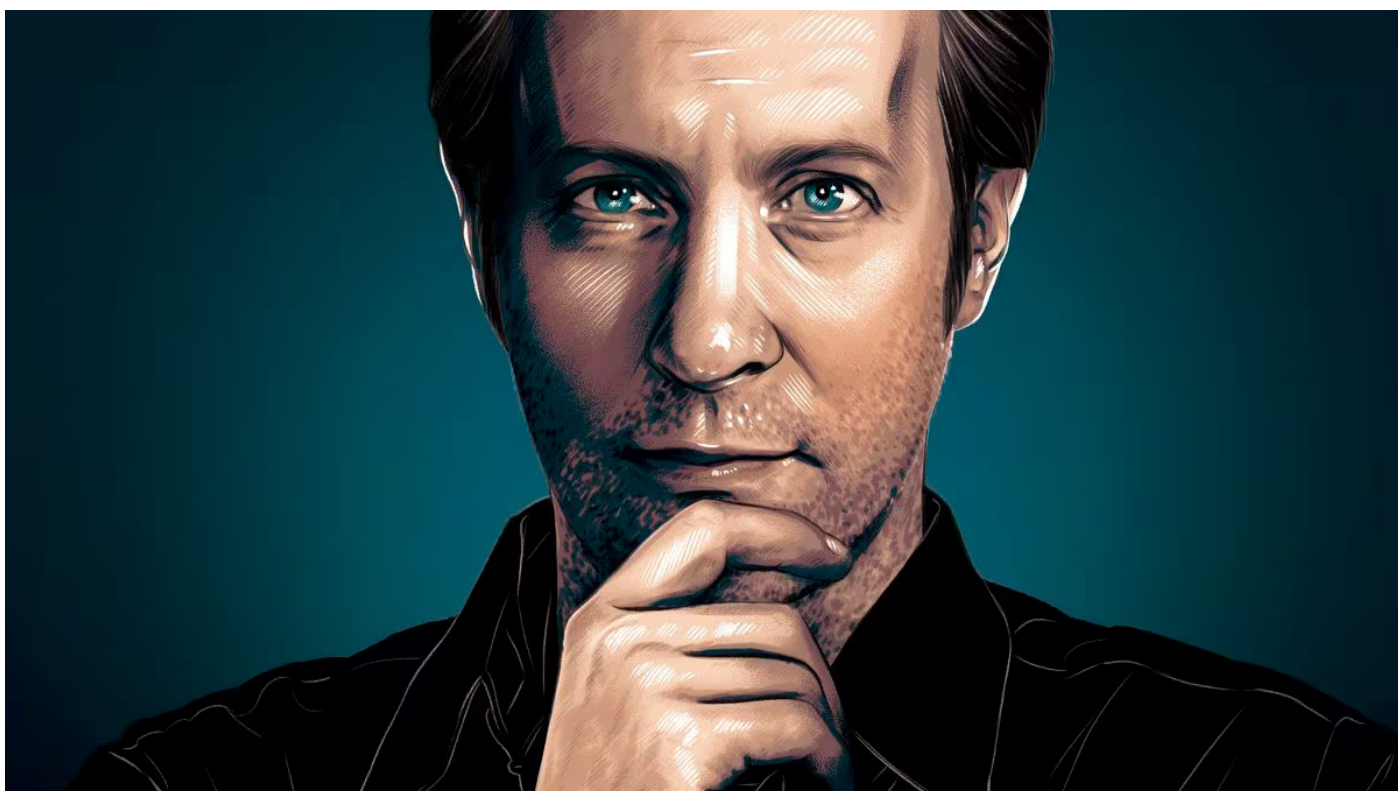


EN PORTADA >

David Eagleman, el investigador de los secretos de nuestro cerebro

El neurocientífico estadounidense, uno de los escritores científicos más interesantes de nuestro tiempo, explica que cada vez que aprendemos algo, nuestras neuronas cambian. Y que si alguien pierde la vista, parte de las células con las que veía le ayudarán, por ejemplo, a que oiga mejor



ALEXANDRA ESPAÑA



JAVIER SAMPEDRO

04 FEB 2024 - 05:30 CET

     9 

Nada del cerebro le es ajeno a David Eagleman, neurocientífico, tecnólogo, empresario y uno de los escritores científicos más

interesantes de nuestro tiempo. Nacido en Nuevo México hace 52 años, investiga en plasticidad cerebral, sinestesia, percepción del tiempo y lo que él llama [neuroderecho, en la intersección entre el conocimiento del cerebro y sus implicaciones legales](#). Su libro [Incógnito. Las vidas secretas del cerebro](#), de 2011, se tradujo a 28 idiomas, incluido el español en Anagrama, y ahora vuelve en la misma editorial con otra obra ambiciosa, [Una red viva](#), que gira en torno a una idea fundamental de la neurociencia actual: que el cerebro está en cambio permanente para adaptarse a la experiencia y al aprendizaje. Su ciencia no solo es de primera clase, sino también de primera mano, pero su escritura brillante y cristalina —un perfecto reflejo de su mente— convierte uno de los asuntos más complejos de la investigación actual en un paseo triunfal para el lector. Conversamos con él por videoconferencia en la primera entrevista que concede a un medio español en la última década.

¿Podría el cerebro de un recién nacido aprender a vivir [en un mundo de cinco dimensiones](#)? “Aún no sabemos cuánto hay de genética y cuánto de experiencia en nuestro cerebro”, responde desde California por videoconferencia. “Si criáramos un bebé en un mundo de cinco dimensiones, lo que seguramente sería un experimento muy poco ético, puede que encontráramos que el niño puede averiguar cómo desenvolverse ahí. [El tema general de la plasticidad cerebral](#) es que todo es más sorprendente de lo que pensábamos, en el sentido de que el cerebro es un diseño capaz de aprender cualquier cosa que el entorno le presente”.

MÁS INFORMACIÓN

Robert J. Sternberg, psicólogo: “Tus habilidades intelectuales cambian, la inteligencia se va aprendiendo”

Eagleman saca de alguna parte un voluminoso cuenco de ensalada, se lleva a la boca lo que cabe en el tenedor y prosigue con su argumento: “Su ejemplo del mundo de cinco dimensiones [es altamente hipotético](#), pero lo que sí sabemos, por supuesto, es que los bebés nacidos en cualquier lugar del mundo, ya sea en una cultura hiperreligiosa o en un país laico, en una economía basada en la agricultura o en otra superdesarrollada tecnológicamente, como aquí en Silicon Valley,

ajustan su cerebro a cualquiera de esos entornos. [Mis hijos se adaptan a usar una tableta o un móvil](#) igual de bien que otros niños, en otros lugares, se adaptan a las herramientas agrícolas. Así que sí, sabemos que nuestro cerebro es realmente muy flexible”.

La distinción entre genética y experiencia, o entre naturaleza y crianza, no es tan nítida como puede parecer. Lo habitual es pensar que los genes construyen el cerebro y que después el entorno toma el relevo modificando la fuerza de las conexiones entre neuronas (sinapsis) o estableciendo nuevos contactos. Pero formar nuevas conexiones y modular las viejas requiere [reactivar los mismos genes que construyeron el cerebro en primer lugar](#). Le planteo esta cuestión.

“La forma en que pensamos sobre la biología del cerebro es que [la experiencia implica cambios a todos los niveles](#), de manera que sí, tiene usted razón, los genes tienen que estar implicados en la plasticidad cerebral. Es cierto que lo habitual es investigar la plasticidad al nivel de las células neuronales y de las sinapsis que forman, refuerzan o debilitan, pero la principal razón es que eso es más fácil de medir. La experiencia modifica el cerebro a todos los niveles, y [la distinción entre sinapsis y genes no es real](#), sino una frontera arbitraria trazada por nosotros, los observadores humanos”, responde.

Un siglo largo de neurología ha establecido que el córtex (o corteza cerebral), la capa exterior que le da al cerebro su aspecto arrugado, está dividida en cientos de áreas especializadas: en ver, en oír, en hablar, en proyectar, [en gestionar las emociones](#) y todo lo demás. Sin embargo, los anatomistas no han encontrado grandes diferencias entre la arquitectura de circuitos de unas áreas y otras, y tampoco se conocen genes específicos de cada zona. ¿Qué quiere decir esto? Una de las lecciones del nuevo libro de Eagleman es que el cerebro es la misma cosa en todas partes. “El córtex usa el mismo truco, la misma arquitectura de circuitos, en cualquier área. La única razón por la que observamos distinciones — [esta área se dedica a la información visual, esta otra a la audición](#)— es porque cada una recibe distintos cables de entrada”, afirma el neurocientífico.

Por ejemplo, la información de los ojos entra por el nervio óptico hasta la

parte de atrás del cerebro, y por eso esa zona se convierte en lo que llamamos córtex visual, pero si te quedas ciego ese mismo córtex se convierte en auditivo, táctil u otras cosas. Así que [no hay nada fundamental en esa compartimentación del cerebro](#), según explica Eagleman: es solo una cuestión de qué cables de entrada están enchufados a un área o a otra, es decir, de qué tipo de información recibe.

Misteriosa evolución

La evolución del cerebro es de momento tan misteriosa como su funcionamiento. [Los seis millones de años que nos separan de los chimpancés](#) son apenas un pestañeo en las escalas evolutivas, y algunos científicos creen que la clave está en el mero aumento del tamaño del córtex, que se ha triplicado respecto a los chimpancés y los australopitecos. Eagleman es uno de ellos. “[Nuestro córtex es mucho mayor que el de cualquiera de nuestros primos animales](#), y esto es una gran parte del *cambio mágico* respecto a ellos. Hay otros cambios paralelos, como una gran laringe que nos permite comunicarnos con rapidez a través del lenguaje hablado, o un pulgar oponible que es una gran ayuda también, pero la diferencia principal es el tamaño de nuestro córtex”.

El científico prosigue: “Eso implica que hay mucho más territorio entre el *input* y el *output* (la información de entrada y de salida), de modo que, cuando uno recibe información sensorial y tiene que emitir una respuesta, en la mayoría de los animales esas dos áreas están muy cerca una de otra, [pero en nuestro caso están más separadas](#). El resultado es que, cuando ves algo, puedes tomar otro tipo de decisiones. Si me pones comida delante, puedo tomar en consideración que estoy a dieta, o que no quiero comer eso ahora porque estoy haciendo ayuno intermitente, o lo que sea, antes de tirarme sobre la comida”.

Eagleman enseña neurociencia en la Universidad de Stanford, California, pero su trabajo como investigador y docente se le queda muy corto a su mente profunda e inquieta. Es el jefe ejecutivo de [Neosensory](#), una empresa que contribuyó a fundar y que se dedica a desarrollar tecnología para que los ciegos y los sordos puedan recuperar parte de

sus facultades [reclutando zonas del cerebro que normalmente se dedican a otras cosas para sustituir al sentido perdido](#). También es el jefe científico de [BrainCheck](#), una plataforma digital para ayudar a los médicos a diagnosticar problemas cognitivos.

Además de escribir unos libros divulgativos de enorme calidad, escribe y presenta la serie de televisión [The Brain with David Eagleman](#) (el cerebro con David Eagleman) y el *podcast* [El cosmos interior con David Eagleman](#). Si hay una columna vertebral de toda esa actividad frenética, es aprovechar el conocimiento del cerebro para ayudar a la medicina de formas innovadoras y creativas.

“La consciencia es el gran misterio no resuelto de la neurociencia. Hay quien cree que solo son algoritmos”

La siguiente pregunta era inevitable. Y sí, Eagleman ha usado ChatGPT. Asegura que la parte fascinante de estos modelos grandes de lenguaje (*large language models*, LLM, el tipo de sistemas al que pertenece ChatGPT) es que ahora mismo [estamos más en un tiempo de descubrimiento que de invención](#). De la mayoría de las cosas que hemos inventado en el pasado —una lavadora o una cafetera— sabemos exactamente cómo funcionan, puesto que las hemos ideado nosotros, afirma. [Pero estos LLM están llenos de sorpresas](#) y hacen cosas que nadie esperaba, ni siquiera sus programadores. “Esto es asombroso. Creo que lo que hacen realmente muy bien es hallar conexiones en las que no habíamos pensado. Los modelos LLM han leído todo lo que se ha publicado en el mundo, tienen una memoria que lo abarca todo y pueden encontrar vínculos insospechados si les haces la pregunta correcta”.

Eagleman cree que esa habilidad de ChatGPT es extremadamente valiosa para la ciencia. “Cada mes salen 30.000 nuevos papers (artículos científicos revisados por pares) y yo no puedo leer todo eso, pero el LLM sí puede. He publicado un trabajo hace poco en el que propongo que hay

descubrimientos científicos de dos niveles. Los de nivel uno consisten en reunir cosas que yo simplemente no sabía, y ChatGPT puede ser útil ahí. Pero eso es diferente de los descubrimientos de nivel dos, que requieren imaginar un modelo que no existe”.

“Albert Einstein”, sigue Eagleman, “se preguntó [cómo vería la luz si estuviera montado en un fotón](#), y ese experimento mental le condujo a la teoría especial de la relatividad. Lo que hizo no fue reunir cosas que ya estaban en la literatura científica, sino imaginar un nuevo modelo y desarrollar una simulación con él. Y eso no estoy tan seguro de que la inteligencia artificial pueda hacerlo de momento. Creo que por eso los científicos tenemos trabajo todavía”.



El científico David Eagleman en su laboratorio en la Baylor College of Medicine, en Houston, Texas, en el año 2009
JOE BARABAN

Pero Eagleman también es un escritor. Y también cree que conservará el trabajo en esa otra faceta suya, porque, mientras que ChatGPT puede

escribir respuestas “sorprendentemente molonas” sobre cuestiones diversas, [como escritor no es particularmente creativo](#). “Usted y yo podemos estructurar los párrafos o hacer una llamada a un pasaje anterior, y ChatGPT está lejos de eso, al menos de momento. [Así que no, tampoco estoy preocupado como escritor](#)”, afirma.

Pero le digo que la IA le puede copiar:

—Usted es un maestro en encontrar analogías, metáforas y ejemplos ilustrativos. Quizá el modelo pueda analizar sus libros e imitarle en todo eso—, le pregunto.

—[Escribir es duro, usted lo sabe](#) —responde cortésmente—. Me asombraría realmente que ChatGPT llegara a escribir un buen libro en 10 años, pero no estoy convencido de que pueda hacerlo. Escribir un buen libro implica poner juntas nuevas ideas, nuevos modelos, pensar sobre ellos y preguntarse: ¿qué tipo de historia puedo contar para empezar este capítulo e introducir ese concepto? ¿Cómo enlazarlo con lo que vendrá luego y luego y luego? Mientras escribo un libro estoy pensando en todos esos niveles a la vez, y en cómo puede ser la experiencia del lector, y cómo hacer una referencia a lo anterior, cómo está sonando el ritmo, [es como si estuvieras componiendo una sinfonía](#). Como ChatGPT y el resto de LLM solo *piensan* en qué palabra escribir a continuación, no pueden pensar en todos los niveles a la vez.

Eagleman explica que lo que resulta tan difícil de imaginar sobre ChatGPT es que “se ha leído todos los libros que han existido, todos los blogs y las páginas web, y recuerda todo eso. Lo que esto ilustra es que, a un nivel fundamental, nosotros somos también algo parecido a una máquina estadística, y que si copias esas estadísticas y sabes qué palabra va cerca de otra en cada texto que ha producido la humanidad, el resultado es mucho mejor de lo que habíamos imaginado”.

Mente/máquina

Una de las áreas de investigación de Eagleman son las interfaces mente/máquina, pequeños paneles de electrodos que se implantan en el

cerebro para ayudar a personas ciegas o sordas. [¿Cómo consiguen estos electrodos conectar con las neuronas correctas?](#)

“Sabemos el área que interesa controlar. Por ejemplo, si se trata de mover un brazo robótico, pinchas la zona del córtex que normalmente controla el brazo, y además no registras una neurona, sino una colección entera. En una situación normal, si alguien te pone un peso en la muñeca, no tardarás en modular tus órdenes cerebrales para mover el brazo en esa nueva situación y no tirar la taza de café, [y con los pacientes pasa lo mismo](#). Nuestro cerebro está acostumbrado a los cambios en el cuerpo, es como si dijera: este es el objetivo que quiero alcanzar, así que ¿cómo llego allí con lo que tengo? Por eso las interfaces mente/máquina no necesitan contactar con las neuronas exactas. La persona sabe que quiere mover un brazo robótico, y encuentra cómo hacerlo”.

El córtex está hecho de unas unidades repetitivas llamadas columnas, diminutas en superficie pero con millones de neuronas en una organización estereotipada de circuitos. Se repiten una y otra vez a lo largo de todo el córtex, dice Eagleman, “como las costillas de una serpiente”. Todo el mundo está interesado en saber qué hace esa columna, esa especie de unidad básica del cerebro. Hay muchos datos, pero no entendemos aún de qué va. Según él, “lo que hace la estructura del córtex es convertir los datos en representaciones que resulten útiles para actuar en el mundo, y ello requiere abstraer los detalles”.

¿Alcanzarán las máquinas una forma de consciencia? [“Todo el mundo tiene una opinión sobre ello](#), pero realmente no lo sabemos. La consciencia es el misterio central no resuelto de la neurociencia. Una idea es la hipótesis computacional, donde todo son algoritmos, y si replicamos esos algoritmos en silicio tendremos consciencia. Otras escuelas piensan que hay algo particular en la biología del cerebro que aún no hemos descubierto”.

Es increíble lo que da de sí una ensalada.