

La inteligencia artificial y la transformación de la publicación científica

On Artificial Intelligence and the Transformation of Scientific Publishing

Ingri G. Camacho-Triana¹, Julian Arcila-Forero², José D. Gutierrez-Mendoza³,
Ian F. Guarnizo-Martinez⁴, Lenin A. Bulla-Cruz⁵ and Sonia C. Mangones M.⁶

La inteligencia artificial (IA) ha dejado de ser una promesa para convertirse en un componente activo del ecosistema editorial. Frente al crecimiento acelerado del volumen de publicaciones, la presión por la visibilidad y las exigencias de eficiencia en los procesos de revisión y edición, esta tecnología se integra cada vez más al flujo de trabajo de los autores, revisores y equipos editoriales. Su adopción, sin embargo, no está exenta de preocupaciones: amplifica debates sobre rigor científico, equidad, transparencia y responsabilidad editorial.

En esta editorial, nos propusimos examinar, desde la experiencia de *Ingeniería e Investigación* (Iel), una revista de acceso abierto diamante, las tensiones que emergen al incorporar IA en procesos clave de la publicación científica. Más que evaluar tecnologías específicas, buscamos comprender los dilemas éticos y conceptuales que plantean: ¿hasta dónde facilitan la democratización del conocimiento y en qué momento comienzan a homogeneizar, sesgar o trivializar contenidos?, ¿qué responsabilidades nuevas surgen para editores y revisores cuando la automatización interviene en decisiones académicas?

Este texto es, por lo tanto, una invitación a reflexionar sobre el lugar que ocupará la IA en la publicación científica contemporánea: ni promesa redentora ni amenaza inminente, sino un conjunto de herramientas que exige criterios claros, alfabetización crítica y una comprensión renovada de la integridad editorial.

El filtro inicial (*desk review*)

En Iel, la revisión preliminar de pertinencia ha evolucionado hacia un modelo en el cual la IA se integra como un apoyo técnico—no decisorio—para optimizar la fase del *desk review*: evaluación de coherencia temática, verificación de

formato y lenguaje, y detección inicial de posibles conflictos de interés. Estudios recientes [1] muestran que el uso de modelos de lenguaje puede automatizar tareas como la revisión preliminar y la verificación de estilo y consistencia, contribuyendo a reducir la carga editorial y mejorar la normalización de los criterios de admisión.

No obstante, la primera tensión radica en que, al asimilar la IA como filtro de pertinencia exige una reflexión crítica sobre los riesgos de sesgo y exclusión. La literatura advierte que las políticas editoriales sobre el uso de IA no siempre están claras [2]. Cuando la IA asiste la evaluación, la confidencialidad, la diversidad epistemológica y la integridad científica dependen de una supervisión humana [3]. Asimismo, las directrices internacionales, como la *Recomendación sobre la ética de la inteligencia artificial* de la UNESCO [4], insisten en que la supervisión humana debe mantenerse como principio rector para evitar exclusiones inadvertidas y preservar la diversidad académica.

Por ello, en Iel, la IA actúa como un ‘agente colaborador’ en la edición, como un instrumento para mejorar la gestión editorial, operando bajo un modelo híbrido, ético y reflexivo. La decisión editorial final sigue siendo humana, deliberativa y evaluada por pares académicos expertos.

En coherencia con estos desafíos, la Revista avanza hacia un modelo de gobernanza editorial, incorporando principios para el uso de IA en la revisión preliminar. Estos principios incluyen (i) transparencia, mediante la documentación del uso de herramientas automatizadas en cada decisión editorial; (ii) supervisión humana obligatoria, garantizando que toda recomendación generada por IA sea verificada por editores con criterio académico; (iii) uso ético, asegurando que la IA se utilice como asistente técnico y no como un mecanismo autónomo de aceptación o rechazo.

¹ Faculty of Engineering Universidad Nacional de Colombia, Bogotá, Colombia. Email: igcamachot@unal.edu.co

² Faculty of Engineering Universidad Nacional de Colombia, Bogotá, Colombia. Email: jfarclaf@unal.edu.co

³ Faculty of Engineering Universidad Nacional de Colombia, Bogotá, Colombia. Email: jodgutierrezme@unal.edu.co

⁴ Faculty of Engineering Universidad Nacional de Colombia, Bogotá, Colombia. Email: iguarnizom@unal.edu.co

⁵ Transport & Roads, Dept. of Civil & Agricultural Engineering, Faculty of Engineering Universidad Nacional de Colombia, Bogotá, Colombia. Email: labullac@unal.edu.co

⁶ Transport & Roads, Dept. of Civil & Agricultural Engineering, Faculty of Engineering Universidad Nacional de Colombia, Bogotá, Colombia. Email: scmangonesm@unal.edu.co

La revisión por pares

El modelo de déficit es una práctica común en la revisión por pares de manuscritos enviados para publicación en revistas científicas [5]. Las premisas del modelo apuntan a que las observaciones de los pares revisores corresponden a fallas o deficiencias del documento. Estas observaciones son tomadas como ciertas, y los autores deben resolverlas hasta eliminar toda discrepancia. Además, al considerar los comentarios como críticas justas, los desacuerdos con los revisores pueden verse como faltas de comprensión de los autores, en lugar de propiciar una discusión académica justificada.

Bajo esta lógica, la integración de herramientas automatizadas plantea un riesgo de exacerbación. En el proceso de revisión por pares, el uso de la IA como asistente puede favorecer la detección de plagio y de errores metodológicos [6]. Sin embargo, si se depende excesivamente de la herramienta, pueden pasarse por alto los aportes originales y la innovación, llevando a que el modelo de déficit se enfoque principalmente en la identificación de fallas de forma, por encima de las deficiencias sustanciales, donde la experticia del revisor es esencial. Por otra parte, ante el sesgo de una profundidad aparente en los manuscritos generados mediante IA, al emplear lenguaje técnico y argumentos circulares con tono académico, existe el riesgo de que el revisor se confunda si la revisión es somera o precipitada.

En consecuencia, la automatización del proceso de revisión por pares puede potenciar prácticas depredadoras, fundamentalmente en relación con la imposibilidad de identificar manuscritos de baja calidad que han sido transformados —o incluso generados— mediante IA, bajo una productividad mecánica que carece de rigor e integridad. Lo anterior crea un bucle indeseado en el que el automatismo extremo elabora y avala documentos faltos de investigación.

En la revisión por pares asistida por IA, el artículo de [7] sobre “los peligros del uso de IA en la revisión por pares” sintetiza preocupaciones que ya recogen COPE, JAMA, Science y múltiples políticas editoriales: la confidencialidad de los manuscritos inéditos, la pérdida de rigor y especificidad en las evaluaciones, la posible representación fraudulenta del rol de la IA y el riesgo de manipulación sistemática de la revisión por pares [8]. A ello se suma evidencia reciente de prácticas como la inserción de instrucciones ocultas en los manuscritos para sesgar revisiones automatizadas, lo que muestra hasta qué punto se pueden explotar las debilidades de los sistemas que sustituyen el juicio humano [9]. Al mismo tiempo, experiencias como el AI Fast Track de NEJM AI ilustran que algunos editores están explorando modelos híbridos en los que la IA cubre aspectos técnicos mientras se preserva la decisión académica en manos humanas [10]. De aquí se desprende nuestra segunda tensión: ganar eficiencia no puede implicar renunciar a la pluralidad de miradas, a la creatividad y a la responsabilidad personal que sostienen la revisión por pares.

La voz del autor y la corrección de estilo

La IA no solo ha transformado profundamente todos los procesos editoriales; también ha permeado por completo la creación de documentos científicos. Hoy en día, este tipo de herramientas acompaña a los autores en procesos como la gestión bibliográfica [11], la corrección gramatical [12] y de forma, y, en los casos más preocupantes, la concepción misma de los artículos [13]. Su uso ha despertado todo tipo de debates, e incluso ha amenazado con automatizar el oficio del corrector de estilo.

Las herramientas de IA suplen una necesidad fundamental en el mundo actual: pueden generar textos gramaticalmente correctos en segundos, un apoyo esencial en el marco de las prácticas de publicación masificada. No obstante, si bien cuentan con algunos de los elementos involucrados en el proceso de corrección de estilo, no son revisores de lenguaje. La razón es muy sencilla: modelos de lenguaje de gran tamaño (LLMs) no tienen las mismas facultades que los humanos en términos de lenguaje, ni lo adquieren de la misma manera. Estos modelos procesan texto con base en corpus enormes de información lingüística y lo generan con base en la ocurrencia (expresada en términos estadísticos) de una palabra tras otra. A diferencia de los humanos, no tienen conciencia de lo que dicen; no lo entienden. No relacionan conceptos ni expresan su realidad por medio del lenguaje. No son capaces de revisar un texto porque no tienen conocimiento accionable de los temas que se tratan ni de lo que se expresa.

Sin embargo, y en el centro del debate, está el hecho de que la IA, usada bajo la supervisión y con la responsabilidad humana, sí sirve para mejorar la calidad, la legibilidad y la corrección de un texto, pues facilita las búsquedas y tiene ‘claras’ las reglas gramaticales—sobre todo las del inglés. Esto acelera el proceso de redacción y, naturalmente, el de la corrección de estilo, que termina centrándose más en coherencia, cohesión y contenido que en aspectos gramaticales y de sintaxis.

No obstante, si bien estas herramientas ayudan a científicos con poco tiempo o pocas nociones del idioma a publicar sus investigaciones, el precio termina siendo la voz individual de cada autor. ChatGPT y DeepL tienen su propio estilo—algo que los correctores han aprendido a ver después de varios años de interactuar con ellos y sus derivados—, y su uso masivo homogeniza el lenguaje empleado en la producción de texto; estamos caminando hacia un mundo en el que la IA habla *por* nosotros, no *como* nosotros.

Algo similar ocurre en el *desk review* y la edición de estilo. Múltiples revistas han adoptado políticas que permiten el uso de IA para tareas limitadas, como corrección de lenguaje o traducción, pero prohíben que estas herramientas generen contenido sustantivo o sean tratadas como autoras, y exigen declarar su uso con transparencia [14]. Esto confirma la tercera y cuarta tensión: la IA puede ayudar a filtrar, homogeneizar formatos y pulir textos, pero, si no se delimita

su papel, puede reforzar sesgos en la selección temprana de manuscritos, estandarizar en exceso las voces y erosionar la trazabilidad de las contribuciones intelectuales.

La representación gráfica

La incorporación de herramientas de IA en la producción visual de las revistas científicas plantea una tensión relevante entre la vulgarización y la democratización. Por un lado, existe el riesgo de que se desarrollen imágenes, infografías o textos de divulgación que, aunque funcionales, presenten recursos visuales simplificados (mas no necesariamente sencillos) que atenúen la complejidad conceptual de algunos contenidos y generen la percepción de una vulgarización del conocimiento. Esta percepción se amplifica cuando el impacto visual adquiere un peso mayor que la profundidad metodológica, creando un desbalance entre forma y sustancia.

En el frente visual, editoriales como Springer Nature y Elsevier han optado por restringir o prohibir el uso de imágenes generadas por IA en artículos científicos, debido a los riesgos legales y de integridad que conlleva, incluso en casos donde la revisión por pares no detecta ilustraciones abiertamente problemáticas [15]. Esta discusión refuerza una tensión relevante: la IA puede democratizar la producción gráfica para revistas con recursos limitados, pero también puede banalizar o distorsionar el contenido científico si se prioriza el impacto visual sobre la precisión conceptual.

Aun así, desconocer el potencial democratizador de estas tecnologías sería un error. En contextos donde las revistas operan con presupuestos limitados y modelos de acceso abierto, la IA permite crear material gráfico y textos de divulgación de forma accesible, rápida y económica. Su uso puede facilitar la comunicación de hallazgos complejos a públicos más amplios, reducir desigualdades entre revistas con distintos niveles de recursos y fortalecer la presencia de las publicaciones en entornos digitales donde la imagen es determinante para la visibilidad.

El reto editorial consiste en evitar que la facilidad de generación visual conduzca a la proliferación de imágenes genéricas o poco matizadas en términos científicos, o a una dependencia acrítica de modelos que optimizan por coherencia estética en lugar de buscar la precisión conceptual. Sin protocolos claros, la calidad de las imágenes y los textos de divulgación generados por IA puede caer en ambigüedades, donde se difuminan las responsabilidades de los autores, editores y revisores, comprometiendo la confiabilidad del contenido publicado.

Por ello, la discusión no debe centrarse en aceptar o rechazar la IA, sino en establecer criterios que garanticen su uso responsable. La supervisión humana experta, la verificación científica de cada material visual o texto de divulgación y la definición de estándares editoriales sólidos son indispensables para convertir estas herramientas en

un recurso de democratización y no en un vehículo de trivialización. El equilibrio entre acceso, claridad y rigor dependerá, en última instancia, de las políticas editoriales que cada revista adopte y de su capacidad para integrar estas tecnologías sin renunciar a la integridad científica.

La infraestructura invisible (formatos y metadatos)

En el contexto actual de la publicación científica, los metadatos y los formatos digitales han dejado de ser componentes secundarios del proceso editorial para consolidarse como elementos estructurales de la comunicación académica. La creciente demanda de formatos como HTML, XML-JATS, EPUB y otros orientados a la accesibilidad responde tanto a los requisitos de indexadores como a la diversificación de dispositivos y modos de acceso al conocimiento.

En este escenario, la IA se incorpora como un apoyo para la creación de una fuente única (semilla) desde la cual puedan derivarse múltiples formatos de publicación. Este enfoque, explorado desde distintos flujos de trabajo basados en LaTeX, Pandoc [16] o esquemas XML-First [17, 18], permite optimizar la producción editorial mediante procesos verificados por código, en los que la supervisión humana se concentra en la validación de resultados y en la revisión de incidencias reportadas por los sistemas de conversión, más que en la corrección manual reiterada de cada versión (como en los modelos tradicionales), lo cual genera, entre otros, una reducción drástica de costos y de tiempos de procesamiento.

Por otra parte, la centralidad de los metadatos se hace aún más evidente al considerar su papel en la interoperabilidad y el descubrimiento automatizado de los contenidos científicos. Los protocolos de cosecha, plataformas como DOAJ y Crossref [19] y los motores de búsqueda académicos dependen de descripciones normalizadas y consistentes para integrar los artículos en sus sistemas de información. En este punto, la IA puede contribuir tanto a la verificación de la calidad y la coherencia de los metadatos como al análisis de su pertinencia para los procesos de indexación y recuperación de información, incluyendo la selección de títulos, descriptores y palabras clave con mayor potencial de visibilidad.

Sin embargo, la progresiva masificación de estas herramientas plantea un nuevo reto editorial: en un entorno donde todos los equipos pueden optimizar sus metadatos con apoyo de la IA, la diferenciación ya no dependerá únicamente de la técnica, sino de una comprensión más profunda de la manera en que los algoritmos de búsqueda y las IA de revisión de literatura jerarquizan y recomiendan el conocimiento científico.

Desde esta perspectiva, la producción científica avanza hacia una concepción del artículo como objeto computacional, más que como un texto aislado [20]. La estandarización de

formatos y el enriquecimiento de metadatos permiten que los contenidos sean cosechados, enlazados y reutilizados en infraestructuras externas, facilitando su integración en redes de conocimiento legibles por máquinas. En este marco, la incorporación de la IA en la edición científica se orienta principalmente a la automatización progresiva de tareas técnicas, mediante validadores, diagnósticos y flujos reproducibles en un entorno caracterizado por la inmediatez de la circulación del conocimiento y el aumento de la obsolescencia de las publicaciones. Esta tensión, menos asociada a dilemas éticos directos y más a la infraestructura editorial, redefine el alcance técnico de la edición académica y abre un espacio de exploración sobre cómo producir y describir conocimiento en un ecosistema mediado por sistemas automatizados.

Finalmente, en el terreno de los formatos y los metadatos, la presión por la interoperabilidad y la cosecha automatizada se ha intensificado con la expansión de la ciencia abierta y los sistemas de indexación y descubrimiento basados en algoritmos [21]. La IA promete mejorar la calidad y consistencia de metadatos legibles por máquinas, pero también desplaza parte del poder de visibilización científica hacia procesos opacos de clasificación y ranking. Esta quinta tensión nos recuerda que no basta con producir conocimiento: hay que asegurarse de que la forma en que lo describimos, etiquetamos y distribuimos no profundice las desigualdades que ya existen entre regiones, lenguas y comunidades científicas.

Conclusión

Las tensiones que hemos planteado no son abstractas ni locales: están en el centro del debate internacional en torno a la integridad editorial. En conjunto, estas tensiones apuntan a una conclusión común: la cuestión central no es si usamos o no IA en la publicación científica, sino bajo qué principios lo hacemos. Para una revista como Iel, esto implica al menos tres compromisos: (i) reconocer explícitamente el uso de IA en la generación, revisión y edición de contenidos, promoviendo políticas claras de declaración para autores, revisores y editores; (ii) acotar su papel a funciones que no sustituyan el juicio crítico ni la responsabilidad intelectual de las personas; y (iii) invertir en alfabetización en IA en toda la comunidad editorial para comprender mejor sus alcances, sus sesgos y sus límites. Solo así, la IA podrá ser una aliada para fortalecer la calidad, la equidad y la relevancia social de lo que publicamos, en lugar de convertirse en un nuevo factor de opacidad en la comunicación científica.

On Artificial Intelligence and the Transformation of Scientific Publishing

Artificial intelligence (AI) has ceased to be a promise; it has transformed into an active active component of the editorial ecosystem. In light of the accelerated increase in the volume of publications, the pressure for visibility, and the

demand for efficiency in reviewing and editing processes, this technology has been increasingly integrated into the workflows of authors, reviewers, and editorial teams. Its adoption, however, is not without concern: it amplifies debates on scientific rigor, equity, transparency, and editorial responsibility.

In this editorial note, we set out to examine, from the experience of *Ingeniería e Investigación* (Iel), a diamond open-access journal, the tensions that emerge when AI is incorporated into the key processes of scientific publishing. Beyond evaluating specific technologies, we sought to understand the ethical and conceptual dilemmas they pose: To what extent do they facilitate the democratization of knowledge and when do they begin to homogenize, bias, or trivialize contents? What are the new responsibilities for editors and reviewers that emerge when automation intervenes in academic decisions?

Therefore, this text issues an invitation to reflect upon the role to be played by AI in contemporary scientific publishing: it should not be a redeeming promise nor an imminent threat, but a set of tools that demands clear criteria, critical literacy, and a renewed conception of editorial integrity.

The initial filter (desk review)

At Iel, the preliminary relevance review has evolved towards a model where AI is incorporated as a technical, non-decisional support in optimizing the desk review phase, i.e., thematic coherence assessment, format and language verification, and the initial detection of potential conflicts of interest. Recent studies [1] show that the use of language models can automate tasks such as preliminary reviews and style and consistency assessments, contributing to the reduction of the editorial burden and improving the normalization of acceptance criteria.

Nevertheless, the first tension lies in the fact that assimilating AI as a relevance filter demands a critical reflection on bias and exclusion risks [2]. When AI aids the assessment, confidentiality, epistemological diversity, and scientific integrity depend on human supervision [3]. Likewise, international guidelines such as UNESCO's *Recommendation on the ethics of artificial intelligence* [4] insist on the fact that human supervision should be maintained as a guiding principle in order to avoid unseen exclusions and preserve academic diversity.

Therefore, at Iel, AI acts as a 'collaborating agent' in editing processes, as an instrument to improve editorial management, operating under a hybrid, ethical, and reflective model. The final editorial decision is still human, deliberative, and evaluated by expert academic peers.

In line with these challenges, the Journal strives towards an editorial governance model by incorporating principles for the use of AI in preliminary reviews. These principles

include (i) transparency, by documenting the use of automated tools in each editorial decision; (ii) mandatory human supervision, guaranteeing that every AI-generated recommendation is verified by editors exercising academic judgement; (iii) ethical use, guaranteeing that AI is used as a technical assistant rather than as an autonomous acceptance or rejection mechanism.

Peer review

The deficit model constitutes a common practice in the peer review of manuscripts submitted for publication in scientific journals [5]. The premises of this model suggest that the observations made by peer reviewers correspond to failures or deficiencies in the document. These observations are taken as true, and authors must address them until all discrepancies have been eliminated. Furthermore, when considering comments as fair criticism, disagreements with reviewers may be seen as lack of comprehension by the authors, instead of sparking a justified academic discussion.

Under this logic, the integration of automated tools poses a risk of exacerbation. In the peer review process, the use of AI as an assistant may favor the detection of plagiarism and methodological errors [6]. However, if one relies excessively on the tool, one may neglect original contributions and innovation, causing the deficit model to mainly focus on the identification of failures in form, instead of on substantial deficiencies, where the reviewer's expertise is essential. On the other hand, when faced with the bias of the apparent depth of AI-generated manuscripts, as they employ technical language and circular argument with an academic tone, there is a risk that the reviewer may be misled if the review is superficial or hasty.

As a consequence, the automation of the peer review process may bolster predatory practices, essentially in relation to the impossibility of identifying low-quality manuscripts that have been transformed—or even generated—using AI, under a mechanical productivity that lacks rigor and integrity. The above creates an undesired loop wherein extreme automatism elaborates and endorses documents lacking research standards.

In AI-assisted peer reviews, the article by [7] on “the dangers of using AI in the peer review process” synthesizes concerns that have already been mentioned by COPE, JAMA, Science, and many editorial policies: the confidentiality of unpublished manuscripts, the loss of rigor and specificity in evaluations, the potential for a fraudulent representation of the role of AI, and the risk of the systematic manipulation of the peer review [8]. In addition, there is recent evidence of practices such as the insertion of hidden instructions in manuscripts, with the aim of biasing automated reviews, showcasing the extent to which the weaknesses of systems that replace the human judgement can be exploited [9]. At the same time, experiences such as AI Fast Track by NEJM AI indicate that some editors are exploring hybrid models in

which AI covers technical aspects while keeping academic decision-making in human hands [10]. This leads to our second tension: gaining efficiency cannot imply renouncing the plurality of perspectives, creativity, and the personal responsibility that sustains the peer review process.

On the author's voice and proofreading

Not only has AI profoundly transformed all editorial processes but has also completely permeated the creation of scientific documents. Nowadays, this type of tool supports authors in processes such as bibliographic management [11], grammar and form correction [12], and, in the most concerning cases, the conception of the article itself [13]. Its use has sparked all manner of debate, even threatening to automate the craft of proofreaders.

AI tools fill a fundamental need in today's world: they can generate grammatically correct texts in seconds, an essential aid within the framework of massified publishing practices. Nevertheless, although they feature some of the elements involved in the proofreading process, they are not language use reviewers. The reason for this is very simple: large language models (LLMs) do not have the same capabilities as humans when it comes to language, nor do they acquire it in the same way. These models process text based on huge corpora containing linguistic information, and they generate it while considering the occurrence (expressed in statistical terms) of one word after another. Unlike humans, they are not aware of what they say; they do not understand it. They do not relate concepts or express their reality through language. They are not capable of revising a text because they do not possess actionable knowledge of the topics addressed or of what is expressed therein.

However, and at the center of the discussion, there is the fact that AI, when used under human supervision and responsibility, is indeed useful for improving the quality, readability, and correctness of a text, as it facilitates searches and has a certain ‘clarity’ of grammar rules—especially those of the English language. This accelerates the writing process and, naturally, that of proofreading, which ends up focusing more on coherence, cohesion, and content than on aspects of grammar and syntax.

Nevertheless, even though these tools can aid scientists with limited time or language proficiency in publishing their research, the price ends up being the individual voice of each author. ChatGPT and DeepL have their own style—which proofreaders have learned to detect after years of interacting with them and their derivatives—, and their massive use homogenizes the language employed in the production of texts; we are on our way to a world where AI speaks *for* us, not *like* us.

Something similar occurs in the desk review process and style editing. Multiple journals have adopted policies that permit the use of AI for limited tasks, such as language use

correction or translation, but they forbid the generation of substantial content using these tools, as well as their recognition as authors, and they demand that their use be declared with transparency [14]. This confirms the third and fourth tension: AI can help to filter, homogenize formats, and polish texts, but, if its role is not expressly delimited, it may reinforce biases in the early selection of manuscripts, excessively standardize voices, and erode the traceability of intellectual contributions.

Graphical representation

The incorporation of AI tools in the visual output of scientific journal poses a relevant tension between vulgarization and democratization. On the one hand, there is a risk of developing images, infographics, or science communication texts that, although functional, present simplified (though not necessarily simple) visual resources which dim the conceptual complexity of some contents and generate a perceived vulgarization of knowledge. This perception is amplified when visual impact becomes more important than methodological depth, creating an imbalance between form and substance.

On the visual front, publishers such as Springer Nature and Elsevier have decided to restrict or prohibit the use of AI-generated images in scientific articles, given the legal and integrity risks it implies, even in cases where the peer review process does not detect explicitly problematic illustrations [15]. This discussion reinforces a relevant tension: AI can democratize graphic production for journals with limited resources, but it can also trivialize or distort scientific contents if visual impact is prioritized over conceptual precision.

Still, not acknowledging the democratizing potential of these technologies would be a mistake. In contexts where journals operate with limited budgets and under open-access models, AI allows creating graphic material and communication texts in an accessible, fast, and economic manner. Its use can facilitate the dissemination of complex findings to broader audiences, reduce the inequalities among journals with different resource levels, and strengthen the presence of publications in digital environments where images play a determining role in visibility.

The editorial challenge lies in preventing the ease of visual generation from leading to a proliferation of images that are generic or poorly nuanced in scientific terms, or to an acritical dependence on models that optimize for aesthetic coherence rather than seeking conceptual precision. Without clear protocols, the quality of AI-generated science communication images and texts may give rise to ambiguities, blurring the responsibilities of authors, editors, and reviewers while compromising the reliability of the published contents.

Therefore, the discussion should not focus on accepting or rejecting AI, but on establishing criteria that ensure a responsible use. Expert human supervision, the scientific verification of every visual material or communication text, and the definition of solid editorial standards are paramount to transform these tools into democratization resources, not into drivers of trivialization. In the end, the balance between access, clarity, and rigor will depend on the editorial policies adopted by each journal and on their ability to integrate these technologies without renouncing their scientific integrity.

Invisible infrastructure (formats and meta-data)

In the current context of scientific publishing, metadata and digital formats have ceased to be secondary components of the editorial process, consolidating as structural elements of academic communication. The growing demand for formats of HTML, XML-JATS, EPUB, and other alternatives oriented towards accessibility responds to both the requirements of indexing services and the diversification of devices and forms of access to knowledge.

In this scenario, AI is incorporated as a support for the creation of a single source (seed) from which multiple publication formats can be derived. This approach, explored from different workflows based on LaTeX, Pandoc [16], or XML-First [17, 18] schemes, allows optimizing editorial production by means of code-verified processes, wherein human supervision is concentrated on validating results and reviewing incidences reported by conversion systems, rather than on the reiterated manual correction of every version (as in traditional models), which generates, among others, a drastic reduction in costs and processing times.

On the other hand, the centrality of metadata becomes even more evident when considering their role in interoperability and in the automated discovery of scientific contents. Harvesting protocols, platforms such as DOAJ and Crossref [19], and academic search engines depend on normalized and consistent descriptions in order to integrate articles into their information systems. At this point, AI can contribute to both the quality and coherence verification of metadata and the analysis of their relevance for information indexation and retrieval processes, including the selection of titles, descriptors, and keywords with greater potential in terms of visibility.

However, the progressive massification of these tools poses a new editorial challenge: in an environment where all teams can optimize their metadata with AI support, differentiation will no longer depend solely on technique, but also on a deeper understanding of the way in which search algorithm and literature review AIs hierarchize and recommend scientific knowledge. From this perspective, scientific production is advancing towards a conception of the article as a computational object rather than as an isolated text [20]. The standardization of formats and the

enrichment of metadata allow contents to be harvested, linked, and reutilized in external infrastructures, facilitating their incorporation into machine-readable knowledge networks. Within this framework, the incorporation of AI into scientific editing is mainly oriented towards the progressive automation of technical tasks, through the use of validators, diagnostics, and reproducible workflows in an environment characterized by the immediate circulation of knowledge and the increasing obsolescence of publications. This tension, which is less associated with explicitly ethical issues and more linked to editorial infrastructure, redefines the technical scope of academic publishing and opens a space for exploring how knowledge can be produced and described within an ecosystem that is mediated by automated systems.

Finally, in the realm of formats and metadata, the pressure for interoperability and automated harvesting has intensified with the expansion of open science and algorithm-based indexation and discovery systems [21]. AI promises to improve the quality and consistency of machine-readable metadata, but it also displaces part of the scientific visibilization potential towards opaque classification and ranking processes. This fifth tension reminds us that producing knowledge is not enough; we must make sure that the way in which we describe, label, and distribute it does not deepen the existing inequalities among regions, languages, and scientific communities.

Conclusions

The tensions that we have established are neither abstract nor local: they are at the center of the international debate regarding editorial integrity. As a whole, these tensions indicate a common conclusion: the central issue is not whether we use AI in scientific publishing, but the principles under which we should do it. For a journal for IeI, this implies at least three commitments: (i) explicitly recognizing the use of AI in generating, reviewing, and editing contents, promoting clear disclosure policies for authors, reviewers, and editors; (ii) limiting their role to functions that do not substitute the critical judgement or the intellectual responsibility of people; and (iii) investing in AI literacy for the whole editorial community to better understand its scope, bias, and limits. Only thus will AI be an ally to strengthen the quality, equity, and social relevance of what we publish, instead of becoming a new factor of obscurity in scientific communication.

References

- [1] S. Ebadi, H. Nejadghanbar, A. R. Salman et al., "Exploring the impact of generative AI on peer review: insights from journal reviewers," *J. Acad. Ethics*, vol. 23, pp. 1383-1397, 2025. <https://doi.org/10.1007/s10805-025-09604-4>
- [2] D. Bhavsar, L. Duffy, H. Jo et al., "Policies on artificial intelligence chatbots among academic publishers: a cross-sectional audit," *Res. Integr. Peer Rev.*, vol. 10, pp. 1, 2025. <https://doi.org/10.1186/s41073-025-00158-y>
- [3] A. da Veiga, "Ethical guidelines for the use of generative artificial intelligence and artificial intelligence-assisted tools in scholarly journal publishing," *Sci. Ed.*, vol. 12, no. 1, pp. 28-34, 2025. <https://doi.org/10.6087/kcse.352>
- [4] UNESCO, "Recomendación sobre la ética de la inteligencia artificial," Documento SHS/BIO/PI/2021/1, París, Francia, 2021.
- [5] M. Reinhart and C. Schendzielorz, "Peer-review procedures as practice, decision, and governance—the road to theories of peer review," *Sci. Public Policy*, vol. 51, no. 3, pp. 543–552, 2024. <https://doi.org/10.1093/scipol/scad089>
- [6] Frontiers, "Unlocking AI's untapped potential: responsible innovation in research and publishing," 2025. [Online]. Available: <https://www.frontiersin.org/documents/unlocking-ai-potential.pdf>
- [7] C. A. de Souza, "Los peligros del uso de IA en la revisión por pares," 2025. [Online]. Available: <https://blog.scielo.org/es/2025/12/05/los-peligros-del-uso-de-ia-en-la-revision-por-pares/>
- [8] Committee on Publication Ethics (COPE), "Position statement on the use of artificial Intelligence tools," 2023. [Online]. Available: <https://publicationethics.org/cope-position-statements/ai-author>
- [9] L. Kugler, "Researchers Are Hiding AI Prompts in Their Papers," *Communications of the ACM*, Nov. 12, 2025. [Online]. Available: <https://cacm.acm.org/news/researchers-are-hiding-ai-prompts-in-their-papers/>
- [10] NEJM AI, "The NEJM AI Fast Track," *NEJM AI*, 2025. [Online]. Available: <https://ai.nejm.org/about/fast-track>
- [11] "Scientific literature reviews have never been this fun," researchrabbit.ai. <https://www.researchrabbit.ai/features> (accessed Dec. 7, 2025)
- [12] "You think big. We'll take care of the details." Grammarly.com. <https://www.grammarly.com> (Accessed Dec. 7, 2025)
- [13] "Your undetectable AI writer," author.com. <https://aithor.com> (Accessed Dec. 7, 2025)
- [14] Science Editorial Office, "Policy on the use of AI tools in manuscript preparation and peer review," American Association for the Advancement of Science (AAAS), 2023. [Online]. Available: <https://www.science.org/content/page/science-journals-editorial-policies>
- [15] N. Else, "Springer Nature and Elsevier address misuse of AI-generated images in scientific articles," *Nature News*, 2024. [Online]. Available: <https://www.nature.com/articles/d41586-024-00659-8>
- [16] A. Krewinkel and R. Winkler, "Formatting open science: agilely creating multiple document formats for academic manuscripts with Pandoc Scholar," *PeerJ Comp. Sci.*, vol. 3, art. e112, 2017. <https://doi.org/10.7717/peerj-cs.112>
- [17] M. Vaughn, M., "Automating JATS XML tagging with ChatGPT," . Minneapolis, MN, USA: Library Publishing Forum, Minneapolis, MN, 2024. <https://doi.org/10.7717/peerj-cs.112>
- [18] J. Arcila-Forero, "Uso de inteligencia artificial para asistir la creación de XML-JATS en revistas científicas," presented at VII Simp. Rev. Cien., Red Sara (virtual), 2024. <https://red-sara.org/actividad/vii-simposio-de-revistas-cientificas/>

- [19] DOAJ & and Crossref, "Best practices for metadata quality and machine readability," Joint Guidelines, 2024. [Online]. Available: <https://www.crossref.org/documentation/principles-practices/best-practices/>
- [20] M. Stocker, M., Snyder, L., Anfuso, M. *et al.*, "Rethinking the production and publication of machine-readable expressions of research findings," *Sci Data*, 12, 677, 2025. <https://doi.org/10.1038/s41597-025-04905-0>
- [21] Z. Lin, "Hidden Prompts in Manuscripts Exploit AI-Assisted Peer Review," arXiv, vol. 2507.06185, Jul. 2025. [Online]. Available: <https://arxiv.org/abs/2507.06185>