

Man in the Middle y los protocolos cuánticos

Miguel Ángel López Muñoz

6 de noviembre de 2007

Índice

1. Introducción	2
2. La teoría cuántica	2
2.1. La mecánica cuántica	2
2.2. La computación cuántica	3
2.3. La información cuántica	10
2.4. La criptografía cuántica	11
3. Conceptos básicos	16
3.1. El qubit	16
3.2. Múltiples qubits	17
3.3. Notación de Dirac	18
3.4. Puertas cuánticas	19
4. Base del protocolo	21
4.1. Indistinguibilidad de qubits	21
4.2. No clonación	22
4.3. Distinguir entre qubits EPR y no EPR	23
4.4. El ataque EPR Man in the Middle	24
5. Descripción del protocolo	25
6. Comentarios de diseño	27
6.1. Acerca del Paso 5	27
6.2. Acerca del Paso 1	28

6.3. Acerca del Paso 3	31
6.4. Acerca del Paso 4	32

1. Introducción

Este trabajo pretende introducir un nuevo protocolo cuántico basado en pares EPR para generar y compartir claves que no está afectado por el ataque llamado EPR Man in the Middle. Este método de criptoanálisis, que posee múltiples variantes incluso en criptografía clásica, ha demostrado ser tremendamente efectivo en muchos de los protocolos basados en pares EPR que se han diseñado hasta ahora. Añadimos también una serie de comentarios respecto a su diseño, además de una introducción a la computación cuántica con los conceptos fundamentales de la misma para facilitar su comprensión.

2. La teoría cuántica

Antes de empezar a hablar de la criptografía cuántica es mejor detenerse primero en la computación cuántica y la información cuántica, dos ramas de estudio recientes que, además de estar muy emparentadas con la criptografía cuántica, ayudarán a mostrar qué puede aportar este nuevo enfoque.

2.1. La mecánica cuántica

La computación cuántica y la información cuántica son el estudio de las tareas de proceso de información que pueden ser llevadas a cabo empleando sistemas mecánicos cuánticos. Esta idea, tan simple en su concepción, está demostrando ser tan profunda que gracias a ella se han obtenido resultados que se resistían a ser descubiertos con las herramientas clásicas de las que se disponía hasta la fecha.

El motivo más importante de que esté ocurriendo esto es la base física sobre la que se sustentan dichos estudios. Desde su descubrimiento a comienzos de los años veinte, la mecánica cuántica no ha dejado de sorprender a los físicos. Su nacimiento fue motivado por el hecho de que la física que existía hasta la fecha, ahora conocida como física clásica, empezó a predecir una serie de fenómenos que sencillamente no podían suceder, como electrones dando vueltas infinitamente alrededor del núcleo del átomo. Al principio esto se solucionó añadiendo una serie de nuevas hipótesis construidas de tal modo que la estructura general mantenía

coherencia, pero a medida que se sabía más acerca del átomo y la radiación esas hipótesis estaban cada vez más comprometidas.

Se puede decir que la mecánica cuántica es un conjunto de reglas para la construcción de teorías físicas que conciernen al mundo de lo microscópico o atómico, para ser más precisos. La mecánica cuántica ha permitido estudiar gran cantidad de fenómenos aparentemente inconexos entre sí, como la fusión nuclear en las estrellas, la estructura del átomo o los superconductores. Como muchos avances revolucionarios en el mundo de la ciencia (como la teoría de la relatividad), las normas por las que se rige son aparentemente sencillas pero a su vez encierran una profunda complejidad debido al hecho de que se oponen a todo lo que hasta entonces ha sido considerado al respecto. En el caso de la mecánica cuántica, además, dan lugar a resultados que van en contra de la intuición.

Por ejemplo, a principios de los años ochenta, se intentó averiguar si era posible emplear resultados de mecánica cuántica para enviar una señal a mayor velocidad que la de la luz, algo que sería imposible según la teoría de la relatividad de Einstein. Un estudio más detallado en esta materia revela que resolver ese problema se reduce a conseguir clonar un estado cuántico desconocido a priori. Sin embargo se sabe que es imposible copiar un estado cuántico arbitrario, un hecho que no deja de ser chocante, ya que tanto en computación clásica como en teoría de la información clásica realizar una copia de información es no sólo posible, sino prácticamente trivial. Este sorprendente resultado revela que un potencial programador cuántico no podría emplear una instrucción tan sencilla como la reasignación de una pieza de información, algo que sin duda resulta extremadamente chocante para todo aquel que ha programado alguna vez en un ordenador. Este resultado acerca de la no clonación de un estado cuántico es uno de los primeros obtenidos en computación e información cuántica.

2.2. La computación cuántica

Para entender bien lo que es la computación cuántica primero hay que pararse a pensar un momento en lo que es la computación clásica. En realidad, ya existían ideas relativas a la computación mucho antes de que se inventara el primer ordenador. Uno de los primeros en introducirse sin saberlo en este área fue Leibnitz. En el año 1686 escribe los llamados Discursos de Metafísica, pero no los llega a publicar nunca, influenciado por la opinión del filósofo Arnauld, que al leer parte de ellos se queda horrorizado. En la sección V y VI de esos discursos, Leibnitz discute una cuestión crucial: ¿Cómo podemos distinguir un mundo que puede ser

explicado mediante la ciencia de otro que no puede ser explicado? ¿Cómo saber que un determinado fenómeno que nos rodea obedece a una ley y cómo saber que se debe a una conducta errática y aleatoria?

Leibnitz usa el ejemplo de las manchas de tinta en un folio. Si alguien mancha un folio con salpicaduras de tinta y determina un conjunto de puntos en la página, incluso aunque los puntos puedan ser descritos de manera completamente precisa en términos matemáticos (usando, por ejemplo, interpolación de Lagrange) eso no quiere decir que obedezcan ni mucho menos a una determinada ley. Es más, Leibnitz razona que la curva obtenida sería tan difícil de estudiar como el propio conjunto de puntos, con lo que el modelo que pudiéramos obtener carece por completo de interés, puesto que no ayuda a simplificar el fenómeno observado.

Hay varios conceptos aquí que actualmente tienen mucho que ver con la teoría de la computación. El primero de ellos es lo que se ha mencionado como 'difícil de estudiar'. En teoría de la computación la palabra más correcta sería 'complejo'. Leibnitz, de manera implícita, está diciendo que un problema (las manchas de tinta) es complejo cuando cualquier explicación a ese problema no simplifica el problema original.

El otro asunto es lo referente a ley y caos. Para Leibnitz algo es aleatorio cuando no admite manera sencilla de ser explicado, de manera similar a como se entiende en teoría de la información algorítmica. Puede que sea aleatorio de verdad o puede que no, pero eso encaja bien en la definición porque desde nuestro punto de vista ambos fenómenos resultan, no sólo igualmente incomprensibles, indistinguibles entre sí.

Lo más importante de todo esto es que Leibnitz no empleó ninguna máquina para llegar a tales conclusiones, sino que estudió la resolución de problemas desde un punto de vista abstracto. Más tarde él mismo fabricó una de las primeras calculadoras existentes, además de ser uno de los primeros científicos en comprender la importancia de la aritmética binaria. Pero el caso es que esas ideas están más allá de las máquinas empleadas para comprobarlas.

Muchísimo más tarde de estas ideas, el matemático Alan Turing publicó un artículo en 1936 en el que describía un aparato extremadamente sencillo llamado la Máquina de Turing. La Máquina de Turing consiste, sin entrar en muchos detalles, en una cinta de longitud infinita en la que es posible escribir y sobrescribir caracteres de un determinado alfabeto, borrar un carácter (que no es más que sobrescribir el carácter 'en blanco') y avanzar o retroceder a lo largo de la cinta. Es fácil comprobar que versiones más complicadas de la Máquina de Turing, como aquellas que emplean más de una cinta, son completamente equivalentes a la

Máquina de Turing original, también conocida como Máquina de Turing Universal. Sin embargo el potencial de la Máquina de Turing va mucho más allá de eso, y su propio creador postuló que cualquier problema que pudiera ser considerado en términos algorítmicos podría ser resuelto utilizando una máquina de Turing, con independencia del tiempo que tardara en hacerlo. Esta afirmación es conocida como la Tesis de Turing-Church.

La Tesis de Turing-Church conecta dos conceptos distintos en computación. Esos conceptos son la manera abstracta de resolver un problema y la manera concreta de resolver un problema. Un algoritmo es una manera de resolver un problema que puede tener distintas variantes según los medios reales de los que dispongamos. Si un algoritmo dice que para hacer una tarta tengo que ir primero al supermercado de las afueras a comprar los ingredientes, el problema se resolverá de distinta manera si disponemos de coche que si tenemos que viajar en transporte público. Por eso no hay que confundir máquina con algoritmo. Sin embargo, lo que Turing y Church afirman es que la Máquina de Turing Universal es la más general de todas las máquinas.

Poco después de esto, John Von Neumann aportó las primeras ideas para la construcción de un ordenador tal y como lo conocemos. Von Neumann pensó en la electrónica como la herramienta indispensable para la construcción del mismo, y de ese modo no tardó en aparecer el chip, que permite codificar ceros o unos según pase o no la corriente.

En este punto es donde entra de manera fundamental la computación cuántica. Porque es en este momento cuando está en nuestra mano llevar a cabo una elección. Von Neumann eligió la electrónica como la teoría física en la que sustentar un ordenador, tal vez del mismo modo en que Leibnitz hubiera hecho de vivir en su misma época. Pero mucho antes de ellos ya existían los ábacos, y un ábaco no es ni más ni menos que un ordenador construido con piedras. Es un dispositivo muy limitado, pero no hay duda de que tiene tanto derecho a ser considerado una máquina de computación como cualquiera de las modernas.

Mucha gente cree erróneamente que el futuro de los ordenadores pasa por construir dispositivos cada vez más rápidos y con mayor capacidad. Si bien es cierto que durante muchos años eso ha permitido grandes avances en el terreno de la informática, se está empezando a atravesar una crisis al respecto. A medida que los componentes electrónicos son más pequeños, empiezan a aparecer problemas relacionados con la transición de la física clásica a la física atómica.

En el año 1965 apareció la llamada Ley de Moore, que postula que el potencial de las computadoras es duplicado, en términos redondos, aproximadamente una

vez cada dos años. Esta ley, sin embargo, puede tocar techo si las limitaciones físicas de los aparatos impiden el crecimiento.

El asunto esencial es que la computación no tiene por qué estar restringida a la concepción actual que se tiene de los ordenadores. Tal vez si en la época de Von Neumann se hubiera dispuesto de un gran conocimiento en, digamos, Mecánica de Fluidos, se hubieran podido construir ordenadores con dispositivos de agua que resolvían los algoritmos con tanta eficiencia como el más actual de los ordenadores. Cuanto más compleja sea la teoría, más posibilidades computacionales ofrece, y no cabe duda de que la Mecánica de Fluidos es una teoría tan legítima como cualquier otra, teniendo como tiene, por ejemplo, las ecuaciones de Navier-Stokes como problema aún abierto.

La premisa básica para el avance real de la computación es la siguiente: Que nuestros computadores actuales no puedan luchar contra un problema no quiere decir que éste no pueda resolverse más deprisa. Ninguna de las computadoras actuales puede simular siquiera con un grado de proximidad muy cercano a la realidad el derramamiento del líquido de un vaso, y sin embargo la naturaleza lo hace constantemente y no tarda una cantidad exponencial de tiempo en hacerlo. Al construir un ordenador tomamos una decisión. Es, a mayor escala, algo parecido a la evolución de los relojes. En el pasado existían relojes de arena, relojes de Sol e incluso relojes de agua. Con el paso del tiempo se inventaron los relojes analógicos, posibles gracias a la mecánica clásica, y la electrónica aportó los relojes digitales. Finalmente, los relojes más perfectos fabricados en la actualidad son los llamados relojes atómicos, y su perfección reside en que el fenómeno más preciso encontrado hasta la fecha para medir el tiempo reside en la física atómica y los estados de excitación de un electrón.

Es por eso que el próximo salto computacional debe ser más cualitativo que cuantitativo. Por otro lado, el uso de la mecánica cuántica como soporte físico a los ordenadores del futuro es una elección nada arbitraria. No sólo porque la mecánica cuántica es una parte de la física que siempre ha resultado muy difícil de simular en ordenadores clásicos. La mecánica cuántica va, como se ha mencionado antes, en contra de la intuición, debido a que en ella aparecen nuevos fenómenos físicos que apenas estamos empezando a entender. Estos fenómenos están ya aportando nuevas y revolucionarias maneras de abordar la resolución de problemas. Muchos investigadores creen de hecho que ninguna mejora en los computadores actuales, por espectacular que sea, podría jamás equipararse a la mejora producida por el salto de ordenadores clásicos a ordenadores cuánticos.

Se puede pensar que en el fondo, por poderoso que sea un ordenador cuántico, el estudio teórico de los límites del mismo no escapa de los resultados que

ya existían antes de su mera concepción. Esta idea se resume con los términos de eficiente e ineficiente. Un algoritmo es eficiente cuando su velocidad de cálculo crece de manera polinomial con respecto a los datos o menor aún. De otro modo un algoritmo es considerado ineficiente. Lo que se postuló es que la Máquina de Turing es en sí misma un simulador perfecto de eficiencia de cualquier ordenador que pudiera ser fabricado en el futuro, esto es, que si un determinado problema puede ser resuelto de manera eficiente en un determinado ordenador, entonces puede ser resuelto eficientemente en una Máquina de Turing. Ésta es la Tesis Fuerte de Turing-Church:

Cualquier proceso algorítmico puede ser simulado eficientemente usando una Máquina de Turing.

Si esta tesis es correcta, implica que con independencia de la máquina que se esté empleando, para analizar la eficiencia de un problema basta con que nos limitemos a estudiarlo en la Máquina de Turing.

Esta Tesis fue puesta en entredicho no mucho más tarde con la aparición de los ordenadores analógicos, que parecían resolver problemas sin solución eficiente en una Máquina de Turing. Sin embargo la presencia de elementos perturbadores (ruido) limitaban su eficiencia y los convertían en igual de eficientes (e ineficientes) que una Máquina de Turing. Las perturbaciones reales, por tanto, deben ser tenidas en cuenta para evaluar la eficiencia de un modelo computacional, una prueba que la computación cuántica ha logrado superar con teoremas relativos a la capacidad de ruido de los canales cuánticos, así como teoremas relativos a la corrección necesaria para evitar errores de transmisión.

Otro problema con el que esta tesis se encontró fue con la invención de los algoritmos aleatorios, que emplean el azar como parte esencial de su código. Estos algoritmos muchas veces ofrecen una respuesta condicionada, por ejemplo, pueden decir si un número es compuesto con certeza pero sólo pueden decir si es primo con un margen de probabilidad. En cualquier caso, al ofrecer una mejor solución que ningún algoritmo antes conocido, parecían violar la Tesis Fuerte de Turing-Church. Lo que se hizo entonces fue introducir el concepto de la Máquina de Turing Probabilística, y la modificación del enunciado de la Tesis:

Cualquier proceso algorítmico puede ser simulado eficientemente usando una Máquina de Turing Probabilística.

El último intento de modificar esta Tesis vino de la mano de David Deutsch en 1985. Deutsch pensaba que las leyes de la física podían ser usadas para encontrar un enunciado más fuerte de esta Tesis, en clara analogía con la naturaleza como

una gran computadora viviente. Esta idea, de hecho, ya fue empleada en el terreno de la ciencia ficción dos años antes en el libro Guía del Autoestopista Galáctico de Douglas Adams:

-No hablo sino del ordenador que me sucederá. Un ordenador cuyos parámetros funcionales no soy digno de calcular, y sin embargo yo lo proyectaré para vosotros. Un ordenador que podrá calcular la Pregunta de la Respuesta Última, un ordenador de tan infinita y sutil complejidad, que la misma vida orgánica formará parte de su matriz funcional. ¡Y hasta vosotros adoptaréis formas nuevas para introducirlos en el ordenador y conducir su programa de diez millones de años! ¡Sí! Os proyectaré ese ordenador. Y también le daré un nombre. Se llamará la Tierra.

Las expectativas de Deutsch eran que la reformulación de la tesis sería tan fuerte como potente fuera la teoría física empleada. De hecho, su idea específica era que ese dispositivo fuera capaz de simular cualquier proceso físico existente. Debido a que la mecánica cuántica se encuentra en la raíz de todas las demás teorías físicas, Deutsch eligió esta rama para considerar esos nuevos dispositivos. Este concepto es el que dio origen a toda la computación cuántica actual.

Aún no se sabe si este Computador Cuántico Universal será tan potente en términos conceptuales como para ser el tope de la capacidad computacional. Podría suceder que existan en el futuro variantes de las teorías existentes o teorías aún muy jóvenes, como la Teoría de Cuerdas, que vayan incluso más allá, otorgando la posibilidad de un modelo más fuerte de computación.

Una vez empezó a poner en marcha sus ideas, lo primero que se planteó Deutsch era si ese computador cuántico que había postulado sería capaz de resolver eficientemente problemas que se resistían a serlo en una Máquina de Turing. De hecho, dando un paso más, se planteó si sería capaz de resolver problemas que eran ineficientes en una Máquina de Turing Probabilística. Lo que hizo fue construir un simple ejemplo que sugería que, en efecto, el potencial de un computador cuántico excedía el de los computadores clásicos.

Muchos investigadores dedicaron sus esfuerzos a mejorar este primer paso de Deutsch, hasta que en 1994 Peter Shor sorprendió al mundo de la computación con un algoritmo, llamado algoritmo de Shor en su nombre, que resolvía de manera eficiente en un computador cuántico el problema de encontrar los factores primos de un número entero, un problema que hasta la fecha se resistía a ser resuelto con eficiencia (es importante no confundir el problema de encontrar los factores primos de un número entero con el problema de saber si un cierto número entero es o no primo. Para este segundo problema sí que existe un algoritmo eficiente, el

AKS, descubierto en 1992 por M. Agrawal, N. Kayal y N. Saxena , y que se puede leer en [7]). Además, Peter Shor propuso un nuevo algoritmo para el problema del logaritmo discreto que lo resolvía también de manera eficiente, otro de los grandes problemas para los que no se pensaba que pudiera llegarse a encontrar dicho algoritmo.

No sólo estos resultados ponían de manifiesto la importancia que pueden llegar a tener los computadores cuánticos, los ordenadores cuánticos también resultaban de gran interés para una idea que ya había sugerido Richard Feynman en 1982. Este científico había hecho notar que pese a que la base teórica de los ordenadores cuánticos estaba lo bastante avanzada como para poder comprenderlos con gran precisión, los ordenadores clásicos eran incapaces de simular el comportamiento de un ordenador cuántico. De ese modo concluyó que posiblemente la única manera de emplear las ideas de computación cuántica en su totalidad era construir ordenadores que se basaran en principios cuánticos.

Del mismo modo que con el ejemplo de eficiencia de Deutsch, esta afirmación acerca de la simulación en ordenadores cuánticos fue el pistoletazo de salida para muchos investigadores que en los noventa mostraron que es posible emplear computadores cuánticos para simular de manera eficiente sistemas que no poseen una simulación eficiente en computadores clásicos. De este modo, tal y como Deutsch imaginaba, parece que una de las mayores aplicaciones de los computadores cuánticos será la simulación de sistemas de mecánica cuántica demasiado complicados para un ordenador clásico, un problema con profundas consecuencias científicas y tecnológicas.

Por desgracia una de las dificultades que tendrá que pasar el mundo de la programación cuántica será nuestra propia adaptación al mismo. Como ya se apuntó antes, al estar nuestra intuición anclada en el mundo clásico, será mucho más difícil para nosotros aprovecharnos de las características únicas que nos ofrecerá un ordenador cuántico a la hora de programar. Y por otro lado, parece lógico que si bien podemos diseñar algoritmos que se parezcan a aquellos que ya conocemos, el verdadero avance debe venir del uso de nuevas reglas antes imposibles de emplear debido a las limitaciones de la máquina que empleábamos. Además, la exigencia sobre un algoritmo de un ordenador cuántico es grande; no sólo debe funcionar bien, debe funcionar mejor que cualquier algoritmo clásico existente empleado para resolver el mismo problema.

2.3. La información cuántica

Ahora nos detendremos en el campo de la teoría de la información, otro campo muy emparentado con la criptografía. El avance de la teoría de la información y la teoría de la computación ha avanzado de manera paralela casi desde su mismo nacimiento, en tanto que la teoría de la información supone el contexto matemático en el que se basa el estudio de la computación.

La teoría de la información trata de comprender el escurridizo concepto de comunicación entre dos fuentes. Por ejemplo, es posible mandar un mensaje de móvil con todos los caracteres a otro conocido y así tener garantía de que el mensaje será comprendido cuando llegue a su destino. Pero por otro lado, puede ocurrir que la limitación de caracteres a emplear nos obligue a simplificar algunos de ellos en detrimento de otros y, sin embargo, la comunicación sea igualmente posible. En este ejemplo aparecen dos ideas centrales. Una de ellas es que aunque el mensaje enviado es distinto, la información enviada es la misma en ambos casos. La otra es que en la mayor parte de los casos realistas tenemos que enfrentarnos a limitaciones en el canal que empleamos para comunicarnos.

Esta idea que la mayoría de nosotros ha usado de manera implícita fue la que Shannon formuló en 1948. Shannon logró dar una definición matemática de información, una definición tan bien escogida que ha demostrado ser muy superior a otras variantes empleadas sobre el mismo tema. Su definición, de hecho, ha dado lugar a una teoría extremadamente útil y rica, además de aplicable a gran cantidad de problemas de comunicación de la vida real.

Como suele ocurrir con los fundadores de una nueva teoría, Shannon se planteó dos preguntas clave acerca de la misma. La primera era qué recursos son necesarios para enviar información a lo largo de un canal de comunicaciones. La segunda era, conociendo la cantidad de ruido (interferencias o también pérdida de información) de un canal, ¿es posible enviar información de modo que no se vea afectada por el ruido de ese canal?

El propio Shannon respondió a estas preguntas con los dos teoremas fundamentales de la teoría de la información. El primero de ellos cuantifica los recursos físicos necesarios para que pueda enviarse información desde una fuente. El segundo teorema cuantifica cuánta información puede llegar de manera fiable a través de un canal de comunicaciones ruidoso, empleando los códigos correctores de errores, que a grandes rasgos lo que hacen es introducir redundancia suficiente como para garantizar que la mayor parte del ruido sólo destruirá piezas de información recuperables al ser recibida la información restante. Se ha trabajado mucho en este segundo teorema con códigos cada vez más perfectos, variados y

sofisticados, y actualmente son empleados en gran cantidad de dispositivos como cds y dvds, computadoras modernas y sistemas de satélites de comunicaciones.

Uno de los grandes logros de la teoría de la información cuántica es que ha avanzado con grandes pasos para resolver problemas importantes como estos. En el año 1995 Ben Schumacher encontró un teorema análogo al teorema de Shannon acerca de los recursos necesarios para enviar información cuántica, mostrando así que en efecto es posible comunicación real a lo largo de un canal de tales características. En el camino, además, se definió el concepto crucial del qubit, que se usará repetidamente en las secciones siguientes, como un análogo cuántico del bit de información, como una fuente real de información. Sin embargo aún no se ha encontrado un análogo completo del segundo teorema relativo a canales ruidosos, aunque sí se han encontrado códigos correctores de errores que permiten a un ordenador cuántico transmitir información a través de canales ruidosos, un bache que los ordenadores analógicos no pudieron superar.

En lo relativo a posibles mejoras que la teoría de la información cuántica puede ofrecer en comparación a la clásica, uno de los descubrimientos más curiosos se efectuó al responderse a la inevitable pregunta de si un canal cuántico sería capaz de enviar información clásica y con qué grado de eficiencia. C. Bennett y S. Wiesner, dos de los padres fundadores de la criptografía cuántica, no sólo demostraron que un canal cuántico haría ese trabajo sin ninguna clase de problemas, de hecho mostraron que era posible enviar dos bits de información clásica con el envío de un solo qubit, algo que se conoce como codificación superdensa.

Otro resultado que alcanza notables mejoras se refiere a ordenadores cuánticos que funcionan formando redes. Si tenemos dos ordenadores cuánticos conectados que están tratando de resolver un mismo problema, la cantidad de comunicación que debe existir entre ellos para lograrlo es exponencialmente menor que la necesaria en el caso clásico. Sin embargo aún hacen falta muchas mejoras técnicas en este caso, así como un problema concreto y de gran importancia en la actualidad que pueda ilustrar acerca de la capacidad de los ordenadores cuánticos en este área.

2.4. La criptografía cuántica

Finalmente, otra de las nuevas teorías que han aparecido como variante de teorías clásicas es la criptografía cuántica. La criptografía, que no siempre ha sido considerada una rama de las matemáticas y que incluso en sus inicios fue más un arte que una ciencia, es en términos básicos el problema de llevar a cabo comuni-

cación entre dos partes que pueden o no confiar entre ellas. La más famosa pero no la única aplicación de la criptografía es la transmisión de mensajes secretos. En la actualidad esta utilidad no sólo es empleada por motivos militares o para defender secretos de estado, está presente en casi cualquier área de la vida moderna debido al auge de las comunicaciones a distancia. Por ejemplo, a la hora de comprar por Internet en portales como Amazon o Ebay, si el comprador usa el número de la tarjeta de crédito para llevar a cabo la transacción espera que haya un método para hacer llegar el número al vendedor sin que se vea expuesto a ser descubierto por algún observador no deseado. Este es un ejemplo de lo que se conoce como un protocolo criptográfico, una de las subsecciones de la criptografía en la que más ha contribuido la criptografía cuántica.

La criptografía clásica moderna se divide en dos grandes partes, la criptografía de clave privada y la criptografía de clave pública. La primera es conocida por la mayoría de la gente. En ella, sólo existe una clave, que es empleada tanto para cifrar como para descifrar el mensaje. Esa clave, por tanto, es compartida tanto por quien envía el mensaje (a quien se suele llamar Alice) como por quien lo recibe (a quien se suele llamar Bob). Esta ha sido la manera usual en la que se han empleado claves a lo largo de siglos, y tiene el crucial defecto de que a veces distribuir las claves es tan difícil como enviar el mensaje de manera segura. Eso se debe a que si un intruso (llamado Eve) desea interceptar el mensaje, entonces le basta con interceptar las claves cuando sean compartidas y luego emplearlas para leer los mensajes que se envíen Alice y Bob.

Sin embargo, la criptografía de clave pública solventa este problema de una manera sorprendente. En este procedimiento, cada uno de los comunicantes posee dos claves, una de ellas pública y la otra privada. Todo el mundo (incluso Eve) conoce la clave pública de Alice, pero sólo Alice conoce su propia clave privada. Si Bob desea enviar un mensaje a Alice, lo que debe hacer es usar la clave pública de Alice para encriptar ese mensaje. La cuestión clave de este sistema es que una vez que Bob ha encriptado un mensaje con dicha clave sólo Alice, con su clave privada, puede desencriptarlo. En un principio puede parecer que hay trampa en la idea, ya que si Bob ha empleado un procedimiento para construir un conjunto de caracteres a partir de otro lo único que parece que debería hacer es repetir ese proceso a la inversa para volver de nuevo al punto de partida, y puesto que la clave pública de Alice es conocida por todos cualquiera podría imitar ese procedimiento. La cuestión es que los procesos de encriptado y desencriptado no son simétricos. Mientras que uno de ellos es muy sencillo, el otro resulta extremadamente complicado salvo que se posea la clave privada de Alice. Estos procesos son contruidos basándose en problemas muy sencillos de resolver en un sentido

pero extremadamente complicados en el otro, como el problema de encontrar los factores primos de un número entero. Mientras que, dados los factores primos de un número entero, encontrar el propio número es un proceso trivial (basta multiplicarlos todos entre sí, ya que suponemos que también poseemos la multiplicidad de cada factor) el problema inverso, hallar los factores dado el número, es bastante más complicado. No es imposible, por supuesto, pero ahí está otra de las ventajas de la criptografía de clave pública: dado que cambiar de clave es muy sencillo para Alice (basta con que genere una nueva clave pública, otra privada, y anuncie la primera), cambiar las claves cada cierto tiempo garantiza que para cuando la clave privada de Alice sea obtenida ésta ya no será de ninguna utilidad.

El ejemplo de los factores primos no es casual. La mayoría de los sistemas de clave pública que se emplean en la actualidad se basan en ese problema. El más conocido es el RSA, que hoy en día sigue siendo muy usado, aunque hay otros sistemas, como los protocolos de Diffie-Hellman, los primeros científicos en llamar la atención sobre este método de encriptación, que se basan en el conocido problema del logaritmo discreto. Como ya se ha comentado, estos problemas no son imposibles de resolver, pero sí muy ineficientes. Esto es, que cuanto mayores sean los datos de entrada el tiempo de ejecución de la computadora crece a velocidades cada vez más insostenibles (en general exponenciales, es decir, el tiempo de ejecución crece con respecto a los datos a la misma velocidad que una colonia de bacterias aumenta con el paso de los segundos, haciendo que haya datos iniciales no muy elevados para los que resolver estos problemas nos llevara tanto como la edad estimada del Universo).

Pero la computación cuántica ha abierto una brecha enorme en estos procedimientos, pues el mencionado algoritmo de Shor, al ser capaz de encontrar los factores primos de un número de manera eficiente, implica que en el momento en que los ordenadores cuánticos entraran en escena este problema ya no sería nunca más un problema intratable, y por tanto el RSA sería fulminantemente derrotado. Hechos como este han hecho que muchos gobiernos se interesen en los avances de computación cuántica, ya que pueden llegar a suponer un asunto de seguridad nacional. Del mismo modo, la fiabilidad de otros sistemas de clave pública está en entredicho, pues puede ser bastante probable que los problemas en los que se base su seguridad sean resueltos de manera eficiente en un computador cuántico del mismo modo que lo ha sido el problema de los factores primos de un número entero.

Sin embargo no es en esto propiamente dicho en lo que consiste la criptografía cuántica. La criptografía cuántica es más profunda que simplemente usar

los conocimientos de computación cuántica para dinamitar los avances de criptografía realizados hasta la fecha.

Como apuntan Shor y Preskill en el artículo [10], la criptografía cuántica difiere de la clásica en el hecho de que los datos permanecen secretos gracias a las peculiares propiedades de la mecánica cuántica en lugar de la aparente dificultad de computar determinados problemas. Esto aporta la ventaja de que podremos estar tan seguros de la eficacia de un protocolo de criptografía cuántica como admitidas y verificadas estén las leyes de mecánica cuántica en las que se fundamenta.

Uno de los primeros avances de criptografía cuántica es el conocido protocolo BB84. Este protocolo, que posee numerosas variantes, se fundamenta en un principio bien conocido a nivel divulgativo, el principio de incertidumbre de Heisenberg. La idea básica, sin entrar en detalles técnicos, es que si Alice envía qubits a Bob y Eve los intercepta y los observa por el camino, entonces, debido a dicho principio, Eve habrá forzado a los qubits a aparecer en un cierto estado, lo que puede ser explotado en su contra para detectar su presencia. El protocolo BB84 permite por tanto saber con certeza total a nivel teórico y casi total a nivel práctico cuándo una transmisión entre Alice y Bob ha sido espiada. Este protocolo ha sido propuesto como protocolo para intercambio seguro de claves entre Alice y Bob, lo que puede traer de vuelta la criptografía de clave privada ya que solventa uno de sus mayores problemas. Asimismo es el primero de una serie de protocolos llamados protocolos QKD (quantum key distribution).

Las primeras ideas sobre criptografía cuántica fueron propuestas por Wiesner en los años sesenta. En concreto, Wiesner había tenido la idea de usar las leyes de la mecánica cuántica para diseñar billetes de banco que resultaran imposibles de falsificar. Sin embargo sus ideas fueron rechazadas para publicación y quedaron olvidadas hasta que Bennett, que conocía a Wiesner desde que eran estudiantes, se asoció con G. Brassard y trataron de llevar adelante la idea de su amigo, obteniendo como resultado el germen de lo que luego sería el protocolo BB84.

Otro asunto muy importante relativo a la criptografía cuántica es que mientras que la computación cuántica aún tiene que demostrar todo su potencial con la construcción real de un computador, la criptografía cuántica ya ha sido empleada con éxito en un experimento realizado por IBM en el año 2000, en el que se empleaban cables de fibra óptica y emisor y receptor estaban a diez kilómetros de distancia.

Sin embargo aún hay mucho que investigar al respecto. La criptografía clásica se encarga de resolver muchos problemas de seguridad distintos, y su versión cuántica no debería ser menos. Entre los problemas de seguridad a resolver, al-

gunos de los cuales se tratarán más tarde y otros han sido ya mencionados, se encuentran:

1. Seguridad física.

Posiblemente uno de los puntos más débiles de la criptografía cuántica. Aquí se engloban todos los asuntos relacionados con el mantenimiento de los soportes físicos de la información. Aquí también se encuentran las medidas contra accidentes como incendios o sobrecargas eléctricas, pero también la prevención de ataques o de manipulación del envío de datos. En el caso de la criptografía cuántica eso puede suceder alterando los dispositivos que manejan qubits (como pueden ser los láseres que polarizan un fotón, una de las maneras que se barajan para manejar qubits de manera práctica).

2. Seguridad de la información.

Preservación de la información contra observadores no autorizados. Los procedimientos pueden cambiar según el computador esté aislado y se quiera acceder desde el exterior o esté en una red de trabajo y se quiera acceder desde otro computador.

3. Seguridad del canal de comunicación.

En general no se puede considerar nunca que un canal es seguro. Rara vez, de hecho, están bajo el control del usuario, ya que suelen pertenecer a terceros. En criptografía clásica resulta muy difícil, si no imposible, asegurarse de que un canal es seguro, en criptografía cuántica existe al menos la posibilidad de diseñar protocolos que permitan un cierto grado de fiabilidad.

4. Problemas de autenticación.

Debido a los potenciales problemas de los canales de comunicación es conveniente asegurarse de que la información que se recibe viene realmente de quien dice enviarla, y que además no ha sido alterada. En criptografía clásica esto suele conseguirse con las llamadas funciones hash o resumen.

5. No repudio.

No sólo es necesario poder identificar al remitente de un mensaje, a veces también es importante que el remitente no pueda negar la responsabilidad de haberlo enviado. Es necesario entonces impedir que un emisor pueda repudiar un mensaje, es decir, negar su autoría sobre él.

6. Anonimato.

En cierto modo es el completo opuesto del anterior. En encuestas o votaciones cada ciudadano tiene derecho a expresar su opinión y tener garantías de poder preservar su intimidad. Por otro lado existen situaciones como la denuncia de violaciones contra los derechos humanos, sobre todo las perpetradas por países con fuertes sistemas dictatoriales, en los que esta aplicación no sólo es necesaria sino que puede resultar vital para aquel que recurre a ella.

7. Poker por teléfono.

Realizar elecciones al azar a distancia sin que ninguno de los involucrados trate de falsificar el resultado es fundamental en muchos protocolos, así como en transacciones o simplemente por motivos de ocio.

3. Conceptos básicos

3.1. El qubit

El concepto esencial de la computación cuántica es el llamado qubit. La palabra qubit es una abreviatura de quantum bit, y como es lógico a partir de su nombre, su papel es análogo al del propio bit en computación clásica, es decir, la unidad básica de información en esta clase de ordenadores.

El qubit, como el bit, es antes de nada un objeto matemático. Aunque sea utilizado a partir de objetos físicos en un contexto real, el asunto de pensar en él como una entidad abstracta nos otorga la ventaja de que, con independencia de la manera que empleemos para utilizar un qubit, ya sea mediante los estados de polarización de un fotón o el nivel de excitación de un electrón, nuestro estudio será el mismo, y la teoría en torno a él no variará.

Del mismo modo en que un bit posee un estado, 0 o 1, un qubit también posee estados. Dos estados análogos para un qubit son los estados $|0\rangle$ y $|1\rangle$. La notación empleada es la notación de Dirac, y más tarde se explica con un poco más de detalle. La diferencia principal entre un bit y un qubit es que éste último puede además existir en una superposición de estos dos estados, lo que se escribe por medio de combinaciones lineales:

$$|\psi\rangle = \alpha|0\rangle + \beta|1\rangle.$$

Los números α y β son números complejos, pero en la mayoría de los casos pueden suponerse como reales. Desde un punto de vista algebraico, un qubit es un elemento de un espacio vectorial de dimensión dos, donde una base de ese espacio son los estados $|0\rangle$ y $|1\rangle$. El hecho de que se pueda definir un producto escalar de manera natural otorga de hecho a este espacio la estructura de espacio de Hilbert. Dicho espacio queda definido si consideramos que los estados $|0\rangle$ y $|1\rangle$ forman no sólo una base, sino una base ortonormal. De ese modo el producto escalar de dos elementos cualesquiera de este espacio queda completamente definido a partir de sus coordenadas.

Otra diferencia con los bits es que, si bien podemos examinar éstos para saber si valen 0 ó 1, eso no ocurre con un qubit. No sólo se trata de que no podamos obtener con seguridad los valores α y β , de hecho la información que obtenemos es muy reducida. Cuando medimos un qubit obtenemos o bien el resultado 0 con probabilidad $|\alpha|^2$ o bien el resultado 1 con probabilidad $|\beta|^2$. Eso obliga, evidentemente, a que se cumpla la identidad $|\alpha|^2 + |\beta|^2 = 1$, ya que las probabilidades deben sumar uno. Eso quiere decir que los qubits son los elementos de nuestro espacio de Hilbert que poseen norma uno.

3.2. Múltiples qubits

La manera de trabajar con sistemas de múltiples qubits se obtiene de nuevo por analogía con el caso de los bits. Cuando tenemos dos bits tenemos cuatro posibles estados, es decir, 00, 01, 10 y 11. También podemos tener esos dos estados en el caso de los qubits, y usando la notación de Dirac, se corresponderían con $|00\rangle$, $|01\rangle$, $|10\rangle$ y $|11\rangle$. Pero es que además ahora, nuevamente, podemos estar en una superposición de estos estados, lo que se modeliza de nuevo como una combinación lineal de los mismos:

$$|\psi\rangle = \alpha_{00}|00\rangle + \alpha_{01}|01\rangle + \alpha_{10}|10\rangle + \alpha_{11}|11\rangle.$$

Como en el caso de un qubit, al medir este sistema de dos qubits los coeficientes son las probabilidades respectivas de obtener el estado al que acompañan, y por ese motivo la suma de todas ellas es de nuevo igual a uno, lo que se puede entender como que este vector, que ahora pertenece a un espacio vectorial de dimensión cuatro, tiene norma uno en ese mismo subespacio.

La herramienta matemática para modelizar un sistema de múltiples qubits es el producto tensorial entre espacios vectoriales, pero dado que no la usaremos más allá de la explicación previa no ahondaremos mucho más en ello. Basta decir

que en el caso general de un sistema de n qubits existen 2^n amplitudes para cada elemento del sistema, una para cada estado de la forma $|x_1 x_2 \dots x_n\rangle$. Un sistema de tales características es tan grande que ningún ordenador clásico podría manejar tal cantidad de datos en la actualidad (para $n = 500$ este número es más grande que el número de átomos que se estima hay en el Universo).

Uno de los estados de dos qubits más importantes que existen, y que tiene un papel clave en el protocolo que vamos a describir, es el llamado *par EPR*

$$\frac{|00\rangle + |11\rangle}{\sqrt{2}}.$$

Este estado, llamado EPR en honor a Einstein, Podolsky y Rosen, es la clave para lo que se conoce como teleportación cuántica. La idea principal es que al medir el primer qubit obtendremos como resultado de medida o bien el valor 0 con probabilidad $1/2$ o bien el valor 1 con probabilidad $1/2$, pero una vez lo hayamos hecho la medida del segundo qubit siempre nos dará el mismo resultado que la medida del primer qubit. Se suele decir que los dos estados están enlazados, y en términos del producto tensorial eso queda claro, ya que ocurre que este sistema de dos qubits no puede ser escrito como el producto tensorial de dos qubits separados, no como, por ejemplo, en el caso del sistema

$$\frac{|00\rangle + |01\rangle}{\sqrt{2}} = |0\rangle \otimes \frac{|0\rangle + |1\rangle}{\sqrt{2}}.$$

Ésta es, de hecho, la manera rigurosa de definir el enlazamiento. El par EPR también es usado en algoritmos, y es la base central de los algoritmos más importantes de computación cuántica, como el de Deutsch, el de Grover y el de Shor.

3.3. Notación de Dirac

La notación de Dirac, muy empleada en el contexto de mecánica cuántica por matemáticos y físicos por igual, no es más que una reescritura de términos básicos del álgebra lineal sobre espacios de dimensión finita.

La notación referente a matrices es prácticamente idéntica a la empleada en álgebra lineal: al conjugado de un número complejo z se le denota z^* , como es usual. También se denota la traspuesta de una matriz como A^T y su conjugada compleja como A^* , que no es más que la matriz cuyos elementos son los conjugados de los de la matriz A . Por último se denota como la matriz hermítica de A a $A^\dagger$, que no es más que $(A^T)^*$.

La notación referente a vectores, como ya se ha visto antes, cambia sustancialmente. Por $|\psi\rangle$ denotaremos a un vector arbitrario, que ya hemos visto siempre se puede escribir como combinación lineal de $|0\rangle$ y $|1\rangle$. Este vector a veces es conocido también como *ket*. El dual de un vector arbitrario se denota como $\langle\psi|$, y es también conocido como *bra*. El producto escalar interno entre los vectores $|\varphi\rangle$ y $|\psi\rangle$ se denota como $\langle\varphi|\psi\rangle$, y es equivalente al producto entre $|\psi\rangle$ y el dual de $|\varphi\rangle$. Este es el motivo de elección de los nombres anteriores, pues el producto escalar interno, con esta notación, se leería *braket*, lo que recuerda a la palabra inglesa bracket (paréntesis).

Aunque no lo usaremos, no está de más decir que el producto tensorial de $|\varphi\rangle$ y $|\psi\rangle$ se denota por $|\varphi\rangle \otimes |\psi\rangle$, o de manera abreviada por $|\varphi\rangle|\psi\rangle$. El producto tensorial de los vectores de la base tiene una notación más abreviada aún, como se vio anteriormente, pues se puede escribir como

$$|i\rangle \otimes |j\rangle = |ij\rangle, i, j \in \{0, 1\}$$

Pór último y no menos importante denotaremos el producto interno de $|\varphi\rangle$ y $A|\psi\rangle$, o equivalentemente el producto interno de $A^\dagger|\varphi\rangle$ y $|\psi\rangle$, como $\langle\varphi|A|\psi\rangle$.

3.4. Puertas cuánticas

La nueva concepción de información en un ordenador cuántico trae también consigo una nueva manera de entender sus procesos internos. Los ordenadores clásicos emplean lo que se conoce como puertas lógicas. Las puertas lógicas son operaciones que efectuamos a la información que introducimos en el ordenador para manipularla y convertirla en aquello que nos interesa. Para ilustrar la diferencia crucial entre puertas cuánticas y lógicas (o clásicas), podemos pensar en puertas lógicas que reciben un bit de entrada y devuelven otro bit de salida. La única puerta lógica no trivial es la conocida como NOT, cuya operación consiste sencillamente en devolver el valor cero cuando ha recibido el uno y viceversa.

Sin embargo, cuando pensamos en una puerta cuántica que recibe un qubit de entrada y devuelve un qubit de salida, las posibilidades son mucho mayores. Para empezar, podemos definir un análogo a la puerta NOT para los qubits, pero hay que hacerlo con cuidado. Ahora no basta con decir lo que ocurre con los estados $|0\rangle$ y $|1\rangle$, pues estamos olvidando que estos estados están superpuestos entre sí gracias a unos coeficientes que determinan la probabilidad de que los obtengamos al medirlos. De hecho esos coeficientes son esenciales para definir nuestro NOT cuántico. Lo que haremos será convertir el qubit

$$\alpha|0\rangle + \beta|1\rangle$$

en un qubit donde las probabilidades han sido intercambiadas,

$$\alpha|1\rangle + \beta|0\rangle.$$

Lo más interesante y esperanzador es que si interpretamos estos coeficientes como las coordenadas del qubit, es decir,

$$\begin{bmatrix} \alpha \\ \beta \end{bmatrix}$$

entonces la acción que hemos descrito puede escribirse como

$$X = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$$

ya que ocurre que

$$X \begin{bmatrix} \alpha \\ \beta \end{bmatrix} = \begin{bmatrix} \beta \\ \alpha \end{bmatrix}.$$

Esta es la puerta cuántica equivalente a NOT, conocida también como q-NOT, y para la que se suele emplear la notación X . Como en este ejemplo, todas las puertas cuánticas definidas de qubit a qubit deben actuar de manera lineal, y por tanto pueden ser definidas y expresadas a partir de una matriz. El motivo de esto reside en las leyes de la mecánica cuántica, ya que si se permitiera un comportamiento no lineal entonces daríamos lugar a paradojas tales como permitir comunicación instantánea, viajes en el tiempo y violación de la segunda ley de la termodinámica.

Por otro lado no sirve cualquier matriz U que nos propongamos. Dado que la salida debe ser un qubit válido tal y como los hemos definido, U debe respetar las normas, es decir, que si el vector entrante tiene norma uno, el vector saliente también. Por tanto la matriz debe ser unitaria, que se puede definir como que $U^\dagger U = I$. Evidentemente la puerta q-NOT cumple este requisito.

Es claro también que si estas son las únicas condiciones que se le exigen a una puerta cuántica que recibe y devuelve un qubit, hay muchas otras posibilidades no triviales, de hecho todo un continuo de ellas. En términos de álgebra lineal, hay

tantas posibilidades como isometrías de un espacio vectorial bidimensional en sí mismo. Existen muchas otras puertas cuánticas importantes, como la puerta de Hadamard H , que tiene la propiedad de mezclar los estados de la base y es crucial en la teleportación cuántica y en los algoritmos anteriormente mencionados, pero dado que nosotros no la usaremos no diremos más al respecto.

4. Base del protocolo

En esta parte veremos brevemente los aspectos que giran en torno al protocolo y en los que se apoya para ser diseñado, como por ejemplo que en general dos qubits no pueden ser distinguidos. También veremos la base del método de criptoanálisis conocido como EPR Man in the Middle.

4.1. Indistinguibilidad de qubits

Uno de los aspectos en los que se basan multitud de protocolos de criptografía cuántica, y en concreto el protocolo BB84, el padre de casi todos ellos, es que dado un qubit $|\varphi\rangle$ arbitrario, si nos preguntan si ese qubit es el qubit $|\psi_1\rangle$ o el qubit $|\psi_2\rangle$ (y sabemos con certeza que se trata de uno de los dos), no podremos responder con seguridad a menos que estos dos qubits sean ortogonales, es decir, que $\langle\psi_1|\psi_2\rangle = 0$.

La demostración rigurosa de este hecho es bastante sencilla, pero escribirla con propiedad implicaría tener que emplear notación y postulados que no hemos introducido, como el de las medidas sobre un qubit. La idea esencial es que al no ser ortogonales, la proyección del segundo qubit sobre el primero dará lugar a una parte no nula. De ese modo el segundo qubit puede descomponerse en suma directa de su proyección sobre el primer qubit y su proyección sobre el ortogonal del primer qubit. Por ese motivo siempre habrá una parte común a ambos que no podrá ser distinguida con certeza, y siempre existirá una probabilidad de cometer un error.

Por ejemplo, si tenemos un qubit arbitrario $|\varphi\rangle$ y nos preguntan si se trata del qubit $|0\rangle$ o del qubit $(|0\rangle + |1\rangle)/\sqrt{2}$, si al medirlo obtenemos el resultado 1 sabemos con certeza que el qubit obtenido es el segundo, ya que nunca podemos obtener ese resultado con el primer qubit ya que la probabilidad de obtenerlo es cero; pero si obtenemos el resultado 0 no existe manera de que sepamos con certeza cuál de los dos qubits es el que poseemos. Nuestra mejor probabilidad de

acertar, como se ha demostrado rigurosamente, es igual a $1/2$, es decir, la misma que si eligiéramos al azar sin haber realizado ningún estudio previo.

4.2. No clonación

Otro de los resultados importantes para el buen funcionamiento de los protocolos cuánticos (y de éste en particular) es que no es posible realizar una copia de un estado cuántico que desconozcamos (y ya hemos visto que es fácil no poder discernir de cuál se trata).

Es posible mostrar una demostración elemental de este hecho con la notación que hemos empleado hasta ahora. Si una máquina clonadora existiera, entonces en esencia funcionaría de la siguiente manera: Introduciríamos en ella un sistema de dos qubits donde el primero es el qubit a clonar y el segundo es un qubit a nuestra completa elección,

$$|\psi\rangle \otimes |s\rangle.$$

Aunque no lo hemos detallado, la sucesión de puertas cuánticas, pese a poder ser toda una cadena muy numerosa de ellas y actuar no sobre un qubit sino sobre un sistema de dos, puede ser modelizada por una sola matriz U cuadrada, unitaria y de dimensión cuatro, aunque la dimensión de U no juega ningún papel en la prueba. De ese modo, al actuar U sobre la entrada de nuestra máquina clonadora, obtenemos lo siguiente:

$$|\psi\rangle \otimes |s\rangle \longrightarrow U(|\psi\rangle \otimes |s\rangle) = |\psi\rangle \otimes |\psi\rangle.$$

Supongamos que esta máquina trabaja en dos estados $|\psi\rangle$ y $|\varphi\rangle$. Entonces obtenemos

$$\begin{aligned} U(|\psi\rangle \otimes |s\rangle) &= |\psi\rangle \otimes |\psi\rangle \\ U(|\varphi\rangle \otimes |s\rangle) &= |\varphi\rangle \otimes |\varphi\rangle \end{aligned}$$

Tomando el producto interno de cada lado de las igualdades, y teniendo en cuenta que el producto interno de un producto tensorial se define de manera natural como el de sus respectivos factores, tenemos que

$$\langle\psi|\varphi\rangle = (\langle\psi|\varphi\rangle)^2,$$

y esta ecuación, donde la incógnita es $x = \langle\psi|\varphi\rangle$, sólo deja dos posibilidades: o bien ambos qubits son el mismo ($x = 1$) o bien son ortogonales ($x = 0$). Por

tanto no existe una máquina clonadora capaz de clonar estados no ortogonales entre sí, y por tanto menos aún de clonar estados arbitrarios.

Esto es muy importante porque implica que un potencial espía del protocolo, si captura qubits de información, debe, para poder sacar partido de ellos, utilizarlos directamente, ya que no puede hacer copias en las que experimentar sus métodos.

Puede parecer que el hecho de no poder clonar qubits resulte extraño, pero en realidad es natural si uno se para a pensarlo. De ser posible podríamos, por ejemplo, transmitir información instantánea y por tanto viajar más rápido que la luz. La idea básica es que si tuviéramos una máquina clonadora en nuestro poder entonces, por medio de estadísticas sobre tantas copias de un qubit como quisiéramos, podríamos librarnos del escollo que supone no poder conocer con certeza el valor del mismo y abrir ante nosotros nuevas posibilidades físicas que sencillamente no están ahí. Podríamos, de hecho, distinguir qubits arbitrarios, otra cosa que ya hemos comentado que no es posible. Los dos problemas, de hecho, son equivalentes: Si uno no es posible el otro tampoco y viceversa.

Estudios recientes han demostrado que ni aunque permitiéramos puertas cuánticas no unitarias (lo que ya es de por sí ser muy permisivo en términos físicos) podríamos construir un aparato de tales características. Por otro lado, si permitimos un cierto margen de error los resultados que se han obtenido son que es posible realizar un aparato clonador que nos de una copia de un qubit con una fidelidad del 83 %, aunque experimentalmente la fidelidad obtenida está en torno al 81 %. Este hecho no afecta en general a ningún protocolo cuántico ni tampoco al nuestro, dado que en ellos la fidelidad necesaria para poder romperlos de manera satisfactoria debe ser del 100 %.

4.3. Distinguir entre qubits EPR y no EPR

Una curiosísima propiedad de las puertas cuánticas es que, dado que pueden ser representadas por matrices unitarias, son todas ellas invertibles, es decir, que toda puerta aplicada a un qubit tiene marcha atrás. Por otro lado las puertas cuánticas, tal y como las hemos modelizado, no poseen memoria, es decir, que dado un qubit que ha sido sometido a varias puertas cuánticas, el resultado sólo depende de la matriz U que representa esa puerta cuántica, con lo que su desconocimiento nos impide regresar atrás con certeza. Eso implica que dado un qubit, no nos es posible saber si forma o no parte de un par EPR, ya que los procesos llevados a cabo para que eso suceda se pueden llevar a cabo con puertas cuánticas, como por ejemplo la de Hadamard.

4.4. El ataque EPR Man in the Middle

El ataque llamado Man in The Middle ya existía mucho antes de la criptografía cuántica. Ha sido usado con éxito para, por ejemplo, atacar el protocolo de Diffie-Hellman de intercambio de claves por medio del logaritmo discreto, y su mecánica es tal y como indica su nombre. Si Alice y Bob desean comunicarse con un protocolo determinado, cuando Eve emplea el ataque Man in The Middle intenta insertarse entre ambos y fingir para cada uno de ellos que es el otro interlocutor. Para ello, por supuesto, Eve posee control casi absoluto del canal. Puede capturar información y puede mandar información propia, lo que se traduce en que Eve puede interceptar un mensaje de cualquiera de los dos interlocutores y cambiarlo por uno propio.

En criptografía clásica la presencia de Man in The Middle es detectada mediante procedimientos de autenticación. Sin embargo eso es tanto más complicado cuanto menor sea la cantidad de información previa compartida entre Alice y Bob, alcanzando el caso límite más complejo, que es que no compartan ninguna información previa en absoluto.

En criptografía cuántica el ataque Man in The Middle posee una versión propia, precedida del término EPR, ya que se usa en protocolos donde estos pares de qubits están presentes. Con mayor o menor grado de sofisticación, los protocolos que usan pares EPR (y en general pares o grupos mayores de qubits enlazados) se basan en la idea central de que uno de los dos qubits está en posesión de Alice y el otro es enviado a Bob (o viceversa). De ese modo, al realizar mediciones, están compartiendo información de manera encubierta. Lo que Eve hace para insertarse en este proceso es capturar el qubit EPR de Alice, crear un par EPR propio y enviar uno de los qubits de ese par a Bob. De ese modo es como si Eve estuviera 'conectada' a ambos interlocutores por medio de dos cuerdas, una que la comunica con Alice y la otra con Bob, y además pudiera 'agarrar' la cuerda que ellos comparten para impedir el flujo de información. Lo único que tiene que hacer es manipular los envíos que la convengan y dejar pasar el resto.

En el artículo [1] se habla de diferentes protocolos que son vulnerables a este ataque. En algunos de ellos existe comunicación clásica y pública entre Alice y Bob, entonces el ataque es exitoso si se supone que Eve controla también este canal. En el protocolo que aquí presentamos, esta hipótesis está relajada. Como veremos, supondremos que Eve controla el canal clásico, pero supondremos también que existe un canal público donde Alice y Bob pueden comunicarse, pero al precio de ser escuchados por cualquiera. Hay varias razones que nos mueven a ello.

Una de ellas es para declarar públicamente la presencia de Eve al otro interlocutor en caso de ser detectada. Esto no es estrictamente necesario dado el carácter simétrico del protocolo. Otra es poder comparar bits en el paso final, algo que veremos con más detalle posteriormente.

El canal clásico, que puede estar controlado por Eve, es para declarar resultados de mediciones sobre qubits. Estas mediciones pueden ser alteradas por Eve, pero si lo hace, entonces será altamente probable que su presencia sea detectada a posteriori.

La idea básica del protocolo para impedir un ataque de tipo EPR Man in The Middle es mezclar todos los resultados anteriores al mismo tiempo: si Eve captura un qubit, no puede saber con certeza cuál es el que está en su poder, y tampoco si forma o no parte de un par EPR. Por otro lado, al no poder clonarlo, no puede realizar más que un intento por averiguarlo.

Finalmente destacar que se aprovecha también el hecho de obligar a Eve a hacer elecciones aleatorias, como qué qubit sustituir por el capturado, antes de que esa información le resulte útil. Esta idea es empleada en un protocolo diseñado por Beige, Englert, Kurtsiefer y Weinfurter [2]. Básicamente, se fundamenta en que Eve no sabe qué qubits de los enviados son de control y cuáles transportan información. En nuestro caso la idea es que todos los qubits pueden tener ambos cometidos a la vez y no tomamos una decisión hasta el final del intercambio de qubits, con lo que Eve no puede secuestrar los qubits y permanecer oculta al mismo tiempo.

5. Descripción del protocolo

Por fin describiremos la manera en que funciona el protocolo. Usaremos la siguiente notación:

$$\psi_1 = \frac{\sqrt{3}}{2}|0\rangle + \frac{1}{2}|1\rangle;$$

$$\psi_2 = X(\psi_1) = \frac{1}{2}|0\rangle + \frac{\sqrt{3}}{2}|1\rangle;$$

$$\psi_{3,4} = \frac{|00\rangle + |11\rangle}{\sqrt{2}};$$

$$\psi_3 = \text{primer qubit de } \psi_{3,4};$$

$$\psi_4 = \text{segundo qubit de } \psi_{3,4}.$$

Usaremos como base de medida $|0\rangle, |1\rangle$.

Como ya dijimos, nuestra hipótesis inicial será que existe un canal clásico público donde Alice y Bob (a los que nos referiremos de forma abreviada como A y B) pueden publicar lo que quieran sin miedo a que sea modificado por Eve (E en abreviado), pero cualquiera puede verlo también.

Protocolo:

Paso 0: A y B establecen, cada uno por su cuenta, un número natural n_A, n_B y un valor TOL_A, TOL_B pequeño.

Paso 1: Con probabilidad $1/2$, A envía un qubit que será o no EPR.

- a) Si es EPR, envía ψ_3 y se queda con ψ_4 .
- b) Si no es EPR, elige con probabilidad $1/2$ enviar ψ_1 o ψ_2 .

Paso 2: B recibe el qubit, lo mide y declara a A el resultado de su medida. B envía un qubit a A procediendo como ella en el Paso 1.

Paso 3: A recibe el qubit de B y el resultado de su medida sobre el qubit que ella envió. Anota dicho resultado.

- a) Si A decidió enviar un qubit EPR, mide el qubit ψ_4 que guardó. Si su medida no coincide con la de B, E es detectado. Lo declara en el canal público. Si coincide, la anota.
- b) Si no envió un qubit EPR, mide uno igual al que mandó y lo anota.

A procede con el qubit de B que ha recibido de manera análoga a B en el Paso 2.

Paso 4: Los Pasos 1 a 3 (que llamaremos un ciclo) se repiten tantas veces como qubits se deseen intercambiar. En el ciclo n_A , A cuenta el número de veces que ha enviado ψ_1 a B, e_1 , y el número de veces que ha enviado ψ_2 a B, e_2 . Calcula $m =$ número de ceros recibidos de B en los pasos en que envió ψ_1 o ψ_2 .

Si $|\frac{3}{4}e_1 + \frac{1}{4}e_2 - m| > TOL_A$, entonces decide que es altamente probable la presencia de E y aborta el protocolo a través del canal público.

B actúa de manera simétrica con n_B y TOL_B .

Paso 5: Una vez el envío de qubits cesa, A y B poseen tres cadenas binarias cada uno. Si denotamos por m_{IJ} = medidas del interlocutor I sobre los qubits de J , entonces

A posee m_{AA}, m_{AB}, m_{BA} .

B posee m_{BB}, m_{AB}, m_{BA} .

Vamos a suponer, sin pérdida de generalidad, que m_{AB} y m_{BA} están compuestas de k bits, donde k es un número par.

A y B comparan públicamente $k/2$ bits de m_{AB} . Si no coinciden, E es detectado.

A y B comparan públicamente $k/2$ bits de m_{BA} . Si no coinciden, E es detectado.

En caso de que las comparativas no hayan arrojado como resultado la detección de E, A y B usan como clave común los restantes $k/2$ bits de m_{AB} con los restantes $k/2$ bits de m_{BA} , concatenados en dicho orden.

6. Comentarios de diseño

6.1. Acerca del Paso 5

El primer comentario, y uno de los más importantes, es que podemos suponer sin pérdida de generalidad que si E captura los resultados de medida que A y B se envían en los pasos 2 y 3 y los altera, entonces es altamente probable que sea detectada en el paso 5. La esencia de que esto ocurra está en el hecho de que E finge ser dos personas al mismo tiempo y que ignora cuáles serán los qubits que efectuarán el control. Siendo más concretos:

Recordando la notación m_{IJ} , supongamos que E está intentando obtener los qubits del protocolo por el procedimiento de EPR Man in the Middle. De ese modo E poseerá sus propios juegos de qubits y sus propios resultados de medida, en concreto $m_{BA}, m_{AB}, m_{EA}, m_{EB}, m_{AE}$ y m_{BE} .

A cree que posee m_{AA}, m_{AB} y m_{BA} , pero en realidad posee m_{AA}, m_{AE} y m_{EA} , y con B ocurre de manera análoga.

Entonces, en el paso 5 del protocolo,

A y B creen que comparan públicamente $k/2$ bits de m_{AB} , pero en un caso serán bits de m_{AE} (los que posee A) y en el otro serán bits de m_{EB} (los que posee B), y en general estos dos números no tienen por qué coincidir en absoluto.

A y B creen que comparan públicamente $k/2$ bits de m_{BA} , y ocurre de manera análoga al caso anterior.

De hecho, y siendo más concretos, dado $\underline{x}, \underline{y}$ cadenas binarias, se define la *distancia de Hamming* como $d(\underline{x}, \underline{y}) = \#\{i/x_i \neq y_i\}$. Ocurre que $d(C_A, C_B)$, donde C_I es la cadena binaria de todos los bits pertenecientes a I que comparan A y B, es cota inferior del número de veces que ha interferido E en los resultados de medida enviados entre A y B.

Otro posible problema es que al mostrarse públicamente esos bits de control puede que estemos dando datos suficientes para que un observador externo logre deducir el valor de la clave. Eso no ocurre ya que la clave está formada sólo por bits que no participan en el proceso de control, y los valores de m_{AB} y m_{BA} son mutuamente independientes dado que son el resultado de mediciones realizadas sobre distintos qubits. Por ese motivo, en el envío del i -ésimo qubit, aunque el qubit que A manda a B sea público porque será usado como qubit de control, eso no afectará en absoluto al qubit que B manda a A en ese mismo momento, que puede o no aparecer en la clave o también ser usado como qubit de control.

Por supuesto, el modo en que se eligen los qubits de control debe ser completamente aleatorio, para no otorgar a E ninguna pista sobre qué qubits serán empleados para detectarla a posteriori.

Otro asunto a discutir es el uso de la cadena binaria compartida. A y B pueden emplearla como clave para algún protocolo de cifrado público o bien para efectuar un cifrado de Vernam con ella. En caso de que ocurriera lo segundo, se puede decir que este protocolo es una variante del protocolo one-time-pad con entrelazamiento mostrado en [4] que además es resistente a un ataque de tipo EPR Man in The Middle.

Por último, cabe la posibilidad de que E decidiera intervenir sólo en algunas mediciones, siendo así más fácil que lograra pasar desapercibido pero obteniendo a cambio sólo unos pocos bits de la clave, información que aunque en principio puede parecer insuficiente podría resultar útil de cara a criptoanálisis con parte de la clave conocida. Podemos minimizar este impacto empleando como clave final una función hash de la obtenida en este protocolo.

6.2. Acerca del Paso 1

De ahora en adelante supondremos, debido a la sección anterior, que E deja pasar intocados los resultados de medida que A envía a B y viceversa. Nos cen-

traremos también, sin pérdida de generalidad, en los qubits que A envía a B, ya que los qubits que B envía a A reciben un tratamiento completamente análogo.

Cabe preguntarse si la elección de los tres tipos de qubits del protocolo es lo suficientemente satisfactoria como para que la información previa que E pueda deducir de ellos sea nula o, al menos, muy reducida.

Podemos dar una medida de la cantidad de información que estamos otorgando por medio de una herramienta de información cuántica llamada la *Cantidad de Holevo*. La Cantidad de Holevo es una cota superior de la entropía que se puede ejercer sobre un sistema de qubits y sus probabilidades condicionada por todas las posibles maneras de medirlos. Sin entrar en muchos detalles, basta decir que cuanto más cercana a cero se encuentre, menos información estamos aportando, y cuanto más cercana a uno esté, más precisa será la posibilidad de distinguirlos. La Cantidad de Holevo vale uno cuando los qubits son ortogonales (ya que se pueden distinguir con exactitud) y cero cuando son iguales (ya que da igual cómo los midamos, nada será mejor que elegir uno de ellos al azar con una distribución equiprobable). La Cantidad de Holevo se denota por χ y se mide en bits.

En nuestro caso, la manera más cómoda de calcular la Cantidad de Holevo es empleando la matriz de densidad de nuestro sistema de qubits. Simplemente nos limitaremos a escribirla, la definición puede encontrarse en [3]. Dado que nuestro sistema consiste en enviar ψ_1 con probabilidad $1/4$, ψ_2 con probabilidad $1/4$ o ψ_3 con probabilidad $1/2$, la matriz de densidad es

$$\rho = \frac{1}{4}\rho_1 + \frac{1}{4}\rho_2 + \frac{1}{2}\rho_3,$$

donde

$$\rho_1 = \begin{bmatrix} 3/4 & \sqrt{3}/4 \\ \sqrt{3}/4 & 1/4 \end{bmatrix}, \rho_2 = \begin{bmatrix} 1/4 & \sqrt{3}/4 \\ \sqrt{3}/4 & 3/4 \end{bmatrix}, \rho_3 = \begin{bmatrix} 1/2 & 1/2 \\ 1/2 & 1/2 \end{bmatrix}.$$

La Cantidad de Holevo se define por

$$\chi = S(\rho) - \sum_i p_i S(\rho_i),$$

donde $\rho = \sum_i p_i \rho_i$ y $S(\rho)$ es la entropía de Von Neumann del estado ρ , que es el análogo cuántico de la entropía de Shannon en teoría clásica de la información, y se puede definir a partir de sus autovectores λ_i como

$$S(\rho) = - \sum_i \lambda_i \log \lambda_i,$$

usando además el convenio de que $0 \log 0 \equiv 0$. En nuestro caso obtenemos que $\chi \simeq 0,2116205364$, lo que es un número bastante bueno de cara al protocolo (en términos generales, dice que la máxima información que E puede obtener en cada envío es alrededor de la quinta parte de un bit).

Siendo autocríticos, una posible estrategia para E puede ser alternar los envíos de los pares EPR $\sqrt{3}/2|00\rangle + 1/2|11\rangle$ y $1/2|00\rangle + \sqrt{3}/2|11\rangle$ con probabilidad $1/2$. Es una buena idea ya que mezcla todos los casos anteriores a la vez, y su probabilidad de acertar a ciegas es mayor, ya que unifica varios de ellos al mismo tiempo. Esto puede solventarse escogiendo nuevos ψ_1, ψ_2 de modo que $X(\psi_1) \neq \psi_2$, es decir, que la probabilidad de obtener cero en ψ_1 no sea la misma que la de obtener uno en ψ_2 y viceversa. Sin embargo, como veremos en los comentarios al Paso 3, esta estrategia conlleva que E está usando constantemente pares EPR, y por tanto sigue estando expuesto a ser descubierto. Además, en los comentarios al Paso 4, veremos que no supone mejora ni empeoramiento hacer dicho cambio en relación con las posibles estrategias llevadas a cabo por E.

Otro importante comentario es si esta Cantidad de Holevo puede ser mejorada, y la respuesta es positiva, pero con ciertos matices. En las siguientes cuentas asumiremos que $\psi_1 = X(\psi_2)$.

Dejando el qubit ψ_3 prefijado, sin pérdida de generalidad (puesto que podemos rotar todo el conjunto si así lo deseáramos) tenemos ahora que

$$\rho_1 = \begin{bmatrix} \cos^2 \alpha & \sin \alpha \cos \alpha \\ \sin \alpha \cos \alpha & \sin^2 \alpha \end{bmatrix},$$

$$\rho_2 = \begin{bmatrix} \sin^2 \alpha & \sin \alpha \cos \alpha \\ \sin \alpha \cos \alpha & \cos^2 \alpha \end{bmatrix},$$

$$\rho_3 = \begin{bmatrix} 1/2 & 1/2 \\ 1/2 & 1/2 \end{bmatrix}.$$

Calculamos los autovalores de la matriz $\rho = 1/4\rho_1 + 1/4\rho_2 + 1/2\rho_3$, y obtenemos que son

$$\lambda_1 = \frac{1}{2} + \frac{\sqrt{a^2 + a + 1/4}}{2}, \lambda_2 = \frac{1}{2} - \frac{\sqrt{a^2 + a + 1/4}}{2},$$

donde $a = \sin \alpha \cos \alpha$.

Un sencillo cálculo demuestra que χ es mínima cuando los autovectores de ρ son respectivamente uno y cero, y eso ocurre cuando $\sqrt{a^2 + a + 1/4} = 1$. Esto a su vez sólo ocurre si $a = 1/2$ ó $a = -3/2$. Finalmente, a la hora de despejar el valor de α , el segundo caso no arroja soluciones reales y por tanto no nos sirve, y el primer caso, $a = 1/2$, nos da como posibles soluciones $\alpha = \pi/4$ y $\alpha = -3\pi/4$.

Ambos valores llevan de hecho a la misma situación (como en este caso, puede ocurrir que distintos qubits generen la misma matriz de densidad) y es que $\rho_1 \equiv \rho_2 \equiv \rho_3$, lo que resulta bastante lógico teniendo en cuenta que si $\chi = 0$ eso quiere decir que los tres estados deben ser absolutamente indistinguibles. Por otro lado, cuanto más próximo esté el valor de a a $-1/2$, mínimo global de la ecuación $a^2 + a + 1/4 = 0$, más fácil será distinguir entre los tres qubits. En nuestro caso concreto tenemos que $\alpha = \pi/6$, y por tanto $a = \sin(\pi/6) \cos(\pi/6) = \sqrt{3}/4$, que es un valor bastante próximo a $1/2$ y por tanto satisfactorio, como ya suponíamos cuando habíamos calculado la Cantidad de Holevo.

6.3. Acerca del Paso 3

Sobre el Paso 2 no hay nada reseñable que decir, pues simplemente describe que B debe declarar su medida a A y nos remite de nuevo al Paso 1. En el Paso 3, sin embargo, ocurre algo interesante, y es que E ya puede ser descubierta, y además con certeza total.

Si nos encontramos en el caso *a*) de dicho apartado, entonces A mandó un qubit EPR a B. Estos qubits actúan como qubits de control de manera pasiva, pues E está forzada a realizar una elección. En el apartado *a*), mande el qubit que mande, si el resultado de medida que B declara a A no coincide con el de A (y recordemos que E debe dejar pasar los resultados de medidas de B para A) entonces E es detectado ya que A sabe que debería estar compartiendo un par EPR común con B y por tanto los resultados de medida no pueden ser distintos.

En el caso *b*), sin embargo, E no puede ser detectado en este mismo paso, ya que al no ser de tipo EPR el qubit enviado a B el hecho de que sus resultados no coincidan no quiere decir que E esté interfiriendo. Sin embargo, las decisiones que E tome también servirán para detectarle, pues si envía qubits EPR cuando no debe hacerlo estará introduciendo un sesgo en la frecuencia de ceros y unos, como veremos en el siguiente paso.

6.4. Acerca del Paso 4

Antes de este paso, A previamente, en el Paso 0, eligió un número natural n_A , que representa el momento en el que considera que puede chequear el sesgo ejercido sobre los qubits que no son EPR al mandar E, sin saberlo, qubits EPR en su lugar. En ese mismo paso eligió también una tolerancia TOL_A , que representa el error máximo que permitimos a dicho valor por efectos de la estadística. En general, sin la presencia de E obtendremos una distribución de probabilidad sobre los ceros de los qubits no EPR distinta que con la presencia de E, ya que éste habrá introducido qubits EPR de manera accidental que alterarán dicha probabilidad. La cuenta que efectuamos en dicho paso no es más que comparar el valor esperado de ceros con el valor real obtenido. Si esa diferencia es muy acusada (es decir, mayor que la tolerancia permitida por A), entonces se concluye que la presencia de E es altamente probable y el protocolo no es seguro, por lo que es mejor abortarlo. En el caso de B, todo funciona de manera simétrica.

Es importante destacar que tanto la tolerancia de A y B como el momento en que realizan el chequeo es conocido sólo por cada uno de ellos, por lo que E no puede aprovechar esa información para intentar discernir en qué momento tratarán de buscarla por este método. Asimismo puede hacerse una variante del protocolo donde este procedimiento se repita en varias ocasiones, pero habrá que tener en cuenta que será tanto más preciso cuanto mayor sea el ciclo (la sucesión de pasos 1 a 3 del protocolo) en que decidamos hacerlo.

Por último hay que tener en cuenta algo importante, y es que si E decide emplear la estrategia de enviar siempre qubits iguales a ψ_3 a B, entonces no está introduciendo sesgo alguno en el promedio de ceros sobre el conjunto de los qubits con los que no acierta la elección. Eso se debe a la simetría de los qubits ψ_1 y ψ_2 elegidos para el protocolo, que hacen que utilizando la ley de la probabilidad total y la probabilidad condicionada, el valor esperado de ceros sea $1/2 \cdot 1/4 + 1/2 \cdot 3/4 = 1/2$, que es exactamente el que obtendría E de emplear constantemente ψ_3 en esos casos. Cabe pensar si introduciendo alguna asimetría en los qubits, y por lo tanto en las distribuciones de probabilidad, no se podría eliminar por completo esta posibilidad. Puede hacerse, pero entonces lo que ocurriría es que, simplemente, aparecerían nuevas posibilidades explotables por E.

Por ejemplo, si dejamos ψ_2 tal cual está pero tomamos en vez de ψ_1 el qubit

$$\varphi_1 = \sqrt{\frac{2}{3}}|0\rangle + \frac{1}{\sqrt{3}}|1\rangle$$

entonces el valor esperado de ceros sería igual a $1/2 \cdot 1/4 + 1/2 \cdot 2/3 = 11/24$. Pero en este caso, si E emplea el par EPR

$$\sqrt{\frac{11}{24}}|00\rangle + \sqrt{\frac{13}{24}}|11\rangle,$$

vuelve a simular de nuevo ese valor esperado.

¿Invalida esto el protocolo? No, pues E aún puede ser atrapada debido al apartado a) del Paso 3. En concreto, para no ser descubierta debería coincidir siempre con B en todas las medidas que haga sobre los qubits que A sabe que son EPR. Teniendo en cuenta que en promedio la mitad de los qubits que A manda a B serán EPR y que la probabilidad de E de coincidir con B en su medida es $2/4 = 1/2$ (pues hay cuatro casos posibles y dos son favorables para E), su probabilidad de pasar desapercibida es de

$$\left(\frac{1}{2}\right)^{\frac{k}{2}} = \left(\frac{1}{\sqrt{2}}\right)^k,$$

donde k es el número de qubits que A envía a B. Esta cifra resulta ser considerablemente pequeña a poco que aumentemos el número de qubits. Debemos además tener en cuenta que esta probabilidad será aún menor debido a que también existirá intercambio de qubits de B para A, ya que sólo estamos considerando un sentido en estos cálculos.

Referencias

- [1] D. Richard Kuhn. Vulnerabilities in Quantum Key Distribution Protocols. Preprint, math. AG/0305076.
- [2] A. Beige, B.-G. Englert, C. Kurtsiefer, H. Weinfurter. Secure communication with a publicly known key. Preprint, math. AG/0111106v2.
- [3] Nielsen and Chuang, *Quantum Information and Quantum Computation*, Cambridge University Press (2000).
- [4] Qing-yu Cai. An one-time-pad key communication protocol with entanglement. Preprint, math. AG/0309108v5.
- [5] Manuel J. Lucena López, *Criptografía y Seguridad en Computadores*, edición electrónica bajo licencia Creative Commons (2007).

- [6] M. Genovese. Photonic realization of quantum information protocols. In School on Theory and Technology in Quantum Information, Communication, Computation and Cryptography (Trieste/Italy, 2006). Pages 28-32.
- [7] M. Agrawal, N. Kayal, N. Saxena, Primes in P, *Annals of Mathematics*, 160 No 02 (2004) 781-793.
- [8] Douglas Adams, *Guía del Autoestopista Galáctico*, Anagrama (2005).
- [9] G. Chaitin, The limits of reason, *Scientific American* Vol 294 No 03 (2006) 74-81.
- [10] P. Shor, J. Preskill. Simple proof of security of the BB84 quantum key distribution protocol. Preprint, math. AG/0003004v2.
- [11] G. Brassard. Brief history of quantum cryptography: a personal perspective. Preprint, math. AG/0604072v1.

Agradecimientos

El autor quiere agradecer especialmente a Ismael Capel por la ayuda prestada. También quiere destacar la gran ayuda que supuso el Curso de Iniciación a la Investigación, impartido como parte del Máster de Investigación Matemática, para la elaboración de este trabajo.