

Landgrebe, Jobst y Smith, Barry, *Why Machines Will Never Rule the World. Artificial Intelligence without Fear*, Routledge Publisher, 2022, 344 páginas. ISBN 9781032309934

Lucio Nontol
Seton Hall University
nontollu@shu.edu

Desde el año 2014 y, tal vez, desde mucho antes, Nick Bostrom, filósofo y fundador del *Future of Humanity Institute* perteneciente a *Oxford University* defiende la idea de que un futuro próximo las operaciones de la mente humana serán postergadas por las súper-capacidades de las máquinas, lo que para algunos parece factible. A esta corriente filosófica se la conoce con el nombre de transhumanismo.¹ Bostrom es demasiado optimista con el potencial de los inventos tecnológicos, cree que mejorarán y posiblemente cambiarán la vida de los seres humanos. Las ofertas que, a juicio de los transhumanistas, la humanidad va a obtener no parecen nada desdeñables:

«Extensión del tiempo de vida, la erradicación de enfermedades, la eliminación de sufrimientos innecesarios y el aumento de las capacidades intelectuales, físicas y emocionales... El enfoque transhumanista no se limita únicamente a gadgets o medicamentos, sino que abarca también diseño de modelos económicos, sociales e institucionales; y el desarrollo de aspectos culturales y habilidades y técnicas psicológicas»².

En el año 2020 Elon Musk pronosticó que dentro de los próximos cinco años la Inteligencia Artificial superaría a la inteligencia humana y se convertiría en un «dictador inmortal» sobre toda la humanidad.

Pareciera que Bostrom y Musk dejan el futuro de la humanidad en manos de la Inteligencia Artificial (en adelante IA), como si la naturaleza humana estuviera inconclusa y necesitara ser moldeada por un «ente» al estilo IA. Además, para estos pensadores, la humanidad ha entrado en una crisis debido a que es incapaz de manejar las máquinas que ella misma ha creado y está siendo derrotada por su propio progreso. El entusiasmo que acoge este movimiento con respecto a las tecnologías no va a la par con el pesimismo que

¹ N. BOSTROM, *Superintelligence. Paths, Dangers, Strategies*, Oxford University Press, Oxford, 2014.

² <https://nickbostrom.com/ethics/values.html>

muestran a la hora de proponer y defender valores éticos esenciales. Pese a que insinúan la necesidad de un enfoque ético de la IA no la desarrollan.

Sin embargo, en medio de toda la euforia tecnológica que rodea a pensadores como Bostrom y Musk entre otros muchos, surge un libro coescrito titulado *Why Machines Will Never Rule the World. Artificial Intelligence without Fear*, (*¿Por qué las máquinas nunca gobernarán el mundo? Inteligencia Artificial sin Miedo*) cuya tesis principal podría resumirse en lo siguiente: «El aprendizaje automático y todas las demás aplicaciones de *software* funcionales están lejos de lograr cualquier cosa que se parezca a la capacidad de los humanos. En términos prácticos, no la acercará a la posibilidad de pleno funcionamiento del cerebro humano».

La demostración de Landgrebe y Smith adquiere relevancia debido a los argumentos interdisciplinarios que los conduce a plantear su propuesta. A partir de un análisis exhaustivo basado en la Física, la Biología, la Lingüística, la filosofía de la mente, incluso considerando los progresos actuales de la IA, afirman que construir una IA requiere emular en *software* el tipo de sistemas que manifiestan inteligencia a nivel humano. La inteligencia a nivel humano es el resultado de un sistema complejo y dinámico que no se puede modelar matemáticamente de manera que permita ser emulado por una computadora. Por lo tanto, lograr una IA, tal como plantean sus entusiastas (al menos a través de las computadoras), es imposible. Antes de seguir ahondando en el libro que estamos reseñando conviene decir algo sobre quiénes son los autores.

Jobst Landgrebe es médico, bioquímico y matemático. Sus estudios de Matemática le sirvieron para llegar a la conclusión de que existe un gran desajuste entre los métodos matemáticos y los datos biológicos. A sus 30 años abandonó la academia y se convirtió en consultor de negocios y empezó a trabajar en sistemas de *software* de IA. Intentaba construir sistemas de IA para imitar lo que los seres humanos pueden hacer, el resultado fue negativo: descubrió el mismo problema que había encontrado años antes en la biología.

Barry Smith es un filósofo estadounidense. Tiene el título de *Julian Park Distinguished Professor* de Filosofía en la Universidad de Buffalo y es investigador en el *Research Scientist* del estado de Nueva York. De 2002 a 2006 fue director del *Institute for Formal Ontology and Medical Information Science* en Leipzig y Saarbrücken, Alemania.

Dos prominentes investigadores con una trayectoria reconocida nos ofrecen una obra inédita. Su propuesta se encamina a enfatizar la importancia de IA sin llegar ni a un optimismo exagerado ni a un pesimismo derrotista. Su enfoque se orienta, más bien, hacia el realismo. Tras esta breve contextualización de obra nuestro objetivo en esta reseña no

es resumir el libro ni tampoco sumergirnos en detalles técnicos. Simplemente daremos una muestra de los argumentos principales de esta obra que está dividida en tres partes precedida por una introducción en la que resume la totalidad. La primera parte: *Propiedades de la mente humana* analiza el concepto de mente, su relación con el cuerpo y la complejidad (lenguaje, ética, ámbito social, etc.) de la mente humana y la IA. La segunda parte: *Límites de los modelos matemáticos* presenta la complejidad de los sistemas matemáticos más importantes de la historia, su relación con los sistemas de la IA y el futuro de dichos modelos. La tercera parte: *Los límites y el potencial de la IA* evidencia los argumentos que sostienen su propuesta: la imposibilidad de una IA que supere a los humanos. Y cuando dicen imposible lo dicen en serio. No se puede encontrar una solución; no por deficiencias en los datos, el *hardware*, el *software* o el cerebro humano, sino más bien por razones *a priori* de las matemáticas.

Smith y Landgrebe ofrecen un examen crítico de las proyecciones injustificadas de la IA tales como las máquinas que se separan de la humanidad o que se autorreplican y se convierten en «agentes plenamente éticos». Afirman, además, que no puede haber una especie de «voluntad de máquina». Cada aplicación de IA se basa en las intenciones de los seres humanos, incluidas las de producir resultados aleatorios.

Esto significa que la Singularidad, un punto en el que la IA se vuelve incontrolable e irreversible (como un momento *Skynet* de la franquicia cinematográfica «Terminator») nunca va a ocurrir. Las afirmaciones descabelladas en sentido contrario solo sirven para inflar el potencial de la IA y distorsionar la comprensión pública de la naturaleza, las posibilidades y los límites de la tecnología.

Traspassando las fronteras de varias disciplinas científicas, Smith y Landgrebe sostienen que la idea de una IA que iguale la inteligencia general de los humanos es imposible debido a los límites matemáticos de lo que se puede modelar y es «computable». Estos límites son aceptados prácticamente por todos los que trabajan en este campo; sin embargo, hasta ahora no han sabido apreciar las consecuencias que tienen sobre lo que una IA puede lograr. Como afirman los autores: «Superar estas barreras requeriría una revolución en las matemáticas que sería de mayor importancia que la invención del cálculo por Newton y Leibniz hace más de 350 años». Atendiendo a los argumentos de Smith y Landgrebe podríamos que decir que por razones matemáticas, la IA no puede imitar la forma en que funciona el cerebro humano. De hecho, los autores dicen que es imposible diseñar una máquina que rivalice con el rendimiento cognitivo humano. El libro insiste en

defender la idea de que no puede haber IA porque está más allá de los límites de lo que, en principio, es posible lograr mediante una máquina.

Finalmente, ¿deberíamos leer *Why Machines Will Never Rule the World*? Si eres un investigador dedicado a la IA o tienes un interés técnico en el tema, es posible que lo disfrutes. Es amplio y está impecablemente investigado, pero también es académico y, en ocasiones, exigente, por lo tanto, estamos delante de una obra interesante y novedosa que conviene leer para continuar el debate planteado por Smith y Landgrebe.