



Arquisur revista

ISSN: 1853-2365

ISSN: 2250-4206

arquisurrevista@fadu.unl.edu.ar

Universidad Nacional del Litoral

Argentina

Cristaldo, Juan Carlos

Una reflexión sobre las publicaciones de arquitectura y el urbanismo en el contexto de la inteligencia artificial

Arquisur revista, vol. 14, núm. 26, 2024, Julio-Diciembre, pp. 28-31

Universidad Nacional del Litoral

Argentina

DOI: <https://doi.org/10.14409/ar.v14i26.14106>

Disponible en: <https://www.redalyc.org/articulo.oa?id=699780064004>

- Cómo citar el artículo
- Número completo
- Más información del artículo
- Página de la revista en redalyc.org

redalyc.org

Sistema de Información Científica Redalyc

Red de revistas científicas de Acceso Abierto diamante

Infraestructura abierta no comercial propiedad de la academia

Una reflexión sobre las publicaciones de arquitectura y el urbanismo en el contexto de la inteligencia artificial

Juan Carlos Cristaldo

Facultad de Arquitectura, Diseño y Arte
Universidad Nacional de Asunción
Paraguay

juan.cristaldo@cidi.fada.una.py

<https://orcid.org/0000-0002-0675-671X>

CÓMO CITAR

Cristaldo, J. C. Una reflexión sobre las publicaciones de arquitectura y el urbanismo en el contexto de la inteligencia artificial. *ARQUISUR Revista*, 14(26), 28-31.
<https://doi.org/10.14409/ar.v14i26.14106>

El lanzamiento de ChatGPT (iniciales de Chat Generative Pre-Trained Transformer), que tuvo lugar el 22 de noviembre de 2022, revolucionó el mundo. Este Modelo Extenso de Lenguaje (Large Language Model o LLM) impactó a la sociedad global al presentar una herramienta que respondía de manera definitiva a la pregunta de si sería posible para los seres humanos desarrollar aplicaciones prácticas de Inteligencia Artificial (IA). Por supuesto, este no fue el primer momento de impacto: Deep Blue ya había derrotado a Kasparov (Deep Blue IBM, s.f.), y AlphaGo había vencido a Lee Sedol, ganador de 18 títulos mundiales (AlphaGo, 30 de octubre de 2024). Pero ChatGPT es diferente en el sentido de que es una herramienta transversal no especializada. Puede escribir prosa, documentos técnicos, poemas, e incluso programar y, por tanto, a pesar de los errores que comete y las «alucinaciones» en que incurre, es aplicable a múltiples campos de las actividades humanas, hasta en la redacción de ensayos académicos —para desesperación y confusión de docentes en todo el mundo.

Aunque ChatGPT está aún muy lejos de ser una IA General (IAG) —un programa con inteligencia similar a la humana y capacidad de aprender de modo independiente— su surgimiento ha desatado un acelerado proceso de adopción tecnológica en los más diversos campos, junto con expectativas —probablemente infundadas— de crecimiento económico infinito y también temores —quizá exagerados, pero es probable que no totalmente carentes de fundamento— que apuntan a posibles desastres apocalípticos.

En definitiva, es innegable que existen enormes potencialidades y desafíos que se derivan hoy de estas tecnologías. Miremos las potencialidades: aplicaciones de Inteligencia Artificial están revolucionando la interpretación de imágenes de resonancia magnética en el contexto del diagnóstico de cánceres de mama (Shelth & Giger, 2019), herramientas de *deep learning* están asistiendo a investigadores en el descubrimiento de nuevos materiales (Merchant *et al.*, 2023), mientras que otros investigadores emplean *machine learning* para descubrir nuevas moléculas que puedan ser utilizadas en el desarrollo de medicinas (Zewe, 17 de junio de 2024).

Al mismo tiempo, existen desafíos de medio o largo plazo, aunque también de corto plazo, que derivan de estas tecnologías. Exploremos primero los problemas de medio o largo plazo. Frecuentemente, la reflexión y discusión sobre este tipo de desafíos se circunscriben a una comunidad muy pequeña de expertos, pero se propone aquí que —por sus implicaciones— deberían ser motivo de consideración y discusión por parte de toda la sociedad. Se citan a continuación algunos de estos desafíos, elegidos precisamente por la relevancia de sus posibles consecuencias:

El problema de la alineación: refiere a cómo garantizar que la IA tenga objetivos alineados con objetivos humanos y no establezca propios, que estén en conflicto con metas humanas. En otras palabras, si potencialmente se desarrolla una inteligencia artificial que alcanza el nivel humano (IAG) o incluso que supera el nivel humano, la duda es cómo garantizar que la herramienta no cree objetivos independientes que sean diferentes o estén en conflicto con metas humanas. Este es el tipo de escenario distópico —pero factible y objeto de preocupación real por parte de especialistas en todo el mundo— en el que «las máquinas deciden controlar el mundo» (Ngo *et al.*, 19 de marzo de 2024).

Respecto del problema de alcanzar el umbral de IAG, diversos especialistas y líderes de la industria tienen perspectivas distintas sobre cuánto tiempo falta para que la humanidad alcance un evento de «singularidad», es decir, el momento de creación de una IAG. Aún más, existen desacuerdos incluso sobre cómo definir una IAG. Jensen Huang, CEO de NVIDIA propone que sería una herramienta capaz de funcionar «8% mejor que la mayoría de la población en tareas específicas, tales como pasar un examen de acreditación o resolver problemas lógicos». Mustafa Suleyman señala que una IAG sería capaz de convertir USD 100 000 en USD 1 000 000 sin recibir instrucciones. Por su parte, el cofundador de Apple, Steve Wozniak, propone una definición más sucinta: una IAG podría hacer una taza de café (*The Economist*, 28 de marzo de 2024). Pero, en definitiva, la expectativa es que se llegue al escenario de que existan IAG en un período de tiempo relativamente acotado, de los 5 años propuestos por Huang de NVIDIA, al rango de entre 5 a 20 años de Geoffrey Hinton (Korinek & Suh, 2024).

De cualquier manera, lo relevante es notar que: (i) en general ya no se discute si una IAG ocurrirá o no, sino que se habla del horizonte temporal en que ocurrirá; (ii) en la literatura, luego de ChatGPT, el horizonte temporal se redujo, siendo que anteriormente se hablaba de 20 a 100 años y ahora se mencionan períodos de entre 5 a 20 años; (iii) todo señala que muy probablemente se alcance la singularidad antes de haber resuelto el problema de la alineación.

Estos desafíos, inmensos, y que parecen ser sacados del reino de la ciencia ficción, son problemas que aún pertenecen al futuro. Discutamos ahora los desafíos que ya tenemos en el presente, al momento de escribir estas líneas. Dividiremos los problemas presentes en dos subgrupos: (i) problemas relacionados con la IA que están «aguas arriba» del usuario, es decir, desafíos que afectan a las compañías que desarrollan IA y a la sociedad en general; y (ii) problemas relacionados con la IA que están «aguas abajo» del usuario, es decir, desafíos que involucran al usuario y a las instituciones que interactúan con el usuario, tales como instituciones educativas o lugares de trabajo.

En cuanto a los problemas presentes que están «aguas arriba» del usuario, y atento al LLM y a la propiedad intelectual, los modelos son entrenados «consumiendo» enormes cantidades de información de la Internet, lo que incluye textos de Wikipedia, noticias, libros digitales, hilos de discusión en Reddit, sin reconocimiento autoral ni compensación (Johri, 6 de junio de 2023). Se estima que para 2021 Open Ai había usado «todas las palabras de Internet en idioma inglés para entrenar a su modelo, pero que necesitaba aún más» (*The Daily*, 16 de abril de 2024), y luego pasó a transcribir audio (como audiolibros o podcasts) y videos (especialmente de YouTube) para crear más datos para entrenar su modelo. Por supuesto, de este enorme conjunto de datos, una gran parte debe estar protegida por derechos autorales. En todo el mundo están surgiendo estudios académicos que se focalizan en esta cuestión, así como disputas legales —de autores de libros, programadores de software, etc.— donde se demanda a empresas como Open Ai y Microsoft, que utilizan datos y obras científicas, literarias y artísticas, para entrenar a los LLM (Dornis & Stober, 2024; Gerken, 27 de diciembre de 2023).

Con relación a los impactos ambientales del consumo de energía necesaria para la IA, en el contexto del cambio climático, las propiedades aparentemente mágicas de esta enmascaran no solo el origen de sus capacidades, basadas en los datos usados para entrenar los modelos. Se hace incluso difícil imaginar o recordar que existe un costo energético asociado a procesar cada *prompt* que es cargado en un LLM. Y sin embargo este costo existe. La demanda energética de la IA solo crecerá en el futuro y puede constituirse en un factor más de incremento de emisiones que acentúan el cambio climático. No hay una respuesta o un modelo predominante de cómo proveer energía a los centros de datos necesarios para sostener los LLM, pero, al momento de redactar este artículo, diversas empresas están impulsando proyectos de energía nuclear para atender la demanda de las herramientas de IA (Da Silva, 15 de octubre de 2024).

Consideremos ahora problemas presentes que están «aguas abajo» del usuario.

¿Cómo se define autoría en el mundo de la IA? En nuestros roles de docentes, autores y colegas, esta es una cuestión tanto práctica como ética para la que no tenemos una respuesta clara. Concretamente, ¿cuál es un uso ético y adecuado de las herramientas de IA en el ámbito académico? ¿Se las puede usar para buscar bibliografía y resumir textos para leer antes de escribir? Probablemente muchos coincidirían en que la respuesta sería sí, porque la herramienta asiste al autor. Ahora bien, ¿se pueden usar herramientas de IA también para escribir artículos enteros? En general, es factible que nuestra intuición ética indique que la respuesta sería no, porque tal escenario borraría cualquier rastro de autoría y de contribución humana al proceso. Y, por el contrario, si aceptamos la noción de que artículos enteros generados por IA son aceptables en publicaciones, ¿eso significaría que convertirse en un científico/académico/artista competente se resume en el mundo contemporáneo a saber escribir *prompts*?

Nuestras herramientas nos construyen: el *homo sapiens* es también *homo faber* (Lawlor & Moulard-Leonard, 2022). Nosotros hacemos nuestras herramientas, pero nuestras herramientas nos construyen. Esto ha sido así desde el descubrimiento del fuego y de la agricultura. En las sociedades contemporáneas, millones de personas hacen ejercicio físico como recreación o como opción de preservación de la salud física o mental. En el escenario concreto de la IA, la pregunta es: ¿cómo afecta ya y cómo afectará en el futuro nuestro desarrollo neurocognitivo y nuestra biología en general el «tercerizar» el esfuerzo de pensar? En el futuro, ¿pensaremos voluntariamente —no por la necesidad profesional de resolver problemas—, del mismo modo que hoy vamos al gimnasio para ejercitar un cuerpo que en rigor no requiere —para su productividad económica— esfuerzo físico?

Por último, nos referiremos a la polarización del discurso público desde los algoritmos, la polarización política derivada de las redes sociales.

Estos algoritmos buscan maximizar la interacción y conexión con determinada plataforma o red social y para esto muestran a los usuarios contenidos que refuerzan sus puntos de vista políticos y sus valores, creando una disociación cognitiva entre sectores de la sociedad que tienen perspectivas ideológicas diferentes (Cho *et al.*, 2020; Conover *et al.*, 2021). Este parece ser un ejemplo concreto de la disquisición teórica del «Maximizador de Clips» (Brindle, 6 de diciembre de 2017; LessWrong, s.f.), donde una IA recibe la tarea aparentemente simple e inocua de optimizar el proceso de creación de clips y, a partir de esta premisa estrecha, termina destruyendo a la humanidad y al Universo al asignar todos los recursos disponibles a su función de utilidad. En el caso de las redes sociales, la función de utilidad es maximizar la interacción y conexión con los usuarios, brindando información y datos que han demostrado interesarle para, con esto, maximizar la información que recogen de los mismos y luego comercializar estos datos o servicios ultrasegmentados de publicidad a terceros. La disgregación del tejido social y de gobiernos democráticos sería, en esta hipótesis, apenas un subproducto desafortunado.¹

¿CÓMO VER LAS PUBLICACIONES ACADÉMICAS EN ESTE CONTEXTO?

Ante los inmensos desafíos del futuro, se propone recuperar los valores fundamentales de las publicaciones científicas y de vanguardia artística que siguen siendo relevantes para la coyuntura contemporánea.

Por un lado, de la historia y praxis de las revistas científicas —que impulsaron innovaciones como la revisión por pares, la periodicidad breve y la discusión pública de resultados y métodos—, se propone rescatar las siguientes aspiraciones: (i) constituir por medio del intercambio, redes y espacios virtuales para compartir métodos y resultados; (ii) establecer una dinámica de discusión pública de los métodos y resultados publicados, entendiendo que el debate público de ideas lleva a la aceleración del ritmo de mejora de los resultados; (iii) sostener un discurso racional, equilibrado y basado en evidencia; (iv) promover relaciones de cooperación y competición entre los miembros de la comunidad académica; y (v) fortalecer el rol de los expertos como veedores y garantes de la calidad del conocimiento publicado.

Por otro lado, de la historia y praxis de las revistas artísticas/arquitectónicas —que desde un lenguaje iconoclasta y transgresor presentaban una visión crítica del presente, valoraciones no culturalmente obvias del pasado y propuestas de cara al futuro—, se destacan los siguientes aspectos: (i) la capacidad de enunciar de modo claro y potente discursos sobre las crisis y problemas del arte y la sociedad actuales; (ii) la voluntad de promover un debate sobre valores o legados histórico culturales que se desean rescatar; (iii) sostener un criterio de potencia e incluso transgresión y poética en el discurso como método de apertura de nuevos caminos culturales; (iv) constituir un espacio virtual para compartir métodos y resultados considerados válidos; y (v) fortalecer una comunidad internacional de artistas que componen una vanguardia.

Por todo lo aquí escrito, está claro que los *journals* como ARQUISUR REVISTA y todos los medios de comunicación tradicionales enfrentan una crisis de resignificación y de reflexión sobre su propio propósito y sentido ante el surgimiento de las herramientas de IA y ante las potencialidades y riesgos que se asocian a dicha tecnología. Sin embargo, se postula que los valores y aspiraciones citados en los párrafos precedentes —que derivan de publicaciones académicas y de revistas de vanguardia artística— siguen siendo válidos para el contexto contemporáneo y brindan un andamiaje intelectual y ético que fortalece la relevancia de las publicaciones que como ARQUISUR REVISTA, se focalizan en la arquitectura y el urbanismo. El futuro requiere de nosotros seguir trabajando en red, de modo riguroso, transparente, colaborativo, formulando conclusiones basadas en evidencia y visiones artísticas poéticas y creativas. 🌱

Msc. Arq. JUAN C. CRISTALDO
Miembro del Comité editorial ARQUISUR REVISTA
Facultad de Arquitectura, Diseño y Arte
Universidad Nacional de Asunción
Paraguay
Noviembre 2024

1. También puede considerarse como un problema contemporáneo el uso intencionalmente negativo de las redes sociales por actores sociales. Por ejemplo, la diseminación intencional de *fake news* o la utilización de IA para crear falsas imágenes o videos, lo que ha dado en denominarse *deep fakes*. El autor considera que, en general, ambos problemas coexisten y se potencian: a saber, la polarización que deriva de las «burbujas de información» creadas por los algoritmos y el uso malintencionado de las redes.

REFERENCIAS BIBLIOGRÁFICAS

- A lphaGo (30 de octubre de 2024). Google DeepMind. <https://deepmind.google/research/breakthroughs/alphago/>
- Brindle, J. (6 de diciembre de 2017). This game about paperclips says a lot about human desire. VICE. <https://www.vice.com/en/article/universal-paperclips-depth-desire/>
- Cho, J.; Ahmed, S.; Hilbert, M.; Liu, B. & Luu, J. (2020). Do search algorithms endanger democracy? An experimental investigation of algorithm effects on political polarization. *Journal of Broadcasting & Electronic Media*, 64(2), 150-172. <https://doi.org/10.1080/08838151.2020.1757365>
- Conover, M.; Ratkiewicz, J.; Francisco, M.; Goncalves, B.; Menczer, F. & Flammini, A. (2021). Political polarization on Twitter. Proceedings of the International AAAI Conference on Web and Social Media, 5(1), 89-96. <https://doi.org/10.1609/icwsm.v5i1.14126>
- Da Silva, J. (15 de octubre de 2024). Google turns to nuclear to power AI data centres. BBC. <https://www.bbc.com/news/articles/c748gn94k950>
- Deep Blue IBM (s. f.). IBM. <https://www.ibm.com/history/deep-blue>
- Gerken, B.T. (27 de diciembre de 2023). New York Times sues Microsoft and OpenAI for 'billions'. BBC. <https://www.bbc.com/news/technology-67826601>
- Johri, S. (6 de junio de 2023). The making of ChatGPT: From data to dialogue. Harvard Kenneth C. Griffin. <https://sitn.hms.harvard.edu/flash/2023/the-making-of-chatgpt-from-data-to-dialogue/>
- Korinek, A. & Suh, D. (2024). Scenarios for the transition to AGI. NBER. https://www.nber.org/system/files/working_papers/w32255/w32255.pdf
- Lawlor, L. & Moulard-Leonard, V. (2022). Henri Bergson. En Zalta, E.N. & Nodelman, U. (Eds.). *The Stanford Encyclopedia of Philosophy*. <https://plato.stanford.edu/archives/win2022/entries/bergson/>
- LessWrong (s.f.). Squiggle Maximizer (formerly «Paperclip maximizer»). LessWrong. <https://www.lesswrong.com/tag/squiggle-maximizer-formerly-paperclip-maximizer>
- Merchant, A.; Batzner, S.; Schoenholz, S.S.; Aykol, M.; Cheon, G. & Cubuk, E. D. (2023). Scaling deep learning for materials discovery. *Nature*, 624(7990), 80-85. <https://doi.org/10.1038/s41586-023-06735-9>
- Ngo, R.; Chan, L. & Mindermann, S. (19 de marzo de 2024). The alignment problem from a deep learning perspective. arXiv. <https://arxiv.org/pdf/2209.00626>
- Peeters Online Journals (s.f.). Peeters. https://poj.peeters-leuven.be/content.php?url=journal&journal_code=JDS
- Sheth, D. & Giger, M. L. (2019). Artificial intelligence in the interpretation of breast cancer on MRI. *Journal of Magnetic Resonance Imaging*, 51(5), 1310-1324. <https://doi.org/10.1002/jmri.26878>
- The Daily (16 de abril de 2024). Los datos de inteligencia artificial y sus desafíos. *The New York Times*. <https://www.nytimes.com/2024/04/16/podcasts/the-daily/ai-data.html>
- The Economist (28 de marzo de 2024). How to define artificial general intelligence. *The Economist*. <https://www.economist.com/the-economist-explains/2024/03/28/how-to-define-artificial-general-intelligence>
- Zewe, A. (17 de junio de 2024). A smarter way to streamline drug discovery. *MIT News*. Massachusetts Institute of Technology. <https://news.mit.edu/2024/smarter-way-streamline-drug-discovery-0617>