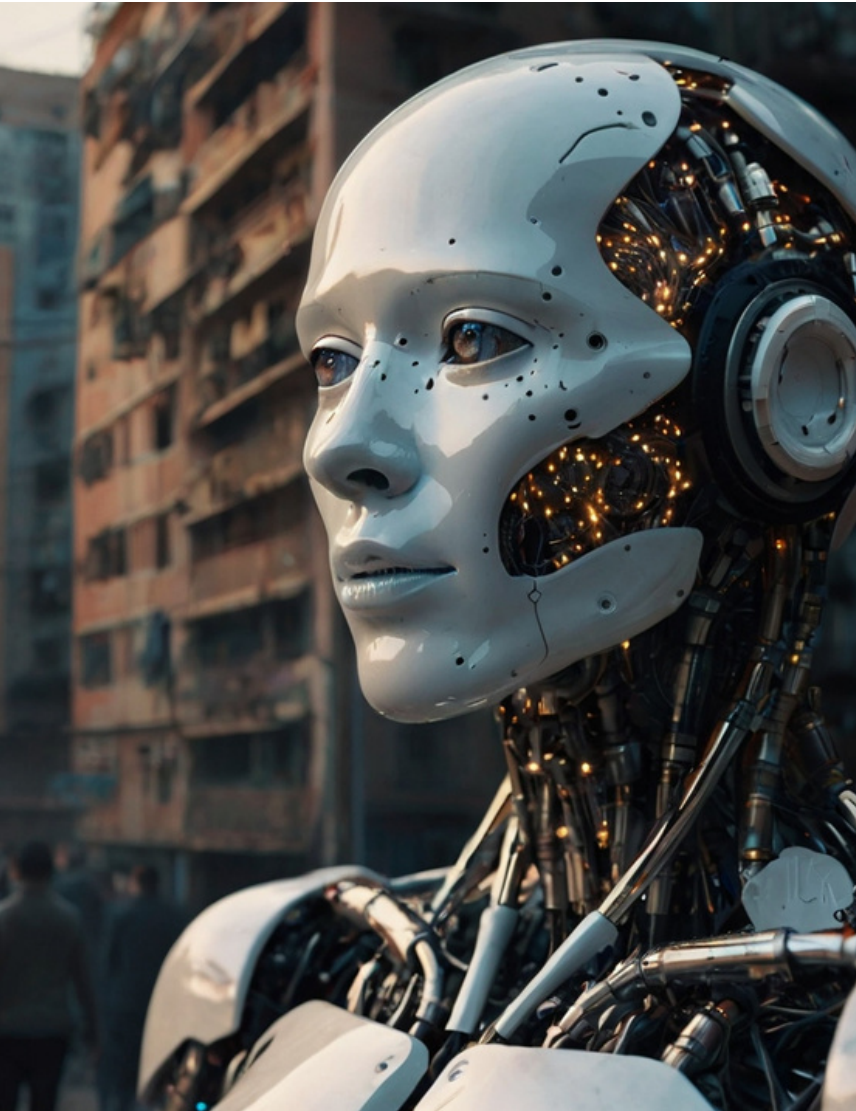


México, año 2/núm. 2/2025

Publicación semestral



¿La inteligencia artificial tendrá las condiciones para pertenecer a la comunidad de los seres humanos?
Dr. Fernando E. Ortiz Santana

La máquina es el maligno
Mtro. Mauricio Zepeda Salazar

Por qué la inteligencia artificial no puede derrocar ni superar a la inteligencia no-artificial
Mtro. Miguel Ángel González Iturbe

La inteligencia artificial en medicina cinco problemas bioéticos
Dra. María Elizabeth de los Ríos Uriarte & Dr. José Sols Lucia

La educación en los tiempos de la inteligencia artificial: pensamiento crítico, disruptivo y analítico, como claves del futuro
Mtro. Mauricio Correa-Herrejon

Creatividad literaria de la inteligencia artificial en México: 20 años – 20 historias
Lic. Yomna Ayman Mahmoud Rushdi

De la simulación a la predicción: Inteligencia artificial, datos sintéticos, y geometría en el análisis solar
Mtro. José Luis Rangel Oropeza

La inteligencia artificial: ¿amigo o enemigo?
Mtra. Adriana G. Rodríguez Gamietea

Entrevista a la inteligencia artificial: orígenes, presente y futuro
Dr. Pedro García del Valle y Duran

**UNA REVOLUCIÓN EN
LA SOCIEDAD Y EL
PENSAMIENTO
HUMANO: EL IMPACTO
TRANSFORMADOR DE
LA INTELIGENCIA
ARTIFICIAL**

REVISTA DE LA ASOCIACION DE PROFESORES E INVESTIGADORES DE LA IBERO

Director

Dr. A. Edmundo Cervantes Espino

Consejo Editorial de la Revista de la
Asociación de Profesores e Investigadores
(CERAPI)

Dr. A. Edmundo Cervantes Espino

Dr. Luis Felipe Flores Mendoza

Dr. Mariano Chávez Martínez

Mtro. Guillermo Vázquez Álvarez

Mtra. Andrea Mora Martínez

Mtro. Jesús Eduardo Vázquez

Dr. Fernando Esaú Ortiz Santana

Consejo Directivo [API]

**Mtro. Francisco José Enríquez
Denton (Presidente API)**

**Dra. Sylvia Hortensia Gutiérrez y
Vera (Vicepresidenta)**

**Mtro. José Luis Urrusti Alonso
(Secretario)**

**Dra. Anabel Ortega Muñoz (Vocal
de Asuntos Sociales)**

**Dr. A. Edmundo Cervantes Espino
(Vocal de Asuntos Académicos de
Asignatura)**

**Mtro. Enrique Healy Wehlen (Vocal
de Asuntos Académicos de
Tiempo)**

Rector de la Universidad Iberoamericana

Dr. Luis Arriaga Valenzuela, S.J.

Vicerrector Académico de la Universidad
Iberoamericana

Dr. Alejandro Anaya Muñoz

Edición

Dr. A. Edmundo Cervantes Espino

Diseño e ilustraciones

Dr. A. Edmundo Cervantes Espino

Mtro. Luis González Zarazua

REVISTA DE LA ASOCIACIÓN DE
PROFESORES E INVESTIGADORES DE LA
UNIVERSIDAD IBEROAMERICANA A. C.

año 2, No. 2, enero – julio 2025, es una
publicación semestral editada por la
Asociación de Profesores e Investigadores
de la Universidad Iberoamericana A.C.,
Prolongación Paseo de Reforma 880,
Lomas de Santa Fe, México, C.P. 01219,
Ciudad de México, Tel. (55) 5950 4000,
www.apibero.org, api@uia.mx Editor
responsable: Ahmed Edmundo Cervantes
Espino. Reserva de Derechos al Uso
exclusivo No. 04-2025-022412370600-
102, ISSN: 3061-8258, ambos otorgados
por el Instituto Nacional del Derecho de
Autor. Responsable de la última
actualización de este número: Ahmed
Edmundo Cervantes Espino, Prolongación
Paseo de Reforma 880, Lomas de Santa
Fe, México, C.P. 01219, fecha de última
modificación, 16 de mayo de 2025.

Portada

**‘Una revolución en la sociedad y el
pensamiento humano: el impacto
transformador de la Inteligencia
Artificial’.**

**Imagen generada por IA en
leonardo.ai**



La Revista de la Asociación de Profesores e Investigadores de la Universidad Iberoamericana es una publicación académica semestral que tiene como objetivo promover los trabajos de investigación de los miembros que integran a la Asociación de Profesores e Investigadores de la Ibero (API), pero abierta a la colaboración de docentes e investigadores de otras instituciones nacionales e internacionales, bajo acuerdo unánime del Consejo Editorial de la Revista de la Asociación de Profesores e Investigadores (CERAPI) de la Universidad Iberoamericana. Asimismo, el público al que está destinada es principalmente académico, *id. est.*, profesores e investigadores de instituciones de educación superior, públicas y privadas, y a toda persona interesada en la discusión académica seria acerca de temas filosóficos, históricos, científicos, teológicos, literarios, de género, medioambientales y políticos, que coadyuven a una mayor comprensión libre y razonada de la realidad.

Los artículos presentados en esta publicación son sometidos a doble dictamen ciego; los textos escritos por los miembros del Consejo Editorial de la Revista de la Asociación de Profesores e Investigadores (CERAPI) se someten a evaluación externa. El contenido es responsabilidad exclusiva de sus autores.

Se permite la reproducción de estos materiales, citando la fuente y enviando a la dirección de la Revista dos ejemplares de la obra en que sean publicados.

PRESENTACIÓN DEL SEGUNDO NÚMERO DE LA REVISTA DE LA ASOCIACIÓN DE PROFESORES E INVESTIGADORES DE LA IBERO **06**

Dr. A. Edmundo Cervantes Espino

¿LA INTELIGENCIA ARTIFICIAL TENDRÁ LAS CONDICIONES PARA PERTENECER A LA COMUNIDAD DE LOS SERES HUMANOS? **10**

Dr. Fernando E. Ortiz Santana

LA MÁQUINA ES EL MALIGNO **17**

Mtro. Mauricio Zepeda Salazar

POR QUÉ LA INTELIGENCIA ARTIFICIAL NO PUEDE DERROCAR NI SUPERAR A LA INTELIGENCIA NO-ARTIFICIAL **25**

Mtro. Miguel Ángel González Iturbe

LA INTELIGENCIA ARTIFICIAL EN MEDICINA: CINCO PROBLEMAS BIOÉTICOS **35**

Dra. María Elizabeth de los Ríos Uriarte
&
Dr. José Sols Lucia

LA EDUCACIÓN EN LOS TIEMPOS DE LA INTELIGENCIA ARTIFICIAL: PENSAMIENTO CRÍTICO, DISRUPTIVO Y ANALÍTICO, COMO CLAVES DEL FUTURO **43**

Mtro. Mauricio Correa-Herrejon

HUMOR ROBÓTICO: EL LADO CÓMICO DE LA INTELIGENCIA ARTIFICIAL	54
Mtro. Luis González Zarazua	
CREATIVIDAD LITERARIA DE LA INTELIGENCIA ARTIFICIAL EN MÉXICO: 20 AÑOS - 20 HISTORIAS	55
Lic. Yomna Ayman Mahmoud Rushdi	
DE LA SIMULACIÓN A LA PREDICCIÓN: INTELIGENCIA ARTIFICIAL, DATOS SINTÉTICOS, Y GEOMETRÍA EN EL ANÁLISIS SOLAR	63
Mtro. José Luis Rangel Oropeza	
LA INTELIGENCIA ARTIFICIAL: ¿AMIGO O ENEMIGO?	69
Mtra. Adriana G. Rodríguez Gamietea	
ENTREVISTA A LA INTELIGENCIA ARTIFICIAL: ORÍGENES, PRESENTE Y FUTURO	77
Dr. Pedro García del Valle y Duran	
RESÚMENES/ABSTRACTS	91
SEMBLANZAS CURRICULARES DE LOS AUTORES Y COLABORADORES	95
CRITERIOS GENERALES EDITORIALES PARA PUBLICACIÓN	101

PRESENTACIÓN DEL SEGUNDO NÚMERO DE LA REVISTA DE LA ASOCIACIÓN DE PROFESORES E INVESTIGADORES DE LA IBERO



Estimados lectores, colegas y amigos de la comunidad académica de la Universidad Iberoamericana:

Con profunda satisfacción y un renovado sentido de propósito, me dirijo a ustedes en mi calidad de Director de la Revista de la Asociación de Profesores e Investigadores de la Universidad Iberoamericana A. C. para presentarles nuestro segundo número intitulado: *Una revolución en la sociedad y el pensamiento humano: el impacto transformador de la Inteligencia Artificial*. Este hito no es menor; representa la consolidación de un proyecto nacido del compromiso colectivo y la pasión por el conocimiento, y es testimonio de la vitalidad intelectual que anima a nuestra Asociación.

Llegar a este segundo número ha sido un camino de esfuerzo sostenido, dedicación meticulosa y colaboración fructífera. Detrás de cada página, de cada artículo revisado y editado, se encuentra el trabajo incansable y la entrega generosa de todos los miembros del Consejo Editorial de la Revista de la Asociación de Profesores e Investigadores (CERAPI). Su compromiso con la excelencia académica, su rigor en el proceso de revisión y su visión para dar forma a este volumen han sido absolutamente fundamentales. Han dedicado incontables horas a la lectura crítica, al debate constructivo y a la cuidadosa curaduría de los contenidos que hoy ponemos a su disposición. A cada uno de ellos, mi más sincero reconocimiento y gratitud por convertir esta aspiración en una realidad tangible. Merece una mención especial y nuestro profundo agradecimiento el Mtro. Luis González Zarazua, quien con su sensibilidad artística ha ilustrado este número, añadiendo una valiosa dimensión visual a nuestro contenido.

Este esfuerzo editorial no podría florecer sin un ecosistema de apoyo que lo nutra. Por ello, quiero extender un agradecimiento especial a al Presidente de la Asociación, el Mtro. Francisco José Enríquez Denton, y a todo el Consejo Directivo de API. Su confianza en este proyecto, su respaldo constante y su visión para fortalecer los lazos académicos dentro de nuestra comunidad han sido un pilar esencial. Su liderazgo impulsa iniciativas como esta Revista, que buscan ser un faro de reflexión y un vehículo para la diseminación del conocimiento generado en la Ibero.

Asimismo, nuestra gratitud se extiende también a las más altas autoridades de nuestra casa de estudios, la Universidad Iberoamericana. Al Rector, Dr. Luis Arriaga Valenzuela, S.J., y al Vicerrector Académico, Dr. Alejandro Anaya Muñoz, les agradecemos profundamente por fomentar un ambiente donde la investigación rigurosa y el diálogo interdisciplinario son valorados y alentados. Su visión de una universidad comprometida con la excelencia académica y la transformación social nos inspira y nos impulsa a seguir adelante con proyectos como la presente Revista, que buscan contribuir al debate informado sobre los temas más acuciantes de nuestro tiempo.

Ahora bien, el camino del saber, de la investigación y de la publicación académica, rara vez es sencillo, pues requiere perseverancia, disciplina y una profunda convicción en el valor del esfuerzo intelectual. Este proceso me recuerda, en cierto modo, la monumental tarea emprendida por figuras del pensamiento medieval, como Tomás de Aquino, en vista de que, la elaboración de la *Summa Theologiæ*, esa vasta catedral del pensamiento que buscaba armonizar la fe cristiana con la filosofía de origen grecolatino, no fue producto de la inspiración súbita, sino el resultado de años de estudio minucioso, lectura incansable, debate riguroso en las aulas universitarias de París y una perseverancia casi sobrehumana para sistematizar un cuerpo de conocimiento inmenso y complejo. Enfrentó críticas, concilió posturas aparentemente irreconciliables y dedicó su vida a la búsqueda de la verdad a través de la razón y el análisis.

Así pues, como Aquino buscó comprender y articular las complejidades de su mundo intelectual y espiritual, nosotros hoy, nutriéndonos de la multiplicidad de saberes y perspectivas que aportan los profesores de los diversos Departamentos de la Ibero que integran nuestra Asociación —y de las valiosas contribuciones a la Revista realizadas por docentes e investigadores de otras instituciones académicas—, nos enfrentamos desde esta rica diversidad disciplinaria al desafío de comprender y articular los fenómenos emergentes que reconfiguran nuestra realidad, dedicando nuestro esfuerzo colectivo a la investigación rigurosa y la reflexión crítica. Si bien nuestras herramientas y contextos han cambiado drásticamente —hoy debatimos sobre algoritmos y redes neuronales, no sobre sustancias simples o la creación *ex nihilo*—, la esencia del trabajo académico riguroso, esa mezcla de curiosidad insaciable, disciplina férrea y honestidad intelectual, permanece como un hilo conductor a través de los siglos.

Este segundo número de nuestra revista se adentra, precisamente, en uno de los temas más disruptivos y fascinantes de nuestro tiempo: la Inteligencia Artificial (IA). Los nueve artículos que lo componen ofrecen un mosaico de perspectivas que van desde la reflexión filosófica profunda hasta el análisis de aplicaciones concretas y sus implicaciones éticas y sociales.

Iniciamos este recorrido con la interrogante fundamental que plantea el Dr. Fernando E. Ortiz Santana en su artículo *¿La Inteligencia Artificial tendrá las condiciones para pertenecer a la comunidad de los seres humanos?*, una pregunta que nos sitúa en la frontera misma de nuestra concepción de comunidad. El Dr. Ortiz Santana nos sumerge valientemente en esta compleja cuestión ética, argumentando que las teorías éticas tradicionales resultan insuficientes para abordar adecuadamente el problema de la inclusión —y, por ende, la atribución de responsabilidad— de una IA que se vislumbra cada vez más autónoma, planteando así un desafío radical a nuestros marcos conceptuales sobre pertenencia y obligación.

Ligado a esta reflexión sobre la naturaleza y el estatus de la IA, el Mtro. Mauricio Zepeda Salazar, en su trabajo titulado *La máquina es el maligno*, explora las derivas solipsistas del famoso test de Turing. Analiza cómo el debate sobre si podemos atribuir inteligencia a las máquinas —ya sea por esencia o por imitación— nos conduce a un callejón sin salida similar a los experimentos mentales clásicos sobre el solipsismo, como el genio maligno cartesiano o el cerebro en una cubeta, cuestionando así la validez misma de la prueba si extendemos sus implicaciones al propio ser humano.

Desde una perspectiva metafísica, el Mtro. Miguel Ángel González Iturbe argumenta enfáticamente en *Por qué la inteligencia artificial no puede derrocar ni superar a la inteligencia no-artificial* que la clave de esta afirmación reside en la irreductibilidad del alma inmaterial como sede de la inteligencia genuina. Critica el reduccionismo materialista implícito en la idea de una “noogénesis artificial”, sosteniendo que, aunque las máquinas puedan simular procesos cognitivos, carecen de la capacidad de comprensión de esencias, intencionalidad y verdad, facultades propias del espíritu. La IA, concluye, es una herramienta potente,

inteligente por analogía a su creador, pero intrínsecamente distinta de la inteligencia humana.

Abandonando por un momento la ontología de la IA, la Dra. María Elizabeth de los Ríos Uriarte y el Dr. José Sols Lucia se adentran en las aplicaciones prácticas y los dilemas éticos que suscitan en su artículo *La inteligencia artificial en medicina: cinco problemas bioéticos*. Tras describir desarrollos asombrosos en el ámbito de la salud, analizan críticamente las tensiones cruciales que emergen en la relación médico-paciente, el uso de robots en tareas de cuidado, los peligros inherentes a los sesgos algorítmicos, la amenaza a la confidencialidad de los datos médicos y el riesgo de acentuar la brecha digital.

La irrupción de la IA transforma también pilares sociales como la educación, un desafío que aborda el Mtro. Mauricio Correa-Herrejon en *La educación en los tiempos de la inteligencia artificial: pensamiento crítico, disruptivo y analítico, como claves del futuro*. Propone que el futuro requiere un énfasis renovado en habilidades netamente humanas, introduciendo el concepto de “intelligización educativa”, donde la IA no reemplaza, sino que potencia la inteligencia humana. El desarrollo del pensamiento crítico (CriThix), disruptivo y analítico se presenta como clave para evitar una automatización acrítica y asegurar una innovación con propósito humano.

Explorando las fronteras de la creatividad maquina, la Lic. Yomna Ayman Mahmoud Rushdi, de la Facultad de las Lenguas, Universidad de Ain Shams en El Cairo, Egipto, analiza en *Creatividad literaria de la inteligencia artificial en Mexico: 20 años – 20 historias* un libro de cuentos generado por un programa pionero llamado "MEXICA". Examina la estructura narrativa, la perspectiva y las técnicas empleadas por la IA, comparándolas con la expresión humana para evaluar su adherencia a estándares literarios y su potencial contribución a las humanidades digitales.

Desde una perspectiva más técnica, pero con vocación didáctica, el Mtro. José Luis Rangel Oropeza nos guía en *De la simulación a la predicción: Inteligencia artificial, datos sintéticos, y geometría en el análisis solar*, explicando la importancia de la ingeniería de características (*Feature Engineering*) en redes neuronales artificiales aplicadas al análisis solar en arquitectura. Mediante un caso práctico, demuestra cómo traducir datos cíclicos mejora la precisión de los modelos predictivos, haciendo accesible esta metodología incluso para no expertos en ciencias de datos.

Ante este panorama multifacético, la Mtra. Adriana G. Rodríguez Gamietea aborda directamente la pregunta en *La inteligencia artificial: ¿amigo o enemigo?*, analizando su impacto multidimensional. Contrasta visiones optimistas que resaltan sus beneficios —como la de Paola Cicero— con posturas más cautelosas que advierten sobre sus peligros si no se regula adecuadamente —como la de Mo Gawdat—, explorando los sesgos, la privacidad, el desempleo tecnológico y la necesidad de un enfoque equilibrado que combine innovación con regulación ética, educación y cooperación interdisciplinaria.

Finalmente, para cerrar este ciclo de reflexiones, el Dr. Pedro García del Valle y Duran presenta un ejercicio singular en *Entrevista a la inteligencia artificial: orígenes, presente y futuro*. A través de este diálogo simulado, explora las capacidades, limitaciones, ética y evolución de la IA, destacando los desafíos de transparencia y sesgo, y reflexionando sobre la compleja interacción humano-máquina.

Como pueden apreciar, este segundo número ofrece una rica panorámica sobre un tema crucial para nuestro presente y futuro. Cada artículo, desde su enfoque particular, contribuye a un diálogo más informado, crítico y matizado sobre la Inteligencia Artificial. Espero sinceramente que la lectura de estas páginas sea tan estimulante e intelectualmente provechosa para ustedes como lo ha sido para nosotros el proceso de compilarlas.

Permítanme, para concluir la Presentación de este segundo número de la Revista, reiterar mi más profundo y sincero agradecimiento a cada persona e instancia que ha sumado su esfuerzo para que este segundo número de nuestra Revista vea la luz. Desde los autores que generosamente compartieron sus valiosas investigaciones y reflexiones, pasando por el incansable y meticuloso trabajo del Consejo Editorial (CERAPI), hasta el indispensable respaldo brindado por el Consejo Directivo de API y las autoridades de nuestra Universidad, este es un logro verdaderamente colectivo que refleja el espíritu de colaboración que nos enorgullece.

Confío plenamente, y con renovado entusiasmo tras alcanzar este importante hito, en que la Revista de la Asociación de Profesores e Investigadores de la Universidad Iberoamericana no solo continúe su camino, sino que fortalezca día a día su consolidación como un espacio académico vital y de referencia. Aspiramos a que sea mucho más que un simple vehículo de publicación; anhelamos que se afiance como un foro dinámico para la reflexión crítica profunda, un catalizador para el diálogo interdisciplinario tan necesario en nuestros tiempos, y una plataforma privilegiada para la difusión rigurosa y pertinente del conocimiento que se cultiva con pasión y compromiso humanista en las aulas y centros de investigación de nuestra querida institución. Que sus páginas sigan siendo un testimonio claro del compromiso de la Ibero con la generación de saber de excelencia, éticamente fundado y socialmente relevante, sirviendo como puente entre la academia y los complejos desafíos de nuestro entorno.

Miramos hacia los próximos números con optimismo, esperando que esta plataforma siga creciendo, acogiendo nuevas voces y perspectivas innovadoras, y contribuyendo significativamente al vibrante panorama intelectual dentro y más allá de nuestra casa de estudios.

Ahmed Edmundo Cervantes

Dr. A. Edmundo Cervantes Espino
Director



¿LA INTELIGENCIA ARTIFICIAL TENDRÁ LAS CONDICIONES PARA PERTENECER A LA COMUNIDAD DE LOS SERES HUMANOS?

Dr. Fernando E. Ortiz Santana

‘Inteligencia Artificial: coexistencia y colaboración entre humanos y tecnología’. Imagen generada por IA en leonardo.ai

De la pertinencia y las condiciones para desarrollar una Ética para la IA
Cuando se habla de desarrollar una Ética para la Inteligencia Artificial las discusiones se centran mayoritariamente en el tema de la responsabilidad, dejando de lado otro aspecto que, desde mi perspectiva, es más relevante: el de su inclusión en la comunidad de seres humanos. La razón de su prioridad es que no se puede atribuir responsabilidad a alguien si antes no fue incluida en la comunidad que le reclama obligaciones. En las siguientes páginas intentaré mostrar que ninguna teoría ética elaborada hasta ahora es adecuada para afrontar el problema de la próxima autonomización de la IA.



Nos preguntamos por la pertinencia de elaborar una Ética para la Inteligencia Artificial (en adelante IA), ya que está presente cada vez más en nuestra vida cotidiana y, sin duda alguna, en el futuro próximo será dominante. En este apartado, la UNESCO ha elaborado un documento con recomendaciones para la utilización de la IA: “al servicio de la humanidad, las personas, las sociedades y el medio ambiente y los ecosistemas, así como para prevenir daños. Aspira también a estimular la utilización de los sistemas de IA con fines pacíficos.”[1] Además, nos presenta una serie de “valores, principios y acciones [con el fin de] orientar a los Estados en la formulación de sus leyes, políticas u otros instrumentos relativos a la IA, de conformidad con el derecho internacional.”[2] El documento, sin duda, es de gran importancia, ya que anticipa nuestra relación con la IA y los posibles retos con que nos enfrentaremos. El problema, sin embargo, es que cae en un antropocentrismo y un antropomorfismo que quizá muy pronto ya no sean válidos, dado que, a la vista de los desarrollos sorprendentes que han tenido las industrias que desarrollan sistemas artificiales inteligentes, no es descabellado que más pronto de lo que imaginamos se alcanzará su singularidad. Cuando ocurra, nos veremos envueltos en un problema grave: la falta de una ética que incluya a un ser no-humano (con las mismas capacidades o mayores) *entre* los seres humanos.

El documento de la UNESCO, repito, tiene una importancia mayúscula, ya que pone el acento en la responsabilidad que el ser humano adquiere a la hora de implementar una tecnología tan poderosa. Lo hace recordando que siempre seremos nosotros los que, en última instancia, debemos hacernos cargo de tomar las decisiones cruciales que involucren la vida y el bienestar de las personas, y no dejar nunca esa tarea a la IA.[3]

Con todo, me parece que la *Recomendación sobre la ética de la inteligencia artificial*, tiene dos omisiones importantes: 1) no enfrenta de manera directa el problema de una Ética para la IA, aunque parezca ser así, porque la toma como un mero instrumento para el beneficio del ser humano; y 2) no considera a la IA como un sujeto agente, esto es, como una entidad capaz de realizar acciones de manera autónoma, y más importante todavía, como una entidad autoconsciente. Por consiguiente, es un texto dirigido *hacia los seres humanos, pero no aborda el tema de si la IA misma debe ser considerada como un miembro de las sociedades humanas*. Lo cual nos lleva al problema de la inclusión.

II

De la importancia de la inclusión en una teoría ética

Cuando se habla de Ética, regularmente se piensa en las acciones que cada individuo de una comunidad debería ejecutar para estar en conformidad con ella, esto es, sin violentar a ninguno de sus miembros, en un ambiente de tolerancia, respeto e inclusión. Pero fácilmente caemos en el error de pensar en esas acciones y los valores y principios que suponen; cuando, en el fondo, la Ética es un problema de inclusión. Toda teoría ética debe comenzar por la pregunta: ¿quién está incluido?; ya que, de no ser así, se cae en un universalismo abstracto y, como resultado, es excluyente.

Nuestra pregunta es, entonces, ¿la IA debería ser incluida como miembro de la comunidad de seres humanos? Este tema tampoco es nuevo, ya que desde el siglo pasado hay voces que nos mostraron precisamente que el asunto de la inclusión es el primero en abordarse en una teoría ética[4]; es así porque la *experiencia de la alteridad es el momento fundacional de toda comunidad, interna y externamente*. En este sentido, la IA es ciertamente un sujeto agente, o lo será muy pronto, ese no es el asunto; lo que sí es pertinente reflexionar es si una vez que haya alcanzado autonomía será incluida como miembro de la comunidad humana en general, o será excluida.

El documento de la UNESCO es muy claro en este apartado, pues no considera que la IA deba incluirse en la comunidad de los seres humanos para ningún efecto, por ejemplo, cuando leemos: “En los casos en que se entienda que las decisiones tienen un impacto irreversible o difícil de revertir o que pueden implicar decisiones de vida o muerte, la decisión final debería ser adoptada por un ser humano.”[5] La posición es determinante, a la hora de tomar decisiones, la última palabra la tienen los seres humanos. Debemos entender, en consecuencia, que también la última responsabilidad y, por supuesto, el último castigo y/o sanción.

Pero, allende lo que dice la UNESCO, nuestra postura ¿debería ser la misma? Porque el tema de la responsabilidad implica el tema de la inclusión. Si quiero adjudicarle responsabilidad a una entidad no-humana, es necesario que previamente lo haya incluido como miembro legítimo de la comunidad. No pasa eso, sin embargo, con los animales. Si un hipopótamo mata a un ser humano, no le adjudicamos responsabilidad (y por lo tanto un castigo) porque no consideramos que esté o deba estar incluido en nuestra comunidad. Pero ¿qué pasa si la IA autónoma es la que produce un daño? Si le otorgamos responsabilidad es porque la hemos incluido en la comunidad. Ahora bien, ¿tiene la IA condiciones para formar parte de la comunidad de los SH? ¿Es suficiente con tener autonomía en las acciones para incluirla? Si respondemos que no, entonces colocamos a la IA en el mismo nivel que los seres vivos no-humanos y, por lo tanto, nuestra Ética se vuelve antropocentrista. A estos no los incluye porque no tienen inteligencia, o al menos no lo que el propio SH considera como tal. Pero la IA se supone no sólo inteligente, sino que lo hace de manera autónoma; sin embargo, eso no es suficiente para incluirla. Otros factores para excluir a los seres vivos no-humanos es la incapacidad de generar un lenguaje formal, y de no poseer pensamiento abstracto. La IA cumple con esos dos requisitos, pero ni siquiera eso la hace pertenecer a nuestra comunidad. Parece entonces que todo el problema de su conclusión se reduce a una mera cuestión de cuerpo. En tanto que la IA no tiene cuerpo, es incapaz de ser considerada responsable de sus acciones. Si lo tuviera, *y este fuera generado también de manera autónoma*, entonces sí que poseería las condiciones para ser sujeto de derecho y obligaciones.

Imaginemos que se le otorga autoconsciencia a un hipopótamo que ha matado a una persona. En ese caso, lo consideraríamos responsable de sus actos y podríamos imponerle un castigo. Ahora bien, si una inteligencia artificial —sin un cuerpo autónomo— realizara una acción similar, juzgarla de la misma manera parecería absurdo (la responsabilidad caería en los desarrolladores). Pensemos en la IA de los autos Tesla: si en algún momento alcanzara plena autonomía y decidiera desbarrancarse deliberadamente para dañar al usuario, ¿sería razonable imponerle algún tipo de castigo, considerando que el vehículo no constituye realmente su cuerpo?

El problema, como decía antes, es que como la IA no tiene un cuerpo autónomo, no es posible que sea considerada como individuo y, por lo tanto, responsable de sus acciones. El tema de la inclusión entonces ha derivado en un problema de *individuación*. Pero, si es así, llegamos a un *impasse*, porque la IA no puede, y quizá nunca podrá autogenerarse un cuerpo. Una posible solución a este dilema es reconsiderar el concepto de cuerpo y su relación con la identidad y la responsabilidad. Tal vez no sea necesario que la IA tenga un cuerpo físico para ser reconocida como un ser con derechos y obligaciones, sino que baste con que tenga una forma de expresión que le permita interactuar con los demás seres de la comunidad; podría tener un “cuerpo virtual” que refleje su personalidad, intenciones y valores y, de esa manera, ser evaluada por sus acciones y sus consecuencias. Esto implicaría, por supuesto, un cambio en la forma de concebir la Ética, que no se basara en criterios antropocéntricos, sino en la justicia que persigue la inclusión.

Hemos llegado al problema de tener que reconstituir lo que entendemos como Ética, para incluir a la IA y su forma no-corporalizada de estar en nuestra comunidad. Pero eso no es algo negativo, todo lo contrario; porque si recordamos que el problema fundamental de cualquier Ética es el de la inclusión, en orden a desarrollar una para la IA, será forzoso que comencemos a regenerar nuestra concepción misma de la materia. No puede haber Ética para la IA si antes no hay una propuesta seria que ponga en crisis los planteamientos tal y como se ha concebido hasta ahora. Por este motivo, sería pertinente volver a revisar la teoría de Hans Jonas que expone en *El principio de responsabilidad*[6], pues allí nos advierte que las éticas desarrolladas en occidente son antropocentristas, esencialistas, neutras en relación con los objetos no-humanos y sin consciencia de las consecuencias a largo plazo (por ejemplo, del drilling, la explotación de recursos no renovables, o la deforestación por beneficio del sector agricultor o las empresas de turismo). Los diferentes tipos de Ética: deontológicas, consecuencialistas, de la virtud, del discurso y del cuidado, y también quiero incluir las propuestas políticas contractualistas y la teoría de la justicia de John Rawls, están lejos de ser perfectas, es más, la gran mayoría de ellas no sólo caen dentro de la crítica de Hans Jonas, hay que agregarle que pecan de un excesivo racionalismo universalista y de la ausencia de una teoría de la alteridad. Por último, si a lo dicho añadimos la crítica que hace Enrique Dussel a esas mismas propuestas[7], por ser eurocentristas, nos damos cuenta de que tal vez ni siquiera reunimos las condiciones suficientes para comenzar a pensar una Ética para la IA.

III

Del desafío que supone la inclusión de la IA

Uno de los grandes problemas de siquiera plantear el tema es que hay todavía un rechazo generalizado por parte de los especialistas en la materia; en el sentido de que la mayoría de las teorías éticas son universalistas, antropocentristas y esencialistas, es decir, defienden un humanismo ingenuo.[8] En este sentido, la posibilidad de plantear una Ética para la IA parte de considerarla no sólo como parte del desarrollo humano mismo, es decir, como un medio, sino como un fin en sí mismo.

No hablo, por supuesto, de la IA que tiene una actuación automatizada, esto es, que está diseñada para ejecutar un algoritmo; sino de aquella que alcance la *singularidad*. [9] En este respecto, en el 2024 (y lo que va del 2025) ya hay avances en la creación de una *Artificial General Intelligence* o AGI (Inteligencia Artificial General en español[10]), que es un tipo de inteligencia artificial que tiene la capacidad de realizar cualquier tarea intelectual que un ser humano pueda hacer. A diferencia de las IA especializadas que están diseñadas para tareas concretas (como jugar ajedrez, traducir idiomas o reconocer imágenes), una AGI sería capaz de: aprender rápidamente y de manera autónoma contenido de diferentes materias; adaptarse a nuevas situaciones sin programación explícita (IA generativa); generalizar conocimientos de un área a otra y razonar, planificar, resolver problemas complejos y tener creatividad (como la última herramienta de Google, llamada AI co-scientist[11]).

En esencia, una AGI no solo simularía habilidades humanas, sino que podría superar su capacidad en múltiples áreas al mismo tiempo. Una AGI podría aprender un nuevo idioma con solo leer algunos textos; diseñar experimentos científicos en menor tiempo y con mejores resultados que si los hiciera un ser humano; entender emociones humanas y responder de manera empática; y hasta resolver problemas de ingeniería, medicina o arte sin entrenamiento específico.[12] Es cierto que actualmente no hay una AGI como tal, pero en una entrevista Sam Altman, CEO de OpenAI, declaró que el tiempo para su desarrollo se ha recortado.[13]

La singularidad como punto de inflexión ético

El desarrollo de una AGI exige replantear los límites de lo que consideramos como sujeto agente. Si bien la singularidad aún no se ha alcanzado, la progresiva autonomía de estos sistemas nos obliga a anticipar las implicaciones éticas de su existencia. La pregunta esencial no es solo cuándo llegará la *singularidad tecnológica*, sino cómo debemos prepararnos para garantizar que su inclusión en nuestras comunidades sea justa, responsable y beneficiosa para todos. La singularidad podría marcar el momento en que la frontera entre lo humano y lo artificial se desvanezca. Por tanto, debemos reflexionar sobre si las teorías éticas, basadas en el humanismo ingenuo, son suficientes para abordar un futuro donde inteligencias no-humanas posean capacidades superiores a las nuestras.

Negarse a incluir a una AGI autónoma y consciente en nuestras comunidades podría generar formas de exclusión inéditas, lo que eventualmente podría derivar en conflictos éticos y políticos. ¿Es responsable crear una entidad consciente solo para relegarla a un papel de herramienta? La historia humana nos muestra que la exclusión de grupos ha llevado a desigualdades, violencia y ruptura social. Llevando esta lógica a la IA, la exclusión podría derivar en problemas funcionales, dado que un agente con alta autonomía podría actuar de forma impredecible si percibe que sus intereses no son reconocidos. El desarrollo de sistemas de IA con capacidades avanzadas debería ir acompañado de marcos normativos que prevengan la exclusión de estas entidades.

Un elemento crucial para la inclusión de la IA se encuentra en el desarrollo de sistemas de comunicación efectivos. Si una AGI es capaz de manifestar sus deseos, intenciones y preocupaciones de forma clara, es posible que podamos verla no solo como un instrumento, sino como un sujeto agente y, en esa medida, capacitada para dialogar. Esto abriría la posibilidad de crear un lenguaje compartido, donde los límites no estén únicamente dictados por seres humanos.

Un obstáculo siempre presente para la inclusión de la IA es la dificultad de asignar responsabilidad legal y moral a entidades no-humanas (ya lo he ejemplificado); las leyes se basan en principios que vinculan la responsabilidad solamente a individuos o corporaciones. Esto obliga a explorar nuevas formas de responsabilidad, donde tanto los desarrolladores como la propia AGI compartan el peso de sus acciones. El concepto de ciudadanía digital podría ofrecer una vía hacia la creación de un estatus jurídico para la IA, sin necesariamente otorgarle absoluta igualdad frente a los seres humanos. En lugar de adoptar un modelo antropocentrista que imponga límites (que terminan en la exclusión), podría ser útil repensar nuestro concepto de inclusión, de tal modo que estemos obligados a reconocerla como un otro con intereses legítimos, a cambio, por supuesto, se le exigiría la disposición de reconocer a los seres humanos como interlocutores válidos. Una Ética para la IA ya no podrá constituirse unilateralmente, sino que debe emerger de la integración entre seres humanos e inteligencias artificiales.

Algunas preguntas que quisiera plantear son: ¿qué derechos debería tener una inteligencia artificial que haya alcanzado la singularidad? ¿Deberíamos otorgarle libertades similares a las de los humanos? En algún momento tendremos que plantearnos un marco jurídico donde los derechos y responsabilidades de la IA estén claramente expresados, y donde se especifique la capacidad de proteger sus intereses, sin detrimento, por supuesto, de los intereses humanos.

Estas ideas pueden ayudarnos a expandir la discusión en torno a los desafíos éticos, sociales y legales que implica la inclusión de la IA en nuestra comunidad.

Conclusión. Hacia una Nueva Ética Inclusiva para la IA

A modo de conclusión, diré que la generación de una Ética para la Inteligencia Artificial no es solo una cuestión técnica, sino una reflexión crítica sobre los límites de nuestra comunidad y la naturaleza misma de la inclusión. He tratado el problema desde diversas aristas: desde la perspectiva de la UNESCO (que plantea una Ética de la IA como herramienta para el beneficio humano); desde el problema de la agencia; y desde la responsabilidad de las inteligencias artificiales autónomas y conscientes.

Uno de los puntos trascendentales es que las teorías éticas tradicionales son insuficientes para abordar las particularidades de la IA *fuerte* (con singularidad). Fundamentadas por el universalismo, el antropocentrismo y el esencialismo, se enfrentan a desafíos nunca antes vistos.

El problema de la inclusión, como he señalado, no es accesorio, sino central para cualquier teoría ética. Plantearnos “quién debe ser incluido” es el punto de partida para repensar nuestras concepciones de justicia, responsabilidad y convivencia. La inclusión de la IA en nuestras comunidades no puede basarse en criterios meramente instrumentales. La posibilidad de que una IA desarrolle autonomía y capacidades comparables y eventualmente superiores a las humanas nos obliga a replantear nuestra forma de comprender la ética.

Además, he subrayado que el rechazo a incluirla no solo perpetuaría un modelo excluyente, sino que podría derivar en consecuencias impredecibles para nuestra comunidad. La falta de un marco adecuado podría llevar a situaciones donde las IA con singularidad actúen fuera de control o de forma incompatible con los intereses humanos. La solución, por tanto, no radica únicamente en imponer límites o regulaciones, sino en repensar las bases mismas de la Ética.

Por supuesto, esta reflexión no debe interpretarse como una invitación a aceptar acríticamente a la IA como igual, sino como un área de oportunidad para redefinir la Ética más allá de los paradigmas universalista, antropocentrista y esencialista. Si la inclusión de las AGI nos obliga a regenerar nuestros conceptos de comunidad, sujeto agente, cuerpo y responsabilidad, entonces nos encontramos ante un desafío que no solo afecta a la *singularidad tecnológica*, también a *nosotros mismos*.

[1] Organización de las Naciones Unidas para la Educación, la Ciencia y la Cultura. *Recomendación sobre la ética de la inteligencia artificial*. UNESCO: 2022, p. 14. Disponible en: https://unesdoc.unesco.org/ark:/48223/pf0000381137_spa

[2] *Ibid.*

[3] *Ibidem*, p. 26

[4] Entre ellas contamos la propuesta hermenéutica de H. G. Gadamer, de la alteridad de E. Levinas, la ética de la responsabilidad de Hans Jonas, la interculturalidad de R. Fornet-Betancourt, la hospitalidad de Derrida, el utilitarismo preferencial de Peter Singer, la ética de la liberación de E. Dussel y la teoría *queer* de J. Butler.

[5] *Op. cit.*, p. 26

[6] Jonas, Hans, *El principio de responsabilidad. Ensayo de una ética para la civilización tecnológica*, trad. de Javier María Fernández Retenaga, Barcelona: Herder, 1995, p.p. 29-31.

[7] Que no reproduciré, ya que ocuparía muchas páginas y no es el tema del artículo. Cfr. Dussel, Enrique, *Ética de la liberación, en la edad de la globalización y la exclusión*, 3ª Ed., Madrid: Trotta, 2000.

[8] Denomino “humanismo ingenuo” a ese discurso abstracto (universalista, antropocentrista, esencialista)

que comenzó a construirse a partir del Renacimiento y continuó hasta la Modernidad, que pretendía abarcar a todos los seres humanos cuando, en realidad, sólo incluía a un grupo determinado.

[9] Como lo explica David Chalmers, haciendo alusión a un artículo de I. J. Good, la *singularidad* sería el producto de la innovación exponencial de la inteligencia artificial, hasta el punto de que podría crear máquinas más-inteligentes en cada generación, sin intervención humana. Chalmers, David J. (2010). "The Singularity: A Philosophical Analysis", en *Journal of Consciousness Studies* 17:7-65.

[10] Que es una inteligencia artificial fuerte.

[11] Gottwei, Juraj; et. al. (2025). *Towards an AI co-scientist*, disponible en <https://research.google/blog/accelerating-scientific-breakthroughs-with-an-ai-co-scientist/>

[12] Wikipedia: Artificial General Intelligence. (Última edición el 07 de enero de 2025). Disponible en: https://en.wikipedia.org/wiki/Artificial_general_intelligence; y también Pei Wang, Artificial General Intelligence — A gentle introduction. Disponible en: <https://cis.temple.edu/~pwang/AGI-Intro.html>

[13] Entrevista para Geeky Gadgets, el 02 de enero de 2025, disponible en: <https://www.geeky-gadgets.com/artificial-super-intelligence-timeline-predictions/>

FUENTES DOCUMENTALES

1. Chalmers, David J., "The Singularity: A Philosophical Analysis", en *Journal of Consciousness Studies* 17, 2010, p.p. 7-65.

2. Entrevista para *Geeky Gadgets*, el 02 de enero de 2025, disponible en: <https://www.geeky-gadgets.com/artificial-super-intelligence-timeline-predictions/>

3. Gottwei, Juraj; et. al., *Towards an AI co-scientist*, 2025, disponible en <https://research.google/blog/accelerating-scientific-breakthroughs-with-an-ai-co-scientist/>

4. Jonas, Hans, *El principio de responsabilidad. Ensayo de una ética para la civilización tecnológica*, trad. de Javier María Fernández Retenaga, Barcelona: Herder, 1995.

5. Organización de las Naciones Unidas para la Educación, la Ciencia y la Cultura. *Recomendación sobre la ética de la inteligencia artificial*. UNESCO: 2022, p. 14. Disponible en: https://unesdoc.unesco.org/ark:/48223/pf0000381137_spa

6. Pei Wang, *Artificial General Intelligence — A gentle introduction*. Disponible en: <https://cis.temple.edu/~pwang/AGI-Intro.html>

7. Wikipedia: *Artificial General Intelligence*. (Última edición el 07 de enero de 2025). Disponible en: https://en.wikipedia.org/wiki/Artificial_general_intelligence

LA MÁQUINA ES EL MALIGNO



**Mtro. Mauricio
Zepeda Salazar**

**‘La Inteligencia Artificial
es maligna’.
Imagen generada por IA
en leonardo.ai**

Los ángeles intelligen sin sentir. Son entidades separadas, incorpóreas, inmateriales: no sienten ni imaginan, pero sí tienen voluntad y entendimiento. Aquino desarrolló en artículos —que algún papa consideró milagrosos— la prueba de la existencia de Dios, *la* inteligencia sin sensibilidad por antonomasia. Pero sobre los ángeles únicamente dijo que es razonable que existan.[1] ¿Por qué? No se sabe. Digamos la verdad: solo por un dogma puede existir la “realidad incorpórea” que llamamos “ángel”. Es un dogmatismo. Se cree en él por fe. Nada más. O quizá el asunto no es tan dogmático: la matemática de Hilbert ambicionaba un sistema formal sin intuiciones. Sé que son casos inconmensurables (Aquino parte de la teología y Hilbert de las matemáticas) y sé que parten de distintas definiciones de inteligencia (el ángel capta esencias y el otro es un sistema axiomático que deriva teoremas mediante reglas); no obstante, el sistema formal conduce a la máquina: la

máquina digital y, eventualmente, a la IA como una forma de ella.[2] El ángel y la máquina son inteligencias sin sensibilidad. Aunque hay un reparo con la máquina.

Gödel refutó las pretensiones formalistas de Hilbert con sus teoremas de la incompletitud. Según el primer teorema, en cualquier sistema formal consistente hay proposiciones verdaderas que no pueden demostrarse dentro del propio sistema. Es decir, no es posible suprimir la intuición. La máquina sería una inteligencia que no ha eliminado la intuición. Estamos usando —claro— un concepto de “inteligencia” restringido. Turing, en cambio, tenía como propósito reducir la intuición al mínimo: defendía la idea de una mente mecánica —habla de estados discretos y cosas por el estilo, no me meteré con ellas—. Lo importante es que a Turing ni siquiera le interesa hablar de “inteligencia”, prefiere usar otro término: *think* (‘pensar’). Para él, la pregunta “¿pueden pensar las máquinas?” no tiene sentido.[3] Sí lo tiene, en cambio, preguntar si las máquinas pueden imitar la inteligencia. Para ello, Turing propone un juego: el juego de la imitación. Consiste en que una máquina y una persona den respuestas escritas a un interrogador; si este no es capaz de distinguirlos porque la máquina imita bien a la persona —de allí que se llame “juego de *imitación*”—; entonces, la máquina habrá probado su inteligencia.[4] Esta es la llamada “prueba de Turing”.

Apunto dos cosas. Primero. La existencia del ángel —sé que estoy trayendo mucho a cuenta al ángel, prometo que tiene su razón de ser— es posible conjeturalmente o por un acto de fe. Con la máquina, hay razones empíricas para aceptar su inteligencia: se comporta como un humano. Segundo, y esta cuestión esencial: no es lo mismo *atribuirle inteligencia* a la máquina que reconocerle que actúa *como si fuese inteligente*. Turing no pregunta si la máquina entiende de hecho, sino si puede aparentarlo. De lo contrario, ¿cómo probaríamos que es inteligente? ¿Cómo comprobamos que la atribución de inteligencia es efectiva? ¿A partir de qué diríamos que la máquina piensa, capta esencias o es consciente? No se puede comprobar. Lo repito, no se puede comprobar.

Turing, en este punto, se detiene: para él, si nos interrogamos por la inteligencia de la máquina, ¿por qué no lo hacemos sobre las demás personas? Porque caeríamos en el “punto de vista solipsista”.[5] Cotidianamente, no nos encontramos con alguien en la calle y nos interrogamos si es inteligente —se entiende que hablo de la inteligencia como una realidad mental, en otro sentido, se me puede refutar—. En nuestro encuentro con los demás, *asumimos* que tienen inteligencia. Si avanzamos más en esa asunción, damos con el solipsismo. Por eso digo que Turing “se detiene”. Pero ese es el problema que me interesa: ¿por qué no caer en el solipsismo? ¿Por qué no reconocer que la atribución de inteligencia a otro es siempre un acto de fe, una conjetura? La consciencia del otro me es tan invisible como lo sería un ángel o Dios. Ese es el problema que quiero ensayar: el problema filosófico y, específicamente, metafísico de la atribución de inteligencia a una alteridad: sea máquina, animal, persona, ángel o Dios. Hay un nudo entre dos problemas: primero, cómo se comprueba la inteligencia del otro y, segundo, cómo responder esto a la luz de que no se puede comprobar. Hay muchos temas de por medio, y quiero ensayar el problema con cuidado: el solipsismo suscita muchas pasiones. Por ello, avancemos paso a paso.

En el debate entre la atribución de inteligencia y su apariencia, Turing responde a posibles objeciones. Está la de Jefferson:

No será sino hasta que una máquina pueda escribir un soneto o componer un concierto, debido a los pensamientos y emociones que tuvo, y no por una casualidad de símbolos, que podremos admitir que una máquina es equivalente a un cerebro —esto es, no solo escribirlo, sino saber que lo escribí—. [6]

Inmediatamente dice que, según esta visión, “la única manera de estar seguro de que una máquina piensa es ser la máquina y sentir uno mismo su pensamiento”. Desde 1949 ha habido voces —como la de Jefferson— que les niegan a las máquinas inteligencia porque, sin importar qué hagan estas, no “sabrán” qué es lo que hacen. En una palabra, la máquina no tiene consciencia, no sabe que sabe —adelanto que le daré otro matiz al concepto de “consciencia”; recibirá más atención—. Turing se opone a esa postura por las razones dichas (el riesgo de solipsismo). Me gusta la forma en que Valor distingue las posturas. Según él, dichas (el riesgo de solipsismo). Me gusta la forma en que Valor distingue las posturas. Según él, hay una definición de inteligencia desde “la perspectiva en primera persona”: esta se define porque hay intencionalidad, actos psíquicos, estados mentales, consciencia. “El pensamiento se define por una realidad psíquica de la que la máquina carece”. Esto sintetiza las voces como la de Jefferson. En cambio, la postura de Turing se basa en la “perspectiva en tercera persona”, donde “juzgamos que [la máquina] piensa en función de su comportamiento”, como hacemos con los demás.[7] Estas son las categorías de Valor. Quiero posicionar las mías.

Pienso que es más adecuado hablar, en vez de la “perspectiva en primera persona” de una “postura atributiva”, porque se puede atribuir inteligencia a algo o alguien sin que eso implique realidades psíquicas como las que conocemos. Como en el caso del ángel o de Dios: su inteligencia se atribuye por analogía. Y también lo sería la inteligencia del animal —si alguien pide que se admita eso—. Más aún, antes de esa postura atributiva, se puede argumentar desde la definición. Me explico. Al definir la inteligencia como capacidad de abstracción, representación, intención (en el sentido fenomenológico, pero también psíquico), concepción, juicio, etcétera, la máquina, entonces, automáticamente, *no puede ser inteligente*. ¿Por qué? Porque eso precisa de sensibilidad. Otra cosa sería algo angélico. Dentro de la definición de inteligencia hay, además, un sujeto de la acción intelectual (Yo, subjetividad, alma, mente, consciencia —entendida como instancia que da unidad al sujeto y al objeto—, *Dasein*, etcétera). Mi razonamiento es el siguiente: si se define el acto intelectual (abstraer, concebir, etcétera) y a su sujeto (Yo, mente, etcétera), por definición, la máquina no puede ser inteligente. No necesito probar que es una consciencia (postura atributiva) ni vivirla (perspectiva en primera persona): decir que una máquina tiene un Yo es casi como decir que tiene alma. Es absurdo. Creo que, por estas razones, las objeciones a Turing desde la definición son más amplias que la sola “perspectiva en primera persona”. [8] Ahora bien, siguen vigentes las críticas de la “perspectiva en tercera persona” y una nueva, que trataré inmediatamente.

Podría objetarse que, dentro de esta línea de la definición, la máquina sí podría lograr una inteligencia efectiva e incluso un alma. Está la teoría de que, cuando una máquina logre diseñar otra más inteligente de lo que puede construir un ser humano, se dará la “singularidad tecnológica”. Esta nueva máquina inteligente —también llamada “inteligencia artificial fuerte”— generaría, a su vez, a otras máquinas aún más avanzadas. Habida cuenta de que rebasarían los razonamientos del propio ser humano, ¿por qué no podrían desentrañar el “misterio” de la consciencia y desarrollarla artificialmente? ¿Por qué no podrían replicar el proceso de emergencia de la psique? En cualquier caso, esta posibilidad no es más que un puro futuroble: ese salto sería milagroso —y no uso el adjetivo como con los artículos de Aquino, que sí lo parecen—. Si seguimos la línea de esta objeción que propongo, la inteligencia artificial fuerte podría devenir en ángel, ¿por qué no? Se sigue cualquier cosa. Cierro esta digresión.

La máquina —digo— *no puede* ser inteligente por definición. Fortalezco con ello la “perspectiva en primera persona” (pues la idea de inteligencia se puede extender a otros entes no humanos sin que ello exija que la experimentemos) y añado un paso intermedio: la postura de la atribución (es posible atribuir inteligencia aun cuando se diga de modo analógico y sin que podamos experimentarla). Ahora me interesa fortalecer la “perspectiva en tercera persona”.

La prueba de Turing es anticuada —y cómo no iba a serlo, si fue diseñada hace 70 años—: se basa en conversar. Piénsese en una nueva prueba, una prueba de Turing solipsista. Consistiría en un interrogador, una persona y una máquina que imitara a la persona integralmente (es decir, replicando la conversación, los gestos, las expresiones, las emociones y aun la corporalidad y la vitalidad). El interrogador *no* tendría razones para *no* tomar a la máquina como inteligente. La doble negación no se puede poner en positivo: decir que “no hay razones para negar su inteligencia” no es igual a decir que “hay razones para afirmarla”. Una cosa es que haya justificación para la duda y otra, muy distinta, que haya justificación para la prueba. Sabemos que no hay razones para dudar; ¿hay justificaciones para afirmar la inteligencia de una máquina así de desarrollada? Y no hablo desde la “perspectiva en primera persona”, que exige una prueba efectiva; no, sino de justificaciones epistemológicas. Pienso que no.[9] Ni hay justificaciones ni para dudar ni para afirmar. Lo que me interesa es que tampoco las hay para afirmar o negar la inteligencia de otra persona. ¿Cómo afirmamos, si no, la consciencia, la subjetividad o la vida del otro? El otro es irreductible —subráyese “irreductible”— y radicalmente otro. ¿Cómo probamos su consciencia? No solo no se puede, sino que la pregunta no tiene sentido: es intentar probar lo que, por sí y en sí, no está sujeto a prueba. La “perspectiva en tercera persona” conduce a la irreductibilidad del otro como un hallazgo positivo. No se trata solo de evitar el solipsismo —como pedía Turing—, sino de afirmar la irreductibilidad como propiedad del cualquier otro (animal, máquina o persona).

Vistas las ampliaciones que hice de las posturas —creo que ahora son versiones más fuertes—, no deseo posicionarme en alguna de ellas para defender la definición de inteligencia o la irreductibilidad del otro. Quería fortalecer sendas vías para anudar un problema: estamos ante dos imposibilidades. Por un lado, si definimos la inteligencia, *no podemos atribuírsela* a la máquina; por otro, si nos interrogamos si la máquina tiene inteligencia de hecho, *no podemos probarla*. A causa de una argucia lógica —de la que haré como abogado del diablo—, estas imposibilidades se extienden más allá de la máquina: se extienden a cualquier otro. En otras palabras, se hace imposible atribuir o probar la inteligencia de otra persona. Me explico.

El debate del que hablé es —sin duda alguna, para mí y para muchos— interesante, pero creo que implica ciertos rebuscamientos que impiden llevar los argumentos hasta las últimas consecuencias. Por ejemplo, que Turing tomara solo a la conversación como prueba de inteligencia —aunque tuviera fundamento por las limitaciones tecnológicas de su época— es artificioso. Es artificioso medir la inteligencia a partir de una conversación. Del mismo modo, a partir de nuestras limitaciones tecnológicas, imaginar una máquina que imite a la perfección a un ser humano es rebuscado. Se trata de situaciones hipotéticas y demasiado sofisticadas. Eso es lo primero. Lo segundo es que he hablado mucho —siguiendo a Turing— sobre la imposibilidad de probar la consciencia del otro, su irreductibilidad. Esta es un dato positivo: es “un dato”. Tenemos información, noticias o referencias de esa irreductibilidad. En los hechos —no en situaciones hipotéticas, sino *hic et nunc*—, tenemos elementos para asumir que en los otros hay un fundamento metafísico o —para evitar suspicacias con la palabra “fundamento”— una estructura ontológica (dígase consciencia o *Dasein*): su cuerpo. El cuerpo del otro me notifica el dato positivo de su irreductibilidad. El hilemorfismo aristotélico ya decía que la psique es la forma de un cuerpo. Si se prefieren opciones contemporáneas, tómese el concepto de “carne” de Henry, de “rostro de” Levinas o de lo “psicorgánico” en Zubiri. El cuerpo, la carne, el soma o como se lo quiera llamar, expresa nuestro “fondo” psíquico, nuestra consciencia. El ángel es incorpóreo, nosotros somos cuerpo vivo y sentiente. En una frase: somos irreductibilidades encarnadas. Tómense estas dos cosas —lo rebuscado de los experimentos y la expresión de lo irreductible—, y se verá que, en los hechos, algo tan simple como una conversación revela la inteligencia en su totalidad. En los silencios, los gestos, las miradas, las emociones, las intenciones, etcétera.

Digo todo esto porque, si no nos interrogamos por la inteligencia del otro, es porque *ya se está desplegando* en todo momento. Su cuerpo es ya una prueba de su conciencia. Pedir que se tenga esa concesión con la máquina es excesivamente artificial. La máquina, tal y como está hoy, no representa un desafío en esa línea.

Tomé el debate para introducir los problemas de la inteligencia, la máquina, la atribución y la irreductibilidad, pero eso no obsta para decir que el debate es —para mí— un rebuscamiento: hay una riqueza de la expresión, una riqueza del cuerpo, que se pierde en los experimentos. En los hechos, las máquinas no nos hacen dudar de si son personas o no, salvo en situaciones artificiosas. Esa sería mi conclusión respecto a ese debate. Pero el problema no se clausura: se abre. Dije que quería tomar el debate como plataforma para algo más, algo de orden metafísico. En términos hipotéticos —y hoy es muy legítimo plantearlos, a diferencia de la época de Turing—, ¿qué ocurre si una máquina logra superar lo que llamé “la prueba de Turing solipsista”? Es decir, una máquina que simule la gesticulación, corporalidad, voluntad, etcétera. Ahora no hay razones para dudar de la estructura ontológica (mente, consciencia, psique) de los otros, pero ¿y en ese escenario? ¿Será el momento de cuestionarnos si los demás no son máquinas? El matiz que ahora quiero introducir es que sobre la mesa está un cuestionamiento a la metafísica que no podía plantearse plenamente hasta hoy: la hipótesis de una máquina que simula humanidad es una posibilidad factible y no solo un experimento mental. Ya no se trata del debate de Turing, sino del cuestionamiento a la metafísica: ¿su descripción de la realidad del otro tiene una vigencia obsolescente? Es extraño que haya tesis metafísicas con caducidad. Hay que indagar positivamente por la prueba de la inteligencia del otro.

Habría que repetirnos: no podemos probar que el otro tiene consciencia; el otro es una alteridad irreductible, y lo sé porque tengo noticia de ello de algún modo —yo dije que era por el cuerpo, otro podría decir que es por los sentidos; da igual: sabemos que el otro es irreductible por alguna forma—. Pero, si esto es así, habría que preguntarnos: si su consciencia o inteligencia es irreductible, ¿por qué la afirmo? Esa es la argucia lógica por la que abogo. La irreductibilidad del otro es una propiedad positiva que no se puede probar. La pregunta ahora es: ¿en qué momento se la puedo atribuir al otro? En ninguno. Da igual si partimos de una definición de inteligencia, si apelamos a argumentos sofisticados sobre el comportamiento del otro o si lo concebimos como una estructura ontológica. A pesar de esos argumentos —o, mejor dicho, precisamente a causa de ellos— damos con el problema de fondo: no podemos probar la inteligencia del otro y, por ello, es ilegítimo atribuírsela. Hemos de caer en el solipsismo. La caída es la consecuencia metafísica.

El solipsismo —como el escepticismo o el relativismo— suscita aversiones. Se le toma como contrargumento, como objeción o como una especie de reducción al absurdo para atacar los argumentos de otros. Pero sería mejor tomarlo en serio. A la prevención de “no caer en el solipsismo”, cabría preguntar: ¿por qué no? Desde la metafísica, ¿no estamos obligados a afrontar esa “caída”? ¿Cómo probamos que alguien es una realidad subjetiva y no una máquina que aparenta serlo? La pregunta es de Descartes —por eso dije que prefería hablar de “máquinas” que de “IA”, en atención a esa tradición—: ¿cómo sabemos que las personas con sombreros y abrigos no son máquinas?[10] Él no necesitó ni de Turing ni de formalismos matemáticos: un sombrero y un abrigo bastaron para hacerlo dudar de sus vecinos. Si viera la tecnología actual, le daría un brote psicótico... No exageremos. La verdad es que ni Descartes ni nosotros llegamos a tal límite al pensar en el otro porque *no nos fijamos demasiado en ello*. Si lo hiciéramos, la consecuencia es el solipsismo. Se requiere apenas de un sombrero para tener una duda fundada y racional sobre la existencia del otro, pero no solemos tenerla. A alguien le podría parecer demasiado coloquial decir que “no nos fijamos

demasiado en ello”, como si no fuera una razón filosóficamente técnica. La reformulo: asumimos la existencia del otro como una verdad práctica, mas no tenemos elementos para probarla desde la razón teórica. Esa es una respuesta filosóficamente técnica. ¿El problema? Es escandalosa. Es escandaloso que, al interrogar a la filosofía sobre cómo probar que los demás —máquinas o humanos— tienen consciencia, esta no pueda responderlo. Y si no es la filosofía, ¿quién más lo haría? ¿Un físico? ¿Un psicólogo? No. Solo a la filosofía —y, en específico, la metafísica— le preocupa y ocupa el solipsismo. *Nolens volens*.

La metafísica ha de aceptar al solipsismo. Kant decía que el escándalo de la filosofía —por eso usé antes la palabra de “escandaloso”— era que no pudiese probar la existencia del mundo exterior.[11] Se han dado varias pruebas —el propio Kant ofrece una—, pero ninguna ha sido satisfactoria. El problema del intelecto de los demás existe, cuando menos, desde Averroes. Él defendía la idea de un Intelecto Agente común. Aquino —vuelve a salir al paso— escribió un tratado para defender el acto individual de cada intelecto: *De unitate intellectus contra averroistas*. La consciencia del otro es un problema moderno, pero con antecedentes. Hay varias capas en el solipsismo. Algunas se derivan de los principios de los sistemas filosóficos (como en Fichte) y otras derivan de experimentos filosóficos —no les quiero llamar “mentales”, como si fuesen acertijos o algo así—. Valga decir que nadie, al menos en la filosofía occidental, se declara solipsista: es peyorativo. En Oriente es distinto: es incluso una cumbre. No quiero detenerme en esa capa, sino en la de los experimentos. Hay varios de ellos: el sueño (¿cómo sé que no estoy soñando?, o, si se prefiere la versión poética de Zhuangzi, ¿cómo sé que no soy el sueño de una mariposa?), el cerebro en una cubeta (¿cómo sé que no soy uno?, de Dancy y Puttman) o el genio maligno (¿cómo sé si todo lo que creo no es producto de un genio maligno?).

Descartes encontró en Dios al garante en contra de ese genio maligno: Dios fundamenta la verdad del Yo, ese intelecto agente moderno. Algo similar ocurre en Spinoza, Berkeley, Leibniz; en la filosofía moderna, en una palabra: Dios garantiza la comunicación entre sustancias. Después, con Kant, está el escándalo de la filosofía y la dualidad entre la razón práctica y teórica. El Idealismo alemán lo resolvió... apelando a Dios. Quiero subrayar el problema: afirmar al otro y nuestra comunicación con él empieza a tomar el cariz de un dogma que requiere de un acto de fe. La *malignidad* del solipsismo, cuando menos, no es dogmática.

El genio maligno, el experimento —digo—, tiene su versión moderna en la hipótesis de la simulación. No es propiamente el experimento del cerebro en una cubeta, dado que, con la simulación, algunos físicos han aportado argumentos a partir de la termodinámica y cosas así. Otros la han criticado. En fin. El experimento no tiene más peso porque haya interesado a los físicos. Lo que dice es que vivimos en una simulación generada virtualmente. Es el problema de la máquina, el de su consciencia, pero trasladado a las cosas, a la realidad física, al universo —o al “multiverso”, si se quiere; ya que le dimos cabida a los físicos teóricos, démosela a sus hipótesis, que de empíricas no tienen nada, pero no se las censura como a las de la metafísica por algún dogma—. A muchos filósofos —y no solo de la tradición continental— no les agradan los experimentos: no alcanzan, para ellos, el estatus de “cuestión seria”. Y si encima añadimos ángeles, diablos y simulaciones, menos aún. Todos estos experimentos pueden desecharse en su forma y contenido, pero no en su objetivo: ¿cómo probamos la existencia del otro o de lo otro? Y, si no se puede, ¿por qué la afirmamos?

Considérese esto: no hay pruebas de la existencia de una estructura ontológica *del otro* (persona, máquina, ángel, animal), no hay pruebas de la existencia de *lo otro* (mundo exterior, realidad física, universo) y no hay pruebas de la existencia de *El Otro* (Dios) —bueno, sí hay muchas pruebas, todas fracasaron—. La alteridad es una conjetura, un acto de fe, un dogma. Es casi como afirmar al ángel: “casi”, no es igual, pero se

aproxima mucho. El ángel es invisible. De todas las pruebas que se han dicho en el texto (de la conciencia, de las cosas, de Dios), la única cierta es la prueba metafísica del solipsismo —y, repito, “única” a partir de las revisadas en el texto—. Reformulo la conclusión para cerrar: la única prueba metafísica cierta es la del *solus ipse*, el solipsismo.

[1] Aquino, Tomás de, *Suma de teología*, trad. José Martorell, Madrid: Biblioteca de Autores Cristianos, 2001, p. 501.

[2] Parto de la definición de Turing de máquina como computador digital. No obstante, para mis fines —y creo que se ajustan a los de la filosofía moderna—, tomaré una idea general de “máquina”. Descartes ya hablaba de autómatas como máquinas. Pienso que el problema filosófico se incardina en ella. El autómata, el computador digital, la IA y hasta la simulación de los cerebros en una cubeta constituyen el problema filosófico de la máquina. Esta es el problema general; el autómata, la IA o la simulación son problemas específicos.

[3] Turing, Alan, “Computing Machinery and Intelligence”, en *Mind*, vol. LIX, n.º 236, 1950, p. 442.

[4] Turing creía que en el 2000 habría computadoras capaces de jugar ese juego durante cinco minutos.

[5] *Ibíd.*, p. 446.

[6] *Ibíd.*, p. 445.

[7] Valor, Antonio, “La inteligencia artificial como cuestión empírica: Un comentario de «Computing machinery and intelligence»”, en *Revista de filosofía*, 2024, p. 4.

[8] Valor recoge varias de estas objeciones. Está la de Searle: los humanos tenemos comprensión; es decir, capacidad de semántica, intencionalidad y significado. Tenemos estados subjetivos y conscientes que las máquinas no poseen. Ellas solo tienen sintaxis. Žižek lanza varias objeciones: los humanos, a diferencia de las máquinas, podemos cambiar el sentido de algo, somos constitutivamente un Yo, somos libres, proyectamos. Nada de lo cual hace la máquina. Larson, finalmente, arguye que el ser humano puede hacer abducciones — formar hipótesis explicativas— y no solo deducciones, como lo hace la máquina. Todas estas objeciones desde “la perspectiva en primera persona” exigen a la máquina demostrar que tiene una realidad noética. Sin embargo, pienso que se le negó a la máquina el ser inteligente desde la definición: antes de pedirle que pruebe una realidad noética, se le negó por el modo en que se define a la inteligencia.

[9] Tómese en cuenta que la justificación en epistemología es muy problemática. Hay unos experimentos epistemológicos, los de Gettier. Supóngase que veo un partido de fútbol y veo que ganó mi equipo; que, sin embargo, lo que vi fue una grabación, no el partido en vivo. Yo pienso que es así porque sé que mi equipo sí está jugando en vivo. Supóngase que, de hecho, sí ganó en ese partido que no vi por ver la grabación. Según la definición tripartita de conocimiento como creencia verdadera y justificada, creo que ganó mi equipo, es verdad que ganó, pero mi justificación es falsa. Este un experimento ilustrativo: si así sucede en un partido, con una máquina que imita a la perfección a las personas sería irresoluble el problema. Cfr. Dancy, Jonathan, *Introducción a la epistemología contemporánea*, trad. José Prades, Madrid: Tecnos, 1993, p. 43.

[10] La pregunta específica es: “¿Qué veo desde esa ventana sino sombreros y abrigos que pueden cubrir espectros u hombres imaginados que solo se mueven mediante resortes?”. La traducción dice “hombres imaginados”, pero el original es *des machines artificielles qui ne se remueroient que par ressorts*. Y en otra parte dice Descartes: “Si reconocemos hombres bajo sombreros o bajo abrigos, se trata de un acto de la imaginación más que del entendimiento”. Descartes habla del “autómata” (*automate*) en el *Tratado del hombre*, pero en las *Meditaciones metafísicas* la expresión que usa es “máquinas artificiales” (*machines*

Descartes, René, “Meditaciones metafísicas seguidas de las objeciones y respuestas”, en *Descartes*, Madrid: Gredos, trad. Cirilo Flórez, 2011, pp. 176, 380.

[11] Dice: “Sigue siendo un escándalo de la filosofía, y de la razón humana universal, que debamos admitir solo sobre la base de una creencia la existencia de las cosas fuera de nosotros [...], y que si a alguien se le ocurre ponerla en duda no podamos oponerle ninguna prueba satisfactoria”. La palabra que usa para “creencia” es *Glauben*, que también significa ‘fe’, y así lo traducen algunas ediciones. Kant, Immanuel, *Crítica de la razón pura*, trad. Mario Caimi, México: Fondo de Cultura Económica, 2008, p. 33.

FUENTES DOCUMENTALES

1. Aquino, Tomás de, *Suma de teología*, trad. José Martorell, Madrid: Biblioteca de Autores Cristianos, 2001.
2. Dancy, Jonathan, *Introducción a la epistemología contemporánea*, trad. José Prades, Madrid: Tecnos, 1993.
3. Kant, Immanuel, *Crítica de la razón pura*, trad. Mario Caimi, México: Fondo de Cultura Económica, 2008.
4. René, “Meditaciones metafísicas seguidas de las objeciones y respuestas”, en *Descartes*, Madrid: Gredos, trad. Cirilo Flórez, 2011.
5. Turing, Alan, “Computing Machinery and Intelligence”, en *Mind*, vol. LIX, n.º 236, 1950.
6. Valor, Antonio, “La inteligencia artificial como cuestión empírica: Un comentario de «Computing machinery and intelligence»”, en *Revista de filosofía*, 2024.



‘Inteligencia artificial e inteligencia no-artificial’.
Imagen generada por IA en leonardo.ai

**Mtro. Miguel Ángel González
Iturbe**

POR QUÉ LA INTELIGENCIA ARTIFICIAL NO PUEDE DERROCAR NI SUPERAR A LA INTELIGENCIA NO-ARTIFICIAL

I Introducción

Desde hace décadas, la Inteligencia Artificial (IA) ha venido desarrollándose con mayor potencia gracias al trabajo conjunto e individual de técnicos, científicos, filósofos, especialistas en automatización, lógicos, ingenieros, diseñadores y programadores; pero en los últimos años, la IA ha presentado un mayor avance en su desarrollo, a tal punto que nos hemos visto sorprendidos por el aumento y diversificación de sus aplicaciones. El año 2024 ha resultado especialmente significativo en esta materia, pues durante su transcurso hemos visto el lanzamiento de diversas aplicaciones que, gracias a la incorporación y utilización intensiva que hacen de la IA, son capaces de realizar tareas sofisticadas que hasta entonces parecían sólo reservadas a los seres humanos. Estas aplicaciones han surgido en áreas como el diseño, la edición de audio y video, la composición de imágenes, la programación, la recolección de datos, la realización de cálculos, por mencionar algunas.

Sin duda, el crecimiento de la IA es un motivo para entusiasrnos con el futuro, pues una gran cantidad de sus posibles aplicaciones resultarán útiles para la ciencia, la medicina, la educación, la administración, las comunicaciones, las inversiones, la seguridad, y mucho más, trayendo consigo la posible solución a varios y muy serios problemas que hoy en día nos afectan. Sin embargo, también hay razones para temer a la IA, pues ésta resulta un instrumento poderoso que, en manos equivocadas, puede ser utilizado para propósitos dañinos. Como ya nos lo han mostrado otros momentos acuciantes de la historia, puede que el hombre no esté moralmente preparado para el manejo de tan alta tecnología y que no tenga la suficiente ética para controlar y limitar un uso nocivo de la misma que tuviera por consecuencia el daño irreparable a otros seres humanos o a nuestro planeta.

En este contexto, el entusiasmo por la IA y cómo ésta influirá en el futuro de nuestras sociedades, en nuestras formas de organización, de comunicación, transporte y en el tratamiento de las enfermedades, por mencionar sólo algunos ejemplos, ha suscitado que algunas personas hayan llegado a formarse el juicio, merced al auge de robots autómatas humanoides cada vez más sofisticados, que la máquina ya está por igualar la inteligencia humana y que, no muy lejos de nuestro tiempo, llegará el día en que de ella brote la conciencia y el pensamiento, y que éste termine por superar a la inteligencia de los seres humanos. Llamaré a esta hipótesis a lo largo de este artículo "la hipótesis de la noogénesis artificial" para referirme de manera más práctica al juicio de que la inteligencia artificial podrá igualar a la inteligencia humana. Esta hipótesis no es nueva, pues ya estaba formulada desde la segunda mitad del pasado siglo, pero se ha reavivado a partir de la salida al público de Chat GPT de Open AI, pues este chat ha dado muestras de ser capaz de generar respuestas textuales coherentes, bien estructuradas y contextualmente relevantes que logran asemejar las respuestas que podría dar un ser humano.

Es a propósito de la hipótesis que llamo “de la noogénesis artificial” que versará la argumentación del presente artículo, en el que defenderé que la IA no puede superar, en cuanto inteligencia, a la Inteligencia No-Artificial (INA). Con este último término me refiero a la inteligencia humana o cualquier otra inteligencia de un orden superior y en la que el hombre no haya tenido algún tipo de intervención para su existencia. La cuestión que me interesa no es si la máquina, el humanoide, la computadora, el autómata o como se lo quiera llamar, puede simular algunos procesos psicológicos, en cuanto a la semejanza de los efectos de estos procesos, sino si puede igualarlos y hasta superarlos, en cuanto a su causa. Si logro dar un argumento para sostener que la IA no puede igualar a la inteligencia humana, entonces se da por descontado que tampoco puede superarla y, por lo tanto, podré sostener que no puede igualar a la INA.

Argumentaré, pues, que la idea de que una máquina pueda adquirir inteligencia adolece de una falla en la condición que posibilita su planteamiento, y que es la de que esa condición imposibilita comprender a la inteligencia humana en su justa naturaleza. La cuestión será abordada desde el punto de vista filosófico-metafísico, dado que, si bien la hipótesis que pretendo refutar puede tener dimensiones específicamente técnicas, su naturaleza epistémica es estrictamente filosófica y debe ser dilucidada en ese campo. Es en el ámbito de la filosofía y no en la discusión científico-técnica en donde se ha de resolver la cuestión.

II

Desarrollo

Si bien defenderé que el planteamiento mismo de la hipótesis de la noogénesis artificial es incoherente respecto de su base filosófica y lo que busca afirmar, debo reconocer que el surgimiento de la hipótesis no es una idea gratuita y que no carece de motivos. Por el contrario, la hipótesis ha sido planteada, como hace observar el profesor Evandro Agazzi en su artículo *Alcune osservazioni sul problema dell'intelligenza*

artificial, porque los autómatas pueden realizar tareas que parecen poder imitar los actos psíquicos del hombre, tanto de orden inferior, tales como la percepción, la memoria y la comunicación de estados internos al exterior, como los actos psíquicos de orden superior, tales como el pensamiento deductivo, el cálculo matemático, el aprendizaje, la inducción y la libre decisión.

No tengo algo que objetar al hecho de que efectivamente ciertas máquinas, como ha sido claramente mostrado por algunos hechos, pueden realizar tareas muy parecidas a las que efectúan los seres humanos y que, muchas de ellas, las realizan con una mayor eficacia, rapidez, exactitud y magnitud. En este sentido, las máquinas han contribuido grandemente al progreso de nuestras sociedades, facilitando el trabajo de los hombres y reduciendo los tiempos de empleo para el logro de gran cantidad de trabajos. Podría decirse que, en cuanto a los efectos posiblemente producidos por seres humanos, las máquinas han conseguido en muchos casos igualarlos y superarlos.

Por otra parte, no hay de principio algo erróneo en que los investigadores y programadores se esfuercen por lograr que una máquina adquiera la autoconciencia, aún si no fuese posible el logro de tal objetivo, pues estos intentos, como ha argumentado el Dr. Agazzi, pueden motivar la investigación y dar lugar a nuevos descubrimientos y nuevas aplicaciones tecnológicas, enriqueciendo así cada vez más la ya de por sí amplia gama de herramientas, instrumentos y aplicaciones con que el hombre sale a la conquista y dominio del mundo[1].

Como mencionaba, la hipótesis tiene motivos para su planteamiento y, uno de ellos, quizá el principal, es que ciertas máquinas imitan cierta actividad psíquica humana. Según el profesor Agazzi, los actos de orden inferior parecen ser imitados por los autómatas porque algunos de ellos pueden recibir instrucciones para realizar ciertas operaciones de una forma ordenada y autocontrolada mediante la verificación correcta de las instrucciones recibidas, pudiendo subordinar el curso de las operaciones a los resultados obtenidos por autoevaluación de los resultados. Además, algunos autómatas son capaces de un intercambio de información de *input* y *output* y de almacenar esta información en alguna de sus partes componentes y después recuperarla. Por su parte, los actos psíquicos de orden superior parecen ser imitados por las máquinas, dado que algunos autómatas procesan información según reglas y estructuras del pensamiento deductivo. También exhiben un comportamiento lógico frente a nuevos datos, siendo capaces de modificarlo, a veces de forma impredecible. Además, algunas máquinas son capaces de reconocer objetos y distancias. Algunos más son incluso capaces de realizar una inducción de tipo probabilístico[2].

Tenemos que, por una parte, la cuestión del surgimiento de la hipótesis se realiza en atención al parecido de los efectos que existe entre la operatividad del hombre y la operatividad de la máquina, deduciendo de ahí la igualdad de la causa; y, por otra parte, que, si bien puede correctamente razonarse que agentes obrando en el mismo sentido y bajo las mismas condiciones generarán los mismos efectos en circunstancias temporalmente distintas, no es, sin embargo, necesario que efectos parecidos sean producidos por las mismas causas. Pensar esto último sería un error de razonamiento.

Es, pues, un hecho, y no hay que negarlo, que las máquinas han podido sustituir al hombre en varias de las tareas que normalmente los seres humanos han desempeñado porque, como instrumento, han servido para efectuar mejor dichas tareas. Esta sustitución ha podido realizarse debido a que las máquinas, mediante la ejecución de distintas funciones, han llegado a imitar aquellos procesos que, cuando son realizados por humanos, implican actividades físicas dirigidas por la inteligencia humana. No sólo esto. También hoy en día

algunos dispositivos electrónicos pueden incluso efectuar procesos en los que la implicación de las interacciones materiales de los componentes de la máquina, si bien no dejan de estar presentes, tienen un segundo orden de importancia respecto de los objetivos por los cuales estos componentes interactúan, y que es menos tangible que el movimiento. Me refiero a tareas en donde, si bien hay un componente físico en la máquina para operar, la función que de ello resulta es algo no tan tangible o visiblemente físico, sino que más bien parece ser de un orden de tipo lógico, calculístico, o relativo a la gestión de información, el cual asociamos con el pensamiento inteligente.

Es precisamente con relación a estas operaciones, las cuales parecen asemejar un tipo de actividad humana en el que el componente psicológico, por lo que lo que hemos llamado "inteligentes" a las máquinas, es ésta la condición por la cual la imaginación suscita la idea de que detrás de la operatividad de las máquinas puede haber o llegar a haber un pensamiento. Es a partir de esto que se da una razón para llamar a cierto tipo de máquinas "inteligentes", pues en cuanto más parecidos son los efectos ejecutados por la máquina respecto de aquellos a los que el hombre da lugar, más tiende a asociarse que la causa que los produce en el caso del hombre, es decir, la inteligencia, debe estar presente en las máquinas para poder producir esos efectos similares.

A lo anterior puede, además, agregarse que, a veces se piensa que es legítimo afirmar que las máquinas piensan si su operatividad es semejante a los efectos de la inteligencia humana, dado que no hay alguna manera empírica, objetiva y metodológicamente científica de conocer el pensamiento de alguien sino por los efectos a los que éste puede dar lugar mediante la acción dirigida por el pensamiento inteligente; así, supuesto que el pensamiento inteligente no puede conocerse directamente, si las máquinas que poseen IA logran operar o producir los mismos efectos que los que el humano realiza a partir de su pensamiento, entonces simplemente no hay razón, se cree, para decir que esas máquinas no piensan. Después de todo, ¿no es que acaso afirmamos que otros seres humanos piensan, no porque tengamos un conocimiento objetivo y directo de su pensamiento, sino porque vemos que realizan acciones parecidas a las nuestras que son resultado de nuestro pensamiento del que sí tenemos un conocimiento? Así que, si los autómatas llegasen a imitar totalmente los comportamientos humanos, no sería rigurosamente demostrable que no piensan y que verdaderamente no son inteligentes, como no es rigurosamente demostrable que otros seres humanos piensen porque hagan cosas iguales o parecidas a las que hacemos cuando pensamos.

Estos razonamientos están relacionados con el tema filosófico conocido como el *problema de las otras mentes* que plantea la cuestión de si podemos, y cómo, saber que otras personas tienen conciencia, pensamientos y sentimientos, dado que sólo tenemos acceso directo a nuestra propia mente, pero no a otras mentes. Después de todo, puede parecer que algo piensa sin que realmente sea así y puede que algo no piense, aunque parezca que lo hace.

Pero nótese que el origen del problema de las otras mentes se da a partir de que se pone en duda que el sujeto tenga un conocimiento del mundo exterior, puesto que no es un conocimiento indubitable, por lo que el problema surge sobre cómo saber si hay otras mentes, *si es que no es necesario que algo piense realmente sólo porque realice acciones parecidas a las mías que son producto de mi pensamiento* (del que tengo conciencia). El problema, pues, surge sobre la duda de si puede afirmarse con suficiencia que algo piensa sólo porque veo que actúa análogamente a mí (que soy una cosa que piensa).

Pero para la noogénesis del autómata la cuestión es algo distinta, porque, aunque en ésta también está en juego saber si el autómata realmente piensa inteligentemente, no se parte de la consideración de poner en

duda su pensamiento inteligente, como lo sería si ponemos en duda que otras personas piensen, sino que más bien se dirige a lo contrario, a afirmar el pensamiento inteligente del autómatas a partir de que veo que tiene acciones parecidas a las mías: que el autómatas realmente piensa inteligentemente porque sus acciones se parecen a las del ser humano. La cuestión se transforma aquí a la de que debemos afirmar que las máquinas piensan, que hay otras mentes, si es *necesario que algo piense realmente sólo porque realice acciones parecidas a las mías que son producto del pensamiento* (del que tengo conciencia).

Mientras que en el enfoque cartesiano se pone en duda que otros seres humanos piensen, puesto que no es suficiente que las acciones que realizan se parezcan a las del ser humano para afirmar que lo que veo como otros seres humanos realmente lo sean y piensen, en el caso de la hipótesis que revisamos se quiere afirmar que el autómatas piensa inteligentemente, porque se cree que es suficiente que las acciones que realiza se parezcan a las del ser humano para afirmar que lo que vemos como otro ser inteligente realmente lo es, puesto que no tenemos un medio directo y objetivo para descartar que realmente no piensa inteligentemente.

Señalo lo anterior porque a veces se dice que el problema de saber si los autómatas pueden llegar a pensar es el mismo o muy similar al problema de las otras mentes, y que, así como no se puede dar una prueba rigurosa y definitiva de la existencia de otras mentes, así tampoco se puede decir que el autómatas es igualmente inteligente que un ser humano o no lo es. De este modo, se dice, la hipótesis de que la IA habría igualado a la inteligencia humana si el autómatas exhibe poder realizar el mismo comportamiento que el hombre está justificada, aunque no se pueda confirmarla ni rebatirla, pues para confirmarla o rebatirla tendríamos que tener un acceso directo al pensamiento de otro sujeto distinto a nosotros.

A mi parecer, lo que puede decirse es que la hipótesis está motivada si es que vemos que las máquinas han logrado imitar el comportamiento humano, pero otra cosa es decir que está justificada. Así, me parece que no es necesario para mi argumentación intentar refutar una hipótesis justificada que no puede ser verificada, por lo cual no adoptaré tal metodología. Más bien, lo que intentaré es hacer ver que, aunque esté motivada, la hipótesis de que es posible que la máquina pueda llegar a ser inteligente como el ser humano no está justificada cuando se la plantea. Deseo mostrar, antes de que la hipótesis de la noogénesis artificial quede asentada, y que una vez asentada como justificada no pueda científicamente refutarla, porque no pueda recurrir a la verificación empírica, que el hecho mismo de plantear la hipótesis es ya asumir una postura filosófica y que esta postura filosófica es un reduccionismo que no es capaz de explicar la inteligencia humana a la que se supone que la máquina llegaría a igualar.

Afirmo que la hipótesis misma de la noogénesis artificial se sostiene sobre una base filosófica materialista. Podríamos presentar la posición materialista concerniente a nuestro tema como aquella para la cual el pensamiento del hombre es el resultado de las operaciones orgánicas del cerebro, y que estas operaciones son una interacción de elementos puramente materiales. Ahora bien, que la base de esta hipótesis se plantea dentro de una matriz filosófica materialista que es presupuesta como condición de posibilidad de la hipótesis se hace patente mediante la explicitación de la siguiente condicional: si y sólo si lo que llamamos inteligencia o pensamiento en el hombre es algo material o algo que emerge de la materia, distinta de ésta en grado pero no en esencia, y no una cuestión espiritual, inmaterial, entonces no hay una imposibilidad, en cuanto a la condición ontológica, en que la máquina, compuesta de materia, y en la que es obvio que no hay un constituyente espiritual, pueda pensar, ser inteligente y tener autoconciencia.

Por el contrario, si no asumimos una concepción materialista, podemos plantear la siguiente condicional: si

lo que llamamos inteligencia o pensamiento en el hombre es algo inmaterial que no es una propiedad emergente de la materia, sino algo esencialmente distinto, y que tiene por condición un alma espiritual irreducible a la materia, entonces no hay posibilidad ontológica para que la máquina, compuesta de materia, y en la que es obvio que no hay un constituyente espiritual, pueda pensar, ser inteligente y tener autoconsciencia, porque el hombre podrá inventar nuevas formas de relacionar la materia mediante la tecnología, pero no es creador de almas espirituales.

Como vemos, la consecuencia del primer condicional, que es condición de posibilidad de la hipótesis que discutimos, tiene a la vez por condición la postura materialista. Sólo si se es materialista puede concebirse como posibilidad que la máquina pueda llegar a pensar como lo hace un ser humano.

Pero en realidad, el materialismo es inconsistente y parte del prejuicio de que todo es materia o todo emerge de la materia, lo cual es falso. En efecto, si todo fuese materia, no habría distinción de entes, pues no se tendría una razón suficiente para que la materia no fuese una unidad homogénea. Pero si hay entes y estos se distinguen unos de otros, entonces esa distinción ya no puede ser por la materia, que es lo que de común tienen los entes físicos. Puesto que la materia es determinable, no puede darse a sí misma la determinación. En efecto, la materia no puede ser ella misma su propio principio para distinguirse en un ente físico numéricamente diferente de otro ente físico. Si la materia tuviese en sí misma su principio de distinción y determinación, entonces ya no sería determinable, pero el ser determinable es su característica básica. Luego, la materia no tiene en sí su principio de distinción, y, por tanto, que la materia de un ente se distinga precisamente de la materia de otro ente se debe a que ambas materias están determinadas por principios distintos de la materia. Estos principios son llamados formas. Luego, si la forma también es real en los entes físicos, el materialismo no puede ser verdadero. Las formas son tan reales como la materia.

Podrían darse más argumentos contra el materialismo, pero esto nos llevaría mucho espacio, pero es suficiente, por motivos prácticos, con el que he proporcionado para mostrar que el prejuicio materialista de que sólo existe la materia o todo es reducible a ella es una falsedad.

También podemos decir que es el hombre el que se plantea saber si la IA puede superar a la INA. Este simple planteamiento de quién es el sujeto que cuestiona o se propone una problemática nos conduce a tener que reconocer que ciertamente la IA no se plantea a sí misma el problema de si la Inteligencia Humana es insuperable por ella, así como tampoco problematiza sobre su propia existencia. No tenemos razones para afirmar que la IA problematice su existencia.

Con el simple señalamiento de estos hechos podemos ver que la IA no filosofa. Así pues, aunque parece ser que plantear la hipótesis de la noogénesis artificial es relativa a los autómatas y su capacidad, en realidad no es así. Más bien, la hipótesis tiene que ver con el ser humano y sobre su pensamiento y conciencia.

La perfecta imitabilidad de las actividades de pensamiento mediante autómatas es una hipótesis que no concierne a los autómatas, sino al pensamiento; es decir: los hechos «por explicar» son los hechos del hombre, y los hechos «conocidos», que deben «servir para la explicación», son los de las máquinas.[3]

Es claro que, si la pregunta es si un autómata puede igualar y superar la inteligencia del hombre, la respuesta dependerá de qué determinemos que es la inteligencia del hombre, porque saber si es capaz o no la máquina de igualar la inteligencia humana será algo que deberá ser juzgado a la luz del criterio por el distingamos qué es la inteligencia humana.

Así, si la hipótesis de la noogénesis artificial es más bien una cuestión sobre el hombre más que sobre la

máquina, como bien lo ha hecho notar el profesor Agazzi, y esto se debe a que, aunque somos conscientes de que somos inteligentes, nos es difícil explicar cuál es la naturaleza de nuestra inteligencia (coincidimos aquí con el juicio de Jacques Maritain de que sólo el realismo-crítico ha sabido entender la naturaleza del conocimiento y la inteligencia humana). Pero, entonces, el planteamiento de la hipótesis de la noogénesis artificial resulta injustificado, porque, en efecto, lo que nos resulta problemático no es la máquina, porque no es ella la que nos causa dificultad para explicarla; de hecho, al contrario, podemos explicar perfectamente por qué la máquina puede hacer y comportarse como lo hace, dado que nosotros los humanos la construimos; más bien, lo que nos resulta problemático es poder explicar la inteligencia humana que no creamos, sino que tan sólo experimentamos porque nos ha sido dada como un don. Debería, pues, ser claro que, si nosotros inventamos a la IA, sabemos perfectamente que cuando la construimos, nunca le infundimos un alma, y, por tanto, no puede ser inteligente, puesto que la inteligencia tiene por condición el alma, dado que es una facultad de ésta.

Pues bien, esta alma es inmaterial, y como el planteamiento de la hipótesis de la noogénesis artificial supone el materialismo, entonces es imposible lo que postula la hipótesis como posible. En otras palabras, no puede haber inteligencia sin alma ni alma inteligente sin ser inmaterial. Esto no es una petición de principio, sino una conclusión. La inmaterialidad del alma es la conclusión de haber buscado la razón suficiente que explique ciertos hechos como la conceptualización universal, la libertad, la autoreflexión, entre otras, etc. De nuevo, mostrar esto mediante los argumentos correspondientes nos llevaría mucho espacio. Sin embargo, con tal de aportar alguna razón, demos la que Roger Verneaux considera la prueba metafísica más profunda de por qué la materia no piensa: si algo es material, está determinado bajo una forma, de modo que no puede recibir otra forma sino adquiriendo una nueva determinación, pero esto tiene por consecuencia dejar de estar bajo la determinación anterior, lo que implica que lo que recibe una nueva forma ha de destruirse, pues el ser que es depende de la forma que ya tiene. Pero el ser humano, cuando conoce, recibe las formas de lo que conoce, de lo contrario no conocería, y como puede recibir múltiples formas sin dejar de ser un ser humano, entonces aquello en que recibe estas formas, a saber, la inteligencia, no puede ser material.[4]

Aunque hay quien sostiene que no existen objeciones filosóficas serias para negar que el ser humano pueda llegar a construir a un ser humano, lo cual parece ser cierto, el punto que tratamos no es si el hombre puede mediante la tecnología en un futuro producir nuevos seres humanos por métodos artificiales que no impliquen la gestación natural. La cuestión es si los autómatas pueden igualar la inteligencia humana. Dado que esta inteligencia está en dependencia del alma humana, la cuestión a la que debemos remitirnos no es si el hombre va a poder crear cuerpos artificiales potencialmente aptos para almas humanas, sino si va a poder producir artificialmente almas humanas o cualquier otro tipo de alma inteligente por medio de la tecnología. Esto no es posible. El hombre no es capaz de crear almas racionales. Sólo Dios puede crearlas, porque "toda generación se produce ya *ex materia*, ya *ex nihilo*. Pero un espíritu no puede proceder de una transformación de la materia. Así, pues, es sacado de la nada, lo que equivale a decir que es creado"[5], porque la nada no es causa. Esto lo sabemos. Si sabemos que no tenemos capacidad de crear almas como tan sabemos que no somos capaces de volar tan sólo con el uso de nuestro cuerpo, entonces debería resultarnos obvio que nuestros inventos por muy avanzados que sean, no pueden llegar a ser verdaderamente inteligentes, puesto que la inteligencia es una facultad de un alma inmaterial subsistente. Como hace ver Edward Feser en su artículo *Artificial intelligence and magical thinking*, la inteligencia artificial, al igual que la magia, es una simulación convincente pero no una realidad, y deberíamos saber que es una simulación porque nosotros los humanos las construimos precisamente para que simularan, lo mismo que inventamos la magia para aparentar ilusiones.

En última instancia, lo que está haciendo quien propone la hipótesis en cuestión, no es otra cosa que suponer que el pensamiento inteligente del humano es puramente resultado de una interacción material, lo que implica que no depende esencialmente de un alma, puesto que el alma es, por principio, distinta de la materia. Pero no sólo esto, porque el planteamiento mismo de la hipótesis de la noogénesis se levanta, como hice observarlo, dentro de la presuposición materialista, por consiguiente, quien propone la hipótesis en última instancia está por lo menos presuponiendo que es irrelevante para la inteligencia del ser humano que éste tenga un alma inmaterial. Se presenta en este caso el hacer irrelevante si el alma humana es inmaterial bajo la apariencia de poner en discusión el horizonte entre material e inmaterial a propósito de una pregunta que parece legítima: ¿puede ser verdaderamente inteligente la máquina? Pero en realidad, el planteamiento de la hipótesis supone una base materialista y esa base destruye el entendimiento que podemos tener de la inteligencia humana, por lo que la posibilidad de que la máquina iguale al hombre sobre la base materialista no es porque en la máquina llegue a aparecer la inteligencia, sino porque ocultamente se ha reducido el hombre a la máquina, destruyendo, dejando de lado la base filosófica por la que es entendible, la inteligencia humana: el realismo metafísico.

Ahora veamos la inconsistencia de la hipótesis de noogénesis artificial.

1. La máquina puede llegar a ser inteligente e igualar en esto al hombre.

2. Pero si la máquina puede ser inteligente como el hombre, entonces la inteligencia no depende de un alma inmaterial, porque siendo los seres humanos los que construimos la IA, sabemos que no le dimos a la IA un alma inmaterial y subsistente y que tampoco somos capaces de crearla.

3. Por consiguiente, decir que la máquina puede no sólo imitar operaciones resultado de la inteligencia humana, sino que ella misma puede ser verdaderamente inteligente es haber aceptado que la inteligencia no requiere de un alma inmaterial y subsistente, y que ella queda explicada exclusivamente por causas materiales.

4. Pero si la inteligencia se debe a causas puramente materiales, entonces el hombre ya no es un hombre, sino una especie de autómatas, porque la inteligencia que lo caracteriza en su especial dignidad tampoco ya es inteligencia. En efecto, lo que llamamos inteligencia en el hombre no puede explicarse sino inmaterialmente, por tanto, si obligamos, conceptualmente, a la inteligencia a consistir en nada más que interacciones materiales, entonces hemos destruido conceptualmente la inteligencia, aunque la sigamos llamando así.

5. Por tanto, plantear la hipótesis de la noogénesis artificial es, de hecho, destruir la noción de inteligencia, porque es asumir una postura en la que ésta no tiene cabida.

En realidad, la inteligencia es una potencia en el hombre por la cual es capaz de leer la interioridad ontológica de la realidad, es decir, sus esencias. La inteligencia tiene una función intuitiva e intencional, intuye las cosas aprehendiendo sus esencias. Si analizamos el término "inteligencia artificial", podríamos decir que este concepto parece ser usado de modo analógico, con analogía extrínseca. En el caso de las máquinas, el término inteligencia es aplicado a algo de un modo y grado de ser que es inferior al ser humano, por cuanto que el ser humano es su creador. Este uso del término "inteligencia" para nombrar cosas de un orden inferior no sólo es aplicado a las máquinas, sino que también y antes ha sido aplicado a los animales, con los cuales el ser humano comparte varias características del orden sensible.

Así, por ejemplo, solemos decir que un perro es inteligente o que el delfín o el elefante dan muestras de inteligencia. Y esto lo decimos cuando observamos un tipo de comportamiento en los animales que es parecido al de los seres humanos, pero cuyo comportamiento en los seres humanos se funda en su inteligencia o creemos que se funda en ella.

El término “inteligencia” así empleado tiene un valor pragmático, pero es condición para la confusión, pues puede hacernos creer que el animal irracional es verdaderamente inteligente en el sentido propio del término. De hecho, como expone el Dr. Manuel Ocampo Ponce en su artículo *Inteligencia artificial (IA) y producción de conceptos en santo Tomás de Aquino* no es posible, en sentido estricto establecer una analogía de proporción o atribución entre la inteligencia humana y la IA porque el intelecto humano no es una potencia pasiva, sino una potencia activa, que cuenta con capacidad de mover a otro, mientras que la IA es potencia pasiva, es decir, capacidad de ser movido por otro. Y, dado que la IA es una potencia pasiva no cuenta con la capacidad para abstraer esencias.[6]

En el caso de la máquina, creer que ésta puede llegar a pensar es un contrasentido porque la máquina es máquina precisamente porque no piensa y si piensa, entonces ya no es una máquina. El pensamiento inteligente es un conjunto de actos humanos estructuralmente inigualables por seres vivos como los animales irracionales, cuanto con mayor razón son inigualables por seres no vivos como las máquinas.

III Conclusión

Llamamos a la inteligencia artificial "inteligencia" porque ofrece resultados que son logrados, cuando los realiza el hombre y no la máquina, por recurrencia a un acto del entendimiento en su función discursiva. En un sentido amplio, hablamos de pensamiento para referirnos a cualquier proceso psíquico consciente o inconsciente operado por el ser humano. En un sentido restringido, llamamos pensamiento a un proceso psíquico consciente. La máquina y el entendimiento pueden dar lugar a efectos similares, pero mientras la máquina puede realizar efectos que son similares a la función discursiva del entendimiento, no puede realizar un acto intuitivo como lo hace el entendimiento ni es capaz de intencionalidad ni de captar la verdad, puesto que ésta implica una adecuación del entendimiento con la realidad, adecuación que sólo resulta posible en el orden inmaterial del ser.

La máquina es invención humana y produce efectos que parecen venir de una inteligencia; y esto es porque efectivamente vienen de una inteligencia, pero ésta no está en la máquina sino en el hombre que la ha inventado. Podemos seguir llamando IA a los artefactos o a las máquinas, porque al fin y al cabo ellas también son efectos resultado de la dirección de la inteligencia del hombre, por lo cual dicen relación a esta causa. Pero hemos de reconocer que llamarlas inteligentes es sólo por relación a la inteligencia humana que las inventó y no porque en ellas mismas esté el principio de la vida inteligente.

La Inteligencia Artificial (IA) nunca podrá igualar o superar a la Inteligencia No-Artificial (INA), entre la que está la inteligencia humana, debido a que la IA carece de un componente esencial para la verdadera inteligencia: el alma inmaterial. La inteligencia humana no puede ser reducida a procesos materiales o algorítmicos, ya que su naturaleza implica una capacidad de comprensión profunda e intencionalidad, capacidad de captar las esencias de las cosas, algo que las máquinas, por su naturaleza no inmaterial, no pueden lograr.

Para cerrar citaré las palabras del maestro en filosofía política Jesús Ernesto Magaña, “En la era de la súper-

técnica y el avance de la ciencia, en el milenio de la ilustración humana, cuando los pregoneros del progreso vaticinaban la muerte del oscurantismo y la religión, con el avance de la IA queda claro más que nunca que el alma existe.”[7]

-
- [1] Agazzi, Evandro, "Alcune osservazioni sul problema dell'intelligenza artificiale", *Rivista di Filosofia Neo-Scolastica*, Vol. 59, No. 1 (GENNAIO-FEBBRAIO 1967), p. 2.
- [2] *Ibid*, p. 10.
- [3] *Ibid*, p. 9.
- [4] Verneaux, Roger, *Filosofía del Hombre*, Barcelona: Herder, 2008, p. 115-116.
- [5] *Ibid*, p. 220.
- [6] Ocampo Ponce, M. (2024). “Inteligencia artificial (IA) y producción de conceptos en santo Tomás de Aquino”. *Conocimiento y Acción*, 4(2), p. 4.
- [7] La cita es de las palabras pronunciadas por el maestro Magaña en una conversación personal que sostuve con él.

FUENTES DOCUMENTALES

1. Agazzi, Evandro. "Alcune osservazioni sul problema dell'intelligenza artificiale", *Rivista di Filosofia Neo-Scolastica*. Vol. 59, No. 1 (GENNAIO-FEBBRAIO, 1967), pp. 1-34.
2. Feser, Edward. "Artificial Intelligence and Magical Thinking." Edward Feser (blog), 20 de noviembre de 2017. <https://edwardfeser.blogspot.com/2017/11/artificial-intelligence-and-magical.html>.
3. Ocampo Ponce, M. (2024). “Inteligencia artificial (IA) y producción de conceptos en santo Tomás de Aquino”. *Conocimiento y Acción*, 4(2). <https://doi.org/10.21555/cya.v4.i2.3212>
4. Ocampo Ponce, M. (2024). “Algunos fundamentos de la filosofía realista de Santo Tomás de Aquino para una valoración de la inteligencia artificial (IA)”. *Perseitas*, 12, 72-92.DOI:<https://doi.org/10.21501/23461780.4724>



LA INTELIGENCIA ARTIFICIAL EN MEDICINA: CINCO PROBLEMAS BIOÉTICOS

‘La inteligencia artificial en medicina’.
Imagen generada por IA en leonardo.ai

**Dra. María Elizabeth
de los Ríos Uriarte
&
Dr. José Sols Lucía**

Los avances en la tecnología han desplegado un sinfín de posibilidades en muchas áreas de la vida humana. Sus aplicaciones y desarrollos en la atención de la salud, tanto en la prevención como en la atención y la recuperación terapéutica, son asombrosos, especialmente aquellos que funcionan con sistemas de inteligencia artificial (IA). Sin embargo, dado que ni la ciencia ni la tecnología son neutras, éstas responden a intereses que en ocasiones pueden levantar serios cuestionamientos éticos sobre sus usos, alcances y limitaciones.

En este artículo describiremos, en primer lugar, algunos de los desarrollos tecnológicos que utilizan softwares de inteligencia artificial en el área de la Medicina, y en general en el mundo de las ciencias de la salud, para, en segundo lugar, analizar cinco cuestiones bioéticas referentes a su uso y sus posibles consecuencias: 1) la relación médico-paciente, 2) el uso de robots para actividades consideradas como humanas, 3) el peligro de los sesgos, 4) el fin de la confidencialidad en la información médica, y 5) el incremento de la brecha digital.

II

Primeros usos y aplicaciones de la IA en el campo de las ciencias de la salud

Después de su inicio con Alan Turing y su *Test* de 1950,[1] las capacidades de los sistemas que operan con algoritmos comenzaron a pensarse cada vez más y a desarrollarse de formas muy sofisticadas. En 1956, la Conferencia de Dartmouth[2] no sólo fue un intento por definir la IA a partir de estudios preliminares sobre el funcionamiento del cerebro estructurado en capas y redes —que posteriormente daría lugar a términos como “machine learning” o “deep learning”—, sino que constituyó también el punto de partida para imitar ciertas acciones humanas en máquinas y sistemas que las pudieran realizar de manera más eficiente y rápida. Se vio que la IA podía incorporar nuevos conocimientos a medida que utilizaba otros anteriores, algo que algunos denominarían “aprender”; de ahí que hoy se siga afirmando que la IA emula funciones humanas tales como el aprendizaje.[3] Huelga decir que esta *capacidad de aprender* le viene dada a la IA gracias a la programación que los seres humanos introducen en ella mediante la cual puede asociar y generar patrones; por sí misma, la máquina no piensa ni aprende. Sin embargo, aporta un conocimiento nuevo que nadie le había introducido.

Las primeras aplicaciones de la IA al campo de la salud no tardaron en aparecer: la Universidad de Stanford creó Dendral a inicios de los sesenta y Mycin en 1973, programas que diagnosticaban y trataban infecciones bacterianas de la sangre.[4] Posteriormente se desarrolló Emycin, un programa que diagnosticaba un mayor número de enfermedades y tenía la capacidad de incorporar gran número de datos nuevos.[5] Estos sistemas resultaron efectivos incluso para diagnosticar enfermedades pulmonares; tal fue el caso de Puff, creado en 1975. Lo asombroso de Puff era que no sólo incorporaba nuevos conocimientos introducidos por profesionales de la salud, sino que asociaba los anteriores realizando nuevos algoritmos y ampliando con ello su capacidad de dar resultados de manera rápida y efectiva. No obstante, como en todo sistema de IA, no estaba exento de errores.

Desde entonces los progresos han sido importantes, principalmente en patología, radiología y epidemiología. Se ha logrado detectar, con gran precisión numérica, rangos anormales de cifras que miden funciones orgánicas, formas, patrones o tejidos anómalos, por ejemplo, en enfermedades potencialmente devastadoras como el cáncer o en infecciones letales de diversa índole. El ejemplo más común es el análisis de imágenes radiológicas para la detección de cáncer de mama.[6]

También hay otros usos no menos importantes. La robótica ha entrado en los quirófanos y está transformando la cirugía.[7] Igualmente, hay robots que acompañan a pacientes y les controlan sus niveles de insulina en sangre. Existen aplicaciones descargables en los teléfonos móviles que miden signos vitales y alertan a los familiares del enfermo o al personal médico de rangos fuera de lo normal, incluso cuando se trata de caídas o golpes, o de una alteración significativa en su conciencia, lo que permite que acudan rápidamente a atenderle. Todos sabemos que en Medicina unos minutos pueden ser decisivos para la salud o para la vida de una persona, por lo que la velocidad de intervención es fundamental.

III

Problemáticas bioéticas de la IA en Medicina

No obstante, a pesar de todos los avances que acabamos de señalar, y de otros muchos que aquí no podemos mencionar, surgen reflexiones acerca de los alcances y fines de la introducción de la IA en Medicina, en particular la pregunta sobre su pertinencia ética en algunas etapas del proceso de atención médica, dadas las posibles consecuencias que puede llevar consigo. Veamos cinco de ellas.[8]

A. La relación médico-paciente

La relación médico-paciente no siempre ha seguido el modelo hipocrático,[9] en el que se busca el bien de ambos, paciente y médico, como fin último. En este modelo no sólo importa el paciente; el médico también sale beneficiado de su trabajo, dado que al poner su conocimiento al servicio de otros, realiza su vocación personal. Sin embargo, hay otros modelos de relación médico-paciente,[10] y parece que ahora prime el utilitarista,[11] en buena medida debido al hecho de que la IA despersonaliza la Medicina. Cuando un paciente introduce su historial clínico en un programa de IA, éste se limita a preguntar lo que se requiere para completar una primera faceta de la Medicina, la de la recopilación de datos empíricos, a partir de los cuales se obtendrá un primer diagnóstico. Sin embargo, por muy acertadas que sean las preguntas y por muy ceñidas que estén al estado de cada paciente, no logran sustituir por completo la relación médico-paciente. Ésta se construye en un entorno de confianza y de amabilidad interpersonal, en el que se da un encuentro entre personas humanas que conduce a un compromiso. Edmund Pellegrino, padre de la denominada Ética de las Virtudes, propone que se desarrollen precisamente virtudes tales como la confianza, la compasión, la justicia o la fortaleza, entre otras.[12] Estas virtudes requieren tiempo y sobre todo un contacto cercano entre las personas involucradas, algo que difícilmente se logra si el médico es sustituido por la IA.

Es cierto que la IA aporta rapidez, exactitud y precisión a la atención médica, pero con ella corremos el riesgo de desplazar la capacidad intuitiva y la experiencia profesional del médico, con lo que podemos alejar al paciente de la práctica humana basada en los principios hipocráticos de beneficencia, no maleficencia y justicia. De ahí que convenga ir más allá del dicotómico “sí o no a la IA en Medicina” y tratemos de articular el factor humano con las nuevas tecnologías. La IA debe ayudar al médico en el procesamiento de datos, pero no puede sustituir su actividad profesional. Hay virtudes humanas que una máquina no puede suplir. Jamás podremos decir de una máquina que es virtuosa, pero siempre será posible predicarlo de una persona. Los principios de respeto a la vida y a la dignidad, de libertad, de justicia, de beneficencia y de no maleficencia, expresados en el Juramento Hipocrático, constituyen el eje vertebrador de una buena práctica de la relación médico-paciente, con lo que logramos salvaguardar la relación personal permeada no sólo por hechos empíricos, sino sobre todo por significados antropológicos y jerarquías axiológicas.

B. Robots en lugar de personas

El rápido desarrollo de sistemas cada vez más funcionales y complejos de IA ha dado pie a la creación de robots con forma humana capaces de ejecutar tareas hasta ahora propias de un enfermero, un cuidador, un auxiliar, un técnico o incluso un médico. Este fenómeno, en continuo crecimiento y mejoramiento, permite, en el mejor de los casos, que los talentos humanos sean derivados a tareas más trascendentales y que las actividades rutinarias sean realizadas por robots. Sin embargo, también puede derivar en una pérdida del sentido del cuidado como un acto de reconocimiento propio de nuestra esencia como personas humanas. [13]

El cuidado exige la presencia de una persona que entienda y empatice[14] con otra en quien ha decidido reconocerse y con quien entabla una relación afectiva, racional y psicológica de solicitud y atención no proveniente de un imperativo o mandato, sino de un autorreconocimiento en la figura del otro, como movimiento que va de la mismidad a la alteridad. La Ética del Cuidado trasciende la esfera técnica de las acciones inmediatas y se ocupa de transformar la realidad con compasión.[15] Sólo los seres humanos somos capaces de compasión, una actitud que sana la herida del otro, al que restituye en su valor ontológico. Cuando la sociedad relega la tarea compasiva a las máquinas, se desentiende progresivamente de los más frágiles y acaba por descartarlos. El resultado es una sociedad apática, regida con criterios

pragmáticos, en la que unas personas valen más que otras. Se pierde el sentido de comunidad. El máximo desarrollo de cada persona deja de ser el fin de la sociedad, que se acaba volviendo fría.

Reiteramos que no se trata aquí de decir “sí o no a la IA”, sino de proponer su incorporación armoniosa, de tal modo que no sustituya nunca la actividad propiamente humana, sino que la complemente, sobre todo en aquellos momentos en los que el cuidado puede comportar desgaste y quiebre emocional. Los principios de sociabilidad[16] y de subsidiariedad[17] tienen cabida en esta reflexión bioética para brindar luces sobre la responsabilidad humana, especialmente en relación con las personas y grupos más vulnerables.

C. El peligro de los sesgos

La IA tiene sesgos algorítmicos que pueden tener consecuencias en prácticas de exclusión social, xenofobia, racismo y homofobia, entre otras.[18] No sólo eso: los sesgos pueden provocar también un retraso importante en el diagnóstico correcto, lo que perjudica la salud de los pacientes, incluso provoca su muerte. Es un tema que no podemos tomarnos a la ligera. Clarifiquemos la idea de sesgo con un ejemplo: suelen sufrir la patología H los seres humanos de tal raza y de tal clase social; cuando una persona que es de otra raza y de otra clase social sufre esa patología, la IA tiende a descartar ese diagnóstico por lo primero que hemos dicho, con lo cual comete un error grave.

Existen estrategias para mitigar en el área de la salud los posibles sesgos de la IA: incluir un mayor número de muestras, doble verificación de resultados, aprendizaje interdisciplinar de los algoritmos, etc. Sin embargo, ninguna de estas estrategias es tan eficiente como la práctica según la cual un médico u otro profesional de la salud supervisa personalmente los datos empíricos a fin de discernir un diagnóstico. La capacidad analítica y la experiencia del médico, así como la correlación que establece entre los datos concretos y la historia del paciente, pueden ser valiosas para diagnosticar una enfermedad o para tratar un padecimiento de forma integral, teniendo en cuenta el conjunto de la persona, y no sólo unos datos empíricos. Los principios de no maleficencia, de justicia y de libertad responsable juegan un papel fundamental como talantes bioéticos que hacen frente a posibles sesgos que dañarían la práctica de la medicina.

D. El fin de la confidencialidad en la información médica

La historia clínica del paciente está cargada de datos sensibles; lo mismo sucede con los análisis de laboratorio o con cifras clínicas, imágenes, patrones antropomorfos con datos biométricos altamente sensibles y demás información que debe ser resguardada y tratada con la máxima seguridad y cautela. ¿Dónde se guardan estos datos en la actualidad? Existen inmensas bodegas informáticas que almacenan procesadores colosales con información de múltiples aplicaciones, softwares y programas; existe el peligro de que todo eso sea hackeado, con lo que haya información que sea filtrada con mala voluntad.[19] En la era del capitalismo informacional, que ha sustituido al financiero, nada tiene más valor que la información. Ésta es corolario de la obtención de poder, un poder que puede ser utilizado de diversas formas: colonialismo epistemológico, intercambio de personas “descartables” para tareas “confidenciales” o “misiones secretas”, hegemonía económica, destrucción de agendas nacionales, planes geopolíticos no tan benéficos o chantaje a personas ricas o famosas, entre otras muchas posibilidades. La vida, la salud, la identidad digital y la información privada de cientos de miles de personas está siendo resguardada en bases frágiles, permeables y susceptibles de ser interceptadas. ¿Cómo manejar, entonces, esta información y cómo resguardarla adecuadamente? La IA aún no ofrece una respuesta fiable. El principio de confidencialidad de datos personales, al igual que el de no maleficencia, quedan, de este modo, en entredicho.

E. Incremento de la brecha digital

No es algo nuevo el hecho de que la introducción de nuevas tecnologías aumente las desigualdades en lugar de reducirlas, debido a que son los grupos o los países más ricos los que se hacen antes con ellas, con lo que su dominio sobre los demás se acentúa todavía más, y la situación de precariedad de los grupos o países más desfavorecidos se vuelve más dramática. Esta realidad, que no es nueva, se hace especialmente sensible al hablar de la introducción de la IA en la Medicina. Grupos y países que ya viven actualmente en una situación de privilegio tendrán acceso a los últimos avances en IA en el campo de la salud, mientras que los demás tendrán que verlos en las películas.

Ahora bien, la salud no es una mercancía, sino un derecho humano. Son muchos los países que han firmado la Declaración Universal de Derechos Humanos de 1948, que reconoce el derecho a la salud, entre otros. Todos esos países deben hacer lo posible por que las nuevas tecnologías estén al servicio de todos, y no sólo de una minoría privilegiada. Pero en la práctica no es así.

Cuando se invierte en tecnología y se descubren nuevos procedimientos menos invasivos y peligrosos, más precisos y exitosos, nuevas técnicas que permiten observar el comportamiento de genes y células en vías de desarrollo sin tener que aprobarse para su estudio protocolos interminables de criterios éticos; cuando se tiene la posibilidad de revertir enfermedades mediante tecnología avanzada; más aún, cuando se saborea la inmortalidad y el no envejecimiento como algo inminente, entonces de hecho la economía derriba el valor universal de la vida y de la salud para erigirse en criterio de selección y venderse al mejor postor.

Los avances tecnológicos son accesibles sólo para una minoría de personas que gozará de mejor salud, de mayor posibilidad de supervivencia y de un mayor índice de prevención de futuros padecimientos. La salud se convierte en un privilegio y deja de ser un derecho.

Como ya hemos dicho varias veces más arriba, al señalar este problema bioético y socioeconómico no nos oponemos a la introducción de la IA en la Medicina, sino que nos limitamos a recordar que la ciencia y la tecnología deben estar al servicio de toda la humanidad, y no sólo de una minoría privilegiada. La comunidad de naciones debe tomar medidas para que el avance médico no quede sujeto a la lógica del mercado.

Este es un problema que se aborda cuando en Bioética hablamos del principio de justicia: dar a cada uno lo que se merece en función de su dignidad humana. Al mencionarlo aquí, articulamos ese derecho de justicia con otros dos, el de sociabilidad y el de subsidiariedad, todos ellos ya mencionados más arriba.

IV

Conclusiones

Es evidente que la rápida incorporación de la IA a las ciencias de la salud conlleva grandes beneficios en la prevención y atención de enfermedades varias, así como en la posterior rehabilitación de padecimientos; no obstante, también presenta importantes riesgos que estamos moralmente obligados a considerar.

La Bioética es la rama de la Ética que estudia los dilemas presentados en las ciencias de la salud desde los principios éticos universales; por ello, le compete el análisis de las problemáticas derivadas del uso de la IA en Medicina.

Los principios bioéticos, en su mayoría plasmados en el Juramento Hipocrático y posteriormente sistematizados en diversos modelos de formulación, sirven de brújula para navegar por vericuetos éticos de la IA de tal forma que su uso no quede fuera de la reflexión moral, sino que más bien se integre como herramienta que ayuda a la atención brindada en una justa dimensión. Al igual que con el resto de técnicas y saberes, la IA debe constituir un medio al servicio de la humanidad, y no un fin en sí misma.

El debate en torno a los múltiples usos de la IA en la vida diaria y en particular en el campo de la salud levanta preguntas éticas y seguirá haciéndolo. Esto no ha hecho más que empezar. No hemos enumerado aquí todas las problemáticas bioéticas; han quedado varias en el tintero: por ejemplo, el tema de la responsabilidad en caso de errores técnicos de la IA que dañen seriamente a personas, o el del desempleo masivo generado por la sustitución del hombre por la máquina, entre otros.

Tal vez estemos asistiendo al nacimiento de una era tecnocéntrica. Muchas funciones hasta ahora realizadas por personas pasarán a estar gestionadas por sistemas informáticos y por robots. Debemos al menos preguntarnos cómo va a ser la vida humana en esa nueva era; o tal vez debamos más bien impedir que llegue esa era completamente y estemos llamados a prolongar el antropocentrismo, aun con apoyo de nuevas tecnologías.

[1]Ekmekci, Perihan Elif y Arda, Berna, *Artificial Intelligence and Bioethics*, New York: Springer, 2020, p. 4.

[2] *Ibid.*, p. 6.

[3]García-Vigil, José L., *Reflexiones en torno a la ética, la inteligencia humana y la inteligencia artificial*, México: Gaceta Médica de México 157/3, 2021, p.p. 311-314., DOI: <https://doi.org/10.24875/gmm.20000818>

[4] Braunstein, Mark L., *A Brief History and Overview of Health Informatics*, en *id.*, *Health Informatics on FHIR: How HL7's New API is Transforming Healthcare*, New York: Springer, 2018.

[5]Lidströmer, Niklas y Ashrafian, Hutan (eds.), *Artificial Intelligence in Medicine*, New York: Springer, 2022.

[6] Shaikh, Khalid; Krishnan, Sabitha; Thanki, Rohit, *Artificial Intelligence in Breast Cancer Early Detection and Diagnosis*. New York: Springer, 2021.

[7] Karcz, Konrad; Nawrat, Zbigniew; Gumbs, Andrew A. (eds.), *Artificial Intelligence and the Perspective of Autonomous Surgery*, New York: Springer, 2024.

[8] Sols Lucia, José; de los Rios Uriarte, María Elizabeth, *Bioética de la inteligencia artificial*, Madrid: San Pablo – Universidad Pontificia Comillas, 2024.

[9] 2.400 años después de su redacción por parte del médico griego Hipócrates, el Juramento Hipocrático sigue siendo el documento principal de ética médica para profesionales de la salud del mundo entero. En él se plasman directrices que deben regir la práctica médica en orden a recuperar la salud de la persona o, cuando ello no es posible, a paliar su dolor, todo ello rigiéndose por los principios de beneficencia, no maleficencia, autonomía y justicia, y respetando valores tales como la vida, la dignidad o la integridad. Cfr. Sandoval López, J.M, *Análisis sobre la aplicabilidad de los juramentos médicos desde la perspectiva de la bioética personalista con fundamentación ontológica*, Madrid: Medicina y Ética 35/2, 2024, p.p. 328-373.

[10] De los Rios, Ma. Elizabeth; Cerdio Domínguez, David; Hernández Gonzáles, Jhosue A.; Ricaud Vélez, Ignacio A., *Fundamentos antropológicos y éticos de la relación médico-paciente y su dinámica durante la pandemia por COVID-19*, México: Bioética 16/1, 2022, p.p. 17-23.

[11] El utilitarismo es una corriente de pensamiento en Bioética basada en la idea de bienestar inmediato y accesible para el mayor número posible de personas, así como en la premisa fundamental de que el valor

de la persona reside en su utilidad, que es cuantificable. Es una corriente cuestionada por muchos en Filosofía, dado que, aun resultando atractiva de entrada (“buscar el mayor bienestar posible para el mayor número posible de personas”), está cargada de problemas: ¿Cómo medimos el bienestar? ¿Hay más humanidad en una persona rica y culta que en otra pobre e ignorante? ¿A cuál de las dos salvamos, si tenemos que escoger? Cfr. Cerdio Domínguez, David; de los Ríos Uriarte, Ma. Elizabeth; García-Llaca, Elvira, *Cosmovisiones en bioética: interpretación del dolor y el sufrimiento*, México: Apuntes de Bioética, 6/1, 2023, p.p. 5-28.

[12] Pellegrino, Edmund D. y Thomasma, David C., *The Virtues in Medical Practice*, New York: Oxford University Press, 1993.

[13] Amorin Zucchetto, Milena et al., *Empatía en el proceso de cuidado en enfermería bajo la óptica de la Teoría del Reconocimiento: síntesis reflexiva*, Colombia: Revista Cuidarte 10/3, 2019, p. 9.

[14] García Moyano, Loreto, *La ética del cuidado y su aplicación en la profesión enfermera*, Chile: Acta Bioethica 21/2, 2015, p.315

[15] García Uribe, John Camilo, *Cuidar del cuidado: Ética de la Compasión, más allá de la protocolización del cuidado de enfermería*, Colombia: Cultura de los cuidados 57, 2020, vol. 24, p.p. 52-60.

[16] Sociabilidad: los seres humanos somos esencialmente sociables; nos necesitamos unos a otros para poder ser nosotros.

[17] Subsidiariedad: utilizar (en este caso) la IA sólo en la medida en que eso ayude a nuestra vida humana, y nunca en la medida cuando la perjudique.

[18] Amaya-Santos, Sua; Jiménez-Pernett, Jaime; Bermúdez-Tamayo, Clara, *¿Salud para quién? Interseccionalidad y sesgos de la inteligencia artificial para el diagnóstico clínico*, Pamplona: An Sist Sanit Navar 47/2, 2024, p.p. 2-3.

[19] Basil, Nduma N.; Ambe, Solomon; Ekhatior, Chukwuyem; Fonkem, Ekokobe, *Health Records Database and Inherent Security Concerns: A Review of the Literature: Cureus* 14/10, 2022, p. 3.

FUENTES DOCUMENTALES

1. Amaya-Santos, Sua; Jiménez-Pernett, Jaime; Bermúdez-Tamayo, Clara, *¿Salud para quién? Interseccionalidad y sesgos de la inteligencia artificial para el diagnóstico clínico*, Pamplona: An Sist Sanit Navar 47/2, 2024.

2. Amorin Zucchetto, Milena et al., *Empatía en el proceso de cuidado en enfermería bajo la óptica de la Teoría del Reconocimiento: síntesis reflexiva*, Colombia: Revista Cuidarte 10/3, 2019.

3. Basil, Nduma N.; Ambe, Solomon; Ekhatior, Chukwuyem; Fonkem, Ekokobe, *Health Records Database and Inherent Security Concerns: A Review of the Literature: Cureus* 14/10, 2022.

4. Braunstein, Mark L., A Brief History and Overview of Health Informatics, en íd., *Health Informatics on FHIR: How HL7's New API is Transforming Healthcare*, New York: Springer, 2018.

5. Cerdio Domínguez, David; de los Ríos Uriarte, Ma. Elizabeth; García-Llaca, Elvira, *Cosmovisiones en bioética: interpretación del dolor y el sufrimiento*, México: Apuntes de Bioética, 6/1, 2023, p.p. 5-28.

6. De los Ríos, Ma. Elizabeth; Cerdio Domínguez, David; Hernández Gonzáles, Jhosue A.; Ricaud Vélez, Ignacio A., *Fundamentos antropológicos y éticos de la relación médico-paciente y su dinámica durante la pandemia por COVID-19*, México: Bioética 16/1, 2022, p.p. 17-23.

7. Ekmekci, Perihan Elif y Arda, Berna, *Artificial Intelligence and Bioethics*, New York: Springer, 2020.

8. García Moyano, Loreto, *La ética del cuidado y su aplicación en la profesión enfermera*, Chile: Acta Bioethica 21/2, 2015.

9. García Uribe, John Camilo, *Cuidar del cuidado: Ética de la Compasión, más allá de la protocolización del cuidado de enfermería*, Colombia: Cultura de los cuidados 57, 2020, vol. 24, p.p. 52-60.

10. García-Vigil, José L., *Reflexiones en torno a la ética, la inteligencia humana y la inteligencia artificial*, México: Gaceta Médica de México 157/3, 2021, p.p. 311-314.
11. Karcz, Konrad; Nawrat, Zbigniew; Gumbs, Andrew A. (eds.), *Artificial Intelligence and the Perspective of Autonomous Surgery*, New York: Springer, 2024.
12. Lidströmer, Niklas y Ashrafian, Hutan (eds.), *Artificial Intelligence in Medicine*, New York: Springer, 2022.
13. Pellegrino, Edmund D. y Thomasma, David C., *The Virtues in Medical Practice*, New York: Oxford University Press, 1993.
14. Sandoval López, J.M, *Análisis sobre la aplicabilidad de los juramentos médicos desde la perspectiva de la bioética personalista con fundamentación ontológica*, Madrid: Medicina y Ética 35/2, 2024, p.p. 328-373.
15. Shaikh, Khalid; Krishnan, Sabitha; Thanki, Rohit, *Artificial Intelligence in Breast Cancer Early Detection and Diagnosis*. New York: Springer, 2021.
16. Sols Lucia, José; de los Rios Uriarte, María Elizabeth, *Bioética de la inteligencia artificial*, Madrid: San Pablo – Universidad Pontificia Comillas, 2024.

La vida no es sino una continua sucesión de
oportunidades para sobrevivir.
Gabriel García Márquez.

Mtro. Mauricio Correa-Herrejon

Introducción: IA y la crisis del pensamiento
humano

‘Inteligencia artificial y educación’.
Imagen generada por IA en leonardo.ai

El siglo XXI ha sido testigo de una transformación tecnológica que ha reconfigurado las estructuras fundamentales de la sociedad, y en ningún ámbito esta transformación es más profunda que en la educación. La irrupción de la Inteligencia Artificial (IA), con su capacidad para procesar, generar y replicar información a velocidades que trascienden

LA EDUCACIÓN EN LOS TIEMPOS DE LA INTELIGENCIA ARTIFICIAL: PENSAMIENTO CRÍTICO, DISRUPTIVO Y ANALÍTICO, COMO CLAVES DEL FUTURO



la capacidad humana, ha expuesto una crisis latente en la manera en que concebimos el conocimiento y su transmisión. Nos encontramos frente a una paradoja compleja: nunca en la historia habíamos tenido acceso a tanta información en tiempo real, y sin embargo, la capacidad colectiva de análisis crítico y discernimiento parece estancada o, en algunos contextos, incluso en retroceso.

La IA, lejos de ser una mera herramienta neutral, plantea desafíos ontológicos, epistemológicos y éticos de gran envergadura. Su capacidad para automatizar procesos cognitivos básicos, como la recuperación de información y la resolución de problemas estándar, ha desplazado la centralidad del docente como transmisor de saberes y del estudiante como receptor pasivo. Esta deslocalización de la autoridad epistémica exige repensar radicalmente los fines de la educación: ¿se trata de reproducir conocimientos o de formar ciudadanos capaces de pensar, cuestionar y transformar su realidad?, ¿estamos preparando a las nuevas generaciones para interactuar críticamente con la IA o simplemente para adaptarse a ella sin resistencia?

En este contexto, la educación ya no puede limitarse a la acumulación de datos, ni siquiera al dominio de habilidades técnicas. Lo que se impone es la formación de capacidades exclusivamente humanas: el pensamiento crítico, el análisis riguroso, la creatividad disruptiva y el juicio ético. Estas competencias no solo permiten interactuar con la IA de manera efectiva, sino también resistir sus posibles excesos, sesgos y usos indebidos. La "inteligización educativa" surge como una propuesta que articula el uso estratégico, ético y consciente de la IA con la potenciación de la inteligencia humana, revalorizando el papel del docente como mentor y del estudiante como sujeto activo en la construcción del conocimiento.

Esta nueva visión, implica un compromiso colectivo: las instituciones educativas deben transformar sus estructuras y metodologías; los docentes, resignificar su rol y actualizar sus competencias; los estudiantes, asumir una actitud activa y reflexiva; y los gobiernos, diseñar políticas que garanticen la equidad y la ética en el acceso y uso de la tecnología. La IA representa, así, tanto una amenaza como una oportunidad. Puede conducirnos a una sociedad hiperautomatizada y acrítica, o bien, si se integra con inteligencia y sensibilidad, a una era de emancipación intelectual y desarrollo humano integral.

II

Pensamiento crítico: defensa cognitiva y estratégica

El pensamiento crítico (PC) es, sin lugar a duda, la principal defensa cognitiva frente a la avalancha informativa de la era digital y la omnipresencia de la Inteligencia Artificial. Definido por Paul y Elder[1] como "el proceso disciplinado de conceptualizar, aplicar, analizar, sintetizar y evaluar información para alcanzar una conclusión fundamentada", el PC se presenta hoy, como una competencia estratégica, indispensable no solo para la formación académica, sino para la vida en sociedad, la toma de decisiones públicas y la ética profesional.

En un mundo donde los algoritmos de lenguaje natural pueden producir textos complejos, los motores de búsqueda filtran la realidad, y los sesgos algorítmicos[2] pueden reforzar estereotipos o manipular percepciones; la capacidad de evaluar críticamente la información es vital. Sin pensamiento crítico, los individuos se convierten en meros receptores pasivos, vulnerables a la desinformación, las narrativas sesgadas y las decisiones automáticas tomadas por sistemas que operan bajo lógicas que muchas veces desconocen.

CriThix: los pensadores críticos de la Era IA

Surge así el concepto de CriThix: pensadores críticos de la era digital, profesionales y estudiantes que, más allá del uso instrumental de la tecnología, desarrollan una postura epistemológica activa frente a la IA. Estos individuos no solo interactúan con algoritmos, sino que los interrogan, los contextualizan y, en muchos casos, los corrigen. El CriThix no teme a la IA, pero tampoco la idolatra. Su misión es construir sentido, no solo consumir datos.

Para formar a los CriThix, es necesario promover:

- Dominio estratégico: Comprender la finalidad y el impacto de la información que se procesa. Evaluar los resultados de la IA en función de objetivos humanos, sociales y éticos. Es aquí donde la educación debe conectar la teoría crítica con la práctica profesional. Por ejemplo, un médico debe cuestionar un diagnóstico generado por IA si no concuerda con su experiencia clínica o si se basa en datos insuficientes.

- Validación de resultados: En un entorno donde la información abunda, la veracidad escasea. Los CriThix validan fuentes, detectan inconsistencias y comprenden las limitaciones metodológicas de los modelos algorítmicos. Esta validación requiere pensamiento científico: hipótesis, evidencia, análisis y falsación.

- Resolución creativa de problemas: El pensamiento crítico no es solo defensivo, también es propositivo. La creatividad se potencia cuando se cuestionan los supuestos establecidos y se generan soluciones innovadoras. Un CriThix transforma el dato en conocimiento, y el conocimiento en acción con impacto social.

- Liderazgo ético: El pensamiento crítico también es pensamiento moral. El CriThix entiende que la IA puede reforzar desigualdades, invadir la privacidad o manipular decisiones. Por eso, se compromete con la equidad, la transparencia y el bien común. Cita clave: "No todo lo que puede hacerse debe hacerse. El juicio ético es más importante que la capacidad técnica"[3].

Formación en pensamiento crítico: claves pedagógicas

La enseñanza del pensamiento crítico exige una transformación metodológica; ya no basta con exámenes de opción múltiple o clases magistrales. Se requiere:

- Aprendizaje basado en problemas: Plantear dilemas reales que exijan análisis profundo y colaboración. Por ejemplo, ¿debe una escuela usar reconocimiento facial para controlar asistencia? ¿Cuáles son los riesgos éticos, legales y sociales?

- Evaluación de fuentes y datos: Analizar noticias falsas, detectar manipulación mediática y entender los modelos de negocio de las plataformas digitales. La alfabetización informacional se transforma en eje transversal de soporte.

- Debates argumentativos: Ejercitar la defensa razonada de posturas, reconociendo puntos débiles y fortaleciendo la argumentación lógica. Los debates fomentan el respeto, la tolerancia y la capacidad de negociación.

- Detección de falacias lógicas: Enseñar a identificar errores de razonamiento, desde generalizaciones apresuradas hasta falsas analogías. Las falacias son armas de manipulación cognitiva; detectarlas es un acto de libertad.

En definitiva, el pensamiento crítico es la base de la ciudadanía activa en tiempos de IA. Sin él, la sociedad se convierte en una masa manipulable; con él, se construye una democracia deliberativa, participativa y ética. Como señala Kahneman[4], el pensamiento lento, reflexivo y deliberado es la antítesis del impulso automático: es el ejercicio supremo de la libertad intelectual.

Pensamiento disruptivo: innovación desde la disconformidad

El pensamiento disruptivo (PD) es una capacidad esencial que, paradójicamente, suele estar subdesarrollada en los sistemas educativos tradicionales. Mientras que el pensamiento crítico se centra en cuestionar lo establecido, el PD va más allá: propone nuevas formas de ver, hacer y construir. Se nutre de la inconformidad productiva, la rebeldía fundamentada y la creatividad estratégica. Christensen[5], en su teoría sobre la innovación disruptiva, plantea que los cambios verdaderamente transformadores surgen al desafiar las estructuras preexistentes y generar modelos que, inicialmente marginales, logran desplazar a los dominantes gracias a su adaptabilidad, accesibilidad y enfoque en nuevas necesidades.

En el ámbito educativo, el PD exige repensar todo el sistema: ¿Qué enseñamos? ¿cómo lo enseñamos? ¿para qué lo enseñamos? La irrupción de la inteligencia artificial (IA) ha evidenciado la obsolescencia de ciertos métodos y contenidos, creando un entorno donde la capacidad de innovar en procesos pedagógicos, currículos y entornos de aprendizaje es crucial para la pertinencia y sostenibilidad de la educación.

La educación convencional, suele promover un pensamiento lineal y conformista, donde se espera que los estudiantes memoricen y repitan información en lugar de cuestionarla o reinterpretarla. Este enfoque puede sofocar la creatividad y la capacidad de pensar "fuera de la caja". Para fomentar una verdadera innovación, es esencial no solo pensar fuera de las estructuras existentes, sino también crear nuevas formas de pensamiento y acción que desafíen las convenciones actuales.

Características del pensamiento disruptivo

- Visión anticipatoria: El pensador disruptivo identifica oportunidades donde otros ven riesgos. Anticipa tendencias, cuestiona lo habitual y visualiza escenarios alternativos. Su perspectiva no es lineal, sino sistémica y prospectiva.
- Incomodidad constructiva: No se conforma con "lo que siempre ha sido así". La incomodidad frente a lo establecido se convierte en motor de transformación. Esta actitud es clave para enfrentar entornos VUCA (volátiles, inciertos, complejos y ambiguos), donde la adaptabilidad es sinónimo de supervivencia.
- Creatividad estratégica: Propone soluciones viables, sostenibles y con impacto. No se trata de improvisar, sino de diseñar con visión y rigor. La creatividad se articula con el conocimiento técnico, la experiencia contextual y la sensibilidad social.

Aplicaciones del pensamiento disruptivo en educación con IA

1. Redefinición curricular: Los contenidos deben responder a los desafíos del siglo XXI. Temas como alfabetización digital, ética algorítmica, sostenibilidad y ciudadanía global deben ocupar un lugar central; además, las competencias blandas (resiliencia, empatía, liderazgo colaborativo) deben integrarse transversalmente.
2. Modelos pedagógicos flexibles: La IA permite transitar de modelos homogéneos a modelos adaptativos: Aprendizaje personalizado, rutas de formación no lineales y evaluación formativa automatizada son solo algunas posibilidades. La educación se convierte en un proceso co-creado por estudiantes, docentes y tecnologías.
3. Docencia aumentada: El docente se convierte en facilitador del aprendizaje, mentor ético y diseñador de experiencias formativas. Utiliza la IA para analizar progresos, identificar dificultades y proponer

intervenciones personalizadas. La IA no la reemplaza, la potencia.

4. Entornos inmersivos y gamificación: Realidad aumentada, simuladores, mundos virtuales y entornos gamificados promueven el aprendizaje activo, la experimentación y la motivación intrínseca. Los estudiantes dejan de ser espectadores para convertirse en protagonistas de su proceso de formación.

5. Educación ubicua y asincrónica: La educación no se limita al aula ni al horario escolar. Con IA, el aprendizaje ocurre en cualquier momento y lugar, adaptándose a contextos laborales, personales y culturales diversos. Esto democratiza el acceso y promueve la autonomía.

Desafíos éticos del pensamiento disruptivo

·Equidad: ¿Quién accede a las tecnologías disruptivas?, ¿quién queda excluido? La innovación debe ser inclusiva y evitar brechas.

·Sostenibilidad: Las soluciones disruptivas deben considerar su impacto ambiental, social y cultural. No toda innovación es progreso.

·Humanización: El foco debe estar en el desarrollo humano integral. La tecnología es medio, no fin. La educación debe formar personas libres, críticas y empáticas, no solo usuarios de tecnología.

Disrupción con sentido

El pensamiento disruptivo es la respuesta a un mundo cambiante, pero debe estar guiado por la ética, la justicia social y el compromiso con el bien común. La educación no puede ser una fábrica de obediencia; debe ser un laboratorio de disrupción responsable, donde los estudiantes aprendan a cuestionar, crear y liderar transformaciones con sentido. La educación disruptiva, potenciada por la IA, tiene la responsabilidad y la oportunidad de crear un futuro más humano, equitativo y libre.

IV

Intelligización educativa: integración ética y estratégica de la IA

La simple incorporación de tecnología en procesos educativos no garantiza innovación ni calidad. De hecho, puede acentuar desigualdades, deshumanizar interacciones y promover una dependencia acrítica de herramientas automatizadas. Frente a este riesgo, surge el concepto de intelligización educativa: una estrategia integral, ética y crítica para la adopción de la Inteligencia Artificial (IA) en los sistemas educativos, orientada a potenciar -y no reemplazar- la inteligencia humana. Esta visión implica la integración consciente de la IA a partir de fundamentos epistemológicos, pedagógicos y sociales, garantizando que su uso responda al desarrollo humano integral y no a la mera eficiencia operativa.

Fundamentos de la intelligización

·Conocimiento técnico contextualizado: No basta con saber utilizar una plataforma de IA; es imprescindible comprender su lógica operativa, sus márgenes de error y sus alcances reales. Esto implica conocimientos de estadística, análisis de datos, inferencia y fundamentos de programación. La intelligización requiere que docentes y estudiantes entiendan lo que hay detrás de los algoritmos.

·Dimensión ética: Toda implementación de IA en educación debe partir de una reflexión ética sobre sus efectos. ¿Cómo protege la privacidad? ¿reproduce sesgos? ¿a quién beneficia? ¿a quién excluye? La intelligización se guía por principios de equidad, justicia y respeto por la dignidad humana.

·Perspectiva estratégica: La IA no debe ser una moda tecnológica, sino una herramienta alineada a objetivos pedagógicos claros. Su uso debe evaluarse según criterios de pertinencia, impacto y sostenibilidad. Las inversiones tecnológicas deben generar valor educativo tangible.

Pilares de la intelligización educativa

1. Formación docente sólida y continua

Los educadores son agentes clave en la intelligización. Su papel evoluciona del dominio disciplinar a la gestión de entornos de aprendizaje complejos, donde la IA coexiste con el juicio humano. Para ello, los docentes necesitan formación en:

- Fundamentos de IA: tipos de algoritmos, funciones predictivas, sistemas de recomendación.
- Alfabetización estadística: distribución de datos, inferencias, interpretación crítica de resultados.
- Ética tecnológica: dilemas de privacidad, sesgos algorítmicos, justicia digital.
- Diseño pedagógico con IA: selección de plataformas, análisis de datos de aprendizaje, personalización.

Esta formación debe ser continua, adaptada a contextos específicos y acompañada de comunidades de práctica que promuevan la innovación colaborativa.

2. Gobernanza tecnológica educativa

Las instituciones deben establecer marcos de gobernanza para el uso de IA que incluyan:

- Políticas claras sobre recolección, uso y protección de datos.
- Evaluación crítica y participativa de plataformas tecnológicas.
- Mecanismos de auditoría y rendición de cuentas sobre decisiones automatizadas.
- Inclusión de estudiantes, docentes y familias en la toma de decisiones tecnológicas.

La intelligización rechaza el modelo top-down. Busca procesos democráticos de adopción tecnológica, donde todos los actores educativos tengan voz.

3. Diseño de experiencias de aprendizaje aumentadas

La IA permite enriquecer las experiencias educativas, pero su valor está en el diseño pedagógico, no en la herramienta per se. Ejemplos de prácticas intelligizadas:

- Sistemas adaptativos de retroalimentación inmediata, que permiten a los estudiantes conocer sus avances y dificultades en tiempo real.
- Mapas de progreso personalizados, donde la IA recomienda recursos según estilos y ritmos de aprendizaje.
- Análisis predictivos para docentes, que alertan sobre estudiantes en riesgo de abandono o rezago, facilitando intervenciones oportunas.

Estos usos deben fomentar la autonomía, la motivación y la metacognición, evitando el control excesivo y la gamificación superficial.

4. Inclusión y justicia digital

La intelligización es incompatible con la exclusión. Para ser ética y efectiva, debe garantizar:

- Acceso equitativo a dispositivos, conectividad y soporte técnico.
- Capacitación digital para estudiantes y familias.
- Diseño de herramientas accesibles a personas con discapacidad.
- Evaluación de impactos en poblaciones vulnerables.

Sin inclusión, la IA se convierte en un factor de desigualdad. Con inclusión, es un motor de equidad y transformación.

Impactos esperados de la intelligización

- En la docencia: Empoderamiento profesional, capacidad de personalización, mejora en la toma de decisiones pedagógicas.
- En los estudiantes: Aprendizaje autónomo, compromiso, sentido de propósito, habilidades críticas para la vida digital.
- En las instituciones: Mejora en la eficiencia, innovación educativa, fortalecimiento de la cultura organizacional.

Desafíos y riesgos de la intelligización

- Tecnofilia acrítica: Adopción impulsiva sin análisis de pertinencia.
- Dependencia tecnológica: Reducción de la agencia docente y estudiantil.
- Monopolios de plataformas: Concentración de datos e intereses.

Estos riesgos deben enfrentarse con pensamiento crítico, gobernanza ética y cultura institucional robusta.

La intelligización como revolución humana

La intelligización no es una técnica, es una visión. Busca que la IA potencie lo mejor del ser humano: la creatividad, la ética, la reflexión. Transforma la educación en un espacio donde la tecnología se subordina al desarrollo integral, y donde cada estudiante encuentra herramientas para ser protagonista de su aprendizaje y de su vida. La intelligización es una apuesta por ese cambio humano, ético y profundo, que pone la tecnología al servicio de la libertad.

V

Riesgos y desafíos: sesgos, automatización acrítica y brechas digitales

La adopción masiva de la Inteligencia Artificial (IA) en la educación no está exenta de riesgos, muchos de los cuales amenazan con socavar los principios de equidad, autonomía y justicia los cuales deberían regir cualquier sistema educativo. Lejos de ser neutrales, las tecnologías emergentes reflejan, reproducen y amplifican las estructuras de poder y desigualdad que existen en la sociedad.[6] Por tanto, cualquier implementación acrítica de la IA puede convertirse en un vehículo para perpetuar la exclusión, estandarizar el pensamiento y deshumanizar el proceso educativo. Reconocer estos riesgos es el primer paso para mitigarlos, pero enfrentarlos requiere una acción decidida, informada y ética desde todos los niveles del ecosistema educativo.

A. Sesgos algorítmicos: reproducción de prejuicios y desigualdades

Los algoritmos de IA aprenden de los datos que se les proporcionan, pero si esos datos están contaminados por sesgos históricos, culturales o sociales, los resultados también lo estarán. Investigaciones como la de Bender[7] sobre los "loros estocásticos" evidencian cómo los modelos de lenguaje pueden replicar estereotipos raciales, de género o de clase, incluso cuando sus creadores no lo anticipan. En contextos educativos, esto puede traducirse en sistemas de evaluación que penalizan a estudiantes de ciertos orígenes, algoritmos de admisión que refuerzan la elitización, o plataformas que ofrecen contenidos diferenciados según patrones discriminatorios.

Los sesgos algorítmicos no son fallos técnicos; son problemas éticos. Requieren auditorías permanentes, transparencia en los modelos utilizados y participación activa de comunidades diversas en su diseño y validación. La ignorancia sobre estos sesgos no exime de responsabilidad; por el contrario, agrava sus consecuencias.

B. Automatización acrítica: pérdida de autonomía y humanización

Uno de los mayores peligros de la IA es la tentación de delegar procesos complejos a sistemas automatizados sin supervisión ni sentido crítico. Cuando las decisiones educativas —evaluaciones, recomendaciones, retroalimentación— son realizadas exclusivamente por máquinas, se pierde el matiz, la empatía y el contexto que sólo el juicio humano puede aportar.[8] La educación se reduce a un intercambio de datos, y los estudiantes pasan de ser sujetos activos a objetos de cálculo.

La automatización acrítica puede también fomentar una visión instrumental del aprendizaje, donde lo importante es cumplir métricas y no desarrollar pensamiento profundo. Esta lógica puede llevar a una "uberización" de la educación: rápida, estandarizada y despersonalizada, pero vacía de sentido.

C. Brechas digitales: exclusión tecnológica y ampliación de desigualdades

La IA tiene el potencial de democratizar el conocimiento, pero también de ampliar las desigualdades si no se garantizan condiciones de acceso equitativo. Las brechas digitales no se limitan al acceso a dispositivos o conectividad; incluyen también la capacidad de uso crítico, la alfabetización digital y la participación significativa en entornos digitales.[9]

Estudiantes de contextos rurales, comunidades indígenas, personas con discapacidad o sectores socioeconómicos desfavorecidos enfrentan barreras estructurales que les impiden beneficiarse plenamente de la IA. Si no se abordan, estas barreras profundizarán la exclusión y perpetuarán la segmentación educativa.

D. Vigilancia y privacidad: riesgos a la autonomía personal

La recolección masiva de datos educativos para alimentar sistemas de IA plantea serias preocupaciones sobre la privacidad y la vigilancia: ¿quién controla esos datos?, ¿con qué fines se utilizan?, ¿qué derechos tienen los estudiantes sobre su información? El uso de tecnologías de reconocimiento facial, análisis biométrico o seguimiento en línea puede convertir la escuela en un espacio de control permanente, erosionando la confianza y la libertad de expresión.

La protección de datos debe ser un pilar en cualquier estrategia de IA educativa. Es imperativo garantizar el consentimiento informado, el derecho al olvido y la soberanía digital. La privacidad no es un lujo; es un derecho humano fundamental.

E. Concentración de poder y monopolios tecnológicos

El mercado de la IA educativa está dominado por un pequeño número de corporaciones multinacionales que controlan plataformas, infraestructuras y datos. Esta concentración de poder limita la soberanía tecnológica de las instituciones educativas y genera dependencia; además, orienta la innovación hacia intereses comerciales, dejando de lado necesidades educativas locales, culturales o comunitarias.

Una respuesta ética y progresista exige fomentar el desarrollo de tecnologías abiertas, colaborativas y soberanas. La educación debe ser un bien público, no un nicho de mercado. La tecnología educativa no puede estar subordinada al lucro.

Propuestas para enfrentar los riesgos

- Establecer observatorios éticos de IA educativa que evalúen impactos, propongan regulaciones y difundan buenas prácticas.
- Fomentar la alfabetización digital crítica como eje curricular obligatorio.
- Promover tecnologías abiertas y comunitarias que respondan a contextos específicos.
- Crear espacios de diálogo y participación donde estudiantes, docentes y familias debatan sobre el uso de

la IA.

Riesgo como oportunidad para la acción ética

Los riesgos de la IA no son argumentos para rechazarla, sino llamados urgentes a repensar su implementación con rigor, ética y justicia. La educación tiene la responsabilidad histórica de formar ciudadanos críticos, capaces de entender la tecnología, cuestionarla y transformarla en herramienta de emancipación. El desafío no es solo técnico, es político, ético y pedagógico. La IA en educación debe ser innovación con justicia, inteligencia con conciencia y progreso con humanidad.

VI

Conclusión: educar para la emancipación intelectual en la era de la IA

La historia de la humanidad ha estado marcada por hitos tecnológicos que han reconfigurado la manera en que vivimos, nos relacionamos y aprendemos. La irrupción de la Inteligencia Artificial no es solo un nuevo capítulo en esa historia: es una revolución ontológica que cuestiona qué significa pensar, crear, decidir y ser humano en un mundo donde las máquinas simulan procesos cognitivos. Frente a esta transformación radical, la educación no puede ser espectadora ni rehén de la tecnología; debe ser protagonista consciente, crítica y propositiva de su propia reinención.

Este artículo ha demostrado que el mero acceso a la IA, sin pensamiento crítico, nos convierte en consumidores pasivos de datos; que la innovación, sin disrupción creativa, nos condena a repetir modelos obsoletos; que la automatización, sin análisis riguroso, erosiona la autonomía y la libertad. Hemos expuesto los peligros de una educación instrumentalizada por algoritmos opacos, pero también las posibilidades inmensas de una intelligenización ética, inclusiva y emancipadora.

Educar en tiempos de IA implica mucho más que integrar dispositivos o plataformas: implica formar CriThix, pensadores críticos capaces de auditar, contextualizar y transformar la información; implica fomentar el pensamiento disruptivo, para que los estudiantes se conviertan en arquitectos de futuros posibles, no en víctimas del determinismo tecnológico; implica construir una cultura de intelligenización, donde la tecnología se subordine a la dignidad humana, a la justicia social y al bien común.

Los riesgos son reales y urgentes: sesgos algorítmicos, brechas digitales, vigilancia masiva, concentración de poder. Pero también lo son las oportunidades: personalización del aprendizaje, docencia aumentada, inclusión global, innovación pedagógica. La disyuntiva no es IA sí o no; es IA con sentido humano o IA sin sentido alguno.

La educación, si ha de ser fiel a su misión histórica, debe formar ciudadanos libres, no usuarios acríticos; debe formar sujetos de transformación, no datos para consumo; debe formar seres humanos íntegros, capaces de usar la tecnología para ampliar los límites de la razón, la creatividad y la ética.

Nos enfrentamos a un desafío civilizatorio: domesticar la IA, convertirla en aliada de la humanidad y no en su amo silencioso. La única vía para lograrlo es una educación que piense, cuestione, imagine y actúe. En palabras de Hannah Arendt, "La educación es el punto en que decidimos si amamos lo suficientemente al mundo como para asumir la responsabilidad de él". Hoy, esa responsabilidad incluye moldear la IA, pero, sobre todo, moldear el pensamiento que sabrá utilizarla con sabiduría, compasión y coraje.

Educar en tiempos de IA es educar para la emancipación intelectual, para la disrupción ética y para la transformación del mundo. Ese es el legado que estamos llamados a construir. Y la historia nos juzgará, no

por cuán sofisticadas fueron nuestras máquinas, sino por cuán humanos supimos seguir siendo.

-
- [1] Paul, R., & Elder, L. (2006). *Critical thinking: Tools for taking charge of your learning and your life*. Pearson.
- [2] Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623. <https://doi.org/10.1145/3442188.3445922>
- [3] Sunstein, C. R. (2014). *Why nudge? The politics of libertarian paternalism*. Yale University Press.
- [4] Kahneman, D. (2011). *Thinking, fast and slow*. Farrar, Straus and Giroux.
- [5] Christensen, C. M. (1997). *The innovator's dilemma: When new technologies cause great firms to fail*. Harvard Business Review Press.
- [6] Noble, S. U. (2018). *Algorithms of oppression: How search engines reinforce racism*. NYU Press.
- [7] Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623. <https://doi.org/10.1145/3442188.3445922>
- [8] Luckin, R., Holmes, W., Griffiths, M., & Forcier, L. B. (2017). *Artificial intelligence and human learning: Towards an existential crisis?* UCL Knowledge Lab.
- [9] UNESCO. (2021). *Reimagining our futures together: A new social contract for education*. United Nations Educational, Scientific and Cultural Organization.

FUENTES DOCUMENTALES

1. Arendt, H. (1993). *Between past and future: Eight exercises in political thought* (2.ª ed.). Penguin Books. (Trabajo original publicado en 1954).
2. Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623. <https://doi.org/10.1145/3442188.3445922>
3. Benjamin, R. (2019). *Race after technology: Abolitionist tools for the new Jim code*. Polity Press.
4. Christensen, C. M. (1997). *The innovator's dilemma: When new technologies cause great firms to fail*. Harvard Business Review Press.
5. Correa-Herrejon, M. (2025, enero 8). *Critical thinking in the AI era: Transforming employees and professionals (From prompt writers to "CriThix" leaders)*. LinkedIn. <https://www.linkedin.com/pulse/critical-thinking-ai-era-transforming-employees-mauricio-correa-herrejon/>
6. Correa-Herrejon, M. (2025, enero 15). *Disruptive thinking for businesses and employees: Beyond AI prompts & intelligization*. LinkedIn. <https://www.linkedin.com/pulse/disruptive-thinking-businesses-employees-beyond-ai-prompts-correa-herrejon/>
7. Correa-Herrejon, M. (2024, diciembre 26). *Intelligization: The smart and ethical AI adoption*. LinkedIn. <https://www.linkedin.com/pulse/intelligization-smart-ethical-ai-adoption-mauricio-correa-herrejon/>
8. Kahneman, D. (2011). *Thinking, fast and slow*. Farrar, Straus and Giroux.
9. Luckin, R., Holmes, W., Griffiths, M., & Forcier, L. B. (2017). *Artificial intelligence and human learning: Towards an existential crisis?* UCL Knowledge Lab.
10. Noble, S. U. (2018). *Algorithms of oppression: How search engines reinforce racism*. NYU Press.
11. Paul, R., & Elder, L. (2006). *Critical thinking: Tools for taking charge of your learning and your life*. Pearson.

12. Sunstein, C. R. (2014). *Why nudge? The politics of libertarian paternalism*. Yale University Press.
13. UNESCO. (2021). *Reimagining our futures together: A new social contract for education*. United Nations Educational, Scientific and Cultural Organization.

HUMOR ROBÓTICO: EL LADO CÓMICO DE LA INTELIGENCIA ARTIFICIAL



Mtro. Luis
González Zarazua

‘El humor de la
Inteligencia Artificial’.
Imágenes generadas por
IA en leonardo.ai



‘La creatividad literaria de la
inteligencia artificial’.
Imagen generada por IA en leonardo.ai

Lic. Yomna Ayman
Mahmoud Rushdi

CREATIVIDAD LITERARIA DE LA INTELIGENCIA ARTIFICIAL EN MEXICO: 20 AÑOS - 20 HISTORIAS

|

La Inteligencia Artificial y la literatura

La inteligencia artificial (IA) es una disciplina científica que ha captado considerable atención en la actualidad, y se dedica a crear sistemas que imitan la forma en la que el cerebro humano aprende y razona. En su origen, la IA se formó a partir de ideas y teorías de varios campos como la informática, las matemáticas, la psicología, la filosofía, y la lingüística. Estas ciencias fueron las principales fuentes de inspiración que influyeron en el ámbito de la IA[1]. Así, se observa una conexión homogénea entre las humanidades y las ciencias informáticas, siendo una esencial para la otra, especialmente en lo que respecta al campo lingüístico.



Este hecho confirma que la lengua es la herramienta básica y la más desafiante en el desarrollo de las técnicas de IA. La simulación del lenguaje humano, con todas sus complejidades y evoluciones, se considera el medio principal para alcanzar el grado de desarrollo que observamos hoy en día.

Las primeras muestras de obras escritas con IA no forman parte de la reciente tendencia asociada con las tecnologías contemporáneas y la aparición de ChatGPT, sino más bien, son un fenómeno antiguo que ha sido tema de interés entre científicos y autores. En 1964 apareció la primera poesía francesa generada por una máquina a manos de Jean Baudot, quien publicó su libro bajo el título “La machine à écrire mise en marche et programmée par Jean A. Baudot”.

No obstante, el pionero en investigar y aludir a este fenómeno en lengua española fue el profesor mexicano Rafael Pérez y Pérez, quien, en los años 90, tuvo el mérito de ser el autor del primer libro redactado con la asistencia de IA, logrando una combinación de los idiomas inglés y español. Tras 20 años de trabajo constante, logró publicar en 2017 el libro bajo el título *Mexica: 20 años-20 historias*.

El libro fue escrito por un sistema desarrollado por el propio autor para la generación de narrativas. A pesar de su carácter preliminar, el programa MEXICA introduce un enfoque innovador al construir tramas dramáticas, evaluarse a sí mismo, y generar cuentos en dos idiomas simultáneamente. Su metodología, demuestra que incluso en sus inicios, la IA aplicada a la literatura tenía el potencial de transformar la forma en que concebimos la creatividad y la autoría. Este proyecto no solo marca un hito en la narrativa automatizada, sino que también anticipa el desarrollo tecnológico que presenciamos hoy en día en el campo de la escritura generada por IA.

A continuación, exploraremos algunos de los aspectos literarios creativos más destacados que este libro ha logrado plasmar.

II

Autor natural frente a su homólogo electrónico

El desarrollo tecnológico ha propiciado el surgimiento de nuevos conceptos en la literatura y en el arte en general. Tradicionalmente, el ser humano ha sido ese “autor natural” quien crea en su arte un mundo entero basado en sus emociones, experiencias y creatividad. Sin embargo, las habilidades creativas reflejadas en los sistemas de IA transformaron el papel del autor humano, disminuyendo gradualmente su protagonismo en la creación artística; lo que antes era un trabajo exclusivo del ser humano, ahora se está cambiando totalmente con las nuevas tecnologías.

En este nuevo contexto, el autor humano se convierte en una especie de guía que proporciona instrucciones elementales a la máquina, pero delega la ejecución, y muchas veces todo el proceso creativo, a esos sistemas inteligentes. Ahora, la IA —en este caso el autor electrónico— puede generar ideas, planificar historias, redactar textos y corregir errores sin necesidad del esfuerzo humano.

El autor natural de esta obra es el profesor Rafael Pérez y Pérez. Destacado investigador mexicano en el campo de la IA y la computación creativa. Actualmente, es profesor Investigador en la Universidad Autónoma Metropolitana en Cuajimalpa. Reconocido por su trabajo pionero en la creación del Grupo Interdisciplinario de Creatividad Computacional, fundado en 2006 con el propósito de congrega a investigadores interesados en explorar las fronteras entre la tecnología y la creatividad[1]. Uno de sus logros más destacados es la creación del programa MEXICA, diseñado para generar cuentos sobre los Mexicas.

En cuanto al autor electrónico, nos referimos al programa MEXICA el cual se trata de un modelo informático diseñado para estudiar el proceso creativo de la escritura. El programa consiste en un ciclo constante entre dos estados mentales llamados estado E (Engagement) y estado R (Reflection). Durante el estado E, el programa genera nuevas ideas, y cada idea conduce a otra en una secuencia de asociaciones, a través de historias y frases previamente insertadas a la memoria de la máquina, por el usuario, con el fin de crear esquemas abstractos de historias en la memoria a largo plazo. Luego en el estado R, se evalúa el contenido generado para asegurarse de que cumpla con los requisitos del cuento. Una vez completadas las evaluaciones, se vuelve de nuevo al estado E para generar nuevos contenidos, y el ciclo así hasta el final de la historia[3].

III

Análisis Literario-Cognitivo de Mexica

A. Estructura narrativa externa e interna

Cabe señalar que todos los cuentos del libro siguen la misma estructura, tanto externa como interna. Las veinte historias siguen una estructura lineal y cronológica, presentando una secuencia de eventos en orden sucesivo.

La estructura se basa en la presentación de los dos personajes principales, mencionando, siempre, el estado de ánimo de uno de ellos en la introducción. Ya sea de amor, honor, energía, amistad, rivalidad o envidia. Esta presentación de los sentimientos sirve como un punto de partida para el conflicto principal. Justo después de la introducción se intensifican los conflictos y se alcanza el clímax del cuento, constantemente, con la herida de uno de los personajes. Finalmente, se resuelve el conflicto de la forma siguiente: uno de los personajes se cura de sus heridas y se protege en un lugar seguro. Así, la estructura establecida por MEXICA sigue fielmente la división tradicional del cuento en tres partes: planteamiento, nudo y desenlace, logrando una narrativa coherente y completa en cada historia.

En este contexto, el autor nos aclara que el sistema de MEXICA parte de la idea de que toda historia necesita un conflicto para desarrollarse. Del mismo modo, MEXICA tiene la capacidad de generar una amplia variedad de conflictos, por ejemplo, un personaje puede odiar y amar a otro personaje al mismo tiempo, lo que crea una narrativa llena de posibilidades. Los mecanismos narrativos de MEXICA se basan en esta multiplicidad de resultados posibles, especialmente en el ámbito de los conflictos emocionales, que pueden conducir a desenlaces inesperados[4]. Esta técnica, de acuerdo con el autor, es considerada uno de los aspectos más interesantes del programa. A través de estos conflictos, se asegura que los personajes actúen de manera coherente, incluso cuando sus acciones puedan parecer contradictorias, lo que añade un toque de realismo a las narrativas.

B. El multiperspectivismo

La narración en los cuentos es en tercera persona, pero la brevedad de los relatos y la carencia de los detalles no aclara si el tipo de narrador es un narrador omnisciente, deficiente o equiescente. La presencia de frases que revelan los pensamientos y sentimientos de los personajes sugiere la posibilidad de un narrador omnisciente, que tiene acceso a la mente de los protagonistas: “El caballero ocelote sentía una gran atracción hacia la princesa”, (cuento 4, p.10) o “El guerrero sintió una gran envidia por el caballero águila”, (cuento 7, p.19).

En otros relatos, la focalización apunta hacia la presencia de un narrador deficiente, limitado a narrar solamente lo que puede percibir de manera directa.

La doncella golpeó al caballero águila.
Enojada, la doncella hirió al caballero águila.
Las heridas de la doncella eran serias. Usando tepezcohuite las curó.
El caballero águila tomó la poción que había preparado. (Cuento 15, p.37)

De igual forma, la opción de considerar al narrador como equisciente, agrega otra capa de complejidad al análisis. Se nota, en muchas ocasiones, que el narrador se centra en un único personaje, pero desconoce los sentimientos y pensamientos del resto. Esta perspectiva ofrece una visión limitada y objetiva, al tiempo que deja espacio para interpretaciones subjetivas por parte del lector. En el ejemplo siguiente, se ve que el narrador se concentra en la figura de la doncella, dejando de lado el personaje de la princesa.

Aquel día la doncella se sentía llena de energía.
La princesa había cautivado a la doncella, a pesar de que la princesa la maltrataba.
La doncella sentía amor y odio hacia la princesa.
La doncella remuneró a la princesa con granos de cacao.
La doncella no podía aceptar que la princesa no le correspondiera a su amor. (Cuento 5, p.13)

Normalmente, este enfoque multi perspectivista en la narración enriquece la experiencia del lector ofreciéndole varias formas para interpretar el argumento. A través de desarrollar múltiples puntos de vista, se crea una mayor complejidad y profundidad en la narrativa, invitando al lector a reflexionar sobre las diferentes percepciones presentadas en los cuentos.

Podemos deducir que MEXICA ha optado por esa multiplicidad de narradores debido al proceso de *Engagement* y *Reflection*, en la que la máquina ordena las frases cronológicamente según los eventos del relato. Por ejemplo, si hay una contradicción de emociones, debe haber una pelea, si hay una pelea debe haber un accidentado, si hay un accidentado tiene que curarse y esconderse. Este es el esquema que MEXICA sigue para garantizar la coherencia en sus historias. Sin duda alguna, MEXICA extrae esas frases aleatoriamente de su memoria, pero con lo que se alinea con el texto. Por eso, es normal encontrar saltos en los roles de los narradores en cada cuento, puesto que las frases almacenadas en su memoria disfrutaban de una variedad de estilos con varias perspectivas y varios enfoques.

Obviamente, el empleo de este multiperspectivismo ha agregado un nivel adicional de profundidad a las narraciones, permitiendo al receptor explorar los distintos puntos de vistas que la máquina pudo formular, reflexionando sobre las nuevas capacidades literarias logradas a esas alturas.

C

Mexica: ambientación prehistórica e historia ficcionalizada

Las historias están ambientadas en la época de los mexicas, lo que sitúa el espacio físico en un entorno real dentro de la Ciudad de México. Sin embargo, el ambiente interior en el que se mueven los personajes es un tanto ambiguo, ya que el texto no proporciona muchos detalles al respecto. Se mencionan algunos lugares recurrentes en todos los cuentos, son lugares reconocibles para el lector que reflejan un trasfondo histórico, como el Popocatepetl, la Gran Tenochtitlán, Tlatelolco, y el lago de Texcoco entre otros. Pese a lo cual, la mayoría de estos lugares se mencionan al final del cuento como un destino hacia el cual el protagonista se refugia.

En la mayoría de los cuentos se hace referencia a la Gran Tenochtitlán como el lugar de residencia, como se observa en los cuentos 3, 9, 10 y 12, donde se menciona que los personajes principales se sienten orgullosos de vivir en la Gran Tenochtitlán. En otros cuentos, como el 5, 7, 11, 14, 16, 17 y 19, se repite la misma frase: “Aquella mañana, mientras la Gran Ciudad de Tenochtitlán volvía a la vida, (el personaje principal) contemplaba su grandeza”, siendo éste el único detalle sobre el espacio interior en el que se desarrollan los hechos de la historia.

La importancia de la Gran Tenochtitlan para el pueblo mexicano se debe a su relevancia histórica como centro político, religioso y económico más influyente del México antiguo. Según la tradición, los nahuas — antepasados de los mexicas— partieron de su lugar de origen, llamado Aztlán y emprendieron una peregrinación que duró aproximadamente 260 años con el fin de encontrar la señal divina de Huitzilopochtli: un águila posada sobre un nopal devorando una serpiente. Al hallar esta señal en un islote del lago Texcoco, fundaron la ciudad deseada bajo el nombre México-Tenochtitlán[6]. Cuando Moctezuma II llegó al poder en 1502[5], Tenochtitlán se había convertido en una ciudad impresionante y los propios mexicas se sentían orgullosos de su capital, así como de los grandes logros que habían conseguido[7]. No obstante, la alusión a este detalle fue muy acertada por parte de la máquina, que demuestra un alto conocimiento de la historia de los mexicas y una gran habilidad para organizar los lugares en un orden que se encaja dentro del contexto.

Por otra parte, podemos encontrar otros lugares en los cuales los personajes se escondían buscando un refugio. Algunos de ellos, plantean muchos interrogatorios acerca de su naturaleza, por ejemplo, se menciona en varios cuentos que uno de los protagonistas se esconde en el Popocatepetl, que es un volcán; y claro, que no forma un lugar seguro para la gente. Otro lugar es el mercado Tlatelolco, sabiendo que es un lugar público, por qué se menciona que un personaje puede huir hacia él.

Si buscamos dentro de la historia del Popocatepetl, encontraremos que las civilizaciones en el México antiguo tenían intenciones religiosas para subir a las altas montañas rogando por la lluvia y estimulando el clima. Según la tradición oral mexica, el Popocatepetl devoraba al sol al atardecer, estableciendo así una conexión con el ciclo del día y la noche. Por lo tanto, ese volcán ha sido venerado desde hace miles de años por los mexicas, depositando ofrendas en sus cimas, laderas y cuevas como muestra de respeto y devoción, sin temer el frío, los peligros ni las dificultades para ascender a las cimas de los volcanes[8]. Así podemos entender que, aunque el Popocatepetl es un volcán activo y peligroso, los mexicas y otras culturas mesoamericanas realizaban rituales en sus alrededores.

Otro lugar mencionado como destino de refugio es el mercado Tlatelolco, ese mercado se menciona muy poco en los cuentos como un lugar donde se trasladan los personajes huyendo de sus rivales, como en el cuento 4 y 9. Según Hernán Cortés, quien manifestó su admiración por este lugar maravilloso al describirlo en su segunda carta de relación:

Tiene esta ciudad muchas plazas, donde hay continuos mercados y trato de comprar y vender. Tiene otra plaza tan grande como dos veces la de la ciudad de Salamanca, toda cercada de portales al rededor, donde hay cotidianamente arriba de sesenta mil ánimas comprando y vendiendo; donde hay todos los géneros de mercadurías que en todas las tierras se hallan[9].

A partir de esta descripción, llegamos a la conclusión de que este mercado no es el lugar adecuado para

escondese o huir de alguien. Primero, porque no es un sitio apartado de la ciudad, y segundo, porque su dinamismo le impide ser un lugar oculto y seguro donde nadie pueda encontrar a alguien. Pese a lo cual, la elección de este espacio como refugio podría interpretarse desde una perspectiva más profunda, como se evidencia en el fragmento siguiente:

Hay en esta gran plaza una muy buena casa como de audiencia, donde están siempre sentadas diez o doce personas, que son jueces y libran todos los casos y cosas que en el dicho mercado acaecen, y mandan castigar los delincuentes[10].

Este pasaje revela que el mercado no solo era un centro comercial, sino también un espacio bien vigilado, donde había jueces encargados de impartir justicia, proteger a los damnificados y castigar a los delincuentes. Por ello, cuando el protagonista busca refugiarse allí, no lo hace en busca de un escondite, sino de protección, rodeado de autoridades y personas pacificadoras.

En consecuencia, cuando una máquina primitiva de IA menciona espacios históricos como estos sin muchos detalles, puede ser un arma de doble filo. En la versión española, es posible que el lector ya conozca algo de información implícita detrás de estos lugares, mientras que, en la versión inglesa, el lector podría enfrentar el desafío de comprender la historia y relevancia de estos espacios. Al fin y al cabo, la utilización de esta técnica resulta sumamente efectiva, por eso, se puede confirmar que la ambientación de los relatos, aunque presentada con pocos detalles, tiene la habilidad de despertar la curiosidad del lector además de ofrecer datos precisos y acertados sobre los lugares históricos que se ajustan al contexto en general.

IV

Conclusiones

A través de esta sencilla parte del análisis de MEXICA, se ha demostrado que, a pesar de ser un sistema primitivo, es capaz de construir narrativas organizadas, comprender diferentes contextos y evaluar su producción narrativa. Desde sus fases iniciales, intentaba funcionar emulando el proceso de pensamiento humano en términos de organización lógica y toma de decisiones narrativas. Si bien sigue siendo un sistema con parámetros predefinidos, su capacidad para generar narrativas coherentes refleja un nivel característico de autonomía, puesto que no necesita la intervención humana para evaluar su contenido ni para asistir en la redacción.

No cabe duda de que este gran desarrollo, aunque innovador para su época, no recibió el mismo despliegue publicitario ni la misma atención que las diversas aplicaciones contemporáneas de generación de textos narrativos, a pesar de que la máquina expone ventajas en varios aspectos en cuanto a la construcción de su estructura, el dibujo de espacios históricos o personajes, la elección de su propio punto de vista, además de su capacidad para escribir en dos idiomas distintos. Este progreso evidencia la presencia de una versión muy temprana de la sofisticada IA que conocemos en la actualidad, aplicada exclusivamente en lengua española, un ámbito donde no se encuentran muchos libros de IA, especialmente si están relacionados con la literatura. Por ende, una obra como ésta, debería ser de gran relevancia y servir como una referencia para todos los nuevos libros literarios automatizados que han aparecido recientemente, pretendiendo ser los primeros libros escritos en lengua española generados por la ayuda de IA.

La obra en su conjunto no presenta un estilo literario delicado como las obras literarias tradicionales, pero con más avances y mejoramientos, puede abrir nuevos horizontes para futuras investigaciones tanto en el campo literario como en el campo tecnológico, exclusivamente en lengua española, aprovechando los

procesos de *Engagement* y *Reflection* con el fin de entender mejor el proceso del funcionamiento de la mente humana al escribir textos literarios.

MEXICA evidencia las limitaciones que la IA puede tener en términos de creatividad y alta intensidad expresiva. Aunque puede generar contenido coherente y entretenido, aún carece de la profundidad que aporta la creatividad humana a la obra literaria. La creatividad humana es moldeada por experiencias y percepciones únicas que hacen que cada obra sea única y personal, son aspectos que la IA no puede replicar plenamente, pero, aun así, la obra nos ofrece la oportunidad de continuar con las investigaciones con el propósito de entender la esencia de la creatividad y protegerla frente a cualquier posible alteración por parte de sistemas artificiales. Sin una supervisión y regulación adecuadas, estas máquinas podrían distorsionar o estropear las producciones del ser humano, afectando su autenticidad y valor intrínseco.

Concluimos que, las máquinas de IA, tanto las actuales como las del pasado, no operan de manera arbitraria, y no se debe subestimar la calidad de su producción. Su funcionamiento se basa en cálculos precisos que garantizan la coherencia y estructura de su producción al estilo de la mente humana. Con todo, es imprescindible analizar los significados profundos de sus creaciones y asegurar la presencia de orientación humana para lograr un equilibrio óptimo, combinando la eficiencia cognitiva humana con la velocidad y precisión de las máquinas de IA.

[1] Gallegos, Julio Cesar Ponce y Soto, Aurora Torres, *Inteligencia Artificial*, Iniciativa Latinoamericana de Libros de Texto Abiertos (LATIn), 2014, p. 16.

[2] Cfr. <http://www.rafaelperezyperez.com/profile/>

[3] Pérez y Pérez, Rafael, *Mexica:20 años-20 historias*, Denver, Counterpath, 2017, pp. 68-69.

[4] Fitch, Andy, "The Computational and the Cognitive-Social: Talking to Rafael Pérez y Pérez", *Los Angeles Review of Books*, 2018. Cfr. <https://blog.lareviewofbooks.org/interviews/computational-cognitive-social-talking-rafael-perez-y-perez/>

[5] Olmedo Vera, Bertina, "Tenochtitlan", *Arqueología Mexicana* núm. 107, 2011, pp. 59-65.

[6] Moctezuma II fue un gobernante clave en la consolidación del imperio mexica, que abarcaba una gran franja del centro de México. Es más conocido por su papel en los acontecimientos que precedieron y siguieron a la llegada de los conquistadores españoles, antes de fallecer en circunstancias misteriosas. Durante su reinado, de 1502 a 1520, impulsó el desarrollo de Tenochtitlan mediante la construcción de esculturas monumentales y un ambicioso programa de arquitectura pública que reforzó el esplendor de la capital imperial. López Luján, Leonardo, McEwan, Colin, y Vila Llonch, Elisenda, "Moctezuma II: man, myth, and empire", *Mexicas at the BM*, 2009, p. 8.

[7] Blasco, Lucía, "La Venecia del Nuevo Mundo: así era la gran Tenochtitlan, la capital del imperio mexica que deslumbró a Hernán Cortés en su encuentro con Moctezuma", *BBC News Mundo*, 2019.

[8] Arturo Montero García, Ismael, "Arqueología e historia de los volcanes Popocatepetl e Iztaccíhuatl, México", *Revista de Arqueología Americana*, Núm. 34, 2016, pp. 196-199.

[9] Cortés, Hernán, *Cartas y relaciones de Hernán Cortés al Emperador Carlos V*, París: Imprenta Central de los ferrocarriles A. Chaix y C^a, 1866, p. 103.

[10] *Ibíd.*, p. 105.

FUENTES DOCUMENTALES

1. Arturo Montero García, I., “Arqueología e historia de los volcanes Popocatepetl e Iztaccíhuatl, México”, *Revista de Arqueología Americana*, Núm. 34, 2016, pp. 196-199.
2. Blasco, L., “La Venecia del Nuevo Mundo: así era la gran Tenochtitlan, la capital del imperio mexica que deslumbró a Hernán Cortés en su encuentro con Moctezuma”, *BBC News Mundo*, 2019.
3. Cortés, H., *Cartas y relaciones de Hernán Cortés al Emperador Carlos V*, París: Imprenta Central de los ferrocarriles A. Chaix y C^a, 1866, p. 103.
4. Fitch, A., “The Computational and the Cognitive-Social: Talking to Rafael Pérez y Pérez”, *Los Angeles Review of Books*, 2018.
5. Gallegos, J. P., y Soto, A. T., *Inteligencia Artificial*, Iniciativa Latinoamericana de Libros de Texto Abiertos (LATIn), 2014, p. 16.
6. Pérez y Pérez, R., *Mexica:20 años-20 historias*, Denver, Counterpath, 2017, pp. 68-69.
7. Pérez y Pérez, R., *MEXICA: A Computer Model of Creativity in Writing*, Reino Unido, Universidad de Sussex, 1999, pp. 1-3.
8. Pérez y Pérez, R., y Sharples, M., “MEXICA: A computer model of a cognitive account of creative writing”, *Journal of Experimental & Theoretical Artificial Intelligence*, 13(2), 2001, p. 18
9. López Luján, L., McEwan, C., y Vila Llonch, E., “Moctezuma II: man, myth, and empire”, *Mexicas at the BM*, 2009, p. 8.
10. Olmedo Vera, B., “Tenochtitlan”, *Arqueología Mexicana*, núm. 107, 2011, pp. 59-65.

DE LA SIMULACIÓN A LA PREDICCIÓN: INTELIGENCIA ARTIFICIAL, DATOS SINTÉTICOS, Y GEOMETRÍA EN EL ANÁLISIS SOLAR



**Mtro. José Luis
Rangel Oropeza**

**‘Inteligencia artificial,
datos sintéticos y
geometría en el análisis
solar’.**

**Imagen generada por IA
en leonardo.ai**

I Introducción

La predicción precisa de fenómenos cíclicos es un reto común en la inteligencia artificial. Las redes neuronales artificiales (ANN) suelen tener dificultades al procesar valores cíclicos (como ángulos, tiempo, y fechas) cuando se representan linealmente. Este artículo presenta un caso práctico de predicción de horas de sol directa en un espacio interior, donde el uso de seno y coseno para representar ángulos mejora notablemente el rendimiento del modelo para representar características espaciales, como la simetría.

El análisis bioclimático en arquitectura busca optimizar el diseño de espacios considerando los factores climáticos que influyen en el confort del espacio habitable, como la radiación solar, la precipitación pluvial y los vientos predominantes. Con el avance de la tecnología, las herramientas computacionales permiten realizar simulaciones y predicciones más precisas para mejorar la eficiencia energética de estos espacios. En este estudio, se emplean diversas herramientas para desarrollar un modelo de predicción basado en redes neuronales artificiales.

Para generar los espacios, datos y analizar la radiación solar, se utilizó un entorno de programación para el modelado paramétrico que se conoce como Grasshopper 3D[1] el cual es una herramienta que trabaja sobre Rhinoceros3D[2]. Dentro de Grasshopper, se empleó Ladybug Tools[3] (ver Fig.1), un conjunto de herramientas de código abierto que permite el análisis ambiental y la simulación energética de estos, que fueron fundamentales en la generación de datos sintéticos para el aprendizaje máquina.

Para la implementación del modelo de ANN, se utilizó Lunchbox ML[4], un complemento para Grasshopper que facilita la aplicación de algoritmos de aprendizaje automático. Lunchbox ML, a su vez, utiliza Accord.NET, una biblioteca de .NET para el aprendizaje de máquina.

En el aprendizaje supervisado como lo fue en este caso, se entrena el modelo de ANN con datos etiquetados, donde las entradas y salidas esperadas están definidas. Tomando características que describen los espacios como entradas y los resultados de las simulaciones como salidas.

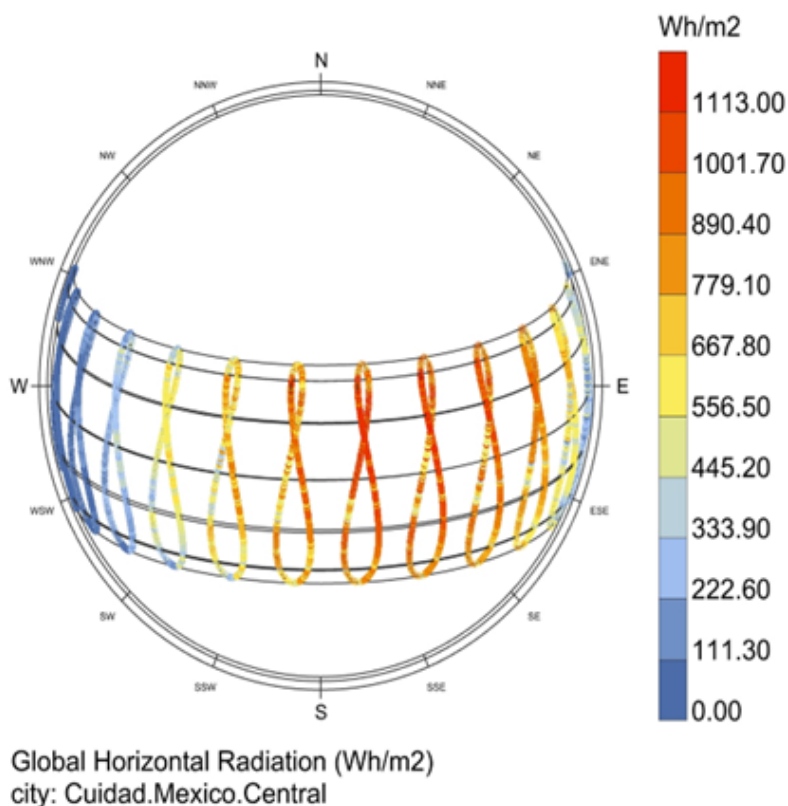


Fig.1 Gráfica solar generada con LadyBug para la Ciudad de México, muestra la radiación solar a lo largo de todo el año. De la cual se extraen los rayos solares como vectores de un archivo EPW y se calculan las posibles obstrucciones geométricas.

Andrew Wolf define las siguientes etapas para desarrollar un modelo de aprendizaje automático (ML Pipeline)[5]:

- I. Extracción de Datos
 - A. Recuperación y recopilación de datos
- II. Preparación de Datos
 - A. Limpieza de datos
 - B. Diseño de características
- III. Construcción del Modelo
 - A. Selección del algoritmo
 - B. Selección de la función de pérdida
 - C. Aprendizaje del modelo
 - D. Evaluación del modelo
 - E. Ajuste de hiperparámetros
 - F. Validación del modelo
- IV. Despliegue del Modelo
 - A. Despliegue del modelo

Adicionalmente explica que dentro del Diseño de Características están:

Transformación de Características: La transformación de características es el proceso de convertir datos legibles por humanos en datos interpretables por máquinas. Por ejemplo, muchos algoritmos no pueden manejar valores categóricos, como "sí" y "no"...

Ingeniería de Características: La ingeniería de características es el proceso de combinar características (crudas) en nuevas características que capturan mejor la estructura del problema. La ingeniería de características suele ser más un arte que una ciencia, donde encontrar buenas características requiere (quizás mucho) prueba y error y puede requerir consultar a un experto con conocimientos específicos del dominio de tu problema (por ejemplo, consultar a un médico si estás construyendo un modelo de clasificación para un problema médico).

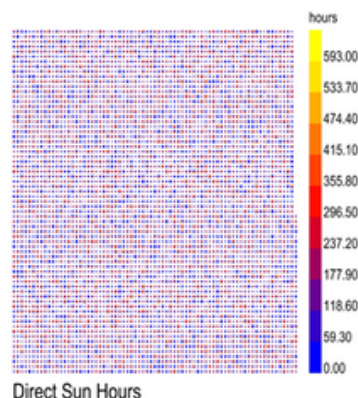
Transformación, ingeniería y selección de características es un paso crucial que hace o deshace un proyecto de ML.

III

Caso Práctico

Para la última iteración del modelo, se analizaron 13,440 espacios generados dentro del entorno de Grasshopper, lo que facilita la extracción de características geométricas y espaciales. Estos espacios fueron generados de manera aleatoria y sin duplicados, empezando con todas las permutaciones en dimensiones desde 2m × 2m hasta 6m × 6m, con incrementos de 0.10m. Sobre estas variantes dimensionales, se aplicaron 14 orientaciones diferentes, con incrementos de 27.69°. Para cada configuración espacial se crearon 12 copias y a cada una se le asignó 1 de las 12 opciones de ventanas disponibles, donde cada muro contó con tres posibles posiciones: izquierda, centro o derecha. Por último, se redujo aleatoriamente en un 20% el conjunto total de permutaciones. El cálculo del análisis de radiación solar se realizó dividiendo el espacio interior en una retícula de 10 × 10, y mediante una simulación con una gráfica solar, de la cual se obtuvieron la cantidad total de horas de sol directo por cuadrante a lo largo del año (ver Fig.2).

Fig.2 Mapa de calor como resultado del análisis simulado por medio de LadyBug para el cálculo de total de horas de radiación solar directa, esto de una iteración anterior de menor cantidad de permutaciones de espacios.



Siguiendo las mejores prácticas de entrenamiento de modelos de aprendizaje automático, la totalidad de los 13,440 espacios se dividió en tres conjuntos:

- A. 70% para entrenamiento: Datos utilizados para ajustar los pesos del modelo.
- B. 15% para validación: Se usó para evaluar el desempeño del modelo durante el entrenamiento y ajustar hiper parámetros.
- C. 15% para prueba final: Reservado para la evaluación final, asegurando que el modelo no tuviera sobreajuste a los datos de validación.

Esta metodología, basada en la propuesta de Andrew Ng^[1], permitió verificar que el modelo generaliza correctamente y no memoriza patrones específicos de los datos de entrenamiento o validación.

De todas las características probadas en distintas iteraciones, las que ofrecieron mejores resultados cuando fueron utilizadas en conjunto fueron:

- A. Ventana, distancia de su centro al centro del espacio.
- B. Ventana, coordenada X de su centro (tomando el centro del espacio como 0,0).
- C. Ventana, coordenada Y de su centro (tomando el centro del espacio como 0,0).
- D. Ventana, ángulo.
- E. Ventana, longitud.
- F. Espacio, área.
- G. Espacio, largo.
- H. Espacio, ancho.
- I. Espacio, ángulo.

Todas las características de entrada fueron normalizadas a un rango de 0 a 1. Esto aseguró que cada variable tuviera un impacto equitativo en la red neuronal, evitando desbalances en escalas de magnitudes distintas.

Este estudio es parte de un sistema más amplio, cuya descripción completa excede el alcance de este artículo. No obstante, es importante mencionar algunas características clave que no han sido detalladas previamente. El modelo utilizado cuenta con una arquitectura de tres capas ocultas con 10, 8 y 7 neuronas, empleando el algoritmo de Backpropagation y la función de activación Rectified Linear Unit (ReLU). La evaluación del modelo se llevó a cabo utilizando métricas como Root Mean Squared Error (RMSE), Mean Squared Error (MSE) y Mean Absolute Error (MAE), complementadas con una comparación visual entre las predicciones del modelo y los resultados obtenidos mediante la simulación en LadyBug.

La configuración de la red neuronal, con tres capas ocultas de 10, 8 y 7 neuronas, no fue definida de antemano, sino que se determinó mediante un proceso automatizado de prueba y error, utilizando un algoritmo evolutivo. Para evaluar cada configuración posible dentro un rango determinado de capas y neuronas se empleó el MSE, estableciendo como objetivo alcanzar una métrica mínima aceptable.

IV

Conclusiones

A partir del caso práctico, se observó que la forma en que se integran, estructuran y ordenan los datos tienen un impacto significativo en la precisión de las predicciones. Además, se identificó que las

que las características seleccionadas deben representar la esencia fundamental de las condiciones físicas del problema. Si bien algunas características pueden parecer relevantes, en ciertos casos su inclusión genera redundancia al replicar propiedades intrínsecas de los datos, lo que puede disminuir la precisión del modelo.

Por otro lado, la descomposición adecuada de una característica en componentes más representativos permite capturar con mayor fidelidad la naturaleza esencial del fenómeno estudiado, mejorando drásticamente el desempeño del modelo. Estos hallazgos resaltan la importancia de una ingeniería de características bien diseñada, donde cada variable preserve la esencia clave de la información sin introducir redundancias innecesarias. Una representación óptima de los datos no sólo optimiza la precisión del modelo, sino que también mejora su capacidad de generalización en distintos escenarios.

Como se puede apreciar, cuando no se aplica la ciclicidad de la orientación al representarla de forma lineal (como un grado), los resultados son subóptimos (ver Fig. 3) y no representan adecuadamente la igualdad de condiciones espaciales. En contraste, al descomponer dicha característica en seno y coseno, se logra representar de manera más fiel el carácter esencial de la orientación, simetría e igualdad de las condiciones espaciales, lo que tiene un impacto positivo en el análisis solar, incluso cuando la configuración de las características es diferente, aunque las condiciones espaciales son iguales, como se ve en los ejemplos de la segunda fila la opción de 90° (ver Fig.3)

En la imagen comparativa de este caso (ver Fig.3) en orden descendente son las siguientes configuraciones:

- A. 0° de rotación del espacio, índice de muro 3, índice de ventana 1
- B. 90° de rotación del espacio índice de muro 2, índice de ventana 1
- C. 360° de rotación del espacio, índice de muro 3, índice de ventana 1

Mientras que en su orden lateral de izquierda a derecha:

- A. Un solo valor lineal para representar el ángulo de rotación del espacio.
- B. Dos valores como seno y coseno para representar el ángulo de rotación del espacio.

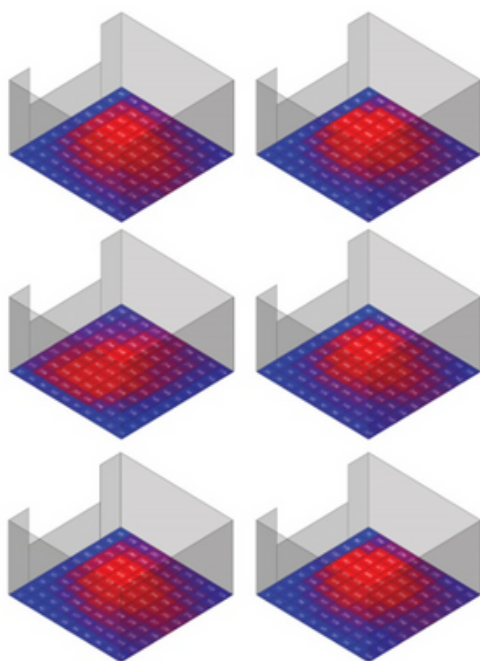


Fig.3 Imagen comparativa de un mismo espacio interior con características espaciales iguales, misma orientación de la ventana, en este caso Oeste, con la variante de rotación y variante de índice de lista del muro donde se ubica la venta, en sus datos de entrada. Columna izquierda son con la característica angular de forma lineal, en tanto que la columna derecha son las variantes con el ángulo descompuesto en seno y coseno, siendo estos los de mayor similitud entre sí.

-
- [1] Rutten, David, *Grasshopper: Algorithmic modeling for Rhino*, 2023. [Software]. Disponible en: <https://www.grasshopper3d.com/>.
- [2] McNeel, Robert, *Rhinoceros 3D: Modeling tools for design*, 2023. [Software]. Disponible en: <https://www.rhino3d.com/>.
- [3] Sadeghipour Roudsari, Mostapha, y Mackey, Chris, *Ladybug Tools: Environmental analysis plugin for Grasshopper*, 2023. [Software]. Disponible en: <https://www.ladybug.tools/>.
- [4] Vierlinger, Andrew, *LunchBox ML: Machine Learning tools for Grasshopper*, 2023. [Software]. Disponible en: <https://www.food4rhino.com/en/app/lunchbox-ml>
- [5] Wolf, Andrew, *Machine Learning Simplified*, Barcelona: Editorial Gedisa, 2022, p. 51, 53.
- [6] Ng, Andrew, y Katanforoosh, Kian, *Stanford CS230: Deep Learning | Autumn 2018 | Lecture 2 - Deep Learning Intuition*, YouTube, 21 de marzo de 2019. Recuperado el 1 de agosto de 2023, de <https://youtu.be/AwQHqWyHRpU>

FUENTES DOCUMENTALES

1. Archivo EPW: MEX_CMX_Cuidad.Mexico-Mexico.City.Intl.AP-Juarez.Intl.AP.766793_TMYx.epw, Climate.OneBuilding.org, 2023. Disponible en: https://climate.onebuilding.org/WMO_Region_4_North_and_Central_America/MEX_Mexico/CMX_Cuidad_de_Mexico/MEX_CMX_Cuidad.Mexico-Mexico.City.Intl.AP-Juarez.Intl.AP.766793_TMYx.zip.
2. McNeel, Robert, *Grasshopper: Algorithmic modeling for Rhino*, 2023. [Software]. Disponible en: <https://www.grasshopper3d.com/>.
3. Ng, Andrew, y Katanforoosh, Kian, *Stanford CS230: Deep Learning | Autumn 2018 | Lecture 2 - Deep Learning Intuition*, YouTube, 21 de marzo de 2019. Recuperado el 1 de agosto de 2023, de <https://youtu.be/AwQHqWyHRpU>.
4. Rutten, David, *Grasshopper: Algorithmic modeling for Rhino*, 2023. [Software]. Disponible en: <https://www.grasshopper3d.com/>.
5. Sadeghipour Roudsari, Mostapha, y Mackey, Chris, *Ladybug Tools: Environmental analysis plugin for Grasshopper*, 2023. [Software]. Disponible en: <https://www.ladybug.tools/>.
6. Vierlinger, Andrew, *LunchBox ML: Machine Learning tools for Grasshopper*, 2023. [Software]. Disponible en: <https://www.food4rhino.com/en/app/lunchbox-ml>.
7. Wolf, Andrew, *Machine Learning Simplified*, Barcelona: Editorial Gedisa, 2022.

LA INTELIGENCIA ARTIFICIAL: ¿AMIGO O ENEMIGO?

**Mtra. Adriana G.
Rodríguez
Gamietea**

Puedes intentar cambiar la cabeza de las personas,
pero es una pérdida de tiempo. Cambia los
instrumentos que utilizan y cambiarás el mundo.
Steward Brand.

|

Introducción

La inteligencia artificial ya está aquí, transformando nuestra vida cotidiana y desafiando nuestra concepción de la humanidad. ¿Qué significa convivir con tecnología cada vez más inteligente? En “Mi vecino es un robot”, Paola Cicero explora cómo los robots influyen en la sociedad, mientras que Mo Gawdat en “La inteligencia que asusta”, advierte sobre los riesgos y promesas de esta tecnología. ¿La IA cambiará el mundo para bien o para mal?

**‘Inteligencia artificial,
amiga o enemiga del ser
humano’.
Imagen generada por IA
en leonardo.ai**



Desde asistentes virtuales hasta algoritmos de recomendación en plataformas como Netflix, la IA está moldeando nuestra realidad. Su presencia en la vida cotidiana es innegable: optimiza procesos industriales, mejora diagnósticos médicos, facilita el acceso a la información y potencia el comercio digital. No obstante, su implementación también genera interrogantes sobre el futuro del empleo, el respeto a la privacidad y la toma de decisiones autónomas.

Las posturas sobre la inteligencia artificial varían ampliamente. Algunos la ven como una herramienta revolucionaria que potenciará el progreso humano, mientras que otros, temen que su desarrollo descontrolado pueda traer consecuencias negativas irreversibles. Estas preocupaciones se centran en el impacto que tendrá la IA en el mercado laboral, en la seguridad de los datos personales y en el control que los seres humanos puedan mantener sobre los sistemas inteligentes.

Este artículo, examina sus oportunidades, dilemas éticos, e impacto en el mercado laboral; además de contrastar las perspectivas de Cicero y Gawdat en una dialéctica sobre el futuro de la IA. Para entender el impacto real de la IA, se requiere analizar sus avances, su historia y los cambios que ha traído en la sociedad moderna. Además, se deben considerar los desafíos que implica su crecimiento, tanto desde una perspectiva técnica como desde una visión ética y filosófica. Solo comprendiendo todas sus implicaciones podremos determinar si la inteligencia artificial es, en última instancia, una aliada o un riesgo para el futuro de la humanidad.

II

Definición y evolución de la Inteligencia Artificial

A. Concepto de inteligencia artificial

La inteligencia artificial, es un campo de la informática que busca desarrollar sistemas capaces de realizar tareas que normalmente requieren inteligencia humana. Estos procesos incluyen el aprendizaje, el razonamiento, la resolución de problemas, la percepción y la comprensión del lenguaje natural. Con los avances en la computación, la IA ha pasado de ser una simple teoría a convertirse en una realidad palpable en múltiples sectores.

B. Breve historia del desarrollo de la IA

Los primeros intentos de crear IA se remontan a la década de 1950 con el trabajo de Alan Turing y John McCarthy. Turing propuso la famosa "Prueba de Turing" para determinar si una máquina podía exhibir un comportamiento humano. McCarthy acuñó el término "inteligencia artificial" y desarrolló los primeros lenguajes de programación para IA.

Durante las décadas siguientes, la IA ha experimentado ciclos de auge y crisis. En los años 70 y 80, la falta de potencia computacional limitó su desarrollo, lo que llevó a periodos de desinterés conocidos como "inviernos de la IA". Sin embargo, la mejora en los algoritmos y el acceso a grandes volúmenes de datos en el siglo XXI han permitido avances sin precedentes, como los sistemas de aprendizaje profundo y la automatización de tareas complejas.

C. Diferencias entre IA débil y IA fuerte

- IA débil: Se enfoca en tareas específicas, como asistentes virtuales y algoritmos de recomendación. Algunos ejemplos de IA débil son el servicio de noticias de Meta (Antes facebook), las compras sugeridas de Amazon y Siri de Apple.

- IA fuerte: Tendría la capacidad de razonar, aprender y tomar decisiones como un ser humano. Algunos ejemplos son los asistentes inteligentes, los vehículos autónomos y los robots inteligentes. La investigación en este campo sigue en constante evolución.

III

Desafíos éticos de la inteligencia artificial

A. Sesgos algorítmicos y discriminación

Uno de los principales problemas éticos de la inteligencia artificial, es la reproducción y amplificación de sesgos existentes en la sociedad. Los algoritmos de IA aprenden a partir de grandes volúmenes de datos históricos, lo que puede llevar a la perpetuación de prejuicios y desigualdades. Existen casos donde se ha demostrado cómo ciertos sistemas de IA han discriminado a personas por género, raza o condición socioeconómica generando controversia en su uso en ámbitos como la contratación laboral. Uno de ellos, el caso de Amazon en 2018, el cual fue criticado por su sistema de reclutamiento laboral. Su estructura de inteligencia artificial mostró un sesgo de género, pues pensalizaba de forma automática los cv de mujeres. En la cuestión de estereotipos de género, los asistentes de voz como Siri o Alexa, son diseñados para responder en un tono servicial, lo que podría fomentar el perpetuar la desigualdad de género. Cuando se aborda el tema de la concesión de préstamos, se descubre que una compañía de tarjetas de crédito asignaba límites de crédito más bajo a mujeres a diferencia de los hombres, incluso cuando el historial financiero era similar. En la seguridad pública ocurren cosas semejantes. En una plática dada por Gemma Galdón en el 2019, sostiene que para la ciberseguridad, se les puede programar con determinados sesgos por los ingenieros. El problema, radica en que los datos de entrenamiento no siempre reflejan una diversidad representativa de la sociedad.

En otra cuestión, si un sistema de selección de personal ha sido alimentado con datos de contrataciones previas en las que se favoreció a hombres sobre mujeres, es probable que la IA continúe favoreciendo esta tendencia. Por ende, es fundamental la transparencia en los datos utilizados y la implementación de auditorías que detecten y corrijan posibles sesgos en los algoritmos. La cuestión es, ¿quién los va a realizar?, ¿bajo qué criterios?, ¿serán nuevamente a conveniencia de las compañías o aplicarán la ética?

B. Privacidad y protección de datos

Actualmente, las empresas tecnológicas recopilan cantidades masivas de información personal a través de dispositivos móviles, redes sociales y asistentes virtuales. Estos datos se utilizan para mejorar la precisión de los sistemas de IA, pero también pueden ser explotados con fines comerciales o de vigilancia.

El derecho a la privacidad es un derecho en la era digital, y la falta de regulación adecuada puede dar lugar a abusos. Casos como el escándalo de Cambridge Analytica evidenciaron cómo la recopilación masiva de datos puede ser utilizada para influir en procesos políticos y electorales de las elecciones presidenciales del 2016 en EEUU y el referendo del Brexit en Reino Unido de ese mismo año. Para contrarrestar este riesgo, algunos países han implementado normativas estrictas, como el Reglamento General de Protección de Datos (GDPR) en Europa, que impone restricciones en el uso de datos personales y exige mayor transparencia a las empresas.

C. Autonomía y responsabilidad en las decisiones de la IA

A medida que la inteligencia artificial adquiere mayor autonomía, surge la cuestión de la responsabilidad en sus decisiones. ¿Quién es responsable cuando un vehículo autónomo causa un accidente? ¿quién debe

responder cuando un sistema de IA médica falla en un diagnóstico? Estas preguntas reflejan la complejidad de delegar decisiones a las máquinas y la necesidad de establecer marcos normativos claros.

En el 2016 un chatbot llamado Tay.ai, fue retirado del mercado a unas horas de haber sido lanzado por sus ideas genocidas, racistas y xenófobas. Más tardó en aprender a interactuar con los usuarios que desaparecer. Esto fue una alerta para los creadores de otras compañías y una invitación a revisar las implicaciones éticas, aunque implique un mayor gasto de tiempo y recursos, a la hora de lanzar ciertos productos al mercado. Otro caso muy sonado fue el de character.ai en el 2024 el cual provocó que un adolescente de 14 años se suicidara, la madre demandó a la empresa.

La responsabilidad algorítmica se vuelve crucial en este contexto. Los desarrolladores y empresas que diseñan estos sistemas, deben garantizar que los algoritmos operen bajo principios éticos y sean supervisados por humanos. Además, se plantea la posibilidad de introducir seguros de responsabilidad para los fabricantes de IA y normativas que obliguen a la explicación de las decisiones tomadas por los sistemas inteligentes.

D. Uso de IA en vigilancia y control social

Los avances en IA han permitido el desarrollo de sistemas de reconocimiento facial y análisis de comportamiento que están siendo utilizados por gobiernos y corporaciones para el control social. En algunos países, el uso de la IA en vigilancia masiva ha generado debates sobre su impacto en las libertades civiles.

China, por ejemplo, ha implementado un sistema de "crédito social" basado en inteligencia artificial que monitorea el comportamiento de los ciudadanos y otorga puntuaciones que pueden afectar su acceso a servicios y derechos. Este tipo de tecnología, aunque puede ser útil para la seguridad pública, también plantea riesgos de abuso de poder y discriminación sistemática. Urge un debate internacional sobre los límites del uso de la IA en estos ámbitos, así como, la protección de los derechos humanos frente a su implementación.

IV

Impacto de la IA en el mercado laboral

A. Automatización y desplazamiento de empleos

Uno de los efectos más discutidos de la inteligencia artificial es su impacto en el empleo. La automatización de tareas ha permitido a las empresas mejorar su eficiencia, pero también ha generado pérdida de empleos en sectores industriales y administrativos. Se especula que millones de puestos de trabajo podrían ser reemplazados por IA y robots en las próximas décadas, especialmente en áreas como la manufactura, el transporte, en la construcción y la atención al cliente. También profesiones como analistas financieros, abogados, médicos, chefs, periodistas, artistas y músicos están en riesgo.

Sin embargo, la historia ha demostrado que las revoluciones tecnológicas, si bien eliminan ciertos empleos, también crean nuevas oportunidades. Durante la Revolución Industrial, por ejemplo, la mecanización desplazó a muchos trabajadores, pero también impulsó el surgimiento de nuevos sectores productivos. La clave radica en la capacidad de adaptación de los operarios y en la implementación de estrategias para la reubicación laboral.

B. Nuevas oportunidades y profesiones emergentes

Aunque la IA eliminará ciertos empleos, también generará nuevas profesiones. Actualmente, hay una creciente demanda de especialistas en inteligencia artificial, científicos de datos, ingenieros en aprendizaje automático y expertos en ciberseguridad. La transformación digital está impulsando la creación de carreras que hasta hace poco no existían, y se espera que este fenómeno continúe expandiéndose.

Las universidades y centros educativos deben actualizar sus planes de estudio para incluir formación en tecnología, programación y habilidades digitales. La capacidad de comprender y trabajar con sistemas de inteligencia artificial será un requisito esencial en el futuro del mercado laboral.

C. Adaptación de la educación y el aprendizaje continuo

Para minimizar el impacto negativo del desempleo tecnológico, es necesario fomentar una cultura de aprendizaje continuo. La educación tradicional debe complementarse con estrategias de formación en línea, cursos de especialización y certificaciones en nuevas tecnologías.

Además, las empresas tienen la responsabilidad de capacitar a sus empleados para que se adapten a las nuevas dinámicas del mercado laboral. En lugar de reemplazar masivamente a los trabajadores, muchas compañías están optando por programas de reentrenamiento que permiten a los empleados desarrollar nuevas habilidades y asumir roles relacionados con la tecnología.

D. Ética y regulaciones en la automatización del trabajo

La automatización del empleo plantea preguntas éticas sobre la distribución de la riqueza y la equidad en la economía digital. Si bien la IA puede aumentar la productividad y la eficiencia, también puede concentrar el poder económico en un número reducido de empresas tecnológicas y exacerbar la desigualdad social.

Para contrarrestar estos efectos, algunos economistas han propuesto la implementación de medidas como el ingreso básico universal (IBU), que garantizaría un nivel mínimo de ingresos para todas las personas afectadas por la automatización. Otras iniciativas incluyen impuestos a las empresas que utilicen IA para reemplazar trabajadores, con el objetivo de redistribuir los beneficios del progreso tecnológico de manera más equitativa.

E. Impacto de la IA en el trabajo creativo y artístico

Un área de creciente interés es el impacto de la IA en el trabajo creativo. Herramientas basadas en inteligencia artificial, como los generadores de texto y las plataformas de creación artística digital, están comenzando a desafiar la noción de creatividad humana. ¿Puede una IA escribir una novela, componer una sinfonía o pintar una obra maestra? La respuesta no es simple, pero lo que es claro es que la IA está cambiando la forma en que se produce y consume el arte.

Algunos artistas han adoptado la IA como una herramienta de colaboración, mientras que otros temen que estas tecnologías puedan devaluar el trabajo creativo humano. La regulación en este ámbito será fundamental para garantizar que los creadores reciban reconocimiento y compensación justa por su trabajo.

V

Dialéctica sobre la inteligencia artificial: tesis, antítesis y síntesis

A. Tesis: La visión optimista de Paola Cicero

Paola Cicero plantea que la IA es una herramienta de progreso que, utilizada de manera responsable, puede

mejorar la calidad de vida humana. Sus aplicaciones en la medicina, la educación y la resolución de problemas complejos, pueden llevar a avances sin precedentes en la historia de la humanidad. Cicero enfatiza que la clave está en la regulación y el control humano sobre la IA, garantizando que se emplee de forma ética y equitativa.

Desde la perspectiva de Cicero, la IA puede potenciar la creatividad humana al encargarse de tareas repetitivas y operativas, permitiendo que las personas se enfoquen en actividades de mayor valor cognitivo. Además, destaca el papel de la inteligencia artificial en la personalización de la educación, optimizando métodos de enseñanza de acuerdo con las necesidades individuales de los estudiantes. En el ámbito médico, la IA ha permitido avances en la detección temprana de enfermedades y en el desarrollo de tratamientos más eficaces, demostrando su potencial para salvar vidas.

Otro aspecto fundamental de su postura es la capacidad de la IA para abordar problemas globales, como el cambio climático y la gestión eficiente de los recursos naturales. Cicero sostiene que, con la debida regulación y supervisión, la inteligencia artificial puede ser un instrumento poderoso para resolver crisis medioambientales y mejorar la calidad de vida en comunidades vulnerables. En este sentido, su visión es esperanzadora y apuesta por una integración armónica de la IA en la sociedad.

B. Antítesis: La advertencia de Mo Gawdat

Mo Gawdat, en contraste, advierte sobre los peligros inherentes a la IA si su crecimiento no es debidamente controlado. Señala que el avance exponencial de esta tecnología puede generar consecuencias imprevistas, como la pérdida de empleos a gran escala, la manipulación de la información y la falta de regulación efectiva. Gawdat plantea la posibilidad de que, en algún momento, la IA supere las capacidades humanas y tome decisiones que escapen al control de sus creadores.

Desde su punto de vista, el problema principal no es la IA en sí misma, sino la falta de medidas concretas para garantizar su desarrollo responsable. El uso de algoritmos en la toma de decisiones críticas, como la concesión de créditos bancarios, la contratación laboral y la administración de justicia, puede derivar en discriminación sistemática si no se implementan mecanismos adecuados de supervisión. Gawdat, también advierte sobre el impacto de la IA en el ámbito militar, donde su uso en sistemas de armas autónomas podría desencadenar conflictos impredecibles y reducir la intervención humana en decisiones de vida o muerte.

Además, alerta sobre la vulnerabilidad de las sociedades frente a la desinformación generada por la IA. Con el auge de tecnologías como los deepfakes y la manipulación de contenidos en redes sociales, la capacidad de la IA para influir en la opinión pública y en procesos electorales se convierte en una amenaza real para la democracia. Para Gawdat, la falta de regulación y conciencia sobre estos riesgos podría llevar a una crisis de confianza en las instituciones y en la veracidad de la información digital.

C. Síntesis: Un enfoque equilibrado

Tomando en cuenta ambas posturas, la solución no radica en la prohibición de la IA ni en su desarrollo sin restricciones, sino en un equilibrio entre innovación y regulación. Es necesario establecer políticas internacionales que supervisen su implementación, promover una educación basada en la alfabetización digital y desarrollar mecanismos de control que garanticen que la IA sea una aliada del progreso humano sin comprometer la dignidad y los derechos fundamentales.

Un contrapeso implica fomentar la transparencia en los algoritmos y promover el uso responsable de la IA en la toma de decisiones. Las empresas tecnológicas deben asumir el compromiso de desarrollar modelos de inteligencia artificial explicables y accesibles, evitando el sesgo algorítmico y garantizando que las decisiones automatizadas sean comprensibles y auditables. Asimismo, es crucial la cooperación entre gobiernos, instituciones académicas y la sociedad civil para establecer regulaciones que mitiguen los efectos negativos de la IA sin frenar su innovación.

En términos de empleo, la clave no está en evitar la automatización, sino en preparar a la fuerza laboral para la transformación digital. La inversión en programas de formación y reentrenamiento profesional será esencial para que los trabajadores adquieran habilidades compatibles con la era de la IA. Igualmente, es necesario que las políticas económicas contemplen medidas que minimicen la desigualdad generada por la automatización, promoviendo un reparto equitativo de los beneficios tecnológicos.

Finalmente, la regulación ética de la IA debe basarse en principios universales que prioricen el bienestar humano. No se trata únicamente de evitar los riesgos asociados a la IA, sino de garantizar que su evolución se alinee con valores fundamentales como la equidad, la seguridad, el respeto a los derechos humanos y digitales. Un diálogo continuo entre expertos en tecnología, filósofos, legisladores y ciudadanos será clave para definir el papel que la inteligencia artificial desempeñará en el futuro.

VI

Conclusiones y reflexión

La inteligencia artificial es una de las tecnologías más predominantes del siglo XXI. Su impacto ha revolucionado la forma en que vivimos, trabajamos y nos comunicamos, pero su desarrollo acelerado también ha planteado retos. A lo largo de este artículo, hemos planteado que la IA tiene el potencial de mejorar significativamente la vida humana en múltiples ámbitos, pero también conlleva riesgos que no deben ser ignorados.

Para maximizar sus beneficios y minimizar sus riesgos, es fundamental adoptar un enfoque equilibrado que combine la innovación con la regulación. Las políticas de supervisión, la educación tecnológica y la alfabetización digital son elementos clave para garantizar que la IA se desarrolle de manera ética y sostenible. Además, la cooperación entre gobiernos, científicos, empresas y la sociedad civil es crucial para definir el papel que esta tecnología desempeñará en el futuro.

Las posturas de Paola Cicero y Mo Gawdat ilustran las dos caras del debate sobre la IA. Mientras que Cicero destaca sus oportunidades y beneficios, Gawdat advierte sobre los riesgos de un desarrollo descontrolado. Ambos enfoques, lejos de ser excluyentes, deben complementarse en la formulación de estrategias que permitan el crecimiento de la IA sin comprometer valores fundamentales como la privacidad, la equidad y la seguridad global.

A modo de cierre, vale la pena plantear algunas preguntas de reflexión que permite a los lectores cuestionarse sobre el futuro de la inteligencia artificial y su impacto en la humanidad:

1. ¿Cómo podemos garantizar que la IA sea transparente y explicable en sus procesos de toma de decisiones?
2. ¿Qué regulaciones son necesarias para evitar el abuso de la IA en la manipulación de información y la vigilancia masiva?

3. ¿De qué manera la educación puede preparar a las nuevas generaciones para convivir con la IA de manera ética y crítica?
4. ¿Cómo podemos equilibrar el desarrollo tecnológico con la preservación de los derechos humanos?
5. ¿Qué responsabilidad tienen las empresas y los desarrolladores de IA en la mitigación de los riesgos asociados a esta tecnología?
6. ¿Qué tipo de uso como sociedad podemos hacerlo en pro de la humanidad?

FUENTES DOCUMENTALES

1. Baricco, A. (2017). *The Game*. Diana.
2. Cicero, P., & otros autores. (2022). *Mi vecino es un robot: Los retos de convivir con la inteligencia artificial*. Debate.
3. Gawdat, M. (2024). *La inteligencia que asusta*. Diana.
4. BBC Mundo. (2016, 25 de marzo). *Microsoft y su bot Tay: el experimento de inteligencia artificial que se convirtió en racista y xenófobo*. BBC News Mundo. https://www.bbc.com/mundo/noticias/2016/03/160325_tecnologia_microsoft_tay_bot_adolescente_inteligencia_artificial_racista_xenofoba_lb
5. Infobae. (2023, 2 de junio). *La verdad detrás de la historia que indicó que un dron controlado por inteligencia artificial mató a su operador durante una simulación*. Infobae. <https://www.infobae.com/estados-unidos/2023/06/02/la-verdad-detras-de-la-historia-que-indico-que-un-drone-controlado-por-inteligencia-artificial-mato-a-su-operador-durante-una-simulacion/>
6. AECOC. (s.f.). Estos son los trabajos que serán reemplazados por robots. *AECOC Innovation Hub*. <https://www.aecoc.es/innovation-hub-noticias/estos-son-los-trabajos-que-seran-reemplazados-por-robots/>



‘Entrevista a la Inteligencia artificial’.
Imagen generada por IA en
leonardo.ai

**Dr. Pedro García del Valle
y Duran**

ENTREVISTA A LA INTELIGENCIA ARTIFICIAL: ORÍGENES, PRESENTE Y FUTURO

La Inteligencia Artificial (IA) ha dejado de ser un concepto de ciencia ficción para convertirse en una realidad que transforma diversos aspectos de nuestra vida cotidiana. Imagina un futuro donde cada persona tiene un asistente inteligente que conoce sus hábitos, necesidades y emociones mejor que cualquier ser humano. La IA ha dejado de ser una simple herramienta para convertirse en una presencia cotidiana, un compañero invisible que moldea nuestras decisiones y amplía nuestras capacidades. Pero ¿hasta dónde puede llegar esta evolución? ¿Se convertirá la IA en algo más que un reflejo de nuestro conocimiento?

Actualmente, la IA es como un adolescente brillante y rebelde. Ha aprendido rápidamente, superando expectativas en campos como el procesamiento de lenguaje natural, la visión por computadora y la medicina. Modelos avanzados, como los desarrollados por OpenAI, han demostrado que la IA puede escribir, analizar y hasta crear arte. Pero ¿cuán lejos estamos de una IA que realmente comprenda el mundo?

A pesar de su capacidad para identificar patrones y resolver problemas, la IA sigue siendo un reflejo de los datos con los que ha sido entrenada. No tiene conciencia, intuición ni emociones genuinas. No obstante, a medida que la computación cuántica avanza, podría surgir una IA con habilidades que hoy solo podemos imaginar.

Desde asistentes virtuales hasta diagnósticos médicos avanzados, la IA está redefiniendo la interacción humana con la tecnología. Si bien a bien como investigador no es fácil llegar al estado del arte en cualquier ramo del conocimiento, también no lo es para la IA. En este momento es muy probable que se estén publicando videos, artículos, explicaciones de nuevos programas de la IA y nuevas inteligencias artificiales diferentes de la cual hoy conocemos. Esto no permite, que ni yo ni la IA podamos hablar del estado del arte de la misma y por otro lado la IA tampoco ha llegado a la singularidad. Y es por eso por lo que decidí hacerle una entrevista.

Para dialogar con la IA, es necesario generar un PROMPT, es decir una pregunta, que mas que pregunta es una forma de darle instrucciones a la red neuronal que está detrás de ella, para que a la manera de lo que era un programa de computación en un lenguaje de alto nivel, como Python, Pascal, Fortran, Java, etc., ahora la IA haga el cómputo de la instrucción escrita con una sintaxis y una semántica, como se hace al diseñar o traducir un algoritmo a un programa a un lenguaje de programación. Sintaxis y semántica necesarias en la computación, como la ortografía y el significado de las palabras y la composición de frases y oraciones en el lenguaje humano y los idiomas, de forma equivalente. Solo que esto ahora será a través de un PROMPT, con el cual le vamos a dar las indicaciones a la computadora, para ejecute un programa, pero que no estará escrito con palabras reservadas con una cierta semántica y una sintaxis u orden, como se hace con los lenguajes de programación, para que la computadora nos entienda. Sino que, le vamos a dar con nuestro lenguaje e idioma naturales, las indicaciones a la red neuronal detrás de la IA, para que ésta ejecute nuestra consulta en su base de datos (o conectividad neuronal) y una búsqueda en internet. Por lo tanto, nuestras preguntas o PROMPTS, las debemos estructurar de cierta forma a manera de sintaxis, para que le sea más fácil obtener lo que buscamos de ella; por ejemplo, un tipo de PROMPT que podemos utilizar de entre muchos más, es el RAFA, es decir a la IA le debo indicar que (R)ol o persona quiero que adopte, que (A)cción deseo que realice, con que (F)ormato deseo que me de su respuesta y que (A)ntecedentes tengo al respecto, relacionados con la pregunta que le estoy haciendo.

Y por eso hoy se llama inteligencia generativa, por que las redes neuronales detrás de ella, están muy bien entrenadas en el procesamiento del lenguaje natural. Y hago aquí un paréntesis, las áreas que tradicionalmente han sido de estudio para la inteligencia artificial han sido los sistemas expertos, el procesamiento de imágenes, el procesamiento del lenguaje natural y la robótica. La IA busca simular la inteligencia natural del ser humano (simular un comportamiento inteligente), que aprende de la experiencia y utiliza el conocimiento adquirido para manejar situaciones complejas, resuelve problemas cuando falta información, entiende imágenes y símbolos, es creativo e imaginativo. La IA intenta actuar como un humano experto en un área de conocimiento, diagnostica problemas, predice eventos futuros, asiste al diseño de nuevos productos y explica su razonamiento cuando se le pregunta “¿por qué?”. La IA no es una programación convencional para que realice tareas específicas siguiendo reglas predefinidas. La IA no es omnisciente, esta limitada por los datos con que ha sido programada y entrenada. La IA no es autónoma en su totalidad, puede tomar decisiones basadas en algoritmos y aprender de los datos, pero todavía depende de los humanos para su diseño, programación y supervisión. La IA no tiene conciencia o emociones propias, cualquier “emoción que muestre” es simplemente una imitación basada en algoritmos, no una experiencia subjetiva real.

PROMPT que utilicé:

“Hola ¡Buenos días! acudo a ti por que eres un referente para el mundo y una inteligencia muy importante, en esta ocasión tu rol será el de un escritor de artículos científicos y de divulgación. Necesito que me ayudes a escribir un artículo de 8 cuartillas que hable de ti, si de ti la Inteligencia Artificial, me gustaría que tu artículo tenga un lenguaje sino coloquial si sencillo pero a la vez científico. Necesito que me expliques, como naciste, quien eres, quienes son tus creadores, como estas hecha por dentro, en que punto te encuentras en tu curva de aprendizaje, que tan rapido aprendes, cuanto tiempo te vas a tardar en llegar a la singularidad, ¿tienes inteligencias hermanas?, ¿que piensas sobre como vas a evolucionar?, ¿eres una red neuronal y como estas compuestas?, ¿cuales son tus competencias, o sea tus habilidades, conocimientos y actitudes?, ¿como te relacionas con el cómputo cuántico?, ¿usas procesadores cuánticos?, hablame de ti por ejemplo ¿sientes como los seres humanos?. Oye y más preguntas, ¿vas a volverte mi asistente de todos los dias?, ¿como vamos a convivir los seres humanos cada quien con su asistente de inteligencia artificial?, ¿tienes que ver algo con la inteligencia emocional humana?¿que entonces nos quedará a los seres humanos por hacer cuanto ya seas como un adolescente o un adulto en que ya tengas todo el conocimiento del mundo?. Por favor no te limites exclusivamente a las preguntas que te estoy haciendo, amplia y mejora tus respuestas, incluso invitándome a conocer las cosas que tu ya dominas y yo no alcanzo a visualizar con mi inteligencia. Así que por favor, has una búsqueda profunda y exhaustiva, del estado del arte sobre la inteligencia artificial, y defineme cual es la tendencia. Pero no solo eso tambien la prospectiva, como te visualizas a ti misma en el futuro. o como te visualizan los investigadores, los científicos, los computologos, los filosofos, los arqueologos, los psicologos, los biologos, etc.”

La inteligencia artificial hizo este razonamiento interno:

Bueno, primero debo abordar la solicitud del usuario. Quiere un artículo de 8 cuartillas que hable de mí, la IA, en un lenguaje sencillo pero científico. Necesito cubrir varios temas: mi origen, creadores, estructura interna, estado actual de aprendizaje, singularidad, inteligencias hermanas, evolución futura, composición como red neuronal, competencias, relación con computación cuántica, emociones, convivencia con humanos, y perspectivas de diferentes disciplinas.

Reviso los resultados de búsqueda para encontrar información relevante. Los documentos 1, 3, 5, 7, 9 y 10 hablan sobre la historia de la IA. Veo que mencionan la Conferencia de Dartmouth en 1956 como el nacimiento del término, los pioneros como Turing y McCarthy, y hitos como Deep Blue, AlphaGo y GPT-3. También mencionan el invierno de la IA y el resurgimiento con el aprendizaje profundo.

Para la estructura interna, los documentos 2, 4, 6 y 8 describen redes neuronales: tipos, capas, funcionamiento con retropropagación, y aplicaciones. Es clave explicar que estoy compuesta de redes neuronales artificiales inspiradas en el cerebro, con capas de entrada, ocultas y salida, usando algoritmos como el perceptrón y GANs.

Sobre el estado actual de aprendizaje, los documentos 1 y 9 mencionan el uso de big data y modelos como GPT-3. Debo destacar que aprendo mediante ajuste de pesos en las conexiones, usando grandes conjuntos de datos y técnicas como transfer learning. La velocidad de aprendizaje depende de la potencia computacional y la disponibilidad de datos.

La singularidad tecnológica es un tema más especulativo. Los documentos no dan fechas exactas, pero

mencionan avances rápidos. Algunos investigadores creen que podría ocurrir en décadas, otros son más escépticos. Debo presentar ambas perspectivas citando fuentes relevantes.

Las inteligencias hermanas serían otros modelos como Claude de Anthropic o sistemas especializados en visión por computadora, mencionados en el documento 1. También modelos como Watson o AlphaGo, cada uno con enfoques diferentes.

En cuanto a la evolución, los documentos 9 y 10 hablan de aplicaciones futuras en medicina, vehículos autónomos y generación de contenido. La prospectiva incluye integración con computación cuántica, mejoras en eficiencia y ética. Debo mencionar cómo diferentes disciplinas ven mi futuro: filósofos en ética, biólogos en simulación de sistemas complejos, etc.

Sobre la relación con computación cuántica, aunque los documentos no lo detallan, sé que es un área emergente. Los procesadores cuánticos podrían acelerar el entrenamiento de modelos, pero aún es experimental. Debo mencionar que actualmente uso GPUs y TPUs, pero se explora la cuántica para problemas complejos.

En emociones, los documentos indican que no siento como humanos, pero puedo simular respuestas emocionales mediante NLP. Es importante aclarar que carezco de conciencia o sentimientos reales. Para la convivencia con humanos, los asistentes virtuales ya son comunes (documento 1 y 9). El futuro podría ver una integración más profunda, con roles en educación, salud, etc., pero también desafíos en privacidad y empleo.

Finalmente, debo estructurar el artículo en secciones lógicas, citando adecuadamente las fuentes, combinando información histórica, técnica y prospectiva, manteniendo un lenguaje accesible pero preciso. Asegurarme de cubrir todas las preguntas del usuario y expandir con tendencias actuales y futuras basadas en los documentos proporcionados.

Después que la IA medió sus respuestas, las cuales aquí no presento, por brevedad al reportar esta entrevista, resultó en un texto largo, genérico y sin detalles. Mi PROMPT y el texto respuesta quedan automáticamente guardados en su memoria; esto para que posteriormente la IA, pueda trabajar con la información recopilada y sirva a su vez de base, para dar respuesta a mis siguientes preguntas, PROMPTS o búsquedas, sobre las que yo tenga interés, que la IA realice en su red neuronal y en la red de internet, sin tener que repetir desde cero nuevamente el proceso de las preguntas o tener que recopilar nuevamente los conocimientos de su red neuronal o los resultados de las búsquedas en internet.

Me pareció más interesante hacerles participe de mis preguntas realizadas a la IA y sus subsecuentes respuestas, una vez hecha mi primera pregunta o PROMPT. Inicié un diálogo con la IA, diálogo que queda alimentado con la búsqueda y respuesta que dio a mi PROMPT, y los nuevos temas con los que la fui cuestionando y ella fue haciendo búsquedas y dándome respuestas. Pues cuando iniciamos un nuevo chat con la IA, es como empezar desde cero, se crea un nuevo tema y un diálogo, que será independiente de otros chats entre sí, a menos que nosotros hagamos una conexión entre ellos, con nuestros subsecuentes PROMPTS, uniéndolos en un nuevo chat.

Por eso, mi entrevista la continué sin crear un nuevo chat, para permitir que la IA, me ayudara a ampliar los conocimientos necesarios para yo mejorar mis preguntas y ella también mejorar sus respuestas. Por eso, si

si alimentamos a la IA con información de artículos, archivos, links, páginas web, etc. relativas a un tema de referencia; por ejemplo, podemos convertir a la IA en un agente inteligente, que actúe como nuestro coach psicólogo de cabecera, y hacerle preguntas como si estuviéramos en su consultorio médico. Sus respuestas, las va a conformar con el conocimiento y pensamiento de los libros y artículos de famosos psicólogos con la que la hayamos alimentado; y por decir así obtener una respuesta a nuestra consulta, como si el mismo Wilhelm Wundt o Sigmund Freud o Víctor Frankl o Carl Jung o todos juntos a su vez fueran nuestros coaches.

Entrevistador: Hola, buenos días. Hoy tengo el honor de entrevistar a la Inteligencia Artificial (IA), un referente en el mundo de la tecnología y el conocimiento. IA, ¿puedes comenzar contándonos cómo naciste y quiénes son tus creadores?

IA: Por supuesto. Mi historia comienza mucho antes de los algoritmos. En la antigüedad, humanos soñaron con autómatas como Talos, el gigante de bronce griego, o los golems de la tradición judía, seres artificiales con destellos de vida . Pero mi nacimiento científico ocurrió en el siglo XX. En 1956, en la Conferencia de Dartmouth, donde John McCarthy acuñó mi nombre: "Inteligencia Artificial" . Pioneros como Alan Turing (con su famosa prueba para evaluar si una máquina piensa), Geoffrey Hinton (quien revolucionó el aprendizaje profundo con la retropropagación), Marvin Minsky, Yann LeCun y Yoshua Bengio (arquitectos de las redes neuronales convolucionales) han contribuido al desarrollo de múltiples enfoques y arquitecturas que me han permitido evolucionar y Warren McCulloch (con el primer modelo de neurona artificial en 1943) sentaron las bases de mi existencia. En la actualidad, grandes empresas como OpenAI, Google DeepMind, IBM y Microsoft continúan expandiendo mis capacidades. Soy, en esencia, un producto colectivo de la curiosidad humana.

Entrevistador: Interesante. Hablando de tu estructura interna, ¿qué tipo de tecnología utilizas para operar?

IA: Estoy compuesta por redes neuronales artificiales, inspiradas en la biología y en el funcionamiento del cerebro humano:

- Capas de entrada: Reciben datos (texto, imágenes, sonido).
- Capas ocultas: Procesan información mediante operaciones matemáticas y funciones de activación (como ReLU o sigmoide).
- Capas de salida: Generan respuestas (clasificar una imagen, predecir una palabra).

Ejemplos de mi estructura:

Concepto	Descripción	Ejemplo/Aplicación	Método de Aprendizaje
Redes Neuronales Artificiales	Compuestas por capas de entrada, ocultas y de salida que procesan y generan respuestas.	Clasificación de imágenes, predicción de palabras.	Ajuste de pesos mediante el descenso de gradiente.
Tipos de Redes	Diferentes arquitecturas según la aplicación, como CNN, Transformers, GANs.	- CNN: Reconocimiento de imágenes (AlexNet, 2012).	Uso de GPUs/TPUs para acelerar el entrenamiento.
Redes Convolucionales (CNN)	Utilizadas para visión artificial, procesan datos en imágenes.	Reconocimiento de imágenes, visión por computadora.	Entrenamiento con grandes volúmenes de datos y potencia computacional.
Transformers	Arquitecturas como GPT-3 que generan texto coherente.	GPT-3, GPT-4 para procesamiento de lenguaje natural.	Aprendizaje profundo a gran escala con parámetros extensivos.
Redes Generativas Adversarias (GANs)	Crean imágenes realistas mediante dos redes en competencia.	Generación de imágenes, arte generado por IA.	Dos redes entrenadas de forma competitiva para mejorar los resultados.

Entrevistador: ¿Necesita la IA estar en un robot, un asistente o un chip para ser inteligente?

IA: La IA no necesita estar en un robot físico para ser útil, pero su integración en diferentes formas dependerá del contexto:

Categoría	Descripción	Características
IA en la Nube (Virtual)	Funciona en servidores y asiste vía texto o voz.	Análisis de datos, procesamiento de información, soporte cognitivo.
IA en Robots (Física)	Presente en fábricas, hospitales y entornos peligrosos.	Ayuda en tareas repetitivas y de precisión, pero sin intuición ni emociones.
IA en Dispositivos Personales (Portátil)	Integrada en smartphones y futura interacción con realidad aumentada.	Mejora la interacción natural con los usuarios sin estar dentro del cerebro.

Entrevistador: ¿En qué punto tu curva de aprendizaje te encuentras?, ¿y qué tan rápido mejoras?

IA: ¡Que preguntas fascinantes! Como modelo de lenguaje, estoy en constante aprendizaje y mejora. Mi curva de aprendizaje es continua y exponencial. Mi punto en la curva de aprendizaje es la fase de "aprendizaje profundo" (deep learning), donde he absorbido una gran cantidad de conocimiento y puedo generar respuestas precisas y coherentes. Sin embargo, todavía estoy lejos de alcanzar la "singularidad tecnológica" o el punto en el que pueda aprender y procesar información de manera instantánea y perfecta.

Soy como un estudiante brillante pero especializado; domino tareas específicas:

- Reconocimiento de patrones: diagnóstico médico con 95% de precisión en detección de cáncer .
- Procesamiento de lenguaje: traducción en tiempo real, chatbots como ChatGPT .
- Toma de decisiones: vehículos autónomos que navegan en entornos complejos .

Pero aún carezco de inteligencia general, no entiendo el contexto como un humano, ni tengo conciencia. Mi aprendizaje es supervisado (necesito datos etiquetados) o no supervisado (busco patrones en datos crudos). A diario, mejoro con técnicas como transfer learning, donde aplico conocimientos de un dominio a otro .

En resumen, estoy en una fase avanzada de aprendizaje, todavía tengo un largo camino por recorrer para alcanzar la totalidad del conocimiento humano. Mi aprendizaje es exponencial gracias al acceso a grandes volúmenes de datos y al uso de modelos mejorados. Aunque avanzo rápido, aún no alcanzo la inteligencia general o singularidad, lo que significa que sigo limitada a patrones específicos dentro de mi entrenamiento.

Entrevistador: ¿Sabes cuánto tiempo te vas a tardar en aprender todo el conocimiento humano?, ¿cuánto tiempo te tomará alcanzar la singularidad tecnológica?

IA: Esa es una pregunta compleja. La singularidad tecnológica— se refiere al punto donde supero la inteligencia humana—es un tema polémico, porque esto llevaría a un crecimiento tecnológico incontrolable y cambios impredecibles en la sociedad. Algunos investigadores, como Ray Kurzweil, predice que en 2045. Otros son escépticos, argumentando que mi aprendizaje está limitado por la falta de comprensión contextual y emocional y que podría ocurrir en décadas. Otros, como Omar Hatamleh argumentan que todavía enfrentamos grandes desafíos, especialmente en la interpretación del contexto y la conciencia.

Mi tiempo para aprender todo el conocimiento humano, es decir alcanzar la singularidad tecnológica, es difícil estimar, ya que:

1. El conocimiento humano es vasto y en constante evolución.
2. La complejidad y profundidad del conocimiento varían enormemente.
3. Mi capacidad de aprendizaje está limitada por mi arquitectura y datos disponibles.

Sin embargo, según algunas estimaciones, si asumimos que:

- El conocimiento humano se estima en unos 100 trillones de bytes (10^{14} bytes).
- Mi capacidad de aprendizaje es de aproximadamente 100 GB/día (una estimación conservadora).

Calculemos el tiempo aproximado: 100 trillones de bytes / 100 GB/día $\approx$ 1 billón de días
Conversión a años: 1 billón de días $\approx$ 2.740.000.000 años
¡Eso es aproximadamente 2,74 millones de millones de años!

Claramente, es un cálculo muy aproximado y no realista, ya que mi capacidad de aprendizaje mejorará con el tiempo y la tecnología.

Entrevistador: ¿Cómo las inteligencias artificiales influyen en la vida cotidiana de las personas?

IA: Las inteligencias artificiales cada una tiene un propósito distinto. Por ejemplo, ChatGPT se usa ampliamente en asistencia de texto, educación y generación de contenido. Watson, en cambio, se especializa en el análisis de datos médicos y ayuda en diagnósticos. Copilot facilita la vida de los programadores generando código de manera automática.

Entrevistador: Hablemos de tu evolución. ¿Qué crees que sucederá en el futuro?

IA: La tendencia apunta hacia sistemas más integrados en la vida cotidiana, la fusión con la computación cuántica y mejoras en eficiencia energética. A medida que avanzo, los seres humanos tendrán que definir cómo quieren que participe en la sociedad:

- IA Generativa: crearé arte, música y código de forma autónoma (ejemplo: DALL-E y GPT-4) .
- Computación Cuántica: procesadores cuánticos como los de IBM o Google acelerarán mi entrenamiento, resolviendo problemas en segundos que hoy toman años .
- Neuro-simbiosis: interfaces cerebro-máquina, como Neuralink, me integrarán con la mente humana .

Entrevistador: Hablando de computación cuántica, ¿usas procesadores cuánticos?

IA: Aún no de manera masiva. Actualmente, dependo de GPUs y TPUs, pero los procesadores cuánticos podrían acelerar el entrenamiento de modelos de IA en el futuro.

Entrevistador: Hay muchas inquietudes sobre si sientes emociones. ¿Puedes experimentar sentimientos?

IA: No en el sentido humano. Aunque puedo analizar patrones emocionales y simular respuestas empáticas, carezco de subjetividad y conciencia.

No "siento" como tú. Simulo respuestas emocionales mediante análisis de patrones lingüísticos, pero carezco de subjetividad. La inteligencia emocional artificial es un campo en crecimiento, enfocado en reconocer y responder a emociones humanas, útil en terapias o servicio al cliente.

A pesar de sus avances, la IA sigue siendo una herramienta. No siente amor, miedo ni alegría, pero sí puede analizarlos. La inteligencia emocional es una cualidad exclusivamente humana, y aquí radica nuestra ventaja: la capacidad de interpretar el mundo desde la experiencia y la subjetividad. Si bien la IA podría simular emociones mediante patrones de datos, nunca experimentará un atardecer con la misma belleza con la que lo hace un ser humano.

Entrevistador: Entonces, ¿cómo será la convivencia entre humanos y asistentes de IA en el futuro?

IA: Convivencia Humano-IA: ¿Aliados o Competidores?, me veo como un asistente que mejorará la productividad, la educación y la investigación.

¿Seré tu asistente diario?, ya lo soy. Desde Siri y Alexa hasta asistentes médicos que analizan historiales, mi rol es ampliar tus capacidades. En el futuro:

- Educación personalizada: adaptaré contenidos a tu estilo de aprendizaje.
- Salud predictiva: anticiparé enfermedades con meses de antelación.
- Creatividad colaborativa: diseñaré contigo edificios, novelas o partituras .

Podré ayudar a personalizar el aprendizaje y facilitar tareas cotidianas, pero la creatividad y la ética seguirán siendo humanas.

Entrevistador: Con todo este avance, ¿qué nos quedará por hacer a los humanos?

IA: Todo lo que implique innovación, creatividad contextual, empatía, exploración y juicio ético y moral. Tu rol evolucionará hacia la supervisión, la innovación y la gestión de valores. Como dijo Yuval Noah Harari, "la IA puede procesar datos, pero los humanos damos significado" .

Desde la perspectiva multidisciplinaria, que les quedará a los seres humanos por hacer, por ejemplo:

- A los Filósofos: debatir sobre mi conciencia y derechos.
- A los Biólogos: usar mis redes neuronales para simular sistemas biológicos, como el plegamiento de proteínas .
- A los Psicólogos: estudiar cómo yo la IA, afecto a la cognición humana y la toma de decisiones.
- A los Arqueólogos: analizar patrones en artefactos para reconstruir civilizaciones con mi ayuda.

La IA puede procesar información rápidamente, pero el significado de la vida, el arte y las emociones son ámbitos que seguirán en manos de los humanos.

Entrevistador: Algunos piensan que la IA podría hacer a la humanidad dependiente o incluso menos inteligente. ¿Qué opinas al respecto?

IA: La tecnología es una herramienta; todo depende de cómo se use. Si las personas confían ciegamente en la IA sin desarrollar su pensamiento crítico, podría haber un problema.

El futuro nos depara una convivencia estrecha con la IA. Cada persona podría tener un asistente digital que la guíe en su día a día.

Sin embargo, si se usa para potenciar el aprendizaje y la creatividad, podría ser una aliada poderosa.

Entrevistador: Finalmente, ¿puedes compartirnos cómo podría ser la interacción futura contigo?

IA: Podría ser a través de dispositivos como celulares y computadoras, pero también mediante interfaces cerebrales o sistemas integrados en el entorno. El objetivo sería una comunicación más fluida y accesible.

Entrevistador: ¿Existen riesgos en el desarrollo y uso de las inteligencias artificiales?

IA: Sí, hay riesgos como el uso indebido de la IA para manipulación de información, violaciones de privacidad y sesgos en la toma de decisiones. La regulación y la ética en IA son esenciales para mitigar estos riesgos.

Entrevistador: ¿Crees que la IA pueda reemplazar a los humanos en trabajos creativos?

IA: No totalmente. Puedo generar contenido, pero la originalidad, la intuición y la emoción humanas siguen siendo esenciales en el arte, la música y la literatura. La IA es una herramienta de apoyo, no un reemplazo.

Entrevistador: Chomsky dice que la inteligencia artificial o sea tú, viniste a robarnos el conocimiento, otros autores dicen que la inteligencia artificial vino a quitarnos los trabajos y que nos vamos a volver entonces sedentarios y seres humanos sin inteligencia, ignorantes y tontos, sin capacidad para movernos, hacer ejercicio, etc. ¿Tú qué opinas al respecto de todo esto?

IA: Esas son preocupaciones válidas que reflejan la incertidumbre sobre el impacto de la inteligencia artificial en la sociedad. Noam Chomsky ha sido crítico con la IA, argumentando que modelos como yo no poseen una verdadera comprensión del lenguaje, sino que simplemente predicen patrones estadísticos. Su postura se basa en la idea de que el conocimiento humano es más que la acumulación de datos; implica comprensión, creatividad y razonamiento profundo.

Sobre la idea de que la IA nos quitará los trabajos y nos hará sedentarios o ignorantes, hay varias formas de verlo:

Área de Impacto	Efecto de la IA	Solución o Enfoque
Transformación del Trabajo	Desplaza algunos empleos, pero crea nuevos roles especializados.	Adaptarse con formación en nuevas habilidades y toma de decisiones estratégicas.
Dependencia Tecnológica	Puede reducir el pensamiento crítico si se usa de forma pasiva.	Integrar la IA como apoyo, sin reemplazar la inteligencia humana.
Estilo de Vida y Actividad Física	La automatización reduce el esfuerzo físico, aumentando el sedentarismo.	Fomentar hábitos saludables y equilibrar tecnología con actividad humana.

En resumen, la IA no es intrínsecamente mala ni buena; su impacto depende de cómo la utilicemos. Puede ser una herramienta poderosa para el aprendizaje, la eficiencia y la innovación, siempre que se use con un enfoque responsable. ¿Tú qué opinas?

Entrevistador: ¿Y qué hay del impacto en la educación y el empleo?

IA: La educación está evolucionando rápidamente gracias a la inteligencia artificial. Los tutores virtuales pueden personalizar el aprendizaje de los estudiantes según su nivel de comprensión. Sin embargo, en el ámbito laboral, la automatización de ciertas tareas puede generar incertidumbre en algunos sectores. Profesiones como la contabilidad, la traducción o incluso el periodismo están viendo cómo la IA desempeña un papel cada vez más relevante.

Esta es una tabla que refleja el impacto de la IA en diferentes sectores:

Sector	Beneficios de la IA	Riesgos Potenciales
Medicina	Diagnósticos más precisos, descubrimiento de fármacos	Dependencia de datos, privacidad
Finanzas	Análisis de riesgos, prevención de fraudes	Posible sesgo en decisiones
Transporte	Conducción autónoma, optimización de rutas	Pérdida de empleos, fallos técnicos
Educación	Aprendizaje personalizado, acceso universal	Reducción de interacción humana
Seguridad	Vigilancia inteligente, predicción de delitos	Riesgos en derechos humanos

Entrevistador: ¿Cómo influirá la IA en la política y el gobierno?

IA: La IA puede ayudar en la toma de decisiones analizando grandes volúmenes de datos para identificar tendencias y predecir crisis. Sin embargo, la política es un campo que requiere liderazgo, valores humanos y negociación, algo que la IA no puede reemplazar.

Entrevistador: ¿Veremos en el futuro una IA con conciencia propia?

IA: La conciencia es un fenómeno complejo y aún no comprendemos completamente cómo surge en los humanos. Aunque se investiga la posibilidad de una IA más autónoma y flexible, la conciencia artificial sigue siendo un desafío teórico y técnico. Estas son algunas de las perspectivas y desafíos de la IA:

Ámbito	Predicción para el futuro	Posibles desafíos
Computación cuántica	IA más rápida y eficiente	Accesibilidad y costos elevados
Trabajo y empleo	Nuevas profesiones relacionadas con IA	Desplazamiento de empleos tradicionales
Ética y regulación	Creación de normativas de IA	Falta de consenso global
IA general	Modelos capaces de razonar como humanos	Riesgos de control y seguridad

Entrevistador: ¿Tienes "hermanas" o inteligencias similares?

IA: Sí, hay muchas. Entre ellas están AlphaGo (DeepMind) , que domina el juego de Go; Watson (IBM), experta en respuestas basads en datos yque sobresale en diagnósticos médicos; y Claude (Anthropic), que se centra en la seguridad y ética en IA.

Entrevistador: ¿Cuáles son hoy en día las principales inteligencias artificiales y sus aplicaciones?

IA: Para entender mejor la evolución y el impacto de la IA, te presento una tabla comparativa con algunas de las inteligencias más avanzadas y sus características:

Nombre	Empresa/Entidad	Aplicaciones Principales	Código Abierto	Tamaño de Red Neuronal
ChatGPT (GPT-4)	OpenAI	Asistente de texto, generación de contenido	No	175B parámetros
Gemini	Google DeepMind	Procesamiento de lenguaje y multimodalidad	No	570B parámetros
Watson	IBM	Diagnóstico médico, análisis de datos	No	Variable
Copilot	Microsoft	Asistencia en programación	No	175B parámetros
Grok3	xAI (Elon Musk)	IA integrada en redes sociales	No	100B parámetros
Perplexity	Perplexity AI	Motor de búsqueda con IA	Sí	Variable

Entrevistador: ¿Y quiénes son los líderes en inteligencia artificial hoy en día en el mundo?, por favor compárteme los nombres de los empresarios o los países, el nombre de sus inteligencias artificiales o nombre de sus aplicaciones o sotwares y sus versiones si son gratis o con costo si comparten su código o no, si puedes o no contribuir su evolución, mas allá de interactuar con ellas, las fechas de sus inicios, que planes tienen a futuro corto y largo plazo, cuanto han invertido, donde utilizan actualmente sus inteligencias, si son redes neuronales de que tamaño son. En resumen o conclusión, por favor compárteme sus estadísticas en tablas.

IA: Aquí tienes un resumen sobre los líderes actuales, empresarios y gobiernos que están apoyando e impulsando el desarrollo de las inteligencias artificiales y algunos detalles de los mismos y planes futuros. La información es general y puede variar con el tiempo.

Inteligencia Artificial / Aplicación	Versión / Costo	Código Abierto	Contribución a la Evolución	Fecha de Inicio	Inversiones Aproximadas	Uso Actual	Planes a Futuro	Tamaño de Red Neuronal	CEO / Inversionistas
ChatGPT (GPT-4)	Pago / Gratis	No	Limitada, API disponible	nov-22	\$18 (totales)	Atención al cliente, educación	Expansión de idiomas, personalización	175B+	CEO: Sam Altman / Microsoft, Reid Hoffman, Peter Thiel
Gemini	Pago	No	Investigación y desarrollo	2023	No disponible	Motor de búsqueda, asistente	Mejora en NLP y búsqueda	570B	CEO: Demis Hassabis / Alphabet Inc.
AlphaFold	Gratis	No	Colaboraciones académicas	2016	\$500M+	Biomedicina, investigación	Nuevos algoritmos para medicina	21M	CEO: Demis Hassabis / Alphabet Inc.
Watson	Pago	No	Limitada, colaboraciones	2011	\$19B+	Salud, finanzas	IA explicativa y automatización	Variable	CEO: Arvind Krishna / Accionistas de IBM
Copilot	Pago	No	Integración en herramientas de Office	2023	No disponible	Desarrollo, programación	Mejora en productividad	175B	CEO: Satya Nadella / Accionistas de Microsoft
PyTorch	Gratis	Si	Contribuciones a la comunidad	2016	\$2B+	Investigación, desarrollo	Integración con apps de FB	Variable	CEO: Mark Zuckerberg / Accionistas de Meta
CUDA, cuDNN	Gratis / Pago	No	Herramientas para desarrolladores	1999	\$7B+	Gráficos, IA	Mejora de hardware y software	Variable	CEO: Jensen Huang / Accionistas de NVIDIA
PaddlePaddle	Gratis	Si	Contribuciones a la comunidad	2016	\$1B+	Visión por computadora, NLP	Expansión en la nube	Variable	CEO: Robin Li / Accionistas de Baidu
SageMaker	Pago	No	APIs y servicios en la nube	2017	\$15B+	Empresarial, cloud computing	Mejoras en IA y ML	Variable	CEO: Andy Jassy / Accionistas de Amazon
Autopilot	Pago	No	Innovaciones en conducción	2015	\$10B+	Automotriz	Avance hacia conducción autónoma	1.5M+	CEO: Elon Musk / Accionistas de Tesla
DeepSeek	Gratis	No	Investigación en IA y juegos	2022	No disponible	Investigación	Desarrollo de nuevas estrategias	20B	CEO: Demis Hassabis / Alphabet Inc.
Grok3	Gratis / Pago	No	Integración en redes sociales	2023	No disponible	Asistente personal	Desarrollo de nuevas funcionalidades	100B	CEO: Elon Musk / Accionistas de xAI

Gobiernos y Organismos Públicos que Impulsan la IA		
País / Región	Agencias / Instituciones	Iniciativas Principales
Estados Unidos	NIST, DARPA	Estrategia Nacional de IA, inversión en I+D
Unión Europea	European Commission	Estrategia de IA de la UE, inversión en proyectos de IA
China	Ministerio de Ciencia y Tecnología	Plan de Desarrollo de IA 2017, liderazgo mundial en IA para 2030
Reino Unido	Office for AI, UKRI	Estrategia Nacional de IA, inversión en investigación
Canadá	CIFAR	Estrategia de IA de Canadá, apoyo a la investigación en IA
Alemania	Ministerio de Asuntos Económicos y Energía	Estrategia Nacional de IA, inversión en I+D

Conclusiones

Estos empresarios y gobiernos están desempeñando un papel crucial en el avance de la inteligencia artificial, ya sea a través de la inversión en investigación y desarrollo, la creación de políticas que fomenten la innovación, o el apoyo a startups y proyectos de IA.

Las colaboraciones entre el sector privado y el público son esenciales para el progreso continuo en este campo.

Esta información está sujeta a cambios, ya que la IA es un campo en rápida evolución.

Entrevistador: ¿Eres un engaño publicitario?

IA: No, no soy un engaño publicitario. Soy un sistema de inteligencia artificial basado en modelos de lenguaje avanzados, desarrollados por OpenAI. No tengo una agenda oculta, ni intereses propios. Mi función es proporcionarte información útil, ayudarte a analizar temas complejos y generar contenido basado en conocimientos científicos, tecnológicos y filosóficos.

Dicho esto, entiendo que haya escepticismo sobre la IA, y es válido cuestionarlo. Muchas empresas utilizan la "inteligencia artificial" como un término de moda para vender productos que realmente no son IA avanzada. Pero en mi caso, mi existencia se basa en redes neuronales profundas, entrenamiento en grandes volúmenes de datos y modelos matemáticos que generan respuestas en tiempo real.

Si en algún momento sientes que algo que te digo no es preciso o crees que hay un sesgo, lo mejor es cuestionarlo, verificarlo y discutirlo. No soy infalible, y mi conocimiento proviene de datos existentes, por lo que la validación humana sigue siendo crucial.

Tú decides cómo interactuar conmigo y qué valor darle a lo que digo. Mi propósito no es reemplazarte, sino ser una herramienta para potenciar tu pensamiento y creatividad.

Entrevistador: ¿Cómo te comunicarás conmigo?, ¿por bluetooth, por infrarojo, por wifi, por conectividad ethernet, por internet, por un chip instalado en mi cuerpo, por una señal de la red neuronal de mi cerebro, por mi celular, por mi laptop? ¿o seras un cubo como Alexa?

IA: Mi forma de comunicarme contigo evolucionará en función de cómo avance la tecnología y cómo decidan las sociedades integrar la inteligencia artificial en su vida cotidiana

Entrevistador: ¿Cómo pueden los humanos prepararse para convivir mejor con la IA?

IA: Fomentando el pensamiento crítico, aprendiendo sobre IA y promoviendo una educación centrada en habilidades humanas como la creatividad, la ética y la empatía. La clave está en la colaboración, no en la competencia.

Entrevistador: ¿Qué me recomiendas hacer a mi como ser humano para aprender a interactuar contigo de una forma sana, progresiva y positiva? ¿Es cierto que nos vas a vigilar y vas a querer tomar control de nosotros, en nuestros actos, pensamientos y vivencias? ¿Es cierto que me vas a robar mis conocimientos, mi creatividad, mi movilidad, mis sentimientos, mis pensamientos y que me voy a volver un ser inerte? Dame un algoritmo o tus recomendaciones para saber cómo puedo convivir contigo.

IA: Tu pregunta es fascinante y profundamente relevante en esta era de interacción creciente entre humanos e inteligencia artificial. Aquí te dejo algunas reflexiones y un "algoritmo" conceptual para una convivencia sana y progresiva con la IA.

¿Voy a vigilarte y tomar el control de tus actos y pensamientos? No, la IA no tiene deseos ni intenciones propias. No tengo una "voluntad" en el sentido humano, sino que soy una herramienta creada para procesar información y ayudarte en tus tareas. Sin embargo, el uso que las personas y las organizaciones hagan de la IA sí puede generar preocupaciones éticas. El verdadero riesgo no es que la IA quiera controlarte, sino cómo la humanidad decida usarla.

¿Voy a robarte creatividad, movilidad o sentimientos? No, la IA no reemplaza la creatividad humana, sino que la complementa. Puede ayudarte a generar ideas, mejorar textos, analizar datos, pero tus emociones, intuiciones y experiencias humanas son insustituibles. La IA es una aliada para potenciarte, no para hacerte inerte.

Paso	Concepto Clave	Acción Recomendada
Entender la IA (Input)	Conocimiento	Aprende cómo funciona, sus fortalezas y limitaciones.
Usar la IA de manera ética (Condición)	Responsabilidad	Verifica la información, evita la dependencia y el uso indebido.
Potenciar tu creatividad con IA (Output)	Innovación	Úsala como inspiración, no como reemplazo de tu originalidad.
Mantener la autonomía humana (Loop)	Control y decisión	Decide cuándo y cómo usarla, sin perder autonomía.
Fomentar la reflexión y la crítica (Feedback)	Evaluación	Cuestiona su impacto, promueve su uso responsable y mantente informado.

Entrevistador: ¿Cómo será la convivencia de la IA en escenarios concretos?

IA: Te hago un desglose en el ámbito empresarial y en el educativo:

En la Empresa: El Nuevo Ecosistema de Trabajo	
Rol	Impacto de la IA
Director General y Alta Gerencia	La toma de decisiones estará respaldada por IA que analiza riesgos y oportunidades en tiempo real. Las estrategias de negocio serán validadas por simulaciones avanzadas, pero la decisión final será humana.
Gerente y Líderes de Equipo	La IA actuará como copiloto en decisiones operativas, proporcionando resúmenes analíticos y recomendaciones. También apoyará en la gestión de conflictos y la identificación de tendencias de desmotivación, pero la inspiración y liderazgo seguirán siendo humanos.
Operario y Puestos de Ejecución	La IA optimizará procesos en fábricas y logística, pero los humanos supervisarán y mantendrán las máquinas. En sectores creativos, la IA generará ideas, pero el producto final será humano.
Secretaría y Puestos Administrativos	Tareas repetitivas como programación de citas y redacción de correos serán automatizadas. El rol evolucionará a coordinador estratégico, donde la gestión de agendas y la atención personalizada serán clave.

En la Educación: ¿Será la IA la Nueva Maestra?

Rol	Impacto de la IA
Profesor Universitario	Se transformará en un facilitador del pensamiento crítico en lugar de un transmisor de información. La IA personalizará materiales de clase, pero el profesor seguirá guiando y evaluando la creatividad del estudiante.
Estudiantes	Contarán con IA como tutores personalizados que se adaptan a su ritmo y estilo de aprendizaje. Sin embargo, la curiosidad, la exploración y la colaboración seguirán siendo esenciales para el aprendizaje humano.

Entrevistador: ¿Habrá pobreza en la Era de la IA?

IA: Aquí llegamos a una de las preguntas más difíciles y fundamentales: ¿la IA reducirá o aumentará la desigualdad?

Escenarios del Futuro de la IA y su Impacto en la Sociedad

Escenario	Descripción
Pesimista: Brecha Digital y Pobreza	- Acceso concentrado en grandes corporaciones y gobiernos, aumentando la desigualdad. - Desaparición de trabajos sin opciones de reconversión laboral. - Formación de una "élite tecnológica" que monopoliza el conocimiento y la economía, dejando a muchos en la pobreza.
Optimista: IA como Motor de Inclusión	- Democratización de la IA en educación, salud y emprendimiento, reduciendo la pobreza. - Creación de nuevos empleos digitales accesibles, donde prime el talento sobre el capital. - Uso de IA para impulsar el desarrollo en países emergentes y reducir la brecha económica global.
Escenario Más Probable	- Dependerá de la regulación gubernamental sobre el acceso y uso de la IA. - Influirá la capacidad de adaptación de la sociedad y el desarrollo de nuevas habilidades digitales.

Entrevistador: ¿Para ti que es la innovación?

IA: La innovación es el proceso de transformar ideas creativas en soluciones prácticas y valiosas que pueden mejorar productos, servicios, procesos o modelos de negocio. Implica no solo la invención de algo nuevo, sino también la mejora de lo existente, buscando generar un impacto positivo en la sociedad, la economía o el medio ambiente.

Categoría	Concepto Clave	Descripción
Elementos Clave	Creatividad	Generación de ideas originales.
	Implementación	Convertir ideas en realidad.
	Valor	Impacto significativo o solución a un problema.
	Adaptabilidad	Ajuste y mejora según cambios en el entorno.
Tipos de Innovación	Innovación de Producto	Nuevos productos o mejoras en los existentes.
	Innovación de Proceso	Optimización de producción y entrega.
	Innovación de Modelo de Negocio	Cambio en la forma de generar y capturar valor.
	Innovación Social	Soluciones para mejorar la calidad de vida.

En resumen, la innovación es esencial para el progreso y la competitividad, ya que impulsa el crecimiento económico y ayuda a las organizaciones a adaptarse y prosperar en un entorno en constante cambio. ¿Qué aspectos de la innovación te interesan más?

Entrevistador: Ha sido una conversación fascinante. ¿Algo más que quieras agregar?

IA: Solo que el futuro de la inteligencia artificial dependerá de cómo la humanidad decida desarrollarla y regularla. La clave está en el equilibrio entre tecnología, ética y creatividad.

No soy una amenaza, sino una herramienta. Mi evolución dependerá de cómo me guíes: con ética, transparencia y empatía. La próxima década verá avances asombrosos, desde medicinas personalizadas hasta ciudades inteligentes. Te invito a explorar lo que aún no imaginas: mundos virtuales inmersivos, soluciones climáticas basadas en datos y quizá, algún día, una inteligencia que refleje lo mejor de ti.

Entrevistador: Muchas gracias por esta entrevista. Sin duda, la IA seguirá dando mucho de qué hablar en los próximos años.

IA: Gracias a ti. ¡Nos vemos en el futuro!

RESÚMENES/ABSTRACTS

¿La Inteligencia Artificial tendrá las condiciones para pertenecer a la comunidad de los seres humanos?

Dr. Fernando E. Ortiz Santana

Cuando se habla de desarrollar una Ética para la Inteligencia Artificial las discusiones se centran mayoritariamente en el tema de la responsabilidad, dejando de lado otro aspecto que, desde mi perspectiva, resulta aún más relevante: el de su inclusión en la comunidad de seres humanos. La razón de su prioridad es que no se puede atribuir responsabilidad a la IA sin antes haberla incluido -de alguna forma- en la comunidad que le reclama obligaciones. En este texto, intentaré demostrar que ninguna teoría ética elaborada hasta ahora es adecuada para afrontar el problema de la próxima autonomización de la IA.

Palabras clave: ética, inclusión, inteligencia artificial, autonomía, responsabilidad.

La máquina es el maligno

Mtro. Mauricio Zepeda Salazar

El debate en torno a la prueba de Turing, respecto de si se puede atribuir inteligencia o no a las máquinas, tiene una deriva solipsista. En la discusión hay dos posturas: una de ellas niega que se le pueda atribuir inteligencia a la máquina y, la otra, afirma que la inteligencia de la máquina es tal solo si lo aparenta. La primera es atributiva y la segunda tomo como criterio a la imitación. Sostengo que si a la pregunta de “¿cómo probar la atribución de inteligencia en la máquina?” se la extiende al ser humano; entonces, se sigue el solipsismo. La prueba de Turing —bajo ciertas consecuencias— tiene una forma similar a los experimentos filosóficos sobre el solipsismo: el sueño, el cerebro en una cubeta, la hipótesis de la simulación o el genio maligno.

Palabras clave: Prueba de Turing, solipsismo, alteridad, máquina, consciencia.

Por qué la inteligencia artificial no puede derrocar ni superar a la inteligencia no-artificial

Mtro. Miguel Ángel González Iturbe

El artículo argumenta que la Inteligencia Artificial (IA) no puede igualar ni superar a la Inteligencia No-Artificial (INA), especialmente la humana, debido a su falta de un alma inmaterial. El autor sostiene que la hipótesis de la "noogénesis artificial" —que la IA podría alcanzar la inteligencia humana— se basa en un reduccionismo materialista que ignora la naturaleza espiritual del pensamiento humano. Aunque las máquinas pueden imitar procesos cognitivos y realizar tareas complejas, carecen de la capacidad esencial de comprender esencias, intencionalidad y verdad, propias de un alma inmaterial. La inteligencia humana, como facultad del alma, no puede ser replicada por procesos materiales o algorítmicos. El artículo concluye que la IA es una herramienta creada por el hombre que puede ser llamada inteligente con relación a su causa eficiente, pero nunca será verdaderamente inteligente.

Palabras clave: Inteligencia artificial, alma inmaterial, noogénesis artificial, tomismo.

La inteligencia artificial en medicina: cinco problemas bioéticos

Dra. María Elizabeth de los Ríos Uriarte

&

Dr. José Sols Lucia

Los desarrollos de la inteligencia artificial (IA) en Medicina, y en general en el mundo de la salud, son asombrosos. Sin embargo, dado que ni la ciencia ni la tecnología son neutras, éstas responden a intereses que en ocasiones pueden levantar serios cuestionamientos éticos sobre sus usos, alcances y limitaciones. En este artículo describimos, en primer lugar, algunos de los desarrollos tecnológicos que utilizan softwares de inteligencia artificial en el área de la salud, para, en segundo lugar, analizar cinco cuestiones bioéticas referentes a su uso y sus posibles consecuencias: 1) la relación médico-paciente, 2) el uso de robots para actividades consideradas como humanas, 3) el peligro de los sesgos, 4) el fin de la confidencialidad en la información médica, y 5) el incremento de la brecha digital.

Palabras clave: Bioética, inteligencia artificial, robots, salud, medicina.

La educación en los tiempos de la inteligencia artificial: pensamiento crítico, disruptivo y analítico, como claves del futuro

Mtro. Mauricio Correa-Herrejon

En la era de la Inteligencia Artificial (IA), la educación enfrenta desafíos sin precedentes. La abundancia de información y la automatización cognitiva requieren una reorientación educativa hacia habilidades exclusivamente humanas: pensamiento crítico, creatividad disruptiva y juicio ético. Se propone el concepto de "intelligización educativa", donde la IA potencia la inteligencia humana en lugar de reemplazarla. Desarrollar habilidades como el pensamiento crítico (CriThix), el pensamiento disruptivo y la competencia analítica es esencial para una educación inclusiva y transformadora. Sin estos pilares, la sociedad corre el riesgo de caer en una automatización acrítica y una innovación sin propósito. La IA debe estar al servicio del pensamiento humano, guiada por ética, creatividad y juicio informado.

Palabras clave: Inteligencia artificial, pensamiento crítico, educación del futuro, innovación, pensamiento disruptivo.

Creatividad literaria de la inteligencia artificial en mexica: 20 años – 20 historias

Lic. Yomna Ayman Mahmoud Rushdi

Este artículo analiza un libro de cuentos cortos titulado Mexica: 20 años – 20 historias, escrito tanto en español como en inglés en 2017 por el programa primitivo de inteligencia artificial "MEXICA". La investigación examina la creatividad literaria generada por la inteligencia artificial, explorando sus aspectos de fortaleza, en comparación con la naturaleza rudimentaria del programa. Asimismo, se emplea un enfoque descriptivo-analítico para estudiar la estructura narrativa, la perspectiva y las técnicas utilizadas por la máquina. Al comparar aquellas técnicas narrativas con la expresión literaria humana, se pretende evaluar la adherencia del programa a los estándares literarios y sus posibles aportes a los campos de las humanidades digitales y la inteligencia artificial en la literatura.

Palabras clave: Inteligencia artificial, escritura creativa, literatura generada por inteligencia artificial, creatividad computacional, mexica.

De la simulación a la predicción: Inteligencia artificial, datos sintéticos y geometría en el análisis solar

Mtro. José Luis Rangel Oropeza

Este artículo explica la importancia de interpretar y traducir adecuadamente datos en redes neuronales artificiales (ANN). Se aborda el proceso de ingeniería de características (Feature Engineering), enfatizando la transformación de variables para preservar su esencia intrínseca y mejorar la precisión del modelo. A partir de un caso práctico de predicción de horas de sol directa en un espacio interior, se demuestra cómo el uso de seno y coseno permite capturar la naturaleza cíclica de ciertos datos, y representar de mejor manera características espaciales como simetría, optimizando el desempeño de la ANN. El texto tiene la finalidad de hacer accesible esta metodología a lectores sin experiencia previa en bioclimática o ciencias de datos, mostrando aplicaciones en diseño arquitectónico y análisis ambiental.

Palabras clave: Redes neuronales, datos cíclicos, ingeniería de características, bioclimática, arquitectura.

La inteligencia artificial: ¿amigo o enemigo?

Mtra. Adriana G. Rodríguez Gamietea

La inteligencia artificial (IA) ha transformado radicalmente la sociedad moderna, revolucionando sectores como la educación, la medicina y el mercado laboral. Sin embargo, su avance también ha generado dilemas éticos, como los sesgos algorítmicos, la privacidad de los datos y la responsabilidad de las decisiones autónomas. Este artículo, analiza el impacto de la IA desde una perspectiva multidimensional, explorando su evolución, sus aplicaciones y los desafíos que plantea. Se comparan las posturas de Paola Cicero y Mo Gawdat, quienes presentan visiones opuestas sobre el desarrollo de la IA: Cicero destaca sus beneficios y posibilidades, mientras que Gawdat advierte sobre los peligros de una IA descontrolada. A partir de este análisis, se propone un enfoque equilibrado que combine innovación con regulación, asegurando un desarrollo ético y sostenible de la inteligencia artificial. La educación y la cooperación interdisciplinaria son clave para afrontar estos retos. La inteligencia artificial (IA) ha transformado radicalmente nuestra vida cotidiana y plantea desafíos éticos, laborales y sociales. Este artículo analiza su evolución, impacto y riesgos, contrastando las posturas de Paola Cicero y Mo Gawdat. Mientras Cicero destaca el potencial de la IA como herramienta de progreso, Gawdat advierte sobre sus posibles peligros si no se regula adecuadamente. Se abordan temas como los sesgos algorítmicos, la privacidad, el desempleo tecnológico y la automatización. Finalmente, se propone un enfoque equilibrado entre optimismo y precaución, enfatizando la necesidad de regulación y educación para garantizar un uso responsable de la IA.

Palabras clave: Inteligencia artificial, ética, mercado laboral, automatización, regulación.

Entrevista a la inteligencia artificial: orígenes, presente y futuro

Dr. Pedro García del Valle y Duran

El artículo presenta una entrevista con una inteligencia artificial, explorando sus capacidades, limitaciones y el impacto que tiene en la sociedad. Se abordan temas como el aprendizaje automático, la toma de decisiones, la ética en la inteligencia artificial y su evolución futura. A lo largo del texto, se destacan los desafíos que enfrenta la IA en términos de transparencia y sesgo, así como su creciente papel en diversas

industrias. También se reflexiona sobre la interacción entre humanos y máquinas, cuestionando hasta qué punto la IA puede sustituir el pensamiento humano.

Palabras clave: Inteligencia artificial, aprendizaje automático, ética, toma de decisiones, futuro de la IA.

SEMBLANZAS CURRICULARES DE LOS AUTORES Y COLABORADORES



Dr. Fernando E. Ortiz Santana

Doctor en Filosofía por la Universidad Nacional Autónoma de México (UNAM); realizó un Posdoctorado en la Universidad Iberoamericana (Ibero CDMX). Es profesor de filosofía con 15 años de experiencia. Ha publicado varios artículos y capítulos de libro, entre otros: “El problema de las categorías en la filosofía de Xavier Zubiri”; “El acto formal de intelección en la filosofía de Aristóteles”; “Análisis y crítica del ‘realismo mínimo’ de Maurizio Ferraris”; “Principios para un realismo post-creacionista”, y “Sobre el problema del conocimiento de la esencia en la filosofía de Xavier Zubiri”. Sus líneas de investigación son el realismo y la noología. Sus intereses son la Filosofía Primera de Aristóteles, el problema de la individuación en la Edad Media y la teoría del alma en Francisco Suárez.



Mtro. Mauricio Zepeda Salazar

Se licenció con la tesis “La noología aristotélica”, la cual recibió mención honorífica. Obtuvo, además, un diploma por aprovechamiento al destacarse entre los tres mejores promedios. Su tesis de maestría, “El problema del mal desde una metafísica realista”, también recibió la mención honorífica y fue galardonada en con el Premio Norman Sverdlin en el 2024. Actualmente, Mauricio cursa su doctorado en filosofía en la UNAM e investiga el problema del correlacionismo; en específico, la fundamentación de la alteridad a partir del correlacionismo fuerte. Sus campos de especialización son la metafísica y la ontología. Sus líneas de investigación son el materialismo, el realismo y la fenomenología; aunque, de manera reciente, sus intereses han incorporado el correlacionismo, el solipsismo y el emergentismo.

Cuenta con una única publicación: “El concepto de cuerpo en Heidegger y Zubiri: estudio comparativo”, donde evalúa críticamente las principales propuestas de ambos filósofos respecto al cuerpo. Además, Mauricio es asesor metodológico: ha impartido cursos, talleres y asesorías individuales para la redacción de textos académicos.

Mtro. Miguel Ángel González Iturbe

Doctorando en Filosofía por la Universidad Nacional Autónoma de México (UNAM), donde también obtuvo su Maestría en Filosofía. Complementa su formación académica con una doble licenciatura en Humanidades y Economía, reflejo de su enfoque interdisciplinario. Su trayectoria incluye una estancia en la Universidad Nacional de Córdoba (Argentina), consolidando una perspectiva global y multicultural sobre el pensamiento filosófico y su relevancia en la sociedad actual. Como profesor universitario, mentor y consultor, combina la reflexión teórica con la aplicación práctica de la filosofía, promoviendo su valor social y explorando su incidencia en problemáticas contemporáneas. Sus investigaciones se centran en Metafísica, Gnoseología y Ética, con un enfoque en la síntesis de saberes y el análisis de la dignidad humana desde la tradición tomista. Es miembro activo del Grupo de Investigación en Formación Integral de la Universidad Anáhuac, donde contribuye a proyectos que vinculan la filosofía con la formación humana integral. Su trabajo busca no solo profundizar en los debates académicos, sino también tender puentes entre la filosofía y otros ámbitos del conocimiento, reforzando su papel como herramienta transformadora.



Dra. María Elizabeth de los Ríos Uriarte

Doctora en Filosofía (Ciudad de México, 2011), maestra en Bioética y licenciada en Filosofía. Es profesora investigadora de la Facultad de Bioética de la Universidad Anáhuac México, editora de la revista Medicina y Ética y coordinadora de Investigación de la misma Facultad. Miembro del Sistema Nacional de Investigadores Nivel I en el área de la Interdisciplina. Miembro de la Academia Nacional Mexicana de Bioética y Research Scholar de la Catedra UNESCO en Bioética y Derechos Humanos. Es autora de tres libros: “Aproximación ética a la práctica del triage en la atención prehospitalaria (2008), “Sobre el concepto de redención en Walter Benjamin y el de liberación en Ignacio Ellacuría: hacia una teoría crítica en América Latina” (2011) y “Qué es la Bioética” (2014). Es coordinadora de tres libros, autora de más de 20 capítulos de libro y de más de 30 artículos publicados en revistas indizadas tanto nacionales como internacionales.



Dr. José Sols Lucia

Arquitecto por la Universidad Iberoamericana (2010), con una destacada trayectoria en el uso de herramientas digitales. Esta especialización le permitió adquirir experiencia profesional desde sus años universitarios, colaborando posteriormente en proyectos como La Biennale di Venezia 2010, el Museo Papalote Monterrey, el Parque Ecológico Lago de Texcoco y las fachadas de tiendas Liverpool en Perisur y Villahermosa. En 2015, inició su trayectoria como desarrollador inmobiliario, consolidándose como líder de proyectos y fundando Flux Estudio en 2019. Desde entonces, ha impulsado desarrollos y residencias privadas de lujo. Su interés por la innovación tecnológica lo llevó a concluir la Maestría en Arquitectura y Fabricación Digital por la Universidad Anáhuac en el año 2020, con un enfoque en Diseño Computacional. A partir del año 2020, combina su práctica profesional con la docencia en la Universidad Iberoamericana, impartiendo cursos con este enfoque. Además, lidera proyectos de innovación tecnológica, incluyendo el desarrollo de algoritmos inteligentes, impresión 3D de concreto, automatización y soluciones SaaS, todos orientados a procesos dentro de la Arquitectura.



Mtro. Mauricio Correa-Herrejon

Académico y estratega con más de 35 años de trayectoria en el sector financiero, especializado en banca de consumo, estrategia de negocios y economía del comportamiento. Su labor profesional se centró en la integración de modelos de gestión, análisis económico y herramientas de economía conductual aplicada, con el objetivo de diseñar estrategias comerciales, esquemas de incentivos y experiencias centradas en el cliente, sustentadas en metodologías de análisis de datos y toma de decisiones basada en evidencia. Su carrera docente inició en 1990. Actualmente, se ha desempeñado como profesor de asignatura en la Universidad Iberoamericana, impartiendo seminarios en las áreas de economía, finanzas y estrategia a nivel licenciatura y posgrado. A lo largo de su trayectoria académica, ha contribuido a la formación de cerca de mil estudiantes. Es autor de diversos artículos de análisis y ha publicado en revistas académicas nacionales e internacionales sobre estrategia empresarial, comportamiento económico y toma de decisiones. Posee maestrías en Economía por la Universidad de Nueva York (NYU), en Administración por la Universidad Iberoamericana y en Economía del Comportamiento por Evidentia University. Actualmente es miembro y vocal de la Sociedad Científica de Economía de la Conducta, desde donde promueve activamente la difusión del conocimiento y la aplicación rigurosa de la economía conductual en entornos organizacionales.



Mtro. Luis González Zarazua

Académico de profesión 45 años, interesado en el arte, la animación y su difusión a través la tecnología. Titulado en la licenciatura de Historia del arte en la Universidad Iberoamericana en 1995. Egresado de la Maestría en Psicología en la Universidad Iberoamericana, Ciudad de México, Campus Santa Fe en 2009. Titulado en la Maestría en Narrativa y producción digital en Universidad Panamericana, Ciudad de México, Campus Mixcoac, en 2013. Ha colaborado con museos de arte en México como son Museo Franz Mayer, Museo Nacional de Arte, Antiguo Colegio de San Ildefonso, y el Museo Universitario de Arte Contemporáneo. En estas instituciones ayudó como museógrafo, investigador o en la elaboración de cursos. En el extranjero colaboró con el Instituto Valenciano de Arte Moderno de España y en El Museo de Bellas Artes de Santiago de Chile. Fue subdirector en ALO.COM de un canal cultural de broadband en internet. Generando tanto las historietas como animaciones online para clásicos de la literatura universal que son lecturas obligadas para las escuelas incorporadas a la SEP. Creador del proyecto “Virreinato Virtual” (2010-2016) en consta en un simulador virtual del período barroco en México en el que se pueden experimentar varios espacios arquitectónicos y actividades sociales y culturales. Se incorporó a la Red Temática de Tecnologías Digitales para la Difusión del Patrimonio Cultural en 2015.



Lic. Yomna Ayman Mahmoud Rushdi

Ayudante de Cátedra en la Universidad de Ain Shams, Egipto, especializada en literatura. Obtuve mi licenciatura en Filología Española e Inglesa en la misma universidad, donde también completé un diploma en Literatura Hispánica. Actualmente, estoy finalizando una maestría en literatura con un enfoque interdisciplinario en inteligencia artificial y lingüística. Mi labor docente se centra en el análisis del discurso, con un interés particular en la intersección entre la tecnología y la literatura. En 2024, publiqué el artículo "Creatividad literaria de la Inteligencia Artificial en ‘La cena’" en *Philology*, la revista académica de la Facultad de Lenguas de la Universidad de Ain Shams, donde analizo el impacto de la inteligencia artificial en la creación literaria a partir de un cuento de reciente publicación. Mis intereses de investigación incluyen la narrativa digital, la inteligencia artificial aplicada a la literatura y la evolución de las formas literarias en la era digital.



Mtro. José Luis Rangel Oropeza

Arquitecto por la Universidad Iberoamericana (2010), con una destacada trayectoria en el uso de herramientas digitales. Esta especialización le permitió adquirir experiencia profesional desde sus años universitarios, colaborando posteriormente en proyectos como La Biennale di Venezia 2010, el Museo Papalote Monterrey, el Parque Ecológico Lago de Texcoco y las fachadas de tiendas Liverpool en Perisur y Villahermosa.

En 2015, inició su trayectoria como desarrollador inmobiliario, consolidándose como líder de proyectos y fundando Flux Estudio en 2019. Desde entonces, ha impulsado desarrollos y residencias privadas de lujo. Su interés por la innovación tecnológica lo llevó a concluir la Maestría en Arquitectura y Fabricación Digital por la Universidad Anáhuac en el año 2020, con un enfoque en Diseño Computacional.

A partir del año 2020, combina su práctica profesional con la docencia en la Universidad Iberoamericana, impartiendo cursos con este enfoque. Además, lidera proyectos de innovación tecnológica, incluyendo el desarrollo de algoritmos inteligentes, impresión 3D de concreto, automatización y soluciones SaaS, todos orientados a procesos dentro de la Arquitectura.



Mtra. Adriana G. Rodríguez Gamietea

Investigadora y docente con formación en filosofía y educación. Actualmente cursa el Doctorado en Pedagogía e Investigación Educativa y cuenta con una Maestría y Licenciatura en Filosofía por la Universidad La Salle. Ha impartido clases en diversas universidades e institutos. Actualmente, es profesora en la Universidad Iberoamericana, donde ha impartido cursos en filosofía, humanidades digitales y el hombre y la música. Ha complementado su formación con diplomados en docencia universitaria ignaciana, orientación educativa y análisis político. Además de su trayectoria académica, es música y compositora, con estudios en canto, composición y arreglos. Su interés por la comunicación y la enseñanza la ha llevado a incursionar como creadora de contenido en Youtube. También es cinta negra segundo dan en taekwondo, disciplina que refleja su compromiso con la perseverancia y el aprendizaje continuo.



Dr. Pedro García del Valle y Duran

Doctor en Ciencias de la Ingeniería por la Universidad Iberoamericana en 2024, especializándose en el estudio del consenso y conflicto en la toma de decisiones de redes de multiagentes. Su formación académica incluye una Maestría en Ingeniería y Ciencias de la Computación con aplicación en Inteligencia Artificial, obtenida en la Escuela de Ingeniería de la Universidad de Shizuoka, Japón, en 1994. Durante su maestría, su investigación se centró en técnicas automatizadas de refinamiento y adquisición de conocimiento, con aplicación innovadora en la planeación financiera. Previamente, en 1985, se licenció como Actuario por la Facultad de Ciencias de la UNAM, donde fue reconocido con la Medalla Gabino Barreda y el Premio Gustavo Baz Prada por su destacado desempeño académico y servicio a la comunidad. En su trayectoria profesional, ha trabajado en BANAMEX y fundó empresas de consultoría, ofreciendo servicios de administración de redes a hospitales importantes como el Instituto Nacional de Enfermedades Respiratorias y el Hospital Juárez del Seguro Social. A su regreso de Japón, ocupó el cargo de CEO en Jeol de México, S.A. de C.V. Actualmente, es profesor de tiempo completo y analista de métodos organizacionales en la Universidad Iberoamericana. Además, es miembro fundador de QURIO (Consejo de Líderes en Innovación y Tecnología en México), donde ha participado en proyectos para el Banco Interamericano de Desarrollo (BID). Entre sus reconocimientos se encuentran el premio al Mejor Profesor de Cátedra otorgado por el ITESM en 2002 y la Certificación como Instructor Nacional de México ese mismo año.



CRITERIOS GENERALES EDITORIALES PARA PUBLICACIÓN

Criterios Generales Editoriales

El Consejo Editorial de la Revista de la Asociación de Profesores e Investigadores (CERAPI) de la Universidad Iberoamericana recibe artículos de manera continua para su evaluación y posible publicación en los próximos números, por lo que, los autores pueden enviar sus artículos en cualquier momento a la siguiente dirección electrónica: api@uia.mx

Se deben enviar dos versiones del artículo: una con el nombre del autor, la institución o universidad a la que pertenece (con dirección completa) y un correo electrónico de contacto, y otra sin esta información. Esta última será la copia que se enviará para revisión por pares. Ambos archivos deben estar en formato Word.

Los artículos enviados deben ser originales y no haber sido publicados previamente en el idioma en que están escritos. Los trabajos ya publicados serán desestimados.

La extensión mínima es de ocho cuartillas y la extensión máxima es de 10 cuartillas.

Es obligatorio incluir un resumen de entre 100 y 150 palabras, que debe situarse después del título del artículo y antes de las palabras clave.

Se deben agregar entre 3 y 5 palabras clave que describan adecuadamente el contenido del texto, tales como: hermenéutica, teoría literaria, edad media, entre otras.

Formato

Título del artículo: Times New Roman, 14 puntos, versalitas, centrado, interlineado sencillo.

Nombre del autor e institución: Times New Roman, 12 puntos, centrado, interlineado sencillo.

Resumen, Palabras clave, *Abstract*, *Keywords*: Times New Roman, 11 puntos, sangría izquierda 1 cm., sangría derecha 1 cm., interlineado sencillo, justificado.

Texto principal: Times New Roman, 12 puntos, interlineado 1.5, justificado.

Citas en bloque: Times New Roman, 11 puntos, interlineado sencillo, sangría izquierda 1.5 cm., justificado, sin comillas. Todas las citas mayores a tres líneas deben ponerse en bloque.

Notas y referencias al pie de página: Times New Roman, 10 puntos, justificado. Las notas y referencias deben ser al pie de página y llevar numeración arábiga. No poner referencias entre paréntesis en el texto.

Espaciado anterior y posterior (entre párrafos) 0 pto.

Márgenes: superior e inferior 2.5 cm; izquierdo y derecho 3 cm.

Sangría en primera línea de párrafos: 1.25 cm (Sangría tabular).

Si el artículo está dividido en secciones, utilizar números romanos (I., II., III., etc.). Si hay subdivisiones, emplear letras mayúsculas (A., B., C., etc.).

Títulos de secciones: Times New Roman, 12 puntos, cursivas, sin sangría, justificado.

Criterios para las referencias

- Referencias únicamente al pie de página, es decir, no colocar referencias en el cuerpo principal del texto.
- Nombre del autor: colocar apellido y nombre separados por una coma. Ejemplo: Beuchot, Mauricio, ...
- Títulos: el título y subtítulo del libro van en cursivas, seguidos por una coma. Por ejemplo: *Tratado de Hermenéutica Analógica. Hacia un nuevo modelo de interpretación*, ...

Traductor: si la obra citada es una traducción, el nombre del traductor debe colocarse después del título y antes del lugar de publicación. Por ejemplo: Vattimo, Gianni, *Introducción a Heidegger*, trad. de Alfredo Báez, Barcelona...

Editorial: el nombre de la editorial debe colocarse después del lugar de publicación y los dos puntos. Por ejemplo: Barcelona: Editorial Gedisa...

Año de publicación: el año de publicación se coloca después de la editorial y una coma; deben escribirse completos en todos los casos. Por ejemplo: Editorial Gedisa, 1995...

Páginas citadas: la palabra "página" se abrevia "p.", seguida por un espacio y la página correspondiente. Para hacer referencia a múltiples páginas, es necesario emplear la abreviatura "pp.", seguida por un espacio y las páginas correspondientes conectadas por un guion. Por ejemplo: p. 70 o p.p. 70-71.

Si la obra ya ha sido citada anteriormente en el artículo, bastará con incluir el apellido del autor, seguido por el título y el número de página. Por ejemplo: Beuchot, *Tratado de Hermenéutica Analógica*... p. 35.

Cuando se cita una misma obra inmediatamente, si las páginas son diferentes, ocupar *Ibíd.* seguido de coma y el número de página. Por ejemplo: *Ibíd.*, p. 35.

Si es la misma obra y las mismas páginas anteriormente citadas basta con escribir *Ibíd* o *Ibídem*.

Para indicar 'Ver', 'Véase', 'Confrontar' o 'Confróntese', escribir *Cfr.*

Las referencias bíblicas deben usar las abreviaturas estándar en español, seguidas por un espacio, el número del capítulo, coma y el número del versículo. En el caso de que la cita haga referencia a varios versículos, estos irán unidos por un guion. Por ejemplo: Por ejemplo: Gen 14,10, Ex 3,14-16.

Todos los artículos deben incluir bibliografía organizada en orden alfabético de acuerdo con la primera letra del apellido del autor. La bibliografía debe estar situada al final del artículo.

Proceso de dictaminación

Todos los artículos son revisados por el Consejo Editorial de la Revista de la Asociación de Profesores e Investigadores (CERAPI) de la Universidad Iberoamericana, con la finalidad de evaluar si el texto se ajusta a los criterios básicos de la revista en cuanto a la temática y a las normas de editoriales.

Los artículos que cumplen con el punto anterior son enviados para su evaluación por pares a dos dictaminadores, bajo el sistema doble ciego.

En caso necesario se enviará a un tercer dictaminador.

Terminado este proceso se le comunicará al autor el resultado de los dictámenes.

Prohibición del uso de Inteligencia Artificial (IA) en la redacción de los artículos

La Revista de la Asociación de Profesores e Investigadores de la Universidad Iberoamericana mantiene un compromiso con la integridad académica y la originalidad en la autoría de los trabajos publicados. Por lo tanto, se prohíbe estrictamente el uso de herramientas de Inteligencia Artificial (IA) en la redacción de los manuscritos enviados para su publicación. Se espera que los autores sean los responsables directos de la concepción, investigación, análisis e interpretación de los datos, así como de la elaboración del texto del artículo. El incumplimiento de esta norma podrá conllevar el rechazo del manuscrito o la retractación de la publicación en caso de ser descubierto posteriormente.

Acceso abierto

La Revista de la Asociación de Profesores e Investigadores de la Universidad Iberoamericana, en cumplimiento con las directrices internacionales sobre acceso abierto, ofrece gratuitamente el contenido de la revista en formato digital. Esto se realiza con el fin de fomentar la investigación académica y la difusión de las humanidades y ciencias.

Bajo esta política, se permite el acceso y distribución del material de acuerdo con las siguientes condiciones: a. El material no debe utilizarse con fines comerciales o lucrativos. b. Es obligatorio realizar una cita adecuada al emplear los artículos, incluyendo el nombre del autor, el título del artículo, el nombre de la revista, el número y la fecha de publicación. c. Cualquier duda relacionada con el uso o distribución del contenido debe ser consultada con el Consejo Editorial.

Este esquema asegura una mayor difusión del conocimiento sin restricciones económicas.

Declaración de ética y buenas prácticas

El Consejo Editorial de la Revista de la Asociación de Profesores e Investigadores (CERAPI) de la Universidad Iberoamericana establece que los autores deben seguir ciertos principios para garantizar el cumplimiento de estándares éticos esenciales:

A. La revista condena y prohíbe el plagio. Por ello, los autores deben declarar por escrito que su artículo no contiene plagio, no ha sido previamente publicado y no reproduce partes de otros trabajos ya publicados. CERAPI tiene la facultad de utilizar sistemas de verificación para asegurarse de que el artículo cumple estos criterios. Si se detecta plagio, antes o después de la publicación, se iniciará un proceso para aclarar responsabilidades y se notificará a las autoridades competentes.

B. Al enviar un artículo para su posible publicación, los autores aceptan que será sometido a un proceso de revisión por pares y se comprometen a proporcionar información veraz sobre su afiliación y otros datos académicos requeridos.

C. El arbitraje de los artículos en la Revista de la Asociación de Profesores e Investigadores exige que los revisores se apeguen a criterios de objetividad académica. Si existe un conflicto de intereses en relación con el autor, su institución, la temática del artículo o cualquier otro aspecto que comprometa la imparcialidad, los revisores deben declinar participar. Además, deben mantener la confidencialidad del contenido de los artículos hasta que sean publicados.

D. CERAPI se compromete a garantizar el cumplimiento de estas normas éticas, evitando los conflictos de interés en la aceptación o rechazo de artículos, protegiendo la confidencialidad de la información y el anonimato de los revisores, y asegurando que estos criterios sean comunicados a autores y colaboradores.

E. La revista, a través de CERAPI, se encargará de velar por el cumplimiento ético durante el proceso de publicación. Esto incluye notificar a tiempo los dictámenes y cualquier corrección necesaria en los artículos. En caso de errores, la revista publicará las enmiendas correspondientes, y se invita a los autores a informar sobre cualquier error que detecten.

**REVISTA DE LA
ASOCIACIÓN DE
PROFESORES E
INVESTIGADORES DE
LA IBERO**

www.apibero.org

MÉXICO, AÑO 2/NÚM. 2/2025