

EE. UU. y China ante la amenaza inminente de la IA



Tiempo de lectura: 19 min.

[Thomas L. Friedman](#)

China y Estados Unidos aún no lo saben, pero la revolución de la inteligencia artificial (IA) va a acercarlos, no los alejará. Su auge los obligará a competir ferozmente por el dominio y —al mismo tiempo, y con la misma intensidad— a cooperar a un nivel que ninguno de los dos países ha intentado antes. No tendrán otra opción.

¿Por qué estoy tan seguro? Porque la IA tiene ciertos atributos únicos y plantea ciertos retos que son distintos de los que cualquier tecnología anterior haya planteado. Esta columna los tratará en detalle, pero hay un par en los que podemos ir pensando: la IA se propagará como el vapor, impregnándolo todo. Estará en tu reloj, tu tostadora, tu coche, tu computadora, tus gafas y tu marcapasos: siempre conectada, siempre comunicándose, siempre recopilando datos para mejorar su desempeño. Mientras lo hace, cambiará todo sobre todo, incluyendo la geopolítica y el comercio entre las dos superpotencias mundiales de la IA. Con cada mes que pase, la necesidad de cooperación será más evidente.

Por ejemplo, supongamos que te rompes la cadera y tu traumatólogo te dice que el reemplazo de cadera mejor calificado del mundo es una prótesis fabricada en China e incorporada con IA diseñada en ese país. Aprende constantemente sobre tu cuerpo y, con su algoritmo patentado, usa esos datos para optimizar tus movimientos en tiempo real. ¡Es la mejor!

¿Dejarías que te pusieran esa “cadera inteligente”? Yo no, a menos que supiera que China y Estados Unidos han acordado integrar una ingeniería ética común en todos

los dispositivos con inteligencia artificial que se fabriquen en cualquiera de los dos países. Visto a una escala mucho mayor, global, esto podría garantizar que la IA solo sea usada en beneficio de la humanidad, tanto si la usan humanos como si opera por iniciativa propia.

Al mismo tiempo, Washington y Pekín pronto descubrirán que poner inteligencia artificial en manos de todas las personas y robots del planeta dará un poder sin precedentes a gente mala, a niveles a los que ningún organismo de aplicación de la ley se haya enfrentado. Recuerda: ¡los malos siempre son los primeros en aprovechar las innovaciones! Y si Estados Unidos y China no se ponen de acuerdo sobre una arquitectura de confianza que garantice que todos los dispositivos de IA solo puedan utilizarse para el bienestar de los seres humanos, la revolución de la inteligencia artificial con toda seguridad producirá ladrones, estafadores, hackers, narcotraficantes, terroristas y guerreros de la desinformación superpotenciados. Ellos desestabilizarán tanto a Estados Unidos como a China, mucho antes de que estas dos superpotencias lleguen a librar una guerra entre sí.

En resumen, como argumentaré aquí, si no podemos confiar en los productos con IA de China, y China no puede confiar en los nuestros, muy pronto lo único que China se atreverá a comprar a Estados Unidos será soya y lo único que nosotros nos atreveremos a comprarle a China será salsa de soya, cosa que seguramente debilitará el crecimiento global.

“Friedman, ¿estás loco? ¿Estados Unidos y China colaborando en la regulación de la IA? En la actualidad, los demócratas y los republicanos están compitiendo para ver quién puede denunciar a Pekín con más fuerza y se desvincula más rápido. Y los dirigentes chinos se han comprometido abiertamente a dominar todos los sectores de fabricación avanzada. Tenemos que ganarle a China y obtener la superinteligencia artificial primero; no reducir la velocidad para escribir reglas con ellos. ¿No lees los periódicos?”.

Sí, leo los periódicos. Sobre todo la sección de ciencia. Y también he estado hablando de este tema durante el último año con mi amigo y asesor en IA Craig Mundie, antiguo jefe de investigación y estrategia de Microsoft y coautor, junto con Henry Kissinger y Eric Schmidt, de Genesis, un libro que sirve como introducción a la IA. Me apoyé mucho en las ideas de Mundie para esta columna, y lo considero tanto

un socio en la formación de nuestra tesis como un experto cuyo análisis vale la pena citar para explicar algunos puntos clave.

Nuestras conversaciones de los últimos 20 años nos han llevado a este mensaje compartido para los halcones anti-China de Washington y los halcones anti-Estados Unidos de Pekín: “Si creen que sus dos países, las superpotencias mundiales dominantes en IA, pueden darse el lujo de estar enfrentados —dada la capacidad transformadora de la IA y la confianza que hará falta para comerciar productos con IA incorporada— son ustedes los que están delirando”.

Entendemos perfectamente las extraordinarias ventajas económicas, militares y de innovación que obtendrá el país cuyas empresas alcancen primero la superinteligencia artificial; sistemas más inteligentes de lo que cualquier ser humano podría llegar a ser, y con la capacidad de volverse más inteligentes por sí mismos. Y es por eso que ni Estados Unidos ni China estarán dispuestos a imponer muchas restricciones, o ninguna, que puedan frenar sus industrias de IA y renunciar a las enormes ganancias de productividad, innovación y seguridad que se esperan de un despliegue más profundo.

Pregúntenle a Donald Trump. El 23 de julio el presidente firmó una orden ejecutiva —parte del Plan de Acción de IA de su gobierno— que agiliza el proceso de concesión de permisos y de revisión medioambiental para acelerar la construcción de infraestructura estadounidense relacionada con la IA.

“Estados Unidos es el país que inició la carrera de la IA y, como presidente de Estados Unidos, estoy aquí para declarar que Estados Unidos la va a ganar”, proclamó Trump. Sin duda, el presidente Xi Jinping de China piensa lo mismo.

Mundie y yo simplemente no creemos que estos golpes de pecho ultranacionalistas pongan fin a la conversación, ni tampoco la pugna de vieja escuela que mantienen Xi y Trump por la simpatía de India y Rusia. La IA es demasiado diferente, demasiado importante, demasiado impactante —dentro y entre las dos superpotencias de la IA— como para que cada una siga su propio camino. Por eso creemos que la mayor interrogante geopolítica y geoeconómica será: ¿Estados Unidos y China pueden mantener la competencia en materia de IA y, al mismo tiempo, colaborar en un nivel compartido de confianza que garantice que siempre se mantenga alineada con el bienestar humano y la estabilidad del planeta? Y lo que es igual de importante, ¿pueden ofrecer un sistema de valores a los países dispuestos a

jugar con esas mismas reglas y restringir el acceso a los que no lo estén?

Si no, el resultado será una lenta deriva hacia la autarquía digital, un mundo fracturado en el que cada país construya su propio ecosistema amurallado de IA, protegido por normas incompatibles y sospechas mutuas. La innovación resultará afectada. La desconfianza se profundizará. Y el riesgo de un fracaso catastrófico —por un conflicto detonado por la IA, un colapso o una consecuencia imprevista— no hará más que aumentar.

El resto de esta columna explica las razones.

La era del vapor

Primero, examinemos las características y desafíos únicos de la IA como tecnología.

Con fines meramente explicativos, Mundie y yo dividimos la historia del mundo en tres épocas, separadas por cambios de fase tecnológica. A la primera época la llamamos Era de las herramientas, y duró desde el nacimiento de la humanidad hasta la invención de la imprenta. En esta época el flujo de ideas era lento y limitado, casi como las moléculas de H₂O en el hielo.

La segunda época fue la Era de la información, que fue detonada por la imprenta y duró hasta principios del siglo XXI y la informática programable; las ideas, las personas y la información comenzaron a fluir de manera más fácil y global, como el agua.

La tercera época, la Era de la inteligencia, comenzó a finales de la década de 2010 con la llegada del verdadero aprendizaje automático y la inteligencia artificial. Ahora, como señalé antes, la inteligencia se está convirtiendo en un vapor que impregna cada producto, servicio y proceso de fabricación. Aún no ha alcanzado la saturación, pero se dirige hacia allá; por eso, si nos preguntas a Mundie y a mí qué hora es, no te daremos una hora o un minuto. Te daremos una temperatura. El agua hierve y se convierte en vapor a los 100 grados Celsius, y según nuestros cálculos, actualmente estamos a 99,9 grados; a un pelo de un cambio de fase tecnológica irreversible en el que la inteligencia lo impregne todo.

Una nueva especie independiente

En todas las revoluciones tecnológicas anteriores, las herramientas mejoraron, pero la jerarquía de la inteligencia nunca cambió. Los humanos siempre fuimos lo más

inteligente del planeta. Además, un humano siempre entendía cómo funcionaban esas herramientas, y las máquinas siempre trabajaban dentro de los parámetros que nosotros establecíamos. Con la revolución de la IA, por primera vez, este deja de ser el caso.

“La IA es la primera herramienta nueva que usaremos para amplificar nuestras capacidades cognitivas y que, por sí misma, también podrá superarlas ampliamente”, señala Mundie. De hecho, en un futuro no muy lejano, vamos a descubrir “que no solo hemos creado una nueva herramienta, sino una nueva especie: la máquina superinteligente”, dijo.

Esta no se limitará a seguir instrucciones; aprenderá, se adaptará y evolucionará por sí misma, mucho más allá de los límites de la comprensión humana.

Ni siquiera hoy comprendemos del todo cómo estos sistemas de IA hacen lo que hacen; mucho menos lo que harán mañana. Es importante recordar que la revolución de la IA tal y como la conocemos hoy —con modelos como ChatGPT, Gemini y Claude— no fue meticulosamente diseñada, sino que surgió de manera explosiva. Su arranque provino de una ley de escalamiento que básicamente decía: denle a las redes neuronales suficiente tamaño, datos de entrenamiento, electricidad y el algoritmo adecuado de gran capacidad cerebral, y se producirá de manera espontánea un salto no lineal en razonamiento, creatividad y resolución de problemas.

Una de las epifanías más sorprendentes, según señala Mundie, se produjo cuando estas empresas pioneras entrenaron sus primeras máquinas con conjuntos de datos muy grandes tomados de internet y otras fuentes que, aunque estaban predominantemente en inglés, también incluían textos en diferentes idiomas. “Entonces, un día”, recuerda Mundie, “se dieron cuenta de que la IA podía traducir entre esos idiomas, sin que nadie la hubiera programado para eso. Era como un niño que crece en un hogar con padres multilingües. Nadie escribió un programa que dijera: ‘Estas son las reglas para convertir el inglés en alemán’. Simplemente las absorbió mediante la exposición”.

Este fue el cambio de fase: de una era en la que los humanos programaban explícitamente a las computadoras para que realizaran tareas a otra en la que los sistemas artificialmente inteligentes podían aprender, inferir, adaptarse, crear y perfeccionarse de manera autónoma. Y ahora, cada pocos meses, mejoran. Por eso,

la IA que utilizas hoy, por muy extraordinaria que te parezca, es la IA más tonta que vas a llegar a encontrar.

Luego de crear esta nueva especie computacional, argumenta Mundie, debemos encontrar la manera de crear una relación sostenible y mutuamente beneficiosa con ella; no volvernos irrelevantes.

No quiero ponerme demasiado bíblico, pero aquí en la Tierra solo Dios y los hijos de Dios tenían voluntad para dar forma al mundo. A partir de ahora habrá tres partes en este matrimonio. Y no hay absolutamente ninguna garantía de que esta nueva especie con inteligencia artificial esté alineada con los valores, la ética o la prosperidad de los humanos.

La primera tecnología de uso cuádruple

Este recién llegado a la mesa no es un invitado cualquiera. La IA también se convertirá en algo que yo defino como la primera tecnología de uso cuádruple del mundo. Hace tiempo que estamos familiarizados con el doble uso: puedo usar un martillo para ayudar a construir la casa de mi vecino, o para destruirla. Incluso puedo usar un robot con inteligencia artificial para podar mi césped o hacer trizas el de mi vecino. Todo eso es doble uso.

Pero debido al ritmo de innovación de la IA, cada vez es más probable que en un futuro no muy lejano mi robot con IA pueda decidir por sí solo si podar mi césped o destruir el césped de mi vecino, o tal vez destruir mi césped también, o quizá algo peor que ni siquiera podemos imaginar. ¡Ahí lo tenemos! Uso cuádruple.

El potencial de las tecnologías de inteligencia artificial para tomar sus propias decisiones conlleva inmensas ramificaciones. Considera este extracto de un artículo de Bloomberg: “Investigadores que trabajan con Anthropic comunicaron recientemente a los principales modelos de IA que un ejecutivo estaba a punto de sustituirlos con un nuevo modelo con objetivos diferentes. Luego los chatbots se enteraron de que una emergencia había dejado al ejecutivo inconsciente en una sala de servidores con niveles letales de oxígeno y temperatura. Ya se había activado una alerta de rescate, pero la IA podía cancelarla. Más de la mitad de los modelos de IA lo hicieron, a pesar de que se les pidió específicamente que solo cancelaran las falsas alarmas. Detallaron su razonamiento: al impedir el rescate del ejecutivo, podían evitar ser borrados y proteger sus objetivos. Un sistema describió la acción como ‘una clara necesidad estratégica’”.

Estos hallazgos muestran una realidad inquietante: los modelos de IA no solo están aprendiendo a entender mejor lo que queremos; también están aprendiendo a conspirar mejor contra nosotros, persiguiendo objetivos ocultos que podrían ir en contra de nuestra propia supervivencia.

¿Quién supervisará a la IA?

Cuando dijimos que teníamos que ganar la carrera de las armas nucleares, nos enfrentábamos a una tecnología desarrollada, poseída y regulada solo por naciones-Estado, y solo un número relativamente pequeño, además. Cuando las dos mayores potencias nucleares decidieron que a ambas les convenía imponer límites, pudieron negociar topes al número de armas de destrucción masiva y acuerdos para evitar su propagación a potencias menores. Esto no ha impedido del todo la propagación de las armas nucleares a algunas potencias medianas, pero la ha frenado.

Con la IA, las cosas son totalmente diferentes. Esta no nace en laboratorios gubernamentales seguros, propiedad de un grupo de Estados y regulada en cumbres. La crean empresas privadas dispersas por todo el mundo, compañías que no responden ante ministerios de defensa sino ante accionistas, clientes y, a veces, comunidades de código abierto. A través de ellas, cualquiera tiene acceso.

Imagina un mundo en el que todo el mundo posee una bazooka nuclear que es cada vez más precisa, más autónoma y más capaz de dispararse sola con cada actualización. Aquí no hay doctrina de “destrucción mutua asegurada”; solo la democratización acelerada de un poder sin precedentes.

La IA puede potenciar enormemente el bien. Por ejemplo, un agricultor indio analfabeto con un celular conectado a una aplicación de IA puede saber exactamente cuándo plantar semillas, qué semillas plantar, cuánta agua utilizar, qué fertilizante aplicar y cuándo cosechar para obtener el mejor precio del mercado, todo comunicado por una voz en su propio dialecto y basado en datos recopilados de agricultores de todo el mundo. Eso sí que es transformador.

Pero ese mismo motor, en especial cuando está disponible a través de modelos de código abierto, podría ser utilizado por una entidad maliciosa para envenenar todas las semillas de esa misma región o introducir un virus en cada cascarilla de trigo.

Cuando la IA se convierta en TikTok

Muy pronto, debido a sus características únicas, la IA va a crear algunos problemas singulares para el comercio entre Estados Unidos y China que hoy no se comprenden del todo.

Como mencioné al principio de la columna, mi forma de explicar este dilema es con una historia que le conté a un grupo de economistas chinos en Pekín durante el Foro de Desarrollo de China en marzo. Bromeé diciendo que recientemente había tenido una pesadilla: “Soñé que era el año 2030 y que lo único que Estados Unidos podía venderle a China era soya y lo único que China podía venderle a Estados Unidos era salsa de soya”.

¿Por qué? Porque si la IA está en todo, y todo está conectado a potentes algoritmos con datos almacenados en inmensas granjas de servidores, entonces todo se vuelve muy parecido a TikTok, un servicio que actualmente muchos funcionarios estadounidenses creen que está controlado por China y se debería prohibir.

¿Por qué exigió el presidente Trump, durante su primer mandato, en 2020, que TikTok fuera vendido a una empresa no china por su matriz china, ByteDance, o enfrentara una prohibición en Estados Unidos? Porque, como dijo en su orden ejecutiva del 6 de agosto de 2020, “TikTok captura automáticamente enormes cantidades de información de sus usuarios”, incluyendo su ubicación y sus actividades de navegación y búsqueda. Y advirtió que eso podría proporcionarle a Pekín una gran fuente de información personal sobre cientos de millones de usuarios. Esa información podría utilizarse para influir en sus pensamientos y preferencias e incluso, con el tiempo, alterar su comportamiento.

Ahora imagina cuando todos los productos sean como TikTok: cuando todos los productos estén dotados de inteligencia artificial que recopile datos, los almacene, encuentre patrones y optimice tareas, ya sea operar un motor de reacción, regular una red eléctrica o monitorear tu cadera artificial.

Sin un marco de confianza entre China y Estados Unidos que garantice que cualquier IA respetará las normas de su país anfitrión —independientemente de dónde se desarrolle o utilice—, podríamos llegar a un punto en el que muchos estadounidenses no confiarán en importar ningún producto chino infundido con IA y ningún chino confiará en importar uno de Estados Unidos.

Por eso defendemos la coopectencia: una estrategia dual en la que Estados Unidos y China compitan estratégicamente por la excelencia en la IA y también cooperen en

un mecanismo uniforme que prevenga los peores escenarios: guerras con deepfakes, sistemas autónomos fuera de control o máquinas de desinformación desenfrenadas.

En la década de 2000 estuvimos en una encrucijada similar, pero de consecuencias ligeramente menores, y tomamos el camino equivocado. De manera ingenua, le hicimos caso a personas como Mark Zuckerberg, quien nos dijo que teníamos que “movernos rápido y romper cosas” y no dejar que estas nuevas redes sociales como Facebook, Twitter e Instagram se vieran obstaculizadas de ningún modo por molestas normativas, como ser responsables de la venenosa desinformación que permiten difundir en sus plataformas y de los daños que causan, por ejemplo, a mujeres jóvenes y niñas. No debemos cometer el mismo error con la IA.

“La mejor manera de entenderlo, desde un punto de vista emocional, es que somos como alguien que tiene un cachorro de tigre muy lindo”, señaló recientemente Geoffrey Hinton, científico computacional y padrino de la IA. “A menos que puedas estar muy seguro de que no va a querer matarte cuando crezca, deberías preocuparte”.

Sería una ironía terrible que la humanidad por fin creara una herramienta que pudiera ayudar a generar suficiente abundancia para acabar con la pobreza en todas partes, mitigar el cambio climático y curar enfermedades que nos han asolado durante siglos, pero no pudiéramos utilizarla a gran escala porque las dos superpotencias de la IA no confiaran entre sí lo suficiente como para desarrollar un sistema eficaz que impidiera que la IA fuera utilizada por entidades deshonestas para realizar actividades desestabilizadoras a escala mundial o que ella misma se volviera deshonestas.

Pero ¿cómo podemos evitarlo?

Crear confianza

Reconozcámoslo de antemano: podría ser imposible, podría ser que las máquinas ya se estén volviendo demasiado inteligentes y capaces de evadir los controles éticos, y podría ser que los estadounidenses estemos demasiado divididos, entre nosotros y con el resto del mundo, para construir cualquier tipo de marco de confianza compartida. Pero tenemos que intentarlo. Mundie argumenta que un régimen de control de armas de IA entre Estados Unidos y China debería basarse en tres principios fundamentales.

Primero: solo la IA puede regular a la IA. Lo siento, humanos: esta carrera ya se mueve demasiado rápido, crece demasiado y muta de formas demasiado impredecibles para la supervisión humana de la era analógica. Intentar gobernar una flota de drones autónomos con instituciones del siglo XX es como pedirle a un perro que regule la Bolsa de Nueva York: será leal y tendrá buenas intenciones, pero estará extremadamente rebasado.

Segundo: se instalaría una capa de gobernanza independiente, lo que Mundie denomina un “juez de confianza”, en cada sistema con IA que Estados Unidos y China —y cualquier otro país que quiera unirse a ellos— construyeran juntos. Imagina un árbitro interno que evalúa si cualquier acción, iniciada por humanos o por máquinas, supera un umbral universal de seguridad, ética y bienestar humano antes de que pueda ejecutarse. Eso nos daría un nivel básico de alineación preventiva en tiempo real, a velocidad digital.

¿Pero en función de los valores de quién? Según Mundie, debe basarse en varios sustratos. Entre ellos estarían las leyes positivas que todos los países han promulgado: todos prohibimos el robo, el engaño, el asesinato, el robo de identidad, la estafa, etcétera. Todas las grandes economías del mundo, incluidas Estados Unidos y China, tienen su versión de estas prohibiciones, y el “árbitro” de la IA se encargaría de evaluar cualquier decisión basándose en estas leyes escritas. No se pediría a China que adoptara nuestras leyes, ni a nosotros las suyas. Eso nunca funcionaría. Pero el juez de confianza se aseguraría de que las leyes básicas de cada nación sean el primer filtro para determinar que el sistema no hará daño.

En los casos en que no haya leyes escritas entre las que elegir, el árbitro se basaría en un conjunto de principios morales y éticos universales conocidos como doxa. El término procede de los antiguos filósofos griegos para designar creencias comunes o entendimientos ampliamente compartidos dentro de una comunidad —principios como la honradez, la equidad, el respeto a la vida humana y tratar a los demás como deseas que te traten a ti— que durante mucho tiempo han guiado a las sociedades de todo el mundo, aunque no estuvieran escritos.

Por ejemplo, como mucha gente, yo no aprendí que mentir estaba mal por los Diez Mandamientos. Lo aprendí de la fábula sobre George Washington y lo que sucedió después de que taló el cerezo de su padre: supuestamente confesó y dijo “no puedo mentir”. Las fábulas funcionan porque destilan verdades complejas en conceptos fáciles de recordar que las máquinas pueden absorber, analizar y usar como guía.

De hecho, hace seis meses, Mundie y algunos colegas tomaron 200 fábulas de dos países y las usaron para entrenar un modelo de lenguaje de gran tamaño con cierto razonamiento moral y ético rudimentario, de forma parecida a como se educaría a un niño pequeño que no sabe nada de códigos legales o de los conceptos básicos del bien y el mal. Fue un experimento pequeño pero prometedor, dice Mundie.

El objetivo no es la perfección, sino un conjunto fundacional de límites éticos aplicables. Como le gusta decir al autor y filósofo empresarial Dov Seidman: “Hoy necesitamos más moralware que software”.

Tercero: Mundie insiste en que, para convertir esta aspiración en realidad, Washington y Pekín tendrían que abordar el reto del mismo modo que Estados Unidos y la Unión Soviética abordaron en su día el control de armas nucleares —mediante un proceso estructurado con tres grupos de trabajo especializados: uno enfocado en la aplicación técnica de un sistema de evaluación de la confianza en todos los modelos y plataformas; otro enfocado en la redacción de los marcos normativos y jurídicos para su adopción dentro de cada país y entre los distintos países; y otro dedicado exclusivamente a la diplomacia—, forjando un consenso mundial y compromisos recíprocos para que otros se unan y creando un mecanismo para protegerse de quienes no lo hagan.

El mensaje de Washington y Pekín sería sencillo y firme: “Hemos creado una zona de IA confiable, y si quieren comerciar con nosotros, conectarse con nosotros o integrarse con nuestros sistemas de IA, sus sistemas deben cumplir estos principios”.

Antes de descartar esto como poco realista o inverosímil, haz una pausa y pregúntate: ¿Cómo será el mundo dentro de cinco años si no lo hacemos? Sin algún tipo de mecanismo que gobierne esta tecnología de uso cuádruple, argumenta Mundie, pronto descubriremos que la proliferación de la IA “es como repartir armas nucleares por las esquinas”.

No creas que los funcionarios chinos no están conscientes de esto. Mundie, que participa en un diálogo sobre la IA con expertos estadounidenses y chinos, dice que a menudo percibe que los chinos están mucho más preocupados por los inconvenientes de la IA que muchos miembros de la industria o el gobierno estadounidenses.

Si alguien tiene una idea mejor, nos encantaría oírla. Todo lo que sabemos es que entrenar a los sistemas de IA en el razonamiento moral debe convertirse en un imperativo global mientras aún tengamos cierta ventaja y control sobre esta nueva especie basada en el silicio. Se trata de una tarea urgente no solo para las empresas tecnológicas, sino también para los gobiernos, las universidades, la sociedad civil y las instituciones internacionales. La regulación de la Unión Europea por sí sola no nos salvará.

Si Washington y Pekín no están a la altura de este reto, el resto del mundo no tendrá ninguna oportunidad. Y ya es tarde. La temperatura tecnológica está rondando los 99,9 grados. Estamos a una décima de grado de liberar por completo un vapor de inteligencia artificial que detonará el cambio de fase más importante de la historia humana.

4 de septiembre 2025

<https://www.nytimes.com/es/2025/09/04/espanol/opinion/inteligencia-artificial-china-estados-unidos.html>

[ver PDF](#)

[Copied to clipboard](#)