

EL NUEVO MIEDO

¿Por qué ahora todo el mundo tiene miedo a la inteligencia artificial?

ECONOMÍA

16_09_2026

**Daniele
Ciacchi**



Como ya hemos tenido ocasión de tratar aquí en la *Brújula Cotidiana* (*La Nuova Bussola Quotidiana* en su edición original en italiano) Jacob Coxon, ingeniero que acaba de dimitir de Anthropic, ha escrito sin rodeos en X que cualquiera que desarrolle

inteligencia artificial alberga en su interior la duda de que esta podría exterminarnos antes de que termine la década. Evan Hubinger, también de Anthropic, ha confirmado una probabilidad superior al 10%. A partir de ahí se ha desatado –y con razón– la campaña mediática con titulares apocalípticos por doquier, una nueva “Guerra de los mundos” de Wells pero con un enemigo de casa, alarma de extinción inminente y, una vez más, la humanidad como verdugo de sí misma (los algoritmos no se han creado solos). El libro de Eliezer Yudkowsky y Nate Soares, *If Anyone Builds It, Everyone Dies*, publicado en septiembre de 2025, entró en la lista de los libros más leídos del *New York Times* con una tesis bastante clara sobre la relación entre la creación de una superinteligencia y la muerte de la humanidad.

Sin embargo, estos rumores no son simples cotilleos de bar sino que provienen de expertos del sector y eso tiene su peso. No obstante, una oleada tan compacta y oportuna suscita sospechas. Es cierto que la IA plantea riesgos reales, pero quizá distintos de cómo los entendemos, y detrás de este catastrofismo se vislumbran al menos tres intereses concretos y una ideología general que sirve de marco.

En mayo de 2023 Sam Altman pidió explícitamente ante el Senado estadounidense una nueva agencia federal con competencias para conceder licencias en materia de IA. El senador Peter Welch y el testigo Gary Marcus advirtieron de inmediato del riesgo de lo que se denominó “captura regulatoria”, que consiste en establecer umbrales de cálculo y requisitos de seguridad adaptados a empresas con miles de millones de dólares, excluyendo de hecho a los competidores más pequeños y a los productores de *open source* (código abierto), consolidando así el duopolio Microsoft/OpenAI. **David Sacks**, ex responsable de IA de la Casa Blanca, ha declarado sin temor a ser desmentido que Anthropic estaría llevando a cabo una estrategia de *captura regulatoria* aprovechando el alarmismo apocalíptico.

Además, existe el clásico mantra de “lo importante es que se hable de ello” tan querido por el mundo de la publicidad. Meredith Whittaker, presidenta de Signal, ha señalado que “las narrativas sobre el riesgo existencial son, de hecho, publicidad para una tecnología que hoy en día solo poseen unas pocas empresas”. El periodista **Brian Merchant** observa que afirmar que un producto es tan potente como para acabar con el mundo equivale a atribuirle un poder desmesurado: es otra forma de marketing.

El *Alpocalypse* desvía el foco de atención hacia un peligro probable, y no inminente, alejándolo del debate público sobre los verdaderos problemas que la IA debería poner hoy en el punto de mira del debate público.

Mientras se debate sobre la extinción para 2030, quedan al margen temas de relevancia social como los sesgos algorítmicos, la vigilancia generalizada de los datos personales, la explotación del trabajo creativo que hay detrás del etiquetado de los modelos y el consumo de agua y energía de los *centros de datos*. Se trata de daños medibles y reales, no de hipótesis futuristas propias de una novela de ciencia ficción. Centrar la atención colectiva en un escenario apocalíptico e indefinido tiene la ventaja, para quienes lo alimentan, de posponer a un futuro lejano las cuentas que deberíamos hacer hoy mismo.

Existe también una razón geopolítica detrás de esta carrera por el miedo. Dario Amodei, en la CBS, ha afirmado que, en el fondo, la carrera hacia la Inteligencia Artificial Generativa se debe al miedo ante **el desarrollo paralelo de China en este mismo ámbito**. Trump ha sido tajante al afirmar que quien gane la carrera por la IA lo ganará todo. Un informe del FBI del 8 de septiembre acusa a los laboratorios chinos de haber “destilado” miles de millones de tokens a partir de los modelos estadounidenses para acortar la distancia, reducida hoy a solo un 2,7% según Stanford. Quienes piden que se frene el avance piden, al mismo tiempo, que se detenga a China. El miedo a la extinción debería discurrir por las vías de la carrera armamentística y de la IA para uso militar.

Más que hablar de una verdadera conspiración, hay que hablar aquí de intereses convergentes. No hace falta un plan coordinado cuando todos comen en la misma mesa. Sin embargo, hay una ideología colectiva que actúa como aglutinante, la del “altruismo eficaz”, que hoy en día impregna los laboratorios y las *organizaciones sin ánimo de lucro* que dictan la agenda de la seguridad de la IA. El mecanismo es muy sencillo: si la extinción es un desenlace posible, cualquier medida política monopolística se convierte en moderada si el objetivo es evitarla. Solicitar una licencia federal ya no es una intromisión en el mercado, sino prudencia. Dar la voz de alarma a nivel mundial es un deber moral, no una estrategia de marketing. Ignorar los daños que la inteligencia artificial causa hoy en día no es negligencia, sino establecer las prioridades correctas frente al riesgo máximo de la desaparición de nuestro planeta. En la práctica, un círculo reducido reivindica una visión correcta y privilegiada del futuro —una especie de nuevo “soviético”—, mientras que las generaciones futuras, en cuyo nombre se habla, no pueden rebatir nada.