

Fecha: 17-04-2026
 Medio: El Mercurio
 Supl.: El Mercurio - Cuerpo B
 Tipo: Noticia general
 Título: Estados Unidos despierta al peligroso poder de la IA

Pág.: 6
 Cm2: 495,6

Tiraje: 126.654
 Lectoría: 320.543
 Favorabilidad: ☐ No Definida

¿Debe confiarse la tecnología nueva más potente del mundo a un puñado de hombres? Cinco *geeks* tan famosos que pueden ser identificados por su nombre de pila —Dario, Demis, Elon, Mark y Sam— ejercen un control casi divino sobre los modelos de inteligencia artificial (IA) que moldearán el futuro. El gobierno de Trump se ha mantenido al margen incluso cuando esos modelos han adquirido capacidades asombrosas, convencido de que la competencia sin trabas entre firmas privadas es la mejor manera de asegurar que Estados Unidos gane la carrera de la IA frente a China.

Hasta ahora. De pronto, el trato desenfadado de Estados Unidos hacia la IA parece estar llegando a su fin. La razón es que el vertiginoso avance de los modelos también representa una amenaza para la propia seguridad nacional estadounidense, inquietando a miembros del gobierno de Trump que antes estaban más inclinados a preocuparse por una sobrerregulación. Al mismo tiempo, el creciente resentimiento entre los votantes estadounidenses está convirtiendo a la IA en un foco de conflicto político. Un enfoque *laissez-faire* ya no es políticamente viable ni estratégicamente sensato.

El punto de inflexión fue el anuncio de Claude Mythos por parte de Anthropic el 7 de abril. La más reciente creación de la firma es tan sorprendentemente buena para encontrar vulnerabilidades de *software* que, en manos equivocadas, pondría en riesgo infraestructura crítica, desde bancos hasta hospitales. Los modelos de IA plantean cada vez más otros riesgos también, desde peligros para la bio-

The Economist:

Estados Unidos despierta al peligroso poder de la IA

Después de Mythos, un enfoque *laissez-faire* ya no es políticamente viable ni estratégicamente sensato.



Dario Amodei,
de Anthropic.



Demis Hassabis,
de Google DeepMind.



Elon Musk,
de xAI.



Mark Zuckerberg,
de Meta.



Sam Altman,
de OpenAI.

seguridad hasta estafas a escala industrial.

El jefe de Anthropic, Dario Amodei, pensó con buen criterio que Mythos era demasiado peligroso para un lanzamiento general. En vez de eso, lo ha reservado para el uso de unas 50 grandes firmas, de los sectores de computación, *software* y finanzas, para que refuercen sus propias defensas. El secretario del Tesoro de Estados Unidos, Scott Bessent, quedó tan inquieto que convocó a los mayores bancos a conversaciones urgentes.

No era la primera vez que el gobierno actuaba. Hace apenas unas semanas, el Pentágono intervino después de que Amodei se negara a permitir que el modelo de Anthropic fuera usado en armas completamente autónomas o para vigilancia masiva dentro del país. También entonces el gobierno de Trump se alarmó por el poder que una sola firma ejercía sobre una tecnología central para la seguridad nacional.

Una reacción adversa entre los votantes aumentará la presión sobre el gobierno para intervenir. Las encuestas de opinión están llevando a cada vez más políticos a pensar que la IA será uno de los grandes temas de las elecciones de 2028. Los estadounidenses son mucho más es-

cepticos respecto de la IA que la gente de otros países. Siete de cada diez creen que la IA perjudicará las oportunidades laborales, un alza pronunciada respecto de hace un año, y eso mucho antes de que exista buena evidencia. La oposición de base a los centros de datos está aumentando con fuerza, aunque la IA tiene poco o nada que ver con el alza de los precios de la electricidad. Como signo de los tiempos, la casa de Sam Altman, jefe de OpenAI, ha sido atacada dos veces en los últimos días.

La historia sugiere que, con una tecnología tan transformadora como la IA, un momento Mythos era inevitable. Desde John D. Rockefeller hasta Henry

Ford, las grandes innovaciones industriales de Estados Unidos fueron lideradas por un pequeño número de hombres que acumularon un enorme poder. Finalmente, los gobiernos del siglo XX intervinieron para domesticar industrias excesivamente poderosas, desde la ofensiva antimonopolio que desmanteló Standard Oil hasta la creación de la Reserva Federal y la división de AT&T. Aquellos tiempos eran al menos tan polarizados y febriles como los actuales. Y nuestros cálculos sugieren que los dioses de la IA todavía no son más dominantes que sus predecesores históricos.

Pero la historia también sugiere que controlar la IA estará lleno de dificultades. En parte, porque lo que está en juego si algo sale mal es enorme. También, porque la IA está evolucionando a velocidad de vértigo.

Las compensaciones son agudas. El crecimiento económico se beneficiará de difundir rápidamente las ventajas de la IA, pero la posible reacción adversa podría llevar fácilmente a una sobrerregulación. No hacer nada podría dejar a Estados Unidos vulnerable a un caos malicioso inducido por la IA, pero un exceso regulatorio aseguraría que China gane la carrera de la IA. Eso hace de este un momento peligroso.

El tiempo es escaso. Hace dos años, durante el gobierno de Biden, las discusiones sobre regulación se centraban en gran medida en los riesgos potenciales de la IA. Hoy, sus capacidades ya son alarmantemente poderosas y se vuelven aún más poderosas con cada lanzamiento. El ritmo de la innovación significa que los debates sobre el papel adecuado del gobierno, que en el pasado se



desarrollaron durante años e incluso décadas, ahora deben resolverse en meses.

Y los obstáculos técnicos para un enfoque más intervencionista son formidables. Las herramientas de control estatal, como la nacionalización, son ineficaces porque los ingenieros talentosos pueden moverse libremente entre compañías y la capacidad de cómputo es una mercancía. Peor aún, los principales desarrolladores de modelos están solo unos meses por delante de sus competidores de código abierto, incluidos los de China. Tarde o temprano, las capacidades de sus modelos estarán disponibles para todos.

Aun así, el momento Mythos podría ser el instante en que comience a tomar forma un esquema viable para controlar la IA. Los usuarios de confianza obtendrían acceso temprano a los nuevos modelos más poderosos: OpenAI está siguiendo a Anthropic al desplegar su última herramienta a un grupo limitado de profesionales de ciberseguridad previamente evaluados. Antes de permitir que estos modelos se comercialicen de manera amplia, el gobierno podría exigir certificación de organismos liderados por la industria que los hayan probado para distintos usos.

Cuidado con los *geeks* cargados de regalos

Esta idea tiene ventajas tanto para los grandes desarrolladores de modelos como para el gobierno. Evita el proceso largo de crear un nuevo regulador. Al permitir solo a unos pocos usuarios *premium*, hace posible que los desarrolladores cobren precios más altos y limiten el uso de una capacidad de cómputo escasa. Mientras tanto, el gobierno puede restringir quién puede usar los modelos más poderosos, reduciendo el riesgo de que China los copie y se ponga al día más rápido.

Pero también adolece de problemas graves. Un lanzamiento limitado reducirá la competen-

cia y aumentará la influencia de las compañías de IA ya consolidadas. Ralentizará la difusión de los beneficios de la IA y creará un sistema de dos niveles dentro de la economía estadounidense, perjudicando a las muchas firmas que se verán repetidamente privadas de un acceso temprano privilegiado a los nuevos modelos más poderosos. ¿Qué pasa si desarrollar defensas para la IA toma mucho tiempo o es imposible? ¿Qué pasa con los modelos de código abierto? ¿Cómo se puede exigir que también sigan estas reglas?

Un sistema regulatorio construido sobre estas bases podría resultar injusto. Los de adentro podrían protegerse frente a amenazas de frontera; los de afuera tendrían que esperar lo mejor. Las oportunidades para el *lobby* y para ganancias desproporcionadas serían inmensas. Eso pondría a prueba la honestidad y la competencia del gobierno más abiertamente corrupto de la era política moderna de Estados Unidos. Y una solución que concentre todavía más poder y riqueza entre el puñado de dioses de la IA corre el riesgo de agravar precisamente la reacción política que está comenzando a inquietar a Washington.

Además, el enfoque Mythos solo puede ser la mitad de la solución. La seguridad de la IA no puede garantizarse a escala nacional. Con el tiempo exigirá cooperación internacional, empezando por China. El nuevo foco en la ciberseguridad también debe ir acompañado de una reflexión urgente sobre los efectos económicos y sociales de la IA. Enfrentar la disrupción sobre el empleo y diseñar un sistema tributario adaptado a la IA que favorezca el trabajo son problemas enormes para los que nadie tiene todavía buenas respuestas. Eso tiene que cambiar. El momento Mythos es una señal de alarma para la seguridad de la IA. También exige una reflexión dura en otras áreas.

Artículo traducido del inglés por "El Mercurio".