

TU CEREBRO NO ES UN ORDENADOR

A pesar de todo lo que sabemos sobre la estructura del cerebro humano, sabemos muy poco de la mente humana.

Las mentes humanas no funcionan como ordenadores

Nuestros cerebros no almacenan palabras, ni imágenes ni crean representaciones de las mismas. Nuestras experiencias no se almacenan en búferes de almacenamiento de memoria a corto plazo y luego se transfieren a un almacenamiento de memoria a largo plazo, donde se pueden recuperar con precisión más adelante. En cambio, nacemos con la capacidad innata de interactuar con el mundo de manera efectiva. Desde el nacimiento, prestamos especial atención a información importante (como mover la cabeza para seguir las voces de nuestras madres), y hacemos conexiones (todo bebé es un rompehielos instantáneo, son capaces de convertir en un adulto afable incluso a la persona más arisca que conozcamos).

Así es como aprendemos, cómo somos capaces de transformarnos para poder interactuar eficazmente en un mundo tan diferente del que nuestros antepasados experimentaron.

Los ordenadores procesan la información. Las mentes humanas no

En un ordenador, cada palabra, imagen y estímulo tiene que ser codificado con un patrón inmutable. Todo lo que hacen los ordenadores se guía por el código, por el algoritmo.

Durante muchos años, los gobiernos, las corporaciones, incluso los científicos han tratado de controlar la forma en la que pensamos y actuamos a través de la ingeniería social, la publicidad y otras tecnologías de este tipo. Pero a diferencia de los ordenadores, las mentes humanas tienden a romper la programación para crear algo nuevo. Los humanos saben provocar al cambio y generar revoluciones.

La mente humana explicada en metáforas

Epstein subraya que, a pesar de todos los estudios científicos lanzados sobre el cerebro humano, todavía estamos muy lejos de entender realmente cómo funciona éste. Dice que la forma en la que entendemos el cerebro humano hoy en día no puede sostenerse, simplemente porque no puede explicar completamente la naturaleza de los seres humanos.

Durante todo este tiempo, nuestra comprensión de la mente humana ha sido a través de metáforas o modelos, donde cada metáfora ha ido cambiando según avanzaba la tecnología y la filosofía de la época.

La que prevalece actualmente, la teoría del Procesamiento de la Información, nos ha impedido ver nuestras capacidades únicas (inherentes), y en este proceso nos ha impedido descubrir quiénes somos realmente como seres humanos.

No hay inmortalidad a través del upload de la mente

Epstein dice que la inmortalidad a través del “upload” de la mente nunca se podrá lograr debido al singular problema que afirma que los estímulos externos cambian el cerebro de los individuos de diferentes maneras. Pese a que todos podemos tener las mismas estructuras básicas en nuestro cerebro, la forma en que se “cablea” (cómo interpreta los estímulos, cómo integra la información) es totalmente independiente de nuestras experiencias anteriores.

Para entender el significado del estado físico actual del cerebro, uno debe entender toda la historia de la vida del dueño de ese cerebro. En cierto sentido, los patrones cerebrales no significan nada fuera del cerebro que produce sus señales. Incluso si el Proyecto Cerebro Humano tuvo éxito en el mapeo de todas las conexiones sinápticas del cerebro, lo que producen sólo es cierto para el único cerebro que han estudiado, y sólo es cierto para este momento, hasta que nuevas experiencias lo vuelven a re-cablear. Además, no existe la garantía de que un cerebro que tiene las mismas conexiones sinápticas exactas que el cerebro modelo tomara las mismas decisiones y elecciones. En resumen, hay algo que sucede en el ser humano que no está predeterminado por su hardware físico.

Implicaciones para la Inteligencia Artificial

Muchos científicos piensan que para resolver el desafío de la alineación de la IA, simplemente se necesita entender cómo funciona el cerebro humano. El artículo de Epstein muestra que, de hecho, es imposible hacerlo, incluso con la mejor tecnología, si sólo vemos al ser humano como un ordenador muy avanzado.

El artículo de Epstein también nos muestra el poder único con el que nacen todos los seres humanos: la capacidad de aprender. Esta capacidad nos permite establecer conexiones, asegurando así que nuestra capacidad para crear nuevos conocimientos sea ilimitada, siempre que se le permitiera a nuestra mente funcionar sin impedimento alguno.

Por último, la promesa de la inmortalidad a través del uploading de la mente se rompe. Continuar por este camino utilizando una perspectiva materialista sería un desperdicio de dinero, tiempo y otros recursos que se podrían utilizar mejor para resolver los problemas sociales más apremiantes del mundo que impiden a los seres humanos alcanzar todo su potencial.

El cerebro vacío

Tu cerebro no procesa información, ni recupera conocimientos ni almacena recuerdos. En resumen: tu cerebro no es un ordenador.

No importa lo que se esfuercen en los intentos, los científicos del cerebro y los psicólogos cognitivos nunca encontrarán una copia de la 5a Sinfonía de Beethoven en el cerebro, o copias de palabras, imágenes, reglas gramaticales o cualquier otro tipo de estímulos ambientales. El cerebro humano no está vacío, por supuesto. Pero no contiene la mayoría de las cosas que la gente piensa que tiene, ni siquiera cosas simples como los “recuerdos”.

Nuestra deficiente visión del cerebro tiene profundas raíces históricas, pero la invención de los ordenadores en la década de 1940 nos ha confundido aún más. Desde hace más de medio siglo, psicólogos, lingüistas, neurocientíficos y otros expertos en comportamiento humano han estado afirmando que el cerebro humano funciona como un ordenador.

Para ver cuán vacua es esta idea, considere los cerebros de los bebés. Gracias a la evolución, los neonatos humanos, como los recién nacidos de todas las demás especies de mamíferos, entran

en el mundo preparados para interactuar con él de manera efectiva. La visión de un bebé es borrosa, pero presta especial atención a las caras, y es rápidamente capaz de identificar la de su madre. Prefiere el sonido de las voces a los sonidos que no son de voz, y puede distinguir un sonido de voz básico de otro. Estamos, sin duda, contruidos para hacer conexiones sociales.

Un recién nacido sano también está equipado con más de una docena de reflejos– reacciones ya hechas para ciertos estímulos que son importantes para su supervivencia. Gira su cabeza en la dirección de algo que acaricia su mejilla y luego chupa lo que entra en su boca. Aguanta la respiración cuando se sumerge en el agua. Agarra las cosas colocadas en sus manos con tanta fuerza que casi puede soportar su propio peso. Tal vez lo más importante, los recién nacidos vienen equipados con poderosos mecanismos de aprendizaje que les permiten cambiar rápidamente para que puedan interactuar cada vez más eficazmente con su mundo, incluso si ese mundo es diferente al que enfrentaron sus antepasados lejanos.

Sentidos, reflejos y mecanismos de aprendizaje – esto es con lo que empezamos, y es bastante, cuando lo piensas. Si carecemos de alguna de estas capacidades al nacer, probablemente tendríamos problemas para sobrevivir.

Pero esto es con lo que no nacemos: información, datos, reglas, software, conocimiento, léxicos, representaciones, algoritmos, programas, modelos, memorias, imágenes, procesadores, subrutinas, codificadores, decodificadores, símbolos o búferes– elementos de diseño que permiten que los ordenadores digitales se comporten de forma algo inteligente. No sólo no nacemos con tales cosas, sino que no las desarrollamos, nunca.

No almacenamos las palabras ni las reglas que nos dicen cómo manipularlas. No creamos representaciones de estímulos visuales, los almacenamos en un búfer de memoria a corto plazo y, a continuación, transferimos la representación a un dispositivo de memoria a largo plazo. No recuperamos información, imágenes o palabras de registros de memoria. Los ordenadores hacen todas estas cosas, pero los organismos vivos no.

Los ordenadores, literalmente, procesan información – números, letras, palabras, fórmulas, imágenes. La información primero tiene que ser codificada en un formato que los ordenadores puedan utilizar, lo que significa patrones de unos y ceros ('bits') organizados en pequeños fragmentos ('bytes'). En mi ordenador, cada byte contiene 8 bits, y un cierto patrón de esos bits significa la letra d, otro la letra o, y otro la letra g. Uno al lado del otro, esos tres bytes forman la palabra perro. Una sola imagen – digamos, la fotografía de mi gato Henry en mi escritorio – está representada por un patrón muy específico de un millón de estos bytes ('un megabyte'), rodeado de algunos caracteres especiales que le dicen al ordenador que espere una imagen, no una palabra.

Los ordenadores, literalmente, mueven estos patrones de un lugar a otro en diferentes áreas de almacenamiento físico grabadas en componentes electrónicos. A veces también copian los patrones, y a veces los transforman de varias maneras – por ejemplo cuando estamos corrigiendo errores en un manuscrito o cuando estamos retocando una fotografía. Las reglas que siguen los equipos para mover, copiar y operar en estas matrices de datos también se almacenan dentro del equipo. Juntos, un conjunto de reglas se llama un 'programa' o un 'algoritmo'. Un grupo de algoritmos que trabajan juntos para ayudarnos a hacer algo (como comprar acciones o encontrar una fecha on-line) se llama una "aplicación", lo que la mayoría de la gente ahora llama una "app".

Perdóneme por esta introducción a la programación, pero necesito ser claro: los ordenadores realmente operan con representaciones simbólicas del mundo. Realmente almacenan y

recuperan. Realmente procesan. Realmente tienen memorias físicas. Realmente son guiados en todo lo que hacen, sin excepción, por algoritmos.

Los humanos, por otro lado, no lo hacen, nunca lo hicieron, nunca lo harán. Dada esta realidad, ¿por qué tantos científicos hablan de nuestra actividad mental como si fuésemos ordenadores?

En su libro *In Our Own Image* (2015), el experto en inteligencia artificial George Zarkadakis describe seis modelos o metáforas diferentes que la gente ha empleado en los últimos 2.000 años para tratar de explicar la inteligencia humana.

En la primera, conservada en la Biblia, los seres humanos se formaron a partir de arcilla o tierra, a los que luego un dios inteligente infundió con su espíritu. Ese espíritu 'explicó' nuestra inteligencia – al menos gramaticalmente.

La invención de la ingeniería hidráulica en el siglo III a. C. condujo a la popularidad de un modelo hidráulico de inteligencia humana, la idea de que el flujo de diferentes fluidos en el cuerpo - los 'humores' - explicaba nuestro funcionamiento físico y mental. La metáfora hidráulica persistió durante más de 1.600 años, retrasando el avance de la medicina durante todo ese periodo.

En el 1.500, se idearon los autómatas alimentados por muelles y engranajes, lo que finalmente inspiró a pensadores líderes como René Descartes a afirmar que los humanos son máquinas complejas. En la década de 1.600, el filósofo británico Thomas Hobbes sugirió que el pensamiento surgió de pequeños movimientos mecánicos en el cerebro. En la década de 1.700, los descubrimientos sobre la electricidad y la química condujeron a nuevas teorías de la inteligencia humana, de nuevo, en gran parte metafóricas en la naturaleza. A mediados de la década de 1.800, inspirado en los recientes avances en las comunicaciones, el físico alemán Hermann von Helmholtz comparó el cerebro con un telégrafo.

Cada metáfora reflejaba el pensamiento más avanzado de la época que la generó. Previsiblemente, pocos años después de los albores de la tecnología informática en la década de 1940, se decía que el cerebro funcionaba como un ordenador, con el papel del hardware físico desempeñado por el propio cerebro y nuestros pensamientos sirviendo como software. El evento histórico que lanzó lo que ahora se llama ampliamente "ciencia cognitiva" fue la publicación de *Lenguaje y Comunicación* (1951) por el psicólogo George Miller. Miller propuso que el mundo mental puede ser estudiado rigurosamente usando conceptos de la teoría de la información, la computación y la lingüística.

Este tipo de pensamiento fue llevado a su máxima expresión en el libro *The Computer and the Brain* (1958), en el que el matemático John von Neumann declaró rotundamente que la función del sistema nervioso humano es 'prima facie digital' ('a primera vista digital'). Aunque reconoció que en realidad se sabía poco sobre el papel que el cerebro desempeñaba en el razonamiento humano y la memoria, dibujó una correspondencia paralela entre los componentes de las máquinas de computación de la época y los componentes del cerebro humano.

Impulsado por los avances posteriores tanto en la tecnología informática como en la investigación cerebral, se desarrolló gradualmente un ambicioso esfuerzo multidisciplinario para entender la inteligencia humana, firmemente arraigado en la idea de que los seres humanos son, como los ordenadores, procesadores de información. Este esfuerzo ahora involucra a miles de investigadores, consume miles de millones de dólares en financiación y ha generado una vasta literatura de artículos y libros tanto técnicos como convencionales. El libro de Ray Kurzweil *How to Create a Mind: The Secret of Human Thought Revealed* (2013), ejemplifica esta perspectiva, especulando sobre los 'algoritmos' del cerebro, cómo el cerebro 'procesa los datos', e incluso cómo se asemeja superficialmente en su estructura a la integración de circuitos.

La metáfora del procesamiento de la información (PI) de la inteligencia humana ahora domina el pensamiento humano, tanto en la calle como en las ciencias. Prácticamente no hay ninguna forma de discurso sobre el comportamiento humano inteligente que proceda sin emplear esta metáfora, del mismo modo que ninguna forma de discurso sobre el comportamiento humano inteligente podría proceder en ciertas épocas y culturas sin referencia a un espíritu o deidad. La validez de la metáfora de la P.I. en el mundo actual se asume generalmente sin dudarlo.

Pero la metáfora de la PI es, después de todo, otra metáfora: una historia que contamos para dar sentido a algo que realmente no entendemos. Y como todas las metáforas que la precedieron, sin duda será echada a un lado en algún momento, ya sea reemplazada por otra metáfora o, al final, reemplazada por el conocimiento real.

Hace poco más de un año, en una visita a uno de los institutos de investigación más prestigiosos del mundo, desafié a los investigadores de allí a dar cuenta de un comportamiento humano inteligente sin referencia a ningún aspecto de la metáfora de la PI y no pudieron hacerlo, y cuando cortésmente planteé el problema en las subsiguientes comunicaciones de correo electrónico, seguían sin tener nada que ofrecer meses más tarde. Vieron el problema. No desestimaron el desafío como trivial. Pero no consiguieron encontrar una alternativa. En otras palabras, la metáfora de la PI es "pegajosa". Inunda nuestro pensamiento con lenguaje e ideas que son tan poderosas que nos cuesta mucho poder pensar de otra manera.

La lógica defectuosa de la metáfora PI es fácilmente demostrable. Se basa en un silogismo defectuoso - uno con dos premisas razonables y una conclusión defectuosa. Premisa razonable #1: todos los ordenadores son capaces de comportarse de forma inteligente. Premisa razonable #2: todos los ordenadores son procesadores de información. Conclusión defectuosa: todas las entidades que son capaces de comportarse de forma inteligente son procesadores de información.

Dejando a un lado el lenguaje formal, la idea de que los seres humanos deben ser procesadores de información sólo porque las computadoras son procesadores de información es simplemente tonto, y cuando, algún día, la metáfora de la PI finalmente se abandone, casi con toda seguridad se verá de esa manera por los historiadores, así como ahora vemos las metáforas hidráulicas y mecánicas como tonterías.

Si la metáfora de la PI es tan tonta, ¿por qué es tan "pegajosa"? ¿Qué nos impide dejarla de lado, así como podríamos echar a un lado una rama que estaba bloqueando nuestro camino? ¿Hay alguna manera de entender la inteligencia humana sin apoyarse en tan endeble muleta intelectual? ¿Y qué precio hemos pagado por apoyarnos tanto en esta muleta en particular durante tanto tiempo? Después de todo, la metáfora de la PI, ha guiado el pensamiento de un gran número de investigadores en múltiples campos durante décadas. ¿A qué precio?

En un ejercicio que he realizado en clase muchas veces a lo largo de los años, empiezo escogiendo a un estudiante para que dibuje una imagen detallada de un billete de dólar, "lo más detallada posible", digo, en la pizarra frente a la sala. Cuando el estudiante termina, cubro el dibujo con una hoja de papel, saco un billete de dólar de mi billetera, lo cuelgo en la pizarra y le pido al alumno que repita la tarea. Cuando él o ella ha terminado, quito la cubierta del primer dibujo, y la clase comenta las diferencias.

Este es su dibujo 'de memoria' que hizo Jinny, una de las estudiantes (fíjese en la metáfora):

Y aquí está el dibujo que posteriormente hizo con el billete de dólar delante de ella:



Jinny estaba tan sorprendida por el resultado como probablemente lo esté usted, pero es típico. Como se puede ver, el dibujo hecho en ausencia del billete de dólar es horrible en comparación con el dibujo hecho con un ejemplar, a pesar de que Jinny ha visto un billete de dólar miles de veces.

¿Cuál es el problema? ¿No tenemos una "representación" del billete de dólar "almacenado" en un "registro de memoria" en nuestro cerebro? ¿No podemos simplemente 'recuperarlo' y usarlo para hacer nuestro dibujo?

Obviamente no, y mil años de neurociencia nunca localizarán una representación de un billete de dólar almacenado dentro del cerebro humano por la sencilla razón de que no está allí para ser encontrado.

Una gran cantidad de estudios cerebrales nos dicen, de hecho, que múltiples y a veces grandes áreas del cerebro a menudo están involucradas incluso en las tareas de memoria más mundanas. Cuando las emociones fuertes están involucradas, millones de neuronas pueden volverse más activas. En un estudio de 2016 del neuropsicólogo de la Universidad de Toronto Brian Levine y otros sobrevivientes de un accidente de avión, recordaron que el accidente aumentó la actividad neuronal en 'la amígdala, el lóbulo temporal medio, la línea media anterior y posterior, y la corteza visual' de los pasajeros.

La idea, desarrollada por varios científicos, de que los recuerdos específicos se almacenan de alguna manera en las neuronas individuales es absurda; si acaso, esa afirmación sólo empuja el problema de la memoria a un nivel aún más desafiante: ¿cómo y dónde, después de todo, se almacena la memoria en la célula?

Entonces, ¿qué está ocurriendo cuando Jinny dibuja el billete de dólar sin tenerlo a la vista? Si Jinny nunca hubiera visto un billete de dólar antes, su primer dibujo probablemente no se habría asemejado al segundo dibujo en absoluto. Después de haber visto billetes de dólar antes, ella cambió de alguna manera. Específicamente, su cerebro fue cambiado de tal manera que le

permitió visualizar un billete de dólar, es decir, volver a experimentar ver un billete de dólar, al menos hasta cierto punto.

La diferencia entre los dos diagramas nos recuerda que visualizar algo (es decir, ver algo en su ausencia) es mucho menos preciso que ver algo en su presencia. Por eso somos mucho mejores reconociendo que recordando. Cuando memoramos algo (del latín *re* 'de nuevo', y *memorari* 'ser conscientes de'), tenemos que tratar de revivir una experiencia; pero cuando reconocemos algo, simplemente debemos ser conscientes del hecho de que hemos tenido esta experiencia perceptiva antes.

Tal vez Ud. no acepte esta demostración. Jinny había visto billetes de dólar antes, pero no había hecho un esfuerzo deliberado para 'memorizar' los detalles. Si lo hubiera hecho, se podría argumentar, que presumiblemente podría haber dibujado la segunda imagen sin que el billete estuviera presente. Incluso en este caso, sin embargo, ninguna imagen del billete de dólar ha sido en ningún sentido "almacenada" en el cerebro de Jinny. Simplemente se ha preparado mejor para dibujarlo con precisión, al igual que, a través de la práctica, un pianista se vuelve más hábil en tocar un concierto sin "inhalar" de alguna manera una copia de la partitura.

De este simple ejercicio, podemos comenzar a construir el marco de una teoría libre de metáforas del comportamiento humano inteligente - uno en el que el cerebro no esté completamente vacío, pero al menos esté vacío de la "carga" de la metáfora PI.

A medida que navegamos por el mundo, vamos cambiando debido a una variedad de experiencias. Las más importantes son de tres tipos: (1) observamos lo que está sucediendo a nuestro alrededor (el comportamiento de otras personas, los sonidos de la música, las instrucciones dirigidas hacia nosotros, las palabras en las páginas, las imágenes en las pantallas); (2) estamos expuestos al emparejamiento de estímulos sin importancia (como el sonido de las sirenas) con estímulos importantes (como la aparición de coches de policía); (3) somos castigados o recompensados por comportarnos de ciertas maneras.

Nos volvemos más eficaces en nuestra vida si vamos cambiando con estas experiencias de manera consecuente – si ahora podemos recitar un poema o cantar una canción, si somos capaces de seguir las instrucciones que se nos dan, si respondemos a los estímulos sin importancia de manera similar a como lo hacemos a estímulos importantes, si nos abstenemos de comportarnos de maneras que fueron castigadas, si nos comportamos con más frecuencia de maneras que fueron recompensadas.

A pesar de los titulares engañosos, nadie tiene realmente ni la más mínima idea de cómo cambia el cerebro después de que hayamos aprendido a cantar una canción o a recitar un poema. Pero ni la canción ni el poema han sido 'almacenados' en él. El cerebro simplemente ha cambiado de una manera ordenada que ahora nos permite cantar la canción o recitar el poema bajo ciertas condiciones. Cuando se sale a actuar, ni la canción ni el poema se 'recuperan' en ningún sentido desde ninguna parte del cerebro, más de lo que los movimientos de mis dedos se 'recuperan' cuando toco con mi dedo en mi escritorio. Simplemente cantamos o recitamos, sin necesidad de recuperación.

Hace unos años, le pregunté al neurocientífico Eric Kandel de la Universidad de Columbia – ganador de un Premio Nobel por identificar algunos de los cambios químicos que tienen lugar en las sinapsis neuronales de la *Aplysia* (un caracol marino) después de que aprende algo – cuánto tiempo piensa que nos llevaría entender cómo funciona la memoria humana. Rápidamente respondió 'Cien años'. No pensé en preguntarle si pensaba que la metáfora de la PI estaba ralentizando la neurociencia, pero algunos neurocientíficos están empezando a pensar lo impensable: que la metáfora de la PI no es indispensable.

Algunos científicos cognitivos, en particular Anthony Chemero de la Universidad de Cincinnati, autor de *Radical Embodied Cognitive Science* (2009), ahora rechazan por completo la opinión de que el cerebro humano funciona como un ordenador. La visión general según la metáfora de la PI es que nosotros, como los ordenadores, damos sentido al mundo realizando cálculos sobre las representaciones mentales del mismo, pero Chemero y otros describen otra forma de entender el comportamiento inteligente, como una interacción directa entre los organismos y su mundo.

Mi ejemplo favorito de la diferencia dramática entre la perspectiva de la PI y lo que algunos llaman ahora la visión "anti-representacional" del funcionamiento humano implica dos formas diferentes de explicar cómo un jugador de béisbol logra atrapar una pelota voladora – bellamente explicado por Michael McBeath, ahora en la Universidad Estatal de Arizona, y sus colegas en un artículo científico de 1995. La perspectiva PI requiere que el jugador formule una estimación de varias condiciones iniciales del vuelo de la pelota – la fuerza del impacto, el ángulo de la trayectoria, ese tipo de cosas – crea y analiza un modelo interno de la trayectoria a lo largo de la cual la pelota probablemente se moverá, y entonces utiliza ese modelo para guiar y ajustar los movimientos de su sistema motor de forma continua en el tiempo con el fin de interceptar la pelota.

Todo sería correcto si funcionásemos como los ordenadores, pero McBeath y sus colegas dieron una explicación más simple: para atrapar la pelota, el jugador simplemente necesita seguir moviéndose de tal manera que mantenga la pelota en una relación visual constante con respecto a la base que se desea alcanzar y el escenario circundante (técnicamente, en una 'trayectoria óptica lineal'). Esto puede sonar complicado, pero en realidad es increíblemente simple, y completamente libre de cálculos, representaciones y algoritmos.

Dos profesores de psicología de la Universidad Leeds Beckett en el Reino Unido, Andrew Wilson y Sabrina Golonka, incluyen el ejemplo del béisbol entre muchos otros que se pueden examinar de forma sencilla y sensata fuera del marco de la PI. Llevan años escribiendo blogs sobre lo que llaman un "enfoque más coherente y naturalizado del estudio científico del comportamiento humano... en contradicción con el enfoque dominante de la neurociencia cognitiva». Sin embargo, esto está lejos de ser un movimiento; las ciencias cognitivas convencionales siguen revolcándose sin crítica en la metáfora de la PI, y algunos de los pensadores más influyentes del mundo han hecho grandes predicciones sobre el futuro de la humanidad que dependen de la validez de dicha metáfora.

Una predicción – hecha por el futurista Kurzweil, el físico Stephen Hawking y el neurocientífico Randal Koene, entre otros – es que, debido a que la conciencia humana es supuestamente como el software informático, pronto será posible descargar las mentes humanas a un ordenador, en los circuitos de los cuales llegaremos a ser inmensamente poderosos intelectualmente y, muy posiblemente, inmortales. Este concepto impulsó la trama de la distópica película *Transcendence* (2014) protagonizada por Johnny Depp como el científico similar a Kurzweil cuya mente fue descargada en Internet, con resultados desastrosos para la humanidad.

Afortunadamente, debido a que la metáfora de la PI ni siquiera es ligeramente válida, nunca tendremos que preocuparnos de que una mente humana se vuelva loca en el ciberespacio; por desgracia, nunca lograremos la inmortalidad a través de la descarga. Esto no es sólo debido a la ausencia del software de conciencia en el cerebro; hay un problema más profundo aquí – llamémoslo '*the uniqueness problem*' (el problema de la singularidad) – que es a la vez inspirador y deprimente.

Porque en el cerebro no existen ni los "bancos de memoria" ni las "representaciones" de estímulos existen en el cerebro, y porque todo lo que se requiere para que funcionemos en el mundo es que el cerebro vaya cambiando de manera ordenada como resultado de nuestras experiencias; no hay razón para creer que cualquiera dos de nosotros _cambiamos de la misma manera tras vivir la misma experiencia. Si tú y yo asistimos al mismo concierto, los cambios que ocurren en mi cerebro cuando escucho la quinta sinfonía de Beethoven casi con toda seguridad serán completamente diferentes de los cambios que ocurren en tu cerebro. Esos cambios, sean cuales sean, se construyen sobre una estructura neuronal única que ya existe, habiéndose desarrollado cada una de las estructuras a lo largo de toda una vida de experiencias únicas.

Es por eso, como demostró Sir Frederic Bartlett en su libro *Remembering* (1932), que no hay dos personas que repitan de la misma manera una historia que hayan escuchado y que además, con el tiempo, sus recitaciones de la historia divergirán cada vez más. Nunca se hace una "copia" de la historia; más bien, cada individuo, al escuchar la historia, cambia hasta cierto punto - lo suficiente como para que cuando se le pregunta por la historia más tarde (en algunos casos, días, meses o incluso años después de que Bartlett les leyera por primera vez la historia) - puedan volver a experimentar la experiencia de escuchar la historia hasta cierto punto, aunque no muy bien (ver el primer dibujo del billete de dólar, arriba).

Esto es inspirador, supongo, porque significa que cada uno de nosotros es verdaderamente único, no sólo en nuestra composición genética, sino incluso en la forma en la que nuestros cerebros cambian con el tiempo. También es deprimente, porque hace que la tarea del neurocientífico sea desalentadora casi más allá de la imaginación. Para cualquier experiencia dada, el cambio ordenado podría involucrar a mil neuronas, a un millón de neuronas o incluso a todo el cerebro, con un patrón de cambio diferente para cada cerebro.

Peor aún, incluso si tuviéramos la capacidad de tomar una instantánea de todas las 86 mil millones de neuronas del cerebro y luego simular el estado de esas neuronas en un ordenador, ese vasto patrón no significaría nada fuera del cuerpo del cerebro que lo produjo. Esta es quizás la forma más atroz en que la metáfora de la PI ha distorsionado nuestro pensamiento sobre el funcionamiento humano. Mientras que los ordenadores almacenan copias exactas de datos – copias que pueden persistir sin cambios durante largos períodos de tiempo, incluso si la electricidad se ha apagado – el cerebro mantiene nuestro intelecto sólo mientras permanezca vivo. No hay interruptor de encendido/apagado. O el cerebro sigue funcionando, o desaparecemos. Además, como señaló el neurobiólogo Steven Rose en *El futuro del cerebro* (2005), una instantánea del estado actual del cerebro no tendría sentido a menos que conociéramos toda la historia de la vida del dueño de ese cerebro – tal vez incluso sobre el contexto social en el que él o ella fue criado.

Piensa en lo difícil que es este problema. Para entender incluso los fundamentos de cómo el cerebro mantiene el intelecto humano, es posible que necesitemos saber no sólo el estado actual de todas las 86 mil millones de neuronas y sus 100 billones de interconexiones, no sólo las diferentes intensidades con las que están conectadas, y no sólo los estados de las más de 1.000 proteínas que existen en cada punto de conexión, sino además cómo la actividad momento a momento del cerebro contribuye a la integridad del sistema. Si añadimos a esto la singularidad de cada cerebro, provocado en parte debido a la singularidad de la historia de la vida de cada persona, la predicción de Kandel comienza a sonar demasiado optimista. (En un reciente artículo en *The New York Times*, el neurocientífico Kenneth Miller sugirió que nos llevará 'siglos' sólo el conocer la conectividad neuronal básica.)

Mientras tanto, se están recaudando grandes sumas de dinero para la investigación del cerebro, basadas en algunos casos en ideas y promesas defectuosas que no se pueden cumplir. El caso

más flagrante de la neurociencia que salió mal, documentado recientemente en un informe en Scientific American, se refiere al Proyecto de Cerebro Humano (Human Brain Project) de 1.300 millones de dólares lanzado por la Unión Europea en 2013. Convencido por el carismático Henry Markram de que podía crear una simulación de todo el cerebro humano en una supercomputadora para el año 2023, y que tal modelo revolucionaría el tratamiento de la enfermedad de Alzheimer y otros trastornos, los organismos de la UE financiaron su proyecto prácticamente sin restricciones. En menos de dos años, el proyecto se convirtió en un "naufregio cerebral", y se le pidió a Markram que renunciara a su puesto.

Somos organismos, no ordenadores. Asumámoslo. Sigamos con la tarea de tratar de entendernos a nosotros mismos, pero sin ser penalizados por una carga intelectual innecesaria. La metáfora de la PI ha tenido un plazo de medio siglo, produciendo pocas ideas, si las hay, en el camino. Ha llegado el momento de pulsar la tecla DELETE.

Robert Epstein

Psicólogo senior de investigación en el Instituto Americano de Investigación
y Tecnología del Comportamiento en California.
Es autor de 15 libros, y ex editor en jefe del *Psychology Today*.

Aportación de Carolina Richter