores reales (positivos) $f \in L^\infty(I, \nu)$ en tanto operadores autoadjuntos del álgebra de von Neumann $L^\infty(I, \nu)$. En realidad esta noción coincide con el concepto de reordenada de una función en un espacio de medida abstracto.

• Sea $P = \{x_1 < \dots < x_{2^n+1}\} \subseteq [\alpha, \beta]$ y sea $\{I_i\}_{i=1}^{2^n}$ una partición de $[\alpha, \beta]$ inducida por P , con $\nu(I_i) = \frac{1}{2^n}$ para $1 \leq i \leq 2^n$. Entonces

$$\int_{I_i} h(t) d\nu(t) = \int_{\frac{2^n-i}{2^n}}^{\frac{2^n-i+1}{2^n}} \mu_h(t) dt \quad \text{para } 1 \leq i \leq 2^n.$$

Sea $g : [\alpha, \beta] \rightarrow \mathbb{R}_{\geq 0}$ otra función creciente y continua a izquierda tal que $h \prec g$. Sea $\mathbf{h} = \mathbf{h}(n) \in \mathbb{R}^{2^n}$ (resp. $\mathbf{g} = \mathbf{g}(n) \in \mathbb{R}^{2^n}$) dado por

$$\mathbf{h}_i = 2^n \int_{I_i} h(t) d\nu(t), \quad 1 \leq i \leq 2^n$$

entonces $\mathbf{h} = \mathbf{h}^\uparrow$ (resp. $\mathbf{g} = \mathbf{g}^\uparrow$) y $\mathbf{h} \prec \mathbf{g}$.

□

En este caso de una medida de probabilidad regular ν con $\text{sop}(\nu) \subseteq [\alpha, \beta]$ y difusa definimos inductivamente una sucesión de particiones de $[\alpha, \beta]$ como sigue: para $n = 0$ consideramos $P_0 = \{\alpha, \beta\}$ y para $n \in \mathbb{N}$ sea $P_n \subseteq P_{n+1} = \{x_i\}_{i=0}^{2^{n+1}}$ tal que, si $\{I_i^{(n+1)}\}_{i=1}^{2^{n+1}}$ es la partición de $[\alpha, \beta]$ en subintervalos inducida por P_{n+1} entonces

$$\nu(I_i^{(n+1)}) = 1/2^{n+1} \quad \text{para } 1 \leq i \leq 2^{n+1}.$$

Notemos que siempre podemos encontrar una tal partición, pues la función de acumulación $g(t) = \nu([\alpha, t])$ es continua; pero no es necesariamente única (esto se debe al hecho de que la función de acumulación puede no ser estrictamente creciente) Decimos que $\{I_i^{(n)}\}_{i=1}^{2^n}$, $n \in \mathbb{N}$ es una partición ν -diádica de $[\alpha, \beta]$.

Sea $\{I_i^{(n)}\}_{i=1}^{2^n}$, $n \in \mathbb{N}$ una partición ν -diádica de $[\alpha, \beta]$. Para cada $n \in \mathbb{N}$ y cada $f \in L^1(\nu)$ definimos la familia de aproximaciones discretas

$$E_n(f)(x) = \sum_{i=1}^{2^n} \left(2^n \int_{I_i^{(n)}} f d\nu \right) \chi_{I_i^{(n)}}(x).$$

Un problema natural que surge en este contexto es el de aproximación de funciones por funciones simples (escalonadas), con respecto a alguna topología. Sucede que si queremos utilizar como base para las funciones simples, las funciones características de particiones ν diádicas entonces, lamentablemente, no podemos hallar en general aproximaciones uniformes ni siquiera de

funciones continuas en I . En lo que sigue consideramos el operador maximal asociado a la sucesión de aproximaciones discretas

$$E^*(f)(x) = \sup_{n \in \mathbb{N}} |E_n(f)(x)|.$$

• Sea $\{I_i^{(n)}\}_{i=1}^{2^n}$, $n \in \mathbb{N}$ una partición ν -diádica de $[\alpha, \beta]$ y sea E_n , $n \in \mathbb{N}$ la familia asociada de aproximaciones discretas. Entonces

* El operador maximal E^* es débil $(1, 1)$.

* Si $f \in L^1(\nu)$ entonces, $\lim_{n \rightarrow \infty} E_n(f)(x) = f(x)$ ν -a.e.

En particular, podemos obtener una descomposición de tipo Calderón-Zygmund: sea $f \in L^1(\nu)$ una función no negativa y $\lambda > 0$. Entonces existe una sucesión ν -diádica de conjuntos $\{I_j\}_{j \in J}$ tal que

* $f(x) \leq \lambda$, ν -a.e. para $x \notin \bigcup_{j \in J} I_j$.

* $\nu(\bigcup_{j \in J} I_j) \leq \frac{\|f\|_1}{\lambda}$.

* $\lambda < \frac{1}{\nu(I_j)} \int_{I_j} f \, d\nu \leq 2\lambda$.

□

Estamos interesados en la siguiente consecuencia de los resultados anteriores

• Sea $\{I_i^{(n)}\}_{i=1}^{2^n}$, $n \in \mathbb{N}$ una partición ν -diádica de $[\alpha, \beta]$ y sea E_n , $n \in \mathbb{N}$ la sucesión de aproximaciones discretas asociada a esta partición ν -diádica. Entonces, para cada función ν -esencialmente acotada f en $[\alpha, \beta]$ tenemos

$$\lim_{n \rightarrow \infty} \|f - E_n(f)\|_1 = 0$$

□

Los resultados anteriores, junto con algunas construcciones dentro de los factores de tipo II_1 nos permiten probar nuestra versión del teorema de Schur-Horn antes descrita.

Capítulo 6: Marcos con operador de marco y normas prefijadas

Sea $\mathbb{M} = \mathbb{N}$ ó $\mathbb{M} = \{1, 2, \dots, m\} := \mathbb{I}_m$, para algún $m \in \mathbb{N}$. Una sucesión $\{f_k\}_{k \in \mathbb{M}}$ es un *marco* para $\mathcal{H}$ (ver la sección 2.2.1 para más detalles) si existen constantes positivas $A, B > 0$ tales que

$$A\|x\|^2 \leq \sum_{k \in \mathbb{M}} |\langle x, f_k \rangle|^2 \leq B\|x\|^2, \quad \text{para cada } x \in \mathcal{H}.$$

Sea $\mathcal{F} = \{f_k\}_{k \in \mathbb{M}}$ un marco para $\mathcal{H}$. El operador

$$S : \mathcal{H} \rightarrow \mathcal{H}, \quad \text{dado por } S(x) = \sum_{k \in \mathbb{M}} \langle x, f_k \rangle f_k, \quad x \in \mathcal{H}.$$

es llamado el *operador de marco* de $\mathcal{F}$. Éste es acotado, positivo e inversible (usamos la notación $S \in \mathcal{GL}(\mathcal{H})^+$).

Un par ordenado $(S, \mathbf{c})$ formado por un operador positivo e inversible $S \in \mathcal{GL}(\mathcal{H})^+ \subseteq L(\mathcal{H})$ y una sucesión uniformemente acotada de números positivos $\mathbf{c} \in \ell^\infty(\mathbb{M})^+$, se dice *admissible* si existe un marco $\mathcal{F} = \{f_k\}_{k \in \mathbb{M}}$ para $\mathcal{H}$ tal que

- $\mathcal{F}$ tiene operador de marco S ,
- $\|f_k\|^2 = c_k$ para cada $k \in \mathbb{M}$.

En este caso, decimos que $\mathcal{F}$ es un $(S, \mathbf{c})$ -marco. Denotamos por $F(S, \mathbf{c})$ el conjunto de todos los $(S, \mathbf{c})$ -marcos para $\mathcal{H}$. Luego, el par $(S, \mathbf{c})$ es admisible si $F(S, \mathbf{c}) \neq \emptyset$.

El problema de hallar condiciones bajo las cuales un par $(S, \mathbf{c})$ es admisible ha sido considerado por numerosos investigadores de la teoría de marcos en espacios de Hilbert por sus implicancias para las aplicaciones de esta teoría. En primer lugar, obtenemos una caracterización de la admisibilidad

- Sea $\mathbf{c} \in \ell^\infty(\mathbb{M})^+$ y sea $S \in \mathcal{G}l(\mathcal{H})^+$. Entonces las siguientes condiciones son equivalentes:
- * El par $(S, \mathbf{c})$ es admisible.
- * Existe una sucesión de vectores unitarios $\{y_k\}_{k \in \mathbb{M}}$ en $\mathcal{H}$ tales que

$$S = \sum_{k \in \mathbb{M}} c_k y_k \otimes y_k ,$$

donde, si $\mathbb{M} = \mathbb{N}$, la suma converge en la topología fuerte de operadores.

- * Existe una extensión $\mathcal{K} = \mathcal{H} \oplus \mathcal{H}_d$ de $\mathcal{H}$ tal que, si denotamos

$$S_1 = \begin{pmatrix} S & 0 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} \mathcal{H} \\ \mathcal{H}_d \end{pmatrix} \in L(\mathcal{K})^+ , \quad \text{entonces} \quad \mathbf{c} \in \mathcal{C}[\mathcal{U}_{\mathcal{K}}(S_1)] .$$

En este caso, existe un marco $\mathcal{F} \in F(S, \mathbf{c})$.

□

Notemos que la equivalencia de los ítems 1 y 3 del resultado anterior muestran que teoremas de tipo Schur-Horn en $L(\mathcal{H})$ son relevantes en el estudio del problema de admisibilidad. El siguiente resultado da una caracterización exacta de la admisibilidad para marcos en espacios de Hilbert finito dimensionales. En lo que sigue vamos a utilizar la notación introducida en la descripción de la primera parte del capítulo anterior.

- Sea $\mathbf{c} \in \ell^\infty(\mathbb{N})^+$ y $S \in \mathcal{G}l_n(\mathbb{C})^+$ una matriz positiva inversible, con autovalores $b_1 \geq b_2 \geq \dots \geq b_n > 0$. Las siguientes condiciones son equivalentes:

- * el par $(S, \mathbf{c})$ es admisible.
- * $\sum_{i=1}^k b_i \geq U_k(\mathbf{c})$, para cada $1 \leq k \leq n-1$, y $\sum_{i=1}^n b_i = \sum_{i \in \mathbb{N}} c_i$.

□

En el caso de dimensión infinita, parece no haber un conjunto de condiciones escritas en términos de la mayorización que caractericen de forma exacta la admisibilidad. Sin embargo obtenemos las siguientes condiciones necesarias. En adelante, los espacios de Hilbert considerados son de dimensión infinita.

- Sea $S \in \mathcal{G}l(\mathcal{H})^+$ un operador positivo e inversible y $\mathbf{c} \in \ell^\infty(\mathbb{N})^+$. Si el par $(S, \mathbf{c})$ es admisible entonces,

$$\sum_{i \in \mathbb{N}} c_i = \infty \quad \text{y} \quad U_k(S) \geq U_k(\mathbf{c}), \quad \text{para cada } k \in \mathbb{N}.$$

En particular, se tiene la condición necesaria

$$\limsup \mathbf{c} \leq \|S\|_e.$$

□

El siguiente resultado provee de condiciones suficientes para la admisibilidad. Estas condiciones se obtienen de reforzar las anteriores condiciones necesarias. En su enunciado usamos la notación $P_2(S) = E[\|S\|_e, \|S\|]$, donde E es la medida espectral de $S \in L(\mathcal{H})^+$ y $\|S\|_e$ denota la norma esencial (es decir la norma cociente de $L(\mathcal{H})$ por el ideal cerrado de los compactos) del operador autoadjunto S .

• Sea $S \in \mathcal{G}l(\mathcal{H})^+$ un operador positivo e inversible y $\mathbf{c} \in l^\infty(\mathbb{N})^+$, tal que $\sum_{i \in \mathbb{N}} c_i = \infty$. Supongamos que se satisface alguna de las siguientes condiciones:

- * a) $\text{tr } P_2(S) = \infty$,
- b) $U_k(S) \geq U_k(\mathbf{c})$ para cada $k \in \mathbb{N}$, y
- c) $\|S\|_e > \limsup(\mathbf{c})$.
- * a) $\text{tr } P_2(S) = r \in \mathbb{N}$,
- b) $U_k(S) \geq U_k(\mathbf{c})$ para $1 \leq k \leq r$,
- c) $U_k(S) > U_k(\mathbf{c})$, para $k > r$, y
- d) $\|S\|_e > \limsup(\mathbf{c})$.

Entonces, el par $(S, \mathbf{c})$ es admisible.

□

Este capítulo termina con una serie de ejemplos que muestran que las condiciones anteriores (según sus respectivos casos) no pueden relajarse más en general, y son de alguna manera condiciones óptimas para asegurar la admisibilidad en el caso general de marcos en espacios de Hilbert infinito dimensionales.

Capítulo 7: MayORIZACIÓN DE MATRICES

Decimos que una matriz de entradas positivas $A \in M_n(\mathbb{R})$ es estocástica por filas, notado $A \in RS(n)$, si las sumas de las entradas de cada una de sus filas es 1. Una matriz de entradas positivas tal que $A \in RS(n)$ y $A^t \in RS(n)$ es llamada doble estocástica y lo notamos $A \in DS(n)$. Para las definiciones básicas referimos al lector a las secciones 2.1.1 y 2.1.2 sobre mayorización de vectores y matrices no negativas.

Dadas $X, Y \in M_{n,m} := M_{n,m}(\mathbb{R})$ matrices rectangulares $n \times m$ a coeficientes reales, consideramos las siguientes nociones de mayorización de matrices:

- Y está mayorizada fuertemente por X , notado $X \succ_f Y$, si existe $D \in DS(n)$ tal que $DX = Y$.
- Y está mayorizada direccionalmente por X , notado $X \succ Y$, si para todo $v \in \mathbb{R}^m$, Xv mayoriza como vector a Yv .

Hemos considerado el siguiente concepto

- Dadas $X, Y \in M_{n,m}$ decimos que Y está *mayorizada débilmente* por X , notado $X \succ_d Y$, si existe $A \in RS(n)$ tal que $AX = Y$. □

Las siguientes propiedades con consecuencia de las definiciones. En lo que sigue $R(X) \subseteq \mathbb{R}^m$ denota el conjunto formado por los vectores fila de $X \in M_{n,m}$.

- Sean $X, Y \in M_{n,m}$ entonces,
- * $X \succ_d Y$ si y solo si $R(Y) \subseteq \text{conv}(R(X))$;
- * Si $X \succ Y$ entonces $X \succ_d Y$.

□

Como consecuencia de estas propiedades deducimos

- Sean $X, Y \in M_{n,m}$. $X \succ_d Y$ si y solo si

$$\max_{1 \leq i \leq n} f(X_i) \geq \max_{1 \leq i \leq n} f(Y_i)$$

para cada función convexa $f : V \rightarrow \mathbb{R}$ donde $V \subseteq \mathbb{R}^m$ es un conjunto convexo que contiene $R(X) \cup R(Y)$. Más aún, si consideramos las funciones lineales $\phi_z : \mathbb{R}^m \rightarrow \mathbb{R}$, $z \in \mathbb{R}^m$ definidas por $\phi_z(x) = \langle x, z \rangle$, $X \succ_d Y$ si y solo si

$$\max_{1 \leq i \leq n} \phi_z(X_i) \geq \max_{1 \leq i \leq n} \phi_z(Y_i)$$

para cada $z \in \mathbb{R}^m$. □

El siguiente resultado caracteriza la mayorización direccional entre matrices en $M_{n,m}$, en términos de $\lfloor \frac{n}{2} \rfloor + 1$ politopos, donde $\lfloor r \rfloor$ es la parte entera de $r \in \mathbb{R}$.

- Sean $X, Y \in M_{n,m}$. $X \succ Y$ si y solo si, para $k = 1, \dots, \lfloor \frac{n}{2} \rfloor$ y $k = n$, el conjunto de promedios de k filas diferentes de Y está incluido en la cápsula convexa del conjunto de promedios de k filas diferentes de X . □

Este método alternativo de verificación de la mayorización direccional es muy útil desde el punto de vista computacional (la verificación por definición puede ser un problema de cálculo complejo).

Con respecto al problema de hallar condiciones bajo las cuales las mayorizaciones direccional o débil impliquen la mayorización fuerte obtenemos los siguientes resultados

- Sean $X, Y \in M_{n,m}$ tales que $X \succ Y$.
- * Si $\text{conv}(R(X))$ tiene solo dos puntos extremos entonces $X \succ_f Y$.
- * Si $1 \leq n \leq 3$ entonces, $X \succ Y$ implica que $X \succ_f Y$.

□

Con respecto a este mismo problema obtenemos las siguientes condiciones generales. En lo que sigue, dada $X \in M_{n,m}$ denotamos por $[X, e] \in M_{n, (m+1)}$ la matriz cuyas primeras m columnas son

iguales a las columnas de X y su última columna es el vector $e = (1, \dots, 1) \in \mathbb{R}^n$. En lo que sigue, $A^\dagger$ denota la pseudo-inversa de Moore-Penrose de una matriz A .

• Sean $X, Y \in M_{n,m}$ y supongamos que $[Y, e][X, e]^\dagger$ tiene entradas no negativas. Si $X \succ_d Y$ y $e^t X = e^t Y$ entonces $X \succ_f Y$.

□

Este resultado tiene como consecuencias

• Sean $X, Y \in M_{n,m}$

* Supongamos que $\text{ran}([X, e]) = \mathbb{R}^n$. Si $X \succ_d Y$ y $e^t X = e^t Y$ entonces $X \succ_f Y$.

* Supongamos que las filas de X son los vértices de un simplex, es decir afínmente linealmente independientes. Si $X \succ_d Y$ y $e^t X = e^t Y$ entonces $X \succ_f Y$.

□

Además hacemos un breve estudio de las relaciones de equivalencia asociadas a estos preórdenes.

En una segunda parte del capítulo, desarrollamos una serie de conceptos nuevos, las llamadas mayorizaciones conjuntas de familias abelianas. Estos desarrollos son mas bien elementales, pero sirven como preparación para el estudio de la mayorización conjunta (fuerte) entre familias abelianas de operadores autoadjuntos en factores de tipo II_1 que consideramos posteriormente.

Una *familia abeliana* es una familia ordenada $(a_i)_{i=1,\dots,m}$ de matrices autoadjuntas en $M_n(\mathbb{C})$ tales que $a_i a_j = a_j a_i$, $i, j = 1, \dots, m$. Para introducir las mayorizaciones conjuntas entre tales familias consideramos los siguientes hechos elementales: si $(a_i)_{i=1,\dots,m}$ y $(b_i)_{i=1,\dots,m}$ son dos familias abelianas en $M_n(\mathbb{C})$ entonces existen matrices unitarias $U, V \in M_n(\mathbb{C})$ tales que

$$U^* a_i U = D_{\lambda(a_i)}, \quad V^* b_i V = D_{\lambda(b_i)}, \quad i = 1, \dots, m,$$

donde D_x denota la matriz diagonal con diagonal principal $x \in \mathbb{R}^n$. En este caso $\lambda(a_i)$ es el vector de autovalores correspondientes a a_i , contados con multiplicidad, en algún orden dependiendo de U . Consideremos las matrices $A, B \in M_{n,m}$ cuyas columnas $C_i(A), C_i(B) \in \mathbb{R}^n$ están dadas por

$$C_i(A) = \lambda(a_i), \quad C_i(B) = \lambda(b_i), \quad i = 1, \dots, m.$$

• Sean $(a_i)_{i=1,\dots,m}, (b_i)_{i=1,\dots,m} \subseteq M_n(\mathbb{C})$ dos familias abelianas y sean $A, B \in M_{n,m}$ definidas como antes. Decimos que la familia $(a_i)_{i=1,\dots,m}$ mayoriza conjuntamente de forma débil (respectivamente mayoriza conjuntamente de forma fuerte, mayoriza conjuntamente de forma direccional) a la familia $(b_i)_{i=1,\dots,m}$ y notamos

$$(a_i)_{i=1,\dots,m} \succ_d (b_i)_{i=1,\dots,m}$$

(respectivamente $(a_i)_{i=1,\dots,m} \succ_f (b_i)_{i=1,\dots,m}$, $(a_i)_{i=1,\dots,m} \succ (b_i)_{i=1,\dots,m}$) si $A \succ_d B$ (respectivamente $A \succ_f B$, $A \succ B$).

□

Obtenemos la siguiente caracterización de las mayorizaciones conjuntas de familias abelianas. En lo que sigue, dada una familia abeliana $(a_i)_{i=1}^n$ entonces $C^*(a_1, \dots, a_n)$ denota la C^* -álgebra (abeliana) generada por esta familia.

• Sean $(a_i)_{i=1,\dots,m}$ y $(b_i)_{i=1,\dots,m}$ dos familias abelianas en $M_n(\mathbb{C})$. Entonces

- * $(a_i) \succ_d (b_i)$ si y solo si existe una transformación positiva y unital

$$T : C^*(a_1, \dots, a_m) \rightarrow C^*(b_1, \dots, b_m)$$

tal que $T(a_i) = b_i$ para cada $i = 1, \dots, m$.

- * $(a_i) \succ (b_i)$ si y solo si, para cada $k = 1, \dots, [\frac{n}{2}]$ y $k = n$ tenemos que

$$(\log[\bigwedge^k \exp(a_i)])_{i=1, \dots, m} \succ_d (\log[\bigwedge^k \exp(b_i)])_{i=1, \dots, m}.$$

- * $(a_i) \succ_f (b_i)$ si y solo si existe una transformación positiva, unital y que preserva la traza

$$T : C^*(a_1, \dots, a_m) \rightarrow C^*(b_1, \dots, b_m)$$

tal que $T(a_i) = b_i$ para cada $i = 1, \dots, m$.

□

Los siguientes resultados son traducciones de las correspondientes propiedades de las mayorizaciones de matrices en este nuevo contexto.

- Sean $(a_i)_{i=1, \dots, m}$ y $(b_i)_{i=1, \dots, m}$ dos familias abelianas. Entonces
- * $(a_i) \succ_d (b_i)$ si y solo si para cada función convexa $f : V \rightarrow \mathbb{R}$ se tiene

$$\|f(a_1, \dots, a_m)\| \geq \|f(b_1, \dots, b_m)\|.$$

- * $(a_i) \succ (b_i)$ si y solo si para $k = 1, \dots, [\frac{n}{2}]$ y $k = n$ tenemos que

$$\left\| f \left(\log[\bigwedge^k \exp(a_1)], \dots, \log[\bigwedge^k \exp(a_m)] \right) \right\| \geq \left\| f \left(\log[\bigwedge^k \exp(b_1)], \dots, \log[\bigwedge^k \exp(b_m)] \right) \right\|$$

para cada función convexa $f : V \rightarrow \mathbb{R}$.

- * $(a_i) \succ_f (b_i)$ si y solo si, para cada función convexa $f : V \rightarrow \mathbb{R}$ se verifica que

$$\text{tr } f(a_1, \dots, a_m) \geq \text{tr } f(b_1, \dots, b_m),$$

donde $V \subseteq \mathbb{R}^m$ es un conjunto convexo que contiene $\sigma(a_1, \dots, a_m) \cup \sigma(b_1, \dots, b_m)$.

□

Este capítulo termina con el estudio de las relaciones de equivalencias asociadas a las mayorizaciones conjuntas entre familias abelianas.

Capítulo 8: Mayorización conjunta en factores factores II_1

Los resultados descritos en este capítulo se desarrollan dentro de los denominados factores de tipo II_1 , es decir álgebras de von Neumann de centro trivial y dotadas de un estado fiel normal finito y tracial de forma que si $\alpha \in [0, 1]$ existe una proyección $p \in \mathcal{M}$ tal que $\tau(p) = \alpha$. Los factores de tipo II_1 son también llamados factores finitos continuos. Una transformación doble estocástica (ver la sección 2.4.4 y el Teorema 2.4.20) en un factor finito $(\mathcal{M}, \tau)$ es una transformación lineal $T : \mathcal{M} \rightarrow \mathcal{M}$ positiva ($T(a) \geq 0$ siempre que $a \geq 0$) y unital ($T(1) = 1$). Al conjunto de tales transformaciones lo denotamos por $DS(\mathcal{M}, \tau)$.

Es conocido el hecho de que estas transformaciones (que son las generalizaciones de las matrices doble estocásticas en este contexto) están íntimamente relacionadas con la mayorización (ver el Teorema 2.4.20). Para poder estudiar la mayorización conjunta de familias abelianas de operadores autoadjuntos en factores de tipo II_1 obtenemos el siguiente teorema de representación de ciertas restricciones de tales transformaciones. De aquí que esta representación describe a las transformaciones doble estocásticas localmente. En su enunciado notamos por $\mathcal{D}(\mathcal{M})$ el semigrupo convexo dado por

$$\mathcal{D}(\mathcal{M}) = \text{conv}\{Ad u : u \in \mathcal{U}(\mathcal{M})\}$$

donde $Ad u(a) = u^* a u$.

• Sean $\mathcal{A}, \mathcal{B} \subseteq \mathcal{M}$ C^* -subálgebras abelianas y separables y sea $T \in DS(\mathcal{M})$. Si $\mathcal{S} = T^{-1}(\mathcal{A}) \cap \mathcal{B}$ entonces, existe una sucesión $(\rho_r)_{r \in \mathbb{N}} \subseteq \mathcal{D}(\mathcal{M})$ tal que para cada $b \in \mathcal{S}$ se tiene

$$\lim_{r \rightarrow \infty} \|T(b) - \rho_r(b)\| = 0.$$

□

Estamos interesados en la siguiente consecuencia del resultado anterior. En lo que sigue, una familia abeliana en $\mathcal{M}_{sa}$ es una n -upla de operadores autoadjuntos de $\mathcal{M}$ tales que conmutan dos a dos.

• Sea $T \in DS(\mathcal{M})$ y sean $(a_i)_{i=1}^n, (b_i)_{i=1}^n \subseteq \mathcal{M}_{sa}$ familias abelianas tales que $T(b_i) = a_i$ para $1 \leq i \leq n$. Entonces existe una sucesión $(\rho_r)_{r \in \mathbb{N}} \subseteq \mathcal{D}(\mathcal{M})$ tal que para $1 \leq i \leq n$

$$\lim_{r \rightarrow \infty} \|a_i - \rho_r(b_i)\| = 0.$$

□

Es un hecho conocido (ver Teorema 2.4.5) que toda C^* -subálgebra de $\mathcal{A} \subseteq \mathcal{M}$ induce una medida espectral E sobre el espectro $\Gamma(\mathcal{A})$ a valores proyecciones de $\mathcal{M}$ que permite describir los elementos de $\mathcal{A}$. Así, decimos que el álgebra es difusa si $E(\{x\}) = 0$ para todo $x \in \Gamma(\mathcal{A})$.

Para obtener el teorema de representación local anterior, consideramos una reducción al caso en que la C^* -subálgebras abelianas y separables $\mathcal{A}, \mathcal{B} \subseteq \mathcal{M}$ son además difusas. El siguiente resultado muestra que tal reducción es siempre posible.

• Sea $\mathcal{A} \subseteq \mathcal{M}$ una C^* -subálgebra abeliana y separable. Entonces existe $a \in \mathcal{M}_{sa}$ tal que $C^*(\mathcal{A}, a)$ es abeliana y difusa.

□

El resultado anterior está relacionado con los resultados de refinamiento para resoluciones espectrales desarrollados en el Capítulo 3, sin embargo son resultados diferentes.

La ventaja de poder reducir el estudio de la forma local de las transformaciones doble estocásticas al caso de álgebras abelianas, separables y difusas radica en el siguiente resultado de aproximación. En lo que sigue una partición de la unidad es una familia finita de proyecciones ortogonales dos a dos $\{p_i\}_{i=1}^n$ y tales que $\sum_{i=1}^n p_i = 1$.

• Sea $\mathcal{B} \subset \mathcal{M}$ una C^* -subálgebra abeliana, separable y difusa. Entonces existe un conjunto no acotado $\mathbb{M} \subseteq \mathbb{N}$ tal que para cada $m \in \mathbb{M}$ existen $k = k(m)$ particiones de la unidad $\{q_i^{t,m}\}_{i=1}^m \subseteq \mathcal{B}' \cap \mathcal{M}$, $1 \leq t \leq k$, con $\tau(q_i^{t,m}) = 1/m$ ($1 \leq i \leq m$, $1 \leq t \leq k$), y tales que para cada $b \in \mathcal{B}$, si notamos $\beta_i^{t,m} = m \tau(b q_i^{t,m})$, entonces

$$\lim_{m \rightarrow \infty} \left\| b - \frac{1}{k} \sum_{t=1}^k \left(\sum_{i=1}^m \beta_i^{t,m} q_i^{t,m} \right) \right\| = 0.$$

□

Además necesitamos extender un resultado conocido de Dixmier con respecto a la esperanza condicional al centro de un álgebra finita de von Neumann al siguientes caso

• Sea $\mathcal{B} \subset \mathcal{M}$ una C^* -subálgebra separable y sea $\{p_i\}_{i=1}^m \subseteq \mathcal{B}' \cap \mathcal{M}$ una partición de la unidad. Entonces existe una sucesión $\{\rho_i\}_{i \in \mathbb{N}} \subset \mathcal{D}(\mathcal{M})$ tal que para cada $b \in \mathcal{B}$, si notamos $\beta_i(b) = \tau(b p_i) / \tau(p_i)$, entonces

$$\lim_{j \rightarrow \infty} \left\| \rho_j(b) - \sum_{i=1}^m \beta_i(b) p_i \right\| = 0.$$

□

Las restricciones anteriores con respecto a las trazas de las proyecciones de las particiones de la unidad son una condición necesaria para poder utilizar el siguiente teorema de tipo Birkhoff

• Sean $\{p_i\}_{i=1}^m, \{q_i\}_{i=1}^m \subseteq \mathcal{M}$ particiones de la unidad tales que $\tau(p_i) = \tau(q_i) = \frac{1}{m}$, y sea $T \in DS(\mathcal{M})$. Entonces existe $\rho \in \mathcal{D}(\mathcal{M})$ tal que si $\beta_1, \dots, \beta_m \in \mathbb{R}$ y $\alpha_i = m \sum_{j=1}^m \beta_j \tau(T(q_j) p_i)$ para $1 \leq i \leq m$, tenemos que

$$\sum_{i=1}^m \alpha_i p_i = \rho \left(\sum_{i=1}^m \beta_i q_i \right).$$

□

Una vez obtenidos estos resultados, los aplicamos en la segunda parte del trabajo al desarrollo de la mayorización conjunta de familias abelianas es decir. A tal fin consideramos las siguientes nociones

• Sea $(X, \mu_X), (Y, \mu_Y)$ dos espacios de probabilidad. Una transformación lineal positiva y unital $\nu : L^\infty(Y, \mu_Y) \rightarrow L^\infty(X, \mu_X)$ es un núcleo doble estocástico si $\int_X \nu(1_\Delta) d\mu_X = \mu_Y(\Delta)$, para cada conjunto μ_Y -medible $\Delta \subseteq Y$.

Sean $\bar{a} = (a_i)_{i=1}^n, \bar{b} = (b_i)_{i=1}^n$ dos familias abelianas en $\mathcal{M}_{sa}$. Decimos que $\bar{a}$ está conjuntamente mayorizado por $\bar{b}$, notado $\bar{a} \prec \bar{b}$, si existe un núcleo doble estocástico $\nu : L^\infty(\sigma(\bar{b}), \mu_{\bar{b}}) \rightarrow L^\infty(\sigma(\bar{a}), \mu_{\bar{a}})$ tal que $\nu(\pi_i) = \pi_i$, para cada $1 \leq i \leq n$.

□

Si $(x_1, \dots, x_n)$ es una familia finita en $\mathcal{M}$, sea $\mathcal{U}_{\mathcal{M}}(x_1, \dots, x_n)$ la órbita unitaria conjunta de la familia con respecto al grupo de operadores unitarios $\mathcal{U}_{\mathcal{M}}$ de $\mathcal{M}$, i.e.

$$\mathcal{U}_{\mathcal{M}}(x_1, \dots, x_n) = \{(u^* x_1 u, \dots, u^* x_n u) : u \in \mathcal{U}_{\mathcal{M}}\} \subseteq \mathcal{M}^n.$$

Consideramos también la cápsula convexa de la órbita unitaria de la familia $(x_i)_{i=1}^n$,

$$\text{conv}(\mathcal{U}_{\mathcal{M}}(x_i)_{i=1}^n) = \{(\rho(x_i))_{i=1}^n, \rho \in \mathcal{D}\} \subseteq \mathcal{M}^n.$$

Denotamos por $\overline{\text{conv}}(\mathcal{U}_{\mathcal{M}}(x_i)_{i=1}^n)$, $\overline{\text{conv}}^w(\mathcal{U}_{\mathcal{M}}(x_i)_{i=1}^n)$ y $\overline{\text{conv}}^1(\mathcal{U}_{\mathcal{M}}(x_i)_{i=1}^n)$ las respectivas clausuras entrada a entrada en la topología de la norma, en la topología débil de operadores, y en la topología L^1 inducida por τ .

El siguiente resultado describe varias caracterizaciones de la mayorización conjunta, y generaliza en el caso de factores II_1 el Teorema 2.4.20.

• Sean $\bar{a} = (a_i)_{i=1}^n$, $\bar{b} = (b_i)_{i=1}^n$ dos familias abelianas en $\mathcal{M}_{sa}$. Entonces los siguientes son equivalentes:

- * $\bar{a}$ está conjuntamente mayorizado por $\bar{b}$.
- * $\bar{a} \in \overline{\text{conv}}(\mathcal{U}_{\mathcal{M}}(\bar{b}))$.
- * $\bar{a} \in \overline{\text{conv}}^1(\mathcal{U}_{\mathcal{M}}(\bar{b}))$.
- * $\bar{a} \in \overline{\text{conv}}^w(\mathcal{U}_{\mathcal{M}}(\bar{b}))$.
- * $\mu_{\bar{a}} \prec \mu_{\bar{b}}$.
- * Existe una transformación completamente positiva $T \in DS(\mathcal{M})$ tal que $a_i = T(b_i)$, $1 \leq i \leq n$.
- * Existe $T \in DS(\mathcal{M})$ tal que $a_i = T(b_i)$, $1 \leq i \leq n$.
- * $\tau(f(a_1, \dots, a_n)) \leq \tau(f(b_1, \dots, b_n))$ para cada función continua y convexa $f : \mathbb{R}^n \rightarrow \mathbb{R}$.

□

Dadas familias $\bar{a} = (a_i)_{i=1}^n$, $\bar{b} = (b_i)_{i=1}^n \subseteq \mathcal{M}$, decimos que $\bar{a}$ y $\bar{b}$ son conjuntamente aproximadamente unitariamente equivalentes en $\mathcal{M}$ si $\bar{a} \in \overline{\text{conv}}(\mathcal{U}_{\mathcal{M}}(\bar{b}))$ es decir, si existe una sucesión de operadores unitarios $(u_n)_{n \in \mathbb{N}} \subseteq \mathcal{M}$ tal que $\lim_{n \rightarrow \infty} \|u_n b_i u_n^* - a_i\| = 0$ para cada $1 \leq i \leq n$. El siguiente resultado caracteriza la relación de equivalencia inducida por el preorden de la mayorización conjunta entre familias abelianas

• Sean $\bar{a} = (a_i)_{i=1}^n$ y $\bar{b} = (b_i)_{i=1}^n \subset \mathcal{M}_{sa}$ familias abelianas. Entonces los siguientes son equivalentes:

- * $\bar{a}$ y $\bar{b}$ son conjuntamente aproximadamente unitariamente equivalentes en $\mathcal{M}$.
- * $\bar{a} \prec \bar{b}$ y $\bar{b} \prec \bar{a}$
- * $\mu_{\bar{a}} = \mu_{\bar{b}}$
- * $\tau(f(a_1, \dots, a_n)) = \tau(f(b_1, \dots, b_n))$ para cada función continua y convexa $f : \mathbb{R}^n \rightarrow \mathbb{R}$.

* $\tau(f(a_1, \dots, a_n)) = \tau(f(b_1, \dots, b_n))$ para cada función continua $f : \mathbb{R}^n \rightarrow \mathbb{R}$.

□

Finalizamos nuestro estudio de la mayorización conjunta estableciendo una comparación entre ésta y un concepto que puede considerarse implícito en el desarrollo de una noción relacionada, desarrollado por P. Alberti y A. Uhlman.

Sea $\mathcal{A}$ una C^* -álgebra abeliana y sea $\mathcal{A}^*$ su espacio dual. En este contexto, una función lineal $V : \mathcal{A}^* \rightarrow \mathcal{A}^*$ es una función estocástica si es positiva y $V(\psi)(1) = \psi(1)$ para cada $\psi \in \mathcal{A}^*$. El conjunto de funciones estocásticas de $\mathcal{A}$ es denotado por $ST(\mathcal{A})$. Notemos que una función estocástica $V \in ST(\mathcal{A})$ mapea el espacio de estados de $\mathcal{A}$ sobre sí mismo. Dado un estado $\varphi \in \mathcal{A}^*$, el conjunto de funciones estocásticas tales que $V(\varphi) = \varphi$ es denotado por $ST_\varphi(\mathcal{A})$.

En [2] se considera la siguiente noción de mayorización: sea $\bar{w}, \bar{u}$ dos n -uplas ordenadas de estados de una C^* -álgebra abeliana $\mathcal{A}$. Entonces $\bar{u}$ está mayorizada por $\bar{w}$, y notamos $\bar{u} \triangleleft \bar{w}$, si existe una función estocástica $V \in ST(\mathcal{A})$ tal que $V\bar{w} = \bar{u}$, $1 \leq i \leq n$.

Sea $a \in \mathcal{M}^+$ con $\tau(a) = 1$. Entonces a induce estado normal $\tau_a \in \mathcal{M}_*$ dado por $\tau_a(x) = \tau(ax)$. De esta forma, podemos considerar la siguiente definición:

- Sean $\bar{a}, \bar{b} \subset \mathcal{M}^+$ dos n -uplas ordenadas tales que $\tau(a_i) = \tau(b_i) = 1$, $1 \leq i \leq n$, y sea $\mathcal{A}$ una subálgebra abeliana de von Neumann de $\mathcal{M}$. Entonces decimos que $\bar{a}$ está $\mathcal{A}$ -mayorizada por $\bar{b}$, y notamos $\bar{a} \triangleleft_{\mathcal{A}} \bar{b}$, si $(\tau_{a_i}) \triangleleft_{\tau} (\tau_{b_i})$ como estados sobre $\mathcal{A}$.

□

A partir de las definiciones anteriores establecemos la siguiente comparación entre la noción de mayorización conjunta anterior y nuestra noción de mayorización conjunta.

- Sea $\bar{a}, \bar{b} \subset \mathcal{M}_{sa}$ familias abelianas tales que $\bar{a} \prec \bar{b}$. Si $\mathcal{B}$ denota el álgebra de von Neumann abeliana generada por la familia $\bar{b}$, entonces $\bar{a} \triangleleft_{\mathcal{B}} \bar{b}$.

□

Sin embargo la mayorización de Alberti y Uhlman no implica la mayorización conjunta en el sentido desarrollado en este trabajo.

Capítulo 2

Preliminares

2.1. Mayorización y submayorización

2.1.1. Mayorización de vectores en $\mathbb{R}^n$

En esta sección presentamos algunos aspectos básicos de la teoría de mayorización. En el libro [44] se puede hallar una exposición detallada sobre esta noción. Dado $\mathbf{b} = (b_1, \dots, b_n) \in \mathbb{R}^n$, denotemos por $\mathbf{b}^\downarrow \in \mathbb{R}^n$ el vector de coordenadas ordenadas de forma decreciente, obtenido por permutación de las coordenadas de $\mathbf{b}$. Si $\mathbf{b}, \mathbf{c} \in \mathbb{R}^n$ entonces decimos que $\mathbf{c}$ está submayorizado por $\mathbf{b}$, y notamos $\mathbf{c} \prec_w \mathbf{b}$, si

$$\sum_{i=1}^k b_i^\downarrow \geq \sum_{i=1}^k c_i^\downarrow \quad k = 1, \dots, n. \quad (2.1)$$

Si $\mathbf{c} \prec_w \mathbf{b}$ y además se verifica que

$$\sum_{i=1}^n b_i = \sum_{i=1}^n c_i \quad (2.2)$$

decimos que $\mathbf{c}$ está mayorizado por $\mathbf{b}$ y lo notamos $\mathbf{c} \prec \mathbf{b}$. Como una consecuencia inmediata de la definición, notemos que tanto la noción de submayorización como la de mayorización son invariantes bajo permutaciones de las coordenadas de los vectores. Por otro lado, si $\mathbf{c} \prec \mathbf{b}$, es decir que valen las ecuaciones (2.1) y (2.2) entonces también vale

$$\sum_{i=1}^k b_i^\uparrow \leq \sum_{i=1}^k c_i^\uparrow \quad k = 1, \dots, n, \quad (2.3)$$

donde $\mathbf{b}^\uparrow$ denota el vector de coordenadas ordenadas en forma creciente obtenido de $\mathbf{b}$ por permutación de sus coordenadas. De hecho, basta que dos cualesquiera de las tres condiciones dadas por estas fórmulas valgan, para que valga la tercera.

Ejemplo 2.1.1. Consideremos el conjunto $\Omega \subseteq \mathbb{R}^n$ dado por

$$\Omega = \{\mathbf{b} \in \mathbb{R}^n : b_i \geq 0, 1 \leq i \leq n, \sum_{i=1}^n b_i = 1\}.$$

Entonces Ω es un $n - 1$ simplex cuyos vértices son los vectores de la base canónica. Es sencillo verificar que si $\mathbf{b} \in \Omega$ entonces

$$\left(\frac{1}{n}, \dots, \frac{1}{n}\right) \prec \mathbf{b} \prec (1, 0, \dots, 0).$$

Así el conjunto Ω posee un mínimo con respecto a la mayorización. Sin embargo posee varios máximos (esto se debe a que la mayorización no es un orden); en efecto, cada vector de la base canónica $\{e_i\}_{i=1}^n$ es un máximo para Ω (notemos que $e_i \prec e_j$ y $e_j \prec e_i$). Sin embargo la mayorización es un preorden.

Observación 2.1.2. Sea $\mathbf{b} \in \mathbb{R}^n$ y denotemos por Φ_k el conjunto formado por todos los subconjuntos de k elementos de $\{1, \dots, n\}$ para $1 \leq k \leq n$. Entonces es evidente que

$$\sum_{i=1}^k b_i^\downarrow = \max_{F \in \Phi_k} \sum_{i \in F} b_i, \quad \text{para } 1 \leq k \leq n. \quad (2.4)$$

Esta sencilla igualdad es de gran utilidad al momento de extender la noción de mayorización al contexto de sucesiones acotadas de números reales, donde el reordenamiento de forma decreciente de una sucesión puede no tener sentido.

A continuación, describimos algunas caracterizaciones de la mayorización. Más adelante consideraremos otras caracterizaciones en términos de las llamadas matrices doble-estocásticas y orto-estocásticas. Comenzamos con algunas nociones básicas: una transformación lineal T en $\mathbb{R}^n$ es una *T-transformación* si existe $0 \leq t \leq 1$ e índices $j \leq k$ tales que

$$Ty = (y_1, \dots, y_{j-1}, ty_j + (1-t)y_k, \dots, (1-t)y_j + ty_k, y_{k+1}, \dots, y_n).$$

Por otro lado, dada una función convexa $\Phi : \mathbb{R}^n \rightarrow \mathbb{R}^m$ decimos que es *compatible por permutaciones* si para cada permutación $\sigma \in \mathbb{S}_n$ existe una permutación $\sigma' \in \mathbb{S}_m$ tal que

$$\Phi(\mathbf{b}_\sigma) = \Phi(\mathbf{b})_{\sigma'}, \quad \forall \mathbf{b} \in \mathbb{R}^n$$

donde $\mathbf{b}_\sigma \in \mathbb{R}^n$ es el vector de coordenadas $(\mathbf{b}_\sigma)_i = b_{\sigma(i)}$.

Teorema 2.1.3. Sean $\mathbf{b}, \mathbf{c} \in \mathbb{R}^n$. Los siguientes son equivalentes:

1. $\mathbf{c} \prec \mathbf{b}$.
2. $\mathbf{c}$ es obtenido a partir de $\mathbf{b}$ por un número finito de T-transformaciones.
3. $\mathbf{c}$ pertenece a la cápsula convexa de los vectores de $\mathbb{R}^n$ obtenidos al permutar las coordenadas de $\mathbf{b}$.
4. Para cada función $\Phi : \mathbb{R}^n \rightarrow \mathbb{R}^m$ convexa y compatible por permutaciones se tiene $\Phi(\mathbf{c}) \prec_w \Phi(\mathbf{b})$.

Los ítems 2 y 3 muestran una primera interpretación intuitiva de la mayorización: si $\mathbf{c} \prec \mathbf{b}$ entonces el vector $\mathbf{c}$ está relativamente “más mezclado” que el vector $\mathbf{b}$ (ver Ejemplo 2.1.1).

Observación 2.1.4. El ítem 4 pone de manifiesto una utilidad interesante de la mayorización (y tal vez poco evidente a partir de la definición original) en relación con funcionales convexos en $\mathbb{R}^n$. Si $\text{tr}(\mathbf{c})$ denota la suma de las coordenadas del vector $\mathbf{c} \in \mathbb{R}^n$ entonces, dada una función Φ convexa compatible por permutaciones, entonces el funcional

$$\phi(\mathbf{c}) := \text{tr}(\Phi(\mathbf{c})) = \sum_{i=1}^m \Phi(\mathbf{c})_i$$

satisface, por el ítem 4,

$$\mathbf{c} \prec \mathbf{b} \Rightarrow \phi(\mathbf{c}) \leq \phi(\mathbf{b}). \quad (2.5)$$

Es decir, el hecho de que $\mathbf{c}$ está mayorizado por $\mathbf{b}$ (que corresponde a verificar n condiciones sobre los vectores reordenados) implica una familia de desigualdades con respecto a esta clase de funcionales convexos. Dos familias significativas de funciones convexas y compatibles por permutaciones y sus correspondientes funcionales son las siguientes:

- a) Sea $f : I \rightarrow \mathbb{R}$ una función convexa de una variable y consideramos $\Phi_f^{(n)} : \mathbb{R}^n \rightarrow \mathbb{R}^n$ dada por $\Phi_f^{(n)}(\mathbf{c}) = (f(c_1), \dots, f(c_n))$ y entonces $\phi_f^{(n)}(\mathbf{c}) = \sum_{i=1}^n f(c_i)$. Más aún, cuando se desea utilizar el ítem 4 para verificar la mayorización de vectores, se puede reemplazar esta condición por la condición más débil: si $\mathbf{b}, \mathbf{c} \in \mathbb{R}^n$ entonces $\mathbf{c} \prec \mathbf{b}$ si y solo si

$$\phi_f^{(n)}(\mathbf{c}) \leq \phi_f^{(n)}(\mathbf{b}) \quad \text{para toda función convexa } f : I \rightarrow \mathbb{R}.$$

- b) Una norma $\Phi : \mathbb{R}^n \rightarrow \mathbb{R}_{\geq 0}$ se dice gauge simétrica si es invariante por permutaciones i.e $\Phi(\mathbf{c}_\sigma) = \Phi(\mathbf{c})$ para toda permutación $\sigma \in \mathbb{S}_n$ y satisface $\Phi(\mathbf{c}) = \Phi(|c_1|, \dots, |c_n|)$. De hecho, para esta clase de funciones se tiene un resultado más fuerte que el dado por la fórmula 2.5: más precisamente

$$\mathbf{c} \prec_w \mathbf{b} \Rightarrow \Phi(\mathbf{c}) \leq \Phi(\mathbf{b}) \quad \text{para } \Phi \text{ gauge simétrica.} \quad (2.6)$$

Es un resultado conocido de von Neumann que estas funciones están relacionadas con las normas unitariamente invariantes en $\mathbb{M}_n(\mathbb{C})$.

Ejemplo 2.1.5. Como ejemplos de funciones correspondientes a los ítems a) y b) de la observación anterior consideremos los siguientes. Sea $f(t) = -t \log(t)$ que es una función convexa en $(0, \infty)$, con lo cual el funcional $-H$ inducido por $\Phi_f^{(n)}$ en $\mathbb{R}_{\geq 0}^n$ es

$$-H(\mathbf{b}) = \sum_{i=1}^n b_i \log(b_i).$$

En física, el funcional H es denominado *entropía*. Como consecuencia de la desigualdad (2.5) y el Ejemplo 2.1.1 concluimos

$$H(1, 0, \dots, 0) \leq H(\mathbf{b}) \leq H\left(\frac{1}{n}, \dots, \frac{1}{n}\right), \quad \mathbf{b} \in \Omega$$

con la notación del Ejemplo 2.1.1; esta es una de las propiedades básicas de la entropía.

Por otro lado, consideremos las funciones gauge simétricas

$$\Phi_{(k)}(\mathbf{b}) = \sum_{i=1}^k |\mathbf{b}_i^\downarrow|$$

denominadas normas de Ky-Fan. Estas normas juegan un papel especial con respecto a la teoría de mayorización de matrices autoadjuntas que veremos más adelante.

2.1.2. Matrices no negativas. Teorema de Schur-Horn

En lo que sigue, $M_{n,m} := M_{n,m}(\mathbb{R})$ denota el espacio vectorial real de las matrices $n \times m$ a coeficientes reales; en particular notamos $M_n := M_{n,n}$. Sea $A = (a_{ij}) \in M_{n,m}$. Decimos que A es no negativa (resp. positiva) si $a_{ij} \geq 0$ (resp. $a_{ij} > 0$) y notamos $A \geq_{(ij)} 0$ (resp. $A >_{(ij)} 0$). Notemos que la condición “ A es no negativa” ($A \geq_{(ij)} 0$) es diferente de la condición “ A es positiva semidefinida” ($A \geq 0$).

Una matriz no negativa $A \in M_n$ es *estocástica por filas* si es tal que para toda fila, la suma de las entradas de la fila es igual a 1. Si denotamos por $e \in \mathbb{R}^n$ el vector con todas sus entradas iguales a 1, el conjunto de matrices estocásticas por fila en M_n es el politopo caracterizado por

$$RS(n) = \{A \in M_n : A \geq_{(ij)} 0, Ae = e\}.$$

Una matriz estocástica por filas $A \in M_n$ con la propiedad de que su traspuesta A^t es también estocástica por filas es *doble estocástica*. El conjunto de matrices doble estocásticas también es un politopo en M_n , caracterizado por

$$DS(n) = \{D \in M_n : D \geq_{(ij)} 0, De = e, D^t e = e\}.$$

Ejemplos de matriz doble estocástica están dados por las matrices de permutación: si $\sigma \in \mathbb{S}_n$ es una permutación de n elementos entonces P_σ es la matriz determinada por $P_\sigma \mathbf{b} = \mathbf{b}_\sigma$. Además vale que $P_\sigma P_\tau = P_{\sigma \circ \tau}$ con lo cual el conjunto de matrices de permutación es un grupo. El siguiente teorema de Birkhoff asegura que la cápsula convexa de las matrices de permutación son todas las matrices doble estocásticas.

Teorema 2.1.6 (Birkhoff). $D \in M_n$ es matriz doble estocástica si y solo si, para algún $k \in \mathbb{N}$, existen matrices de permutación $P_1, \dots, P_k \in M_n$ y escalares no negativos $\alpha_1, \dots, \alpha_k \in \mathbb{R}$ tales que $\alpha_1 + \dots + \alpha_k = 1$ y

$$D = \sum_{j=1}^k \alpha_j P_j.$$

Como consecuencia del teorema de Birkhoff 2.1.6 y el Teorema 2.1.3 deducimos el siguiente

Teorema 2.1.7. Sean $\mathbf{b}, \mathbf{c} \in \mathbb{R}^n$. Los siguientes son equivalentes:

1. $\mathbf{c} \prec \mathbf{b}$.
2. Existe $D \in DS(n)$ tal que $D\mathbf{b} = \mathbf{c}$.

Dentro del semigrupo compacto de las matrices doble estocásticas se encuentra un subconjunto importante para nosotros, que es el formado por las denominadas matrices orto-estocásticas. Una matriz $O \in M_n$ es orto-estocástica si existe una matriz unitaria $U \in M_n(\mathbb{C})$ tal que $O_{ij} = |U_{ij}|^2$ para $1 \leq i, j \leq n$. Es inmediato verificar que toda matriz orto-estocástica es doble-estocástica.

Ejemplo 2.1.8 (Schur). Sea $\mathbf{b} \in \mathbb{R}^n$ y denotemos por $M_{\mathbf{b}} \in M_n(\mathbb{C})$ a la matriz diagonal con respecto a la base canónica $\{e_i\}_{i=1}^n$ de $\mathbb{C}^n$ que tiene a $\mathbf{b}$ en su diagonal principal. Por otro lado, sea $\mathcal{C} : M_n(\mathbb{C}) \rightarrow \mathbb{C}^n$ la compresión a la diagonal según la base canónica, i.e.

$$\mathcal{C}(A) = (\langle Ae_i, e_i \rangle)_{i=1}^n.$$

Sea $U \in \mathcal{M}_n(\mathbb{C})$ una matriz unitaria y notemos que en este caso se tiene que $\mathcal{C}(U^* M_{\mathbf{b}} U) \in \mathbb{R}^n$ pues $M_{\mathbf{b}}$ es autoadjunta. Más aún

$$\mathcal{C}(U^* M_{\mathbf{b}} U) = O\mathbf{b}, \quad \text{donde } O_{ij} = |U_{ij}|^2, \quad 1 \leq i, j \leq n. \quad (2.7)$$

Luego, del hecho de que las matrices orto-estocásticas son doble-estocásticas y del Teorema 2.1.7 deducimos que $\mathcal{C}(U^* M_{\mathbf{b}} U) \prec \mathbf{b}$, para toda matriz unitaria $U \in \mathcal{M}_n(\mathbb{C})$.

El siguiente teorema establece la recíproca del ejemplo anterior. En su enunciado utilizamos la notación del Ejemplo 2.1.8.

Teorema 2.1.9 (Schur-Horn). Sean $\mathbf{b}, \mathbf{c} \in \mathbb{R}^n$. Entonces, $\mathbf{c} \prec \mathbf{b}$ si y solo si existe una matriz unitaria $U \in M_n(\mathbb{C})$ tal que

$$\mathcal{C}(U^* M_{\mathbf{b}} U) = \mathbf{c}.$$

□

Como un corolario sencillo del teorema de Schur-Horn, la fórmula (2.7) y el Teorema 2.1.7 deducimos que $\mathbf{c} \prec \mathbf{b}$ si y solo si existe una matriz doble-estocástica D y una matriz orto-estocástica O tales que $D\mathbf{b} = O\mathbf{b} = \mathbf{c}$. Sin embargo no debe confundirse este hecho con la afirmación de que toda matriz doble-estocástica sea orto-estocástica, que en realidad es falsa. Por ejemplo, la matriz

$$A = \frac{1}{2} \begin{pmatrix} 1 & 1 & 0 \\ 1 & 0 & 1 \\ 0 & 1 & 1 \end{pmatrix}$$

es claramente doble-estocástica. Sin embargo es un ejercicio sencillo probar que no es orto-estocástica.

2.1.3. MayORIZACIÓN de matrices autoadjuntas

Sea $A \in M_n(\mathbb{C})$ una matriz y consideremos $|A| = (A^* A)^{1/2}$ su módulo (el cálculo funcional en álgebras más generales será descrito más adelante). El vector de valores singulares de A , notado $\mu_A \in \mathbb{R}^n$, es el vector de autovalores de $|A|$ contados con multiplicidad y ordenado de forma decreciente. En particular, si A es una matriz autoadjunta, μ_A es el vector de módulos de los autovalores de A (contados con multiplicidad) ordenados de forma decreciente. En la década del 30, J. von Neumann estableció una dualidad interesante entre las funciones gauge simétricas (ver ítem b) de la Observación 2.1.4) y las llamadas normas unitariamente invariantes, es decir las

normas $||| \cdot |||$ en $M_n(\mathbb{C})$ tales que $|||UAV||| = |||A|||$ siempre que U, V sean matrices unitarias. Más concretamente, von Neumann notó que si Φ es una función gauge simétrica, entonces la función

$$|||A||| = \Phi(\mu_A), \quad A \in M_n(\mathbb{C})$$

determina una norma unitariamente invariante y recíprocamente, dada una norma unitariamente invariante $||| \cdot |||$ entonces se obtiene una función gauge simétrica

$$\tilde{\Phi}(\mathbf{b}) = |||M_{\mathbf{b}}|||, \quad \text{de forma que} \quad |||A||| = \tilde{\Phi}(\mu_A).$$

De esta forma, si $|||A|||(k) = \Phi_{(k)}(\mu_A)$ denotan las normas inducidas por los funcionales de Ky-Fan definidos en el Ejemplo 2.1.5 entonces se tiene

$$|||A|||(k) \leq |||B|||(k) \Rightarrow |||A||| \leq |||B|||, \quad \text{para toda norma unitariamente invariante,}$$

que no es más que otra forma de expresar la implicación $\mu_A \prec_w \mu_B \Rightarrow \Phi(\mu_A) \leq \Phi(\mu_B)$ para toda función gauge simétrica (ver ecuación (2.6)). Debido a que la mayorización provee de un método para probar desigualdades generales con respecto a normas de operadores, esta noción es de gran utilidad.

En el artículo [5] Ando extiende la mayorización de vectores al caso de la mayorización entre matrices autoadjuntas, probablemente debido a la importancia de la mayorización en el análisis matricial descrita en el párrafo anterior. Si $A, B \in M_n(\mathbb{C})_{sa}$ son matrices autoadjuntas entonces se dice que A submayoriza a B y se nota $B \prec_w A$ si $\lambda(B) \prec_w \lambda(A)$, donde $\lambda(A)$ denota el vector de autovalores de A , contados con multiplicidad. Esta noción puede considerarse como una extensión “no conmutativa” de la mayorización de vectores reales al caso del espacio vectorial real de las matrices autoadjuntas. Notemos por otro lado que, en tanto la mayorización se define en términos del espectro de las matrices, esta noción corresponde más bien a una comparación entre los operadores lineales determinados por estas matrices con respecto a alguna (cualquier) base ortonormal.

A continuación describimos algunas caracterizaciones de la mayorización entre matrices autoadjuntas, algunas de las cuales son traducciones de las obtenidas en el Teorema 2.1.3 para el caso de la mayorización de vectores. A tal fin consideramos las siguientes nociones. Sea $A \in M_n(\mathbb{C})$ y denotemos por $\mathcal{U}(n)$ el grupo compacto de matrices unitarias de $M_n(\mathbb{C})$. Entonces, la órbita unitaria de A , notada $\mathcal{U}(n)(A)$ está definida por

$$\mathcal{U}(n)(A) = \{U^*AU : U \in \mathcal{U}(n)\}. \quad (2.8)$$

Por otro lado, dado un conjunto X en un espacio vectorial real entonces $\text{conv}(X)$ denota su cápsula convexa. Decimos que una transformación lineal T en $M_n(\mathbb{C})$ es doble estocástica si es unital ($T(1) = 1$), positiva ($T(A) \geq 0$ siempre que $A \geq 0$) y preserva la traza $\text{tr}(\cdot)$ de $M_n(\mathbb{C})$ ($\text{tr}(T(A)) = \text{tr}(A)$). Las transformaciones doble estocásticas son los análogos de las matrices doble estocásticas.

Teorema 2.1.10. Sean $A, B \in M_n(\mathbb{C})_{sa}$ matrices autoadjuntas. Los siguientes son equivalentes:

1. $A \prec B$.
2. $A \in \text{conv}(\mathcal{U}(n)(B))$.

3. $\text{tr}(f(A)) \leq \text{tr}(f(B))$ para toda función convexa $f : I \rightarrow \mathbb{R}$ tal que $\sigma(A) \cup \sigma(B) \subseteq I$.
4. Existe una transformación doble estocástica T en $\mathcal{M}_n(\mathbb{C})$ tal que $T(B) = A$.

En este contexto, el ítem 3. tal vez sea el más significativo para las aplicaciones, pues provee de una familia de desigualdades para funcionales convexos definidos sobre el espacio vectorial real de las matrices autoadjuntas.

2.1.4. Mayorización en $\ell_{\mathbb{R}}^{\infty}(\mathbb{N})$ y el teorema de Schur-Horn

En esta sección describiremos una extensión de la mayorización de vectores en $\mathbb{R}^n$ al caso de vectores en $\ell_{\mathbb{R}}^{\infty}(\mathbb{N})$, el espacio vectorial real de las sucesiones reales uniformemente acotadas, desarrollada por Neumann en [56]. Esta generalización está en el espíritu de la noción original de la mayorización (más adelante veremos extensiones “no conmutativas” en dimensión infinita). Curiosamente, la motivación de Neumann provino de la geometría (simpléctica), más precisamente del estudio de generalizaciones del teorema de Konstant. Si bien Neumann desarrolló versiones del teorema de Schur-Horn en el contexto de operadores autoadjuntos acotados en un espacio de Hilbert separable, la formulación general que obtuvo no es la adecuada para el estudio de la admisibilidad de marcos en espacios de Hilbert que realizaremos en el Capítulo 3. Es por eso que desarrollamos algunos aspectos de la teoría de mayorización de Neumann que utilizaremos para deducir una versión del teorema de Schur-Horn para operadores autoadjuntos más adecuada a nuestros fines.

Mayorización en $\ell_{\mathbb{R}}^{\infty}(\mathbb{N})$. Si $\mathbf{a} \in \ell_{\mathbb{R}}^{\infty}(\mathbb{N})$ entonces puede no tener sentido considerar el reordenamiento de $\mathbf{a}$ de forma decreciente, que es el punto de partida de la noción de mayorización en $\mathbb{R}^n$. En efecto, consideremos el vector

$$v = (1, 0, 1, 0, 1, \dots) \in \ell_{\mathbb{R}}^{\infty}(\mathbb{N})$$

que tiene una infinidad de coordenadas iguales a 0 y 1 respectivamente. En este caso, si queremos considerar un reordenamiento en forma decreciente de este vector, es decir el vector $v^{\downarrow}$ obtenido por una permutación de las coordenadas de v de forma que las coordenadas de $v^{\downarrow}$ estén ordenadas de forma decreciente, nos encontraríamos con el problema de ubicar los infinito ceros “al final de una infinidad de unos”. Sin embargo Neumann encontró una forma de resolver este problema. En efecto, consideremos Φ_k el conjunto formado por todos los subconjuntos de $\mathbb{N}$ de k elementos; entonces dada $\mathbf{a} \in \ell_{\mathbb{R}}^{\infty}(\mathbb{N})$ definimos

$$U_k(\mathbf{a}) = \sup_{F \in \Phi_k} \sum_{i \in F} a_i \quad \text{y} \quad L_k(\mathbf{a}) = \inf_{F \in \Phi_k} \sum_{i \in F} a_i = -U_k(-\mathbf{a}) \quad (2.9)$$

De la Observación 2.1.2 vemos que si consideramos la inmersión natural de $\mathbb{R}^n$ en $\ell_{\mathbb{R}}^{\infty}(\mathbb{N})$ entonces las funciones definidas arriba reemplazan las sumas que aparecen en las ecuaciones (2.1) y (2.3). En este contexto, general análogos de la ecuación (2.2) tienen poco interés debido al hecho de que las series determinadas por $\mathbf{a} \in \ell_{\mathbb{R}}^{\infty}(\mathbb{N})$ son en general divergentes. A continuación damos una breve descripción de algunos de los resultados de [56] relacionados con mayorización en $\ell_{\mathbb{R}}^{\infty}(\mathbb{N})$, pero antes consideramos las siguientes nociones.

Definición 2.1.11. Sea Π el conjunto de todas las funciones biyectivas en $\mathbb{N}$ y, para cada $k \in \mathbb{N}$, denotemos por $\Pi_k \subseteq \Pi$ el conjunto de permutaciones σ tales que $\sigma(n) = n$ para cada $n > k$. Dado $\mathbf{a} \in \ell_{\mathbb{R}}^{\infty}(\mathbb{N})$ y $\sigma \in \Pi$, notamos:

1. $\mathbf{a}_\sigma = (a_{\sigma(1)}, a_{\sigma(2)}, \dots)$.
2. $\Pi \cdot \mathbf{a} = \{\mathbf{a}_\sigma, \sigma \in \Pi\}$, la órbita de $\mathbf{a}$ bajo la acción de Π .
3. $\text{conv}(\Pi \cdot \mathbf{a})$, la cápsula convexa de la órbita de $\mathbf{a}$. □

En lo que sigue $\text{cl}_{\|\cdot\|_\infty}(X) = \overline{X}^{\|\cdot\|_\infty}$ denota la clausura en norma $\|\cdot\|_\infty$ de $X \subseteq \ell_\mathbb{R}^\infty(\mathbb{N})$.

Teorema 2.1.12. Sean $\mathbf{a}, \mathbf{b} \in \ell_\mathbb{R}^\infty(\mathbb{N})$. Los siguientes son equivalentes:

1. $\mathbf{a} \in \text{cl}_{\|\cdot\|_\infty}(\text{conv}(\Pi \cdot \mathbf{b}))$.
2. $U_k(\mathbf{b}) \geq U_k(\mathbf{a})$, $L_k(\mathbf{b}) \leq L_k(\mathbf{a})$ para $k \geq 1$.

Si además $\mathbf{b} \in \ell_\mathbb{R}^1(\mathbb{N})$ entonces los siguientes son equivalentes:

3. $\mathbf{a} \in \text{cl}_{\|\cdot\|_1}(\text{conv}(\Pi \cdot \mathbf{b}))$.
4. $U_k(\mathbf{b}) \geq U_k(\mathbf{a})$, $L_k(\mathbf{b}) \leq L_k(\mathbf{a})$ para $k \geq 1$ y $\sum_{i=1}^\infty b_i = \sum_{i=1}^\infty a_i$. □

Es interesante notar que, si $\mathbf{a} \in \ell_\mathbb{R}^\infty(\mathbb{N})$ y definimos

$$a_i^+ = \max\{a_i - \limsup \mathbf{a}, 0\} \quad \text{y} \quad a_i^- = \min\{a_i - \liminf \mathbf{a}, 0\}, \quad i \in \mathbb{N}, \quad (2.10)$$

entonces, $\lim_{i \rightarrow \infty} a_i^+ = 0$, $\lim_{i \rightarrow \infty} a_i^- = 0$ y para cada $k \in \mathbb{N}$ vale que

$$U_k(\mathbf{a}) = U_k(\mathbf{a}^+) + k \limsup \mathbf{a} \quad \text{y} \quad L_k(\mathbf{a}) = L_k(\mathbf{a}^-) + k \liminf \mathbf{a}. \quad (2.11)$$

En particular deducimos que

$$\mathbf{a} \in \text{cl}_{\|\cdot\|_\infty}(\text{conv}(\Pi \cdot \mathbf{b})) \Rightarrow \limsup \mathbf{b} \geq \limsup \mathbf{a} \quad \text{y} \quad \liminf \mathbf{b} \leq \liminf \mathbf{a}.$$

Definición 2.1.13. Si $\mathbf{a}, \mathbf{b} \in \ell_\mathbb{R}^\infty(\mathbb{N})$ entonces decimos que $\mathbf{b}$ mayoriza a $\mathbf{a}$ y notamos $\mathbf{a} \prec \mathbf{b}$ si se verifica alguna de las condiciones 1. ó 2. del Teorema 2.1.12; si además $\mathbf{b} \in \ell_\mathbb{R}^1(\mathbb{N})$ entonces decimos que $\mathbf{b}$ mayoriza a $\mathbf{a}$ (en el sentido ℓ^1) si se verifica alguna de las condiciones 3. ó 4. del Teorema 2.1.12.

El teorema de Schur-Horn en $L(\mathcal{H})$. En lo que sigue introducimos las notaciones necesarias para enunciar el teorema de Schur-Horn en la versión infinito dimensional de Neumann.

Sea $\mathbb{M}$ un conjunto, sea $\mathcal{K}$ un espacio de Hilbert con $\dim \mathcal{K} = |\mathbb{M}|$ y sea $\mathcal{B} = \{e_n\}_{n \in \mathbb{M}}$ una base ortonormal de $\mathcal{K}$

1. Para cualquier $\mathbf{a} = (a_n)_{n \in \mathbb{M}} \in \ell^\infty(\mathbb{M})$, denotemos por $M_{\mathcal{B}, \mathbf{a}} \in L(\mathcal{K})$ el operador diagonal dado por $M_{\mathcal{B}, \mathbf{a}} e_n = a_n e_n$, $n \in \mathbb{M}$. Cuando sea claro del contexto qué base se está usando abreviaremos $M_{\mathcal{B}, \mathbf{a}} = M_{\mathbf{a}}$.
2. Si $\mathbf{a} \in \mathbb{C}^n$, haciendo abuso de notación, denotamos por $M_{\mathbf{a}} \in \mathcal{M}_n(\mathbb{C})$ la matriz diagonal (con respecto a la base canónica de $\mathbb{C}^n$) que tiene las entradas de $\mathbf{a}$ en su diagonal.

3. La compresión diagonal $\mathcal{C}_{\mathcal{B}} : L(\mathcal{K}) \rightarrow L(\mathcal{K})$ asociada a la base $\mathcal{B}$, está definida por $\mathcal{C}_{\mathcal{B}}(T) = M_{\mathcal{B}, \mathbf{a}}$, donde $\mathbf{a} = (\langle Te_n, e_n \rangle)_{n \in \mathbb{N}}$. $\square$

Definición 2.1.14. Sea $\mathcal{H}$ un espacio de Hilbert, $S \in L(\mathcal{H})$ y $\mathcal{B}$ una base ortonormal de $\mathcal{H}$. Entonces,

1. $\mathcal{U}_{\mathcal{H}}(S) = \{U^* S U : U \in \mathcal{U}(\mathcal{H})\}$.

2. $\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)] = \{\mathbf{c} \in \ell^\infty(\mathbb{N}) : M_{\mathcal{B}, \mathbf{c}} \in \mathcal{C}_{\mathcal{B}}(\mathcal{U}_{\mathcal{H}}(S))\}$. $\square$

Observación 2.1.15. Dado $S \in L(\mathcal{H})$, la definición de $\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]$ no depende de la base ortonormal $\mathcal{B}$. De hecho, si $\mathcal{B}'$ es otra base ortonormal de $\mathcal{H}$, $U \in \mathcal{U}(\mathcal{H})$ manda $\mathcal{B}$ sobre $\mathcal{B}'$, y $\mathbf{c} \in \ell^\infty(\mathbb{N})^+$ satisface $M_{\mathcal{B}, \mathbf{c}} = \mathcal{C}_{\mathcal{B}}(T)$ para algún $T \in \mathcal{U}_{\mathcal{H}}(S)$, entonces

$$M_{\mathcal{B}', \mathbf{c}} = U M_{\mathcal{B}, \mathbf{c}} U^* = U \mathcal{C}_{\mathcal{B}}(T) U^* = \mathcal{C}_{\mathcal{B}'}(U T U^*) \in \mathcal{C}_{\mathcal{B}'}(\mathcal{U}_{\mathcal{H}}(S)) .$$

Con lo cual $\{\mathbf{c} \in \ell^\infty(\mathbb{N}) : M_{\mathcal{B}', \mathbf{c}} \in \mathcal{C}_{\mathcal{B}'}(\mathcal{U}_{\mathcal{H}}(S))\} = \mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]$. Notemos que $\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]$ es el conjunto de todas las posibles diagonales de S en sus representaciones matriciales con respecto a las bases ortonormales de $\mathcal{H}$. $\square$

Teorema 2.1.16 (Schur-Horn para operadores diagonales [56]). Sea $S \in L(\mathcal{H})$ un operador autoadjunto tal que $S = M_{\mathcal{B}, \mathbf{a}}$ para algún $\mathbf{a} \in \ell^\infty(\mathbb{N})$ y alguna base ortonormal $\mathcal{B}$, entonces

$$\text{cl}_{\|\cdot\|_\infty}(\text{conv}(\Pi \cdot \mathbf{a})) = \text{cl}_{\|\cdot\|_\infty}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]) . \quad (2.12)$$

$\square$

De esta forma, dado $\mathbf{b} \in \ell^\infty(\mathbb{N})$ entonces $\mathbf{a}$ está mayorizado (según la Definición 2.1.13) por $\mathbf{b}$ si y sólo si para cada base ortonormal existe una sucesión de operadores unitarios $(U_n)_{n \geq 1} \subseteq \mathcal{U}_{\mathcal{H}}$ tal que

$$\lim_{n \rightarrow \infty} \|\mathcal{M}_{\mathcal{B}, \mathbf{a}} - \mathcal{C}_{\mathcal{B}}(U_n^* M_{\mathcal{B}, \mathbf{b}} U_n)\|_\infty = 0,$$

que resulta la generalización del Teorema 2.1.9.

Para describir la versión de [56] para operadores autoadjuntos no diagonalizables, introducimos las siguientes nociones. Denotamos por $L_0(\mathcal{H})$ el *-ideal cerrado de los operadores compactos y consideremos el cociente $\mathcal{A}(\mathcal{H}) = L(\mathcal{H})/L_0(\mathcal{H})$, que es una C^* -álgebra unital, conocida como el álgebra de Calkin. Dado $T \in L(\mathcal{H})$, el *espectro esencial* de T , notado $\sigma_e(T)$, es el espectro de la clase $T + L_0(\mathcal{H})$ en el álgebra $\mathcal{A}(\mathcal{H})$. La norma esencial $\|T\|_e = \inf\{\|T + K\| : K \in L_0(\mathcal{H})\}$ de T es la norma (cociente) de $T + L_0(\mathcal{H})$ en $\mathcal{A}(\mathcal{H})$. Notemos que $\sigma_e(T)$ es un subconjunto compacto de $\sigma(T)$. Si $S \in L(\mathcal{H})_h$ definimos

$$\alpha^+(S) = \max \sigma_e(S) = \|S\|_e \quad \text{y} \quad \alpha_-(S) = \min \sigma_e(S) , \quad (2.13)$$

y si $S = \int_{\sigma(S)} t \, dE(t)$ es la representación espectral de S con respecto a la medida espectral E , consideramos los siguientes operadores compactos:

$$\begin{aligned} S^+ &= \int_{[\alpha^+(S), \|S\|]} (t - \alpha^+(S)) dE(t) , \quad \text{y} \\ S_- &= \int_{[-\|S\|, \alpha_-(S)]} (t - \alpha_-(S)) dE(t) . \end{aligned} \quad (2.14)$$

Notemos que $S_- \leq 0 \leq S^+$ y que si S es no diagonalizable entonces $\alpha^+(S) > \alpha_-(S)$. De esta forma se tiene

Teorema 2.1.17 (Scur-Horn para operadores no diagonales). Sea $S \in L(\mathcal{H})$ un operador autoadjunto no diagonalizable. Entonces

$$\overline{\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]}^{\|\cdot\|_{\infty}} = \text{cl}_{\|\cdot\|_{\infty}}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S^+)]) + [\alpha_-(S), \alpha^+(S)]^{\mathbb{N}} + \text{cl}_{\|\cdot\|_{\infty}}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S_-)]),$$

donde $\alpha^+(S)$, $\alpha_-(S)$, S^+ , S_- son definidos por las ecuaciones (2.13) y (2.14). $\square$

Como ya hemos comentado, esta versión del teorema de Schur-Horn no es la conveniente para nuestro estudio del problema de admisibilidad de marcos en espacios de Hilbert. Sin embargo, una formulación del teorema de Schur-Horn basada en los teoremas 2.1.12 y 2.1.16 resulta ser adecuada. El hecho que el teorema 2.1.16 sea válido solo para operadores diagonales no representa en realidad una limitación, en virtud del siguiente

Teorema 2.1.18 (Weyl-Berg-von Neumann). Sea $S \in L(\mathcal{H})_h$ un operador autoadjunto. Entonces existe un operador diagonal $D \in \overline{\mathcal{U}_{\mathcal{H}}(S)}^{\|\cdot\|_{\infty}}$ con respecto a alguna (toda) base ortonormal $\mathcal{B}$ de $\mathcal{H}$.

En realidad el teorema anterior es la versión moderna del teorema de Weyl-von Neumann, y es un resultado conocido en la teoría de álgebras de operadores.

Observación 2.1.19. Dos operadores $S, T \in L(\mathcal{H})$ se dicen *aproximadamente unitariamente equivalentes* si existe una sucesión $(U_n)_{n \geq 1} \in \mathcal{U}_{\mathcal{H}}$ tal que

$$\lim_{n \rightarrow \infty} \|T - U_n^* S U_n\| = 0. \quad (2.15)$$

Esta es una relación de equivalencia, que es estudiada en la teoría de álgebras de operadores. Por ejemplo, se prueba en el libro de Davidson [28] (II.4.4) que $S, T \in L(\mathcal{H})_h$ son aproximadamente unitariamente equivalentes si y solo si $\sigma_e(S) = \sigma_e(T)$ y $\dim \ker(S - \lambda I) = \dim \ker(T - \lambda I)$ para cada $\lambda \notin \sigma_e(S)$. $\square$

2.2. Marcos en espacios de Hilbert

Introducimos algunos hechos básicos acerca de los marcos en espacios de Hilbert. Para una descripción completa de la teoría de marcos y sus aplicaciones se recomiendan los trabajos de Daubechies, Grossmann y Meyer [27], Aldroubi [3], la monografía de Heil y Walnut [38] o los libros de Young [70] y Christensen [25].

2.2.1. Definiciones básicas

Antes de definir los marcos consideramos los siguientes conceptos básicos relacionados con ellos. Sea $\mathcal{H}$ un espacio de Hilbert separable y consideremos $\mathbb{M} = \mathbb{N}$ ó $\mathbb{M} = \{1, \dots, m\} := I_m$ para algún $m \in \mathbb{N}$. Una sucesión $\{f_n\}_{n \in \mathbb{M}} \subseteq \mathcal{H}$ es una sucesión *Bessel* si existe una constante $\beta > 0$ tal que

$$\sum_{n \in \mathbb{M}} |\langle f, f_n \rangle|^2 \leq \beta \|f\|^2, \quad \text{para cada } f \in \mathcal{H}. \quad (2.16)$$

Observación 2.2.1. De esta forma, si $\{f_n\}_{n \in \mathbb{M}}$ es una sucesión Bessel entonces la ecuación (2.16) muestra que el operador dado por $\mathcal{H} \ni f \mapsto (\langle f, f_n \rangle)_{n \in \mathbb{M}}$ está bien definido y resulta acotado entre los espacios $\mathcal{H}$ y $\ell^2(\mathbb{M})$. Es sencillo verificar que el adjunto del operador antes descrito está dado por $\ell^2(\mathbb{M}) \ni (a_i)_{i \in \mathbb{M}} \mapsto \sum_{n \in \mathbb{M}} a_n f_n$; notemos que este operador es acotado por ser el adjunto de un operador acotado.

El siguiente resultado es un ejercicio interesante de análisis funcional.

Teorema 2.2.2. Sea $\mathcal{H}$ un espacio de Hilbert separable y sea $\{f_n\}_{n \in \mathbb{M}} \subseteq \mathcal{H}$ una sucesión Bessel. Los siguientes son equivalentes:

1. El operador $T : \ell^2(\mathbb{M}) \rightarrow \mathcal{H}$ determinado por $T(e_n) = f_n$ para todo $n \in \mathbb{M}$ es epimorfismo acotado.
2. El operador $T^* : \mathcal{H} \rightarrow \ell^2(\mathbb{M})$ dado por $T^*(f) = (\langle f, f_n \rangle)_{n \in \mathbb{M}}$ es monomorfismo acotado y de rango cerrado.
3. Existen constantes $\alpha, \beta > 0$ tales que

$$\alpha \|f\|^2 \leq \sum_{n \in \mathbb{M}} |\langle f, f_n \rangle|^2 \leq \beta \|f\|^2, \quad \text{para cada } f \in \mathcal{H}. \quad (2.17)$$

4. Existen constantes $\alpha, \beta > 0$ tales que

$$\alpha I \leq TT^* \leq \beta I. \quad (2.18)$$

En este caso, $\alpha > 0$ (resp $\beta > 0$) verifica la fórmula (2.17) si y solo si verifica (2.18).

Definición 2.2.3. Sea $\mathcal{H}$ un espacio de Hilbert separable, y $\mathcal{F} = \{f_n\}_{n \in \mathbb{M}}$ una sucesión Bessel en $\mathcal{H}$. $\mathcal{F}$ es un *marco* para $\mathcal{H}$ si satisface alguna de las condiciones del Teorema 2.2.2.

En el caso en que $\{f_n\}_{n \in \mathbb{M}}$ resulta un marco para $\mathcal{H}$ entonces se debe tener $\dim \mathcal{H} \leq |\mathbb{M}|$, donde $|\mathbb{M}|$ denota el cardinal del conjunto $\mathbb{M}$. En la siguiente observación describimos una serie de objetos asociados a un marco para $\mathcal{H}$:

Observación 2.2.4. Sea $\mathcal{F} = \{f_n\}_{n \in \mathbb{M}}$ un marco para $\mathcal{H}$.

1. El epimorfismo $T \in L(\ell^2(\mathbb{M}), \mathcal{H})$ determinado por $T(e_n) = f_n$ para $n \in \mathbb{M}$ (ver Teorema 2.2.2) es llamado el operador de *síntesis* (o de premarco) para $\mathcal{F}$.
2. El adjunto $T^* \in L(\mathcal{H}, \ell^2(\mathbb{M}))$ de T , está dado por

$$T^*(f) = \sum_{n \in \mathbb{M}} \langle f, f_n \rangle e_n, \quad f \in \mathcal{H},$$

donde $\{e_n\}_{n \in \mathbb{M}}$ denota la base canónica de $\ell^2(\mathbb{M})$. T^* es el llamado *operador de análisis* para $\mathcal{F}$. El rango de este operador, que es un subespacio de $\ell^2(\mathbb{M})$, es el espacio de coeficientes admisibles.

3. Las constantes óptimas α, β en la ecuación (2.17) se llaman las *cotas de marco* para $\mathcal{F}$. El marco $\mathcal{F}$ se dice *ajustado* si $\alpha = \beta$, y *Parseval* si $\alpha = \beta = 1$. Los marcos Parseval también son llamados marcos normalizados y ajustados .
4. El operador $S = TT^* \in L(\mathcal{H})^+$, llamado el *operador de marco* de $\mathcal{F}$, satisface

$$Sf = \sum_{n \in \mathbb{M}} \langle f, f_n \rangle f_n = \left(\sum_{n \in \mathbb{M}} f_n \otimes f_n \right) f, \quad \text{para cada } f \in \mathcal{H}. \quad (2.19)$$

Se sigue del Teorema 2.2.2 que $\alpha I \leq S \leq \beta I$ y en particular, $S \in \mathcal{G}l(\mathcal{H})^+$. $\square$

2.2.2. El exceso de un marco

Sea $\{f_n\}_{n \in \mathbb{M}}$ un marco para $\mathcal{H}$. Entonces, el operador de síntesis T es un epimorfismo es decir, para cada $f \in \mathcal{H}$ existe una sucesión de coeficientes $(a_n)_{n \in \mathbb{M}} \in \ell^2(\mathbb{M})$ de forma que

$$f = T((a_n)_{n \in \mathbb{M}}) = \sum_{n \in \mathbb{M}} a_n f_n.$$

Esto muestra que $\{f_n\}_{n \in \mathbb{M}}$ es un sistema de generadores; más aún, los coeficientes de generación pueden ser elegidos dentro del espacio de Hilbert $\ell^2(\mathbb{M})$. Sin embargo esta es una propiedad compartida con las bases ortonormales, que poseen como rango de coeficientes admisibles todo $\ell^2(\mathbb{M})$. Entonces ¿qué interés tienen los marcos? En realidad, a diferencia de las bases ortonormales, los coeficientes de reconstrucción no son únicos, lo que resulta de interés en las aplicaciones. En efecto, esta redundancia producto de la sobre-determinación propia de los marcos se traduce en una cierta propiedad de robustez.

Ejemplo 2.2.5. Imaginemos el siguiente procedimiento simplificado de transmisión de datos: en este caso, los datos son vectores de un espacio de Hilbert $\mathcal{H}$; éstos son representados con respecto a una base ortonormal y se transmiten sus coeficientes de Fourier. Si alguno de estos coeficientes “se altera” durante la transmisión (digamos que por ruido en el canal de transmisión) entonces no hay forma, ni siquiera, de detectar esta alteración. Sin embargo, si consideramos la descomposición de los datos con respecto a marco, podemos utilizar las relaciones lineales que hay entre los coeficientes de reconstrucción (conocidas a priori) para detectar una alteración de los datos, y tal vez para realizar una corrección de tal alteración. En efecto, en el caso de los marcos, el rango del operador de análisis puede ser un subespacio propio de $\ell^2(\mathbb{M})$ de forma que si los coeficientes “recibidos” en el modelo de transmisión, no corresponden a este espacio debe haber una alteración de los datos. Más aún, podemos pensar en proyectar estos datos alterados sobre el espacio de coeficientes admisibles con estos datos “corregidos”. Como veremos más adelante, todo vector se puede reconstruir a partir de sus coeficientes con respecto a un marco.

La siguiente noción da cuenta del nivel de redundancia de la información en un marco.

Definición 2.2.6. Sea $\mathcal{F} = \{f_n\}_{n \in \mathbb{M}}$ un marco en $\mathcal{H}$. El número cardinal

$$e(\mathcal{F}) = \dim \left\{ (c_n)_{n \in \mathbb{M}} \in \ell^2(\mathbb{M}) : \sum_{n \in \mathbb{M}} c_n f_n = 0 \right\},$$

es llamado el exceso del marco. Holub [43] y Balan, Casazza, Heil y Landau [13] probaron que

$$e(\mathcal{F}) = \sup\{|I| : I \subseteq \mathbb{M} \text{ y } \{f_n\}_{n \notin I} \text{ es un marco en } \mathcal{H}\}.$$

Esta caracterización justifica el nombre “exceso de $\mathcal{F}$ ”. El marco $\mathcal{F}$ es llamado *base de Riesz* si $e(\mathcal{F}) = 0$, i.e., si el operador de síntesis de $\mathcal{F}$ es inversible. $\square$

Si $\{f_n\}_{n \in \mathbb{M}}$ es un marco para $\mathcal{H}$ notemos que el exceso de $\{f_n\}_{n \in \mathbb{M}}$ coincide con la dimensión del núcleo del operador de síntesis T asociado a este marco. Por otro lado, $\ker(T) = R(T^*)^\perp$ de forma que el exceso de $\{f_n\}_{n \in \mathbb{M}}$ está dando una medida del tamaño del complemento ortogonal del espacio de coeficientes admisibles. En la medida que el exceso sea mayor, entonces el espacio de coeficientes es “más pequeño”, es decir, los vectores en él están sujetos a una cantidad mayor de restricciones. Esto se traduce en una mayor capacidad de detectar “errores de transmisión.”^{en} nuestro modelo simplificado de transmisión del Ejemplo 2.2.5.

2.2.3. El problema de la reconstrucción de marcos

Como hemos mencionado, los marcos son sistemas de generadores (a coeficientes en $\ell^2(\mathbb{M})$) posiblemente linealmente dependientes. Un problema natural dentro de la teoría de marcos es el siguiente: dado un marco $\{f_n\}_{n \in \mathbb{M}}$ para $\mathcal{H}$, y dado $f \in \mathcal{H}$ hallar coeficientes $(a_n)_{n \in \mathbb{M}} \in \ell^2(\mathbb{M})$ de forma que

$$f = \sum_{n \in \mathbb{M}} a_n f_n.$$

Este procedimiento se denomina “reconstrucción” de $f \in \mathcal{H}$ a partir del marco $\{f_n\}_{n \in \mathbb{M}}$. De la posible sobredeterminación de $\{f_n\}_{n \in \mathbb{M}}$ notamos que la elección anterior puede no ser única; se busca entonces un método sistemático que permita, dado $f \in \mathcal{H}$ hallar coeficientes para la reconstrucción de f que tengan alguna propiedad de *optimalidad*.

Las siguientes observaciones elementales proveen de un tal método. Si $\{f_n\}_{n \in \mathbb{M}}$ es un marco para $\mathcal{H}$ entonces $S = TT^* \in \mathcal{G}l(\mathcal{H})^+$ es un operador positivo e inversible. Sea $S^{-1} \in L(\mathcal{H})$ el inverso de S y notemos que las identidades $SS^{-1} = S^{-1}S = I$ se traducen (ver la fórmula (2.19)) en

$$f = \sum_{n \in \mathbb{M}} \langle f, S^{-1}f_n \rangle f_n = \sum_{n \in \mathbb{M}} \langle f, f_n \rangle S^{-1}f_n. \quad (2.20)$$

Notemos que la sucesión $\{S^{-1}f_n\}_{n \in \mathbb{M}}$ es un marco para $\mathcal{H}$: en efecto, basta notar que su operador de síntesis determinado por $\tilde{T}(e_n) = S^{-1}f_n = S^{-1}Te_n$ se extiende a un epimorfismo de $\mathcal{H}$, donde T denota el operador de síntesis de $\{f_n\}_{n \in \mathbb{M}}$. El marco $\{S^{-1}f_n\}_{n \in \mathbb{M}}$ se llama *marco dual canónico* del marco $\{f_n\}_{n \in \mathbb{M}}$.

De la primera igualdad vemos que la sucesión $(\langle f, S^{-1}f_n \rangle)_{n \in \mathbb{M}}$ proporciona coeficientes para la reconstrucción de f a partir de $\{f_n\}_{n \in \mathbb{M}}$. Por otro lado, el hecho de que $\{S^{-1}f_n\}_{n \in \mathbb{M}}$ es un marco para $\mathcal{H}$ implica que la dependencia de los coeficientes de reconstrucción $f \mapsto (\langle f, S^{-1}f_n \rangle)_{n \in \mathbb{M}}$ es lineal y continua. Además posee la siguiente propiedad

Proposición 2.2.7. Sea $\{f_n\}_{n \in \mathbb{M}}$ un marco para $\mathcal{H}$. Entonces la sucesión de coeficientes

$$(\langle f, S^{-1}f_n \rangle)_{n \in \mathbb{M}} \in \ell^2(\mathbb{M})$$

tiene norma mínima entre todas las sucesiones $(a_n)_{n \in \mathbb{M}} \in \ell^2(\mathbb{M})$ que satisfacen

$$f = \sum_{n \in \mathbb{M}} a_n f_n.$$

Hasta aquí, el problema de reconstrucción de marcos tiene una solución canónica. Sin embargo, esta solución de carácter teórico resulta en algunos casos concretos impracticable: en efecto, involucra el invertir un operador en un espacio de dimensión (posiblemente) infinita. Es por eso que el problema de la reconstrucción (práctica) sigue siendo material de investigación en la actualidad.

Sin embargo, hay una clase de marcos para los cuales el problema de la inversión de su operador de marco puede llevarse a cabo, de forma elemental. En efecto, recordemos que un marco $\{f_n\}_{n \in \mathbb{M}}$ se dice ajustado si sus cotas óptimas de marco son iguales. Esto se traduce en el hecho (ver ítem 4 de la Observación 2.2.4) de que su operador de marco $S = \alpha \cdot I$ es un múltiplo (no nulo) de la identidad. En este caso el marco dual resulta $\{\frac{f_n}{\alpha}\}_{n \in \mathbb{M}}$; más aún, si $\alpha = 1$ (cuando el marco es Parseval) entonces el marco coincide con su dual y la fórmula de reconstrucción (2.20) se convierte en

$$f = \sum_{n \in \mathbb{M}} \langle f, f_n \rangle f_n$$

que es formalmente igual a la fórmula de reconstrucción para una base ortonormal. Es por esta razón que los marcos Parseval son particularmente importantes para las aplicaciones: su proceso de reconstrucción es formalmente el de una base ortonormal, pero tienen la ventaja de brindar una reconstrucción más estable (en el sentido del Ejemplo 2.2.5). En ciertos casos es incluso deseable considerar un marco Parseval con propiedades adicionales; por ejemplo, de forma que las normas de los elementos del marco tengan una cierta distribución a priori $(c_n)_{n \in \mathbb{M}}$, es decir $\|f_n\| = c_n$ para todo $n \in \mathbb{M}$. Más adelante vamos a considerar este problema, denominado *problema de la admisibilidad de marcos*.

2.3. Álgebras de von Neumann

En esta sección describimos algunos aspectos básicos de la teoría de álgebras de von Neumann. Nuestra introducción de esta vasta área de la matemática se circunscribirá estrictamente a aquellos resultados necesarios para desarrollos ulteriores. Más aun, algunos teoremas son enunciados de forma simplificada para no entorpecer la intención didáctica de la exposición.

2.3.1. Una introducción abstracta a las álgebras de von Neumann

Comenzamos con una descripción abstracta de esta clase de álgebras. Más adelante, indicaremos las relaciones entre esta presentación y las correspondientes álgebras de operadores (álgebras de von Neumann concretas) en espacios de Hilbert.

Un álgebra $C^* \mathcal{M}$ sobre el cuerpo de números complejos $\mathbb{C}$ (o más brevemente una C^* -álgebra) es un álgebra de Banach compleja con unidad y dotada de una involución isométrica $a \mapsto a^*$ es decir, para $\alpha, \beta \in \mathbb{C}$ y $a, b \in \mathcal{M}$ vale

$$\text{i) } (\alpha a + \beta b)^* = \bar{\alpha} a^* + \bar{\beta} b^*$$

$$\text{ii) } (ab)^* = b^* a^*$$

$$\text{iii) } (a^*)^* = a$$

$$\text{iv) } \|a^*\| = \|a\|$$

que satisface la condición adicional

$$\|a^*a\| = \|a\|^2.$$

Naturalmente, en tanto álgebras de Banach complejas, en las C^* -álgebras se pueden desarrollar las nociones de espectro y radio espectral; más aún, se pueden definir elementos normales y autoadjuntos con respecto a su involución. Un primer ejemplo de C^* -álgebra es $L(\mathcal{H})$, es decir la $*$ -álgebra de todos los operadores lineales y acotados en H para un espacio de Hilbert complejo, con su estructura de álgebra usual y con la involución dada por la adjunción usual de operadores. Evidentemente, cualquier subálgebra de $L(\mathcal{H})$ cerrada en norma y bajo la adjunción también cumplirá con las condiciones de las C^* -álgebras. El siguiente teorema, famoso en este contexto, afirma que las subálgebras de $L(\mathcal{H})$ cerradas en norma y bajo la adjunción son el modelo de toda C^* -álgebra. Antes necesitamos la siguiente

Definición 2.3.1. Dadas dos C^* -álgebras $\mathcal{M}, \mathcal{N}$ un $*$ -morfismo $\pi : \mathcal{M} \rightarrow \mathcal{N}$ es un morfismo de álgebras que preserva la involución. En el caso en que $\mathcal{N} = L(\mathcal{H})$ decimos que π es una representación de $\mathcal{M}$.

Teorema 2.3.2 (Gelfand-Naimark-Segal). Sea $\mathcal{M}$ una C^* -álgebra. Entonces existe un espacio de Hilbert $\mathcal{H}$ y una representación isométrica $\pi : \mathcal{M} \rightarrow L(\mathcal{H})$.

El Teorema 2.3.2 permite pensar toda C^* -álgebra como una C^* -subálgebra de $L(\mathcal{H})$, para algún espacio de Hilbert $\mathcal{H}$. En este sentido, las C^* -álgebras son álgebras de operadores. Estamos en condiciones de definir las álgebras de von Neumann

Definición 2.3.3. Un álgebra de von Neumann (a.v. Neumann) $\mathcal{M}$ es una C^* -álgebra para la cual existe un espacio de Banach E de forma que $\mathcal{M}$ es, en tanto espacio de Banach, el espacio dual de E .

De esta forma, las álgebras de von Neumann son “álgebras de operadores” dotadas de cierta estructura adicional: la topología débil*, que admiten como espacios duales a espacios de Banach.

Ejemplo 2.3.4. Un ejemplo importante de álgebra de von Neumann es el mismo $L(\mathcal{H})$. En efecto, es bien sabido que si $L^1(\mathcal{H})$ denota el espacio de Banach de los operadores traza dotados de la norma $\|\cdot\|_1$ dada por $\|A\|_1 = \text{tr}(|A|)$ entonces $(L^1(\mathcal{H}))^* = L(\mathcal{H})$ como espacio de Banach.

El siguiente resultado, que es una mejora del Teorema 2.3.2 en el caso de las álgebras de von Neumann .

Teorema 2.3.5. Sea $\mathcal{M}$ un álgebra de von Neumann . Entonces existe un espacio de Hilbert $\mathcal{H}$, una subálgebra de von Neumann $\tilde{\mathcal{M}} \subseteq L(\mathcal{H})$ y una representación $\pi : \mathcal{M} \rightarrow L(\mathcal{H})$ isométrica tal que $\pi(\mathcal{M}) = \tilde{\mathcal{M}}$ y tal que π es continua con respecto a las topologías débiles* de $\mathcal{M}$ y $\tilde{\mathcal{M}}$.

Recordemos que la relativización de la topología débil* a la bola cerrada de $\mathcal{M}$ resulta compacta. Si bien la introducción anterior de las a.v. Neumann es más técnica que la usual, pone de manifiesto los elementos más importantes de esta clase de álgebras. En lo que sigue recordamos algunas topologías conocidas de $L(\mathcal{H})$ y damos descripciones alternativas de las álgebras de von Neumann.

2.3.2. Topologías en $L(\mathcal{H})$. Álgebras de von Neumann concretas

Dado un espacio de Hilbert complejo $\mathcal{H}$, describimos las siguientes tres topologías en $L(\mathcal{H})$. Sea $a \in L(\mathcal{H})$ y consideremos en primer lugar, la *topología de la norma*, dada por:

$$\|a\| = \sup_{\xi \in \mathcal{H}} \frac{\|a\xi\|}{\|\xi\|}.$$

Si consideramos el caso $\mathcal{H} = L^2([0, 1])$ con la medida de Lebesgue, el álgebra $L^\infty([0, 1])$ actúa en $\mathcal{H}$ por multiplicación punto a punto y la norma de $f \in L^\infty([0, 1])$ es el supremo esencial de $|f|$. Entonces las funciones continuas $C([0, 1])$ forman una subálgebra cerrada en norma de $L^\infty([0, 1])$ en H y bajo la adjunción (esto vale en general para cualquier espacio compacto y cualquier medida finita).

La *topología fuerte* en $L(\mathcal{H})$ es aquella definida por las seminormas $a \mapsto \|a\xi\|$ cuando ξ recorre $\mathcal{H}$. Entonces una sucesión (o una red) a_n converge a a si y sólo si $a_n\xi$ converge a $a\xi$ en $\mathcal{H}$ para todo $\xi \in \mathcal{H}$. La topología fuerte es mucho más débil que la de la norma. De hecho se puede ver, en el ejemplo de $L^\infty([0, 1])$ y $C([0, 1])$ actuando en $L^2([0, 1])$, que $C([0, 1])$ es densa en $L^\infty([0, 1])$ en la topología fuerte. Para ver una sucesión que converge en la topología fuerte pero no en norma, tomemos x_n la función característica de $[0, 1/n]$ vista como elemento de $L^\infty([0, 1])$. Entonces a_n tiende a cero en la topología fuerte, pero $\|a_n - a_m\| = 1$ para $n \neq m$.

La *topología débil* en $L(\mathcal{H})$ es aquella definida por las seminormas $a \mapsto |\langle a\xi, \eta \rangle|$ cuando ξ y η recorren $\mathcal{H}$. La desigualdad de Cauchy-Schwartz muestra que la topología fuerte es más fuerte que la débil (y de ahí los nombres). De hecho la topología débil lo es tanto que la bola unitaria de $L(\mathcal{H})$ es débilmente compacta, lo que con frecuencia es muy útil. Para ver un ejemplo de sucesión que converge en la topología débil pero no en la fuerte, consideremos en $\ell^2(\mathbb{N})$ el operador “shift”

$$S((x_0, x_1, x_2, \dots)) = (0, x_0, x_1, x_2, \dots).$$

Se puede comprobar mediante cálculos directos que la sucesión $\{S^n\}_{n \in \mathbb{N}}$ converge en la topología débil pero no en la fuerte, porque $\|S^n\xi\| = \|\xi\|$ y además $\langle S^n\xi, \eta \rangle \rightarrow 0$ para todo $\xi, \eta \in \mathcal{H}$.

En lo que sigue introducimos algunas notaciones que necesitamos para enunciar un resultado que resume algunas caracterizaciones de las álgebras de von Neumann. Si $S \subseteq L(\mathcal{H})$

$$S' = \{x \in L(\mathcal{H}) : xs = sx \text{ para todo } s \in S\}$$

es el *conmutante* de S , y $S'' = (S')'$.

Teorema 2.3.6 (del doble conmutante). Sea S una subálgebra de $L(\mathcal{H})$ que contiene la identidad y que es cerrada bajo la involución de $L(\mathcal{H})$. Entonces

1. S es un álgebra de von Neumann si sólo si $S = S''$.
2. S'' es un álgebra de von Neumann y coincide con las clausuras de S en las topologías fuerte y débil de operadores.

Este teorema muestra que dos nociones, una analítica (clausura en la topología fuerte) y otra algebraica (ser igual al doble conmutante de un conjunto), coinciden para *-subálgebras de $L(\mathcal{H})$ que contienen al 1. Del teorema anterior vemos que una C^* -subálgebra de $L(\mathcal{H})$ es un álgebra de von Neumann si sólo si es fuertemente cerrada o, equivalentemente, si es débilmente cerrada.

Ejemplos 2.3.7. De álgebras de von Neumann

- i) El álgebra $L^\infty([0, 1])$ actuando en $L^2([0, 1])$ es su propio conmutante, por lo que claramente es un álgebra de von Neumann. Es un ejemplo además de álgebra de von Neumann *abeliana maximal*.
- ii) Si G es un grupo y $g \mapsto u_g$ es una representación unitaria de G , entonces el conmutante $\{u_g\}'$ es un álgebra de von Neumann.
- iii) Si H es de dimensión finita, se verifica que un álgebra de von Neumann $\mathcal{M}$ es simplemente una suma directa de álgebras de matrices correspondientes a una cierta descomposición ortogonal $H = H_1 \oplus \dots \oplus H_k$.

El centro $\mathcal{Z}(\mathcal{M})$ de un álgebra de von Neumann es abeliano. En dimensión finita será una suma directa de copias de $\mathbb{C}$, una por cada sumando de la descomposición. En general, usando el teorema espectral, von Neumann definió ([69]), en el caso separable, una noción de “integral directa” de espacios de Hilbert $\int_X^\oplus \mathcal{H}(\lambda) d\mu(\lambda)$ de manera que, por ejemplo, para $L^\infty([0, 1])$ en $L^2([0, 1])$ la correspondiente descomposición de $L^2([0, 1])$ sería $\int_{[0, 1]}^\oplus \mathcal{H}(\lambda) dx(\lambda)$ donde $\mathcal{H}(\lambda) \equiv \mathbb{C}$. El álgebra $\mathcal{M}$ respeta esta descomposición, y se llega a una noción de integral directa de álgebras de von Neumann: $\mathcal{M} = \int_X^\oplus \mathcal{M}(\lambda) d\lambda$ en $\int_X^\oplus \mathcal{H}(\lambda) d\lambda$, siendo la descomposición completa esencialmente única. Los individuos $\mathcal{M}(\lambda)$ tienen centro trivial (esto se puede intuir estudiando el caso finito dimensional). Entonces toda álgebra de von Neumann es una integral directa de álgebras de von Neumann de centro trivial, lo que hace que las álgebras de von Neumann de centro trivial sean especialmente interesantes.

Definición 2.3.8. Un álgebra de von Neumann $\mathcal{M}$ cuyo centro consiste de múltiplos escalares de la identidad es llamada un *factor*.

Como ejemplos de factores, citemos el más elemental, que es el mismo $L(\mathcal{H})$. Veamos brevemente un ejemplo no trivial de factor construido por Murray y von Neumann que no es $L(\mathcal{H})$. Sea Γ un grupo discreto tal que todas sus clases de conjugación salvo la de la identidad son infinitas (por ejemplo, el grupo libre en dos generadores). A estos grupos se los llama i.c.c., por “infinite conjugacy classes”. Consideremos $\ell^2(\Gamma)$ el espacio de Hilbert cuya base está indexada por Γ , sea $\gamma \mapsto u_\gamma$ la representación regular a izquierda de Γ , es decir, consideremos al grupo Γ representado en $\mathcal{L}(\ell^2(\Gamma))$ como operadores u_γ , donde actuar con u_γ es multiplicar los índices por γ (o sea, $u_\gamma(\{\lambda_g\}_{g \in \Gamma}) = \{\lambda_{\gamma g}\}_{g \in \Gamma}$).

Vistos como matrices en $\ell^2(\Gamma)$, con respecto a la base canónica indexada por $\gamma \in \Gamma$, los u_γ , y por ende todas las combinaciones lineales de los $\{u_\gamma\}$, son de la forma $x_{\gamma, \nu} = f(\gamma^{-1}\nu)$ para alguna función de soporte finito en Γ , es decir que son matrices con todas las diagonales constantes. Lo mismo vale para límites débiles de tales operadores salvo que f ya no será de soporte finito. Aplicando uno de tales operadores al elemento de la base correspondiente a la identidad vemos que $f \in \ell^2$. En efecto, imaginemos por comodidad que el vector correspondiente a la identidad es el vector $v_1 = (1, 0, 0, 0, \dots)$; entonces

$$(x_{\gamma, \nu}) \cdot v_1 = (f(\gamma_1), f(\gamma_2), \dots),$$

y como este vector tiene que estar en $\ell^2(\Gamma)$, debe ser

$$\sum_{\gamma \in \Gamma} f(\gamma)^2 = \|(x_{\gamma, \nu}) \cdot v_1\| < \infty.$$

Es por tanto conveniente y preciso escribir a los elementos de $\mathcal{M} = \{u_\gamma\}''$ como sumas

$$\sum_{\gamma \in \Gamma} f(\gamma)u_\gamma,$$

donde $f \in \ell^2$ (aunque no todas las funciones de ℓ^2 definen elementos de $\mathcal{M}$). No prestaremos atención al sentido de convergencia de la suma para simplificar ideas. En cualquier caso, para que $\sum_{\gamma \in \Gamma} f(\gamma)u_\gamma$ esté en el centro de $\mathcal{M}$, debe conmutar con u_ν para todo $\nu \in \Gamma$, lo que implica que $f(\nu\gamma\nu^{-1}) = f(\gamma)$, es decir que f es constante en las clases de conjugación. Pero f está en ℓ^2 y todas las clases de conjugación no triviales son infinitas, de manera que si f fuera no nula en un elemento no trivial, no podría estar en ℓ^2 . Entonces el soporte de f es la identidad, lo que implica que el centro de $\mathcal{M}$ es trivial, o sea que $\mathcal{M}$ es un factor. Lo llamaremos $\text{vN}(\Gamma)$.

Se puede ver que este factor no es isomorfo a ningún $L(\mathcal{H})$ observando que la función lineal $\text{tr}(\sum f(\gamma)u_\gamma) = f(1)$ tiene la propiedad $\text{tr}(ab) = \text{tr}(ba)$ y no es idénticamente cero. Es un hecho standard que no existe función con esa propiedad en $L(\mathcal{H})$ a menos que $\dim \mathcal{H} < \infty$. Como veremos más adelante, este es un ejemplo de factor de tipo II_1 .

2.3.3. Comparación de proyecciones y Clasificación de factores

La herramienta que utilizaron Murray y von Neumann para clasificar los factores fue la teoría de comparación de proyecciones que desarrollaron ellos mismos. Si bien estamos interesados en los factores de tipo II_1 , describimos algunos aspectos generales de esta teoría para poner en contexto tales factores.

Dada un álgebra de von Neumann $\mathcal{M}$, llamaremos $\mathcal{P}(\mathcal{M})$ al conjunto

$$\mathcal{P}(\mathcal{M}) = \{p \in \mathcal{M} : p^2 = p \text{ y } p^* = p\} \quad (2.21)$$

de proyecciones ortogonales de $\mathcal{M}$.

Definición 2.3.9 ([68]). Sean $p, q \in \mathcal{P}(\mathcal{M})$.

- a) Decimos que p y q son equivalentes, y lo notamos $p \sim q$, si existe una isometría parcial $u \in \mathcal{M}$ tal que $u^*u = p$ y $uu^* = q$.
- b) Decimos que q es mayor que p , y lo notamos $p \preceq q$, si existe $p_1 \in \mathcal{P}(\mathcal{M})$ tal que $p \sim p_1 \leq q$.

Se verifica que $\sim$ es una relación de equivalencia en $\mathcal{P}(\mathcal{M})$ y que la validez de $p \preceq q$ se mantiene si se reemplaza p y q por proyecciones equivalentes.

Notemos que la isometría parcial u debe estar en $\mathcal{M}$, de manera que la noción de comparación depende de $\mathcal{M}$. Si $\mathcal{M} = L(\mathcal{H})$, entonces dos proyecciones son equivalentes si y sólo si sus imágenes tienen la misma dimensión. A raíz de esto surgió la idea de que las clases de equivalencia de proyecciones constituyen una noción abstracta de dimensión para un factor arbitrario. El siguiente resultado sobre comparación de proyecciones confirma esto:

Teorema 2.3.10. Si $\mathcal{M}$ es un factor y p, q son proyecciones en $\mathcal{M}$ entonces $p \preceq q$ ó $q \preceq p$.

Definición 2.3.11. Una proyección $p \in \mathcal{P}(\mathcal{M})$ se dice *finita* si para todo $p_0 \in \mathcal{P}(\mathcal{M})$ que cumpla $p \sim p_0 \leq p$ resulta $p_0 = p$. En caso contrario, p se dice *infinita*. Si p es infinita y no tiene subproyecciones finitas, entonces decimos que p es *propiamente infinita*.

En el caso particular en que $\mathcal{M} = L(\mathcal{H})$, la noción de proyección finita o infinita coincide con que su rango sea respectivamente un subespacio de dimensión finita o infinita.

El orden de las proyecciones respeta el comportamiento que uno esperaría intuitivamente (en el sentido de la teoría de conjuntos) respecto de la finitud o infinitud:

Proposición 2.3.12. Sean $p, q \in \mathcal{P}(\mathcal{M})$. Si $p \preceq q$ y q es finito, entonces p es finito.

En particular si existe p infinito, resulta $1 (= Id_H)$ infinito.

Notemos que la forma en que hemos definido el concepto de proyección infinita es en realidad análoga a lo que se hace en teoría de conjuntos, es decir que una proyección es infinita si es equivalente a una subproyección propia.

Definición 2.3.13. Un álgebra de von Neumann $\mathcal{M}$ se dice *finita*, *infinita*, *propiamente infinita* si $1 \in \mathcal{M}$ es una proyección respectivamente finita, infinita o propiamente infinita.

Definición 2.3.14. Se dice que $p \in \mathcal{P}(\mathcal{M})$ es *minimal* si $p \neq 0$ y si $q \in \mathcal{P}(\mathcal{M})$, $q \leq p$ implica $q = 0$, o $q = p$.

Finalmente estamos en condiciones de enunciar la clasificación de los factores de Murray y von Neumann:

Definición 2.3.15. Un factor $\mathcal{M}$ se dice de tipo I, II o III si satisface las respectivas propiedades:

1. Tipo I: hay proyecciones minimales.
 - a) Tipo I_n , $n < \infty$: hay n proyecciones minimales equivalentes que suman 1.
 - b) Tipo I_∞ : hay infinitas proyecciones minimales equivalentes que suman 1.
2. Tipo II_1 : no hay proyecciones minimales y toda proyección es finita.
3. Tipo II_∞ : no hay proyecciones minimales y hay tanto proyecciones finitas como infinitas.
4. Tipo III: no hay proyecciones finitas no nulas.

La clasificación puede extenderse a álgebras de von Neumann arbitrarias. Omitiremos el enunciado de esta clasificación: en realidad la propiedad que utilizaremos con referencia al tipo de un álgebra es la existencia o no de estados y pesos traciales, de acuerdo con lo expresado en 2.3.25.

Un teorema importante de Murray y von Neumann dice que toda álgebra de von Neumann admite proyecciones P_{I_n} ($n \in \mathbb{N}$), P_{I_∞} , P_{II_1} , P_{II_∞} , P_{III} centrales, ortogonales dos a dos, que suman 1 y tales que $P_k \mathcal{M}$ es de tipo k para cada uno de los proyectores mencionados.

Observación 2.3.16. Se verifica que si un factor $\mathcal{M}$ tiene una traza tr (ver Definición 2.3.19) con $\text{tr}(x^*x) > 0$ para $x \neq 0$, entonces $\mathcal{M}$ es finito dimensional o de tipo II_1 . También son hechos conocidos que si el factor $\mathcal{M}$ es de tipo I entonces es de la forma $L(\mathcal{H}) \otimes \text{id}$ en $\mathcal{H} \otimes \mathcal{K}$, y que todo factor II_∞ es un producto tensorial de un factor II_1 y un factor infinito de tipo I.

2.3.4. Funcionales y pesos

En esta sección, como antes, $\mathcal{M}$ denotará un álgebra de von Neumann. Llamaremos $\mathcal{M}_+$ al conjunto de elementos positivos de $\mathcal{M}$. Los funcionales de un álgebra de von Neumann se comenzaron a estudiar originalmente porque se los puede ver como la generalización de la noción de integral. Los pesos serían desde este punto de vista las integrales “impropias”, es decir el caso en que la masa total del espacio es infinita.

Definición 2.3.17. Dada un funcional lineal ϕ sobre $\mathcal{M}$, diremos que es

- i) *positiva*, si $\phi(x^*x) \geq 0$ para todo $x \in \mathcal{M}$.
- ii) un *estado*, si es positiva y además $\phi(1) = 1$.

Definición 2.3.18. Un *peso* en un álgebra de von Neumann $\mathcal{M}$ es una función $\phi : \mathcal{M}_+ \longrightarrow [0, \infty]$, tal que $\phi(\lambda x + y) = \lambda\phi(x) + \phi(y)$ si $x, y \in \mathcal{M}_+$ y $\lambda \in [0, \infty)$ con la convención que $\lambda + \infty = \infty$ y $\lambda \cdot \infty = \infty$ si $\lambda > 0$, mientras que $0 \cdot \infty = 0$.

Definición 2.3.19. Un funcional lineal (o un peso) ϕ se dice

- i) *fiel*, si $0 \neq x \in \mathcal{M}_+ \implies \phi(x) > 0$.
- ii) *normal*, si cada vez que $x_i \nearrow x \in \mathcal{M}$ en la topología fuerte para $x_i \in \mathcal{M}_+$, resulta $\phi(x) = \sup_i \phi(x_i)$.
- iii) *tracial* o *traza*, si $\phi(x^*x) = \phi(xx^*)$, $\forall x \in \mathcal{M}$ (equivalente a que $\phi(xy) = \phi(yx)$, $\forall x, y \in \mathcal{M}$ en el caso en que ϕ es un funcional).
- iv) Un peso se dice *finito* si $\phi(x) < \infty$ para todo $x \in \mathcal{M}_+$.

Observación 2.3.20. Los funcionales normales de un álgebra de von Neumann $\mathcal{M}$, que notaremos con $\mathcal{M}_*$, forman un espacio de Banach con la norma usual. Se puede probar que $(\mathcal{M}_*)^* = \mathcal{M}$ (como dual de espacio de Banach). Esto justifica el nombre de *predual* de $\mathcal{M}$ para el espacio $\mathcal{M}_*$. Esta propiedad de tener un predual único permite caracterizar a las álgebras de von Neumann en abstracto, es decir sin tener que recurrir a una representación.

Proposición 2.3.21. Para un peso ϕ en $\mathcal{M}$ es equivalente:

- i) ser finito,
- ii) $\phi(1) < \infty$,
- iii) que exista un funcional positivo φ en $\mathcal{M}$ tal que $\varphi|_{\mathcal{M}_+} = \phi$.

Para un peso ϕ en $\mathcal{M}$, como no está definido en toda el álgebra, es útil definir los siguientes espacios asociados:

$$\mathcal{D}_\phi = \{x \in \mathcal{M}_+ : \phi(x) < \infty\}; \quad (2.22)$$

$$\mathcal{N}_\phi = \{x \in \mathcal{M} : \phi(x^*x) < \infty\}; \quad (2.23)$$

$$\mathcal{M}_\phi = \mathcal{N}_\phi^* \mathcal{N}_\phi = \left\{ \sum_{i=1}^n x_i^* y_i : x_i, y_i \in \mathcal{N}_\phi, n = 1, 2, \dots \right\}. \quad (2.24)$$

Donde no haya confusión posible, el subíndice ϕ será omitido. Notamos que $\mathcal{M}_\phi$ resulta una subálgebra autoadjunta de $\mathcal{M}$ (aunque no necesariamente $1 \in \mathcal{M}$, ni es cerrada para alguna topología particular)

Definición 2.3.22. Un peso ϕ en M se dice *semifinito* si $\mathcal{M}_\phi$ es débilmente denso en $\mathcal{M}$.

Semifinitud para un peso significa que hay “suficientes” elementos donde es finito. Puede probarse que para toda álgebra de von Neumann existe un peso normal, fiel y semifinito.

Proposición 2.3.23. Las siguientes condiciones son equivalentes para un peso ϕ :

- i) ϕ es semifinito
- ii) $1 = \sup\{e \in \mathcal{P}(\mathcal{M}) : \phi(e) < \infty\}$
- iii) existe una red creciente $\{x_i\}$ en $\mathcal{D}$ tal que $\|x_i\| < 1$ para todo i y $x_i \nearrow 1$.

El siguiente resultado de Haagerup [35] que caracteriza la “normalidad” de un peso.

Proposición 2.3.24. Para un peso ϕ en $\mathcal{M}$, las siguientes condiciones son equivalentes:

- i) ϕ es normal;
- ii) existe una red monótona creciente $\{\psi_i : i \in I\}$ de funcionales normales positivas tales que $\psi_i(x) \nearrow \phi(x)$ para todo $x \in \mathcal{M}_+$;
- iii) existe una familia $\{\psi_i : i \in J\}$ de funcionales positivas normales tales que $\phi(x) = \sum_{i \in J} \psi_i(x)$ para todo $x \in \mathcal{M}_+$ (donde la suma se interpreta como el límite de la red de sumas finitas);

La existencia de un peso semifinito tracial está relacionada con la clasificación de las álgebras de von Neumann.

Definición 2.3.25. Un álgebra de von Neumann se dice

1. *semifinita*, si existe un peso fiel, normal, semifinito y tracial;
2. *finita*, si existe un estado fiel, normal y tracial.
3. *infinita*, si no hay ningún peso fiel, normal, semifinito y tracial.

Las álgebras de von Neumann semifinitas son todas las de tipo I y II. Las finitas son las de tipo I_n , $n < \infty$, y II_1 . En este último caso el hecho no trivial, pero ya probado por Murray y von Neumann, es la existencia de un estado tracial en toda álgebra finita (tomando como definición de finita el hecho de que no hay proyecciones infinitas, como en la Definición 2.3.11).

Observación 2.3.26. Estamos particularmente interesados en los factores de tipo II_1 es decir, álgebras de von Neumann de centro trivial, dotadas de un estado fiel y tracial τ tal que para todo $\alpha \in [0, 1]$ existe una proyección $p \in \mathcal{P}(\mathcal{M})$ de forma que $\tau(p) = \alpha$ (en este caso pensamos a α como la dimensión del rango de p con respecto a $\mathcal{M}$). Más aún, dos proyecciones $p, q \in \mathcal{P}(\mathcal{M})$ son equivalentes en el sentido de Murray-von Neumann en $\mathcal{M}$ si y solo si $\tau(p) = \tau(q)$. Este último hecho resultará ser de importancia para nuestros desarrollos.

2.3.5. Esperanzas Condicionales

El nombre de “esperanza condicional” proviene de la Teoría de Probabilidades. La noción de esperanza condicional en álgebras de von Neumann extiende esta construcción al caso no abeliano. Las esperanzas condicionales han jugado un papel central en la teoría de las álgebras de von Neumann desde la década de 1950, y su importancia y utilidad se acrecentó con los trabajos de Takesaki y Connes a principios de la década de 1970, para luego en la década de 1980 convertirse en uno de los conceptos centrales de la teoría del índice.

Definición 2.3.27. Sean $\mathcal{N} \subseteq \mathcal{M}$ álgebras de von Neumann. Una proyección $\mathcal{E} : \mathcal{M} \rightarrow \mathcal{N}$ que cumple:

1. $\mathcal{E} \geq 0$; es decir, $x \in \mathcal{M}_+ \implies \mathcal{E}(x) \in \mathcal{N}_+$;
2. $\mathcal{E}(nxm) = n\mathcal{E}(x)m$, si $x \in \mathcal{M}$, $n, m \in \mathcal{N}$;
3. $(\mathcal{E}x)^*(\mathcal{E}x) \leq \mathcal{E}(x^*x)$, para todo $x \in \mathcal{M}$.

es una *esperanza condicional* de $\mathcal{M}$ en $\mathcal{N}$. Además,

- a) Una esperanza condicional $\mathcal{E}$ se dice *normal* si $\mathcal{E}(\sup_i x_i) = \sup_i \mathcal{E}(x_i)$, para toda red creciente $\{x_i\}_{i \in I} \subset \mathcal{M}$.
- b) Una esperanza condicional $\mathcal{E}$ se dice *fiel* si $\mathcal{E}(x^*x) = 0 \implies x = 0$ para todo $x \in \mathcal{M}$.
- c) Dado un peso fiel, normal y semifinito ϕ en $\mathcal{M}$, una esperanza condicional $\mathcal{E}$ se dice *ϕ -compatible* si $\phi \circ \mathcal{E} = \phi$ (es decir, $\phi(\mathcal{E}(x)) = \phi(x)$, $\forall x \in \mathcal{M}_+$).

Ejemplos 2.3.28. Mostraremos aquí algunos ejemplos elementales de esperanzas condicionales:

1. Cualquier estado φ de un álgebra de von Neumann $\mathcal{M}$ es una esperanza condicional $\varphi : \mathcal{M} \rightarrow \mathbb{C}$;
2. Si $\mathcal{M} = M_n(\mathbb{C})$, $\mathcal{N} = \mathbb{C}^n$ visto como las matrices diagonales, y $p_1, \dots, p_n$ son los proyectores minimales de $\mathcal{N}$, entonces $\mathcal{E} : \mathcal{M} \rightarrow \mathcal{N}$ dada por $\mathcal{E}(x) = \sum_{i=1}^n p_i x p_i$, es una esperanza condicional;
3. Si G es un grupo compacto de automorfismos de un álgebra de von Neumann $\mathcal{M}$,

$$\mathcal{E}(x) = \int_G \alpha_g(x) dg,$$

donde dg es la medida de Haar, es una esperanza condicional.

Teorema 2.3.29 (Takesaki). Sean $\mathcal{N} \subseteq \mathcal{M}$ álgebras de von Neumann, y sea ϕ un peso fiel normal y tracial en $\mathcal{M}$. Entonces existe una esperanza condicional $\mathcal{E} : \mathcal{M} \rightarrow \mathcal{N}$, que es ϕ -compatible.

El resultado anterior garantiza, en el caso de factores semifinitos, suficientes esperanzas condicionales que conmutan con la traza del factor.

2.4. Hechos básicos sobre C^* -álgebras abelianas

En esta sección describimos algunos resultados básicos acerca de C^* -álgebras abelianas.

2.4.1. Presentaciones de las álgebras abelianas

En esta sección veremos una forma conveniente de presentar las C^* -álgebras abelianas unitales, i.e. como álgebras de funciones continuas sobre un espacio topológico compacto Hausdorff.

Sea $\mathcal{A}$ una C^* -álgebra abeliana y unital. Consideramos el conjunto $\Gamma(\mathcal{A})$ formado por todos los $*$ -homomorfismos no nulos de $\mathcal{A}$ en $\mathbb{C}$ (como álgebras unitales). Notemos que si $\gamma \in \Gamma(\mathcal{A})$ entonces $\ker(\gamma)$ es un ideal maximal de $\mathcal{A}$ (puesto que $\mathcal{A}/\ker(\gamma) \approx \mathbb{C}$, i.e. el cociente es un cuerpo) y en consecuencia $\ker(\gamma)$ es cerrado. En este caso es bien sabido que γ es acotado. Recíprocamente, todo ideal maximal $\mathcal{I}$ es necesariamente cerrado e induce un carácter, que es la proyección al cociente $\rho : \mathcal{A} \rightarrow (\mathcal{A}/\mathcal{I}) \approx \mathbb{C}$. Más aún, a partir de lo anterior es claro que para todo $\gamma \in \Gamma(\mathcal{A})$ se tiene $\|\gamma\| = 1$. Notemos que entonces $\Gamma(\mathcal{A}) \subseteq (\mathcal{A})_1^*$, donde $(\mathcal{A})_1^*$ denota la bola unitaria cerrada del dual, y es sencillo verificar que este conjunto es cerrado bajo la topología débil* con lo que resulta compacto Hausdorff con la topología relativa.

A partir de estas observaciones se prueba el siguiente resultado sobre la transformada de Gelfand.

Teorema 2.4.1. Sea $\mathcal{A}$ una C^* -álgebra abeliana y unital y sea $\Gamma(\mathcal{A})$ su espacio de caracteres. Entonces

1. Si $a \in \mathcal{A}$ entonces $\sigma(a) = \{\gamma(a) : \gamma \in \Gamma(\mathcal{A})\}$.
2. La transformación $\Gamma : \mathcal{A} \rightarrow C(\Gamma(\mathcal{A}))$ de forma que $\Gamma(a)(\gamma) = \gamma(a)$ es un $*$ -isomorfismo de C^* -álgebras.

La transformación Γ del álgebra en el álgebra de funciones continuas sobre el espacio compacto Hausdorff $\Gamma(\mathcal{A})$ es llamada la transformada de Gelfand.

Ejemplos 2.4.2. Consideremos los siguientes ejemplos:

- a) En el caso particular en que $\mathcal{A} = C(X)$ donde X es un espacio topológico compacto Hausdorff se puede ver que el espacio de caracteres de $\mathcal{A}$ es homeomorfo a X vía la asociación $X \ni x \mapsto e_x$ donde e_x denota el morfismo “evaluación en x ”.
- b) Si $a \in L(\mathcal{H})_h$ es un operador autoadjunto entonces la C^* -álgebra generada por a , notada $C^*(a)$ es abeliana y el espacio de caracteres de $C^*(a)$ es homeomorfo a $\sigma(a)$. De hecho, si $\gamma \in \Gamma(\mathcal{A})$ entonces $\gamma(\gamma(a)I - a) = \gamma(a) - \gamma(a) = 0$ lo que implica que $\gamma(a)I - a$ no es un elemento inversible del álgebra $C^*(a)$. Recíprocamente, si $\lambda \in \sigma(a)$ entonces $\lambda I - a$ no es inversible con lo cual existe algún ideal maximal $\mathcal{I}$ (Zorn) que lo contiene. En este caso, si γ denota el carácter inducido por morfismo al cociente $\rho : C^*(a) \rightarrow C^*(a)/\mathcal{I} \approx \mathbb{C}$ entonces $\gamma(\lambda I - a) = 0$, pero entonces $\gamma(a) = \lambda$. De esta forma, la función continua $\gamma \mapsto \gamma(a)$ es el homeomorfismo buscado. Así que obtenemos una versión del cálculo funcional continuo $C^*(a) \approx C(\sigma(a))$ de forma que el operador a se corresponde con la función identidad sobre $\sigma(a)$.

La descripción anterior muestra que los resultados sobre las álgebras de funciones continuas se traducen a resultados sobre C^* -álgebras abelianas. En particular, si X, Y son espacio topológicos compactos, es conocido en hecho de que todo $*$ -homomorfismo $\Phi : C(X) \rightarrow C(Y)$ le corresponde una función continua $\phi : Y \rightarrow X$ de forma $\Phi(f) = f \circ \phi$. Más aún, Φ es monomorfismo (resp. epimorfismo) si y solo si ϕ es suryectiva (resp. inyectiva). De esta forma deducimos

Corolario 2.4.3. Sea $\Phi : \mathcal{A} \rightarrow \mathcal{B}$ un $*$ -morfismo entre álgebras abelianas y unitales. Entonces Φ induce una función continua $\phi : \Gamma(\mathcal{B}) \rightarrow \Gamma(\mathcal{A})$ de forma que $\gamma(\Phi(a)) = \phi(\gamma)(a)$ para cada $a \in \mathcal{A}$ y $\gamma \in \Gamma(\mathcal{B})$. Más aún, Φ es monomorfismo (epimorfismo) si y solo si ϕ es suryectiva (resp. inyectiva).

2.4.2. Medidas espectrales y representaciones

Sea $(X, \mathfrak{N})$ un conjunto dotado de una σ -álgebra de subconjuntos, y $L(\mathcal{H})$ el espacio de los operadores lineales acotados sobre el espacio de Hilbert $\mathcal{H}$. Una *medida espectral* E en $(X, \mathfrak{N})$ a valores en $L(\mathcal{H})$ es una función $E : \mathfrak{N} \rightarrow \mathcal{P}(L(\mathcal{H}))$ de conjuntos medibles a valores proyectores en $L(\mathcal{H})$ de forma que satisface

1. $E(\emptyset) = 0, E(X) = 1$
2. $E(\bigcup_{i=1}^{\infty} \Delta_i) = \sum_{i=1}^{\infty} E(\Delta_i)$ (convergencia en la topología fuerte de operadores) para toda sucesión $(\Delta_i)_{i=1}^{\infty} \subseteq \mathfrak{N}$ dos a dos disjuntos.

Notemos que de la condición 2 deducimos que $E(\Delta_1)E(\Delta_2) = 0$ si $\Delta_1 \cap \Delta_2 = \emptyset$.

Observación 2.4.4. Supongamos que $f : X \rightarrow \mathbb{C}$ es una función simple $f = \sum_{i=1}^n \alpha_i \chi_{\Delta_i}$ en X (donde estamos suponiendo que $\{\Delta_i\}_{i=1}^n$ forma una partición de X). Entonces podemos considerar la integral de f con respecto a la medida espectral E de forma que

$$\int_X f dE = \sum_{i=1}^n \alpha_i E(\Delta_i) \in L(\mathcal{H}), \quad \left\| \int_X f dE \right\| = \max_{1 \leq i \leq n} |\alpha_i|. \quad (2.25)$$

En el caso en que X sea un espacio topológico compacto entonces vemos que podemos definir la integral de una función continua $f \in C(X)$ como límite de integrales de funciones simples (en el sentido de Lebesgue, partiendo la imagen acotada de f de forma diádica y considerando las funciones simples aproximantes).

El siguiente teorema es “el teorema espectral” para álgebras abelianas. Contiene, como caso particular, el teorema espectral para operadores normales y será utilizado en varias partes de nuestro trabajo.

Teorema 2.4.5. Sea $\mathcal{A}$ una C^* -álgebra abeliana y sea $\Gamma(\mathcal{A})$ su espacio de caracteres. Si $\pi : \mathcal{A} \rightarrow L(\mathcal{H})$ es una representación de $\mathcal{A}$ entonces existe una medida espectral E definida sobre la σ -álgebra de conjuntos Borelianos en $\Gamma(\mathcal{A})$ de forma que

1. $E(\Delta) \in \pi(\mathcal{A})''$ (el álgebra de von Neumann generada por $\pi(\mathcal{A})$).
2. $E(\Delta) = \bigwedge \{E(\mathcal{O}) : \Delta \subseteq \mathcal{O}, \mathcal{O} \text{ es abierto}\}.$

3. Sea $a \in \mathcal{A}$ y sea $f \in C(\Gamma(\mathcal{A}))$ la función asociada a a por el isomorfismo (Teorema 2.4.1) $\mathcal{A} \approx C(\Gamma(\mathcal{A}))$. Entonces

$$\pi(a) = \int_{\Gamma(\mathcal{A})} f \, dE.$$

4. Para cada $\xi \in \mathcal{H}$, $\Delta \mapsto \langle E(\Delta)\xi, \xi \rangle := \mu_\xi(\Delta)$ es una medida de forma que, si a y f son como en el ítem anterior entonces

$$\langle \pi(a)\xi, \xi \rangle = \int_{\Gamma(\mathcal{A})} f \, d\mu_x.$$

Observación 2.4.6. Con las notaciones del teorema anterior, sea $f : \Gamma(\mathcal{A}) \rightarrow \mathbb{C}$ una función medible (Borel) y uniformemente acotada. Entonces tiene sentido (ver la Observación 2.4.4) la integral

$$\int_{\Gamma(\mathcal{A})} f \, dE \in \pi(\mathcal{A})''.$$

Además, este proceso de integración tiene la siguiente propiedad (de continuidad) denominada σ -normalidad: si $(f_n)_{n \in \mathbb{N}}$ es una sucesión de funciones medibles, que está uniformemente acotada y tal que las funciones f_n convergen de forma creciente a una función medible (y uniformemente acotada) f entonces

$$\int_{\Gamma(\mathcal{A})} f_n \, dE \xrightarrow{\text{SOT}} \int_{\Gamma(\mathcal{A})} f \, dE$$

y de forma creciente, con respecto al orden usual de operadores.

Ejemplo 2.4.7. Sea $\mathcal{M} \subseteq L(\mathcal{H})$ un álgebra de von Neumann y sea $a \in \mathcal{M}_h$ un operador autoadjunto. Si consideramos $C^*(a)$ la C^* -álgebra unital generada por a entonces se tiene la inclusión $C^*(a) \subseteq \mathcal{M} \subseteq L(\mathcal{H})$. Por el Teorema 2.4.5, el morfismo “inclusión” ($C^*(a) \subseteq L(\mathcal{H})$) induce una medida espectral E sobre (los conjuntos medibles Borel de) $\sigma(a)$ (por el Ejemplo 2.4.2 ítem b), en este caso se tiene $\Gamma(C^*(a)) = \sigma(a)$ a valores en $\mathcal{P}(\mathcal{M})$, pues $C^*(a)'' \subseteq \mathcal{M}$. De esta forma, la medida espectral satisface

$$a = \int_{\sigma(a)} x \, dE.$$

En el caso especial de operadores autoadjuntos $a \in \mathcal{M}_h \subseteq L(\mathcal{H})$ en un álgebra de von Neumann, existe una presentación espectral de a distinta a la del ejemplo anterior, basada en su resolución espectral.

Definición 2.4.8. Sea $\mathcal{M}$ un álgebra de von Neumann. Una familia de proyecciones $\{P_\lambda\}_{\lambda \in \mathbb{R}} \subseteq \mathcal{P}(\mathcal{M})$ es una resolución espectral a derecha si satisface

1. $P_\lambda \leq P_\mu$ siempre que $\lambda \geq \mu$.
2. $P_\lambda = \lim_{\lambda' \rightarrow \lambda^+} P_{\lambda'}$ en la topología fuerte de operadores.
3. $\lim_{\lambda \rightarrow -\infty} P_\lambda = 0$ y $\lim_{\lambda \rightarrow \infty} P_\lambda = I$.

Si existe $a > 0$ tal que $P_{-a} = 0$ y $P_a = 1$ entonces decimos que la resolución es acotada.

Es evidente que en el caso en que la resolución espectral $\{P_\lambda\}_{\lambda \in \mathbb{R}}$ resulte acotada, entonces basta considerar un truncamiento compacto $\{P_\lambda\}_{\lambda \in I}$ para un cierto intervalo compacto $I \subseteq \mathbb{R}$. En adelante consideraremos sólo este truncamiento.

Ejemplo 2.4.9. Siguiendo con la notación del Ejemplo 2.4.7, si $a \in \mathcal{M}_h$ es un operador autoadjunto no nulo y E denota la medida espectral sobre $\sigma(a) \subseteq \mathbb{R}$ entonces si definimos

$$P_\lambda = E(\lambda, \infty), \quad \lambda \in \mathbb{R}$$

entonces resulta que $\{P_\lambda\}_{\lambda \in \mathbb{R}}$ es una resolución acotada (aquí podemos poner como cota $\|a\| > 0$) a derecha.

Observación 2.4.10. De forma dual, una familia $\{P_\lambda\}_{\lambda \in \mathbb{R}} \subseteq \mathcal{P}(\mathcal{M})$ es una resolución a izquierda (acotada) si la familia $\{I - P_\lambda\}_{\lambda \in \mathbb{R}}$ es una resolución a derecha (acotada).

Dada una resolución espectral a izquierda y acotada $\{P_\lambda\}_{\lambda \in [\alpha, \beta]} \subseteq \mathcal{M}$ en un álgebra de von Neumann $\mathcal{M} \subseteq L(\mathcal{H})$ entonces, las sumas de Riemann

$$\sum_{i=0}^n x_i (P_{x_{i+1}} - P_{x_i})$$

para $\alpha = x_0 < \dots < x_n = \beta$, convergen en norma de operadores a un operador autoadjunto $a \in \mathcal{M}_h$, cuando $\max_i |x_{i+1} - x_i| \rightarrow 0$. Más aún, se puede verificar que en este caso, si E denota la medida espectral inducida sobre $\sigma(a)$ por la representación “inclusión” (ver Ejemplo 2.4.7) entonces $P_\lambda = E(\lambda, \infty)$.

El siguiente resultado resume la relación que hay entre resoluciones espectrales a izquierda de un álgebra de von Neumann $\mathcal{M}$ y sus elementos autoadjuntos.

Teorema 2.4.11. Sea $\mathcal{M}$ un álgebra de von Neumann. Entonces las resoluciones espectrales a izquierda (a derecha) y acotadas en $\mathcal{M}$ están en correspondencia biunívoca con el conjunto de operadores autoadjuntos de $\mathcal{M}$. (Esta correspondencia está descrita en el Ejemplo 2.4.9 y la Observación 2.4.10).

2.4.3. Cálculo funcional en varias variables

En el desarrollo de algunos de los tópicos del trabajo vamos a utilizar el cálculo funcional en una o varias variables. En lo que sigue hacemos una primera descripción de este cálculo, en el caso básico de las C^* -álgebras abstractas finitamente generadas.

Sea $\mathcal{A}$ una C^* -álgebra unital y sean $a_1, \dots, a_n$ elementos de la parte autoadjunta de $\mathcal{A}$, notada $\mathcal{A}_{sa}$ que conmutan entre sí. Si $\mathcal{B} = C^*(a_1, \dots, a_n)$ denota la C^* -subálgebra unital de $\mathcal{A}$ generada por estos elementos entonces $\mathcal{B}$ es una C^* -álgebra abeliana finitamente generada. Entonces existe un espacio compacto Hausdorff X tal que $\mathcal{B}$ es $*$ -isomorfa a $C(X)$. Es sabido que X es (salvo homeomorfismos) el espacio de caracteres de $\mathcal{B}$, i.e. el conjunto de homomorfismos $\gamma : \mathcal{B} \rightarrow \mathbb{C}$, dotados de la topología débil*.

En el caso de un operador $a \in \mathcal{A}_{sa}$, X es homeomorfo a $\sigma(a)$. En general, los caracteres γ del álgebra $\mathcal{B}$ están asociados de forma continua e inyectiva a n -uplas $\gamma \mapsto (\gamma(a_1), \dots, \gamma(a_n)) \in \prod_{i=1}^n \sigma(a_i)$ por restricción a las C^* -subálgebras abelianas de $\mathcal{B}$ generadas por los a'_i s. Así, X es

homeomorfo a su imagen bajo esta función continua, que es denominada *el espectro conjunto* de la familia $(a_1, \dots, a_n)$ y lo denotamos $\sigma(a_1, \dots, a_n)$. A partir de lo anterior se tiene un $*$ -isomorfismo natural $\mathcal{B} \approx C(\sigma(a_1, \dots, a_n))$

Ejemplo 2.4.12. Sea $\mathcal{A}$ una C^* -álgebra unital y $a, b \in \mathcal{A}_{sa}$ tales que $ab = ba$. Aunque el espectro conjunto $\sigma(a, b)$ es un subconjunto cerrado de $\sigma(a) \times \sigma(b)$ puede ser, en general, bastante más pequeño que el producto de los espectros. Por ejemplo $\sigma(a, b)^0 = \emptyset$. De hecho se $b = f(a)$ para una función continua $f : \sigma(a) \rightarrow \mathbb{R}$ entonces es sencillo ver que $\sigma(a, b) = \{(x, f(x)), x \in \sigma(a)\} = \text{Graph}(f)$. Notemos que en este caso $C^*(a) = C^*(a, b)$ y $\sigma(a)$ es homeomorfo a $\sigma(a, b)$ de forma obvia. $\square$

Podemos describir ahora el cálculo funcional en varias variables. Sea $f : \sigma(a_1, \dots, a_n) \rightarrow \mathbb{R}$ una función continua definida en el espectro conjunto de los a_i 's. Entonces existe un elemento, denotado por $f(a_1, \dots, a_n)$, que corresponde a la función continua f bajo el $*$ -isomorfismo $\mathcal{B} \approx C(\sigma(a_1, \dots, a_n))$. Notemos que por el teorema de extensión de Tietze podemos considerar funciones definidas en $\prod_{i=1}^n \sigma(a_i) \subseteq \mathbb{R}^n$ sin pérdida de la generalidad (sin embargo en este caso el álgebra $C(\prod_{i=1}^n \sigma(a_i))$ puede ser más grande que $\mathcal{B}$.) De esta forma la asociación $f \mapsto f(a_1, \dots, a_n)$ es un $*$ -epimorfismo de $C(\prod_{i=1}^n \sigma(a_i))$ sobre $\mathcal{B}$. En el caso de familias compuestas por un operador, la construcción anterior es el llamado cálculo funcional continuo (en su versión abstracta).

Ejemplo 2.4.13. Si $(a_i)_{i=1, \dots, m}$ es una familia abeliana de matrices autoadjuntas en $M_n(\mathbb{C})$ entonces existe una matriz unitaria (posiblemente no única) $U \in M_n(\mathbb{C})$ tal que

$$U^* a_i U = D_{\lambda(a_i)}, \quad i = 1, \dots, m,$$

donde D_x denota la matriz diagonal con diagonal principal $x \in \mathbb{R}^n$. En este caso $\lambda(a_i)$ es el vector de autovalores correspondientes a a_i , contados con multiplicidad, en algún orden dependiendo de U . Consideremos la matriz $A \in M_{n,m}$ dada por

$$C_i(A) = \lambda(a_i), \quad C_i(B) = \lambda(b_i), \quad i = 1, \dots, m.$$

Notemos que si U, W son dos matrices unitarias que diagonalizan la familia $(a_i)_{i=1, \dots, m}$ simultáneamente, entonces existe una matriz de permutación Q tal que

$$U^* a_i U = Q^* W^* a_i W Q, \quad i = 1, \dots, m.$$

Así, si $A' \in M_{n,m}$ denota la matriz cuyas columnas $C_i(A')$ son las diagonales principales de las matrices $W^* a_i W$, entonces $A = Q A'$. Esto muestra que el conjunto de filas $R(A)$ no depende de la matriz unitaria U . Este conjunto es justamente el espectro conjunto de la familia $\sigma(a_1, \dots, a_m)$ en este caso. Más aún, si $f : V \rightarrow \mathbb{C}$ es tal que $\sigma(a_1, \dots, a_m) \subseteq V$ entonces

$$f(a_1, \dots, a_m) = U D_x U^*$$

donde D_x es la matriz diagonal con diagonal principal $x = (f(R_1(A)), \dots, f(R_n(A))) \in \mathbb{C}^n$. Notemos que $f(a_1, \dots, a_m) \in M_n(\mathbb{C})$ no depende de la matriz unitaria U .

2.4.4. Valores singulares, preorden espectral y (sub)mayorización

Sea $(\mathcal{M}, \tau)$ un álgebra de von Neumann semifinita. Los τ -valores singulares (nuestra referencia para esta noción es el trabajo de Fack [32]) $\mu_x(t)$ de $x \in \mathcal{M}$ están definidos para cada $t \in \mathbb{R}_0^+$, donde $\mathbb{R}_0^+$ son los números reales no negativos, por la ecuación

$$\mu_x(t) = \inf\{\|xe\| : e \in \mathcal{P}(\mathcal{M}), \tau(1-e) \leq t\}.$$

Para cada $x \in \mathcal{M}$, la función $\mu_x : \mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+$ es decreciente y continua a derecha. Si $x, y \in \mathcal{M}$ entonces

$$|\mu_x(t) - \mu_y(t)| \leq \|x - y\|$$

lo que muestra una dependencia continua de los valores singulares con respecto a la norma de operadores. Además se tiene la siguiente dependencia continua con respecto a las normas $\|\cdot\|_1$ respectivas

$$\|\mu_a - \mu_b\|_1 \leq \|a - b\|_1. \quad (2.26)$$

Si $a \in \mathcal{M}^+$ es un operador positivo entonces tenemos

$$\mu_a(t) = \min\{s \in \mathbb{R}_0^+ : \tau(p^a(s, \infty)) \leq t\}, \quad (2.27)$$

donde $p^a(s, \infty)$ denota la proyección espectral de a correspondiente al intervalo (s, ∞) . Esta última caracterización de los valores singulares de los operadores positivos muestra una propiedad interesante: si $a, b \in \mathcal{M}^+$ son tales que $\tau(p^a(s, \infty)) = \tau(p^b(s, \infty))$ entonces, $\mu_a = \mu_b$. En particular, vemos que $\mu_a = \mu_{uau^*}$ para cada operador unitario $u \in \mathcal{M}$. Más aún, de este hecho y la dependencia continua de los valores singulares en la norma vemos que $\mu_a = \mu_b$, siempre que exista una sucesión $\{u_n\}_{n \in \mathbb{N}}$ de operadores unitarios en $\mathcal{M}$ tal que $\|b - u_n a u_n^*\| \rightarrow 0$. Es decir, los valores singulares son invariantes bajo equivalencia aproximadamente unitaria. Si asumimos además que $\mathcal{M}$ es un factor finito y que τ es la traza normalizada, fiel y normal en $\mathcal{M}$ entonces, Kamei probó en [49] que existe un recíproco de este hecho. Recogemos estas observaciones en la siguiente proposición, pero primero fijamos alguna notación:

1. Dado $a \in \mathcal{M}$, denotamos por $\mathcal{U}_{\mathcal{M}}(a) = \{uau^* : u \in \mathcal{U}_{\mathcal{M}}\}$ la órbita unitaria de a en $\mathcal{M}$.
2. Denotamos por $\overline{\mathcal{U}_{\mathcal{M}}(a)} = \overline{\mathcal{U}_{\mathcal{M}}(a)}^{\|\cdot\|_\infty}$ la clausura en norma de $\mathcal{U}_{\mathcal{M}}(a)$. Luego, $b \in \mathcal{M}$ es aproximadamente unitariamente equivalente a a si $b \in \overline{\mathcal{U}_{\mathcal{M}}(a)}$ o, equivalentemente, si $\overline{\mathcal{U}_{\mathcal{M}}(a)} = \overline{\mathcal{U}_{\mathcal{M}}(b)}$.

Proposición 2.4.14. Sea $a, b \in \mathcal{M}^+$. Entonces

1. Si $b \in \overline{\mathcal{U}_{\mathcal{M}}(a)}$, entonces $\mu_a = \mu_b$.
2. Si $(\mathcal{M}, \tau)$ es un factor finito entonces, si $\mu_a = \mu_b$ se tiene que $b \in \overline{\mathcal{U}_{\mathcal{M}}(a)}$.

Notemos que el ítem 2 de la proposición anterior no es cierto en general en un factor semifinito, pero no finito (ver Ejemplo 3.1.14). Finalizamos esta sección con las definiciones de tres preórdenes diferentes que vamos a considerar más adelante.

Definición 2.4.15. Sean $a, b \in \mathcal{M}^+$. Decimos que b *domina espectralmente* a a , si se verifica que $\mu_a(t) \leq \mu_b(t)$, para todo $t \geq 0$. En este caso escribimos $a \preceq b$.

Por ejemplo, es bien sabido que si $a \leq b$ entonces $a \preceq b$. El siguiente resultado que aparece en [19] describe algunas caracterizaciones del preorden espectral en el caso en que $\mathcal{M}$ es un factor.

Teorema 2.4.16. Sea $(\mathcal{M}, \tau)$ un factor semifinito y sean $a, b \in \mathcal{M}^+$ operadores positivos. Entonces los siguientes son equivalentes:

1. b domina espectralmente a a .
2. $P^a(s, \infty) \preceq P^b(s, \infty)$ para $s \in \mathbb{R}_0^+$.
3. $\tau(P^a(s, \infty)) \leq \tau(P^b(s, \infty))$ para $s \in \mathbb{R}_0^+$.
4. $\tau(g(a)) \leq \tau(g(b))$ para toda función continua y creciente $g : \mathbb{R} \rightarrow \mathbb{R}$.

Definición 2.4.17. Sean $a, b \in \mathcal{M}^+$. Decimos que a está *submayorizado* por b , y notamos $a \prec_w b$, si

$$\int_0^s \mu_a(t) dt \leq \int_0^s \mu_b(t) dt, \quad \text{para cada } s \geq 0.$$

El siguiente resultado que aparece en [49] establece una caracterización de la submayorización en factores semifinitos que vamos a utilizar más adelante (desigualdades de tipo Jensen)

Teorema 2.4.18. Sea $(\mathcal{M}, \tau)$ un factor semifinito y sean $a, b \in \mathcal{M}^+$ operadores positivos. Entonces $a \prec_w b$ si y solo si

$$\tau(f(a)) \leq \tau(f(b))$$

para toda función convexa creciente $f : \mathbb{R} \rightarrow \mathbb{R}$.

Si $a \prec_w b$ y además vale que $\tau(a) = \tau(b)$ entonces decimos que a está *mayorizado* por b y lo notamos $a \prec b$. Así, si $a, b \in \mathcal{M}^+$ tenemos que $a \leq b \Rightarrow a \preceq b \Rightarrow a \prec_w b$. El siguiente resultado que aparece en [41] tiene relación con el contenido del capítulo 8. Pero primero consideramos la siguiente noción: sea $(\mathcal{M}, \tau)$ un álgebra de von Neumann semifinita y $T : \mathcal{M} \rightarrow \mathcal{M}$ una transformación lineal. Decimos que T es una transformación doble estocástica (resp subestocástica) si es positiva i.e. $T(a) \geq 0$ siempre que $a \geq 0$, unital i.e. $T(1) = 1$, y preserva trazas i.e. $\tau(T(a)) = \tau(a)$, $a \in \mathcal{M}$ (resp reduce trazas i.e. $\tau(T(a)) \leq \tau(a)$ para todo $a \in \mathcal{M}^+$). Ya hemos considerado la noción de transformación doble estocástica para un factor finito de tipo I_n (ver antes del teorema 2.1.10).

Teorema 2.4.19. Sea $(\mathcal{M}, \tau)$ un factor semifinito y sean $a, b \in \mathcal{M}^+$ operadores positivos. Los siguientes son equivalentes:

1. $a \prec_w b$.
2. Existe T doble subestocástica en $(\mathcal{M}, \tau)$ tal que $T(b) = a$.
3. $\tau(f(a)) \leq \tau(f(b))$ para toda función convexa continua y creciente $f : \mathbb{R} \rightarrow \mathbb{R}$.

Análogamente, se tiene el siguiente teorema que caracteriza a la mayorización

Teorema 2.4.20. Sea $(\mathcal{M}, \tau)$ un factor semifinito y sean $a, b \in \mathcal{M}^+$ operadores positivos. Los siguientes son equivalentes:

1. $a \prec b$.
2. $a \in \overline{\text{conv}}(\mathcal{U}_{\mathcal{M}}(b))$.
3. Existe $T \in DS(\mathcal{M}, \tau)$ tal que $T(b) = a$.
4. $\tau(f(a)) \leq \tau(f(b))$ para toda función convexa continua $f : \mathbb{R} \rightarrow \mathbb{R}$.

Tanto el preorden espectral, así como la submayorización y la mayorización surgen naturalmente en varios contextos de la teoría de operadores (ver[7, 11, 12, 19, 33, 32], HiaiN); por otra parte implican una familia de desigualdades traciales, que son de interés para las aplicaciones. Algunos ejemplos recientes son el estudio de desigualdades de tipo Young [33] y de tipo Jensen [7, 12].

Vamos a necesitar el siguiente resultado de Hiai y Nakamura [42], relativo a funciones en espacios de medida finita (X, ν) . En este caso, las funciones uniformemente acotadas son consideradas como operadores actuando por multiplicación en el espacio de Hilbert $L^2(X, \nu)$ y los valores singulares de las funciones están definidos con respecto a la traza fiel y normal inducida por ν en el álgebra de von Neumann finita $L^\infty(X, \nu)$.

Proposición 2.4.21. Sea (X, ν) un espacio de probabilidad y sean $f, g \in L^\infty(X, \nu)$ funciones positivas. Entonces $f \prec_w g$ si y solo si existe $h \in L^\infty(X, \nu)$ tal que $f \leq h \prec g$.

2.4.5. Hechos básicos de la teoría de Choquet

En esta sección describimos algunos hechos básicos relacionados con la teoría de integrales de borde de Choquet para el caso particular de $\mathbb{R}^n$. Nuestra referencia es el libro de Takesaki [66]. En lo que sigue, $K \subseteq \mathbb{R}^n$ denotará un conjunto convexo compacto. Recordemos que una función $f : K \rightarrow \mathbb{R}$ es convexa si se verifica

$$f(\lambda x + (1 - \lambda)y) \leq \lambda f(x) + (1 - \lambda)f(y), \quad \text{para } x, y \in K. \quad (2.28)$$

Denotamos por $\mathcal{B}(K)$ el conjunto de las funciones continuas.

Proposición 2.4.22. El conjunto $\{f - g : f, g \in \mathcal{B}(K)\}$ es un espacio vectorial uniformemente denso en $C_{\mathbb{R}}(K)$, el álgebra de las funciones continuas sobre K a valores reales.

Notamos por $M_+^\sim(\mathbb{R}^n)$ el espacio de medidas de Borel regulares, finitas y positivas μ de soporte compacto. Si $\mu \in M_+^\sim(\mathbb{R}^n)$ es una medida y $f : \mathbb{R}^n \rightarrow \mathbb{C}$ es una función μ -integrable entonces notamos

$$\mu(f) = \int_{\mathbb{R}^n} f(\eta) d\mu(\eta).$$

Si $\mu, \nu \in M_+^\sim(\mathbb{R}^n)$ son tales que $\nu(\pi_i) = \mu(\pi_i)$ para $1 \leq i \leq n$ y $\nu(\mathbb{R}^n) = \mu(\mathbb{R}^n)$ entonces notamos $\nu \sim \mu$.

La siguiente es una noción que será de importancia en nuestro desarrollo de la mayorización conjunta en factores Π_1 .

Definición 2.4.23. Decimos que ν mayoriza a μ , y notamos $\mu \prec \nu$, si para cada $\nu_1, \dots, \nu_m \in M_+^\sim(\mathbb{R}^n)$ con $\sum_{i=1}^m \nu_i = \nu$ existen $\mu_1, \dots, \mu_m \in M_+^\sim(\mathbb{R}^n)$ tales que $\sum_{i=1}^m \mu_i = \mu$, $\nu_i(1) = \mu_i(1)$ y $\nu_i(\pi_j) = \mu_i(\pi_j)$ para $1 \leq i \leq m$, $1 \leq j \leq n$.

Sean $\mu, \nu \in M_+^\sim(\mathbb{R}^n)$ y tales que están soportadas en $\mathcal{K}$. En la teoría de integrales de borde, el resultado anterior se interpreta como la afirmación de si $\mu \prec \nu$ entonces ν está concentrada más cerca del borde (extremal) del convexo compacto K .

La relación $\prec$ de la definición 2.4.23 no es usualmente llamada “mayorización” en la literatura, pero parece ser una terminología adecuada en este contexto debido al Teorema 8.2.5.

La siguiente caracterización de la mayorización de medidas será de utilidad para nuestro estudio de la mayorización conjunta en factores de tipo II_1 .

Teorema 2.4.24. Sea $\mu, \nu \in M_+^\sim(\mathbb{R}^n)$. Entonces $\mu \prec \nu$ si y solo si $\mu(f) \leq \nu(f)$ para cada función continua convexa $f : \mathbb{R}^n \rightarrow \mathbb{R}$.

Observación 2.4.25. Notemos que como consecuencia de la Proposición 2.4.22 y el Teorema 2.4.24 decimos que, si $\mu, \nu \in M_+^\sim(\mathbb{R}^n)$ son tales que $\mu \prec \nu$ y $\nu \prec \mu$ entonces $\mu(f) = \nu(f)$ para toda $f \in C(K)$, es decir que $\mu = \nu$. Así la mayorización entre medidas es un orden.

Capítulo 3

Refinamientos de resoluciones espectrales

Para una descripción conceptual de los resultados incluidos dentro de este capítulo, así como la relación entre éstos y otros resultados ver la sección “Contexto general del trabajo” del Capítulo 1.

Organización del capítulo En la sección 3.1.1 introducimos la noción de refinamiento entre resoluciones espectrales a derecha y acotadas en factores II_1 arbitrarios. Enunciamos uno de los resultados principales al respecto, que será utilizado en la sección siguiente; su demostración, más bien técnica, se desarrolla en la sección 3.2.

En la sección 3.1.2 usamos el resultado anterior sobre refinamientos para desarrollar el modelado de operadores y para el estudio de la dominación espectral y de la submayorización en factores II_1 . En la sección 3.1.3 hacemos algunas consideraciones sobre el caso semifinito.

En la sección 3.2 probamos el teorema de refinamientos adelantado en la sección 3.1.1.

Notaciones Sea $\mathcal{H}$ un espacio de Hilbert y sea $L(\mathcal{H})$ el álgebra de todos los operadores acotado en $\mathcal{H}$. En esta sección, el par $(\mathcal{M}, \tau)$ denota un álgebra de von Neumann semifinita con traza τ fiel normal y semifinita (f.n.s.). En particular, si $\mathcal{M}$ es un factor finito, entonces τ denota la única traza f.n.s. tal que $\tau(I) = 1$. El espacio real de operadores autoadjuntos en $\mathcal{M}$ es denotado por $\mathcal{M}_{sa}$ y el cono de operadores positivos es denotado por $\mathcal{M}^+$. $\mathcal{P}(\mathcal{M}) \subseteq \mathcal{M}_{sa}$ denota el reticulado de proyecciones ortogonales en $\mathcal{M}$ dotado de la topología fuerte de operadores. Si $a \in \mathcal{M}_{sa}$ entonces $P^a(\Delta)$ denota la proyección espectral correspondiente al conjunto medible $\Delta \subseteq \mathbb{R}$; sin embargo, para simplificar la notación, escribimos $P^a(\alpha, \beta)$ en lugar de $P^a((\alpha, \beta))$. Si $a \in \mathcal{M}$ entonces $R(a)$ denota su rango y $P_{\overline{R(a)}} \in \mathcal{P}(\mathcal{M})$ denota la proyección ortogonal sobre la clausura de su rango. Denotamos por $\mathcal{U}_{\mathcal{M}}$ el grupo de operadores unitarios de $\mathcal{M}$. Finalmente, $\mathbb{R}_0^+$ denota el conjunto de números reales no negativos.

3.1. Refinamientos y modelado de operadores

3.1.1. Refinamientos de resoluciones espectral

Sea $I = [\alpha, \beta] \subseteq \mathbb{R}$ un intervalo cerrado, y recordemos que $\mathcal{P}(L(\mathcal{H}))$ denota el reticulado de proyecciones ortogonales en $L(\mathcal{H})$ dotado de la topología fuerte de operadores. Si $p \in \mathcal{P}(L(\mathcal{H}))$,

decimos que una función $E : I \rightarrow \mathcal{P}(L(\mathcal{H}))$ es una *resolución espectral a derecha y acotada de p* (y abreviamos REDA de p) si E es decreciente y continua a derecha, $E(\beta) = 0$ y $E(\alpha) = p$. Si $p = 1$ entonces esta noción concuerda con la definición usual de REDA en $L(\mathcal{H})$ (ver Definición 2.4.8). Por ejemplo, si $a \in M^+$ es un operador positivo entonces induce una REDA de $p = P_{\overline{R(a)}}$, dada por

$$E_\lambda = P^a(\lambda, \infty), \quad \lambda \in [0, \|a\|] \quad (3.1)$$

Si E es una REDA (de $E(\alpha)$) entonces, identificamos E con la familia $\{E_\lambda\}_{\lambda \in I}$, donde $E_\lambda = E(\lambda)$ para cada $\lambda \in I$. Si el conjunto I está claro del contexto, la REDA es denotada por $\{E_\lambda\}$. Si $\mathcal{N} \subseteq L(\mathcal{H})$ es un álgebra de von Neumann, decimos que $\{E_\lambda\}_{\lambda \in I}$ es una REDA de p en $\mathcal{N}$, si es una REDA de p tal que $E_\lambda \in \mathcal{P}(\mathcal{N})$ para $\lambda \in I$. Por ejemplo, si $a \in \mathcal{N}^+$ es un operador positivo, entonces la REDA inducida por a está en $\mathcal{N}$.

En adelante vamos a considerar la siguiente noción de *refinamiento* entre resoluciones espectrales.

Definición 3.1.1. Sean $\{E_\lambda\}_{\lambda \in I}$ y $\{E'_\lambda\}_{\lambda \in I'}$ REDAs con $I = [\alpha, \beta]$, $I' = [\alpha', \beta']$.

1. Dada una función $f : I \rightarrow I'$, decimos que el par $(\{E'_\lambda\}, f)$ es un *refinamiento* de $\{E_\lambda\}$ si

- a) f es creciente, continua a derecha y $f(\beta) = \beta'$;
- b) $E_\lambda = E'_{f(\lambda)}$ para cada $\lambda \in I$.

decimos que $(\{E'_\lambda\}, f)$ es un *refinamiento fuerte* de $\{E_\lambda\}$ si f también satisface

- (c) $f(\lambda) \geq \lambda$ para cada $\lambda \in I$, y
- (d) $f(\lambda) - f(\mu) \geq \lambda - \mu$, para cada $\lambda > \mu \in I$. □

Es evidente que el refinamiento entre REDAs definido arriba es un preorden. Sea $I = [\alpha, \beta]$ y sea $\{E_\lambda\}_{\lambda \in I}$ una REDA de $p \in \mathcal{P}(L(\mathcal{H}))$. Decimos que $\lambda_0 \in (\alpha, \beta]$ es un *átomo* para $\{E_\lambda\}$, si la resolución no es continua a la derecha en λ_0 . Si $p \neq 1$ entonces α es considerado un átomo. El conjunto de átomos de $\{E_\lambda\}_{\lambda \in I}$ es denotado por $\text{At}(\{E_\lambda\})$; es claro que este conjunto es vacío si y solo la resolución es continua. Decimos que un operador positivo $a \in \mathcal{M}^+$ tiene *distribución continua* si la resolución inducida por $a \in \mathcal{M}^+$ (ver fórmula (3.1)) es continua.

El siguiente resultado es una herramienta fundamental para el desarrollo de este trabajo. Su prueba será desarrollada en la Sección 4. En lo que sigue abreviamos subálgebras abelianas maximales por “masa”.

Teorema 3.1.2. Sea $(\mathcal{M}, \tau)$ un factor II_1 . Dada una REDA $\{E_\lambda\}_{\lambda \in I}$ de p en $\mathcal{M}$, existe una REDA continua $\{E'_\lambda\}_{\lambda \in I'}$ en $\mathcal{M}$ y una función $f : I \rightarrow I'$ tal que $(\{E'_\lambda\}, f)$ es un refinamiento fuerte de $\{E_\lambda\}$. Más aún, si $\mathcal{A} \subseteq \mathcal{M}$ es una masa y $\{E_\lambda\}$ es una REDA en $\mathcal{A}$ entonces podemos elegir a $\{E'_\lambda\}$ también en $\mathcal{A}$.

3.1.2. Modelado de operadores

Comenzamos con los siguientes lemas sobre funciones. Las pruebas son elementales y las omitimos.

Lema 3.1.3. Sean $I = [\alpha, \beta]$, $J = [\alpha', \beta'] \subseteq \mathbb{R}$ intervalos compactos, $g : J \rightarrow [0, 1]$ una función decreciente y continua a derecha, y $h : I \rightarrow [0, 1]$ una función decreciente y continua tal que $h(\alpha) \geq h(\alpha')$ y $h(\beta) \leq g(\beta')$. Si definimos $\tilde{g} : J \rightarrow I$ dada por

$$\tilde{g}(x) = \max\{t \in I : g(x) = h(t)\}$$

entonces, $\tilde{g}$ es una función creciente continua a derecha tal que $g = h \circ \tilde{g}$.

Lema 3.1.4. Sean $I = [\alpha, \beta]$, $J = [\alpha', \beta'] \subseteq \mathbb{R}$ intervalos compactos y sea $f : I \rightarrow I'$ una función creciente continua a derecha tal que $f(\beta') = \beta$. Si $f^\dagger : I \rightarrow J$ es la función dada por

$$f^\dagger(\lambda) = \min\{t \in I : \lambda \leq f(t)\}$$

entonces es creciente y continua a izquierda y tal que, para cada $t \in I$

$$\{\lambda \in I' : f^\dagger(\lambda) > t\} = \{\lambda \in I' : \lambda > f(t)\}. \quad (3.2)$$

Más aún, si $\tilde{J} \subseteq J$ y $g : \tilde{J} \rightarrow I'$ es tal que $f(t) \geq g(t)$ para cada $t \in \tilde{J}$ entonces $g^\dagger \geq f^\dagger$.

Lema 3.1.5. Sean $I = [\alpha, \beta]$, $J = [\alpha', \beta'] \subseteq \mathbb{R}$ y sea $f : I \rightarrow J$ una función creciente y continua a izquierda tal que $f(\alpha) = \alpha'$. Si $f_\dagger : J \rightarrow I$ es la función definida por

$$f_\dagger(\lambda) = \max\{t \in I : \lambda \geq f(t)\}$$

entonces es creciente y continua a derecha y tal que, para cada $t \in I$ vale

$$\{\lambda \in J : \lambda < f(t)\} = \{\lambda \in J : f_\dagger(\lambda) < t\}. \quad (3.3)$$

En lo que sigue consideramos los valores singulares de los operadores positivos de un factor II_1 , así como las relaciones de preorden espectral y (sub)mayorización entre éstos. Estas nociones se describen en la sección 2.4.4 de los Preliminares (ver las definiciones 2.4.15 y 2.4.17).

Teorema 3.1.6. Sea $(\mathcal{M}, \tau)$ un factor II_1 y sea $a \in \mathcal{M}^+$. Entonces, existe un operador $a' \in \mathcal{M}^+$ con distribución continua tal que

1. La resolución espectral de a' refina la resolución espectral de a .
2. Si $b \in \mathcal{M}^+$ entonces, existe una función h_b creciente y continua a izquierda tal que, si $\tilde{b} = h_b(a')$ entonces $\mu_b = \mu_{\tilde{b}}$. Más aún, la resolución espectral de a' refina la resolución espectral de b si y solo si $h_b(a') = b$.
3. Si $c^+ \in \mathcal{M}$ entonces $c \precsim b$ (resp $c \prec_w b$, $c \prec b$) si y solo si $\tilde{c} \leq \tilde{b}$ (resp $\tilde{c} \prec_w \tilde{b}$, $\tilde{c} \prec \tilde{b}$).

Más aún, si $\mathcal{A} \subseteq \mathcal{M}$ es una masa y $a \in \mathcal{A}^+$ entonces podemos elegir a $a' \in \mathcal{A}^+$.

Demostración. Sea $a \in \mathcal{M}^+$ y consideremos la resolución espectral de $p = P_{\overline{R(a)}}$ inducida por a , $\{E_\lambda\}_{\lambda \in I}$ donde $I = [0, \|a\|]$, definida como en la ecuación (3.1). Entonces, por el Teorema 3.1.2 existe un refinamiento fuerte y continuo $(\{E'_\lambda\}_{\lambda \in I'}, f)$, donde $I' = [0, \alpha] \subseteq \mathbb{R}^+$ es un intervalo cerrado. Sea

$$a' = \int_{I'} \lambda dE'_\lambda$$

y notemos que $P^{a'}(0, \infty) = E'_0 = 1$, pues $\{E'_\lambda\}$ es continua. Luego $a' \in \mathcal{M}^+$ es un operador con distribución continua y tal que $P^a(s, \infty) = P^{a'}(f(s), \infty)$ para cada $s \in I$. Sea $h : [0, \|a'\|](= I') \rightarrow [0, 1]$ la función definida por $h(s) = \tau(P^{a'}(s, \infty))$ y notemos que h es continua, decreciente y tal que $h(0) = 1$ y $h(\|a'\|) = 0$.

Sea $b \in \mathcal{M}^+$ y sea $g : [0, \|b\|](= J) \rightarrow [0, 1]$ la función decreciente y continua a derecha definida por $g(s) = \tau(P^b(s, \infty))$. Por el Lema 3.1.3, existe una función creciente y continua a derecha $\tilde{g} : J \rightarrow I'$, tal que $g = h \circ \tilde{g}$ es decir,

$$\tau(P^b(s, \infty)) = \tau(P^{a'}(\tilde{g}(s), \infty)), \quad s \in J. \quad (3.4)$$

Por el Lema 3.1.4 existe $\tilde{g}^\dagger := h_b : I' \rightarrow J$ creciente (y luego uniformemente acotada y medible) continua a derecha tal que

$$\{\lambda \in I' : h_b(\lambda) > s\} = \{\lambda \in I' : \lambda > \tilde{g}(s)\}, \quad s \in J. \quad (3.5)$$

Sea $\tilde{b} = \int_{I'} h_b(\lambda) dE'_\lambda$ y notemos que $P^{\tilde{b}}(s, \infty) = P^{a'}(\tilde{g}(s), \infty)$, que se deduce de la ecuación (3.5). Entonces, por la ecuación (3.4) tenemos que

$$\tau(P^{\tilde{b}}(s, \infty)) = \tau(P^{a'}(\tilde{g}(s), \infty)) = \tau(P^b(s, \infty)), \quad s \in J. \quad (3.6)$$

Luego, b y $\tilde{b}$ tienen los mismos valores singulares.

Supongamos que la resolución espectral de $b \in \mathcal{M}^+$ está refinada por la resolución espectral de a' . Sea $\tilde{b} = h_b(a')$ y notemos que $P^{\tilde{b}}(s, \infty) = P^{a'}(\tilde{g}(s), \infty)$ y $P^b(s, \infty) = P^{a'}(f'(s), \infty)$ para alguna función creciente y continua a derecha $f' : [0, \|b\|] \rightarrow I'$. Entonces $P^{\tilde{b}}(s, \infty) \leq P^b(s, \infty)$ ó $P^b(s, \infty) \leq P^{\tilde{b}}(s, \infty)$ y

$$\tau(P^b(s, \infty)) = \tau(P^{\tilde{b}}(s, \infty)) \Rightarrow P^b(s, \infty) = P^{\tilde{b}}(s, \infty), \quad s \in J.$$

Por otra parte, si $b = h(a')$ para alguna función creciente y continua a izquierda $h : I' \rightarrow [0, \|b\|]$ entonces es sencillo verificar que a' refina a b .

Finalmente, supongamos que $c \in \mathcal{M}^+$ es tal que $\mu_c \leq \mu_b$ o, equivalentemente, que $\tau(P^c(s, \infty)) \leq \tau(P^b(s, \infty))$ para todo $s \geq 0$. Sea $k(s) = \tau(P^c(s, \infty))$, $\tilde{k}$ obtenida de k como en el Lema 3.1.3, y $\tilde{k}^\dagger = h_c$ obtenida de $\tilde{k}$ como en el Lema 3.1.4. Entonces, $h_c \leq h_b$; de hecho, $\tau(P^c(s, \infty)) \leq \tau(P^b(s, \infty))$ implica que $\tilde{g} \leq \tilde{k}$ y, por el Lema 3.1.4, concluimos que $\tilde{k}^\dagger \leq \tilde{g}^\dagger$. El resto del teorema es consecuencia de las igualdades $\mu_c = \mu_{\tilde{c}}$ y $\mu_b = \mu_{\tilde{b}}$. □

Observación 3.1.7. Con las notaciones de la prueba del teorema anterior, decimos que $\tilde{b} \in \mathcal{M}^+$ es un *modelo* de $b \in \mathcal{M}^+$. Como una consecuencia inmediata del ítem 2. en la Proposición 2.4.14, concluimos que el modelo $\tilde{b}$ es aproximadamente unitariamente equivalente a b en $\mathcal{M}$.

La siguiente aplicación del Teorema 3.1.6 provee nuevas caracterizaciones del preorden espectral y la submayorización entre operadores positivos en factores Π_1 . Estas caracterizaciones dan una respuesta parcial afirmativa a un problema propuesto en [33] (ver la sección 3.1.3 para una discusión detallada de este problema). Notemos que estas reformulaciones involucran desigualdades con respecto al orden usual de operadores.

Teorema 3.1.8. Sea $(\mathcal{M}, \tau)$ un factor II_1 y sean $a, b \in \mathcal{M}^+$. Entonces

1. b domina espectralmente a a si y solo si existe

$$c \in \overline{\mathcal{U}_{\mathcal{M}}(b)} \quad \text{tal que} \quad a \leq c \quad (3.7)$$

o, equivalentemente, si existe

$$\overline{\mathcal{U}_{\mathcal{M}}(a)} \ni d \quad \text{tal que} \quad d \leq b. \quad (3.8)$$

Más aún, podemos asumir que a y c conmutan y, b y d conmutan.

2. b submayoriza a a si y solo si existe $c \in \mathcal{M}^+$ tal que

$$a \leq c \prec b. \quad (3.9)$$

Más aún, podemos asumir que a y c conmutan.

Observación 3.1.9. Recordemos que para operadores positivos $a, b \in \mathcal{M}^+$, $a \leq b$ implica que $a \preceq b$. Entonces, en el caso general de un factor semifinito, la existencia de operadores c, d que verifiquen las fórmulas (3.7) ó (3.8), implican la dominación espectral. Análogamente, la existencia de un operador c tal que la fórmula (3.9) se verifique, implica la submayorización. Así, el objetivo del Teorema 3.1.8 es mostrar que las implicaciones recíprocas son ciertas en el caso de un factor finito.

Demostración. Probamos la primera parte del ítem 1. Sea $a \in \mathcal{M}^+$ y notemos que, por el Teorema 3.1.6, existe $a' \in \mathcal{M}^+$ con las propiedades de los ítems 1., 2., y 3. en ese teorema. Sea $b \in \mathcal{M}^+$ tal que $a \preceq b$ es decir, $\mu_a \leq \mu_b$. Entonces, nuevamente por el Teorema 3.1.6, existe una función creciente y continua a izquierda h_b tal que $\tilde{b} = h_b(a')$ tiene los mismos valores singulares que b i.e., $\mu_b = \mu_{\tilde{b}}$. Por el ítem 2. en la Proposición 2.4.14, esto implica que $\tilde{b} \in \overline{\mathcal{U}_{\mathcal{M}}(b)}$. Como por hipótesis $\mu_a \leq \mu_b$ entonces, por los ítems 2. y 3. en el Teorema 3.1.6, tenemos $\tilde{b} = h_b(a') \geq h_a(a') = a$. Así, obtenemos la fórmula (3.7) con $c = \tilde{b}$. La prueba de la segunda parte sigue un argumento similar, considerando el modelo de a con respecto a un refinamiento de b .

Para probar el segundo ítem, sean a y a' como en la primera parte de la prueba. Sea $b \in \mathcal{M}^+$ tal que $a \prec b$ y sea $\nu = \nu_{a'}$ la medida inducida por la traza y la medida espectral de a' en $I' = [0, \|a'\|]$ es decir, $\nu(\Delta) = \tau(\chi_{\Delta}(a'))$. Entonces, si h_a, h_b son como en el Teorema 3.1.6 se tiene $h_a \prec h_b$ en el espacio de medida (I', ν) . Luego, por el ítem 2. en la Proposición 2.4.21 existe $h \in L^\infty(I', \nu)$ tal que $h_a \leq h \prec h_b$. Sea $c = h(a')$ y notemos que, por construcción $c \prec b$. El resultado es una consecuencia de este último hecho y de que $a = h_a(a')$.

□

El siguiente corolario es una consecuencia inmediata del Teorema 3.1.8 y del hecho de que, para operadores positivos $a, b \in \mathcal{M}^+$, $a \preceq b$ implica $a \prec_w b$.

Corolario 3.1.10. Con las notaciones del teorema, si b domina espectralmente a a entonces existe $c \in \mathcal{M}^+$ tal que $a \leq c \prec b$. Más aún, podemos suponer que a y c conmutan.

En el siguiente corolario agregamos un ítem a la caracterización de la dominación espectral dada en el Teorema 3.1.8.

Corolario 3.1.11. Sean $a, b \in \mathcal{M}^+$ entonces, los siguientes son equivalentes:

1. b domina espectralmente a a
2. Existe $c \in \overline{\mathcal{U}_{\mathcal{M}}(b)}$ tal que $a \leq c$.
3. Existe $d \in \overline{\mathcal{U}_{\mathcal{M}}(a)}$ tal que $d \leq b$.
4. Existe una REDA $\{E_\lambda\}_{\lambda \in I}$, donde $I = [0, \|a\|]$ tal que $\tau(E_\lambda) = \tau(P^a(\lambda, \infty))$ para cada $\lambda \in I$

y

$$\lambda E_\lambda \leq E_\lambda b E_\lambda, \quad \forall \lambda \geq 0. \quad (3.10)$$

Demostración. La equivalencia de los ítems 1, 2 y 3 se probó en el Teorema 3.1.8. Supongamos que existe una sucesión $(v_n^* a v_n)_n \subseteq \mathcal{U}_{\mathcal{M}}(a)$ tal que $\lim_{n \rightarrow \infty} \|d - v_n^* a v_n\| = 0$ y $d \leq b$ para algún $d \in \mathcal{M}^+$. Notemos que, para cada $\lambda \geq 0$ tenemos $\tau(P^a(\lambda, \infty)) = \tau(P^d(\lambda, \infty))$, pues $\lim_{n \rightarrow \infty} P^{v_n^* a v_n}(\lambda, \infty) = P^d(\lambda, \infty)$ σ -WOT. Más aún ,

$$\lambda P^d(\lambda, \infty) \leq P^d(\lambda, \infty) d \leq P^d(\lambda, \infty) b P^d(\lambda, \infty).$$

Entonces, si definimos $E_\lambda = P^d(\lambda, \infty)$, $\{E_\lambda\}_{\lambda \in [0, \|a\|]}$ es una REDA con las propiedades requeridas.

Finalmente, asumamos que existe una REDA $\{E_\lambda\}_{\lambda \in [0, \|a\|]}$ como en el ítem 4. Dado $\epsilon > 0$, sea $b_\epsilon = b + \epsilon I$ y notemos que

$$\lambda E_\lambda < E_\lambda b_\epsilon E_\lambda$$

y luego, $P^{E_\lambda b_\epsilon E_\lambda}(\lambda, \infty) = E_\lambda$. En [32] Fack probó la siguiente desigualdad de tipo entrelace: para cada proyección ortogonal $p \in \mathcal{M}$, $p b p \preceq b$. Entonces

$$\tau(P^a(\lambda, \infty)) = \tau(E_\lambda) = \tau(P^{E_\lambda b_\epsilon E_\lambda}(\lambda, \infty)) \leq \tau(P^{b_\epsilon}(\lambda, \infty)).$$

La desigualdad anterior muestra que $\mu_a \leq \mu_{b_\epsilon}$ para cada $\epsilon > 0$. Entonces el resultado es una consecuencia del hecho de que $\lim_{\epsilon \rightarrow 0^+} \mu_{b_\epsilon}(t) = \mu_b(t)$ para cada $t \geq 0$.

□

Observación 3.1.12. En la última parte de la prueba del Corolario 3.1.11 no hemos usado el hecho de que estamos considerando un factor $(\mathcal{M}, \tau)$ sino mas bien un álgebra de von Neumann semifinita . Luego, la existencia de una resolución espectral $\{E_\lambda\}_{\lambda \in I}$ que satisface las hipótesis del ítem 4. implica la dominación espectral en un álgebra de von Neumann semifinita arbitraria. Este último hecho ha sido probado en [33].

Seguidamente, aplicamos los resultados obtenidos a las siguientes desigualdades de tipo Young obtenidas por Farenick y Manjegani en [33].

Corolario 3.1.13. Sea $(\mathcal{M}, \tau)$ un factor II_1 , sean $x, y \in \mathcal{M}$ y $p, q > 1$ índices conjugados. Entonces, existen sucesiones $(u_n)_{n \in \mathbb{N}}, (v_n)_{n \in \mathbb{N}} \subseteq \mathcal{M}$ de operadores unitarios tales que

$$|xy^*| \leq \lim_{n \rightarrow \infty} u_n^* (p^{-1}|x|^p + q^{-1}|y|^q) u_n$$

y

$$\lim_{n \rightarrow \infty} v_n^* |xy^*| v_n \leq p^{-1}|x|^p + q^{-1}|y|^q$$

Demostración. En [33] probaron que si p, q, x, y son como en el enunciado, entonces $|xy^*| \lesssim p^{-1}|x|^p + q^{-1}|y|^q$. El resultado es una consecuencia de este hecho y del Teorema 3.1.8. $\square$

Esta reformulación de las desigualdades de tipo Young, que son en realidad equivalentes al resultado obtenido en [33] son formalmente similares a las obtenidas por T. Ando en el contexto de las álgebras de matrices.

3.1.3. Algunas observaciones sobre el caso semifinito no finito

En el caso de un factor semifinito arbitrario $\mathcal{M}$ actuando en un espacio de Hilbert separable $\mathcal{H}$, Farenick y Manjegani [33] preguntaron si (una instancia particular de) la desigualdad $a \lesssim b$ implica la existencia de un automorfismo Θ de $\mathcal{M}$, o al menos una isometría $v \in \mathcal{M}$ tal que $\Theta(b) \geq a$ o $v^*bv \geq a$. El teorema 3.1.8 da una respuesta parcial afirmativa a este problema en el caso particular de un factor de tipo II_1 , mostrando la existencia no de un unitario sino de una sucesión de unitarios que cumplen aproximadamente esta condición.

En esta sección mostramos que si consideramos $\mathcal{M} = L(\mathcal{H})$ lo anterior no es cierto, para operadores arbitrarios $a, b \in L(\mathcal{H})^+$ tales que $a \lesssim b$.

Recordemos que todo automorfismo de $L(\mathcal{H})$ es interior i.e, para cada automorfismo Θ existe un operador unitario $u \in L(\mathcal{H})$ tal que $\Theta(x) = u^*xu$ para todo $x \in L(\mathcal{H})$. En el siguiente ejemplo exhibimos dos operadores $a, b \in L(\mathcal{H})^+$ tales que $a \lesssim b$ y tal que no existe una isometría $v \in L(\mathcal{H})$ tal que $v^*bv = a$. En particular no existe ningún automorfismo Θ de $L(\mathcal{H})$ tal que $\Theta(b) = a$.

Ejemplo 3.1.14. Sea $\{e_n\}_{n \in \mathbb{N}}$ una base ortonormal del espacio de Hilbert (separable) $\mathcal{H}$ y sean $a, b \in L(\mathcal{H})^+$ dados por

$$a(e_n) = \begin{cases} 0 & \text{si } n = 1 \\ \frac{1}{2^{n-1}} e_n & \text{si } n \geq 2 \end{cases} \quad \text{y} \quad b(e_n) = \frac{1}{2^n} e_n, \quad \forall n \geq 1. \quad (3.11)$$

Notemos que entonces $a \lesssim b$, pues $\text{tr}(P^a(\lambda, \infty)) = \text{tr}(P^b(\lambda, \infty))$ para cada $\lambda \geq 0$. Si suponemos que existe una isometría $v \in L(\mathcal{H})$ tal que $v^*bv \geq a$, entonces

$$a(e_2) = \frac{1}{2} e_2 \Rightarrow \langle bv(e_2), v(e_2) \rangle \geq \frac{1}{2},$$

y entonces $v(e_2) = e_1$. Análogamente, como v es una isometría entonces $v(e_3) \perp v(e_2)$ y como $\langle bv(e_3), v(e_3) \rangle \geq \frac{1}{4}$ concluimos que $v(e_3) = e_2$. Finalmente, siguiendo un argumento inductivo, vemos que $v(e_n) = e_{n-1}$ para cada $n \geq 2$. Luego, $v(e_1) \perp e_i$ para cada $i \geq 1$, con lo cual $v(e_1) = 0$, lo que contradice el hecho de que v es una isometría de $\mathcal{H}$.

Sin embargo siguiente teorema da una respuesta parcial afirmativa al problema anterior.

Teorema 3.1.15. Sea $\mathcal{H}$ un espacio de Hilbert. Dados a, b operadores positivos de $L(\mathcal{H})$ tales que a es compacto, b es diagonalizable y $a \lesssim b$ entonces, existe una isometría parcial $u \in L(\mathcal{H})$ con espacio inicial $\overline{R(a)}$ tal que

$$uau^* \leq b \quad \text{y} \quad (uau^*)b = b(uau^*).$$

Más aún, si $\mathcal{H}$ tiene dimensión finita, la misma desigualdad vale para $a, b \in L_{sa}(\mathcal{H})$ y la isometría puede cambiarse por un unitario.

Demostración. Sea $\{\xi_n\}_{n \in \mathbb{N}}$ una base ortonormal de $\overline{R(a)}$ que consiste de autovectores asociados a una sucesión decreciente de autovalores positivos $\{\lambda_n\}_{n \in \mathbb{N}}$ de a , contados con multiplicidad.

Por otro lado, sea $\{\eta_m\}_{m \in M}$ una base ortonormal de $\overline{R(b)}$ que consiste de autovectores asociados a una sucesión de autovalores positivos $\{\mu_m\}_{m \in M}$ de b , contados con multiplicidad.

Dado $\alpha \geq 0$ sean $\lambda_1, \dots, \lambda_{n_\alpha}$ los autovalores de a estrictamente mayores que α . De forma similar, sea $M(\alpha) = \{m \in M, \mu_m > \alpha\}$. Si m_α denota el cardinal de $M(\alpha)$ entonces notemos que $0 \leq m_\alpha \leq \infty$ y si $\alpha \geq \beta \geq 0$ entonces $M(\alpha) \subseteq M(\beta)$.

Por hipótesis, $P^a(\alpha, +\infty)$ es equivalente a una subproyección de $P^b(\alpha, +\infty)$. Como $P^a(\alpha, +\infty)$ es la proyección (ortogonal) sobre la clausura lineal de $\{\xi_n\}_{n=1}^{n_\alpha}$ y $P^b(\alpha, +\infty)$ es la proyección sobre la clausura lineal de $\{\eta_m\}_{m \in M(\alpha)}$, entonces vemos que $n_\alpha \leq m_\alpha$.

Entonces, podemos definir una inyección $\psi : \mathbb{N} \rightarrow M$ tal que para todo $\alpha \geq 0$ se verifica que $\psi(k) \in M(\alpha)$, $k = 1, \dots, n_\alpha$. Pero entonces, definiendo $u : \overline{R(a)} \rightarrow \overline{R(b)}$ por

$$u(\xi_n) = \eta_{\psi(n)}$$

y extendiendo esta definición a $\mathcal{H}$ como cero en $\ker(a)$, obtenemos que $uau^* \leq b$ y $(uau^*)b = b(uau^*)$, ya que el conjunto $(\eta_m)_{m \in M}$ es un sistema de autovectores para uau^* y b , que es completo para $\overline{R(uau^*)} = R(u)$, y $\ker(uau^*) = R(u)^\perp$ es un subespacio invariante para b .

Finalmente, si $\dim \mathcal{H} = n < \infty$ y $a, b \in L_{sa}(\mathcal{H})$, sea $\xi_1, \dots, \xi_n$ una base de $\mathcal{H}$ que consiste de autovectores de a asociados a una sucesión (finita) decreciente de autovalores $\lambda_1, \dots, \lambda_n$ y sea $\eta_1, \dots, \eta_n$ una base de $\mathcal{H}$ que consiste de autovectores de b asociados a una sucesión (finita) decreciente de autovalores $\mu_1, \dots, \mu_n$. Entonces, si definimos $u : \mathcal{H} \rightarrow \mathcal{H}$ dado por $u(\xi_m) = \eta_m$, $1 \leq m \leq n$, el mismo argumento que antes muestra que $uau^* \leq b$ y $(uau^*)b = b(uau^*)$. Por construcción, u es un operador unitario. □

Observamos que, en general, si $c, d \in L(\mathcal{H})$ son tales que $c \preceq d$ entonces puede no existir una sucesión de operadores unitarios $(u_n)_{n \in \mathbb{N}} \subseteq L(\mathcal{H})$ tales que $\lim_{n \rightarrow \infty} u_n^* d u_n \rightarrow e \geq c$. De hecho, consideremos los operadores $c, d \in L(\mathcal{H} \oplus \mathcal{H})$ dados por

$$c = a + I \oplus I \quad \text{y} \quad d = a + I \oplus 0,$$

donde $a \in L(\mathcal{H})$ es el operador definido por la fórmula (3.11). Entonces, es inmediato verificar que $\tau(P^c(\lambda, \infty)) = \tau(P^d(\lambda, \infty))$ para cada $\lambda \geq 0$. Pero, por otra parte, si existe $e \geq c$ tal que e y d son aproximadamente unitariamente equivalentes entonces (ver II.4.4 en [28]) $\sigma_e(e) = \sigma_e(d)$, donde $\sigma_e(d)$ denota el espectro esencial de d . Pero $0 \in \sigma_e(d) = \sigma_e(e)$ y $e \geq c$, lo que contradice el hecho de que $c \geq I \oplus I$ es un operador inversible.

3.2. Refinamientos de resoluciones espectrales en factores II_1

En esta sección, el par $(\mathcal{M}, \tau)$ denota a un factor II_1 y las REDAs están $\mathcal{M}$. Dada una REDA $\{E_\lambda\}_{\lambda \in [\alpha, \beta]}$ de $p \in \mathcal{P}(\mathcal{M})$, con $0 \leq \alpha < \beta$, decimos que está *uniformemente distribuida* si

$$\tau(E_\lambda) = \frac{\tau(p)(\beta - \lambda)}{\beta - \alpha}, \quad \lambda \in [\alpha, \beta]. \quad (3.12)$$

Lema 3.2.1. Sea $p \in \mathcal{P}(\mathcal{M})$ una proyección en $\mathcal{M}$ y sea $0 \leq \alpha < \beta \in \mathbb{R}$. Entonces, existe una REDA de p en $\mathcal{M}$ $\{E_\lambda\}_{\lambda \in [\alpha, \beta]}$ uniformemente distribuida. Más aún, si $\mathcal{A} \subseteq \mathcal{M}$ es una masa y $p \in \mathcal{A}$ entonces podemos elegir $\{E_\lambda\}$ en $\mathcal{A}$.

Demostración. Sea $p \in \mathcal{P}(\mathcal{M})$. Notemos que si construimos una REDA $\{E_\lambda\}_{\lambda \in [0, s]}$ para algún $s > 0$ y tal que

$$\tau(E_\lambda) = \frac{\tau(p)(s - \lambda)}{s}, \quad \lambda \in [0, s] \quad (3.13)$$

luego podemos reparametrizar esta resolución y obtener una REDA con las propiedades deseadas.

Fijemos $s = \tau(p)$ y consideremos una descomposición $p = p_1^1 + p_2^1$ con $\tau(p_i^1) = s/2$ para $i = 1, 2$. Definimos inductivamente una sucesión $\{p_i^k\}_{i=1}^{2^k}$, $k \in \mathbb{N}$ de la siguiente forma :

$$p_{2i-1}^{k+1} + p_{2i}^{k+1} = p_i^k \quad \text{y} \quad \tau(p_{2i-1}^{k+1}) = \tau(p_{2i}^{k+1}) = s/2^{k+1}, \quad 1 \leq i \leq k+1.$$

Para cada $k \in \mathbb{N}$ y $1 \leq i \leq 2^k$ sea $\beta_i^k = si/2^k$ y consideremos la medida espectral F_k en $[0, s]$ dada por

$$F_k(\Delta) = \sum_{i=1}^{2^k} \chi_\Delta(\beta_i^k) p_i^k + \chi_\Delta(0)(1 - p)$$

para conjuntos de Borel $\Delta \subset [0, s]$. Notemos que $F_k(\Delta) \geq F_{k+1}(\Delta)$ de forma que existe

$$F(\Delta) = \lim_{k \rightarrow \infty} F_k(\Delta) = \bigwedge_{k \in \mathbb{N}} F_k(\Delta), \quad (3.14)$$

donde la convergencia es en la topología fuerte de operadores. Afirmamos que $F : B([0, s]) \rightarrow \mathcal{P}(\mathcal{M})$ es una medida espectral con soporte $[0, s]$. Basta mostrar la σ -aditividad de F pues el resto de las propiedades de F son evidentes; de hecho, si $\{\Delta_i\}_{i \in \mathbb{N}} \subseteq [0, s]$ es una sucesión de conjuntos de Borel dos a dos disjuntos entonces, como $F(\Delta_i) \leq F_k(\Delta_i)$ para cada k , $i \in \mathbb{N}$, $\sum_{i=1}^\infty F(\Delta_i) \leq \sum_{i=1}^\infty F_k(\Delta_i)$ para cada $k \in \mathbb{N}$ y $F(\Delta_i)F(\Delta_j) = 0$ para $i \neq j$. Entonces, tenemos

$$\sum_{i=1}^\infty F(\Delta_i) \leq \lim_{k \rightarrow \infty} \sum_{i=1}^\infty F_k(\Delta_i) = \lim_{k \rightarrow \infty} F_k\left(\bigcup_{i=1}^\infty \Delta_i\right) = F\left(\bigcup_{i=1}^\infty \Delta_i\right).$$

Por otra parte,

$$\begin{aligned} \tau\left(\sum_{i=1}^\infty F(\Delta_i)\right) &= \sum_{i=1}^\infty \tau(F(\Delta_i)) = \sum_{i=1}^\infty \lim_{k \rightarrow \infty} \tau(F_k(\Delta_i)) \\ &= \lim_{k \rightarrow \infty} \sum_{i=1}^\infty \tau(F_k(\Delta_i)) = \tau\left(F\left(\bigcup_{i=1}^\infty \Delta_i\right)\right), \end{aligned}$$

donde el cambio del orden de los símbolos de límite y sumación es válido por el teorema de la convergencia dominada de Lebesgue en $l^1(\mathbb{N})$. Entonces $F(\bigcup_{i=1}^\infty \Delta_i) = \sum_{i=1}^\infty F(\Delta_i)$; más aún, de la fórmula (3.14) vemos que

$$\tau(F((\lambda, s])) = \lim_{k \rightarrow \infty} F_k((\lambda, s]) = \frac{\tau(p)(s - \lambda)}{s}, \quad \lambda \in [0, s].$$

Si $\{E_\lambda\}_{\lambda \in [0, s]}$ está dada por $E_\lambda = F((\lambda, s])$ entonces es una REDA que satisface la ecuación (3.13).

Supongamos que $\mathcal{A} \subseteq \mathcal{M}$ es una masa y $p \in \mathcal{A}$. Entonces $\mathcal{A}$ es un álgebra abeliana de von Neumann sin átomos. Así, para cada $0 \leq \alpha < \tau(p)$ hay una proyección $p \geq q \in \mathcal{A}$ tal que $\tau(q) = \alpha$, pues de otra forma p tendría una sub-proyección que es minimal en $\mathcal{A}$, contradiciendo nuestra hipótesis sobre $\mathcal{A}$. Luego, en este caso es posible construir la REDA uniformemente distribuida en $\mathcal{A}$. □

Aunque técnico, el Lema 3.2.1 tiene algunas consecuencias interesantes. El siguiente corolario generaliza la proposición 5.2 en [11]. Nuestra prueba, basada en el Lema 3.2.1, es más sencilla.

Corolario 3.2.2. Sea $\mathcal{A} \subseteq \mathcal{M}$ una subálgebra de von Neumann abeliana sin proyecciones minimales. Si μ es una medida de probabilidad de soporte compacto en la recta real, entonces existe $b \in \mathcal{A}$ tal que $\tau(P^b(s, \infty)) = \mu((s, \infty))$ para cada $s \in \mathbb{R}$.

Demostración. Como $1 \in \mathcal{A}$ entonces por el Lema 3.2.1 existe una REDA en $\mathcal{A}$ uniformemente distribuida de $1 \in \mathcal{A}$, $\{E_\lambda\}_{\lambda \in [0, 1]}$, tal que

$$\tau(E_\lambda) = 1 - \lambda, \quad \text{para cada } \lambda \in [0, 1]. \quad (3.15)$$

Sea $a = \int_{[0, 1]} \lambda dE_\lambda \in \mathcal{A}^+$. Si $F(s) = \mu((s, \infty))$ entonces, $F : \mathbb{R} \rightarrow [0, 1]$ es una función continua a derecha y decreciente. Para $x \in [0, 1]$, sea $F_+(x) = \min\{s : F(s) \leq x\}$ que está bien definida. Entonces $F_+ : [0, 1] \rightarrow \mathbb{R}$ es una función decreciente (y luego uniformemente acotada y medible) tal que, para cada $s \in \text{sop}(\mu)$ tenemos

$$\{x : s < F_+(x)\} = \{x : F(s) > x\}. \quad (3.16)$$

Sea $b = F_+(a) \in \mathcal{A}$ y notemos que si $s \in \text{sop}(\mu)$ entonces

$$\begin{aligned} P^b(s, \infty) &= P^a(\{x : s < F_+(x)\}) = \\ &= P^a(\{x : F(s) > x\}) = P^a(\{x : F(s) \geq x\}), \end{aligned} \quad (3.17)$$

que se deduce de las ecuaciones (3.15) y (3.16). El hecho de que $\tau(P^b(s, \infty)) = \mu((s, \infty))$ para cada $s \in \text{sop}(\mu)$ (y entonces para cada $s \in \mathbb{R}$) ahora se deduce de las ecuaciones (3.15) y (3.17) y del hecho de que $P^a((\lambda, \infty)) = E_\lambda$ para cada $\lambda \in [0, 1]$. □

Sea $I = [\alpha, \beta]$ y sea $\{E_\lambda\}_{\lambda \in I}$ una REDA de una proyección $p \in \mathcal{P}(\mathcal{M})$. Recordemos que $\lambda_0 \in (\alpha, \beta]$ es un átomo de $\{E_\lambda\}$, si la resolución no es continua a la izquierda en λ_0 i.e, si

$$\lim_{\lambda \rightarrow \lambda_0^-} E_\lambda = E_{\lambda_0} + p(\lambda_0), \quad p(\lambda_0) \neq 0. \quad (3.18)$$

En este caso $p(\lambda_0)$ es la *proyección salto* de $\{E_\lambda\}$ en λ_0 . Si $p \neq 1$ entonces α es considerado un átomo y definimos la proyección salto en este caso como $p(\alpha) = 1 - E_\alpha$. El conjunto de átomos de $\{E_\lambda\}_{\lambda \in I}$ es denotado por $\text{At}(\{E_\lambda\})$; este conjunto es vacío si y solo si la resolución es continua. Notemos que el conjunto de átomos $\text{At}(\{E_\lambda\})$ es numerable. De hecho, si $\lambda_0, \lambda_1 \in \text{At}(\{E_\lambda\})$ entonces es sencillo ver que $p(\lambda_0)p(\lambda_1) = 0$ i.e, $p(\lambda_0)$ y $p(\lambda_1)$ son proyecciones ortogonales. Luego

$$\mathcal{J}(\{E_\lambda\}) := \sum_{\lambda \in \text{At}(\{E_\lambda\})} \tau(p(\lambda)) = \tau\left(\sum_{\lambda \in \text{At}(\{E_\lambda\})} p(\lambda)\right) \leq 1 \quad (3.19)$$

y esto implica que $\text{At}(\{E_\lambda\})$ es numerable. Al número real $\mathcal{J}(\{E_\lambda\})$ lo llamamos el *salto total* de la resolución.

Lema 3.2.3. Sean $\{E_\lambda\}_{\lambda \in I}$, $\{E'_\lambda\}_{\lambda \in I'} \subseteq \mathcal{M}$ REDAs, donde $(\mathcal{M}, \tau)$ es un factor Π_1 . Si existe una función continua a derecha f tal que $(\{E'_\lambda\}, f)$ refina a $\{E_\lambda\}$ entonces $\mathcal{J}(\{E_\lambda\}) \geq \mathcal{J}(\{E'_\lambda\})$.

Demostración. Sea $\lambda_0 \in \text{At}(\{E'_\lambda\})$ y consideremos $\mu_0 = \min\{\mu \in I : f(\mu) \geq \lambda_0\}$ que está bien definido pues f es creciente y continua a derecha. Entonces, por definición de μ_0 $f(\mu_0) \geq \lambda_0$ y $f(\mu) < \lambda_0$ si $\mu < \mu_0$. Así tenemos

$$\lim_{\mu \rightarrow \mu_0^-} E_\mu - E_{\mu_0} = \lim_{\mu \rightarrow \mu_0^-} E'_{f(\mu)} - E'_{f(\mu_0)} \geq \lim_{\lambda \rightarrow \lambda_0^-} E'_\lambda - E'_{\lambda_0} \neq 0,$$

ya que λ_0 es un átomo de $\{E'_\lambda\}$. Luego $\mu_0 \in I$ es un átomo de la resolución $\{E_\lambda\}$ y tenemos

$$\lim_{\mu \rightarrow \mu_0^-} \tau(E_\mu) = \lim_{\mu \rightarrow \mu_0^-} \tau(E'_{f(\mu)}) > \tau(E'_{\lambda_0}) \geq \tau(E'_{f(\mu_0)}) = \tau(E_{\mu_0}) \quad (3.20)$$

pues $f(\mu) \rightarrow \lambda_1 \leq \lambda_0$ cuando $\mu \rightarrow \mu_0^-$ y $\mu_0 \in \text{At}(\{E_\lambda\})$. Consideramos la siguiente relación en $\text{At}(\{E'_\lambda\})$: si $\lambda_1, \lambda_2 \in \text{At}(\{E'_\lambda\})$ entonces $\lambda_1 \approx \lambda_2$ si y solo si existe $\mu_0 \in \text{At}(\{E_\lambda\})$ tal que

$$\tau(E'_{\lambda_1}), \tau(E'_{\lambda_2}) \in \left[\tau(E_{\mu_0}), \lim_{\mu \rightarrow \mu_0^-} \tau(E_\mu) \right).$$

La primera parte de la prueba muestra que esta relación es reflexiva. Por otra parte, es claramente simétrica. Si $\mu_1 < \mu_2$ entonces $\lim_{\mu \rightarrow \mu_2^-} \tau(E_\mu) \leq \tau(E_{\mu_1})$ y

$$[\tau(E_{\mu_2}), \lim_{\mu \rightarrow \mu_2^-} \tau(E_\mu)) \cap [\tau(E_{\mu_1}), \lim_{\mu \rightarrow \mu_1^-} \tau(E_\mu)) = \emptyset.$$

Así, si $\lambda_1 \approx \lambda_2$ entonces existe un único $\mu_0 \in \text{At}(\{E_\lambda\})$ tal que vale la ecuación anterior, y en particular $\approx$ es una relación de equivalencia. Entonces, para cada clase de equivalencia Q , existe un único átomo $\mu_Q \in \text{At}(\{E_\lambda\})$ tal que $\lambda \in (\lim_{\mu \rightarrow \mu_Q^-} f(\mu), f(\mu_Q)]$ para todo $\lambda \in Q$. Sea $Q \in \Pi = \text{At}(\{E'_\lambda\}) / \approx$ una clase de equivalencia y notemos que, si $\lambda_1 < \dots, \lambda_n \in Q$ entonces

$$\begin{aligned} \sum_{i=1}^n \tau(p'(\lambda_i)) &= \sum_{i=1}^n (\lim_{\lambda \rightarrow \lambda_i^-} \tau(E'_\lambda) - \tau(E'_{\lambda_i})) \leq \lim_{\lambda \rightarrow \lambda_1^-} \tau(E'_\lambda) - \tau(E'_{\lambda_n}) \\ &\leq \lim_{\mu \rightarrow \mu_Q^-} \tau(E'_{f(\mu)}) - \tau(E'_{f(\mu_Q)}) = \tau(p(\mu_Q)) \end{aligned}$$

donde $p'(\lambda)$ es la proyección salto de la resolución $\{E'_\lambda\}$ en λ y $p(\mu_Q)$ es la proyección salto de la resolución $\{E_\lambda\}$ en μ_Q . Tomando límite sobre n si fuera necesario, entonces tenemos que $\sum_{\lambda \in Q} \tau(p'(\lambda)) \leq \tau(p(\mu_Q))$. Concluimos que

$$\mathcal{J}(\{E'_\lambda\}) = \sum_{Q \in \Pi} \sum_{\lambda \in Q} \tau(p'(\lambda)) \leq \sum_{Q \in \Pi} \tau(p(\mu_Q)) \leq \mathcal{J}(\{E_\lambda\})$$

ya que se trata de series de términos positivos.

□

Lema 3.2.4. Sea $\{E_\lambda\}_{\lambda \in [\alpha, \beta]} \subseteq \mathcal{M}$ una REDA de $p \in \mathcal{P}(\mathcal{M})$, donde $(\mathcal{M}, \tau)$ es un factor II_1 , y sea $\lambda_0 \in \text{At}(\{E_\lambda\})$. Entonces, existe un refinamiento fuerte $(\{E'_\lambda\}_{\lambda \in I'}, f)$ de $\{E_\lambda\}$, donde $I' = [\alpha, \beta + \tau(p(\lambda_0))]$ tal que $\mathcal{J}(\{E'_\lambda\}) = \mathcal{J}(\{E_\lambda\}) - \tau(p(\lambda_0))$ y para cada $\lambda \in I'$ tenemos $f(\lambda) - \lambda \leq \tau(p(\lambda_0))$. Más aún, si $\mathcal{A} \subseteq \mathcal{M}$ es una masa y $\{E_\lambda\}$ es una REDA en $\mathcal{A}$ entonces podemos elegir $\{E'_\lambda\}$ en $\mathcal{A}$.

Demostración. Por simplicidad, asumimos que $I = [0, \beta]$ ($\alpha = 0$). El caso general se deduce de este por reparametrización. Sea $\lambda_0 \in \text{At}(\{E_\lambda\})$, $p_0 = p(\lambda_0)$ la proyección de salto en λ_0 y $\alpha_0 = \tau(p_0)$. Usando el Lema 3.2.1 existe una REDA uniformemente distribuida de p_0 , $\{U_\lambda\}_{\lambda \in [0, \alpha_0]}$. Sea

$$E'_\lambda = \begin{cases} E_\lambda & \text{si } 0 \leq \lambda < \lambda_0 \\ E_{\lambda_0} + U_{\lambda - \lambda_0} & \text{si } \lambda_0 \leq \lambda \leq \lambda_0 + \alpha_0 \\ E_{\lambda - \alpha_0} & \text{si } \lambda_0 + \alpha_0 < \lambda \leq \alpha_0 + \beta. \end{cases}$$

Es inmediato verificar que $\{E'_\lambda\}_{\lambda \in I'}$ es una REDA de p , donde $I' = [0, \beta + \alpha_0]$. En el caso en que $\{E_\lambda\} \subseteq \mathcal{A}$ para una masa $\mathcal{A} \subseteq \mathcal{M}$ entonces $p(\lambda_0) \in \mathcal{A}$ y podemos elegir la resolución $\{U_\lambda\}$ también en $\mathcal{A}$. La función creciente y continua a derecha $f : I \rightarrow I'$ dada por

$$f(\lambda) = \begin{cases} \lambda & \text{si } 0 \leq \lambda < \lambda_0 \\ \lambda + \alpha_0 & \text{si } \lambda_0 \leq \lambda \leq \beta + \alpha_0 \end{cases} \quad (3.21)$$

es tal que $E_\lambda = E'_{f(\lambda)}$, para $\lambda \in [0, \beta]$. Más aún

$$\text{At}(\{E'_\lambda\}) = f(\text{At}(\{E_\lambda\}) \setminus \{\lambda_0\})$$

y $p(\lambda) = p'(f(\lambda))$ para cada $\lambda \in \text{At}(\{E_\lambda\}) \setminus \{\lambda_0\}$, donde $p'(\lambda)$ es proyección de salto de $\{E'_\lambda\}$ en $\lambda \in \text{At}(\{E'_\lambda\})$. Luego

$$\mathcal{J}(\{E'_\lambda\}) = \sum_{\lambda \in \text{At}(\{E_\lambda\}) \setminus \{\lambda_0\}} \tau(p(f(\lambda))) = \mathcal{J}(\{E_\lambda\}) - \tau(p_0).$$

El resto de las propiedades de f se deducen de la ecuación (3.21). □

Introducimos la siguiente notación para enunciar el Teorema 3.2.6. Si $\{\alpha_k\}_{k \in \mathbb{N}} \in l^1(\mathbb{R}^+)$ es una sucesión de números positivos, decimos que la sucesión $(\{E_\lambda^k\}_{\lambda \in I_k})_{k \in \mathbb{N}} \subseteq \mathcal{M}$ de REDAs de p es $\{\alpha_k\}_{k \in \mathbb{N}}$ -compatible si valen las siguientes condiciones:

1. $\exists 0 \leq \alpha < \beta \in \mathbb{R}$ tal que $I_k = [\alpha, \beta + \sum_{i=1}^k \alpha_i]$ para cada $k \in \mathbb{N}$.
2. $(\{E_\lambda^{k+1}\}, f_k)$ es un refinamiento fuerte de $\{E_\lambda^k\}$ para cada $k \in \mathbb{N}$.
3. $f_k(\lambda) - \lambda \leq \alpha_k$, para cada $\lambda \in I_k$ y cada $k \in \mathbb{N}$.

El siguiente lema elemental será utilizado en la prueba del Teorema 3.2.6.

Lema 3.2.5. Sea $(a_{n,k})_{n,k} \subseteq \mathbb{R}$ tal que $0 \leq a_{n,k} \leq 1$ y $a_{n,k} \leq a_{n',k'}$, siempre que $n \leq n'$ y $k \leq k'$. Entonces

$$\lim_{n \rightarrow \infty} \lim_{k \rightarrow \infty} a_{n,k} = \lim_{k \rightarrow \infty} \lim_{n \rightarrow \infty} a_{n,k} (= \lim_{n \rightarrow \infty} a_{n,n}).$$

Teorema 3.2.6. Sea $(\mathcal{M}, \tau)$ un factor II_1 y sea $\{\alpha_k\}_{k \in \mathbb{N}} \in l^1(\mathbb{R}^+)$. Si $(\{E_\lambda^k\}_{\lambda \in I_k})_{k \in \mathbb{N}} \subseteq \mathcal{M}$ es $\{\alpha_k\}_{k \in \mathbb{N}}$ -compatible entonces existe una REDA $\{E_\lambda\}_{\lambda \in I}$ de p en $\mathcal{M}$ tal que $(\{E_\lambda\}, h_k)$ es un refinamiento fuerte de $\{E_\lambda^k\}$, para cada $k \in \mathbb{N}$. Más aún, si $\mathcal{A} \subseteq \mathcal{M}$ es una masa y cada miembro de la sucesión es una REDA en $\mathcal{A}$ entonces podemos elegir $\{E_\lambda\}$ en $\mathcal{A}$.

Demostración. Por simplicidad asumimos que $\alpha = 0$. El caso general se deduce de este por reparametrización. Sea $I = [0, \beta + \sum_{i=1}^{\infty} \alpha_i]$ y notemos que, ya que $f_k(\lambda) \geq \lambda$ para $\lambda \in I_k$

$$E_\lambda^k = E_{f_k(\lambda)}^{k+1} \leq E_\lambda^{k+1}.$$

Luego, para cada $\lambda \in I$ la sucesión $\{E_\lambda^k\}_{k \in \mathbb{N}}$ es creciente, donde fijamos $E_\lambda^k = 0$ si $\lambda \notin I_k$. Definimos

$$E_\lambda = \bigvee_{k \in \mathbb{N}} E_\lambda^k = \lim_{k \rightarrow \infty} E_\lambda^k \in \mathcal{P}(\mathcal{M}), \quad \lambda \in I \quad (3.22)$$

donde el límite es en la topología fuerte de operadores. Notemos que en el caso en que $E_\lambda^k \in \mathcal{A}$ para una masa $\mathcal{A} \subseteq \mathcal{M}$ entonces $E_\lambda \in \mathcal{A}$ pues $\mathcal{A}$ es cerrada en la topología fuerte de operadores. Para ver que $\{E_\lambda\}_{\lambda \in I}$ es una REDA de p notemos que $E_{\lambda_0} \geq E_\lambda$ si $\lambda_0 \leq \lambda$. Así, $\exists \lim_{\lambda \rightarrow \lambda_0^+} E_\lambda \leq E_{\lambda_0}$ en la topología fuerte de operadores. Si $\{\lambda_n\} \subseteq I$ es una sucesión decreciente tal que $\lim_{n \rightarrow \infty} \lambda_n = \lambda_0$ entonces

$$\begin{aligned} \tau\left(\lim_{n \rightarrow \infty} E_{\lambda_n}\right) &= \lim_{n \rightarrow \infty} \tau(E_{\lambda_n}) = \lim_{n \rightarrow \infty} \lim_{k \rightarrow \infty} \tau(E_{\lambda_n}^k) \\ &= \lim_{k \rightarrow \infty} \lim_{n \rightarrow \infty} \tau(E_{\lambda_n}^k) = \tau\left(\bigvee_{k \in \mathbb{N}} E_{\lambda_0}^k\right) = \tau(E_{\lambda_0}) \end{aligned}$$

donde el cambio del orden de los límites iterados es válido por el Lema 3.2.5. Luego $\{E_\lambda\}_{\lambda \in I}$ es una REDA de p .

Fijamos $k \in \mathbb{N}$ y consideramos la sucesión $\{u_n : I_k \rightarrow I_{k+n}\}_{n \in \mathbb{N}}$ de funciones crecientes y continuas a derecha, dadas inductivamente por $u_1 = f_k$ y $u_n = f_{k+n-1} \circ u_{n-1}$ para $n \geq 2$. Entonces, por inducción, es inmediato verificar que, para cada $n \geq 1$ se tiene

1. $E_\lambda^k = E_{u_n(\lambda)}^{k+n}$,
2. $u_{n+1} \geq u_n$, $\|u_{n+1} - u_n\|_\infty \leq \alpha_{n+1}$,
3. $u_n(\lambda) - u_n(\mu) \geq \lambda - \mu$ si $\lambda, \mu \in I_k$ y $\lambda \geq \mu$.

Sea $h_k : I_k \rightarrow I$ el límite uniforme de la sucesión creciente $\{u_n\}$. Entonces h_k es creciente, continua a derecha, $h_k(\lambda) \geq \lambda$ ($u_1 = f_k$) y $h_k(\lambda) - h_k(\mu) \geq \lambda - \mu$ si $\lambda, \mu \in I_k$.

Sea $\lambda_0 \in I_k$ y notemos que $E_{\lambda_0}^k = E_{u_n(\lambda_0)}^{k+n} \geq E_{h_k(\lambda_0)}^{k+n}$, pues $u_n(\lambda) \leq h_k(\lambda)$. Luego

$$E_{\lambda_0}^k \geq \lim_{n \rightarrow \infty} E_{h_k(\lambda_0)}^{k+n} = E_{h_k(\lambda_0)}. \quad (3.23)$$

Para ver que vale la igualdad en la fórmula (3.23) consideramos

$$\lambda_n := \min\{\lambda \in I_k : u_n(\lambda) \geq h_k(\lambda_0)\}.$$

Entonces $\lambda_n \geq \lambda_{n+1} \geq \lambda_0$, pues $\{u_n\}$ es una sucesión creciente y $\lambda_n \rightarrow \lambda_0^+$. De hecho, si $\lambda > \lambda_0$ y $\lambda - \lambda_0 = \epsilon$ entonces $h_k(\lambda) \geq h_k(\lambda_0) + \epsilon$ y existe $n \in \mathbb{N}$ tal que $u_n(\lambda) > h_k(\lambda_0)$, que implica que $\lambda_0 \leq \lambda_n \leq \lambda$. Finalmente, tenemos

$$E_{h_k(\lambda_0)} \geq E_{u_n(\lambda_n)} \geq E_{u_n(\lambda_n)}^{k+n} = E_{\lambda_n}^k, \quad \forall n \in \mathbb{N}$$

que implica que $E_{h_k(\lambda_0)} \geq \lim_{n \rightarrow \infty} E_{\lambda_n}^k = E_{\lambda_0}^k$. □

Demostración del Teorema 3.1.2. Por simplicidad, asumimos que $I = [0, \beta]$. El caso general se deduce de este por reparametrización. Sea $\{\lambda_n\}_{n \in \mathbb{N}}$ una enumeración del conjunto $\text{At}(\{E_\lambda\})$, donde $N \subseteq \mathbb{N}$ es un segmento inicial y sea $\alpha_n = \tau(p(\lambda_n)) > 0$. Entonces, por la ecuación (3.19), tenemos que $\sum_{n \in N} \alpha_n \leq 1$. Sea $I = I_1$, $\{E_\lambda\} = \{E_\lambda^1\}$ y sea $(\{E_\lambda^2\}_{\lambda \in I_2}, f_1)$ un refinamiento fuerte obtenido como en el Lema 3.2.4, a partir de $\{E_\lambda^1\}_{\lambda \in I_1}$ y del átomo λ_1 . Recordemos que en este caso $I_2 = [0, \beta + \tau(p_1)]$ y definamos $f_1 = g_2 : I_1 \rightarrow I_2$.

Procedemos de forma inductiva, como sigue: supongamos que para $1 \leq t \leq k-1$ tenemos una REDA $\{E_\lambda^t\}_{\lambda \in I_t}$ de p , donde $I_t = [0, \beta + \sum_{j=1}^{t-1} \alpha_j]$ y para $1 \leq i \leq k-2$ funciones crecientes y continua a derecha $f_i : I_i \rightarrow I_{i+1}$ tales que $(\{E_\lambda^{i+1}\}, f_i)$ refina fuertemente a $\{E_\lambda^i\}$ y tales que $f_i(\lambda) - \lambda \leq \alpha_i$ para $\lambda \in I_i$. Supongamos además que para $2 \leq l \leq k-1$ existen funciones crecientes y continuas a derecha $g_l : I \rightarrow I_l$ tales que $E_\lambda = E_{g_l(\lambda)}^l$ para $\lambda \in I$ y tales que

$$\text{At}(\{E_\lambda^l\}) = g_l(\text{At}(\{E_\lambda\}) \setminus \{\lambda_1, \dots, \lambda_{l-1}\})$$

y

$$\mathcal{J}(\{E_\lambda^l\}) = \mathcal{J}(\{E_\lambda\}) - \sum_{j=1}^{l-1} \alpha_j.$$

Consideramos la construcción del Lema 3.2.4 aplicada a la REDA $\{E_\lambda^{k-1}\}_{\lambda \in I_{k-1}}$ y el átomo $g_{k-1}(\lambda_{k-1})$. Luego, obtenemos $\{E_\lambda^k\}_{\lambda \in I_k}$ una REDA de p , donde $I_k = [0, \beta + \sum_{j=1}^{k-1} \alpha_j]$, y una función creciente y continua a derecha $f_{k-1} : I_{k-1} \rightarrow I_k$ tal que $(\{E_\lambda^k\}, f_{k-1})$ es un refinamiento fuerte de $\{E_\lambda^{k-1}\}$; en este caso tenemos $f_{k-1}(\lambda) - \lambda \leq \alpha_{k-1}$. Luego,

$$E_\lambda^{k-1} = E_{f_{k-1}(\lambda)}^k \Rightarrow E_\lambda = E_{g_{k-1}(\lambda)}^{k-1} = E_{f_{k-1}(g_{k-1}(\lambda))}^k \quad (3.24)$$

para $\lambda \in I$. Si definimos $g_k = f_{k-1} \circ g_{k-1} : I \rightarrow I_k$ entonces g_k es una función creciente y continua a derecha tal que $E_\lambda = E_{g_k(\lambda)}^k$ por la ecuación (3.24). Entonces tenemos

$$\text{At}(\{E_\lambda^k\}) = f_{k-1}(\text{At}(\{E_\lambda^{k-1}\}) \setminus \{g_{k-1}(\lambda_k)\}) = g_k(\text{At}(\{E_\lambda\}) \setminus \{\lambda_1, \dots, \lambda_k\}).$$

Más aún, $\mathcal{J}(\{E_\lambda^k\}) = \mathcal{J}(\{E_\lambda^{k-1}\}) - \alpha_{k-1} = \mathcal{J}(\{E_\lambda\}) - \sum_{i=1}^{k-1} \alpha_i$.

De esta forma obtenemos una sucesión de resoluciones $\{E_\lambda^k\}_{\lambda \in I_k}$ donde $I_k = [0, \beta + \sum_{j=1}^{k-1} \alpha_j]$, y funciones crecientes y continuas a derecha $\{f_k : I_k \rightarrow I_{k+1}\}$ para $k \in \mathbb{N}$ en las hipótesis del Teorema 3.2.6. Así, existen una REDA de p $\{E'_\lambda\}_{\lambda \in I'}$ y una sucesión de funciones crecientes y continuas a derecha $h_k : I_k \rightarrow I'$ tales que, para cada $k \in \mathbb{N}$ tenemos $E_\lambda^k = E'_{h_k(\lambda)}$, para cada $\lambda \in I_k$. Por el Lema 3.2.3 entonces $\mathcal{J}(\{E'_\lambda\}) \leq \mathcal{J}(\{E_\lambda^k\})$ para cada $k \in \mathbb{N}$ y luego $\mathcal{J}(\{E'_\lambda\}) = 0$ i.e., $\{E'_\lambda\}$ es continua. Por otra parte, si tomamos $f = h_1 : I_1 (= I) \rightarrow I'$ entonces $E_\lambda = E_\lambda^1 = E'_{h_1(\lambda)} = E'_{f(\lambda)}$ para cada $\lambda \in I$. Entonces, $(\{E'_\lambda\}, f)$ es un refinamiento fuerte y continuo de $\{E_\lambda\}$. □

Capítulo 4

Desigualdades de tipo Jensen

Para una descripción conceptual de los resultados incluidos dentro de este capítulo, así como la relación entre éstos y otros resultados ver la sección “Contexto general del trabajo” del Capítulo 1.

Organización del capítulo En la sección 4.1 probamos desigualdades de tipo Jensen con respecto al preorden espectral, para la clase de funciones convexas monótonas. En el caso especial de los factores de tipo II_1 , obtenemos reformulaciones de estas desigualdades a partir de las caracterizaciones del preorden espectral obtenidas en el capítulo anterior.

En la sección 4.2 obtenemos desigualdades de tipo Jensen con respecto a esperanzas condicionales para funciones convexas arbitrarias en varias variables. En el caso particular de los factores II_1 obtenemos desigualdades de con respecto a la submayorización para funciones convexas arbitrarias en varias variables. Al final de esta sección obtenemos algunos resultados necesarios para el estudio de teoremas de tipo Schur-Horn que realizamos en el capítulo siguiente.

En la sección 4.3 traducimos los resultados obtenidos en el caso del álgebra de matrices $M_n(\mathbb{C})$.

Notaciones Sea $\mathcal{H}$ un espacio de Hilbert y sea $L(\mathcal{H})$ el álgebra de todos los operadores acotado en $\mathcal{H}$. En esta sección, el par $(\mathcal{M}, \tau)$ denota un factor finito con traza f.n.s. tal que $\tau(I) = 1$. Las letras $\mathcal{A}, \mathcal{B}$ denotan C^* -álgebras unitales. $\mathcal{A}_{sa}$ denota el espacio real de los operadores autoadjuntos y $\mathcal{A}^+$ el cono de operadores positivos de $\mathcal{A}$. $\mathcal{P}(\mathcal{M}) \subseteq \mathcal{M}_{sa}$ denota el reticulado de proyecciones ortogonales en $\mathcal{M}$ dotado de la topología fuerte de operadores. Además, si $p, q \in \mathcal{P}(\mathcal{M})$ entonces $p \wedge q$ y $p \vee q$ denotan las proyecciones sobre la intersección y la clausura del subespacio generado por los rangos de estas proyecciones respectivamente. Si $a \in \mathcal{M}$ entonces $R(a)$ denota su rango y $P_{\overline{R(a)}} \in \mathcal{P}(\mathcal{M})$ denota la proyección ortogonal sobre la clausura de su rango y $\ker(a)$ denota su núcleo. Denotamos por $\mathcal{U}_{\mathcal{M}}$ el grupo de operadores unitarios de $\mathcal{M}$.

4.1. Funciones monótonas convexas-preorden espectral

Recordemos que una función real f definida en un segmento (a, b) , donde $-\infty \leq a < b \leq \infty$, es **convexa** si la desigualdad

$$f(\lambda x + (1 - \lambda)y) \leq \lambda f(x) + (1 - \lambda)f(y), \quad (4.1)$$

vale siempre que $a < x < b$, $a < y < b$ y $0 \leq \lambda \leq 1$. Una función g es cóncava si $-g$ es convexa.

Es inmediato verificar que la ecuación (4.1) es equivalente a la condición

$$\frac{f(t) - f(s)}{t - s} \leq \frac{f(u) - f(t)}{u - t}, \quad (4.2)$$

siempre que $a < s < t < u < b$.

Teorema 4.1.1. Sea $\mathcal{A}$ una C^* -álgebra, $\varphi : \mathcal{A} \rightarrow \mathbb{C}$ un estado y f una función convexa definida en algún intervalo (α, β) ($-\infty \leq \alpha < \beta \leq \infty$). Entonces tenemos

$$f(\varphi(a)) \leq \varphi(f(a)),$$

para cada $a \in \mathcal{A}_{sa}$ cuyo espectro está contenido en (α, β) .

Demostración. Sea $t = \varphi(a)$ y notemos que entonces $\alpha < t < \beta$. Si γ es el supremo de los cocientes de la izquierda de la desigualdad (4.2), donde $\alpha < s < t$, entonces γ no es mayor que ninguno de los cocientes de la derecha de la misma desigualdad, para $\mu \in (t, \beta)$. Entonces

$$f(\sigma) \geq f(t) + \gamma(\sigma - t) \quad (\alpha < \sigma < \beta).$$

Luego

$$f(a) - f(t) + \gamma(a - t) \geq 0.$$

Si aplicamos φ a ambos miembros de esta desigualdad, y considerando que φ es un estado, obtenemos la desigualdad requerida. $\square$

En esta sección consideramos funciones monótonas convexas y cóncavas. El siguiente resultado de Brown y Kosaki ([19]) indica que el orden apropiado para establecer una desigualdad del tipo Jensen para esta clase de funciones es el preorden espectral. Para la definición y propiedades elementales de este orden ver la sección 2.4.4, Definición 2.4.15 y Teorema 2.4.16).

Sea $\mathcal{A}$ una álgebra de von Neumann semifinita y sea $v \in \mathcal{A}$ una contracción; entonces, para cada operador positivo $a \in \mathcal{A}$ y cada función convexa, continua y monótona f definida en $[0, +\infty)$ y tal que $f(0) = 0$, Brown y Kosaki probaron que

$$v^* f(a) v \preceq f(v^* a v).$$

El siguiente teorema es una generalización del resultado anterior, en términos de transformaciones positivas y uniales, y funciones convexas monótonas.

Teorema 4.1.2. Sea $\mathcal{A}$ una C^* -álgebra unital, $\mathcal{B}$ un álgebra de von Neumann y $\phi : \mathcal{A} \rightarrow \mathcal{B}$ una transformación positiva y unital. Entonces, para cada función convexa y monótona f , definida en algún intervalo I , y para cada operador autoadjunto $a \in \mathcal{A}$ cuyo espectro está contenido en I se tiene

$$f(\phi(a)) \preceq \phi(f(a)) \quad (4.3)$$

Demostración. De acuerdo con la definición, dado $\alpha \in \mathbb{R}$ debemos probar que existe una proyección $q \in \mathcal{A}$ tal que

$$P^{f(\phi(a))}(\alpha, +\infty) = P^{\phi(a)}\{f > \alpha\} \sim q \leq P^{\phi(f(a))}(\alpha, +\infty).$$

Afirmamos que $P^{\phi(a)}\{f > \alpha\} \wedge P^{\phi(f(a))}(-\infty, \alpha] = 0$. De hecho, tomamos

$$\bar{\eta} \in R\left(P^{\phi(a)}\{f > \alpha\}\right), \quad \|\bar{\eta}\| = 1.$$

Como f es monótona tenemos que $\alpha < f(\langle \phi(a)\bar{\eta}, \bar{\eta} \rangle)$, y usando el Teorema 4.1.1 obtenemos

$$\alpha < \langle \phi(f(a))\bar{\eta}, \bar{\eta} \rangle.$$

Por otro lado, para cada $\bar{\xi} \in R\left(P^{\phi(f(a))}(-\infty, \alpha]\right)$, $\|\bar{\xi}\| = 1$

$$\alpha \geq \langle \phi(f(a))\bar{\xi}, \bar{\xi} \rangle.$$

Teniendo esto en consideración y usando la fórmula de Kaplansky deducimos que

$$\begin{aligned} P^{f(\phi(a))}(\alpha, +\infty) &= P^{f(\phi(a))}(\alpha, +\infty) - \left(P^{f(\phi(a))}(\alpha, +\infty) \wedge P^{\phi(f(a))}(-\infty, \alpha]\right) \\ &\sim \left(P^{f(\phi(a))}(\alpha, +\infty) \vee P^{\phi(f(a))}(-\infty, \alpha]\right) - P^{\phi(f(a))}(-\infty, \alpha] \\ &\leq I - P^{\phi(f(a))}(-\infty, \alpha] = P^{\phi(f(a))}(\alpha, +\infty). \end{aligned}$$

□

Observaciones 4.1.3.

1. Notemos que del teorema anterior no podemos inferir un resultado similar para funciones cóncavas y monótonas pues no es cierto que $a \lesssim b \Rightarrow -b \lesssim -a$. Sin embargo, usando un argumento similar se puede deducir la correspondiente desigualdad de Jensen para tales funciones.
2. Si $0 \in I$ y $f(0) \leq 0$ el mismo resultado vale para transformaciones positivas y contractivas. La prueba sigue las mismas líneas, pero debemos usar el Corolario 4.2.5, que probamos en la sección que sigue, en lugar del Teorema 4.1.1. □

El siguiente resultado es una reformulación del Teorema 4.1.2 utilizando los Teoremas 3.1.15 y 3.1.8 en cada caso correspondiente.

Corolario 4.1.4. Sea $\mathcal{A}$ una C^* -álgebra unital y $\phi : \mathcal{A} \rightarrow \mathcal{M}$ una transformación positiva y unital en un factor semifinito $(\mathcal{M}, \tau)$. Entonces, para cada función convexa y monótona f definida en $[0, +\infty)$ se tiene:

1. Si $\mathcal{M} = L(\mathcal{H})$ con $\mathcal{H}$ un espacio de Hilbert separable, f tal que $f(0) = 0$, y para cada operador positivo a de $\mathcal{A}$ tal que $\phi(a)$ y $\phi(f(a))$ son compactos, entonces existe una isometría parcial $u \in L(\mathcal{H})$ con espacio inicial $\overline{R(f(\phi(a)))}$ tal que:

$$uf(\phi(a))u^* \leq \phi(f(a)) \quad \text{y} \quad \left(uf(\phi(a))u^*\right)\phi(f(a)) = \phi(f(a))\left(uf(\phi(a))u^*\right).$$

2. Si $(\mathcal{M}, \tau)$ es un factor II_1 entonces existen sucesiones $(u_n)_{n \in \mathbb{N}}, (v_n)_{n \in \mathbb{N}} \in \mathcal{M}$ de operadores unitarios tales que

$$f(\phi(a)) \leq \lim_{n \rightarrow \infty} u_n^* \phi(f(a)) u_n$$

y

$$\lim_{n \rightarrow \infty} v_n^* f(\phi(a)) v_n \leq \phi(f(a)).$$

□

4.2. Funciones convexas arbitrarias-submayorización

En esta sección obtenemos desigualdades de tipo Jensen para funciones convexas de varias variables. Resultados relacionados con los que aparecen en esta sección pueden hallarse en [37] y [59].

El siguiente ejemplo de J. S. Aujla y F. C. Silva muestra que la conclusión del Teorema 4.1.2 puede no valer si la función f no es monótona.

Ejemplo 4.2.1. Consideramos la transformación positiva $\phi : \mathcal{M}_4 \rightarrow \mathcal{M}_2$ dada por

$$\phi \left(\begin{pmatrix} A_{11} & A_{12} \\ A_{21} & A_{22} \end{pmatrix} \right) = \frac{A_{11} + A_{22}}{2}$$

Tomamos $f(t) = |t|$ y sea A la matriz

$$A = \left(\begin{array}{cc|cc} -2 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ \hline 0 & 0 & 1 & 1 \\ 0 & 0 & 1 & 1 \end{array} \right)$$

entonces

$$\phi(f(A)) = \begin{pmatrix} 3/2 & 1/2 \\ 1/2 & 1/2 \end{pmatrix} \quad y \quad f(\phi(A)) = \begin{pmatrix} 1/\sqrt{2} & 0 \\ 0 & 1/\sqrt{2} \end{pmatrix}.$$

Cálculos sencillos muestran que

$$\text{rank}(P^{\phi(f(A))}(0,5,+\infty)) = 1 < 2 = \text{rank}(P^{f(\phi(A))}(0,5,+\infty)).$$

□

Sin embargo, se tiene una desigualdad de tipo Jensen con respecto al orden usual para cada función convexa, si la transformación ϕ toma valores en un álgebra conmutativa $\mathcal{B}$. Más generalmente, basta asumir que los operadores $\phi(f(a))$ y $f(\phi(a))$ de $\mathcal{B}$ conmutan. Más aún, nuestros métodos permiten deducir los siguientes resultados para funciones de varias variables.

Definición 4.2.2. Sea U un subconjunto convexo de $\mathbb{R}^n$. Una función $f : U \rightarrow \mathbb{R}$ es convexa si para todo $x, y \in U$ y para todo $0 \leq \lambda \leq 1$ se verifica que

$$f(\lambda x + (1 - \lambda)y) \leq \lambda f(x) + (1 - \lambda)f(y).$$

La siguiente proposición enuncia un hecho elemental bien conocido acerca de las funciones convexas. Omitiremos su demostración.

Proposición 4.2.3. Sea f una función convexa definida en un conjunto abierto y convexo $U \subseteq \mathbb{R}^n$ y sea $K \subseteq U$ compacto. Entonces, existe una familia numerable de funciones lineales $\{f_i\}_{i \geq 1}$ tal que para cada $x \in K$ se verifica

$$f(x) = \sup_{i \geq 1} f_i(x).$$

Teorema 4.2.4. Sean $\mathcal{A}$ y $\mathcal{B}$ C^* -álgebras unitales, $\phi : \mathcal{A} \rightarrow \mathcal{B}$ una transformación positiva y unital, y $f : U \rightarrow \mathbb{R}$ una función convexa. Sea $(a_1, \dots, a_n) \subseteq \mathcal{A}_{sa}$ una familia abeliana $\prod_{i=1}^n \sigma(a_i) \subseteq U$ y tal que $(\phi(a_1), \dots, \phi(a_n), \phi(f(a_1, \dots, a_n))) \subseteq \mathcal{B}_{sa}$ resulta también una familia abeliana. Entonces tenemos

$$f(\phi(a_1), \dots, \phi(a_n)) \leq \phi(f(a_1, \dots, a_n)). \quad (4.4)$$

Más aún, si $\tilde{0} = (0, \dots, 0) \in U$ y $f(\tilde{0}) \leq 0$ entonces la ecuación (4.4) vale si ϕ es positiva y contractiva.

Demostración. Sea $\widehat{\mathcal{B}}$ la C^* -subálgebra abeliana de $\mathcal{B}$ generada por $\phi(a_1), \dots, \phi(a_n)$ y $\phi(f(a_1, \dots, a_n))$. Por otra parte, sea $\{f_i\}_{i \geq 1}$ la sucesión de funciones lineales dadas por la Proposición 4.2.3, tal que para cada $(x_1, \dots, x_n) \in \prod_{i=1}^n \sigma(a_i)$

$$f(x_1, \dots, x_n) = \sup_{i \geq 1} f_i(x_1, \dots, x_n).$$

Como $f \geq f_i$ ($i \geq 1$) tenemos que $f(a_1, \dots, a_n) \geq f_i(a_1, \dots, a_n)$ y entonces obtenemos

$$\phi(f(a_1, \dots, a_n)) \geq \phi(f_i(a_1, \dots, a_n)) = f_i(\phi(a_1), \dots, \phi(a_n)), \quad (4.5)$$

donde la última igualdad vale pues f_i es lineal. Como $f_i(\phi(a_1), \dots, \phi(a_n))$ pertenecen a la C^* -álgebra abeliana $\widehat{\mathcal{B}}$ para cada $i \geq 1$, y $\phi(f(a_1, \dots, a_n))$ también pertenece a la C^* -álgebra abeliana $\widehat{\mathcal{B}}$ tenemos que

$$\phi(f(a_1, \dots, a_n)) \geq \max_{1 \leq i \leq n} [f_i(\phi(a_1), \dots, \phi(a_n))] = \max_{1 \leq i \leq n} [f_i](\phi(a_1), \dots, \phi(a_n)).$$

Pero como $\max_{1 \leq i \leq n} [f_i] \rightarrow f$ uniformemente en conjuntos compactos por el teorema de Dini, entonces

$$\phi(f(a_1, \dots, a_n)) \geq f(\phi(a_1), \dots, \phi(a_n)).$$

Si ϕ es contractiva y $f(0, \dots, 0) \leq 0$, las funciones f_i satisfacen $f_i(0) \leq 0$ y podemos reemplazar (4.5) por:

$$\phi(f(a)) \geq \phi(f_i(a)) \geq f_i(\phi(a)).$$

Repitiendo el mismo argumento llegamos a la desigualdad requerida. □

Corolario 4.2.5. Sean $\mathcal{A}, \mathcal{B}$ C^* -álgebras y supongamos que $\mathcal{B}$ es conmutativa. Sea $f : U \rightarrow \mathbb{R}$ una función convexa definida en algún conjunto abierto convexo U y $\phi : \mathcal{A} \rightarrow \mathcal{B}$ positiva y unital. Entonces

$$f(\phi(a_1), \dots, \phi(a_n)) \leq \phi(f(a_1, \dots, a_n)), \quad (4.6)$$

para cada familia abeliana $(a_1, \dots, a_n) \subseteq \mathcal{A}_{sa}$ tal que $\prod_{i=1}^n \sigma(a_i) \subseteq U$. Más aún, si $(0, \dots, 0) \in U$ y $f(0, \dots, 0) \leq 0$ entonces la ecuación (4.6) vale para ϕ contractiva.

Sean $\mathcal{C} \subseteq \mathcal{B}$ C^* -álgebras. Recordemos que una esperanza condicional (ver sección 2.3.5 de los Preliminares, Definición 2.3.27) $\mathcal{E} : \mathcal{B} \rightarrow \mathcal{C}$ es una proyección $\mathcal{C}$ -lineal positiva de $\mathcal{B}$ sobre $\mathcal{C}$ de norma

1. Por ejemplo, los estados son esperanzas condicionales. El centralizador de $\mathcal{E}$ es la C^* -subálgebra de $\mathcal{B}$ definida por:

$$\mathcal{B}^{\mathcal{E}} = \{b \in \mathcal{B} : \mathcal{E}(ba) = \mathcal{E}(ab), \forall a \in \mathcal{B}\}$$

Notemos que, para cada $b \in \mathcal{B}^{\mathcal{E}}$, $\mathcal{E}(b) \in \mathcal{Z}(\mathcal{C})$; donde $\mathcal{Z}(\mathcal{C}) = \{c \in \mathcal{C} : cb = bc, \forall b \in \mathcal{C}\}$ es el centro de $\mathcal{C}$.

En [36] Hansen y Pedersen consideraron el espacio vectorial $C_b(X, \mathcal{A})$ de todas las funciones continuas y uniformemente acotadas de un espacio localmente compacto y Hausdorff X , dotado de una medida de Borel μ , a valores en un C^* -álgebra unital $\mathcal{A}$. Es bien conocido el hecho de que este espacio es una C^* -álgebra con la suma puntual, multiplicación puntual, involución puntual y la norma

$$\|g\|_{\infty} = \sup_X \|g(x)\|.$$

Si la función $x \rightarrow \|g(x)\|$ es integrable, la función g se dice μ -integrable y podemos considerar la integral de Bochner

$$\int_X g(x) d\mu(x).$$

Una función $d \in C_b(X, \mathcal{A})$ es llamada densidad siempre que d^*d sea integrable y que se verifique

$$\int_X d^*(x) d(x) d\mu(x) = I.$$

A cada densidad d le está asociada una transformación positiva y unital, $\phi : C_b(X, \mathcal{A}) \rightarrow \mathcal{A}$, definida por

$$\phi(g) = \int_X d^*(x) g(x) d(x) d\mu(x).$$

En este contexto, dado un estado φ de $\mathcal{A}$ y una función convexa f definida en un intervalo abierto I , Hansen y Pedersen probaron ([36]) la siguiente desigualdad de tipo Jensen

$$\varphi \left(f \left(\int_X d^*(x) g(x) d(x) d\mu(x) \right) \right) \leq \varphi \left(\int_X d^*(x) f(g(x)) d(x) d\mu(x) \right) \quad (4.7)$$

para cada $g \in C_b(X, \mathcal{A})_{sa}$ tal que $h = \int_X d^*(x) g(x) d(x) d\mu(x) \in \mathcal{A}^{\varphi}$ y el espectro de $g(x)$ está contenido en I para cada $x \in X$.

El Teorema 4.2.7 es una mejora del resultado anterior (4.7), que nos permite conectar los resultados de Hansen y Pedersen con la mayorización en factores finitos. Pero primero desarrollamos el siguiente

Lema 4.2.6. Sea $\mathcal{B}$ una C^* -álgebra, φ un estado definido en $\mathcal{B}$, y $(b_1, \dots, b_n) \subseteq \mathcal{B}^{\varphi}$ una familia abeliana. Entonces, existe una medida de Borel μ definida en $K := \sigma(b_1, \dots, b_n)$ y una transformación positiva y unital $\Psi : \mathcal{B} \rightarrow L^{\infty}(K, \mu)$ tal que:

i. $\Psi(f(b_1, \dots, b_n)) = f$ para cada $f \in C(K)$.

ii. $\varphi(x) = \int_K \Psi(x)(t) d\mu(t)$ para cada $x \in \mathcal{B}$.

Demostración. Notemos que para cada función continua g definida en $\sigma(b_1, \dots, b_n)$ la asociación

$$g \mapsto \varphi(g(b_1, \dots, b_n)),$$

determina un funcional lineal acotado en $C(\sigma(b_1, \dots, b_n))$. Entonces, por el teorema de representación de Riesz, existe una medida de Borel μ definida en los subconjuntos de Borel de $\sigma(b_1, \dots, b_n)$, tal que para cada función continua g en $\sigma(b_1, \dots, b_n)$,

$$\varphi(g(b_1, \dots, b_n)) = \int_{\sigma(b_1, \dots, b_n)} g(t) d\mu(t).$$

Dado $x \in \mathcal{B}^+$, definimos el siguiente funcional en $C(\sigma(b_1, \dots, b_n))$

$$\varphi_x(g) = \varphi(xg(b_1, \dots, b_n)).$$

Sea $C^*(b_1, \dots, b_n) \subseteq \mathcal{B}^\varphi$ la C^* -álgebra abeliana generada por $b_1, \dots, b_n$. Como para cada operador positivo $y \in C^*(b_1, \dots, b_n)$ se tiene $\varphi(xy) = \varphi(y^{1/2}xy^{1/2}) \leq \|x\|\varphi(y)$, φ_x no solo es positivo y acotado sino que está dominado por φ . Entonces existe una función h_x de $L^\infty(\sigma(b_1, \dots, b_n), \mu)$ tal que, para cada $g \in C(\sigma(b_1, \dots, b_n))$,

$$\varphi(xg(b_1, \dots, b_n)) = \int_{\sigma(b_1, \dots, b_n)} g(t) h_x(t) d\mu(t).$$

La asociación $x \mapsto h_x$ extendida por linealización, define una transformación lineal, positiva y unital $\Psi : \mathcal{B} \rightarrow L^\infty(\sigma(b_1, \dots, b_n), \mu)$ que satisface la condición (i), pues

$$\varphi(f(b_1, \dots, b_n)g(b_1, \dots, b_n)) = \varphi(fg(b_1, \dots, b_n)) = \int_{\sigma(b_1, \dots, b_n)} g(t)f(t) d\mu(t).$$

Para probar (ii), notemos que

$$\varphi(x) = \varphi(x \cdot 1(b_1, \dots, b_n)) = \int_{\sigma(b_1, \dots, b_n)} 1 \Psi(x)(t) d\mu(t) = \int_{\sigma(b_1, \dots, b_n)} \Psi(x)(t) d\mu(t).$$

□

Teorema 4.2.7. Sean $\mathcal{A}, \mathcal{B}$ C^* -álgebras unitales y $\phi : \mathcal{A} \rightarrow \mathcal{B}$ una transformación positiva y unital. Supongamos que existe una esperanza condicional $\mathcal{E} : \mathcal{B} \rightarrow \mathcal{C}$, de $\mathcal{B}$ sobre la C^* -subálgebra $\mathcal{C}$. Entonces para cada función convexa $f : U \rightarrow \mathbb{R}$ definida en algún conjunto abierto y convexo $U \subseteq \mathbb{R}^n$ se verifica que

$$\mathcal{E}(g[f(\phi(a_1), \dots, \phi(a_n))]) \leq \mathcal{E}(g[\phi(f(a_1, \dots, a_n))]). \quad (4.8)$$

donde $(a_1, \dots, a_n) \in \mathcal{A}_{sa}$ es una familia abeliana tal que $\prod_{i=1}^n \sigma(a_i) \subseteq U$, $\phi(a_1), \dots, \phi(a_n) \in \mathcal{B}^\mathcal{E}$ también forman una familia abeliana, y $g : I \rightarrow \mathbb{R}$ es una función convexa y creciente definida en algún intervalo abierto, con $\text{Im}(f) \subseteq I$.

Demostración. Supongamos que $\mathcal{E}$ es un estado. Definimos $b_i = \phi(a_i)$ para $1 \leq i \leq n$ y como $b_i \in \mathcal{B}^\varphi$, por el lema anterior existe una medida de Borel μ definida en los subconjuntos borelianos de $\sigma(b_1, \dots, b_n)$ y una transformación positiva y unital $\Psi : \mathcal{B} \rightarrow L^\infty(\sigma(b_1, \dots, b_n), \mu)$ tal que:

i. $\Psi(f(b_1, \dots, b_n)) = f$ para cada $f \in C(\sigma(b_1, \dots, b_n))$.

ii. $\varphi(x) = \int_{\sigma(b_1, \dots, b_n)} \Psi(x)(t) d\mu(t)$ para cada $x \in \mathcal{B}$.

Consideramos la transformación $\Phi : C(\sigma(a_1, \dots, a_n)) \rightarrow L^\infty(\sigma(b_1, \dots, b_n), \mu)$ definida por $\Phi(h) = \Psi(\phi(h(a_1, \dots, a_n)))$. Entonces, Φ es positiva y unital. Más aún,

$$\varphi(\phi(h(a_1, \dots, a_n))) = \int_{\sigma(b_1, \dots, b_n)} \Psi(\phi(h(a_1, \dots, a_n)))(t) d\mu(t) = \int_{\sigma(b_1, \dots, b_n)} \Phi(h)(t) d\mu(t).$$

Sea g una función convexa creciente. Entonces, usando el Teorema 4.2.4,

$$\begin{aligned} \varphi(g[f(\phi(a_1, \dots, a_n))]) &= \varphi(g \circ f(b_1, \dots, b_n)) = \int_{\sigma(b_1, \dots, b_n)} g \circ f(t) d\mu(t) \\ &= \int_{\sigma(b_1, \dots, b_n)} g[f(\Phi(\pi_1), \dots, \Phi(\pi_n))] d\mu(t) \\ &\leq \int_{\sigma(b_1, \dots, b_n)} g[\Phi(f(a_1, \dots, a_n))(t)] d\mu(t) \\ &= \int_{\sigma(b_1, \dots, b_n)} g[\Psi(\phi(f(a_1, \dots, a_n)))(t)] d\mu(t) \\ &\leq \int_{\sigma(b_1, \dots, b_n)} \Psi(g[\phi(f(a_1, \dots, a_n))])(t) d\mu(t) = \varphi(g[\phi(f(a_1, \dots, a_n))]). \end{aligned}$$

El caso general se deduce del anterior, componiendo la esperanza condicional con los estados de la C^* -álgebra $\mathcal{C}$. $\square$

A través del Teorema 4.2.7, podemos obtener una desigualdad de tipo Jensen para funciones convexas arbitrarias con respecto a la submayorización en factores finitos. La submayorización ha sido descrita en la sección 2.4.4 de los Preliminares; vamos a considerar el Teorema 2.4.19 que caracteriza esta noción.

Teorema 4.2.8. Sea $\mathcal{A}$ una C^* -álgebra unital, $(\mathcal{M}, \tau)$ un factor finito y $\phi : \mathcal{A} \rightarrow \mathcal{M}$ una transformación positiva y unital. Entonces para cada función convexa $f : U \rightarrow \mathbb{R}$ definida en algún conjunto abierto y convexo $U \subseteq \mathbb{R}^n$ se verifica que

$$f(\phi(a_1), \dots, \phi(a_n)) \prec_w \phi(f(a_1, \dots, a_n)). \quad (4.9)$$

donde $(a_1, \dots, a_n) \in \mathcal{A}_{sa}$ es una familia abeliana tal que $\prod_{i=1}^n \sigma(a_i) \subseteq U$ y $\phi(a_1), \dots, \phi(a_n) \in \mathcal{M}$ también forman una familia abeliana.

Demostración. Por el Teorema 4.2.7, considerando como esperanza (al centro del factor) el estado tracial del factor finito $\mathcal{M}$ se tiene que

$$\tau[g(f(\phi(a_1), \dots, \phi(a_n)))] \leq \tau[g(\phi(f(a_1, \dots, a_n)))]$$

para cada función convexa y creciente g . Pero entonces el teorema es una consecuencia del Teorema 2.4.19 y la desigualdad anterior. $\square$

El siguiente corolario, que es una consecuencia inmediata de los Teoremas 3.1.8 y 4.2.8 es en realidad equivalente al Teorema 4.2.8.

Corolario 4.2.9. Con las notaciones del teorema anterior, existe un operador $c \in \mathcal{M}_h$ tal que

$$f(\phi(a_1), \dots, \phi(a_n)) \leq c \prec \phi(f(a_1, \dots, a_n)). \quad (4.10)$$

Más aún, podemos asumir que $f(\phi(a_1), \dots, \phi(a_n))$ y c conmutan.

Los siguientes resultados correspondientes al caso de una variable serán necesarios para nuestro estudio de un teorema de tipo Schur-Horn en factores II_1 .

Corolario 4.2.10. Sea $\mathcal{A} \subseteq \mathcal{M}$ una masa del factor II_1 $(\mathcal{M}, \tau)$ y sea $E_{\mathcal{A}}$ la esperanza condicional que preserva la traza sobre $\mathcal{A}$. Entonces, si $b \in \mathcal{M}$ es un operador autoadjunto tenemos que

$$E_{\mathcal{A}}(b) \prec b.$$

Demostración. Por el Teorema 4.2.8 se tiene que para cualquier función convexa $f : I \rightarrow \mathbb{R}$ tal que $\sigma(b) \subseteq I$ vale que

$$f(E_{\mathcal{A}}(b)) \prec_w E_{\mathcal{A}}(f(b)) \Rightarrow \tau(f(E_{\mathcal{A}}(b))) \leq \tau(E_{\mathcal{A}}(f(b))) = \tau(f(b)),$$

donde hemos usado el hecho de que $E_{\mathcal{A}}$ preserva trazas. Entonces el corolario es una consecuencia de este hecho y del Teorema 2.4.20. □

En [11] se puede hallar otra prueba del Corolario 4.2.10.

Corolario 4.2.11. Sea $\mathcal{A} \subseteq \mathcal{M}$ una masa del factor II_1 $(\mathcal{M}, \tau)$ y sea $E_{\mathcal{A}}$ la esperanza condicional que preserva la traza sobre $\mathcal{A}$. Entonces, si $b \in \mathcal{M}$ es un operador autoadjunto tenemos

$$\|E_{\mathcal{A}}(b)\|_1 \leq \|b\|_1.$$

Demostración. Notemos que por el Corolario 4.2.10 tenemos que $E_{\mathcal{A}}(b) \prec b$. Si consideramos la función convexa $f(x) = |x|$, entonces por el Teorema 2.4.20 deducimos que

$$\|E_{\mathcal{A}}(b)\|_1 = \tau(f(E_{\mathcal{A}}(b))) \leq \tau(E_{\mathcal{A}}(f(b))) = \|b\|_1. \quad \square$$

Notemos que hemos probado que la esperanza condicional reduce la norma $\|\cdot\|_1$ solo en el caso de operadores autoadjuntos.

4.3. El caso finito dimensional

Tanto el preorden espectral como la mayorización tienen caracterizaciones bien conocidas en el caso de dimensión finita. El objetivo de esta sección es re-escribir las desigualdades de Jensen previamente obtenidas en términos de estas caracterizaciones de los preórdenes anteriores. A lo largo de esta sección identificamos el espacio $L(\mathbb{C}^n)$ con el espacio de matrices complejas $\mathcal{M}_n(\mathbb{C})$, el espacio vectorial real $L_{sa}(\mathcal{H})$ con el espacio vectorial real $\mathcal{M}_n(\mathbb{C})_h$ de matrices autoadjuntas y el cono positivo $L(\mathcal{H})^+$ con el cono positivo de matrices $\mathcal{M}_n(\mathbb{C})^+$. Dada una matriz autoadjunta A , $(\lambda_1(A), \dots, \lambda_n(A))$ denota los autovalores de A contados con multiplicidad y ordenados de forma decreciente

Comencemos recordando las descripciones del preorden espectral y la (sub)mayorización en $\mathcal{M}_n(\mathbb{C})_h$.

Observación 4.3.1. Sean $A, B \in \mathcal{M}_n(\mathbb{C})_h$.

1. Usando el Teorema 3.1.15 es sencillo ver que las siguientes condiciones son equivalentes

- a) $A \preceq B$.
- b) Existe una matriz unitaria U tal que $(UAU^*)B = B(UAU^*)$ y $UAU^* \leq B$.
- c) $\lambda_i(A) \leq \lambda_i(B)$ ($1 \leq i \leq n$).

2. Cálculos simples muestran que dada una matriz autoadjunta C , las funciones $\mu_C(t)$ consideradas en la definición de la submayorización tienen la siguiente forma

$$\mu_C(t) = \begin{cases} \lambda_1(C) & \text{si } 0 \leq t < \frac{1}{n} \\ \vdots & \\ \lambda_n(C) & \text{si } \frac{n-1}{n} \leq t < 1 \end{cases}$$

Luego, se tiene que

$$\int_0^\alpha \mu_A(t) dt \leq \int_0^\alpha \mu_B(t) dt \quad (\forall \alpha \in \mathbb{R}^+) \iff \sum_{i=1}^k \lambda_i(A) \leq \sum_{i=1}^k \lambda_i(B) \quad (k = 1, \dots, n).$$

En la siguiente proposición, hacemos un breve resumen de las diferentes versiones de la desigualdad de Jensen obtenidas en las secciones anteriores usando las descripciones de la Observación 4.3.1

Proposición 4.3.2. Sea $\mathcal{A}$ una C^* -álgebra unital y $\phi : \mathcal{A} \rightarrow \mathcal{M}_n(\mathbb{C})$ una transformación positiva y unital. Supongamos que $a \in \mathcal{A}_{sa}$ y que $f : I \rightarrow \mathbb{R}$ es una función cuyo dominio es un intervalo que contiene el espectro de a . Entonces

1. Si f es una función monótona convexa entonces, para cada $i \in \{1, \dots, n\}$

$$\lambda_i(f(\phi(a))) \leq \lambda_i(\phi(f(a))).$$

2. Si f es una función convexa

$$\sum_{i=1}^k \lambda_i(f(\phi(a))) \leq \sum_{i=1}^k \lambda_i(\phi(f(a))).$$

Más aún, si $0 \in I$ y $f(0) \leq 0$ entonces las desigualdades de arriba siguen valiendo para transformaciones contractivas ($\phi(I) \leq I$).

Ejemplo 4.3.3. Dadas dos matrices $n \times n$, $A = (a_{ij})$ y $B = (b_{ij})$, denotamos por $A \circ B = (a_{ij}b_{ij})$ su producto de Schur. Es un hecho conocido que la transformación $A \mapsto A \circ B$ es positiva para cada matriz positiva B , y que si además $I \circ B = I$ entonces esta transformación también es unital. Entonces, las desigualdades de arriba se aplican al caso de $\phi : \mathcal{M}_n(\mathbb{C}) \rightarrow \mathcal{M}_n(\mathbb{C})$ dada por $\phi(A) = A \circ B$ donde B tiene las propiedades mencionadas antes. $\square$

Ejemplo 4.3.4. Consideremos la transformación $\phi : \mathcal{M}_m \rightarrow \mathcal{M}_n(\mathbb{C})$ dada por

$$\phi(A) = \sum_{i=1}^r W_i^* A W_i$$

donde $W_1, \dots, W_r$ son matrices $m \times n$ tales que

$$\sum_{i=1}^r W_i^* W_i = I.$$

Estas transformaciones son positivas y unitales como resultado de la condición anterior. Como antes, las conclusiones de la Proposición 4.3.2 se aplican también a estas transformaciones. $\square$

Ejemplo 4.3.5. Como una aplicación final consideremos para cada $\alpha \in (0, 1)$ la transformación positiva y unital $\phi_\alpha : \mathcal{M}_{2n} \rightarrow \mathcal{M}_n(\mathbb{C})$ dada por

$$\phi \left(\begin{pmatrix} A & B \\ C & D \end{pmatrix} \right) = \alpha A + (1 - \alpha)D.$$

Aplicando la Proposición 4.3.2 a estas transformaciones y tomando matrices diagonales por bloques obtenemos que para cada función monótona y convexa f

$$\lambda_i(f(\alpha A + (1 - \alpha)D)) \leq \lambda_i(\alpha f(A) + (1 - \alpha)f(D)) \quad (i = 1, \dots, n)$$

y para funciones convexas arbitrarias f

$$\sum_{i=1}^r \lambda_i(f(\alpha A + (1 - \alpha)D)) \leq \sum_{i=1}^r \lambda_i(\alpha f(A) + (1 - \alpha)f(D)) \quad (r = 1, \dots, n)$$

donde A y D son matrices autoadjuntas cuyo espectro está contenido en el dominio de f . Estas desigualdades han sido probadas por J. S. Aujla y F. C. Silva en [12]. $\square$

Capítulo 5

Teoremas de tipo Schur-Horn

Para una descripción conceptual de los resultados incluidos dentro de este capítulo, así como la relación entre éstos y otros resultados ver la sección “Contexto general del trabajo” del Capítulo 1.

Organización del capítulo En la sección 5.1 desarrollamos una versión del teorema de Schur-Horn para operadores autoadjuntos en $L(\mathcal{H})$. Nuestra versión está basada en una reducción al caso obtenido por A. Neumann mediante el uso del teorema de Weyl-von Neumann. Además obtenemos una versión especial del teorema de Schur-Horn para el caso de los operadores de tipo traza.

En la sección 5.2 obtenemos una versión del teorema de Schur-Horn dentro del contexto de los factores de tipo II_1 , similar a la conjeturada por Arveson y Kadison en [11]. A tal fin, en la sección 5.2.1 consideramos algunas propiedades de las reordenadas de funciones en espacios de medida incluidos en la recta real y algunos argumentos de aproximación, para el caso de medidas difusas de soporte compacto en la recta real. En la sección 5.2.2 usamos los resultados anteriores y algunas construcciones de los factores de tipo II_1 para probar un teorema de tipo Schur-Horn.

Notaciones. Sea $\mathcal{H}$ un espacio de Hilbert separable, $L(\mathcal{H})$ el álgebra de operadores lineales acotados en $\mathcal{H}$, $L_0(\mathcal{H})$ el ideal de operadores compactos, $\mathcal{G}l(\mathcal{H})$ el grupo de operadores inversibles, $L(\mathcal{H})_h$ el espacio real de operadores hermitianos, $L(\mathcal{H})^+$ el cono de operadores semidefinidos positivos, $\mathcal{U}(\mathcal{H})$ el grupo de operadores unitarios, y $\mathcal{G}l(\mathcal{H})^+$ el conjunto de operadores (definidos) positivos inversibles. Denotamos por $L^1(\mathcal{H})$ el ideal de operadores traza en $L(\mathcal{H})$. Además notamos $L^1(\mathcal{H})_h = L^1(\mathcal{H}) \cap L(\mathcal{H})_h$ y $L^1(\mathcal{H})^+ = L^1(\mathcal{H}) \cap L(\mathcal{H})^+$.

El espacio de Banach de las sucesiones complejas absolutamente sumables es $\ell^1(\mathbb{N})$ y $\ell^1_{\mathbb{R}}(\mathbb{N})$ (resp. $\ell^1(\mathbb{N})^+$) es el subconjunto de sucesiones reales (resp. no negativas). Análogamente usamos las notaciones $\ell^\infty(\mathbb{N})$, $\ell^\infty_{\mathbb{R}}(\mathbb{N})$ y $\ell^\infty(\mathbb{N})^+$ para sucesiones acotadas.

Dado un operador $A \in L(\mathcal{H})$, $R(A)$ denota al rango de A , $\ker A$ el núcleo de A , $\sigma(A)$ el espectro de A , A^* el adjunto de A , $\rho(A)$ el radio espectral de A , y $\|A\|$ la norma espectral de A . Decimos que A es una isometría (resp. co-isometría) si $A^*A = I$ (resp. $AA^* = I$).

Consideramos también el cociente $\mathcal{A}(\mathcal{H}) = L(\mathcal{H})/L_0(\mathcal{H})$, que es una C^* -álgebra unital, conocida como el álgebra de Calkin. Dado $T \in L(\mathcal{H})$, el *espectro esencial* de T , notado $\sigma_e(T)$, es el espectro de la clase $T + L_0(\mathcal{H})$ en el álgebra $\mathcal{A}(\mathcal{H})$. La norma esencial $\|T\|_e = \inf\{\|T + K\| : K \in L_0(\mathcal{H})\}$ de T es la norma (cociente) de $T + L_0(\mathcal{H})$ en $\mathcal{A}(\mathcal{H})$. Notemos que $\sigma_e(T)$ es un subconjunto compacto de $\sigma(T)$. Si $S \in L(\mathcal{H})_h$ definimos

$$\alpha^+(S) = \max \sigma_e(S) = \|S\|_e \quad \text{y} \quad \alpha_-(S) = \min \sigma_e(S), \quad (5.1)$$

y si $S = \int_{\sigma(S)} t \, dE(t)$ es la representación espectral de S con respecto a la medida espectral E , consideramos los siguientes operadores compactos:

$$\begin{aligned} S^+ &= \int_{[\alpha^+(S), \|S\|]} (t - \alpha^+(S)) dE(t) , \quad y \\ S_- &= \int_{[-\|S\|, \alpha_-(S)]} (t - \alpha_-(S)) dE(t) . \end{aligned} \quad (5.2)$$

Notemos que $S_- \leq 0 \leq S^+$.

Dado un subconjunto $\mathcal{M}$ de un espacio de Banach $(\mathcal{X}, \|\cdot\|)$, su clausura es denotada por $\overline{\mathcal{M}}$ o $\text{cl}_{\|\cdot\|}(\mathcal{M})$, y su cápsula convexa $\text{conv}(\mathcal{M})$. Por otro lado, dado un subespacio cerrado $\mathcal{S}$ de $\mathcal{H}$, denotamos por $P_{\mathcal{S}}$ la proyección ortogonal (i.e. autoadjunta) sobre $\mathcal{S}$. Si $B \in L(\mathcal{H})$ satisface $P_{\mathcal{S}} B P_{\mathcal{S}} = B$, en algunos casos usaremos la compresión de B a $\mathcal{S}$, (i.e. la restricción de B a $\mathcal{S}$ como un operador de $\mathcal{S}$ en $\mathcal{S}$), y decimos que *consideramos a B actuando en $\mathcal{S}$* .

Finalmente, cuando $\dim \mathcal{H} = n < \infty$, identificamos $\mathcal{H}$ con $\mathbb{C}^n$, $L(\mathcal{H})$ con $\mathcal{M}_n(\mathbb{C})$, y usamos las siguientes notaciones: $\mathcal{M}_n(\mathbb{C})_h$ para $L(\mathcal{H})_h$, $\mathcal{M}_n(\mathbb{C})^+$ para $L(\mathcal{H})^+$, $\mathcal{U}(n)$ para $\mathcal{U}(\mathcal{H})$, y $\mathcal{G}l_n(\mathbb{C})$ para $\mathcal{G}l(\mathcal{H})$.

5.1. Teoremas de tipo Schur-Horn en $L(\mathcal{H})$

En esta sección presentamos una versión diferente del “teorema de Schur-Horn en dimensión infinita” obtenido por A. Neumann en [56]. Nuestra formulación evita algunas distinciones técnicas entre el caso diagonalizable y no diagonalizable. Por otra parte, esta versión se aplica directamente al problema de la admisibilidad para marcos en el caso de dimensión infinita.

Dada una sucesión $\mathbf{a} \in \ell_{\mathbb{R}}^{\infty}(\mathbb{N})$, Neumann [56] definió:

$$U_k(\mathbf{a}) = \sup_{|F|=k} \sum_{i \in F} a_i \quad y \quad L_k(\mathbf{a}) = \sup_{|F|=k} \sum_{i \in F} a_i.$$

Estas nociones generalizan las sumas parciales que aparecen en la definición de mayorización. En la primera parte de esta sección extendemos estas funciones al caso de operadores autoadjuntos arbitrarios en un espacio de Hilbert $\mathcal{H}$. Denotemos por $\mathcal{P}_k$ el conjunto de proyecciones ortogonales sobre subespacios k -dimensionales de $\mathcal{H}$.

Definición 5.1.1. Dado $S \in L(\mathcal{H})_h$, definimos para cualquier $k \in \mathbb{N}$,

$$U_k(S) = \sup_{P \in \mathcal{P}_k} \text{tr}(SP) \quad y \quad L_k(S) = \inf_{P \in \mathcal{P}_k} \text{tr}(SP) = -U_k(-S) .$$

□

Observación 5.1.2. Es sencillo ver que U_k y L_k satisfacen las siguientes propiedades:

1. Para cada $k \in \mathbb{N}$, la función U_k es convexa.
2. Para cada $k \in \mathbb{N}$, las funciones U_k y L_k son unitariamente invariantes, i.e. $U_k(S) = U_k(U^* S U)$, para cada $U \in \mathcal{U}(\mathcal{H})$ y $S \in L(\mathcal{H})_h$. □

Sea $\mathbb{M}$ un conjunto, sea $\mathcal{K}$ un espacio de Hilbert con $\dim \mathcal{K} = |\mathbb{M}|$ y sea $\mathcal{B} = \{e_n\}_{n \in \mathbb{M}}$ una base ortonormal de $\mathcal{K}$. Consideramos las siguientes notaciones:

- a) Para cualquier $\mathbf{a} = (a_n)_{n \in \mathbb{M}} \in \ell^\infty(\mathbb{M})$, denotemos por $M_{\mathcal{B}, \mathbf{a}} \in L(\mathcal{K})$ el operador diagonal dado por $M_{\mathcal{B}, \mathbf{a}} e_n = a_n e_n$, $n \in \mathbb{M}$. Cuando sea claro del contexto qué base se está usando abreviaremos $M_{\mathcal{B}, \mathbf{a}} = M_{\mathbf{a}}$.
- b) Si $\mathbf{a} \in \mathbb{C}^n$, haciendo abuso de notación, denotamos por $M_{\mathbf{a}} \in \mathcal{M}_n(\mathbb{C})$ la matriz diagonal (con respecto a la base canónica de $\mathbb{C}^n$) que tiene las entradas de $\mathbf{a}$ en su diagonal.
- c) La compresión diagonal $\mathcal{C}_{\mathcal{B}} : L(\mathcal{K}) \rightarrow L(\mathcal{K})$ asociada a la base $\mathcal{B}$, está definida por $\mathcal{C}_{\mathcal{B}}(T) = M_{\mathcal{B}, \mathbf{a}}$, donde $\mathbf{a} = (\langle T e_n, e_n \rangle)_{n \in \mathbb{M}}$. $\square$

El siguiente resultado afirma que la Definición 5.1.1 extiende las funciones definidas por Neumann para operadores diagonalizables. En este caso identificamos (como espacios reales normados) el álgebra diagonal con respecto a una base fija con $\ell_{\mathbb{R}}^\infty(\mathbb{N})$.

Proposición 5.1.3. Sea $\mathcal{B} = \{e_n\}_{n \in \mathbb{N}}$ una base ortonormal de un espacio de Hilbert $\mathcal{H}$. Si $\mathbf{a} \in \ell_{\mathbb{R}}^\infty(\mathbb{N})$ entonces, para cada $k \in \mathbb{N}$

$$U_k(M_{\mathcal{B}, \mathbf{a}}) = U_k(\mathbf{a}).$$

Para probar esta Proposición necesitamos los siguientes resultados técnicos.

Lema 5.1.4. Sea $S \in L_0(\mathcal{H})^+$, y denotemos por $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n \geq \dots$ los autovalores positivos de S , con multiplicidad (si $\dim R(S) < \infty$, completamos esta sucesión con ceros). Entonces, para cada $k \in \mathbb{N}$,

$$U_k(S) = \sum_{i=1}^k \lambda_i.$$

Más aún, si $P \in \mathcal{P}_k$ es la proyección sobre el subespacio generado por un conjunto ortonormal de autovectores de $\lambda_1, \dots, \lambda_k$, entonces $U_k(S) = \text{tr}(SP)$.

Demostración. Fijemos $k \in \mathbb{N}$. Basta mostrar que $\text{tr}(SQ) \leq \text{tr}(SP) = \sum_{i=1}^k \lambda_i$ para cada $Q \in \mathcal{P}_k$. Esto es una consecuencia del teorema de Schur (la diagonal está mayorizada por los autovalores, ver ejemplo 2.1.8), que también vale en contexto. $\square$

En [56], Neumann probó el siguiente resultado (Lema 2.17): si $\mathbf{a} \in \ell_{\mathbb{R}}^\infty(\mathbb{N})$,

$$a_i^+ = \max\{a_i - \limsup \mathbf{a}, 0\} \quad \text{y} \quad a_i^- = \min\{a_i - \liminf \mathbf{a}, 0\}, \quad i \in \mathbb{N}, \quad (5.3)$$

entonces, para cada $k \in \mathbb{N}$,

$$U_k(\mathbf{a}) = U_k(\mathbf{a}^+) + k \limsup \mathbf{a} \quad \text{y} \quad L_k(\mathbf{a}) = L_k(\mathbf{a}^-) + k \liminf \mathbf{a}. \quad (5.4)$$

$\square$

El próximo resultado extiende la ecuación (5.4) al caso de un operador autoadjunto. Este hecho es necesario para probar la Proposición 5.1.3, pero también es una herramienta útil para manejar a las funciones U_k y L_k .

Proposición 5.1.5. Sea $S \in L(\mathcal{H})_h$. Entonces, para cada $k \in \mathbb{N}$,

1. $U_k(S) = U_k(S^+) + k \alpha^+(S)$
2. $L_k(S) = L_k(S_-) + k \alpha_-(S)$

donde $\alpha^+(S)$, $\alpha_-(S)$, S^+ , S_- están definidos en las ecuaciones (5.1) y (5.2). En particular,

$$\lim_{k \rightarrow \infty} \frac{U_k(S)}{k} = \alpha^+(S) = \|S\|_e \quad \text{y} \quad \lim_{k \rightarrow \infty} \frac{L_k(S)}{k} = \alpha_-(S) . \quad (5.5)$$

Demostración. Notemos $\alpha^+ = \alpha^+(S) (= \|S\|_e)$, y definamos

$$P_2 = P_2(S) = E[\|S\|_e, \|S\|] = E[\alpha^+, \|S\|] , \quad (5.6)$$

donde E es la medida espectral de S . Recordemos que

$$S^+ = \int_{[\alpha^+, \|S\|]} (t - \alpha^+) dE(t) = (S - \alpha^+)P_2 .$$

Entonces $S - S^+ = S(I - P_2) + \alpha^+P_2 \leq \alpha^+I$. Así, para cada $k \in \mathbb{N}$ y $Q \in \mathcal{P}_k$,

$$\text{tr}(SQ) = \text{tr}(S^+Q) + \text{tr}((S - S^+)Q) \leq U_k(S^+) + k\alpha^+ , \quad (5.7)$$

que implica que $U_k(S) \leq U_k(S^+) + k\alpha^+$ para cada $k \in \mathbb{N}$.

Para probar la otra desigualdad, supongamos primero que $\text{tr } P_2 = +\infty$. Sean $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n \geq \dots$ los autovalores de S^+ , elegidos como en el Lema 5.1.4. Sea $Q_k \in \mathcal{P}_k$ la proyección sobre el subespacio generado por un conjunto ortonormal de autovectores de $\lambda_1, \dots, \lambda_k$. Luego $Q_k \leq P_2$ y, por el Lema 5.1.4,

$$\text{tr}(SQ_k) = \text{tr}(S^+Q_k) + \text{tr}((S - S^+)Q_k) = \sum_{i=1}^k \lambda_i + k\alpha^+ = U_k(S^+) + k\alpha^+ .$$

Concluimos que $U_k(S) = U_k(S^+) + k\alpha^+$. Si asumimos que $\text{tr } P_2 = r < \infty$ entonces, si $k \leq r$, el argumento anterior muestra que $U_k(S) = U_k(S^+) + k\alpha^+$. Sea $k > r$ y tomemos $\varepsilon > 0$; como $P_\varepsilon = E[\alpha^+ - \varepsilon, \alpha^+)$ tiene rango infinito (de otra forma $\|S\|_e \leq \alpha^+ - \varepsilon$), podemos tomar $Q \leq P_\varepsilon$ una proyección de rango $k - r$. Si $Q_k = Q + P_2$, entonces

$$\begin{aligned} U_k(S) &\geq \text{tr}(SQ_k) = \text{tr}(SP_2) + \text{tr}(SQ) \\ &= \text{tr}(S^+) + r\alpha^+ + \text{tr}(SP_\varepsilon Q) \\ &\geq \text{tr}(S^+) + r\alpha^+ + (k - r)(\alpha^+ - \varepsilon) \\ &= U_k(S^+) + k\alpha^+ - \varepsilon(k - r) . \end{aligned}$$

Como ε es arbitrario, $U_k(S) = U_k(S^+) + k\alpha^+$. La fórmula para $L_k(S)$ se deduce de aplicar el ítem 1 a $-S$. Como $S^+ \in L_0(\mathcal{H})^+$, entonces sus autovalores convergen a cero. Luego, por el Lema 5.1.4, obtenemos que $\lim_{k \rightarrow \infty} \frac{U_k(S^+)}{k} = 0$ y de forma análoga para $L_k(S_-)$. De esta manera, la ecuación (5.5) resulta evidente. □

Demostración de la Proposición 5.1.3. Se deduce usando en Lema 5.1.4, Proposición 5.1.5, ecuación (5.4) y de las siguiente identidades: si $S = M_{\mathcal{B}, \mathbf{a}}$, entonces

$$1. \alpha^+(S) = \limsup \mathbf{a}, \text{ y } \alpha_-(S) = \liminf \mathbf{a}.$$

$$2. S^+ = M_{\mathcal{B}, \mathbf{a}^+} \text{ y } S_- = M_{\mathcal{B}, \mathbf{a}^-},$$

donde $\mathbf{a}^+$ y $\mathbf{a}^-$ están definidas por la ecuación (5.3). $\square$

Para enunciar la versión infinito dimensional del teorema de Schur-Horn, introducimos las siguientes nociones:

Definición 5.1.6. Sea $\mathcal{H}$ un espacio de Hilbert, $S \in L(\mathcal{H})$ y $\mathcal{B}$ una base ortonormal de $\mathcal{H}$. Entonces,

$$1. \mathcal{U}_{\mathcal{H}}(S) = \{U^*SU : U \in \mathcal{U}(\mathcal{H})\}.$$

$$2. \mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)] = \{\mathbf{c} \in \ell^\infty(\mathbb{N}) : M_{\mathcal{B}, \mathbf{c}} \in \mathcal{C}_{\mathcal{B}}(\mathcal{U}_{\mathcal{H}}(S))\}.$$

$\square$

Observación 5.1.7. Dado $S \in L(\mathcal{H})$, la definición de $\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]$ no depende de la base ortonormal $\mathcal{B}$. De hecho, si $\mathcal{B}'$ es otra base ortonormal de $\mathcal{H}$, $U \in \mathcal{U}(\mathcal{H})$ manda $\mathcal{B}$ sobre $\mathcal{B}'$, y $\mathbf{c} \in \ell^\infty(\mathbb{N})^+$ satisface $M_{\mathcal{B}, \mathbf{c}} = \mathcal{C}_{\mathcal{B}}(T)$ para algún $T \in \mathcal{U}_{\mathcal{H}}(S)$, entonces

$$M_{\mathcal{B}', \mathbf{c}} = UM_{\mathcal{B}, \mathbf{c}}U^* = U\mathcal{C}_{\mathcal{B}}(T)U^* = \mathcal{C}_{\mathcal{B}'}(UTU^*) \in \mathcal{C}_{\mathcal{B}'}(\mathcal{U}_{\mathcal{H}}(S)).$$

De esta forma $\{\mathbf{c} \in \ell^\infty(\mathbb{N}) : M_{\mathcal{B}', \mathbf{c}} \in \mathcal{C}_{\mathcal{B}'}(\mathcal{U}_{\mathcal{H}}(S))\} = \mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]$. $\square$

Dado un operador diagonal $M_{\mathbf{a}}$, Neumann mostró que, si $\mathbf{a}, \mathbf{c} \in \ell_{\mathbb{R}}^\infty(\mathbb{N})$, los siguiente enunciados son equivalentes (ver el Corolario 2.18 y el Teorema 3.13 de [56]):

$$1. \mathbf{c} \in \text{cl}_{\|\cdot\|_\infty}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(M_{\mathbf{a}})]).$$

$$2. U_k(\mathbf{a}) \geq U_k(\mathbf{c}) \text{ y } L_k(\mathbf{a}) \leq L_k(\mathbf{c}), k \in \mathbb{N}.$$

Nuestro objetivo es probar la generalización de esta equivalencia para operadores $S \in L(\mathcal{H})_h$ (mediante una reducción al caso diagonalizable).

Lema 5.1.8. Sea $S, T \in L(\mathcal{H})_h$, entonces $S \in \text{cl}_{\|\cdot\|}(\mathcal{U}_H(T))$ si y solo si

$$\text{cl}_{\|\cdot\|}(\mathcal{U}_H(S)) = \text{cl}_{\|\cdot\|}(\mathcal{U}_H(T)).$$

En este caso $U_k(S) = U_k(T)$ y $L_k(S) = L_k(T)$ para cada $k \in \mathbb{N}$.

Demostración. Si $\{V_n\}_{n \in \mathbb{N}}$ es una sucesión en $\mathcal{U}(\mathcal{H})$ tal que $\|V_nTV_n^* - S\| \xrightarrow{n \rightarrow \infty} 0$, entonces

$$\|V_n^*SV_n - T\| = \|V_n(S - V_nTV_n^*)V_n^*\| = \|V_nTV_n^* - S\| \xrightarrow{n \rightarrow \infty} 0.$$

Así $\text{cl}_{\|\cdot\|}(\mathcal{U}_H(S)) = \text{cl}_{\|\cdot\|}(\mathcal{U}_H(T))$. Por la Observación 5.1.2 tenemos que

$$U_k(V_nTV_n^*) = U_k(T) \text{ y } L_k(V_nTV_n^*) = L_k(T), k \in \mathbb{N}.$$

Supongamos que $S \in \text{cl}_{\|\cdot\|}(\mathcal{U}_H(T))$, con lo cual existe una sucesión $(V_n)_{n \in \mathbb{N}} \subseteq \mathcal{U}_{\mathcal{H}}$ tal que $V_nTV_n^* \xrightarrow{n \rightarrow \infty} S$, fijemos $k \in \mathbb{N}$ y tomemos $P \in \mathcal{P}_k$. Entonces

$$\text{tr } SP = \lim_{n \rightarrow \infty} \text{tr } V_nTV_n^*P \leq \lim_{n \rightarrow \infty} U_k(V_nTV_n^*) = U_k(T).$$

De esta forma $U_k(S) \leq U_k(T)$. Análogamente $L_k(S) \geq L_k(T)$. Las desigualdades reversas se deducen del hecho de que $V_n^*SV_n \xrightarrow{n \rightarrow \infty} T$. $\square$

Teorema 5.1.9. Sea $S \in L(\mathcal{H})_h$ y $\mathbf{c} \in \ell_{\mathbb{R}}^{\infty}(\mathbb{N})$. Luego, las siguientes condiciones son equivalentes:

1. $\mathbf{c} \in \text{cl}_{\|\cdot\|_{\infty}}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)])$.
2. $U_k(S) \geq U_k(\mathbf{c})$ y $L_k(S) \leq L_k(\mathbf{c})$ para cada $k \in \mathbb{N}$.

En este caso,

$$\max \sigma_e(S) \geq \limsup \mathbf{c} \quad \text{y} \quad \min \sigma_e(S) \leq \liminf \mathbf{c}. \quad (5.8)$$

Demostración. Como hemos mencionado, el caso diagonalizable fue probado por Neumann. Si S no es diagonalizable, por el Teorema 2.1.18, existe un operador diagonalizable $D \in \text{cl}_{\|\cdot\|}(\mathcal{U}_{\mathcal{H}}(S))$. Por el Lema 5.1.8,

$$U_k(D) = U_k(S), \quad L_k(D) = L_k(S) \quad \text{para cada } k \in \mathbb{N},$$

y $\text{cl}_{\|\cdot\|}(\mathcal{U}_{\mathcal{H}}(D)) = \text{cl}_{\|\cdot\|}(\mathcal{U}_{\mathcal{H}}(S))$. Esto implica que

$$\text{cl}_{\|\cdot\|_{\infty}}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(D)]) = \text{cl}_{\|\cdot\|_{\infty}}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]),$$

pues la función $T \mapsto \mathcal{C}_{\mathcal{B}}(T)$ es continua, para cada base ortonormal $\mathcal{B}$. Así, el caso general se reduce al caso diagonalizable. La ecuación (5.8) se sigue del hecho de que

$$\limsup \mathbf{c} = \lim_{k \rightarrow \infty} \frac{U_k(\mathbf{c})}{k} \quad \text{y} \quad \liminf \mathbf{c} = \lim_{k \rightarrow \infty} \frac{L_k(\mathbf{c})}{k}, \quad (5.9)$$

y la ecuación (5.5). □

Un resultado similar vale para operadores autoadjuntos en $L^1(\mathcal{H})$. Para enunciarlo consideramos las siguientes nociones: sea Π el conjunto de todas las funciones biyectivas en $\mathbb{N}$ y, para cada $k \in \mathbb{N}$, denotemos por $\Pi_k \subseteq \Pi$ el conjunto de permutaciones σ tales que $\sigma(n) = n$ para cada $n > k$. Dado $\mathbf{a} \in \ell_{\mathbb{R}}^{\infty}(\mathbb{N})$ y $\sigma \in \Pi$, notamos:

1. $\mathbf{a}_{\sigma} = (a_{\sigma(1)}, a_{\sigma(2)}, \dots)$.
2. $\Pi \cdot \mathbf{a} = \{\mathbf{a}_{\sigma}, \sigma \in \Pi\}$, la órbita de $\mathbf{a}$, bajo la acción de Π .
3. $\text{conv}(\Pi \cdot \mathbf{a})$, la cápsula convexa de la órbita de $\mathbf{a}$. □

5.1.10. Si $\mathbf{b}, \mathbf{a}$ son sucesiones en $\ell_{\mathbb{R}}^1(\mathbb{N})$, Neumann [56] (ver también la sección 2.1.12 de los Preliminares) probó que los siguientes enunciados son equivalentes:

1. $\mathbf{b} \in \text{cl}_{\|\cdot\|_1}(\text{conv}(\Pi \cdot \mathbf{a}))$
2. a) $U_k(\mathbf{a}) \geq U_k(\mathbf{b})$ y $L_k(\mathbf{a}) \leq L_k(\mathbf{b})$ para todo $k \in \mathbb{N}$.
b) $\sum_{k=1}^{\infty} b_k = \sum_{k=1}^{\infty} a_k$. □

Vamos a extender este resultado en el siguiente teorema de tipo Schur-Horn de mayorización en el sentido ℓ^1 .

Teorema 5.1.11. Sea $S \in L^1(\mathcal{H})_h$, y $\mathbf{b} \in \ell_{\mathbb{R}}^1(\mathbb{N})$. Entonces, los siguiente son equivalentes

1. $\mathbf{b} \in \text{cl}_{\|\cdot\|_1}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)])$.

2. $U_k(S) \geq U_k(\mathbf{b})$, $L_k(S) \leq L_k(\mathbf{b})$ para cada $k \in \mathbb{N}$, y $\sum_{k=1}^{\infty} b_k = \text{tr } S$.

Demostración. $1 \rightarrow 2$. Evidentemente $\text{cl}_{\|\cdot\|_1}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]) \subseteq \text{cl}_{\|\cdot\|_{\infty}}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)])$. Luego, por la Proposición 5.1.9,

$$U_k(S) \geq U_k(\mathbf{b}) \text{ y } L_k(S) \leq L_k(\mathbf{b}) \text{ para cada } k \in \mathbb{N}.$$

La igualdad $\sum_{k=1}^{\infty} b_k = \text{tr } S$ vale siempre que $\mathbf{b} \in \mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]$. El caso general se sigue de la $\ell^1(\mathbb{N})$ -continuidad de la función $\mathbf{b} \mapsto \sum_{k=1}^{\infty} b_k$.

$2 \rightarrow 1$. Sea $\mathbf{a} \in \ell_{\mathbb{R}}^1(\mathbb{N})$ y $\mathcal{B} = \{e_k\}_{k \in \mathbb{N}}$ una base ortonormal de $\mathcal{H}$ tal que $S = M_{\mathcal{B}, \mathbf{a}}$. Por 5.1.10 y la Proposición 5.1.3, basta mostrar que

$$\text{cl}_{\|\cdot\|_1}(\text{conv}(\Pi \cdot \mathbf{a})) \subseteq \text{cl}_{\|\cdot\|_1}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]).$$

Notemos que $\text{cl}_{\|\cdot\|_1}(\text{conv}(\Pi \cdot \mathbf{a})) = \text{cl}_{\|\cdot\|_1}(\text{conv}(\Pi_0 \cdot \mathbf{a}))$, donde $\Pi_0 = \bigcup_{k \in \mathbb{N}} \Pi_k$. De hecho, basta ver que $\Pi \cdot \mathbf{a} \subseteq \text{cl}_{\|\cdot\|_1}(\text{conv}(\Pi_0 \cdot \mathbf{a}))$. Dado $\sigma \in \Pi$, $\mathbf{a}_{\sigma} \in \Pi \cdot \mathbf{a}$, y $\varepsilon > 0$, consideremos $N \in \mathbb{N}$ tal que $\sum_{k > N} |a_k| < \frac{\varepsilon}{2}$ y $N_0 \in \mathbb{N}$ tal que $\sigma^{-1}(\mathbb{I}_N) \subseteq \mathbb{I}_{N_0}$. Definimos $\sigma_0 \in \Pi_{N_0}$ de forma que $\sigma(k) = \sigma_0(k)$ para cada $k \in \mathbb{I}_{N_0}$ tal que $\sigma(k) \in \mathbb{I}_N$. Entonces tenemos

$$\begin{aligned} \|\mathbf{a}_{\sigma} - \mathbf{a}_{\sigma_0}\|_1 &= \sum_{\sigma(k) \notin \mathbb{I}_N} |a_{\sigma(k)} - a_{\sigma_0(k)}| \\ &\leq \sum_{\sigma(k) \notin \mathbb{I}_N} |a_{\sigma(k)}| + \sum_{\sigma_0(k) \notin \mathbb{I}_N} |a_{\sigma_0(k)}| < \varepsilon. \end{aligned}$$

Si consideramos $\mathbf{b} \in \text{conv}(\Pi_0 \cdot \mathbf{a})$ entonces, existe $n \in \mathbb{N}$ tal que $\mathbf{b} \in \text{conv}(\Pi_n \cdot \mathbf{a})$. Esto significa que las primeras n entradas de $\mathbf{b}$ son una combinación convexa de permutaciones de las primeras n entradas de $\mathbf{a}$, y $b_k = a_k$ para todo $k > n$. De aquí que $(b_1, \dots, b_n) \prec (a_1, \dots, a_n)$. Denotamos por $\mathcal{B}_n = \{e_k : k \leq n\}$ y $\mathcal{H}_n = \text{span}\{\mathcal{B}_n\}$. Entonces, por el teorema de Schur-Horn 2.1.9 para mayorización de vectores, existe un unitario $U_0 \in L(\mathcal{H}_n)$ tal que

$$M_{\mathcal{B}, \mathbf{b}}|_{\mathcal{H}_n} = \mathcal{C}_{\mathcal{B}_n}(U_0^* M_{\mathcal{B}, \mathbf{a}}|_{\mathcal{H}_n} U_0).$$

Si tomamos

$$U = \begin{pmatrix} U_0 & 0 \\ 0 & I \end{pmatrix} \begin{matrix} \mathcal{H}_n \\ \mathcal{H}_n^{\perp} \end{matrix} \in \mathcal{U}(\mathcal{H}),$$

obtenemos que $M_{\mathcal{B}, \mathbf{b}} = \mathcal{C}_{\mathcal{B}}(U^* M_{\mathcal{B}, \mathbf{a}} U)$, y $\mathbf{b} \in \mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]$. Así,

$$\text{cl}_{\|\cdot\|_1}(\text{conv}(\Pi \cdot \mathbf{a})) = \text{cl}_{\|\cdot\|_1}(\text{conv}(\Pi_0 \cdot \mathbf{a})) \subseteq \text{cl}_{\|\cdot\|_1}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]),$$

que completa la prueba. □

Observación 5.1.12. Comparando 5.1.10 con el Teorema 5.1.11 vemos que, si $S = M_{\mathcal{B}, \mathbf{a}}$ para alguna $\mathbf{a} \in \ell_{\mathbb{R}}^1(\mathbb{N})$ y alguna base ortonormal $\mathcal{B}$ de $\mathcal{H}$, entonces

$$\text{cl}_{\|\cdot\|_1}(\text{conv}(\Pi \cdot \mathbf{a})) = \text{cl}_{\|\cdot\|_1}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)]).$$

En particular, $\text{cl}_{\|\cdot\|_1}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)])$ es un conjunto convexo.

Por otro lado, como las funciones U_k son convexas y las funciones L_k son cóncavas, $k \in \mathbb{N}$, del Teorema 5.1.9 podemos concluir que el conjunto $\text{cl}_{\|\cdot\|_{\infty}}(\mathcal{C}[\mathcal{U}_{\mathcal{H}}(S)])$ es convexo, para cada $S \in L(\mathcal{H})_h$. En realidad, este hecho es conocido y puede ser deducido de los resultados de Neumann (ver los Teoremas 2.1.16 y 2.1.17 de los Preliminares).

5.2. Un teorema de tipo Schur-Horn para factores II_1

Vamos a reescribir el teorema de Schur-Horn para mayorización de vectores en términos relacionados con el factor finito de tipo I_n $M_n(\mathbb{C})$. Sea $\mathcal{B} = \{e_i\}_{i=1}^n$ una base ortonormal de $\mathbb{C}^n$ y denotemos por $\mathcal{D}_{\mathcal{B}}$ el álgebra diagonal con respecto a esta base, es decir las matrices $D \in M_n(\mathbb{C})$ tales que $De_i = \lambda_i e_i$ para ciertos $\lambda_i \in \mathbb{C}$ con $1 \leq i \leq n$. Esta álgebra tiene la siguiente propiedad: si $A \in M_n(\mathbb{C})$ es tal que conmuta con el álgebra $\mathcal{D}$ es decir $AD = DA$ para toda $D \in \mathcal{D}$ entonces $A \in \mathcal{D}$. Esta propiedad implica que no existe un álgebra abeliana $\mathcal{D}'$ que contiene de forma estricta a $\mathcal{D}$. De esta forma, $\mathcal{D}$ es un subálgebra abeliana maximal (que abreviamos masa) del factor $M_n(\mathbb{C})$. Recíprocamente, si $\mathcal{D} \subseteq M_n(\mathbb{C})$ es una masa entonces existe una base ortonormal $\{e_i\}_{i=1}^n$ de forma que $\mathcal{D}$ es el álgebra diagonal con respecto a esta base. Dada una masa $\mathcal{D}$ que es el álgebra diagonal con respecto a la base ortonormal $\{e_i\}_{i=1}^n$ recordemos (ver notaciones antes de la Proposición 5.1.3) que la compresión a $\mathcal{D}$, $\mathcal{C}_{\mathcal{B}} : M_n(\mathbb{C}) \rightarrow \mathcal{D}$ está dada por

$$\mathcal{C}_{\mathcal{B}}(A) = M_{\mathcal{B}, \mathbf{a}}$$

donde $\mathbf{a} = (\langle Ae_i, e_i \rangle)_{i=1}^n$. Esta transformación resulta una proyección positiva, unital y preserva la traza de $M_n(\mathbb{C})$ en el sentido $\text{tr}(A) = \text{tr}(\mathcal{C}_{\mathcal{B}}(A))$ y es la única transformación lineal con estas propiedades.

Si $x, y \in \mathbb{R}^n$ son tales que $x \prec y$ entonces el teorema de Schur-Horn (ver 2.1.9) es equivalente a la existencia de una matriz unitaria $U \in M_n(\mathbb{C})$ tal que

$$\mathcal{C}_{\mathcal{B}}(U^* M_y U) = M_{\mathcal{B}, x}.$$

Esta última versión puede leerse de la siguiente manera: si $D \in \mathcal{D}$ es una matriz autoadjunta (y diagonal con respecto a la base $\mathcal{B}$) y $A \in M_n(\mathbb{C})$ es una matriz autoadjunta tal que $D \prec A$ entonces existe una matriz unitaria tal que

$$\mathcal{C}_{\mathcal{B}}(U^* A U) = D$$

con lo cual se tiene

$$\mathcal{C}_{\mathcal{B}}(\mathcal{U}(n)(A)) = \{D \in \mathcal{D} : D \prec A\}$$

donde $\mathcal{U}(n)(A) = \{U^* A U, U \in \mathcal{U}(n)\}$ es la órbita unitaria de A , siendo $\mathcal{U}(n)$ el grupo compacto de unitarios de $M_n(\mathbb{C})$. Más aún, siguiendo ideas similares, el Teorema de tipo schur-Horn 5.1.9 puede re-escribirse como la igualdad

$$\text{cl}_{\|\cdot\|_{\infty}}(\mathcal{C}_{\mathcal{B}}(\mathcal{U}_L(\mathcal{H})(S))) = \{D \in \mathcal{D}_{\mathcal{B}} : D \prec S\}$$

donde $\mathcal{D}_{\mathcal{B}}$ denota el álgebra diagonal con respecto a la base ortonormal $\{e_i\}_{i \in \mathbb{N}}$ es decir

$$\mathcal{D}_{\mathcal{B}} = \{M_{\mathcal{B}, \mathbf{a}} : \mathbf{a} \in \ell^{\infty}(\mathbb{N})\}$$

y $\mathcal{C}_{\mathcal{B}}$ denota la (única) proyección positiva, unital y que preserva la traza sobre $\mathcal{D}_{\mathcal{B}}$. Más aún, el Teorema 5.1.11 (Schur-Horn de tipo ℓ^1) puede re-escribirse como la igualdad

$$\text{cl}_{\|\cdot\|_1}(\mathcal{C}_{\mathcal{B}}(\mathcal{U}_{L(\mathcal{H})}(S))) = \{D \in L^1(\mathcal{D}_{\mathcal{B}}) : D \prec S\}$$

donde

$$L^1(\mathcal{D}_{\mathcal{B}}) = \{M_{\mathcal{B}, \mathbf{a}} : \mathbf{a} \in \ell^1(\mathbb{N})\}.$$

En el caso del teorema de Schur-Horn en su versión para operadores en $L(\mathcal{H})$, fijando una base ortonormal $\mathcal{B} = \{e_k\}_{k \in \mathbb{N}}$ de $\mathcal{H}$ entonces la compresión a la diagonal determinada por $\mathcal{B}$, es decir $\mathcal{C}_{\mathcal{B}}(A) = M_{\mathcal{B}, \mathbf{a}}$ donde $M_{\mathcal{B}, \mathbf{a}}$ es el operador diagonal con respecto a la base $\mathcal{B}$ dado por $M_{\mathcal{B}, \mathbf{a}} e_n = \langle A e_n, e_n \rangle e_n$, es una esperanza condicional sobre el “álgebra diagonal determinada por $\mathcal{B}$ ” que conmuta con la traza usual de $L(\mathcal{H})$ es decir $\text{tr}(A) = \text{tr}(\mathcal{C}_{\mathcal{B}}(A))$. Esta álgebra diagonal tiene la siguiente propiedad: si $A \in L(\mathcal{H})$ es tal que $AD = DA$ para todo D en el álgebra diagonal determinada por $\mathcal{B}$ entonces A es un operador diagonal con respecto a la base $\mathcal{B}$. Esto implica que no hay álgebras abelianas en $L(\mathcal{H})$ que contengan estrictamente al álgebra diagonal determinada por $\mathcal{B}$. De aquí que esta álgebra diagonal es una subálgebra abeliana maximal de $L(\mathcal{H})$ o más brevemente una *masa* de $L(\mathcal{H})$.

En el caso de factores de tipo II_1 se conoce la existencia de masas en ellos con distintas propiedades: en particular, dada una masa $\mathcal{A} \subseteq \mathcal{M}$ de un factor II_1 $(\mathcal{M}, \tau)$ entonces existe una única esperanza condicional $E_{\mathcal{A}}$ que conmuta con la traza en el sentido anterior (que es un ejemplo especial de transformación doble estocástica en $(\mathcal{M}, \tau)$). Una cuestión natural es la de caracterizar el conjunto de posibles diagonales con respecto a esta álgebra de un operador autoadjunto b es decir el conjunto

$$E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(b)).$$

En el trabajo [11] Arveson y Kadison han conjeturado que vale la igualdad

$$E_{\mathcal{A}}(\text{cl}_{\|\cdot\|_{\infty}}(\mathcal{U}_{\mathcal{M}}(b))) = \{a \in \mathcal{A} : a \prec b\}.$$

En lo que sigue desarrollamos las herramientas necesarias para probar la caracterización

$$\text{cl}_{\|\cdot\|_1}(E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(b))) = \{a \in \mathcal{A} : a \prec b\}.$$

En tanto caracterización del conjunto $\{a \in \mathcal{A} : a \prec b\}$ es más débil que la conjeturada por Arveson y Kadison en el sentido de que a priori

$$E_{\mathcal{A}}(\text{cl}_{\|\cdot\|_{\infty}}(\mathcal{U}_{\mathcal{M}}(b))) \subseteq \text{cl}_{\|\cdot\|_1}(E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(b))).$$

Sin embargo notemos que en el contexto de los factores de tipo II_1 todos los operadores son de tipo traza (pues aquí la traza es finita en todo operador acotado). De esta forma, nuestro resultado está más en el espíritu del teorema de Schur-Horn de tipo ℓ^1 descrito anteriormente.

5.2.1. Reordenamientos y aproximaciones de funciones

En lo que sigue ν denotará una medida de probabilidad regular, difusa y de Borel soportada en el conjunto compacto $\text{sop}(\nu) \subseteq [\alpha, \beta]$ y $h : [\alpha, \beta] \rightarrow \mathbb{R}_{\geq 0}$ denotará una función creciente y continua

a izquierda. Más aún, $\mu_h(\cdot)$ son los valores singulares de h con respecto a la medida ν (i.e. como elemento del álgebra de von Neumann $L^\infty(\nu)$).

Reordenamientos

Dada una función $f : X \rightarrow \mathbb{C}$ definida en un espacio de medida (X, ν) definimos su función de distribución $d_f : (0, \infty) \rightarrow [0, \infty]$ por

$$d_f(t) = \nu(\{x \in X : |f(x)| > t\}).$$

A partir de la función de distribución de f definimos la reordenada de f (ver [31]), denotada $f^* : [0, \nu(X)] \rightarrow [0, \infty)$ dada por

$$f^*(s) = \inf\{t : d_f(t) \leq t\}.$$

La reordenada de una función es una función decreciente y continua a derecha que tiene la misma distribución que la función original. Esta noción juega un papel importante en el estudio de aquellas nociones que solo dependen de la función de distribución; por ejemplo, si $\phi : [0, \infty) \rightarrow [0, \infty)$ es una función creciente y diferenciable y tal que $\phi(0) = 0$ entonces

$$\int_X \phi(|f(x)|) d\nu(x) = \int_0^\infty \phi'(t) d_f(t) dt$$

lo cual muestra que la integral de la izquierda solo depende de la función de distribución de f . En particular, si consideramos $\phi(t) = t^p$ ($p \geq 1$) entonces vemos que la norma $\|\cdot\|_p$ de una función sólo depende de la función de distribución.

Consideremos un operador positivo $a \in \mathcal{M}^+$ donde $(\mathcal{M}, \tau)$ es un factor de tipo II_1 . En este caso, definimos ν_a como la medida escalar asociada al operador positivo a y la traza de forma que $\nu_a(\Delta) = \tau(P^a(\Delta))$ para $\Delta \subseteq \mathbb{R}$ un conjunto medible Borel. En este caso ν_a está soportada en el compacto $\sigma(a)$. Más aún se tiene el *-isomorfismo (ver sección 2.4.1, Ejemplo 2.4.2 ítem b)) $C^*(a) \approx C(\sigma(a))$ de forma que a se identifica con la función identidad $f(x) = x$ en $\sigma(a)$. Bajo esta representación espectral de a (bajo la identificación anterior), si consideramos la medida ν_a en $\sigma(a)$ entonces la reordenada de la función identidad f coincide con la función que determina los valores singulares de a . En efecto, basta notar que

$$d_f(t) = \nu_a(\{x \in \sigma(a) : x > t\}) = \tau(P^a(t, \infty))$$

de forma que

$$f^*(s) = \inf\{t \geq 0 : \tau(P^a(t, \infty)) \leq t\} = \mu_a(s)$$

por la caracterización (2.27). Más aún, siguiendo un argumento similar al anterior se prueba que, si g es una función positiva definida en $\sigma(a)$ entonces la reordenada de la función g en el espacio de medida $(\sigma(a), \nu_a)$ coincide con los valores singulares del operador $g(a)$. De esta forma, en lo que sigue denotamos a la reordenada de una función g (en las condiciones anteriores) en un espacio de medida (I, ν) por μ_g , pues ésta coincide con los valores singulares del operador $g \in L^\infty(I, \nu)$.

Lema 5.2.1. Sean $h : [\alpha, \beta] \rightarrow \mathbb{R}_{\geq 0}$ y ν como antes y sea $[\gamma, \delta] \subseteq [\alpha, \beta]$, $\tilde{h} = h \cdot 1_{[\gamma, \delta]}$. Entonces, si definimos $r = \nu([\delta, \beta])$ tenemos que

$$\mu_{\tilde{h}}(t) = \begin{cases} \mu_h(t + r) & \text{si } t \in [0, \nu([\gamma, \delta])] \\ 0 & \text{si } t \in [\nu([\gamma, \delta]), 1]. \end{cases} \quad (5.10)$$

Demostración. Para $t \in [0, 1]$, definimos

$$k(t) = \inf\{s \in [\alpha, \beta] : \nu([\alpha, s]) \geq 1 - t\}.$$

Es evidente que k es decreciente y continua. Afirmamos que $\mu_h(t) = h(k(t))$. En efecto, por continuidad tenemos que $\nu([\alpha, k(t)]) \geq 1 - t$ con lo cual $\mu_h(t) \leq h(k(t))$. Si $\mu_h(t) < h(k(t))$ entonces existe $X \in [\alpha, \beta]$ con $\nu(X) \geq 1 - t$ y $\sup_X h < h(k(t))$. Entonces, del hecho de que h es creciente, existe δ tal que $X \subset [\alpha, k(t) - \delta]$. Pero entonces $k(t) - \delta$ sería una cota inferior para el conjunto $\inf\{s \in [\alpha, \beta] : \nu([\alpha, s]) \geq 1 - t\}$, que es una contradicción.

Denotemos por $Y = [\gamma, \delta]$ y consideremos

$$\mu_{\tilde{h}}(t) = \inf\{\|\tilde{h} 1_X\| : \nu(X) \geq 1 - t\}.$$

Si $t \geq \nu(Y)$, entonces $\nu([\alpha, \beta] \setminus Y) \geq 1 - t$. Si definimos $X = [\alpha, \beta] \setminus Y$, entonces $\nu(X) \geq 1 - t$ y $\tilde{h} 1_X = 0$. Así, en este caso, $\mu_{\tilde{h}}(t) = 0$.

Asumamos entonces que $t < \nu(Y)$, con lo que $1 - t > \nu([\alpha, \beta] \setminus Y)$. Sea

$$X = ([\alpha, \beta] \setminus Y) \cup [\gamma, k'(t)],$$

donde elegimos

$$k'(t) = \inf\{s \in [\gamma, \delta] : \nu([\gamma, s]) \geq \nu(Y) - t\}.$$

Con este conjunto X , $\|\tilde{h} 1_X\| = h(k'(t))$. Para cualquier otra elección X' (salvo diferencia de medida 0) con $\nu(X') \geq 1 - t$ habrá una intersección de medida positiva con $(k'(t), \beta]$, y entonces $\|\tilde{h} 1_{X'}\| > k'(t)$. Luego $\mu_{\tilde{h}}(t) = h(k'(t))$.

Si sumamos $\nu([\alpha, \gamma])$ a ambos lados de la desigualdad que define $k'(t)$, vemos que la desigualdad se transforma en $\nu([\alpha, k'(t)]) \geq \nu([\alpha, \delta]) - t = 1 - (t + \nu([\delta, \beta]))$. Esto muestra que $k'(t) = k(t + \nu([\delta, \beta]))$. $\square$

Observación 5.2.2. En lo anterior el hecho de que la medida ν sea difusa así como la continuidad a izquierda de h no son esenciales. En el caso de que ν tenga átomos entonces $k(t)$ resulta continua a derecha y los cálculos se realizan de la misma forma.

Recordemos que para dado operador positivo $a \in \mathcal{M}^+$ de un álgebra de von Neumann finita $(\mathcal{M}, \tau)$ (con $\tau(1) = 1$) entonces

$$\tau(a) = \int_0^1 \mu_a(t) dt.$$

En particular, usando el Lema 5.2.1 tenemos

$$\begin{aligned} \int_{\gamma}^{\delta} h(t) d\nu(t) &= \int_0^1 h(t) \cdot \chi_{[\gamma, \delta]}(t) d\nu(t) = \int_0^1 \mu_{\tilde{h}}(t) dt \\ &= \int_0^{\nu([\gamma, \delta])} \mu_h(t + r) dt \end{aligned}$$

Si hacemos el cambio de variable $u = t + r$ ($du = dt$) entonces tenemos

$$\int_{\gamma}^{\delta} h(t) d\nu(t) = \int_r^{r+\nu([\gamma, \delta])} \mu_h(t) dt$$

donde $r = \nu([\delta, \beta])$ con lo que $r + \nu([\gamma, \delta]) = \nu([\gamma, \beta])$. De esta forma hemos probado la siguiente

Lema 5.2.3. Sea $P = \{x_1 < \dots < x_{2^n+1}\} \subseteq [\alpha, \beta]$ y sea $\{I_i\}_{i=1}^{2^n}$ la partición de $[\alpha, \beta]$ inducida por P , con $\nu(I_i) = \frac{1}{2^n}$ para $1 \leq i \leq 2^n$. Entonces

$$\int_{I_i} h(t) d\nu(t) = \int_{\frac{2^{n-i}}{2^n}}^{\frac{2^n-i+1}{2^n}} \mu_h(t) dt \quad \text{para } 1 \leq i \leq 2^n. \quad (5.11)$$

Teorema 5.2.4. Sea ν una medida de probabilidad de soporte compacto $I = [\alpha, \beta]$ y difusa, $P = \{x_1 < \dots < x_{2^n+1}\} \subseteq [\alpha, \beta]$ y sea $\{I_i\}_{i=1}^{2^n}$ la partición de $[\alpha, \beta]$ inducida por P , con $\nu(I_i) = \frac{1}{2^n}$ para $1 \leq i \leq 2^n$. Sean $g, h : [\alpha, \beta] \rightarrow \mathbb{R}_{\geq 0}$ funciones crecientes y continuas a izquierda tales que $h \prec g$. Sea $\mathbf{h} = \mathbf{h}(n) \in \mathbb{R}^{2^n}$ (resp. $\mathbf{g} = \mathbf{g}(n) \in \mathbb{R}^{2^n}$) dado por

$$\mathbf{h}_i = 2^n \int_{I_i} h(t) d\nu(t), \quad 1 \leq i \leq 2^n$$

entonces $\mathbf{h} = \mathbf{h}^\uparrow$ (resp. $\mathbf{g} = \mathbf{g}^\uparrow$) y $\mathbf{h} \prec \mathbf{g}$.

Demostración. Notemos que si $1 \leq i < j \leq 2^n$ entonces, por la ecuación (5.11), tenemos

$$\frac{\mathbf{h}_i}{2^n} = \int_{\frac{2^{n-i}}{2^n}}^{\frac{2^n-i+1}{2^n}} \mu_h(t) dt \leq \int_{\frac{2^{n-i}}{2^n}}^{\frac{2^n-i+1}{2^n}} \mu_h\left(t + \frac{(i-j)}{2^n}\right) dt = \int_{\frac{2^{n-j}}{2^n}}^{\frac{2^n-j+1}{2^n}} \mu_h(u) du = \frac{\mathbf{h}_j}{2^n}$$

donde hemos usado que $\mu_h(t) \leq \mu_h(t + \frac{(i-j)}{2^n})$, el cambio de variable $u = t + \frac{(i-j)}{2^n}$ y el hecho de que la medida de Lebesgue es invariante por traslaciones. Entonces $\mathbf{h} = \mathbf{h}^\uparrow$.

Nuevamente, por la ecuación (5.11) entonces vemos que

$$\frac{\sum_{i=1}^k \mathbf{h}_i}{2^n} = \int_{\frac{n-k}{n}}^1 \mu_h(t) dt \geq \int_{\frac{n-k}{n}}^1 \mu_g(t) dt = \frac{\sum_{i=1}^k \mathbf{g}_i}{2^n}, \quad 1 \leq i \leq 2^n - 1$$

por hipótesis, pues $h \prec g$ (la desigualdad se invierte pues estamos considerando las sumas correspondientes a los valores singulares más pequeños). Finalmente

$$\frac{\sum_{i=1}^{2^n} \mathbf{h}_i}{2^n} = \int_{\alpha}^{\beta} h(t) d\nu(t) = \int_{\alpha}^{\beta} g(t) d\nu(t) = \frac{\sum_{i=1}^{2^n} \mathbf{g}_i}{2^n}$$

□

Aproximaciones por funciones simples

Vamos a considerar algunos resultados de aproximación de funciones por funciones escalonadas o simples. Incidentalmente, obtenemos una descomposición de tipo Calderón-Zygmund para funciones positivas e integrables en un espacio de medida $([\alpha, \beta], \nu)$, donde ν es de probabilidad, regular, difusa y de soporte incluido en $[\alpha, \beta]$. Comenzamos recordando algunos resultados de aproximación en espacios de medida abstractos.

Sean (X, μ) y (Y, ν) espacios de medida y sea T un operador de $L^p(X, \mu)$ en el espacio de funciones medibles de Y en $\mathbb{C}$. Decimos que T es débil (p, q) ($q < \infty$) si existe una constante $C < \infty$ tal que, para $\lambda > 0$ entonces

$$\nu(\{y \in Y : |T(f)(y)| > \lambda\}) \leq \left(\frac{C \|f\|_p}{\lambda} \right)^q.$$

Si $\{T_t\}_{t \in R}$ es una familia de operadores lineales en $L^p(X, \mu)$ entonces el operador maximal asociado a esta familia está definido por

$$T^*(f)(x) = \sup_{t \in R} |T_t(f)(x)|.$$

Una prueba del siguiente teorema puede hallarse en [31].

Teorema 5.2.5. Sea $\{T_t\}_{t \in R}$ una familia de operadores lineales en $L^p(X, \mu)$ tales que el operador maximal es débil (p, q) entonces el conjunto

$$\{f \in L^p(X, \mu) : \lim_{t \rightarrow t_0} T_t(f)(x) = f(x), \nu - \text{a.e.}\}$$

es cerrado en $L^p(X, \mu)$.

En el caso de una medida de probabilidad regular ν con $\text{sop}(\nu) \subseteq [\alpha, \beta]$ y tal que no tiene átomos, definimos inductivamente una sucesión de particiones de $[\alpha, \beta]$ como sigue: para $n = 0$ consideramos $P_0 = \{\alpha, \beta\}$ y para $n \in \mathbb{N}$ sea $P_n \subseteq P_{n+1} = \{x_i\}_{i=0}^{2^{n+1}}$ tal que, si $\{I_i^{(n+1)}\}_{i=1}^{2^{n+1}}$ es la partición de $[\alpha, \beta]$ en subintervalos inducida por P_{n+1} entonces

$$\nu(I_i^{(n+1)}) = 1/2^{n+1} \quad \text{para } 1 \leq i \leq 2^{n+1}.$$

Notemos que siempre podemos encontrar una tal partición, pues la función de acumulación $g(t) = \nu([\alpha, t])$ es continua; pero no es necesariamente única (esto se debe al hecho de que la función de acumulación puede no ser estrictamente creciente) Decimos que $\{I_i^{(n)}\}_{i=1}^{2^n}$, $n \in \mathbb{N}$ es una partición ν -diádica de $[\alpha, \beta]$.

Observación 5.2.6. Sea ν una medida difusa de soporte compacto en $\mathbb{R}^n$. Entonces, puede no haber una partición del soporte en cubos como en el caso de una medida soportada en la recta: esto se debe al hecho de que la función de acumulación

$$g(x_1, \dots, x_n) = \nu(\{(a_1, \dots, a_n) \in \mathbb{R}^n : a_i \leq x_i, 1 \leq i \leq n\})$$

puede no ser ni siquiera separadamente continua en general.

Definición 5.2.7. Sea $\{I_i^{(n)}\}_{i=1}^{2^n}$, $n \in \mathbb{N}$ una partición ν -diádica de $[\alpha, \beta]$. Para cada $n \in \mathbb{N}$ y cada $f \in L^1(\nu)$ definimos la familia de aproximaciones discretas

$$E_n(f)(x) = \sum_{i=1}^{2^n} \left(2^n \int_{I_i^{(n)}} f \, d\nu \right) \chi_{I_i^{(n)}}(x). \quad (5.12)$$

Observación 5.2.8. Sean E_n las aproximaciones discretas inducidas por partición ν -diádica $\{I_i^{(n)}\}_{i=1}^{2^n}$, $n \in \mathbb{N}$ de $[\alpha, \beta]$. Entonces, para cada $n \in \mathbb{N}$, E_n es un operador lineal que reduce las normas $\|\cdot\|_1$ y $\|\cdot\|_\infty$.

Teorema 5.2.9. Sea $\{I_i^{(n)}\}_{i=1}^{2^n}$, $n \in \mathbb{N}$ una partición ν -diádica de $[\alpha, \beta]$ y sea E_n , $n \in \mathbb{N}$ la familia asociada de aproximaciones discretas. Entonces

- i) El operador maximal E^* es débil $(1, 1)$.
- ii) Si $f \in L^1(\nu)$ entonces, $\lim_{n \rightarrow \infty} E_n(f)(x) = f(x)$ ν -a.e.

iii) Si $g \in L^\infty(\nu)$ entonces $\lim_{n \rightarrow \infty} \|g - E_n(g)\|_1 = 0$.

Demostración. i) Consideramos una función no negativa f ; el caso general se deduce de este. Sea $n \in \mathbb{N}$, $\lambda > 0$, y sea

$$\Omega_n = \{x \in [\alpha, \beta] : E_n(f)(x) > \lambda \text{ y } E_k(f)(x) \leq \lambda \text{ si } k < n\}.$$

Entonces es claro que Ω_n es una unión de intervalos de $\{I_i^{(n)}\}_{i=1}^{2^n}$, $\Omega_n \cap \Omega_m = \emptyset$ si $n \neq m$ y

$$\{x \in [\alpha, \beta] : E^*(f)(x) > \lambda\} = \cup_{n \in \mathbb{N}} \Omega_n.$$

Entonces tenemos

$$\begin{aligned} \nu(\{x \in [\alpha, \beta] : E^*(f)(x) > \lambda\}) &= \sum_{n \in \mathbb{N}} \nu(\Omega_n) \\ &\leq \sum_{n \in \mathbb{N}} \frac{1}{\lambda} \int_{\Omega_n} E_n(f) d\nu \\ &\leq \frac{1}{\lambda} \|f\|_1. \end{aligned}$$

ii) Observemos que de la regularidad de la medida ν se deduce que la clase de funciones continuas en $[\alpha, \beta]$ es densa en $L^1([\alpha, \beta], \nu)$. Entonces, por la primera parte del teorema y por el Teorema 5.2.5 vemos que para probar $\lim_{n \rightarrow \infty} E_n(g)(x) = g(x)$ ν -a.e. para toda $g \in L^1([\alpha, \beta], \nu)$ basta probarlo para g continua.

Para cada $x \in [\alpha, \beta]$ y $n \in \mathbb{N}$ sea $I_n(x)$ el único subintervalo en $\{I_i^{(n)}\}_{i=1}^{2^n}$ tal que $x \in I_n(x)$ y sea $a_n(x)$ la longitud de $I_n(x)$. Entonces, $a_n(x) \geq a_{n+1}(x)$ de forma que existe $l(x) := \lim_{n \rightarrow \infty} a_n(x)$. Si g es una función continua en $[\alpha, \beta]$, afirmamos que

$$\lim_{n \rightarrow \infty} E_n(g)(x) = g(x) \text{ siempre que } l(x) = 0.$$

De hecho, sea $x \in [\alpha, \beta]$ tal que $l(x) = 0$ y sea $\varepsilon > 0$. Entonces, por continuidad uniforme existe $\delta > 0$ tal que $|g(x) - g(y)| \leq \varepsilon$ siempre que $|x - y| \leq \delta$. Sea $n \in \mathbb{N}$ tal que $a_n(x) \leq \delta$ y notemos que

$$\begin{aligned} |E_n(g)(x) - g(x)| &= |2^n \int_{I_n(x)} g(t) d\nu(t) - g(x)| \\ &\leq 2^n \int_{I_n(x)} |g(t) - g(x)| d\nu(t) \leq \varepsilon \end{aligned}$$

pues $|t - x| \leq \delta$ para cada $t \in I_n(x)$, que demuestra nuestra afirmación. De lo anterior, basta mostrar que

$$\nu(\{x : l(x) > 0\}) = 0. \quad (5.13)$$

A tal fin, sea $\varepsilon > 0$ y sea $n \in \mathbb{N}$ tal que $(\beta - \alpha)/2^n \leq \varepsilon$. Entonces,

$$\{x : l(x) > \varepsilon\} \subseteq \{x : a_{2n}(x) > \frac{(\beta - \alpha)}{2^n}\} = \bigcup_{i \in K} I_i^{(2n)} \subseteq [\alpha, \beta] \quad (5.14)$$

pues $a_{2n}(x) \geq l(x)$, donde $K \subseteq \{1, \dots, 2^{2n}\}$. Sea $x_i \in I_i^{(2n)}$ para cada $i \in K$ y notemos que

$$(\beta - \alpha) \geq \sum_{i \in K} a_{2n}(x_i) \geq \sum_{i \in K} \frac{(\beta - \alpha)}{2^n} = |K| \frac{(\beta - \alpha)}{2^n}$$

donde $|K|$ denota el cardinal del conjunto K . Este último hecho implica que $|K| \leq 2^n$ y en este caso tenemos que

$$\nu(\{x : a_{2n}(x) > \frac{(\beta - \alpha)}{2^n}\}) = \sum_{i \in K} \nu(I_i^{(2n)}) = \frac{|K|}{2^{2n}} \leq \frac{1}{2^n} \quad (5.15)$$

De las fórmulas (5.14) y (5.15) vemos que $\nu(\{x : l(x) > \varepsilon\}) = 0$ y como $\varepsilon > 0$ es arbitrario, deducimos que se verifica la ecuación (5.13).

iii). Notemos que por la Observación 5.2.8 se tiene que $\|E_n(g)\|_\infty \leq \|g\|_\infty$ y por el ítem anterior $\lim_{n \rightarrow \infty} E_n(g)(x) = g(x)$ ν -a.e. Como ν es una medida finita, entonces el corolario es una consecuencia del teorema de la convergencia dominada de Lebesgue. $\square$

La siguiente consecuencia del Teorema 5.2.9 es una descomposición de tipo Calderón-Zygmund para una función no negativa e integrable f .

Corolario 5.2.10. Sea $f \in L^1(\nu)$ una función no negativa. Entonces, para cada $\lambda > 0$ existe una sucesión ν -diádica de conjuntos $\{I_j\}_{j \in J}$ tal que

1. $f(x) \leq \lambda \nu$ -a.e. para $x \notin \bigcup_{j \in J} I_j$.
2. $\nu(\bigcup_{j \in J} I_j) \leq \frac{\|f\|_1}{\lambda}$.
3. $\lambda < \frac{1}{\nu(I_j)} \int_{I_j} f \, d\nu \leq 2\lambda$.

5.2.2. Un teorema de tipo Schur-Horn en factores II_1

Corolario 5.2.11. Sea $b \in \mathcal{M}^+$ y sea $c \in \mathcal{U}_{\mathcal{M}}(b)$. Entonces $E_{\mathcal{A}}(c) \prec b$, con lo que tenemos

$$\overline{E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|_1} \subseteq \{a \in \mathcal{A} : a \prec b\} := \Omega_{\mathcal{A}}(b).$$

Demostración. Por el Corolario 4.2.10 tenemos $E_{\mathcal{A}}(b) \prec b$. Si $c \in \mathcal{U}_{\mathcal{M}}(b)$ entonces $\mu_c = \mu_b$ con lo que deducimos que $E_{\mathcal{A}}(c) \prec b$. Finalmente, el corolario se deduce de la dependencia continua de los valores singulares con respecto a la norma $\|\cdot\|_1$ (ver la desigualdad (2.26)). $\square$

Teorema 5.2.12. Sea $b \in \mathcal{M}^+$ y $a \in \mathcal{A}$ tal que $a \prec b$. Entonces $a \in \overline{E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|_1}$, con lo cual tenemos

$$\overline{E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|_1} = \Omega_{\mathcal{A}}(b).$$

Demostración. Sea $a' \in \mathcal{A}$ como en el Teorema 3.1.6, sea $\sigma(a') \subseteq [\alpha, \beta]$ y sea $\nu(\Delta) = \tau(P^{a'}(\Delta))$ la medida escalar asociada a a' i.e. $\nu(\Delta) = \tau(P^{a'}(\Delta))$ para $\Delta \subseteq \mathbb{R}$ medible Borel. Sea $g := h_b$ la función creciente continua a izquierda como en el ítem 2 en el Teorema 3.1.6, $\tilde{b} = g(a')$ y notemos que

$$b \in \overline{\mathcal{U}_{\mathcal{M}}(\tilde{b})}^{\|\cdot\|_1}$$

pues $\mu_b = \mu_{\tilde{b}}$. Como $E_{\mathcal{A}}$ es continua en norma y la convergencia en norma implica $\|\cdot\|_1$ -convergencia entonces

$$\overline{E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|_1} = \overline{E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(\tilde{b}))}^{\|\cdot\|_1}$$

De esta forma podemos asumir que $b = g(a') \in \mathcal{A}$. Sea $\{I_i\}_{i=1}^n$ una partición de $[\alpha, \beta]$ dada por intervalos tales que $\nu(I_i) = \frac{1}{n}$ (esto se puede hacer en este caso, pues ν es una medida difusa de forma que la función $m(t) = \nu([\alpha, t])$ es continua). Sea $h := h_a$ tal que $h(a') = a$ y notemos que $h \prec g$ en $L^\infty(\nu)$.

Sea $\epsilon > 0$ y notemos que, por el Teorema 5.2.9 existe $n \in \mathbb{N}$ tal que

$$\|b - \sum_{i=1}^n \mathbf{g}_i p_i\|_1 \leq \epsilon, \quad \|a - \sum_{i=1}^n \mathbf{h}_i p_i\|_1 \leq \epsilon \quad (5.16)$$

donde los vectores $\mathbf{h}, \mathbf{g} \in \mathbb{R}^n$ están dados por

$$\mathbf{h}_i = n \int_{I_i} h(t) d\nu(t), \quad \mathbf{g}_i = n \int_{I_i} g(t) d\nu(t), \quad \text{para } 1 \leq i \leq n.$$

Entonces, como en el Teorema 5.2.4 $\mathbf{h} = \mathbf{h}^\uparrow$, $\mathbf{g} = \mathbf{g}^\uparrow$ y $\mathbf{h} \prec \mathbf{g}$. Por el teorema de Schur-Horn para vectores, existe una matriz unitaria $U \in \mathbb{R}^{n \times n}$ tal que $\mathcal{C}(UM_{\mathbf{g}}U^*) = \mathbf{h}$ donde $\mathcal{C}$ es la compresión a la diagonal y $M_{\mathbf{g}}$ es la matriz diagonal, ambas con respecto a la base canónica de $\mathbb{R}^n$. Del hecho de que $\tau(p_i) = \frac{1}{n}$ se sigue que, para $1 \leq j < i \leq n$ existe una isometría parcial v_{ij} tal que

$$v_{ij}v_{ij}^* = p_i \quad \text{y} \quad v_{ij}^*v_{ij} = p_j.$$

Entonces, si $1 \leq j < i \leq n$ sea $v_{ji} = v_{ij}^*$; de esta forma tenemos un sistema de matrices reales es decir, una base del espacio vectorial $\mathbb{R}^{n \times n}$ visto dentro del factor $(\mathcal{M}, \tau)$. Sea

$$u = \sum_{i,j=1}^n U_{ij} v_{ij}.$$

Es inmediato verificar que $u \in \mathcal{U}_{\mathcal{M}}$ es un operador unitario del factor $\mathcal{M}$. Más aún

$$\begin{aligned} p_i(u(\sum_{j=1}^n \mathbf{g}_j p_j)u^*)p_i &= (p_i \sum_{l,k=1}^n U_{lk} v_{lk})(\sum_{j=1}^n \mathbf{g}_j p_j)u^*p_i = \\ (\sum_{k=1}^n U_{ik} v_{ik})(\sum_{j=1}^n \mathbf{g}_j p_j)(\sum_{l=1}^n \overline{U_{il}} v_{li}) &= \sum_{j=1}^n \mathbf{g}_j \sum_{k,l} U_{ik} \overline{U_{il}} v_{ik} p_j v_{li} \\ &= (\sum_{j=1}^n |U_{ij}|^2 \mathbf{g}_j) p_i = \mathbf{h}_i p_i \end{aligned}$$

De lo anterior se deduce que

$$\begin{aligned} E_{\mathcal{A}}(u(\sum_{j=1}^n \mathbf{g}_j p_j)u^*) &= E_{\mathcal{A}}(u(\sum_{j=1}^n \mathbf{g}_j p_j)u^*)(\sum_{i=1}^n p_i) \\ &= \sum_{i=1}^n p_i E_{\mathcal{A}}(u(\sum_{j=1}^n \mathbf{g}_j p_j)u^*) p_i = E_{\mathcal{A}}(p_i (u(\sum_{j=1}^n \mathbf{g}_j p_j)u^*) p_i) \\ &= E_{\mathcal{A}}(\sum_{i=1}^n \mathbf{h}_i p_i) = \sum_{i=1}^n \mathbf{h}_i p_i. \end{aligned}$$

Por la fórmula (5.16) vemos que

$$\|u(b - \sum_{i=1}^n \mathbf{g}_i p_i)u^*\|_1 \leq \epsilon.$$

Por el Corolario 4.2.11 los cálculos previos concluimos

$$\|E_{\mathcal{A}}(ubu^*) - a\|_1 \leq \|E_{\mathcal{A}}(u(b - \sum_{i=1}^n \mathbf{g}_i p_i)u^*)\| + \|a - \sum_{i=1}^n \mathbf{h}_i p_i\|_1 \leq 2\epsilon$$

□

Corolario 5.2.13. $\overline{E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|_1} = E_{\mathcal{A}}(\overline{\text{conv}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|_1})$ es un conjunto convexo.

Demostración. Como caso particular de la teoría de mayorización conjunta en factores finitos tenemos que

$$\overline{\text{conv}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|} = \overline{\text{conv}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|_1} = \overline{\text{conv}(\mathcal{U}_{\mathcal{M}}(b))}^{\sigma\text{-WOT}} = \Omega_{\mathcal{M}}(b) \quad (5.17)$$

donde $\Omega_{\mathcal{M}}(b)$ denota los operadores positivos $a \in \mathcal{M}$ tales que $a \prec b$. Por el Corolario 4.2.10 y por transitividad de la mayorización vemos que en general

$$E_{\mathcal{A}}(\overline{\text{conv}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|_1}) = \Omega_{\mathcal{A}}(b). \quad (5.18)$$

Pero entonces, el corolario es una consecuencia del Teorema 5.2.12. □

Corolario 5.2.14. $\Omega_{\mathcal{A}}(b)$ es σ -WOT compacto.

Demostración. Por las ecuaciones (5.17) y (5.18) tenemos que, en general,

$$E_{\mathcal{A}}(\overline{\text{conv}(\mathcal{U}_{\mathcal{M}}(b))}^{\sigma\text{-WOT}}) = \Omega_{\mathcal{A}}(b) \quad (5.19)$$

Pero el lado izquierdo de la ecuación (5.19) es un conjunto compacto pues $E_{\mathcal{A}}$ es σ -WOT continua (normal) y $\overline{\text{conv}(\mathcal{U}_{\mathcal{M}}(b))}^{\sigma\text{-WOT}}$ es σ -WOT compacto. □

Observación 5.2.15. En realidad, el hecho de que $\overline{E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|_1}$ sea convexo es equivalente al Teorema 5.2.12; en efecto, asumamos la convexidad de dicho conjunto; entonces, por el Corolario 4.2.11 y por la ecuación (5.18) tenemos

$$\Omega_{\mathcal{A}}(b) = E_{\mathcal{A}}(\overline{\text{conv}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|_1}) \subseteq \overline{E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|_1} \quad (5.20)$$

mientras que la otra inclusión está dada en Corolario 5.2.11. Siguiendo un argumento similar al anterior, para probar la afirmación (más fuerte)

$$\Omega_{\mathcal{A}}(b) = \overline{E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|}$$

basta con probar que $\overline{E_{\mathcal{A}}(\mathcal{U}_{\mathcal{M}}(b))}^{\|\cdot\|}$ es un conjunto convexo.

Capítulo 6

Marcos con operador de marco y normas prefijadas

Para una descripción conceptual de los resultados incluidos dentro de este capítulo, así como la relación entre éstos y otros resultados ver la sección “Contexto general del trabajo” del Capítulo 1.

Organización del capítulo. En la sección 2 probamos algunas extensiones del teorema de Schur-Horn para un operador autoadjunto en un espacio de Hilbert $\mathcal{H}$. En la sección 3 reformulamos el problema de la admisibilidad de marcos para poder aplicar la teoría de mayorización a este problema. Obtenemos así condiciones equivalentes para la admisibilidad, en el caso de dimensión finita (tanto para sucesiones $\mathbf{c}$ finitas como infinitas). En la sección 4 estudiamos el caso general, mostrando condiciones necesarias y suficientes de forma separada. En la sección 5 se presentan varios ejemplos para los casos límites en términos de las condiciones halladas previamente y mostramos que, en general, éstas no pueden debilitarse más. También estudiamos diferentes tipos de marcos en $F(S, \mathbf{c})$, en términos de sus excesos.

Notaciones. Adoptamos las notaciones del capítulo anterior.

6.1. Reformulación de la admisibilidad de marcos

Sea $\mathbb{M} = \mathbb{N}$ ó $\mathbb{M} = \{1, 2, \dots, m\} := \mathbb{I}_m$, para algún $m \in \mathbb{N}$. Recordemos que una sucesión $\{f_k\}_{k \in \mathbb{M}}$ es un *marco* para $\mathcal{H}$ (ver la sección 2.2.1 para más detalles) si existen constantes positivas $A, B > 0$ tales que

$$A\|x\|^2 \leq \sum_{k \in \mathbb{M}} |\langle x, f_k \rangle|^2 \leq B\|x\|^2, \quad \text{para cada } x \in \mathcal{H}.$$

Sea $\mathcal{F} = \{f_k\}_{k \in \mathbb{M}}$ un marco para $\mathcal{H}$. El operador

$$S : \mathcal{H} \rightarrow \mathcal{H}, \quad \text{dado por } S(x) = \sum_{k \in \mathbb{M}} \langle x, f_k \rangle f_k, \quad x \in \mathcal{H}.$$

es llamado el *operador de marco* de $\mathcal{F}$. Éste es acotado, positivo e inversible (usamos la notación $S \in \mathcal{G}l(\mathcal{H})^+$).

Un par ordenado $(S, \mathbf{c})$ formado por un operador positivo en inversible $S \in \mathcal{G}l(\mathcal{H})^+ \subseteq L(\mathcal{H})$ y una sucesión uniformemente acotada de números positivos $\mathbf{c} \in \ell^\infty(\mathbb{M})^+$, se dice *admisibile* si existe un marco $\mathcal{F} = \{f_k\}_{k \in \mathbb{M}}$ para $\mathcal{H}$ tal que

- $\mathcal{F}$ tiene operador de marco S ,
- $\|f_k\|^2 = c_k$ para cada $k \in \mathbb{M}$.

En este caso, decimos que $\mathcal{F}$ es un $(S, \mathbf{c})$ -marco. Denotamos por $F(S, \mathbf{c})$ el conjunto de todos los $(S, \mathbf{c})$ -marcos para $\mathcal{H}$. Luego, el par $(S, \mathbf{c})$ es admisibile si $F(S, \mathbf{c}) \neq \emptyset$.

Consideramos las siguiente formulaciones equivalentes de la admisibilidad, que ponen de manifiesto su relación con el teorema de Schur-Horn de la teoría de mayorización.

Proposición 6.1.1. Sea $\mathbf{c} \in \ell^\infty(\mathbb{M})^+$ y sea $S \in \mathcal{G}l(\mathcal{H})^+$. Entonces las siguientes condiciones son equivalentes:

1. El par $(S, \mathbf{c})$ es admisibile.
2. Existe una sucesión de vectores unitarios $\{y_k\}_{k \in \mathbb{M}}$ en $\mathcal{H}$ tales que

$$S = \sum_{k \in \mathbb{M}} c_k y_k \otimes y_k ,$$

donde, si $\mathbb{M} = \mathbb{N}$, la suma converge en la topología fuerte de operadores.

3. Existe una extensión $\mathcal{K} = \mathcal{H} \oplus \mathcal{H}_d$ de $\mathcal{H}$ tal que, si denotamos

$$S_1 = \begin{pmatrix} S & 0 \\ 0 & 0 \end{pmatrix} \begin{matrix} \mathcal{H} \\ \mathcal{H}_d \end{matrix} \in L(\mathcal{K})^+ , \quad \text{entonces} \quad \mathbf{c} \in \mathcal{C}[\mathcal{U}_{\mathcal{K}}(S_1)] . \quad (6.1)$$

En este caso, existe un marco $\mathcal{F} \in F(S, \mathbf{c})$ con $e(\mathcal{F}) = \dim \mathcal{H}_d$.

Demostración. $1 \leftrightarrow 2$. Sea $\mathcal{F} = \{f_n\}_{n \in \mathbb{M}}$ un marco para $\mathcal{H}$. Entonces notemos que el operador de marco S se puede escribir como

$$S(x) = \sum_{n \in \mathbb{M}} \langle x, f_n \rangle f_n = \sum_{n \in \mathbb{M}} (f_n \otimes f_n) x = \sum_{n \in \mathbb{M}} \|f_n\|^2 \left(\frac{f_n}{\|f_n\|} \otimes \frac{f_n}{\|f_n\|} \right) x$$

donde la convergencia es en la norma de $\mathcal{H}$. De esta forma, si $\mathcal{F}$ es un $(S, \mathbf{c})$ -marco se tiene una descomposición como en 2. Por otra parte, si S admite una descomposición como en 2. entonces es claro que el marco $\{\sqrt{c_n} y_n\}_{n \in \mathbb{M}}$ es un marco para $\mathcal{H}$ con operador de marco S y tal que $\|\sqrt{c_n} y_n\|^2 = c_n$ para $n \in \mathbb{M}$.

$1 \leftrightarrow 3$. Supongamos que $\mathcal{F} = \{f_k\}_{k \in \mathbb{M}} \in F(S, \mathbf{c})$. Sea $(T_0, \mathcal{K}_0, \mathcal{B}_0)$ un operador de síntesis para $\mathcal{F}$. Consideramos la descomposición polar $T_0 = U|T_0|$, donde $U : \mathcal{K}_0 \rightarrow \mathcal{H}$ es una coisometría con espacio inicial $(\ker T_0)^\perp$ y rango $\mathcal{H}$. Notemos que U^* mapea isométricamente $\mathcal{H}$ sobre $\ker T_0^\perp$. Denotemos $\mathcal{H}_d = \ker T_0$, y $\mathcal{K} = \mathcal{H} \oplus \mathcal{H}_d$. Sea $V : \mathcal{K} \rightarrow \mathcal{K}_0$ el operador unitario dado por

$$V(\xi_1, \xi_2) = U^* \xi_1 + \xi_2 , \quad \text{para} \quad (\xi_1, \xi_2) \in \mathcal{H} \oplus \mathcal{H}_d = \mathcal{K} .$$

Sea $\mathcal{B} = V^*(\mathcal{B}_0)$ una base ortonormal de $\mathcal{K}$, y $T = T_0V \in L(\mathcal{K}, \mathcal{H})$. Entonces $(T, \mathcal{K}, \mathcal{B})$ es otro operador de síntesis para $\mathcal{F}$, con $\ker T = \mathcal{H}_d$. Sea $T_1 \in L(\mathcal{K})$ dado por $T_1\xi = T\xi \oplus 0_{\mathcal{H}_d}$, $\xi \in \mathcal{K}$. Luego $T_1^*T_1 = T^*T$, $|T_1| = |T|$, y

$$T_1T_1^* = \begin{pmatrix} TT^* & 0 \\ 0 & 0 \end{pmatrix} \begin{matrix} \mathcal{H} \\ \mathcal{H}_d \end{matrix} = \begin{pmatrix} S & 0 \\ 0 & 0 \end{pmatrix} = S_1 .$$

If $T_1 = U_1|T_1| = U_1|T|$ es la descomposición polar de T_1 , entonces U_1 actúa sobre $\mathcal{H} = (\ker T_1)^\perp$ como un operador unitario y vale $V = U_1 + P_{\mathcal{H}_d} \in \mathcal{U}(\mathcal{K})$. Como $T_1 = V|T|$,

$$S_1 = T_1T_1^* = V|T|^2V^* = V(T^*T)V^* \implies V^*S_1V = T^*T .$$

Por otro lado, si $\mathcal{B} = \{e_k\}_{k \in \mathbb{N}}$ entonces $\langle T^*Te_k, e_k \rangle = \langle Te_k, Te_k \rangle = \|f_k\|^2 = c_k$, para cada $k \in \mathbb{M}$. De esta forma

$$\mathcal{C}_{\mathcal{B}}(V^*S_1V) = \mathcal{C}_{\mathcal{B}}(T^*T) = M_{\mathcal{B}, \mathbf{c}} \implies \mathbf{c} \in \mathcal{C}[\mathcal{U}_{\mathcal{K}}(S_1)] .$$

Recíprocamente, supongamos que existe una extensión $\mathcal{K} = \mathcal{H} \oplus \mathcal{H}_d$ de $\mathcal{H}$ y $V \in \mathcal{U}(\mathcal{K})$ tal que $M_{\mathcal{B}, \mathbf{c}} = \mathcal{C}_{\mathcal{B}}(V^*S_1V)$, para alguna base ortonormal $\mathcal{B} = \{e_k\}_{k \in \mathbb{N}}$ de $\mathcal{K}$. Sea $T = S_1^{1/2}V$. Como S es inversible, entonces $R(T) = \mathcal{H}$ y $\dim \ker T = \dim \mathcal{H}_d$. Con lo cual $\mathcal{F} = \{Te_k\}_{k \in \mathbb{M}}$ es un marco para $\mathcal{H}$, con operador de marco $TT^*|_{\mathcal{H}} = S_1|_{\mathcal{H}} = S$. Ya que $T^*T = V^*S_1V$ y $\mathcal{C}_{\mathcal{B}}(V^*S_1V) = M_{\mathcal{B}, \mathbf{c}}$, entonces $\|Te_k\|^2 = \langle T^*Te_k, e_k \rangle = c_k$, para cada $k \in \mathbb{M}$ y $\mathcal{F} \in F(S, \mathbf{c})$ con $e(\mathcal{F}) = \dim \mathcal{H}_d$. $\square$

6.2. El caso finito dimensional

En esta sección supondremos que $\mathcal{H}$ es de dimensión finita. Consideramos separadamente los casos de marcos de una cantidad finita e infinita de elementos.

Supongamos que $S \in \mathcal{M}_n(\mathbb{C})^+$ y $|\mathbb{M}| = m < \infty$. En este caso, el teorema de Schur-Horn 2.1.9 brinda una caracterización completa de la admisibilidad para el par $(S, \mathbf{c})$.

Teorema 6.2.1. Sea $\mathbf{c} \in \mathbb{R}_{\geq 0}^m$ y sea $S \in \mathcal{G}l_n(\mathbb{C})^+$, con autovalores $b_1 \geq b_2 \geq \dots \geq b_n > 0$. Entonces, el par $(S, \mathbf{c})$ es admisible si y solo si

$$\sum_{i=1}^k b_i \geq \sum_{i=1}^k c_i \quad \text{para } 1 \leq k \leq n-1, \quad \text{y} \quad \sum_{i=1}^n b_i = \sum_{i=1}^m c_i .$$

Es decir, si $\mathbf{c} \prec (b_1, \dots, b_n, 0, \dots, 0) \in \mathbb{R}^m$. $\square$

Este resultado fue obtenido en [21] y [50], desde el análisis matricial y la teoría de operadores. Incluso, las pruebas que se desarrollan en estos trabajos pueden ser adaptadas para probar el teorema de Schur-Horn de formas más simples que las de la literatura.

En lo que sigue, consideramos la admisibilidad para sucesiones infinitas en espacios de Hilbert finito dimensionales.

Teorema 6.2.2. Sea $\mathbf{c} \in \ell^\infty(\mathbb{N})^+$ y $S \in \mathcal{G}l_n(\mathbb{C})^+$, con autovalores $b_1 \geq b_2 \geq \dots \geq b_n > 0$. Las siguientes condiciones son equivalentes:

1. el par $(S, \mathbf{c})$ es admisible.
2. $\sum_{i=1}^k b_i \geq U_k(\mathbf{c})$, para cada $1 \leq k \leq n-1$, y $\sum_{i=1}^n b_i = \sum_{i \in \mathbb{N}} c_i$.

Demostración. Sea $\mathbf{b} = (b_1, \dots, b_n, 0, \dots, 0, \dots) \in \ell^\infty(\mathbb{N})^+$.

$2 \rightarrow 1$: Sea $\mathcal{H}$ espacio de Hilbert separable infinito dimensional, y consideremos

$$S_1 = \begin{pmatrix} S & 0 \\ 0 & 0 \end{pmatrix} \in L(\mathbb{C}^n \oplus \mathcal{H}) .$$

Luego, existe una base ortonormal $\mathcal{B} = \{e_k\}_{k \in \mathbb{N}}$ de $\mathcal{K} = \mathbb{C}^n \oplus \mathcal{H}$ tal que $S_1 = M_{\mathcal{B}, \mathbf{b}}$. Por la Proposición 5.1.3,

$$U_k(S_1) = \sum_{i=1}^k b_i \geq U_k(\mathbf{c}), \quad \text{para cada } k \in \mathbb{N}.$$

Notemos además que $L_k(S_1) = 0 \leq L_k(\mathbf{c})$ para cada $k \in \mathbb{N}$ y $\sum_{i=1}^n b_i = \sum_{i \in \mathbb{N}} c_i$. Entonces, por el Teorema 5.1.11, existe una sucesión $\{V_m\}_{m \in \mathbb{N}}$ en $\mathcal{U}(\mathcal{K})$ tal que

$$\mathcal{C}_{\mathcal{B}}(V_m^* S_1 V_m) \xrightarrow[m \rightarrow \infty]{\|\cdot\|_1} M_{\mathbf{c}} ,$$

donde $\|A\|_1 = \text{tr}|A|$. Por la Proposition 6.1.1, existe una sucesión acotada de epimorfismos $T_m : \mathcal{K} \rightarrow \mathbb{C}^n$ tal que $T_m T_m^* = S$ para todo $m \in \mathbb{N}$, y $(\|T_m(e_i)\|^2)_{i \in \mathbb{N}} \xrightarrow[m \rightarrow \infty]{\ell^1(\mathbb{N})} \mathbf{c}$. Entonces, por un argumento diagonal estándar, podemos asegurar la existencia de una subsucesión, que por simplicidad notamos $\{T_m\}_{m \in \mathbb{N}}$, tal que

$$T_m(e_i) \xrightarrow[m \rightarrow \infty]{} x_i \in \mathbb{C}^n, \quad \text{con } \|x_i\| = \sqrt{c_i} \quad \text{para cada } i \in \mathbb{N}.$$

Sea $T_0 : \text{span}\{\mathcal{B}\} \rightarrow \mathbb{C}^n$ el único operador (densamente definido) tal que $T_0(e_i) = x_i$ para cada $i \in \mathbb{N}$. Notemos que T_0 es acotado pues, si $x = \sum_{i=1}^r \alpha_i e_i$ y $C = \sum_{i \in \mathbb{N}} c_i = \text{tr } S$, entonces

$$\begin{aligned} \|T_0(x)\| &= \left\| \sum_{i=1}^r \alpha_i x_i \right\| \leq \sum_{i=1}^r |\alpha_i| \|x_i\| \\ &\leq \left(\sum_{i=1}^r c_i \right)^{1/2} \left(\sum_{i=1}^r |\alpha_i|^2 \right)^{1/2} \leq C^{1/2} \|x\| . \end{aligned}$$

Denotamos por $T \in L(\mathcal{K}, \mathbb{C}^n)$ la extensión acotada de T_0 a $\mathcal{K}$.

Notemos que $\|T_m - T\| \xrightarrow[m \rightarrow \infty]{} 0$. En efecto, sea $\varepsilon > 0$ y $i_0 \in \mathbb{N}$ tal que $\sum_{i=i_0}^\infty c_i < \varepsilon$. Entonces, existe $m_1 \in \mathbb{N}$ tal que

$$\sum_{i=i_0}^\infty \|T_m(e_i)\|^2 \leq \varepsilon, \quad \text{para cada } m \geq m_1 . \quad (6.2)$$

Esto es una consecuencia del hecho de que $(\|T_m(e_i)\|^2)_{i=i_0}^\infty \xrightarrow[m \rightarrow \infty]{\ell^1(\mathbb{N})} (c_i)_{i=i_0}^\infty$. Por otro lado, existe $m_2 \geq m_1$ tal que

$$\sum_{i=1}^{i_0-1} \|T_m(e_i) - T(e_i)\| \leq \varepsilon, \quad \text{para cada } m \geq m_2 . \quad (6.3)$$

Sea $m \geq m_2$ y $x = \sum_{i=1}^r \alpha_i e_i$. De las ecuaciones (6.2) y (6.3) se deduce que

$$\begin{aligned} \|(T_m - T)(x)\|^2 &\leq \left(\sum_{i=1}^r |\alpha_i|^2 \right) \left(\sum_{i=1}^r \|(T_m - T)(e_i)\|^2 \right) \\ &\leq \|x\|^2 \left(\sum_{i=1}^{i_0-1} \|(T_m - T)(e_i)\|^2 + 2 \sum_{i=i_0}^{\infty} (\|T_m(e_i)\|^2 + \|T(e_i)\|^2) \right) \\ &\leq 5\varepsilon \|x\|^2, \end{aligned}$$

Entonces $TT^* = \lim_{m \rightarrow \infty} T_m T_m^* = S$. Así hemos probado que el marco $\mathcal{F} = \{x_i\}_{i \in \mathbb{N}} \in F(S, \mathbf{c})$.

1 $\rightarrow$ 2: Se sigue del Teorema 5.1.9, aplicado a S_1 y $\mathbf{c}$, y de la Proposición 6.1.1. □

Observación 6.2.3. El Teorema 6.2.2 puede ser reformulado en términos de operadores de rango finito y sucesiones en $\ell^1(\mathbb{N})$ de la siguiente forma: sea $\mathcal{K}$ un espacio de Hilbert separable e infinito dimensional. Sea $S_1 \in L(\mathcal{K})^+$ tal que $\dim R(S_1) < \infty$. Entonces $\mathcal{C}[\mathcal{U}_{\mathcal{K}}(S_1)]$ es cerrado, como subconjunto de $\ell^1(\mathbb{N})$.

De hecho, supongamos que $S_1 \neq 0$ (el caso $S_1 = 0$ es trivial). Entonces, existe

$$\mathbf{b} = (b_1, \dots, b_m, 0, \dots, 0, \dots) \in \ell^1(\mathbb{N})^+$$

con $b_m > 0$, y una base ortonormal $\mathcal{B} = \{e_n\}_{n \in \mathbb{N}}$ de $\mathcal{K}$ tal que $S_1 = M_{\mathcal{B}, \mathbf{b}}$. Sea $\mathbf{c} \in \ell^1(\mathbb{N})^+$. Por el Teorema 5.1.11, la condición 2 del Teorema 6.2.2 significa que $\mathbf{c} \in \text{cl}_{\|\cdot\|_1}(\mathcal{C}[\mathcal{U}_{\mathcal{K}}(S_1)])$. Pero, por la Proposición 6.1.1, la condición 1 del Teorema 6.2.2 significa que $\mathbf{c} \in \mathcal{C}[\mathcal{U}_{\mathcal{K}}(S_1)]$.

Observemos que, aunque

$$\text{cl}_{\|\cdot\|_1}(\text{conv}(\Pi \cdot \mathbf{b})) = \text{cl}_{\|\cdot\|_1}(\mathcal{C}[\mathcal{U}_{\mathcal{K}}(S_1)]) = \mathcal{C}[\mathcal{U}_{\mathcal{K}}(S_1)],$$

como se muestra en la Observación 5.1.12, no vale que $\text{conv}(\Pi \cdot \mathbf{b})$ sea cerrado como subconjunto de $\ell^1(\mathbb{N})^+$. Por ejemplo, si $\mathbf{b} = (1, 0, 0, \dots)$ entonces, por el Teorema 5.1.11,

$$\mathbf{c} = \left(\frac{1}{2^n} \right)_{n \in \mathbb{N}} \in \text{cl}_{\|\cdot\|_1}(\mathcal{C}[\mathcal{U}_{\mathcal{K}}(e_1 \otimes e_1)]) = \text{cl}_{\|\cdot\|_1}(\text{conv}(\Pi \cdot \mathbf{b})) .$$

Pero $\mathbf{c} \notin \text{conv}(\Pi \cdot \mathbf{b})$, pues cada sucesión en $\text{conv}(\Pi \cdot \mathbf{b})$ tiene una cantidad finita de entradas no nulas. En este caso, $\mathbf{c} = \mathcal{C}_{\mathcal{B}}(x \otimes x) \in \mathcal{C}[\mathcal{U}_{\mathcal{K}}(e_1 \otimes e_1)]$, donde $x = \sum_{n \in \mathbb{N}} 2^{-\frac{n}{2}} e_n$. □

6.3. La admisibilidad en el caso de dimensión infinita

A lo largo de esta sección $\mathcal{H}$ denota un espacio de Hilbert separable infinito dimensional. Nuestro primer resultado brinda condiciones necesarias para la admisibilidad:

Teorema 6.3.1. Sea $S \in \mathcal{G}l(\mathcal{H})^+$ y $\mathbf{c} \in \ell^\infty(\mathbb{N})^+$. Si el par $(S, \mathbf{c})$ es admisible entonces,

$$\sum_{i \in \mathbb{N}} c_i = \infty \quad \text{y} \quad U_k(S) \geq U_k(\mathbf{c}), \quad \text{para cada } k \in \mathbb{N}.$$

En particular, se tiene la condición necesaria

$$\limsup \mathbf{c} \leq \|S\|_e.$$

Demostración. Supongamos que existe un marco $\mathcal{F} \in F(S, \mathbf{c})$. Entonces, por la Proposición 6.1.1, existe una extensión $\mathcal{K} = \mathcal{H} \oplus \mathcal{H}_d$ de $\mathcal{H}$ tal que, si denotamos

$$S_1 = \begin{pmatrix} S & 0 \\ 0 & 0 \end{pmatrix} \begin{matrix} \mathcal{H} \\ \mathcal{H}_d \end{matrix} \in L(\mathcal{K})^+, \quad \text{entonces} \quad \mathbf{c} \in \mathcal{C}[\mathcal{U}_{\mathcal{K}}(S_1)] .$$

Luego, $\sum_{i \in \mathbb{N}} c_i = \text{tr } M_{\mathbf{c}} = \text{tr } S_1 = \infty$. Por la Proposición 5.1.5, $U_k(S) = U_k(S_1)$ para cada $k \in \mathbb{N}$. Entonces, podemos aplicar el Teorema 5.1.9, a partir del cual se deduce el teorema. $\square$

Observación 6.3.2. Sea $S \in \mathcal{G}l(\mathcal{H})^+$ y $\mathbf{c} \in \ell^\infty(\mathbb{N})^+$. Entonces, por el Teorema 5.1.9 y la Proposición 6.1.1, las siguientes condiciones son equivalentes :

1. $U_k(S) \geq U_k(\mathbf{c})$ para cada $k \in \mathbb{N}$.
2. Existe una sucesión $\mathcal{F}_k = \{f_{ik}\}_{i \in \mathbb{N}}$, $k \in \mathbb{N}$ de marcos para $\mathcal{H}$, tal que S es el operador de marco de cada $\mathcal{F}_k$ y $\|f_{ik}\| \xrightarrow[k \rightarrow \infty]{} \sqrt{c_i}$ uniformemente para $i \in \mathbb{N}$.

Notemos que las desigualdades que involucran las funciones L_k , $k \in \mathbb{N}$, siempre se cumplen si consideramos una extensión suficientemente grande $\mathcal{H} \oplus \mathcal{H}_d$ de $\mathcal{H}$. En este caso, $\limsup \mathbf{c} \leq \|S\|_e$. $\square$

Notemos que las condiciones del Teorema 6.3.1 no son suficientes para asegurar que el par $(S, \mathbf{c})$ es admisible, como muestra el Ejemplo 6.4.1. Es decir, en general no podemos remover las barra de clausura en las fórmulas del Teorema 5.1.9, como ya ha sido mencionado en [56], para el caso diagonalizable.

En [50] se probó el siguiente resultado que da condiciones suficientes que aseguran que el par $(S, \mathbf{c})$ es admisible:

Teorema 6.3.3 (Larson-Korleson). Sea $S \in \mathcal{G}l(\mathcal{H})^+$ y $\mathbf{c} \in \ell^\infty(\mathbb{N})^+$. Supongamos que $\sum_{i \in \mathbb{N}} c_i = \infty$ y $\|\mathbf{c}\|_\infty < \|S\|_e$. Entonces el par $(S, \mathbf{c})$ es admisible. $\square$

El siguiente resultado, que generaliza el Teorema 6.3.3, refuerza las condiciones necesarias para la admisibilidad dadas por el Teorema 6.3.1 para obtener condiciones suficientes. Recordemos la notación $P_2(S) = E[\|S\|_e, \|S\|]$, donde E es la medida espectral de $S \in L(\mathcal{H})^+$.

Teorema 6.3.4. Sea $S \in \mathcal{G}l(\mathcal{H})^+$ y $\mathbf{c} \in \ell^\infty(\mathbb{N})^+$, tal que $\sum_{i \in \mathbb{N}} c_i = \infty$. Supongamos que se satisface alguna de las siguientes condiciones:

1. a) $\text{tr } P_2(S) = \infty$,

- b) $U_k(S) \geq U_k(\mathbf{c})$ para cada $k \in \mathbb{N}$, y
 - c) $\|S\|_e > \limsup(\mathbf{c})$.
- 2.
- a) $\text{tr } P_2(S) = r \in \mathbb{N}$,
 - b) $U_k(S) \geq U_k(\mathbf{c})$ para $1 \leq k \leq r$,
 - c) $U_k(S) > U_k(\mathbf{c})$, para $k > r$, y
 - d) $\|S\|_e > \limsup(\mathbf{c})$.

Entonces, el par $(S, \mathbf{c})$ es admisible.

Demostración. Por la Proposición 6.1.1, para probar que existe un marco $\mathcal{F} \in F(S, \mathbf{c})$, basta mostrar que existe una sucesión de vectores unitarios $\{x_k\}_{k \in \mathbb{N}}$ tal que

$$S = \sum_{k \in \mathbb{N}} c_k x_k \otimes x_k .$$

Supongamos que vale la primera condición. De las ecuaciones (5.5) y (5.9), la condición $\|S\|_e > \limsup(\mathbf{c})$ implica que $U_m(S) - U_m(\mathbf{c}) \xrightarrow{m \rightarrow \infty} \infty$. Entonces, existe $m_0 \in \mathbb{N}$ y $\varepsilon > 0$ tal que

$$c_m \leq \|S\|_e - \varepsilon \quad \text{para } m \geq m_0, \quad \text{y} \quad \|S\|_e \leq U_{m_0}(S) - U_{m_0}(\mathbf{c}) .$$

Sea $\mu_1 \geq \mu_2 \geq \dots \geq \mu_n \geq \dots$ la sucesión de autovalores de S^+ , elegidos como en el Lema 5.1.4. Sea $\{y_n\}_{n \in \mathbb{N}}$ un sistema ortonormal tal que $S^+ y_n = \mu_n y_n$. Si $\lambda_n = \mu_n + \|S\|_e$, $n \in \mathbb{N}$ entonces $\|S\| \geq \lambda_n \geq \|S\|_e$, y $S y_n = \lambda_n y_n$, $n \in \mathbb{N}$. Por la Proposición 5.1.5, para cada $k \in \mathbb{N}$,

$$\sum_{i=1}^k \lambda_i y_i \otimes y_i \leq S, \quad \text{y} \quad U_k(S) = \sum_{i=1}^k \lambda_i .$$

sea n_0 el primer número entero tal que $\sum_{i=1}^{n_0} c_i > \sum_{i=1}^{m_0} \lambda_i$. Luego $n_0 \geq m_0 + 1$, and

$$h = \sum_{i=1}^{n_0} c_i - \sum_{i=1}^{m_0} \lambda_i \leq c_{n_0} < \|S\|_e \leq \lambda_{m_0+1} .$$

Sea $\mathbf{c}_0 = (c_1, \dots, c_{n_0})$. Como

$$\sum_{i=1}^k \lambda_i = U_k(S) \geq U_k(\mathbf{c}) \geq U_k(\mathbf{c}_0), \quad 1 \leq k \leq m_0 ,$$

entonces $\mathbf{c}_0 \prec (\lambda_1, \dots, \lambda_{m_0}, h, 0, \dots, 0) \in \mathbb{R}^{n_0}$. Denotemos por

$$S_1 = h y_{m_0+1} \otimes y_{m_0+1} + \sum_{i=1}^{m_0} \lambda_i y_i \otimes y_i \leq S ,$$

y $S_2 = S - S_1$. Entonces, el par $(S_1, \mathbf{c}_0)$, actuando en $\text{span}\{y_1, \dots, y_{m_0+1}\}$, satisface las condiciones del Teorema 6.2.1. De esta forma, existe un conjunto de vectores unitarios $\{x_1, \dots, x_{n_0}\}$ tales que

$$\sum_{i=1}^{n_0} c_i x_i \otimes x_i = S_1.$$

Observemos que $S_2 \geq 0$, $R(S_2)$ es cerrado (por teoría de Fredholm), y $\|S_2\|_e = \|S\|_e$. Podemos aplicar entonces el Teorema 6.3.3 al par $(S_2, \{c_i\}_{i>n_0})$, actuando en $R(S_2)$ y deducir la existencia de vectores unitarios x_k , para $k > n_0$, tales que

$$S_2 = \sum_{i=n_0+1}^{\infty} c_i x_i \otimes x_i.$$

Con lo cual obtenemos una descomposición de $S = \sum_{i \in \mathbb{N}} c_i x_i \otimes x_i$ en términos de operadores de rango uno.

Supongamos la condición 2. Entonces, existe $\delta > 0$ tal que

1. $U_{r+k}(S) - \delta > U_{r+k}(\mathbf{c})$, para cada $k \in \mathbb{N}$.
2. Existe $m_0 \geq 1$ tal que $c_m \leq \|S\|_e - \delta$ para $m \geq m_0$.

Sea $m_1 = \max\{m_0, r+1\}$, $\mu_1 \geq \dots \geq \mu_r$ los autovalores más grandes de S^+ , y sea $\{y_1, \dots, y_r\}$ un conjunto ortonormal de autovectores asociados a los autovalores anteriores. Denotemos por

$$\lambda_i = \mu_i + \|S\|_e, \quad 1 \leq i \leq r, \quad \text{y} \quad \lambda_i = \|S\|_e - \frac{\delta}{2m_1}, \quad r+1 \leq i \leq m_1+1.$$

Entonces, por la Proposición 5.1.5,

1. $U_k(S) = \sum_{i=1}^k \lambda_i$, para $1 \leq k \leq r$.
2. $U_k(\mathbf{c}) \leq U_k(S) - \delta \leq \sum_{i=1}^k \lambda_i$, para $r+1 \leq k \leq m_1+1$.

Como $Q = E(\|S\|_e - \delta/2m_1, \|S\|_e)$ tiene rango infinito, existe un conjunto ortonormal $\{y_{r+1}, \dots, y_{m_1+1}\} \subseteq R(Q)$. En consecuencia

$$\sum_{i=1}^{m_1+1} \lambda_i y_i \otimes y_i \leq S.$$

Seguimos un argumento similar al anterior. Sea n_0 el primer número entero tal que $\sum_{i=1}^{n_0} c_i > \sum_{i=1}^{m_1} \lambda_i$. Luego $n_0 \geq m_1+1$ y

$$h = \sum_{i=1}^{n_0} c_i - \sum_{i=1}^{m_1} \lambda_i \leq c_{n_0} \leq \|S\|_e - \delta \leq \lambda_{m_1+1}.$$

Sea $\mathbf{c}_0 = (c_1, \dots, c_{n_0})$. Del hecho de que

$$\sum_{i=1}^k \lambda_i = U_k(S) \geq U_k(\mathbf{c}) \geq U_k(\mathbf{c}_0), \quad 1 \leq k \leq r, \quad \text{y}$$

$$\sum_{i=1}^k \lambda_i \geq U_k(S) - \delta \geq U_k(\mathbf{c}) \geq U_k(\mathbf{c}_0), \quad r+1 \leq k \leq m_1,$$

deducimos que $\mathbf{c}_0 \prec (\lambda_1, \dots, \lambda_{m_1}, h, 0, \dots, 0) \in \mathbb{R}^{n_0}$. Entonces, por el Corolario 6.2.1, existe un conjunto de vectores unitarios $\{x_1, \dots, x_{n_0}\}$ tal que

$$S_1 = \sum_{i=1}^{m_1} \lambda_i y_i \otimes y_i + h y_{m_0+1} \otimes y_{m_0+1} = \sum_{i=1}^{n_0} c_i x_i \otimes x_i.$$

Además $S_1 \leq \sum_{i=1}^{m_1+1} \lambda_i y_i \otimes y_i$, con lo cual $S_2 = S - S_1 \geq 0$ y $\|S_2\|_e = \|S\|_e$. Como antes, aplicamos el Teorema 6.3.3 al par $(S_2, \{c_i\}_{i>n_0})$, actuando en $R(S_2)$, y obtenemos una descomposición

$$S_2 = \sum_{i=n_0+1}^{\infty} c_i x_i \otimes x_i.$$

De esta forma obtenemos una descomposición $S = \sum_{i \in \mathbb{N}} c_i x_i \otimes x_i$ en términos de operadores de rango 1. □

El Ejemplo 6.4.2 muestra que la condición 2 (c) del Teorema 6.3.4 no puede despreciarse en general.

Corolario 6.3.5. Sea $0 < A \in \mathbb{R}$ y $\mathbf{c} \in \ell^\infty(\mathbb{N})^+$ tal que $0 < c_i \leq A$, $i \in \mathbb{N}$. Denotemos por $J = \{i \in \mathbb{N} : c_i = A\}$. Supongamos que

$$\sum_{i \notin J} c_i = \infty \quad \text{y} \quad \limsup_{i \notin J} c_i < A \quad (\text{o equivalentemente} \quad \sup_{i \notin J} c_i < A).$$

Entonces el par $(AI, \mathbf{c})$ es admisible. □

6.4. Algunos ejemplos

En el siguiente ejemplo veremos que

$$U_k(S) > U_k(\mathbf{c}), \quad k \in \mathbb{N} \quad \text{y} \quad \|S\|_e = \limsup(\mathbf{c}) \not\equiv F(S, \mathbf{c}) \neq \emptyset.$$

Ejemplo 6.4.1. Sea $S = I \in L(\mathcal{H})$ y $a \in (0, 1)$. Sea $\mathbf{c} \in \ell^\infty(\mathbb{N})^+$ dado por $c_1 = p \in (0, 1)$ y

$$c_k = \begin{cases} a^k & \text{si } k \neq 1 \text{ es impar,} \\ 1 - a^k & \text{si } k \text{ es par.} \end{cases}$$

Entonces, $0 < c_k < 1$ para $k \in \mathbb{N}$, $\sum_k c_k = \infty = \sum_k (1 - c_k)$, y $\limsup \mathbf{c} = 1 = \|S\|_e$. Supongamos que existe un marco $\mathcal{F} = \{f_k\}_{k \in \mathbb{N}} \in F(S, \mathbf{c})$. Entonces

$$\|x\|^2 = \sum_{k \in \mathbb{N}} |\langle x, f_k \rangle|^2, \quad \text{para cada } x \in \mathcal{H}.$$

En particular tenemos que, para cada $j \in \mathbb{N}$,

$$\|f_j\|^2 = \sum_{k \in \mathbb{N}} |\langle f_j, f_k \rangle|^2 = \|f_j\|^4 + \sum_{k \neq j} |\langle f_j, f_k \rangle|^2.$$

Así, si $j \neq 1$ vale la desigualdad

$$|\langle f_1, f_j \rangle|^2 = |\langle f_j, f_1 \rangle|^2 \leq \sum_{k \neq j} |\langle f_j, f_k \rangle|^2 = \|f_j\|^2 - \|f_j\|^4 = c_j(1 - c_j).$$

De esta forma se tiene

$$\begin{aligned} p = \|f_1\|^2 &\leq \|f_1\|^4 + \sum_{j \neq 1} c_j(1 - c_j) \\ &= p^2 + \sum_{j \neq 1} a^j(1 - a^j) = p^4 + \sum_{j \neq 1} a^j - \sum_{j \neq 1} a^{2j} \\ &= p^2 + \frac{1}{1 - a} - \frac{1}{1 - a^2} = p^2 + \frac{a}{1 - a^2} \end{aligned} \quad (6.4)$$

Si tomamos $p = \frac{1}{2}$ y $a \in (0, 1)$ tal que $\frac{a}{1 - a^2} < \frac{1}{4}$ vemos que $p > p^2 + \frac{a}{1 - a^2}$, que contradice la ecuación (6.4). En este caso, $F(S, \mathbf{c}) = \emptyset$. Pero notemos que el par $(S, \mathbf{c})$ satisface todas las condiciones necesarias del Teorema 6.3.1, pues el hecho de que $c_j < 1$, $j \in \mathbb{N}$ implica $U_k(S) = k > U_k(\mathbf{c})$ para cada $k \in \mathbb{N}$. $\square$

En el segundo ejemplo veremos que en general,

$$U_k(S) \geq U_k(\mathbf{c}), \quad k \in \mathbb{N} \quad \text{y} \quad \|S\|_e > \limsup(\mathbf{c}) \not\Rightarrow F(S, \mathbf{c}) \neq \emptyset.$$

Ejemplo 6.4.2. Sea $S = M_{\mathbf{s}}$ el operador diagonal con respecto a una base ortonormal de $\mathcal{H}$ dado por $\mathbf{s} = \{1 - (i + 1)^{-1}\}_{i \in \mathbb{N}}$ y sea $(c_i)_{i \in \mathbb{N}}$ dado por $c_1 = 1$ y $c_i = 1/2$ para cada $i \geq 2$. Notemos que

1. $1 = \|S\|_e > 1/2 = \limsup(\mathbf{c})$,
2. $U_1(S) = U_1(\mathbf{c})$,
3. $U_k(S) = k > 1 + (k - 1)/2 = U_k(\mathbf{c})$ para cada $k \geq 2$.

Sin embargo, $F(S, \mathbf{c}) = \emptyset$. En efecto, supongamos por el contrario que existe $\mathcal{F} \in F(S, \mathbf{c})$. Entonces, por la Proposición 6.1.1 existe una extensión $\mathcal{K} = \mathcal{H} \oplus \mathcal{H}_d$ de $\mathcal{H}$ tal que, si

$$S_1 = \begin{pmatrix} S & 0 \\ 0 & 0 \end{pmatrix} \begin{matrix} \mathcal{H} \\ \mathcal{H}_d \end{matrix} \in L(\mathcal{K})^+, \quad \text{entonces} \quad \mathbf{c} \in \mathcal{C}[\mathcal{U}_{\mathcal{K}}(S_1)].$$

Sea $V \in \mathcal{U}(\mathcal{K})$ tal que, en la base ortonormal $\mathcal{B} = \{e_k\}_{k \in \mathbb{N}}$, $M_{\mathbf{c}} = \mathcal{C}_{\mathcal{B}}(V^* S_1 V)$. sea $x = P_{\mathcal{H}} V e_1$. Entonces se tiene $\|x\| \leq 1$ y $\langle Sx, x \rangle = \langle M_{\mathbf{c}} e_1, e_1 \rangle = c_1 = 1$, mientras que $\|S\| = 1$. Luego $Sx = x$, y 1 es un autovector de S , pero esto es absurdo. Notemos que en este caso la condición 2 (c) del Teorema 6.3.4 no vale, pues $\|S\| = \|S\|_e$, con lo cual $r = \text{tr } P_2(S) = 0$; pero $U_1(S) = 1 = U_1(\mathbf{c})$. Además, $\sum_k c_k = \infty = \sum_k (1 - c_k)$. $\square$

6.4.1. El exceso de marcos en $F(S, \mathbf{c})$

Sea $S \in \mathcal{GL}(\mathcal{H})^+$ y $\mathbf{c} = (c_i)_{i \in \mathbb{M}} \in \ell^\infty(\mathbb{M})^+$ tal que el par $(S, \mathbf{c})$ es admisible. Entonces, podemos encontrar diversos tipos de marcos $\mathcal{F} \in F(S, \mathbf{c})$. Consideremos en conjunto

$$\text{Nul}(S, \mathbf{c}) = \{ e(\mathcal{F}) : \mathcal{F} \in F(S, \mathbf{c}) \} .$$

En el siguiente ejemplo, mostramos que este conjunto puede abarcar todos los números naturales. Por otro lado, este ejemplo también muestra que pueden existir pares admisibles $(S, \mathbf{a})$, que satisfacen tan solo las condiciones necesarias del Teorema 6.3.1, es decir $U_k(S) = U_k(\mathbf{a})$, $k \in \mathbb{N}$, y $\limsup \mathbf{a} = \|S\|_e$.

Ejemplo 6.4.3. Sea $\mathcal{H}$ un espacio de Hilbert con base ortonormal $\mathcal{B} = \{x_n\}_{n \in \mathbb{N}}$. Sea

$$\mathbf{a} = \left(\frac{1}{2}, 1, \frac{1}{2}, 1, \frac{1}{2}, \dots \right) \in \ell^\infty(\mathbb{N})^+, \quad \text{y} \quad S = M_{\mathcal{B}, \mathbf{a}} \in \mathcal{GL}(\mathcal{H})^+ .$$

Entonces, el marco (base de Riesz) $\mathcal{F}_0 = \{a_n^{1/2} x_n\}_{n \in \mathbb{N}}$ tiene operador de marco S , de forma que $\mathcal{F}_0 \in F(S, \mathbf{a})$. Por otro lado, sea

$$\mathcal{F}_1 = \left\{ \frac{1}{\sqrt{2}} x_2, x_4, \frac{1}{\sqrt{2}} x_2, x_6, \frac{1}{\sqrt{2}} x_1, x_8, \frac{1}{\sqrt{2}} x_3, x_{10}, \dots \right\} .$$

Se puede ver que también vale $\mathcal{F}_1 \in F(S, \mathbf{a})$, pero $e(\mathcal{F}_1) = 1$. Análogamente,

$$\mathcal{F}_2 = \left\{ \frac{1}{\sqrt{2}} x_2, x_4, \frac{1}{\sqrt{2}} x_2, x_6, \frac{1}{\sqrt{2}} x_8, x_{10}, \frac{1}{\sqrt{2}} x_8, x_{12}, \frac{1}{\sqrt{2}} x_1, \dots \right\} \in F(S, \mathbf{a}) ,$$

con $e(\mathcal{F}_2) = 2$. De forma similar, podemos definir marcos $\mathcal{F}_k \in F(S, \mathbf{a})$ con $e(\mathcal{F}_k) = k$, para cada $k \in \mathbb{N} \cup \{\infty\}$. Notemos que

$$\mathcal{F}_\infty = \left\{ \frac{1}{\sqrt{2}} x_1, x_4, \frac{1}{\sqrt{2}} x_2, x_8, \frac{1}{\sqrt{2}} x_2, x_{12}, \frac{1}{\sqrt{2}} x_3, x_{16}, \frac{1}{\sqrt{2}} x_6, x_{20}, \frac{1}{\sqrt{2}} x_6, \dots \right\} .$$

es decir, $\mathcal{F}_\infty$ es el marco inducido por el operador $T : \ell^2(\mathbb{N}) \rightarrow \mathcal{H}$ dado por

$$T(e_n) = \begin{cases} x_{4k} & \text{si } n = 2k \\ \frac{1}{\sqrt{2}} x_{2k-1} & \text{si } n = 6k - 5 \\ \frac{1}{\sqrt{2}} x_{4k-2} & \text{si } n = 6k - 3 \\ \frac{1}{\sqrt{2}} x_{4k-2} & \text{si } n = 6k - 1 . \end{cases}$$

Así se tiene $\text{Nul}(S, \mathbf{a}) = \mathbb{N} \cup \{0, \infty\}$. □

Proposición 6.4.4. Sea $S \in \mathcal{GL}(\mathcal{H})^+$ y $\mathbf{c} \in \ell^\infty(\mathbb{N})^+$. Supongamos que el par $(S, \mathbf{c})$ es admisible y $\liminf \mathbf{c} < \min \sigma_e(S)$. Entonces $\text{Nul}(S, \mathbf{c}) = \{\infty\}$.

Demostración. Sea $\mathcal{F} = \{f_n\}_{n \in \mathbb{N}} \in F(S, \mathbf{c})$, con $e(\mathcal{F}) = d$. Por la Proposición 6.1.1 existe una extensión $\mathcal{K} = \mathcal{H} \oplus \mathcal{H}_d$ de $\mathcal{H}$ tal que, si denotamos

$$S_1 = \begin{pmatrix} S & 0 \\ 0 & 0 \end{pmatrix} \begin{matrix} \mathcal{H} \\ \mathcal{H}_d \end{matrix} \in L(\mathcal{K})^+, \quad \text{entonces} \quad \mathbf{c} \in \mathcal{C}[\mathcal{U}_{\mathcal{K}}(S_1)] .$$

Por el Teorema 5.1.9, $\min \sigma_e(S_1) \leq \liminf \mathbf{c}$. Pero, si $\dim \mathcal{H}_d = e(\mathcal{F}) < \infty$, entonces $\sigma_e(S_1) = \sigma_e(S)$. □

Observación 6.4.5. La Proposición 6.4.4 tiene la siguiente aplicación para los marcos Parseval: Sea $\mathcal{F} = \{f_n\}_{n \in \mathbb{N}}$ marco Parseval para $\mathcal{H}$ (i.e. que tiene operador de marco $S = I_{\mathcal{H}}$). Si $\liminf_{n \in \mathbb{N}} \|f_n\| < 1$, entonces $e(\mathcal{F}) = \infty$. □

Ejemplo 6.4.6. Sea $\mathcal{H}$ un espacio de Hilbert con base ortonormal $\mathcal{B} = \{x_i\}_{i \in \mathbb{N}}$. Sea

$$\mathbf{a} = (1, 2, 1, 2, \dots) , \quad S = M_{\mathcal{B}, \mathbf{a}} \in \mathcal{G}l(\mathcal{H})^+ \quad \text{y} \quad \mathbf{c} = \left(\frac{3}{2}, \frac{3}{2}, \frac{3}{2}, \dots\right).$$

Mostraremos que $\text{Nul}(S, \mathbf{c}) = \mathbb{N} \cup \{0, \infty\}$. En este caso,

$$\alpha_-(S) = 1 < \liminf \mathbf{c} = \frac{3}{2} = \limsup \mathbf{c} < 2 = \|S\|_e .$$

En efecto, tomemos la base de Riesz $\mathcal{F}_0 = \{f_n\}_{n \in \mathbb{N}}$ dada por

$$f_n = \begin{cases} \frac{x_n}{\sqrt{2}} + x_{n+1} & \text{si } n \text{ es par} \\ \frac{-x_{n-1}}{\sqrt{2}} + x_n & \text{si } n \text{ es impar} \end{cases} .$$

Es sencillo ver que $\mathcal{F}_0 \in F(S, \mathbf{c})$. Usando que

$$\left(\frac{3}{2}, \frac{3}{2}, \frac{3}{2}, \frac{3}{2}\right) \prec (2, 2, 2, 0) ,$$

entonces se puede intercalar un número arbitrario de paquetes de cuatro vectores con normas $\sqrt{\frac{3}{2}}$ asociados a los paquetes de tres entradas pares de la diagonal de S en la construcción anterior.

Cada uno de estos paquetes le agrega exceso 1 al sistema completo. De esta forma pueden encontrarse marcos $\mathcal{F}_k \in F(S, \mathbf{c})$ con $e(\mathcal{F}_k) = k$ para cada $k \in \mathbb{N} \cup \{\infty\}$. □

Observación 6.4.7. Sea $S \in \mathcal{G}l_n(\mathbb{C})^+$ y $\mathbf{c} \in \ell^\infty(\mathbb{M})^+$. Si el par $(S, \mathbf{c})$ es admisible, entonces $\text{Nul}(S, \mathbf{c}) = \{|\mathbb{M}| - n\}$. Sin embargo, si $k > n$, $\mathbf{c} = (1, \dots, 1) \in \mathbb{C}^k$ y $S = \frac{k}{n}I \in \mathcal{M}_n(\mathbb{C})$, entonces $F(S, \mathbf{c})$ es el conjunto de marcos ajustados esféricos de k elementos en $\mathbb{C}^n$. Dykema, Freeman, Korleson, Larson, Ordower y Weber [30] mostraron que, en este caso, $F(S, \mathbf{c})$ tiene una estructura geométrica importante, que contiene órbitas de elementos cualitativamente distintos. □

Capítulo 7

Mayorización de matrices

Para una descripción conceptual de los resultados incluidos dentro de este capítulo, así como la relación entre éstos y otros resultados ver la sección “Contexto general del trabajo” del Capítulo 1.

Organización del capítulo. A continuación damos una descripción de nuestro estudio acerca de las distintas mayorizaciones de matrices. La mayorización débil de matrices tiene una interpretación geométrica simple. Esto nos permite obtener un procedimiento efectivo para verificar esta propiedad. Es bien conocido el hecho de que la mayorización fuerte de matrices implica la mayorización direccional; probaremos que la mayorización direccional de matrices implica la mayorización débil de matrices y daremos ejemplos mostrando que, en general, las implicaciones recíprocas no son ciertas. Sin embargo, estudiamos condiciones adicionales bajo las cuales estas implicaciones son ciertas; este problema ha sido considerado en varios artículos, por ejemplo [62], [63], [45].

Usamos algunos hechos elementales de la teoría de convexidad para obtener nuevas caracterizaciones de las distintas mayorizaciones de matrices. En particular, obtenemos un criterio simple y efectivo para determinar si $X \succ Y$. Además, damos otras descripciones de las distintas mayorizaciones de matrices, que involucran la comparación de trazas de diferentes familias de matrices. También consideramos las relaciones de equivalencias inducidas por estos preórdenes y hallamos las matrices minimales.

Sea $M_n(\mathbb{C})$ el álgebra de matrices $n \times n$ con entradas complejas. Una *familia abeliana* es una familia ordenada de matrices autoadjuntas en $M_n(\mathbb{C})$ que conmutan entre sí. Introducimos tres mayorizaciones diferentes entre familias abelianas que llamamos *mayorizaciones conjuntas*. Varios de los resultados obtenidos para las mayorizaciones matriciales se traducen en este contexto de mayorizaciones entre familias abelianas, los cuales nos permiten deducir caracterizaciones de estas últimas.

Notaciones Denotamos por $M_{n,m} = M_{n,m}(\mathbb{R})$ (resp. $M_n = M_n(\mathbb{R})$) el espacio vectorial real de matrices $n \times m$ con entradas reales (resp. $n \times n$) y $M_{n,m}(\mathbb{C})$ (resp. $M_n(\mathbb{C})$) el espacio vectorial complejo $n \times m$ de matrices (resp. $n \times n$) con entradas complejas. $GL(n)$ denota el grupo de matrices $n \times n$ inversibles (con entradas reales) y $\mathbb{S}_n$ el grupo de permutaciones de orden n .

Los vectores en $\mathbb{R}^n$ (o $\mathbb{C}^n$) son considerados como vectores columna. Sin embargo, algunas veces los describimos como $v = (v_1, \dots, v_n) \in \mathbb{C}^n$. Los elementos de la base canónica son denotados por $e_1, \dots, e_n \in \mathbb{R}^n$. Dado $x \in \mathbb{R}^n$, $\mathbb{R}x$ denota el subespacio real generado por x y $\mathbb{C}x$ es el subespacio complejo generado por x . Además consideramos el vector $e = (1, \dots, 1) \in \mathbb{R}^n$.

Si $x = (x_1, \dots, x_n)$, $y = (y_1, \dots, y_n) \in \mathbb{C}^n$ entonces $\langle x, y \rangle$ denota su producto interno i.e., $\langle x, y \rangle = \sum_{i=1}^n x_i \bar{y}_i$. Dada $A \in M_n(\mathbb{C})$, decimos que A es positiva semidefinida si $\langle Ax, x \rangle \geq 0$ para cada $x \in \mathbb{C}^n$. La traza canónica en $M_n(\mathbb{C})$ es denotada por tr .

Si $v_1, \dots, v_k \in \mathbb{C}^n$ entonces denotamos por

$$v_1 \wedge \dots \wedge v_k \in \bigwedge^k \mathbb{C}^n$$

su producto antisimétrico. Dado $A \in M_n(\mathbb{C})$, $\bigwedge^k A \in M_t(\mathbb{C})$ denota la k -ésima potencia antisimétrica de A . Es sabido que

1. $(\bigwedge^k A)(\bigwedge^k B) = \bigwedge^k(AB)$
2. $(\bigwedge^k A)^* = \bigwedge^k(A^*)$

para $A, B \in M_n(\mathbb{C})$ (ver por ejemplo [16]).

Para $X \in M_{n,m}$, $R_i(X)$ (o más brevemente X_i) denota la i -ésima fila de X y $C_i(X)$ denota la i -ésima columna de X . Consideramos también los conjuntos de filas y columnas de X

$$R(X) = \{R_i(X) : i = 1, \dots, n\} \quad \text{y} \quad C(X) = \{C_i(X) : i = 1, \dots, m\}.$$

Dada $X \in M_{n,m}(\mathbb{C})$, $X^t \in M_{m,n}(\mathbb{C})$ denota su transpuesta, $X^* \in M_{m,n}(\mathbb{C})$ denota su adjunta y $X^\dagger \in M_{m,n}(\mathbb{C})$ es la pseudoinversa de *Moore-Penrose* de X . La dimensión del rango columna de X es notada $\text{rank}(X)$.

Dado $S \subseteq \mathbb{R}^n$, $\text{conv}(S)$ denota la cápsula convexa de S , i.e. el conjunto de combinaciones convexas de elementos de S . Vamos a usar la siguiente terminología: llamamos *politopo* a la cápsula convexa de un número finito de puntos en $\mathbb{R}^n$. Un politopo generado por puntos afinmente independientes es llamado *símples*.

7.1. Mayorización de matrices

En lo que sigue vamos a utilizar las propiedades de la mayorización de vectores desarrolladas en los preliminares, así como algunos hechos de las matrices positivas allí descritos. Adoptamos también las notaciones introducidas en los preliminares.

Dadas $X, Y \in M_{n,m}$ consideramos las siguientes nociones de mayorización de matrices:

- Y está *mayorizada fuertemente* por X , notado $X \succ_f Y$, si existe $D \in DS(n)$ tal que $DX = Y$.
- Y está *mayorizada direccionalmente* por X , notado $X \succ Y$, si para todo $v \in \mathbb{R}^m$, $Xv \succ Yv$.

Observación 7.1.1. En [54] A. W. Marshall y I. Olkin definen, dadas matrices $X, Y \in M_{n,m}$, Y está mayorizada por X si existe $D \in DS(m)$ tal que $XD = Y$. Esta noción fue llamada en [26] y [14] *mayorización multivariada*. La noción de mayorización fuerte introducida arriba corresponde a la mayorización multivariada de las matrices *transpuestas*. Por otro lado, la mayorización direccional ha sido considerada en [45], [52] y [63], por ejemplo.

□

Cuando $X, Y \in M_{n,1}$, i.e. X e Y son vectores en $\mathbb{R}^n$, las mayorizaciones fuerte y direccional de matrices coinciden con la mayorización de vectores. En este caso, el teorema de Schur-Horn (ver la sección de Preliminares ó [44]) establece que la mayorización fuerte de matrices (mayorización direccional) es equivalente a la existencia de una matriz unitaria $U \in M_n(\mathbb{C})$ tal que $(U \circ \bar{U})^t X = Y$, donde “ $\circ$ ” denota el producto de Schur de matrices. Pero en general, dado $X, Y \in M_{n,m}$, es sabido que la existencia de una matriz doble estocástica $D \in DS(n)$ tal que $DX = Y$ no implica que $(U \circ \bar{U})^t X = Y$ para alguna matriz unitaria $U \in M_n(\mathbb{C})$ (ver la página 431 de [54]).

7.1.1. Mayorización débil de matrices

Introducimos la siguiente definición de mayorización de matrices.

Definición 7.1.2. Dadas $X, Y \in M_{n,m}$ decimos que Y está *mayorizada débilmente* por X , notado $X \succ_d Y$, si existe $A \in RS(n)$ tal que $AX = Y$.

Hemos considerado matrices estocásticas por filas, pero también tiene sentido considerar matrices estocásticas por filas no cuadradas. Digamos que $A \in M_{n,m}$ es estocástica por filas si A es no negativa y todas las sumas de las entradas de cada una de sus filas es 1. De esta forma, la noción de mayorización débil puede extenderse a pares de matrices con el mismo número de columnas pero diferente número de filas como sigue: sea $X \in M_{n,m}$ y $Y \in M_{p,m}$ entonces $X \succ_d Y$ si existe una de matriz estocástica por filas $A \in M_{p,n}$ tal que $AX = Y$. Sin embargo no desarrollaremos esta noción más general.

Observación 7.1.3. Dadas $X, Y \in M_{n,m}$ consideremos las dos m -uplas de vectores $(x_i)_{i=1\dots m}$ y $(y_i)_{i=1\dots m}$ en $\mathbb{R}^n$ definidos por

$$x_i = C_i(X), \quad y_i = C_i(Y), \quad i = 1, \dots, m.$$

Entonces, es sencillo verificar las siguientes equivalencias:

1. $X \succ_d Y$ si y solo si existe $A \in RS(n)$ tal que $Ax_i = y_i$ para cada $i = 1, \dots, m$.
2. $X \succ Y$ si y solo si $\sum_{j=1}^m \alpha_j x_j \succ \sum_{j=1}^m \alpha_j y_j$ para cualquier m -upla de escalares $(\alpha_1, \dots, \alpha_m) \in \mathbb{R}^m$.
3. $X \succ_f Y$ si y solo si existe $D \in DS(n)$ tal que $Dx_i = y_i$ para cada $i = 1, \dots, m$.

Luego, cada mayorización de matrices puede considerarse como una relación entre la m -upla (ordenada) de vectores columnas $(x_i)_{i=1\dots m}$ y $(y_i)_{i=1\dots m}$. □

Es inmediato de las definiciones y el Teorema 2.1.7 que la mayorización fuerte implica la mayorización direccional. Ahora damos una caracterización de la mayorización débil que nos permite verificar que es más débil que la mayorización direccional.

Proposición 7.1.4. Sean $X, Y \in M_{n,m}$ entonces,

- (i) $X \succ_d Y$ si y solo si $R(Y) \subseteq \text{conv}(R(X))$;
- (ii) Si $X \succ Y$ entonces $X \succ_d Y$.

Demostración. (i) Sean $X, Y \in M_{n,m}$ y $A = (a_{ij}) \in M_n$. Entonces $AX = Y$ si y solo si

$$R_i(Y) = \sum_{k=1}^n a_{ik} R_k(X) \quad i = 1, \dots, n.$$

Luego, si existe $A \in RS(n)$ tal que $AX = Y$ verificamos que $R(Y) \subseteq \text{conv}(R(X))$. Por otro lado, si $R(Y) \subseteq \text{conv}(R(X))$ entonces, por la ecuación anterior, podemos construir las filas de una matriz $A \in RS(n)$ tal que $AX = Y$.

(ii) Sean $X, Y \in M_{n,m}$ tal que $X \succ Y$ y supongamos que existe $1 \leq i \leq n$ tal que $R_i(Y) \notin \text{conv}(R(X))$. En este caso, existe un hiperplano que separa $R_i(Y)$ de $\text{conv}(R(X))$ i.e., existen $v \in \mathbb{R}^m$ y $t > 0$ tales que

$$\langle R_i(Y), v \rangle \geq t \quad \text{y} \quad \langle R_j(X), v \rangle < t \quad \text{para todo } j = 1, \dots, n.$$

Pero esto contradice la mayorización de los vectores $Xv \succ Yv$ pues

$$((Yv)^\downarrow)_1 \geq (Yv)_i = \langle R_i(Y), v \rangle \geq t > ((Xv)^\downarrow)_1.$$

Luego, $X \succ_d Y$. □

Observación 7.1.5. Como consecuencia de la Proposición 7.1.4 obtenemos un método eficiente para verificar si es válida la relación $X \succ_d Y$. De hecho, por el ítem i), solo tenemos que verificar si cada fila en Y puede ser escrita como una combinación convexa de las filas de X . Para esto se puede resolver un problema de programación lineal, cuyas variables son los coeficientes para las combinaciones convexas a encontrar. Más aún, para matrices de dimensiones pequeñas esto se puede estudiar de forma gráfica (ver Observación 7.1.16).

Notemos que, aunque la mayorización débil de las matrices $X \succ_d Y$, para $X, Y \in M_{n,m}$, puede considerarse como una relación algebraica entre las *columnas* de X y de Y (ver Observación 7.1.3), en la Proposición 7.1.4 obtenemos una caracterización geométrica de esta relación en términos de las *filas* de X e Y . □

Los siguientes ejemplos muestran que, en general, las diferentes mayorizaciones de matrices no son equivalentes.

Ejemplo 7.1.6. $X \succ_d Y$ no implica $X \succ Y$.

Sean

$$X = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \quad \text{y} \quad Y = \begin{pmatrix} 1 & 0 \\ 1/2 & 1/2 \end{pmatrix}.$$

Entonces, si tomamos $A = Y \in RS(n)$, es claro que $AX = Y$ con lo cual $X \succ_d Y$. Sin embargo, si consideramos $v = (2, 1)$ entonces $Xv \not\succ Yv$. Así se tiene que $X \not\succ Y$. □

Ejemplo 7.1.7. $X \succ Y$ no implica $X \succ_f Y$.

Este es un hecho conocido. En [62] se puede encontrar un ejemplo de A. Horn. Nuestro ejemplo usa matrices más pequeñas. En realidad, veremos en el Corolario 7.1.23 que las matrices de este

ejemplo tienen el mínimo número de filas y columnas que hacen falta para que la implicación sea falsa. Sean

$$X = \begin{pmatrix} 0 & 3 & -3 & 0 \\ 0 & -2 & -2 & 4 \end{pmatrix}^t \quad \text{y} \quad Y = \begin{pmatrix} 2 & -2 & 0 & 0 \\ 0 & 0 & -2 & 2 \end{pmatrix}^t.$$

Entonces $X \succ Y$ pero $X \not\succ_f Y$. Una prueba de este hecho aparece en la Observación 7.1.16. $\square$

Proposición 7.1.8. Sean $X, Y \in M_{n,m}$ y supongamos que $\text{rank}(X) = n$. Los siguientes son equivalentes:

- (i) $X \succ_d Y$.
- (ii) $YX^\dagger \in RS(n)$ y $\ker(X) \subseteq \ker(Y)$.

Demostración. Supongamos que $X \succ_d Y$, i.e. existe $A \in RS(n)$ tal que $AX = Y$. Del hecho de que $\text{rank}(X) = n$, vemos que $XX^\dagger = I_n$ y $A = AXX^\dagger = YX^\dagger$. La ecuación $Y = AX$ implica que $\ker(X) \subseteq \ker(Y)$ de forma evidente.

Recíprocamente, si $YX^\dagger \in RS(n)$ y $\ker(X) \subseteq \ker(Y)$, luego $X^\dagger X$ es la proyección ortogonal sobre $\ker X^\perp \supseteq \ker Y^\perp$ y $YX^\dagger X = Y$, es decir $X \succ_d Y$. $\square$

En lo que sigue, consideramos la mayorización débil de matrices cuando $X, Y \in M_n$, en particular cuando $X \in GL(n)$. El siguiente corolario es una consecuencia de la Proposición 7.1.8.

Corolario 7.1.9. Supongamos que $X, Y \in M_n$ y $X \in GL(n)$. Entonces, $X \succ_d Y$ si y solo si $YX^{-1} \in RS(n)$.

Proposición 7.1.10. Sean $X, Y \in M_n$. Si $X \succ_d Y$ entonces $|\det(X)| \geq |\det(Y)|$. Más aún, si $X \succ_d Y$ y $|\det(X)| = |\det(Y)| \neq 0$ entonces existe una matriz de permutación $P \in M_n$ tal que $Y = PX$.

Demostración. Sea $S_X = \text{conv}(R(X) \cup \{0\})$ (resp. $S_Y = \text{conv}(R(Y) \cup \{0\})$) el politopo generado por $R(X) \cup \{0\}$ (resp. $R(Y) \cup \{0\}$). Luego, $|\det(X)|$ y $|\det(Y)|$ son los volúmenes de S_X y S_Y respectivamente. Si $X \succ_d Y$ tenemos, por la Proposición 7.1.4, que $S_Y \subseteq S_X$. Luego $|\det(X)| \geq |\det(Y)|$.

Si suponemos además que $X \succ_d Y$ y $|\det(X)| = |\det(Y)| \neq 0$ entonces, usando la terminología descrita en los Preliminares, S_X y S_Y son simplejos con vértices $R(X) \cup \{0\}$ y $R(Y) \cup \{0\}$ respectivamente. Como $S_Y \subseteq S_X$ y $|\det(X)| = |\det(Y)| \neq 0$, entonces deben coincidir. En particular, estos conjuntos tienen los mismos vértices, es decir que X e Y tienen las mismas filas salvo orden. $\square$

7.1.2. Convexidad y mayorización de matrices

Comenzamos esta sección recordando algunas caracterizaciones conocidas de la mayorización fuerte en términos de funciones convexas. Una prueba de este resultado puede hallarse en [26].

Teorema 7.1.11. Sean $X, Y \in M_{n,m}$. Entonces $X \succ_f Y$ si y solo si, para cada función convexa $f : V \rightarrow \mathbb{R}$ vale que

$$\sum_{j=1}^n f(X_j) \geq \sum_{j=1}^n f(Y_j)$$

donde $V \subseteq \mathbb{R}^m$ es un conjunto convexo tal que $R(X) \cup R(Y) \subseteq V$.

Observación 7.1.12. En lo que sigue, utilizamos los siguientes hechos elementales acerca de conjuntos y funciones convexas:

(i) Dados $z, w_i \in \mathbb{R}^m$, $i = 1, \dots, n$,

$$z \in \text{conv}(\{w_i : i = 1, \dots, n\}) \text{ si y solo si } \max_{1 \leq i \leq n} \langle w_i, v \rangle \geq \langle z, v \rangle \text{ para todo } v \in \mathbb{R}^m.$$

(ii) Dados dos conjuntos convexas V_1 y V_2 , $V_1 \subset V_2$ si y solo si

$$\max_{x \in V_1} f(x) \leq \max_{x \in V_2} f(x)$$

para cada función convexa f definida sobre $V_1 \cup V_2$.

□

Como consecuencia de la Proposición 7.1.4 y la Observación 7.1.12 concluimos el siguiente:

Corolario 7.1.13. Sean $X, Y \in M_{n,m}$. $X \succ_d Y$ si y solo si

$$\max_{1 \leq i \leq n} f(X_i) \geq \max_{1 \leq i \leq n} f(Y_i)$$

para cada función convexa $f : V \rightarrow \mathbb{R}$ donde $V \subseteq \mathbb{R}^m$ es un conjunto convexo que contiene $R(X) \cup R(Y)$. Más aún, si consideramos las funciones lineales $\phi_z : \mathbb{R}^m \rightarrow \mathbb{R}$, $z \in \mathbb{R}^m$ definidas por $\phi_z(x) = \langle x, z \rangle$, $X \succ_d Y$ si y solo si

$$\max_{1 \leq i \leq n} \phi_z(X_i) \geq \max_{1 \leq i \leq n} \phi_z(Y_i)$$

para cada $z \in \mathbb{R}^m$.

El siguiente teorema caracteriza la mayorización direccional entre matrices en $M_{n,m}$, en términos de $\lceil \frac{n}{2} \rceil + 1$ politopos, donde $\lceil r \rceil$ es la parte entera de $r \in \mathbb{R}$.

Teorema 7.1.14. Sean $X, Y \in M_{n,m}$. $X \succ Y$ si y solo si, para $k = 1, \dots, \lceil \frac{n}{2} \rceil$ y $k = n$, el conjunto de promedios de k filas diferentes de Y está incluido en la cápsula convexa del conjunto de promedios de k filas diferentes de X .

Demostración. Sean $X, Y \in M_{n,m}$, y supongamos que el conjunto de promedios de k filas diferentes de Y está incluido en la cápsula convexa del conjunto de promedios de k filas diferentes de X . Sea $v \in \mathbb{R}^m$ y $1 \leq k \leq \lceil \frac{n}{2} \rceil$. Entonces, existe una permutación $\sigma \in \mathbb{S}_n$ tal que

$(Yv)_{\sigma(1)} \geq \dots \geq (Yv)_{\sigma(n)}$, donde $(Yv)_i$ es la i -ésima coordenada de $Yv \in \mathbb{R}^n$. Por hipótesis, existe una familia $(c_\mu)_{\mu \in \mathbb{S}_n} \subseteq \mathbb{R}_0^+$ tal que $\sum_{\mu \in \mathbb{S}_n} c_\mu = 1$ y

$$\frac{1}{k} \sum_{j=1}^k Y_{\sigma(j)} = \sum_{\mu \in \mathbb{S}_n} c_\mu \left(\frac{1}{k} \sum_{j=1}^k X_{\mu(j)} \right).$$

Luego vemos que

$$\begin{aligned} \sum_{j=1}^k (Yv)_j^\downarrow &= k \left\langle \frac{1}{k} \sum_{j=1}^k Y_{\sigma(j)}, v \right\rangle = k \left\langle \sum_{\mu \in \mathbb{S}_n} c_\mu \left(\frac{1}{k} \sum_{j=1}^k X_{\mu(j)} \right), v \right\rangle = \\ &= k \sum_{\mu \in \mathbb{S}_n} c_\mu \left\langle \frac{1}{k} \sum_{j=1}^k X_{\mu(j)}, v \right\rangle \leq k \max_{\mu \in \mathbb{S}_n} \left\langle \frac{1}{k} \sum_{j=1}^k X_{\mu(j)}, v \right\rangle = \\ &= \sum_{j=1}^k (Xv)_j^\downarrow. \end{aligned}$$

Notemos que la hipótesis para $k = n$ implica que

$$\sum_{j=1}^n Y_j = \sum_{j=1}^n X_j.$$

Sea $\lfloor \frac{n}{2} \rfloor < k < n$, y $\tau \in \mathbb{S}_n$ una permutación tal que $(Yv)_{\tau(1)} \leq \dots \leq (Yv)_{\tau(n)}$. Nuevamente, por hipótesis, existe $(d_\mu)_{\mu \in \mathbb{S}_n} \subseteq \mathbb{R}_0^+$, $\sum_{\mu \in \mathbb{S}_n} d_\mu = 1$, tal que

$$\frac{1}{n-k} \sum_{j=1}^{n-k} Y_{\tau(j)} = \sum_{\mu \in \mathbb{S}_n} d_\mu \left(\frac{1}{n-k} \sum_{j=1}^{n-k} X_{\mu(j)} \right),$$

ya que $1 \leq n-k \leq \lfloor \frac{n}{2} \rfloor$. De aquí que

$$\begin{aligned} \sum_{j=1}^k (Yv)_j^\downarrow &= \left\langle \sum_{j=1}^n Y_j, v \right\rangle - \sum_{j=1}^{n-k} \langle Y_{\tau(j)}, v \rangle = \\ &= \left\langle \sum_{j=1}^n X_j, v \right\rangle - (n-k) \left\langle \sum_{\mu \in \mathbb{S}_n} d_\mu \left(\frac{1}{n-k} \sum_{j=1}^{n-k} X_{\mu(j)} \right), v \right\rangle \leq \\ &\leq \left\langle \sum_{j=1}^n X_j, v \right\rangle - (n-k) \min_{\mu \in \mathbb{S}_n} \left\langle \frac{1}{n-k} \sum_{j=1}^{n-k} X_{\mu(j)}, v \right\rangle = \\ &= \sum_{j=1}^k (Xv)_j^\downarrow, \end{aligned}$$

Luego, $Xv \succ Yv$. Como $v \in \mathbb{R}^m$ era arbitrario entonces $X \succ Y$. Por otro lado, si supongamos que $X \succ Y$ entonces, dada $\mu \in \mathbb{S}_n$,

$$\max_{\sigma \in \mathbb{S}_n} \left\langle \sum_{i=1}^k X_{\sigma(i)}, v \right\rangle = \sum_{i=1}^k (Xv)_i^\downarrow \geq \sum_{i=1}^k (Yv)_i^\downarrow \geq \left\langle \sum_{i=1}^k Y_{\mu(i)}, v \right\rangle \quad \text{para todo } v \in \mathbb{R}^n.$$

y por la Observación 7.1.12, tenemos que $\frac{1}{k} \sum_{i=1}^k Y_{\mu(i)}$ pertenece a la cápsula convexa de

$$\left\{ \frac{1}{k} \sum_{i=1}^k X_{\sigma(i)} : \sigma \in \mathbb{S}_n \right\}.$$

□

El Corolario 7.1.13 y el Teorema 7.1.14 implican la siguiente descripción de la mayorización direccional en términos de la mayorización débil.

Corolario 7.1.15. Sean $X, Y \in M_{n,m}$. $X \succ Y$ si y solo si $\overline{X}(k) \succ_d \overline{Y}(k)$ para $k = 1, \dots, \lfloor \frac{n}{2} \rfloor$ y $k = n$, donde $\overline{X}(k)$ (respectivamente $\overline{Y}(k)$) es la matriz de $\frac{n!}{k!(n-k)!}$ filas, que son todos los posibles promedios de k filas diferentes de X (respectivamente de Y).

Como consecuencia del Corolario 7.1.15 y la Observación 7.1.5 obtenemos un criterio efectivo para verificar la relación $X \succ Y$. En efecto, con la notación de arriba, solo tenemos que verificar si $\overline{X}(k) \succ_d \overline{Y}(k)$ para $k = 1, \dots, \lfloor \frac{n}{2} \rfloor$ y $k = n$ (i.e., $\lfloor \frac{n}{2} \rfloor + 1$ instancias de mayorización débil de matrices). Para un tal k , podemos recurrir a la programación lineal (como se comentó en la Observación 7.1.5) para verificar si se tiene $\overline{X}(k) \succ_d \overline{Y}(k)$.

Observación 7.1.16. Sean X, Y las matrices del Ejemplo 7.1.7. Para determinar que $X \succ Y$, solo tenemos que verificar que $\overline{X}(k) \succ_d \overline{Y}(k)$ para $k = 1, 2, 4$, por el Corolario 7.1.15.

Figura 7.1: Polígonos correspondientes a $k = 1$ y $k = 2$

En primer lugar, $\overline{X}(4) = (0, 0) = \overline{Y}(4) \in M_{1,2}$, de forma que $\overline{X}(4) \succ_d \overline{Y}(4)$. Más aún,

$$\overline{X}(2) = \begin{pmatrix} 3/2 & -3/2 & 0 & 0 & 3/2 & -3/2 \\ -1 & -1 & 2 & -2 & 1 & 1 \end{pmatrix}^t \quad \text{y} \quad \overline{Y}(2) = \begin{pmatrix} 0 & 1 & 1 & -1 & -1 & 0 \\ 0 & -1 & 1 & -1 & 1 & 0 \end{pmatrix}^t.$$

Los gráficos en la Figura 1 muestran la inclusión de los polígonos que prueban $\overline{X}(k) \succ_d \overline{Y}(k)$ para $k = 1, 2$. De lo anterior deducimos que $X \succ Y$. Por otra parte, la función convexa $f(x, y) = \max\{-y, \frac{y}{2} + x, \frac{y}{2} - x\}$ y el Teorema 7.1.11 muestran que $X \not\succ_f Y$ en el Ejemplo 7.1.7.

□

7.1.3. Cuando la mayorización débil implica la mayorización fuerte

En esta sección estudiamos condiciones bajo las cuales la mayorización débil o direccional de matrices implica la mayorización fuerte de matrices. Este problema es de interés y ha sido considerado en los trabajos de investigación [45], [62], [63].

Proposición 7.1.17. Sean $X, Y \in M_{n,m}$ tales que $X \succ Y$. Supongamos que $\text{conv}(R(X))$ tiene solo dos puntos extremales. Entonces $X \succ_f Y$.

Demostración. Notemos que, como $\text{conv}(R(X))$ tiene solo dos puntos extremales, los puntos en $R(X)$ están contenidos en una recta de $\mathbb{R}^m$. Luego, por hipótesis, los puntos de $R(Y)$ también deben pertenecer a esta recta. Sea $Z \in M_{n,m}$ la matriz cuyas filas son todas iguales a $R_1(X)$. Es sencillo verificar que $X \succ Y$ (resp. $X \succ_f Y$) si y solo si $X - Z \succ Y - Z$ (resp. $X - Z \succ_f Y - Z$).

Luego, podemos suponer que $\text{rank } X \leq 1$ y $\text{rank } Y \leq 1$. Si $X = 0$ el resultado es inmediato. Si $Y = 0$ y $\text{rank } X = 1$ supongamos que $Xe_1 \neq 0$ y consideremos la matriz $D \in DS(n)$ tal que $D(Xe_1) = Ye_1 = 0$. Entonces $DX = 0 = Y$, pues cada columna de X es un múltiplo real de $C_1(X) = Xe_1$. Si $\text{rank } Y = \text{rank } X = 1$, sean $x_1, y_1 \in \mathbb{R}^n$ y $x_2, y_2 \in \mathbb{R}^m$ tales que $X = x_1 x_2^t$ y $Y = y_1 y_2^t$. Más aún, como $\mathbb{R}y_2 = \text{ran}(Y^t) = \text{ran}(X^t) = \mathbb{R}x_2$, podemos asumir que $y_2 = x_2$. Notemos que $Xx_2 = \langle x_2, x_2 \rangle x_1$ y $Yx_2 = \langle x_2, x_2 \rangle y_1$.

Como $X \succ Y$, entonces $x_1 \succ y_1$ y existe $D \in DS(n)$ tal que $Dx_1 = y_1$. Luego

$$DX = Dx_1 x_2^t = y_1 x_2^t = Y$$

y vemos entonces que $X \succ_f Y$. □

Dada $X \in M_{n,m}$ denotamos por $[X, e] \in M_{n, (m+1)}$ la matriz cuyas primeras m columnas son iguales a las columnas de X y su última columna es el vector $e = (1, \dots, 1) \in \mathbb{R}^n$. En [45], S.-G. Hwang y S.-S. Pyo probaron el siguiente teorema.

Teorema. Sean $X, Y \in M_{n,m}$ tales que $[Y, e][X, e]^\dagger$ tiene entradas no negativas. Entonces $X \succ Y$ si y solo si $X \succ_f Y$.

Extendemos este resultado reemplazando $X \succ Y$ por $X \succ_d Y$ mas la condición $e^t X = e^t Y$. Notemos que si $X \succ Y$ entonces $X \succ_d Y$ y $e^t X = e^t Y$ (ver Corolario 7.1.15), pero la otra implicación no vale (ver Observación 7.1.20). Más aún, usando la noción de mayorización débil de matrices nuestra prueba es más sencilla.

Teorema 7.1.18. Sean $X, Y \in M_{n,m}$ y supongamos que $[Y, e][X, e]^\dagger$ tiene entradas no negativas. Si $X \succ_d Y$ y $e^t X = e^t Y$ entonces $X \succ_f Y$.

Para probar este teorema vamos a utilizar el siguiente lema cuya demostración es elemental y omitimos.

Lema 7.1.19. Sean $X, Y \in M_{n,m}$ entonces

$$\begin{aligned} X \succ_d Y &\text{ si y solo si } [X, e] \succ_d [Y, e] \\ X \succ Y &\text{ si y solo si } [X, e] \succ [Y, e] \\ X \succ_f Y &\text{ si y solo si } [X, e] \succ_f [Y, e] \\ e^t X = e^t Y &\text{ si y solo si } e^t [X, e] = e^t [Y, e] \end{aligned}$$

Demostración del Teorema 7.1.18. Sean $Z = [X, e]$ y $W = [Y, e]$. Aplicando en Lema 7.1.19 solo tenemos que probar que, si $WZ^\dagger$ tiene entradas no negativas, entonces $Z \succ_d W$ y $e^t Z = e^t W$ implica $Z \succ_f W$.

Supongamos que $Z \succ_d W$, entonces existe una matriz estocástica por filas A tal que $W = AZ$. Multiplicando ambos lados de la ecuación por $Z^\dagger$ obtenemos:

$$WZ^\dagger = AZZ^\dagger = AP$$

donde P es la proyección ortogonal sobre el rango de Z . Como $APZ = AZ = W$, entonces tendremos que $Z \succ_f W$ en tanto verifiquemos que AP es doble estocástica. Por hipótesis sabemos que $AP = WZ^\dagger$ tiene entradas no negativas con lo cual basta mostrar que $APe = e$ y $e^t AP = e^t$. Dado que $Z = [X, e]$, entonces e está en la imagen de Z y $Pe = e$. Así tenemos que

$$APe = Ae = e$$

pues A es estocástica por filas. Por hipótesis, $e^t Z = e^t W$, con lo cual

$$e^t AP = e^t WZ^\dagger = e^t ZZ^\dagger = e^t P = e^t$$

lo que muestra que AP es doble estocástica, $(AP)Z = W$, y por el Lema 7.1.19 también vale que $(AP)X = Y$. □

Observación 7.1.20. La condición $X \succ_w Y$ y $e^t X = e^t Y$ del Teorema 7.1.18 es más débil que las hipótesis $X \succ Y$ del Teorema de Hwang y Pyo. De hecho, sean $X, Y \in M_{6,2}$ las siguientes matrices

$$X = \begin{pmatrix} 0 & 1 & 1 & -1 & -1 & 0 \\ 1 & 1 & -1 & -1 & 1 & 1 \end{pmatrix}^t \quad \text{y} \quad Y = \begin{pmatrix} 2/3 & 2/3 & 1 & -1 & -2/3 & -2/3 \\ 1 & 1 & -1 & -1 & 1 & 1 \end{pmatrix}^t.$$

Es sencillo ver que $X \succ_d Y$ y $e^t X = (0, 2) = e^t Y$. Sin embargo, la Figura 2 muestra que $\overline{X}(2) \not\succ_d \overline{Y}(2)$ (donde $\blacksquare$ representan las filas de $\overline{X}(2)$ y $\blacktriangle$ representan las filas de $\overline{Y}(2)$). Entonces, por el Corolario 7.1.15, $X \not\succ Y$.

Figura 7.2: $\text{conv}(\overline{Y}(2)) \not\subseteq \text{conv}(\overline{X}(2))$

□

Los siguientes resultados son consecuencias del Teorema 7.1.18.

Corolario 7.1.21. Sean $X, Y \in M_{n,m}$ y supongamos que $\text{ran}([X, e]) = \mathbb{R}^n$. Si $X \succ_d Y$ y $e^t X = e^t Y$ entonces $X \succ_f Y$.

Demostración. Se deduce de la Proposición 7.1.8 y el Teorema 7.1.18. □

Corolario 7.1.22. Sean $X, Y \in M_{n,m}$ tales que las filas de X son los vértices de un símplex. Si $X \succ_d Y$ y $e^t X = e^t Y$ entonces $X \succ_f Y$.

Demostración. El hecho de que las filas de X generen un *símplex* es equivalente a que el conjunto $\{R_2(X) - R_1(X), \dots, R_n(X) - R_1(X)\}$ sea linealmente independiente. Entonces, la dimensión del rango de la matriz

$$Z = \begin{pmatrix} 0 \\ R_2(X) - R_1(X) \\ \vdots \\ R_n(X) - R_1(X) \end{pmatrix}$$

es $n - 1$. Luego, el subespacio $\mathcal{S}$ generado por las columnas de Z también tiene dimensión $n - 1$ y $e \notin \mathcal{S}$. Usando el hecho de que $C_i(Z) = C_i(X) - x_{1i}e$, $1 \leq i \leq m$, concluimos que el conjunto $\{C_1(X), \dots, C_m(X), e\}$ genera $\mathbb{R}^n$. Por el Corolario 7.1.21, obtenemos que $X \succ_f Y$. □

Corolario 7.1.23. Sean $X, Y \in M_{n,m}$ con $1 \leq n \leq 3$. Entonces, $X \succ Y$ implica que $X \succ_f Y$.

Demostración. Sean $X, Y \in M_{n,m}$, con $1 \leq n \leq 3$, tales que $X \succ Y$. Si $\text{conv}(R(X))$ es un segmento, el resultado se sigue de la Proposición 7.1.17. En otro caso, $n = 3$ y tenemos que $\text{conv}(R(X))$ es un triángulo contenido en $\mathbb{R}^m$ con vértices $X_i = R_i(X)$, $i = 1, 2, 3$. Aplicando el Corolario 7.1.22 en este caso, vemos que $X \succ_f Y$. □

Observación 7.1.24. Sean $X, Y \in M_{n,m}$. Notemos que $X \succ Y$ es equivalente a las desigualdades $f(Xv) \geq f(Yv)$ para cada $v \in \mathbb{R}^m$ y cada función convexa simétrica $f : \mathbb{R}^n \rightarrow \mathbb{R}$ (ver Teorema 2.1.3). Por otra parte, si consideramos las funciones convexas y simétricas $f_\infty(z_1, \dots, z_n) = \max(z_1, \dots, z_n)$ y $f_1(z_1, \dots, z_n) = z_1 + \dots + z_n$, entonces $X \succ_d Y$ es equivalente a $f_\infty(Xv) \geq f_\infty(Yv)$ para cada $v \in \mathbb{R}^m$, mientras que $e^t X = e^t Y$ es equivalente a $f_1(Xv) \geq f_1(Yv)$ para cada $v \in \mathbb{R}^m$.

Supongamos que $\text{ran}([X, e]) = \mathbb{R}^n$. El Corolario 7.1.21 establece que si $X \succ_d Y$ y $e^t X = e^t Y$ entonces $X \succ_f Y$. Podemos re-escribir este resultado como sigue: si $f_\infty(Xv) \geq f_\infty(Yv)$ y $f_1(Xv) \geq f_1(Yv)$ para cada $v \in \mathbb{R}^m$ entonces, $f(Xv) \geq f(Yv)$ para cada $v \in \mathbb{R}^m$ y cada función convexa simétrica $f : \mathbb{R}^n \rightarrow \mathbb{R}$. Esta reformulación del resultado es similar al siguiente teorema de interpolación : si $A \in M_{n,m}$ es tal que $\|Av\|_\infty \leq \|v\|_\infty$ y $\|Av\|_1 \leq \|v\|_1$ para cada $v \in \mathbb{R}^m$ entonces $\|Av\| \leq \|v\|$ para cada $v \in \mathbb{R}^m$ y cada norma gauge simétrica. □

7.1.4. Relaciones de equivalencias asociadas a las mayorizaciones de matrices

Como ya hemos mencionado, las mayorizaciones de matrices previamente consideradas son relaciones de preorden. Como $X \succ_d Y$ si y solo si $R(Y) \subseteq \text{conv}(R(X))$, es claro que $X \succ_d Y$ y $Y \succ_d X$ es equivalente a la igualdad de conjuntos $\text{conv}(R(X)) = \text{conv}(R(Y))$. El siguiente teorema describe la relación de equivalencia asociada a las mayorizaciones direccional y fuerte de matrices.

Teorema 7.1.25. Sean $X, Y \in M_{n,m}$. Entonces los siguientes son equivalentes

- i) Existe una matriz de permutación $Q \in M_n$ tal que $QX = Y$.
- ii) $X \succ_f Y$ y $Y \succ_f X$.

iii) $X \succ Y$ y $Y \succ X$.

Antes de probar este resultado, consideramos la siguiente propiedad de la mayorización direccional de matrices.

Lema 7.1.26. Sean $X, Y \in M_{n,m}$ tales que $X \succ Y$ y $R_{i_0}(X) = R_{j_0}(Y)$. Si $\tilde{X} \in M_{(n-1),m}$ (resp. $\tilde{Y} \in M_{(n-1),m}$) denota la matriz que se obtiene de borrar la i_0 -ésima fila de X (resp. la j_0 -ésima fila de Y), entonces $\tilde{X} \succ \tilde{Y}$.

Demostración. Es consecuencia del siguiente hecho (ver los Preliminares): si $x, y \in \mathbb{R}^r$ entonces, para cada $\lambda \in \mathbb{R}$,

$$x \succ y \iff (x_1, \dots, x_r, \lambda) \succ (y_1, \dots, y_r, \lambda).$$

□

Demostración del Teorema 7.1.25 Las implicaciones $i) \Rightarrow ii) \Rightarrow iii)$ son evidentes. Solo tenemos que probar la implicación $iii) \Rightarrow i)$. Usamos inducción en el número de filas de X e Y . Si $n = 1$ es inmediato: notemos que si $X, Y \in M_{1,m}$ entonces $X \succ Y$ implica $X = Y$.

En el caso en que $n > 1$, si $X \succ Y$ y $Y \succ X$ entonces, $X \succ_d Y$ y $Y \succ_d X$. Luego la cápsula convexa de $R(X)$ coincide con la de $R(Y)$ y en particular tienen los mismos puntos extremales. Si z es un punto extremal de $\text{conv}(R(X)) = \text{conv}(R(Y))$ entonces, $z = R_{i_0}(X) = R_{j_0}(Y)$ con $1 \leq i_0, j_0 \leq n$.

Si $\tilde{X}, \tilde{Y} \in M_{(n-1),m}$ son como en el lema anterior, entonces tenemos que $\tilde{X} \succ \tilde{Y}$ y $\tilde{Y} \succ \tilde{X}$. Por la hipótesis inductiva, las filas de $\tilde{X}$ son un reordenamiento de las filas de $\tilde{Y}$. Luego las filas de X son un reordenamiento de las filas de Y , lo que implica $i)$.

□

El siguiente corolario, que es una consecuencia de la Proposición 7.1.10, es un análogo del Teorema 7.1.25 para la mayorización débil de matrices, en el caso particular en que $X, Y \in GL(n)$.

Corolario 7.1.27. Sean $X, Y \in M_n$ con $Y \in GL(n)$. Entonces los siguientes son equivalentes

- i) Existe una matriz de permutación $Q \in M_n$ tal que $QX = Y$.
- ii) $X \succ_d Y$ y $Y \succ_d X$.

Seguidamente, determinamos el conjunto de matrices minimales con respecto a los preórdenes que hemos considerado hasta ahora. En este contexto, un elemento *minimal* con respecto a un preorden $\ll$ en un conjunto P es un elemento $m \in P$ tal que, dado $n \in P$, si $n \ll m$ entonces $m \ll n$.

Proposición 7.1.28. $X \in M_{n,m}$ es minimal con respecto a cualquiera de los preórdenes $\succ_d, \succ$ ó $\succ_f$ si y solo si $X_1 = \dots = X_n$, es decir, todas las filas de X son iguales.

Demostración. Si $R(X) = \{v\}$, para algún $v \in \mathbb{R}^m$, entonces $\text{conv}(R(X)) = \{v\}$. Entonces, si $X \succ_d Y$ es claro que $X = Y$. Por otro lado, sea $X \in M_{n,m}$ una matriz con al menos dos filas diferentes. Entonces $R(X)$ contiene dos puntos diferentes (en $\mathbb{R}^m$). Si $D \in DS(n)$ es la matriz con todas sus entradas iguales a $1/n$ tenemos que $Y = DX \prec_f X$. Más aún, como $R_1(Y) = R_2(Y) = \dots = R_n(Y)$, entonces $\text{conv}(R(Y))$ contiene solo un punto, con lo cual $R(X) \not\subset \text{conv}(R(Y))$. Luego $Y \not\succ_d X$ y X no es minimal con respecto a ninguna de las mayorizaciones de matrices.

□

7.2. Mayorizaciones conjuntas

En [5] T. Ando introdujo la relación de mayorización entre matrices autoadjuntas de la siguiente forma: si $a, b \in H(n)$, el espacio real de matrices autoadjuntas en $M_n(\mathbb{C})$ y $\lambda(a), \lambda(b) \in \mathbb{R}^n$ denotan los vectores de autovalores de a y b respectivamente, contados con multiplicidad, a *mayoriza* a b (en el sentido de Ando) si $\lambda(a) \succ \lambda(b)$; en este caso escribimos $a \succ b$ (ver la sección 2.1.3 de los Preliminares).

El objeto de esta sección es introducir y caracterizar algunas posibles extensiones de la mayorización en $H(n)$, que llamamos *mayorizaciones conjuntas*. Algunos de los resultados de esta sección están basados en el material previamente obtenido acerca de la mayorización de matrices.

7.2.1. Mayorización conjunta entre familias abelianas en $H(n)$

Una *familia abeliana* es una familia ordenada $(a_i)_{i=1,\dots,m}$ de matrices autoadjuntas en $M_n(\mathbb{C})$ tales que

$$a_i a_j = a_j a_i, \quad i, j = 1, \dots, m.$$

Para introducir las mayorizaciones conjuntas consideramos los siguientes hechos elementales: si $(a_i)_{i=1,\dots,m}$ y $(b_i)_{i=1,\dots,m}$ son dos familias abelianas en $M_n(\mathbb{C})$ entonces existen matrices unitarias $U, V \in M_n(\mathbb{C})$ tales que

$$U^* a_i U = D_{\lambda(a_i)}, \quad V^* b_i V = D_{\lambda(b_i)}, \quad i = 1, \dots, m,$$

donde D_x denota la matriz diagonal con diagonal principal $x \in \mathbb{R}^n$. En este caso $\lambda(a_i)$ es el vector de autovalores correspondientes a a_i , contados con multiplicidad, en algún orden dependiendo de U . Consideremos las matrices $A, B \in M_{n,m}$ dadas por

$$C_i(A) = \lambda(a_i), \quad C_i(B) = \lambda(b_i), \quad i = 1, \dots, m.$$

Definición 7.2.1. Sean $(a_i)_{i=1,\dots,m}, (b_i)_{i=1,\dots,m} \subseteq M_n(\mathbb{C})$ dos familias abelianas y sean $A, B \in M_{n,m}$ definidas como antes. Decimos que la familia $(a_i)_{i=1,\dots,m}$ *mayoriza conjuntamente de forma débil* (respectivamente *mayoriza conjuntamente de forma fuerte, mayoriza conjuntamente de forma direccional*) a la familia $(b_i)_{i=1,\dots,m}$ y notamos

$$(a_i)_{i=1,\dots,m} \succ_d (b_i)_{i=1,\dots,m}$$

(respectivamente $(a_i)_{i=1,\dots,m} \succ_f (b_i)_{i=1,\dots,m}$, $(a_i)_{i=1,\dots,m} \succ (b_i)_{i=1,\dots,m}$) si $A \succ_d B$ (respectivamente $A \succ_f B$, $A \succ B$).

Notemos que si U, W son dos matrices unitarias que diagonalizan la familia $(a_i)_{i=1,\dots,m}$ simultáneamente, entonces existe una matriz de permutación Q tal que

$$U^* a_i U = Q^* W^* a_i W Q, \quad i = 1, \dots, m.$$

Así, si $A' \in M_{n,m}$ denota la matriz cuyas columnas $C_i(A')$ son las diagonales principales de las matrices $W^* a_i W$, entonces $A = Q A'$. Es decir, la definición anterior no depende de la matriz unitaria U . Esto también muestra que el conjunto de de filas $R(A)$ no depende de la matriz unitaria

U . Este conjunto es llamado el *espectro conjunto* de la familia y lo denotamos por $\sigma(a_1, \dots, a_m)$. Más aún, si $f : V \rightarrow \mathbb{C}$ es tal que $\sigma(a_1, \dots, a_m) \subseteq V$ entonces consideramos

$$f(a_1, \dots, a_m) = UD_x U^*$$

donde D_x es la matriz diagonal con diagonal principal $x = (f(R_1(A)), \dots, f(R_n(A))) \in \mathbb{C}^n$. Notemos que $f(a_1, \dots, a_m) \in M_n(\mathbb{C})$ no depende de la matriz unitaria U . Decimos que la matriz $f(a_1, \dots, a_m)$ se obtiene de la familia $(a_i)_{i=1, \dots, m}$ por cálculo funcional.

De aquí en adelante, siempre que el contexto lo permita, evitaremos escribir los subíndices correspondientes a las familias de matrices y notamos $(a_i) \succ_d (b_i)$ (resp. $(a_i) \succ_f (b_i)$, $(a_i) \succ (b_i)$).

Proposición 7.2.2. Sean $(a_i)_{i=1, \dots, m}$ y $(b_i)_{i=1, \dots, m}$ dos familias abelianas en $M_n(\mathbb{C})$. Entonces

1. $(a_i) \succ_d (b_i)$ si y solo si $\text{conv}(\sigma(b_1, \dots, b_m)) \subseteq \text{conv}(\sigma(a_1, \dots, a_m))$.
2. $(a_i) \succ (b_i)$ si y solo si, para cada $\gamma_1, \dots, \gamma_m \in \mathbb{R}$ vale que

$$\gamma_1 a_1 + \dots + \gamma_m a_m \succ \gamma_1 b_1 + \dots + \gamma_m b_m \text{ (en el sentido de Ando).}$$

3. $(a_i) \succ_f (b_i)$ si y solo si existen $k \in \mathbb{N}$, matrices unitarias $W_1, \dots, W_k \in M_n(\mathbb{C})$ números no negativos $\mu_1, \dots, \mu_k$, $\sum_{j=1}^k \mu_j = 1$, tales que

$$b_i = \sum_{j=1}^k \mu_j W_j^* a_i W_j, \quad \text{para } 1 \leq i \leq m. \quad (7.1)$$

Demostración. Los ítems 1. y 2. son consecuencias de las definiciones, por lo cual omitimos la prueba. La prueba del ítem 3. es postpuesta hasta la prueba del Teorema 7.2.5. □

7.2.2. Caracterizaciones de las mayorizaciones conjuntas

En esta subsección consideramos algunas caracterizaciones de las diferentes nociones de mayorización conjunta previamente introducidas.

Recordemos que una transformación lineal $T : S \rightarrow M_n(\mathbb{C})$ de un subespacio unital y autoadjunto $S^* = S \subseteq M_n(\mathbb{C})$, $I \in S$, es llamada *unital* si $T(I) = I$, donde I denota la matriz identidad; *positiva* si $T(a)$ es positiva semidefinida siempre que a sea positiva semidefinida, y *que preserva la traza* si $\text{tr}(T(a)) = \text{tr}(a)$ para cada $a \in S$.

Lema 7.2.3. Sea $\mathcal{D}$ el álgebra diagonal en $M_n(\mathbb{C})$ y sea $T : \mathcal{D} \rightarrow \mathcal{D}$ positiva y unital. Entonces existe $E \in RS(n)$ tal que

$$T(D_x) = D_{Ex}. \quad (7.2)$$

Si además, T preserva la traza entonces $E \in DS(n)$. Por otra parte, si T es como en (7.2) para alguna matriz $E \in RS(n)$ (respectivamente $E \in DS(n)$), entonces T positiva y unital (resp. es positiva, unital y preserva la traza).

Demostración. Fijamos la base canónica $\{e_i\}_{i=1,\dots,n}$ de $\mathbb{C}^n$ e identificamos $\mathcal{D}$ con $\mathbb{C}^n$ como espacios vectoriales por la función $D_x \mapsto x$, donde D_x es la matriz diagonal con diagonal principal $x \in \mathbb{C}^n$. Luego, bajo esta identificación, T induce una transformación lineal $\tilde{T}$ en $\mathbb{C}^n$ dada por $\tilde{T}x = \sum_{i=1}^n T(D_x)_{ii} e_i$. Sea E la matriz de $\tilde{T}$ con respecto a la base canónica. Entonces E satisface $Ee = e$ y $Ex \geq 0$ siempre que $x \geq 0$, donde $y \geq 0$ significa que todas las coordenadas de $y \in \mathbb{R}^n$ son no negativas. Luego $E \in RS(n)$ y $T(D_x) = D_{Ex}$. Más aún, si T preserva la traza entonces $\text{tr}(Ex) = \text{tr}(x)$, donde $\text{tr}(y) = y_1 + \dots + y_n$. Esto implica que $E \in DS(n)$. La recíproca es evidente de lo anterior. $\square$

En lo que sigue haremos uso de los siguientes resultados.

Lema 7.2.4. Sea $\mathcal{A} \subseteq M_n(\mathbb{C})$ una $*$ -subálgebra unital de $M_n(\mathbb{C})$. Entonces existe una transformación lineal positiva, unital y que preserva la traza, $\Psi : M_n(\mathbb{C}) \rightarrow \mathcal{A}$, tal que $\Psi(a) = a$ para todo $a \in \mathcal{A}$.

Dada $(a_i)_{i=1,\dots,m} \subseteq M_n(\mathbb{C})$, $C^*(a_1, \dots, a_m)$ denota la $*$ -subálgebra unital generada por los a_i 's, es decir, la $*$ -subálgebra $\mathcal{A}$ más pequeña de $M_n(\mathbb{C})$ con la propiedad de que $a_i \in \mathcal{A}$, $i = 1, \dots, m$.

Notemos que si $(a_i)_{i=1,\dots,m}$ es una familia abeliana en $M_n(\mathbb{C})$ y U es una matriz unitaria que diagonaliza simultáneamente esta familia entonces U diagonaliza cualquier $a \in C^*(a_1, \dots, a_m)$ i.e., $U^*aU = D_x$ para algún $x \in \mathbb{C}^n$. Luego, $C^*(a_1, \dots, a_m) \subseteq UDU^* = \{UD_xU^* : x \in \mathbb{C}^n\}$.

Teorema 7.2.5. Sean $(a_i)_{i=1,\dots,m}$ y $(b_i)_{i=1,\dots,m}$ dos familias abelianas en $M_n(\mathbb{C})$. Entonces

1. $(a_i) \succ_d (b_i)$ si y solo si existe una transformación positiva y unital

$$T : C^*(a_1, \dots, a_m) \rightarrow C^*(b_1, \dots, b_m)$$

tal que $T(a_i) = b_i$ para cada $i = 1, \dots, m$.

2. $(a_i) \succ (b_i)$ si y solo si, para cada $k = 1, \dots, [\frac{n}{2}]$ y $k = n$ tenemos que

$$(\log[\bigwedge_{i=1,\dots,m}^k \exp(a_i)])_{i=1,\dots,m} \succ_d (\log[\bigwedge_{i=1,\dots,m}^k \exp(b_i)])_{i=1,\dots,m}.$$

3. $(a_i) \succ_f (b_i)$ si y solo si existe una transformación positiva, unital y que preserva la traza

$$T : C^*(a_1, \dots, a_m) \rightarrow C^*(b_1, \dots, b_m)$$

tal que $T(a_i) = b_i$ para cada $i = 1, \dots, m$.

Demostración. Sean $U, V \in M_n(\mathbb{C})$ matrices unitarias tales que

$$U^*a_iU = D_{\lambda(a_i)}, \quad V^*b_iV = D_{\lambda(b_i)}, \quad i = 1, \dots, m$$

donde $\lambda(a_i), \lambda(b_i) \in \mathbb{R}^n$. Como antes, si $a \in C^*(a_1, \dots, a_m)$ entonces $U^*aU \in \mathcal{D}$.

1. Supongamos que existe una transformación positiva y unital $T : C^*(a_1, \dots, a_m) \rightarrow C^*(b_1, \dots, b_m)$ tal que $T(a_i) = b_i$ para cada $i = 1, \dots, m$. Sea $\tilde{T} : \mathcal{D} \rightarrow \mathcal{D}$ definida por

$$\tilde{T}(D_x) = V^*T(\Psi(UD_xU^*))V,$$

donde $\Psi : M_n(\mathbb{C}) \rightarrow C^*(a_1, \dots, a_m)$ es la transformación obtenida en el Lema 7.2.4. Notemos que $\tilde{T}$ es positiva, unital y tal que $\tilde{T}(D_{\lambda(a_i)}) = D_{\lambda(b_i)}$, $i = 1, \dots, m$. Por el Lema 7.2.3 existe $E \in RS(n)$ tal que $E\lambda(a_i) = \lambda(b_i)$, $i = 1, \dots, m$. De aquí que $(a_i) \succ_d (b_i)$ (ver Observación 7.1.3).

Por otro lado, si $(a_i) \succ_d (b_i)$, sea $E \in RS(n)$ tal que $E\lambda(a_i) = \lambda(b_i)$, $i = 1, \dots, m$. Sea $T : U\mathcal{D}U^* \rightarrow C^*(b_1, \dots, b_m)$ definida por

$$T(UD_xU^*) = \Phi(VD_{E_x}V^*),$$

donde $\Phi : M_n(\mathbb{C}) \rightarrow C^*(b_1, \dots, b_m)$ es la transformación obtenida como en el Lema 7.2.4. Basta considerar entonces la restricción de T a $C^*(a_1, \dots, a_m)$.

2. Notemos que $\bigwedge^k U$ es una matriz unitaria que diagonaliza la familia $(\bigwedge^k a_i)$. Sea $A \in M_{n,m}$ tal que para $1 \leq i \leq m$, $C_i(A) = \lambda(a_i)$. Para $1 \leq k \leq n$, sea $A_{(k)}$ la matriz $\frac{n!}{k!(n-k)!} \times m$ cuyas columnas son $C_i(A_{(k)}) = \lambda(i, k)$, donde

$$\bigwedge^k U^* \left(\log \left[\bigwedge^k \exp(a_i) \right] \right) \bigwedge^k U = D_{\lambda(i,k)}, \quad 1 \leq i \leq m.$$

Vamos a mostrar que, para $1 \leq k \leq n$, $A_{(k)} = k \cdot \bar{A}(k)$ (salvo una permutación de filas) donde $\bar{A}(k)$ es como en el Corolario 7.1.15. Sea $1 \leq j \leq n$ y denotemos por $u_j = C_j(U)$ las columnas de U . Entonces, para calcular las filas de $A_{(k)}$ basta notar que, para cada $1 \leq i \leq m$ y cualquier elección de $1 \leq l_1 < \dots < l_k \leq n$, $u_{l_1} \wedge \dots \wedge u_{l_k}$ es un autovector de $\log[\bigwedge^k \exp(a_i)]$ correspondiente al autovalor

$$\log \left[\prod_{j=1}^k \exp(\lambda(a_i)_{l_j}) \right] = \lambda(a_i)_{l_1} + \dots + \lambda(a_i)_{l_k},$$

donde $\lambda(a_i)_{l_j}$ es la l_j -ésima entrada del vector $\lambda(a_i)$. La igualdad $A_{(k)} = k \cdot \bar{A}(k)$ es una consecuencia inmediata de estas observaciones. Pero notemos que entonces, el resultado es consecuencia de las hipótesis $(A_{(k)}) \succ_d B_{(k)}$ para $k = 1, \dots, [\frac{n}{2}]$ y $k = n$ y el Corolario 7.1.15.

3. La prueba dada para la primera parte del teorema puede extenderse de forma evidente para probar este ítem. □

Ejemplo 7.2.6. Un sistema de proyecciones (o partición de la unidad) en $M_n(\mathbb{C})$ es una familia $\{P_i\}_{i=1, \dots, k}$ de proyecciones ortogonales tales que $\sum_{i=1}^k P_i = I$. Dada una tal familia, la compresión asociada, $\mathcal{C} : M_n(\mathbb{C}) \rightarrow M_n(\mathbb{C})$ está dada por

$$\mathcal{C}(A) = \sum_{i=1}^k P_i A P_i.$$

$\mathcal{C}$ es un ejemplo de transformación positiva unital que preserva la traza. En particular, si P_i es la proyección ortogonal sobre $\mathbb{C}e_i$, $i = 1, \dots, n$, entonces la compresión asociada a este sistema de proyecciones es llamada la compresión diagonal y notada $\mathcal{C}_0$.

En las páginas 331-332 de [54], A. W. Marshall y I. Olkin dan un ejemplo de mayorización multivariada que reformulamos en términos de mayorización conjunta fuerte (en este contexto, es

una consecuencia del ítem 3. del Teorema 7.2.5): Sea $(a_i)_{i=1,\dots,m}$ una familia abeliana en $M_n(\mathbb{C})$ y sea $\mathcal{C}_0$ la compresión diagonal. Entonces,

$$(a_i) \succ_f (\mathcal{C}_0(a_i)). \quad (7.3)$$

Es conveniente notar aquí que, dada una familia abeliana $(a_i)_{i=1,\dots,m}$, el resultado anterior no es válido para una compresión arbitraria $\mathcal{C}$ puesto que la familia $(\mathcal{C}(a_i))_{i=1,\dots,m}$ puede no ser abeliana. $\square$

Prueba de la Proposición 7.2.2. Supongamos que existen $k \in \mathbb{N}$, matrices unitaria $W_1, \dots, W_k \in M_n(\mathbb{C})$ y números no negativos $\mu_1, \dots, \mu_k$, $\sum_{j=1}^k \mu_j = 1$, tales que vale la ecuación (7.1). Definamos $T : C^*(a_1, \dots, a_m) \rightarrow C^*(b_1, \dots, b_m)$ por

$$T(a) = \Phi \left(\sum_{j=1}^k \mu_j W_j^* a W_j \right),$$

donde $\Phi : M_n(\mathbb{C}) \rightarrow C^*(b_1, \dots, b_m)$ es la transformación obtenida como en el Lema 7.2.4. Es evidente que T preserva la traza, es positiva y unital. Luego, por el Teorema 7.2.5 concluimos que $(a_i) \succ_f (b_i)$.

Por otra parte, si $(a_i) \succ_f (b_i)$ sean $U, V, \lambda(a_i), \lambda(b_i)$ ($1 \leq i \leq m$) como en la prueba del Teorema 7.2.5. Entonces, existe $E \in DS(n)$ tal que $E\lambda(a_i) = \lambda(b_i)$ para $1 \leq i \leq m$ (ver Observación 7.1.3). Por el teorema Birkhoff existen $k \in \mathbb{N}$, matrices de permutación $P_1, \dots, P_k \in DS(n)$ y números no negativos $\mu_1, \dots, \mu_k \in \mathbb{R}$, $\sum_{j=1}^k \mu_j = 1$ tales que $E = \sum_{j=1}^k \mu_j P_j$. Entonces, para $1 \leq i \leq m$ tenemos

$$\begin{aligned} b_i &= V^* D_{\lambda(b_i)} V = V^* D_{E\lambda(a_i)} V = V^* \left(\sum_{j=1}^k \mu_j P_j^t D_{\lambda(a_i)} P_j \right) V \\ &= \sum_{j=1}^k \mu_j (UP_j V)^* a_i (UP_j V) \end{aligned}$$

$\square$

7.2.3. Mayorizaciones conjuntas y funciones convexas

En esta sección consideramos caracterizaciones de las mayorizaciones conjuntas en términos del cálculo funcional descrito antes de la Proposición 7.2.2. Dada una familia arbitraria de matrices cuadradas $(a_i)_{i=1,\dots,m} \subseteq M_n(\mathbb{C})$, el (*primer*) *rango numérico conjunto* (ver [51]) está definido por

$$W(a_1, \dots, a_m) = \{(v^* a_1 v, \dots, v^* a_m v) : v \in \mathbb{C}^n, v^* v = 1\}.$$

Vamos a relacionar el rango numérico conjunto $W(a_1, \dots, a_m)$ con el espectro conjunto $\sigma(a_1, \dots, a_m)$ de una familia abeliana.

Lema 7.2.7. Sea $(a_i)_{i=1,\dots,m}$ una familia abeliana. Entonces,

$$W(a_1, \dots, a_m) = \text{conv}(\sigma(a_1, \dots, a_m)).$$

Demostración. $W(a_1, \dots, a_m)$ es invariante bajo conjugaciones unitarias de los a_i 's por una matriz unitaria $U \in M_n$ fija. Por lo tanto, podemos suponer que $a_i = D_{\lambda(a_i)}$, $i = 1, \dots, m$. Si $v^*v = 1$ tenemos

$$\begin{aligned} (v^*a_1v, \dots, v^*a_mv) &= \left(\sum_{j=1}^n |v_j|^2 \lambda_j(a_1), \dots, \sum_{j=1}^n |v_j|^2 \lambda_j(a_m) \right) \\ &= \sum_{j=1}^n |v_j|^2 (\lambda_j(a_1), \dots, \lambda_j(a_m)) \end{aligned}$$

donde $\sum_{j=1}^n |v_j|^2 = 1$. El lema es una consecuencia inmediata de estos hechos. $\square$

Proposición 7.2.8. Sean $(a_i)_{i=1, \dots, m}$ y $(b_i)_{i=1, \dots, m}$ dos familias abelianas. Entonces, los siguientes son equivalentes:

1. $(a_i) \succ_d (b_i)$.
2. $W(b_1, \dots, b_m) \subseteq W(a_1, \dots, a_m)$.
3. Para cada función convexa $f : V \rightarrow \mathbb{R}$ se tiene

$$\|f(a_1, \dots, a_m)\| \geq \|f(b_1, \dots, b_m)\|.$$

donde $V \subseteq \mathbb{R}^m$ es un conjunto convexo que contiene $\sigma(a_1, \dots, a_m) \cup \sigma(b_1, \dots, b_m)$.

Demostración. $1. \Leftrightarrow 2.$ se sigue del Lema 7.2.7 y del ítem 1. de la Proposición 7.2.2. Por otra parte, $1. \Leftrightarrow 3.$ es una consecuencia del Corolario 7.1.13. $\square$

La siguiente Proposición es una consecuencia del ítem 2. del Teorema 7.2.5 y la Proposición 7.2.8.

Proposición 7.2.9. Sean $(a_i)_{i=1, \dots, m}$ y $(b_i)_{i=1, \dots, m}$ dos familias abelianas. Entonces, $(a_i) \succ (b_i)$ si y solo si, para $k = 1, \dots, \lfloor \frac{n}{2} \rfloor$ y $k = n$ tenemos que

$$\left\| f \left(\log \left[\bigwedge^k \exp(a_1) \right], \dots, \log \left[\bigwedge^k \exp(a_m) \right] \right) \right\| \geq \left\| f \left(\log \left[\bigwedge^k \exp(b_1) \right], \dots, \log \left[\bigwedge^k \exp(b_m) \right] \right) \right\|$$

para cada función convexa $f : V \rightarrow \mathbb{R}$, donde $V \subseteq \mathbb{R}^m$ es un conjunto convexo que contiene a $\sigma(a_1, \dots, a_m) \cup \sigma(b_1, \dots, b_m)$.

La siguiente Proposición es una reformulación del Teorema 7.1.11 en este contexto.

Proposición 7.2.10. Sean $(a_i)_{i=1, \dots, m}$ y $(b_i)_{i=1, \dots, m}$ dos familias abelianas. Entonces, $(a_i) \succ_f (b_i)$ si y solo si, para cada función convexa $f : V \rightarrow \mathbb{R}$ se verifica que

$$\text{tr } f(a_1, \dots, a_m) \geq \text{tr } f(b_1, \dots, b_m),$$

donde $V \subseteq \mathbb{R}^m$ es un conjunto convexo que contiene $\sigma(a_1, \dots, a_m) \cup \sigma(b_1, \dots, b_m)$.

7.2.4. Relaciones de equivalencias asociadas a las mayorizaciones conjuntas

Las mayorizaciones conjuntas hasta aquí consideradas son preórdenes en el conjunto de familias abelianas en $M_n(\mathbb{C})$ de una cierta cantidad de miembros fija. El siguiente teorema describe las relaciones de equivalencias asociadas a estos preórdenes.

Teorema 7.2.11. Sean $(a_i)_{i=1,\dots,m}$ y $(b_i)_{i=1,\dots,m}$ dos familias abelianas en $M_n(\mathbb{C})$.

a) Los siguientes son equivalentes:

- i) $(a_i) \succ_d (b_i)$ y $(b_i) \succ_d (a_i)$.
- ii) $W(a_1, \dots, a_m) = W(b_1, \dots, b_m)$.

b) Los siguientes son equivalentes:

- i) Existe una matriz unitaria $W \in M_n$ tal que $W^* a_i W = b_i$ para cada $i = 1, \dots, m$.
- ii) $(a_i) \succ_f (b_i)$ y $(b_i) \succ_f (a_i)$.
- iii) $(a_i) \succ (b_i)$ y $(b_i) \succ (a_i)$.

Demostración. a). Es una consecuencia inmediata de la Proposición 7.2.8.

b). Notemos que el automorfismo interior $\alpha : C^*(a_1, \dots, a_m) \rightarrow C^*(b_1, \dots, b_m)$ inducido por W preserva la traza, es positivo y unital. Luego i) implica ii). Evidentemente ii) $\Rightarrow$ iii). Por otro lado, si $(a_i) \succ (b_i)$ y $(b_i) \succ (a_i)$, por el Teorema 7.1.25 existe una matriz de permutación $Q \in M_n$ tal que

$$V(Q^t(U^* a_i U)Q)V^* = b_i, \quad i = 1, \dots, m$$

donde $U, V \in M_n$ son como en la prueba del Teorema 7.2.5. Basta tomar entonces $W = UQV^*$. □

Capítulo 8

Mayorización conjunta en factores factores II_1

Para una descripción conceptual de los resultados incluidos dentro de este capítulo, así como la relación entre éstos y otros resultados ver la sección “Contexto general del trabajo” del Capítulo 1.

Organización del capítulo. El capítulo está organizado de la siguiente manera. En la sección 8.0.5 recordamos algunos hechos acerca de C^* -subálgebras abelianas de un factor II_1 . En la sección 8.1, después de describir algunos resultados técnicos, enunciamos y probamos nuestro resultado principal sobre la estructura local de las transformaciones doble estocásticas. En la sección 8.2 introducimos y desarrollamos la noción de mayorización conjunta entre familias abelianas de operadores autoadjuntos en un factor II_1 y obtenemos varias caracterizaciones de esta relación. Finalmente, en la sección 8.3 probamos los resultados adelantados en la sección 8.1.

Notaciones A lo largo de este capítulo, $\mathcal{M}$ denotará un factor II_1 con traza fiel normal y normalizada τ . Las C^* -subálgebras de $\mathcal{M}$ consideradas siempre serán unitales. El subespacio de operadores autoadjuntos de $\mathcal{M}$ es denotado por $\mathcal{M}_{sa}$, y consideramos familias abelianas $(a_1, \dots, a_n) = (a_i)_{i=1}^n$ en $\mathcal{M}_{sa}$, es decir familias finitas de operadores autoadjuntos que conmutan entre sí. Si $(a_i)_{i=1}^n \subseteq \mathcal{M}_{sa}$ es una familia abeliana entonces $C^*(a_1, \dots, a_n)$ denota la C^* -subálgebra (unital) abeliana y separable de $\mathcal{M}$ generada por los elementos de la familia. Si $\mathcal{A}$ es C^* -subálgebra abeliana arbitraria de $\mathcal{M}$ entonces $\Gamma(\mathcal{A})$ denota su espacio de caracteres, i.e. el conjunto de $*$ -homomorfismos $\gamma : \mathcal{A} \rightarrow \mathbb{C}$, dotados de la topología débil- $*$. $\Gamma(\mathcal{A})$ es un conjunto compacto y $\mathcal{A} \simeq C(\Gamma(\mathcal{A}))$, donde $C(\Gamma(\mathcal{A}))$ denota la C^* -álgebra de funciones continuas en $\Gamma(\mathcal{A})$ (ver sección 2.4.1 de los Preliminares).

8.0.5. Medidas espectrales conjuntas y distribuciones conjuntas

En lo que sigue haremos uso de los resultados descritos en las secciones 2.4.1, 2.4.2 y 2.4.3. Si $\mathcal{A} \subseteq \mathcal{M}$ es una C^* -subálgebra separable entonces $\Gamma(\mathcal{A})$ es metrizable y la representación $C(\Gamma(\mathcal{A})) \simeq \mathcal{A} \subseteq \mathcal{M}$ induce una medida espectral $E_{\mathcal{A}}$ que toma valores en el reticulado $\mathcal{P}(\mathcal{M})$ de proyecciones de $\mathcal{M}$. Sea $\mu_{\mathcal{A}}$ la medida regular de Borel en $\Gamma(\mathcal{A})$ definida por

$$\mu_{\mathcal{A}}(\Delta) = \tau(E_{\mathcal{A}}(\Delta)).$$

La regularidad de $\mu_{\mathcal{A}}$ se sigue del hecho de que cada conjunto abierto es σ -compacto [60, 2.18]. La función $\Lambda : L^\infty(\Gamma(\mathcal{A}), \mu_{\mathcal{A}}) \rightarrow \mathcal{M}$ dada por

$$\Lambda(h) = \int_{\Gamma(\mathcal{A})} h dE_{\mathcal{A}}$$

es un $*$ -monomorfismo normal (notemos que en este caso la topología débil- $*$ de $L^\infty(\Gamma(\mathcal{A}), \mu_{\mathcal{A}})$ restringida a la bola unitaria es metrizable) y tenemos que

$$\tau(\Lambda(h)) = \int_{\Gamma(\mathcal{A})} h d\mu_{\mathcal{A}}, \quad \forall h \in L^\infty(\Gamma(\mathcal{A}), \mu_{\mathcal{A}}). \quad (8.1)$$

Consideramos entonces el álgebra de von Neumann $L^\infty(\mathcal{A}) := \Lambda(L^\infty(\Gamma(\mathcal{A}), \mu_{\mathcal{A}})) \subseteq \mathcal{M}$.

Cuando $\mathcal{A} = C^*(a_1, \dots, a_n)$, $E_{\bar{a}} := E_{\mathcal{A}}$ y $\mu_{\bar{a}} := \mu_{\mathcal{A}}$ son la medida espectral conjunta y la distribución conjunta de la familia abeliana $\bar{a}$ y el isomorfismo normal definido arriba es notado $\Lambda_{\bar{a}} : L^\infty(\Gamma(\bar{a}), \mu_{\bar{a}}) \rightarrow L^\infty(\mathcal{A})$. Es inmediato verificar que $\Lambda_{\bar{a}}(\pi_i) = a_i$, $1 \leq i \leq n$, y escribimos $h(a_1, \dots, a_n) := \Lambda_{\bar{a}}(h)$.

En el caso particular en que $x \in \mathcal{M}$ es un operador normal, las partes real e imaginaria de x son operadores autoadjuntos que conmutan entre sí. Identificando el plano complejo con $\mathbb{R}^2$ en la forma usual, es sencillo ver que el espectro de x como operador normal coincide con el espectro conjunto del par abeliano $(\operatorname{Re}(x), \operatorname{Im}(x))$, y que la medida espectral de x coincide con la medida espectral conjunta del par $(\operatorname{Re}(x), \operatorname{Im}(x))$.

8.1. Forma local de las transformaciones doble estocásticas

Recordemos que una transformación lineal $\Phi : \mathcal{M} \rightarrow \mathcal{M}$ es doble estocástica [39] si es unital, positiva, y preserva la traza. En [39] se probó que en un factor finito cada transformación doble estocástica es normal. Denotamos al conjunto de todas las transformaciones doble estocásticas en $\mathcal{M}$ por $\operatorname{DS}(\mathcal{M})$. Estas transformaciones juegan un papel importante en la teoría de mayorización entre operadores autoadjuntos (ver la sección de Preliminares o [1, 2, 39, 40]); de esta forma, el estudio de su estructura es un desarrollo natural dentro de esta teoría.

En esta sección obtenemos una descripción conveniente de lo que denominamos la forma local de las transformaciones doble estocásticas que operan en un factor II_1 . Utilizamos esta estructura local para el estudio de la mayorización conjunta de familias abelianas en la sección 8.2.

Comenzamos con la presentación de cierta terminología y los enunciados del Teorema 8.1.1, la Proposición 8.1.2, y el Lema 8.1.4. Las demostraciones de estos resultados será desarrollada hacia el final del capítulo en la sección 8.3. Si bien son resultados técnicos son herramientas útiles para considerar problemas de aproximación.

Sea $\mathcal{A} \subseteq \mathcal{M}$ una C^* -subálgebra abeliana, y sean $E_{\mathcal{A}}$ y $\mu_{\mathcal{A}}$ la medida espectral y la medida de distribución de $\mathcal{A}$ definidas como en la sección 8.0.5. Si $x \in \Gamma(\mathcal{A})$ es tal que $E_{\mathcal{A}}(\{x\}) \neq 0$, decimos que x es un átomo para $E_{\mathcal{A}}$. El conjunto de átomos de $E_{\mathcal{A}}$ es denotado por $\operatorname{At}(E_{\mathcal{A}})$. Como $\mu_{\mathcal{A}} = \tau \circ E_{\mathcal{A}}$, la fidelidad de la traza implica que $\operatorname{At}(\mu_{\mathcal{A}}) = \operatorname{At}(E_{\mathcal{A}})$. Decimos que $\mathcal{A}$ es difusa si $\operatorname{At}(E_{\mathcal{A}}) = \emptyset$. El siguiente teorema establece que la medida espectral de una C^* -subálgebra separable $\mathcal{A}$ de un factor II_1 puede ser refinada de forma coherente.

Teorema 8.1.1. Sea $\mathcal{A} \subseteq \mathcal{M}$ una C^* -subálgebra abeliana y separable. Entonces existe $a \in \mathcal{M}_{sa}$ tal que $C^*(\mathcal{A}, a)$ es abeliana y difusa.

Como los átomos de $E_{\mathcal{A}}$ están en correspondencia con el conjunto de proyecciones minimales de $L^\infty(\mathcal{A})$, el teorema 8.1.1 provee una forma de sumergir $\mathcal{A}$ en una C^* -subálgebra separable $\tilde{\mathcal{A}} = C^*(\mathcal{A}, a)$ tal que $L^\infty(\tilde{\mathcal{A}})$ no tiene proyecciones minimales (ver la Observación 8.3.4 para más detalles).

Proposición 8.1.2. Sea $\mathcal{B} \subset \mathcal{M}$ una C^* -subálgebra abeliana, separable y difusa. Entonces existe un conjunto no acotado $\mathbb{M} \subseteq \mathbb{N}$ tal que para cada $m \in \mathbb{M}$ existen $k = k(m)$ particiones de la unidad $\{q_i^{t,m}\}_{i=1}^m \subseteq \mathcal{B}' \cap \mathcal{M}$, $1 \leq t \leq k$, con $\tau(q_i^{t,m}) = 1/m$ ($1 \leq i \leq m$, $1 \leq t \leq k$), y tales que para cada $b \in \mathcal{B}$, si notamos $\beta_i^{t,m} = m \tau(b q_i^{t,m})$, entonces

$$\lim_{m \rightarrow \infty} \left\| b - \frac{1}{k} \sum_{t=1}^k \left(\sum_{i=1}^m \beta_i^{t,m} q_i^{t,m} \right) \right\| = 0.$$

Observación 8.1.3. Para m fijo y $\{q_i^t\}_{i=1}^m$, $1 \leq t \leq k$, particiones de la unidad, es inmediato verificar que la transformación lineal

$$b \mapsto \frac{1}{k} \sum_{t=1}^k \left(\sum_{i=1}^m m \tau(b q_i^t) q_i^t \right)$$

es una contracción con respecto a la norma de operadores.

El siguiente lema está relacionado con el teorema de Dixmier sobre esperanzas condicionales a valores centrales en un álgebra de von Neumann finita. Notamos por $\mathcal{D}(\mathcal{M})$ el semigrupo convexo dado por

$$\mathcal{D}(\mathcal{M}) = \text{conv}\{Ad u : u \in \mathcal{U}(\mathcal{M})\}$$

donde $Ad u(a) = u^* a u$.

Lema 8.1.4. Sea $\mathcal{B} \subset \mathcal{M}$ una C^* -subálgebra separable y sea $\{p_i\}_{i=1}^m \subseteq \mathcal{B}' \cap \mathcal{M}$ una partición de la unidad. Entonces existe una sucesión $\{\rho_i\}_{i \in \mathbb{N}} \subset \mathcal{D}(\mathcal{M})$ tal que para cada $b \in \mathcal{B}$, si notamos $\beta_i(b) = \tau(b p_i)/\tau(p_i)$, entonces

$$\lim_{j \rightarrow \infty} \left\| \rho_j(b) - \sum_{i=1}^m \beta_i(b) p_i \right\| = 0.$$

En el siguiente lema establecemos un análogo del teorema de Birkhoff para operadores de espectro discreto en factores Π_1 .

Lema 8.1.5. Sean $\{p_i\}_{i=1}^m, \{q_i\}_{i=1}^m \subseteq \mathcal{M}$ particiones de la unidad tales que $\tau(p_i) = \tau(q_i) = \frac{1}{m}$, y sea $T \in DS(\mathcal{M})$. Entonces existe $\rho \in \mathcal{D}(\mathcal{M})$ tal que si $\beta_1, \dots, \beta_m \in \mathbb{R}$ y $\alpha_i = m \sum_{j=1}^m \beta_j \tau(T(q_j) p_i)$ para $1 \leq i \leq m$, tenemos que

$$\sum_{i=1}^m \alpha_i p_i = \rho \left(\sum_{i=1}^m \beta_i q_i \right). \quad (8.2)$$

Demostración. Sea $\gamma_{i,j} = m \tau(T(q_j) p_i) \geq 0$; entonces es inmediato verificar que $(\gamma_{i,j}) \in \mathbb{R}^{m \times m}$ es una matriz doble estocástica y, más aun, que $\alpha_i = \sum_{j=1}^m \gamma_{i,j} \beta_j$ para cada $i = 1, \dots, m$. Por el teorema de Birkhoff 2.1.6, la matriz doble estocástica $(\gamma_{i,j})$ puede escribirse como una combinación convexa de matrices de permutación, i.e. $(\gamma_{i,j}) = \sum_{\sigma \in \mathbb{S}_m} \eta_\sigma P_\sigma$, donde $\eta_\sigma \geq 0$, $\sum_{\sigma \in \mathbb{S}_m} \eta_\sigma = 1$ y P_σ es la matriz de permutación $m \times m$ inducida por $\sigma \in \mathbb{S}_m$. Entonces tenemos

$$\alpha_i = \sum_{j=1}^m \gamma_{i,j} \beta_j = \sum_{\sigma \in \mathbb{S}_m} \eta_\sigma \beta_{\sigma(i)} \quad 1 \leq i \leq m. \quad (8.3)$$

El hecho de que $\mathcal{M}$ es un factor II_1 y que los elementos de las particiones $\{p_i\}_i, \{q_i\}_i$ tienen las mismas trazas implica la existencia de operadores unitarios u_σ tales que $u_\sigma q_{\sigma(i)} (u_\sigma)^* = p_i$, $1 \leq i \leq m$, para cada $\sigma \in \mathbb{S}_m$. De hecho, si $\sigma \in \mathbb{S}_m$, de la igualdad $\tau(q_{\sigma(i)}) = \tau(p_i)$ deducimos que existen isometrías parciales $v_{i,\sigma} \in \mathcal{M}$ tales que $v_{i,\sigma} v_{i,\sigma}^* = p_i$ y $v_{i,\sigma}^* v_{i,\sigma} = q_{\sigma(i)}$ para $i = 1, \dots, m$. Entonces $u_\sigma = \sum_{i=1}^m v_{i,\sigma} \in \mathcal{M}$ son los unitarios requeridos. De la ecuación (8.3), y definiendo $\rho(\cdot) = \sum_{\sigma \in \mathbb{S}_m} \eta_\sigma u_\sigma(\cdot) u_\sigma^*$, tenemos

$$\begin{aligned} \sum_{i=1}^m \alpha_i p_i &= \sum_{i=1}^m \left(\sum_{\sigma \in \mathbb{S}_m} \eta_\sigma \beta_{\sigma(i)} \right) p_i = \sum_{\sigma \in \mathbb{S}_m} \eta_\sigma \left(\sum_{i=1}^m \beta_{\sigma(i)} u_\sigma q_{\sigma(i)} (u_\sigma)^* \right) \\ &= \sum_{\sigma \in \mathbb{S}_m} \eta_\sigma u_\sigma \left(\sum_{i=1}^m \beta_i q_i \right) (u_\sigma)^* = \rho \left(\sum_{i=1}^m \beta_i \right). \end{aligned}$$

□

Seguidamente, establecemos y probamos el resultado sobre la forma local de las transformaciones doble estocásticas. De forma resumida, éste describe la relación entre operadores normales y sus imágenes bajo la acción de $T \in DS(\mathcal{M})$, cuando estos operadores pertenecen a cierto sistema de operadores determinado por dos C^* -subálgebras (separables) abelianas de $\mathcal{M}$. Observamos que, si consideramos diferentes álgebras abelianas podemos obtener diferentes descripciones de T . Más aun, también es posible que ninguna de estas descripciones sea válida globalmente, i.e. para todos los operadores en $\mathcal{M}$ (ver Observaciones 8.2.7).

Teorema 8.1.6 (Forma local). Sean $\mathcal{A}, \mathcal{B} \subseteq \mathcal{M}$ C^* -subálgebras abelianas y separables y sea $T \in DS(\mathcal{M})$. Si $\mathcal{S}$ es el subsistema de operadores de $\mathcal{B}$ dado por $\mathcal{S} = T^{-1}(\mathcal{A}) \cap \mathcal{B}$ entonces, existe una sucesión $(\rho_r)_{r \in \mathbb{N}} \subseteq \mathcal{D}(\mathcal{M})$ tal que $\lim_{r \rightarrow \infty} \|T(b) - \rho_r(b)\| = 0$ para cada $b \in \mathcal{S}$.

Demostración. Notemos que basta probar el teorema para subálgebras abelianas, separables y difusas de $\mathcal{M}$; de hecho, supongamos que este resultado vale para tales álgebras y sean $\mathcal{A}, \mathcal{B} \subseteq \mathcal{M}$ C^* -subálgebras separables y abelianas arbitrarias. Entonces, por el Teorema 8.1.1 existen subálgebras abelianas, separables y difusas $\tilde{\mathcal{A}}$ y $\tilde{\mathcal{B}}$ de $\mathcal{M}$ tales que $\mathcal{A} \subseteq \tilde{\mathcal{A}}$ y $\mathcal{B} \subseteq \tilde{\mathcal{B}}$. De esta forma obtenemos una sucesión $(\rho_r)_{r \in \mathbb{N}}$ tal que $\lim_{r \rightarrow \infty} \|T(b) - \rho_r(b)\| = 0$ para cada $b \in T^{-1}(\mathcal{A}) \cap \mathcal{B} \subseteq T^{-1}(\tilde{\mathcal{A}}) \cap \tilde{\mathcal{B}}$. Así que suponemos además que $\mathcal{A}$ y $\mathcal{B}$ son difusas.

Por la Proposición 8.1.2 existe un conjunto no acotado $\mathbb{M}$ y, para cada $m \in \mathbb{M}$, $k = k(m)$ particiones de la unidad $\{q_i^{j,m}\}_{i=1}^m \subseteq \mathcal{B}' \cap \mathcal{M}$ y $\{p_i^{j,m}\}_{i=1}^m \subseteq \mathcal{A}' \cap \mathcal{M}$, $1 \leq j \leq k$ (para simplificar la notación no escribimos el super-índice m en las proyecciones) tales que, para cada $b \in T^{-1}(\mathcal{A}) \cap \mathcal{B}$

y cada $r \in \mathbb{N}$ existe $m_0 = m_0(r, b)$ tal que si $m \geq m_0$ entonces

$$\left\| b - \frac{1}{k} \sum_{j=1}^k \left(\sum_{i=1}^m \beta_i^j q_i^j \right) \right\| < \frac{1}{r} \quad (8.4)$$

y

$$\left\| T(b) - \frac{1}{k} \sum_{j=1}^k \left(\sum_{i=1}^m \alpha_i^j p_i^j \right) \right\| < \frac{1}{r} \quad (8.5)$$

donde $\beta_i^j = m \tau(b q_i^j)$, $\alpha_i^j = m \tau(T(b) p_i^j)$, $\tau(p_i^j) = \tau(q_i^j) = 1/m$, (de la construcción de tales particiones es evidente que podemos suponer que ambas tienen el mismo conjunto $\mathbb{M}$ y para cada $m \in \mathbb{M}$ tienen el mismo $k(m)$). Como $\|T\| = 1$, entonces por la ecuación (8.4) tenemos que

$$\left\| T(b) - \frac{1}{k} \sum_{j=1}^k \left(\sum_{i=1}^m \beta_i^j T(q_i^j) \right) \right\| \leq \frac{1}{r}. \quad (8.6)$$

Aplicando la ecuación (8.6) y el hecho de que la transformación lineal en la Observación 8.1.3 es contractiva (con $\{p_i^j\}$ como las particiones de la unidad), obtenemos

$$\left\| \frac{1}{k} \sum_{j=1}^k \sum_{i=1}^m \alpha_i^j p_i^j - \frac{1}{k^2} \sum_{j=1}^k \left(\sum_{t=1}^k \sum_{i=1}^m \alpha_i^{j,t} p_i^t \right) \right\| \leq \frac{1}{r}, \quad (8.7)$$

donde $\alpha_i^{j,t} = m \sum_{l=1}^m \beta_l^j \tau(T(q_l^j) p_i^t)$. Por el Lema 8.1.5 existen $\rho_{j,t}^m \in \mathcal{D}(\mathcal{M})$ tales que

$$\sum_{i=1}^m \alpha_i^{j,t} p_i^t = \rho_{j,t}^m \left(\sum_{l=1}^m \beta_l^j q_l^j \right), \quad 1 \leq j, t \leq k. \quad (8.8)$$

De (8.5), (8.7), y (8.8) deducimos que

$$\left\| T(b) - \frac{1}{k^2} \sum_{j=1}^k \sum_{t=1}^k \rho_{j,t}^m \left(\sum_{l=1}^m \beta_l^j q_l^j \right) \right\| \leq \frac{2}{r}, \quad (8.9)$$

Por el Lema 8.1.4 existen sucesiones $(\tilde{\rho}_n^j)_{n \in \mathbb{N}} \subseteq \mathcal{D}(\mathcal{M})$, $1 \leq j \leq k$, (independientes de b) tales que para cada $r \in \mathbb{N}$ existe $n_0 = n_0(r, b)$ tal que, si $n \geq n_0$ entonces

$$\left\| \sum_{l=1}^m \beta_l^j q_l^j - \tilde{\rho}_n^j(b) \right\| \leq \frac{1}{r}, \quad 1 \leq j \leq k. \quad (8.10)$$

De (8.9) y (8.10), junto con el hecho de que cada $\rho \in \mathcal{D}(\mathcal{M})$ es contractiva obtenemos, para cada $n \geq n_0(r, b)$

$$\left\| T(b) - \frac{1}{k^2} \sum_{j=1}^k \sum_{t=1}^k \rho_{j,t}^m(\tilde{\rho}_n^j(b)) \right\| \leq \frac{3}{r}, \quad (8.11)$$

Consideremos un subconjunto denso numerable $\{b_1, b_2, \dots\}$ de $\mathcal{B}$. Notemos que de las estimaciones anteriores, la desigualdad (8.11) sigue valiendo para b_j siempre que $n \geq n_0(r, b_j)$, $m \geq m_0(r, b_j)$. Definamos

$$n(r) = \max\{n_0(r, b_1), \dots, n_0(r, b_r)\}, \quad m(r) = \max\{m_0(r, b_1), \dots, m_0(r, b_r)\}$$

y luego consideremos

$$\rho_r(\cdot) := \frac{1}{k^2} \sum_{j=1}^k \sum_{t=1}^k \rho_{j,t}^{m(r)} \circ \tilde{\rho}_{n(r)}^j(\cdot) \in \mathcal{D}(\mathcal{M}),$$

donde $k = k(m(r))$. Por las estimaciones anteriores tenemos que $\|T(b_j) - \rho_r(b_j)\| \leq \frac{3}{r}$ siempre que $1 \leq j \leq r$. Sea $b \in \mathcal{B}$, $\epsilon > 0$ y $l \in \mathbb{N}$ tal que $\|b - b_l\| < \epsilon/3$. Si $r > \max\{l, 9/\epsilon\}$, entonces $\|T(b_l) - \rho_{r,n_r}(b_l)\| < \epsilon/3$, y

$$\|T(b) - \rho_r(b)\| \leq \|T(b - b_l)\| + \|T(b_l) - \rho_r(b_l)\| + \|\rho_r(b_l - b)\| < \epsilon.$$

□

Corolario 8.1.7. Sea $T \in DS(\mathcal{M})$ y sean $(a_i)_{i=1}^n, (b_i)_{i=1}^n \subseteq \mathcal{M}_{sa}$ familias abelianas tales que $T(b_i) = a_i$ para $1 \leq i \leq n$. Entonces existe una sucesión $(\rho_r)_{r \in \mathbb{N}} \subseteq \mathcal{D}$ tal que para $1 \leq i \leq n$ $\lim_{r \rightarrow \infty} \|a_i - \rho_r(b_i)\| = 0$.

Demostración. Consideremos $\mathcal{A} = C^*(a_1, \dots, a_n)$ y $\mathcal{B} = C^*(b_1, \dots, b_n)$, que son C^* -subálgebras abelianas y separables de $\mathcal{M}$. Aplicando el Teorema 8.1.6 con respecto a estas álgebras obtenemos una sucesión $(\rho_r)_{r \in \mathbb{N}} \subseteq \mathcal{D}$ tal que $\lim_{r \rightarrow \infty} \|T(b) - \rho_r(b)\| = 0$ para cada $b \in T^{-1}(\mathcal{A}) \cap \mathcal{B}$. Pero notemos que, por nuestra elección, $b_i \in T^{-1}(\mathcal{A}) \cap \mathcal{B}$ y entonces,

$$\|T(b_i) - \rho_r(b_i)\| = \|a_i - \rho_r(b_i)\| \xrightarrow{r} 0.$$

□

8.2. Núcleos doble estocásticos

La noción de núcleo doble estocástico que desarrollamos en esta sección surgió de forma natural, a partir de una familia de ejemplos (ver Ejemplo 8.2.2) que asemeja a las matrices doble estocásticas en este contexto.

Definición 8.2.1. Sea $(X, \mu_X), (Y, \mu_Y)$ dos espacios de probabilidad. Una transformación lineal positiva y unital $\nu : L^\infty(Y, \mu_Y) \rightarrow L^\infty(X, \mu_X)$ es un núcleo doble estocástico si $\int_X \nu(1_\Delta) d\mu_X = \mu_Y(\Delta)$, para cada conjunto μ_Y -medible $\Delta \subseteq Y$.

Ejemplo 8.2.2. Sean X y Y espacios topológicos compactos y sean μ_X y μ_Y medidas de probabilidad, regulares y de Borel en X y Y respectivamente. Consideremos $D \in L^1(\mu_X \times \mu_Y)$ y sea $\nu(f)(x) = \int_Y D(x, y) f(y) d\mu_Y(y)$. Entonces $\nu : L^\infty(X, \mu_X) \rightarrow L^\infty(Y, \mu_Y)$ es un núcleo doble estocástico si y solo si

1. $D(x, y) \geq 0$ $(\mu_X \times \mu_Y)$ -a.e.
2. $\int_X D(x, y) d\mu_X(x) = 1$ μ_Y -c.t.p.

$$3. \int_Y D(x, y) d\mu_Y(y) = 1 \text{ } \mu_X\text{-c.t.p.}$$

En particular, si $\mu_X = \mu_Y$ es una medida con soporte finito $\{x_i\}_{i=1}^m$ y tal que $\mu_X(\{x_i\}) = \frac{1}{m}$ para $1 \leq i \leq m$ entonces D es núcleo doble estocástico si y solo si la matriz $m \times m$ $(D(x_i, x_j))_{i,j}$ es doble estocástica.

La siguiente proposición establece la relación entre (la localización de) una transformación doble estocástica y los núcleos doble estocásticos.

Proposición 8.2.3. Sean $\bar{a} = (a_i)_{i=1}^n$, $\bar{b} = (b_i)_{i=1}^n \subseteq \mathcal{M}_{sa}$ familias abelianas. Los siguiente enunciados son equivalentes:

1. Existe $T \in DS(\mathcal{M})$ tal que $T(b_i) = a_i$, $1 \leq i \leq n$.
2. Existe $\nu : L^\infty(\sigma(\bar{b}), \mu_{\bar{b}}) \rightarrow L^\infty(\sigma(\bar{a}), \mu_{\bar{a}})$ núcleo doble estocástico tal que $\nu(\pi_i) = \pi_i$, $1 \leq i \leq n$.

Demostración. Supongamos que $T(b_i) = a_i$, $1 \leq i \leq n$, con $T \in DS(\mathcal{M})$. Sean $\mathcal{A} = C^*(a_1, \dots, a_n)$, $\mathcal{B} = C^*(b_1, \dots, b_n)$. Como $\mathcal{M}$ es un álgebra de von Neumann finita, existe una esperanza condicional $\mathcal{E}_{\mathcal{A}} : \mathcal{M} \rightarrow L^\infty(\mathcal{A})$ que conmuta con τ . Entonces $\nu = \Lambda_{\bar{a}}^{-1} \circ \mathcal{E}_{\mathcal{A}} \circ T \circ \Lambda_{\bar{b}}$ es el núcleo doble estocástico que buscamos, pues

$$\nu(\pi_i) = \Lambda_{\bar{a}}^{-1}(\mathcal{E}_{\mathcal{A}}(T(\Lambda_{\bar{b}}(\pi_i)))) = \Lambda_{\bar{a}}^{-1}(\mathcal{E}_{\mathcal{A}}(T(b_i))) = \Lambda_{\bar{a}}^{-1}(\mathcal{E}_{\mathcal{A}}(a_i)) = \Lambda_{\bar{a}}^{-1}(a_i) = \pi_i,$$

y para cualquier subconjunto Boreliano $\Delta \subseteq \sigma(\bar{b})$

$$\int_{\sigma(\bar{a})} \nu(1_\Delta) d\mu_{\bar{a}} = \tau(\Lambda_{\bar{a}}(\nu(1_\Delta))) = \tau(\mathcal{E}_{\mathcal{A}} \circ T \circ \Lambda_{\bar{b}}(1_\Delta)) = \tau(\Lambda_{\bar{b}}(1_\Delta)) = \mu_{\bar{b}}(\Delta).$$

Recíprocamente, supongamos que existe ν como en 2. Sea $\mathcal{E}_{\mathcal{B}} : \mathcal{M} \rightarrow L^\infty(\mathcal{B})$ la esperanza condicional sobre $L^\infty(\mathcal{B})$ que conmuta con τ . Entonces definamos $T = \Lambda_{\bar{a}} \circ \nu \circ \Lambda_{\bar{b}}^{-1} \circ \mathcal{E}_{\mathcal{B}}$. Evidentemente T es positiva, unital y conmuta con τ . Además,

$$T(b_i) = \Lambda_{\bar{a}} \circ \nu \circ \Lambda_{\bar{b}}^{-1} \circ \mathcal{E}_{\mathcal{B}}(b_i) = \Lambda_{\bar{a}} \circ \nu \circ \Lambda_{\bar{b}}^{-1}(b_i) = \Lambda_{\bar{a}}(\nu(\pi_i)) = \Lambda_{\bar{a}}(\pi_i) = a_i, \quad 1 \leq i \leq n.$$

□

En lo que sigue, introducimos y desarrollamos la noción de mayorización conjunta entre familias abelianas en un factor Π_1 .

Definición 8.2.4. Sean $\bar{a} = (a_i)_{i=1}^n$, $\bar{b} = (b_i)_{i=1}^n$ dos familias abelianas en $\mathcal{M}_{sa}$. Decimos que $\bar{a}$ está conjuntamente mayorizado por $\bar{b}$, notado $\bar{a} \prec \bar{b}$, si existe un núcleo doble estocástico $\nu : L^\infty(\sigma(\bar{b}), \mu_{\bar{b}}) \rightarrow L^\infty(\sigma(\bar{a}), \mu_{\bar{a}})$ tal que $\nu(\pi_i) = \pi_i$, para cada $1 \leq i \leq n$.

La siguiente terminología será utilizada en el enunciado del teorema de caracterización de la mayorización conjunta 8.2.5: si $(x_1, \dots, x_n)$ es una familia finita en $\mathcal{M}$, sea $\mathcal{U}_{\mathcal{M}}(x_1, \dots, x_n)$ la órbita unitaria conjunta de la familia con respecto al grupo de operadores unitarios $\mathcal{U}_{\mathcal{M}}$ de $\mathcal{M}$, i.e.

$$\mathcal{U}_{\mathcal{M}}(x_1, \dots, x_n) = \{(u^* x_1 u, \dots, u^* x_n u) : u \in \mathcal{U}_{\mathcal{M}}\}.$$

Vamos a considerar también la cápsula convexa de la órbita unitaria de la familia $(x_i)_{i=1}^n$,

$$\text{conv}(\mathcal{U}_{\mathcal{M}}(x_i)_{i=1}^n) = \{(\rho(x_i))_{i=1}^n, \rho \in \mathcal{D}\}.$$

Denotamos por $\overline{\text{conv}}(\mathcal{U}_{\mathcal{M}}(x_i)_{i=1}^n)$, $\overline{\text{conv}}^w(\mathcal{U}_{\mathcal{M}}(x_i)_{i=1}^n)$ y $\overline{\text{conv}}^1(\mathcal{U}_{\mathcal{M}}(x_i)_{i=1}^n)$ las respectivas clausuras entrada a entrada en la topología de la norma, en la topología débil de operadores, y en la topología L^1 inducida por τ .

Teorema 8.2.5. Sean $\bar{a} = (a_i)_{i=1}^n$, $\bar{b} = (b_i)_{i=1}^n$ dos familias abelianas en $\mathcal{M}_{sa}$. Entonces los siguientes son equivalentes:

1. $\bar{a}$ está conjuntamente mayorizado por $\bar{b}$.
2. $\bar{a} \in \overline{\text{conv}}(\mathcal{U}_{\mathcal{M}}(\bar{b}))$.
3. $\bar{a} \in \overline{\text{conv}}^1(\mathcal{U}_{\mathcal{M}}(\bar{b}))$.
4. $\bar{a} \in \overline{\text{conv}}^w(\mathcal{U}_{\mathcal{M}}(\bar{b}))$.
5. $\mu_{\bar{a}} \prec \mu_{\bar{b}}$.
6. Existe una transformación completamente positiva $T \in DS(\mathcal{M})$ tal que $a_i = T(b_i)$, $1 \leq i \leq n$.
7. Existe $T \in DS(\mathcal{M})$ tal que $a_i = T(b_i)$, $1 \leq i \leq n$.
8. $\tau(f(a_1, \dots, a_n)) \leq \tau(f(b_1, \dots, b_n))$ para cada función continua y convexa $f : \mathbb{R}^n \rightarrow \mathbb{R}$.

Dada una subálgebra von Neumann abeliana $\mathcal{N} \subseteq \mathcal{M}$ sea $\mathcal{E}_{\mathcal{N}}$ la (única) esperanza condicional sobre $\mathcal{N}$ que preserva la traza. Entonces, $\mathcal{E}_{\mathcal{N}}$ es una transformación doble estocástica (completamente positiva) de $\mathcal{M}$.

Corolario 8.2.6. Sea $\mathcal{N} \subseteq \mathcal{M}$ una subálgebra abeliana von Neumann y sea $(b_i)_{i=1}^n \subseteq \mathcal{M}_{sa}$ una familia abeliana. Entonces $(\mathcal{E}_{\mathcal{N}}(b_i))_{i=1}^n \prec (b_i)_{i=1}^n$.

Observaciones 8.2.7.

1. Sea $x \in \mathcal{M}$ un operador normal. Recordemos que hay una forma natural para identificar la medida espectral usual de x con la medida espectral conjunta del par abeliano $(\text{Re}(x), \text{Im}(x))$ (ver el último párrafo de la sección 8.0.5). Si $T \in DS(\mathcal{M})$ entonces $T(x) = y$ si y solo si $T(\text{Re}(x)) = \text{Re}(y)$ y $T(\text{Im}(x)) = \text{Im}(y)$, ya que T es positivo.

De estos hechos y del Teorema 8.2.5, vemos que si $x, y \in \mathcal{M}$ son operadores normales entonces $y \prec x$ en el sentido de [40] si y solo si $(\text{Re}(y), \text{Im}(y)) \prec (\text{Re}(x), \text{Im}(x))$ en el sentido de la definición 8.2.4. Es decir, la mayorización conjunta (bajo esta identificación) extiende la mayorización entre operadores normales en el caso de factores II_1 .

2. Si $p \in \mathcal{M}$ es una proyección con traza $1/2$, existe una isometría parcial $v \in \mathcal{M}$ con $v^*v = p$, $vv^* = 1 - p$. Sea $a_1 = 1$, $a_2 = v^* - v$, $b_1 = 1$, $b_2 = v - v^*$. Definamos la transformación lineal $T : C^*(p, v) \rightarrow C^*(p, v)$ por $T(p) = p$, $T(1) = 1$, y $T(v) = v^*$, $T(v^*) = v$ (el álgebra $C^*(p, v)$ puede pensarse como $M_2(\mathbb{C})$ y T como la transposición). Entonces T no es una transformación completamente positiva, pero puede ser extendida a $T \in DS(\mathcal{M})$. La equivalencia entre 6 y 7 en el Teorema 8.1.6 muestra que existe una transformación completamente positiva

$T' \in DS(\mathcal{M})$ con $T'(b_1) = a_1$, $T'(b_2) = a_2$. De hecho se puede describir una transformación T' con estas características, por ejemplo $T' = u \cdot u^*$, con $u = 2p - 1$. De esta forma vemos que la forma “local” de la transformación puede diferir substancialmente de la forma “global”.

En el resto de la sección probamos las implicaciones necesarias para demostrar el Teorema 8.2.5. El siguiente lema generaliza un resultado de [40].

Lema 8.2.8. Sean $\bar{a} = (a_i)_{i=1}^n$, $\bar{b} = (b_i)_{i=1}^n \subseteq \mathcal{M}_{sa}$ familias abelianas. Si $\bar{a} \in \overline{\text{conv}}^w(\mathcal{U}_{\mathcal{M}}(\bar{b}))$ entonces existe una transformación completamente positiva $T \in DS(\mathcal{M})$ tal que $a_i = T(b_i)$, $1 \leq i \leq n$.

Demostración. Sea

$$\{(b_1^j, \dots, b_n^j)\}_{j \in J} \subseteq \text{conv}(\mathcal{U}_{\mathcal{M}}(b_1, \dots, b_n)), \quad b_i^j \xrightarrow[j]{\text{WOT}} a_i,$$

para $1 \leq i \leq n$. Entonces existe una sucesión $(\rho_j)_{j \in J} \subseteq \mathcal{D}$ tal que

$$(b_1^j, \dots, b_n^j) = (\rho_j(b_1), \dots, \rho_j(b_n)),$$

para cada $j \in J$. Notemos que ρ_j es una transformación completamente positiva y doble estocástica y la red $\{\rho_j\}_{j \in J}$ está acotada en norma. Entonces existe una subred (que, haciendo abuso de notación, denotamos $\{\rho_j\}_{j \in J}$) que converge en la topología BW de Arveson [10], i.e. existe una transformación completamente positiva $T : \mathcal{M} \rightarrow \mathcal{M}$ tal que $\rho_j(x) \xrightarrow[j]{\text{WOT}} T(x)$ si $x \in \mathcal{M}$. Por normalidad de la traza, T preserva la traza, es positiva y unital. Como

$$\rho_j(b_i) = b_i^j \xrightarrow[j]{\text{WOT}} a_i \Rightarrow T(b_i) = a_i, \quad 1 \leq i \leq n.$$

□

Lema 8.2.9. Sean $\bar{a} = (a_i)_{i=1}^n$, $\bar{b} = (b_i)_{i=1}^n \subseteq \mathcal{M}_{sa}$ familias abelianas. Si $\bar{a} \prec \bar{b}$, entonces $\mu_{\bar{a}} \prec \mu_{\bar{b}}$.

Demostración. Por hipótesis $\bar{a} \prec \bar{b}$ es decir, existe un núcleo doble estocástico

$$\nu : L^\infty(\sigma(\bar{b}), \mu_{\bar{b}}) \rightarrow L^\infty(\sigma(\bar{a}), \mu_{\bar{a}})$$

tal que $\nu(\pi_i) = \pi_i$, $1 \leq i \leq n$. Sea $\nu_1, \dots, \nu_m \in M_+^\infty(\mathbb{R}^n)$ tales que $\sum_{j=1}^m \nu_j = \mu_{\bar{a}}$. Definamos medidas ν'_j por $\nu'_j(\Delta) = \nu_j(\nu(1_\Delta))$. Por continuidad de ν , $\nu'_j(f) = \nu_j(\nu(f))$ para cada $f \in L^\infty(\sigma(\bar{b}), \mu_{\bar{b}})$. De esta forma,

$$\nu'_j(\pi_i) = \nu_j(\nu(\pi_i)) = \nu_j(\pi_i), \quad 1 \leq i \leq n, \quad 1 \leq j \leq m$$

y análogamente $\nu_j(1) = \nu'_j(1)$, con lo que $\nu_j \sim \nu'_j$, para $1 \leq j \leq m$. Finalmente,

$$\sum_{j=1}^m \nu'_j(\Delta) = \sum_{j=1}^m \nu_j(\nu(1_\Delta)) = \mu_{\bar{a}}(\nu(1_\Delta)) = \mu_{\bar{b}}(\Delta).$$

Entonces $\sum_{j=1}^m \nu'_j = \mu_{\bar{b}}$, lo que muestra que $\mu_{\bar{a}} \prec \mu_{\bar{b}}$.

□

Lema 8.2.10. Sean $\bar{a} = (a_i)_{i=1}^n$, $\bar{b} = (b_i)_{i=1}^n \subset \mathcal{M}_{sa}$ familias abelianas. Si $\mu_{\bar{a}} \prec \mu_{\bar{b}}$, entonces existe $T \in DS(\mathcal{M})$ tal que $T(b_i) = a_i$ para $1 \leq i \leq n$.

Demostración. Por compacidad, podemos considerar particiones $\{\Delta_j^r\}_{j=1}^{m(r)}$ de $\sigma(\bar{a})$ con $\text{diam}(\Delta_j^r) < 1/r$ para cada $1 \leq j \leq m$. Fijamos puntos $x_1^r, \dots, x_{m(r)}^r$ con $x_j^r \in \Delta_j^r$ y definimos medidas μ_j^r por $\mu_j^r(\cdot) = \mu_{\bar{a}}(\cdot \cap \Delta_j^r)$. Entonces evidentemente $\sum_j \mu_j^r = \mu_{\bar{a}}$. Como $\mu_{\bar{a}} \prec \mu_{\bar{b}}$ por hipótesis, existen medidas ν_j^r con $\nu_j^r \sim \mu_j^r$ y $\sum_j \nu_j^r = \mu_{\bar{b}}$. Sean g_j^r las derivadas de Radon-Nikodym $g_j^r = d\nu_j^r/d\mu_{\bar{b}}$. Notemos que $\sum_j g_j^r = 1$ ($\mu_{\bar{b}}$ - a.e.). Definimos funciones $D_r : \sigma(\bar{a}) \times \sigma(\bar{b}) \rightarrow \mathbb{R}$ dadas por

$$D_r(s, t) = \sum_{j=1}^{m(r)} \frac{g_j^r(t)}{\mu_{\bar{a}}(\Delta_j^r)} 1_{\Delta_j^r}(s).$$

Definamos $\nu_r : L^\infty(\sigma(\bar{b}), \mu_{\bar{b}}) \rightarrow L^\infty(\sigma(\bar{a}), \mu_{\bar{a}})$ por

$$\nu_r(b)(s) = \int_{\sigma(\bar{b})} b(t) D_r(s, t) d\mu_{\bar{b}}(t).$$

Las transformaciones ν_r son núcleos doble estocásticos, lo cual puede verse de las equivalencias $\mu_j^r \sim \nu_j^r$. Por la Proposición 8.2.3 tenemos la sucesión de transformaciones doble estocásticas asociadas estos núcleos $\{T_r\}_r \subset DS(\mathcal{M})$ tal que

$$T_r(b_i) = \int_{\sigma(\bar{a})} \nu_r(\pi_i) dE_{\bar{a}} \in L^\infty(\mathcal{A}), \quad 1 \leq i \leq n.$$

La red acotada $\{T_r\}_{r \in \mathbb{N}}$ tiene una subred $\{T_k\}_{k \in K}$ que converge hacia un punto de acumulación $T \in DS(\mathcal{M})$ en la topología BW. Como esta red está acotada, $T(b_i) = \lim_{k \in K} T_k(b_i) \in L^\infty(\mathcal{A})$ en la topología débil de operadores. Afirmamos que $T(b_i) = a_i$, $1 \leq i \leq n$. Para ver esto, como la red $\{T_k(b_i)\}_{k \in K}$ está acotada, basta probar que

$$\lim_k \tau(x T_k(b_i)) = \tau(x a_i), \quad 1 \leq i \leq n, \quad \forall x \in \mathcal{A}.$$

Equivalentemente, basta mostrar que para cada función continua $f \in C(\sigma(\bar{a}))$ y cada $i = 1, \dots, n$,

$$\lim_k \int_{\sigma(\bar{a})} f(s) \left(\int_{\sigma(\bar{b})} D_k(s, t) \pi_i(t) d\mu_{\bar{b}}(t) \right) d\mu_{\bar{a}}(s) = \int_{\sigma(\bar{a})} f(s) \pi_i(s) d\mu_{\bar{a}}(s).$$

Esto puede verse por argumento de aproximación estándar, usando la continuidad uniforme de f , el hecho de que los diámetros de Δ_j^r tienden a 0 cuando r crece, y la equivalencia $\mu_j^r \sim \nu_j^r$. $\square$

El siguiente lema es una consecuencia directa del Teorema 2.4.24 y la identidad en la ecuación (8.1).

Lema 8.2.11. Sean $\bar{a} = (a_i)_{i=1}^n$, $\bar{b} = (b_i)_{i=1}^n \subset \mathcal{M}_{sa}$ dos familias abelianas. Entonces $\mu_{\bar{a}} \prec \mu_{\bar{b}}$ si y solo si $\tau(f(a_1, \dots, a_n)) \leq \tau(f(b_1, \dots, b_n))$ para cada función continua convexa $f : \mathbb{R}^n \rightarrow \mathbb{R}$.

Demostración del Teorema 8.2.5. La Proposición 8.2.3 muestra la equivalencia (7) $\Leftrightarrow$ (1) y el Corolario 8.1.7 es (7) $\Rightarrow$ (2). La implicación (2) $\Rightarrow$ (3) $\Rightarrow$ (4) es trivial. El Lema 8.2.8 muestra que (4) $\Rightarrow$ (6), y es evidente que (6) $\Rightarrow$ (7). Los Lemmas 8.2.9 y 8.2.10 prueban la equivalencia (5) $\Leftrightarrow$ (1). De esta forma tenemos que (1)-(7) son equivalentes. Finalmente, el Lema 8.2.11 establece que (5) $\Leftrightarrow$ (8). $\square$

Dadas familias $\bar{a} = (a_i)_{i=1}^n$, $\bar{b} = (b_i)_{i=1}^n \subseteq \mathcal{M}$, decimos que $\bar{a}$ y $\bar{b}$ son conjuntamente aproximadamente unitariamente equivalentes en $\mathcal{M}$ si $\bar{a} \in \overline{\mathcal{U}_{\mathcal{M}}(\bar{b})}$ es decir, si existe una sucesión de operadores unitarios $(u_n)_{n \in \mathbb{N}} \subseteq \mathcal{M}$ tal que

$$\lim_{n \rightarrow \infty} \|u_n b_i u_n^* - a_i\| = 0, \quad 1 \leq i \leq n.$$

Es evidente que esta es una relación de equivalencia. Más aún, si $\bar{a}$ y $\bar{b}$ son conjuntamente aproximadamente unitariamente equivalentes en $\mathcal{M}$ entonces $\bar{a}$ es una familia abeliana si y solo si $\bar{b}$ lo es. En [11] se dio una caracterización de esta relación de equivalencia entre operadores autoadjuntos en términos de las distribuciones espectrales de los operadores. El siguiente resultado muestra una lista de caracterizaciones de esta relación para el caso de familias abelianas en factores Π_1 .

Teorema 8.2.12. Sean $\bar{a} = (a_i)_{i=1}^n$ y $\bar{b} = (b_i)_{i=1}^n \subset \mathcal{M}_{sa}$ familias abelianas. Entonces los siguientes son equivalentes:

1. $\bar{a}$ y $\bar{b}$ son conjuntamente aproximadamente unitariamente equivalentes en $\mathcal{M}$.
2. $\bar{a} \prec \bar{b}$ y $\bar{b} \prec \bar{a}$
3. $\mu_{\bar{a}} = \mu_{\bar{b}}$
4. $\tau(f(a_1, \dots, a_n)) = \tau(f(b_1, \dots, b_n))$ para cada función continua y convexa $f : \mathbb{R}^n \rightarrow \mathbb{R}$.
5. $\tau(f(a_1, \dots, a_n)) = \tau(f(b_1, \dots, b_n))$ para cada función continua $f : \mathbb{R}^n \rightarrow \mathbb{R}$.

Demostración. Por el Teorema 8.2.5 tenemos (1) $\Rightarrow$ (2) y (2) $\Leftrightarrow$ (4). Más aun, (4) es equivalente a $\mu_{\bar{a}}(f) = \mu_{\bar{b}}(f)$ para cada función continua y convexa f . Entonces $\mu_{\bar{a}}(f) = \mu_{\bar{b}}(f)$ para cada función continua f [4, Proposición I.1.1], y esto implica que $\mu_{\bar{a}} = \mu_{\bar{b}}$. De esta forma, (4) $\Rightarrow$ (5) $\Rightarrow$ (3). Nuevamente, por el Teorema 8.2.5 (3) $\Rightarrow$ (2) y entonces (2)-(5) son equivalentes. Finalmente, probamos que (3) $\Rightarrow$ (1). Supongamos que $\mu_{\bar{a}} = \mu_{\bar{b}}$ y notemos que entonces, $\sigma(\bar{a}) = \text{sop } \mu_{\bar{a}} = \text{sop } \mu_{\bar{b}} = \sigma(\bar{b})$ y para cada conjunto de Borel Δ en $\sigma(\bar{a})$ tenemos

$$\tau(E_{\bar{a}}(\Delta)) = \mu_{\bar{a}}(1_{\Delta}) = \mu_{\bar{b}}(1_{\Delta}) = \tau(E_{\bar{b}}(\Delta)). \quad (8.12)$$

sea $\epsilon > 0$. Por compacidad, elegimos $B_1, \dots, B_m$ un cubrimiento finito disjunto de $\sigma(\bar{a}) = \sigma(\bar{b})$, y puntos $x_j \in B_j$ con la propiedad de que $|\pi_i(\lambda) - \pi_i(x_j)| < \epsilon/2$ para cada $\lambda \in B_j$, $1 \leq i \leq n$, $1 \leq j \leq m$. Entonces obtenemos, usando el teorema espectral,

$$\left\| a_i - \sum_{j=1}^m \pi_i(x_j) E_{\bar{a}}(B_j) \right\| < \frac{\epsilon}{2}, \quad \left\| b_i - \sum_{j=1}^m \pi_i(x_j) E_{\bar{b}}(B_j) \right\| < \frac{\epsilon}{2}.$$

para $1 \leq i \leq n$. De la ecuación (8.12) vemos que $\tau(E_{\bar{a}}(B_j)) = \tau(E_{\bar{b}}(B_j))$ para cada $1 \leq j \leq m$. Como en la prueba del Lema 8.1.5, obtenemos operadores unitarios $w_\epsilon \in \mathcal{U}(\mathcal{M})$ tales que $w_\epsilon^* E_{\bar{b}}(B_j) w_\epsilon = E_{\bar{a}}(B_j)$ para cada j . Entonces

$$w_\epsilon^* \left(\sum_{j=1}^m \pi_i(x_j) E_{\bar{b}}(B_j) \right) w_\epsilon = \sum_{j=1}^m \pi_i(x_j) E_{\bar{a}}(B_j).$$

Luego, para cada $1 \leq i \leq n$ tenemos

$$\|w_\epsilon^* b_i w_\epsilon - a_i\| \leq \left\| w_\epsilon^* \left(b_i - \sum_{j=1}^m \pi_i(x_j) E_{\bar{b}}(B_j) \right) w_\epsilon \right\| + \frac{\epsilon}{2} < \epsilon.$$

□

Como una consecuencia inmediata del Teorema 8.2.12 vemos que la clausura en norma de la órbita unitaria conjunta de una familia abeliana es cerrada en la topología fuerte en factores II_1 . Esto generaliza resultados de [61, teorema 5.4 (4)]. Dada $\bar{x} = (x_i)_{i=1}^n \subseteq \mathcal{M}$ entonces $\overline{\mathcal{U}_{\mathcal{M}}(\bar{x})}^s$ denota la clausura entrada a entrada en la topología fuerte de operadores.

Corolario 8.2.13. Sea $\bar{a} = (a_i)_{i=1}^n \subseteq \mathcal{M}_{sa}$ una familia abeliana. Entonces se tiene la igualdad

$$\overline{\mathcal{U}_{\mathcal{M}}(\bar{a})}^{\|\cdot\|} = \overline{\mathcal{U}_{\mathcal{M}}(\bar{a})}^s.$$

Demostración. Sea $\bar{b} = (b_i)_{i=1}^n \in \overline{\mathcal{U}_{\mathcal{M}}(\bar{a})}^s$. Entonces existe una red $(a_1^j, \dots, a_n^j)_{j \in J} \subseteq \mathcal{U}_{\mathcal{M}}(\bar{a})$ tal que a_i^j converge fuertemente hacia b_i para cada $i = 1, \dots, n$. Sea $f : \mathbb{R}^n \rightarrow \mathbb{R}$ una función continua. Entonces $\tau(f(a_1^j, \dots, a_n^j)) = \tau(f(a_1, \dots, a_n))$ para cada j . Así, por [66, Lema II.4.3] concluimos que $\tau(f(b_1, \dots, b_n)) = \tau(f(a_1, \dots, a_n))$. Notemos que (5) del Teorema 8.2.12 implica que $\bar{b} \in \overline{\mathcal{U}_{\mathcal{M}}(\bar{a})}$. Por otra parte, la otra inclusión es trivial. □

La unicidad de la traza y resultados similares de doble mayorización de operadores autoadjuntos [49] permiten verificar que cualquier automorfismo de un factor II_1 es puntualmente aproximadamente (en norma) interior. El siguiente corolario es una mejora de este resultado.

Corolario 8.2.14. Sea Θ un automorfismo de $\mathcal{M}$. Entonces $\Theta|_{\mathcal{A}}$ es aproximadamente interior para cada C^* subálgebra abeliana y separable $\mathcal{A} \subset \mathcal{M}$.

Demostración. La unicidad de la traza asegura que Θ preserva la traza. Por ser multiplicativo, el rango de una subálgebra abeliana es una subálgebra abeliana. Es decir, Θ es una transformación DS que mapea toda familia abeliana de $\mathcal{M}$ en otra familia abeliana de $\mathcal{M}$. Consideremos un subconjunto denso y numerable $\{a_i\}$ de $\mathcal{A}$, y usemos el Teorema 8.2.12 para obtener operadores unitarios u_n para cada subconjunto finito $\{a_1, \dots, a_n\}$. Un argumento de tipo $\epsilon/3$ muestra entonces que la sucesión $\{\text{Ad } u_n\}$ aproxima Θ en toda $\mathcal{A}$. □

8.2.1. Comparación con la mayorización conjunta de Alberti y Uhlmann

En [2] Alberti y Uhlmann consideraron una noción general de mayorización conjunta entre n -uplas de estados de una C^* -álgebra abeliana fija. En esta sección mostramos como aplicar esta noción general en nuestro caso particular de una C^* -subálgebra abeliana de un factor II_1 y la comparamos con la mayorización conjunta que hemos desarrollado hasta aquí.

Sea $\mathcal{A}$ una C^* -álgebra abeliana y sea $\mathcal{A}^*$ su espacio dual. En este contexto, una función lineal $V : \mathcal{A}^* \rightarrow \mathcal{A}^*$ es una función estocástica si es positiva y $V(\psi)(1) = \psi(1)$ para cada $\psi \in \mathcal{A}^*$. El conjunto de funciones estocásticas de $\mathcal{A}$ es denotado por $ST(\mathcal{A})$. Notemos que una función estocástica $V \in ST(\mathcal{A})$ mapea el espacio de estados de $\mathcal{A}$ sobre sí mismo. Dado un estado $\varphi \in \mathcal{A}^*$, el conjunto de funciones estocásticas tales que $V(\varphi) = \varphi$ es denotado por $ST_\varphi(\mathcal{A})$.

Definición 8.2.15 ([2]). Sea $\bar{w}, \bar{u}$ dos n -uplas ordenadas de estados de una C^* -álgebra abeliana $\mathcal{A}$. Entonces $\bar{u}$ está mayorizada por $\bar{w}$, y notamos $\bar{u} \triangleleft \bar{w}$, si existe una función estocástica $V \in ST(\mathcal{A})$ tal que $V\bar{w} = \bar{u}$, $1 \leq i \leq n$.

En realidad, la definición 8.2.15 es una formulación equivalente de la definición de mayorización dada en [2] (que está dada en términos de funcionales generalizados de Ky-Fan, ver Teorema 3.2.3 en [2]). Si en la definición de arriba pedimos que $V \in ST_\varphi(\mathcal{A})$ para un estado fijo $\varphi \in \mathcal{A}^*$, decimos que $\bar{u}$ está φ -mayorizada por $\bar{w}$, y notamos $\bar{u} \triangleleft_\varphi \bar{w}$, y esta noción puede ser considerada como una mayorización *ponderada* (de forma que φ es el peso).

Sea $a \in \mathcal{M}^+$ con $\tau(a) = 1$. Entonces a induce estado normal $\tau_a \in \mathcal{M}_*$ dado por $\tau_a(x) = \tau(ax)$. De esta forma, podemos considerar la siguiente definición:

Definición 8.2.16. Sean $\bar{a}, \bar{b} \subset \mathcal{M}^+$ dos n -uplas ordenadas tales que $\tau(a_i) = \tau(b_i) = 1$, $1 \leq i \leq n$, y sea $\mathcal{A}$ una subálgebra abeliana de von Neumann de $\mathcal{M}$. Entonces decimos que $\bar{a}$ está $\mathcal{A}$ -mayorizada por $\bar{b}$, y notamos $\bar{a} \triangleleft_{\mathcal{A}} \bar{b}$, si $(\tau_{a_i}) \triangleleft_\tau (\tau_{b_i})$ como estados sobre $\mathcal{A}$.

Aclaremos que la definición dada arriba no está contenida dentro del desarrollo en [2], aunque bien puede considerarse implícita en este trabajo.

Observación 8.2.17. Notemos que las familias no son necesariamente abelianas. Pero si $\mathcal{E}_{\mathcal{A}}$ denota la esperanza condicional sobre $\mathcal{A}$ que conmuta con τ , entonces para cada $a, x \in \mathcal{M}$ tenemos $\tau(ax) = \tau(\mathcal{E}_{\mathcal{A}}(a)x)$. Es decir, estamos comparando en realidad las familias abelianas $(\mathcal{E}_{\mathcal{A}}(b_i))$, $(\mathcal{E}_{\mathcal{A}}(a_i))$ de operadores en el álgebra abeliana $\mathcal{A}$. Por otro lado, esta noción depende fuertemente del álgebra abeliana $\mathcal{A}$, y en este sentido esta noción de mayorización puede considerarse como *local*.

La siguiente Proposición establece una dualidad que nos permite comparar nuestra noción de mayorización conjunta con la de Alberti y Uhlman.

Lema 8.2.18. Si $T \in DS(\mathcal{M})$ entonces existe una única $T' \in DS(\mathcal{M})$ tal que $\tau(T(a)b) = \tau(aT'(b))$ para todos $a, b \in \mathcal{M}$.

Demostración. Sea $b \in \mathcal{M}^+$. Consideramos el funcional positivo $\phi_b \in \mathcal{M}_*$ dado por $\phi_b(x) = \tau(bT(x))$. Como para cada $x \in \mathcal{M}^+$, $0 \leq \phi_b(x) \leq \|b\|\tau(x)$, podemos aplicar el [48, teorema 7.3.13] para encontrar $c \in \mathcal{M}^+$, $\|c\| \leq \|b\|$, tal que $\phi_b(x) = \tau(cx)$. Del hecho de que la traza es fiel concluimos que este c es único, y lo denotamos por $c = T'(b)$. Como cada $a \in \mathcal{M}$ es una combinación lineal de (cuatro) operadores positivos, entonces T' se extiende por linealidad a todo $\mathcal{M}$ en tal forma que $\tau(T(a)b) = \tau(aT'(b))$, para todos $a, b \in \mathcal{M}$. Usando nuevamente el hecho de

que la traza es fiel, deducimos que T' es una transformación lineal positiva y unital. Finalmente, $\tau(T'(b)) = \tau(1T'(b)) = \tau(bT(1)) = \tau(b)$, con lo que $T' \in DS(\mathcal{M})$. $\square$

La siguiente proposición muestra una comparación entre la noción de mayorización conjunta desarrollada en este capítulo y la definición 8.2.15.

Proposición 8.2.19. Sea $\bar{a}, \bar{b} \subset \mathcal{M}_{sa}$ familias abelianas tales que $\bar{a} \prec \bar{b}$. Si $\mathcal{B}$ denota el álgebra de von Neumann abeliana generada por la familia $\bar{b}$, entonces $\bar{a} \triangleleft_{\mathcal{B}} \bar{b}$.

Demostración. Sea $\mathcal{E}_{\mathcal{B}}$ la esperanza condicional sobre $\mathcal{B}$ que conmuta con la traza y sea $T \in DS(\mathcal{M})$ tal que $T(b_i) = a_i$, para $1 \leq i \leq n$ (que existe por el Teorema 8.2.5). Por el Lema 8.2.18 existe $T' \in DS(\mathcal{M})$ tal que $\tau(T(a)b) = \tau(aT'(b))$ para todos $a, b \in \mathcal{M}$. Entonces $\mathcal{E}_{\mathcal{B}} \circ T'|_{\mathcal{B}}$ es una transformación positiva y unital que preserva la traza de $\mathcal{B}$ en sí mismo. Consideremos $V : \mathcal{B}^* \rightarrow \mathcal{B}^*$ dado por $V(\psi)(b) = \psi \circ \mathcal{E}_{\mathcal{B}} \circ T'(b)$; y de aquí es inmediato que V es una función estocástica tal que $V(\tau) = \tau$. Por otro lado, notemos que

$$V(\tau_{b_i})(x) = \tau(b_i \mathcal{E}_{\mathcal{B}}(T'(x))) = \tau(b_i T'(x)) = \tau(T(b_i)x) = \tau(a_i x) = \tau_{a_i}(x).$$

De esto concluimos que $\bar{a} \triangleleft_{\mathcal{B}} \bar{b}$. $\square$

El hecho de que $\bar{a} \triangleleft_{\mathcal{D}} \bar{b}$ con respecto a alguna subálgebra de von Neumann abeliana $\mathcal{D} \in \mathcal{M}$ no implica que $\bar{a} \prec \bar{b}$, aún si $\mathcal{D} = \mathcal{B}$ y las familias tienen un solo elemento. En efecto, consideremos el siguiente ejemplo (que puede realizarse dentro de cualquier factor II_1):

Ejemplo 8.2.20. Sean $a, b \in \mathcal{M}_2(\mathbb{C})_{sa}$ dadas por

$$a = \begin{pmatrix} 1/2 & 1/2 \\ 1/2 & 1/2 \end{pmatrix}, \quad b = \begin{pmatrix} 1/4 & 0 \\ 0 & 3/4 \end{pmatrix}.$$

La $*$ -álgebra $\mathcal{B}$ generada por b es el álgebra diagonal con respecto a la base canónica de $\mathbb{C}^2$. De esta forma $\mathcal{E}_{\mathcal{B}}(a) = \frac{1}{2} I_2$. Entonces vemos que $\mathcal{E}_{\mathcal{B}}(a) \triangleleft_{\mathcal{B}} b$, que a su vez muestra que $a \triangleleft_{\mathcal{B}} b$. Por otro lado $a \not\prec b$, pues si suponemos $a \prec b$ entonces se tendría $(1, 0) \prec (3/4, 1/4)$ como vectores en $\mathbb{R}^2$, que es falso.

8.3. Algunas herramientas técnicas

En esta sección probamos los resultados presentados al comienzo de la de sección 8.1. Comenzamos con la prueba de que cualquier C^* -subálgebra abeliana y separable de $\mathcal{M}$ puede ser sumergida en una C^* -subálgebra abeliana, separable y difusa. Luego de esto, probamos algunos resultados de aproximación que valen para C^* subálgebras abelianas, separables y difusas de $\mathcal{M}$.

8.3.1. Refinamientos de medidas espectrales

Un problema técnico que surge cuando consideramos medidas espectrales está relacionado con los átomos de las medidas. En general, los problemas relativos a medidas totalmente atómicas o totalmente difusas se estudian mediante métodos definidos según el caso. Sin embargo estos métodos no son útiles en el caso de medidas que tienen parte difusa y atómica no trivial. En esta sección mostramos que en nuestro contexto es posible eliminar la parte atómica de las medidas espectrales sumergiendo el álgebra en un ambiente más grande (pero no mucho más grande). Si bien nuestro

interés es aplicar estos resultados para el estudio de las transformaciones doble estocásticas y la mayorización conjunta, estas técnicas y posteriores refinamientos nos permitirán el estudio de otros problemas naturales dentro de los factores Π_1 .

Comenzamos observando algunos hechos elementales acerca de las inclusiones de C^* -subálgebras abelianas en factores Π_1 que necesitamos para nuestro estudio. Si $\mathcal{A} \subseteq \mathcal{B}$ son C^* -álgebras unitales, entonces la función $\Phi : \Gamma(\mathcal{B}) \rightarrow \Gamma(\mathcal{A})$ dada por $\Phi(\gamma) = \gamma|_{\mathcal{A}}$ es una suryección continua sobre $\Gamma(\mathcal{A})$. Supongamos además que $\mathcal{A} \subseteq \mathcal{B} \subseteq \mathcal{M}$ son separables y sean $E_{\mathcal{A}}, E_{\mathcal{B}}$ sus medidas espectrales, entonces $E_{\mathcal{A}} = E_{\mathcal{B}} \circ \Phi^{-1}$ y $\mu_{\mathcal{A}} = \mu_{\mathcal{B}} \circ \Phi^{-1}$. Por otro lado si $\mathcal{A} \subseteq \mathcal{B} \subseteq \mathcal{C}$ y $\Phi : \Gamma(\mathcal{B}) \rightarrow \Gamma(\mathcal{A})$, $\Psi : \Gamma(\mathcal{C}) \rightarrow \Gamma(\mathcal{B})$ son las suryecciones continuas asociadas a las inclusiones anteriores entonces $\Phi \circ \Psi : \Gamma(\mathcal{C}) \rightarrow \Gamma(\mathcal{A})$ es la suryección asociada a la inclusión $\mathcal{A} \subseteq \mathcal{C}$. La propiedad anterior pone de manifiesto el hecho de que la asociación $(\mathcal{A} \subseteq \mathcal{B}) \mapsto (\Phi : \Gamma(\mathcal{B}) \rightarrow \Gamma(\mathcal{A}))$ es funtorial entre las categorías correspondientes.

Por otro lado, como $\mu_{\mathcal{A}} = \tau \circ E_{\mathcal{A}}$ entonces, por fidelidad de la traza, $\text{At}(\mu_{\mathcal{A}}) = \text{At}(E_{\mathcal{A}})$. Sea $\sum_{x \in \text{At}(E_{\mathcal{A}})} \mu_{\mathcal{A}}(\{x\})$ la masa atómica total de $E_{\mathcal{A}}$. Como $\mu_{\mathcal{A}}$ es finita, el conjunto de átomos es numerable. Notemos que $\mathcal{A}$ es difusa justamente cuando su masa atómica total es cero. Finalmente, remarcamos el hecho de que los elementos de $\text{At}(E_{\mathcal{A}})$ están en correspondencia biunívoca con el conjunto de proyecciones minimales de $L^\infty(\mathcal{A})$.

Lema 8.3.1. Con la notación de arriba, si $x \in \text{At}(E_{\mathcal{B}})$ entonces $\Phi(x) \in \text{At}(E_{\mathcal{A}})$, y la masa atómica total de $E_{\mathcal{B}}$ es menor que la masa atómica total de $E_{\mathcal{A}}$.

Demostración. Sea $x \in \text{At}(E_{\mathcal{B}})$ y notemos que $0 \neq \mu_{\mathcal{B}}(\{x\}) \leq \mu_{\mathcal{B}}(\Phi^{-1}(\Phi(\{x\}))) = \mu_{\mathcal{A}}(\Phi(\{x\}))$, con lo que $\Phi(x) \in \text{At}(E_{\mathcal{A}}) = \text{At}(\mu_{\mathcal{A}})$. Consideramos la relación de equivalencia en $\text{At}(E_{\mathcal{B}})$ inducida por Φ , i.e. $x \sim y$ si $\Phi(x) = \Phi(y)$. Si $Q \in \mathcal{Q} := \text{At}(E_{\mathcal{B}})/\sim$ es tal que $\Phi(x) = x_Q$ para cada $x \in Q$, entonces usando que Q es numerable obtenemos que

$$\sum_{x \in Q} \mu_{\mathcal{B}}(\{x\}) = \mu_{\mathcal{B}}(Q) \leq \mu_{\mathcal{B}}(\Phi^{-1}(\{x_Q\})) = \mu_{\mathcal{A}}(\{x_Q\}).$$

Entonces concluimos que

$$\sum_{x \in \text{At}(E_{\mathcal{B}})} \mu_{\mathcal{B}}(\{x\}) = \sum_{Q \in \mathcal{Q}} \sum_{x \in Q} \mu_{\mathcal{B}}(\{x\}) \leq \sum_{Q \in \mathcal{Q}} \mu_{\mathcal{A}}(x_Q) \leq \sum_{x \in \text{At}(E_{\mathcal{A}})} \mu_{\mathcal{A}}(\{x\}),$$

pues se trata de series de términos positivos. □

El Lema 8.3.1 muestra que es posible reducir la masa atómica total agrandando el álgebra $\mathcal{A}$. En la Proposición 8.3.3 damos una formulación más específica de este último hecho. El siguiente lema es una consecuencia inmediata del Lema 3.2.1

Lema 8.3.2. Sea $p \in \mathcal{P}(\mathcal{M})$ y sean $\alpha, \beta \in \mathbb{R}$ con $0 < \alpha < \beta$. Entonces existe $a \in \mathcal{M}_{sa}$ tal que $[\alpha, \beta] \subseteq \sigma(a) \subseteq [\alpha, \beta] \cup \{0\}$, $P_{\overline{R(a)}} = p$, y

$$\mu_a((x, \beta]) = \frac{\tau(p)(\beta - x)}{\beta - \alpha}, \quad \alpha \leq x \leq \beta.$$

En particular, $\text{At}(E_a) \subseteq \{0\}$.

En el siguiente resultado mostramos que podemos “borrar” un átomo de un álgebra abeliana, sumergiéndola en el álgebra obtenida de agregar un generador a la original.

Proposición 8.3.3. Con las notaciones de arriba, sea $x_0 \in \Gamma(\mathcal{A})$ un átomo de $E_{\mathcal{A}}$ y sean $\alpha, \beta \in \mathbb{R}$ con $0 < \alpha < \beta$. Entonces existe $a \in \mathcal{A}' \cap \mathcal{M}_{sa}$ con $[\alpha, \beta] \subseteq \sigma(a) \subseteq [\alpha, \beta] \cup \{0\}$, $P_{\overline{R(a)}} = E_{\mathcal{A}}(\{x_0\})$, y tal que $E_{\mathcal{B}}$ no tiene átomos en la fibra $\Phi^{-1}(x_0)$, donde $\mathcal{B} = C^*(\mathcal{A}, a)$.

Demostración. Sea $p = E_{\mathcal{A}}(\{x_0\})$ y sea $a \in \mathcal{M}_{sa}$ como en el Lema 8.3.2, de forma que a satisface las condiciones sobre el $\sigma(\bar{a})$ y $P_{\overline{R(a)}}$. Como p es una proyección minimal en $L^\infty(\mathcal{A})$, tenemos $pb = pbp = \lambda_b p$ para cada $b \in \mathcal{A}$ y entonces $ab = apb = \lambda_b a = bpa = ba$. De esto se sigue que $a \in \mathcal{A}' \cap \mathcal{M}$.

Sea $\mathcal{B} = C^*(\mathcal{A}, a)$ y sea $\Phi : \Gamma(\mathcal{B}) \rightarrow \Gamma(\mathcal{A})$, $\Psi : \Gamma(\mathcal{B}) \rightarrow \Gamma(C^*(a))$ las suryecciones continuas inducidas por las inclusiones $\mathcal{A} \subseteq \mathcal{B}$ y $C^*(a) \subseteq \mathcal{B}$.

Notemos que la restricción $\Psi|_{\Phi^{-1}(x_0)}$ es inyectiva. De hecho, sean $x, y \in \Phi^{-1}(x_0)$ tal que $\Psi(x) = \Psi(y)$, i.e. las restricciones de los caracteres a $C^*(a)$ coinciden. Pero como $\Phi(x) = \Phi(y) (= x_0)$ los caracteres también coinciden en $\mathcal{A}$ y luego son iguales como caracteres en $\mathcal{B}$, ya que $\mathcal{B}$ está generada por $\mathcal{A}$ y $C^*(a)$.

Por otra parte, si $x \in \Gamma(\mathcal{B})$ es tal que $x(a) \neq 0$, entonces $\Phi(x) = x_0$. En efecto, supongamos que $\Phi(x) \neq x_0$ y sea $f \in C(\Gamma(\mathcal{A}))$ con $f(\Phi(x)) = 0$ y $f(x_0) = 1$. Entonces $f \circ \Phi \geq 1_{\Phi^{-1}(x_0)}$, pero

$$\int_{\Gamma(\mathcal{B})} f \circ \Phi dE_{\mathcal{B}} \geq \int_{\Gamma(\mathcal{B})} 1_{\Phi^{-1}(x_0)} dE_{\mathcal{B}} = E_{\mathcal{B}}(\Phi^{-1}(x_0)) = E_{\mathcal{A}}(\{x_0\}) = p.$$

Notemos que si $0 \in \sigma(a)$ entonces es un punto aislado, y en cualquier caso tenemos $p \in C^*(a) \subseteq \mathcal{B}$. Entonces $0 = f \circ \Phi(x) \geq x(p) \geq 0$, y así $x(p) = 0$. Como $0 \leq a \leq \beta p$, $x(a) = 0$.

Sea $z \in \Phi^{-1}(x_0)$. Entonces o bien $z(a) = 0$ o bien $z(a) \neq 0$. Veamos que en cualquier caso obtenemos $E_{\mathcal{B}}(\{z\}) = 0$. Si $z(a) \neq 0$, de los párrafos anteriores deducimos que $\Psi^{-1}(\Psi(z)) = \{z\}$ y consecuentemente $E_{\mathcal{B}}(\{z\}) = E_a(\{\Psi(z)\}) = 0$ ya que $\Psi(z)(a) \neq 0$ y $\text{At}(E_a) \subseteq \{0\}$. Si $z(a) = 0$ (hay a lo sumo un tal carácter) entonces, de los párrafos anteriores tenemos

$$\begin{aligned} \{z\} &= \Phi^{-1}(x_0) \setminus \{x \in \Phi^{-1}(x_0) : x(a) \neq 0\} \\ &= \Phi^{-1}(x_0) \setminus \Psi^{-1}(\{x \in \Gamma(C^*(a)) : x(a) \neq 0\}) \end{aligned}$$

y

$$\begin{aligned} E_{\mathcal{B}}(\Psi^{-1}(\{x \in \Gamma(C^*(a)) : x(a) \neq 0\})) &= E_a(\{x \in \Gamma(C^*(a)) : x(a) \neq 0\}) \\ &= E_{\mathcal{A}}(\{x_0\}) = E_{\mathcal{B}}(\Phi^{-1}(x_0)). \end{aligned}$$

De esto concluimos que $E_{\mathcal{B}}(\{z\}) = 0$. □

Prueba del Teorema 8.1.1. Recordemos que el conjunto $\text{At}(E_{\mathcal{A}})$ de átomos de $E_{\mathcal{A}}$ es numerable. Si $\text{At}(E_{\mathcal{A}}) = \emptyset$ entonces $E_{\mathcal{A}}$ ya es difusa. En otro caso, enumeramos $\text{At}(E_{\mathcal{A}}) = \{x_i : 1 \leq i \leq r\}$, donde $r \in \mathbb{N} \cup \{\infty\}$. Para $1 \leq i \leq r$, sea $I_i = [\alpha_i, \beta_i]$ con $\alpha_i = 1 + \frac{1}{2n} < \beta_i = 1 + \frac{1}{2n-1}$. Entonces $I_i \cap \overline{\bigcup_{1 \leq j \leq r, j \neq i} I_j} = \emptyset$ y $\bigcup_{i=1}^r I_i \subseteq [1, 2]$. Para cada $i = 1, \dots, r$ existe, por la Proposición 8.3.3, $a_i \in \mathcal{A}' \cap \mathcal{M}_{sa}$ tal que $P_{\overline{R(a_i)}} = E_{\mathcal{A}}(\{x_i\})$, $I_i \subseteq \sigma(a_i) \subseteq I_i \cup \{0\}$, y tal que $E_{\mathcal{A}_i}$ no tiene átomos en la fibra $\Phi_i^{-1}(x_i)$, donde $\Phi_i : \Gamma(\mathcal{A}_i) \rightarrow \mathcal{A}$ denota la suryección continua inducida por la inclusión $\mathcal{A} \subseteq \mathcal{A}_i := C^*(\mathcal{A}, a_i)$. Sea $a = \sum_{i=1}^r a_i \in \mathcal{A}' \cap \mathcal{M}_{sa}$ (esta suma converge en la topología fuerte de

operadores pues los rangos de los operadores a_i son ortogonales y $\|a_i\| \leq 2$ para cada i). Entonces $\mathcal{B} := C^*(\mathcal{A}, a)$ es una subálgebra abeliana de $\mathcal{M}$.

Sea $\mathcal{B} = C^*(\mathcal{A}, a)$. Afirmamos que la medida espectral $E_{\mathcal{B}}$ de $\mathcal{B}$ no tiene átomos. De hecho, notemos que $1_{I_i} \in C(\cup_{1 \leq j \leq r} I_j)$ es una función continua (pues la distancia entre los conjuntos I_i y $\cup_{i \neq j} I_j$ es positiva); entonces, como $1_{I_i}(a) = a_i$, se sigue que $\mathcal{A}_i \subset \mathcal{B}$ para cada $1 \leq i \leq r$. Supongamos que $x \in \text{At}(\Gamma(\mathcal{B}))$ y sea $\Phi : \Gamma(\mathcal{B}) \rightarrow \Gamma(\mathcal{A})$ como antes. Entonces por el Lema 8.3.1 existe $i \in \{1, \dots, r\}$ tal que $\Phi(x) = x_i \in \text{At}(E_{\mathcal{A}})$. Como $\Phi = \Phi_i \circ \Psi_i$, donde $\Psi_i : \Gamma(\mathcal{B}) \rightarrow \Gamma(\mathcal{A}_i)$ es la suryección inducida por la inclusión $\mathcal{A}_i \subseteq \mathcal{B}$, concluimos que $\Psi_i(x) \in \Phi_i^{-1}(x_i)$ es un átomo de la medida $E_{\mathcal{A}_i}$, nuevamente por el Lema 8.3.1. Pero esto es una contradicción ya que por construcción no hay átomos en la fibra $\Phi_i^{-1}(x_i)$. $\square$

Observación 8.3.4. Dada una C^* subálgebra abeliana $\mathcal{A} \subset \mathcal{M}$, una forma directa para hallar una C^* -subálgebra abeliana $\mathcal{A} \subseteq \tilde{\mathcal{A}} \subset \mathcal{M}$ con medida espectral difusa es considerar una MASA en $\mathcal{M}$ que contenga a $\mathcal{A}$. La información adicional que obtenemos del Teorema 8.1.1 es que $\tilde{\mathcal{A}}$ puede elegirse separable (como C^* -álgebra) siempre que $\mathcal{A}$ sea separable. En este último caso, el espacio de caracteres de $\tilde{\mathcal{A}}$ es metrizable, un hecho que es importante para nuestros cálculos.

8.3.2. Aproximaciones discretas en álgebras abelianas, separables y difusas

Dado un espacio métrico compacto siempre es posible hallar, usando la continuidad uniforme, aproximaciones discretas de funciones continuas por combinaciones lineales de funciones características de ciertos conjuntos $\{Q_i\}_{i=1}^m$. Pero si consideramos una medida en este espacio y requerimos medidas iguales para estos conjuntos, puede que no haya tal aproximación uniforme basada en funciones características. La Proposición 8.1.2 establece, en lenguaje de las álgebras de von Neumann, que hay una solución intermedia de este problema, que se obtiene promediando sobre un conjunto adecuado de particiones del espacio métrico. Esta solución está inspirada en la prueba del Lema 4.1 de [40] y es una herramienta fundamental para nuestros desarrollos. Antes consideramos una propiedad elemental de las medidas de probabilidad difusas (es decir sin átomos) de soporte compacto.

Lema 8.3.5. Sea $K \subset \mathbb{R}^n$ un compacto con la topología usual de $\mathbb{R}^n$ y sea μ una medida de probabilidad, regular, de Borel y difusa en K . Entonces, para cada $\alpha \in (0, 1)$ existe un conjunto medible $S \subset K$ tal que $\mu(S) = \alpha$.

Prueba de la Proposición 8.1.2. $\Gamma(\mathcal{B})$ es un espacio métrico compacto, con métrica d (que induce la topología w^*). Sea $r \in \mathbb{N}$; por compacidad existe una partición $\{\tilde{Q}_i\}_{i=1}^{k_0}$ de $\Gamma(\mathcal{B})$ con $\text{diam}_d(\tilde{Q}_i) < \frac{1}{r}$. Entonces $\sum_{i=1}^{k_0} \mu_{\mathcal{B}}(\tilde{Q}_i) = 1$. Sea $m = m(r)$ tal que

$$1/m \leq \min\{\mu_{\mathcal{B}}(\tilde{Q}_j)^2 : 1 \leq j \leq k_0\}$$

y notemos que $\lim_{r \rightarrow \infty} m(r) = \infty$, pues la medida $\mu_{\mathcal{B}}$ es difusa (esta condición lo suficientemente débil de forma que nos permite ajustar el m a condiciones adicionales, por ejemplo en el caso de dos particiones involucradas). Entonces para cada $1 \leq j \leq k_0$ existe $k_j \in \mathbb{N}$ tal que $\mu_{\mathcal{B}}(\tilde{Q}_j) = k_j/m + \delta_j$ con $0 \leq \delta_j < 1/m$. Si definimos $\tilde{k} = \tilde{k}(r) = \min_j\{k_j\}$ entonces

$$\tilde{k} \geq \max\{\mu_{\mathcal{B}}(\tilde{Q}_j)^{-1} - 1, 1 \leq j \leq k_0\}$$

con lo que $\lim_{r \rightarrow \infty} \tilde{k}(r) = \infty$.

Para $1 \leq j \leq k_0$, elijamos $\tilde{k}$ particiones $\{\tilde{Q}_{j,s}^t\}_{s=0}^{k_j}$ de cada $\tilde{Q}_j$ ($1 \leq t \leq \tilde{k}$), con $\mu_{\mathcal{B}}(\tilde{Q}_{j,s}^t) = 1/m$ si $1 \leq s \leq k_j$ y $\mu_{\mathcal{B}}(\tilde{Q}_{j,0}^t) = \delta_j$, de forma tal que $\tilde{Q}_{j,0}^t \subset \tilde{Q}_{j,t}^1$, $2 \leq t \leq \tilde{k}$. Siempre podemos hacer tal elección: usando el Lema 8.3.5 elegimos $\tilde{Q}_{j,0}^t \subseteq \tilde{Q}_{j,t}^1$ con $\mu_{\mathcal{B}}(\tilde{Q}_{j,0}^t) = \delta_j < 1/m$, y entonces tomamos una partición $\{\tilde{Q}_{j,s}^t\}_{s=1}^{k_j}$ de $\tilde{Q}_j \setminus \tilde{Q}_{j,0}^t$ lo que puede hacerse por el Lema 8.3.5 (notemos que $\mu_{\mathcal{B}}(\tilde{Q}_j \setminus \tilde{Q}_{j,0}^t) = k_j/m$). De nuestra elección, $\tilde{Q}_{j,0}^t \cap \tilde{Q}_{j,0}^{t'} = \emptyset$ si $t \neq t'$. Este último hecho es importante para nuestros cálculos.

Para cada $1 \leq t \leq \tilde{k}$, sea $\tilde{Q}_{0,0}^t = \cup_{j=1}^{k_0} \tilde{Q}_{j,0}^t$. Entonces $\mu_{\mathcal{B}}(\tilde{Q}_{0,0}^t) = 1 - \sum_j k_j/m = (m - \sum_{j=1}^{k_0} k_j)/m$. Finalmente, tomamos particiones de cada conjunto $\tilde{Q}_{0,0}^t$ de $n_1 = m - \sum_j k_j$ subconjuntos $\{\tilde{Q}_i^t\}_{i=1}^{n_1}$ de medida $1/m$. Renombrando las $\tilde{k}$ particiones $\{\tilde{Q}_{j,s}^t\}_{j,s} \cup \{\tilde{Q}_i^t\}_i$, terminamos con $\tilde{k}$ particiones $\{Q_i^{t,m}\}_{i=1}^m$, para $1 \leq t \leq \tilde{k}$, tales que:

1. $\mu_{\mathcal{B}}(Q_i^{t,m}) = 1/m$, para cada $i \in \{1, \dots, m\}$, $t \in \{1, \dots, \tilde{k}\}$,
2. $\text{diám}_d(Q_i^{t,m}) \leq 1/r$, si $i > n_1$,
3. Si $1 \leq i, i' \leq n_1$ entonces $Q_i^{t,m} \cap Q_{i'}^{t',m} = \emptyset$ si $i \neq i'$ ó $t \neq t'$.

La construcción de las k particiones $\{Q_i^{t,m}\}_{i=1}^m$ fue hecha de tal forma que los subconjuntos que no tienen diámetros pequeños son disjuntos, aún para particiones diferentes.

Sea $\mathbb{M} = \{m(r), r \geq 1\}$ y para cada $m = m(r) \in \mathbb{M}$ sea $k(m) = \tilde{k}(r)$ como antes y para i, t, m , sea $q_i^{t,m} = E_{\mathcal{B}}(Q_i^{t,m})$. El conjunto $\mathbb{M}$ es no acotado y, para cada $t = 1, \dots, k$, $\{q_i^{t,m}\}_{i=1}^m \subset \mathcal{B}' \cap \mathcal{M}$ es una partición de la unidad.

Sea $b \in \mathcal{B}$, $\epsilon > 0$, y sea $f \in C(\Gamma(\mathcal{B}))$ tal que $b = \int_{\Gamma(\mathcal{B})} f dE_{\mathcal{B}}$. Luego, por compacidad, existe $\delta > 0$ tal que si $Q \subseteq \Gamma(\mathcal{B})$ con $\text{diam}_d(Q) < \delta$ entonces $\text{diam}(f(Q)) < \epsilon$. Sea $r \in \mathbb{N}$ tal que $1/r < \delta$ y $2\|b\|/k(r) \leq \epsilon$ y sea $m = m(r) \in \mathbb{M}$. Si

$$\beta_i^{t,m} = m \tau(b q_i^{t,m}) = m \int_{Q_i^{t,m}} f d\mu_{\mathcal{B}}$$

entonces, las propiedades 1-3 implican las siguientes

- 1'. $\tau(q_i^{t,m}) = 1/m$, para cada $i \in \{1, \dots, m\}$, $t \in \{1, \dots, k\}$,
- 2'. si $i > n_1$, entonces $|f(x) - \beta_i^{t,m}| \leq \epsilon$, $\forall x \in Q_i^{t,m}$.
- 3'. si $1 \leq i, i' \leq n_1$ entonces $q_i^{t,m} \perp q_{i'}^{t',m}$ si $i \neq i'$ ó $t \neq t'$.

Pero entonces tenemos

$$\begin{aligned} \left\| b - \frac{1}{k} \sum_{t=1}^k \sum_{i=1}^m \beta_i^{t,m} q_i^{t,m} \right\| &= \left\| \frac{1}{k} \sum_{t=1}^k \left(b - \sum_{i=1}^m \beta_i^{t,m} q_i^{t,m} \right) \right\| \\ &\leq \left\| \frac{1}{k} \sum_{t=1}^k \sum_{i=1}^{n_1} \int_{Q_i^{t,m}} (f - \beta_i^{t,m}) dE_{\mathcal{B}} \right\| + \epsilon \\ &\leq \left\| \frac{2\|b\|}{k} \sum_{t=1}^k \sum_{i=1}^{n_1} q_i^{t,m} \right\| + \epsilon = \frac{2\|b\|}{k} + \epsilon \leq 2\epsilon \end{aligned}$$

donde la primer desigualdad es una consecuencia de 2' y la última igualdad se sigue de 3'. $\square$

Prueba del Lema 8.1.4. Fijamos un conjunto numerable $B = (b_j)_{j \in \mathbb{N}} \subseteq \mathcal{B}$ que es denso en norma. En las construcciones previas al teorema de Dixmier [48, 8.3.4], se muestra que para cada j , existe una sucesión $\{\rho_j^n\}_{n \in \mathbb{N}} \subseteq \mathcal{D}(\mathcal{M})$ tal que para cada $1 \leq h \leq j$,

$$\|\rho_j^n(b_h) - \tau(b_h)I\| \xrightarrow{n} 0.$$

Para cada $j \in \mathbb{N}$, sea $n_0 = n_0(j) \in \mathbb{N}$ tal que si $n \geq n_0$ entonces

$$\|\rho_j^n(b_h) - \tau(b_h)I\| \leq 1/j, \quad 1 \leq h \leq j.$$

Si definimos $\rho_j = \rho_j^{n_0(j)}$ para $j \in \mathbb{N}$ entonces obtenemos que $\|\rho_j(b_h) - \tau(b_h)I\| \xrightarrow{j} 0$ para cada $h \in \mathbb{N}$. Sea $b \in \mathcal{B}$ y $\varepsilon > 0$; por hipótesis existe $b_h \in B$ tal que $\|b - b_h\| < \varepsilon/3$, y $j_0 \in \mathbb{N}$ tal que, si $j \geq j_0$, entonces $\|\rho_j(b_h) - \tau(b_h)I\| < \varepsilon/3$. Luego tenemos que

$$\|\rho_j(b) - \tau(b)I\| \leq \|\rho_j(b - b_h)\| + \|\rho_j(b_h) - \tau(b_h)I\| + |\tau(b_h - b)| < \varepsilon$$

de lo que vemos que $\lim_j \|\rho_j(b) - \tau(b)I\| = 0$ para cada $b \in \mathcal{B}$.

Para cada $i = 1, \dots, m$, consideramos el factor $\Pi_1 p_i \mathcal{M} p_i$ con traza normalizada dada por $\tau_i(p_i x) = \tau(x p_i) / \tau(p_i)$. Por la propiedad de aproximación de Dixmier mencionada en el primer párrafo, aplicada a la C^* -subálgebra separable $p_i \mathcal{B}$ en el factor finito $p_i \mathcal{M} p_i$, existe una sucesión $\{\rho_j^i\}_{j \in \mathbb{N}} \in \mathcal{D}(p_i \mathcal{M} p_i)$ tal que

$$\lim_{j \rightarrow \infty} \|\rho_j^i(p_i b) - \tau_i(p_i b) p_i\| = 0, \quad \forall b \in \mathcal{B}.$$

Para cada $\rho \in \mathcal{D}(p_i \mathcal{M} p_i)$, podemos considerar una extensión $\tilde{\rho} \in \mathcal{D}(\mathcal{M})$ como sigue: si $\rho(p_i b) = \sum_{h=1}^k \lambda_h u_h b u_h^*$, con $u_h \in \mathcal{U}(p_i \mathcal{M} p_i)$, definimos $\tilde{\rho} \in \mathcal{D}(\mathcal{M})$ por $\tilde{\rho}(b) = \sum_{h=1}^k \lambda_h \tilde{u}_h b \tilde{u}_h^*$, donde $\tilde{u}_h = u_h + (1 - p_i) \in \mathcal{U}(\mathcal{M})$. Consideremos estas extensiones $\{\tilde{\rho}_j^i\}_{j \in \mathbb{N}} \in \mathcal{D}(\mathcal{M})$, $1 \leq i \leq m$ y definamos $\rho_j = \prod_{i=1}^m \tilde{\rho}_j^i$ para $j \geq 1$. Es inmediato verificar que si $1 \leq i \leq m$ entonces $\rho_j(b p_i) = \tilde{\rho}_j^i(b p_i)$ para cada $b \in \mathcal{B}$. Luego, si $b \in \mathcal{B}$,

$$\begin{aligned} \left\| \rho_j(b) - \sum_{i=1}^m \beta_i(b) p_i \right\| &= \left\| \sum_{i=1}^m \rho_j(b p_i) - \tau_i(b p_i) p_i \right\| \\ &= \left\| \sum_{i=1}^m \tilde{\rho}_j^i(b p_i) - \tau_i(b p_i) p_i \right\| \xrightarrow{j \rightarrow \infty} 0. \end{aligned}$$

$\square$

Bibliografía

- [1] P.M. Alberti, A. Uhlmann, *Stochasticity and partial order. Doubly stochastic maps and unitary mixing*. Mathematische Monographien, **18**. VEB Deutscher Verlag der Wissenschaften, Berlin, 1981.
- [2] P.M. Alberti, A. Uhlmann, *Dissipative motion in state spaces*. Teubner-Texte zur Mathematik, **33**. BSB B. G. Teubner Verlagsgesellschaft, Leipzig, 1981.
- [3] A. Aldroubi, *Portraits of frames*, Proc. Amer. Soc. 123 (1995) 1661-1668.
- [4] E.M. Alfsen, *Compact Convex Set and Boundary Integrals*, Springer-Verlag, New York, NY 1971.
- [5] T. Ando, *Majorization, doubly stochastic matrices and comparison of eigenvalues*, Lecture Note, Hokkaido Univ., 1982.
- [6] Tsuyoshi Ando y Fumio Hiai, Hölder type inequalities for matrices, Math. inequalities and Appl. 1 (1998) 1-30.
- [7] J. Antezana, P. Massey, y D. Stojanoff, *Jensen's Inequality and Majorization*, preprint (<http://xxx.lanl.gov/abs/math.FA/0411442>).
- [8] J. Antezana, P. Massey, M. Ruiz y D. Stojanoff, *The Schur-Horn theorem for operators and frames with prescribed norms and frame operator*, to appear in Illinois J. of Math. <http://xxx.lanl.gov/abs/math.FA/0508646>
- [9] M. Argerami y P. Massey, *The local form of doubly stochastic maps and joint majorization in II_1 factors*, preprint.
- [10] W. Arveson, *Subalgebras of C^* -algebras* Acta Math. **123** (1969), 141–224.
- [11] W. Arveson, R. Kadison, *Diagonals of self-adjoint operators*, preprint (<http://xxx.lanl.gov/abs/math.OA/0508482>).
- [12] Jaspal Singh Aujla y Fernando C. Silva, *Weak majorization inequalities and convex functions*, Linear Algebra Appl. 369 (2003), 217-233.
- [13] R. Balan, P. Casazza, C. Heil y Z. Landau, *Deficits and excesses of frames*, Adv. Comp . An.,**18**: 93-116 (2003).

- [14] L. Beasley, S. Lee, *Linear operators preserving multivariate majorization*, Linear Algebra and Appl. 304 (2000), 141-159.
- [15] Julius Bendat y Seymour Sherman, Monotone and convex operator functions, Transactions de the American Mathematical Society 79 (1955), 58-71
- [16] Rajendra Bhatia, *Matrix Analysis*, Berlin-Heidelberg-New York, Springer 1997.
- [17] Rajendra Bhatia y Chandler Davis, More Operator Versions of the Schwarz inequality, Communication in Mathematical Physics, Springer Verlag 215 (2000), 239-244.
- [18] G. Birkhoff, *Three observations on linear algebra. (Spanish)* Univ. Nac. Tucumán. Revista A. 5, (1946). 147-151.
- [19] Lawrence G. Brown y Hideki Kosaki, Jensen's inequality in semi-finite von Neumann algebras, Journal of Operator Theory 23 (1990), 3-19.
- [20] P. Casazza, *The art of frame theory*, Taiwanese J. Math. 4 (2000), 129-201.
- [21] P. Casazza y M. Leon, *Frames with a given frame operator* preprint.
- [22] P. Casazza y M. Leon, *Existence and construction of finite tight frames* preprint.
- [23] G.-S. Cheon, Y.-H. Lee, *The doubly stochastic matrices of a multivariate majorization*, J. Korean Math. Soc. 32 (4) (1995), 857-867.
- [24] J.B. Conway, *A course in functional analysis*, Springer-Verlag, New York, NY 1990.
- [25] O. Christensen, *An introduction to frames and Riesz bases*, Birkhäuser, Boston, 2003.
- [26] G. Dahl, *Matrix majorization*, Linear Algebra and Appl. 288 (1999), 53-73.
- [27] I. Daubechies, A. Grossmann and Y. Meyer, *Painless nonorthogonal expansions*, J. Math. Phys. 27 (1986) 1271-1283.
- [28] K. R. Davidson, *C^* -algebras by example*. Fields Institute Monographs, 6. AMS, Providence, RI, 1996.
- [29] R.J. Duffin y A.C. Schaeffer, *A class of nonharmonic Fourier series*, Trans. Amer. Math. Soc. 72 (1952), 341-366.
- [30] K. Dykema, D. Freeman, K. Korleson, D. Larson, M. Ordower y E. Weber, *Ellipsoidal tight frames and projection decomposition of operators*, Illinois J. Math. 48 (2004), 477-489.
- [31] J. Duandikoetxea, *Introduction to Fourier Analysis*
- [32] T. Fack, *Sur la notion de valeur caractéristique*, J. Operator Theory (1982), 307-333.
- [33] D.R. Farenick y S.M. Manjegani, *Young's Inequality in Operator Algebras*, to appear in J. of the Ramanujan Math. Soc. (<http://xxx.lanl.gov/abs/math.OA/0303318>).

- [34] M. Fujii y I. Kasahara, *A remark on the spectral order of operators*, Proc. Japan Acad., 47 (1971), 986-988.
- [35] U. Haagerup, *Normal Weights in W^* algebras*, J. of Func. An. 19 (1975) 302-318.
- [36] Frank Hansen y Gert K. Pedersen, *Jensen's operator inequality*, Bulletin de London Mathematical Society 35 (2003), 553-564.
- [37] Frank Hansen y Gert K. Pedersen, *Jensen's trace inequality in several variables*, Internat. J. Math. 14 (2003), 667-681.
- [38] C. E. Heil y D.F. Walnut, *Continuous and discrete wavelet transforms*, SIAM Rev. 31 (1989), 628-666.
- [39] F. Hiai, *Majorization and Stochastic maps in von Neumann algebras*, J. Math. Anal. Appl. **127** (1987), no. 1, 18-48.
- [40] F. Hiai, *Spectral majorization between normal operators in von Neumann algebras*, Operator algebras and operator theory (Craiova, 1989), 78-115, Pitman Res. Notes Math. Ser., 271, Longman Sci. Tech., Harlow, 1992.
- [41] F. Hiai, Y. Nakamura, *Distance between unitary orbits in von Neumann algebras*, Pacific J. of Math, vol 138 (1989) 259-294.
- [42] F. Hiai, Y. Nakamura, *Closed Convex Hulls of Unitary Orbits in von Neumann Algebras*, Trans. Amer. Math. Soc. **323** (1991), 1-38.
- [43] J. R. Holub, *Pre-frame operators, Besselian frames and near-Riesz bases in Hilbert spaces*, Proc. Amer. Math. Soc. **122** (1994) 779-785.
- [44] R. Horn, C. Johnson, *Matrix Analysis*, Cambridge Univ. Press, 1985.
- [45] S.-G. Hwang, S.-S. Pyo, *Matrix majorization via vector majorization*, Linear Algebra and Appl. 332-334 (2001), 15-21.
- [46] R. Kadison, *The Pythagorean theorem I: the finite case*, Proc. N.A.S. (USA), 99(7):4178-4184, 2002.
- [47] R. Kadison, *The Pythagorean theorem II: the infinite discrete case*, Proc. N.A.S. (USA), 99(8):5217-5222, 2002.
- [48] Kadison y Ringrose, *Fundamentals of the Theory of Operator Algebras, Vol II*, Academic Press, Orlando, Florida, 1986.
- [49] Eizaburo Kamei, *Majorization in finite factors*, Math. Japonica 28, No. 4 (1983), 495-499.
- [50] K. A. Kornelson, D. R. Larson, *Rank-one decomposition of operators and construction of frames*, Wavelets, frames and operator theory, 203-214, Contemp. Math., **345**, Amer. Math. Soc., Providence, RI, 2004.

- [51] C.-K. Li, Y.-T. Poon, *Convexity of the Joint Numerical Range*, SIAM J. Matrix Analysis Appl. 21 (1999), 668-678.
- [52] C.-K. Li, E. Poon, *Linear Operators Preserving Directional Majorization*, Linear Algebra and Appl. 325 (2001), no. 1-3, 141-146.
- [53] E. H. Lieb y M. B. Ruskai, Some operator inequalities of the Schwarz type, Advances in Math. 12 (1974), 269-273.
- [54] A. W. Marshall, I. Olkin, *Inequalities: Theory of Majorization and its Applications*, Academic Press, New York, 1979.
- [55] F.D. Martínez Pería, P. Massey, L. Silvestre, *Weak matrix majorization*. Linear Algebra Appl. 403 (2005), 343-368.
- [56] A. Neumann, *An infinite-dimensional version of the Schur-Horn convexity theorem*, J. Funct. Anal. 161 (1999), 418-451.
- [57] Vern I. Paulsen, Completely bounded transformaciones y dilations, Notes in Mathematics Series 146, Pitman, New York, 1986.
- [58] G.K. Pedersen, *C^* -algebras and their automorphism groups*. London Mathematical Society Monographs, 14. Academic Press, Inc., London-New York, 1979.
- [59] G. K. Pedersen, Convex trace functions of several variables on C^* -algebras. J. operador Theory 50 (2003), 157-167.
- [60] W. Rudin, *Real and complex analysis*. Third edition. McGraw-Hill Book Co., New York, 1987.
- [61] D. Sherman, *Unitary orbits of normal operators in von Neumann algebras*, preprint.
- [62] S. Sherman, *A correction to "On a conjecture concerning doubly stochastic matrices"*, Proc. Amer. Math. Soc. 5 (1954), 998-999.
- [63] S. Shreiber, *On a result of Sherman concerning doubly stochastic matrices*, Proc. Amer. Math. Soc. 9 (1958) 350-353.
- [64] B. Simon, Trace ideals and their applications, London Mathematical Society Lecture Note Series, 35, Cambridge University Press, Cambridge-New York, 1979.
- [65] Allan M. Sinclair y Roger R. Smith, Hoshschild Cohomology de von Newmann Algebras, London Math. Soc. Lectures Notes Series 203, Cambridge Univ. Press, 1995.
- [66] M. Takesaki, *Theory of Operator Algebras I*, Encyclopaedia of Mathematical Sciences V, Springer Verlag, 2nd printing of the First Edition 1979.
- [67] J.A.Tropp, I.S.Dhillon, R.W.Heath Jr. y T. Strohmer *Designing structred tight frames via an alternating projecion method*, IEEE Transactions on information theory, Vol. 51 N°1 (2005).
- [68] F.J. Murray, J. von Neumann, *On rings of operators*, Ann. Math. 37 (1936), 116-229.

- [69] F.J. Murray, J. von Neumann, *On rings of operators II*, Trans. A.M.S. 41 (1937) 208-248.
- [70] R. M. Young, *An introduction to nonharmonic Fourier series* (revised first edition) Academic Press, San Diego, 2001.