



RIDAA
Repositorio Institucional
Digital de Acceso Abierto de la
Universidad Nacional de Quilmes



Universidad
Nacional
de Quilmes

Vanni, Leonardo

Los problemas de la medición cuántica sin decoherencia



Esta obra está bajo una Licencia Creative Commons Argentina.
Atribución - No Comercial - Sin Obra Derivada 2.5
<https://creativecommons.org/licenses/by-nc-nd/2.5/ar/>

Documento descargado de RIDAA-UNQ Repositorio Institucional Digital de Acceso Abierto de la Universidad Nacional de Quilmes de la Universidad Nacional de Quilmes

Cita recomendada:

Vanni, L. (2012). *Los problemas de la medición cuántica sin decoherencia. (Tesis de doctorado). Universidad Nacional de Quilmes, Bernal, Argentina. Disponible en RIDAA-UNQ Repositorio Institucional Digital de Acceso Abierto de la Universidad Nacional de Quilmes* <http://ridaa.unq.edu.ar/handle/20.500.11807/103>

Puede encontrar éste y otros documentos en: <https://ridaa.unq.edu.ar>

Los problemas de la medición cuántica sin decoherencia

TESIS DOCTORAL

Leonardo Vanni

idaeos@gmail.com

Resumen

En la teoría cuántica de la medición existen dos problemas centrales: el problema de la lectura definida y el problema de la base preferida. El primer problema consiste en el hecho de que la teoría no puede dar cuenta de valores bien definidos registrados como resultados en la medición, puesto que predice que el estado final del sistema y aparato es una superposición sin valor definido. El segundo problema consiste en el hecho de que bajo ciertas circunstancias (referidas a la preparación del estado del sistema a medir) la teoría no puede dar cuenta, además, de una base bien definida de estados a la cual pertenece el resultado obtenido en la medición.

Históricamente la pretendida resolución a estos problemas viene de la mano del formalismo de decoherencia, el cual se ha desarrollado como un intento de encontrar el límite clásico de la mecánica cuántica.

El objetivo principal de esta tesis consiste en brindar una respuesta a los mencionados problemas sin apelar a decoherencia alguna. Respecto del problema de la base preferida presente en las correlaciones del estado final de la medición, mostraremos que, en lugar de resolverse apelando a una interacción posterior con el entorno, tal como en la teoría de la decoherencia, puede ser resuelto analizando el proceso previo a dicho estado final sin la necesidad del entorno. La esencia del procedimiento consiste en reconsiderar la medición, no como la mera correlación que se manifiesta en el estado final, sino como el proceso capaz de establecerla. Desde esta nueva perspectiva es posible diferenciar distintos procesos que conducen a distintas correlaciones, independientemente de que éstas puedan vincularse matemáticamente mediante un cambio de base. La eliminación de la ambigüedad de la base se logra al determinar qué proceso es caracterizado por el operador de evolución que conecta el estado inicial de la medición con el final. Estrictamente hablando, el problema no queda resuelto, sino más bien disuelto, al considerar que la medición involucra todo el proceso que desemboca en una correlación final, y ello sin la necesidad de agregar interacciones posteriores con sistemas como el entorno.

Respecto del problema de la lectura definida nuestro objetivo es un poco más modesto. No vamos a resolver el problema estrictamente sino que, en el marco de ese problema, brindaremos una respuesta a la incompatibilidad que establece el postulado del colapso respecto de las evoluciones unitarias. La estrategia aquí es incorporar los aparatos de medición en el cálculo de las probabilidades condicionales establecidas para una secuencia de dos mediciones. Asumiendo valores definidos para la primera, demostramos que el postulado del colapso sobre el sistema puede ser derivado del propio formalismo de la teoría cuántica, aun cuando el sistema compuesto que incluye a los aparatos nunca abandona la evolución unitaria que predice la ecuación de Schrödinger

Prólogo

Introducción

Capítulo 1. Nociones básicas de la mecánica cuántica

- 1.1 Observables y propiedades**
- 1.2 Estados y probabilidades**
- 1.3 Estados puros y estados mezcla**
- 1.4 Estados reducidos de sistemas compuestos**
- 1.5 El principio de superposición**
- 1.6 El principio de incerteza**
- 1.7 La evolución de los estados cuánticos**

Capítulo 2. Los problemas de la medición cuántica

- 2.1 La teoría cuántica de la medición**
 - 2.1.2 Primera etapa: preparación**
 - 2.1.3 Segunda etapa: interacción**
- 2.2 El problema de la lectura definida**
 - 2.2.1 Copenhague y la hipótesis del colapso**
 - 2.2.2 Las interpretaciones modales**
 - 2.2.3 La interpretación subjetiva de Wigner**
 - 2.2.4 Everett y la multiplicidad de mundos**
- 2.3 El problema de la base preferida**

Capítulo 3. El formalismo de decoherencia

- 3.1 Idea básica y antecedentes**
- 3.2 La teoría de la decoherencia**
- 3.3 Algunas críticas a la decoherencia**

Capítulo 4. Valores definidos sin decoherencia

- 4.1 Mediciones ideales**
- 4.2 Mediciones no ideales**

Capítulo 5. Base preferida sin decoherencia

- 5.1 Condición de macroscopicidad en el experimento de Stern y Gerlach**
- 5.2 El aparato de cuatro luces**
- 5.3 ¿Detección de entidades lingüísticas?**
- 5.4 ¿Cuál juego de salidas?**

Capítulo 6. Conclusiones

Apéndice. La matemática de la mecánica cuántica

- A.1 Campo de complejos**
- A.2 Espacios y subespacios vectoriales**
- A.3 Producto interno**
- A.4 Operadores lineales**
- A.5 Producto tensorial**

Bibliografía

Prólogo

Toda teoría científica se compone de dos partes fundamentales: un formalismo y una interpretación. Es por medio de la interpretación que se establecen las conexiones conceptuales entre los elementos del formalismo y los elementos de la realidad. Con esas conexiones la interpretación da sentido a la formalización y provee la semántica de la teoría. Esto es de capital importancia para la ciencia, una importancia subestimada por parte de la mayoría de los científicos actuales, y en especial los físicos. Frente a la exactitud de las elaboraciones matemáticas, los análisis conceptuales son a veces dejados de lado como si se tratara del ejercicio de otra disciplina, más propio de la filosofía, cuyas especulaciones son muchas veces consideradas innecesarias para el desarrollo y la aplicación de la física. Esto es un error grave. Ciencia y filosofía se conjugan en la interpretación de la teoría, y un cambio de interpretación, inclusive de una sola fórmula, puede abrir la puerta a nuevas conexiones conceptuales y esto a su vez, a desarrollar toda una nueva teoría. Hay sobrada evidencia histórica de este proceso en la física. La interpretación que ofrece Einstein de las formulas de Lorentz para dar inicio a la relatividad, o la interpretación de Maxwell de la ley de inducción de Faraday al considerar la existencia de un campo eléctrico como el causante de la corriente inducida y completar así las famosas ecuaciones del electromagnetismo, son algunos de los muchos ejemplos. Todos tienen en común que en un punto, el avance de la teoría se basa en la interpretación, en las conexiones conceptuales introducidas al formalismo. Esto no puede ser desestimado. Sólo así será posible trascender los límites ficticios de lo que actualmente es definido como ciencia y como filosofía, de modo de producir una integración más fructífera para ambas.

El presente trabajo intenta ser una contribución en este aspecto, pues hemos brindado respuestas conceptuales como solución a problemas físicos concretos. Aunque modesto, buscamos ofrecer un aporte a la integración entre ciencia y filosofía, bajo el supuesto de que ambas pueden considerarse formas fundamentales a través de las cuales se manifiesta el pensamiento humano.

Este trabajo tuvo su génesis hace ya unos años atrás. Por distintas razones personales que no detallare aquí, quedó postergado hasta el presente año, donde pude brindarle la dedicación adecuada para concretarlo. En el camino he aprendido mucho y en esto le debo agradecer a la Dra. Olimpia Lombardi.

Leonardo Vanni, agosto 2012

Introducción

Desde los orígenes de la física aristotélica en la Antigua Grecia hasta nuestros días, no ha habido un campo de trabajo tan fructífero para la filosofía de la física como el que se presentó en los últimos 100 años. Esto se debe al advenimiento, durante el siglo pasado, de la llamada "física moderna", que incluye esencialmente a la relatividad y la mecánica cuántica. Aunque ambas teorías desafían las nociones newtonianas y nuestra percepción del mundo, es en la mecánica cuántica donde la especulación filosófica ha alcanzado un nivel de riqueza tal que en su accionar ha contribuido notablemente al acercamiento casi perdido entre ciencia y filosofía.

Las reformulaciones que introdujo la mecánica cuántica tuvieron consecuencias tan profundas que han llegado al nivel de repercutir en la propia estructura lógica con la que pensamos el universo. Es por esto que se habla de "lógica cuántica", y hasta de un "mundo cuántico", nada parecido en otra teoría. Pese a los aspectos revolucionarios de la relatividad, rara vez se habla de algo como "mundo relativista", y nunca de una "lógica relativista".

La mecánica cuántica es verdaderamente una teoría muy misteriosa. No sólo por las peculiaridades conceptuales que presenta, sino porque a pesar de eso es una de las más exitosas. Actualmente existe elevado consenso respecto de considerarla universalmente válida, hasta llegar a considerar al mundo clásico como derivado de ella. En esto consiste el llamado límite clásico de la mecánica cuántica, que a diferencia del límite clásico de la relatividad, es un terreno mucho más oscuro y complejo.

Los intentos por encontrar el límite clásico de la mecánica cuántica sólo han conseguidos éxitos parciales y sujetos a restricciones dudosas, o al menos no totalmente generales como se pretende. Esto aplica al formalismo de decoherencia, en particular en su enfoque ortodoxo, Decoherencia Inducida por el Ambiente (EID, por sus siglas en inglés). Como su nombre lo indica, la decoherencia intenta explicar como la "coherencia", palabra que remite a la relación de fase de los componentes de un estado superposición, puede perderse de modo de recuperarse características clásicas para dicho estado, al menos en su interpretación. La idea

básica consiste en asumir al sistema y aparato, entre los cuales se produce la medición, como un sistema abierto en contacto con un tercer sistema, el entorno (environment), típicamente con muchos grados de libertad y capaz de interactuar con los dos primeros. La injerencia de un tercer sistema tiene consecuencias sobre la superposición final del sistema compuesto por sistema + aparato + entorno, de modo que en el límite hipotético de considerar infinitos grados de libertad del entorno, la coherencia cuántica en el sistema resulte eliminada.

Aunque con mucho reconocimiento a nivel práctico, en verdad el formalismo está plagado de dificultades conceptuales y excepciones teóricas, algunas de las cuales profundizaremos en la tesis. De este modo, la formalización de límite clásico por parte de la decoherencia es actualmente más pretendida que real. No obstante, la decoherencia ha tenido gran relevancia histórica como intento de brindar una respuesta conjunta a los llamados problemas de la teoría de la medición cuántica.

En la teoría cuántica de la medición, existen dos problemas centrales: el problema de la lectura definida y el problema de la base preferida. El primer problema consiste en el hecho de que la teoría no puede dar cuenta de valores bien definidos registrados como resultados en la medición, puesto que se predice, como consecuencia de la evolución unitaria, que el estado final del compuesto sistema y aparato es una superposición sin valor definido. El segundo problema consiste en el hecho de que bajo ciertas circunstancias (referidas a la preparación del estado del sistema a medir) la teoría no puede dar cuenta, además, de una base bien definida de estados a la cual pertenece el resultado obtenido en la medición.

La formulación del primer problema, el de la lectura definida, data de los orígenes mismos de la teoría cuántica de la medición, la cual se debe esencialmente a los trabajos de von Neumann en la década del '50 del siglo pasado. Desde entonces ha habido un sinnúmero de propuestas para resolver el problema, algunas apelando a supuestos algo inverosímiles, como la posibilidad de la injerencia de la mente humana en el proceso de medición en la interpretación subjetiva de Wigner, o la posibilidad de múltiples universos en la teoría de muchos mundos de Everett. De todas éstas, la respuesta más sólida, y hoy día más aceptada, consiste en adoptar el llamado postulado del colapso. Este postulado afirma que, después de la evolución unitaria que predice la ecuación de Schrödinger y que establece la correlación

entre sistema y aparato a través de un estado superposición, existe una evolución adicional, no unitaria, que consiste en una proyección a uno de los componentes de la superposición. De este modo, la medición se completa al colapsar el estado a un componente de valor definido.

Si bien el postulado del colapso, con el agregado de la regla de Born de la probabilidad, permite una descripción probabilística de los fenómenos cuánticos que está perfectamente acorde con los resultados experimentales, dicho postulado es altamente problemático desde un punto de vista conceptual. En primer lugar, porque es un postulado adicional, no derivado de la teoría, y en segundo, porque está en contradicción con lo que la teoría predice. Una evolución de tipo colapso viola la ecuación de Schrödinger.

El aporte de la decoherencia al problema de la lectura definida consiste en demostrar que, luego de la interacción con el entorno, y como consecuencia de los muchos grados de libertad de este último, el estado final del sistema objeto puede ser descripto como una mezcla de estados que admite una interpretación por ignorancia. La interpretación por ignorancia concibe el estado mezcla como una mezcla estadística, representando la ignorancia acerca de cuál estado es aquel en el que verdaderamente se encuentra el sistema. El problema de la lectura definida se resolvería si pudiera mostrarse que el sistema efectivamente está en uno de los estados de la mezcla, y con valor definido, sólo que el observador desconoce cuál. Como veremos, la interpretación por ignorancia no es trivialmente aplicable a estados mezclas cuánticos y, por tanto, esta solución no es completamente aceptable.

La formulación del segundo problema, el de la base preferida, se debe prácticamente al formalismo de la decoherencia, y su supuesta solución se presenta como uno de sus logros. Tal como los principales seguidores del formalismo lo presentan, parece un problema diseñado para ser abordado y resuelto por la misma decoherencia. Como en el problema anterior, una vez más el recurso consiste en apelar a la interacción del sistema y aparato con el entorno, de modo que la descomposición del triple sistema compuesto formado por sistema + aparato + entorno elimine la ambigüedad del cambio de base. El fundamento de la demostración descansa en el hecho matemático por el cual una descomposición triortonormal de un vector de estado, a diferencia de la biortonormal, sí es única. Esa base es la única que permite mantener la correlación previamente establecida entre sistema y aparato, de modo

que el registro de la medición no se pierda aun cuando al final se desee ignorar los grados de libertad del entorno.

Estos son, esencialmente, los aportes de la decoherencia a los problemas de la medición.

El objetivo principal de esta tesis consiste en brindar una respuesta a los mencionados problemas sin apelar a decoherencia alguna. Respecto del problema de la base preferida presente en las correlaciones del estado final de la medición, mostraremos que, en lugar de resolverse apelando a una interacción posterior con el entorno, tal como en la teoría de la decoherencia, puede ser resuelto analizando el proceso previo a dicho estado final sin la necesidad del entorno. La esencia del procedimiento consiste en reconsiderar la medición, no como la mera correlación que se manifiesta en el estado final, sino como el proceso capaz de establecerla. Desde esta nueva perspectiva es posible diferenciar distintos procesos que conducen a distintas correlaciones, independientemente de que éstas puedan vincularse matemáticamente mediante un cambio de base. La eliminación de la ambigüedad de la base se logra al determinar qué proceso es caracterizado por el operador de evolución que conecta el estado inicial de la medición con el final. Estrictamente hablado, el problema no queda resuelto, sino más bien disuelto, al considerar que la medición involucra todo el proceso que desemboca en una correlación final, y ello sin la necesidad de agregar interacciones posteriores con sistemas como el entorno. Respecto del problema de la lectura definida nuestro objetivo es un poco más modesto. No vamos a resolver el problema estrictamente sino que, en el marco de ese problema, brindaremos una respuesta a la incompatibilidad que establece el postulado del colapso respecto de las evoluciones unitarias. La estrategia aquí es incorporar los aparatos de medición en el cálculo de las probabilidades condicionales establecidas para una secuencia de dos mediciones. Asumiendo valores definidos para la primera, demostramos que el postulado del colapso sobre el sistema puede ser derivado del propio formalismo de la teoría cuántica, aun cuando el sistema compuesto que incluye a los aparatos nunca abandona la evolución unitaria que predice la ecuación de Schrödinger. Con estos dos núcleos centrales como objetivos, pasamos ahora a detallar cómo se organiza la tesis.

En el Capítulo 1 brindaremos una presentación concisa pero completa de los contenidos conceptuales más relevantes de la mecánica cuántica a modo de introducción formal que capacitará al lector, aun con conocimientos mínimos de la teoría, comprender el desarrollo de los siguientes capítulos. Podemos decir que toda la tesis reposa en las bases conceptuales de este primer capítulo. Asimismo, este primer capítulo reposa a su vez en las bases matemáticas más fundamentales que hemos desarrollado en el Apéndice. Creemos que en dicho apéndice se presenta el compendio matemático más reducido posible necesario para entender completamente cualquier consideración cuántica desarrollada en la tesis. Quien posea inquietudes de justificación más formales puede recurrir a este apéndice en cualquier momento, y a lo largo de toda la tesis haremos referencia a él cada vez que consideremos necesaria una justificación matemática más rigurosa.

En el Capítulo 2 presentaremos la teoría cuántica de la medición como es entendida en nuestros días, sobre la base originalmente desarrollada por von Neumann, de modo de brindar así el marco adecuado para presentar los dos problemas de la medición que ya hemos mencionado. Al final de este capítulo repasaremos las propuestas más destacadas presentadas respecto del problema de la lectura definida. En este repaso podríamos haber mencionado a la decoherencia, pero debido a su peso propio y relevancia histórica vinculada a los dos problemas de la medición, hemos decidido dedicarle un capítulo propio, el próximo.

El Capítulo 3, como acabamos de decir, lo dedicaremos enteramente a la teoría de la decoherencia. Partiendo de algunos antecedentes históricos, explicaremos el formalismo y luego desarrollaremos cómo es aplicado al proceso de medición, para derivar finalmente el modo en que este enfoque pretende dar respuestas a los problemas que hemos mencionado. Al final del capítulo repasaremos algunas de las críticas que se le han formulado y sus debilidades más fundamentales.

En los Capítulos 4 y 5 se corporiza la propuesta central de la tesis.

En el Capítulo 4, en el marco del problema de la lectura definida, brindaremos una solución a la incompatibilidad entre el postulado del colapso y la evolución unitaria que predice la ecuación de Schrödinger. Sin pretender resolver el problema de la lectura definida, es decir, justificar por qué existen valores bien definidos, demostraremos que el postulado del colapso,

en realidad, no es tal, sino que puede ser derivado dentro de la misma teoría y en perfecta concordancia con las evoluciones unitarias, pero consideradas en el sistema completo que incluya todas las partes en interacción.

En el Capítulo 5, daremos respuesta completa al problema de la base preferida. Como ya hemos indicado, no lo resolveremos, sino que más bien lo disolveremos en virtud de considerar la medición como un proceso, y no sólo como la correlación que dicho proceso puede establecer. Para poder rastrear qué proceso puede ser o no involucrado en el establecimiento de la correlación del estado final de la medición, indagaremos la existencia de un hipotético aparato, que hemos llamado "aparato de cuatro luces", que resulta cuánticamente imposible. En virtud de consideraciones sobre este aparato, elaboraremos argumentos que permitirán seleccionar una base como preferida sin apelar a la decoherencia.

Finalmente, en el Capítulo 6 brindaremos las conclusiones de la tesis, donde intentamos repasar los aspectos conceptuales más destacados en relación al camino transitado.

Capítulo 1. Nociones básicas de la mecánica cuántica

Todo sistema cuántico queda caracterizado en un tipo especial de espacio vectorial llamado espacio de Hilbert, en el que se pueden representar matemáticamente los dos aspectos centrales de su descripción: los observables y los estados. Las particularidades técnicas de los espacios de Hilbert y sus características como espacios de vectoriales se desarrollan en el Apéndice A.2. Como en el caso especial de la mecánica cuántica se requiere de la operación de espacios de Hilbert sobre números complejos, puede resultar útil la adicional lectura del Apéndice A.1.

En lo que sigue desarrollaremos cómo pueden ser representados los observables y los estados cuánticos, y los alcances que esta representación adquiere en la elaboración de las descripciones de la teoría.

1.1 Observables y propiedades

En general, ya sea para la mecánica clásica o cuántica, los observables de un sistema son entendidos como aquellas propiedades caracterizadas por magnitudes físicas que están sujetas a evaluación experimental. Diremos que son las “propiedades-tipo” del sistema. Por ejemplo, la posición X , la energía E , etc. Cada una de estas propiedades-tipo se corresponde con un conjunto de “propiedades-caso” o “propiedades de valor”, que son los posibles valores que puede tomar la propiedad tipo al ser medida (para la distinción entre propiedades-tipo y propiedades-caso, bajo la denominación de ‘determinables’ y ‘determinados’, ver Sanford 2011). Si medimos el observable posición como propiedad-tipo, por ejemplo, podemos obtener el valor $x = 2mt$ como propiedad-caso. Pero otra medición del mismo observable puede arrojar otro valor como caso.

El conjunto de valores posibles para un observable es llamado espectro de ese observable. En lo que sigue, consideraremos observables con espectro discreto y finito¹, por lo que sus valores pueden etiquetarse con un índice natural. Diremos que un observable A

¹ La generalización del tratamiento que se sigue en estas páginas para el caso de observables con valores continuos es inmediata, aunque no exenta de sutiles técnicas de carácter matemático las cuales no intervienen en lo tratado en esta tesis.

tiene un espectro de N elementos compuesto por el conjunto de los posibles valores $\{a_i\}$, con $i = 1 \cdots N$.

En una medición de un dado observable, se registrará uno y sólo uno de sus posibles valores. Por lo tanto, cada medición puede ser pensada como una posible pregunta sobre el valor de la magnitud. La respuesta a esa pregunta tiene que ser única. Carecería de sentido físico la medición si, al medir la posición por ejemplo, obtuviéramos $x = 2mt$ y $x = 4mt$. Cada elemento del espectro de un observable A determina una propiedad de valor p_i^A correspondiente a la afirmación 'el observable A tiene valor a_i '. Considerada como pregunta, podemos pensar que esta expresión determina una cuestión experimental, cuya respuesta se obtiene de los resultados de la medición (Hughes 1989, p. 60). La respuesta será afirmativa si el valor obtenido es a_i , y negativa en caso contrario. Por medio de subconjuntos² de valores del espectro es posible especificar propiedades de valor más generales. Diremos que el subconjunto Δ del espectro de un observable A determina la propiedad de valor p_Δ^A correspondiente a la afirmación 'el valor de A pertenece al conjunto Δ '. Como antes, el sistema posee la propiedad si la cuestión experimental puede ser respondida afirmativamente luego de efectuar una medición, y no la posee en caso contrario.

Estas consideraciones valen para cualquier teoría física, son formulaciones sobre propiedades físicas de un sistema sujeto a evaluación experimental. Veremos ahora qué entidades matemáticas representan esas propiedades en el caso que nos interesa: la teoría cuántica. Pues bien, uno de los postulados básicos de la mecánica cuántica es que cada observable A de un sistema físico se puede representar con un operador lineal hermítico A , y cada uno de los posibles valores que puede tomar A en una medición se corresponde con alguno de los autovalores $\{a_i\}$ del correspondiente operador A (ver Apéndice A.4). Por simplicidad, en todo lo que sigue, y como la mayoría de la literatura hace, identificaremos el observable A , con el operador A que lo representa en la teoría, aunque está claro que el

² Los posibles subconjuntos incluyen al vacío y al espectro mismo, los cuales corresponden a la propiedad nula y la universal respectivamente. La primera nunca se cumple, puesto que en toda medición algún valor se obtiene. Por eso mismo, la propiedad universal es la que se cumple siempre.

primero es una propiedad de un sistema físico y el segundo es elemento formal de la teoría matemática que utiliza la mecánica cuántica.

Como se señala en el Apéndice A.4, el hecho de que el operador sea hermítico asegura que sus autovalores sean reales, como se requiere para el valor de una magnitud física. Cada autovalor a_i de un observable (operador) A se corresponde a un subespacio unidimensional generado por el autovector correspondiente $|a_i\rangle$ en el espacio de Hilbert, y dicho subespacio a su vez se corresponde con el proyector $\Pi_{a_i} = |a_i\rangle\langle a_i|$ que participa en la descomposición espectral de A (ver Apéndice A.4). Es esto lo que nos da la pauta de cómo pueden ser representadas las propiedades de valor del observable A . La propiedad p_i^A correspondiente a la cuestión experimental 'el observable A tiene valor a_i ', será representada en el formalismo de la teoría cuántica por el proyector $\Pi_{a_i} = |a_i\rangle\langle a_i|$ correspondiente al subespacio generado por $|a_i\rangle$. En forma más general, la propiedad p_Δ^A correspondiente a la cuestión 'el valor de A pertenece al conjunto Δ ', donde Δ es un subconjunto de valores del espectro, podrá ser representada por el proyector $\Pi_\Delta = \sum_{a_i \in \Delta} |a_i\rangle\langle a_i|$ correspondiente al subespacio generado por el conjunto de autovectores $\{|a_i\rangle\}$, con $a_i \in \Delta$.

La capacidad predictiva de esta forma de representar propiedades quedará evidenciada más claramente con la noción de estado, noción que permite asignar probabilidades a dichas propiedades.

1.2 Estados y probabilidades

Si los valores para los observables determinan las cuestiones experimentales del sistema físico, ya sea para la mecánica clásica o cuántica, el estado de dicho sistema es entendido como el conjunto de respuestas que pueden ser asignadas a esas posibles cuestiones experimentales. Debido a la naturaleza determinista de la mecánica clásica (Hughes 1989, Cap. 3), cada cuestión experimental puede ser respondida por sí o no en esa teoría. Esto significa que, en el caso clásico, un estado es una función que asigna uno o cero a cada cuestión experimental: uno si la respuesta es afirmativa, y cero si no lo es. Debido a ello, desde el punto de vista clásico, el estado determina con certeza todas las propiedades de

valor del sistema en un cierto instante, e inversamente, el conjunto de todas las propiedades de valor en un cierto instante determina el estado del sistema.

En la mecánica cuántica esto no es así (Hughes 1989, Caps. 2 y 3). Es necesario recurrir a una descripción probabilística donde el estado es entendido como una función o medida de probabilidad, que asigna a cada cuestión experimental un valor entre cero y uno. Cuánticamente hablando, el estado no determina con certeza todas las propiedades de valor de las propiedades-tipo del sistema.

Para estados llamados puros³, la entidad matemática capaz de determinar esta función de probabilidad es un vector en el espacio de Hilbert que caracteriza al sistema (ver Apéndice A.2). Por requerimientos de normalización de la probabilidad (la exigencia de que la suma de las probabilidades de todas las posibilidades de una medición sume uno), se requiere que los vectores de estado estén normalizados a la unidad. De la definición de norma de un vector (Apéndice A.3), y considerando el caso general de un espacio de Hilbert sobre los complejos (Apéndice A.2), el requerimiento de normalización significa que dos vectores que difieren sólo en una fase (es decir en un número complejo de módulo uno) representaran el mismo estado cuántico. Esto puede comprenderse a partir de la propia definición de la probabilidad cuántica, tal como sigue. Cada vector normalizado en el espacio de Hilbert, que simbolizamos como $|v\rangle$, induce una función de probabilidad para cada propiedad p_Δ representada por el proyector Π_Δ dada por

$$P_v(p_\Delta) = \langle v | \Pi_\Delta | v \rangle = \langle v | \Pi_\Delta \Pi_\Delta | v \rangle = \|\Pi_\Delta | v \rangle\|^2 \quad (1.1)$$

donde se ha utilizado el hecho que un proyector es idempotente en la segunda igualdad, y la definición de norma de un vector, en la tercera (Apéndice A.4). El vector $|v\rangle$ normalizado es llamado vector de estado.

De la acción de un proyector sobre un vector (Apéndice A.4), es posible asociar un significado geométrico a la última ecuación: un vector de estado determina una función de probabilidad para una propiedad, que es el cuadrado en norma de la proyección de dicho

³ Entre los estados de la mecánica cuántica se distinguen estados puros y estados mezcla. De estos últimos hablaremos en una sección más adelante.

vector sobre el subespacio asociado a la propiedad. Esto puede ser mejor entendido con el siguiente ejemplo. Por simplicidad consideremos un observable A no degenerado en un espacio de Hilbert sobre los reales de dimensión 2. Existirán dos propiedades p_i^A correspondientes a los autovalores a_i representadas por los proyectores que llamaremos Π_i , con $i=1,2$. Estos proyectores proyectan sobre subespacios unidimensionales (rectas) generados por cada uno de los correspondientes autovectores $\{|a_i\rangle\}$. Por tratarse de un operador hermítico, estos autovectores son ortogonales y determinan una base del espacio (ver Figura 1).

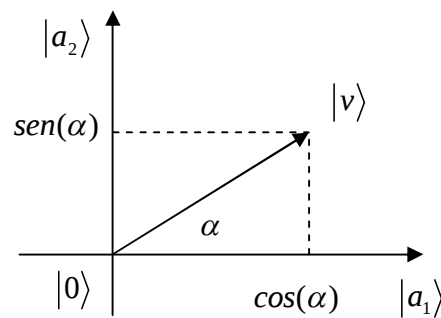


Figura 1

Supongamos que el vector $|v\rangle$ forma un ángulo α respecto de la recta horizontal, que en la figura asumimos asociada al proyector Π_1 , es decir, al subespacio generado por $|a_1\rangle$. Por ortogonalidad, la vertical corresponderá al subespacio generado por $|a_2\rangle$. De la ecuación (1.1), para cada una de las propiedades elementales p_i^A , con $i=1,2$ tenemos

$$\begin{aligned} P_v(p_1^A) &= \|\Pi_1 |v\rangle\|^2 = \cos^2(\alpha) \\ P_v(p_2^A) &= \|\Pi_2 |v\rangle\|^2 = \sin^2(\alpha) \end{aligned} \quad (1.2)$$

Como el vector $|v\rangle$ tiene norma igual a uno, las funciones $\sin(\alpha)$ y $\cos(\alpha)$ son los componentes de sus proyecciones sobre los ejes coordenados. Es fácil ver que la definición dada por la ecuación (1.1) cumple con los axiomas de Kolmogorov que se esperan para una función de probabilidad (Hughes 1989, p. 88).

En el caso particular en el que el vector de estado coincide con uno de los autovectores del observable A , por ejemplo $|v\rangle = |a_1\rangle$, es claro de la ecuación (1.2) que

$$P_v(p_1^A) = 1$$

$$P_v(p_2^A) = 0$$

es decir, la probabilidad es uno para la propiedad asociada al autovalor cuyo autovector es vector de estado, y cero en cualquier otro caso. Dicho de otro modo, si el sistema es preparado de tal manera que su estado coincida con uno de los autovectores de A , entonces existe la certeza de obtener el correspondiente autovalor al efectuar una medición de A . Esto justifica el llamado vínculo autovector-autovalor (eigenstate-eigenvalue link, ver Wilce 2006).

Cabe señalar que en este ejemplo hemos considerado cuestiones experimentales referidas a las propiedades de valor de un solo experimento: la medición del observable A . Pese a esto, la expresión dada por la fórmula (1.1) asigna una probabilidad a todas las cuestiones experimentales de todos los observables posibles en el espacio de Hilbert. Sin embargo, se puede demostrar que sobre este espacio muestral más general, que incluye propiedades de valor de distintos observables, esta fórmula no cumple con los axiomas de Kolmogorov, por lo que, estrictamente desde un punto de vista matemático, no es una medida de probabilidad legítima (Mittelstaedt 1998, p. 92; Griffiths 2002).

1.3 Estados puros y estados mezcla

Los estados con los que hemos tratado hasta ahora son estados puros; son los estados determinados por vectores normalizados el espacio de Hilbert. No obstante, estos estados no agotan todos los posibles estados que podemos considerar. Es posible establecer medidas de probabilidad para todas las cuestiones experimentales del sistema que no sean inducidas por un vector. Éstas se logran combinando adecuadamente funciones de probabilidad provenientes de distintos vectores de estado.

Consideremos la medición de un observable A cuyo espectro viene dado por los autovalores $\{a_i\}$ que determinan las cuestiones experimentales referidas a las propiedades

p_Δ^A , siendo Δ un subconjunto del espectro de A . Las propiedades elementales asociadas a cada autovalor a_i son representadas por los subespacios generados por los respectivos autovectores $|a_i\rangle$. Según la ecuación (1.1), cada uno de estos vectores, correspondientes a estados puros, especifica la medida de probabilidad dada por $P_i(p_\Delta^A) = \langle a_i | \Pi_\Delta | a_i \rangle$. Se puede demostrar que la combinación de estas medidas, efectuada de la siguiente manera,

$$P(p_\Delta^A) = \sum_i b_i P_i(p_\Delta^A), \quad \text{con } b_i \geq 0 \text{ y } \sum_i b_i = 1, \quad (1.3)$$

es también una medida de probabilidad, es decir, es una función de probabilidad que cumple con los axiomas de Kolmogorov (Ballentine 1998, Cap. 2). Es fácil ver que esta probabilidad no viene de un vector, por lo que debe ser pensada como proveniente un nuevo tipo de estado. Si proviniera de un vector, tendría que existir algún $|v\rangle$ tal que

$$P(p_\Delta^A) = \sum_i b_i \langle a_i | \Pi_\Delta | a_i \rangle = \langle v | \Pi_\Delta | v \rangle$$

Pues bien, de la propia sumatoria vemos que, a menos que $|v\rangle$ sea uno de los $|a_i\rangle$, no es posible satisfacer esta última ecuación.

Los estados que determinan una función de probabilidad como la dada en la ecuación (1.3), es decir, formada por una combinación de probabilidades provenientes de estados puros, son llamados estados mezcla (Hughes 1989, p. 93; Ballentine 1998, Cap. 2), y la combinación específica indicada en esa ecuación es llamada combinación convexa (Hughes 1989, p. 143; Ballentine 1998, p. 37).

Con todo esto en mente, el problema inmediato que se presenta es determinar con qué elemento matemático representar este nuevo tipo de estados. Pues bien, hay una formulación en términos de operadores de estados que es sencilla, elegante, y que además unifica el tratamiento que hemos abordado para estados puros. Partamos de considerar la probabilidad $P_i(p_\Delta^A) = \langle a_i | \Pi_\Delta | a_i \rangle$ sobre cualquier cuestión experimental referida a la medición del observable A inducida por el estado puro $|a_i\rangle$ que es autovector de A . Consideremos además al proyector Π_{a_j} asociado al subespacio generado por el vector $|a_j\rangle$. Como los

vectores $|a_i\rangle$ son ortonormales, tenemos que $\Pi_{a_j}|a_i\rangle$ es igual a $|a_i\rangle$ si $i = j$, y cero si $i \neq j$. Si en este punto introducimos la función delta de Kronecker, dada por

$$\delta_{ij} = \begin{cases} 1 & \text{si } i = j \\ 0 & \text{si } i \neq j \end{cases}$$

se obtiene que $\Pi_{a_i}|a_j\rangle = \delta_{ij}|a_j\rangle$, por lo que

$$\begin{aligned} P_i(p_\Delta^A) &= \langle a_i | \Pi_\Delta | a_i \rangle \\ &= \sum_j \langle a_j | \Pi_\Delta \delta_{ij} | a_j \rangle \\ &= \sum_j \langle a_j | \Pi_\Delta \Pi_{a_i} | a_j \rangle = \text{Tr}[\Pi_\Delta \Pi_{a_i}] \end{aligned} \quad (1.4)$$

Hemos conseguido así una nueva fórmula para calcular la probabilidad de la propiedad p_Δ^A cuando el sistema está en el estado puro $|a_i\rangle$. Como antes el operador Π_Δ representa la propiedad p_Δ^A , pero la información del estado puro $|a_i\rangle$ esta ahora contenida en el operador Π_{a_i} .

Esta estrategia nos permite considerar una nueva representación de los estados puros en términos de operadores: los estados puros quedarán representados por los operadores proyectores sobre el subespacio generado por el vector que representa el estado puro. La utilidad de este tipo de representación de estados en términos de operadores reside en su posibilidad de generalización para estados mezclas. Para poner de manifiesto lo dicho, hagamos uso de la ecuación (1.4) en la fórmula (1.3) para calcular la probabilidad inducida por un estado mezcla

$$\begin{aligned} P(p_\Delta^A) &= \sum_i b_i P_i(p_\Delta^A) \\ &= \sum_i b_i \text{Tr}[\Pi_\Delta \Pi_{a_i}] \\ &= \text{Tr} \left[\Pi_\Delta \sum_i b_i \Pi_{a_i} \right] \end{aligned}$$

donde hemos hecho uso de la propiedad para la traza de una suma de operadores (Apéndice A.4). Esta última ecuación tiene exactamente la misma forma que la última ecuación que hemos usado en (1.4) para calcular la probabilidad de estados puros, si definimos el operador de estado como

$$\rho = \sum_i b_i \Pi_{a_i}$$

Este operador es una combinación convexa de operadores correspondientes a estados puros, y así la medida de probabilidad que induce es también una combinación convexa de las medidas de probabilidades inducidas por estados puros. Por tratarse de este tipo de combinación, todo operador de estado resulta hermítico; además se cumple $Tr[\rho] = 1$, que es la condición de normalización que se requiere para la probabilidad (Ballentine 1998, Cap. 2).

Ya sea ρ un estado puro correspondiente a un proyector, o un estado mezcla correspondiente a una combinación convexa de proyectores, se obtiene la fórmula general para la probabilidad de una propiedad p_Δ^A en el estado ρ dada por

$$P_\rho(p_\Delta^A) = Tr[\Pi_\Delta \rho]$$

La construcción realizada para definir los estados mezcla sugiere la siguiente interpretación que, como veremos, resulta cuestionable: los estados mezcla representan una mezcla estadística de los posibles estados puros en los que puede encontrarse el sistema. Los distintos coeficientes con los que se pesa cada estado puro brindan una medida de la ignorancia que poseemos acerca del verdadero estado en el que se encuentra el sistema.

Si esto fuera así, un estado mezcla nos permitiría calcular probabilidades que resultarían ser producto de una indeterminación meramente gnoseológica y no ontológica (Cohen 1989, p. 41), pues no surgirían de una regularidad subyacente sino de una limitación de nuestro conocimiento. Sus valores cambiarían si se agregara conocimiento sobre el sistema. Esta interpretación es llamada interpretación por ignorancia (Hughes 1989, p. 96; Mittelstaedt 1998, p. 16; Ballentine 1998, p. 39), y si bien es como clásicamente se interpretarían las probabilidades calculadas en la mecánica estadística, en el caso cuántico tal interpretación

presenta severas limitaciones que tienen que ver con la indeterminación del conjunto de estados puros con los que se puede expresar un mismo estado mezcla. En efecto, la descomposición de un estado mezcla en términos de una combinación convexa de estados puros de la forma $\rho = \sum b_i \Pi_i$ en general no es única. Existen infinitas descomposiciones distintas de la forma $\rho = \sum b'_i \Pi'_i$ por medio de Π'_i no ortogonales (Mittelstaedt 1998, p. 79; Ballentine 1998, p. 39). Pero aun en el caso en que decidamos privilegiar sólo las combinaciones ortogonales para hacer uso del teorema de la descomposición espectral sobre ρ , todavía es posible encontrarse con una falta de unicidad en el desarrollo de ρ si éste es degenerado (Mittelstaedt 1998, p. 79). En otras palabras, si existe al menos un b_i igual a un b_j para $i \neq j$, entonces es posible reemplazar $\Pi_i + \Pi_j$ por la suma de otros proyectores $\Pi'_i + \Pi'_j$ también ortogonales, tal que $\Pi'_i + \Pi'_j = \Pi_i + \Pi_j$ (Hughes 1989, p. 139). En general, se tiene $\rho = \sum b_i \Pi_i = \sum b'_i \Pi'_i$. Esta falta de unicidad establece una indeterminación en el conjunto de estados puros que componen un estado mezcla.

Estas propiedades formales tienen consecuencias en la pretendida interpretación por ignorancia. Debido a la falta de unicidad mencionada, un estado mezcla no sólo representaría la ignorancia sobre cuál estado es aquel en el que verdaderamente se encuentra el sistema, sino una ignorancia mucho mayor, que involucra el desconocimiento del conjunto de estados posibles que podemos asignar al sistema.

Otra objeción para la interpretación por ignorancia, aún en el caso de estados mezcla no degenerados, tiene que ver con los estados de las partes de un sistema compuesto, lo cual será abordado en la siguiente subsección.

1.4 Estados reducidos de sistemas compuestos

El formalismo de operadores de estado es sumamente útil en la descripción de los estados de subsistemas de un sistema compuesto. Primero hacemos notar que si tenemos un sistema S_A caracterizado por un espacio de Hilbert H_A , y otro sistema S_B caracterizado por un espacio de Hilbert H_B , entonces el sistema $S = S_A + S_B$ compuesto por S_A y S_B queda caracterizado por un espacio Hilbert $H = H_A \otimes H_B$ que es el producto tensorial de H_A y H_B (Apéndice A.5). Adicionalmente, si A es un operador que representa un observable, por

ejemplo, en el sistema S_A , entonces el operador $A \otimes I_B$ es el operador que lo representa en el sistema compuesto (Hughes 1989, p. 149), siendo I_B el operador identidad en el sistema S_B . Por consiguiente, si $A = \sum a_i \Pi_{a_i}^A$ es la descomposición espectral de A en el subsistema S_A , entonces $A \otimes I_B = \sum_i a_i \Pi_{a_i}^A \otimes I_B$ es la descomposición espectral del correspondiente operador del sistema compuesto. De este modo, cada cuestión experimental del sistema compuesto, pero referida a la medición del operador A en el sistema S_A correspondiente a la propiedad p_Δ^A , 'el valor de A pertenece al conjunto Δ ', estará dada por el proyector $\Pi_\Delta^A \otimes I_B$ donde, como ya hemos considerado, Δ es algún subconjunto del espectro de A .

Si el sistema compuesto $S = S_A + S_B$ es preparado en un estado ρ_{AB} , puro o no, se puede demostrar que cada uno de los subsistemas por separado queda descrito por uno de los siguientes estados

$$\begin{aligned}\rho_A &= Tr_B[\rho_{AB}] \\ \rho_B &= Tr_A[\rho_{AB}]\end{aligned}$$

Esto es así porque la probabilidad para cada cuestión experimental referida a la medición de A en el sistema compuesto, calculada por medio de $\Pi_\Delta^A \otimes I_B$ en el estado ρ_{AB} , es la misma probabilidad calculada por medio Π_Δ^A con el estado ρ_A del subsistema S_A . Es decir, se demuestra que

$$\begin{aligned}P(p_\Delta^A) &= Tr_A[\rho_A \Pi_\Delta^A] \\ &= Tr[\rho_{AB} \Pi_\Delta^A \otimes I_B]\end{aligned}$$

Supongamos ahora que el sistema compuesto es preparado en un estado puro que puede adoptar la forma

$$|\Psi_{AB}\rangle = \sum_i c_i |a_i\rangle |b_i\rangle$$

lo cual sigue de la descomposición de Schmidt (Apéndice A.5). Este vector corresponde a un estado puro cuyo operador de estado es $\rho_{AB} = |\Psi_{AB}\rangle \langle \Psi_{AB}|$. No es difícil mostrar que, en este caso,

$$\rho_A = \text{Tr}_B[\rho_{AB}] = \sum_i |c_i|^2 |a_i\rangle\langle a_i|$$

$$\rho_B = \text{Tr}_A[\rho_{AB}] = \sum_i |c_i|^2 |b_i\rangle\langle b_i|$$

Estas ecuaciones manifiestan una peculiaridad interesante que permite formular otra objeción que se suma a las ya planteadas a la interpretación por ignorancia para estados mezcla. Si la interpretación por ignorancia fuera válida, los estados mezcla involucrarían menor información que los estados puros, pues en un estado mezcla no sabemos en cuál de los posibles estados puros que componen la mezcla se encuentra el sistema. Bajo esta consideración, las últimas ecuaciones nos muestran que los estados de las partes tienen menos información que el estado de la totalidad. Dicho de otra manera, especificar los estados de todas las partes no basta para obtener el estado del sistema compuesto por esas partes. Si la interpretación por ignorancia fuera correcta, entonces el subsistema A estaría en alguno de los estados representados por $\rho_{a_i} = |a_i\rangle\langle a_i|$. Asimismo, el subsistema B estaría en algunos de los $\rho_{b_k} = |b_k\rangle\langle b_k|$. Por consiguiente, el estado compuesto debería estar en una mezcla de estados formada con $\rho_{a_i} \otimes \rho_{b_k}$. Pero esto conduce a una contradicción, pues partimos del supuesto de que el sistema compuesto se encontraba en el estado puro representado por $\rho_{AB} = |\Psi_{AB}\rangle\langle\Psi_{AB}|$. Este argumento demuestra que la interpretación por ignorancia no es admisible para interpretar estados mezcla en la mecánica cuántica.

1.5 El principio de superposición

El principio de superposición de estados pone de manifiesto uno de los aspectos fundamentales de la mecánica cuántica que la diferencia de la mecánica clásica, y consiste en su carácter netamente probabilístico. El principio se puede enunciar de la siguiente manera.

Si $|v_1\rangle$ y $|v_2\rangle$ son vectores de estado, entonces la combinación lineal dada por

$|v\rangle = a_1|v_1\rangle + a_2|v_2\rangle$, de modo tal que $\| |v\rangle \| = 1$, también es un vector de estado.

Este principio tiene consecuencias en el carácter probabilístico de la teoría. Supongamos una vez más el ejemplo del experimento consistente en la medición de un observable A no degenerado en un espacio de Hilbert de dimensión dos, y cuyo espectro está dado por los autovalores $\{a_i\}$ correspondientes a los autovectores $\{|a_i\rangle\}$, con $i=1,2$. Cada uno de estos autovectores representan estados puros que determinan las medidas de probabilidad $P_{a_i}(p_\Delta^A) = \langle a_i | \Pi_\Delta | a_i \rangle$ para las posibles propiedades p_Δ^A referidas a las cuestiones experimentales correspondientes a la medición de A . Como ya hemos indicado, asumimos Δ como algún subconjunto del espectro de A , de modo que las propiedades p_Δ^A quedan representadas por

$$\Pi_\Delta = \sum_{a_i \in \Delta} |a_i\rangle \langle a_i|$$

Consideremos el vector $|v\rangle = \alpha_1 |a_1\rangle + \alpha_2 |a_2\rangle$, con $\|v\| = 1$. Entonces, por la manera en que están formados los proyectores Π_Δ , se obtiene

$$\begin{aligned} P_v(p_\Delta^A) &= \langle v | \Pi_\Delta | v \rangle \\ &= (\alpha_1^* \langle a_1 | + \alpha_2^* \langle a_2 |) \Pi_\Delta (\alpha_1 |a_1\rangle + \alpha_2 |a_2\rangle) \\ &= |\alpha_1|^2 \langle a_1 | \Pi_\Delta | a_1 \rangle + |\alpha_2|^2 \langle a_2 | \Pi_\Delta | a_2 \rangle \\ &= |\alpha_1|^2 P_{a_1}(p_\Delta^A) + |\alpha_2|^2 P_{a_2}(p_\Delta^A) \end{aligned}$$

Como $\|v\| = 1$, vale que $\sum |\alpha_i|^2 = 1$. Esto significa que $P_v(p_\Delta^A)$ se obtiene como una combinación convexa de las probabilidades $P_{a_i}(p_\Delta^A)$. Pero por el principio de superposición se tiene que, si $|a_1\rangle$ y $|a_2\rangle$ son estados puros posibles, entonces $|v\rangle = \alpha_1 |a_1\rangle + \alpha_2 |a_2\rangle$ también lo es. Por lo tanto, tenemos un estado puro cuya probabilidad puede obtenerse como combinación convexa de probabilidades provenientes de otros estados puros.

Esto no tiene antecedente clásico. Clásicamente, combinaciones convexas de probabilidades provenientes de estados puros es siempre una probabilidad inducida por un estado mezcla y, por consiguiente, con valores entre cero y uno.

Una teoría donde vale el principio de superposición no puede ser una teoría determinista, ya que un estado puro que puede determinar valores ceros y unos para ciertas cuestiones

experimentales puede ser visto como una combinación convexa de estados asociados a otras cuestiones experimentales cuyas funciones de probabilidad (por tratarse de una combinación convexa) no arrojan certezas, sino valores entre cero y uno. Desde el punto de vista cuántico, al valer el principio de superposición, por más que tratemos con estados puros, éstos no pueden determinar certezas para todas las cuestiones experimentales.

1.6 El principio de incerteza

Enunciado originalmente por Heisenberg, el principio de incerteza (o principio de indeterminación) establece la imposibilidad de determinar simultáneamente y con una exactitud arbitraria, valores obtenidos en la medición de observables incompatibles (Hughes 1989, pp. 265-270; Ballentine 1998, pp. 165-187). En el marco de la mecánica cuántica, este principio queda formalizado mediante la llamada relación de incerteza, que impone una cota al producto de las incertezas de los observables medidos (Hughes 1989, p. 263; Ballentine 1998, p. 166). Aunque en muchos textos no se discute, se presenta aquí el problema de considerar a qué refiere exactamente la palabra ‘incerteza’ en este contexto (Hughes 1989, p. 266). La lectura usual, y de la que proviene la mencionada denominación del principio, consiste en atribuirle un significado estadístico, es decir considerando a la incerteza como la varianza⁴ de la magnitud medida, esto es, como la incerteza estadística que se presenta como resultado de limitaciones propias de las mediciones cuánticas. Otra lectura le atribuye un significado ontológico, y considera a la incerteza como una indefinición intrínseca en los valores simultáneos que pueden tomar los observables incompatibles, independientemente de que éstos sean medidos o no. Cualquiera sea la lectura elegida, es posible demostrar que la relación de incerteza es consecuencia de la existencia de observables incompatibles en la teoría cuántica (Hughes 1989, pp. 265-270; Ballentine 1998, pp. 165-187; Sakurai 1985, p. 35).

Desde un punto de vista formal, dos observables A y B se dicen incompatibles si sus correspondientes operadores A y B no conmutan. Que los operadores no conmutan significa que $AB - BA \neq 0$, lo cual se simboliza como $[A, B] \neq 0$. Si dos operadores no conmutan,

⁴ También llamada “dispersión cuadrática media”.

entonces algunos de los proyectores que participan en sus desarrollos espectrales tampoco lo hacen. Esto nos permite hablar de subespacios incompatibles. Diremos que dos subespacios son incompatibles si sus proyectores asociados no conmutan.

Geoméricamente, es fácil ver que dos subespacios son incompatibles si no son ni paralelos ni ortogonales. En subespacios unidimensionales, dos subespacios son incompatibles si sus rectas son oblicuas. Aunque más difícil de ver, la idea permanece para subespacios de más dimensiones (Hughes 1989, p. 104). Identificando, como ya hemos señalado, los observables con sus correspondientes operadores, tenemos que dos observables A y B son incompatibles si existe al menos un autovector $|a_i\rangle$ de A que genera un subespacio oblicuo al subespacio generado por algún autovector $|b_i\rangle$ de B . De este modo el proyector $\Pi_{a_i} = |a_i\rangle\langle a_i|$ no conmutará con $\Pi_{b_i} = |b_i\rangle\langle b_i|$ y, por supuesto, tampoco A conmutará con B .

Sin apelar a la demostración general, la existencia de una relación de incerteza entre observables incompatibles puede ser fácilmente visualizada si analizamos dos observables en un espacio de Hilbert de dimensión dos. Supongamos que $\{|a_i\rangle\}$, con $i=1,2$, son los autovectores de un observable A que generan los subespacios asociados a las propiedades de valor $\{a_i\}$, y $\{|b_i\rangle\}$, con $i=1,2$, son los autovectores de otro observable B que generan los subespacios asociados a las propiedades de valor $\{b_i\}$. Si A y B son incompatibles, entonces los subespacios generados por los $\{|a_i\rangle\}$ son oblicuos a los generados por $\{|b_i\rangle\}$. Esto es representado en la Figura 2.

Bajo esta condición, tenemos que si el sistema es preparado, por ejemplo, en $|b_1\rangle$, de acuerdo con la proyección del vector de estado $|b_1\rangle$ sobre los $\{|b_i\rangle\}$, habrá certeza (probabilidad igual a cero o uno) para la propiedad de valor b_1 referida a la medición de B pero, de acuerdo con la proyección del mismo vector $|b_1\rangle$ sobre los $\{|a_i\rangle\}$, habrá indeterminación (probabilidades entre cero y uno) para las propiedades de valor a_i referidas a la medición de A . Si, por otro lado, preparamos el sistema en algunos de los $\{|a_i\rangle\}$ de modo de tener certeza para las propiedades de valor de A , entonces tenemos indeterminación en las propiedades de B . Por tener autovectores que generan subespacios oblicuos, no existe un estado que sea paralelo o perpendicular a todos los subespacios considerados y que, por

consiguiente, determine certezas para las propiedades de valor de A y B simultáneamente. Esta es la esencia geométrica del principio de incerteza.

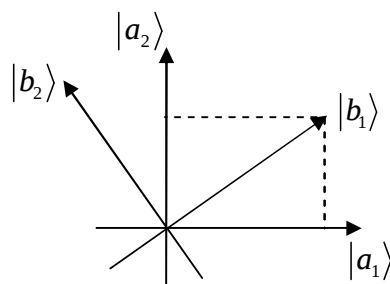


Figura 2

Cabe señalar que el principio de incerteza es independiente del principio de superposición. No obstante, sin el principio de incerteza, es decir, en el caso en que todos los observables fueran compatibles, el principio de superposición no tendría interpretación física (Hughes 1989, pp. 108-111). Esto se debe a que, en ese caso, los estados puros provenientes de superposiciones no podrían distinguirse de estados mezcla, en el sentido que producirían las mismas probabilidades que estados mezclas.

En efecto, es en virtud de la existencia de observables incompatibles, es decir de observables con propiedades de valor asociadas a subespacios oblicuos, que se pueden diferenciar estados puros, provenientes de una superposición, de estados mezcla, provenientes de una combinación convexa. Para explicitar esta afirmación, volvamos al ejemplo de dos observables incompatibles A y B en un espacio de dimensión dos, cuyas propiedades de valor corresponden a los subespacios indicados por los ejes de la Figura 2. Como consecuencia de tal disposición se tiene, por ejemplo, que los $\{|b_i\rangle\}$ pueden expresarse como combinaciones lineales no triviales de los $\{|a_i\rangle\}$, es decir $|b_i\rangle = c_{i1}|a_1\rangle + c_{i2}|a_2\rangle$, $i = 1, 2$, con todos los coeficientes c_{ij} distintos de cero, lo cual asegura que cada $|b_i\rangle$ no es paralelo ni ortogonal a alguno de los $|a_i\rangle$. Por lo tanto, $|b_i\rangle$ es un estado puro superposición de los $\{|a_i\rangle\}$. Para toda propiedad p_A^A referida a la medición de A y representada por el proyector

$$\Pi_{\Delta}^A = \sum_{a_i \in \Delta} |a_i\rangle\langle a_i|$$

la probabilidad $P_{|b_i\rangle}(p_{\Delta}^A) = \langle b_i | \Pi_{\Delta}^A | b_i \rangle$ inducida por $|b_i\rangle$, que es superposición de los $\{|a_i\rangle\}$, no puede distinguirse de la probabilidad proveniente de una mezcla de esos estados $\{|a_i\rangle\}$. No es difícil demostrar que, siendo $|b_i\rangle$ como se indicó arriba, la probabilidad $P_{|b_i\rangle}(p_{\Delta}^A)$ puede calcularse

$$P_{|b_i\rangle}(p_{\Delta}^A) = |c_{i1}|^2 P_{|a_1\rangle}(p_{\Delta}^A) + |c_{i2}|^2 P_{|a_2\rangle}(p_{\Delta}^A)$$

Es sólo por medio del cálculo de las probabilidades de las propiedades p_{Δ}^B referidas a la medición de B y representadas por los proyectores $\Pi_{\Delta}^B = \sum_{b_i \in \Delta} |b_i\rangle\langle b_i|$, que puede demostrarse

$$P_{|b_i\rangle}(p_{\Delta}^B) \neq |c_{i1}|^2 P_{|a_1\rangle}(p_{\Delta}^B) + |c_{i2}|^2 P_{|a_2\rangle}(p_{\Delta}^B)$$

Es decir que si no existiese B incompatible con A , de modo tal que sus propiedades correspondieran a subespacios oblicuos, la fórmula de la probabilidad calculada con un estado puro coincidiría con la calculada con uno mezcla. Esto significa que no sería necesario suministrar un contenido conceptual y empírico a la superposición, pues toda superposición podría ser tratada como mezcla.

1.7 La evolución de los estados cuánticos

Hasta ahora hemos hablado de los estados de la mecánica cuántica en su carácter sincrónico, es decir, en su capacidad para determinar las situaciones experimentales en un dado instante (Bub 1997, p. 13). Por su puesto, esto no basta para las descripciones físicas. Es necesario también establecer su carácter diacrónico, es decir, su capacidad para determinar el desarrollo de tales situaciones en el tiempo. Es en este aspecto que hablamos de la evolución dinámica de los sistemas. No hay física verdaderamente si no podemos hablar del contenido dinámico del estado, capaz de establecer la evolución temporal de los sistemas y ser reflejado en las situaciones experimentales.

La ecuación dinámica para los estados en la mecánica cuántica está gobernada, al igual que en la mecánica clásica, por el llamado hamiltoniano, que se asocia a la energía del sistema (Ballentine 1998, pp. 67-68; Hughes 1989, p. 198, p. 77). No obstante, la energía, como toda magnitud cuántica, queda representada por un operador llamado hamiltoniano, H , que es el operador cuyos autovalores son los valores de la energía del sistema. El hamiltoniano determina la tasa de cambio del vector de estado $|\varphi\rangle$ de un sistema de la siguiente manera

$$i\hbar \frac{\partial}{\partial t} |\varphi\rangle = H |\varphi\rangle$$

Esta ecuación es la llamada ecuación de Schrödinger (Ballentine 1998, pp. 67-68; Hughes 1989, p. 77). Si H es independiente del tiempo, esta ecuación se puede integrar fácilmente haciendo uso de la teoría de ecuaciones diferenciales, obteniéndose

$$|\varphi\rangle = e^{-i/\hbar H(t-t_0)} |\varphi_0\rangle$$

Donde $|\varphi_0\rangle$ es el estado a un tiempo t_0 . El operador $U(t) = e^{-i/\hbar H(t-t_0)}$, que aplicado al estado al tiempo inicial da como resultado el estado a un tiempo posterior, es llamado operador de evolución del sistema (Ballentine 1998, pp. 67-68; Hughes 1989, p. 77). Como el hamiltoniano es hermítico, se puede probar que $U(t)$ es unitario, lo cual significa que $U^{-1}(t) = U^\dagger(t)$.

Dado el hamiltoniano, queda establecido el operador de evolución, y esto determina el cambio del vector de estado, de modo que $|\varphi\rangle = U(t)|\varphi_0\rangle$. Si se conoce cómo cambia un vector de estado, es posible saber cómo cambia el correspondiente operador de estado haciendo uso de la notación de Dirac, al dar cuenta que si $|\varphi\rangle = U(t)|\varphi_0\rangle$, entonces $\langle\varphi| = \langle\varphi_0|U^\dagger(t)$. Por lo tanto, el operador $\rho_0 = |\varphi_0\rangle\langle\varphi_0|$, correspondiente al vector $|\varphi_0\rangle$, cambia como

$$\rho = U(t)\rho_0 U^\dagger(t)$$

Es interesante señalar que, según la ecuación de Schrödinger, el estado evoluciona en forma determinista desde un estado conocido a un tiempo inicial. Esto no contradice lo que hemos afirmado acerca del carácter netamente probabilístico de la mecánica cuántica si recordamos que, cuánticamente hablando, el estado no determina las respuestas a todas las cuestiones experimentales, sino la probabilidad de esas respuestas. Por consiguiente, si bien el estado queda completamente determinado para todo tiempo posterior al inicial, no es posible determinar qué propiedades del sistema se obtendrán en los experimentos realizados en tiempos posteriores a ese tiempo inicial. Esto significa que, al contrario que la mecánica clásica, la evolución del estado no es la evolución “verdadera” del sistema. La evolución del sistema viene dada experimentalmente por la específica sucesión de propiedades obtenidas en las mediciones. Esto da pie a mencionar los llamados formalismos de historias cuánticas, en los cuales se estudia la evolución de los sistemas cuánticos como una sucesión ordenada de propiedades cuánticas (Griffiths 1984; Griffiths y Omnés 1999; Laura y Vanni 2009).

Capítulo 2. Los problemas de la medición cuántica

Con las bases conceptuales provistas en el capítulo anterior, estamos en condiciones de presentar los dos problemas centrales de la medición cuántica, los cuales serán abordados por separado en capítulos posteriores para presentar, como parte central de esta tesis, una respuesta a cada uno de ellos. Previamente a ello, en lo que sigue se expondrán las bases de lo que se entiende como teoría cuántica de la medición.

2.1 La teoría cuántica de la medición

En toda teoría, la descripción del proceso de medición es de importancia central, no sólo para la corroboración o no de dicha teoría, sino también para ajustar una correcta interpretación de la misma. Esto es así porque cualquier interpretación debería proveer relaciones entre las expresiones teóricas y los resultados experimentales, y es el proceso de medición el que logra establecer este vínculo.

En la mecánica cuántica, la descripción del proceso de medición ha sido históricamente problemática. Actualmente, hay consenso respecto de considerar la física cuántica universalmente válida (Peres y Zurek 1982); sin embargo, aún persisten fuertes desacuerdos en su interpretación. Si una teoría es universalmente válida, debe contener los medios necesarios para su propia corroboración, es decir, debe ser capaz de describir los procesos de medición que testean sus resultados. Las condiciones mediante las cuales la mecánica cuántica puede o no cumplir este requisito de validez universal fueron inicialmente estudiadas por von Neumann (1955), quien desarrolló lo que hoy conocemos como teoría cuántica de la medición.

El concepto básico de la teoría de la medición en mecánica cuántica consiste en tratar como cuánticos tanto al objeto que se quiere medir, como al aparato que lo mide. Bajo esta condición, el proceso de medición es concebido como una interacción entre dos sistemas que se encuentran en pie de igualdad: el sistema objeto y el sistema aparato que mide al primero. La utilidad de la medición reside en el hecho de que, después de la interacción, estos dos sistemas quedan correlacionados, de modo tal que las propiedades de uno de ellos, por

ejemplo los valores q_i de un observable Q que se desea medir en el sistema objeto, se corresponden con propiedades del otro sistema, por ejemplo los valores a_i de un observable A que caracteriza la variable indicadora (pointer) del aparato. La variable indicadora es la variable sobre la cual se leen los resultados de la medición. Podemos decir que la variable indicadora es una suerte de representante en el mundo macroscópico de la variable medida en un sistema del mundo microscópico. Rara vez la variable medida es la variable registrada en la medición. Es justamente la medición el proceso que se encarga de establecer una correspondencia entre sistema y aparato de modo que la variable medida quede correlacionada con la registrada.

Esta correspondencia se establece por medio de la llamada función del pointer, que vincula los q_i con los a_i . Más exactamente, diremos que la función del pointer f es tal que $q_i = f(a_i)$. Así, el conocimiento de un valor del aparato de medición nos revela información sobre un valor del sistema.

Ahora bien, al tratarse de una interacción, la medición puede afectar la evolución del sistema objeto, por lo que el concepto clásico de "observación pasiva", así como la atribución de propiedades objetivas a dicho sistema, se ven seriamente comprometidos. Desde un punto de vista cuántico, no podemos "observar" al sistema objeto aislado: lo que observamos es siempre el sistema más la acción producida por la propia observación. No obstante, dicha acción queda completamente determinada por el operador de evolución construido con el hamiltoniano de la interacción para el sistema compuesto: objeto + aparato. Este operador, en el cual reside la dinámica de la medición, deberá satisfacer ciertos requerimientos para dar sentido a lo que llamamos medición. Esto se analizará a continuación, al estudiar las diferentes etapas de la medición (Mittelstaedt 1998).

2.1.2 Primera etapa: preparación

Supongamos que el sistema objeto S es un sistema cuántico representado mediante un espacio de Hilbert H_S , y el sistema aparato M es otro sistema cuántico representado mediante un espacio de Hilbert H_M . En el sistema S se desea medir el observable $Q = \sum q_i |q_i\rangle\langle q_i|$. Antes de la medición, S es preparado en un estado general que puede ser

expresado en términos de la autobase de Q de modo que $|\psi\rangle = \sum_i c_i |q_i\rangle$. El aparato es preparado en un estado $|a_0\rangle$, que es un autoestado con valor de referencia a_0 del observable indicador $A = \sum a_i |a_i\rangle\langle a_i|$. En esa etapa los sistemas S y M son completamente independientes; no obstante, son considerados como un sistema compuesto en el estado $|\Psi_0\rangle = |\psi\rangle |a_0\rangle \in H_S \otimes H_M$.

2.1.3 Segunda etapa: interacción

Los sistemas S y M son puestos en interacción durante un intervalo Δt . La dinámica de la evolución será gobernada por el hamiltoniano de interacción del sistema compuesto $S + M$ dado por $H_{SM} \in H_S \otimes H_M$, a través del operador de evolución $U \equiv U(\Delta t) = e^{-i/\hbar H_{SM} \Delta t}$.

El tipo de interacción es determinada, o elegida, de acuerdo con el observable que se quiere medir en el sistema S , y con qué observable en el aparato M se lo quiera vincular. Bajo estas consideraciones es esperable asumir que, elegido el observable que se desea medir en el sistema y el aparato para hacerlo (es decir, el arreglo experimental), a cada proceso de medición le corresponde un operador unitario U que lo caracteriza.

Los requisitos que debe cumplir este operador U para que la interacción tenga sentido como medición constituyen la calibración del proceso y del aparato. En la calibración se requerirá que, cuando el sistema es preparado en un autoestado $|q_i\rangle$ del observable Q , la medición arroje con certeza un único valor a_i para el observable del pointer, el cual se corresponderá con el valor $q_i = f(a_i)$ para la variable del sistema.

En el marco de las llamadas mediciones ideales⁵, si el sistema es preparado en un autoestado $|q_i\rangle$ de la variable a medir, U deberá ser tal que conduzca a la siguiente evolución

$$|\Psi_0\rangle = |q_i\rangle |a_0\rangle \rightarrow |\Psi_1\rangle = |q_i\rangle |a_i\rangle \quad (2.1)$$

Un proceso de medición de este tipo se encuentra en concordancia con una visión realista de la mecánica cuántica, porque el hecho de preservar el estado del sistema nos concede la posibilidad de atribuirle propiedades objetivas, ya sea de estado o de valor. En este caso, la lectura a_i del pointer indica que la medición para la variable Q ha arrojado como resultado el

⁵Las mediciones ideales son a veces llamadas también mediciones estándar.

valor $q_i = f(a_i)$, el cual puede ser atribuido al sistema como una propiedad determinada. En ese sentido, la función del pointer vincula el valor leído del observable del aparato después de la medición, con el valor del observable que efectivamente "posee" el sistema.

Las mediciones ideales son también mediciones repetibles, porque tienen la propiedad de dejar inalterado el estado del sistema si éste se encuentra en un autoestado de la variable que se mide. Sin embargo, un concepto más general de la medición lo único que exige es que, después de la interacción, el aparato preserve algún rastro del estado del sistema antes de la medición, rastro necesario para identificar qué se ha medido. Así, en el caso de las mediciones que llamaremos no ideales, la evolución resulta

$$|\Psi_0\rangle = |q_i\rangle |a_0\rangle \rightarrow |\Psi_1\rangle = |\eta_i\rangle |a_i\rangle \quad (2.2)$$

con $|\eta_i\rangle \neq |q_i\rangle$. El hecho de que el estado inicial del sistema resulte alterado no invalida la medición, ya que el valor del observable indicador permanece correlacionado al valor inicial del observable del sistema. Aunque este tipo de interacción modifica el estado del sistema, aun así en esta etapa todavía podemos adoptar una perspectiva realista al afirmar que la función del pointer en este caso vincula el valor leído en el aparato, no con el valor que tiene el sistema, sino con el que "tenía" antes de efectuar la medición (Mittelstaedt 1998).

La atribución realista de propiedades al sistema tiene mucho que ver con la posibilidad de obtener certezas en los resultados, o al menos certeza en algo que caracterice a dichos resultados como sucede en la calibración, ya que calibrar es de alguna manera ajustar las certezas que definen la interpretación de lo que se mide. Sin ello, no habría una medición en general.

Para mediciones ideales se puede demostrar (Mittelstaedt 1998) que, si se desea medir una variable Q en el sistema objeto, y si P_A es la variable canónica conjugada⁶ a la variable indicadora A del aparato, entonces

$$H_{SM} = -\lambda \hbar / \Delta t (Q \otimes P_A)$$

⁶Dos variables P_A y A se dicen canónicas conjugadas si $[P_A, A] = i\hbar$.

es el hamiltoniano de interacción adecuado que asegura que $U = e^{-i/\hbar H_{SM} \Delta t}$ cumple los requisitos de la calibración como fueran indicados más arriba. Para mediciones más generales, la forma explícita del hamiltoniano, cuyo correspondiente operador de evolución cumpla lo requerido en la calibración, puede ser difícil de hallar. No obstante, se puede demostrar que tal hamiltoniano tiene siempre que ser función de la variable Q que se pretende medir.

Dejando ya de lado las condiciones del postulado de calibración, si el sistema es preparado en una combinación lineal de autoestados del observable Q , esto es, $|\psi\rangle = \sum c_i |q_i\rangle$, como consecuencia de la linealidad del operador de evolución U la evolución resulta, para mediciones ideales,

$$|\Psi_0\rangle = |\psi\rangle |a_0\rangle = \sum_i c_i |q_i\rangle |a_0\rangle \rightarrow |\Psi_1\rangle = U |\Psi_0\rangle = \sum_i c_i |q_i\rangle |a_i\rangle \quad (2.3)$$

y en el caso de mediciones no ideales

$$|\Psi_0\rangle = |\psi\rangle |a_0\rangle = \sum_i c_i |q_i\rangle |a_0\rangle \rightarrow |\Psi_1\rangle = U |\Psi_0\rangle = \sum_i c_i |\eta_i\rangle |a_i\rangle \quad (2.4)$$

Ya sea el caso de mediciones ideales o no, la correlación entre sistema y aparato queda expresada por el estado final de la medición, el cual tiene la forma de una descomposición de Schmidt (Apéndice A.5), cuya característica central es ser una suma de índice único. Ese índice es el vínculo entre el valor de la variable del sistema y el valor obtenido como resultado de la medición en la variable indicadora.

Ahora bien, si el estado final del aparato es una suma de autoestados de la variable A , sólo sabemos que el resultado de la medición será alguno de los valores a_i , pero no sabemos cuál exactamente. La mecánica cuántica no puede predecir cuál será el resultado de la medición. Sólo puede predecir cuál será la distribución de probabilidad sobre los resultados posibles. De este modo se pierde la certeza. No obstante, se asume (y, de hecho, así se comprueba) que la medición tendrá un valor definido relativo al pointer, es decir, que aunque no sepamos cuál, se registrará sólo uno de los posibles valores a_i de la variable indicadora, y tal valor será atribuido al aparato. Ahora bien, si la atribución de valor corresponde a la

atribución de estado, entonces el hecho experimental de que la lectura del pointer conduzca siempre a un valor a_i bien definido conduce a una contradicción al asumirse que tal lectura se realiza sobre el estado final $|\Psi_1\rangle$, puesto que este estado constituye una superposición de distintos autoestados de la variable A , cada uno correspondiente a un distinto valor a_i . Esta dificultad nos introduce en los llamados problemas de la medición en mecánica cuántica.

Dentro de los variados problemas conceptuales de la teoría, los problemas de la medición son sin duda los más discutidos y constituyen el marco sobre el cual se elabora la presente tesis. Actualmente podemos decir que estos problemas se centran alrededor de dos núcleos principales (Schlosshauer 2004)

- *El problema de la lectura definida:* La teoría no es capaz de explicar por qué, si el estado final de la medición es una superposición de estados de la variable indicadora con distintos valores posibles, dicha medición arroja siempre un resultado bien definido consistente en uno y sólo uno de esos valores.
- *El problema de la base privilegiada:* Bajo ciertas condiciones en el desarrollo del estado inicial del sistema antes de la medición, se demuestra que el estado final no es único, sino que puede expresarse por medio de distintas correlaciones poniendo en evidencia una ambigüedad de la teoría en la determinación de qué variable está siendo medida.

Si bien históricamente el primer problema, el de la lectura definida, fue el primero en evidenciarse y en recibir la mayor cantidad de propuestas para solucionarlo, el segundo problema, el de la base preferida, es tan grave o más que el primero. En lo que sigue entraremos en el detalle de la definición de cada uno de estos dos problemas

2.2 El problema de la lectura definida

Comencemos por considerar un proceso de medición ideal de una variable Q por medio de un aparato con variable indicadora A . Como puede verse de la ecuación (2.3), aunque el estado inicial antes de la medición es un estado factorizado entre el sistema y el aparato, el

estado final resulta ser una superposición de estados del sistema y del aparato que impide distinguir uno del otro.

Como se ha considerado una base ortonormal tanto en el sistema como en el aparato, estos estados son excluyentes en el sentido de que su producto interno da cero. El problema se presenta al considerar la variable indicadora del aparato, la cual se pretende sea aquella que registre los resultados de la medición en el mundo macroscópico de modo que el aparato sirva como tal. Si esto se cumple, entonces el estado final representa una superposición de estados excluyentes y asociados a resultados macroscópicamente distinguibles. La mecánica cuántica presenta aquí un problema, porque las superposiciones de estados excluyentes no tienen sentido en el mundo macroscópico. Pero más grave aún es que los datos de la experiencia indican que esto nunca se observa, de modo que lo que predice la teoría cuántica de la medición por alguna razón es incompleto. Algún otro mecanismo participa de la medición y no es contemplado por la expresión dada por la ecuación (2.3). La teoría cuántica de la medición es incapaz de explicar por qué siempre se registran valores bien definidos de la variable indicadora, aun cuando la teoría predice una superposición de estados para dicha variable.

Mucho se ha discutido en el intento de encontrar una respuesta adecuada a las dificultades que plantea este problema (Peres 1980; Ballentine 1998; Ballentine 1990; Schlosshauer 2004; Laura y Vanni 2008a), e incluso algunas interpretaciones han sido diseñadas específicamente para suministrar una solución conceptualmente admisible. Presentaremos en lo que sigue las principales propuestas abordadas para tratar este problema.

2.2.1 Copenhague y la hipótesis del colapso

No es sencillo caracterizar con precisión qué se entiende por interpretación de Copenhague, en la medida en que no constituye una doctrina única y definida sino, más bien, un conjunto de ideas emparentadas pero no siempre consistentes entre sí. Incluso podría decirse que hay tantas versiones de la interpretación de Copenhague como defensores de tal perspectiva, entre los cuales se cuentan Niels Bohr, Werner Heisenberg, Wolfgang Pauli y

Max Born. No obstante, es posible formular ciertos aspectos centrales de tal interpretación⁷. En primer lugar, se considera que los conceptos físicos básicos no pertenecen al ámbito microscópico sino al macroscópico. Desde esta perspectiva, la mecánica cuántica es un conjunto de reglas destinadas a establecer correlaciones entre eventos macroscópicos; por lo tanto, adquieren una importancia central las correlaciones entre los resultados de las mediciones sobre el sistema cuántico y la preparación del dispositivo experimental.

En cuanto al modo de evolución de los sistemas cuánticos, la interpretación de Copenhague postula un dualismo: mientras se encuentra aislado, el sistema evoluciona según ley determinista dada por la ecuación de Schrödinger; pero en el instante de la medición, y sólo en la medición, el sistema está sujeto a una ley indeterminista, el postulado de proyección (o de colapso) según el cual, ante la medición de un observable Q , el vector de estado se proyecta –la función de onda “colapsa”– abruptamente e indeterminísticamente sobre uno de los autovectores del observable Q , resultando así un valor definido de la medición correspondiente al autovalor asociado al autovector sobre el cual colapsa el sistema (Ghirardi 2011).

En el marco de la llamada teoría cuántica de la medición, en 1932 von Neumann introdujo formalmente el postulado del colapso como una hipótesis de partida de la teoría (von Neumann 1955). De este modo, el colapso se supone como un hecho irreductible que no puede ser explicado en términos de una evolución continua gobernada por la ecuación de Schrödinger (Jammer 1974, pp. 474-479).

En la formulación original de von Neumann, a veces llamada de colapso fuerte, la evolución por medio del colapso es tal que si el sistema es preparado en un estado puro $|\psi\rangle$, superposición de autoestados $|q_i\rangle$ de la variable Q que se quiere medir (es decir, $|\psi\rangle = \sum c_i |q_i\rangle$), entonces en la medición de Q el estado del sistema se proyecta (colapsa) a uno de los autoestados $|q_i\rangle$ con una probabilidad dada por $|c_i|^2$, arrojando el valor q_i como resultado de la medición.

⁷ Una breve pero muy buena síntesis de las tesis centrales de la interpretación de Copenhague puede encontrarse en Loewer (1998, pp. 317-318).

En una postulación posterior, conocida como colapso débil, se asume que si el sistema es preparado en $|\psi\rangle = \sum c_i |q_i\rangle$, luego de la medición del observable Q el sistema queda descrito por un estado mezcla $\rho_\psi = \sum |c_i|^2 |q_i\rangle\langle q_i|$. En este caso no hay una proyección, pero el hecho de que el estado luego de la medición sea un estado mezcla habilitaría el uso de la interpretación por ignorancia, al suponer que el sistema efectivamente está en alguno de los $|q_i\rangle$, tal como si hubiera ocurrido una proyección, aunque no sabemos en cuál exactamente. Ya hemos analizado las dificultades que se presentan en la interpretación por ignorancia, por lo que esta manera de suponer el estado después de una medición está seriamente comprometida.

No obstante, lo más grave es que, ya sea se trate del colapso fuerte –que si bien supone una evolución de un estado puro a otro puro, el segundo es una proyección del primero– o del colapso débil –que supone una evolución de un estado puro a uno mezcla–, en ningún caso este tipo de evoluciones puede ser producido por un operador de evolución unitario como se requiere en la ecuación Schrödinger.

Esta situación es conceptualmente inadmisibles pues, en la medición, el sistema a medir, que está compuesto de átomos y partículas gobernados por la mecánica cuántica, interactúa con otro sistema, el aparato, que también está compuesto de átomos y partículas gobernadas por la mecánica cuántica. ¿Por qué, entonces, en la medición se produce una interacción como el colapso, que no puede ser descrita por la ecuación de Schrödinger?

En esta tesis mostraremos que, a través de la condicionalización de resultados conocidos en una primera medición, el postulado del colapso puede reconciliarse con una evolución unitaria para mediciones posteriores. Si bien no brindaremos una respuesta al hecho de la obtención de valores definidos en las mediciones, presentaremos una solución a la contradicción que supone la existencia de valores definidos en presencia de evoluciones unitarias.

2.2.2 Las interpretaciones modales

En algún sentido, las interpretaciones modales brindan la respuesta más sencilla al problema de la lectura definida, pues lo hacen renunciando al vínculo normalmente aceptado

entre la atribución de estado y la de valor. El problema se disuelve al aceptar que un valor bien definido como resultado de la medición no tiene por qué implicar una pérdida de la superposición del estado sobre el cual dicho valor es obtenido.

Actualmente puede afirmarse que existe una familia de interpretaciones modales, aunque casi todas tienen su origen en los trabajos de van Fraassen de la década del '70 (van Fraassen 1972, 1973, 1974). La idea central en todas ellas consiste en asumir que los estados cuánticos, a diferencia de los estados clásicos, codifican información sobre las posibilidades más que sobre lo actual. Para ello se recurre a diferenciar las atribuciones de estado de las atribuciones de valor. La atribución de estado determina las propiedades que el sistema puede poseer, con sus correspondientes probabilidades como son dadas en el formalismo usual de la teoría cuántica. La atribución de valor, en cambio, determina la propiedad que realmente el sistema posee. Se presentan así dos ámbitos o dominios de aplicación, cada uno evolucionando independientemente: el ámbito de las posibilidades, representado por la evolución de los estados cuánticos de acuerdo con la ecuación de Schrödinger, y el ámbito de lo actual, que viene dado por el conjunto de propiedades que se actualizan en el sistema con las mediciones realizadas. De acuerdo con las interpretaciones modales, entonces, un observable puede tener un valor específico, sin que esto signifique que el estado del sistema sea el correspondiente autoestado del observable. Esto brinda una solución al problema de la lectura definida sin apelar al postulado del colapso (Dieks 1994; 2007b). Los observables pueden tener valor definido aun preservando las evoluciones unitarias en los estados.

No obstante, al relajarse el vínculo directo entre estado y valor, la relación entre lo posible y lo actual es menos trivial de lo que se supone con la interpretación usual. Para determinar cuáles propiedades pueden actualizarse de acuerdo con las posibilidades establecidas por el estado, la interpretación debe introducir reglas adicionales que se añaden al formalismo. Estas reglas son a veces llamadas "reglas de actualización" (Lombardi y Castagnino 2008). Las diferentes variantes de las interpretaciones modales difieren esencialmente en el modo en que eligen la regla de actualización, pero todas ellas comparten la interpretación del estado como la base de las características modales del sistema, y no de las actuales (Dieks 2007a).

2.2.3 La interpretación subjetiva de Wigner

Esta interpretación surge como un intento de identificar qué tiene de particular el proceso de medición, de modo que su evolución sea gobernada por el colapso y no por la ecuación de Schrödinger como se esperaría de una interacción cuántica entre un sistema objeto y otro sistema llamado aparato.

Si bien no el primero, Eugene Wigner (1961) fue quien más difundió la tesis según la cual la particularidad del acto de medición se encuentra en la intervención de un observador humano⁸. Según esta interpretación, en toda medición no existe en realidad una sola interacción, sino una cadena de interacciones que culmina en última instancia en un registro macroscópico capaz de ser reconocido por la conciencia humana. Es la conciencia humana el elemento particular que está presente en la medición, y no en otro tipo de interacciones.

Wigner llega a proponer la situación conocida como “la paradoja del amigo de Wigner”, que se puede presentar de la siguiente manera. Como hemos visto, en la teoría cuántica de la medición, si el sistema S interactúa con un aparato cuántico M , el estado del sistema compuesto $S + M$ se describe como una superposición; pero si M se reemplaza por un amigo O al cual se le pregunta por el estado de $S + O$, la idea de una superposición se torna insostenible, pues “implica que el amigo se encontraba en un estado de animación suspendida antes de responder mi pregunta. De aquí se sigue que, en mecánica cuántica, un ser con conciencia debe cumplir un papel diferente al de un dispositivo de medición inanimado” (Wigner 1961, p. 294). La conciencia influye sobre los sistemas físicos de un modo análogo al modo en que recibe su influencia.

Como decíamos, Wigner no fue el único en adherir a este tipo de interpretación para la medición. Ya von Neumann, años antes que Wigner, había propuesto que la medición consiste en una cadena de interacciones que concluye en la observación humana. Según von Neumann, el colapso es el resultado de la intervención de la mente sobre los eventos físicos⁹.

⁸ Esta interpretación puede legítimamente denominarse “subjetiva”, en la medida en que hace intervenir, de un modo esencial, la conciencia del sujeto cognoscente individual.

⁹ Max Jammer (1974, p. 482) compara la solución propuesta por von Neumann con la doctrina de Anaxágoras, según la cual la materia forma un todo indiferenciado y sólo la acción de la mente (*Nous*) introduce una separación de las cosas entre sí.

John Wheeler aventura una hipótesis aún más especulativa sobre la base de la posibilidad de experiencias de acción retardada (ver Davies y Brown 1986): por ejemplo, en la experiencia de las dos rendijas, es posible demorar la elección de qué observable será medido hasta después de que la partícula ha atravesado la pantalla. A partir de ello, Wheeler especula que la conciencia es responsable de la actualización retroactiva de lo real, donde esta acción retroactiva puede alcanzar incluso el pasado más remoto.

La interpretación subjetiva reinstala un dualismo ontológico en el corazón de la propia física: el mundo físico, material, considerado aisladamente evoluciona según la ecuación de Schrödinger y responde a un estricto determinismo; en el ámbito de lo mental, en cambio, reina la indeterminación. A su vez, lo mental interactúa con lo físico a través del proceso de medición, introduciendo los eventos indeterministas a los que alude la mecánica cuántica. Tal vez este carácter dualista sea el principal motivo por el cual la interpretación subjetiva, salvo contadas excepciones, no ha sido seriamente considerada por la comunidad científica¹⁰.

2.2.4 Everett y la multiplicidad de mundos

Las soluciones al problema de la lectura definida basadas en la hipótesis del colapso son dualistas, en el sentido de que postulan dos modos de evolución completamente diferentes en los sistemas cuánticos y suponen la descomposición del mundo en sistema y dispositivo de medición. Sin embargo, el formalismo de la mecánica cuántica no brinda indicación alguna acerca del criterio para efectuar tal descomposición.

Es en este contexto que Hugh Everett III (1957) formula la primera interpretación monista de la mecánica cuántica, posteriormente ampliada por Bryce De Witt (1970), llamada interpretación de “muchos mundos” (many worlds). Esta interpretación se basa estrictamente en el formalismo de la teoría sin el agregado de ningún postulado adicional: existe un único modo de evolución cuántica, regido por la ecuación de Schrödinger; por lo tanto, la superposición de estados nunca colapsa reduciéndose a una de sus componentes. A fin de reconciliar este supuesto con la experiencia –que adjudica un único valor definido al

¹⁰Una de estas excepciones es Roger Penrose (1989), quien especula que en el futuro contaremos con una teoría más general, que incluirá a la mecánica cuántica y que, a la vez, será capaz de tratar la mente en términos precisos.

observable luego de la medición– los autores asumen que el universo completo se ramifica en diversas copias de sí mismo, una por cada componente de la superposición, que jamás vuelven a interactuar. Los propios observadores humanos se multiplican en copias, lo cual explica que no puedan registrar o percibir que la bifurcación ha ocurrido. De este modo se resuelve el problema de la medición: medimos una única componente de la superposición pues sólo tenemos acceso a la que permanece en nuestra “rama” del universo. El vector de estado representa efectivamente la realidad física pero describiendo no uno, sino todos los “mundos”; es nuestra observación la que se encuentra confinada a una parte limitada de tal realidad.

Dentro de esta interpretación el concepto de probabilidad pierde significado (Healey 1984, p. 502) puesto que, si luego de una interacción el universo entero se despliega en copias de sí mismo, todas igualmente actuales, ¿qué sentido tiene adjudicar cierta probabilidad a cada una de ellas?¹¹. La interpretación de muchos mundos constituye una lectura totalmente actualista de la mecánica cuántica: las probabilidades desaparecen; todo aquello que en otras versiones es considerado posible aquí es considerado actual.

Si bien esta interpretación resuelve muchas de sus dificultades conceptuales, ha sido muy criticada por introducir elementos no testeables empíricamente y resultar tremendamente antieconómica así como extravagante desde un punto de vista metafísico: cada interacción cuántica, de las innumerables que suceden sin cesar, produciría una bifurcación global del espacio-tiempo (Sonego 1993, p. 12). John Earman subraya las consecuencias de esta interpretación en lo que se refiere a la noción de espacio-tiempo, señalando que, si la idea de bifurcación del universo debe ser tomada literalmente, entonces implica una ramificación de su estructura espacio-temporal; pero es muy difícil concebir “un mecanismo causal por el cual una medición del spin de un electrón causa una bifurcación global del espacio-tiempo” (Earman 1986, p.224).

¹¹Healey (1984, p. 503) señala también la dificultad adicional que introduce el carácter t -invariante de la Ecuación de Schrödinger: si la descripción es t -invariante y ninguna restricción introduce una irreversibilidad de facto, entonces el proceso de permanente bifurcación del universo es reversible; la descripción temporalmente invertida consistiría en múltiples universos en constante fusión.

Otras críticas apuntan a la vaguedad de la noción de bifurcación: si la bifurcación del universo es un evento físico, debe tener una descripción precisa dentro de la teoría; pero el formalismo cuántico no brinda indicación acerca del modo y del instante en que la bifurcación ocurre.

Finalmente, desde un punto de vista más técnico, Richard Healey (1984) señala otro problema al que se enfrenta la interpretación de muchos mundos y que tiene que ver con el otro gran problema de la medición, el problema de la base preferida, al que nos referiremos detalladamente en la próxima sección. El origen del problema reside en que, bajo ciertas condiciones en el desarrollo de una superposición de estados, es posible efectuar un cambio de base de modo que la misma superposición pueda ser representada como compuesta de otros estados. En este caso se presenta una ambigüedad adicional: además de existir un mecanismo indefinido capaz de producir una bifurcación del universo, una por cada componente de la superposición, queda también indefinido el conjunto de las posibles bifurcaciones. Es decir, no sólo existe una indeterminación en la bifurcación tomada, sino también en la colección de las posibles bifurcaciones a tomar.

El problema de la base preferida tiene consecuencias que exceden a la interpretación de muchos mundos, y brindar una solución a tal problema constituye uno de los núcleos centrales de la presente tesis.

2.3 El problema de la base preferida

Si bien no tan conocido como el problema de la lectura definida, el problema de la base preferida es conceptualmente tan grave como aquél. De hecho, incluso lo empeora, porque si en virtud del problema de la lectura definida no podemos explicar el colapso a un autoestado del observable que se mide, debido al problema de la base preferida ni siquiera podremos explicar en qué base se produce dicho colapso.

El problema de la lectura definida y el problema de la base preferida se encuentran estrechamente vinculados, y ambos constituyen los dos grandes problemas que se presentan en la teoría cuántica de la medición. La presente tesis intenta dar una respuesta de carácter netamente conceptual a estos problemas, tarea que se abordará en los capítulos siguientes.

Habiendo caracterizado el problema de la lectura definida en la sección anterior, prosigamos ahora con una presentación detallada del problema de la base preferida.

Consideremos un proceso de medición ideal de la variable Q por medio de un aparato con variable indicadora A . Si inicialmente el sistema es preparado en el estado $|\psi\rangle = \sum c_i |q_i\rangle$, el proceso de medición, que queda determinado por el operador de evolución U , puede representarse como indica la ecuación (2.3) de modo tal que

$$|\Psi_0\rangle = |\psi\rangle |a_0\rangle \rightarrow |\Psi_1\rangle = U |\Psi_0\rangle = \sum_i c_i |q_i\rangle |a_i\rangle \quad (2.5)$$

El estado final después de la medición, dado por $|\Psi_1\rangle = \sum c_i |q_i\rangle |a_i\rangle$, es una superposición de estados del sistema y el aparato que pone de manifiesto la correlación que la medición establece entre la variable $Q = \sum q_i |q_i\rangle \langle q_i|$ del sistema y la variable $A = \sum a_i |a_i\rangle \langle a_i|$ del aparato, de modo tal que, por medio de la función del pointer, el valor a_i leído en el aparato corresponde al valor q_i atribuible al sistema.

Ahora bien, bajo la condición de que los coeficientes c_i que desarrollan el estado inicial del sistema sean iguales en módulo (Zurek 1981), es posible realizar un cambio de base $\{|q_i\rangle\} \rightarrow \{|q_i'\rangle\}$ y $\{|a_i\rangle\} \rightarrow \{|a_i'\rangle\}$, de manera que el estado final después de la medición se pueda escribir como la combinación $|\Psi_1\rangle = \sum c_i |q_i'\rangle |a_i'\rangle$. De este modo, el mismo proceso de medición dado por (2.5) puede ser representado también como

$$|\Psi_0\rangle = |\psi\rangle |a_0\rangle \rightarrow |\Psi_1\rangle = U |\Psi_0\rangle = \sum_i c_i |q_i'\rangle |a_i'\rangle \quad (2.6)$$

Damos cuenta aquí que el proceso de medición no cambió, sólo hemos cambiado la forma matemática del desarrollo del estado final después de la medición. El problema de la base preferida surge al interpretar el estado final después de la medición. Si en (2.5) lo interpretábamos como representado la medición del observable $Q = \sum q_i |q_i\rangle \langle q_i|$ por medio de un aparato con variable indicadora $A = \sum a_i |a_i\rangle \langle a_i|$, entonces en (2.6) debería interpretarse como representando la medición del observable $Q' = \sum q_i' |q_i'\rangle \langle q_i'|$ por medio de un aparato con variable indicadora $A' = \sum a_i' |a_i'\rangle \langle a_i'|$. Sin embargo, no hemos cambiado el proceso de medición, sólo hemos cambiado la forma del desarrollo del estado final.

Esto pone de manifiesto una ambigüedad teórica grave. Bajo la condición de coeficientes de igual módulo, la teoría cuántica de la medición no nos permite en principio determinar desde el formalismo qué observable fue medido (Zurek 1981, 1982; Schlosshauer 2004). En otras palabras, aun con la descripción de un único proceso de interacción, es decir, empleando un único operador de evolución U , la teoría permitiría la ambigüedad de especificar la correlación establecida por otra medición, producto de otra interacción. Sin algún elemento adicional no podemos determinar una base privilegiada, esto es, cuál de las correlaciones es la que realmente está presente entre el sistema y el aparato. En términos del colapso, no podríamos saber en cuál de las superposiciones se producirá el colapso (Zurek 1981). No sabríamos qué observable se midió y tampoco cuál es el pointer del aparato de medición, esto es, qué “juego” de lecturas se efectivizará en la medición. Este argumento se generaliza al afirmar: “En el mundo real, incluso cuando no conocemos el resultado de una lectura, sabemos cuáles son las alternativas y podemos actuar confiadamente como si sólo una de ellas hubiera ocurrido. Pero, ¿cómo un observador, que aún no ha observado el detector, puede expresar su ignorancia acerca de la lectura sin contar con la certeza acerca del “menú” de posibilidades?” (Zurek 1991, p. 38).

De hecho, es posible demostrar que existen infinitos cambios de bases del tipo $\{|q_i\rangle\} \rightarrow \{|q_i'\rangle\}$ y $\{|a_i\rangle\} \rightarrow \{|a_i'\rangle\}$, de modo que la ambigüedad se extiende sobre un conjunto infinito de posibles observables. El problema se torna más grave aun al tomar en cuenta que, dentro de esos infinitos observables, es posible encontrar observables Q y Q' que no conmutan, es decir, que son incompatibles: en completa contradicción con la teoría, la ambigüedad del cambio de base permitiría que fueran medidos por un mismo proceso de medición.

Tanto el problema de la lectura definida como el de la base preferida ponen de manifiesto una aparente incompletitud por parte de la teoría cuántica. En el primer caso se manifiesta la incapacidad de la teoría para explicar el resultado que predice el colapso. En el segundo, su incapacidad para determinar, además, el observable en cuya autobase se produce el colapso. Aunque menos difundido, el problema de la base preferida resulta ser en realidad más grave, puesto que no sólo manifestaría una incompletitud de la mecánica cuántica, sino que además

presentaría una inconsistencia en la propia teoría al permitir la posibilidad de medición de observables incompatibles.

A diferencia del problema de la lectura definida, que ha sido abordado con un sinnúmero de propuestas para su solución, algunas de las cuales hemos repasado brevemente en secciones previas, el problema de la base preferida sólo ha sido encarado con un único enfoque suficientemente sólido como para ser efectivo. Se trata del formalismo de decoherencia. De hecho, algunos autores proclaman que este formalismo sería capaz de solucionar ambos problemas de la medición a la vez. Por esta razón, y por su importancia en el tema de esta tesis, le dedicaremos por completo el siguiente capítulo.

Capítulo 3. El formalismo de decoherencia

El formalismo de decoherencia es actualmente uno de los enfoques más prominentes para dar cuenta del vínculo entre las descripciones cuántica y clásica que se presentan en el proceso de medición (Ludwig 1950; Zeh 1970; Zurek 1981, 1982, 1991; Bacciagaluppi y Hemmo 1994; Joos 2000; Zurek 2003; Paz y Zurek 2002; Schlosshauer 2004). El formalismo pretende explicar de qué manera las superposiciones cuánticas a nivel microscópico son eliminadas en el proceso de medición para dar lugar a la emergencia del mundo clásico al nivel macroscópico. Esto no sólo brindaría una respuesta satisfactoria al primer problema de la medición, el problema de la lectura definida, sino que, como muchos autores consideran (Zurek 1981; Zurek 1982; Schlosshauer 2004), resolvería también el problema de la base preferida.

3.1. Idea básica y antecedentes

Como su nombre indica, la decoherencia es el proceso mediante el cual se pierde coherencia de un estado puro que es superposición de otros estados. Un estado se dice coherente cuando las relaciones de fase entre los miembros que componen la superposición se mantienen constantes. Por ejemplo, consideremos el siguiente estado puro

$$|\psi\rangle = \alpha_1 |q_1\rangle + \alpha_2 |q_2\rangle$$

donde α_1 y α_2 son números complejos cualesquiera que, en sus formas exponenciales, pueden expresarse como $\alpha_1 = Ae^{i\delta_1}$ y $\alpha_2 = Be^{i\delta_2}$. La coherencia se traduce en una relación constante entre las fases δ_1 y δ_2 . Ahora bien, visto como operador, el estado $|\psi\rangle$ adopta la forma

$$\begin{aligned}\rho_\psi &= |\psi\rangle\langle\psi| \\ &= (\alpha_1 |q_1\rangle + \alpha_2 |q_2\rangle)(\alpha_1^* \langle q_1| + \alpha_2^* \langle q_2|) \\ &= |A|^2 |q_1\rangle\langle q_1| + Ae^{i\delta_1} Be^{-i\delta_2} |q_1\rangle\langle q_2| + Ae^{-i\delta_1} Be^{i\delta_2} |q_2\rangle\langle q_1| + |B|^2 |q_2\rangle\langle q_2|\end{aligned}$$

que en su representación matricial se expresa como

$$\rho_{\psi} = \begin{pmatrix} |A|^2 & AB e^{i(\delta_1 - \delta_2)} \\ AB e^{i(\delta_2 - \delta_1)} & |B|^2 \end{pmatrix}$$

Tal como se observa en esta última expresión, en el operador de estado la coherencia se manifiesta en los términos fuera de la diagonal en la correspondiente matriz, y a través de la diferencia $\pm(\delta_1 - \delta_2)$. Si, por alguna razón, la relación entre δ_1 y δ_2 fluctuara al azar con el paso del tiempo, la diferencia $\pm(\delta_1 - \delta_2)$ podría tomar cualquier valor. No obstante, bajo la consideración de observar el resultado promedio de esas fluctuaciones, se obtendrían términos nulos fuera de la diagonal y así la coherencia del estado se perdería. Esto se debe a los factores $e^{\pm i(\delta_1 - \delta_2)}$ que, por ser complejos, se distribuyen en el círculo de radio unidad sumando en promedio cero. Bajo estos supuestos, el operador de estado se torna diagonal

$$\bar{\rho}_{\psi} = \begin{pmatrix} |A|^2 & 0 \\ 0 & |B|^2 \end{pmatrix}$$

La importancia de este resultado reside en el hecho de que este último operador consiste en una mezcla de estados puros que se puede escribir como

$$\bar{\rho}_{\psi} = |A|^2 |q_1\rangle\langle q_1| + |B|^2 |q_2\rangle\langle q_2|$$

la cual, aunque con todas las objeciones que ya hemos visto, suele ser interpretada por ignorancia, interpretación que se ajusta a la presunción clásica de asumir que el sistema está en algunos de los estados $|q_1\rangle$ o $|q_2\rangle$ que componen la mezcla. De este modo la decoherencia es el proceso que elimina las superposiciones cuánticas, transformándolas en mezclas que supuestamente pueden interpretarse en términos clásicos.

Existen varios tipos de formalismos de decoherencia, pero la mayoría tiene su origen en consideraciones estadísticas sobre el gran número de grados de libertad de los instrumentos de medición al ser estos considerados macroscópicos, ya que en principio todo instrumento

debe ser capaz de dejar registro en el mundo de nuestra percepción para ser considerado como tal.

El formalismo tiene sus primeros antecedentes en trabajos desarrollados en la década del '50 como intentos de resolver los problemas de la medición. Se buscaba dar cuenta de la evolución no unitaria que describe el colapso como un proceso termodinámico irreversible, considerando la medición como un proceso que ocurre dentro de un sistema cerrado macroscópico (Ludwig 1950; van Kampen 1954; Daneri, Loinger y Prosperi 1962, 1966). Desde esta perspectiva, la mecánica cuántica era interpretada como una teoría que estudia procesos termodinámicamente irreversibles que se producen en objetos macroscópicos, causados por objetos microscópicos. El carácter indeterminista de la teoría sería producto de la imposibilidad de preparar sistemas microscópicos que produzcan los mismos efectos en los instrumentos de medición macroscópicos.

Más tarde, pero sobre la base de estas ideas, se intentó tratar la irreversibilidad como producto de considerar la medición como un proceso que ocurre en un sistema abierto en interacción con un entorno, y que es capaz de desembocar en la pérdida de coherencia en las superposiciones cuánticas (Zeh 1970). No obstante, estos primeros enfoques tuvieron muchas falencias formales, y las principales críticas que se les dirigieron provienen del hecho de que existen ejemplos donde la pérdida de coherencia, y por lo tanto los efectos cuánticos de la superposiciones, se presenta en sistemas que no interactúan con ningún entorno, ni con instrumentos de medición macroscópicos (Jauch, Wigner y Yanase 1967). De hecho, existen ejemplos donde, en virtud de una configuración específica para el estado de un fotón, el colapso se produce antes de la interacción con el instrumento de medición que lo detectará (Jammer 1974). Otras críticas apuntan al hecho de tratar la medición mediante consideraciones termodinámicas, lo cual no hace más que trasladar al ámbito de la mecánica cuántica el problema adicional de la mecánica estadística vinculado a la fundamentación de la irreversibilidad.

Si bien muchos de estos ejemplos siguen siendo cuestión de debate, no fue hasta la década de los '80 que Zurek y sus colaboradores desarrollaron la idea central sobre la base de la cual la decoherencia es entendida actualmente, al menos en su enfoque ortodoxo (Zurek

1981, 1982). Este enfoque suele denominarse Decoherencia Inducida por el Entorno (Environment Induced Decoherence, EID), debido al papel central que cumple el entorno en el proceso de medición cuántica. En la próxima sección presentaremos las bases teóricas de este enfoque.

3.2 La teoría de la decoherencia

Consideremos un sistema S de dos estados, al que se le desea medir un observable $Q = \sum q_i |q_i\rangle\langle q_i|$, con $i = 1, 2$. Para ello se dispone de un instrumento M , con variable indicadora $A = \sum a_i |a_i\rangle\langle a_i|$. Si el estado inicial del sistema S es una superposición $|\psi_S\rangle = c_1 |q_1\rangle + c_2 |q_2\rangle$, como hemos visto en el Capítulo 2, luego de la interacción entre S y M el sistema compuesto se encuentra en el estado correlacionado dado por

$$|\psi_{SM}\rangle = c_1 |q_1\rangle \otimes |a_1\rangle + c_2 |q_2\rangle \otimes |a_2\rangle$$

Según se entiende en la teoría cuántica de la medición (von Neumann 1955), es en este punto donde la medición termina, y se presentan los ya tratados problemas de la lectura definida y de la base preferida. Sin embargo, para el formalismo de decoherencia ésta es solo la primera etapa de la medición, donde se establece la correlación entre el sistema S y el aparato M . Esta etapa es llamada 'pre-medición' (Zurek 1981). En la pre-medición no concluye la medición ya que el sistema compuesto $S + M$ nunca se encuentra aislado, sino en interacción con un entorno E de un enorme número de grados de libertad¹². Es recién después de la posterior interacción con el entorno que la medición cobra sentido como tal y se produce el registro irreversible de su resultado. Para el formalismo de decoherencia, es sólo en ese punto donde tiene sentido buscar la solución a los problemas de la medición.

Si suponemos que el entorno E está conformado por un número N de sistemas de dos estados (por ejemplo, partículas de spin $\frac{1}{2}$) que no interactúan entre sí, el estado del sistema compuesto $S + M + E$ antes de la interacción de $S + M$ con el entorno E resulta ser

¹²Para todos los fines prácticos, el entorno consiste en los demás grados de libertad de propio aparato que no intervinieron en la pre-medición.

$$|\Psi_{SME}\rangle = (c_1|q_1\rangle \otimes |a_1\rangle + c_2|q_2\rangle \otimes |a_2\rangle) \otimes \prod_{k=1}^N (\gamma_k |\uparrow\rangle_k + \delta_k |\downarrow\rangle_k)$$

donde $|\uparrow\rangle_k$ y $|\downarrow\rangle_k$ son los autoestados de la componente en la dirección k del spin correspondiente a la partícula k . Por simplicidad, se considera que los hamiltonianos del aparato y del entorno son nulos, de modo que la evolución del sistema compuesto $S + M + E$ estará regida por un hamiltoniano de interacción H_{int} del tipo

$$H_{\text{int}} = (|q_1\rangle\langle q_1| \otimes |a_1\rangle\langle a_1| + |q_2\rangle\langle q_2| \otimes |a_2\rangle\langle a_2|) \otimes \sum_k g_k (|\uparrow\rangle\langle\uparrow| + |\downarrow\rangle\langle\downarrow|)_k$$

Se puede probar que, bajo la influencia de este H_{int} , el estado $|\Psi_{SME}\rangle$ evolucionará como

$$|\Psi_{SME}(t)\rangle = c_1|q_1\rangle \otimes |a_1\rangle \otimes |\varepsilon_1(t)\rangle + c_2|q_2\rangle \otimes |a_2\rangle \otimes |\varepsilon_2(t)\rangle \quad (3.1)$$

donde

$$\begin{aligned} |\varepsilon_1(t)\rangle &= \prod_k (\delta_k e^{ig_k t} |\uparrow\rangle_k + \delta_k e^{-ig_k t} |\downarrow\rangle_k) \\ |\varepsilon_2(t)\rangle &= \prod_k (\delta_k e^{-ig_k t} |\uparrow\rangle_k + \delta_k e^{ig_k t} |\downarrow\rangle_k) \end{aligned}$$

De este modo, el estado del sistema $S + M$ queda entrelazado (*entangled*) con el estado del entorno E . El operador $\rho_{SME}(t)$ correspondiente a este estado entrelazado será

$$\begin{aligned} \rho_{SME}(t) = & |\Psi_{SME}(t)\rangle\langle\Psi_{SME}(t)| = |c_1|^2 |q_1\rangle\langle q_1| \otimes |a_1\rangle\langle a_1| \otimes |\varepsilon_1(t)\rangle\langle\varepsilon_1(t)| \\ & + c_1 c_2^* |q_1\rangle\langle q_2| \otimes |a_1\rangle\langle a_2| \otimes |\varepsilon_1(t)\rangle\langle\varepsilon_2(t)| \\ & + c_1^* c_2 |q_2\rangle\langle q_1| \otimes |a_2\rangle\langle a_1| \otimes |\varepsilon_2(t)\rangle\langle\varepsilon_1(t)| \\ & + |c_2|^2 |q_2\rangle\langle q_2| \otimes |a_2\rangle\langle a_2| \otimes |\varepsilon_2(t)\rangle\langle\varepsilon_2(t)| \end{aligned}$$

donde los términos “cruzados” (términos fuera de la diagonal en la representación matricial) están dados por el segundo y tercer renglón y, como hemos visto, son los que manifiestan las correlaciones cuánticas que impiden una interpretación en términos clásicos.

En cada instante t la descripción del sistema $S + M$ viene dada por el operador reducido que se obtiene trazando sobre los grados de libertad del entorno, esto es

$$\rho_{SM}(t) = Tr_E [\rho_{SME}(t)] = |c_1|^2 |q_1\rangle\langle q_1| \otimes |a_1\rangle\langle a_1| + c_1 c_2^* r(t) |q_1\rangle\langle q_2| \otimes |a_1\rangle\langle a_2| \\ + c_1^* c_2 r^*(t) |q_2\rangle\langle q_1| \otimes |a_2\rangle\langle a_1| + |c_2|^2 |q_2\rangle\langle q_2| \otimes |a_2\rangle\langle a_2|$$

donde el factor $r(t)$ resulta ser

$$r(t) = \langle \varepsilon_1(t) | \varepsilon_2(t) \rangle = \prod_k \left[\cos(2g_k t) + i \left(|\gamma_k|^2 - |\delta_k|^2 \right) \sin(2g_k t) \right] \quad (3.2)$$

Este factor, llamado '*factor de correlación*', determina en cada instante el valor de los términos fuera de la diagonal. Zurek y sus colaboradores demuestran que, para N suficientemente grande, a medida que el tiempo transcurre, los estados $|\varepsilon_1(t)\rangle$ y $|\varepsilon_2(t)\rangle$ tienden hacia la ortogonalidad, de modo que, en un intervalo llamado *tiempo de decoherencia*, se obtiene

$$r(t) = \langle \varepsilon_1(t) | \varepsilon_2(t) \rangle \approx 0 \quad (3.3)$$

Con esta condición, entonces, se tiene que el operador reducido $\rho_{SM}(t)$ se aproxima a la forma dada por

$$\bar{\rho}_{SM} = |c_1|^2 |q_1\rangle\langle q_1| \otimes |a_1\rangle\langle a_1| + |c_2|^2 |q_2\rangle\langle q_2| \otimes |a_2\rangle\langle a_2| \quad (3.4)$$

donde los términos fuera de la diagonal han desaparecido, pero sin perderse la correlación entre el sistema objeto S y el aparato M . En virtud de la interacción con el entorno, De acuerdo con el formalismo de decoherencia, el sistema $S + M$, que al finalizar la pre-medición se encontraba en un estado puro, luego del tiempo de decoherencia puede ser descripto, para todo fin práctico, como encontrándose en un estado mezcla. Según los teóricos de la decoherencia, en este punto ambos problemas de la medición habrían sido resueltos.

El problema de la lectura definida quedaría resuelto porque, al converger a un estado mezcla, las superposiciones cuánticas del sistema $S + M$ desaparecen. Por lo tanto, dicho

estado sería susceptible de ser interpretado por ignorancia, asumiendo que luego de la medición el aparato M se encontrará en uno y sólo uno de sus posibles estados $|a_1\rangle$ o $|a_2\rangle$ y, por consiguiente, su variable indicadora registrará uno y sólo uno de sus valores a_1 o a_2 , sin problemas de superposición macroscópica.

El problema de la base preferida quedaría resuelto porque, debido a (3.3), los estados del entorno se hacen ortogonales. Con ello, la superposición en (3.1) se convierte en una descomposición triortonormal que, a diferencia de la biortonormal, es única (Andrew y Bub 1994; Auletta 2004). La ambigüedad presente en la descomposición biortonormal (en el caso en que los coeficientes sean iguales) ha desaparecido. Entre los infinitos cambios de base que introducía el problema de la base preferida en el nivel de la pre-medición, existe una que es privilegiada en virtud de la presencia del entorno: es la base de autoestados del hamiltoniano de interacción entre el aparato y el entorno. Sólo en esa base se preserva la información sobre la correlación establecida en la pre-medición entre sistema y aparato, a pesar de la imperfecta aislación del aparato respecto del entorno. Por eso, luego de trazar sobre el entorno no se pierde la correlación entre el sistema S y el aparato M , tal como lo pone de manifiesto el estado de la ecuación (3.4). Esta base privilegiada es la llamada 'base del puntero' (pointer basis. Zurek 1981, 1982), que depende del entorno particular que esté actuando en cada caso. En nuestro ejemplo es la base formada por los autoestados $\{|a_i\rangle\}$.

Con estos resultados, los teóricos de la decoherencia no sólo están convencidos de que su formalismo resuelve los problemas de la medición, sino que, adicionalmente, es el mecanismo capaz de dar cuenta de la emergencia del mundo clásico en virtud de su eficacia para explicar cómo las superposiciones cuánticas son eliminadas en las interacciones con sistemas macroscópicos.

No obstante, existen aún muchos problemas conceptuales en el mecanismo de la decoherencia, y una gran cantidad de críticas se han señalado a este respecto. En la próxima sección presentaremos algunas de ellas.

3.3 Algunas críticas a la decoherencia

Es de utilidad analizar las debilidades que presenta el formalismo de decoherencia como fuera presentado en la sección anterior. La primera que podemos mencionar, y quizás una de

las más importantes, consiste en señalar que el formalismo de la decoherencia inducida por el entorno sólo puede ser aplicado a sistemas abiertos en contacto con un entorno. Esto presenta de inmediato el problema de su capacidad para dar cuenta de la clasicidad observada en el universo como un todo que, por definición, es en un sistema cerrado, sin entorno (Zurek 1991; Pessoa 1998). Pese a que en los primeros estadios cosmológicos del universo los efectos cuánticos no son despreciables, las observaciones actuales parecen concordar perfectamente con el límite clásico tomado de la relatividad general sin involucrar en absoluto aspectos cuánticos. De acuerdo con la decoherencia inducida por el ambiente, y al contrario de todas las observaciones corroboradas, tendríamos que estar viviendo en un universo con características cuánticas a nivel cosmológico, porque sin entorno el universo no podría haber alcanzado el límite clásico.

La estrategia general para abordar esta inconsistencia consiste en dividir el universo en distintos grados de libertad mediante los cuales sea posible distinguir un sistema de interés respecto de ciertos otros grados de libertad que cumplen el papel de entorno. Por ejemplo, en aplicaciones de la teoría cuántica de campos, la división puede hacerse en términos de un “campo perturbativo” en relación al campo de “fondo” del universo (Calzetta, Hu y Mazzitelli 2001). Pero esta estrategia no sólo presenta la limitación de explicar la decoherencia sólo en una parte del universo, sino que además pone en escena otro problema grave del enfoque. En general, y ya no sólo al nivel cosmológico, no existe un criterio formal que determine la división objetiva entre sistema y entorno (Zurek 1998). En cada caso de interés esta división termina efectuándose *ad hoc* y de hecho pueden introducirse distintas divisiones, arrojando distintos resultados, aun para un mismo caso de estudio (Fortin 2012).

Si bien estas objeciones son graves y no pueden desestimarse, existen otras críticas más vinculadas al objetivo de esta tesis que presentaremos a continuación. Éstas se relacionan con el hecho de que, como hemos visto en la sección anterior, la decoherencia parece capaz de solucionar los problemas de la medición, pero de forma sólo aproximada. La correlación con estados ortogonales del entorno, que conduce a la anulación de los términos no diagonales, nunca se logra en forma exacta. El tiempo de decoherencia es meramente el tiempo a partir del cual los estados del entorno se vuelven aproximadamente ortogonales, y

para todos los fines prácticos los resultados obtenidos no pueden distinguirse del límite ideal asociado a la pérdida total de coherencia, como supone el límite clásico. Esto no sería grave bajo el supuesto de que toda teoría física sólo debe dar cuenta de nuestras percepciones sensibles (las “apariencias”). Como sostiene d’Espagnat, “la decoherencia explica las recién mencionadas apariencias, y éste es su resultado más importante” (d’Espagnat 2000). El propio Zurek encuadra este tipo de consideraciones dentro su interpretación existencial de la decoherencia, según la cual el problema de la emergencia de la clasicidad se reduce a explicar nuestra *percepción* de la clasicidad (Zurek 1998). Según este autor, la clasicidad obtenida por medio del mecanismo de decoherencia no es estrictamente objetiva, sino sólo “operacionalmente” objetiva.

Pero estas consideraciones encierran otro problema, aun más grave. El factor de correlación dado por la expresión (3.2), además de ser aproximado como cero en el límite clásico, resulta ser una función cuasiperiódica (Zurek 1981, 1982), por lo cual, estrictamente hablando, su límite para tiempo infinito ni siquiera está definido. Por el contrario, existe un tiempo de recurrencia para el cual los estados del entorno dejan de ser ortogonales, y la función de correlación se aparta de cero, de hecho se hace arbitrariamente cercana a uno (Zurek 1982). Así, la proclamada solidez de la correlación establecida por el entorno se debilita. Al perder la ortogonalidad de los estados del entorno, la descomposición (3.1) deja de ser única y la ambigüedad del cambio de base se presenta nuevamente. Una vez más, es sólo a los efectos prácticos que se puede considerar un límite efectivo para el factor de correlación, porque cuando el número N de grados de libertad del entorno se hace enormemente grande, el tiempo de recurrencia se hace tan grande o más que la edad del universo (Zurek 1982), consiguiéndose así una irreversibilidad al menos efectiva. Es por este motivo que en el marco del programa de decoherencia se presenta la necesidad crucial de considerar el entorno como un sistema con un enorme número de grados de libertad. De lo contrario, no sólo el problema de la base preferida no quedaría resuelto, ni siquiera en forma efectiva, sino que además se predecirían oscilaciones apreciables en la emergencia de lo clásico, lo cual estaría en completa oposición con lo observado.

Pero aquí hay más todavía: es en relación al otro problema de la medición, al de la lectura definida, que se pone en evidencia un error muy común en la comprensión del proceso de decoherencia. La convergencia al mundo clásico no es resultado exclusivo de la particular evolución temporal sujeta a la interacción con un entorno, sino de una adicional operación de “grano grueso”, que consiste de un recorte de ciertos grados de libertad que pueden ser discriminados en el sistema global completo (Castagnino, Lombardi y Vanni 2008). Este grano grueso se introduce a través de la operación de trazado sobre los grados de libertad del entorno luego de establecerse la correlación dada por la ecuación (3.1). En el marco de la explicación de la decoherencia, es esta operación lo único capaz de transformar el estado puro inicial en un estado final mezcla, y así pretender resolver el problema de la lectura definida. La necesidad de esta operación resulta evidente a partir de los principios mismos de la mecánica cuántica, ya que la evolución de los estados cuánticos no deja de ser unitaria por más que intervenga un sistema particular como el entorno. Así, en la proclamada solución del problema de la lectura definida por parte de la decoherencia interviene en forma esencial una operación que es independiente de la mayoría de las consideraciones que suelen hacerse sobre este enfoque, puesto que se trata de una mera operación algebraica.

Terminamos esta sección mencionando una última objeción respecto del tan mencionado límite clásico de la mecánica cuántica. Se ha dicho que este límite se obtiene como resultado de la pérdida de las superposiciones cuánticas que permite pasar de estados puros a estados mezcla. En general, el problema de la emergencia de la clasicidad puede ser puesto en términos de cómo la mecánica cuántica conduce a la ocurrencia de eventos que registramos como clásicos en nuestra experiencia. El mero hecho de perder las superposiciones por medio de la decoherencia, y obtener así estados mezcla, que pese a las advertencias mencionadas en el Capítulo 1, son forzados a una interpretación por ignorancia, no basta para justificar la ocurrencia de registros clásicos. El carácter diagonal del estado mezcla, que como hemos visto ni siquiera es producto del formalismo en sí sino de una operación de trazado, no basta para explicar qué evento particular, de los posibles interpretados en la mezcla, se actualizará en el proceso de medición (Adler 2003).

A la luz de todo lo dicho, parecen inevitables serias dudas sobre la efectividad del programa de la decoherencia inducida por el entorno para obtener el límite clásico y resolver los problemas de la medición. Respecto de los problemas de la medición, podemos decir, resumiendo, que la decoherencia resuelve el problema de la base privilegiada sólo de una manera efectiva, ya que estrictamente hablando, aun en presencia del entorno, la ambigüedad del cambio de base no se elimina esencialmente, y que resuelve el problema de la lectura definida por medio de la aplicación de una operación de grano grueso, la cual nada tiene que ver con la presencia de un entorno, y de un modo aún más cuestionable, a través de una interpretación por ignorancia del estado mezcla que no es admisible en el caso cuántico.

Capítulo 4. Valores definidos sin decoherencia

Es claro que, si se acepta el postulado del colapso como característica de la teoría, entonces queda justificada la existencia de salidas bien definidas en las mediciones, tal como indican los datos de la experiencia. Ahora bien, ¿qué sucede con la recíproca? Es decir, si se acepta la existencia de salidas bien definidas, ¿vale el colapso como una característica de la teoría?

En este capítulo demostraremos que, asumiendo valores definidos como nociones primitivas, el postulado del colapso puede ser deducido del formalismo cuántico sin entrar en contradicción con la existencia de evoluciones unitarias. Efectuaremos esta demostración dentro del esquema habitual de la teoría cuántica de la medición y sin apelar a ninguna noción de decoherencia.

Ahora bien, no vamos a encontrar una justificación al hecho de obtener valores definidos en las mediciones, sino una justificación de que la presencia de valores definidos no se contradice con el carácter unitario de la evolución de los estados cuánticos. Para ello, adicionalmente haremos uso de la noción la probabilidad condicional. Demostraremos que, asumiendo conocido el resultado en una primera medición, el cálculo de las probabilidades para una medición posterior, condicionalizadas al primer resultado, coincide con el cálculo de probabilidades que se obtendría de asumir el colapso en la primera medición (Laura y Vanni 2008a, 2008b). La clave de la demostración consistirá en incorporar a los aparatos de medición en las descripciones, y seguir la prescripción de predicar siempre sobre sus variables indicadoras, de modo que sea posible articular las conjunciones que participan en la fórmula de la probabilidad condicional, aun cuando correspondan a la conjunción de variables incompatibles en el sistema medido.

4.1 Mediciones ideales

Consideramos la medición ideal sucesiva de dos observables en un sistema S : en primer lugar, del observable Q por medio de un aparato M_A con variable indicadora A , y luego del observable R por medio del aparato M_B cuya variable indicadora es B . En el marco de la formulación original de la teoría cuántica de la medición, tal como fuera presentada en el Capítulo 2, hacemos uso de la evolución (2.3) que rige las mediciones ideales

$$|\Psi_0\rangle = |\psi\rangle|z_0\rangle = \sum_i c_i |q_i\rangle |a_0\rangle \rightarrow |\Psi_1\rangle = U|\Psi_0\rangle = \sum_i c_i |q_i\rangle |a_i\rangle \quad (4.1)$$

La diferencia ahora es que tendremos dos aparatos; por lo tanto, tratamos con la evolución del sistema compuesto $S + M_A + M_B$. Supongamos el estado inicial de este sistema compuesto viene dado por $|\psi\rangle|a_0\rangle|b_0\rangle$, donde $|\psi\rangle$ es un estado genérico del sistema S , y $|a_0\rangle$ y $|b_0\rangle$ son los estados de referencia de las variables indicadoras de los aparatos M_A y M_B respectivamente. Operando con la ecuación (4.1), obtenemos que la evolución del sistema compuesto resulta

$$\begin{aligned} |\Psi_0\rangle &= |\psi\rangle|a_0\rangle|b_0\rangle \\ &= \sum_i c_i |q_i\rangle|a_0\rangle|b_0\rangle \rightarrow |\Psi_1\rangle = \sum_i c_i |q_i\rangle|a_i\rangle|b_0\rangle \\ &= \sum_i c_i \sum_j r_{ij} |r_j\rangle|a_i\rangle|b_0\rangle \rightarrow |\Psi_2\rangle = \sum_{ij} c_i r_{ij} |r_j\rangle|a_i\rangle|b_j\rangle \end{aligned} \quad (4.2)$$

Antes de la primera medición, para poder aplicar la ecuación (4.1), hemos desarrollado el estado $|\psi\rangle$ en términos de los autoestados de la variable Q , es decir, consideramos $|\psi\rangle = \sum_i c_i |q_i\rangle$. En esta primera medición se produce una interacción entre S y M_A que correlaciona Q con la variable A del primer aparato. Antes de la segunda medición, y para volver a utilizar la ecuación (4.1), hemos desarrollado cada $|q_i\rangle$ en términos de los autoestados de la variable R , considerando $|q_i\rangle = \sum_j r_{ij} |r_j\rangle$. En esta segunda medición se produce una interacción entre S y M_B y así R queda correlacionada con B .

Haciendo uso de la fórmula habitual para la probabilidad condicional, podemos calcular las probabilidades para los diferentes valores r_j del observable R en la segunda medición, condicionadas respecto de un valor conocido q_i del observable Q en la primera. Efectuaremos el cálculo en términos de los correspondientes valores de las variables indicadoras A y B de los aparatos, las cuales quedan asociadas a las variables Q y R del sistema en virtud de las correlaciones. Este punto es crucial para asegurar que la conjunción de la fórmula de la probabilidad condicional esté bien definida. La conjunción entre propiedades de las variables A y B es legítima, ya que estas variables resultan trivialmente compatibles debido al hecho de que pertenecen a sistemas cuánticos distintos. La

probabilidad mencionada es entonces traducida en términos de los registros de los aparatos a la probabilidad del valor b_i en la variable B en el segundo aparato, condicionada al valor a_j en la variable A del primero. Trabajando en el sistema compuesto $S + M_A + M_B$ que incluye los dos aparatos, considerando que el estado final de ese sistema compuesto luego de completada la secuencia de mediciones viene dado por $\rho_2 = |\Psi_2\rangle\langle\Psi_2|$, haciendo uso de los proyectores que representan las propiedades de valor como fuera desarrollado en el Capítulo 1, y luego de ciertos cálculos algo largos pero directos (Laura y Vanni 2008a, 2008b), la fórmula de la probabilidad buscada resulta ser

$$\begin{aligned}
 P(b_j / a_i) &= \frac{P(b_j \wedge a_i)}{P(a_i)} \\
 &= \frac{\text{Tr}[\rho_2 (\mathbf{I}_S \otimes \Pi_{a_i} \otimes \Pi_{b_j})]}{\text{Tr}[\rho_2 (\mathbf{I}_S \otimes \Pi_{a_i} \otimes \mathbf{I}_B)]} \quad (4.3) \\
 &= \text{Tr}[\rho_{q_i} \Pi_{r_j}]
 \end{aligned}$$

donde $\rho_{q_i} = |q_i\rangle\langle q_i|$ y $\Pi_{r_j} = |r_j\rangle\langle r_j|$. La última ecuación reproduce la probabilidad de acuerdo con el postulado del colapso, porque es la probabilidad de obtener r_j como si el sistema S hubiera evolucionado en forma no unitaria al componente $|q_i\rangle$ de su superposición original. Sin embargo, el sistema compuesto, que incluye a los aparatos, nunca deja de evolucionar según la ecuación de Schrödinger, tal como se muestra en la ecuación (4.2).

Tenemos, entonces, que la probabilidad condicional aplicada en términos de los registros a_i y b_j de los aparatos en un estado ρ_2 que viene de una evolución unitaria del sistema compuesto $S + M_A + M_B$, como indica la expresión en el segundo renglón de (4.3), coincide con la probabilidad aplicada en términos de los registros r_i de la variable del sistema S sobre un estado ρ_{q_i} , que resulta ser el que se obtendría del postulado del colapso en la medición de Q con resultado q_i , como indica la el último renglón de (4.3).

De esta manera, el postulado del colapso puede ser entendido como la regla necesaria para compensar el efecto de restringir las descripciones al sistema objeto S ignorando los aparatos, lo cual implica el recorte de una parte del sistema completo $S + M_A + M_B$, que es el que verdaderamente evoluciona en forma unitaria en la medición. Cuando se considera la

totalidad del sistema que interviene en la interacción, la regla del colapso puede considerarse "espuria" e innecesaria. Para utilizar una analogía geométrica: sería algo así como la regla en los ángulos para compensar la perspectiva cuando nos restringimos a una representación plana de un objeto tridimensional. Este efecto, que se denomina contracción proyectiva, es en realidad un mecanismo agregado, "espurio" respecto de la verdadera descripción, pero indispensable para lograr el escorzo cuando se pretende realizar la descripción en una dimensión menos. En este sentido, el estado que proviene del postulado del colapso es una suerte de "sombra" sobre un espacio de menor dimensión del verdadero estado que es una superposición en un espacio mayor.

En esta idea simbólica motivada por nuestras deducciones se encuentra implícita la noción de proyección que, como ya sabemos, está asociada al postulado del colapso, porque éste implica la proyección de un vector de estado sobre un subespacio del espacio de Hilbert total del sistema. Sin embargo, a la luz de lo recién desarrollado, podemos decir que el postulado del colapso trata en realidad con un tipo menos trivial de proyección: la que proviene de proyectar el sistema compuesto "sobre" una de sus partes (el sistema objeto), condicionada al valor asignado en otras de las partes (los aparatos).

Hablamos de proyectar sobre una parte, a la que llamamos sistema y diferenciamos de las otras que llamamos aparatos, porque es la parte que deseamos medir. En rigor de verdad, los aparatos están aquí en pie de igualdad respecto del sistema que se quiere medir. Tanto el sistema objeto como los aparatos son subsistemas del sistema total que evoluciona según la ecuación de Schrödinger. No hay ningún privilegio de macroscopicidad que pueda diferenciar a los aparatos del sistema. La diferenciación entre éstos proviene del hecho que el sistema es la única subparte que interactúa con todas las otras que caracterizamos como aparatos. El colapso queda reproducido en la subparte del sistema compuesto que llamamos sistema objeto, y como consecuencia de la condicionalización aplicada sobre los registros de los demás subsistemas que se ponen en interacción con el primero, y que llamamos aparatos.

4.2 Mediciones no ideales

En el caso de mediciones no ideales aplicamos el mismo razonamiento de la sección anterior, pero ahora haciendo uso de la evolución (2.4) correspondiente a ese tipo de medición:

$$|\Psi_0\rangle = |\psi\rangle |a_0\rangle = \sum_i c_i |q_i\rangle |a_0\rangle \rightarrow |\Psi_1\rangle = U |\Psi_0\rangle = \sum_i c_i |\eta_i\rangle |a_i\rangle \quad (4.4)$$

Como en el caso anterior, utilizando la fórmula de la probabilidad condicional sobre el registro de los aparatos, y considerando la evolución unitaria del sistema compuesto $S + M_A + M_B$, es posible reducir la probabilidad para una segunda medición, conocido el resultado en la primera, en términos de un estado sobre el espacio del sistema S . Sin embargo, en este caso no reproduce exactamente el colapso, sino una versión especial de él que corresponde a una reducción a uno de los estados $|\eta_i\rangle$ que desarrollan el estado final dado en (4.4) (Laura y Vanni 2008a, 2008b). Concretamente, se obtiene que

$$P(b_j / a_i) = \frac{P(b_j \wedge a_i)}{P(a_i)} = \text{Tr} \left[\rho_{\eta_i} \Pi_{r_j} \right]$$

donde $\rho_{\eta_i} = |\eta_i\rangle \langle \eta_i|$. La diferencia con el caso anterior puede entenderse si consideramos la evolución de las mediciones no ideales en la etapa de calibración, tal como indica la ecuación (2.2). Al no preservarse el autoestado inicial del observable medido, el efecto de la condicionalización después de la primera medición se establece en relación con algunos de los componentes que participan en el desarrollo del estado posterior a tal medición, es decir, con algunos de los $|\eta_i\rangle$.

Como en el caso de las mediciones ideales, podemos decir de forma analógica, que el estado reducido que se obtiene en el espacio del sistema objeto S , estado que no corresponde a una evolución unitaria, es la "sombra" sobre el sistema objeto proyectada por el verdadero estado del sistema compuesto, que continúa siendo una superposición producto de los efectos de una evolución unitaria.

En el marco de una descripción que incluya a los aparatos, las demostraciones que hemos presentado tornan explícita la relación, normalmente aceptada en la teoría cuántica, entre un

autovalor y su correspondiente autovector (vínculo autovector-autovalor). En el postulado del colapso se considera que, al colapsar el sistema a un autoestado de la variable medida, esta variable asume el correspondiente autovalor como resultado en la medición. Nuestro trabajo demuestra que asumir un valor definido como resultado de la medición de una variable no necesariamente implica que el sistema se encuentre en el correspondiente autovector. Esta conclusión puede considerarse un fuerte respaldo para la hipótesis fundamental del gran número de interpretaciones modales, las cuales en esencia se basan en el hecho de relajar el mencionado vínculo autovector-autovalor tan utilizado en la interpretación ortodoxa de la mecánica cuántica.

Capítulo 5. Base preferida sin decoherencia

Habiendo delineado los conceptos básicos de la mecánica cuántica, presentado los problemas de la medición, y explicado la postura del formalismo de decoherencia, abordaremos ahora la propuesta de esta tesis respecto al problema de la base preferida.

Sin recurrir explícitamente a la decoherencia, muchos han intentado encontrar respuesta a las dificultades que introduce la medición recurriendo al carácter macroscópico de los instrumentos. Se apela a una condición de macroscopicidad que define y caracteriza todo aparato de medición, pero en rigor de verdad dicha condición nunca es formulada rigurosamente (Vanni y Laura 2005). Pese a que no fue hasta el advenimiento de la decoherencia que se articuló cierta noción de macroscopicidad referida al enorme número de grados de libertad de los instrumentos, se ha recurrido a dicha noción en forma reiterada a la hora de explicar muchas de las características de la medición cuántica, en particular en ejemplos fundacionales como el experimento de Stern y Gerlach (Stern y Gerlach 1922). En esta línea, no es difícil imaginar un argumento que recurra a la macroscopicidad para privilegiar una base de resultados posibles y que permita franquear el problema de la base preferida.

En este capítulo presentamos el argumento de la macroscopicidad en el experimento de Stern y Gerlach, y luego demostraremos que dicho argumento no resulta suficiente para dar una respuesta completa al problema de la base preferida. Finalmente presentaremos nuestra propuesta, que consiste en un análisis lógico del proceso de medición, sin hacer uso en absoluto del proceso de decoherencia.

5.1 Condición de macroscopicidad en el experimento de Stern y Gerlach

El experimento de Stern y Gerlach es una de las primeras experiencias que permitieron poner en evidencia una propiedad inherentemente cuántica de las partículas, llamada spin. El experimento consiste esencialmente en dirigir partículas a través de un imán capaz de generar un campo magnético no uniforme en una dirección privilegiada del espacio. En el marco de la teoría cuántica se puede demostrar que, por medio de su spin, una partícula puede interactuar

con el campo magnético externo de modo tal que al atravesar el imán es deflectada en alguna de dos direcciones en virtud del carácter bivaluado de su proyección de spin (en el caso de spin $1/2$). Dicho de otro modo, si se coloca una pantalla frente al imán y a este último se lo orienta en la dirección vertical, la dirección z , se obtienen dos posibles resultados como salidas del dispositivo: “mancha arriba” y “mancha abajo”, las cuales se corresponden con las zonas de detección de las partículas que abandonan el imán.

Decimos que el aparato de Stern y Gerlach mide la proyección del spin en z , magnitud que indicaremos como S_z , porque es capaz de establecer una correlación entre esa variable y la componente del momento en la misma dirección. La partícula cambiará su momento de acuerdo con el signo de la componente del spin, y así será desviada hacia arriba o hacia abajo, lo cual conducirá a su vez a una correlación con su ubicación superior o inferior en la pantalla. Partículas con proyección de spin positiva quedarán con sus funciones de onda localizadas en la zona superior de la pantalla (se detectará mancha arriba), y partículas con proyección de spin negativa quedarán con sus funciones de onda localizadas en la zona inferior de la pantalla (se detectará mancha abajo).

Éste es un ejemplo de medición donde la magnitud microscópica medida (proyección de spin) queda naturalmente correlacionada con una magnitud macroscópicamente distinguible (ubicación arriba o abajo) sin ningún mecanismo de amplificación ulterior en el aparato, y sin ninguna consideración referida a los muchos grados de libertad del aparato por tratarse de un dispositivo macroscópico.

Si indicamos con $|+\rangle$ y $|-\rangle$ a los autoestados de S_z , correspondientes a la proyección del spin en la dirección z positiva y negativa respectivamente, entonces en la situación más general una partícula antes de atravesar el imán se encontrará en una superposición dada por $\alpha|+\rangle + \beta|-\rangle$. No obstante, para ubicarnos dentro de las condiciones del problema del cambio de base, consideremos coeficientes iguales: $\alpha = \beta = 1/\sqrt{2}$. Con esta preparación como entrada del dispositivo de Stern y Gerlach, es posible demostrar se producirá la evolución (Ballentine 1998)

$$|\Psi_0\rangle = \frac{1}{\sqrt{2}}(|+\rangle + |-\rangle)|\Phi_0\rangle \rightarrow |\Psi_1\rangle = \frac{1}{\sqrt{2}}|+\rangle|\Phi_+\rangle + \frac{1}{\sqrt{2}}|-\rangle|\Phi_-\rangle$$

donde $|\Phi_+\rangle$ y $|\Phi_-\rangle$ son los estados de la parte espacial de la función de onda correspondientes a las ubicaciones “mancha arriba” y “mancha abajo” respectivamente. El estado $|\Phi_0\rangle$ corresponde a la ubicación central de la partícula antes de entrar al aparato, es decir, a “mancha central” en el caso de que no se deflectara.

En virtud de la distinguibilidad macroscópica asociada a $|\Phi_+\rangle$ y $|\Phi_-\rangle$, es posible elaborar un argumento a modo de solución al problema de la base. Este argumento se basa en considerar como privilegiada la base asociada a salidas localizadas y distinguibles. Para desarrollarlo en detalle consideremos el cambio de base dado por

$$\begin{aligned} |+\rangle &\rightarrow |\square\rangle = \frac{1}{\sqrt{2}}(|+\rangle + |-\rangle) & |\Phi_+\rangle &\rightarrow |\Phi_\square\rangle = \frac{1}{\sqrt{2}}(|\Phi_+\rangle + |\Phi_-\rangle) \\ |-\rangle &\rightarrow |\otimes\rangle = \frac{1}{\sqrt{2}}(|+\rangle - |-\rangle) & |\Phi_-\rangle &\rightarrow |\Phi_\otimes\rangle = \frac{1}{\sqrt{2}}(|\Phi_+\rangle - |\Phi_-\rangle) \end{aligned} \quad \begin{matrix} (5.1) \\ (5.2) \end{matrix}$$

Con estas transformaciones, el estado final después de la medición puede ser expresado como

$$|\Psi_1\rangle = \frac{1}{\sqrt{2}}|+\rangle|\Phi_+\rangle + \frac{1}{\sqrt{2}}|-\rangle|\Phi_-\rangle = \frac{1}{\sqrt{2}}|\square\rangle|\Phi_\square\rangle + \frac{1}{\sqrt{2}}|\otimes\rangle|\Phi_\otimes\rangle \quad (5.3)$$

Los estados $|\square\rangle$ y $|\otimes\rangle$ en las ecuaciones (5.1) resultan ser los autoestados de la variable S_x , correspondientes a la proyección positiva y negativa, respectivamente, del spin en el eje x . En esta situación se esperaría que $|\Phi_\square\rangle$ y $|\Phi_\otimes\rangle$ fueran los estados correspondientes a las ubicaciones “mancha derecha” y “mancha izquierda” respectivamente¹³. Si así fuera el caso, en virtud de la ambigüedad del cambio de base como viene dada en (5.3), no podríamos determinar si el Stern y Gerlach arrojará alguna de las salidas “arriba” o “abajo”, de modo de asegurar que se midió S_z , o si por el contrario arrojará alguna de las salidas “derecha” o “izquierda”, de modo de asegurar que se midió S_x .

¹³ Asociamos las dos direcciones del eje x con “derecha” e “izquierda”, así como hemos asociado las dos direcciones del eje z con “arriba” y “abajo”.

Pues bien, por tratarse $|\Phi_+\rangle$ y $|\Phi_-\rangle$ de estados distinguibles macroscópicamente, el cambio de base dado por las ecuaciones (5.1) y (5.2), aunque legítimo desde el punto de vista matemático, carece de sentido físico. Esto es así porque combinaciones de estados macroscópicos asociados a localizaciones “arriba” y “abajo” no da como resultado estados asociados a localizaciones “derecha” e “izquierda”. Los estados $|\Phi_\square\rangle$ y $|\Phi_\otimes\rangle$, si bien son los correctos para obtener el segundo desarrollo en la ecuación (5.3), no corresponden a salidas “derecha” e “izquierda” legítimas del aparato. Así, cuando el campo magnético del aparato de Stern y Gerlach se orienta en la dirección z , la partícula que llega es desviada macroscópicamente hacia arriba o hacia abajo, pero nunca a derecha o izquierda. En esa situación las únicas salidas posibles son “mancha arriba” y “mancha abajo” porque son las únicas asociadas a estados bien localizados macroscópicamente. De este modo, el aparato de Stern y Gerlach en la dirección z sólo puede medir S_z . Para medir S_x se necesitará rotar el campo magnético hacia la dirección x . Sólo en ese caso el imán producirá desviaciones macroscópicas hacia la derecha o hacia la izquierda, y entonces las salidas vinculadas a los estados $|\Phi_\square\rangle$ y $|\Phi_\otimes\rangle$ son las que cobrarán sentido físico, mientras que las vinculadas a $|\Phi_+\rangle$ y $|\Phi_-\rangle$ lo perderán.

Este argumento determinaría un criterio capaz de privilegiar naturalmente una de las infinitas bases que permite el cambio de base, resolviendo así la ambigüedad. La base privilegiada sería aquella asociada a las salidas que, por la propia preparación del dispositivo experimental, se requerirá devengan macroscópicamente distinguibles, de modo que dicho dispositivo opere como aparato de medición y deje así registro en nuestra experiencia sensible.

De este modo, el diseño del aparato mismo determina cuál es la base preferida. La respuesta a la pregunta acerca de qué está midiendo el observador se responde por la simple elección de lo que el observador desea medir, elección que implementará por medio del dispositivo experimental necesario para hacer manifiesta la discriminación de los valores de la magnitud que eligió medir. Este argumento, aunque no determina explícitamente cuál es la base privilegiada, al menos parece en principio suficiente para eliminar formalmente la ambigüedad del cambio de base. En el experimento de Stern y Gerlach su implementación

parece ser clara y satisfactoria. No obstante, en mediciones más generales, como es en el caso de las “mediciones bit-a-bit” (esto es, sistemas de spin $1/2$ que “miden” sistemas de spin $1/2$), se torna insuficiente.

Veremos que el problema de la base preferida no está vinculado a la posibilidad de distinción macroscópica, distinción que corresponde a un plano epistemológico, sino a un hecho más profundo vinculado con la capacidad de distinción entre entidades en un plano ontológico. En la próxima sección, a la vez que elaboraremos el argumento que nos permitirá finalmente presentar nuestra propuesta, demostraremos que la macroscopicidad no resulta suficiente para la solución del problema de la base preferida.

5.2 El aparato de cuatro luces

En esta sección demostraremos que el problema de la base preferida descansa sobre el supuesto del funcionamiento de un aparato capaz de activar más de un juego de salidas, y que dicho aparato no puede existir por requerimientos de la propia mecánica cuántica, independientemente de cualquier condición de macroscopicidad. En el experimento de Stern y Gerlach, si el imán es orientado en la dirección z , se tiene el juego de salidas “arriba” - “abajo”. Si, en cambio, es orientado en la dirección x , se tiene el juego “derecha” - “izquierda”. El aparato de Stern y Gerlach, sin importar cuál sea la preparación de entrada, sólo puede activar un juego de salidas a la vez, lo cual está determinado por la disposición experimental relacionada con la dirección del imán.

Esta disposición experimental determina sólo un juego de salidas como legítima, pero la macroscopicidad no es esencial para ello. A fin de indagar con mayor profundidad esta tesis, analizaremos las llamadas “mediciones bit-a-bit” (“bit-a-bit measurement”) (Peres 1974; Paz y Zurek 2002). En este tipo de mediciones se tiene un sistema de spin $1/2$ que mide otro sistema de spin $1/2$. Supongamos que $|+\rangle$ y $|-\rangle$ son los autoestados de la proyección de spin S_z que se quiere medir en uno de esos sistemas de spin $1/2$, al que consideraremos sistema objeto. Supongamos que $|\uparrow\rangle$ y $|\downarrow\rangle$ son los autoestados de misma variable pero en el otro sistema de spin $1/2$, al que consideraremos aparato. El estado general del sistema objeto será $\alpha|+\rangle + \beta|-\rangle$, donde tomaremos $\alpha = \beta = 1/\sqrt{2}$ según la condición del problema de

cambio de base. Si asumimos $|\uparrow\rangle$ como estado inicial del sistema aparato, se puede demostrar que existe un hamiltoniano de interacción entre los dos sistemas de spin $1/2$ tal que produce la siguiente evolución (Peres 1974):

$$|\Psi_0\rangle = \frac{1}{\sqrt{2}}(|+\rangle + |-\rangle)|\uparrow\rangle \rightarrow |\Psi_1\rangle = \frac{1}{\sqrt{2}}|+\rangle|\uparrow\rangle + \frac{1}{\sqrt{2}}|-\rangle|\downarrow\rangle \quad (5.4)$$

cuyo estado final, ante el siguiente cambio de base

$$\begin{aligned} |+\rangle &\rightarrow |\square\rangle = \frac{1}{\sqrt{2}}(|+\rangle + |-\rangle) & |\uparrow\rangle &\rightarrow |\rightarrow\rangle = \frac{1}{\sqrt{2}}(|\uparrow\rangle + |\downarrow\rangle) \\ |-\rangle &\rightarrow |\otimes\rangle = \frac{1}{\sqrt{2}}(|+\rangle - |-\rangle) & |\downarrow\rangle &\rightarrow |\leftarrow\rangle = \frac{1}{\sqrt{2}}(|\uparrow\rangle - |\downarrow\rangle) \end{aligned}$$

puede ser expresado como

$$|\Psi_1\rangle = \frac{1}{\sqrt{2}}|+\rangle|\uparrow\rangle + \frac{1}{\sqrt{2}}|-\rangle|\downarrow\rangle = \frac{1}{\sqrt{2}}|\square\rangle|\rightarrow\rangle + \frac{1}{\sqrt{2}}|\otimes\rangle|\leftarrow\rangle \quad (5.5)$$

Esta expresión es análoga al estado final del Stern y Gerlach dado en (5.3) pero, a diferencia de tal caso, ahora no tenemos ningún argumento de macroscopicidad para privilegiar algún juego de salidas del sistema aparato. En el experimento de Stern y Gerlach con el campo magnético en la dirección z , el desarrollo en términos de los estados $|\Phi_{\square}\rangle$ y $|\Phi_{\otimes}\rangle$ se descarta porque, si bien estos estados eran combinaciones de $|\Phi_{+}\rangle$ y $|\Phi_{-}\rangle$ asociados a mancha arriba y abajo, no representaban verdaderas salidas macroscópicas derecha e izquierda como las que produce el aparato cuando el imán es rotado en la dirección x .

En las mediciones bit a bit no sucede lo mismo. El segundo desarrollo que muestra (5.5), en términos de los estados $\{|\rightarrow\rangle, |\leftarrow\rangle\}$, es tan legítimo como el primero en términos de los estados $\{|\uparrow\rangle, |\downarrow\rangle\}$. Los estados $\{|\rightarrow\rangle, |\leftarrow\rangle\}$ no sólo son combinaciones de $\{|\uparrow\rangle, |\downarrow\rangle\}$, sino que quedan asociados a un verdadero juego de salidas, al menos tan legítimo como el juego

de salidas asociadas a $\{|\uparrow\rangle, |\downarrow\rangle\}$, supuesto el segundo sistema de spin $1/2$ como aparato. El argumento de macroscopicidad no es aplicable aquí.

Se podría objetar que las mediciones bit a bit no son “verdaderas” mediciones puesto que, si bien establecen una correlación entre sistema y aparato, no producen registro a nivel macroscópico (Ballentine 1998, p. 174). Es aquí donde los partidarios de la decoherencia señalan la importancia del entorno para “completar la medición”. Como ya hemos desarrollado en el Capítulo 3, es a través de una posterior interacción del aparato con el entorno, con muchos grados de libertad, que la correlación establecida entre el sistema y aparato es capaz de trasladarse a un estado sin superposiciones cuánticas compatible con un registro macroscópico, y que además corresponde a una base particular privilegiada en el proceso. Diversas objeciones al enfoque de la decoherencia ya fueron discutidas en el mencionado capítulo, y el experimento de Stern y Gerlach podría incluirse entre ellas. Dicho experimento es muy particular porque produce correlaciones macroscópicas y privilegia una base de manera natural sin que medie ningún otro sistema como entorno. La decoherencia es incapaz de brindar una respuesta en este caso, al menos la versión de la decoherencia inducida por el entorno. No negamos la relevancia de la decoherencia para solucionar los problemas de la medición y en particular el de la base preferida. Lo que pretendemos mostrar aquí es que este último problema puede ser reducido a algo más fundamental sin apelar a ningún mecanismo de decoherencia.

El ejemplo de la medición bit a bit será el inicio de nuestro argumento. Como fuera señalado, en ese caso ambos juegos de salidas involucradas están en pie de igualdad, y esto establecería una ambigüedad entre bases cuya distinción macroscópica, veremos, no es esencial al argumento.

El problema del cambio de base se presentaría como problema si existiese un aparato capaz de habilitar la actualización de los distintos juegos de salidas, independientemente del hecho de que algún mecanismo de amplificación posterior transformara esa actualización en un registro de nuestra percepción. Sean $\{\uparrow, \downarrow\}$ las salidas asociadas a los estados $\{|\uparrow\rangle, |\downarrow\rangle\}$, y $\{\rightarrow, \leftarrow\}$ las salidas asociadas a los estados $\{|\rightarrow\rangle, |\leftarrow\rangle\}$ en la medición bit a bit arriba analizada. Podríamos imaginar algún mecanismo de amplificación macroscópica para cada

uno de estos juegos de salidas, por ejemplo, que condujeran al encendido de luces de distintos colores. Llamemos a este hipotético dispositivo “aparato de cuatro luces”, una para cada salida: $\{\uparrow, \downarrow, \rightarrow, \leftarrow\}$. Como la medición bit a bit deja en pie de igualdad los distintos juegos de salidas $\{\uparrow, \downarrow\}$ y $\{\rightarrow, \leftarrow\}$, este aparato de cuatro luces estaría trasladando la ambigüedad del cambio de base a nivel macroscópico. Como consecuencia del problema del cambio de base, no podríamos determinar desde la teoría cuál juego de luces veremos encendidas y, por lo tanto, qué mide semejante aparato.

Pues bien, en la próxima sección demostraremos que este aparato no puede existir, independientemente de la transformación al mundo macroscópico que promueven las luces de colores.

5.3 ¿Detección de entidades lingüísticas?

Para desarrollar nuestro argumento aquí sólo necesitamos del caso de mediciones ideales. La aplicación al caso más general de mediciones es inmediata y no aporta diferencias significativas. Consideremos la medición de la variable Q con autoestados $\{|q_i\rangle\}$ en un sistema objeto, mediante un aparato de variable indicadora A con autoestados $\{|a_i\rangle\}$. Supongamos inicialmente el sistema preparado en el estado $|\psi\rangle = \sum c_i |q_i\rangle$. Bajo las condiciones del problema del cambio de base, donde los coeficientes c_i que participan en el desarrollo del estado inicial son todos iguales, ante el cambio de base dado por $\{|q_i\rangle\} \rightarrow \{|q'_i\rangle\}$ y $\{|a_i\rangle\} \rightarrow \{|a'_i\rangle\}$ tenemos las dos posibles evoluciones

$$|\Psi_0\rangle = \sum_i c_i |q_i\rangle |a_0\rangle \rightarrow |\Psi_1\rangle = U |\Psi_0\rangle = \begin{cases} \sum_i c_i |q_i\rangle |a_i\rangle & \text{(a)} \\ \sum_i c_i |q'_i\rangle |a'_i\rangle & \text{(b)} \end{cases} \quad (5.6)$$

Como ya fuera analizado en la Sección 2.3, es la igualdad de los estados finales, aquí representados en las expresiones (a) y (b) de la última ecuación, lo que manifiesta el problema

de cambio de base, porque con el mismo proceso de medición no sabemos si hemos medido Q por medio de la variable indicadora A , estableciéndose la correlación del estado (a), o si, por el contrario, hemos medido Q' por medio de la variable indicadora A' , estableciéndose la correlación del estado (b). Y esto aun cuando Q y Q' puedan resultar variables incompatibles en el sistema objeto.

Pues bien, esta situación resultaría verdaderamente un problema si existiese un aparato capaz de habilitar, sin cambiar nada en él, los dos juegos de salidas $\{a_i\}$ y $\{a'_i\}$. Decimos 'sin cambiar nada en él' porque, si fuera necesaria la implementación de un cambio en el dispositivo experimental para habilitar uno de los juegos o el otro, entonces sería posible determinar qué se está midiendo de acuerdo con ese cambio, como sucede en el experimento de Stern y Gerlach desarrollado en una sección anterior. En ese experimento el dispositivo de medición es tal que la rotación del imán de la dirección z a la dirección x para habilitar el juego de salidas "derecha"- "izquierda, por ejemplo, nos permite diferenciar que en el primer caso medimos S_z y en el segundo caso medimos S_x . Sólo con la existencia de un aparato que permitiera habilitar los dos juegos de salidas posibles (el aparato que hemos llamado de cuatro luces), es que las dos correlaciones (a) y (b) en (5.6) cobrarían sentido físico independiente, y su igualdad implicaría una ambigüedad seria en la determinación de qué se midió.

Procedamos ahora por el absurdo. Supongamos que efectivamente existe tal aparato de cuatro luces, de modo que (a) y (b) en (5.6) efectivamente impliquen mediciones distintas con dicho aparato. Supongamos, auxiliamente y por simplicidad, que las variables A y A' que participan en esas mediciones distintas comparten el mismo valor de referencia, de modo que $|a_0\rangle = |a'_0\rangle$. Entonces, como producto de la calibración de la medición en (a), si preparamos el sistema en el estado inicial $|q_i\rangle$, se activará la salida de valor a_i asociado a $|a_i\rangle$, puesto que, procediendo según la ecuación (2.1), tenemos

$$|q_i\rangle|a_0\rangle \rightarrow |q_i\rangle|a_i\rangle \quad (5.7)$$

Por otro lado, de la calibración de la medición en (b), si preparamos el sistema en el estado inicial $|q_i'\rangle$ se activara la salida de valor a_i' asociado a $|a_i'\rangle$,

$$|q_i'\rangle|a_0\rangle \rightarrow |q_i'\rangle|a_i'\rangle \quad (5.8)$$

Ahora bien, por ser $\{|q_i\rangle\}$ y $\{|q_i'\rangle\}$ bases del espacio de Hilbert del sistema, tenemos que cada $|q_i'\rangle$ puede ser desarrollado en términos de los $|q_i\rangle$. De hecho, es el cambio de base aquello que las vincula. Consideremos en particular el siguiente vínculo: $|q_i'\rangle = \sum_k n_k |q_k\rangle$. Pero entonces, de la medición cuya calibración viene dada por (5.7), tenemos que, como consecuencia de la linealidad de la evolución temporal,

$$|q_i'\rangle|a_0\rangle = \sum_k n_k |q_k\rangle|a_0\rangle \rightarrow \sum_i n_i |q_i\rangle|a_i\rangle \quad (5.9)$$

Consideremos ahora las ecuaciones (5.8) y (5.9). La ecuación (5.9) supone la activación de la salida asociada al valor a_i ; en cambio, la ecuación (5.8) supone la activación de alguna de las salidas asociadas a los a_i' . Con total independencia de cualquier condición de macroscopicidad que permita una "distinción humana" entre las salidas $\{a_i\}$ y las $\{a_i'\}$, si éstas corresponden a bases distintas, deben representar realidades físicas distintas ya que refieren a observables que diferentes y, por lo tanto, debe ser posible distinguirlas. Pero si las salidas $\{a_i\}$ y $\{a_i'\}$ pueden ser distinguidas, entonces las ecuaciones (5.8) y (5.9) no pueden ser simultáneamente válidas (en el sentido que se habiliten simultáneamente las salidas $\{a_i\}$ y las $\{a_i'\}$), puesto que si el aparato no cambió, la interacción no cambió y el estado inicial no cambió, entonces las salidas del aparato deberían ser las mismas. De lo contrario se podría distinguir (5.8) y (5.9) como procesos distintos, lo cual no puede ser posible si no se ha cambiado nada; nada físico al menos, pues la única diferencia entre (5.8) y (5.9) es que desarrollamos el estado inicial del sistema $|q_i'\rangle$ en distintas bases matemáticas. La posibilidad de registrar ya sea alguno de los $\{a_i\}$ o de los $\{a_i'\}$ nos permitiría distinguir, según sea el

caso, si hemos expresado el estado inicial como $|\psi\rangle = |q_i'\rangle$ en virtud de (5.8), o si por el contrario, lo hemos expresado como $|\psi\rangle = \sum n_k |q_k\rangle$ en virtud de (5.9).

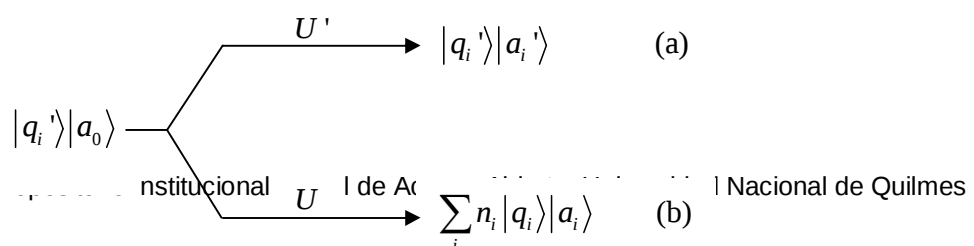
Esta conclusión no tiene sentido físico, puesto que $|q_i'\rangle$ y $\sum n_k |q_k\rangle$ son estrictamente idénticos, en el sentido de identidad lógica. Sólo son dos expresiones matemáticas, dos "nombres" para designar la misma entidad extralingüística. La existencia del aparato supuesto, el aparato de cuatro luces, implicaría la posibilidad de un aparato capaz de "medir" entidades lingüísticas, lo cual es conceptualmente absurdo.

Por consiguiente, dada una cierta interacción, un mismo aparato cuántico no puede activar más que un juego de salidas distintas. De este modo, sólo una de las correlaciones (a) o (b) expresadas en (5.6) tiene sentido físico, independientemente de la igualdad matemática entre dichas ecuaciones. Matemáticamente el cambio de base es correcto. Es por un argumento lógico en el contexto de la estructura de la medición cuántica, y con independencia de cualquier distinción macroscópica –que sólo es necesaria para nuestra percepción–, que concluimos que o bien se activan las salidas $\{a_i\}$, o bien las $\{a_i'\}$, pero no los dos juegos de salidas a la vez sin cambiar algo en el aparato.

Este argumento elimina la ambigüedad del problema del cambio de base respecto de la posibilidad de medir más de un observable. Como uno y sólo uno de los posibles juegos de salidas podrán habilitarse como resultado del proceso en (5.6), dicho proceso está representando la medición de un único observable. Falta determinar ahora cuál es precisamente ese observable.

5.4 ¿Cuál juego de salidas?

En la sección anterior hemos mostrado que, en el marco de la teoría cuántica de la medición, todo aparato sólo puede activar un único juego de salidas asociado a una determinada base. Intentemos ahora especificar cuál es dicha base. Usemos las ecuaciones (5.8) y (5.9) pero de la siguiente manera más esquemática



(5.10)

Si se activa uno de los $\{a_i\}$, eso significa que se produjo la evolución representada por la flecha inferior de (5.10), regida por el operador de evolución U . Si, en cambio, se activa uno de los $\{a'_i\}$, eso significa que se produjo la evolución representada por la flecha superior de (5.10), regida por el operador U' .

Ambas evoluciones parten de un mismo estado inicial, pero conducen a estados diferentes, uno vinculado a la activación de los $\{a_i\}$, y el otro a la activación de los $\{a'_i\}$. Pero, dado el carácter determinista de las evoluciones cuánticas, la única manera de que esto suceda es si $U \neq U'$. Esto significa que, aunque $\{|q_i\rangle\}$ y $\{|q'_i\rangle\}$ participen en el cambio de base de modo de establecer la ambigüedad representada por la igualdad de los estados en (5.6), pero expresados como distintas correlaciones, dichas correlaciones tienen que ser producto de procesos diferentes caracterizados por operadores de evolución diferentes, los cuales pueden ser discriminados al nivel de la calibración. Dicho de otro modo, aunque se presente la ambigüedad de distintas correlaciones con un mismo estado final –y como consecuencia de que sólo un juego de salidas puede activarse–, cuando las posibles mediciones correspondientes a esas correlaciones se analizan al nivel de la calibración, se comprueba que dichas correlaciones no pueden ser producidas por un mismo proceso físico.

Debido a que el operador de evolución que caracteriza la medición de un observable es función unívoca de dicho observable, se desprende que la especificación de U conduce a la especificación del observable que se mide. Así la ambigüedad teórica que presenta el problema de la base preferida es en realidad una falsa ambigüedad. Desde la propia teoría se puede especificar perfectamente la descripción de qué observable es medido por medio de la especificación del operador de evolución utilizado para conducir el estado inicial a alguna de las correlaciones del estado final (Lombardi y Vanni 2010). Aunque el estado final pueda ser el mismo expresado por medio de distintas correlaciones, el proceso necesario para producir una u otra es único, y queda determinado por el particular operador de evolución implicado.

Así, la respuesta a la pregunta acerca de qué se mide en el proceso (5.6) es, en algún sentido, contestada tautológicamente: se mide el observable que hemos elegido medir, es

decir, el observable que participa en el operador de evolución U incorporado en la descripción teórica del aparato de medición.

Esto elimina completamente la aparente ambigüedad del cambio de base. Con el argumento de la sección anterior eliminamos la ambigüedad respecto de la posibilidad de la medición de más de un observable. De los infinitos observables vinculados a las respectivas posibles correlaciones con las que se puede especificar el estado final de la medición, sólo se mide uno. Con el argumento de esta sección, hemos eliminado la indeterminación de cuál de ellos es el observable efectivamente medido.

Como puede observarse, nuestra estrategia elimina las ambigüedades del cambio de base presentes en el estado final de la medición, al diferenciar aspectos del proceso que participan antes de ser establecido dicho estado final. Este enfoque se distingue netamente de la solución brindada por el mecanismo de decoherencia. El formalismo de decoherencia intenta eliminar la ambigüedad del estado final en la medición apelando a un proceso posterior al establecimiento de la correlación, la interacción con el entorno, lo cual conduce a que el proceso finalmente estudiado sea diferente del original considerado en la teoría de la medición. En el presente trabajo, por el contrario, hemos eliminado la ambigüedad apelando al proceso previo al establecimiento de la correlación. Es dicho proceso el que determina la correlación, aunque matemáticamente pueda estar igualada a expresiones de otras correlaciones.

Capítulo 6. Conclusiones

Comenzamos esta tesis presentando, en el Capítulo 1, una introducción a la teoría cuántica que, si bien concisa, suponemos organizada de manera consistente para recorrer los aspectos conceptuales más relevantes de la teoría y algunas de sus peculiaridades.

En el Capítulo 2 procedimos a presentar lo que actualmente se conoce como la teoría cuántica de la medición, de modo de establecer el marco necesario para exponer los dos problemas nucleares de la teoría a este respecto. Éstos son el problema de la lectura definida y el problema de la base preferida.

Si bien el primero es históricamente presentado como "el" problema de la medición, pese a su mucha menor difusión el segundo es tan importante como el primero. De hecho, ya lo hemos dicho, lo empeora, porque si el problema de la lectura definida establece una indeterminación teórica sobre la definición del único valor que se registra en la medición, el problema de la base preferida supone una indeterminación teórica sobre la base de componentes que fija el conjunto de los valores posibles. Dicho de otro modo, el primer problema implica una incapacidad de la teoría en la determinación del valor específico que se registra en la medición, el segundo problema implica una incapacidad en la determinación del particular "menú" de posibilidades al que pertenece dicho valor.

Respecto del problema de la lectura definida, la respuesta tradicional consiste en apelar al postulado del colapso, que combinado con la regla de Born de la probabilidad es suficiente para elaborar una descripción probabilística de la mecánica cuántica. El postulado es todo lo efectivo que se requiere para las predicciones probabilísticas; no obstante, como fue discutido, resulta conceptualmente inadmisibile. Se trata de un postulado ad hoc que, de hecho, viola la ley de evolución fundamental de la mecánica cuántica: la evolución unitaria que establece la ecuación de Schrödinger. Actualmente cierta justificación al problema de la lectura definida se obtiene en el contexto del formalismo de decoherencia, porque puede reproducir, dentro de ciertos límites, una versión débil del colapso. Otras propuestas, aunque más extravagantes, han sido ensayadas pero sin mucho éxito; hemos efectuado un breve recorrido por ellas al final de la Sección 2.2.

Respecto del problema de la base preferida, la respuesta tradicional y, debemos agregar, la presentación misma del problema, aparece con la difusión del formalismo de decoherencia. Debido a la relevancia histórica de este enfoque respecto de ambos problemas de la medición, le hemos dedicado un capítulo entero, el Capítulo 3. En él hemos presentado la idea básica –al menos en su enfoque ortodoxo (decoherencia inducida por el entorno)–, la cual se reduce a explotar los efectos de la interacción con un sistema adicional, el entorno, cuando el proceso de medición es considerado sobre un sistema abierto. Si bien el enfoque ha tenido muchos éxitos y actualmente se considera la ortodoxia en el tema, se encuentra realmente inmerso en problemas conceptuales, algunos de los cuales hemos presentado al final del Capítulo 3.

El Capítulo 4 se dedicó al problema de la lectura definida, más precisamente, al análisis del colapso como solución al problema de la lectura definida. En el presente capítulo no brindamos una respuesta que permita resolver el problema de la existencia de valores bien definidos, sino que presentamos un argumento que, partiendo de ellos como nociones primitivas, sirve para otro propósito. Es claro que, aceptando al colapso como postulado adicional en la teoría, se deriva la existencia de valores definidos en la medición. Nosotros demostramos la recíproca de esta afirmación, es decir, hemos mostrado que, aceptando valores definidos, es posible derivar el colapso dentro de la teoría sin abandonar las evoluciones unitarias.

La demostración es completamente independiente del formalismo de decoherencia, y se basa en la estrategia crucial de incorporar los aparatos en la descripción de una secuencia de mediciones cuánticas. Calculando la probabilidad para una segunda medición condicionalizada sobre los valores obtenidos en la primera, y bajo la prescripción de predicar únicamente en términos de las variables de los aparatos, hemos demostrado que el resultado coincide con el que se obtendría por medio del postulado del colapso cuando se predica sobre los registros del sistema como si éste estuviera aislado.

Sobre esta base, hemos mostrado que el colapso, que es inconsistente con una evolución unitaria, es consecuencia de restringir la descripción de una evolución unitaria en el sistema compuesto que incluye a los aparatos al espacio del sistema objeto. Este resultado expresa, de una manera particular y en el contexto de la teoría cuántica de la medición, el carácter

holístico de los fenómenos cuánticos. En efecto, recortar las descripciones solamente a una parte del todo, aun cuando se contemplen los efectos de la interacción con las otras partes, tiene consecuencias menos triviales que en el caso clásico. La descripción cuántica de sólo una parte de un sistema compuesto presenta obstáculos y complicaciones conceptuales insospechadas desde el punto de vista clásico, pero que la misma teoría puede superar si se describe al sistema compuesto como un todo, es decir, si no se excluye ninguna de las partes correlacionadas producto de alguna interacción. El colapso es un ejemplo de lo dicho en el marco de la teoría de la medición: una regla espuria, estrictamente ilegítima ya que es consecuencia de ignorar una parte relevante del proceso, esto es, los aparatos de medición. No obstante, como regla instrumental para aplicar sobre el sistema de interés, el colapso es efectivo para conducir a los resultados esperados por la teoría. Si, por razones prácticas, se desea prescindir de los aparatos, entonces el postulado del colapso puede considerarse una manera efectiva y conveniente para calcular probabilidades sobre el sistema; pero no debe olvidarse que, estrictamente, el sistema completo, que abarca todas las partes en interacción, nunca abandona su evolución unitaria.

Al final del Capítulo 4 se discutió la posibilidad de considerar al colapso como una suerte de regla de proyección, pero no del vector de estado a uno de sus componentes como normalmente se lo entiende, sino como la proyección de aspectos del todo sobre una parte. Algo así como la regla del escorzo en geometría proyectiva, que permite reconstruir aspectos de un objeto a partir de su sombra sobre un plano de dimensión menor. El colapso, en este sentido simbólico, es la "sombra" sobre una parte, el sistema de interés, de una verdadera evolución unitaria que rige el sistema completo, sistema de interés más aparatos. Aspectos más formales de esta analogía están actualmente bajo estudio, y la exploración de sus alcances queda para futuros trabajos.

Algo más que podemos mencionar en relación a lo desarrollado en el Capítulo 4 es que, al tratar con valores definidos aun sin abandonar evoluciones unitarias, hemos naturalmente escindido el vínculo tradicional asumido entre autovectores y autovalores. Ésta es una característica central de las interpretaciones modales, común a todas ellas. Por consiguiente,

se abren aquí expectativas respecto de incorporar presente trabajo en el marco de alguna particular interpretación de tipo modal.

El Capítulo 5 se dedicó por completo al problema de la base preferida. En la primera sección de ese capítulo indagamos sobre la posibilidad de eliminar la ambigüedad del problema del cambio de base apelando una preferencia macroscópica que puede manifestarse, según alguna particular disposición experimental, en un conjunto particular de salidas posibles del aparato. Analizamos esta posibilidad en el famoso experimento de Stern y Gerlach porque en él se presenta, por el tipo de interacción involucrada, la particular característica de evidenciar salidas macroscópicamente distinguibles sin necesidad de ningún mecanismo de amplificación posterior, ni ningún otro sistema que funcione como entorno para que se produzca la decoherencia. Dada una particular orientación del campo magnético, sólo interviene un juego de salidas que involucran distinguibilidad macroscópica. Combinaciones lineales de estados asociados a esas salidas no corresponden a nuevas salidas legítimas del aparato, puesto que combinaciones de funciones de onda macroscópicamente localizadas en una dirección no dan como resultado funciones de onda macroscópicamente localizadas en otra dirección. Esto parece presentarse como un recurso para privilegiar una base sobre las demás: la base privilegiada sería aquella cuyas salidas son macroscópicamente distinguibles.

El argumento anterior, si bien muy sólido en el ejemplo presentado, no basta para mediciones más generales donde intervienen correlaciones entre sistemas microscópicos, como es el caso de las mediciones bit-a-bit. Es en este tipo de mediciones donde se pone en juego el formalismo de decoherencia como la respuesta al problema más aceptada y difundida en nuestros días. Debido a que las interacciones que establecen correlaciones entre sistemas microscópicos no tienen registro a nivel macroscópico, en el marco del formalismo no son consideradas como verdaderas mediciones o como mediciones completas. Desde el punto de vista del enfoque ortodoxo de la decoherencia, la mera interacción entre dos sistemas cuánticos se concibe como una "pre-medición". La verdadera medición se completa posteriormente, cuando se completa la interacción con un tercer sistema, el entorno. Considerado como un sistema de un enorme número grados de libertad, el entorno permite expresar las correlaciones establecidas en la pre-medición en términos de una

descomposición triortonormal para sistema + aparato + entorno, descomposición que, a diferencia de la biortonormal, es única. Este resultado supuestamente resuelve el problema de la base preferida, y se proclama como un logro de la decoherencia.

Nuestra estrategia ha sido totalmente opuesta en el siguiente sentido. En lugar de eliminar la ambigüedad apelando a un proceso posterior al establecimiento de la correlación entre sistema y aparato, hemos eliminado la ambigüedad apelando a un proceso anterior. Para ello, en primer lugar señalamos que el problema del cambio de base resultaría un verdadero problema sólo si existiese un aparato que fuera capaz de habilitar más de un juego de salidas, uno por cada base involucrada en el cambio de base. Sólo en ese caso las distintas correlaciones cobrarían sentido físico, estableciéndose como verdaderas mediciones diferentes con un mismo aparato. La especulación sobre la existencia de semejante aparato, que hemos llamado "aparato de cuatro luces", nos condujo a analizar cómo se comportaría al nivel de la calibración. La conclusión fue que, por la naturaleza misma de la evolución unitaria propia de la mecánica cuántica, distinguiendo sobre qué juego de salidas se obtiene un resultado definido, con tal aparato podríamos distinguir la base en la cual fue expresado el estado inicial del sistema en el caso general. Como el desarrollo en una base es una mera forma matemática, una expresión que adopta el vector de estado, poder distinguir un desarrollo de otro implicaría distinguir entre entidades lingüísticas, discriminar el nombre con el que designamos al vector de estado en las operaciones de la teoría. Tal aparato no puede existir, pues se trataría de un aparato capaz de "medir" el nombre de las entidades físicas, y esto no tiene sentido físico alguno.

La distinguibilidad entre diferentes juegos de salidas considerada aquí es producto de la diferenciación de estados más allá de que éstos puedan conducir, mediante algún otro mecanismo, la activación de salidas distinguibles mediante nuestros sentidos. La condición que define al aparato de cuatro luces es permitir más de un juego de salidas. Si hablamos de diferentes juegos, es porque pueden ser reconocidos y, por lo tanto, diferenciados en la teoría. Con este significado esencial de distinguibilidad, cualquier condición de macroscopicidad es irrelevante, y todo lo dicho vale para cualquier aparato cuántico, sea que tenga o no salidas macroscópicas.

De este modo, podemos agregar que no es por la condición de macroscopicidad que se activa un solo juego de salidas, sino que porque se activa un solo juego de salidas es posible la ulterior macroscopicidad de sus resultados. Así, si el aparato es cuántico, de modo de encontrarse regido por una evolución unitaria, uno y sólo uno de los infinitos juegos de salidas puestos en pie de igualdad por un cambio de base se puede hacer efectivo.

Lo dicho hasta aquí elimina la ambigüedad del cambio de base, aunque no elimina la indeterminación sobre cuál es la base y, por lo tanto, cuál es el observable medido. Un último argumento, también referido a la calibración, nos permitió concluir que el proceso para producir una determinada correlación es único, porque al nivel de la calibración es único el operador unitario que conecta el estado inicial con el final. Si bien con independencia de la calibración existen infinitas correlaciones con las que se puede escribir el mismo estado final de la medición, el proceso necesario para producir cada una de ellas es distinto, y viene discriminado por el operador de evolución que interviene en el proceso. Como el operador de evolución es función unívoca del observable medido, especificar el operador es especificar el observable medido. Esto contesta la pregunta sobre cuál la base y por lo tanto cuál es el observable objeto de la medición.

En algún sentido, la respuesta a qué está midiendo el observador se responde tautológicamente por la simple elección de lo que el observador desea medir, elección que implementará por medio del dispositivo experimental adecuado para ello, y que describirá en la teoría por medio de un determinado operador de evolución que le permitirá hacer los cálculos de las correlaciones.

En este punto alguien podría preguntarse si nuestra manera de resolver el problema y determinar la base preferida brinda una respuesta que coincide con la ofrecida por el formalismo de decoherencia. Pues bien, en rigor de verdad la pregunta no aplica, es impropia, pues se trata de problemas diferentes. Al abordar la medición como un proceso sobre un sistema abierto, la respuesta que brinda la decoherencia con la incorporación del entorno se refiere a la base preferida en procesos que involucran sistemas de ese tipo. Nosotros aquí presentamos y dimos respuesta al problema original, que es respecto a la elección de la base que afecta al compuesto sistema + aparato considerado como un todo aislado.

✚

Apéndice. La matemática de la mecánica cuántica

No es posible hablar de mecánica cuántica sin una comprensión mínima de su estructura matemática. En este apéndice ofrecemos un recorrido sintético sobre los conceptos matemáticos más básicos de la mecánica cuántica que, si bien abordados desde una perspectiva introductoria, brindarán los elementos mínimos necesarios para entender por completo los problemas y peculiaridades que se presentan en el desarrollo de esta tesis.

A.1 Campo de complejos

Antes de hablar de espacios vectoriales, debemos definir la noción de campo¹⁴ (MacLane y Birkhoff 1979, Cap. 3 y 8). Un campo $C = \langle C, +, \cdot, 0, 1 \rangle$ es una estructura algebraica formada por un conjunto C de elementos, entre los cuales se definen las operaciones suma (+) y producto ($\cdot$), siendo el producto distributivo respecto a la suma, y tales que para todo elemento $x \in C$ existe un elemento $-x \in C$ tal que $x + (-x) = 0$, y para todo elemento $x \in C$ distinto del elemento 0, existe $x^{-1} \in C$ tal que $x \cdot x^{-1} = 1$.

El elemento $-x$ es el inverso aditivo y el elemento x^{-1} el inverso multiplicativo de x . El 0 es el elemento neutro para la suma, y 1 es el elemento neutro para la multiplicación. Los elementos cualesquiera de un campo son llamados escalares.

Un ejemplo de campo lo forman los números reales, $\mathbb{R}$, con la suma y multiplicación habitual. Otro ejemplo es el campo de los números complejos, $\mathbb{C}$, los cuales intervienen en forma intrínseca en las descripciones de la teoría cuántica, motivo por el cual le dedicaremos las próximas líneas.

Un número complejo w puede ser representado por un par de números reales, a, b , que en su forma binómica (Rojo 1981, p. 347; Hughes 1989, p. 30) adopta la expresión

$$w = a + ib$$

¹⁴ A veces también se lo denomina *cuerpo* (Rojo 1981, p. 278)

siendo i la unidad imaginaria, definida como aquella tal que $i^2 = -1$. El número real a (que multiplica la unidad real 1) se denomina componente real del complejo; el número b (que multiplica la unidad imaginaria i) se denomina componente imaginario. Si se grafica el componente real en el eje horizontal de un sistema de ejes cartesianos, y el imaginario en el vertical, un número complejo puede interpretarse geoméricamente como un segmento orientado en el plano con origen en el centro del sistema de ejes (Hughes 1989, p. 28). Ver Figura 3.

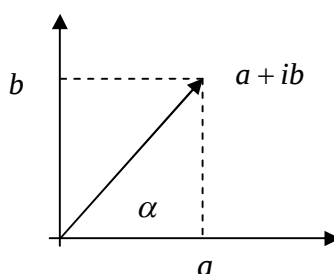


Figura 3

La suma de dos números complejos se define en forma directa sumando componente a componente. Así, si $w = a + ib$ y $z = b + id$, entonces

$$\begin{aligned} z + w &= (a + ib) + (b + id) \\ &= (a + b) + i(b + d) \end{aligned}$$

Claramente, con esta definición la suma entre complejos resulta asociativa y conmutativa. La multiplicación puede obtenerse fácilmente como resultado de la propiedad distributiva y haciendo uso de que $i^2 = -1$, de modo que

$$\begin{aligned} z \cdot w &= (a + ib) \cdot (b + id) \\ &= ab + iad + ibd + i^2 bd \\ &= (ab - bd) + i(ad + bd) \end{aligned}$$

El elemento neutro para la suma es el cero complejo, dado por $0 + i0$. El elemento neutro para la multiplicación es el número uno complejo dado por $1 + i0$.

De acuerdo con lo dicho, y con el modo en que fuera definida la suma, es claro que el inverso aditivo para el complejo $w = a + ib$ es $-w = -a - ib$, de modo que

$$w + (-w) = (a + ib) + (-a - ib) = 0 + i0$$

Para encontrar el inverso multiplicativo de un complejo, necesitamos definir previamente la operación conjugación y el concepto de módulo. Dado un número complejo $w = a + ib$, se define como el conjugado de w al número $w^* = a - ib$. Se puede demostrar fácilmente que $(w^*)^* = w$, que $(w \cdot z)^* = w^* \cdot z^*$, y que $(w + z)^* = w^* + z^*$. Por medio de la conjugación es posible encontrar una expresión para la parte real e imaginaria de un complejo, pues $w + w^* = 2\text{Real}(w)$, y $w - w^* = 2i \text{Imag}(w)$.

Si multiplicamos un complejo por su conjugado, se obtiene

$$\begin{aligned} w \cdot w^* &= (a + ib) \cdot (a - ib) \\ &= a^2 + b^2 \end{aligned}$$

Esta cantidad resulta ser siempre un número real mayor o igual a cero, siendo cero si y solo si w es cero. Esto nos permite definir el módulo de w como

$$|w| = \sqrt{w \cdot w^*} = \sqrt{a^2 + b^2}$$

Geométricamente, el módulo resulta ser igual a la longitud del segmento que representa al complejo en el plano (Hughes 1989, p. 28).

Volviendo a la multiplicación de un complejo por su conjugado, en términos de su modulo obtenemos que $w \cdot w^* = |w|^2$, de lo que se desprende que $w \cdot w^* / |w|^2 = 1$. Esto nos permite encontrar la expresión que identifica el inverso multiplicativo de cualquier complejo w (distinto de cero). En efecto

$$w^{-1} = \frac{1}{w} = \frac{w^*}{|w|^2}$$

El conjunto de los números complejos, con la suma y el producto según han sido definidos, forman un campo como fuera definido más arriba, el campo de los complejos $\mathbb{C}$.

Para terminar esta breve introducción a los números complejos, presentaremos la notación exponencial, que es muy utilizada en un sinnúmero de aplicaciones físicas, en particular en el contexto de la mecánica cuántica. Para usar la notación exponencial debemos definir primero la exponencial compleja. A fin de preservar ciertas propiedades básicas de la potenciación, se define $e^{i\alpha} = \cos(\alpha) + i\sin(\alpha)$ (Rojo 1981, p. 366). Si ahora pensamos al parámetro α como el ángulo (también llamado fase) de w respecto del eje real según la construcción geométrica de la Figura 3, obtenemos que $a = |w|\cos(\alpha)$ y que $b = |w|\sin(\alpha)$, siendo $|w|$ el módulo de w . Utilizando este resultado, vemos que cualquier número complejo de la forma $w = a + ib$ puede escribirse en notación exponencial como

$$w = |w|e^{i\alpha}$$

La notación exponencial facilita notablemente las operaciones algebraicas. Así, si tenemos w , de módulo $|w|$ y fase α , y z , de módulo $|z|$ y fase β , la multiplicación resulta ser el complejo obtenido por el producto de los módulos y la suma de las fases:

$$w \cdot z = |w|e^{i\alpha} \cdot |z|e^{i\beta} = |w||z|e^{i(\alpha+\beta)}$$

Si $z = |z|e^{i\beta}$, entonces en notación exponencial la inversa de z se obtiene simplemente como $z^{-1} = e^{-i\beta}/|z|$. Esto permite expresar la división de forma simétrica a la multiplicación, y resulta ser el complejo obtenido con la división de los módulos y la diferencia de las fases:

$$w \cdot z^{-1} = |w|e^{i\alpha} \cdot \frac{e^{-i\beta}}{|z|} = \frac{|w|}{|z|}e^{i(\alpha-\beta)}$$

A.2 Espacios y subespacios vectoriales

Ahora que disponemos de la estructura del campo de los complejos, estamos en condiciones de especificar la estructura de espacio vectorial en forma suficientemente general como para aplicar estas nociones a la mecánica cuántica.

Un espacio vectorial $V = \langle V, +, \cdot, |0\rangle \rangle$ sobre un campo $C = \langle C, +, \cdot, 0, 1 \rangle$ es una estructura algebraica formada por un conjunto V cuyos elementos indicamos con el símbolo $|x\rangle$, entre los cuales está definida una operación interna suma (+), con elemento neutro $|0\rangle$, y una operación externa ($\cdot$) consistente en un producto por los elementos del campo C (producto por un escalar), tales que, para todo $|u\rangle, |v\rangle$ y $|w\rangle$ pertenecientes a V y para todo a, b perteneciente a C , se cumple:

1. $(|u\rangle + |v\rangle) + |w\rangle = |u\rangle + (|v\rangle + |w\rangle)$
2. $|u\rangle + |v\rangle = |v\rangle + |u\rangle$
3. $|v\rangle + |0\rangle = |v\rangle$
4. $a \cdot (b \cdot |v\rangle) = (a \cdot b) \cdot |v\rangle$
5. $(a + b) \cdot |v\rangle = a \cdot |v\rangle + b \cdot |v\rangle$
6. $a \cdot (|v\rangle + |u\rangle) = a \cdot |v\rangle + a \cdot |u\rangle$
7. $0 \cdot |v\rangle = |0\rangle$
8. $1 \cdot |v\rangle = |v\rangle$

Dentro de un espacio vectorial existen subconjuntos de particular importancia para la física cuántica, que son llamados subespacios vectoriales, a los cuales introduciremos a continuación. Un subespacio S en un espacio vectorial $V = \langle V, +, \cdot, |0\rangle \rangle$ sobre el campo $C = \langle C, +, \cdot, 0, 1 \rangle$ es un subconjunto de vectores pertenecientes a V tales que

1. $|0\rangle \in S$
2. Si $|v\rangle, |u\rangle \in S \Rightarrow |v\rangle + |u\rangle \in S$
3. Si $|v\rangle \in S$ y $a \in C \Rightarrow a \cdot |v\rangle \in S$

En otras palabras, un subespacio es un conjunto de vectores que contiene el cero y que es cerrado bajo las operaciones de suma y producto por un escalar.

En todo espacio o subespacio vectorial es posible determinar un conjunto de vectores mediante los cuales expresar cualquier otro vector. Este particular conjunto es llamado base y debe cumplir los requisitos que se detallan a continuación. Una base de un espacio o subespacio vectorial sobre el campo $C = \langle C, +, \cdot, 0, 1 \rangle$ es un conjunto de vectores $\{|v_i\rangle\}$ tales que, para todo $|x\rangle$ en el espacio o subespacio, resulta

1. $|x\rangle = \sum_i^N c_i |v_i\rangle, \quad c_i \in C$
2. $\sum_i^N c_i |v_i\rangle = 0 \Leftrightarrow c_i = 0 \quad \forall i$

La primera condición significa que cualquier vector puede ser generado por combinaciones lineales de los elementos $\{|v_i\rangle\}$ de la base, es decir, por combinaciones de los vectores multiplicados por escalares. La segunda significa que los elementos de la base son linealmente independientes (Rojo 1980, p. 53). El número N es el número de elementos de la base del espacio o subespacio, y se define como la dimensión de dicho espacio o subespacio. Los coeficientes c_i se conocen como los coeficientes del vector $|x\rangle$ en la base $\{|v_i\rangle\}$.

Un ejemplo clásico de espacio vectorial es el determinado por el conjunto de puntos del plano que llamamos $\mathbb{R}^2$. Cada punto del plano puede ser especificado al dar un par de números reales, que indican sus coordenadas al especificar un sistema de ejes cartesianos. Éste es un espacio de dimensión 2, y sus elementos pueden ser representados por siguiente arreglo en forma de matriz columna

$$|v\rangle = \begin{pmatrix} v_x \\ v_y \end{pmatrix}$$

Se trata de un espacio vectorial sobre el campo de los reales, con elemento neutro dado por

$$|0\rangle = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$$

y con las operaciones de suma y producto por un escalar definidas de modo que

$$\text{Si } |v\rangle = \begin{pmatrix} v_x \\ v_y \end{pmatrix} \text{ y } |w\rangle = \begin{pmatrix} w_x \\ w_y \end{pmatrix}, \text{ entonces se tiene que } |w\rangle + |v\rangle = \begin{pmatrix} v_x + w_x \\ v_y + w_y \end{pmatrix}$$

$$\text{Si } a \in \mathbb{R} \text{ y } |v\rangle = \begin{pmatrix} v_x \\ v_y \end{pmatrix}, \text{ entonces se tiene que } a \cdot |v\rangle = \begin{pmatrix} av_x \\ av_y \end{pmatrix}$$

Los elementos del espacio $\mathbb{R}^2$ pueden ser geoméricamente representados mediante segmentos orientados que parten del origen de un sistema cartesiano y que apuntan a un punto específico¹⁵. El origen es precisamente el elemento $|0\rangle$, como puede verse en la Figura 4.

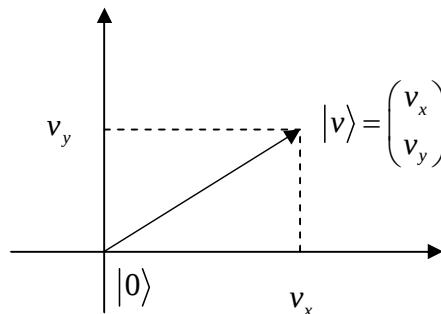


Figura 4

De forma similar es posible definir las operaciones de suma y producto por un escalar entre elementos con n componentes reales. Ésta es una generalización inmediata de $\mathbb{R}^2$, y se trata del espacio vectorial $\mathbb{R}^n$. Claramente, cuando el producto por un escalar se realiza

¹⁵ Existe una clara analogía entre un espacio vectorial de dimensión 2 con el cuerpo de los complejos, pero no la indagaremos aquí.

considerando complejos, se tiene un espacio vectorial sobre el campo de los complejos, y esto define $\mathbb{C}^n$.

A.3 Producto interno

Además de las operaciones que definen un espacio vectorial sobre un campo, es posible definir un producto entre vectores llamado producto interno (Rojo 1980, p. 214; Hughes 1989, p. 43). A veces también es llamado producto escalar, porque el producto interno entre vectores sobre un campo es siempre un escalar perteneciente al campo. En lo que sigue, asumiremos el producto interno entre vectores sobre el campo de los complejos $\mathbb{C}$.

El producto interno permite introducir la noción de métrica en un espacio, noción que lleva a consecuencias geométricas importantes como las definiciones de longitud, ortogonalidad y ángulo entre vectores. En el caso de la mecánica cuántica, permite la definición de un concepto teórico central: una medida de probabilidad.

Sea $V = \langle V, +, \cdot, |0\rangle \rangle$ un espacio vectorial sobre el campo complejo $\mathbb{C}$, el producto interno entre dos vectores cualesquiera en V , $|v\rangle$ y $|u\rangle$, que simbolizamos como $\langle v | \cdot | u \rangle \equiv \langle v | u \rangle$, cumple las siguientes condiciones

1. $\langle v | v \rangle \geq 0$ y $\langle v | v \rangle = 0$ si y solo si $|v\rangle = |0\rangle$, siendo 0 el cero complejo
2. $\langle u | v \rangle = \langle v | u \rangle^*$
3. $\langle u | (a \cdot |v\rangle) \rangle = a \langle u | v \rangle$
4. Si $|w\rangle$ también pertenece a V , $\langle u | (|v\rangle + |w\rangle) \rangle = \langle u | v \rangle + \langle u | w \rangle$

Si $|v\rangle \in V$ puede representarse en el caso especial de $\mathbb{C}^n$ como una matriz columna, entonces $\langle v |$ adquiere sentido como el mismo vector pero dispuesto en una matriz fila y con sus elementos conjugados, de modo que $\langle v | u \rangle$ se reduce a la multiplicación habitual de fila por columna matricial. En general, el conjunto de los $\{\langle v | \}$ se llama espacio dual de V (Isham 1995, p. 35). En algunos libros, según la llamada notación de Dirac (Isham 1995, p. 35), se habla de $|v\rangle$ como un “ket”, y de su dual $\langle v |$ como un “bra”, de modo que con su producto

interno se forma un “bra-ket”. Es importante notar que si a pertenece al campo de los complejos, entonces el dual de $a|v\rangle$ es $\langle v|a^*$. Con todas estas consideraciones, si tenemos

$$|v\rangle = \begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_n \end{pmatrix}, \quad |u\rangle = \begin{pmatrix} u_1 \\ u_2 \\ \vdots \\ u_n \end{pmatrix} \in \mathbb{C}^n$$

entonces el producto interno entre los vectores $|v\rangle$ y $|u\rangle$ viene dado por

$$\langle v| \cdot |u\rangle \equiv \langle v|u\rangle = (v_1^*, v_2^*, \dots, v_n^*) \cdot \begin{pmatrix} u_1 \\ u_2 \\ \vdots \\ u_n \end{pmatrix} = \sum_i^n v_i^* u_i$$

Como se puede observar en esta última fórmula, el producto interno entre vectores es siempre un escalar perteneciente al campo de los complejos. Este producto induce la norma de un vector, dada por

$$\| |v\rangle \| \equiv \langle v|v\rangle = \sum_i^n |v_i|^2$$

la cual define una noción de distancia en el espacio vectorial. La norma de un vector da una medida de su “longitud”; por consiguiente, un vector $|v\rangle$ se dice normalizado si $\| |v\rangle \| = 1$. En el caso de un espacio vectorial sobre el campo real, donde los vectores admiten la representación geométrica como segmentos orientados, la norma coincide con la longitud del segmento que lo representa. Esto es claro de la Figura 4.

El producto interno entre dos vectores permite definir un ángulo entre ellos (*Isham* 1995), que en el caso un espacio vectorial sobre el campo real se puede entender geométricamente como el ángulo entre los segmentos orientados que representa dichos vectores. Se puede demostrar (Rojo 1981, p. 302) que

$$\langle v|u\rangle = \|v\| \|u\| \cos(\theta_{vu})$$

donde θ_{vu} es el ángulo entre $|v\rangle$ y $|u\rangle$. Es clarísimo de esta última fórmula que dos vectores son perpendiculares si y solo si el producto interno da cero, pues en ese caso θ_{vu} es un ángulo recto. Cuando sucede esto se dicen que los vectores son ortogonales. Dos o más vectores se dicen ortonormales si están normalizados, y además son ortogonales.

Un espacio vectorial dotado de un producto interno es llamado también espacio de Hilbert.

Con las nociones de espacio vectorial y producto interno, podremos adentrarnos en el concepto de operador, específicamente en el de operador lineal, lo cual será presentado en la próxima subsección.

A.4 Operadores lineales

Un operador A sobre un espacio vectorial $V = \langle V, +, \cdot, |0\rangle \rangle$ es una función que, aplicada sobre un vector $|v\rangle \in V$ ¹⁶, da otro vector, esto es, $A|v\rangle = |u\rangle \in V$. Dicho de otra manera, un operador es una aplicación que mapea vectores en vectores (*Isham* 1995, p. 45; *Hughes* 1989, p. 42). Un operador lineal es una clase especial de operador que cumple las siguientes propiedades

1. $A(|v\rangle + |w\rangle) = A|v\rangle + A|w\rangle$
2. $A(a|v\rangle) = a(A|v\rangle)$

para todo $|v\rangle, |w\rangle \in V$. Es posible definir la suma y producto entre operadores por medio de su aplicación sobre vectores. Así, dados dos operadores A y B actuando sobre un espacio vectorial $V = \langle V, +, \cdot, |0\rangle \rangle$, se define la suma y el producto de modo tal que, para todo $|v\rangle \in V$

1. $(A + B)|v\rangle = A|v\rangle + B|v\rangle$

¹⁶ Si bien se trata de una función sobre un vector, es convencional notar la aplicación del operador sobre el vector de forma $A|v\rangle$ y no con la notación típicamente asociada a funciones $A(|v\rangle)$.

$$2. (AB)|v\rangle = A(B|v\rangle)$$

De este modo resulta que, si A y B son lineales, los operadores $A+B$ y AB también son lineales.

Así como cualquier vector perteneciente a $\mathbb{C}^n$ puede representarse con un arreglo en forma de columna con n números complejos, cualquier operador actuando sobre el espacio vectorial $\mathbb{C}^n$ puede representarse como un arreglo en forma de una matriz de $n \times n$ de números complejos.

Como ejemplo de aplicación concreta, volvamos al espacio vectorial real de dimensión 2 que hemos llamado $\mathbb{R}^2$. En ese espacio los operadores quedarán representados por matrices de 2×2 de números reales. Así, si A es un operador sobre $\mathbb{R}^2$, tenemos

$$A = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$$

La aplicación del operador sobre un elemento del espacio vectorial está definida según la multiplicación de matrices (Hughes 1989, p.17), de modo que si

$$|v\rangle = \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} \in \mathbb{R}^2$$

entonces

$$A|v\rangle = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = \begin{pmatrix} a_{11}v_1 + a_{12}v_2 \\ a_{21}v_1 + a_{22}v_2 \end{pmatrix} = \begin{pmatrix} w_1 \\ w_2 \end{pmatrix} = |w\rangle$$

Esta multiplicación matricial puede, a su vez, ser puesta en términos de los componentes del operador A y los componentes del vector $|v\rangle$. Si a_{ij} es el elemento ubicado en la columna

i fila j en el arreglo matricial que representa a A ¹⁷, y v_j es el componente j de $|v\rangle$, entonces el componente w_i del vector $|w\rangle$ es dado por

$$w_i = \sum_j a_{ij} v_j$$

Un concepto de importancia crucial respecto de los operadores lineales utilizados en mecánica cuántica es el de hermiticidad. Para definirlo, introducimos primero el concepto de hermítico conjugado de un operador A , que simbolizamos con $A^\dagger$, como el operador cuya matriz resulta de intercambiar fila por columna de la matriz de A y conjugar cada uno de sus elementos. Un operador A se dice hermítico si $A = A^\dagger$. La hermiticidad de un operador es importante, porque las magnitudes físicas en mecánica cuántica, llamadas observables, serán representadas sólo por operadores hermíticos.

Otro concepto crucial asociado a un operador es el de sus autovalores y autovectores. Para un dado operador A , si existe un vector $|a_i\rangle$ tal que

$$A|a_i\rangle = a_i|a_i\rangle$$

siendo a_i un escalar, se dice que $|a_i\rangle$ es un autovector de A , y a_i es su correspondiente autovalor. Se puede probar que todo operador hermítico tiene sus autovalores reales, y que si dos autovalores son distintos, entonces sus correspondientes autovectores son ortogonales (Hughes 1989, p. 49). Así, si A es un operador hermítico que actúa sobre un espacio de dimensión N , y tiene N autovalores distintos, sus correspondientes autovectores forman una base ortogonal del espacio. Se habla entonces de la autobase de A .

Si el operador A que actúa sobre un espacio de dimensión N tiene N autovalores distintos, cada autovalor corresponde a un único autovector que genera un subespacio de dimensión uno. Si el operador en cuestión tiene M autovalores distintos, con $M < N$, cada autovalor corresponde no a un único autovector, sino a un conjunto de autovectores, y así el subespacio asociado a ese autovalor es de dimensión mayor que uno. No obstante, si el

¹⁷ El elemento a_{ij} se llama *elemento de matriz* del operador A

operador es hermítico, siempre es posible encontrar una base ortonormal del espacio formada con autovectores (Hughes 1989, p. 51). Cuando un operador tiene más de un autovector para cada autovalor, se dice que es degenerado.

Existe una clase especial de operadores lineales cuyo papel es importantísimo en las descripciones de las propiedades cuánticas. Son los llamados proyectores. En términos técnicos, un operador Π es un proyector si es hermítico y además es idempotente¹⁸.

La importancia de los proyectores es que están asociados a subespacios de un espacio vectorial. Se puede demostrar que, para cada subespacio S en un espacio vectorial, existe un proyector Π_S tal que, para todo $|v\rangle$ en el espacio, $\Pi_S |v\rangle = |v_S\rangle$, siendo $|v_S\rangle$ el componente del vector $|v\rangle$ sobre el subespacio S (Hughes 1989, p. 47). Se dice que $|v_S\rangle$ es el vector proyección de $|v\rangle$ sobre el subespacio S y que Π_S proyecta sobre S .

En el caso de un subespacio de dimensión uno es fácil calcular cuál es el vector $|v_S\rangle$ proyectado. Esto puede verse en la Figura 5. Si el vector $|s\rangle$ es base de S y tiene norma igual a uno, haciendo uso de las funciones trigonométricas elementales, tenemos que la proyección de cualquier vector $|v\rangle$ sobre S está dada por $|v_S\rangle = \|v\| \cos(\theta_{sv}) |s\rangle$, donde θ_{sv} es el ángulo entre el vector $|v\rangle$ y la recta asociada al subespacio S .

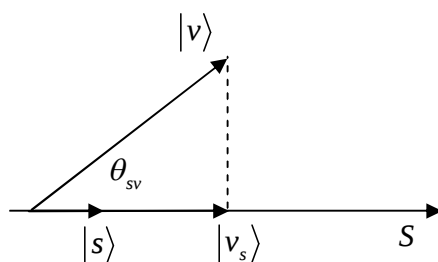


Figura 5

Los proyectores pueden ser representados por medio de la notación de Dirac en forma muy conveniente para las operaciones vectoriales. Lo dicho es fácil de ver si consideramos un

¹⁸ Un operador es idempotente si su cuadrado es igual a sí mismo, es decir, si $\Pi^2 = \Pi$.

subespacio S de dimensión uno. Si el vector $|s\rangle$ es de norma uno y genera el subespacio S , entonces el proyector sobre S puede ser simbólicamente representado por la expresión $\Pi_s = |s\rangle\langle s|$ si entendemos que para cualquier vector $|v\rangle \in V$, dicha expresión opera sobre dicho vector de la siguiente manera

$$\Pi_s |v\rangle = (|s\rangle\langle s|) |v\rangle = |s\rangle (\langle s|v\rangle)$$

La expresión es correcta, pues de la definición del producto interno entre $|s\rangle$ y $|v\rangle$, y de la ortonormalidad de $|s\rangle$, tenemos que $\langle s|v\rangle = \|v\| \cos(\theta_{sv})$, como ya hemos visto en la construcción de la Figura 5. Así la ecuación de aquí arriba cumple

$$\Pi_s |v\rangle = |s\rangle \|v\| \cos(\theta_{sv}) = |v_s\rangle$$

como se requiere de Π_s . En forma más general, si el proyector Π_s proyecta sobre un subespacio S generado por los vectores $\{|s_i\rangle\}$, entonces se tiene que

$$\Pi_s = \sum_i |s_i\rangle\langle s_i|$$

Un resultado importante muy usado en mecánica cuántica es el llamado teorema de la descomposición espectral (Hughes 1989, p. 50). Este teorema asegura que, si A es un operador hermítico que opera sobre un espacio dimensión finita, con M autovalores distintos $\{a_i\}$, $i = 1 \dots M$, entonces se lo puede descomponer en términos de los proyectores $\Pi_{a_i} = |a_i\rangle\langle a_i|$ asociados a los subespacios generados por los correspondientes autovectores $\{|a_i\rangle\}$, de la siguiente y única manera

$$A = \sum_i^M a_i |a_i\rangle\langle a_i|$$

Esta expresión para operadores es muy útil en los manejos algebraicos.

Para terminar la sección definimos una operación muy importante que nos permitirá una representación general para los estados en la mecánica cuántica. Nos referimos a la traza de

un operador (Hughes 1989, Cap. 5). Para operadores con espectro discreto, la traza se define como la suma de los elementos de la diagonal de la matriz representativa. Resulta que la suma de los elementos de la diagonal es igual a la suma de los autovalores de A . Así, si a_{ij} son los elementos de matriz del operador A , y a_j son sus autovalores, entonces la traza es dada por

$$Tr[A] = \sum_i a_{ii} = \sum_j a_j$$

Se puede demostrar, además, que para cualquier base $\{|v_i\rangle\}$ de vectores ortonormales del espacio de Hilbert sobre el cual aplica A , la traza viene dada también por

$$Tr[A] = \sum_i \langle v_i | A | v_i \rangle$$

La aplicación de la traza es imprescindible para obtener una generalización de la noción de estado que requiere del uso de operadores. Asociado a ello son relevantes las siguientes propiedades básicas de la traza: para cualesquiera operadores A y B , y siendo a un escalar cualquiera, se cumple que

$$Tr[AB] = Tr[BA]$$

$$Tr[A + B] = Tr[A] + Tr[B]$$

$$Tr[aA] = aTr[A]$$

A.5 Producto tensorial

Teniendo distintos espacios de Hilbert, es posible construir un espacio de Hilbert compuesto por medio del producto tensorial entre los primeros. Valiéndonos del producto tensorial es posible caracterizar un sistema cuántico compuesto de otros subsistemas que forman parte de él.

Si tenemos un espacio de Hilbert H_A de dimensión N_A con una base dada por $\{|v_i^A\rangle\}$, $i=1\cdots N_A$, y otro espacio de Hilbert H_B de dimensión N_B con una base $\{|v_j^B\rangle\}$, $j=1\cdots N_B$, se define el producto tensorial entre H_A y H_B como el espacio de Hilbert compuesto $H = H_A \otimes H_B$ de dimensión $N_A N_B$ que contiene al conjunto de vectores formados por la combinación

$$|v\rangle = \sum_{ij} c_{ij} |v_i^A\rangle \otimes |v_j^B\rangle$$

donde $|v_i^A\rangle \otimes |v_j^B\rangle$ viene dado por la concatenación $|v_i^A\rangle |v_j^B\rangle$ (Hughes 1989, p. 148; Ballentine 1998, p. 160). Por definición, el conjunto de los $N_A N_B$ vectores dados por $\{|v_i^A\rangle |v_j^B\rangle\}$ forman una base del producto tensorial $H_A \otimes H_B$.

Es interesante mencionar que, si bien cualquier vector $|v\rangle \in H_A \otimes H_B$ puede ser escrito como combinación lineal de $|v_i^A\rangle \in H_A$ y $|v_j^B\rangle \in H_B$, no todo $|v\rangle \in H_A \otimes H_B$ puede ser escrito como una única concatenación de la forma $|v^A\rangle |v^B\rangle$ con $|v^A\rangle \in H_A$ y $|v^B\rangle \in H_B$ (Hughes 1989, p. 149). Esto significa que el producto tensorial $H_A \otimes H_B$ no es simplemente el producto cartesiano entre H_A y H_B , sino que lo contiene.

Un resultado importante respecto a la descomposición de vectores en un espacio compuesto $H_A \otimes H_B$ es la llamada descomposición de Schmidt o descomposición biortonormal (Mittelstaedt 1998, p. 26), la cual asegura que todo $|v\rangle \in H_A \otimes H_B$ puede descomponerse con una suma de índice único del siguiente modo

$$|v\rangle = \sum_i c_i |v_i^A\rangle \otimes |v_i^B\rangle$$

Teniendo caracterizado el tratamiento de vectores en un espacio compuesto por productos tensoriales, podemos definir cómo se aplican los operadores sobre él. Si A es un operador sobre H_A y B es un operador sobre H_B , es posible construir el operador $A \otimes B$ sobre $H_A \otimes H_B$ haciendo uso del hecho que cualquier vector en $H_A \otimes H_B$ puede ser escrito como combinación de la concatenación entre vectores de la base de H_A (sobre los que aplica A) y vectores de la base de H_B (sobre los que aplica B), de modo que

$$(A \otimes B) |v_i^A\rangle |v_j^B\rangle = A |v_i^A\rangle \otimes B |v_j^B\rangle$$

Es importante notar que, si bien hemos construido $A \otimes B$ que actúa sobre $H_A \otimes H_B$, no todo operador O sobre $H_A \otimes H_B$ es de esa forma.

Es posible definir el producto interno entre dos vectores en $H_A \otimes H_B$ al notar que

$$\left(\langle v_i^A | \otimes \langle v_j^B | \right) \cdot \left(|v_{i'}^A\rangle \otimes |v_{j'}^B\rangle \right) = \langle v_i^A | \cdot |v_{i'}^A\rangle \langle v_j^B | \cdot |v_{j'}^B\rangle$$

Finalmente definimos la traza parcial de un operador O sobre el espacio compuesto $H_A \otimes H_B$ al sumar sobre los vectores de la base de uno de los espacios H_A o H_B

$$Tr_A[O] = \sum_i \langle v_i^A | O | v_i^A \rangle$$

$$Tr_B[O] = \sum_j \langle v_j^B | O | v_j^B \rangle$$

Una traza parcial no es un escalar. En el primer caso, al trazar sobre H_A nos da un operador que actúa sobre H_B . En el segundo, al trazar sobre H_B nos da un operador que actúa sobre H_A .

Bibliografía

- Adler, S. (2003), "Why decoherence has not solved the measurement problem: a response to P. W. Anderson", *Studies in History and Philosophy of Modern Physics*, **34**: 135-142.
- Andrew, E. y Bub, J. (1994), "Triorthogonal uniqueness theorem and its relevance to the interpretation of quantum mechanics", *Physical Review A (Atomic, Molecular, and Optical Physics)*, **49**: 4213-4216.
- Auletta, G. (2004), "Decoherence and triorthogonal decomposition", *International Journal of Theoretical Physics*, **43**: 2263-2274.
- Ballentine, L. (1990), "Limitations of the projection postulate", *Foundations of Physics*, **20**: 1329-1343.
- Ballentine, L. (1998), *Quantum Mechanics. A Modern Development*, Singapore: World Scientific.
- Bub, J. (1997), *Interpreting the Quantum World*, Cambridge: Cambridge University Press.
- Bacciagaluppi, G. y Hemmo, M. (1994), "Making sense of approximate decoherence", *Proceedings of the Philosophy of Science Association*, **1**: 345-354.
- Calzetta, E. A., Hu, B. L. y Mazzitelli, F. D. (2001), "Coarse-grained effective action and renormalization group theory in semiclassical gravity and cosmology", *Physics Reports*, **352**: 459-520.
- Castagnino, M., Lombardi, O. y Vanni, L. (2008), "Hacia una interpretación ontológicamente pluralista de la mecánica cuántica", en R. de Andrade Martins, C. C. Silva y L. A. P. Martins (eds.), *Filosofia e História da Ciência no Cone Sul. Seleção de Trabalhos do 5º Encontro*, Campinas: AFHIC, 321-330.
- Cohen, D. W. (1989), *An Introduction to Hilbert Space and Quantum Logic*, New York: Springer-Verlag.
- d'Espagnat, B. (2000), "A note on measurement", *Physics Letters A*, **282**: 133-137.

- Daneri, A., Loinger, A. y Prosperi, G. M. (1962), "Quantum theory of measurement and ergodicity conditions", *Nuclear Physics*, **33**: 297-319 (También incluido en Wheeler y Zurek 1983).
- Daneri, A., Loinger, A. y Prosperi, G. M. (1966), "Further remarks on the relations between statistical mechanics and quantum theory of measurement", *Nuovo Cimento*, **44B**: 119-128.
- Davies, P. C. W. y Brown, J. (eds) (1986), *The Ghost in the Atom*, Cambridge: Cambridge University Press.
- Dieks, D. (1994), "The modal interpretation of quantum mechanics, measurement and macroscopic behavior", *Physical Review D*, **49**: 2290-2300.
- Dieks, D. (2007a), "Probability in modal interpretations of quantum mechanics", *Studies in History and Philosophy of Modern Physics*, **38**: 292-310.
- Dieks, D. (2007b), "Probability in non-collapse interpretations of quantum mechanics", en Th. M. Nieuwenhuizen, V. Spicka, B. Mehmani, M. Jafar-Aghdami y A. Yu Khrennikov (eds.), *Beyond the Quantum*, Singapore: *World Scientific*, 345-355.
- Earman, J. (1986), *A Primer on Determinism*. Dordrecht: Reidel Publishing Company.
- Fortin, S. (2012), "Una interpretación de la decoherencia y el límite clásico desde la perspectiva de los sistemas cerrados", Tesis Doctoral, Universidad Nacional Tres de Febrero.
- Ghirardi, G. (2011), "Collapse theories", en E. N. Zalta (ed.), *The Stanford Encyclopedia of Philosophy*, (Winter 2011 Edition), URL = <<http://plato.stanford.edu/archives/win2011/entries/qm-collapse/>>
- Griffiths, R. (1984), "Consistent histories and the interpretation of quantum mechanics", *Journal of Statistical Physics*, **36**: 219-272.
- Griffiths, R. y Omnes, R. (1999), "Consistent histories and quantum measurements", *Physics Today*, **52**: 26-31.
- Griffiths, R. (2002), "Probabilities and quantum reality: Are there correlata?", *Foundations of Physics*, **33**: 1423-1459.

- Healey, R. A. (1984), "How many worlds?", *Noûs*, **18**: 591-616.
- Hughes, R. I. G. (1989), *The Structure and Interpretation of Quantum Mechanics*, Cambridge MA: Harvard University Press.
- Isham, C. J. (1995), *Lectures on Quantum Theory: Mathematical and Structural Foundations*, London: Imperial College Press.
- Jammer, M. (1974), *The Philosophy of Quantum Mechanics*, New York: John Wiley & Sons.
- Jauch, J. M., Wigner, E. P. y Yanase, M. M. (1967), "Some comments concerning measurements in quantum mechanics", *Nuovo Cimento*, **48B**: 144-151.
- Joos, E. (2000), "Elements of environmental decoherence", en P. Blanchard, D. Giulini, E. Joos, C. Kiefer y I.-O. Stamatescu (eds.), *Decoherence: Theoretical, Experimental, and Conceptual Problems, Lecture Notes in Physics 538*, Heidelberg-Berlin: Springer, pp. 1-17.
- Laura, R. y Vanni, L. (2008a), "Probabilidades condicionales y colapso en las mediciones cuánticas", *Filosofia e Historia da Ciência no Cone Sul (AFHIC), Seleccion de Trabalhos do 5o Encontro, Florianópolis: AFHIC*, 383-391.
- Laura, R. y Vanni, L. (2008b), "Conditional probabilities and collapse in quantum measurements", *International Journal of Theoretical Physics*, **47**: 2382-2392.
- Laura, R. y Vanni, L. (2009), "Time translation of quantum properties", *Foundations of Physics*, **39**: 160-173.
- Lombardi, O. y Castagnino, M. (2008), "A modal-Hamiltonian interpretation of quantum mechanics", *Studies in History and Philosophy of Modern Physics*, **39**: 380-443.
- Lombardi, O. y Vanni, L. (2010), "Medición cuántica y decoherencia: ¿qué medimos cuando medimos?", *Scientiae Studia*, **8**: 273-291.
- Ludwig, G. (1954), *Die Grundlagen der Quantenmechanik*, Berlin-Göttingen-Heidelberg: Springer.
- MacLane, S. y Birkhoff, G. (1979), *Algebra*, New York: Macmillan.

- Mittelstaedt, P. (1998), *The Interpretation of Quantum Mechanics and The Measurement Process*, Cambridge: Cambridge University Press.
- Paz, J. P. y Zurek, W. H. (2002), "Environment-induced decoherence and the transition from quantum to classical", en D. Heiss (ed.), *Fundamentals of Quantum Information: Quantum Computation, Communication, Decoherence and All That*, Heidelberg-Berlin: Springer, 77-148.
- Penrose, R. (1989), "The Emperor's New Mind". Oxford: Oxford University Press.
- Peres, A. (1974), "Quantum measurements are reversible", *American Journal of Physics*, **42**: 886-891.
- Peres, A. (1980), "Can we undo quantum measurements?", *Physical Review D*, **22**: 879-883.
- Peres, A. y Zurek, W. H. (1982), "Is quantum theory universally valid?", *Journal of Physics*, **50**: 807-810.
- Pessoa Jr., O. (1998), "Can the decoherence approach help to solve the measurement problem?", *Synthese*, **113**: 323-346.
- Rojo, A. O. (1980), *Algebra II*, Buenos Aires: El Ateneo.
- Rojo, A. O. (1981), *Algebra I*, Buenos Aires: El Ateneo.
- Sakurai, J. J. (1985), *Modern Quantum Mechanics*, New York: Addison-Wesley Publishing.
- Sanford, D. H. (2011), "Determinates and Determinables", en E. N. Zalta (ed.), *The Stanford Encyclopedia of Philosophy*, (Spring 2011 Edition), URL = <<http://plato.stanford.edu/archives/spr2011/entries/determinate-determinables/>>.
- Schlosshauer, M. (2004), "Decoherence, the measurement problem, and interpretations of quantum mechanics", *Reviews of Modern Physics*, **76**: 1267-1305.
- Sonego, S. (1993), "Conceptual foundations of quantum theory: a map of the land", *Université Libre de Bruxelles* (Preprint).

- Stern, O. y Gerlach, W. (1922), "Der experimentelle Nachweis des magnetischen Moments des Silberatoms", *Zeitschrift für Physik*, **8**: 110.
- van Fraassen, B. C. (1972), "A formal approach to the philosophy of science", en R. Colodny (ed.), *Paradigms and Paradoxes: The Philosophical Challenge of the Quantum Domain*, Pittsburgh: University of Pittsburgh Press, 303-366.
- van Fraassen, B. C. (1973), "Semantic analysis of quantum logic", en C. A. Hooker (ed.), *Contemporary Research in the Foundations and Philosophy of Quantum Theory*, Dordrecht: Reidel, 80-113.
- van Fraassen, B. C. (1974), "The Einstein-Podolsky-Rosen paradox", *Synthese*, **29**: 291-309.
- van Kampen, N. (1954), "Quantum statistics of irreversible processes", *Physica*, **20**: 603-622.
- Vanni, L. y Laura, R. (2005), "Mediciones cuánticas sin decoherencia", *Anales AFA*, **17**: 28-31.
- von Neumann, J. (1955), *Mathematical Foundations of Quantum Mechanics*, Princeton: Princeton University Press.
- Wheeler, J. A. y Zurek, W. H. (eds) (1983), *Quantum Theory and Measurement*, Princeton: Princeton University Press.
- Wilce, A. (2006), "Quantum Logic and Probability Theory", en E. N. Zalta (ed.), *The Stanford Encyclopedia of Philosophy*, (2006 Edition), URL = <http://plato.stanford.edu/entries/qt-quantlog/>.
- Zurek, W. H. (1981), "Pointer basis of quantum apparatus: Into what mixture does the wave packet collapse?", *Physical Review D*, **24**: 1516-1525.
- Zurek, W. H. (1982), "Environment-induced superselection rules", *Physical Review D*, **26**: 1862-1880.
- Zurek, W. H. (1991), "Decoherence and the transition from quantum to classical", *Physics Today*, **44**: 36-44.

- Zurek, W. H. (1998), "Decoherence, einselection, and the existential interpretation," *Philosophical Transactions of the Royal Society*, **A356**: 1793-1821 (Los Alamos National Laboratory, quant-ph/9805065).
- Zurek, W. H. (2003), "Decoherence, einselection, and the quantum origins of the classical", *Reviews of Modern Physics*, **75**: 715-775.
- Zeh, D. (1970), "On the interpretation of measurement in quantum theory", *Foundations of Physics*, **1**: 69-76.

Para citar este documento

Vanni, Leonardo (2015). Los problemas de la medición cuántica sin decoherencia (Tesis de posgrado). Universidad Nacional de Quilmes, Bernal, Argentina: Repositorio Institucional Digital de Acceso Abierto. Disponible en: <http://ridaa.demo.unq.edu.ar>

