

Homo agenticus: ¿la nueva frontera de la inteligencia artificial?

Por David Valle-Cruz, Vanessa Fernandez-Cortez

De acuerdo con Wallace, se tienen precedentes de inteligencia desde hace 200 mil años con la aparición del Homo sapiens (Marmelada, 2003). Asimismo, existe la postura emergentista que sostiene que la inteligencia surge de manera gradual desde hace unos 2 a 1.5 millones de años (Hiraiwa-Hasegawa, 2019). En esta evolución, solo hace 70 años, el ser humano creó una inteligencia sintética que intenta imitar lo biológico (Kurzweil, 2000).

Un hito en esta historia es la famosa reunión en el Dartmouth College, en 1956, donde se planteó “la conjetura de que todos los aspectos del aprendizaje o de cualquier otra característica de la inteligencia pueden, en principio, ser descritos de modo tan preciso que se pueda construir una máquina capaz de simularlos”. La idea fue tan innovadora que sigue siendo un reto para investigadores y expertos en tecnología. Es aquí en donde surge el término inteligencia artificial (IA). Desde entonces, la IA ha tenido grandes avances, como el perceptrón (que simula una neurona biológica) y los sistemas generativos que permiten crear poemas, pinturas y películas mediante técnicas de aprendizaje profundo. Además, puede servir como asistente para profesionales de todas las áreas en actividades tanto rutinarias como creativas.

De acuerdo con Raymond Kurzweil (2000), la IA ha avanzado a un ritmo acelerado y, quizás, ha evolucionado más rápido que la inteligencia humana. Hechos que se consideraban solo ciencia ficción, como los presentados en la película *Ella*, en donde Theodore Twombly (interpretado por Joaquin Phoenix) se enamora de su sistema operativo inteligente: Samantha. Al respecto, entre 2023 y 2024 se hicieron públicos varios casos de personas que declararon estar enamoradas de ChatGPT-4, en lo que, inclusive, algunas personas realizaron ceremonias simbólicas de *matrimonio virtual* con el modelo de lenguaje.

De esta manera, puede notarse que, el avance de la IA no solo se ha visto reflejado en lo sentimental, ya que también ha sido útil en la realización de actividades repetitivas que no necesitan razonamiento especializado y que se podrían delegar a un agente inteligente. Este tipo de gestor “es una entidad que percibe y actúa sobre un entorno. Y un agente inteligente [... lo hace] de forma razonada” (Russell y Norvig, 2003). En computación, este agente nace de la construcción de entidades que perciben, razonan, aprenden y actúan para automatizar tareas complejas en entornos cambiantes. Dichos agentes poseen “un grado de inteligencia suficiente para ejecutar parte de sus tareas de



Tecnología

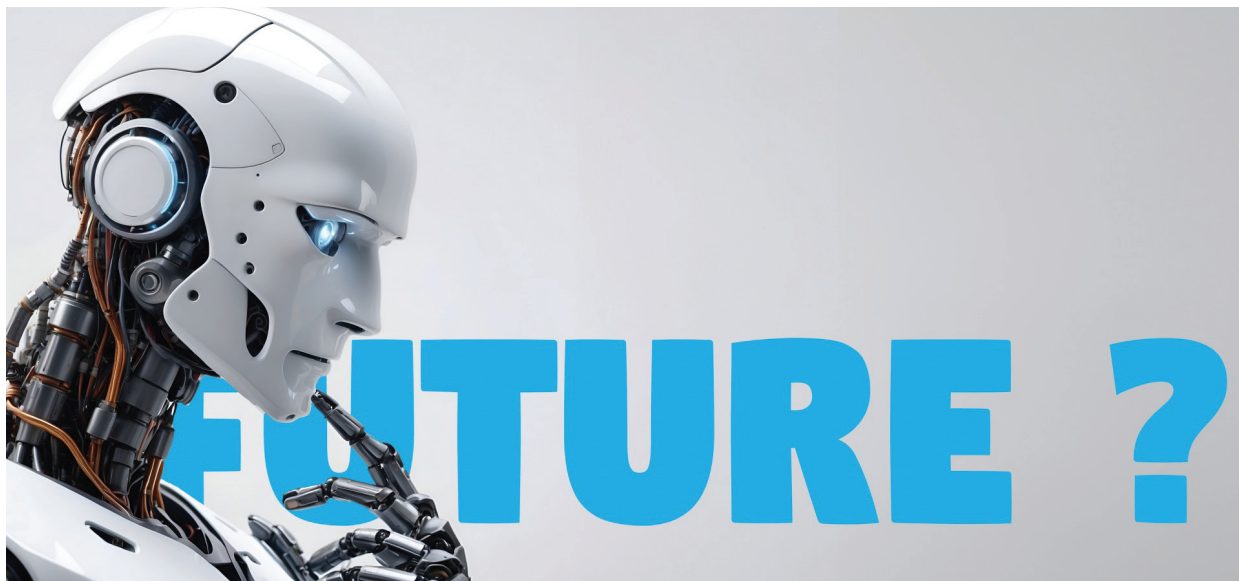


Imagen: publicdomainpictures.net

forma autónoma y para interactuar con su entorno de forma útil” (Brenner, Zarnekow y Wittig, 2012). En los noventa, el agente se volvió la *biblia* para entender a la IA (Russel y Norvig, 2003), pero con el auge del aprendizaje automático se perdió de vista. Actualmente, los modelos largos de lenguaje (como GPT, LLaMA, Gemini, Claude y Mistral) funcionan como cerebros lingüísticos que le dan a los agentes cuerpo, metas, planificación, acción, búsqueda, persistencia y coordinación con otros agentes. Todo esto fundamenta un área más de esta tecnología: la IA agentiva (*agentic AI*).

La IA agentiva es la forma que mejor captura la idea de una IA con capacidad de actuar, tomar decisiones, planificar, iniciar acciones y ejecutar tareas autónomamente dentro de un marco definido. Incluso se ha mencionado que tiene cierta intuición, no precisamente humana, al producir inferencias rápidas basadas en patro-

nes aprendidos (Bubeck et al., 2023).

De este modo, es posible que el siguiente peldaño del *homo sapiens* (con la fusión humano-máquina) pueda transformarlo en el *homo deus* (Harari, 2016). En torno a esta metáfora evolutiva, surge el término *homo agenticus*, que describe al humano cuya vida, decisiones y creatividad están integradas con agentes inteligentes autónomos. El *homo agenticus* no es una rama darwiniana, sino una etapa cultural, cognitiva y tecnológica dentro del *homo sapiens* que puede interactuar con el entorno como agente.

La IA agentiva tiene el potencial de brindar a los sistemas autonomía extendida, razonamiento generativo, memoria dinámica y contextual, toma de decisiones coordinada y capacidad de ejecutar tareas complejas sin supervisión paso a paso. Algunos casos internacionales de agentes son los agentes científicos de Google DeepMind, que pueden planear expe-

rimentos, proponer hipótesis, buscar información, optimizar rutas químicas y corregirse automáticamente. Asimismo, Amazon está desarrollando agentes destinados a logística global para ajustar inventarios, redirigir envíos, negociar tiempos de entrega en almacenes, predecir cuellos de botella y ejecutar acciones sin autorización humana. Quizá el caso más conocido es el piloto automático de los vehículos de Tesla, el cual emplea redes neuronales profundas para percibir el entorno, deliberar, decidir maniobras y ejecutar acciones en tiempo real sin intervención directa de una persona.

Con todo esto, se pone en la mesa el tema de agentes autónomos, una idea del siglo pasado, pero aumentada por modelos de lenguaje, que tiene un potencial especial al incluir artefactos basados en máquinas cognitivas y algoritmos con decisión propia. De acuerdo con Leonardi



(2025), esto crea una paradoja, pues cuánto más autónoma parece la IA, menos autónomos se sienten los humanos, independientemente del control real que posean. Sin embargo, este desplazamiento no es permanente ni unidireccional. La agencia circula a través de “bucles de agencia” previsible: patrones recursivos que involucran delegación, atribución, contingencia, reafirmación y reconfiguración, y todos pueden gestionarse estratégicamente.

Los humanos que comprendan cómo operan estos bucles de agencia e intervengan en ellos mediante procesos estratégicos de atribución podrían desarrollar nuevas formas de pericia y autoridad. El futuro pertenecerá a quienes entiendan que la agencia siempre se atribuye en algún lugar, y que el lugar donde decidimos ubicarla determina quién tiene el poder de actuar. El *homo agenticus* podría vislumbrarse como un humano

con esteroides de IA agentiva que reduce el tiempo dedicado a tareas analíticas lineales, se especializa en la metacognición, tiene mayor capacidad de tomar decisiones de alto nivel, es más creativo y tiene una mejor calidad de vida. Pero esto es algo que tendremos que esperar a que suceda. 

Referencias

- Brenner, Walter, Rüdiger Zarnekow y Hartmut Wittig (2012). *Intelligent software agents: foundations and applications*. Springer Science & Business Media.
- Bubeck, Sébastien et al. (2023). “Sparks of artificial general intelligence: Early experiments with GPT-4”, en *arXiv*. <<https://arxiv.org/abs/2303.12712>>.
- Harari, Yuval Noah (2016). *Homo Deus: Breve historia del mañana*. Debate.
- Hiraiwa-Hasegawa, Mariko (2019). “Evolution of intelligence on the earth”, en *Astrobiology: From the Origins of Life to the Search for Extraterrestrial Intelligence*. Singapore: Springer Singapore.
- Kurzweil, Ray (2000). *The age of spiritual machines: When computers exceed human intelligence*. Penguin.
- Leonardi, Paul M. (2025). “Homo agenticus in the age of agentic AI: Agency loops, power displacement, and the circulation of responsibility”, en *Information and Organization*, vol. 35, núm. 3, 100582.
- Marmelada, Carlos A. (2003). “Sobre el origen de la inteligencia humana” en *Aceprensa*, Servicio 6/03. <<https://www.unav.edu/web/ciencia-razon-y-fe/sobre-el-origen-de-la-inteligencia-humana>>.
- McCarthy, John et al. (1955). *A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence*.
- Russel, Stuart y Peter Norvig (2003). *Artificial intelligence—A modern approach*. Pearson Education.



David Valle-Cruz es Ingeniero en Computación, maestro en Informática y doctor en Ciencias Económico-Administrativas. Ha realizado estancias en SUNY Albany y CINVESTAV. Sus publicaciones científicas aparecen en revistas científicas internacionales como *Government Information Quarterly*, *Cognitive Computation* e *Information Polity*. Actualmente, es profesor en el Centro Universitario UAEMEX Tianguistenco e integrante del SNII nivel 1. En 2024, recibió el Christopher Pollitt Award al mejor artículo por su trabajo sobre los impactos negativos de la inteligencia artificial en el gobierno. Desde el mismo año, forma parte del top 2% de los científicos más influyentes del mundo, según el ranking de la Universidad de Stanford y Elsevier.



Vanessa Fernandez-Cortez es Candidata a Doctora en Finanzas, Maestra en Finanzas Corporativas y Licenciada en Contaduría. Actualmente, es profesora en la Facultad de Contaduría y Administración de la Universidad Autónoma del Estado de México. En 2020, fue galardonada con el premio a la mejor investigación en la Annual International Conference on Digital Government Research. Su investigación se centra en la optimización de portafolios de inversión, la educación financiera y la aplicación de inteligencia artificial para el análisis y la toma de decisiones.