

Transición cuántico clásica, la mecánica cuántica en el espacio de fases, la decoherencia y el problema de la medición

Juan Pablo Paz

9 de mayo de 2025

Índice

1. La transición cuántico-clásica y la decoherencia en el espacio de fases	2
2. La mecánica cuántica en el espacio de fases	4
2.1. Definición de la función de Wigner	4
2.2. Propiedades de $W(\alpha)$	6
3. Representación de Gatos de Schrödinger en el espacio de fases. La interferencia cuántica y la negatividad de la función de Wigner	8
4. Evolución de la función de Wigner	10
4.1. El impacto de la decoherencia sobre un gato de Schrödinger . . .	13
5. Medición de la función de Wigner	16
6. El problema de la medición	19
7. El proceso de medición según Von Neumann	21
8. Los problemas de la descripción de Von Neumann	24
9. Decoherencia y la aparición de una base privilegiada en los estados del aparato	26
10. El proceso de decoherencia y la frontera entre el mundo clásico y el cuántico	27
11. ¿Cuál es la base de estados punteros?	30

1. La transición cuántico-clásica y la decoherencia en el espacio de fases

Como mencionamos a lo largo del texto, la mecánica cuántica tiene aspectos anti intuitivos que impiden su adaptación a nuestro sentido común. Uno de ellos es que, de acuerdo a los principios que hemos discutido, existen estados cuánticos que jamás son observados para sistemas macroscópicos, siendo que en principio nada impediría lo contrario. El ejemplo más famoso de estos estados fue presentado en 1935 por Erwin Schrödinger apelando a un ser vivo, un gato (que de ahí en más pasó a denominarse como “el gato de Schrödinger”): mediante un dispositivo conceptualmente simple, la mecánica cuántica permitiría preparar estados que son superposición de dos alternativas macroscópicamente distinguibles (en uno de ellos el gato está vivo y en el otro está muerto). Es decir, la mecánica cuántica admite estados en los que el gato está vivo y muerto a la vez. Más rigurosamente, está en una superposición de los dos estados cuánticos asociados a esas dos alternativas. En cambio, la física clásica no admite este tipo de estados sino que, por el contrario, solamente podría convivir con situaciones en las cuales preparamos un conjunto de muchos objetos de modo tal que cada uno de ellos se encuentra en uno de dos estados posibles. En el caso analizado por Schrödinger, este sería un conjunto de gatos que están vivos o muertos. El conjunto está caracterizado por las probabilidades que podemos asociar a estos estados. Es decir, con cierta probabilidad, en este caso, encontraremos al gato vivo y con otra probabilidad estará muerto. Pero cada uno de los gatos estará en alguno de estos estados, que son mutuamente excluyentes. Esa es, entonces, la diferencia crucial entre el caso clásico y el cuántico: en el primero cada gato está vivo o muerto mientras que en el segundo, cada uno de ellos está en un estado donde las dos alternativas son exploradas simultáneamente.

La preparación de estados del tipo “gato de Schrödinger” para objetos realmente macroscópicos es, aún hoy, un desafío para la ciencia y las tecnologías cuánticas modernas. Sin embargo, en las últimas décadas se avanzó mucho en este sentido. Uno de los primeros logros se produjo en 1997 en el laboratorio de Serge Haroche, en el que (como se ha discutido en el capítulo destinado a la electrodinámica cuántica en cavidades -o cavity QED-) se logró preparar estados “tipo gatos” consistentes en superposiciones de dos estados coherentes centrados, cada uno de ellos, alrededor de un punto diferente del espacio de cuadraturas del campo electromagnético (que es totalmente análogo al espacio de fases de un oscilador mecánico). Si bien estos estados no son realmente macroscópicos, ya que contienen un número medio de fotones relativamente bajo -de alrededor de $n = 7$ -, estos experimentos abrieron la puerta a una generación de realizaciones que, usando diversas técnicas similares, han permitido crear gatos cada vez más grandes (entre las tecnologías utilizadas se destaca, por ejemplo, la que se denomina “electrodinámica cuántica de circuitos” que permite manipular dispositivos superconductores acoplados con resonadores de radio frecuencia que actúan de manera muy similar a las cavidades espejadas usadas en los experimentos del grupo de Haroche). No repetiremos aquí la descripción detallada del

procedimiento por el cual es posible preparar en la cavidad un estado del tipo

$$|\Psi\rangle_{\text{cat}} = N(|\alpha_1\rangle + |\alpha_2\rangle), \quad (1)$$

donde $|\alpha_1\rangle$ y $|\alpha_2\rangle$ son dos estados coherentes y N es una constante de normalización (donde $N = \sqrt{2 + 2\text{Re}(\langle\alpha_1|\alpha_2\rangle)}$). Resumidamente, el proceso involucra la preparación de un estado coherente en la cavidad para luego inyectar un átomo preparado en una superposición de dos estados (habitualmente denominados $|e\rangle$ y $|g\rangle$) con los que el campo interactúa de manera no resonante (o dispersiva) introduciéndose un desfase entre ellos que es proporcional al número de fotones. Lo crucial aquí es que, gracias a la interacción dispersiva, un estado coherente inicial evoluciona en otro estado coherente que depende del estado interno del átomo. Al inyectar un átomo en una superposición de estos dos estados, el átomo y el campo se entrelazan, lo cual resulta un ingrediente crucial en todo el proceso. En esta situación, luego de medir el estado del átomo a la salida de la cavidad (haciéndolo atravesar previamente una zona de Ramsey donde se induce una rotación en el espacio interno del átomo) la cavidad queda preparada en un estado gato, en el que la separación entre los estados coherentes es proporcional al tiempo de interacción y al desfase que la interacción dispersiva genera para cada fotón. En este punto, se recomienda al lector el repaso de los principios en los que se basa este procedimiento.

En definitiva, después de un proceso de preparación, el estado de la cavidad es del tipo de los que describimos más arriba. En su experimento Haroche puso a prueba la física de un proceso que naturalmente induce la transición entre el mencionado estado (donde el gato está vivo y muerto a la vez) y otro en el cual el felino está vivo o muerto con cierta probabilidad (que en este caso es la misma para las dos alternativas). El proceso se denomina “decoherencia” y la virtud del experimento de Haroche fue que constituyó la primera observación de la dinámica del mencionado proceso en tiempo real: el gato pasa de una situación que solo se puede explicar con la mecánica cuántica a otra en la que la explicación clásica es la que se impone. La relevancia del proceso de decoherencia en el contexto del estudio de la frontera entre el régimen cuántico y el régimen clásico de un sistema físico había sido destacada desde principios de los años 1980, por científicos como Dieter Zeh, Wojciech Zurek y muchos otros, pero recién pudo ser observada (como proceso dinámico) en el experimento de Haroche. El origen físico de la decoherencia es muy simple: este proceso aparece inevitablemente siempre que un sistema cuántico interactúe -aunque sea muy débilmente- con su entorno. El entorno actúa, en efecto, como un objeto que registra el estado cuántico del sistema e induce el colapso de un estado puro a otro mixto. En esta sección describiremos la dinámica del proceso de decoherencia para un sistema simple, un oscilador armónico (que, como mencionamos, es completamente análogo al sistema estudiado por Haroche: un modo del campo electromagnético almacenado en una cavidad superconductora). Describiremos entonces la evolución de un oscilador armónico amortiguado preparado inicialmente en un estado tipo “gato de Schrödinger” y usaremos una ecuación maestra (con la forma de Lindblad) para modelar el proceso disipativo inducido por la absorción de fotones por los espejos de la cavidad. Nos enfocaremos en el estudio

de la desaparición de los efectos de “interferencia cuántica” que son aquellos que nos obligan a aceptar que el gato está vivo y muerto a la vez observando cómo rápidamente el estado se transforma en otro que puede describirse como una mezcla estadística de gatos que están vivos o muertos.

Utilizaremos una metodología muy poderosa para analizar la evolución del estado cuántico del oscilador: representaremos al operador densidad ρ mediante una función $W(\alpha)$ que está definida en el espacio de fases. Este método fue originalmente presentado por Eugene Wigner y, por ese motivo, la función $W(\alpha)$ se denomina, precisamente, la función de Wigner. En lo que sigue, describiremos en detalle el método de la función de Wigner y luego analizaremos la evolución de “gatos de Schrödinger” utilizándola intensivamente.

2. La mecánica cuántica en el espacio de fases

El espacio de las fases es el escenario natural para la física clásica. En ese caso los estados son descriptos mediante distribuciones de probabilidad: es decir, funciones positivas que integradas sobre todo el espacio son iguales a la unidad. Wigner demostró que la física cuántica puede representarse también en el espacio de fases asociando a cada estado una función $W(\alpha)$ que toma valores reales (utilizaremos la notación $\alpha = (x, p)$ para identificar a un punto en el espacio de las fases). Esta función, como veremos, tiene muchas de las propiedades de las distribuciones de probabilidad pero, a diferencia de ellas, puede ser negativa. La negatividad de la función de Wigner es lo que distingue a los estados cuánticos de los clásicos y por eso nos permite identificar cuál es la huella del carácter cuántico de un estado. Podemos afirmar que la función de Wigner es una distribución de quasi-probabilidad reflejando con ese nombre su carácter no positivo. Pese a eso, tiene una estrecha relación con las distribuciones de probabilidad predichas por la mecánica cuántica: cuando se integra la función de Wigner a lo largo de cualquier recta en el espacio de fases se obtiene un número real y positivo que es idéntico a la probabilidad de observar un cierto valor en la medición de la cuadratura asociada a la mencionada recta (seremos más precisos al respecto más abajo). Veremos ahora cómo definir esta función y cuáles son sus propiedades más importantes.

2.1. Definición de la función de Wigner

Mostraremos aquí dos formas equivalentes de definir la función de Wigner. En primer lugar lo haremos siguiendo la definición original presentada por el propio Wigner y luego mediante un método equivalente pero mucho más poderoso. La definición original propuesta por Wigner para el caso de una partícula en una sola dimensión (la generalización es inmediata) es la siguiente:

$$W(x, p) = \int du \langle x - u/2 | \rho | x + u/2 \rangle \exp(iup) / 2\pi. \quad (2)$$

De esta definición surge de manera inmediata que esta función es real (ya que al tomar su compleja conjugada se obtiene la misma función, lo que se demuestra

cambiando la variable de integración de u a $-u$ y usando el carácter hermitico del estado ρ). También surge inmediatamente que la integral de $W(x, p)$ sobre todo el espacio de fases es igual a la unidad (lo que se demuestra integrando primero en la variable de momento, usando la representación de la delta de Dirac como integral y recordando que la traza de ρ es igual a la unidad). Por último, también surge que la integral sobre el momento no es otra cosa más que la densidad de probabilidad $\rho(x, x)$ (lo que se demuestra simplemente usando que la integral en momento da lugar a la aparición de la delta de Dirac $\delta(u)$). Asimismo, la integral en posición es idéntica a $\rho(p, p)$ (lo que se demuestra insertando la representación de la identidad en la base de autoestados del momento y realizando manipulaciones simples de la ecuación resultante). Menos evidente es la propiedad de que $W(x, p)$ provee una descripción completa del estado ρ (es decir que conociendo $W(x, p)$ en todo el espacio de fases es posible reconstruir el estado ρ) y que el producto interno entre estados puede calcularse a partir de la integral del producto de las correspondientes funciones de Wigner. Pero, como dijimos más arriba, hay otra forma equivalente de definir la función de Wigner que no solo permite demostrar con relativa facilidad sus propiedades más importantes sino que también es útil para calcular estas funciones (en particular, es muy útil para los gatos de Schrödinger que analizaremos más abajo).

En efecto, la función de Wigner puede definirse también como el valor medio de un operador hermitico:

$$W(\alpha) = \text{Tr}(\rho A(\alpha)), \quad (3)$$

donde $A(\alpha)$, que suele denominarse como “operador de punto” en el espacio de fases, se define como

$$A(\alpha) = \frac{D(\alpha) R D^\dagger(\alpha)}{\pi}. \quad (4)$$

En esta expresión, el operador de desplazamiento en el espacio de las fases es

$$D(\alpha) = \exp(\alpha a^\dagger - \alpha^* a), \quad (5)$$

y el operador de reflexión R está definido en la base de posición (o de momento) como aquel que satisface $R|x\rangle = |-x\rangle$ (o $R|p\rangle = |-p\rangle$). Resulta ilustrativo demostrar explícitamente la equivalencia entre las dos definiciones de la función de Wigner que dimos más arriba. El valor medio del operador de punto es

$$W(\alpha) = \text{Tr}(\rho D(\alpha) R D^\dagger(\alpha) / \pi) = \int dq \langle q | \rho D(\alpha) R D^\dagger(\alpha) | q \rangle / \pi. \quad (6)$$

Factorizando los operadores de desplazamiento como el producto de una traslación en momento seguida por otra en posición y usando la forma del operador de reflexión se obtiene

$$W(x, p) = \int dq \langle q | \rho | 2q - x \rangle \exp(2ip(x - q)/\hbar) / \pi. \quad (7)$$

Cambiando la variable de integración a $u = 2(x - q)$, que implica que $q = x - u/2$, se obtiene

$$W(x, p) = \int du \langle x - u/2 | \rho | x + u/2 \rangle \exp(ipu/\hbar) / 2\pi, \quad (8)$$

que es la expresión original de Wigner que presentamos más arriba.

2.2. Propiedades de $W(\alpha)$

En lo que sigue, demostraremos algunas de las propiedades más importantes de $W(\alpha)$, que son heredadas de propiedades de los operadores de punto $A(\alpha)$.

Propiedad (i) $W(\alpha)$ es una función real, cuya integral sobre todo el espacio de las fases es igual a la unidad y que provee una descripción completa del estado cuántico ρ . Esto surge del hecho de que los operadores de punto $A(\alpha)$ son hermíticos y forman una base ortonormal y completa del espacio conjunto de operadores. En efecto, se verifica (queda para el lector comprobar estas identidades) que:

$$A^\dagger(\alpha) = A(\alpha), \quad \int d\alpha A(\alpha) = I, \quad \text{y} \quad \text{Tr}(A(\alpha)A(\beta)) = \delta(\alpha - \beta)/2\pi. \quad (9)$$

De estas propiedades de $A(\alpha)$ surge el carácter real de $W(\alpha)$, el hecho de que su integral sobre todo el espacio es igual a la unidad y que el estado ρ (así como también cualquier otro operador) puede descomponerse como combinación lineal de operadores de punto de la siguiente manera:

$$\rho = 2\pi \int d\alpha W(\alpha) A(\alpha). \quad (10)$$

Esto implica que la función de Wigner $W(\alpha)$ no es otra cosa más que la proyección del estado ρ sobre cada operador de punto. También esto implica que conociendo la función de Wigner, el estado ρ queda completamente determinado. Cabe mencionar que, tal como lo adelantamos más arriba, cualquier operador O puede descomponerse como combinación lineal de operadores de punto. Los coeficientes de esta descomposición (que definen la representación de Weyl-Wigner del operador) son

$$O_W(\alpha) = \text{Tr}(OA(\alpha)), \quad (11)$$

y entonces

$$O = 2\pi \int d\alpha O_W(\alpha) A(\alpha). \quad (12)$$

Propiedad (ii) La integral de $W(\alpha)$ sobre cualquier recta en el espacio de fases es la densidad de probabilidad de obtener un dado resultado en la medición de la cuadratura asociada a dicha recta.

Para demostrar esto es útil notar que la traza de cualquier producto de operadores puede obtenerse como la integral de las respectivas representaciones de Weyl-Wigner. En efecto, la ortonormalidad de los operadores de punto implica que

$$\text{Tr}(\rho_1 \rho_2) = 2\pi \int d\alpha W_1(\alpha) W_2(\alpha), \quad (13)$$

y en general que

$$\text{Tr}(O_1 O_2) = 2\pi \int d\alpha O_{1W}(\alpha) O_{2W}(\alpha). \quad (14)$$

Naturalmente, estos resultados pueden usarse para calcular la probabilidad de obtener un cierto resultado en una dada medición ya que la misma es siempre idéntica al valor medio del proyector asociado a ese resultado. Por lo tanto, las probabilidades siempre pueden obtenerse como integrales de la función de Wigner multiplicadas por la representación de Weyl-Wigner del proyector correspondiente.

Consideremos el operador Q definido como $Q = aX - bP$. Como el espectro del operador Q es continuo, los autovalores son números reales tanto positivos como negativos. Denotaremos como P_Q al proyector sobre el autoestado del operador Q con autovalor c . Por lo visto anteriormente, si calculamos la representación de Weyl-Wigner de este proyector podremos calcular la probabilidad de obtener cualquier resultado en la medición de Q (como el espectro de Q es continuo obtendremos, más rigurosamente, una densidad de probabilidad). Este cálculo es relativamente sencillo y el resultado es notablemente simple: si denotamos como $P_{QW}(x, p)$ a la representación de Weyl-Wigner del mencionado proyector, demostraremos que

$$P_{QW}(x, p) = \delta(ax - bp - c), \quad (15)$$

con lo cual, la integral de la función de Wigner a lo largo de la línea definida por la ecuación $ax - bp = c$ será la densidad de probabilidad de obtener el resultado c a partir de la medición del observable $Q = ax - bp$.

Para demostrar esto usaremos que el proyector P_Q puede escribirse siempre en función de los operadores posición (X) y momento (P) como:

$$P_Q = \delta(aX - bP - c). \quad (16)$$

Usando esta expresión (y la representación de la delta de Dirac como exponencial) es simple calcular la representación de Weyl-Wigner de este proyector. En efecto,

$$P_{QW}(x, p) = \int du \langle x - u/2 | \delta(aX - bP - c) | x + u/2 \rangle \exp(iup/\hbar)/2\pi. \quad (17)$$

En esta expresión resulta conveniente escribir la delta de Dirac como la integral de una función exponencial, que a su vez puede descomponerse en el producto de dos operadores como:

$$\delta(aX - bP - c) = \int d\lambda \exp(-ib\lambda P) \exp(i\lambda aX) \exp(iab\lambda^2/2)/2\pi. \quad (18)$$

Reemplazando esta expresión en la identidad que define a $P_{QW}(x, p)$ y operando con los desplazamientos en posición y momento obtenemos

$$P_{QW}(x, p) = \int d\lambda \int du \exp(i\lambda^2 ab/2) \exp(iup + i\lambda a(x + u/2)) \delta(u + \lambda b)/(2\pi)^2. \quad (19)$$

Ahora resulta sencillo realizar la integral sobre u , debido a la presencia de la delta de Dirac y obtener (evaluando la exponencial en $u = -\lambda b$)

$$P_{QW}(x, p) = \int d\lambda \exp(i\lambda^2 ab/2) \exp(-i\lambda bp) \exp(i\lambda ax) \exp(-i\lambda^2 ab/2) \exp(-i\lambda c) / (2\pi)^2. \quad (20)$$

Finalmente, cancelando las dos contribuciones que contienen a λ^2 y usando la representación de la delta de Dirac como integral, obtenemos

$$P_{QW}(x, p) = \delta(ax - bp - c). \quad (21)$$

Este es el resultado que buscábamos, del que se deduce que es posible calcular todas las distribuciones de probabilidad asociadas a la medición de cualquier cuadratura como una integral de la función de Wigner sobre una cierta recta.

3. Representación de Gatos de Schrödinger en el espacio de fases. La interferencia cuántica y la negatividad de la función de Wigner

Consideremos un estado tipo “gato de Schrödinger” construido como superposición de dos estados coherentes:

$$|\Psi\rangle_{\text{cat}} = N_{\text{cat}}(|\beta_1\rangle + |\beta_2\rangle), \quad (22)$$

donde $N_{\text{cat}} = \frac{1}{\sqrt{2+2\text{Re}(\langle\beta_1|\beta_2\rangle)}}$ es una constante que asegura la normalización del estado y $\beta_{1,2}$ son dos números complejos que caracterizan a los estados coherentes que participan de la superposición. La matriz densidad de este estado es, obviamente,

$$\rho_{\text{cat}} = |\Psi_{\text{cat}}\rangle\langle\Psi_{\text{cat}}| = N_{\text{cat}}^2 (|\beta_1\rangle\langle\beta_1| + |\beta_2\rangle\langle\beta_2| + |\beta_1\rangle\langle\beta_2| + |\beta_2\rangle\langle\beta_1|). \quad (23)$$

Calcularemos ahora la función de Wigner de este estado. Para eso debemos obtener expresiones para elementos de matriz del operador de punto $A(\alpha)$ entre dos estados coherentes arbitrarios (y usarlas para calcular las funciones de Wigner asociadas a cada uno de los cuatro términos que aparecen en la expresión de arriba). Sin pérdida de generalidad podemos enfocar el cálculo del término

$$W_{ij}(\alpha) = \langle\beta_i|A(\alpha)|\beta_j\rangle, \quad (24)$$

para luego aplicar el resultado eligiendo convenientemente los números β_i y β_j . Para realizar el cálculo conviene recordar que

$$|\beta_j\rangle = D(\beta_j)|0\rangle \quad (25)$$

y que para componer dos operadores de desplazamiento podemos utilizar la expresión

$$D(\gamma_1)D(\gamma_2) = D(\gamma_1 + \gamma_2) \exp(\gamma_1 \wedge \gamma_2/2), \quad (26)$$

donde $\gamma_1 \wedge \gamma_2 = \gamma_1 \gamma_2^* - \gamma_2 \gamma_1^* = \frac{2i(x_1 p_2 - x_2 p_1)}{\hbar}$, operación que tiene un aspecto muy similar a la que se realiza para calcular el producto vectorial de dos vectores en dimensión 2. En efecto, la fase que aparece en la expresión anterior, que involucra la composición de dos traslaciones en el espacio de fases, puede identificarse con el área del paralelogramo formado por los vectores γ_1 y γ_2 .

Usando estos resultados podemos obtener la siguiente expresión simple para $W_{ij}(\alpha)$:

$$W_{ij}(\alpha) = \exp(-2|\alpha - (\beta_i + \beta_j)/2|) \exp(\alpha \wedge (\beta_i - \beta_j)) \exp(\beta_i \wedge \beta_j/2). \quad (27)$$

Esta expresión es el ingrediente necesario para obtener finalmente la función de Wigner de un estado tipo “gato de Schrödinger” compuesto por la superposición de dos estados coherentes $|\beta_i\rangle$ y $|\beta_j\rangle$. Consideraremos solo el estado formado por la suma de ambos y dejaremos como ejercicio el cálculo de la función de Wigner de aquellos estados del tipo $|\Psi_\phi\rangle = N_\phi(|\beta_1\rangle + \exp(i\phi)|\beta_2\rangle)$, donde ϕ es una fase arbitraria. En efecto, cuando $\phi = 0$ la función de Wigner es

$$W(\alpha) = \left(\frac{N^2}{\pi} \right) \left(\exp(-|\alpha - \beta_1|^2) + \exp(-2|\alpha - \beta_2|^2) \right. \\ \left. + 2 \exp(-2|\alpha - (\beta_1 + \beta_2)/2|^2) \cos(2\text{Im}(\alpha \wedge (\beta_1 - \beta_2))) \right). \quad (28)$$

Es decir, la función de Wigner tiene dos picos Gaussianos, uno centrado en β_1 y el otro centrado en β_2 . Pero además contiene un término de interferencia que está modulado por otro pico Gaussiano centrado esta vez en el punto medio entre los dos anteriores. Es decir, en el punto $(\beta_1 + \beta_2)/2$. Este término de interferencia contiene un factor oscilante que puede reescribirse en términos de las coordenadas posición y momento ya que

$$\cos(\text{Im}(\alpha \wedge \Delta\beta)) = \cos\left(\frac{2x_\alpha \Delta p_\beta}{\hbar} - \frac{2p_\alpha \Delta x_\beta}{\hbar}\right), \quad (29)$$

donde para simplificar la notación definimos $\Delta\beta = \beta_1 - \beta_2$.

Examinando esta expresión es simple demostrar que las oscilaciones siempre son paralelas al vector $\beta_1 - \beta_2$, es decir que la función de Wigner oscila a lo largo del eje perpendicular al vector que une los dos picos Gaussianos. Podemos ver esto en un ejemplo sencillo. Para evitar complicaciones algebraicas innecesarias, analizaremos el caso $\beta_1 = \beta = -\beta_2$. Es decir, supondremos que los estados coherentes están centrados en puntos opuestos en el espacio de las fases. Obviamente, las consideraciones anteriores nos permiten afirmar que el caso general es completamente análogo a este, trasladando el origen del espacio de fases al punto medio entre los dos estados. Entonces, para el estado $|\Psi_\beta\rangle = N_\beta(|\beta\rangle + |-\beta\rangle)$, la función de Wigner es

$$W(\alpha) = \left(\frac{N_\beta^2}{\pi} \right) \left(\exp(-2|\alpha - \beta|^2) + \exp(-2|\alpha + \beta|^2) \right. \\ \left. + 2 \exp(-2|\alpha|^2) \cos\left(\frac{2x_\alpha p_\beta - 2p_\alpha x_\beta}{\hbar}\right) \right). \quad (30)$$

Es evidente que la longitud de onda de las oscilaciones disminuye cuando p_β y x_β aumentan (son inversamente proporcional a estas magnitudes). La funcionalidad de Wigner en el origen toma un valor que es el doble que el alcanzado en cada uno de los picos Gaussianos y esto refleja la interferencia constructiva entre los estados (es fácil comprobar que un gato “impar”, construido como la resta de los estados anteriores) tendrá una funcionalidad de Wigner negativa en el origen. En cualquier caso, la negatividad de la función de Wigner para este tipo de estados es evidente y puede observarse con claridad en la figura que se incluye abajo.

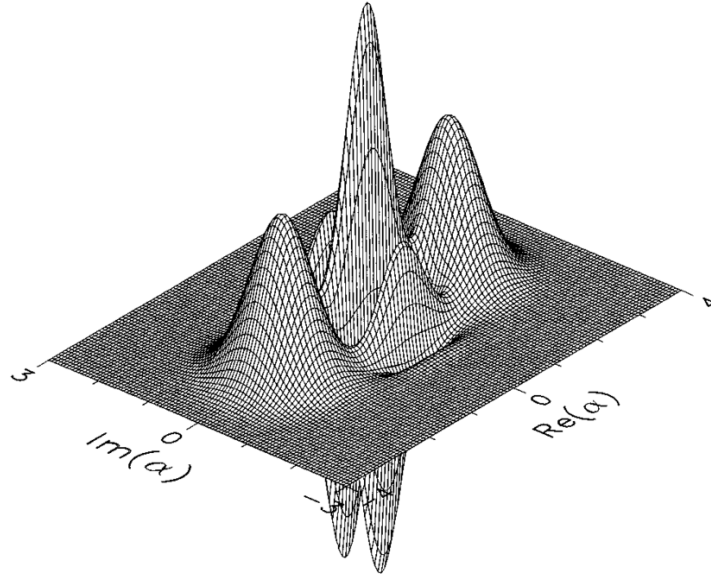


Figura 1: Función de Wigner de un estado tipo gato de Schrödinger que consiste en dos estados coherentes ubicados a lo largo del eje real del espacio de fases. Entre las campanas gaussianas de los estados coherentes individuales se pueden ver las franjas emergentes de la interferencia cuántica.

4. Evolución de la función de Wigner

Para estudiar la dinámica del proceso de decoherencia inducido por la interacción entre el sistema y su entorno analizaremos la evolución de la función de Wigner. Supondremos que la evolución de la matriz densidad está determinada por una ecuación maestra con la forma de Linblad y derivaremos a partir de ella una ecuación para la evolución de $W(\alpha)$. Comencemos por analizar el caso en el que el sistema evoluciona de manera unitaria y la matriz densidad obedece

una ecuación de la forma

$$\dot{\rho} = \frac{1}{i\hbar} [H, \rho].$$

Supondremos que el Hamiltoniano es de la forma $H = \frac{p^2}{2m} + V(x)$, pero luego nos concentraremos solamente en el caso en el que el potencial depende cuadráticamente de la coordenada del sistema (el oscilador armónico). Sin embargo, es interesante examinar la evolución de la función de Wigner en este caso, más general, en primer lugar. La ecuación para la función $W(\alpha)$ se obtiene a partir de la ecuación maestra simplemente multiplicando por el operador de punto $A(\alpha)$ y tomando la traza. De ese modo obtenemos

$$\dot{W}(\alpha) = \frac{1}{i\hbar} \text{Tr}([H, \rho]A(\alpha)) = \dot{W}_K(\alpha) + \dot{W}_V(\alpha),$$

donde $\dot{W}_K$ y $\dot{W}_V$ contienen las contribuciones al lado derecho de la ecuación que se originan, respectivamente, en el término cinético y el término del potencial en el Hamiltoniano.

La contribución proveniente de la energía cinética puede calcularse fácilmente. En efecto,

$$\dot{W}_K(\alpha) = \int du \exp\left(\frac{ipu}{\hbar}\right) \langle x - \frac{u}{2} | (P^2 \rho - \rho P^2) | x + \frac{u}{2} \rangle \frac{1}{4\pi i \hbar m}.$$

Usando la forma explícita de la acción del operador momento en la representación posición, se puede obtener

$$\dot{W}_K(\alpha) = \hbar^2 \partial_x \int du \exp\left(\frac{ipu}{\hbar}\right) \partial_u \langle x - \frac{u}{2} | \rho | x + \frac{u}{2} \rangle \frac{1}{2\pi i \hbar m}.$$

Integrando por partes se obtiene una expresión sencilla para este término que resulta ser

$$\dot{W}_K(\alpha) = -\frac{p}{m} \partial_x W(\alpha).$$

Por otra parte, la contribución de la energía potencial puede obtenerse también de manera sencilla ya que

$$\dot{W}_V(\alpha) = \int du \exp\left(\frac{ipu}{\hbar}\right) \langle x - \frac{u}{2} | (V(x)\rho - \rho V(x)) | x + \frac{u}{2} \rangle \frac{1}{2\pi i \hbar}.$$

Desarrollando el potencial en serie alrededor del punto x resulta evidente que solamente contribuyen las potencias impares de la variable u , y que cada una de ellas puede expresarse en términos de derivadas respecto del momento p . De este modo se obtiene

$$\dot{W}_V(\alpha) = \sum_{k \geq 0} \frac{-1)^k \hbar^{2k} \partial_x^{2k+1} V(x)}{(2k+1)! 2^{2k}} \partial_p^{2k+1} W(\alpha).$$

Es interesante notar que el término con $k = 0$ en la expresión anterior, combinado con la contribución asociada a la energía cinética, no es otra cosa más

que el corchete de Poisson entre el Hamiltoniano y $W(\alpha)$, que da cuenta de la dinámica clásica asociada al flujo Hamiltoniano. Este término simplemente hace que la función de Wigner se propague siguiendo las trayectorias clásicas que pasan por cada punto del espacio de las fases. Los términos con potencias más altas son todas contribuciones de naturaleza cuántica ya que vienen acompañadas con potencias crecientes de $\hbar$. Es decir, la evolución de la función de Wigner, generada por la evolución unitaria de la mecánica cuántica es tal que

$$\dot{W}(\alpha) = \{H, W(\alpha)\}_{PB} + \sum_{k>0} (-1)^k \frac{\hbar^{2k} \partial_x^{2k+1} V(x) \partial_p^{2k+1} W(x, p)}{(2k+1)! 2^{2k}},$$

donde el símbolo $\{H, W\}_{PB}$ denota al corchete de Poisson entre H y la función de Wigner. El término de la derecha en la ecuación anterior suele denominarse también “corchete de Moyal” e incluye todas las potencias de $\hbar^2$ (el número de términos en esta suma está limitado, obviamente, por el número de derivadas no nulas que tiene el potencial $V(x)$). Es claro que solamente en el caso donde el potencial depende cuadráticamente de la coordenada la evolución de $W(\alpha)$ será idéntica a la clásica. En otros casos, las derivadas superiores de $W(\alpha)$ respecto del momento jugarán un rol crucial y serán responsables, en general, de generar zonas donde la función de Wigner se vuelva negativa aun en el caso de que la misma sea inicialmente positiva en todos los puntos del espacio de las fases. De hecho, es posible demostrar que los únicos estados con funciones de Wigner no negativa son los estados Gaussianos. Como cualquier potencial no cuadrático genera estados no Gaussianos a partir de estados que inicialmente tienen esa característica, cualquier potencial no cuadrático dará lugar a negatividades en la función de Wigner.

En lo que sigue nos restringiremos a considerar la dinámica de un oscilador armónico y, por consiguiente, el término unitario de la evolución será muy sencillo: para el oscilador armónico, la función de Wigner evoluciona (en el caso unitario) siguiendo a un flujo Hamiltoniano. Por ejemplo, si preparamos un estado tipo gato de Schrödinger como el que se muestra en la Figura xx (con dos picos Gaussianos y oscilaciones en la zona ubicada entre esos dos picos), entonces la evolución temporal será simplemente la asociada a una rotación alrededor del origen, con una frecuencia de rotación que es independiente de la distancia alrededor mencionado origen (que es, precisamente, lo que caracteriza al movimiento oscilatorio armónico, en el cual la frecuencia es independiente de la amplitud del movimiento). Por eso, la evolución unitaria no afecta en nada a las franjas de interferencia que caracterizan a la función de Wigner de un estado “gato”. No desaparecen ni se crean franjas, simplemente las que existen rotan y se reubican periódicamente. En particular, el valor de la función de Wigner en el origen ($W(\alpha = 0)$) se mantiene constante (en el caso de la Figura, donde los estados coherentes que participan de la superposición están ubicados de manera simétrica respecto al origen). Esto puede verse de manera muy simple: es la causa de que el corchete de Poisson, que tiene un término proporcional a p y otro proporcional a x , se anula en el origen y, por lo tanto, $\dot{W}(0) = 0$.

4.1. El impacto de la decoherencia sobre un gato de Schrödinger

Estudiemos ahora qué sucede cuando el oscilador armónico interactúa con un entorno y enfoquemos nuestro análisis a discutir en particular qué es lo que sucede con las franjas de interferencia que caracterizan al estado inicial. Modelaremos la interacción entre el oscilador y su entorno utilizando una ecuación para la evolución de la función de Wigner que se obtiene de una de las ecuaciones maestras que ya analizamos anteriormente. En efecto, utilizaremos la ecuación que incluye dos operadores de Lindblad. El primero será $L_1 = \sqrt{\gamma(1+n)}a$ y el segundo será $L_2 = \sqrt{\gamma n}a^\dagger$. Esta ecuación, como vimos en el capítulo anterior, permite describir la aproximación al equilibrio del oscilador: el estado asintótico es el de equilibrio, donde el valor medio del operador número es

$$\langle a^\dagger a \rangle = n \left(\frac{\hbar\omega}{k_B T} \right),$$

con $n(x) = \frac{1}{\exp(x)-1}$ es la distribución de Planck. La frecuencia γ es la que define el tiempo de relajación ya que, como vimos, todos los valores medios relevantes en el problema alcanzan sus valores de equilibrio para tiempos mayores que el tiempo de relajación t_{rel} , donde $t_{\text{rel}} = \frac{1}{\gamma}$. Como veremos, la interacción con el entorno hace que las franjas de interferencia que caracterizan a la función de Wigner decaigan también, pero en una escala de tiempos mucho menor que t_{rel} .

Para estudiar el impacto de la parte no unitaria de la evolución de la función de Wigner conviene reescribir la ecuación maestra para el estado ρ en el caso que consideramos. Como vimos, en términos de los operadores posición (X) y momento (P) esta ecuación resulta ser

$$\dot{\rho} = \frac{1}{i\hbar} [H, \rho] + \frac{\gamma}{i\hbar} ([x, \{p, \rho\}] - [p, \{x, \rho\}]) + \gamma(1+2n) \left(\frac{m\omega}{\hbar} \right) \left([x, [x, \rho]] + \frac{[p, [p, \rho]]}{m^2\omega^2} \right).$$

A partir de esta ecuación es posible deducir la que gobierna la evolución de la función de Wigner. Para eso, procedemos del mismo modo en que lo hicimos para estudiar el régimen unitario: multiplicamos la ecuación maestra por el operador de punto $A(\alpha)$ y tomamos la traza a ambos lados de la ecuación resultante. De ese modo, el lado izquierdo será simplemente la derivada temporal de la función de Wigner. En la sección anterior examinamos cómo describir el efecto del término unitario en la ecuación de evolución de ρ . Ese término estará nuevamente presente en la ecuación para $W(\alpha)$ y será idéntico al corchete de Poisson entre el Hamiltoniano y $W(\alpha)$ (Recordemos que, como estamos analizando la evolución de un oscilador armónico, la ecuación inducida por la parte unitaria es idéntica a la clásica). Dejamos como ejercicio para el lector la deducción del efecto de los términos no unitarios sobre $W(\alpha)$ y presentamos aquí el resultado final al que podemos arribar después de una manipulación relativamente simple de las integrales que aparecen en la definición de $W(\alpha)$. En efecto, la ecuación de evolución de la función de Wigner resulta ser

$$\dot{W}(\alpha) = \{H, W(\alpha)\}_{PB} + \gamma (\partial_x(xW(\alpha)) + \partial_p(pW(\alpha))) + \gamma(1+2n)\hbar m\omega \left(\partial_p^2 W(\alpha) + \frac{\partial_x^2 W(\alpha)}{m^2\omega^2} \right).$$

La forma de esta ecuación de evolución es realmente simple: en primer lugar aparece el corchete de Poisson, cuyo origen ya analizamos y que se debe al primer término de la ecuación maestra. En segundo lugar aparece un término proporcional a γ que es totalmente simétrico ante intercambios de posición y momento. Su efecto es netamente disipativo y está asociado al amortiguamiento en la evolución de los valores medios (tanto de la posición, como del momento y de la energía). Por cierto, la evolución de los valores medios se calcula a partir de la ecuación de evolución de $W(\alpha)$ de manera muy similar a como se lo hace a partir de la ecuación maestra para ρ : si multiplicamos la ecuación de evolución de $W(\alpha)$ por la posición x e integramos sobre todo el espacio de las fases obtenemos del lado izquierdo la derivada del valor medio de la posición. En ese caso es simple demostrar que, además del corchete de Poisson, la única contribución no nula proviene del factor proporcional a γ y que contiene la derivada respecto a la posición (y lo mismo sucede con la evolución del valor medio del momento, la única contribución no nula se origina en el factor que contiene la derivada primera respecto a p). En tercer lugar aparece un término difusivo con derivadas segundas respecto al momento y a la posición, nuevamente el término difusivo es simétrico frente a intercambios de momento y posiciones, tal como es natural. El término difusivo contiene un prefactor numérico que depende tanto de la tasa de relajación γ como de la temperatura del entorno (que entra en escena a través de la función $n(\hbar\omega/(k_B T))$). La constante de Planck solo aparece en el término difusivo. Vale la pena notar que existe un régimen interesante en el cual la constante $\hbar$ desaparece completamente de la ecuación para $W(\alpha)$: es el régimen de altas temperaturas, en el cual suponemos que $\hbar\omega/k_B T$ es mucho menor que la unidad. En esa circunstancia es posible aproximar el factor $1 + 2n$ como

$$1 + 2n(\hbar\omega/k_B T) = \coth(\hbar\omega/(2k_B T)) \approx \frac{2k_B T}{\hbar\omega} + \dots,$$

donde despreciamos todos los términos que decaen para temperaturas mucho mayores que $\hbar\omega/k_B$. En ese régimen la ecuación maestra es

$$\dot{W} = \{H, W\}_{PB} + 2\gamma (\partial_x(xW) + \partial_p(pW)) + \gamma m k_B T \left(\partial_p^2 W + \frac{\partial_x^2 W}{m^2 \omega^2} \right).$$

Esta es la ecuación de Fokker-Planck, utilizada para describir la evolución de una función de distribución clásica en presencia de disipación y ruido térmico. Podemos decir que es una ecuación dinámica totalmente clásica, pero no debemos perder de vista que la hemos obtenido como caso límite de la evolución de un sistema cuántico. Fuera del régimen de altas temperaturas, la ecuación para $W(\alpha)$ tiene también el aspecto de una ecuación de Fokker-Planck con un coeficiente difusivo diferente, que no es directamente proporcional a la temperatura y que depende de la constante de Planck.

Analicemos ahora el impacto de esta ecuación en la evolución de las franjas de interferencia. Es bastante evidente que la escala de tiempos para la cual el impacto del término disipativo (aquel en el cual solamente aparece γ) y aquella asociada al término difusivo, pueden ser muy diferentes y que dependerán

fuertemente de la naturaleza de la función de distribución $W(\alpha)$. En efecto, veremos que el término difusivo es responsable del decaimiento de las franjas de interferencia en una escala de tiempos que es típicamente mucho menor que la fijada por la tasa de relajación γ .

Podríamos resolver la ecuación de evolución para $W(\alpha)$ (tanto de manera analítica como numérica) pero la presentación de esos resultados requeriría un detalle innecesario para nuestro propósito. Bastará con estimar el tiempo necesario para que la función de Wigner en el origen (que denotamos simplemente como $W(0)$) decaiga. Recordemos que inicialmente su valor es el doble del alcanzado en los picos ubicados en cada uno de los estados coherentes que intervienen en la superposición. Con ese fin, podemos evaluar el efecto del término difusivo sobre las franjas de interferencia y notar que las derivadas segundas son tales que

$$\left(\partial_p^2 + \frac{\partial_x^2}{(m\omega)^2}\right) \cos\left(\frac{2xp_\beta}{\hbar} - \frac{2px_\beta}{\hbar}\right) = 4\left(\frac{x_\beta^2}{\hbar^2} + \frac{p_\beta^2}{(m\omega\hbar)^2}\right) \cos\left(\frac{xp_\beta}{\hbar} + \frac{px_\beta}{\hbar}\right).$$

Entonces el efecto del término difusivo en la evolución de la función de Wigner en el origen será

$$\dot{W}(0) = -\gamma(1 + 2n)4|\beta|^2 W(0),$$

donde hemos despreciado el efecto del término disipativo (lo que tiene sentido en el régimen en el que γ es muy pequeño pero el producto de $\gamma(1 + 2n)|\beta|^2$ es finito). En este caso, la función de Wigner en el origen decae exponencialmente rápido con una tasa, que denominaremos “tasa de decoherencia” Γ_{dec} , que es

$$\Gamma_{\text{dec}} = \gamma(1 + 2n)4|\beta|^2.$$

La dependencia con la temperatura es tal que en el régimen de temperaturas altas Γ_{dec} es proporcional a $k_B T$. Por otra parte, la tasa de decaimiento de la interferencia está determinada por la separación entre los paquetes que intervienen en la superposición. Esta dependencia viene dada por el factor $|\beta|^2$. Es decir, dado que la energía media asociada al estado coherente $|\beta\rangle$ es proporcional a $|\beta|^2$ y esta energía es el número medio de fotones contenidos en el estado N_β , podríamos reescribir la expresión anterior de modo tal que

$$\Gamma_{\text{dec}} = \gamma(1 + 2n)4N_\beta,$$

lo que implica decir que el tiempo de decoherencia es aproximadamente igual al tiempo necesario para que un único fotón sea absorbido por el entorno (este tiempo es $\frac{1}{\gamma N_\beta}$, ya que $\frac{1}{\gamma}$ es el tiempo necesario para perder todos los fotones). Es decir, la escala de tiempos involucrada en la decoherencia es notablemente más pequeña que la involucrada en la pérdida de energía. Este hecho permite considerar la posibilidad de responsabilizar a la decoherencia inducida por la interacción con el entorno del comportamiento de todos los sistemas que se comportan clásicamente, aun aquellos que conservan su energía. Esto a primera vista parece sorprendente ya que atribuir el comportamiento clásico a un proceso

asociado a la disipación acarrea el riesgo de asociar dicho comportamiento con un régimen donde el comportamiento pone de manifiesto cierta irreversibilidad. La enorme separación de las escalas de tiempo de relajación y de decoherencia, que hemos puesto de manifiesto aquí (en un ejemplo muy sencillo pero no trivial) da esperanzas para que la decoherencia permita predecir la existencia de un régimen cuántico donde la reversibilidad se mantenga.

Este modelo, basado en la ecuación de Linblad que utilizamos, describe bastante bien los resultados de los experimentos realizados por el grupo de Haroche en los que el proceso de decoherencia fue observado instante a instante. En efecto, en el link LINK! Puede verse la película que muestra la evolución de la función de Wigner medida en el laboratorio como función del tiempo. Naturalmente, la mayor virtud de esta película es que no se obtiene a partir de una simulación numérica sino que es un resultado experimental, lo cual constituye una verdadera hazaña científica y tecnológica. Animaciones similares, generadas numéricamente, se obtuvieron un tiempo antes que las observadas en el laboratorio y pueden verse en la Figura yyy. Allí se observa la evolución de la función de Wigner para dos estados iniciales: uno en el cual los estados coherentes están separados en posición (y $p_\beta = 0$) y otro en el cual la separación es en momentos (y $x_\beta = 0$). La ecuación maestra que usamos aquí predice que ambos estados deberían decaer del mismo modo. Sin embargo, en la simulación que se muestra, esto no sucede. El motivo es simple e interesante a la vez: la simulación no corresponde a la solución de una ecuación de Linblad sino que se corresponde con la solución de la evolución de un oscilador acoplado con un baño de infinitos osciladores a través de una interacción que involucra la posición del sistema (y no su momento). Este sistema, usualmente denominado como el “movimiento Browniano cuántico” puede resolverse exactamente (un esbozo de la solución se presenta en el capítulo anterior) y tiene la virtud de distinguir, para tiempos suficientemente cortos, entre estados separados en posición y en momento. En cambio, la ecuación de Linblad que usamos, trata a la posición y al momento de manera totalmente simétrica (lo cual no es una buena aproximación en muchas situaciones).

En la película experimental de Haroche se muestra la evolución de la función de Wigner en cada punto del espacio de fases y en función del tiempo. Cabe preguntarse cómo es que dicha función puede medirse. En la próxima sección analizaremos esta pregunta.

5. Medición de la función de Wigner

La función de Wigner, como vimos, es el valor medio de un operador hermitico (un observable) y por lo tanto puede medirse experimentalmente. La pregunta que responderemos aquí es: ¿cómo hacerlo? (Como vimos la definición de $W(\alpha)$ es:

$$W(\alpha) = \text{Tr}(D(\alpha)RD^\dagger(\alpha)\rho),$$

es decir, $W(\alpha)$ es el valor medio de un operador que no solamente es hermitico sino que también es unitario. Para medir este valor medio se puede apelar a

un algoritmo cuántico que es notablemente simple pero que es conceptualmente muy rico. Lo denominamos habitualmente: el algoritmo de scattering y está ilustrado en la Figura zzz. Involucra a un sistema que inicialmente está preparado en un estado ρ (que en el caso que analizamos será el oscilador armónico, pero que puede asociarse a cualquier otro sistema físico, con un espacio de estados de dimensión arbitraria). Haremos interactuar a este sistema con un qubit, que actúa como un objeto “de prueba”, que interactúa con el sistema y luego es sometido a una observación. En este algoritmo no realizamos ninguna medición sobre el sistema cuyo estado inicial es ρ sino que solamente observamos el estado del qubit de prueba. El algoritmo involucra la preparación de muchos sistemas en el mismo estado ρ y la repetición del proceso de interacción, tras el cual se mide (tal como describiremos más abajo) el valor medio de alguna magnitud asociada al qubit de prueba. De esa medición extraeremos información no solamente sobre ρ sino sobre el operador involucrado en la interacción.

En el algoritmo se prepara inicialmente el qubit en un estado $|0\rangle$ y luego se lo somete a una compuerta de Hadamard que lo transforma en el estado $\frac{|0\rangle+|1\rangle}{\sqrt{2}}$. En esta situación, el estado del conjunto formado por el qubit y el otro sistema está descrito por una matriz densidad que denominaremos ρ_{QS} y que es

$$\rho_{QS} = \frac{1}{2} (|0\rangle\langle 0| \otimes \rho + |0\rangle\langle 1| \otimes \rho + |1\rangle\langle 0| \otimes \rho + |1\rangle\langle 1| \otimes \rho).$$

Tras preparar este estado inicial se somete a ambos sistemas a una interacción del tipo control-U. Esto se corresponde con un operador unitario que aplica el operador U al sistema si el estado del qubit es $|1\rangle$ mientras que aplica el operador identidad si el qubit está en el estado $|0\rangle$. El operador control-U puede escribirse explícitamente como

$$\text{ctrl-U} = |0\rangle\langle 0| \otimes I + |1\rangle\langle 1| \otimes U.$$

Aplicando este operador para hacer evolucionar al estado completo ρ_{QS} se obtiene que

$$\rho'_{QS} = \text{ctrl-U} \rho_{QS} \text{ctrl-U}^\dagger.$$

Es simple aplicar este operador sobre el estado completo obteniendo que

$$\rho'_{QS} = \frac{1}{2} (|0\rangle\langle 0| \otimes \rho + |0\rangle\langle 1| \otimes \rho U^\dagger + |1\rangle\langle 0| \otimes U \rho + |1\rangle\langle 1| \otimes U \rho U^\dagger).$$

Dado que no realizaremos ninguna medición sobre el sistema inicialmente preparado en el estado ρ , podemos calcular la matriz densidad reducida del qubit, que resulta ser

$$\rho_Q = \frac{1}{2} (|0\rangle\langle 0| + |1\rangle\langle 1| + |0\rangle\langle 1| \text{Tr}(\rho U^\dagger) + |1\rangle\langle 0| \text{Tr}(U \rho)).$$

Esta matriz densidad puede escribirse en términos de la identidad y las matrices de Pauli del siguiente modo

$$\rho_Q = \frac{1}{2} (I + \text{Re}(\text{Tr}(U \rho)) \sigma_x + \text{Im}(\text{Tr}(U \rho)) \sigma_y).$$

En particular, si U es no solo unitario sino también hermítico (como es el caso del operador de punto) la parte imaginaria de $\text{Tr}(U\rho)$ es nula y por lo tanto solamente tenemos una combinación lineal de la identidad y σ_x . Es evidente que si medimos el valor medio de σ_x después de la interacción (tal como se indica en la Figura) obtendremos que

$$\langle\sigma_x\rangle = \text{Re}(\text{Tr}(U\rho)).$$

La medición de la función de Wigner puede realizarse de esta manera: podemos ejecutar muchas veces este algoritmo eligiendo operadores U de la forma $U = D(\alpha)RD^\dagger(\alpha)$ midiendo en cada α la polarización del qubit a lo largo de la componente x . De este modo mediremos la función de Wigner $W(\alpha)$ en cada punto del espacio de fases (a menos de una constante de normalización ya que $\langle\sigma_x\rangle = \text{Tr}(U\rho) = \pi W(\alpha)$).

Cabe mencionar que la medición de la función de Wigner en realidad puede hacerse aún más fácilmente. Por cierto, no es necesario generar esta interacción en la que el operador U depende del punto α . Por cierto, la alternativa más simple es aplicar sobre el sistema inicialmente preparado en el estado ρ un operador de desplazamiento $D^\dagger(\alpha)$ (o $D(-\alpha)$) y luego ejecutar el algoritmo de scattering eligiendo siempre $U = R$. De este modo obtenemos el mismo resultado que antes. La polarización en x del qubit dará como resultado el valor medio del operador de reflexión R en el estado desplazado, que no es otra cosa más que la función de Wigner $W(\alpha)$.

Esta propuesta de medición de $W(\alpha)$ fue presentada por Luis Davidovich (un notable colega brasileño) y su estudiante Luis Lutterbach en 1996. El algoritmo de scattering fue extendido e interpretado en el contexto de la tomografía y la espectroscopia unos años más tarde (en 2002) por una colaboración entre físicos argentinos y norteamericanos. En ese trabajo (titulado “Tomography and spectroscopy as dual forms of quantum computation”) se resalta el hecho de que el algoritmo de scattering puede ser visto de dos maneras distintas y complementarias. Por un lado, si dejamos fijo ρ (preparamos siempre el sistema de la misma forma) y tenemos capacidad para variar el operador U de manera controlada, ejecutando el algoritmo para una base completa del espacio de operadores, tenemos una potente herramienta topográfica. Cada valor medio nos dará una componente del estado ρ a lo largo de un operador que integra una base completa. La otra visión, complementaria a esta, es aquella en la que no variamos (o no podemos variar) el operador U que interviene en la interacción pero, en cambio, podemos variar el estado ρ . Si ejecutamos el algoritmo preparando un conjunto completo (una base) del espacio de operadores, entonces tendremos información sobre el operador U , información espectroscópica. Por ejemplo, si se prepara el estado más sencillo, el máximamente mixto, el algoritmo permitirá medir $\text{Tr}(U)$ que es la suma de los autovalores de ese operador (su traza), que contiene información importante y no es un objeto sencillo de medir de otro modo.

Finalmente, vale la pena resaltar que la operación control- R puede ejecutarse en el contexto de la implementación de electrodinámica cuántica en cavidades

usando un átomo que atraviesa la cavidad (a modo de qubit) e interactúa con el campo almacenado en la misma de manera dispersiva induciendo un desfase igual a π por cada fotón (lo cual en el contexto de cavity QED es un desafío difícil de alcanzar pero no imposible). Experimentos de este tipo fueron los que permitieron confeccionar la película que muestra la evolución de la función de Wigner realizada por el grupo de Serge Haroche y que puede encontrarse en el mencionado LINK.

6. El problema de la medición

Con más de cien años de edad, la mecánica cuántica sigue dando lugar a debates conceptuales sobre su interpretación y sus alcances. Uno de estos debates suele darse bajo el nombre del “problema de la medición” y a él nos referiremos en este capítulo. Si bien el debate es realmente importante, hay una paradoja que está vinculada con estas discusiones. Las conclusiones a las que se arriben pueden ser radicalmente distintas en lo que se refiere a la visión del Universo que la física cuántica nos provee, pero las predicciones cuantitativas que ella hace respecto a los experimentos que se realizan en todos los laboratorios del mundo son las mismas, independientemente de la interpretación que se adopte. En este sentido, la física cuántica no cumple con un principio que, según el parecer de muchos científicos y filósofos, debería satisfacer: la teoría sugiere (o induce) su propia interpretación. Con la cuántica esto no es así y, como consecuencia, desde los años 1920 tuvieron lugar intensas discusiones entre las cuales se destacan las protagonizadas por Niels Bohr y Albert Einstein sobre esta cuestión.

En resumidas cuentas, podemos describir al problema de la medición de la siguiente manera. De acuerdo a los postulados de la mecánica cuántica que expusimos en el Capítulo 4 y 6 (en este último caso, el postulado vinculado a la evolución temporal), los sistemas físicos parecen comportarse de una manera cuando no son observados por un aparato de medición y de otra, totalmente distinta, cuando lo son. Cuando no son observados, los sistemas evolucionan unitariamente y, como vimos, su dinámica puede asociarse con un operador de evolución U (que cumple que $U^\dagger U = U U^\dagger = I$). Pero según lo postulamos en el Capítulo 4, cuando se realiza una medición de un cierto observable A , se pueden detectar los resultados a_n (que forman el espectro del operador A) y tras ello el estado del sistema cambia y pasa a ser aquel definido por la proyección del estado inicial sobre el subespacio asociado al autovalor detectado. Esta evolución es lo que en la literatura se conoce como el “colapso de la función de onda” y no está descripta por un operador unitario sino de una manera radicalmente diferente. Pero el problema de la medición abarca también otra cuestión importante: ¿hasta dónde podemos aplicar la mecánica cuántica? Si la respuesta es que la cuántica puede aplicarse a cualquier sistema cerrado (a todo el Universo, por ejemplo), entonces el propio aparato debería ser descripto mediante un cierto estado, un vector en un espacio de Hilbert. Y la interacción entre el aparato y el sistema observado debería ser pasible de una descripción simple, que obviamente debe depender de la naturaleza física del sistema y del

instrumento de medición. En definitiva, si realmente creemos en la universalidad de la mecánica cuántica, deberíamos ser capaces de describir la interacción entre un sistema y un aparato tratándolos a ambos como dos partes de un sistema compuesto, que interactúan por vía de algún Hamiltoniano que está diseñado con el fin de generar correlaciones entre ciertas propiedades del sistema y otras del aparato. Pero ambos protagonistas, sistema y aparato, deberían ser descriptos de acuerdo a lo que establece la física cuántica: su estado debe asociarse a un vector en un espacio de Hilbert.

Debemos hacer aquí una confesión, una suerte de mea culpa, ya que eludimos totalmente hasta aquí la discusión de este tema apelando a un recurso que puede ser cuestionado (y la intención de este capítulo es precisamente esa: cuestionar una actitud que se refleja en muchos textos de mecánica cuántica donde se respeta la orden de muchos de los padres fundadores de esta teoría ante la necesidad de terminar con debates infinitos, “shut up and calculate!”). Hasta aquí afirmamos que el estado de un sistema es información (información en manos de un observador, que intenta describir el estado de situación en su laboratorio). Apelando a esta descripción, afirmamos que no resulta sorprendente que al adquirir información al realizar una medición, el observador debe actualizar el estado del sistema. En ese sentido, el colapso de la función de onda surge como algo natural. Es el proceso de actualización de la mencionada información. Pero, debemos confesar que en nuestra opinión, esta visión de alguna manera “esconde la basura debajo de la alfombra” ya que deja de lado la pregunta: ¿qué es la información? Y aunque esta pregunta suena demasiado pretenciosa, tiene una respuesta (parcial, pero parte de la respuesta completa, en todo caso) ya que no hay información en un sentido abstracto sino que siempre la información tiene una representación física (en el campo de la hoy muy desarrollada “información cuántica”, Wojciech Zurek propuso una consigna que alude a la historia de su país adoptivo, los EEUU: “there is no information without representation”, parafraseando la consigna que levantaban muchos de los colonos en la etapa previa a la independencia de ese país y que decía “no taxation without representation”). Entonces, el problema que describíamos antes surge nuevamente si queremos aplicar la mecánica cuántica al objeto material utilizado para recopilar, o registrar, la información sobre el sistema medido. Ahí la basura vuelve a estar a la vista y el problema de la medición se vuelve evidente nuevamente.

Está claro que este problema no existiría si aceptáramos que el universo se compone de ciertos objetos que evolucionan de acuerdo a lo establecido por la física cuántica y por otros que se comportan como “aparatos de medición”. Pero obviamente esto atenta contra la visión que considera a la cuántica como una teoría universal y a la vez induce una pregunta: ¿dónde está la frontera entre el mundo cuántico y el clásico? ¿Tiene que ver con el carácter macroscópico de los objetos de medición? ¿Con alguna otra propiedad? Estas preguntas no encuentran respuesta alguna en el marco de la física que hoy conocemos, y no consideramos razonable esta división, con una frontera rígida, entre el mundo clásico y el cuántico (sobre todo en una época en la cual cada vez resulta evidente que la física cuántica puede aplicarse a objetos cada vez más grandes” como el centro de masa de una nanopartícula que puede llevarse a su estado fundamental

pese a que la partícula tiene alrededor de 10^{11} átomos).

El primer intento serio de formular el problema de la medición apelando a una descripción cuántica del sistema y el aparato fue propuesto por el genial John von Neumann en su famoso libro sobre mecánica cuántica, en el que presenta la formulación axiomática más o menos en los términos en los que hoy la estudiamos. El tratamiento de von Neumann es tan claro y sistemático que puede utilizarse como punto de partida para exponer uno a uno los componentes fundamentales que lo constituyen. Comenzaremos por la descripción de la medición según von Neumann y luego pasaremos a la discusión de los avances realizados en las últimas décadas para resolver algunos de los problemas detectados por él.

7. El proceso de medición según Von Neumann

Consideremos un sistema S y un aparato A y describamos a ambos de acuerdo a los postulados de la mecánica cuántica. Es decir, el espacio de estados del sistema compuesto por S y A será el producto tensorial $\mathcal{H}_{SA} = \mathcal{H}_S \otimes \mathcal{H}_A$.

Consideremos que en $\mathcal{H}_S$ hay una base ortonormal formada por los estados $|\phi_j\rangle$ $j = 1, \dots, D_S$, donde D_S es la dimensión del espacio $\mathcal{H}_S$. Elegiremos a esta base como la de los autoestados del operador $\hat{A}$, que es el que intentamos medir. Estos estados son tales que $\hat{A}|\phi_j\rangle = a_j|\phi_j\rangle$. Consideremos también que el aparato tiene estados que podemos denominar $|A_j\rangle$ que también forman una base de su espacio de estados (en realidad, la dimensión del espacio de estados del sistema y el aparato no tiene necesariamente que ser la misma, el aparato puede tener una dimensión mucho mayor que la del sistema, como veremos). Los estados $|A_j\rangle$ del aparato son aquellos en los que el mismo debería encontrarse si el sistema tiene un valor a_j para el observable $\hat{A}$ y suelen denominarse “estados puntero”. Como veremos más abajo, la idea de Von Neumann es describir un proceso mediante el cual podemos asegurar que si el estado del aparato es $|A_j\rangle$ entonces el valor de la propiedad medida es a_j y el estado del sistema es $|\phi_j\rangle$. En muchas descripciones, el aparato suele considerarse como una partícula que se mueve en una dimensión (es decir, tiene un espacio de Hilbert de dimensión infinita y continua). En ese caso, la propiedad a_j se asociará a un intervalo en el cual podremos encontrar la “aguja”, o el puntero.

El proceso de medición descrito por Von Neumann puede analizarse en cuatro etapas.

Etapas 1: Pre-medición

En primer lugar, preparamos un estado del sistema en una superposición arbitraria de la base $|\phi_j\rangle$ y al aparato en algún estado de referencia $|A_0\rangle$. Es decir, el estado cuántico del conjunto sistema-aparato es

$$|\Psi_{SA}\rangle = \left(\sum_n c_n |\phi_n\rangle \right) |A_0\rangle, \quad (31)$$

donde los coeficientes c_n son números complejos que, según los postulados de la mecánica cuántica que presentamos anteriormente, están destinados a definir las probabilidades $p_n = |c_n|^2$ para cada uno de los resultados posibles de la medición. Pero esto vendrá más adelante.

Etapa 2: Interacción de medición

Después de preparar el estado inicial del conjunto S - A , se los hace interactuar mediante algún Hamiltoniano que debe ser diseñado para implementar una medición. Esto es, si el estado inicial fuera $|\phi_n\rangle|A_0\rangle$, este estado debería evolucionar en un estado de la forma $|\phi_n\rangle \otimes |A_n\rangle$. De este modo, cada estado del $|\phi_n\rangle$ del sistema debe correlacionarse con un estado $|A_n\rangle$ del aparato. Esto, naturalmente, siempre puede implementarse mediante una interacción unitaria. Podemos reflexionar un poco sobre la naturaleza de esta interacción (que podríamos denominar interacción de medición, por ejemplo). Su característica principal es que el estado del sistema determina el futuro del estado del aparato, lo cual es una suerte de generalización de lo que sucede con el operador unitario que implementa la compuerta control-not. En ese caso, el estado del qubit de control determina el estado final del blanco (al que se le aplica el operador σ_x solo si el control está en el estado $|1\rangle$). En este caso, en la interacción de mediciones el sistema tiene más de dos estados posibles, pero para cada uno de ellos el aparato pasa a estar en un estado que se obtiene del inicial aplicando un operador (una traslación generalizada) que depende de cada $|\phi_n\rangle$. En efecto, estamos imponiendo la condición que establece que

$$U_{\text{medición}}|\phi_n\rangle \otimes |A_0\rangle = |\phi_n\rangle \otimes |A_n\rangle. \quad (32)$$

Cuando el aparato está compuesto por una partícula que se mueve en una dimensión (la punta de una aguja, en una descripción analógica de una medición), el operador de evolución temporal asociado a esta etapa podría ser

$$U_{\text{medición}} = \exp(-i\lambda\hat{A} \otimes \hat{P}), \quad (33)$$

donde $\hat{P}$ es el momento del aparato y λ es una constante que mide la intensidad de la interacción. Esta interacción, como resulta evidente, genera una traslación en la posición de la aguja del aparato cuya magnitud depende del valor que adopte el observable $\hat{A}$ del sistema. Naturalmente, si el operador de evolución temporal satisface esta propiedad, el estado inicial (aquel que preparamos en la primera etapa del proceso de medición) se transforma en

$$|\Psi_1\rangle = U_{\text{medición}}|\Psi_0\rangle = \sum_n c_n |\phi_n\rangle \otimes |A_n\rangle. \quad (34)$$

En esta etapa se establecen correlaciones entre el sistema y el aparato y, evidentemente, gracias a ellas sabemos que si el estado del aparato es $|A_n\rangle$ entonces el del sistema es $|\phi_n\rangle$. En ese sentido, el aparato está detectando el valor del observable $\hat{A}$. Si bien esto parece muy natural y razonable, veremos

que para que esto conduzca a una medición correcta del observable $\hat{A}$ hay que hacer una hipótesis que el razonamiento original de Von Neumann soslaya y que fue puesta en cuestión en la década de 1980. Veremos estos cuestionamientos más adelante ya que ellos conducen a, y son resueltos por, el proceso que se conoce como “decoherencia”.

Etapa 3: Primer colapso

Hasta aquí, el proceso de medición que hemos descrito involucra solamente la interacción entre dos sistemas cuánticos (S y A) mediante un operador de evolución unitario que cumple con propiedades simples y fácilmente realizables. En esta tercera etapa, Von Neumann postula que debe haber un proceso no unitario, que caracteriza a la primera etapa del colapso de la función de onda, que lleve el estado puro $|\Psi_1\rangle$ a un estado mixto de la forma

$$\rho_3 = \sum_n |c_n|^2 |\phi_n\rangle\langle\phi_n| \otimes |A_n\rangle\langle A_n|. \quad (35)$$

En esta etapa, la evolución que conduce al estado mixto ρ_3 está fuera de aquellas descritas por la mecánica cuántica tal como la formulamos. El carácter no unitario de este tipo de evolución se hace evidente si notamos que el estado al final de la etapa 2 es puro (el vector $|\Psi_2\rangle$) mientras que el operador densidad ρ_3 corresponde a un estado mixto. En consecuencia, la evolución no preserva la pureza y, por lo tanto, no puede ser unitaria. Lo importante de la etapa 3 es que al pasar de un estado puro a otro mixto, los coeficientes c_n definen inequívocamente ciertas probabilidades. Esto es así ya que el estado ρ_3 es idéntico al estado asociado al ensemble que asocia la probabilidad $p_n = |c_n|^2$ al estado $|\phi_n\rangle \otimes |A_n\rangle$. Es decir, el estado ρ_3 describe una mezcla estadística en los que la propiedad A del sistema toma el valor a_n con probabilidad p_n , tal como lo establece la mecánica cuántica.

Etapa 4: Segundo colapso

Para finalizar su descripción del proceso de medición, Von Neumann se ve obligado a postular una nueva etapa de colapso mediante una evolución no unitaria. Esta evolución es la que lleva al estado del conjunto formado por S y A de la mezcla estadística ρ_3 en la que todos los valores del observable $\hat{A}$ están presentes, cada uno con una probabilidad p_n , a otra en la que una única respuesta, asociada a un único autovalor de $\hat{A}$, está presente. Es decir, si de la medición se obtiene el resultado a_j , el estado al finalizar la etapa 4 será

$$\rho_4 = |\phi_j\rangle\langle\phi_j| \otimes |A_j\rangle\langle A_j|, \quad (36)$$

que describe una situación en la que el sistema queda en el estado asociado al autovalor medido y el aparato en el estado correspondiente al correlacionado con dicho autovalor. Nuevamente, pasamos de una mezcla estadística a un estado puro y, por lo tanto, la evolución no puede ser unitaria. Esta cuarta, y última,

etapa resulta innecesaria si nos contentamos con una descripción del proceso de medición en términos de probabilidades asociadas a cada una de las respuestas, pero resulta necesaria si, de alguna manera u otra, aspiramos a obtener de la mecánica cuántica una descripción del estado final en el que podamos identificar unívocamente el valor de la propiedad medida.

En lo que sigue, presentaremos algunos de los problemas conceptuales que la descripción de Von Neumann contiene y discutiremos en qué medida podemos resolver estos cuestionamientos.

8. Los problemas de la descripción de Von Neumann

En la descripción anterior hay algunos problemas evidentes, que en modo alguno fueron ocultos desde el primer momento por el autor de la misma. Por ejemplo, los procesos físicos que podrían estar asociados a las etapas 3 y 4 son completamente desconocidos y no es muy plausible que se generen a partir de las leyes conocidas de la física. Son una descripción ad-hoc ya que ni siquiera motivaron la búsqueda de este tipo de procesos, aunque en años recientes se han propuesto algunos, conocidos como “modelos de colapso espontáneo”, y que involucran una ampliación de las leyes de la física que hoy conocemos (por ejemplo, los modelos de colapso espontáneo propuestos por Girardi, Rimini y Weber apelan a un colapso generado por algún proceso estocástico fundamental que se hace cada vez más importante cuanto más macroscópico es el sistema, mientras que aquellos propuestos por Penrose involucran algún proceso -desconocido hasta ahora- en el que interviene la gravitación). Pero antes de llegar a la etapa 3 pasamos por la etapa 2 que, a simple vista, parece no presentar ninguna dificultad conceptual. Sin embargo, la tiene y su resolución es crucial. Como mencionamos, en la etapa 2 se generan correlaciones entre el sistema y el aparato, los estados $|s_n\rangle$ se correlacionan con los $|A_n\rangle$. Esto es lo que nos permite decir que el aparato está midiendo el observable $\hat{A}$, que es diagonal en la base de los autoestados $|\phi_n\rangle$. Sin embargo, esto no es tan evidente y, como veremos, a menos que hagamos una hipótesis muy fuerte, la existencia de correlaciones no nos permite decir ni siquiera cuál es el observable que el aparato está midiendo.

Como vimos, el estado al final de la etapa 2 es

$$|\Psi_2\rangle = \sum_n c_n |\phi_n\rangle \otimes |A_n\rangle. \quad (37)$$

Pero resulta evidente que este estado puede escribirse de muchas otras maneras en las que no aparezcan los estados $|\phi_n\rangle$ ni los estados $|A_n\rangle$ sino otras bases de los espacios de estados del sistema y el aparato. Veamos esto con un ejemplo sencillo, que puede generalizarse fácilmente. Consideremos el caso en el cual el sistema tiene dimensiones $D_S = 2$ y supongamos que todas las amplitudes son iguales (es decir, $c_0 = c_1 = 1/\sqrt{2}$). Es decir, el estado anteriormente presentado

es

$$|\Psi_2\rangle = \frac{1}{\sqrt{2}}(|\phi_0\rangle \otimes |A_0\rangle + |\phi_1\rangle \otimes |A_1\rangle). \quad (38)$$

Ahora bien, este estado del conjunto S - A puede reescribirse en función de los siguientes

$$|\phi_{\pm}\rangle = \frac{|\phi_0\rangle \pm |\phi_1\rangle}{\sqrt{2}} \quad \text{y} \quad |A_{\pm}\rangle = \frac{|A_0\rangle \pm |A_1\rangle}{\sqrt{2}}. \quad (39)$$

En efecto, en términos de estos estados tenemos que

$$|\Psi_2\rangle = \frac{1}{\sqrt{2}}(|\phi_+\rangle \otimes |A_+\rangle + |\phi_-\rangle \otimes |A_-\rangle). \quad (40)$$

Así escrito, el estado pone de manifiesto correlaciones entre el sistema y el aparato, pero estas correlaciones involucran a los estados $|\phi_{\pm}\rangle$ y $|A_{\pm}\rangle$. Cuando lo escribimos de esta forma, el aparato parece estar midiendo otra propiedad del sistema (aquella que es diagonal en la base formada por los estados $|\phi_{\pm}\rangle$): podríamos decir que si el aparato se encuentra en el estado $|A_-\rangle$ entonces el sistema estará en $|\phi_-\rangle$, etc. Entonces, la etapa 2 no describe bien el proceso de medición de un observable sino que simplemente un proceso en el cual se generan correlaciones entre S y A . Este problema en la formulación de Von Neumann fue discutido ampliamente durante la década de 1980 por Wojciech Zurek, Dieter Zeh y otros. En algún sentido, la etapa 2 no nos dice qué es lo que el aparato está midiendo a menos que hagamos una hipótesis fundamental. Esta hipótesis es que el aparato siempre se encontrará en alguno de los estados punteros, aquellos asociados a la base que denominamos $|A_n\rangle$. En algún sentido, a menos que aceptemos que esta es una base privilegiada y que debemos utilizarla, entonces no podemos decir que durante la etapa 2 se establecen correlaciones que implican la medición del observable $\hat{A}$. El aparato de medición es, entonces, un sistema que no es tan cuántico como lo planteamos. Si solamente puede existir (o ser preparado) en estados punteros entonces no todos los estados de su espacio de Hilbert son estados posibles: el aparato no cumple, entonces, con el principio de superposición. A menos que aceptemos este hecho, tampoco podríamos avanzar hacia la etapa 3 (el primer colapso) ya que, nuevamente, deberíamos preguntarnos en qué base debemos escribir al estado mixto ρ_3 ? O, dicho de otro modo, ¿cuál es el motivo por el que podemos escribir la expresión antes mencionada para ρ_3 y no otra en la que los estados del aparato que aparezcan sean los estados $|A_{\pm}\rangle$ y los del sistema aquellos que denominamos $|\phi_{\pm}\rangle$.

Repasemos un poco las implicancias de lo que acabamos de expresar: Para que el aparato sirva para medir un dado observable es necesario que posea una base privilegiada tal que los estados puntero $|A_n\rangle$ puedan ser preparados y existir de manera estable mientras que sus superposiciones deben de estar excluidas del menú de estados posibles. ¿Cuál podría ser el origen del comportamiento de un sistema cuántico que dé lugar a una propiedad como esta? La respuesta a esta pregunta fue uno de los hallazgos de la década de 1990 y se conoce bajo el nombre de un proceso físico muy simple: la decoherencia.

9. Decoherencia y la aparición de una base privilegiada en los estados del aparato

A finales de la década de 1980 se tomó conciencia de que un proceso físico muy simple podía servir para explicar la aparente paradoja planteada en la etapa 2 del proceso de medición descrito por Von Neumann y a la vez explicar dinámicamente la etapa 3 (es decir, la transición de un estado puro del sistema S - A a otro mixto). Se trata del proceso de decoherencia, inducido por la interacción del aparato con otro sistema físico que lo rodea, formando un entorno en el cual este aparato vive. Es decir, para resolver el problema de la medición (hasta el final de la etapa 3) es necesario ampliar nuestro universo: en lugar de estar limitado al conjunto formado por el sistema S y el aparato A , debemos incluir a otra componente, un entorno E formado por un conjunto muy grande de grados de libertad a los que, en general, no tenemos acceso o no podemos controlar. El triplete S - A - E constituye el universo y el entorno E interactúa con el aparato A . De esta interacción, podemos ver que se deriva el comportamiento clásico del aparato y la aparición de una base privilegiada de estados, que se comportan clásicamente. Los estados $|A_n\rangle$ serán relativamente inmunes a la interacción con el entorno (o, en algún sentido preciso, serán los más inmunes de su espacio de estados) mientras que sus superposiciones serán inestables y decaerán rápidamente en mezclas estadísticas del conjunto S - A . En el marco de esta idea, el objeto que debemos analizar es la matriz densidad reducida del conjunto S - A , que se obtiene tomando la traza parcial sobre el entorno E . La interacción entre A y E produce correlaciones entre ellos, que en muchos casos corresponden a genuino entrelazamiento. Dependiendo del tipo de interacción que exista entre S y E (así como también de sus propios Hamiltonianos, que definen la dinámica interna de estas componentes) la base preferida $|A_n\rangle$ tendrá diversos tipos de propiedades. Analicemos el caso más sencillo de todos: uno en el cual el estado inicial del entorno es puro (lo natural, en la mayoría de los casos, es considerar entornos térmicos pero podemos comenzar describiendo el caso de temperatura nula). La dinámica combinada de la interacción propia de la etapa 2 y la interacción entre A y su entorno, darán lugar a un estado del conjunto S - A - E de la forma

$$|\Psi_3\rangle = \sum_n c_n |\phi_n\rangle \otimes |A_n\rangle \otimes |E_n\rangle, \quad (41)$$

donde los estados $|E_n\rangle$ pertenecen al espacio de estados del entorno. Esta interacción puede describirse realizando diversas aproximaciones que muchas veces pueden ser razonables y otras no tanto. Por ejemplo, la aproximación típica consiste en suponer que la interacción asociada a la medición domina completamente la evolución durante un intervalo de tiempo muy corto. En ese momento, se puede despreciar la interacción con el entorno pero antes y después es esa la interacción que domina la evolución del aparato. En cualquier caso, de ese modo podremos obtener una expresión para el estado del conjunto de tres subsistemas y, tomando la traza parcial sobre el entorno, se determina el estado del conjunto

formado por S y A . De acuerdo a las reglas usuales de la mecánica cuántica (es decir, sin hipótesis adicionales de ningún tipo) este estado es

$$\rho_{SA} = \sum_{nm} c_n c_m^* |\phi_n\rangle \langle \phi_m| \otimes |A_n\rangle \langle A_m| \langle E_m | E_n \rangle \quad (42)$$

El proceso de decoherencia tiene lugar cuando los estados $|E_n\rangle$ que se correlacionan con los correspondientes punteros del aparato son ortogonales entre sí (al menos, en forma aproximada). En ese caso, el solapamiento entre estos estados será muy pequeño, $\langle E_m | E_n \rangle \approx 0$. En ese caso, la doble suma colapsa en una sola y el estado del conjunto S - A es precisamente igual al que Von Neumann requiere al finalizar la Etapa 3 del proceso de medición. Es decir, la decoherencia viene, en este caso, a resolver dos problemas de manera simultánea. Por un lado, explica que solamente los estados punteros del aparato (o las mezclas estadísticas de ellos) pueden ser estables. Las superposiciones de estados punteros decaen, al correlacionarse cada posición del puntero con un estado del entorno que es ortogonal al que se correlaciona con otra posición del puntero, en una mezcla estadística. La decoherencia se manifiesta, entonces, como una tendencia a la diagonalización de la matriz densidad del aparato en una cierta base, que es seleccionada por la interacción con su entorno. Pero, por otra parte, al seleccionar a los estados punteros, la decoherencia explica también la primera Etapa del colapso descrita por Von Neumann en su Etapa 3. No dice nada, en cambio, respecto a la Etapa 4 y genera, entonces, un límite clásico que siempre es de carácter estadístico, lo que para muchos autores es enteramente satisfactorio pero para algunos otros no lo es. La decoherencia no dice nada sobre cómo se selecciona una dada posición del aparato respecto a otra sino que describe un escenario donde ambas alternativas están presentes con distintas probabilidades.

10. El proceso de decoherencia y la frontera entre el mundo clásico y el cuántico

Analizaremos aquí algunos detalles de modelos realistas del proceso de decoherencia que es inducido por un entorno cuando este interactúa con algún objeto cuántico. En las discusiones usuales que se encuentran en la literatura sobre el proceso de decoherencia solo aparecen dos protagonistas: un sistema cuántico S y su entorno (también cuántico) E . En la sección anterior presentamos este proceso en el contexto de la discusión del problema de la medición y por eso aparecieron tres protagonistas, el sistema S (cuya propiedad $\hat{A}$ queríamos medir), el aparato de medición A y el entorno E , que interactuaba con el aparato. El entorno, en ese contexto, es el responsable de inducir el comportamiento clásico del aparato A . En lo que sigue, dejaremos de lado al tercero (al sistema S) y nos concentraremos en la discusión del proceso de decoherencia inducido sobre un sistema A (preservaremos la letra A para denominar al sistema para evitar confusiones con las secciones anteriores).

La decoherencia es un proceso que está presente en numerosos sistemas fisi-

cos. De hecho, todos aquellos que se comportan de manera clásica en nuestros laboratorios (¡que son muchísimos!) deben su carácter clásico a algún tipo de decoherencia con un entorno u otro, que depende de la física del sistema. Atribuir la clasicidad a la decoherencia es algo que podemos hacer y que es un salto conceptual importante. El comportamiento clásico de los sistemas físicos no es atribuido, según este punto de vista, al carácter macroscópico de los objetos ni a nada por el estilo sino a su apertura a la interacción con un entorno. Cualquier sistema que interactúe, aunque sea muy débilmente, con un entorno suficientemente complejo, sufrirá decoherencia y solo podrá existir de manera estable en un conjunto muy pequeño de estados cuánticos. Los demás estados, la mayoría de los que constituyen el espacio de estados del sistema, se vuelven inestables al proceso de decoherencia y decaen rápidamente en mezclas estadísticas. Veremos aquí la razón por la cual en los modelos relevantes para el proceso de medición, por ejemplo, la decoherencia ocurre en una escala de tiempos muy pequeña, que suele ser la más pequeña de todas las escalas presentes en el problema.

Hemos visto cómo describir la dinámica de un sistema cuántico abierto, que evoluciona acoplado continuamente con un entorno. Varios de los modelos que estudiamos en ese capítulo son apropiados para describir el proceso de decoherencia. Antes de hacerlo, conviene resaltar una característica del proceso de decoherencia: es muy difícil estudiar este proceso de manera controlada y recién en las últimas décadas se han conseguido realizar experimentos en los que se observa claramente el decaimiento de la coherencia cuántica (el pasaje de un estado que es superposición de varias alternativas microscópicamente distinguibles) a una mezcla. Por cierto, lo que resulta muy difícil de lograr es contar con un sistema A que interactúe con un entorno E de modo tal que la interacción entre ellos tenga una intensidad que sea controlable desde afuera. Es decir, es muy difícil contar con una “perilla” que aumente o disminuya la intensidad de dicha interacción de modo tal de aumentar o disminuir la decoherencia. El primer experimento en el que se observó la dinámica del proceso de decoherencia, el decaimiento en tiempo real, fue realizado por el grupo de Serge Haroche en París y consistió en preparar un estado tipo “gato de Schrödinger” del campo electromagnético en una cavidad y observar su decaimiento a lo largo del tiempo. Hemos discutido en detalle el proceso de preparación de estados “gato” en los experimentos de CQED en el Capítulo XX. Asimismo, hemos presentado una ecuación maestra que tiene la forma de Lindblad y que describe la evolución de un oscilador armónico amortiguado (con un único operador de Lindblad que es proporcional al operador de destrucción a). Es decir, estamos en condiciones de describir la evolución del gato de Schrödinger preparado por Haroche y verificar la concordancia de nuestra descripción teórica con el experimento. En efecto, la ecuación de Lindblad resultante puede resolverse por distintos métodos numéricos y puede comprobarse que la coherencia cuántica decae tan rápidamente, con un tiempo característico que igual al tiempo que tarda un fotón almacenado en la cavidad en absorberse en las paredes que la conforman. Es decir, apenas se absorbe un fotón el sistema pierde coherencia, aunque la variación de su energía media sea muy pequeña. Esta diferencia entre la escala de tiempos que caracteriza al decaimiento de la energía y de la coherencia cuántica es

una característica de muchos sistemas y da lugar a la existencia de un régimen en el cual el sistema se comporta clásicamente debido a la interacción con su entorno pero todavía es reversible desde el punto de vista de su dinámica ya que la energía se conserva aproximadamente en esa escala de tiempos. Como parte del material complementario de este capítulo presentamos la solución de la ecuación de Lindblad para el gato de Schrödinger creado por Haroche, pero aquí nos enfocaremos en la discusión de otro caso, de características similares, que es más relevante en el contexto del problema de la medición.

También vimos en la sección correspondiente la solución del problema del movimiento Browniano cuántico. Este caso es el que analizaremos en detalle aquí ya que provee un escenario un poco más realista que el anterior en cuanto a lo que se refiere a la identificación de un mecanismo microscópico de interacción entre el aparato y su entorno. En efecto, en el escenario provisto por el MBC la interacción entre el sistema y su entorno es por vía de un Hamiltoniano que es proporcional al operador posición del puntero del aparato y también es proporcional a los operadores posición de los infinitos osciladores que conforman el entorno (en efecto, como vimos en el Capítulo XX, $H_{\text{int}} = \sum_k \lambda_k q_k \cdot x$, donde x es la coordenada del sistema A y q_k las de los osciladores del entorno). Este modelo admite una solución exacta y ella fue usada para estudiar la dinámica del proceso de decoherencia. Usaremos aquí una versión aproximada de la ecuación maestra que describe la evolución de la matriz densidad reducida de A (donde ya hemos tomado la traza sobre el entorno) y que se obtiene en el límite de altas temperaturas. Es decir, supondremos que el entorno está en un estado inicial de equilibrio térmico a una cierta temperatura T y que la escala de energías que esta provee ($k_B T$) es la más grande de todas las presentes en el problema. En este límite (que, como mencionamos, debe utilizarse con cuidado) la ecuación maestra puede escribirse como

$$\dot{\rho} = \frac{1}{i\hbar}[H, \rho] + \gamma[x, \{p, \rho\}] - D[x, [x, \rho]]. \quad (43)$$

El coeficiente γ define la tasa de pérdida de energía, es decir, es el que determina la tasa de amortiguamiento del movimiento y aparece, por ejemplo, en la ecuación de evolución de los valores medios que en el caso del momento resulta ser $\langle \dot{p} \rangle = -\gamma \langle p \rangle - m\omega^2 \langle x \rangle$. El coeficiente D define la tasa de difusión del sistema y puede escribirse como $D = \gamma \frac{mk_B T}{\hbar^2}$, o equivalentemente $D = \gamma / \lambda_T^2$ donde λ_T es la longitud térmica de de Broglie que se define como $\lambda_T = \hbar / \sqrt{mk_B T}$. Como mencionamos, esta ecuación no tiene la forma de Lindblad (en particular, el término de disipación no respeta esa estructura). Pero la ecuación puede utilizarse para evolucionar estados en los que la dispersión en posición no sea mucho menor que λ_T (en esos casos la ecuación anterior conduce a una violación inicial de la positividad ya que el término disipativo domina la evolución). Eso es lo que haremos aquí.

Estudiaremos qué es lo que sucede si consideramos un estado inicial del sistema A formado por la superposición de dos estados coherentes de la forma

$$|\Psi_0\rangle_S = N(|\alpha\rangle + |-\alpha\rangle), \quad (44)$$

donde N es una constante de normalización y elegiremos α de modo tal que el valor medio de la posición sea

$$\langle x \rangle = L \quad (45)$$

(o sea, $\alpha = L/\sigma\sqrt{2}$). La matriz densidad reducida inicial puede escribirse como la suma de cuatro términos

$$\rho_S(0) = N^2(|\alpha\rangle\langle\alpha| + |-\alpha\rangle\langle-\alpha| + |\alpha\rangle\langle-\alpha| + |-\alpha\rangle\langle\alpha|). \quad (46)$$

Los dos primeros términos son casi diagonales en la base de posición ya que la función de onda del estado $|\alpha\rangle$ es una Gaussiana centrada en $x = L$ y con ancho σ (que se supone que es mucho menor que L). Los términos no diagonales son los dos últimos y contienen dos picos, uno de ellos ubicado en $x = L$ y el otro en $x = -L$. En el límite de alta temperatura y muy baja disipación, podemos despreciar el término proporcional a γ en la ecuación maestra y considerar que el sistema evoluciona inicialmente dominado por la difusión. En ese caso se obtiene una solución simple de la ecuación maestra, que es válida para tiempos cortos (comparados con $1/\gamma$). Esta resulta ser

$$\rho_{\text{ND}}(t) = \rho_{\text{ND}}(0) \exp\left(-\gamma t \frac{4L^2}{\lambda_T^2}\right), \quad (47)$$

donde ρ_{ND} denota la parte no diagonal de ρ_A en la base de autoestados de la posición. Entonces, como consecuencia de la interacción con el entorno, la matriz densidad del sistema A tiende a diagonalizarse en la base de autoestados de la posición. La tasa de decaimiento de los términos no diagonales define a la tasa de decoherencia, que resulta ser en este caso

$$\gamma_{\text{deco}} = \gamma \left(\frac{2L}{\lambda_T}\right)^2. \quad (48)$$

Para entender este resultado conviene analizarlo un poco con algunos valores concretos para los parámetros. Si reemplazamos distancias $2L$ del orden de 1 centímetro, temperatura de 300K y masas de 1 gramo, entonces la tasa de decoherencia resulta ser 30 órdenes de magnitud más baja que la disipación. Es decir, en ese caso $\left(\frac{2L}{\lambda_T}\right) \approx 10^{-15}$. Más allá de que un resultado como este no es rigurosamente aplicable, pone de manifiesto la propiedad antes anunciada: la tasa de decoherencia es mucho mayor que la de disipación, lo que da lugar al surgimiento de un régimen en el cual el sistema se comporta clásicamente debido a la interacción con el entorno, siendo que el entorno no ha tenido tiempo todavía como para afectar significativamente a la energía del problema.

11. ¿Cuál es la base de estados punteros?

Como mencionamos más arriba, el proceso de decoherencia induce el comportamiento clásico de un aparato siempre y cuando seleccione dinámicamente

un conjunto de estados $|A_j\rangle$ que son preferidos. Estos estados deben ser relativamente estables mientras que sus superposiciones deben decaer en mezclas estadísticas de los mismos. La naturaleza de los estados puntero fue discutida abundantemente en la literatura científica durante la década de 1990 y se lograron identificar distintos casos significativos, que describiremos aquí.

A. El entorno que registra una huella del aparato

El primer caso analizado fue aquel en el cual la dinámica del sistema A (el aparato de mediciones) no es relevante y su Hamiltoniano puede suponerse nulo (o sea, los cambios generados por su propia dinámica ocurren en tiempos que son infinitos para aquellos que caracterizan la situación que estudiamos). En ese caso, el Hamiltoniano total es la suma del del entorno (que podríamos modelar como un baño de osciladores, por ejemplo) sumado a la interacción con el aparato, que podemos considerar que es del mismo tipo que aquella analizada en el movimiento Browniano cuántico (en la que aparece la posición de A y la de los osciladores del entorno). Si suponemos que cada uno de los osciladores del entorno está en un estado inicial coherente (con una cierta amplitud que, como veremos, no resulta relevante para lo que sigue) es fácil encontrar la solución exacta de este problema. En efecto, cada oscilador del entorno sentirá una fuerza constante cuya magnitud depende de la coordenada x del aparato (el operador que aparece en el Hamiltoniano de interacción, que es lineal también en q_k). Como vimos en el Capítulo XX, dedicado al estudio del oscilador armónico, en este caso el estado coherente inicial se desplaza en una magnitud que depende de x y del tiempo y que puede calcularse fácilmente (de hecho, el cálculo explícito está presentado en el capítulo correspondiente). Realizando este cálculo podemos obtener el solapamiento entre el estado del entorno que se obtiene cuando la coordenada del sistema es x y aquel que surge cuando dicha coordenada es x' . El resultado es, simplemente

$$\langle \Phi_{eX}(t) | \Phi_{eX'}(t) \rangle = \exp \left(-\frac{(x-x')^2}{\hbar} \sum_k \lambda_k^2 \frac{2 \cos^2(\omega_k t/2)}{m_k \omega_k^3} \right). \quad (49)$$

Entonces, estos estados se vuelven ortogonales a medida que la diferencia entre la posición de la aguja del aparato se hace cada vez más grande. Obviamente, el hecho de que el mecanismo de interacción haya involucrado al observable x del aparato y no a algún otro no resulta relevante y el resultado se puede generalizar fácilmente reemplazando la diferencia $(x-x')$ por la diferencia entre dos autovalores del observable que aparece en el Hamiltoniano de interacción. Por ese motivo, podemos decir que los estados punteros en este caso serán autoestados del operador posición, que por otra parte no se verán afectados por la evolución generada por el Hamiltoniano $H_{\text{tot}} = H_{\text{entorno}} + H_{\text{int}}$. Este régimen, en el cual como dijimos, los estados punteros son autoestados del Hamiltoniano de interacción entre el aparato y su entorno, fue el primer caso estudiado por Wojciech Zurek y, a raíz de eso, por varios años se definió a los estados punteros como aquellos autoestados del Hamiltoniano de interacción. Pero pronto se

tomó conciencia clara de que esto sucedía bajo la aproximación crucial de que la dinámica propia del aparato A era irrelevante en el problema. Esta puede ser una buena aproximación para un aparato (que está quieto esperando ser utilizado para medir alguna magnitud física) pero no es nada razonable para otros sistemas. Por cierto, si queremos explicar la clasicidad de todos los sistemas físicos como una propiedad derivada del proceso de decoherencia entonces es vital analizar otros casos en los que H_A no se anula.

Entonces, se hizo necesario contar con una manera de definir los estados punteros que sea independiente del modelo o del mecanismo de interacción y que pueda ser utilizada en un contexto más general. Fue así, que surgió la idea de definirlos mediante lo que se dio en llamar “la criba de la predictibilidad” (predictability sieve). La idea es que los estados punteros deben ser aquellos que son mínimamente afectados por el entorno y, consecuentemente, su evolución es predecible. Eso es lo que sucede en el caso anterior: cuando la evolución del aparato A es despreciable, los estados punteros son seleccionados por el entorno y son inmutables, no cambian. Con esta idea se buscaron varias alternativas para definir a los estados punteros y se adoptó una que tiene la gran virtud de capturar algo de la idea esencial antes expresada y a la vez puede ser calculada en algunos ejemplos sencillos que rápidamente se volvieron emblemáticos. En efecto, se propuso cuantificar el efecto del entorno sobre el sistema apelando al estudio de la pérdida de pureza definida como

$$\chi = 1 - \text{Tr}(\rho^2). \quad (50)$$

Como dijimos, esta no es la única forma de estudiar y cuantificar el efecto del entorno sobre el sistema pero resulta sumamente simple analizarla si se cuenta con una ecuación maestra que determina la evolución de la matriz densidad ρ . Dado que $\dot{\chi} = -2 \text{Tr}(\rho \dot{\rho})$, la pérdida de pureza está cuantificada por el valor medio de la parte derecha de la ecuación maestra (excluyendo al término del conmutador con el Hamiltoniano, que no da lugar a pérdida de coherencia alguna ya que en general resulta que $\text{Tr}(\rho[H, \rho]) = 0$, para todo H (lo cual no resulta sorprendente ya que la evolución unitaria, generada por el conmutador de ρ con el Hamiltoniano, preserva la pureza). La criba de predictibilidad fue usada en varios contextos y con ella se identificaron dos nuevos regímenes, que están asociados a dos situaciones físicas completamente distintas de aquella vinculada al proceso de medición (en la que el aparato A no evoluciona). Por una parte se estudió el límite opuesto, en el cual las escalas de tiempo involucradas en la evolución del entorno son mucho mayores que aquellas asociadas a la dinámica del aparato A . Es el régimen del “entorno lento” que en ciertas circunstancias es apropiado para describir sistemas de espines acoplados con un baño formado también por espines. Por otra parte, se estudió el caso intermedio, en el cual ni el sistema ni el entorno son lentos sino que interactúan entre sí, afectándose mutuamente su estado a la vez que cambian en el tiempo. En este caso, el modelo paradigmático en el que se analizó este régimen es el del movimiento Browniano cuántico. En ambos regímenes, como veremos, los estados punteros son muy diferentes entre sí y también son distintos de los que emergen en el

régimen estudiado por Zurek, que está dominado por la interacción de A con su entorno. Veremos estos dos casos por separado.

B. Los estados coherentes son estados punteros

Analicemos el caso en el que nuestra partícula, que tiene una coordenada x , interactúa con un baño térmico de osciladores armónicos, que es el modelo que venimos analizando reiteradamente. Sabemos que en el límite de altas temperaturas la ecuación maestra es la ecuación xx que vimos más arriba. Con ella, podemos calcular la derivada de la pureza como

$$\dot{\chi} = 2\gamma\chi - D\text{Tr}(\rho x \rho x - x^2 \rho^2). \quad (51)$$

Claramente, la ecuación no podría aplicarse si el segundo término fuera mucho menor que el primero ya que el primero hace aumentar la pureza, que no puede ser mayor que la unidad (y por lo tanto conduciría a una violación de la positividad de los estados). Por cierto, como ya hemos dicho reiteradamente, la ecuación solo puede utilizarse para evolucionar estados en los que la dispersión en posición no sea demasiado chica (comparada con la longitud térmica de de Broglie). En efecto, para estados puros en los que $\rho^2 = \rho$, el lado derecho de la ecuación anterior se simplifica y resulta ser

$$\dot{\chi} = 2\gamma\chi - D\Delta^2 x, \quad (52)$$

donde $\Delta^2 x = \langle x^2 \rangle - \langle x \rangle^2$. Como una primera aproximación a la solución de la ecuación anterior podemos considerar que el estado es casi puro y mantener la ecuación anterior como

$$\dot{\chi} = 2\gamma\chi - 2D\Delta^2 x. \quad (53)$$

En el régimen donde la disipación es despreciable (γ tiende a cero, pero D se mantiene finito) el primer término del lado derecho puede ser despreciado. Nuevamente, si despreciamos la dinámica del aparato A , los estados que preservan la pureza son aquellos para los que $\Delta x = 0$ (los autoestados de la posición). Pero, en cambio, si tomamos en cuenta la evolución del aparato A , cuyo Hamiltoniano también es el de un oscilador armónico con una cierta frecuencia Ω , debemos reemplazar del lado derecho la expresión de la dispersión en función del tiempo, lo cual es una tarea sencilla: usando la expresión para los operadores de Heisenberg de un oscilador armónico vemos que Δx evoluciona en el tiempo convirtiéndose en Δp tras un cuarto de periodo y retornando a Δx tras otro cuarto de periodo. Obviamente, podemos preguntarnos cuál será el estado que minimice la velocidad con la que disminuye la pureza en un dado instante, pero en ese caso obtendremos una respuesta diferente para cada instante: el estado en cuestión será aquel que minimice la dispersión $\Delta^2 x(t)$ en ese instante y en general no será ni un autoestado de momento ni un autoestado de la posición. Pero la pregunta que resulta más relevante es si existe un conjunto de estados que dan lugar al mínimo valor de $\dot{\chi}$ para un conjunto de instantes que sea relevante para la dinámica del sistema. Es decir, si existe un conjunto robusto de

estados que minimizan $\dot{\chi}$ al menos para una escala relevante de tiempos. Para eso es útil razonar de la siguiente manera: Podemos analizar cómo es la forma de $\dot{\chi}$ cuando promediamos el lado derecho de la ecuación anterior para un periodo de la evolución. O sea, cuál es el estado que minimiza $\dot{\chi}$, el promedio temporal de la velocidad de caída de la pureza. En ese caso, tomando el promedio sobre un periodo, la ecuación para el valor medio de $\dot{\chi}$ es:

$$\dot{\chi} = -D(\Delta^2 x + \Delta^2 p/m^2 \Omega^2), \quad (54)$$

donde las dispersiones que aparecen aquí son las del estado inicial (el estado puntero, que queremos encontrar). Entonces, de acuerdo a este argumento, los estados punteros en este régimen dinámico, en el cual el aparato A evoluciona en el tiempo, serán aquellos que minimicen la suma de los cuadrados de las dispersiones, que puede reescribirse como

$$\dot{\chi} = -D(\langle H \rangle - \bar{H})/m\Omega^2, \quad (55)$$

donde $\bar{H}$ es la energía asociada a los valores medios de la posición y del momento, es decir $\bar{H} = \langle p \rangle^2/2m + m\Omega^2 \langle x \rangle^2/2$. Por consiguiente, los estados que minimizan el promedio temporal de $\dot{\chi}$ son aquellos que hacen mínima la diferencia entre el valor medio de la energía y la energía de los valores medios. Como vimos en el Capítulo XX, esto define unívocamente a los estados coherentes (los autoestados del operador de aniquilación). Por este motivo, podemos afirmar que los estados coherentes son seleccionados como estados punteros en este célebre y relevante modelo del movimiento Browniano cuántico.

C. Los autoestados de la energía como estados punteros

Un nuevo régimen donde se genera un tipo drásticamente distinto de estados punteros surge cuando las escalas de tiempo que caracterizan la evolución del entorno son mucho más largas que aquellas que aparecen en la dinámica del sistema. Es el caso del “entorno lento”. Es bastante simple obtener la principal característica de los estados punteros en este caso mediante un análisis que no involucra demasiada complejidad: En efecto, si el Hamiltoniano total es la suma de tres términos, el correspondiente al aparato A , el que es propio del entorno E y el que determina la interacción entre ambos

$$H_{\text{TOT}} = H_A + H_{\text{int}} + H_E, \quad (56)$$

podemos pasar a una representación de interacción donde tomamos $H_0 = H_A$. En este caso, el Hamiltoniano en esa representación será simplemente

$$H_I = H_{\text{int}}^I(t) + H_E. \quad (57)$$

En el Hamiltoniano de interacción, en esta representación de interacción, aparecen las frecuencias de Bohr del aparato A que, como afirmamos en primer lugar, son mucho mayores que cualquier otra escala de frecuencias en el problema. Por consiguiente, la interacción evoluciona (oscila) mucho más rápidamente de

lo que cualquier evolución temporal generada por H_E . Si tomamos un promedio temporal sobre la escala de tiempos más cortas, la dependencia temporal de H_{int} se promediará a cero salvo cuando en el Hamiltoniano de interacciones aparezca un observable del aparato A que conmute con su Hamiltoniano. Es decir, cuando en H_{int} aparezca una constante de movimiento para el aparato A . Solo en este caso, la interacción entre el aparato y su entorno podrá seleccionar estados punteros para A . Naturalmente, estos estados serán los autoestados del H_{int} que, son también autoestados de H_A , ya que estos operadores conmutan entre sí. Por eso, en este caso los estados punteros son autoestados de la energía de A .

D. La decoherencia y los problemas que resuelve

Resumamos aquí lo que hemos presentado hasta ahora. Como vimos, la formulación habitual del problema de la medición presenta dos dificultades importantes. Por un lado, si describimos al aparato y al sistema medido cuánticamente no podemos realmente asegurar que el aparato y el sistema se correlacionen de modo tal que el observable medido sea $\hat{A}$, el que deseamos medir. Esto ocurre por un motivo simple: si el aparato de medición puede ser preparado y existir establemente en cualquiera de sus estados cuánticos entonces las correlaciones entre S y A no seleccionan unívocamente un observable del sistema S que tiene una base de autoestados que se correlaciona unívocamente con un conjunto de estados punteros del aparato (esta es la crítica a la formulación de la Etapa 2 hecha por Von Neumann). Por otra parte, el análisis tradicional del problema de la medición no brinda ninguna explicación dinámica para la primera (ni para la segunda) etapa del colapso, que está descrita por la Etapa 4 de la descripción de Von Neumann. La decoherencia generada sobre A por su interacción con un entorno E resuelve estos dos cuestionamientos de manera simultánea. Por un lado, la decoherencia selecciona dinámicamente un conjunto de estados “puntero” del aparato A que son los únicos en los que este dispositivo puede existir de manera estable. Estos estados, que son mínimamente perturbados por la interacción con el entorno, pueden ser determinados operacionalmente mediante un criterio del tipo de la “Criba de predictabilidad” (predictability sieve). Pero al mismo tiempo, el proceso de decoherencia explica dinámicamente el motivo por el cual el estado reducido del aparato se vuelve diagonal en la base de estados puntero, resolviendo el problema de la Etapa 3.

Algo sobre Darwinismo cuántico y el OI