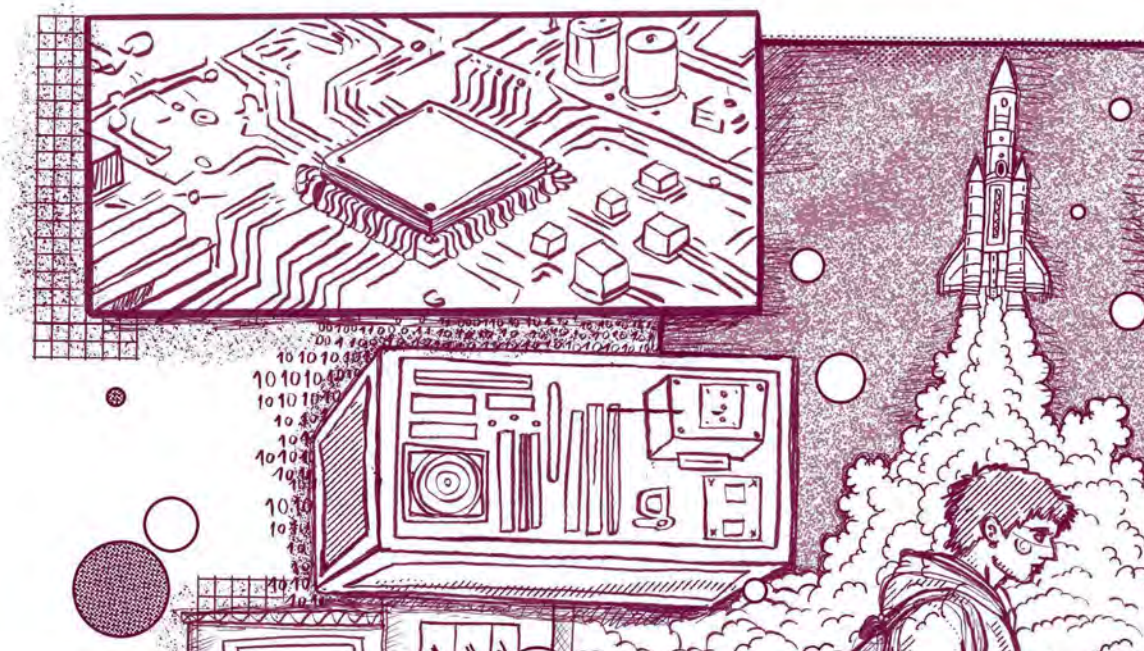




Espacio conceptual

La máquina mutante: breve historia de la IA

¿A qué llamamos Inteligencia Artificial? ¿Cómo funcionan estas máquinas? ¿Qué modelos entraron en disputa durante su desarrollo? El autor reconstruye la historia de esta tecnología e identifica el tablero mundial en el que se juegan sus últimos avances.

Imaginemos viajar por una autopista en un automóvil que cada diez kilómetros se transforma en otro auto: cambia el tablero, el motor, la carrocería... Ahora imaginemos que el conductor de ese automóvil conduce cada vez a mayor velocidad. Ese automóvil es la mal llamada “inteligencia artificial”, ese conductor somos nosotros, sus usuarios. Bajarse del auto es muy difícil, disminuir la velocidad suena poco tentador, confiar en que a la larga aprenderemos a manejar semejante máquina con destreza puede costarnos caro. Quizás si miramos la historia, la lógica y los impactos externos de esta máquina mutante, podamos manejarla mejor.

Las dos vidas de la IA

Por una paradoja de la historia, hoy llamamos “inteligencia artificial” al tipo de tecnología que fue pensada como lo contrario de la “inteligencia artificial”. En los años 40, Warren McCulloch y Walter Pitts intentaron replicar mediante una red de nodos eléctricos el funcionamiento de las neuronas cerebrales. El experimento era doble: por un lado, intentar que una máquina reprodujera algo similar al pensamiento humano; por otro, explicar el funcionamiento del pensamiento humano como si fuera una máquina. Las redes neuronales artificiales de Pitts y McCulloch

atrajeron el interés de matemáticos, biólogos, psicólogos y también del Departamento de Defensa norteamericano que, en el contexto de la Segunda Guerra Mundial y la posterior Guerra Fría, estaba interesado en desarrollar máquinas de guerra autónomas. El mismo año en que Pitts y McCulloch empezaron a experimentar con sus redes neuronales artificiales, Arturo Rosenbluth, Norbert Wiener y Julian Bigelow imaginaron una máquina autocorrectora capaz de modificar sus metas finales según sus propios errores. Wiener llamó “cibernética” al principio de retroalimentación de información que permitiría desarrollar máquinas autónomas que se adaptaran inductivamente a las señales del entorno, sin necesidad de poseer inteligencia interna.

El experimento cibernético contó con financiamiento del Ejército norteamericano hasta los años 60. La inversión era muy grande, pero los logros limitados: era imposible reponer artificialmente la cantidad de neuronas que tiene un cerebro humano, y aún más difícil reproducir su funcionamiento.

A su vez, a fines de los años 50, un grupo de científicos congregados en Dartmouth, encabezados por John McCarthy y Marvin Minsky, acuñaron el término “inteligencia artificial” para oponerse precisamente al experimento cibernético. Su propuesta era dotar a las máquinas de una lógica interna capaz de darles un objetivo distinto al de ajustarse adaptativamente al entorno. Solo se trataba de diseñar un software con todas las reglas lógicas necesarias para que la máquina supiera responder a cualquier problema que se le presentara en el mundo exterior.

Durante años, este enfoque captó la financiación estatal y el interés social. Hasta que también le llegó su invierno en la década del 70: las máquinas aprendían a manejarse con reglas lógicas dentro del software, como si fuera una maqueta, un mundo de juguete que replicaba simbólicamente al mundo real. Pero al confrontarla con el mundo real, aparecían variables no previstas y fenómenos complejos que las reglas lógicas no podían resolver.

Así, la historia de la inteligencia artificial es la historia de una disputa entre dos paradigmas distintos, dos formas de concebir cómo piensa la mente. Los cibernéticos adhieren a un paradigma conexionista: pensar es un conjunto de funciones elementales distribuidas a través de una red neuronal, cuyo comportamiento y significado sencillamente emerge de interacciones elementales, como una caja negra que incorpora información por un lado y produce una respuesta por otro. Los lógicos adhieren a un paradigma simbólico: pensar consiste en calcular símbolos, el cerebro es una máquina que manipula representaciones del mundo exterior.

El segundo invierno de la inteligencia artificial llegó entre 1987 y 1993, con el colapso comercial de los llamados “sistemas expertos” que reemplazaban los mundos de juguete por largas listas de conocimiento especializado con reglas del tipo “SI: FIEBRE, ENTONCES: BUSCAR INFECCIÓN”. Pero eran demasiado costosos de mantener y demasiado rígidos para adaptarse a situaciones imprevistas. La lección que dejaron ambos inviernos es que el mundo real no cabe en una lista de reglas, por larga que sea.

Mientras tanto, pasaron cosas: en 1990 surgió la web, que facilitó y masificó el tráfico de internet mediante un sistema de protocolos e interfaces que permitían

Num. 18, – ISSN 2683-7129 (en línea)

navegar entre diferentes páginas. A principios del siglo XXI se desarrollaron las plataformas, infraestructuras digitales que permiten la interacción intensiva de los usuarios y captan los datos de esa interacción. En menos de una década, el volumen de datos disponibles había alcanzado una escala sin precedentes en la historia de la humanidad. Y la idea de replicar el funcionamiento cerebral con una red de nodos que permitieran el cálculo distribuido y paralelo volvió a parecer factible. Ya no era necesario reponer la misma cantidad de neuronas que tiene un cerebro: bastaba con crear capas de redes neuronales artificiales por las cuales un flujo masivo de datos fuera subiendo, y un algoritmo que lo reenviara al principio si encontraba un error: la llamada “retropropagación”. Nacía el *deep learning*, el aprendizaje maquínico profundo que le permite a una máquina combinar cantidades masivas de datos hasta dar con una respuesta probabilística: no era exacta, pero nunca se quedaba sin respuesta. Como el término “inteligencia artificial” ya había quedado instalado, a esta tecnología bisnieta de las redes neuronales de Pitts y McCulloch le tocó el nombre de su enemiga histórica: la IA.



Una foto de participantes de la “Conferencia de Dartmouth” (en inglés Dartmouth Summer Research Project on Artificial Intelligence) en 1956. La foto pertenece a la familia de Marvin Minsky.

Los algoritmos que aprenden

Los algoritmos existen desde que el matemático persa Muhammad al-Juarismi formalizó el cálculo numérico de base decimal en el año 825. Cuando Europa adoptó el sistema decimal para facilitar la contabilidad comercial, adoptó los

procedimientos de cálculo de al-Juarismi y latinizó su nombre como *algorismus*. Durante doce siglos la humanidad usó algoritmos cada vez más complejos para mecanizar manualmente cálculos. El primer salto se produjo en el siglo XX, cuando el sistema binario permitió automatizar esos cálculos empleando electricidad: el estado de una corriente eléctrica puede ser encendido o apagado, y eso se representa perfectamente con 0 y 1. Sobre esa base, se desarrollaría la teoría matemática de la información en 1938, la computadora EDVAC en 1952 y el diseño de microchips cada vez más pequeños y poderosos.

Hoy asistimos al tercer salto de esta historia: algoritmos que aprenden. Un algoritmo clásico aplica reglas rígidas sobre datos pasivos: los datos entran, las reglas se aplican, el resultado sale. Un algoritmo de aprendizaje automático, en cambio, modifica sus propias reglas internas según los datos que recibe. Los datos dejan de ser pasivos y se vuelven activos: le dicen a la máquina cómo tiene que comportarse. Ese es el auto que se transforma a medida que manejamos, una tecnología sumamente inestable que hace muy difícil cualquier proyección, planificación o legislación. Y sin embargo, debemos planificar y legislar esta máquina antes de que siga acelerando.

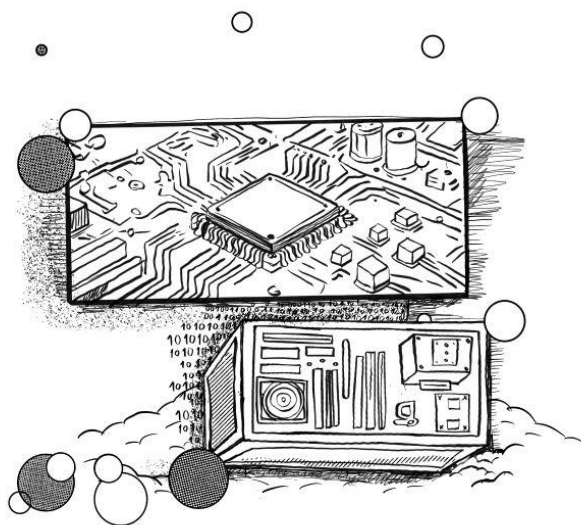
Percepción en lugar de pensamiento

Llegados a este punto, comienzan los debates: ¿es inteligencia?, ¿piensa?, ¿puede reemplazarnos?, ¿puede superarnos? Quizás convenga volver al principio: el *deep learning* fue desarrollado y presentado en sociedad en 2012 como una tecnología de reconocimiento de imágenes. Más adelante se desarrollaría la inteligencia artificial generativa, capaz de generar imágenes, y los *LLM* (grandes modelos de lenguaje) capaces de reconocer y generar texto. Pero el origen fue el reconocimiento de texto y, de alguna manera, el *deep learning* no deja de funcionar como un instrumento óptico aplicado al conocimiento: el flujo de información es como un haz de luz proyectado por los datos de entrenamiento, comprimido por el algoritmo y difractado hacia el mundo por la lente del modelo estadístico. Y como cualquier instrumento óptico, no solo amplifica sino también deforma y comprime. Toda compresión de información produce pérdida y difracción.

Yann LeCun, uno de los padres del *deep learning*, afirma que los sistemas de IA actuales no son versiones sofisticadas de la cognición, sino de la percepción. La IA no es un autómatas pensante; es un algoritmo que reconoce patrones. El algoritmo no percibe una imagen: calcula píxeles, valores numéricos de brillo y proximidad. No aprende: mapea una distribución estadística de valores numéricos y establece una función matemática. El aprendizaje profundo es la aplicación de técnicas de reconocimiento de patrones a toda clase de datos no visuales: imagen, movimiento, forma, estilo y hasta decisión ética pueden describirse como distribuciones estadísticas.

Para poder aplicar los patrones de reconocimiento visual a datos no visuales, la IA necesita granularlos, reducirlos a su mínima expresión, con el menor contexto posible, y después etiquetarlos para que la máquina los reconozca. En este proceso, los datos no solo se desvirtúan, sino que quedan expuestos a distintos

tipos de sesgos. Están, por un lado, los sesgos del mundo, las formas de prejuicio ya existentes que los etiquetadores y desarrolladores portan al momento de diseñar sus modelos de IA. Luego está el sesgo de datos, que surge de la preparación y el etiquetado de los datos de entrenamiento por parte de operadores humanos, donde taxonomías viejas y conservadoras producen una mirada distorsionada del mundo. Finalmente, el sesgo algorítmico es la amplificación y reproducción automática y masiva de los sesgos anteriores producida por los propios algoritmos de aprendizaje, dado que la compresión de información genera difracción y pérdida, como ya vimos.



La Guerra Fría de la IA

Si la historia interna de la IA es la historia de dos paradigmas en disputa, la historia externa es la vieja historia de la lucha por la hegemonía. En 2022, el gobierno de Estados Unidos amplió los controles de exportación de chips hacia China con el objetivo de debilitar el desarrollo de herramientas de IA por empresas de ese país. En 2025, la empresa china DeepSeek lanzó DeepSeek-R1, una aplicación de chatbot gratuita para iOS y Android con un costo por token 27 veces menor al de OpenAI, con un consumo energético entre diez y cuarenta veces inferior, según datos de la propia empresa.

Mientras tanto, el sur global ocupa un lugar perfectamente definido en esta geopolítica: el de proveedor de insumos y trabajo barato. Es también el campo de pruebas beta de sistemas de IA sin regulación: desde las operaciones de escaneo de iris de Worldcoin en Argentina, Colombia, Perú y Chile hasta la anotación de datos en campos de refugiados en Kenia y Líbano. El monopolio de los marcos regulatorios en manos de Estados Unidos y China restringe el margen de soberanía digital de los países del sur global. Sin embargo, Latinoamérica cuenta con una larga tradición de no alineamiento, búsqueda de independencia y una cosmovisión humanista capaz de terciar entre los dos polos. Un ejemplo de ello es la iniciativa LatamGPT, dirigida por el CENIA de Chile: un proyecto abierto que usa *fine-tuning*

Num. 18, – ISSN 2683-7129 (en línea)

para adaptar el comportamiento de modelos de IA ya existentes a contextos, lenguajes y valores latinoamericanos, con infraestructura diversificada entre hiperescaladores estadounidenses y tecnología china. A todo bipolarismo le llega su Movimiento de No Alineados.

Este artículo de la revista Scholé es producto del trabajo colaborativo entre los autores y los diferentes equipos de producción del ISEP.

Autoría: Alejandro Galliano

Equipo de producción de la revista

Dirección: Laura Percaz

Producción de contenidos: Lila Far

Edición: Martín Schuliaquer

Diseño e ilustración: Facundo Fernández

Maquetación: Daniel Wolovelsky

Desarrollo web: Santiago Rubiolo

Coordinación de producción: Matías Pardo

Coordinación de infraestructura digital: Paula Fernández

Autoridades

Martín Llaryora | Gobernador

Myriam Prunotto | Vicegobernadora

Horacio Ademar Ferreyra | Ministro de Educación

Victoria Flores | Secretaria General de Ambiente, Economía Circular y Biocidadanía

Luis Sebastián Franchi | Secretario de Educación

Gabriela Cristina Peretti | Secretaria de Innovación, Desarrollo Profesional y Tecnologías en Educación

María Eugenia Rotondi | Directora General de Innovación Educativa

Nora Esther Bedano | Secretaria de Coordinación Territorial

Claudia Amelia Maine | Secretaria de Fortalecimiento Institucional y Educación Superior

Lucía Escalera | Subsecretaria de Administración

Equipo de gestión del ISEP

Adriana Fontana | Directora

Lucía Cugini | Secretaria Académica

Victoria Farina | Secretaria de Organización Institucional

Laura Percaz | Secretaria de Contenidos Educativos

Cómo citar este material

Galliano, A. (2026). La máquina mutante: breve historia de la IA. *Scholé* n.º 18. Para el Instituto Superior de Estudios Pedagógicos, Ministerio de Educación de la Provincia de Córdoba.

ISSN: 2683-7129

Este material está bajo una licencia Creative Commons ([CC BY-NC 4.0](https://creativecommons.org/licenses/by-nc/4.0/))

