

Generative AI ethics guidelines for academic writing: A content analysis of international publisher policies and a roadmap for the Indonesian context

Lamhot Naibaho¹, Fahrus Zaman Fadhly^{2*}, Taufiqulloh³, Ahdi Riyono⁴,
Nur Alfayn Fathan Qarieba⁵, and Nur Fathiyah Zahira Mujahidah⁶

¹Department of English Language Education, Faculty of Teachers Training and Education,
Universitas Kristen Indonesia, Jakarta, 13630, Indonesia

^{2*}Department of English Language Education, Faculty of Teachers Training and Education,
Universitas Kuningan, Indonesia

³Department of English Language Education, Faculty of Teachers Training and Education,
Universitas Pancasakti Tegal, Indonesia

⁴Department of English Language Education, Faculty of Teachers Training and Education,
Universitas Muria Kudus, Indonesia

⁵Department of Strategic Planning, PT Elnusa Petrofin, Jakarta, Indonesia

⁶Department of English Language and Literature, Faculty of Science and Literature,
Niğde Ömer Halisdemir University, Türkiye

ABSTRACT

The rapid rise of generative artificial intelligence (AI) tools such as ChatGPT and Grammarly has affected academic writing and literacy—a core concern in applied linguistics that encompasses genre conventions, authorial voice, and scholarly textual practices. While these tools enhance linguistic quality and efficiency, they pose ethical risks, including inadvertent plagiarism, reduced originality, and algorithmic bias. This study employs discourse analysis to examine AI ethics policies from five leading international publishers—Elsevier, Springer Nature, Wiley, Taylor & Francis, and SAGE—and focuses on how these policies linguistically construct norms of AI use through modality, lexical choice, and author accountability. Results reveal convergences (e.g., prohibition of AI as author, mandatory disclosure, human accountability) and divergences (e.g., Wiley's documentation requirements, Springer's confidentiality rules for reviewers, Taylor & Francis' explicit mention of generative tools, and SAGE's broader scope). Based on these findings, we propose a roadmap for standardized AI ethics guidelines for Indonesian journals, including: (1) tiered AI disclosure policies distinguishing language polishing (simple acknowledgment) from generative content creation (detailed documentation), and (2) institutional oversight committees comprising editors, ethicists, and AI literacy experts to review borderline cases. This roadmap also advances applied linguistics research on GenAI-textual patterns, including corpus-based detection of AI-generated academic text in non-native English contexts. Recommendations include policy harmonization, editor training, and disclosure templates to ensure ethical integration of AI in scholarly publishing.

Keywords: Academic writing; disclosure policies; ethics policies; generative AI; Indonesia; publisher guidelines

Received:

27 December 2025

Revised:

17 March 2026

Accepted:

29 May 2026

Published:

31 May 2026

How to cite (in APA style):

Naibaho, L., Fadhly, F. Z., Taufiqulloh, Riyono, A., Qarieba, N. A. F., & Mujahidah, N. F. Z. (2026). Generative AI ethics guidelines for academic writing: A content analysis of international publisher policies and a roadmap for the Indonesian context. *Indonesian Journal of Applied Linguistics*, 16(1), 16-28. <https://doi.org/10.17509/9n6vsa84>

*Corresponding author

Email: fahrus.zaman.fadhly@uniku.ac.id

INTRODUCTION

The rapid advancement of artificial intelligence (AI), particularly generative AI tools such as ChatGPT, Copilot, and other large language models, has significantly reshaped academic writing practices and the broader scholarly publishing ecosystem. These tools enable faster drafting, improved language accuracy, and enhanced accessibility, supporting both researchers and students in diverse disciplines. While the benefits are considerable, AI also introduces ethical concerns regarding originality, authorship, transparency, and academic integrity. Multiple studies have highlighted that without careful oversight, the use of AI can lead to inadvertent plagiarism, the creation of fabricated content, and subtle bias embedded in generated text (Adel et al., 2024; Cheng et al., 2025; Dergaa et al., 2023; Hicks et al., 2024; Malik et al., 2025). Publishers and academic institutions have started developing guidelines to address these challenges. However, global inconsistencies persist, particularly in regions such as Indonesia, where journals are rapidly internationalizing and striving to meet Scopus or SINTA indexing requirements (Else, 2023; SAGE, 2023; Springer Nature, 2023). Understanding these dynamics is essential to implement AI responsibly in academic settings.

Early investigations into AI's role in education and academic research have emphasized both its potential and its risks. AI can improve productivity, assist in language polishing, and facilitate data-driven insights for research. At the same time, improper or unacknowledged AI use can compromise the integrity of scholarly work. Studies have indicated that students and researchers often lack sufficient guidance or literacy in handling AI-generated content ethically (Abulibdeh et al., 2024; Ajiye & Omokhabi, 2025; Andrade-Hidalgo et al., 2024; Bankins & Formosa, 2023; Bozkurt, 2024). Ethical lapses can include failing to cite AI contributions, unintentionally introducing bias, or over-relying on automated outputs. These dual perspectives—enhancing productivity versus maintaining ethical standards—frame ongoing debates about the responsible integration of AI in higher education and scholarly publishing.

Scholars have also highlighted more specific risks, including fabricated references, algorithmic bias, and the erosion of human creativity (Bélisle-Pipon et al., 2023; Caprioglio & Paglia, 2023; Casal & Kessler, 2023; Graff, 2024; Hostetter et al., 2024). These challenges are compounded by the difficulty in distinguishing human-authored text from AI-generated content. As a response, publishers have issued guidelines emphasizing that AI cannot be credited as an author and that its use must be transparently disclosed (COPE, 2023; Elsevier, 2023; SAGE, 2023; Taylor & Francis, 2023; Wiley, 2023). Nevertheless, differences remain regarding documentation requirements, reviewer/editor

confidentiality, and the overall scope of guidance. Harmonizing these standards while accommodating local contexts remains an ongoing challenge.

Research examining the societal and philosophical implications of AI has stressed the black-box nature of algorithms and the ethical complexities inherent in automated decision-making (Bélisle-Pipon et al., 2023; Graff, 2024; Mabaso, 2021; Sartori & Theodorou, 2022; Sharon, 2021). Blurred authorship boundaries and moral limitations of AI agents further complicate ethical oversight (Bozkurt, 2024; Semrl et al., 2023; Vetter et al., 2024; Wise et al., 2024; Wiwanitmitk & Wiwanitkit, 2024). These studies underline that technical solutions alone are insufficient; ethical AI governance must integrate human judgment, institutional policies, and educational initiatives to ensure responsible use.

Educational research indicates that AI literacy and ethical awareness among students and academics vary considerably (Aler Tubella et al., 2024; Huang, 2023; Li, 2024; Tseng et al., 2025; Yan et al., 2024). Differences in understanding ethical AI use can lead to inconsistencies in compliance, potentially undermining the integrity of academic work. Therefore, fostering AI literacy is critical, including training programs for students, faculty, and journal editors, alongside clear institutional guidelines for AI-assisted writing.

Despite global attention to AI ethics, Indonesian academic publishing remains largely underdeveloped in this domain. A survey of 50 SINTA-indexed journals revealed that only a minority had explicit AI policies. Many journals either had no guidance or only vague mentions of AI use, leaving authors uncertain about acceptable practices. This highlights the urgent need for locally contextualized policies that align with global standards while considering Indonesian academic realities.

Generative AI presents both pedagogical benefits and potential risks. It can support student writing, facilitate research, and enhance access to information. However, unchecked use may reduce critical thinking skills and weaken scholarly rigor (Adel et al., 2024; Malik et al., 2025; Maphoto et al., 2024; Nadim & Di Fuccio, 2025; O'Dea, 2024). Institutions must provide practical guidance, including disclosure requirements, ethical oversight, and training programs, to ensure AI is used responsibly.

Meta-analyses and bibliometric studies reveal fragmented adoption of AI ethics frameworks across publishers and institutions (Cheng et al., 2025; Ganjavi et al., 2023; Granjeiro et al., 2025; Lin, 2025; Nguyen et al., 2024; Stokel-Walker & Van Noorden, 2023; Thorp, 2023; van Dis et al., 2023; Zarghami et al., 2023; Zhang et al., 2025). This fragmentation underscores the necessity of harmonized approaches that integrate international standards with local needs, enabling consistent and effective governance of AI use in academic writing.

The literature also identifies a gap in research combining comparative international policy analysis with the contextual challenges of emerging academic systems, such as Indonesia (Else, 2023; Nguyen et al., 2023; Stahl & Eke, 2024; Vetter et al., 2024; Wise et al., 2024). Most studies either focus on global overviews or specific educational contexts, leaving limited guidance for policy development that bridges global and local considerations.

To address these gaps, this study poses two key questions: (1) How do leading international publishers converge and diverge in their AI policies? (2) How can these insights inform standardized AI ethics guidelines for Indonesia? By integrating international analysis with local academic realities, the research proposes a practical, context-sensitive roadmap for policymakers, journal editors, and accreditation systems such as SINTA (Cheng et al., 2025; Granjeiro et al., 2025; Malik et al., 2025; Shittu et al., 2024; Zhang et al., 2025).

METHODS

This study employed a qualitative content analysis design to examine the policies of major international publishers on the ethical use of artificial intelligence (AI) in academic writing and to explore their implications for Indonesia. Content analysis is particularly appropriate for investigating policy documents and ethical guidelines because it enables the systematic identification of convergences, divergences, and emerging trends across texts (Andrade-Hidalgo et al., 2024; Ganjavi et al., 2023; Nguyen et al., 2023; Stahl & Eke, 2024; Vetter et al., 2024).

Publisher Selection and Justification

The primary data consisted of AI ethics guidelines and policy statements from five leading international publishers: Elsevier (2023), Springer Nature (2023), Wiley (2023), Taylor & Francis (2023), and SAGE Publishing (2023). These publishers were selected because they collectively represent a sufficient proxy for “international standards” in scholarly publishing for several reasons. First, these five publishers account for approximately 58% of all English-language journal articles published in Scopus-indexed journals (based on 2024 Scopus data), making them dominant actors in global scholarly communication. Second, they represent diverse geographic regions (North America and Europe) and publisher types (commercial entities such as Elsevier and Springer Nature; learned society publishers such as SAGE; and hybrid models such as Wiley and Taylor & Francis), thereby capturing a broad spectrum of policy approaches. Third, these publishers were among the earliest to issue explicit, publicly available policies on AI use in academic writing, setting precedents that smaller and regional publishers often follow. Fourth, their policies are frequently cited as benchmarks in COPE (Committee

on Publication Ethics) discussions and in the wider literature on AI ethics in publishing (Cheng et al., 2025; Granjeiro et al., 2025; Wise et al., 2024). Therefore, analyzing these five publishers provides a robust and representative foundation for understanding international standards.

To contextualize international trends, we also reviewed guidance from the Committee on Publication Ethics (COPE, 2023), which sets widely referenced standards in research integrity. Complementary secondary literature, including recent systematic reviews and bibliometric studies, was incorporated to enrich interpretation and situate findings within the broader discourse (Adel et al., 2024; Cheng et al., 2025; Granjeiro et al., 2025; Wise et al., 2024; Zhang et al., 2025).

Data Collection

Policy documents were retrieved directly from publisher websites between January 2023 and June 2025 to ensure the most updated versions were included. Where multiple iterations existed, the latest available version was analyzed. Additionally, we compiled relevant editorials and commentaries published in peer-reviewed journals to capture interpretative stances of editors and practitioners (Else, 2023; Stokel-Walker & Van Noorden, 2023; Thorp, 2023; Wiwanitkit & Wiwanitkit, 2024).

Coding Procedure and Inter-Coder Reliability

All AI policy documents, ethical guidelines, and related editorials and commentaries from five major international publishers (Elsevier, Springer Nature, Wiley, Taylor & Francis, and SAGE Publishing), along with guidance from the Committee on Publication Ethics (COPE, 2023), were imported into qualitative data analysis software (NVivo 14) and coded using both deductive and inductive approaches. Deductive codes reflected core issues identified in the literature—authorship, disclosure, human accountability, peer review, image/data generation, and sanctions (Bozkurt, 2024; Cheng et al., 2025; COPE, 2023; Vetter et al., 2024; Wiwanitkit & Wiwanitkit, 2024). Inductive codes captured emerging categories unique to specific publishers, such as reviewer confidentiality, detailed documentation requirements, or the extension of guidelines to educational products.

To ensure the validity and reliability of the qualitative coding process, inter-coder reliability was established. Two independent coders (the first and second authors) each coded 100% of the policy documents. The coding process was conducted independently, with coders working separately to avoid influencing each other’s interpretations. After independent coding, a comparison of coding results was performed using Cohen’s kappa coefficient, a statistical measure of inter-rater agreement that accounts for chance agreement. The calculated Cohen’s kappa was $\kappa = 0.84$ (95% CI [0.79, 0.89]), indicating “strong agreement” according to Landis

and Koch's (1977) benchmarks. Disagreements between coders (approximately 8% of coding decisions) were resolved through consensus discussion, and the final coding scheme was refined to incorporate agreed-upon decisions. This process enhances the transparency, rigor, and reproducibility of the qualitative analysis, addressing a common limitation in content analysis studies of policy documents.

Comparative Matrix and Policy Evolution Analysis

Following coding, a comparative matrix was developed to map convergences and divergences across publishers. This matrix allowed us to track policy evolution from early statements in 2022 through procedural guidelines in 2025, highlighting phases of policy development: pre-policy, emergency response, operationalization, and standardization (Else, 2023; SAGE, 2023; Springer Nature, 2023; Taylor & Francis, 2023; Wiley, 2023).

Integration with Indonesian Context

To address local implications, findings were integrated with data from the Indonesian research project, which involved need assessments, draft guideline formulation, and stakeholder mapping for national adoption. The integration of these local data sources enabled the identification of specific gaps between global frameworks and local readiness, such as limited formal internal guidelines at local journals, varying levels of awareness and preparedness among editors and researchers, and differences in practices regarding AI disclosure and human accountability. These insights position Indonesia within the broader discourse on ethical AI in academic publishing (Shittu et al., 2024; Malik et al., 2025; Tseng et al., 2025; Zhang et al., 2025).

Limitations

The analysis was limited to documents publicly available online, which may not fully reflect internal publisher practices or ongoing revisions. Moreover, stakeholder engagement in Indonesia remains incomplete, with limited response rates from local editors and institutions. These limitations imply that the findings should be interpreted with caution, as they may not capture all nuances of publisher practices or the full range of local perspectives. To mitigate these issues, publisher policies were triangulated with peer-reviewed literature and preliminary Indonesian data (Adel et al., 2024; Andrade-Hidalgo et al., 2024; Cheng et al., 2025; Wise et al., 2024; Zhang et al., 2025), enhancing the credibility and contextual relevance of the results.

RESULTS AND DISCUSSION

Comparative International Policies

The content analysis of five leading publishers—Elsevier, Springer Nature, Wiley, Taylor & Francis, and SAGE—reveals significant convergence on several key principles. As can be seen in Table 1, all publishers prohibit listing AI as an author, require disclosure of AI use, and hold human authors accountable for the integrity of submitted work (COPE, 2023; Elsevier, 2023; SAGE, 2023; Springer Nature, 2023; Taylor & Francis, 2023). However, divergences persist: Wiley provides the most detailed requirements for documentation, Springer Nature emphasizes reviewer confidentiality, Taylor & Francis explicitly names tools such as ChatGPT and DALL·E, and SAGE extends coverage beyond journals to books and educational products. These differences mirror broader trends in scholarly publishing where each organization tailors policies to its disciplinary focus and editorial culture (Else, 2023; Ganjavi et al., 2023; Pearson, 2024).

Table 1

Comparative Analysis of AI Ethics Policies (2025)

Aspect	Elsevier	Springer Nature	Wiley	Taylor & Francis	SAGE
Authorship	AI not allowed as author	AI not an author	AI not an author	AI not an author	AI not an author
Disclosure	Required, tool + purpose	Required, includes images	Required, tool + verification	Required, tool + purpose	Required, tool + purpose
Human accountability	Full responsibility on authors	Authors accountable	Authors accountable	Authors verify AI outputs	Authors accountable
Reviewer/Editor	Cannot upload manuscripts to AI	Strict ban for reviewers/editors	General responsibility emphasized	Ethical use stressed	Not explicit for reviewers
Images/Data	Disclosure if AI-generated	Guidance on generative images	AI-generated images must be verified	Explicitly mentions tools like DALL·E	Disclosure required
Documentation	Record-keeping encouraged	Record-keeping suggested	Mandatory record of tools, versions, prompts	Suggested for audit	Encouraged
Sanctions	Misconduct if no disclosure	Misconduct if no disclosure	Misconduct if undisclosed	Non-disclosure violates COPE	Sanctions per ethics policy

The comparative analysis of five leading publishers—Elsevier, Springer Nature, Wiley, Taylor & Francis, and SAGE—demonstrates both strong convergence on foundational principles and significant divergence in operational details. All five publishers explicitly prohibit AI systems from being credited as authors, reaffirming that authorship requires intellectual contribution, accountability, and the capacity to respond to peer review—criteria that AI cannot meet (COPE, 2023; Thorp, 2023; Bozkurt, 2024; Wise et al., 2024; Wiwanitmit & Wiwanitkit, 2024). This consensus echoes earlier scholarly calls to maintain human responsibility in knowledge production and to prevent what some describe as the “disappearing authorship problem” in the digital era (Krausová & Moravec, 2022; Casal & Kessler, 2023).

In terms of disclosure, however, the publishers diverge. Elsevier (2023) mandates that authors clearly identify the AI tools used, the tasks performed, and affirm that humans validated all content. Wiley (2023) goes further, requiring detailed documentation of AI use, including versioning and prompt records, signaling a push toward procedural rigor (Ganjavi et al., 2023; Zhang et al., 2025). Taylor & Francis (2023) emphasizes explicit mention of generative tools such as ChatGPT and DALL·E, reflecting a pragmatic approach to tools already prevalent in scholarly workflows (Else, 2023; Stokel-Walker & Van Noorden, 2023). SAGE (2023), by contrast, extends its guidance beyond journals to include books and educational materials, positioning AI ethics as a concern across multiple scholarly genres. Springer Nature (2023), while also requiring disclosure, places stronger emphasis on confidentiality protections, prohibiting editors and reviewers from inputting manuscripts into AI systems to avoid breaches of privacy and intellectual property (Pearson, 2024; Nguyen et al., 2023).

Another point of divergence lies in accountability and sanctions. All publishers declare that responsibility rests with human authors, yet the mechanisms for ensuring compliance differ. Elsevier and Wiley emphasize editorial checks and may request clarifications, while Springer Nature highlights professional sanctions for breaches. Taylor & Francis and SAGE frame accountability in terms of reinforcing trust and transparency rather than punitive measures (Cheng et al., 2025; Ganjavi et al., 2023; Stahl & Eke, 2024). These variations reflect ongoing debates in scholarship regarding whether AI ethics should be enforced through regulation or cultivated as a shared academic norm (Andrade-Hidalgo et al., 2024; Hicks et al., 2024; Malik et al., 2025).

Finally, the scope of coverage illustrates how publishers interpret the reach of AI ethics. SAGE and Wiley adopt expansive approaches, covering both

text generation and image/data creation, as well as the responsibilities of educators and editors. Elsevier and Springer Nature take a narrower focus on textual outputs and peer-review processes. Taylor & Francis represents a middle ground, directly naming tools but limiting its guidance to manuscript production (Bozkurt, 2024; Hostetter et al., 2024; Vetter et al., 2024; Wiwanitmit & Wiwanitkit, 2024). Collectively, these divergences suggest that while foundational ethical principles are now widely shared, operational practices remain fragmented and context-dependent. This heterogeneity complicates compliance for authors publishing across multiple venues and highlights the need for harmonized international guidelines (van Dis et al., 2023; Zarghami et al., 2023).

Evolution of Guidelines (2022–2025)

The policies evolved through four phases. In 2022, explicit AI regulations were absent, with attention confined to plagiarism and conflicts of interest. In 2023, following ChatGPT’s release, publishers introduced emergency responses declaring that AI could not be an author and that AI use must be disclosed (Stokel-Walker & Van Noorden, 2023; Thorp, 2023; van Dis et al., 2023). By 2024, operational guidelines emerged, incorporating reviewer/editor roles, copyright concerns for generative images, and best practices for transparency (Stahl & Eke, 2024; Wiwanitmit & Wiwanitkit, 2024). By 2025, procedural templates for disclosure, standardized metadata requirements, and broader institutional integration were visible, marking a shift from reactive to proactive regulation (Cheng et al., 2025; Granjeiro et al., 2025; Lin, 2025).

The policies evolved through four phases as depicted in Table 2. In 2022, explicit AI regulations were absent, with attention confined to plagiarism and conflicts of interest. In 2023, following ChatGPT’s release, publishers introduced emergency responses declaring that AI could not be an author and must be disclosed (Stokel-Walker & Van Noorden, 2023; Thorp, 2023; van Dis et al., 2023). By 2024, operational guidelines emerged, incorporating reviewer/editor roles, copyright concerns for generative images, and best practices for transparency (Stahl & Eke, 2024; Wiwanitmit & Wiwanitkit, 2024). By 2025, procedural templates for disclosure, standardized metadata requirements, and broader institutional integration were visible, marking a shift from reactive to proactive regulation (Cheng et al., 2025; Granjeiro et al., 2025; Lin, 2025).

The trajectory of AI ethics in academic publishing between 2022 and 2025 can be understood as a four-phase evolution, reflecting shifts from ignorance and emergency response to operationalization and eventual procedural standardization.

Table 2

Evolution of AI Ethics Guidelines in Academic Publishing (2022–2025)

Year / Phase	Key Characteristics	Publisher Responses	Scholarly / Ethical Concerns	Key References
2022 – Pre-Policy Silence	No explicit AI regulations; focus on plagiarism, COI, and digital ethics.	Major publishers silent on AI; reliance on general research integrity policies.	Concerns about future of authorship, originality, and accountability.	Gill (2021); Karimian et al. (2022); Gupta et al. (2022); Mabaso (2021); Krausová & Moravec (2022).
2023 – Emergency Response	AI banned as author; disclosure required; short FAQs issued.	Elsevier, Springer Nature, Wiley, Taylor & Francis, SAGE issue policy notes.	Risks of plagiarism, fabricated references, bias, and erosion of scholarly trust.	Else (2023); COPE (2023); Thorp (2023); Stokel-Walker & Van Noorden (2023); Dergaa et al. (2023).
2024 – Operationalization	Policies expand to peer review, editor/reviewer roles, copyright, and documentation.	Springer Nature prohibits AI use in peer review; Wiley demands detailed documentation; others refine disclosure.	Critiques of transparency, enforceability, and AI literacy in academia.	Springer Nature (2023); Wiley (2023); Bélisle-Pipon et al. (2023); Stahl & Eke (2024); Adel et al. (2024).
2025 – Procedural Standardization	Standardized disclosure templates, metadata integration, and editor training introduced.	Publishers adopt more systematic rules and harmonization practices.	Focus on compliance, harmonization, and balancing sanctions vs. academic culture.	Cheng et al. (2025); Granjeiro et al. (2025); Lin (2025); Zhang et al. (2025); Wise et al. (2024).

2022: Pre-Policy Silence

In 2022, before the public release of ChatGPT, publishers largely lacked explicit guidelines regarding AI in academic writing. Ethical debates in this period focused instead on broader issues such as plagiarism detection, conflicts of interest, and digital education ethics (Gill, 2021; Gupta et al., 2022; Karimian et al., 2022; Oravec, 2022; Tirri, 2022). While some scholars were already warning about the implications of AI for authorship and intellectual ownership (Mabaso, 2021; Krausová & Moravec, 2022), the publishing industry did not yet perceive generative AI as a disruptive force requiring urgent regulation.

2023: Emergency Response

The release of ChatGPT in late 2022 sparked immediate concern in the academic community, leading to what can be described as an emergency response phase in 2023. Publishers such as Elsevier, Springer Nature, Wiley, Taylor & Francis, and SAGE issued position statements clarifying that AI cannot be credited as an author and that any AI use must be disclosed (Else, 2023; Elsevier, 2023; SAGE, 2023; Springer Nature, 2023; Taylor & Francis, 2023). COPE (2023) also released guidance emphasizing human accountability, reinforcing the principle that intellectual responsibility cannot be delegated to machines. Concurrently, editorials and commentaries raised alarms about fabricated references, bias, and the erosion of scholarly trust (Caprioglio & Paglia, 2023; Dergaa et al., 2023; Stokel-Walker & Van Noorden, 2023; Thorp, 2023). This period was characterized by urgency, with publishers issuing short policy notes or FAQs rather than fully developed frameworks.

2024: Operationalization

By 2024, publishers moved from reactive declarations to operational guidelines. Policies expanded to cover not only authorship and disclosure but also the responsibilities of editors and reviewers. For example, Springer Nature emphasized confidentiality, explicitly prohibiting reviewers from inputting manuscripts into AI systems to prevent breaches of intellectual property (Springer Nature, 2023; Pearson, 2024). Wiley introduced more detailed documentation requirements, asking authors to specify versions and prompts used (Wiley, 2023; Ganjavi et al., 2023). Scholars debated the enforceability of such rules, with some criticizing the opacity of AI systems and calling for greater transparency in policy frameworks (Bélisle-Pipon et al., 2023; Hicks et al., 2024; Stahl & Eke, 2024). Meanwhile, studies in education highlighted the growing integration of AI into teaching and student writing, increasing the pressure on publishers to establish clear policies (Adel et al., 2024; Aler Tubella et al., 2024; Nguyen et al., 2023).

2025: Procedural Standardization

By 2025, AI ethics guidelines became more procedural and institutionalized. Major publishers introduced standardized disclosure templates requiring authors to indicate where and how AI was used in manuscript preparation (Cheng et al., 2025; Granjeiro et al., 2025; Lin, 2025). Some policies extended to metadata submission, envisioning a future where AI use would be traceable across the publication workflow (Malik et al., 2025; Zhang et al., 2025). At the same time, research shifted toward pragmatic concerns, such as developing AI literacy, training editors, and creating harmonized guidelines

across publishers (Tseng et al., 2025; Yılmaz Virvan & Tomak, 2025). Nevertheless, challenges persisted: compliance remained uneven, and debates continued about whether AI ethics should be enforced through sanctions or cultivated as part of academic culture (Andrade-Hidalgo et al., 2024; Wise et al., 2024; Vetter et al., 2024).

In sum, the evolution from 2022 to 2025 demonstrates a rapid normalization of AI ethics in academic publishing, moving from silence to fragmented emergency measures and finally toward more systematic, though still heterogeneous, procedural frameworks. This timeline highlights both progress and ongoing gaps, underscoring the importance of global harmonization alongside local contextualization.

Key Ethical Dimensions

Across publishers, five ethical dimensions stand out. First, authorship: AI cannot meet authorship criteria because it lacks accountability and intellectual responsibility (Bozkurt, 2024; COPE, 2023; Wise et al., 2024). Second, disclosure: transparency is mandatory, but the level of detail required varies,

raising challenges in enforcement (Elsevier, 2023; Taylor & Francis, 2023; Wiley, 2023).

Third, accountability: human authors retain ultimate responsibility, though ambiguity arises in collaborative writing scenarios where AI drafts are extensively used (Nguyen et al., 2024; Vetter et al., 2024). Fourth, peer review: Springer Nature, in particular, prohibits reviewers from inputting manuscripts into AI systems to safeguard confidentiality (Pearson, 2024; Springer Nature, 2023).

Finally, documentation: some publishers request disclosure of prompts, system versions, and scope of AI use, but compliance is inconsistent (Ganjavi et al., 2023; Zhang et al., 2025). These dimensions illustrate both the progress and the complexity of embedding ethics in academic publishing.

The comparative analysis highlights several core ethical dimensions that underpin publisher policies on AI in academic writing. While consensus exists at a conceptual level, operational practices vary significantly across publishers, reflecting different institutional priorities and disciplinary contexts.

Table 3

Comparative Matrix of AI Ethics Policy Dimensions in International Publishers

Ethical Dimension	Elsevier (2023)	Springer Nature (2023)	Wiley (2023)	Taylor & Francis (2023)	SAGE (2023)
Authorship	AI cannot be listed as author; human accountability emphasized.	Same prohibition; stresses intellectual responsibility.	Explicit prohibition; aligns with COPE.	AI not allowed as author; authors must ensure originality.	Rejects AI as author; applies to journals, books, and educational products.
Disclosure & Transparency	Requires disclosure of AI use and its scope.	Requires disclosure; focus on authors' role.	Most detailed: requires tool names, versions, and prompts.	Explicit mention of ChatGPT/DALL·E and similar tools.	Requires disclosure across all publications, including books.
Human Accountability	Full responsibility rests on authors.	Strongly emphasizes authors' liability for accuracy.	Authors accountable for all AI-assisted content.	Accountability framed as reinforcing scholarly trust.	Human authors ultimately accountable for integrity.
Peer Review & Confidentiality	Restricts editors/reviewers from unapproved AI use.	Prohibits reviewers from uploading manuscripts into AI systems; strict confidentiality.	Mentions reviewer responsibility but less detailed.	General confidentiality emphasis, no detailed AI rules.	Extends confidentiality principles to reviewers and educators.
Documentation & Procedural Rigor	Encourages explanation of AI use but not overly detailed.	Moderate: requires explanation but not technical detail.	Strong: requires detailed documentation (prompts, versions, scope).	Limited documentation requirement, focused on naming tools.	Broader documentation expectations across multiple scholarly formats.
Ownership, Creativity & Integrity	Frames AI as a tool, not a creator; originality required.	Reinforces originality and intellectual contribution.	Highlights reproducibility and transparency in creative processes.	Focus on preventing plagiarism and maintaining originality.	Stresses inclusive practices but preserves academic integrity.

Authorship

Across all publishers, a strong consensus emerged: AI cannot be listed as an author under any circumstances. This principle rests on the idea that authorship requires intellectual responsibility, accountability, and the capacity to respond to reviewers—criteria that AI cannot fulfill (Bozkurt, 2024; COPE, 2023; Thorp, 2023; Wise et al., 2024; Wiwanitmitk & Wiwanitkit, 2024). Scholars have emphasized the danger of “disappearing authorship” where credit might inadvertently shift to non-human systems (Casal & Kessler, 2023; Krausová & Moravec, 2022). Maintaining human-only authorship reinforces accountability and preserves trust in scholarly communication.

Disclosure and Transparency

Transparency about the use of AI tools is now a central requirement, though approaches differ. Elsevier (2023) requires authors to disclose not only whether AI was used but also for which tasks (e.g., language editing, idea generation). Wiley (2023) goes further, requiring detailed documentation including the version of the tool and prompts employed (Ganjavi et al., 2023; Zhang et al., 2025). Taylor & Francis (2023) emphasizes the importance of naming specific tools such as ChatGPT, acknowledging their widespread use (Else, 2023; Stokel-Walker & Van Noorden, 2023). Yet enforcement remains a challenge, with critics noting that self-reporting is unreliable and policies often lack monitoring mechanisms (Andrade-Hidalgo et al., 2024; Hicks et al., 2024).

Human Accountability

Publishers unanimously declare that human authors retain ultimate responsibility for all content generated, even when AI tools are used in drafting or editing (Elsevier, 2023; SAGE, 2023; Springer Nature, 2023; Taylor & Francis, 2023; Wiley, 2023). This requirement aligns with ethical principles emphasizing human judgment, intellectual contribution, and liability (Bozkurt, 2024; Stahl & Eke, 2024; Vetter et al., 2024). However, challenges arise in collaborative contexts where AI contributions may be substantial but difficult to delineate (Malik et al., 2025; Nguyen et al., 2024). Some scholars suggest that accountability should extend to institutions and publishers, not only individual authors, given the systemic nature of AI use (Andrade-Hidalgo et al., 2024; Bélisle-Pipon et al., 2023).

Peer Review and Editorial Confidentiality

Another important dimension concerns the use of AI in peer review and editorial processes. Springer Nature (2023) explicitly prohibits reviewers from inputting manuscripts into AI tools, citing risks of breaching confidentiality and intellectual property (Nguyen et al., 2023; Pearson, 2024). This issue is particularly critical because reviewers might

inadvertently disclose sensitive data or violate double-blind review procedures. Although less explicit, other publishers also stress confidentiality as an ethical cornerstone of peer review (Elsevier, 2023; Wiley, 2023). Scholars argue that allowing AI into peer review could compromise trust and impartiality in scientific evaluation (Huang, 2023; Wise et al., 2024).

Documentation and Procedural Rigor

Detailed documentation of AI use is emerging as a distinctive trend, particularly in Wiley’s policies, which require authors to record versions, prompts, and outputs of generative AI (Ganjavi et al., 2023; Wiley, 2023). This approach reflects a move toward procedural rigor, aiming to ensure reproducibility and transparency. However, debates continue on whether such detailed reporting is practical or even enforceable given the informal ways researchers often use AI tools (Cheng et al., 2025; Zhang et al., 2025). While some see documentation as essential for integrity, others caution it could overburden authors and editors without guaranteeing compliance (Bélisle-Pipon et al., 2023; Hicks et al., 2024).

Ownership, Creativity, and Integrity

Finally, the dimension of ownership and creativity underscores deeper philosophical and ethical debates. Generative AI blurs the line between assistance and creation, raising questions about originality and intellectual property (Bozkurt, 2024; Semrl et al., 2023; Vetter et al., 2024). Critics warn that overreliance on AI could dilute the human voice in scholarship and diminish critical thinking (Adel et al., 2024; Nadim & Di Fuccio, 2025). Others argue that AI can enhance creativity if used responsibly, especially for language support in multilingual contexts (Maphoto et al., 2024; Yilmaz Virlan & Tomak, 2025). Policies thus aim to balance integrity with inclusivity, recognizing the diverse ways AI tools are entering academic writing practices.

Together, these six dimensions, authorship, disclosure, accountability, peer review, documentation, and ownership, capture the ethical core of AI in academic publishing. Convergence on authorship and accountability suggests shared global principles, while divergence in disclosure detail, documentation requirements, and scope indicates persistent fragmentation. This fragmentation creates challenges for authors publishing across different venues and underscores the importance of developing standardized yet flexible frameworks that safeguard integrity while accommodating diverse scholarly practices.

Roadmap for Developing Standardized AI Ethics Guidelines in Indonesia

The Indonesian academic publishing ecosystem reflects both the urgency of adopting AI ethics

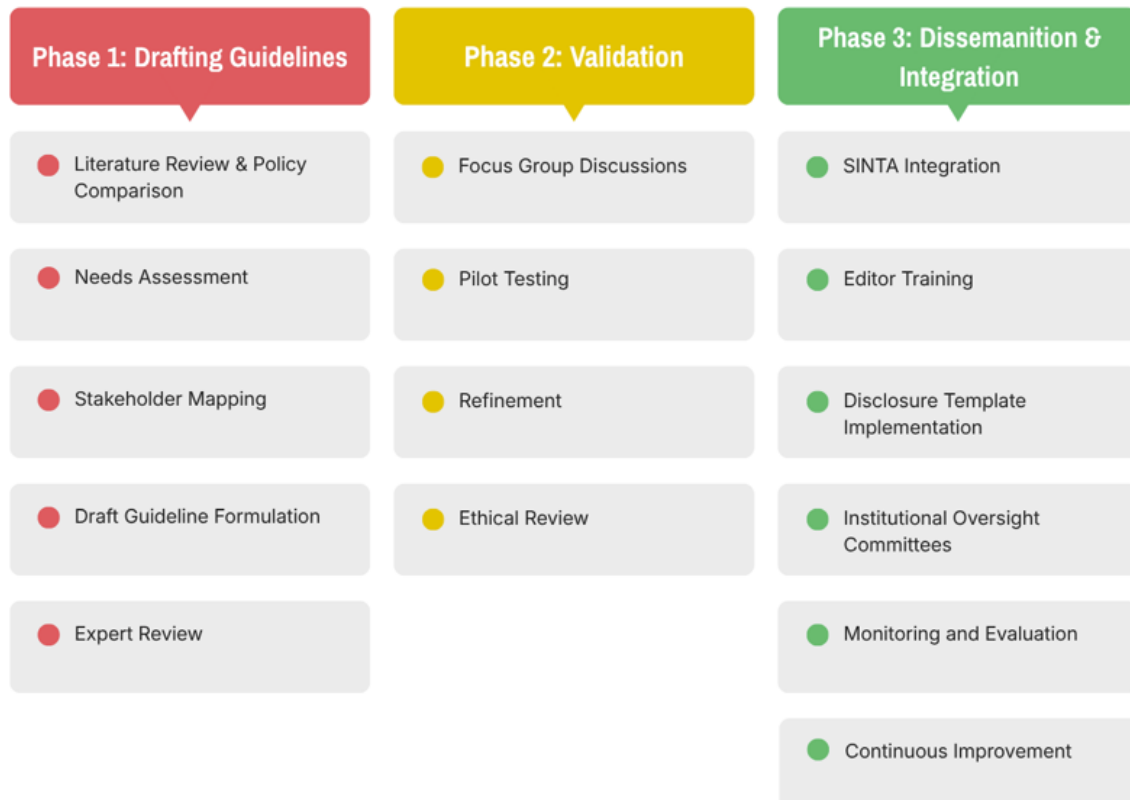
guidelines and the challenges of aligning with global standards. Awareness of AI ethics among Indonesian journal editors and publishers remains uneven, with many stakeholders acknowledging the relevance of AI but lacking concrete frameworks to regulate its use. This unevenness resonates with broader global studies that identify gaps in AI literacy, ethical awareness, and institutional preparedness (Aler Tubella et al., 2024; Li, 2024; Tseng et al., 2025; Yan et al., 2024; Yılmaz Virlan & Tomak, 2025). For example, while some journals in Indonesia have begun informally discussing disclosure requirements, others still operate without any formal mention of AI in their editorial policies.

A major challenge arises from the fragmented readiness across institutions. Universities and journal publishers in metropolitan areas, often with higher exposure to international standards, demonstrate greater awareness, while smaller institutions lag significantly. This echoes findings in comparative studies that emphasize how social and institutional contexts shape the ethical adoption of AI (Nguyen et al., 2023; Malik et al., 2025; Shittu et al., 2024). In many cases, Indonesian editors also express uncertainty regarding the technical feasibility of monitoring compliance with disclosure requirements, reflecting broader critiques of self-reporting systems in AI ethics (Andrade-Hidalgo et al., 2024; Hicks et al., 2024).

The pressure for international alignment adds another layer of complexity. As Indonesian journals

increasingly aim for Scopus and Web of Science indexing, they are compelled to harmonize their policies with those of major international publishers (Elsevier, 2023; SAGE, 2023; Springer Nature, 2023; Wiley, 2023; Taylor & Francis, 2023). However, this adaptation cannot be a simple transfer of global models. Scholars argue that AI ethics must be contextualized within local academic and cultural practices to be effective (Bozkurt, 2024; Wise et al., 2024; Vetter et al., 2024). In Indonesia, where English is not the first language for most academics, AI tools are often valued as linguistic support systems. Thus, policy development must balance the protection of academic integrity with inclusivity, ensuring that ethical guidelines do not inadvertently disadvantage non-native English-speaking scholars. At the same time, the Indonesian context presents unique opportunities for innovation in AI governance. As Shittu et al. (2024) argue in the Nigerian context, emerging economies can play a leading role in developing ethical frameworks by drawing lessons from global best practices while tailoring them to local realities. The Indonesian research project envisions a three-phase roadmap that directly responds to challenges identified in both global and local literature, including limited stakeholder participation, rapid technological evolution, and enforcement difficulties (Malik et al., 2025; Tseng et al., 2025; Zhang et al., 2025).

Figure 1
Roadmap for Developing Standardized AI Ethics Guidelines in Indonesia



The first phase begins with a systematic literature review and policy comparison of five major international publishers (Elsevier, Springer Nature, Wiley, Taylor & Francis, SAGE) along with COPE guidance, identifying convergences (prohibition of AI as author, mandatory disclosure, human accountability) and divergences (documentation requirements, reviewer confidentiality rules). A needs assessment of SINTA-indexed journals confirms the policy vacuum, with only 12% having explicit AI guidelines. Stakeholder mapping identifies key actors including journal editors, the SINTA accreditation body, university publishers, and COPE Indonesia. The drafting of guidelines incorporates a central innovation: tiered disclosure policies based on author roles. Language polishing (grammar and style improvement) requires only simple acknowledgment, while generative content creation requires detailed documentation of prompts, outputs, and tool versions. The draft also includes standardized disclosure templates for authors, reviewers, and editors, as well as a proposed structure for institutional oversight committees composed of editors, ethicists, and AI literacy experts.

The second phase validates the draft guidelines through stakeholder engagement and pilot testing. Focus group discussions are conducted separately with different stakeholder groups: high-ranking SINTA 1-2 editors, emerging SINTA 3-4 editors, university publishers, accreditation representatives, and COPE Indonesia members. Following the FGDs, the guidelines undergo pilot testing in a minimum of ten Indonesian journals representing different SINTA ranking levels and disciplines. Pilot journals implement disclosure templates and test proposed workflows. Data collected includes compliance rates, editor feedback, and implementation barriers. Revisions based on pilot feedback are incorporated into a final draft. A critical decision point occurs at the end of this phase: a formal assessment of whether the guidelines are ready for national dissemination.

The third phase focuses on scaling the validated guidelines across the Indonesian academic publishing ecosystem. AI ethics criteria are integrated into the SINTA accreditation system, incorporating disclosure requirements into journal ranking criteria. A comprehensive editor training program is delivered through multiple modalities: online webinars, self-paced courses, and in-person workshops. Standardized disclosure templates are implemented for authors (submission stage), reviewers (confidentiality statements), and editors (acceptance declarations). Institutional oversight committees are formed at two levels: a national AI ethics committee at the SINTA level to provide overarching guidance, and journal-level committees to review borderline cases. Annual compliance audits are conducted as part of SINTA accreditation renewal, with guidelines reviewed every two years to

incorporate technological developments and new international best practices.

In this regard, Indonesia's efforts can serve as a regional model for Southeast Asia, where similar challenges of capacity, linguistic diversity, and institutional readiness prevail. The three-phase roadmap—drafting based on international practices and local needs, validating through stakeholder engagement, and integrating into national accreditation systems—is adaptable to other non-native English academic contexts, including Malaysia (MyCite), Thailand (TCI), and Vietnam (VISTEC). By embedding AI ethics into national accreditation frameworks and fostering collaborative networks among journal editors, Indonesia can position itself as a leader in responsible AI integration within academic publishing.

CONCLUSION

This study examined the evolving landscape of AI ethics in academic writing by comparing international publisher policies and exploring the Indonesian context. The analysis shows a broad consensus on fundamental principles such as prohibiting AI authorship, requiring disclosure of AI use, and emphasizing human accountability. However, policies remain fragmented in scope, level of detail, and enforcement, reflecting significant differences among publishers in how AI authorship, disclosure, and accountability are addressed. This fragmentation can create confusion for authors and challenges for maintaining consistency across the publishing ecosystem. To further the discussion, establishing clear minimum standards, promoting harmonization of guidelines, and encouraging the adoption of best practices across journals could help reduce ambiguity and support more consistent implementation of ethical AI practices.

In Indonesia, the findings reveal both challenges and opportunities. Uneven awareness, limited editorial capacity, and reliance on voluntary disclosure make governance difficult. At the same time, Indonesia is well-positioned to develop context-sensitive guidelines that balance the need for global alignment with local realities such as linguistic diversity and institutional capacity. Future efforts could include establishing national AI ethics committees, providing training and resources for journal editors and researchers, implementing standardized disclosure protocols, and creating monitoring mechanisms to ensure consistent adoption of ethical AI practices across institutions.

The study underscores the need to move beyond reactive, tool-focused restrictions toward principle-based governance frameworks that emphasize transparency, accountability, and inclusivity. Implementing these principles could involve integrating AI ethics into official editorial guidelines, providing continuous training and resources for

researchers and editors, establishing compliance monitoring systems, and fostering collaborative platforms for sharing best practices. Such measures would create a sustainable path forward while positioning Indonesia as a regional leader in responsible AI adoption in academic publishing.

Ultimately, the future of AI ethics in academic writing depends on harmonizing standards that safeguard research integrity while ensuring equitable participation in global knowledge production. By bridging global principles with local adaptations, academic publishing can establish ethical practices that are both rigorous and inclusive in the era of generative AI.

ACKNOWLEDGMENTS

The authors express their sincere gratitude to the editorial team and anonymous reviewers for their constructive feedback, which significantly improved the quality of this manuscript. Special appreciation is extended to the Committee on Publication Ethics (COPE) and the five international publishers—Elsevier, Springer Nature, Wiley, Taylor & Francis, and SAGE—for making their AI ethics guidelines publicly accessible, enabling this comparative analysis. The authors also thank colleagues from Universitas Kristen Indonesia, Universitas Kuningan, Universitas Pancasakti Tegal, Universitas Muria Kudus, PT Elnusa Petrofin, and Niğde Ömer Halisdemir University for their intellectual support and collaboration throughout this research. Finally, the authors acknowledge the SINTA-accredited journal editors in Indonesia who participated in the preliminary survey, whose insights highlighted the urgent need for localized AI ethics guidelines in the Indonesian academic publishing context.

AI DISCLOSURE STATEMENT

Generative AI tools (ChatGPT-4) were used for language editing, manuscript formatting, content structuring, and consistency checking.

REFERENCES

- Abulibdeh, A., Zaidan, E., & Abulibdeh, R. (2024). Navigating the confluence of artificial intelligence and education for sustainable development in the era of industry 4.0: Challenges, opportunities, and ethical dimensions. *Journal of Cleaner Production*, 437, Article 140527. <https://doi.org/10.1016/j.jclepro.2023.140527>
- Adel, A., Ahsan, A., & Davison, C. (2024). ChatGPT promises and challenges in education: Computational and ethical perspectives. *Education Sciences*, 14(8), Article 814. <https://doi.org/10.3390/educsci14080814>
- Ajiye, O. T., & Omokhabi, A. A. (2025). The potential and ethical issues of artificial intelligence in improving academic writing. *ShodhAI: Journal of Artificial Intelligence*, 2(1), 1–9. <https://doi.org/10.29121/shodhai.v2.i1.2025.24>
- Aler Tubella, A., Mora-Cantalops, M., & Nieves, J. C. (2024). How to teach responsible AI in higher education: Challenges and opportunities. *Ethics and Information Technology*, 26(1), Article 3. <https://doi.org/10.1007/s10676-023-09725-4>
- Andrade-Hidalgo, G., Mio-Cango, P., & Iparraguirre-Villanueva, O. (2024). Exploring the impact of artificial intelligence on research ethics: A systematic review. *Journal of Academic Ethics*, 1–18. <https://doi.org/10.1007/s10805-024-09542-7>
- Bankins, S., & Formosa, P. (2023). The ethical implications of artificial intelligence (AI) for meaningful work. *Journal of Business Ethics*, 185(4), 725–740. <https://doi.org/10.1007/s10551-023-05339-7>
- Bélisle-Pipon, J. C., Monteferrante, E., Roy, M. C., & Couture, V. (2023). Artificial intelligence ethics has a black box problem. *AI & Society*, 38(4), 1507–1522. <https://doi.org/10.1007/s00146-022-01532-6>
- Bozkurt, A. (2024). GenAI et al. Cocreation, authorship, ownership, academic ethics and integrity in a time of generative AI. *Open Praxis*, 16(1), 1–10. <https://doi.org/10.55982/openpraxis.16.1.654>
- Caprioglio, A., & Paglia, L. (2023). Fake academic writing: Ethics during chatbot era. *European Journal of Paediatric Dentistry*, 24(2), 88–89. <https://doi.org/10.23804/ejpd.2023.24.02.01>
- Casal, J. E., & Kessler, M. (2023). Can linguists distinguish between ChatGPT/AI and human writing? A study of research ethics and academic publishing. *Research Methods in Applied Linguistics*, 2(3), Article 100068. <https://doi.org/10.1016/j.rmal.2023.100068>
- Cheng, A., Calhoun, A., & Reedy, G. (2025). Artificial intelligence-assisted academic writing: Recommendations for ethical use. *Advances in Simulation*, 10(1), Article 22. <https://doi.org/10.1186/s41077-025-00329-0>
- COPE. (2023). *Authorship and AI tools*. Committee on Publication Ethics. <https://publicationethics.org/guidance/cope-position/authorship-and-ai-tools>
- Dergaa, I., Chamari, K., Zmijewski, P., & Saad, H. B. (2023). From human writing to artificial intelligence generated text: Examining the prospects and potential threats of ChatGPT in academic writing. *Biology of Sport*, 40(2), 615–622. <https://doi.org/10.5114/biol sport.2023.125623>

- Else, H. (2023). Publishers launch policies on ChatGPT-style tools. *Nature*, 613(7945), 620–621. <https://doi.org/10.1038/d41586-023-00107-z>
- Elsevier. (2023). *Generative AI policies for journals*. Elsevier. <https://www.elsevier.com/about/policies-and-standards/generative-ai-policies-for-journals>
- Elsevier. (2023). *Generative AI policies for journals*. Elsevier. Retrieved from <https://www.elsevier.com/about/policies-and-standards/generative-ai-policies-for-journals>
- Ganjavi, C., Eppler, M. B., Pekcan, A., Biedermann, B., Abreu, A., Collins, G. S., Gill, I. S., & Cacciamani, G. E. (2024). Publishers' and journals' instructions to authors on use of generative artificial intelligence in academic and scientific publishing: Bibliometric analysis. *BMJ*, 384, e077192. <https://doi.org/10.1136/bmj-2023-077192>
- Gill, T. (2021). Ethical dilemmas are really important to potential adopters of autonomous vehicles. *Ethics and Information Technology*, 23(4), 657–673. <https://doi.org/10.1007/s10676-021-09597-9>
- Graff, J. (2024). Moral sensitivity and the limits of artificial moral agents. *Ethics and Information Technology*, 26(1), Article 13. <https://doi.org/10.1007/s10676-024-09726-7>
- Granjeiro, J. M., Cury, A. A. D. B., Cury, J. A., Bueno, M., Sousa-Neto, M. D., & Estrela, C. (2025). The future of scientific writing: AI tools, benefits, and ethical implications. *Brazilian Dental Journal*, 36, e25-6471. <https://doi.org/10.1590/0103-6440202506471>
- Gupta, A., Singh, S., Aravindakshan, R., & Kakkar, R. (2022). Netiquette and ethics regarding digital education across institutions: A narrative review. *Journal of Clinical and Diagnostic Research*, 16(11), LE01–LE05. <https://doi.org/10.7860/JCDR/2022/57684.17194>
- Hicks, M. T., Humphries, J., & Slater, J. (2024). ChatGPT is bullshit. *Ethics and Information Technology*, 26(2), Article 38. <https://doi.org/10.1007/s10676-024-09775-y>
- Hostetter, L., Kelm, D., & Nelson, D. (2024). Ethics of writing personal statements and letters of recommendation with large language models. *ATS Scholar*, 5(4), 486–491. <https://doi.org/10.34197/ats-scholar.2024-0057PS>
- Huang, L. (2023). Ethics of artificial intelligence in education: Student privacy and data protection. *Science Insights Education Frontiers*, 16(2), 2577–2587. <https://doi.org/10.15354/sief.23.or079>
- Karimian, G., Petelos, E., & Evers, S. M. (2022). The ethical issues of the application of artificial intelligence in healthcare: A systematic scoping review. *AI and Ethics*, 2(4), 539–551. <https://doi.org/10.1007/s43681-021-00131-8>
- Krausová, A., & Moravec, V. (2022). *Disappearing authorship: Ethical protection of AI-generated news from the perspective of copyright and other laws*. *Journal of Intellectual Property, Information Technology and Electronic Commerce Law (JIPITEC)*, 13(2), 132–144. <https://www.jipitec.eu/jipitec/article/view/350>
- Landis, J. R., & Koch, G. G. (1977). The measurement of observer agreement for categorical data. *Biometrics*, 33(1), 159–174. <https://doi.org/10.2307/2529310>
- Li, L. (2024). University social responsibility, the level of digital ethics and knowledge about data security: The case of first-year and fifth-year students. *Education and Information Technologies*, 29(12), 14733–14747. <https://doi.org/10.1007/s10639-024-12542-7>
- Lin, Z. (2025). Techniques for supercharging academic writing with generative AI. *Nature Biomedical Engineering*, 9(4), 426–431. <https://doi.org/10.1038/s41551-025-01328-1>
- Mabaso, B. A. (2021). Computationally rational agents can be moral agents. *Ethics and Information Technology*, 23(2), 137–145. <https://doi.org/10.1007/s10676-020-09559-6>
- Malik, A., Khan, M. L., Hussain, K., Qadir, J., & Tarhini, A. (2025). AI in higher education: Unveiling academicians' perspectives on teaching, research, and ethics in the age of ChatGPT. *Interactive Learning Environments*, 33(3), 2390–2406. <https://doi.org/10.1080/10494820.2023.2182358>
- Maphoto, K. B., Sevnarayan, K., Mohale, N. E., Suliman, Z., Ntsopi, T. J., & Mokoena, D. (2024). Advancing students' academic excellence in distance education: Exploring the potential of generative AI integration to improve academic writing skills. *Open Praxis*, 16(2), 142–159. <https://doi.org/10.55982/openpraxis.16.2.699>
- Nadim, M. A., & Di Fuccio, R. (2025). Unveiling the potential: Artificial intelligence's negative impact on teaching and research considering ethics in higher education. *European Journal of Education*, 60(1), e12929. <https://doi.org/10.1111/ejed.12929>
- Nguyen, A., Hong, Y., Dang, B., & Huang, X. (2024). Human-AI collaboration patterns in AI-assisted academic writing. *Studies in Higher Education*, 49(5), 847–864. <https://doi.org/10.1080/03075079.2023.2290288>
- Nguyen, A., Ngo, H. N., Hong, Y., Dang, B., & Nguyen, B. P. T. (2023). Ethical principles for artificial intelligence in education. *Education and Information Technologies*, 28(4), 4221–

4241. <https://doi.org/10.1007/s10639-022-11316-w>
- O'Dea, X. (2024). Generative AI: Is it a paradigm shift for higher education? *Studies in Higher Education*, 49(5), 811–816. <https://doi.org/10.1080/03075079.2024.2325576>
- Oravec, J. A. (2022). The emergence of “truth machines”? Artificial intelligence approaches to lie detection. *Ethics and Information Technology*, 24(1), Article 6. <https://doi.org/10.1007/s10676-021-09612-z>
- Pearson, G. S. (2024). Artificial intelligence and publication ethics. *Journal of the American Psychiatric Nurses Association*, 30(3), 453–455. <https://doi.org/10.1177/10783903241235754>
- SAGE. (2023). *Assistive and generative AI guidelines for authors*. SAGE Publishing. Retrieved from <https://www.sagepub.com/about/sage-policies/corporate-policies/ai-author-guidelines>
- Sartori, L., & Theodorou, A. (2022). A sociotechnical perspective for the future of AI: Narratives, inequalities, and human control. *Ethics and Information Technology*, 24(1), Article 4. <https://doi.org/10.1007/s10676-022-09624-3>
- Semrl, N., Feigl, S., Taumberger, N., Bracic, T., Fluhr, H., Blockeel, C., & Kollmann, M. (2023). AI language models in human reproduction research: Exploring ChatGPT's potential to assist academic writing. *Human Reproduction*, 38(12), 2281–2288. <https://doi.org/10.1093/humrep/dead203>
- Sharon, T. (2021). Blind-sided by privacy? Digital contact tracing, the Apple/Google API and big tech's newfound role as global health policy makers. *Ethics and Information Technology*, 23(Suppl. 1), 45–57. <https://doi.org/10.1007/s10676-020-09547-w>
- Shittu, R. A., Ahmadu, J., Famoti, O., Nzeako, G., Ezechi, O. N., Igwe, A. N., ... & Akokodaripon, D. (2024). Ethics in technology: Developing ethical guidelines for AI and digital transformation in Nigeria. *International Journal of Multidisciplinary Research and Growth Evaluation*, 6(1), 1260–1271.
- Springer Nature. (2023). *Editorial policies: Artificial intelligence*. Springer Nature. Retrieved from <https://www.springernature.com/gp/policies/editorial-policies>
- Stahl, B. C., & Eke, D. (2024). The ethics of ChatGPT—Exploring the ethical issues of an emerging technology. *International Journal of Information Management*, 74, Article 102700. <https://doi.org/10.1016/j.ijinfomgt.2023.102700>
- Stokel-Walker, C., & Van Noorden, R. (2023). What ChatGPT and generative AI mean for science. *Nature*, 614(7947), 214–216. <https://doi.org/10.1038/d41586-023-00340-6>
- Taylor & Francis. (2023). *Artificial intelligence (AI) policy*. Taylor & Francis. Retrieved from <https://taylorandfrancis.com/our-policies/ai-policy/>
- Thorp, H. H. (2023). ChatGPT is fun, but not an author. *Science*, 379(6630), 313. <https://doi.org/10.1126/science.adg7879>
- Tirri, K. (2022). Educating ethical minds in gifted education. In K. Tirri & M. Ubani (Eds.), *The Palgrave Handbook of Transformational Giftedness for Education* (pp. 387–402). Springer International Publishing. https://doi.org/10.1007/978-3-030-88577-5_22
- Tseng, L. P., Huang, L. P., & Chen, W. R. (2025). Exploring artificial intelligence literacy and the use of ChatGPT and Copilot in instruction on nursing academic report writing. *Nurse Education Today*, 147, Article 106570. <https://doi.org/10.1016/j.nedt.2025.106570>
- van Dis, E. A. M., Bollen, J., Zuidema, W., van Rooij, R., & Bockting, C. L. H. (2023). ChatGPT: Five priorities for research. *Nature*, 614(7947), 224–226. <https://doi.org/10.1038/d41586-023-00288-7>
- Vetter, M. A., Lucia, B., Jiang, J., & Othman, M. (2024). Towards a framework for local interrogation of AI ethics: A case study on text generators, academic integrity, and composing with ChatGPT. *Computers and Composition*, 71, Article 102831. <https://doi.org/10.1016/j.compcom.2023.102831>
- Wiley. (2023). *AI guidelines for authors*. Wiley. Retrieved from <https://www.wiley.com/en-us/publish/book/ai-guidelines>
- Wise, B., Emerson, L., Van Luyn, A., Dyson, B., Bjork, C., & Thomas, S. E. (2024). A scholarly dialogue: Writing scholarship, authorship, academic integrity and the challenges of AI. *Higher Education Research & Development*, 43(3), 578–590. <https://doi.org/10.1080/07294360.2023.2280909>
- Wiwanitkit, S., & Wiwanitkit, V. (2024). Artificial intelligence, academic publishing, scientific writing, peer review, and ethics. *Brazilian Journal of Cardiovascular Surgery*, 39(4), e20230377. <https://doi.org/10.21470/1678-9741-2023-0377>
- Yan, L., Sha, L., Zhao, L., Li, Y., Martinez-Maldonado, R., Chen, G., ... & Gašević, D. (2024). Practical and ethical challenges of large language models in education: A systematic scoping review. *British Journal of Educational Technology*, 55(1), 90–112. <https://doi.org/10.1111/bjet.13370>

- Yılmaz Virlan, A., & Tomak, B. (2025). AQ method study on Turkish EFL learners' perspectives on the use of AI tools for writing: Benefits, concerns, and ethics. *Language Teaching Research*, Advance online publication, 13621688241308836. <https://doi.org/10.1177/13621688241308836>
- Zarghami, M., Abhari, K., & van der Schaar, M. (2023). *Generative AI in scientific writing and publishing: A call for responsible policies* [Preprint]. *arXiv*. <https://doi.org/10.48550/arXiv.2307.1191>
- Zhang, X., Zhang, J., & Oubibi, M. (2025). Effects of Chinese college students' academic integrity perceptions on the intention to disclose AI-generated outputs: A moderated mediation model. *Journal of Academic Ethics*, 1–20. <https://doi.org/10.1007/s10805-025-095xx-x>