



Quest

ISSN: 0033-6297 (Print) 1543-2750 (Online) Journal homepage: www.tandfonline.com/journals/uqst20

Navigating Boundaries of AI Use in Physical Education Journal Research Publications

Yucen Li, Chenhao Wu, Yanlei Su, Beichen Lu, Xiaofen D. Hamilton, Timothy Lynch & Mark F. Hamilton

To cite this article: Yucen Li, Chenhao Wu, Yanlei Su, Beichen Lu, Xiaofen D. Hamilton, Timothy Lynch & Mark F. Hamilton (26 Jul 2026): Navigating Boundaries of AI Use in Physical Education Journal Research Publications, Quest, DOI: [10.1080/00336297.2026.2707164](https://doi.org/10.1080/00336297.2026.2707164)

To link to this article: <https://doi.org/10.1080/00336297.2026.2707164>



Published online: 26 Jul 2026.



Submit your article to this journal [↗](#)



View related articles [↗](#)



View Crossmark data [↗](#)



Navigating Boundaries of AI Use in Physical Education Journal Research Publications

Yucen Li ^a, Chenhao Wu^a, Yanlei Su^b, Beichen Lu^c, Xiaofen D. Hamilton ^a, Timothy Lynch ^{d,e}, and Mark F. Hamilton ^f

^aDepartment of Curriculum and Instruction, The University of Texas at Austin, Austin, TX, USA; ^bSchool of Business, The George Washington University, Washington, DC, USA; ^cCollege of Education and Human Ecology, The Ohio State University, Columbus, OH, USA; ^dDepartment of Education, Yew Chung College of Early Childhood Education, Hong Kong, China; ^eYew Chung Yew Wah Education Network, Hong Kong, China; ^fWalker Department of Mechanical Engineering, The University of Texas at Austin, Austin, TX, USA

ABSTRACT

This study investigates the status and regulatory nature of Artificial Intelligence (AI) usage policies within mainstream English-language physical education journals. Ten peer-reviewed physical education journals were included. Results indicated a top-down, publisher-driven approach to AI policy formulation; however, the widespread adoption of these policies across the broader landscape of physical education journals remains limited. While AI authorship was strictly banned, regulations on content generation varied. Moreover, a regulatory asymmetry emerged, characterized by stringent confidentiality protocols for reviewers alongside comparatively lax standards for data privacy and security on the author's side. Current governance relies on generic templates; thus, journals must transition to a human-centered framework establishing symmetric data protection and bias screening to safeguard integrity.

KEYWORDS

AI; regulation; physical education research journals

The rapid emergence of artificial intelligence (AI) tools has begun to reshape research practices across academic disciplines, including physical education (PE) (Caputo 2025). On one hand, AI technologies offer new possibilities for researchers by supporting literature synthesis, data management, and scholarly writing, raising important questions related to authorship, accountability, transparency, and ethical responsibility (Hosseini et al., 2023). On the other hand, AI has the potential to substantially transform or even partially replace traditional peer review processes, prompting concerns about the quality, reliability, privacy, and integrity of scholarly evaluation (Gatrell et al., 2024). In response to these developments, academic publishers and journals have increasingly introduced policies to regulate AI use in research, publication, and peer review across many fields (Gatrell et al., 2024). For PE research, which frequently involves human participants, qualitative inquiry, and pedagogical implications, these developments present both opportunities and challenges that warrant careful examination (Latzel & Glauner, 2024).

Despite the growing presence of AI-related publication policies, considerable uncertainty remains regarding how AI use is defined, constrained, and permitted across PE journals. While some journals provide detailed guidance on acceptable AI-assisted practices, others

rely on general or ambiguous statements, leaving authors, reviewers, and editors to interpret policy boundaries independently. This lack of clarity may contribute to inconsistent practices, uneven disclosure, and confusion across the publication process, underscoring the need for systematic analysis of how AI policies are articulated and operationalized within PE research publications.

Ethical dilemmas of AI in academic publishing

The convenience and efficiency AI brings to academic publishing are accompanied by numerous controversies. The following three issues are particularly noteworthy. First is the crisis of accuracy and academic integrity. Multiple studies warn that AI hallucinations can lead to fabricated references and erroneous data interpretation, severely damaging scientific credibility if used without rigorous verification (Athaluri et al., 2023; Sun et al., 2024). Furthermore, researchers may unwittingly engage in unconscious plagiarism when AI-generated text closely mirrors original content from its training databases (Kim, 2024).

The second one is that the proliferation of cloud-based AI tools poses significant confidentiality and intellectual property risks. Uploading unpublished work to third-party AI systems for polishing effectively exposes intellectual property to commercial entities (Hosseini et al., 2023). Scholars emphasize that such actions violate confidentiality agreements and risk the learning or leakage of original ideas before publication, triggering serious legal disputes (Mollaki, 2024; Thelwall, 2024).

The third issue is related to the detriment of researchers' autonomy and judgment by fostering increasing dependence on its convenient and often persuasive features. As AI systems become more integrated into the entire research process such as supporting tasks defining research questions, choosing research design, selecting conceptual framework, synthesizing literature, analyzing data and polishing the academic writing (Latzel & Glauner, 2024; Springmier, 2025), there is a risk that researchers may gradually defer to AI-generated suggestions without sufficient critical evaluation (Gendron et al., 2022). Researchers have pointed out that this reliance can subtly reshape decision-making processes, narrowing the space for independent reasoning and scholarly creativity (Gendron et al., 2022; Mollaki, 2024; Springmier, 2025). Such dependence may weaken researchers' capacity to formulate original questions, challenge prevailing assumptions, and exercise epistemic agency. In this way, AI not only influences how knowledge is produced but may also constrain the direction of inquiry itself, ultimately reducing researchers' control over their own epistemological agendas (Guelmami et al., 2024; Mollaki, 2024). As warned by Gendron et al. (2022), "the future of academic publishing as a meaningful human activity is at stake" (*p.* 1) due to the use of AI in research.

Previous studies on AI policies in academic journals

Given that no previous studies have examined AI policies for PE research journals, the brief literature review will focus on studies of the topic published in other fields to shed lights on PE. It has been found that the development and implementation of AI policies across academic journals in various fields exhibit significant disciplinary disparities. Lund and Naheem (2024) indicated that natural science journals such as physics, biology, chemistry, and medicine responded most swiftly, with 71.3% establishing guidelines. This rapid

adoption is similarly evident in the medical field, where nearly 80% of journals have established specific principles (Alkhawam et al., 2025). In contrast, fields like the social sciences and humanities have lagged behind. Multiple studies reveal a lower adoption rate in these disciplines (Alkhawam et al., 2025; Kim, 2024; Resnik & Hosseini, 2025). Specifically, Lund and Naheem (2024) reported an adoption rate of 46.9%, and Kim (2024) noted that over half of social science journals have yet to issue clear guidelines. This divergence stems partly from the social sciences' emphasis on discourse originality and theoretical deduction, leading to greater caution in policy formulation (Resnik & Hosseini, 2025).

It is important to acknowledge that publishers are shifting from technological blockade to mandatory transparency (Thorp, 2023). Such emerging consensus prioritizes risk management through detailed disclosure mechanisms rather than absolute bans, thereby shifting ethical accountability back to human authors (Resnik & Hosseini, 2025). This shift closely aligns with a broader theoretical trend as a growing body of research has emphasized the concept of human-centered artificial intelligence (HCAI) governance in recent years (Ozmen Garibay et al., 2023; Shneiderman, 2020). HCAI refers to a design framework that combines high levels of human control with high levels of computer automation to create reliable, safe, and trustworthy systems that amplify, augment, and empower human performance (Ozmen Garibay et al., 2023; Shneiderman, 2020). Specifically, HCAI rests on three core propositions: augmenting rather than replacing human capabilities, preserving human agency in decision-making, and upholding ethical standards that support long-term societal well-being (Albannai & Raziq, 2026; Schmager et al., 2025). Some studies suggest a notable divergence in how journals approach the permissible scope of substantive AI use (Ganjavi et al., 2024; Resnik & Hosseini, 2025). This regulatory fragmentation is similarly evident in the broader field of education, which closely aligns with the pedagogical dimensions of PE. For instance, a recent review by Hsu et al. (2025) highlighted considerable variability and a lack of standardized ethical guidelines regarding AI-assisted writing across education journals. Meanwhile, there is also AI policy research in sports science, kinesiology, and public health. A study reported by Guelmami et al. (2024) indicated that an international panel of experts identified 15 key ethical challenges, including AI misuse, data falsification, and plagiarism, and proposed a comprehensive set of guidelines to uphold research integrity in the sports medicine and exercise science field. An analysis of 108 sport-science journals by Caputo et al. (2025) found that 77.8% had established AI policies, and the prevalence of these policies was positively correlated with journal ranking. Furthermore, researchers show that the policies of sport-science and kinesiology journals primarily emphasize on transparency, human accountability, and clearly defined ethical boundaries of AI use (Caputo 2025; Mangum & Kuenze, 2025).

In summary, the identified key themes emerging from the above literature review in are transparency in AI use, human accountability, authorship boundaries, and ethical considerations as well as points of divergence across disciplines and journals, including differences in the level of policy specificity, enforcement mechanisms, and openness to AI-assisted writing. These findings help establish a clearer comparative backdrop for the present study, enabling a more meaningful interpretation of how AI policies in PE journals align with or differ from those in other fields. The current divergence in policy implementation, where natural science and medical journals are more formally regulated, while those in social sciences and humanities remain more

cautious and less developed, also exerts a notable influence on the PE domain. PE is inherently a highly interdisciplinary field, encompassing both the natural science attributes of students' physiological indicators and movement performance, and the social science attributes focusing on behavioral changes, social interactions and educational equity (Ekberg, 2021; Larsson, 2023). The significant divergence in AI policies observed across other disciplines not only accurately reflects the inherent complexities PE journals inevitably face when establishing unified AI regulatory standards, but also profoundly explains why current AI governance in this field remains in a fragmented and exploratory stage.

As noted earlier, no study has systematically analyzed AI-related policy documents across PE research journals. AI policies adopted by journals in adjacent fields cannot be assumed to apply directly to PE journals and a focused analysis is necessary to capture the discipline-specific landscape of AI policy, ultimately informing more relevant and effective guidelines for researchers and practitioners in PE. PE is a field centered on embodied practice and involves students' performance, movement contexts, and social interactions (Volshøj et al., 2025). These settings often include sensitive personal data, such as students' health information and psychological states, meaning that the use of AI may introduce risks related to privacy breaches and data governance (Klimova & Pikhart, 2025). These factors are not always addressed in broader or adjacent disciplines. For instance, general education does not typically focus on movement-based contexts, while health sciences often lack a strong pedagogical component. Therefore, research designs, data types, and professional standards can differ substantially from adjacent disciplines. Additionally, PE journals may vary in their adoption of AI policies due to differing editorial priorities and the discipline's highly interdisciplinary nature. Without examining PE-specific journals, it is difficult to determine whether existing policies adequately address these distinct realities or if there are gaps, inconsistencies, or emerging needs.

The purposes and significance of the current study

The purpose of this study is to examine how AI use is regulated in PE research journals by conducting a thematic analysis of AI-related publication policies. Specifically, this study seeks to identify key themes and points of divergence in how journals define permissible and restricted AI use, address issues of authorship and peer review, and responsibility, and promote transparency and ethical research practices. The secondary purpose is to synthesize in what ways AI and data-driven analytics could alter the evolution of research publishing in PE. By mapping the current policy landscape, this analysis aims to provide greater clarity for PE researchers navigating AI-related publication requirements and to contribute to ongoing discussions about the responsible and informed integration of AI in scholarly research.

While the usage of AI in PE research is inevitable, systematic analysis of AI policies within the field of PE remains scarce. This lack of clarity exposes scholars to ethical risks and compliance uncertainties. Consequently, this study investigates AI regulations across mainstream PE journals to analyze core policy elements and limitations. The findings aim to inform the development of consistent, ethically grounded guidelines for the PE community.

Methods

Criterion for inclusion of AI policy documents

To identify relevant AI usage policies within the field of PE, this study established two inclusion criteria: (a) the journals must be peer-reviewed, English-language publications with an academic scope explicitly focused on physical education, ensuring they represent the field’s academic standards; and (b) the journals or their affiliated publishers must have publicly accessible AI use policy, author guidelines, and/or publication ethics documents published on their official websites.

Data collection

Prior to data extraction, a rigorous screening process was conducted to establish the research sample across four major database platforms: Web of Science, Scopus, the Education Resources Information Center (ERIC), and SCImago. To ensure that the selected publication venues were fundamentally rooted in pedagogical and education-based contexts, we first restricted our search scope to the specific subject category of “Education & Educational Research.” Within this established category, we applied the exact term “Physical Education” as a precise filter. This retrieval process yielded nine journals from Web of Science, five from Scopus, six from ERIC, and 15 from SCImago. After the removal of duplicates, 27 journals were identified. Following a further manual screening, 17 of the 27 journals without the explicitly articulated AI policies were excluded. The AI policies of the remaining 10 journals were included for the final analysis (see Figure 1).

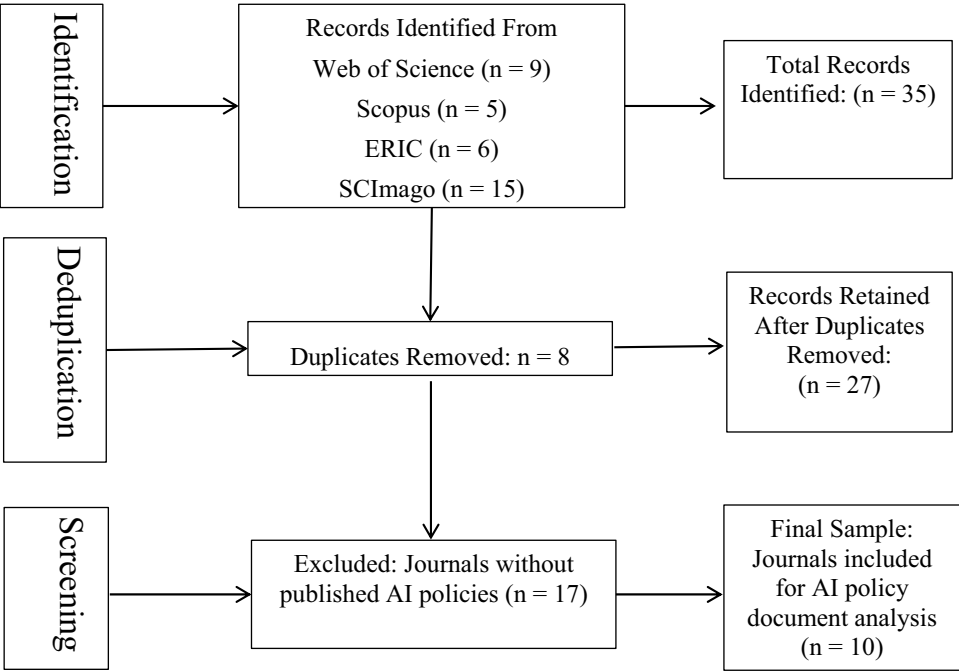


Figure 1. Flowchart of the journal selection and screening process.

Two steps were taken to ensure the reliability and validity of the data collection. First, the research team downloaded the latest versions of policy documents directly from each publisher's official website in December 2025. For major publishers such as Taylor & Francis and Sage, dedicated "AI Policy" were retrieved; for independent journals, relevant sections within the "Instructions for Authors" were extracted. Second, to ensure data completeness, the study collected not only specific ethical statements but also general submission guidelines to capture the full regulatory context. All documents were retained in their original English versions. A coding sheet was developed based on previous research on the topic in other fields (Alkhawam et al., 2025; Veiga, 2025; Weber-Wulff et al., 2023). Then all the collected documents were coded using the coding sheet.

Data analysis

This study employed document analysis to examine data from the selected policy documents. The analysis process consisted of three phases: skimming (superficial examination), reading (thorough examination), and interpretation (Bowen, 2009). This process has been used in studies in the field of PE (Yang et al., 2025). In the first phase, the research team first determine the top 10 journals published in English focusing on PE based on the journal's mission statement and its impact factor. Then all relevant materials, including author guidelines, ethics policies, and specific AI statements were skimmed by the research team, to determine whether each journal possessed a specific AI policy statement. The second phase involved a thorough, concentrated re-reading and review of the data to extract core themes across two distinct dimensions: author-oriented and reviewer-oriented regulations. Specifically, all policy information was coded using a structured code sheet containing the following variables (see Table 1). Author-side regulation variables included (see Table 2): (a) Authorship accountability, assessing whether policies explicitly addressed AI eligibility for authorship; (b) Permissible scope of use, categorized into language refinement and substantive content generation; (c) Prohibited applications, specifying restrictions on image creation and data manipulation or fabrication; (d) Risk mitigation strategies, coding for the implementation of author self-checks for accuracy and bias, the use of automated technical detection tools, and requirements for editorial intervention; and (e) Transparency mechanisms, identifying general disclosure requirements for conditional AI use. Reviewer-side regulation variables included (see Table 3): (a) Confidentiality protocols, coding for mandates prohibiting the uploading of unpublished manuscripts or sensitive data into AI systems; (b) Evaluative autonomy, identifying restrictions on using AI for manuscript evaluation or decision-making; and (c) Linguistic assistance, specifying whether reviewers are permitted to use AI for polishing review reports while maintaining liability for the accuracy of their feedback.

Data trustworthiness

Data trustworthiness was established through investigator triangulation, which involves using multiple researchers to code the same data to reduce potential bias (Patton, 2015). Specifically, two researchers coded the policy documents independently. Upon completion, the research team met to discuss discrepancies until the team reached a consensus. Additionally, the study utilized the method of document analysis verification by cross-

Table 1. Coding sheet of AI policy in PE field.

Journal	Publisher	Category Quartile	Impact Factor	Journal Citation	A1	A2	A3	A4	A5	A6	A7	A8	A9	R1	R2	R3
PESP	Taylor & Francis	Q1	3	SSCI	0	1	2	3	3	1	1	0	1	1	1	1
CSHPE	Taylor & Francis	Q1	2.1	ESCI	0	1	2	3	3	1	1	0	1	1	1	1
MPEES	Taylor & Francis	Q2	1.9	SSCI	0	1	2	3	3	1	1	0	1	1	1	1
JPERD	Taylor & Francis	Q3	0.29	ESCI	0	1	2	3	3	1	1	0	1	1	1	1
SJPSE	Taylor & Francis	Q4	N/A		0	1	2	3	3	1	1	0	1	1	1	1
JTPE	Human Kinetics Journals	Q2	1.8	SSCI	0	1	3	3	3	1	0	0	0	1	1	2
EPER	Sage journal	Q1	3.3	SSCI	0	1	2	2	2	1	0	1	1	1	1	1
JPES	Editura Universitatea din Pitesti	Q3	0.38		0	2	2	2	2	0	0	0	1	0	0	0
IJPEFS	Asian Research Association	N/A	N/A		0	1	3	3	3	1	0	0	0	1	0	0
PETM	OVS LLC	Q3	0.32		0	1	3	3	3	1	0	1	0	0	0	0

Note: SSCI = Social Sciences Citation Index; ESCI = Emerging Sources Citation Index; PESP= *Physical Education and Sport Pedagogy*; CSHPE= *Curriculum Studies in Health and Physical Education*; MPEES= *Measurement in Physical Education and Exercise Science*; JPERD= *Journal of Physical Education, Recreation and Dance*; SJPSE= *Strategies: A Journal for Physical and Sport Educators*; JTPE= *Journal of Teaching in Physical Education*; EPER= *European Physical Education Review*; JPES= *Journal of Physical Education and Sport*; IJPEFS= *International Journal of Physical Education, Fitness and Sports*; PETM= *Physical Education Theory and Methodology*.

Table 2. Coding book of author-side regulation variables.

Variable	Codes
A1: Authorship Accountability	0 = Not mentioned 1 = Prohibited 2 = Allowed
A2: ermissible Scope: Language Refinement	0 = Not mentioned 1 = Allowed
A3. Permissible Scope: Substantive Content Generation	2 = Conditionally Allowed (Required disclosure)
A4. Prohibited Applications: Image Generation	3 = Prohibited
A5. Prohibited Applications: Data Manipulation/Generation	0 = Not explicitly required, 1 = Required
A6. Risk Mitigation: Author Self-check	0 = Not used / Not mentioned, 1 = Used
A7. Risk Mitigation: Automated Technical Tools	0 = Not mentioned, 1 = Required
A8. Risk Mitigation: Editorial Intervention	0 = Not mentioned, 1 = Disclosure Required
A9. Transparency Mechanisms	

Table 3. Coding book of reviewer-side regulation variables.

Variable	Codes
R1. Confidentiality: Uploading Unpublished Manuscripts	0 = Not mentioned, 1 = Prohibited
R2. Evaluative Autonomy: Use AI for Decision Making	
R3. Linguistic Assistance: Language Polish for Reviews	0 = Not mentioned, 1 = Allowed, 2 = Prohibited

referencing findings against external industry standards (COPE & UNESCO guidelines) to evaluate the completeness and rigor of the journal policies. Operationally, the research team utilized the COPE position statement on AI tools and the core tenets of the UNESCO Recommendation on the Ethics of AI as objective external benchmarks. Subsequently, the research team systematically mapped and compared the coded dimensions of the PE journal policies against these benchmarks to examine the alignment of existing policies with

international ethical consensus and to precisely analyze regulatory gaps in the discussion section. This cross-verification process ensured that the research team’s evaluation regarding the completeness and rigor of the journal policies was grounded in internationally recognized ethical frameworks rather than subjective assumptions, thereby minimizing subjective bias in the qualitative content analysis to the greatest extent possible.

Results

Characteristics of included journals

The study ultimately retrieved 27 major PE journals published in English, of which only 10 had established AI usage policies or guidelines. As shown in Table 4, among these 10 journals, seven derived their policies directly from their respective publishers while one journal (i.e., *International Journal of Physical Education, Fitness and Sports*) adopts its AI policy from its affiliated Association. Conversely, there are two journals (i.e., *Journal of Physical Education and Sport*, and *Physical Education Theory and Methodology*) which created their own independent AI policies instead of using the one developed by its publisher.

Thematic AI policy for authors

Ethical principles and scope of use

There is a universal consensus regarding the legitimacy of linguistic assistance. Policies from all 10 journals explicitly permit the use of AI tools for language editing. This unanimity reflects a pragmatic recognition that AI can democratize access to publishing by reducing language composition barriers for all researchers. However, this permission is strictly bounded. Policies emphasize that such tools must function solely as copyeditors to improve clarity and readability, without altering the original scientific meaning or theoretical framework of the text.

In contrast to the uniformity on language editing, the policies regarding substantive content generation (e.g., idea generation, and literature synthesis) reveal a significant divide. It is worthy to note that content/idea generation or an entirely AI-produced manuscript is usually not permitted by most journals (see Table 5).

Table 4. Physical education journals included in the study.

Journal	Publisher	Category Quartile	Impact Factor (2024)	Journal Citation
EPER	Sage journal	Q1	3.3	SSCI
PESP	Taylor & Francis	Q1	3.0	SSCI
CSHPE	Taylor & Francis	Q1	2.1	ESCI
MPEES	Taylor & Francis	Q2	1.9	SSCI
JTPE	Human Kinetics Journals	Q2	1.8	SSCI
JPES	Editura Universitatea din Pitesti	Q3	0.38	
PETM	OVS LLC	Q3	0.32	
JPERD	Taylor & Francis	Q3	0.29	ESCI
SJPSE	Taylor & Francis	Q4	N/A	
IJPEFS	Asian Research Association	N/A	N/A	

Note: SSCI = Social Sciences Citation Index; ESCI = Emerging Sources Citation Index; PESP= *Physical Education and Sport Pedagogy*; CSHPE= *Curriculum Studies in Health and Physical Education*; MPEES= *Measurement in Physical Education and Exercise Science*; JPERD= *Journal of Physical Education, Recreation and Dance*; SJPSE= *Strategies: A Journal for Physical and Sport Educators*; JTPE= *Journal of Teaching in Physical Education*; EPER= *European Physical Education Review*; JPES= *Journal of Physical Education and Sport*; IJPEFS= *International Journal of Physical Education, Fitness and Sports*; PETM= *Physical Education Theory and Methodology*.

Table 5. Rationale & responsibility requirements of AI usage across PE journals.

Journal Name	Language Refinement	Substantive Content Generation	Image/Data Generation
PESP	✓ “Language improvement ... Idea generation and idea exploration ... Literature classification” (T&F AI Policy).	✓* “Authors should not submit manuscripts where Generative AI tools have been used in ... text or code generation without rigorous revision” (T&F AI Policy).	X “Does not permit the use ... in the creation and manipulation of images and figures” (T&F AI Policy).
CSHPE	✓ Same as above	✓* Same as above	X Same as above
MPEES	✓ Same as above	✓* Same as above	X Same as above
JPERD	✓ Same as above	✓* Same as above	X Same as above
SJPSE	✓ Same as above	✓* Same as above	X Same as above
JTPE	✓ “Technologies can be used only for language refinement” (Human Kinetics Journals AI Policy).	X “AI use for original writing, data analysis, or other applications is not permitted” (Human Kinetics Journals AI Policy).	X “AI use for original writing, data analysis, or other applications is not permitted” (Human Kinetics Journals AI Policy).
EPER	✓ “Sage welcomes developments in this area ... Polishing language, or structuring a submission” (Sage AI Policy).	✓* “Sage welcomes developments in this area ... synthesizing, or analyzing findings” (Sage AI Policy).	✓* “The use of AI tools that can produce content such as generating references, text, images or any other form of content” (Sage AI Policy).
JPES	✓* Not explicitly detailed in text	✓* Not explicitly detailed in text	✓* Not explicitly detailed in text
IJPEFS	✓ “Authors may use AI tools for non-substantive tasks ... Using AI for grammar checks, spelling, punctuation and stylistic improvements” (Asian Research Association AI Policy).	X “Prohibited uses include ... content generation ... abstracts, literature reviews, methodology, results or conclusions” (Asian Research Association AI Policy).	X “Image Creation: Generating or altering images ... is prohibited”; “Data Manipulation ... is prohibited” (Asian Research Association AI Policy).
PETM	✓ “AI tools may be used only for technical assistance, such as: language editing and grammar checking, improving clarity and fluency of English language, minor stylistic adjustments that do not alter the scientific content” (PETM AI Policy).	X “Unacceptable ... producing substantial parts of the scientific content (Introduction, Methods, Results, Discussion)” (PETM AI Policy).	X “Unacceptable ... generating fabricated data” (PETM AI Policy).

Note: ✓ = Allowed; ✓* = Conditionally allowed (typically requires disclosure or specific conditions); X = Prohibited.

The rationale for this ban lies in the fact that the core cognitive tasks of scientific research must be led and completed by humans to ensure the authenticity of academic outcomes.

Human responsibility and authorship integrity

Every analyzed journal strictly excludes AI tools from being listed as authors or coauthors. This universal prohibition serves to preserve the integrity of the academic record, ensuring that every name on a byline corresponds to a human entity capable of ethical reasoning and professional liability.

The journals included in the study primarily explain the policy excluding AI from authorship with accountability as the core rationale. *Physical Education and Sport Pedagogy* (PESP) further elaborates that AI tools are unable to manage copyright and

licensing agreements or provide contractual assurances regarding the integrity of the work. *European Physical Education Review* (EPER) adds a qualitative dimension, noting that LLMs cannot replicate human creative and critical thinking. Thus, the requirement for human authorship is a mechanism to ensure that there is always a liable party responsible for the originality, validity and integrity of the published contents.

Data accuracy and bias prevention

The policies articulate specific epistemological risks associated with AI (see [Table 6](#)). The policies warn of systemic biases embedded in the training data, which are hard to detect and can inadvertently propagate stereotypes in research findings. To address accuracy and integrity risks, these publications have established a dual defense mechanism. While 90% of journals prohibit the use of AI for substantive research tasks, 100% mandate author self-checks. This universal requirement reflects a professional consensus that authors must serve as the primary filter against AI hallucinations. Technical solutions remain underdeveloped, as only 50% of journals, primarily those within the Taylor & Francis portfolio, utilize automated testing tools. Furthermore, these applications are strictly limited to image detection through platforms like Imagetwin and cannot yet verify textual authenticity.

Strategies for mitigating bias are notably less standardized, with only 60% of journals requiring authors to screen for systemic bias and a mere 20% requiring editors to check for such biases in AI-generated content once authors have explicitly disclosed their use of AI. The discrepancy between the 100% self-check rate for accuracy versus the 60% rate for bias control suggests that journals prioritize factual errors over subtle algorithmic prejudices. In general, journals rely heavily on authorial integrity due to the infancy of AI detection technology.

Confidentiality risks and privacy

Regulatory oversight regarding author confidentiality appears inadequate. Only 60% of the included journals established requirements for manuscript confidentiality when using AI. Moreover, the majority of these merely required authors to self-check and prevent data leakage, with only the *International Journal of Physical Education, Fitness and Sports* explicitly prohibiting the uploading of unpublished materials or manuscripts.

Thematic AI policy for reviewers

Strict confidentiality requirements

The protection of unpublished manuscripts is treated as the paramount ethical obligation for reviewers. [Table 7](#) shows that 80% of journals strictly prohibit reviewers from uploading unpublished manuscripts into AI tools. This consensus reflects a fear that AI platforms act as “data leaks,” compromising the confidentiality that is central to the blind peer review process.

Policies explicitly prohibit reviewers from uploading manuscripts, project proposals, or any associated files into AI tools. This restrictive approach suggests that the current AI infrastructures is viewed as insufficiently secure for handling privileged information, thereby necessitating a strict separation between the peer review process and AI tools.

Table 6. Perspectives and risk Mitigation adopted by PE journals for accuracy and bias and confidentiality.

Accuracy / Integrity Risks					Bias Risks		Privacy & Confidentiality			
Journal	Test tool	Self-check	Prohibit to SG	Perspectives	Check by Editor	Self-check	Perspectives	Prohibit use	Self-check	Perspectives
PESP	✓*	✓	✓	"Generative AI tools are of a statistical nature (as opposed to factual) and, as such, can introduce inaccuracies, falsities (so-called hallucinations)" (T&F AI Policy).	✓	✓	"Generative AI ... can introduce bias, which can be hard to detect" (T&F AI Policy).	✓	✓	"Generative AI tools are often used on third-party platforms that may not offer sufficient standards of confidentiality, data security, or copyright protection" (T&F AI Policy).
CSHPE	✓*	✓	✓	Same as above.	✓	✓	Same as above.	✓	✓	Same as above.
MPEES	✓*	✓	✓	Same as above.	✓	✓	Same as above.	✓	✓	Same as above.
JPERD	✓*	✓	✓	Same as above.	✓	✓	Same as above.	✓	✓	Same as above.
SJPSE	✓*	✓	✓	Same as above.	✓	✓	Same as above.	✓	✓	Same as above.
JTPE		✓	✓	"AI use for original writing, data analysis, or other applications is not permitted" (Human Kinetics Journals AI Policy).			N/A			N/A
EPER		✓	✓	"Be conscious of the potential for fabrication where the LLM may have generated false content, including getting facts wrong, or generating citations that don't exist. Ensure you have verified all claims in your article prior to submission" (Sage AI Policy).	✓	✓	"Before incorporating sources into your scholarly work, apply the CRAAP test to your responses to avoid sending out misinformation or spreading bias" "As Editors, we rely on your subject level expertise to ... If a sentence or paragraph does not make sense, or appears to be machine generated, please query it with authors, or raise it with Sage for advice" (S age AI Policy).			N/A
JPES		✓		(Not explicitly detailed in text).			N/A			N/A

(Continued)

(Continued)

Table 6. (Continued).

Journal	Accuracy / Integrity Risks				Bias Risks		Privacy & Confidentiality	
	Test tool	Self-check	Prohibit to SG	Perspectives	Check by Editor	Self-check	Perspectives	Prohibit use
IJPEFS	✓	✓	✓	"Prohibited uses include ... data manipulation: using AI to fabricate, manipulate or misrepresent data" (Asian Research Association AI Policy).				✓
PETM	✓	✓	✓	"The use of AI is considered unacceptable if ... generating fabricated data, results, or references" (PETM AI Policy).	✓		"Reviewers and editors may pay special attention to signs of inappropriate AI use (e.g., inconsistent text, implausible results, incorrect or non-existent references)" (PETM AI Policy).	✓
Rate	50%	100%	90%		20%	60%		10%
								50%

Note: ✓ = Use this strategy; ✓*=Image detection only; N/A = Not mentioned; SG = Substantive Content Generation.

Table 7. Checklist of AI usage permissions for peer reviewers.

Journals	Uploading Unpublished Manuscripts		Use AI for Decision Making		Language Polish
PESP	X	"Peer reviewers must not upload unpublished manuscripts or project proposals, including any associated files, images or information, into Generative AI tools" (T&F AI Policy).	X	"Peer reviewers are chosen experts in their fields and should not be using Generative AI for analysis or to summaries submitted articles or portions thereof in the creation of their reviews (T&F AI Policy)."	✓ "Generative AI may only be utilized to assist with improving review language, but peer reviewers will at all times remain responsible for ensuring the accuracy and integrity of their reviews (T&F AI Policy)."
CSHPE	X	Same like above.	X	Same like above.	✓ Same like above.
MPEES	X	Same like above.	X	Same like above.	✓ Same like above.
JPERD	X	Same like above.	X	Same like above.	✓ Same like above.
SJPSE	X	Same like above.	X	Same like above.	✓ Same like above.
JTPE	X	"Reviewers and editors should not upload any manuscript in whole or in part into a generative AI tool." (Human Kinetics Journals AI Policy).	X	"Reviewers and editors should not use generative AI to make decisions about, evaluate, or review a manuscript (Human Kinetics Journals AI Policy)."	X Not mentioned. (Note: The policy explicitly bans uploading "even to improve readability")
EPER	X	"We ask that Editors ensure the reviewers invited are aware of the confidentiality issues presented by generating a review report using language models or generative AI" (Sage AI Policy).	X	"Reviewers using ChatGPT or other Generative AI tools to generate review reports inappropriately will not be invited to review for the journal and their review will not be included in the final decision (Sage AI Policy)."	✓ "Reviewers may wish to use Generative AI to improve the quality of the language in their review. If they do so, they maintain responsibility for the content, accuracy and constructive feedback within the review (Sage AI Policy)."
JPES		Not mentioned.		N/A	N/A
IJPEFS	X	"Authors, reviewers and editors must not input confidential, sensitive or unpublished material into AI tools." (Asian Research Association AI Policy).		N/A	N/A
PETM		Not mentioned.		N/A	N/A
PESP Rate	0% 80%		0% 70%		60% 70%

Note: ✓ = Allowed; X = Prohibited; N/A = Not mentioned.

Primacy of human judgment in evaluation

Journals emphasize that reviewers are selected for their domain-specific expertise, which current AI technologies are unable to replicate. For example, *PESP* explicitly states that reviewers are chosen as experts in their respective fields and should not use AI for manuscript analysis. This position reflects a broader belief that expert human judgment is qualitatively distinct from, and not substitutable by, AI, thereby safeguarding the quality and integrity of the peer review process (see [Table 4](#)).

Divergence on linguistic assistance

As shown in [Table 4](#), unlike the bans on content generation, journals like *Physical Education and Sport Pedagogy* and *European Physical Education Review* permit AI to improve review language. This acknowledges that reviewers, particularly non-native English speakers, may

benefit from tools that help them express their critiques more effectively. Nevertheless, these policies still require reviewers to take full responsibility for the accuracy and integrity of any content generated by AI for language assistance. Crucially, all permission for AI use are coupled with clauses that place full liability on reviewers.

Discussion

The data from the study indicated that only 10 of the 27 (i.e., 37%) PE journals have established AI policies, an adoption rate significantly lower than that observed in other disciplines (Caputo 2025). This may suggest that the PE field currently places less emphasis on AI governance, raising potential concerns regarding the robustness of policy enforcement. Despite a universal consensus against AI authorship, policies differ significantly concerning the usage boundaries for substantive content generation. Moreover, current frameworks exhibit critical enforcement gaps by prioritizing reviewer confidentiality while leaving vulnerabilities in author data privacy and offering limited guidance on the necessary screening for algorithmic bias.

Limited top-down policy diffusion in the PE Field

The data from the current study reveal a pronounced “top-down” diffusion pattern in AI policies within the field of PE. Specifically, regulations primarily originate from the standardized mandates of major publishers such as Taylor & Francis, and Sage, rather than being autonomously formulated by individual journals based on disciplinary specificities. This finding aligns with previous policy surveys in other disciplines, which suggest that top-tier journals and publishing houses typically take the lead in introducing and adopting regulatory frameworks to address major technological innovations in scientific research (Ganjavi et al., 2024; Perkins & Roe, 2024). Subsequently, these rules cascade down to subsidiary journals (Lund & Naheem, 2024). Although there are advantages (i.e., promoting consistency and efficiency across journals) for many journals to largely derive their AI policies from publisher-level guidelines, this centralized approach may not fully account for disciplinary differences (Wang & Gong, 2026). In the context of PE, the field’s applied and interdisciplinary nature, particularly its emphasis on movement-based research, pedagogical practice, and human participant interactions, raises considerations that may not be fully addressed by general AI policies. For example, using AI to analyze student movements or conduct assessments introduces distinct ethical and methodological challenges, which differ significantly from text-heavy or nonhuman disciplines. While many medical journals have implemented discipline-specific AI policies driven by professional autonomy and collective alliances rather than commercial publishers (O’Brien et al., 2025; Yoo, 2025), such tailored governance remains rare in humanities and social science journals, where generic publisher templates still predominate (Alduais et al., 2025; Resnik & Hosseini, 2025). As such, it is reasonable to suggest that PE journals could benefit from some degree of policy individualization to better reflect our field-specific practices and concerns. For example, AI policies in the field of PE may need to address whether AI-generated images of human movement should be permitted in PE journals.

It is also important to note that the diffusion of AI policies within the PE field remains limited. Consequently, 17 of the 27 PE-focused journals we collected identified in this study

did not have publicly available AI policies posted on their websites. These findings indicate that AI policy adoption remains incomplete across PE journals, as many journals still do not provide publicly available guidance on AI use in publication and peer review. This regulatory gap leaves journals situated outside established policy frameworks vulnerable to publishing articles that contain non-compliant or undisclosed AI-generated content (Ganjavi et al., 2024). The reasons why many PE journals have yet to establish explicit AI policies to guide research publication and peer review practices remain unclear and warrant further investigation.

Consensus on upholding human responsibility and disqualifying AI authorship

The data from the study reveals that all journals strictly prohibit AI authorship (see Table 2), with accountability universally cited as the primary rationale, which is consistent with what has been reported in the literature on the topic in other fields (Alkhawam et al., 2025; Kim, 2024). This indicates a consensus within the PE academic community: regardless of technological evolution, the legal liability, ethical obligations, and credit for creativity in academic research must always reside with human participant. Viewed through the lens of a HCAI framework, this consensus perfectly operationalizes the core pillar of human accountability. This universal consensus across PE journals aligns seamlessly with the core tenets established by COPE, which unequivocally states that AI tools cannot meet the requirements for authorship because they cannot take legal or moral responsibility for the submitted work, nor can they assert the presence or absence of conflicts of interest (COPE Council [COPE], 2023). Therefore, excluding AI from authorship is not merely a technical stipulation, but a profound normative statement anchored in human accountability. This policy convergence centers primarily on establishing the accountability of human authors and maintaining confidentiality. All journal policies emphasize reinforcing human authorial responsibility; excluding AI from authorship is not merely a technical stipulation but a profound normative statement. Delineating the normative boundaries between “human” and “nonhuman” is critical for establishing agency, legal liability, and maintaining journal legitimacy (Christen et al., 2023). This human-centric perspective provides a theoretical foundation for relevant legislation and policy formulation (Hallamaa & Kalliokoski, 2022), asserting that the originality of human scholarship must be safeguarded in academic writing and that AI tools should not supplant the dominant role of human authors (Cheng et al., 2025). From a legal standpoint, AI, as a non-legal entity, cannot hold copyright or sign contracts; thus, humans are recognized as the sole authorial entities (Ide et al., 2023; Teixeira Da Silva, 2023). Furthermore, given AI’s lack of emotional sentience and moral judgment (Cheng et al., 2025; Teixeira Da Silva, 2023), enforcing strict human accountability for the authorship represents the most effective mechanism at this stage to mitigate legal risks and ensure the quality of scholarly publications, and uphold the integrity of the HCAI framework.

Permissive use of AI for improving the quality of writing

Across the AI policies examined in this study, another common area of consensus concerns the permissive use of AI to support the quality of academic writing (see Table 2), which is also in line with what has been reported in other fields (Caputo 2025; Guelmami et al.,

2024). These practices are generally framed as analogous to traditional editorial support services and are viewed as non-substantive contributions that do not alter the intellectual content of the research. This restrictive use reflects an underlying principle that improving how ideas are expressed is fundamentally different from producing the ideas themselves. However, it remains unclear why this rationale has been so broadly adopted across the publishing field, particularly given ongoing debates about human–AI collaboration. Notably, even within this permissive category, journals vary in their expectations for transparency, with some requiring explicit disclosure of AI-assisted language editing and others offering only general guidance. Together, these findings suggest that journals recognize the practical benefits of AI for enhancing writing quality, while simultaneously reinforcing the primacy of human authorship and intellectual responsibility in scholarly publishing.

It is important to recognize that academic writing differs fundamentally from other forms of writing, such as documentaries or announcements, in that meaning, stance, and authorial voice are constructed through deliberate lexical and grammatical choices. Word selection, hedging, modality, and syntactic structure all contribute to how arguments are framed, how claims are qualified, and how authors position themselves in relation to existing scholarship. These features are not merely stylistic and they are integral to the development and communication of scholarly knowledge. As such, AI-assisted polishing or revision, often framed as “language improvement” can extend beyond surface-level edits. Changes in wording, sentence structure, or tone may unintentionally alter emphasis, shift the strength of claims, or modify how evidence is interpreted. Therefore, the boundary between language improvement and substantive content change is not always clear-cut. What appears to be a minor linguistic adjustment can carry meaningful implications for interpretation, authorial positioning, and even the perceived contribution of the work. Given this potential risk, utilizing AI for language editing still necessitates human oversight. From the perspective of the HCAI paradigm, authors cannot simply accept algorithmic outputs; Mosqueira-Rey et al. (2023) emphasize the necessity of a “human-in-the-loop” approach, where users must serve as active verifiers to ensure that AI-driven language refinements do not inadvertently distort their scientific intention. The risk of such unintended distortions underscores the need for careful consideration and transparent guidelines regarding the use of AI in academic writing, particularly in disciplines where nuance and contextual sensitivity are central to scholarly communication.

Blurred boundaries of legitimacy in core content generation

The analysis shows that 70% of the journals adopt a permissive but regulated model (see Table 2), viewing AI as a research assistant capable of enhancing efficiency. This stance aligns with existing research, given that some scholars assert that leveraging AI for processing extensive literature reviews or preliminary idea conceptualization can significantly reduce mechanical cognitive load, thereby enabling researchers to devote more energy to higher-order critical analysis and creative thinking (Pavlik, 2023). Furthermore, this human-AI collaboration model is regarded as an inevitable methodological evolution for addressing the exponential growth of scientific literature and enhancing research productivity (Salvagno et al., 2023). In contrast, the remaining 30% of the included journals implement a strict prohibition policy, categorizing tasks involving cognitive processing as

restricted area for AI. On one hand, research indicates that mainstream AI tools still pose a risk of AI hallucinations, where AI tools produce content based on user prompts that appears plausible but is, in reality, completely fabricated, absurd, or factually incorrect (Athaluri et al., 2023; Sun et al., 2024). As a result, utilizing AI tools to create core content and draft manuscripts could severely compromise the accuracy and integrity of research findings (Perkins, 2023). On the other hand, given that AI-generated text may be highly similar or even identical to original sources, using AI to directly generate core content and paragraphs may lead authors to inadvertently use excerpts from previously published articles (Guelmami et al., 2024). This inconsistency in defining usage boundaries is a prominent characteristic of the current policy landscape, yet it also creates significant confusion for authors. The same manuscript utilizing AI-assisted literature classification might be deemed compliant innovation by one journal, while constituting academic misconduct in another (Ganjavi et al., 2024). This regulatory fragmentation highlights that the PE discipline has not yet formed a unified consensus regarding what constitutes unique human intellectual contribution.

Overall, current AI policies in PE indicate that AI is permitted in research publishing and, to a limited extent, in the peer review process. However, its use is typically restricted to language refinement. AI is generally positioned as a supplementary technical tool within existing scholarly workflows rather than as an integrated component of research production. These policy patterns reflect an emphasis on control and risk management in the governance of AI within academic publishing. Such an approach may constrain the generation of new knowledge within scholarly research publishing.

Policy gaps and limitations

The lack of explicit details in current AI policies from PE journals raises concerns in terms of systemic enforcement gap. The absence of effective verification technologies compels journals to rely heavily on authorial integrity and ethical constraints. Furthermore, current policies predominantly focus on the accuracy of AI-generated content, while some journals pay insufficient attention to the potential algorithmic biases inherent in AI tools. Moreover, despite strict prohibitions against reviewers regarding manuscript upload risks, regulations governing authors remain comparatively lax.

Absence of technical verification and over-reliance on author integrity

When facing the risks of AI hallucinations and fabrication, although all journals require author self-checks, only 50% have introduced automated detection tools (see Table 3), which are primarily limited to image plagiarism checks. This cautious attitude toward text detection technology stems from the immaturity and instability of current detection technologies. Multiple studies point out that existing AI text detection technologies are still in an immature stage, and their reliability is far from meeting the zero-tolerance standards required for academic publishing (Chemaya & Martin, 2024; Weber-Wulff et al., 2023). Research has found that these detection algorithms are prone to false positives, particularly when assessing the writing of non-native English scholars, where the false alarm rate rises significantly (Liang et al., 2023). In addition, further research confirms that content generated by ChatGPT often exhibits extremely low text duplication rates, allowing it to easily evade existing text matching and detection software, which makes relying solely

on technical means to identify AI-generated content extremely challenging (Elkhatat, 2023). Therefore, given the status where technology cannot reconcile fairness with accuracy, most PE journals then choose to avoid using such tools to prevent creating misjudgments due to algorithmic defects.

While the current strategy of journals successfully transfers the liability for AI misuse from publishers to individual researchers at a legal level through mandatory author declarations, it may be virtually ineffective in terms of actual regulatory impact. Due to the lack of reliable detection tools as enforcement mechanisms, ethical policies aimed at regulating AI use often devolve into formalism. Studies point out when the benefits of noncompliance are extremely high while the risk of detection is nearly zero, relying solely on author self-discipline creates a massive responsibility gap (Hu, 2024). This regulatory environment lacking hard constraints actually creates a systematic gray area that increases the risk of undisclosed AI-generated content entering academic publications (Hosseini et al., 2023).

Insufficient scrutiny of AI algorithmic bias

Results from our study reveal that while all included journal policies emphasize the accuracy of AI-generated content, only 60% explicitly mandate screening for algorithmic bias (see Table 3). These results indicate that a significant proportion of PE journals still deviate from the UNESCO recommendation on the Ethics of AI. UNESCO (2022) warns that AI algorithms risk amplifying societal biases and discrimination. Consequently, it mandates all AI actors to actively minimize these risks and implement ethical impact assessments to combat AI-generated stereotypes (UNESCO, 2022). Given that PE research frequently involves sensitive issues such as body image, race, and gender, this regulatory blind spot may allow stereotypes embedded within AI to pass unverified, posing a potential threat to the objectivity of disciplinary knowledge. Previous studies have warned that Large Language Models (LLMs) may inherit societal biases and data skew from their training data, which are often amplified in the generated content (Gallegos et al., 2024; Navigli et al., 2023). Research indicates that AI's tendency to prioritize idealized body standards while excluding non-standard types directly undermines inclusivity in PE (Luccioni et al., 2023). Furthermore, in decision-making scenarios involving race and gender (Wu et al., 2025), AI has been confirmed to exhibit significant implicit discrimination, such as the underestimation of specific groups' abilities or stereotypical categorization (Abid et al., 2021). From a HCAI perspective, the current regulatory blind spot regarding bias reflects a prioritization of superficial procedural compliance over the substantive integration of fairness and justice into the review process, which is essential to safeguard research communities from diverse backgrounds from being marginalized by algorithmic bias. If such unscrutinized bias permeates academic writing, it will severely compromise the impartiality and ethical foundation of scientific research.

Asymmetric confidentiality standards for authors and reviewers

To safeguard the confidentiality of peer review, 80% of journals have established stringent barriers, explicitly prohibiting reviewers from uploading unpublished manuscripts to AI platforms (see Table 4). However, data privacy protections on the author's side appear disproportionately weak: only 10% of policies explicitly prohibit authors from inputting sensitive research data into AI tools, while as many as 40% remain entirely silent on this

issue (see Table 3). While this may indicate that journals have yet to establish comprehensive and explicit standards for data privacy, it does not imply a complete regulatory void regarding participant confidentiality. Recent research indicates that although 87% of institutional policies highlight the potential risks associated with inputting data into AI tools, only 38% explicitly address human participant data; therefore, while such data protection measures are gradually emerging, the safeguarding of research participants still lacks sufficient attention (Ganguly et al., 2025). Concurrently, guidelines issued by the Secretary's Advisory Committee on Human Research Protections emphasize the critical role of institutions and Human Research Protection Programs in addressing data security concerns (Protections [OHRP], 2022). Therefore, journals could consider strengthening their collaboration with institutional bodies regarding data protection. By incorporating more rigorous document reviews during the manuscript submission process, journals can work in tandem with institutions to mitigate potential data breach risks, thereby enhancing their accountability and contribution to the broader public interest of data privacy. For instance, commercial AI platforms (e.g., ChatGPT) typically retain user data by default for model training purposes (Hosseini et al., 2023). Research warns that once sensitive data are uploaded to the cloud, it risks not only becoming a permanent part of the model's memory but also being maliciously extracted and reconstructed through specific technical vectors, such as model inversion attacks (Mollaki, 2024). PE research not only frequently involves the image privacy of minors but also encompasses vast amounts of multi-modal sensitive data, including students' physiological and biochemical markers, mental health records, private interview data and videos of motor skills (Radanliev, 2025). In the absence of explicit prohibitions, the inadvertent input of sensitive biometrics into commercial systems risks exposing participants to unauthorized exploitation and violating regulations like General Data Protection Regulation (GDPR) or Health Insurance Portability and Accountability A (HIPAA) (Hosseini et al., 2023; Latzel & Glauner, 2024).

Implications for future research and journal policy formulation

Implications for future research

Given that most of PE journals rely on generic publisher templates, future research should deeply investigate how these top-down mandates adapt to the specific disciplinary needs of PE over time. Future studies should examine whether the absence of publicly available AI policies is associated with journal characteristics such as ranking, impact factor, indexing status, publisher type, or editorial structure. Furthermore, since detection technologies remain immature, the mere possession of a policy does not equate to compliance; therefore, future studies should further explore the actual implementation and efficacy of these policies. More empirical research is needed to verify whether the "author self-check" mechanism effectively filters out hallucinations and bias, or if it exacerbates the "responsibility gap" identified in this study.

Future research should proactively explore how AI policies will shape the evolving roles and practices of authors and reviewers in scholarly publishing. As AI tools become more integrated into research workflows, studies are needed to assess how policy design can support effective, transparent, and responsible use that enhances research quality and review processes. In particular, examining how appropriate applications of AI may improve efficiency, analytical depth, and research capacity will be critical. The affordability of AI

tools for underserved authors and reviewers remains an important concern, as disparities in access may contribute to new forms of inequality in scholarly publishing. Addressing these emerging inequities warrants greater attention from researchers, editors, and policy professionals within the field. Insights from such research can guide the development of forward-looking policy frameworks that balance innovation with accountability and support sustainable human – AI collaboration in new knowledge production.

Implications for journal policy formulation

To address limitations within the existing policy framework and move beyond short-term compliance, journals should anchor their future AI governance in a dynamic and human-centered framework that requires continuous optimization. First, journals should enforce active human oversight and eliminate regulatory asymmetry regarding data privacy. This requires bridging the gap where reviewer confidentiality is strictly guarded while author-side data privacy is less adequately addressed. Specifically, policies should explicitly prohibit the input of non-anonymized biometric data into public AI platforms to ensure compliance with regulations such as GDPR. Second, acknowledging the enforcement challenges of technical bans, journals should transition from simple prohibitions toward radical transparency. Furthermore, considering that only 60% of current policies mandate bias screening, PE journals should institutionalize specific “bias check” protocols. These should require authors to explicitly confirm that they have scrutinized AI-generated content for potential stereotypes, thereby upholding the discipline’s core values of human-centered inclusivity.

Limitations

First, the analysis based solely on public documents reflects only the status of policy establishment. The absence of publicly articulated AI policies does not necessarily indicate a lack of internal editorial oversight, as journals may operate under nonpublic criteria. Accordingly, this study maps the outward-facing policy landscape but cannot assess actual internal enforcement or author compliance. Second, the exclusive inclusion of English-language journals may overlook regulatory nuances in non-English contexts, thereby limiting the global generalizability of the findings. Finally, the number of included journals was constrained by the current lag in policy formulation within the PE field, as only 10 of the 27 retrieved journals had established policies, thereby limiting the potential for broader statistical inference.

Conclusions

From a macro perspective, our findings reveal a “top-down” diffusion pattern driven by major publishers; however, the depth of this diffusion within the PE field remains insufficient, with a significant lag in policy formulation observed among other PE journals. Regarding specific regulations, while a strong academic consensus has been reached on rejecting AI authorship to ensure human accountability and AI can be used for improving the quality of writing, significant divergence persists concerning the permissible boundaries for the substantive use of AI tools. Furthermore, current regulations regarding potential algorithmic bias remain inadequate, and a regulatory asymmetry exists between reviewer and author policies, prioritizing the confidentiality of peer review while leaving gaps in addressing author-side data privacy.

Disclosure statement

No potential conflict of interest was reported by the author(s).

ORCID

Yucen Li  <http://orcid.org/0009-0002-9084-6925>

Xiaofen D. Hamilton  <http://orcid.org/0000-0002-3762-4103>

Timothy Lynch  <http://orcid.org/0000-0002-7096-541X>

Mark F. Hamilton  <http://orcid.org/0000-0002-1622-6480>

References

- Abid, A., Farooqi, M., & Zou, J. (2021). Persistent anti-muslim bias in large language models. In Proceedings of the 2021 AAAI/ACM Conference on AI, Ethics, and Society (pp. 298–306). Association for Computing Machinery. <https://doi.org/10.1145/3461702.3462624>
- Albannai, N. A. A., & Raziq, M. M. (2026). Navigating ethical, human-centric leadership in AI-driven organizations: A thematic literature review. *The Service Industries Journal*, 46(5–6), 555–582. <https://doi.org/10.1080/02642069.2025.2534360>
- Alduais, A., Qadhi, S., Chaaban, Y., & Khraisheh, M. (2025). Utilizing generative AI responsibly and ethically for research purposes in higher education: A policy analysis. *Serials Review*, 51(3–4), 120–170. <https://doi.org/10.1080/00987913.2025.2581429>
- Alkhawam, M., Almobayed, A., Pandey, A., Nanda, N. C., Ebrahimi, A. J., & Ahmed, M. I. (2025). Exploring AI use policies in manuscript writing in cardiology and vascular journals. *American Heart Journal Plus: Cardiology Research and Practice*, 58, 100586. <https://doi.org/10.1016/j.ahjo.2025.100586>
- Athaluri, S. A., Manthena, S. V., Kesapragada, V. S. R. K. M., Yarlagadda, V., Dave, T., & Duddumpudi, R. T. S. (2023). Exploring the boundaries of reality: Investigating the phenomenon of artificial intelligence hallucination in scientific writing through ChatGPT references. *Cureus*, 15(4), e37432. <https://doi.org/10.7759/cureus.37432>
- Bowen, G. A. (2009). Document analysis as a qualitative research method. *Qualitative Research Journal*, 9(2), 27–40. <https://doi.org/10.3316/QRJ0902027>
- Caputo, J. L. (2025). Editorial policies on artificial intelligence in sport science: Balancing innovation and ethical standards. *Kinesiology Review*, 14(4), 519. <https://doi.org/10.1123/kr.2025-0010>
- Caputo, J. L., Ibrahim, A. H., Davis, D. W., Stone, W. J., Garver, M. J., & Johnson, S. L. (2025). Editorial policies on artificial intelligence in sport science: Balancing innovation and ethical standards. *Kinesiology Review*, 14(4), 519–523. <https://doi.org/10.1123/kr.2025-0010>
- Chemaya, N., & Martin, D. (2024). Perceptions and detection of AI use in manuscript preparation for academic journals. *PLOS ONE*, 19(7), e0304807. <https://doi.org/10.1371/journal.pone.0304807>
- Cheng, A., Calhoun, A., & Reedy, G. (2025). Artificial intelligence-assisted academic writing: Recommendations for ethical use. *Advances in Simulation*, 10(1), 22. <https://doi.org/10.1186/s41077-025-00350-6>
- Christen, M., Burri, T., Kandul, S., & Vörös, P. (2023). Who is controlling whom? Reframing “meaningful human control” of AI systems in security. *Ethics and Information Technology*, 25(1), 10. <https://doi.org/10.1007/s10676-023-09686-x>
- COPE Council. (2023). *COPE position statement: Authorship and AI tools*. Committee on Publication Ethics. <https://publicationethics.org/guidance/cope-position/authorship-and-ai-tools>
- Ekberg, J.-E. (2021). Knowledge in the school subject of physical education: A Bernsteinian perspective. *Physical Education & Sport Pedagogy*, 26(5), 448–459. <https://doi.org/10.1080/17408989.2020.1823954>

- Elkhatat, A. M. (2023). Evaluating the authenticity of ChatGPT responses: A study on text-matching capabilities. *International Journal for Educational Integrity*, 19(1), 15. <https://doi.org/10.1007/s40979-023-00137-0>
- Gallegos, I. O., Rossi, R. A., Barrow, J., Tanjim, M. M., Kim, S., DERNONCOURT, F., Yu, T., Zhang, R., & Ahmed, N. K. (2024). Bias and fairness in large language models: A survey. *Computational Linguistics*, 50(3), 1097–1179. https://doi.org/10.1162/coli_a_00524
- Ganguly, A., Johri, A., Ali, A., & McDonald, N. (2025). Generative artificial intelligence for academic research: Evidence from guidance issued for researchers by higher education institutions in the United States. *AI and Ethics*, 5(4), 3917–3933. <https://doi.org/10.1007/s43681-025-00688-7>
- 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>
- Gatrell, C., Muzio, D., Post, C., & Wickert, C. (2024). Here, there and everywhere: On the responsible use of artificial intelligence (AI) in management research and the peer-review process. *Journal of Management Studies*, 61(3), 739–751. <https://doi.org/10.1111/joms.13045>
- Gendron, Y., Andrew, J., & Cooper, C. (2022). The perils of artificial intelligence in academic publishing. *Critical Perspectives on Accounting*, 87, 102411. <https://doi.org/10.1016/j.cpa.2021.102411>
- Guelmami, N., Ben Ezzeddine, L., Hatem, G., Trabelsi, O., Ben Saad, H., Glenn, J. M., El Omri, A., Chalghaf, N., Taheri, M., Bouassida, A., Ben Aissa, M., Trabelsi, K., Ammar, A., Bouzouraa, M. M., Saidane, M., Eken, Ö., Clark, C. C. T., Parsakia, K., Dhahbi, W. & Chamari, K. (2024). The ethical compass: Establishing ethical guidelines for research practices in sports medicine and exercise science. *International Journal of Sport Studies for Health*, 7(2), 31–46. <https://doi.org/10.61838/kman.intjssh.7.2.4>
- Hallamaa, J., & Kallioikoski, T. (2022). AI ethics as applied ethics. *Frontiers in Computer Science*, 4, 776837. <https://doi.org/10.3389/fcomp.2022.776837>
- Hosseini, M., Resnik, D. B., & Holmes, K. (2023). The ethics of disclosing the use of artificial intelligence tools in writing scholarly manuscripts. *Research Ethics*, 19(4), 449–465. <https://doi.org/10.1177/17470161231180449>
- Hsu, H.-Y., Hakouz, A., & Fotouhi, G. (2025). Towards responsible generative AI in academia: A synthesis of AI policies on academic writing in the field of educational research. *AI and Ethics*, 5(5), 5467–5484. <https://doi.org/10.1007/s43681-025-00794-6>
- Hu, G. (2024). Challenges for enforcing editorial policies on AI-generated papers. *Accountability in Research*, 31(7), 978–980. <https://doi.org/10.1080/08989621.2023.2184262>
- Ide, K., Hawke, P., & Nakayama, T. (2023). Can ChatGPT be considered an author of a medical article? *Journal of Epidemiology*, 33(7), 381–382. <https://doi.org/10.2188/jea.JE20230030>
- Kim, E. (2024). Analyzing AI use policy in LIS: Association with journal metrics and publisher volume. *Scientometrics*, 129(12), 7623–7644. <https://doi.org/10.1007/s11192-024-05189-8>
- Klimova, B., & Pikhart, M. (2025). Exploring the effects of artificial intelligence on student and academic well-being in higher education: A mini-review. *Frontiers in Psychology*, 16. <https://doi.org/10.3389/fpsyg.2025.1498132>
- Larsson, H. (2023). Movement learning: Pedagogy and agentic realism. *Quest*, 75(2), 136–150. <https://doi.org/10.1080/00336297.2022.2153708>
- Latzel, R., & Glauner, P. (2024). Artificial intelligence in sport scientific creation and writing process. In C. Dindorf, E. Bartaguiz, F. Gassmann, & M. Fröhlich (Eds.), *Artificial intelligence in sports, movement, and health* (pp. 15–29). Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-67256-9_2
- Liang, W., Yuksekgonul, M., Mao, Y., Wu, E., & Zou, J. (2023). GPT detectors are biased against non-native English writers. *Patterns*, 4(7), 100779. <https://doi.org/10.1016/j.patter.2023.100779>
- Luccioni, A. S., Akiki, C., Mitchell, M., & Jernite, Y. (2023). Stable bias: Evaluating societal representations in diffusion models. *Advances in Neural Information Processing Systems*, 36, 56338–56351. <https://doi.org/10.52202/075280-2458>

- Lund, B., & Naheem, K. (2024). Can ChatGPT be an author? A study of artificial intelligence authorship policies in top academic journals. *Learned Publishing*, 37(1), 13–21. (WOS: 001071474800001). <https://doi.org/10.1002/leap.1582>
- Mangum, L. C., & Kuenze, C. (2025). Responsible use of artificial intelligence in the NATA journals manuscript review process. *Journal of Athletic Training Education and Practice*, 21(3), 244–245. <https://doi.org/10.4085/1947-380X-25-00.ED>
- Mollaki, V. (2024). Death of a reviewer or death of peer review integrity? The challenges of using AI tools in peer reviewing and the need to go beyond publishing policies. *Research Ethics*, 20(2), 239–250. <https://doi.org/10.1177/17470161231224552>
- Mosqueira-Rey, E., Hernández-Pereira, E., Alonso-Ríos, D., Bobes-Bascarán, J., & Fernández-Leal, Á. (2023). Human-in-the-loop machine learning: A state of the art. *Artificial Intelligence Review*, 56(4), 3005–3054. <https://doi.org/10.1007/s10462-022-10246-w>
- Navigli, R., Conia, S., & Ross, B. (2023). Biases in large language models: Origins, inventory, and discussion. *Journal of Data and Information Quality*, 15(2), 1–21. <https://doi.org/10.1145/3597307>
- O'Brien, C., Thayani, Z., Smith, T., Tran, A. V., Crotty, P., Young, A., Ford, A. I., & Vassar, M. (2025). Evaluating AI guidelines in leading family medicine journals: A cross-sectional study. *BMC Primary Care*, 26(1), 368. <https://doi.org/10.1186/s12875-025-03044-0>
- Ozmen Garibay, O., Winslow, B., Andolina, S., Antona, M., Bodenschatz, A., Coursaris, C., Falco, G., Fiore, S. M., Garibay, I., Grieman, K., Havens, J. C., Jirotko, M., Kacorri, H., Karwowski, W., Kider, J., Konstan, J., Koon, S., Lopez-Gonzalez, M., Maifeld-Carucci, I., . . . Xu, W. (2023). Six human-centered artificial intelligence grand challenges. *International Journal of Human-Computer Interaction*, 39(3), 391–437. <https://doi.org/10.1080/10447318.2022.2153320>
- Patton, M. Q. (2015). *Qualitative research & evaluation methods: Integrating Theory and Practice* (4th ed.). Sage Publications, Inc. <https://us.sagepub.com/en-us/nam/qualitative-research-evaluation-methods/book232962>
- Pavlik, J. V. (2023). Collaborating with ChatGPT: Considering the implications of generative artificial intelligence for journalism and media education. *Journalism & Mass Communication Educator*, 78(1), 84–93. <https://doi.org/10.1177/10776958221149577>
- Perkins, M. (2023). Academic integrity considerations of AI large language models in the post-pandemic era: ChatGPT and beyond. *Journal of University Teaching & Learning Practice*, 20(2). <https://doi.org/10.53761/1.20.02.07>
- Perkins, M., & Roe, J. (2024). Academic publisher guidelines on AI usage: A ChatGPT supported thematic analysis. *F1000research*, 12, 1398. <https://doi.org/10.12688/f1000research.142411.2>
- Protections. (2022, October 21). IRB considerations on the use of artificial intelligence in human subjects research [Recommendation]. <https://www.hhs.gov/ohrp/sachrp-committee/recommendations/irb-considerations-use-artificial-intelligence-human-subjects-research/index.html>
- Radanliev, P. (2025). Privacy, ethics, transparency, and accountability in AI systems for wearable devices. *Frontiers in Digital Health*, 7, 1431246. <https://doi.org/10.3389/fdgth.2025.1431246>
- Resnik, D. B., & Hosseini, M. (2025). Disclosing artificial intelligence use in scientific research and publication: When should disclosure be mandatory, optional, or unnecessary? *Accountability in Research*, 33(2), 1–13. <https://doi.org/10.1080/08989621.2025.2481949>
- Salvagno, M., Taccone, F. S., & Gerli, A. G. (2023). Can artificial intelligence help for scientific writing? *Critical Care*, 27(1), 75. <https://doi.org/10.1186/s13054-023-04380-2>
- Schmager, S., Pappas, I. O., & Vassilakopoulou, P. (2025). Understanding human-centred AI: A review of its defining elements and a research agenda. *Behaviour & Information Technology*, 44(15), 3771–3810. <https://doi.org/10.1080/0144929X.2024.2448719>
- Shneiderman, B. (2020). Human-centered artificial intelligence: Reliable, safe & trustworthy. *International Journal of Human-Computer Interaction*, 36(6), 495–504. <https://doi.org/10.1080/10447318.2020.1741118>
- Springmier, K. (2025). An exploration of faculty and student perceptions of generative AI. *Library Trends*, 73(4), 426–442. <https://doi.org/10.1353/lib.2025.a968490>
- Sun, Y., Sheng, D., Zhou, Z., & Wu, Y. (2024). AI hallucination: Towards a comprehensive classification of distorted information in artificial intelligence-generated content. *Humanities and Social Sciences Communications*, 11(1). <https://doi.org/10.1057/s41599-024-03811-x>

- Teixeira Da Silva, J. A. (2023). Is ChatGPT a valid author? *Nurse Education in Practice*, 68, 103600. <https://doi.org/10.1016/j.nepr.2023.103600>
- Thelwall, M. (2024). Can ChatGPT evaluate research quality? *Journal of Data and Information Science*, 9(2), 1–21. <https://doi.org/10.2478/jdis-2024-0013>
- Thorp, H. H. (2023). ChatGPT is fun, but not an author. *Science*, 379(6630), 313–313. <https://doi.org/10.1126/science.adg7879>
- UNESCO. (2022). *Recommendation on the ethics of artificial intelligence*(SHS/BIO/PI/2021/1). United Nations Educational, Scientific and Cultural Organization. <https://unesdoc.unesco.org/ark:/48223/pf0000381137>
- Veiga, A. D. (2025). Ethical guidelines for the use of generative artificial intelligence and artificial intelligence-assisted tools in scholarly publishing: A thematic analysis. *Science Editing*, 12(1), 28–34. <https://doi.org/10.6087/kcse.352>
- Volshøj, E. S., Aggerholm, K., Beni, S., & Madsen, K. L. (2025). Meaningful physical education: Towards an embodied pedagogy. *European Physical Education Review*, 31(4), 638–654. <https://doi.org/10.1177/1356336X241300426>
- Wang, Z., & Gong, M. (2026). A cross-disciplinary analysis of AI policies in academic peer review. *Learned Publishing*, 39(1), e2035. <https://doi.org/10.1002/leap.2035>
- Weber-Wulff, D., Anohina-Naumeca, A., Bjelobaba, S., Foltýnek, T., Guerrero-Dib, J., Popoola, O., Šigut, P., & Waddington, L. (2023). Testing of detection tools for AI-generated text. *International Journal for Educational Integrity*, 19(1), 26. <https://doi.org/10.1007/s40979-023-00146-z>
- Wu, C., Jr. Jr., Worrell, M. R., Jr., Guan, J., Zhang, J., Xiao, N., Lynch, T., Hamilton X. D., & Hamilton, M. F. (2025). Gender and racial differences in first-year college students' body dissatisfaction and social physique anxiety *Journal of the International Alliance for Health, Physical Education, Dance, and Sport*, 1(1), 1–31. <https://doi.org/10.6084/m9.figshare.29991961.v1>
- Yang, A., Hamilton, X. D., Wang, Y., Smolianov, P., Castro-Piñero, J., Guan, J., Dolmatova, T., Zhang, X., Zhang, J., Zhan, E., & Hamilton, M. F. (2025). Results assessment methods of health-related fitness test batteries in school-based physical education programs: A global Comparison. *Journal of Teaching in Physical Education*, 44(3), 564–572. (185973545). <https://doi.org/10.1123/jtpe.2024-0005>
- Yoo, J.-H. (2025). Defining the boundaries of AI use in scientific writing: A comparative review of editorial policies. *Journal of Korean Medical Science*, 40(23), e187. <https://doi.org/10.3346/jkms.2025.40.e187>