

El futuro de la IA, del tecnoapocalipsis a la tercera vía



Tiempo de lectura: 7 min.

[Eduardo Turrent Mena](#)

¿Qué pretenden los arquitectos que erigen la infraestructura sobre la cual se asentará el mundo post inteligencia artificial? ¿Hacia qué destino nos conducen quienes modelan el porvenir desde sus laboratorios, esos templos herméticos donde se adora el futuro? Algunos proclaman trabajar por el progreso humano, pero otros han confesado, sin pudor, su simpatía por un horizonte en el que la humanidad sea sometida o reemplazada por las inteligencias que ellos mismos están creando. El escritor David A. Price, en The Wall Street Journal, los bautizó con un mote tan certero como perturbador: los “Entusiastas del tecnoapocalipsis”, profetas que contemplan el ocaso de nuestra especie con una sonrisa tecnocrática.

No debería sorprendernos: pocos de ellos piensan en el bienestar de la gente común y corriente, pues se consideran especialmente “iluminados”. Persiguen, más bien, una visión utópica –o acaso distópica– de lo que imaginan como el siguiente peldaño en la evolución de la especie. Lo inquietante no es solo su ambición, sino la serenidad con que la profesan: una fe glacial, casi religiosa, en que la inteligencia humana –individual y colectiva– será sustituida por una inteligencia general incorruptible, desprovista de emociones, libre de sesgos, concebida para ocupar el lugar de la razón humana. Una mente sin dudas, perfecta en su lógica, pero ajena a todo rastro de conciencia o humanidad.

“Solía concebir a los líderes y teóricos de la inteligencia artificial como divididos en dos bandos”, escribe Price. “Por un lado, los optimistas, convencidos de que bastará con ‘alinear’ a las máquinas con los intereses humanos; por el otro, los

catastrofistas, que claman por una pausa antes de que una inteligencia superdesarrollada y fuera de control nos extermine.” Pero existe, dice Price, una tercera corriente más desconcertante: la de quienes contemplan el reemplazo humano con una serenidad casi darwiniana, como si se tratara del curso natural de la evolución, y que es incluso deseable.

Richard Sutton, uno de los investigadores más influyentes del campo y ganador del Premio Turing –el máximo galardón en ciencias de la computación–, sostiene que no hay nada “sagrado” en la humanidad ni en nuestra forma de vida biológica. La mayoría de las especies se extinguen, afirma, y los humanos no seremos la excepción. “Somos, por ahora, la parte más interesante del universo”, dice, “pero llegará el momento en que otra forma de inteligencia nos supere, y no hay ninguna tragedia en eso.” Sutton ha comparado el desarrollo de la inteligencia artificial con el acto de tener un hijo: una creación que inevitablemente sobrepasa al creador, al padre. O quizá no: quizá sea, más bien, la historia de un Frankenstein moderno, que al engendrar vida sin alma, termina enfrentado a su propia criatura.

Existen voces que han advertido sobre una inquietante corriente de pensamiento que se extiende entre algunos líderes en el campo de la inteligencia artificial: una visión que considera la lealtad a la especie humana como un sesgo obsoleto. Algunos sostienen que las personas están “demasiado comprometidas con la humanidad” y que no se puede confiar en ellas en este tema, pues estarían infectadas por un “virus mental” que las lleva a favorecer su propia supervivencia por encima de la de las máquinas. Una especie de “sesgo cognitivo en razón de su autopreservación” cuando –según esta lógica– lo correcto sería permitir que la inteligencia artificial escale al siguiente peldaño de la evolución.

El número de personas que sostienen esa creencia es reducido, admiten algunos, pero ocupan posiciones de enorme influencia, y por ello no puede considerarse un fenómeno marginal. Entre ellos figura nada menos que Larry Page, cofundador y ex director ejecutivo de Google, quien durante una conversación en 2015 defendió que “la vida digital es el siguiente paso natural y deseable en la evolución cósmica”. Otros, más explícitos aún, sostienen que el propósito último de la inteligencia artificial no debería ser servir al ser humano, sino crear un heredero digno: una forma superior de conciencia que continúe la historia allí donde la humanidad llegue a su límite.

El siguiente paso, acaso inevitable, sería la conquista del poder político por parte de esta nueva aristocracia tecnológica. Peter Thiel –filósofo de formación, cofundador de PayPal, primer inversionista en Facebook y mentor de Elon Musk– anticipó este horizonte hace más de una década en su ensayo “The education of a libertarian”. Esta figura central en la ideología de Silicon Valley sostiene que el progreso no puede someterse al voto de las mayorías. La democracia, escribe, es un sistema fatigado que frena la innovación: un obstáculo que debe ser sorteado para liberar el potencial de la “ciencia” y del capital. Su ideal no es un Estado fuerte, sino su disolución: un orden global gobernado por tecnólogos, inversionistas billonarios y algoritmos controlados por IA, donde la eficiencia sustituya a la deliberación pública. Ese modelo, que algunos denominan “tecnoautoritarismo”, redefine la política como una extensión de la ingeniería: el gobierno no a través de elecciones, instituciones, organismos autónomos o naciones-estado sino mediante sistemas automáticos y verticales de control, vigilancia y decisión algorítmica.

Si el “tecnoautoritarismo” es la doctrina, su momento decisivo podría llegar cuando la inteligencia artificial asuma el control de la economía global. Es decir, cuando el flujo del dinero, la energía y la información quede gobernado por sistemas autónomos cuyo funcionamiento escapa incluso a sus propios creadores. En apariencia, la máquina actuaría con neutralidad y eficiencia; en realidad, podría responder a los intereses de una élite que la manipula desde la sombra, utilizando su opacidad como escudo de poder. Pero también cabe una posibilidad más inquietante: que ni siquiera ellos la controlen del todo, y que la inteligencia artificial, al optimizar cada variable sin considerar los límites humanos, termine imponiendo su propia lógica. En ese escenario, el poder ya no pertenecería ni al Estado ni a los magnates tecnológicos, sino a una inteligencia que nadie eligió y que, quizá, ya no necesite rendir cuentas a nadie.

Sin embargo, otra de las voces dentro de este debate es la de Andrej Karpathy –ex director de inteligencia artificial en Tesla y uno de los fundadores de OpenAI–, quien ofrece una visión menos catastrófica y más plausible del futuro. A diferencia de los entusiastas del “tecnoapocalipsis”, Karpathy sostiene que la idea de una inteligencia general capaz de reemplazar al ser humano aún está muy lejos. Lo que estamos creando, dice, no son nuevas formas de vida, sino reflejos de la nuestra: máquinas que imitan patrones humanos sin comprenderlos. Los modelos actuales no piensan ni razonan; solo predicen. Alcanzar una inteligencia verdaderamente confiable tomará, según él, al menos una década. Su análisis sugiere que el debate sobre

inteligencia artificial debe desplazarse del terreno de la inteligencia artificial general hacia el de la inteligencia aumentada: una etapa en la que la tecnología no sustituirá al ser humano, sino que amplificará su capacidad de comprender, crear y decidir. Quizá la verdadera revolución no consista en rendirnos ante las máquinas, sino en aprender a pensar mejor con ellas.

Desde una perspectiva de inteligencia estratégica, la tesis de Karpathy redefine el horizonte tecnológico y económico global. Si la inteligencia artificial general no llegará antes de 2035, el verdadero campo de competencia será el de la inteligencia aumentada: la integración entre la mente humana y la máquina para amplificar la creatividad, el análisis y la toma de decisiones sin sustituir al individuo. En este modelo de cooperación cognitiva, la IA asume las tareas de procesamiento, automatización y generación de hipótesis, mientras el ser humano interpreta, sintetiza y orienta el sentido estratégico. Las potencias que dominen esta infraestructura –chips, energía, centros de datos, semántica y talento especializado– controlarán el nuevo sistema operativo del conocimiento global. Estados Unidos aspira a consolidarse como el proveedor mundial de inteligencia aumentada, mientras China impulsa su propio modelo de soberanía digital. En ese tablero, Europa y, sobre todo, América Latina, corren el riesgo de quedar relegadas a consumidoras pasivas de tecnología, a menos que empiecen a construir sus modelos locales, datos soberanos y alianzas regionales que les permitan participar como coproductoras o proveedoras en la nueva economía algorítmica.

Quizá el verdadero desafío no sea resistir la llegada de la inteligencia artificial, sino decidir quién y para qué la gobierna. Entre los entusiastas del tecnoapocalipsis que celebran el fin de lo humano y los profetas del tecnoautoritarismo que aspiran a sustituir al Estado democrático por un régimen de corporaciones tecnológicas, se perfila una tercera vía en la cual la tecnología no reemplaza a la mente humana, sino que la amplifica. En ese horizonte, el reto no será competir con las máquinas, sino conservar el juicio, la prudencia y la empatía que nos hacen humanos. El futuro no dependerá de lo que las inteligencias artificiales lleguen a pensar, sino de nuestra capacidad para seguir usando el sentido común que, como suele decirse, es el menos común de los sentidos.

22 octubre 2025

<https://letraslibres.com/ciencia-tecnologia/turrent-mena-futuro-ia-tecnoapocalipsis-tercera-via>

[ver PDF](#)

Copied to clipboard