

Review article

A Review of the TITAN Guideline: Advancing Transparency and Responsible Artificial Intelligence Use in Medical Research and Scholarly Publishing

Arzoo Nazir¹ and Shah Zeb^{2,*}

¹ College of Veterinary Medicine, Institute of Comparative Medicine, Yangzhou University, 225009 Yangzhou, China

² Department of Chemical Biology, Institute of Physical Chemistry, Polish Academy of Sciences, Kasprzaka 44/52, 01-224 Warsaw, Poland

*Correspondence: shahzeb1267@gmail.com

Article information:

Received: 17-07-2026

Revised: 13-08-2026

Accepted: 13-08-2026

Published: 14-08-2026

Academic Editor:

Dr. Naveed Ahmed

Keywords:

Scientific misconduct

TITAN guidelines

Artificial intelligence

Research ethics

Publishing ethics

Abstract: Artificial Intelligence (AI) is increasingly being integrated into medical research and scholarly publishing, supporting activities such as literature searching, data analysis, medical imaging, manuscript preparation, and peer review. Despite these opportunities, AI use introduces concerns related to hallucinations, bias, privacy, confidentiality, copyright, reproducibility, and scientific accountability. Existing guidance provides important principles for responsible AI use, but reporting practices remain variable. This article reviews the TITAN guidelines as a practical framework for improving transparency and responsible reporting of AI use in medical research and scholarly publishing. The TITAN guidelines provide a proportionate and technology neutral approach to AI reporting. It distinguishes minor uses, such as language assistance, from substantive applications involving research design, analysis, interpretation, or scientific content. The framework emphasizes that human researchers retain responsibility for evaluating AI outputs and the integrity of published work. It also provides a basis for authors, reviewers, editors, and publishers to incorporate standardized AI reporting into scholarly workflows. TITAN offers a practical framework for documenting meaningful AI involvement while supporting transparency, reproducibility, and human accountability. Its flexible structure can accommodate emerging AI technologies, including multimodal and agentic systems. Periodic revision and coordinated adoption by medical journals and research communities will be important to maintain its relevance as AI capabilities continue to evolve.

Article citation: Nazir, A. & Zeb, S. A Review of the TITAN Guideline: Advancing Transparency and Responsible Artificial Intelligence Use in Medical Research and Scholarly Publishing. *Electron J Med Res.* 2026. 2(3): 59-76. <https://doi.org/10.64813/ejmr.2026.126>

Introduction

Artificial Intelligence (AI) has rapidly become an important part of biomedical research and scholarly publishing. Advances in machine learning, deep learning, natural language processing, computer vision, and generative AI have created new opportunities to support activities that previously required substantial human time and expertise [1]. AI-assisted tools are now being explored for literature searching, study planning, statistical analysis, medical

imaging, manuscript preparation, and peer review. These developments can improve research efficiency, facilitate the handling of large and complex datasets, and support researchers in navigating an increasingly information-intensive scientific environment [2].

The growing use of AI, however, also introduces important scientific and ethical challenges. One major concern is the generation of inaccurate or fabricated information, commonly described as AI hallucination. Generative AI systems can produce convincing statements, references, explanations, or

interpretations that lack reliable supporting evidence. In medical research, such errors may be particularly consequential because scientific findings can influence subsequent research, clinical practice, and health-related decisions. AI-generated content therefore requires appropriate human verification rather than being accepted solely because it appears plausible or professionally written [3].

Bias is another important concern. AI systems may reproduce or amplify biases present in their training data, algorithms, or implementation. In biomedical research, this can affect AI performance and interpretation across different populations, institutions, and clinical settings. Models developed using limited or unrepresentative datasets may not perform consistently when applied to other populations. Researchers must therefore consider whether AI-assisted outputs are appropriate for the specific population and research context in which they are being used [4].

The use of AI also raises questions about research integrity, privacy, confidentiality, copyright, and accountability. Researchers may interact with AI systems using unpublished findings, clinical information, proprietary datasets, or other sensitive materials. Without appropriate safeguards, these practices may risk patient confidentiality and research data protection. Similarly, AI-generated text, images, code, and other materials may raise questions regarding originality, attribution, and intellectual property. These concerns demonstrate that responsible AI use involves more than technical accuracy; it requires consideration of the wider ethical and scientific environment [5].

Transparency is therefore essential when AI contributes to medical research or scholarly communication. Readers, reviewers, and editors should be able to understand whether AI was used, for what purpose, and how its outputs were evaluated. A general statement that AI was used may not be sufficient when the technology has contributed to substantive aspects of research design, analysis, or interpretation, or manuscript development. Information such as the AI tool used, its intended purpose, relevant prompts or inputs, human oversight, and validation procedures may be necessary to understand the role of AI and assess its potential influence on the final work [6].

Existing organizations and publishers have introduced recommendations addressing several of these concerns. Guidance from bodies such as the International Committee of Medical Journal Editors (ICMJE), Committee on Publication Ethics (COPE),

World Association of Medical Editors (WAME), Nature, Elsevier, and Springer Nature has emphasized principles including disclosure, human accountability, confidentiality, and responsible AI use [7]. These initiatives provide an important foundation, but differences in scope and level of detail remain. In particular, researchers may lack a standardized framework for reporting the practical characteristics of AI-assisted research in a consistent and reproducible manner [8].

In this context, the TITAN Guideline is proposed as a practical framework for advancing transparency and responsible AI use in medical research and scholarly publishing. Rather than restricting AI use, TITAN seeks to make meaningful AI involvement visible, assessable, and accountable. It focuses on key reporting domains including AI disclosure, tool identification, intended purpose, relevant prompts and inputs, human oversight, verification, ethical safeguards, bias, and reproducibility [9].

The underlying principle of TITAN is that human researchers remain responsible for the scientific integrity of their work, regardless of the extent to which they use AI. By encouraging proportionate and structured reporting, the guideline aims to help authors, reviewers, editors, and publishers understand the role of AI within the research process. Such transparency can strengthen reproducibility, support responsible innovation, and preserve confidence in the medical literature as AI becomes increasingly integrated into scientific research and publishing [10].

Artificial Intelligence in Medical Research

AI is increasingly being incorporated into different stages of medical research, from identifying research questions to analyzing data and communicating scientific findings. Its applications include machine learning, deep learning, natural language processing, computer vision, and generative AI [11]. Although these technologies differ in their methods and capabilities, they share the ability to process large volumes of information and identify patterns that may be difficult to detect through conventional approaches. As biomedical research becomes increasingly data-intensive, AI has emerged as a potentially valuable research support technology [12].

One important application of AI is literature searching and evidence discovery. Biomedical researchers often need to evaluate large numbers of publications across multiple databases. AI-assisted systems can support semantic searching, article classification, citation mapping, and extraction of relevant information. Natural language processing can

identify relationships between concepts even when different terminology is used across publications [13]. Generative AI can also assist researchers in developing search strategies and summarizing retrieved information. However, AI-generated literature recommendations should not replace systematic database searches because outputs may be incomplete or inaccurate. References suggested by generative AI should also be independently verified, as fabricated or incorrectly described citations remain a recognized limitation of such systems [14].

AI can also contribute to study design and research planning. Researchers may use AI to explore potential research questions, identify variables, generate hypotheses, and examine patterns in preliminary datasets. Machine learning can identify complex relationships that may help researchers develop hypotheses for subsequent testing [15]. However, AI-generated hypotheses should be considered exploratory rather than established scientific conclusions. Automated systems may identify associations caused by confounding, data quality problems, or sampling limitations. Human researchers must therefore determine whether proposed relationships are biologically and methodologically plausible [16].

A major area of application is statistical analysis and prediction. Machine learning models can process high-dimensional datasets and identify nonlinear relationships among variables. Such approaches have been investigated for disease prediction, clinical risk assessment, survival analysis, epidemiology, genomics, and drug development. AI can be particularly useful when datasets contain numerous interacting variables. Nevertheless, predictive performance can be affected by overfitting, data leakage, missing information, class imbalance, and differences between development and validation populations. Consequently, high model accuracy alone does not establish clinical validity or generalizability [17].

Medical imaging represents another major field of AI research. Deep learning and computer vision methods have been applied to computed tomography, magnetic resonance imaging, ultrasound, mammography, and digital pathology [18]. AI systems may assist with image classification, lesion detection, segmentation, and quantitative measurements. These applications can reduce repetitive workloads and support standardized image assessment. However, model performance may change across institutions

because of differences in patient populations, imaging equipment, acquisition protocols, and disease prevalence. External validation is therefore important before conclusions about broader clinical applicability are made [19].

The emergence of generative AI has significantly expanded AI-assisted manuscript preparation. Researchers may use generative systems for language editing, translation, summarization, outlining, and drafting support. These applications can improve readability and reduce the time required for routine writing tasks [20]. However, fluent language does not guarantee scientific accuracy. AI-generated text may contain unsupported claims, incorrect interpretations, or fabricated references. Researchers must therefore verify substantive content against the underlying , and evidence ensures that AI-assisted editing does not alter the intended scientific meaning [21].

AI is also being explored in peer review and editorial workflows. Potential applications include manuscript screening, reference checking, identification of reporting problems, reviewer recommendation, and summarization of submitted work [22]. Such tools could reduce administrative workload, but peer review requires contextual scientific judgment and ethical responsibility. Confidentiality is a particular concern because manuscripts may contain unpublished data, proprietary information, or intellectual contributions. Reviewers and editors should therefore avoid transferring confidential material to AI systems unless appropriate safeguards and authorization are available [23].

These applications demonstrate that AI can influence almost every stage of the medical research workflow. The significance of its use, however, varies considerably. AI used only for spelling or grammar correction has a different scientific impact from AI used for statistical analysis, image interpretation, or generation of scientific conclusions. This distinction supports the need for proportionate AI reporting, in which the level of disclosure reflects the importance of AI to the research process [24].

Overall, AI can provide substantial benefits across medical research, but its value depends on appropriate human oversight and critical evaluation. AI should be regarded as a research support technology rather than an independent source of scientific authority. Clear documentation of its role is therefore necessary to distinguish responsible assistance from unverified or inappropriate reliance on automated outputs [25].

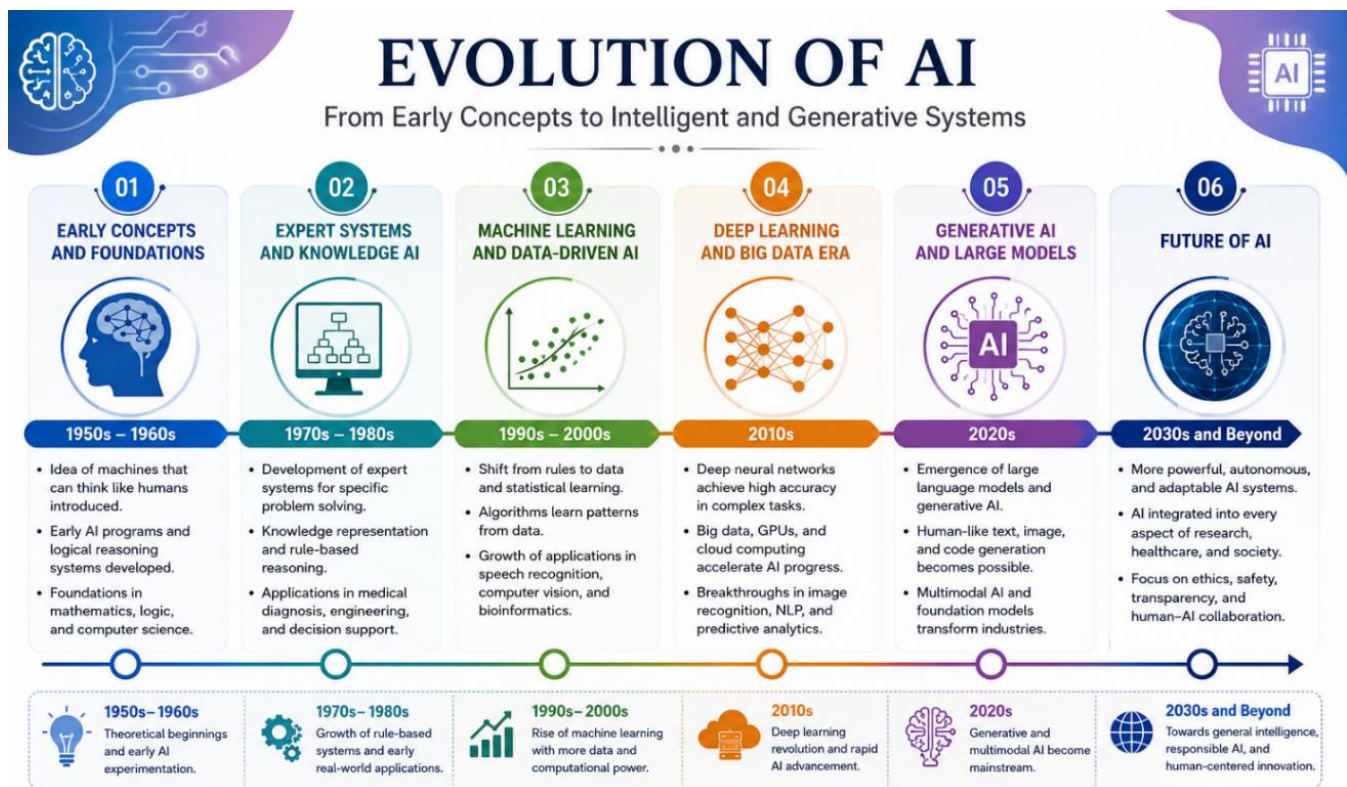


Figure 1. Evolution of AI in medical research and scholarly publishing.

Table 1. Applications of Artificial Intelligence Across the Medical Research Workflow.

Research stage	Examples of AI use	Potential benefit	Major concern
Literature search	Semantic search, citation mapping, summarization	Faster evidence discovery	Incomplete or fabricated information
Study design	Hypothesis and variable exploration	Supports research planning	Unsupported associations
Data analysis	Prediction and pattern recognition	Handles complex datasets	Overfitting and bias
Medical imaging	Detection, classification, segmentation	Efficient image assessment	Limited generalizability
Manuscript preparation	Editing, translation, drafting	Improved efficiency and readability	Hallucinations and altered meaning
Peer review	Screening and reviewer support	Reduced workload	Confidentiality and bias
Literature search	Semantic search, citation mapping, summarization	Faster evidence discovery	Incomplete or fabricated information
Study design	Hypothesis and variable exploration	Supports research planning	Unsupported associations

Existing AI Guidance Before TITAN

The rapid adoption of artificial intelligence in medical research has encouraged professional organizations, medical editors, and major academic publishers to develop policies for its responsible use. These initiatives have addressed several important issues, including disclosure, authorship, human accountability, confidentiality, research integrity, and the verification of AI-generated content. Although these policies differ in scope, they provide an

important foundation for developing more structured approaches to AI reporting [26].

The International Committee of Medical Journal Editors (ICMJE) has emphasized that artificial intelligence tools should not be considered authors because they cannot accept responsibility for the accuracy, integrity, and originality of published work [27]. Human authors remain responsible for reviewing AI-assisted content and ensuring that the final manuscript meets appropriate scientific standards. ICMJE also supports disclosure of AI-assisted

technologies when they have been used in the preparation of a manuscript. This approach establishes an important distinction between technological assistance and scientific accountability. However, disclosure alone may not explain the specific purpose, extent, or methodological influence of AI use [28].

The Committee on Publication Ethics (COPE) has addressed AI within the broader context of publication ethics. Its recommendations emphasize transparency, responsible use, authorship, confidentiality, and the need for human oversight. COPE has also highlighted concerns related to the use of AI during peer review, particularly because submitted manuscripts may contain confidential and unpublished information. The organization's approach reinforces the principle that AI should support, rather than replace, human responsibility in editorial and research activities [29].

The World Association of Medical Editors (WAME) has similarly emphasized responsible AI use in scholarly publishing. Its guidance recognizes potential benefits of AI-assisted writing while drawing attention to inaccurate information, fabricated references, plagiarism, and unclear accountability [30]. WAME emphasizes that authors must verify AI-generated material and remain responsible for the accuracy of their published work. This position is particularly relevant to medical publishing, where inaccurate information may have consequences beyond the academic literature [31].

Major publishers have also introduced specific AI policies. Nature has established guidance distinguishing acceptable AI assistance from uses that may affect authorship, scientific integrity, or the authenticity of research outputs. Its policies emphasize disclosure of substantive AI involvement and recognize particular concerns surrounding AI-generated or AI-modified images. These issues are important because images in scientific publications may represent experimental or clinical evidence rather than merely illustrative material [32].

Elsevier has developed policies emphasizing transparency, human oversight, and author responsibility when generative AI is used. Its approach recognizes that AI may be useful for language- and productivity-related tasks, but that substantive AI-generated content requires careful human evaluation. Confidentiality and responsible handling of information are also relevant when AI tools are used within editorial workflows [33].

Springer Nature has likewise emphasized that AI systems cannot be authors and that human researchers remain responsible for the content of submitted manuscripts. Its policies encourage

disclosure of substantive AI use and careful verification of AI-generated material. Protecting confidential information during editorial and peer review activities is another important consideration [34].

Despite these developments, existing guidance leaves several areas that could benefit from greater standardization. Most policies emphasize whether AI was used and whether its use should be disclosed, but they do not always provide a common structure for reporting the specific characteristics of AI-assisted research. Information such as the tool or model used, version, intended purpose, relevant prompts, input materials, human modifications, validation procedures, and potential influence on research decisions may be inconsistently reported [35].

Reproducibility presents an additional challenge. AI systems can produce different outputs across interactions, and commercial models may be updated without providing researchers with complete technical information. Consequently, simply naming an AI platform may not be sufficient to understand how a particular output was produced. A more structured approach is needed to document the aspects of AI use that are practically relevant to scientific evaluation [6].

The existing guidance therefore provides strong ethical principles but leaves an opportunity for a dedicated reporting framework that connects disclosure with methodological transparency. Such a framework should remain flexible enough to accommodate. This need forms the basis for the proposed TITAN Guideline [36].

The TITAN Guideline

The increasing integration of AI into medical research and scholarly publishing creates a need for a reporting framework that translates broad principles of responsible AI use into practical requirements [37]. Existing guidance has established important expectations regarding disclosure, human accountability, confidentiality, and verification, but researchers may still lack a consistent method for describing how AI contributed to their work. The TITAN Guideline is proposed to address this need by providing a structured framework for transparent and responsible reporting of AI-assisted research and publishing activities [38].

TITAN is designed as a reporting framework rather than a restriction on AI use. Its purpose is to enable researchers to benefit from AI technologies while ensuring that their involvement remains visible and scientifically accountable [3]. The framework recognizes that AI may be used for very different purposes, ranging from language editing and literature

searching to data analysis, image interpretation, research planning, and generation of scientific content. Because these applications have different levels of influence on research outcomes, reporting should reflect the nature and significance of the AI contribution [39].

A central objective of TITAN is to improve transparency. Readers, reviewers, and editors should be able to identify whether AI contributed to a research process and understand the purpose for which it was used. A meaningful disclosure should provide more information than a general statement that “AI was used.” Where relevant, researchers should identify the AI system, its intended function, and the extent to which its outputs influenced the final work [40].

TITAN also emphasizes human accountability. AI systems cannot independently accept responsibility for research findings, ethical decisions, or published conclusions. Human researchers must therefore remain responsible for evaluating AI-generated information, correcting errors, and ensuring that the final work accurately reflects the available evidence. The framework treats human oversight as an essential component of responsible AI-assisted research [41].

Another objective is to strengthen reproducibility. AI-assisted processes can be difficult to reproduce because outputs may vary between interactions, models may be updated, and some commercial systems provide limited technical information. TITAN encourages researchers to document relevant details, such as model or tool identification, access information, important prompts or instructions, and validation procedures when these details are relevant to understanding the research process [42].

The framework is intended for a broad range of stakeholders. Authors can use TITAN to determine what AI-related information should be disclosed during manuscript preparation. Reviewers can use reported information to assess whether AI involvement raises methodological or ethical concerns [43]. Editors and publishers can incorporate the framework into journal policies, submission forms, disclosure statements, and editorial workflows. Institutions may also use its principles when developing internal guidance for responsible AI use [30].

TITAN follows a proportionate reporting philosophy. Not every use of AI requires the same level of documentation. Basic language correction may require limited disclosure, whereas AI used for

statistical analysis, image interpretation, or scientific decision-making requires substantially more information [44]. This approach aims to prevent unnecessary reporting burdens while ensuring that meaningful AI contributions remain transparent. The framework is also intended to remain technology-neutral. By focusing on purpose, contribution, oversight, validation, ethics, and reproducibility rather than specific technologies, TITAN can remain applicable as AI develops. Its core principle is straightforward: when AI plays a meaningful role in producing or communicating medical research, that role should be reported clearly enough for others to understand and evaluate it [45].

Core Components of TITAN

The effectiveness of TITAN depends on converting its principles of transparency and accountability into practical reporting elements. The guideline therefore focuses on key domains that allow readers, reviewers, and editors to understand how AI was used, its influence on the research process, and how its outputs were evaluated. The level of detail should remain proportional to the importance of AI within the study [46].

AI Disclosure and Tool Identification

Authors should clearly disclose the use of AI when it contributes meaningfully to research or manuscript preparation. The disclosure should identify the general purpose of AI use and, where possible, the name of the tool, developer or provider, model, version, and date of access. Such information provides context because different AI systems can produce different outputs, and model updates may affect performance [27]. Disclosure should distinguish substantive scientific use from minor assistance. For example, language correction may require a brief statement, whereas AI used for data analysis or interpretation requires more detailed reporting [47].

Intended Purpose and Prompts

The specific purpose of AI use should be reported because its scientific significance depends on the task performed. Relevant purposes may include literature searching, hypothesis generation, statistical assistance, image analysis, coding, translation, manuscript drafting, or editorial support [48].

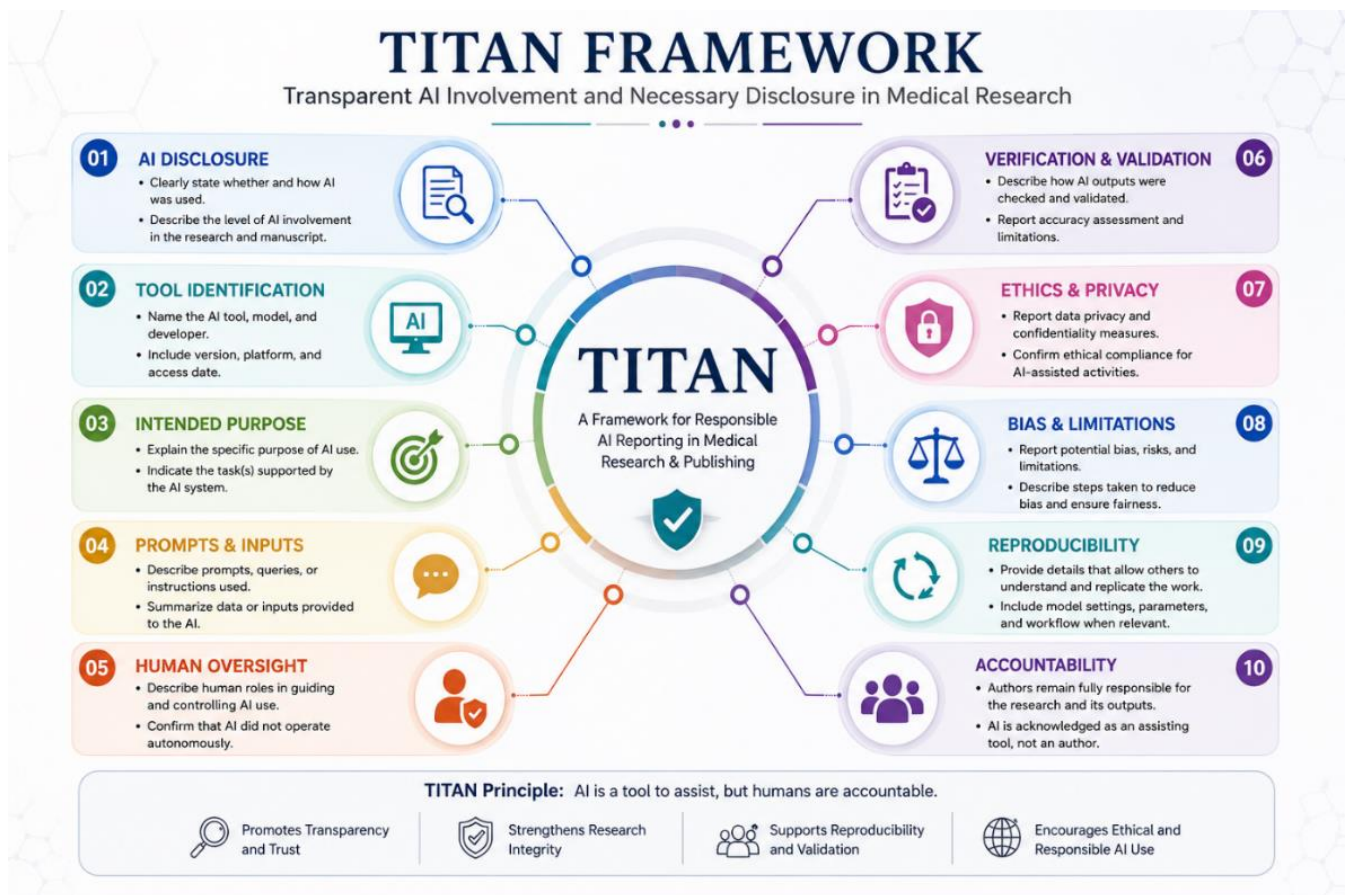


Figure 2. The TITAN framework for transparent and responsible AI reporting.

For generative AI, prompts or instructions can substantially influence outputs. When AI-generated material contributes directly to research decisions or scientific content, relevant prompts or their structure should be documented where feasible. Important input materials should also be described, particularly when they influence the generated output [49].

Human Oversight

Human oversight is a central requirement of TITAN. Researchers should review AI-generated outputs rather than accepting them without critical evaluation. The extent of oversight should reflect the potential consequences of the task [50].

For low-risk language assistance, human review may involve checking clarity and meaning. For statistical, clinical, or methodological applications, oversight should include appropriate expert evaluation and independent assessment of the AI output. The final scientific judgment must remain with the human research team [51].

Verification and Validation

AI-generated information should be verified before it is incorporated into scientific work. References should be checked against reliable databases, factual claims against appropriate evidence, and analytical outputs against validated statistical procedures. Where AI is used to develop predictive or diagnostic models, appropriate internal and external validation should be reported according to established methodological standards [28]. Verification is particularly important for generative AI because apparently credible outputs may contain factual errors, fabricated references, or unsupported conclusions [52].

Ethics, Privacy, and Bias

AI-assisted research should comply with relevant ethical and data protection requirements. Researchers should consider whether patient information, unpublished findings, proprietary datasets, or confidential documents are being entered into external AI systems. Appropriate safeguards should be used before processing sensitive information [53]. Potential bias should also be considered. AI systems may reflect limitations in their training data or design

and may perform differently across populations. Researchers should assess whether the selected AI system is appropriate for the study population and report relevant measures used to identify or reduce bias [54].

Reproducibility and Accountability

AI-assisted workflows should be documented sufficiently to allow others to understand how AI contributed to the research. Depending on the application, this may include the model or tool, version, access date, prompts, relevant inputs, settings, generated outputs, and human modifications [55].

Exact replication may not always be possible because AI systems can change over time. Nevertheless, documenting the available information improves methodological transparency.

Human authors retain final responsibility for the accuracy, integrity, originality, and ethical acceptability of the published work. AI cannot assume authorship or accountability [56].

These domains are intended to operate as an integrated framework rather than as isolated requirements. Their application can vary according to the type of AI involved. A generative language model may require detailed information about prompts and human editing, whereas a machine learning model may require greater emphasis on training data, validation, performance, and external testing [57]. TITAN therefore provides a flexible structure that can be adapted to different research designs while maintaining consistent expectations for transparency, human oversight, and responsible AI use [58].

Table 2. Core Reporting Domains of the TITAN Guideline.

TITAN domain	Key information to report	Primary purpose
AI disclosure	Whether and where AI was used	Transparency
Tool identification	Tool, provider, model/version, access date	Traceability
Intended purpose	Specific task performed by AI	Context
Prompts and inputs	Relevant instructions and input type	Reproducibility
Human oversight	Review, modification, or rejection of outputs	Accountability
Verification	Methods used to check outputs	Accuracy
TITAN domain	Key information to report	Primary purpose
AI disclosure	Whether and where AI was used	Transparency
Tool identification	Tool, provider, model/version, access date	Traceability

Implications for Authors

Authors are the primary users of AI in medical research and therefore have a central responsibility for ensuring that its use remains transparent, appropriate, and scientifically defensible. The TITAN Guideline does not prohibit authors from using artificial intelligence; instead, it provides a structured approach for documenting its contribution and maintaining human responsibility for the final research output [59].

A first responsibility is to use AI for an appropriate and clearly defined purpose. Authors should consider whether an AI system is suitable for the task and whether its use introduces risks that could affect the validity of the research. AI may be useful for literature organization, language improvement, coding assistance, or exploratory analysis, but its outputs

should not automatically be treated as established scientific evidence [26].

Authors should also provide appropriate disclosure. When AI contributes meaningfully to research or manuscript preparation, the relevant use should be reported according to the journal's requirements. The disclosure should identify the general purpose of AI use and, where appropriate, provide information about the tool or model. Greater detail should be provided when AI has influenced research methods, analysis, interpretation, or substantive manuscript content [30].

Verification is a fundamental author responsibility. AI-generated information should be checked before it is included in a manuscript. This includes factual statements, references, numerical information, interpretations, code, and other substantive outputs. Researchers should verify AI-assisted content against

appropriate scientific sources, datasets, analytical procedures, or expert judgment. Particular caution is necessary with generative AI because fluent and confident language does not guarantee factual accuracy [52].

Authors must also maintain human accountability. AI systems cannot assume responsibility for scientific conclusions, ethical decisions, or the integrity of published work. Researchers must review AI-assisted material, correct errors, and ensure that conclusions are supported by the evidence. AI should therefore function as a supporting technology rather than an independent scientific contributor [15].

Privacy and confidentiality require additional attention. Authors should avoid entering identifiable patient information, confidential manuscripts, unpublished research findings, or proprietary data into AI platforms unless appropriate authorization and safeguards are available. The data handling practices of external AI systems should be considered before sensitive information is processed [60].

Finally, authors should preserve relevant information about their AI-assisted workflow when reproducibility may depend on it. Tool names, model versions, important prompts, access dates, outputs, and subsequent human modifications may be useful records. Although exact replication may not always be possible, appropriate documentation allows readers and reviewers to better understand how AI influenced the research [61].

Implications for Reviewers

The increasing use of artificial intelligence in medical research also affects the responsibilities of peer reviewers. Reviewers may encounter manuscripts in which AI has contributed to literature searching, data analysis, image interpretation, coding, manuscript preparation, or other research activities. The TITAN Guideline provides reviewers with a framework for assessing whether such use has been reported adequately and whether appropriate human oversight and verification have been applied [22].

A primary responsibility of reviewers is to assess transparency. When AI appears to have played a meaningful role in the research, reviewers should consider whether the manuscript provides sufficient information about its use. The relevant disclosure should identify the general purpose of AI and, when appropriate, provide details about the tool, model, or process. Reviewers do not need to demand extensive technical information when AI was used only for minor language assistance, but substantive methodological use should receive greater attention [62].

Reviewers should also consider methodological validity. If AI has contributed to data analysis, prediction, medical image interpretation, coding, or other research procedures, the manuscript should provide enough information to understand how the AI-assisted method was implemented and evaluated. Reviewers should examine whether appropriate validation was performed and whether limitations such as overfitting, dataset bias, or limited generalizability have been considered [63].

The possibility of AI-generated inaccuracies requires particular attention. Reviewers should remain alert to unsupported claims, incorrect references, implausible interpretations, or inconsistencies between reported methods and results. However, reviewers should not assume that the presence of AI automatically indicates poor quality research. Evaluation should remain evidence-based and focused on the scientific quality of the manuscript [64].

Confidentiality is another critical responsibility. Reviewers should not upload confidential manuscripts, unpublished data, or proprietary information into external AI systems unless the journal explicitly permits such use and appropriate safeguards are in place. AI tools should never compromise the confidentiality expected within peer review [65].

Reviewers should also retain independent scientific judgment. AI may assist with administrative or analytical tasks, but it should not replace the reviewer's assessment of study quality, originality, methodological rigor, or clinical relevance. If reviewers use AI-supported tools where permitted, they remain responsible for the final review and should follow the journal's disclosure requirements [51]. TITAN therefore positions reviewers as an important safeguard between AI-assisted research and the published scientific record. Their role is not to police AI use but to ensure that meaningful AI involvement is transparent, appropriately validated, and consistent with accepted standards of scientific integrity [32].

Implications for Editors and Publishers

Editors and publishers have an important role in establishing how artificial intelligence is used, disclosed, and evaluated within scholarly publishing. As AI becomes integrated into manuscript preparation, peer review, editorial screening, and publication workflows, journals need policies that promote transparency without unnecessarily limiting appropriate technological innovation. The TITAN Guideline can provide a practical foundation for developing such policies [66].

A key responsibility of journals is to establish clear AI disclosure requirements. Author instructions and submission systems should explain when AI use must be reported and what information should be provided. Disclosure requirements should distinguish minor language assistance from substantive AI involvement in study design, data analysis, interpretation, image generation, or manuscript development. Standardized disclosure fields can reduce ambiguity and improve consistency across submissions [67].

Editors can also integrate TITAN into the manuscript submission and screening process. Submission forms may ask authors to indicate whether AI was used, identify its purpose, and provide additional information when the technology contributed substantially to the research. Such information can help editors determine whether further clarification is required before peer review [68].

Reviewer guidance should also address AI use. Journals should clearly state whether reviewers are permitted to use AI tools during peer review and under what conditions. Where AI-assisted review is permitted, policies should protect manuscript confidentiality and ensure that human reviewers remain responsible for their evaluations. Reviewers should not submit confidential manuscript content to external AI systems without explicit authorization and appropriate safeguards [69].

Editors may also use AI for administrative or screening tasks, such as identifying incomplete submissions, checking references, or supporting reviewer selection. However, automated systems should not replace editorial judgment. AI-assisted recommendations should be reviewed by qualified editors, particularly when decisions may affect acceptance, rejection, or the interpretation of scientific findings [70].

Publishers should provide training and policy updates because AI technologies and their associated risks are changing rapidly. Editorial teams need sufficient knowledge to identify common problems such as hallucinated references, inappropriate AI-generated content, confidentiality risks, and undisclosed substantive AI use. Training should also emphasize that AI-generated outputs require human verification [71]. TITAN can further support workflow standardization across journals. A common reporting structure would allow publishers to develop consistent disclosure forms, editorial checklists, and author

guidance. This could reduce differences between journals and make AI reporting easier for researchers who submit to multiple publications [72].

Implementation should remain proportionate. Requiring extensive documentation for every minor use of AI could create unnecessary administrative burdens and discourage legitimate applications. Journal policies should therefore focus greater scrutiny on AI uses that may influence research methods, results, interpretation, or scientific conclusions. Through clear policies, standardized disclosure procedures, reviewer safeguards, and appropriate training, editors and publishers can help ensure that AI strengthens rather than weakens the integrity of scholarly communication. TITAN offers a practical structure for achieving these goals while allowing journals to adapt requirements to their individual scope and research communities [73].

Strengths of TITAN

The TITAN Guideline has several strengths that may support its use in medical research and scholarly publishing. Its primary strength is its focus on transparency. Rather than treating AI use as a simple yes-or-no disclosure, TITAN encourages researchers to describe the purpose and significance of AI involvement. This provides readers, reviewers, and editors with a clearer understanding of how technology contributed to the research process [74].

A second strength is its emphasis on reproducibility. AI-assisted research can be difficult to reproduce because outputs may vary between interactions and models may change over time. By encouraging documentation of relevant tools, versions, prompts, inputs, and workflow details, TITAN can improve other researchers' ability to understand and assess AI-assisted methods. It does not assume that every AI output can be reproduced exactly, but it encourages researchers to preserve information that can reasonably be documented [75].

Human accountability is another central strength. TITAN clearly separates technological assistance from scientific responsibility. AI systems may support researchers with specific tasks, but human authors remain responsible for the accuracy, integrity, interpretation, and ethical acceptability of published work. This principle is particularly important when AI contributes to substantive scientific content [76].



Figure 3. Editorial workflow for integrating AI reporting into medical publishing.

The framework also offers flexibility. AI technologies are developing rapidly, and a guideline based on specific software products could quickly become outdated. TITAN instead focuses on reporting principles such as purpose, oversight, verification, bias, privacy, and reproducibility. These domains apply to generative AI, machine learning, medical imaging systems, multimodal models, and future technologies.

Another advantage is its stakeholder-oriented approach. TITAN is relevant not only to authors but also to reviewers, editors, publishers, and potentially research institutions. This broader perspective recognizes that responsible AI use extends throughout the scholarly communication process [77].

Finally, TITAN supports proportionate reporting. Not every AI application has the same scientific impact. A brief language editing interaction should not necessarily require the same documentation as AI-assisted statistical analysis or clinical interpretation. By allowing reporting requirements to reflect the significance of AI involvement, the framework can promote transparency without creating unnecessary burdens. These characteristics position TITAN as a practical framework for integrating AI reporting into existing research and publishing practices while preserving scientific accountability and methodological integrity [78].

Challenges

Despite its potential value, implementation of the TITAN Guideline may face several challenges. The most important is the rapid evolution of artificial intelligence. New models, platforms, and capabilities are being introduced frequently, while existing systems may change through updates [79]. A reporting framework must therefore remain sufficiently stable to provide consistent standards while also allowing adaptation to new forms of AI. Excessive dependence on specific technologies could make reporting requirements outdated. AI hallucination and reliability represent another challenge. Generative AI can produce information that appears credible but is inaccurate, incomplete, or fabricated. This creates difficulties when AI is used for literature searches, scientific writing, interpretation, or evidence synthesis. Human verification can reduce these risks, but the required level of verification may vary by task. Journals and researchers may also differ in their ability to evaluate technically complex AI outputs [80].

Privacy and confidentiality are particularly important in medical research. AI systems may process patient information, clinical records, unpublished findings, or proprietary datasets. Researchers may not always know how external platforms store or process submitted information. This creates potential risks involving patient confidentiality, data protection, and

unauthorized reuse of research material. TITAN can encourage reporting of safeguards, but broader institutional and legal policies are still required to manage these risks [81].

Bias and fairness present additional concerns. AI systems may reflect biases within their training data, development processes, or implementation environments. Medical datasets can reflect differences in demographic characteristics, healthcare access, geographic location, and disease prevalence. A model that performs well in one population may produce less reliable results in another. Reporting potential bias and mitigation measures can improve transparency, but identifying all sources of algorithmic bias remains difficult [82].

Copyright and intellectual property are also evolving areas. Researchers may use AI to generate or transform text, images, code, or other materials, creating questions about originality, ownership, attribution, and permitted use. The legal status of AI-generated material continues to develop across jurisdictions. A reporting guideline can promote transparency around AI involvement but cannot independently resolve legal questions that may vary between countries and institutions [83].

Another challenge is reproducibility. AI outputs may differ, depending on the model version, system settings, prompt structure, context, and date of use. Commercial platforms may also restrict access to model architecture or training information. Even when researchers document their interactions carefully, exact replication may not always be possible. TITAN should therefore be understood as improving process transparency rather than guaranteeing identical AI outputs [84].

There may also be concerns about reporting burden. If journals require extensive documentation for every minor use of AI, researchers may face unnecessary administrative work. Excessive requirements could discourage legitimate uses such as language editing or routine organizational assistance. Proportionate reporting is therefore essential, with more detailed disclosure reserved for AI applications that meaningfully influence research methods, results, or interpretation [85].

Finally, inconsistent journal policies may complicate implementation. Different journals may define acceptable AI use differently or require different disclosure formats. Researchers submitting the same work to different journals may therefore face varying expectations. Standardized principles such as TITAN could reduce this inconsistency, but successful implementation would require cooperation among

researchers, editors, publishers, institutions, and professional organizations [30].

These challenges indicate that responsible AI reporting cannot depend on disclosure alone. Effective implementation requires continuing policy development, researcher education, technical awareness, and periodic revision as AI capabilities and associated risks evolve [86].

Future Perspectives

The role of artificial intelligence in medical research is likely to expand as AI systems become more capable, interactive, and integrated with research software and digital infrastructure. Future developments may change not only how researchers use AI but also how scientific work is documented and evaluated. Reporting frameworks such as TITAN will therefore need to remain adaptable while preserving their core principles of transparency, human oversight, and accountability [87].

One important development is agentic AI. Unlike systems designed primarily to respond to individual prompts, agentic systems can potentially plan and execute multiple connected tasks with limited human intervention. In research, such systems may assist with literature searches, data organization, coding, analysis, and preparation of research outputs within a single workflow. Greater autonomy could increase efficiency but may also make it more difficult to identify where individual decisions were made. Future reporting standards may therefore need to document AI actions, decision points, tool interactions, and the extent of human intervention [88].

Multimodal AI is another important area of development. These systems can process different types of information, including text, images, audio, and structured data. In medical research, multimodal systems may support integrated analysis of clinical notes, laboratory results, medical images, and other patient information. Such applications could create new opportunities for research but also introduce complex issues involving validation, data governance, bias, and reproducibility [89].

AI-assisted peer review and editorial decision-making may also become more common. Automated systems could help identify methodological issues, reporting deficiencies, relevant literature, or potential reviewers. However, scientific judgment and ethical responsibility should remain with human reviewers and editors. Future policies may need to define which editorial activities can appropriately be supported by AI and what level of disclosure is required [90].

The rapid development of AI also means that reporting guidelines will require periodic revision. New models may introduce capabilities that are not adequately covered by current reporting categories. Future versions of TITAN may therefore need to incorporate additional domains related to autonomous decision-making, model provenance, multimodal inputs, external tool use, and interaction histories [91].

Future research should also evaluate whether AI reporting guidelines actually improve transparency and reproducibility. Empirical studies could examine how consistently researchers disclose AI use, whether reviewers can identify important risks from reported information, and which reporting elements provide the greatest value [92].

The future of responsible AI reporting will ultimately depend on maintaining a balance between innovation and scientific accountability. As AI becomes more deeply integrated into medical research, transparent documentation should evolve alongside the technology rather than being treated as a fixed requirement [93].

Recommendations for Medical Journals

Medical journals can play a central role in promoting transparent and responsible use of artificial intelligence by incorporating standardized AI reporting practices into their existing publication systems. The TITAN Guideline can provide a practical foundation for this process [94].

First, journals should establish clear AI disclosure policies within their author instructions. Authors should be asked to identify whether AI was used, the purpose of its use, and the extent to which it contributed to the research or manuscript. Reporting requirements should remain proportionate, so minor language assistance is not treated the same as AI involvement in analysis, interpretation, or study design [66].

Second, journals should incorporate standardized AI disclosure fields into electronic submission systems. Structured questions can improve consistency and help editors identify manuscripts that require additional information. Where AI has contributed substantially to scientific methods or content, authors should be encouraged to provide relevant details concerning the tool, model, prompts, human oversight, and validation [95].

Third, journals should provide AI-specific training for editors and reviewers. Training should address hallucinations, fabricated references, algorithmic bias, confidentiality risks, and appropriate evaluation of AI-

assisted research. Reviewers should also receive clear instructions concerning whether AI tools may be used during peer review and how confidential material must be protected [66].

Fourth, journals should develop standardized editorial workflows for handling AI disclosures. Editors should be able to determine when additional clarification, methodological review, or ethical assessment is necessary. AI-assisted editorial screening should support, not replace, professional editorial judgment [96].

Finally, journal policies should be reviewed regularly. Because AI technologies and associated risks are changing rapidly, reporting requirements should be updated periodically in consultation with researchers, editors, publishers, and relevant professional organizations.

Adoption of consistent AI reporting practices can reduce ambiguity across journals and strengthen transparency throughout medical publishing. TITAN provides a flexible structure that journals can adapt to their specific scope while maintaining common principles of disclosure, verification, accountability, and reproducibility [97].

Conclusion

Artificial intelligence is becoming increasingly integrated into medical research and scholarly publishing, offering important opportunities to improve research efficiency, data analysis, scientific communication, and evidence discovery. At the same time, its use raises concerns about hallucinations, bias, privacy, confidentiality, copyright, reproducibility, and scientific accountability. These challenges make transparent reporting of AI-assisted activities increasingly important. Existing guidance from major medical and publishing organizations has established important principles concerning disclosure, human responsibility, confidentiality, and responsible AI use. However, differences in scope and reporting detail indicate a need for a practical framework that connects these principles with specific information about how AI contributes to research. The TITAN Guideline provides such a framework by focusing on AI disclosure, tool identification, intended purpose, relevant prompts and inputs, human oversight, verification, ethics, bias, reproducibility, and final accountability. Its proportionate and technology-neutral approach allows the framework to accommodate different levels of AI involvement and future technological developments. TITAN does not seek to restrict responsible innovation. Instead, it aims to ensure that meaningful AI involvement remains visible and

scientifically accountable. Authors remain responsible for verifying AI-assisted content, while reviewers, editors, and publishers have complementary responsibilities for maintaining transparency and protecting confidentiality.

As AI technologies continue to evolve, periodic refinement of reporting standards will be necessary. By promoting consistent documentation and human oversight, TITAN can support responsible AI integration while strengthening the integrity, transparency, and trustworthiness of medical research and scholarly publishing.

Author Contributions: Conceptualization, S.Z.; methodology, A.N.; validation, S.Z.; resources, A.N.; writing—original draft preparation, A.N.; writing—review and editing, S.Z.; supervision, S.Z.; project administration, A.N. Both authors have read and agreed to the published version of the manuscript.

Funding: None.

Acknowledgments: The authors thank Yangzhou University, China, for providing facilities to write this informative review.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. Agha, R.A., Mathew, G., et al. Transparency in the reporting of Artificial Intelligence—the TITAN guideline. *Methods*. **2025**. 1(2).
2. Agha, R., Mathew, G., et al. Updated reporting guidelines in the age of Artificial Intelligence: SCARE, PROCESS, STROCSS and TITAN 2025. *Prem J Sci*. **2025**. 10: 100081.
3. Resnik, D.B. and Hosseini, M. The ethics of using artificial intelligence in scientific research: new guidance needed for a new tool. *AI and Ethics*. **2025**. 5(2): 1499–1521.
4. Ferrara, E. Fairness and bias in artificial intelligence: A brief survey of sources, impacts, and mitigation strategies. *Sci*. **2023**. 6(1): 3.
5. Al-Kfairy, M., Mustafa, D., et al. Ethical challenges and solutions of generative AI: An interdisciplinary perspective. in *Informatics*. 2024. MDPI.
6. Haibe-Kains, B., Adam, G.A., et al. Transparency and reproducibility in artificial intelligence. *Nature*. **2020**. 586(7829): E14–E16.
7. Al-Qudimat, A.R., Fares, Z.E., et al. Advancing medical research through artificial intelligence: progressive and transformative strategies: a literature review. *Health science reports*. **2025**. 8(2): e70200.
8. Hernandez-Boussard, T., Lee, A.Y., et al. Promoting transparency in AI for biomedical and behavioral research. *Nature medicine*. **2025**. 31(6): 1733–1734.
9. Zhou, S., Feng, J., et al., The rise of artificial intelligence in medicine: a new challenge for Chinese healthcare. **2026**, LWW. p. 2246–2249.
10. Blau, W., Cerf, V.G., et al., Protecting scientific integrity in an age of generative AI. **2024**, National Academy of Sciences. p. e2407886121.
11. Hamet, P. and Tremblay, J. Artificial intelligence in medicine. *metabolism*. **2017**. 69: S36–S40.
12. Malik, P., Pathania, M., and Rathaur, V.K. Overview of artificial intelligence in medicine. *Journal of family medicine and primary care*. **2019**. 8(7): 2328.
13. Jin, Q., Leaman, R., and Lu, Z. PubMed and beyond: biomedical literature search in the age of artificial intelligence. *EBioMedicine*. **2024**. 100.
14. Li, Y., Datta, S., et al. Enhancing systematic literature reviews with generative artificial intelligence: development, applications, and performance evaluation. *Journal of the American Medical Informatics Association*. **2025**. 32(4): 616–625.
15. Khalifa, M. and Albadawy, M. Using artificial intelligence in academic writing and research: An essential productivity tool. *Computer methods and programs in biomedicine update*. **2024**. 5: 100145.
16. Filetti, S., Fenza, G., and Gallo, A. Research design and writing of scholarly articles: New artificial intelligence tools available for researchers. *Endocrine*. **2024**. 85(3): 1104–1116.
17. Wilson, A. and Anwar, M.R. The future of adaptive machine learning algorithms in high-dimensional data processing. *International Transactions on Artificial Intelligence*. **2024**. 3(1): 97–107.
18. Panayides, A.S., Amini, A., et al. AI in medical imaging informatics: current challenges and future directions. *IEEE journal of biomedical and health informatics*. **2020**. 24(7): 1837–1857.
19. Obuchowicz, R., Strzelecki, M., and Piórkowski, A., Clinical applications of artificial intelligence in medical imaging and image processing—A review. **2024**, MDPI. p. 1870.
20. Kumar, I., Yadav, N., and Verma, A. Navigating artificial intelligence in scientific manuscript writing: Tips and traps. *Indian Journal of Radiology and Imaging*. **2025**. 35(S 01): S178–S186.

21. Pividori, M. and Greene, C.S. A publishing infrastructure for Artificial Intelligence (AI)-assisted academic authoring. *Journal of the American Medical Informatics Association*. **2024**. 31(9): 2103–2113.
22. Perlis, R.H., Christakis, D.A., et al. Artificial intelligence in peer review. *JAMA*. **2025**. 334(17): 1520–1522.
23. Checco, A., Bracciale, L., et al. AI-assisted peer review. *Humanities and social sciences communications*. **2021**. 8(1): 25.
24. Kousha, K. and Thelwall, M. Artificial intelligence to support publishing and peer review: A summary and review. *Learned Publishing*. **2024**. 37(1): 4–12.
25. Mazloum, T. and Shahin, B. AI and Research: Benefits and Challenges. *Clinical Oral Science and Dentistry*. **2024**. 6(3): 1–10.
26. Organization, W.H., Ethics and governance of artificial intelligence for health: WHO guidance. **2021**: World Health Organization.
27. Hosseini, M., Resnik, D.B., and Holmes, K. The ethics of disclosing the use of artificial intelligence tools in writing scholarly manuscripts. *Research ethics*. **2023**. 19(4): 449–465.
28. Celik, S.U. Integrating artificial intelligence into scientific writing: a narrative review for clinical and surgical researchers. *The American Journal of Surgery*. **2025**: 116657.
29. Pearson, G.S., Artificial intelligence and publication ethics. **2024**, SAGE Publications Sage CA: Los Angeles, CA. p. 453–455.
30. Yoo, J.-H. Defining the boundaries of AI use in scientific writing: a comparative review of editorial policies. *Journal of Korean Medical Science*. **2025**. 40(23).
31. Joo, H., Radnaabaatar, M., and Jung, J. Large language models and the future of peer-reviewed medical publishing: challenges, editorial responses, and a framework for responsible integration. *Journal of Korean Medical Science*. **2026**. 41(29): e306.
32. Lendvai, G.F. and Petra, A. Artificial intelligence in academic practices and policy discourses across 'Big 5' publishers. *Research Evaluation*. **2026**. 35: rvag004.
33. Hsu, H.-Y., Hakouz, A., and Fotouhi, G. Towards responsible generative AI in academia: A synthesis of AI policies on academic writing in the field of educational research. *AI and Ethics*. **2025**. 5(5): 5467–5484.
34. Sattari-Ardabili, F. and De Hoyos-Guevara, A.J. Artificial intelligence in scientific research: an analysis from the ethics and responsibility perspective. *Sophia Research Review*. **2026**. 3(1): 50–56.
35. BaHammam, A.S., The transparency paradox: why researchers avoid disclosing AI assistance in scientific writing. **2025**, Taylor & Francis. p. 2569–2574.
36. Leung, T.I., de Azevedo Cardoso, T., et al., Best practices for using AI tools as an author, peer reviewer, or editor. **2023**, JMIR Publications Toronto, Canada. p. e51584.
37. Kirkham, A.A., Mackey, J.R., et al. TITAN trial: a randomized controlled trial of a cardiac rehabilitation care model in breast cancer. *JACC: Advances*. **2023**. 2(6): 100424.
38. Rivera, A., Abhari, K., and Xiao, B. Responsible AI design: the authenticity, control, transparency theory. *Journal of the Association for Information Systems*. **2025**. 26(5): 1337–1389.
39. Van Zoonen, W., Tursunbayeva, A., and Morgan-Thomas, A. Beyond AI disclosure: Claim accountability and responsible research in scholarly publishing. *European Management Journal*. **2026**.
40. Bear, C., Mukherjee, A., and Sonnino, R. Challenges for policymakers in enabling transparency for sustainable, healthy and safe food systems. **2025**.
41. Chatila, R., Dignum, V., et al., Trustworthy ai, in *Reflections on artificial intelligence for humanity*. **2021**, Springer. p. 13–39.
42. Cheimarios, N. Scientific software development in the AI era: reproducibility, MLOps, and applications in soft matter physics. *Frontiers in Physics*. **2025**. 13: 1711356.
43. Ruttenberg, D. The AIR framework for research transparency: a critical analysis of stage-specific AI disclosure in the context of accessibility and research integrity. *AI & SOCIETY*. **2026**: 1–14.
44. Raji, I.D., Smart, A., et al. Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing. in *Proceedings of the 2020 conference on fairness, accountability, and transparency*. 2020.
45. McIntosh, T.R., Susnjak, T., et al. From google gemini to openai gpt-4o: A survey on reshaping the generative artificial intelligence (ai) research landscape. *Technologies*. **2025**. 13(2): 51.
46. Bear, C., Mukherjee, A., and Cifuentes, M.L. Transparency solutions for healthy, sustainable and safe food systems in existing policy. **2023**.

47. Avery, J.J., Abril, P.S., and Del Riego, A. Attributing AI authorship: towards a system of icons for legal and ethical disclosure. *Nw. J. Tech. & Intell. Prop.* **2024**. 22: 1.
48. Liu, X., Rivera, S.C., et al. Reporting guidelines for clinical trial reports for interventions involving artificial intelligence: the CONSORT-AI extension. *The Lancet Digital Health*. **2020**. 2(10): e537–e548.
49. Ng, J.Y. A structured framework for effective and responsible generative artificial intelligence chatbot prompt engineering throughout the scientific process: a comprehensive guide for the health and medical researcher. *Frontiers in Artificial Intelligence*. **2026**. 9: 1745928.
50. Gaube, S., Langer, M., et al. Keeping an Eye on AI: A Framework for Effective Human Oversight of AI Systems. *arXiv preprint arXiv:2605.16278*. **2026**.
51. Gong, E.J., Bang, C.S., and Shin, Y.S. Applications of large language models in medical research: from systematic reviews to clinical studies. *Bioengineering*. **2026**. 13(3): 365.
52. Boretti, A. Hallucinations in generative AI: A threat to scholarly integrity and the urgent need for publisher-led academically supervised verification. *Energy Research & Social Science*. **2026**. 136: 104720.
53. Stahl, B.C. and Wright, D. Ethics and privacy in AI and big data: Implementing responsible research and innovation. *IEEE Security & Privacy*. **2018**. 16(3): 26–33.
54. Li, Z. Ethical frontiers in artificial intelligence: navigating the complexities of bias, privacy, and accountability. *International Journal of Engineering and Management Research*. **2024**. 14(3): 109–116.
55. White, M., Haddad, I., et al. The model openness framework: Promoting completeness and openness for reproducibility, transparency, and usability in artificial intelligence. *arXiv preprint arXiv:2403.13784*. **2024**.
56. Albertoni, R., Colantonio, S., and Stefanowski, J. Reproducibility of machine learning: Terminology, recommendations and open issues. *arXiv preprint arXiv:2302.12691*. **2023**.
57. Trujillo, J., Lavalle, A., et al. An integrated requirements framework for analytical and AI projects. *Data & Knowledge Engineering*. **2025**: 102493.
58. Ranjitsingh, L.M. and Rao, T.S. Establish legal and regulatory standards for the testing and validation of AI systems to ensure their reliability and safety in operational environments. *International Journal of System Assurance Engineering and Management*. **2025**. 16(10): 3338–3353.
59. Zhang, J. and Zhang, Z.-m. Ethics and governance of trustworthy medical artificial intelligence. *BMC medical informatics and decision making*. **2023**. 23(1): 7.
60. Mehrtabar, S., Marey, A., et al. Ethical considerations in patient privacy and data handling for AI in cardiovascular imaging and radiology. *Journal of Imaging Informatics in Medicine*. **2025**: 1–22.
61. Chawla, R., Maheshwari, A., et al. RSSDI IJDDC position statement on embracing artificial intelligence responsibly in medical research and scientific publication: consensus recommendations for authors, editors, and peer reviewers. *International Journal of Diabetes in Developing Countries*. **2026**: 1–12.
62. Gatrell, C., Muzio, D., et al., Here, there and everywhere: On the responsible use of artificial intelligence (AI) in management research and the peer-review process. **2024**, Wiley Online Library. p. 739–751.
63. Nagendran, M., Chen, Y., et al. Artificial intelligence versus clinicians: systematic review of design, reporting standards, and claims of deep learning studies. *bmj*. **2020**. 368.
64. Giray, L., Sevnarayan, K., et al. AI slop in academic publishing: History, characteristics, manifestations, causes, and mitigation strategies. *Internet Reference Services Quarterly*. **2026**. 30(2): 221–244.
65. Treviño, M. and Arias-Carrión, O. Artificial intelligence in systematic reviews: Overcoming reproducibility, bias and validation challenges. *Preprints*. **2025**. 2025061895.
66. Ganjavi, C., Eppler, M.B., et al. Publishers' and journals' instructions to authors on use of generative artificial intelligence in academic and scientific publishing: bibliometric analysis. *bmj*. **2024**. 384.
67. da Veiga, A. Ethical guidelines for the use of generative artificial intelligence and artificial intelligence-assisted tools in scholarly publishing: a thematic analysis. *Science Editing*. **2025**. 12(1): 28–34.
68. Al-Jarf, R. To publish or not to publish AI-generated research articles in scholarly journals: A perspective from editors and publishers. in *International Conference on Artificial Intelligence and its Applications in the Age of Digital Transformation*. 2025. Springer.

69. Flanagin, A., Kendall-Taylor, J., and Bibbins-Domingo, K. Guidance for authors, peer reviewers, and editors on use of AI, language models, and chatbots. *Jama*. **2023**. 330(8): 702–703.
70. Mann, S.P., Aboy, M., et al. AI and the future of academic peer review. *arXiv preprint arXiv:2509.14189*. **2025**.
71. Vadner, D.J., Brown, A.N., and Gold, M.H. Artificial Intelligence in medical research and publishing: progress, risks, and future perspectives. *La Presse Médicale*. **2026**: 104371.
72. Benítez-Hidalgo, A., Barba-González, C., et al. TITAN: A knowledge-based platform for Big Data workflow management. *Knowledge-Based Systems*. **2021**. 232: 107489.
73. Marijan, D. and Gotlieb, A. Industry-Academia research collaboration in software engineering: The Certus model. *Information and software technology*. **2021**. 132: 106473.
74. Laubach, M., Cheers, G.M., et al. Clinical translation of 3D-printed patient-specific bone implants: a consensus statement. *International Journal of Surgery (London, England)*. **2025**. 111(11): 7497.
75. Sciences, N.A.o., Medicine, et al., Reproducibility and replicability in science. **2019**: National Academies Press.
76. Kaminski, M.E. Binary governance: Lessons from the GDPR's approach to algorithmic accountability. *S. Cal. L. Rev*. **2018**. 92: 1529.
77. Tardin, M.G., Perin, M.G., et al. Organizational sustainability orientation: A review. *Organization & Environment*. **2024**. 37(2): 298–324.
78. Alnaimat, F., Al-Halaseh, S., and AlSamhori, A.R.F. Evolution of research reporting standards: adapting to the influence of artificial intelligence, statistics software, and writing tools. *Journal of Korean Medical Science*. **2024**. 39(32).
79. Cheng, K., Wu, S., et al. An artificial intelligence-enhanced coaching mode. *International Journal of Surgery (London, England)*. **2025**. 111(9): 6469.
80. Sun, Y., Sheng, D., et al. AI hallucination: towards a comprehensive classification of distorted information in artificial intelligence-generated content. *Humanities and Social Sciences Communications*. **2024**. 11(1): 1278.
81. Momani, A. Implications of artificial intelligence on health data privacy and confidentiality. *arXiv preprint arXiv:2501.01639*. **2025**.
82. Koçak, B., Ponsiglione, A., et al. Bias in artificial intelligence for medical imaging: fundamentals, detection, avoidance, mitigation, challenges, ethics, and prospects. *Diagnostic and interventional radiology*. **2025**. 31(2): 75.
83. Cross, J.L., Choma, M.A., and Onofrey, J.A. Bias in medical AI: implications for clinical decision-making. *PLOS digital health*. **2024**. 3(11): e0000651.
84. Miyakawa, T. No raw data, no science: another possible source of the reproducibility crisis. *Molecular brain*. **2020**. 13(1): 24.
85. Perkins, M. and Roe, J. Academic publisher guidelines on AI usage: A ChatGPT supported thematic analysis. *F1000Research*. **2024**. 12: 1398.
86. Labkoff, S., Oladimeji, B., et al. Toward a responsible future: recommendations for AI-enabled clinical decision support. *Journal of the American Medical Informatics Association*. **2024**. 31(11): 2730–2739.
87. Hanna, M.G., Pantanowitz, L., et al. Future of artificial intelligence—machine learning trends in pathology and medicine. *Modern Pathology*. **2025**. 38(4): 100705.
88. Bandi, A., Kongari, B., et al. The rise of agentic ai: A review of definitions, frameworks, architectures, applications, evaluation metrics, and challenges. *Future Internet*. **2025**. 17(9): 404.
89. Acosta, J.N., Falcone, G.J., et al. Multimodal biomedical AI. *Nature medicine*. **2022**. 28(9): 1773–1784.
90. Huang, X., Dou, S., and Yin, Z. The dual-edged sword: artificial intelligence's evolving role in academic peer review. *Science China Information Sciences*. **2025**. 68(11).
91. Collins, G.S., Moons, K.G., et al. TRIPOD+ AI statement: updated guidance for reporting clinical prediction models that use regression or machine learning methods. *bmj*. **2024**. 385.
92. Hernandez-Boussard, T., Bozkurt, S., et al. MINIMAR (MINimum Information for Medical AI Reporting): Developing reporting standards for artificial intelligence in health care. *Journal of the American Medical Informatics Association*. **2020**. 27(12): 2011–2015.
93. Lai, H., Tong, B., et al. Development methodologies and methodological quality of artificial intelligence reporting guidelines in healthcare: a scoping review. *AI Evidence-Based Medicine*. **2026**. 1(1): 42–52.
94. Kolbinger, F.R., Veldhuizen, G.P., et al. Reporting guidelines in medical artificial intelligence: a systematic review and meta-analysis. *Communications Medicine*. **2024**. 4(1): 71.

95. Scuderi, G.R., Taunton, M.J., et al. The challenges with artificial intelligence in scientific writing. *The Journal of Arthroplasty*. **2026**. 41(2): 299–303.
96. Chigwada, J. and Ngulube, P. Use of artificial intelligence tools in the publishing process: expectations from publishers through author guidelines. *Frontiers in Research Metrics and Analytics*. **2026**. 11: 1740510.
97. Huh, S. Recent issues in medical journal publishing and editing policies: adoption of artificial intelligence, preprints, open peer review, model text recycling policies, best practice in scholarly publishing 4th version, and country names in titles. *Neurointervention*. **2023**. 18(1): 2–8.

Disclaimer/Publisher's Note: The statements, opinions, and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of RES Publishers and/or the editor(s). RES Publishers and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.