



# Research Integrity Challenges in Generative Artificial Intelligence–Assisted Scholarship

**Dr. Sophia Williams**

Associate Professor, University of Melbourne, Melbourne, Australia

## Abstract

Generative artificial intelligence has rapidly become embedded in scholarly workflows, supporting literature exploration, idea development, language refinement, coding, data interpretation, manuscript preparation, reference discovery, and editorial processes. These capabilities can improve accessibility and research efficiency, particularly for multilingual researchers and scholars working with complex information environments. At the same time, generative AI introduces research-integrity risks that differ from conventional plagiarism or authorship disputes because inaccurate, fabricated, biased, confidential, or inadequately attributed content can be produced at scale while retaining a persuasive scholarly appearance.

Current guidance from the Committee on Publication Ethics, International Committee of Medical Journal Editors, and World Association of Medical Editors maintains that AI systems cannot assume authorship responsibility and that human authors remain accountable for generated material. ICMJE further requires disclosure of AI-assisted technologies used in manuscript production, while publisher policies increasingly emphasize verification, confidentiality, transparency, and human oversight.

This study develops a research-integrity governance framework for generative AI–assisted scholarship using conceptual analysis and a transparent simulation-based risk assessment across six stages of the research lifecycle. Six integrity dimensions are examined: factual reliability, citation integrity, authorship accountability, methodological transparency, confidentiality and data protection, and intellectual contribution.

The simulated analysis indicates that citation and reference management presents the highest illustrative risk score, followed by literature review, manuscript drafting, peer review, and data analysis. The paper argues that responsible AI-assisted scholarship should not be governed through blanket prohibition or unrestricted acceptance. Instead, institutions, researchers, reviewers, and journals require a proportional framework combining disclosure, source verification, traceable human decision-making, secure data practices, reproducibility, authorship responsibility, and editorial quality assurance. The proposed model distinguishes legitimate AI assistance from practices that obscure scholarly contribution or compromise the reliability of the research record.

**Keywords:** generative artificial intelligence, research integrity, scholarly publishing, academic authorship, citation hallucination, research ethics, peer review, AI disclosure, scientific misconduct, scholarly communication



## 1. Introduction

Generative artificial intelligence has altered the practical environment in which scholarly knowledge is produced. Large language models and related systems can generate prose, summarize documents, suggest analytical approaches, translate technical material, create programming code, reorganize arguments, and support reference discovery. Khalifa and Albadawy's systematic review of AI in academic writing describes potential applications spanning idea generation, organization of information, literature-related tasks, data management, and writing support. These capabilities can lower linguistic barriers and reduce repetitive academic work, yet they also redistribute cognitive tasks that traditionally provided evidence of a researcher's intellectual contribution.

Research integrity becomes particularly important because generative systems do not operate as conventional databases that simply retrieve verified knowledge. Their outputs may combine accurate information, plausible synthesis, missing context, fabricated details, or nonexistent references in language that appears authoritative. Recent scholarship has documented growing concern regarding hallucinated citations, fabricated scholarly records, and AI-generated manuscripts containing misattributed or unverifiable information. Resnik and colleagues specifically examined whether hallucinated citations produced through generative AI can become sufficiently serious to intersect with established concepts of research misconduct, while Spinellis demonstrated how AI-generated and falsely attributed scholarly material can expose weaknesses in publishing verification systems.

The integrity challenge is wider than factual hallucination. Generative AI can influence research questions, literature selection, analytical interpretation, methodological description, scholarly argument, peer review, and editorial decision support. Bjelobaba and colleagues identified ethical concerns across multiple research phases, including transparency deficiencies, fabrication, bias, privacy, copyright, and inadequate disclosure. Vasconcelos similarly argues that the growth of generative AI requires research-integrity governance to reconsider the relationship between technological assistance and human agency. The central concern is therefore not simply whether researchers use AI, but whether they can remain epistemically and professionally responsible for work partly mediated by systems whose outputs may be difficult to reconstruct.

Leading publishing organizations have begun to establish clearer boundaries. COPE states that AI tools cannot qualify as authors because they cannot assume responsibility for submitted work. ICMJE requires authors to disclose AI-assisted technologies used in manuscript preparation and maintains that human authors remain responsible for accuracy, originality, attribution, and integrity. WAME likewise rejects chatbot authorship and emphasizes transparency in the use of generative systems. Nature Portfolio currently requires human accountability and disclosure while imposing confidentiality requirements on reviewer use, and Elsevier requires declarations for substantive generative-AI assistance in manuscript preparation.

## 2. Objectives of the Study

### 2.1 To Identify Research-Integrity Risks Across the Scholarly Lifecycle

The first objective is to identify how generative AI may affect research integrity from problem formulation to publication. The study examines risks arising during research design, literature review, data analysis, manuscript drafting, peer review, and reference management.



Particular attention is given to fabricated information, selective or distorted literature synthesis, unverified citations, methodological opacity, inappropriate AI authorship, undisclosed assistance, confidentiality breaches, synthetic data presented as observed evidence, and excessive delegation of scholarly judgment.

## 2.2 To Distinguish Legitimate AI Assistance from Integrity-Compromising Use

The second objective is to develop criteria for differentiating acceptable AI assistance from practices that undermine intellectual accountability.

The analysis rejects the simplistic assumption that all AI-assisted writing constitutes misconduct. Language refinement, translation, coding assistance, structured brainstorming, or carefully verified literature support may be legitimate where permitted and transparently disclosed. The integrity threshold is crossed when AI use conceals authorship responsibility, creates unverifiable evidence, compromises confidential information, misrepresents methods, fabricates citations or results, or substitutes automated output for the scholarly judgment claimed by human authors.

## 2.3 To Develop an Operational Governance Framework

The third objective is to formulate a practical governance framework for researchers, universities, journals, reviewers, and publishers.

### The framework emphasizes:

- human accountability;
- transparent disclosure;
- source-level verification;
- reproducibility;
- methodological documentation;
- secure handling of unpublished information;
- authorship responsibility;
- proportional editorial screening; and
- correction of the scholarly record where AI-related errors are discovered.

## 3. Review of Literature

### 3.1 Generative AI, Authorship, and Scholarly Accountability

Authorship in academic publishing is not merely recognition for producing text. It establishes intellectual responsibility for research design, interpretation, accuracy, originality, conflicts of interest, and responses to questions regarding integrity.

This distinction explains why major publication-ethics bodies reject AI authorship. COPE states that AI tools cannot meet authorship requirements because they cannot assume responsibility for submitted work or manage conflicts of interest.

ICMJE similarly states that authors should not list or cite AI-assisted technologies as authors and should disclose how such technologies were used. Its guidance places responsibility for checking AI-generated material, including attribution and potential plagiarism, on the human authors.



WAME's recommendations follow the same principle. Chatbots cannot satisfy conventional authorship criteria because authorship requires accountability that can be assigned to human contributors.

A system may produce substantial quantities of manuscript text without possessing responsibility for:

- why a research question was selected;
- whether evidence was interpreted correctly;
- whether references exist;
- whether participant rights were protected;
- whether reported data are authentic;
- whether competing interests exist; or
- whether correction or retraction is required.

The challenge becomes more complicated when extensive AI assistance blurs the boundary between language support and intellectual contribution. If researchers provide a dataset and prompt an AI system to select hypotheses, interpret findings, generate theoretical arguments, formulate conclusions, and construct the majority of a manuscript, a disclosure stating merely that AI was used for "writing assistance" may inadequately describe the true intellectual workflow.

### 3.2 Hallucinated Evidence, Citation Integrity, and Reproducibility

Citation integrity has emerged as one of the most visible weaknesses associated with generative scholarly tools.

Large language models may generate references that appear formally credible yet contain nonexistent article titles, incorrect author combinations, false journal names, inaccurate publication years, invalid identifiers, or distorted descriptions of genuine studies.

Resnik and colleagues' 2026 discussion of hallucinated citations emphasizes that fabricated references can have consequences beyond ordinary citation mistakes because they may introduce nonexistent evidence into scholarly argument.

Nature reported in April 2026 that hallucinated citations are increasingly contaminating scientific literature and that large numbers of publications may contain invalid AI-generated references. This is particularly important because citation errors are cumulative: once a nonexistent or distorted reference enters published literature, subsequent researchers may reproduce it without verifying the original record.

Spinellis's case study similarly showed how AI-generated and misattributed scholarly articles can challenge established mechanisms such as DOI and author identification, illustrating that apparently plausible bibliographic information may bypass conventional editorial review.

Elsevier's current generative-AI guidance explicitly warns that AI-generated references may be incorrect, hallucinated, or fabricated and therefore requires authors to verify reference-related outputs. It also expects AI-supported literature selection or editing to be appropriately described when applicable.

**Citation hallucination is only one component of a wider reproducibility problem. AI may also:**

- invent numerical values;
- misstate sample characteristics;
- propose statistical procedures not actually performed;
- generate fictitious quotations;



- describe nonexistent supplementary material;
- modify methodological terminology;
- produce plausible but incorrect code; or
- infer causal interpretations unsupported by the study design.

Pellegrina and colleagues argue that AI presents a dual research-integrity role: it can introduce errors and misconduct risks, but AI tools can also be developed to detect statistical problems, citation errors, image manipulation, and other integrity concerns.

### 3.3 Confidentiality, Peer Review, and Editorial Integrity

Peer review introduces a different category of AI risk because manuscripts under review are generally confidential materials.

Uploading unpublished manuscripts, grant proposals, identifiable data, proprietary methods, or confidential reviewer information into external generative-AI systems can create privacy, intellectual-property, and confidentiality concerns.

Nature Portfolio's current AI policy states that reviewers remain responsible for independent expert evaluation, should not upload manuscript content to unsecured or public AI tools, and should not delegate scholarly judgment to AI.

Springer Nature likewise advises peer reviewers not to upload manuscripts into generative-AI tools because submitted material can contain sensitive or proprietary information and because AI output may be false, biased, or incomplete.

Elsevier adopts an even more restrictive position for standard external generative-AI use in peer review, emphasizing that critical evaluation and original judgment remain human responsibilities and that confidentiality must be protected.

The problem is not limited to privacy. If a reviewer asks an AI system to generate the substantive critique of a manuscript and submits that output with minimal independent evaluation, peer review becomes procedurally misleading. The editor assumes the report reflects an expert's disciplinary judgment when substantial reasoning may have been outsourced.

## 4. Research Methodology

### 4.1 Research Design and Analytical Approach

**Third, six major stages of the research lifecycle were selected for comparative assessment:**

1. research design;
2. literature review;
3. data analysis;
4. manuscript drafting;
5. peer review; and
6. citation and reference management.

Fourth, illustrative risk values were assigned using a weighted model to demonstrate a governance instrument that could later be tested with empirical evidence.



The framework is designed for reproducibility. A journal or university could replace the simulated values with observed indicators derived from manuscript audits, surveys, integrity investigations, reviewer assessments, or institutional records.

## 4.2 Integrity-Risk Dimensions and Weighting

Six dimensions were used to evaluate AI-assisted scholarly activity.

**Table 1: Proposed Research-Integrity Risk Assessment Framework**

| <b>Integrity Dimension</b>          | <b>Weight</b> | <b>Core Question</b>                                                                  | <b>Illustrative Indicators</b>                                       |
|-------------------------------------|---------------|---------------------------------------------------------------------------------------|----------------------------------------------------------------------|
| Factual Reliability                 | 20%           | Could AI-generated claims be inaccurate, fabricated, or contextually misleading?      | Hallucinations, unsupported claims, numerical inaccuracies           |
| Citation Integrity                  | 20%           | Are references authentic, accurately represented, and independently verified?         | Nonexistent references, false DOI, citation distortion               |
| Authorship Accountability           | 15%           | Can identifiable human researchers accept responsibility for the intellectual work?   | Undisclosed delegation, AI authorship, unclear contribution          |
| Methodological Transparency         | 20%           | Can readers reconstruct how AI affected the research process?                         | Missing AI disclosure, hidden prompts, undocumented transformations  |
| Confidentiality and Data Protection | 15%           | Could sensitive, unpublished, or personal information be exposed?                     | Manuscript uploads, participant data, proprietary information        |
| Intellectual Contribution           | 10%           | Does the published work accurately represent the researcher's scholarly contribution? | Excessive automated drafting, uncritical synthesis, ghost generation |
| <b>Total</b>                        | <b>100%</b>   | <b>Overall research-integrity exposure</b>                                            | <b>Lifecycle risk assessment</b>                                     |

**Source:** Author-developed conceptual framework informed by current publication-ethics and AI-governance guidance.



F = factual reliability risk C = citation-integrity risk A = authorship-accountability risk M = methodological-transparency risk D = confidentiality and data-protection risk I = intellectual-contribution risk

Each dimension is normalized on a 0–100 scale.

## 4.3 Risk Classification and Governance Thresholds

The proposed model uses four interpretive levels.

- A score from 0 to 39 represents low integrity risk. AI assistance may generally be acceptable subject to ordinary verification and journal policy.
- A score from 40 to 59 represents moderate integrity risk. The activity should require disclosure, independent checking, and documentation of human oversight.
- A score from 60 to 74 represents high integrity risk. Strong verification, explicit disclosure, reproducibility documentation, and editorial review should be required.
- A score of 75 or above represents very high integrity risk. AI-generated outputs should not be incorporated without independent source-level verification, clear human responsibility, and documented safeguards.

These thresholds are analytical tools rather than legal or universal publication standards.

## 5. Analysis

### 5.1 Simulated Research-Integrity Risk Across the Research Lifecycle

The simulation indicates that generative AI does not create an equal level of integrity risk across every scholarly activity.

**Table 2: Simulated Integrity-Risk Scores for Generative AI–Assisted Research Activities**

| Research Stage    | Simulated Risk Score | Risk Level | Principal Integrity Concern                                | Recommended Control                                          |
|-------------------|----------------------|------------|------------------------------------------------------------|--------------------------------------------------------------|
| Research Design   | 58                   | Moderate   | Automated framing may introduce hidden assumptions or bias | Researcher justification and documented conceptual decisions |
| Literature Review | 78                   | Very High  | Selective synthesis, invented studies, distorted evidence  | Database verification and source-by-source checking          |
| Data Analysis     | 69                   | High       | Incorrect code, statistical                                | Reproducible scripts and                                     |

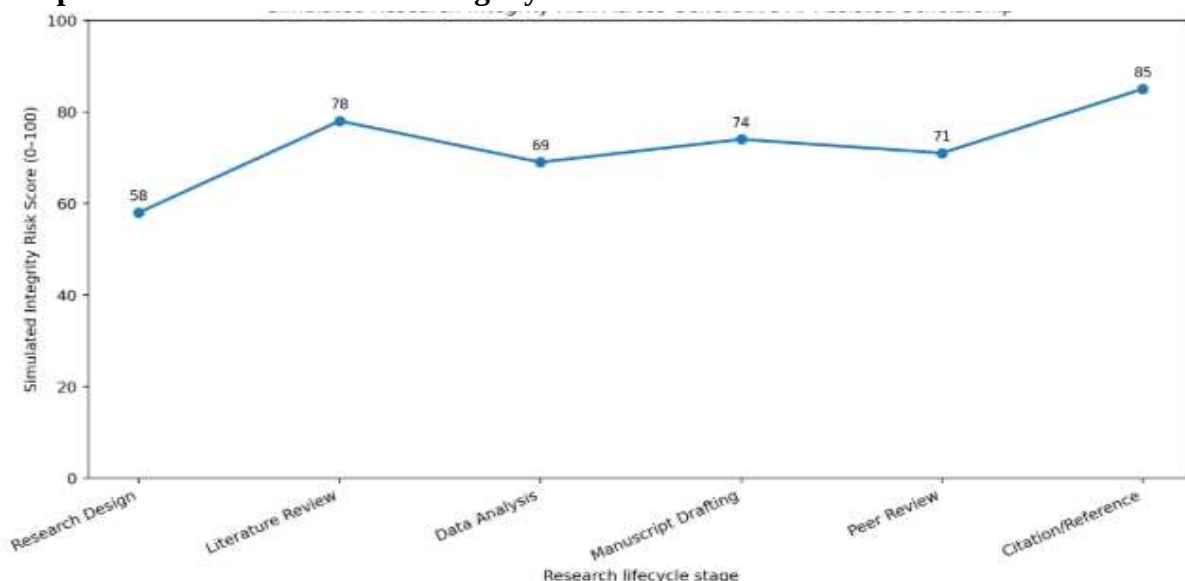


| Research Stage                | Simulated Risk Score | Risk Level | Principal Integrity Concern                            | Recommended Control                                           |
|-------------------------------|----------------------|------------|--------------------------------------------------------|---------------------------------------------------------------|
|                               |                      |            | interpretation, invented values                        | independent validation                                        |
| Manuscript Drafting           | 74                   | High       | Undisclosed intellectual delegation and factual error  | AI-use declaration and line-level author verification         |
| Peer Review                   | 71                   | High       | Confidentiality loss and delegation of expert judgment | Secure tools only; reviewer remains independently responsible |
| Citation/Reference Management | 85                   | Very High  | Hallucinated, corrupted, or misrepresented references  | Mandatory bibliographic verification                          |

**Note:** Scores are simulated solely to demonstrate the proposed methodology and must not be interpreted as measured prevalence rates.

## 5.2 Graphical Analysis

**Graph 1: Simulated Research-Integrity Risk Across Generative AI-Assisted Scholarship**





**Source:** Author-generated simulation.

The line graph demonstrates a nonuniform integrity-risk pattern. Risk rises considerably from research design to literature review, declines during data analysis, increases again in manuscript preparation, remains elevated during peer review, and reaches its maximum at citation and reference management.

The graph should not be interpreted to mean that citation hallucination is empirically 85% likely or that research design is 58% risky. The values represent an analytical scenario used to test governance priorities.

Key analytical finding: The closer an AI-supported activity comes to making verifiable claims about evidence, sources, or scholarly judgment, the stronger the need for independent human verification.

### 5.3 Proposed Human-Accountability Governance Model

The analysis supports a seven-control framework.

#### **Control 1: Declare Meaningful Generative-AI Use**

Disclosure should identify:

- the AI system used;
- the research stage in which it was used;
- the purpose of use;
- whether outputs were incorporated into the manuscript;
- whether generated references or analyses were independently verified.

ICMJE expects authors to disclose AI-assisted technologies, while current publisher policies increasingly require manuscript-level declarations for substantive generative-AI use.

#### **Control 2: Apply a Human Verification Rule**

No factual claim, statistic, quotation, citation, methodological description, or interpretation should enter the final scholarly record solely because an AI system generated it.

The responsible researcher should verify the statement against primary evidence.

#### **Control 3: Verify Every AI-Suggested Reference**

Verification should include:

- author names;
- article title;
- journal or publisher;
- publication year;
- DOI or persistent identifier;
- existence of the source;
- consistency between the cited source and the claim attributed to it.

This control addresses the documented problem of hallucinated citations.



## **Control 4: Preserve Research Provenance**

Researchers should document substantial AI involvement in analytical workflows.

Where an AI tool generated or modified code, transformed qualitative data, assisted in classification, or influenced methodological decisions, sufficient documentation should be retained to permit reconstruction and verification.

## **Control 5: Protect Confidential and Unpublished Material**

Unpublished manuscripts, participant-level data, grant proposals, patent-sensitive material, identifiable health information, confidential peer-review documents, and protected institutional data should not be entered into systems lacking appropriate security and data-governance safeguards.

Nature Portfolio and Springer Nature specifically emphasize reviewer confidentiality and restrictions on uploading manuscripts into unsecured AI environments.

## **Control 6: Preserve Human Intellectual Responsibility**

Researchers should be able to explain and defend:

- the research question;
- theoretical argument;
- methodological choices;
- data-analysis strategy;
- interpretation;
- conclusions; and
- cited evidence.

If an author cannot explain a substantive passage because it was generated automatically, the problem is not stylistic; it is a failure of scholarly responsibility.

## **6. Discussion**

### **6.1 Research Integrity Should Focus on Accountability, Not AI Detection Alone**

The rapid spread of generative artificial intelligence has increased interest in automated AI-text detection systems. Although these tools may provide useful signals for preliminary editorial screening, research integrity cannot be reduced to identifying whether particular sentences appear to have been generated by a machine.

The more important questions concern whether the evidence is authentic, whether the reported methods were actually performed, whether references were accurately verified, whether authors understand and can defend the arguments presented, whether AI assistance was disclosed when required, whether confidential information was protected, and whether the research can be independently scrutinized.

A manuscript written entirely by humans can still contain fabricated data, plagiarism, manipulated findings, or unverifiable references, whereas an AI-assisted manuscript may remain ethically responsible when researchers maintain intellectual control over the work. Responsible use requires authors to verify substantive claims, critically evaluate AI-generated material, disclose meaningful assistance, protect sensitive information, and comply with relevant publisher and



institutional requirements. Therefore, AI detection should be treated, at most, as one component of a broader research-integrity assessment.

## 6.2 Disclosure Must Be Specific Enough to Be Meaningful

Transparency is frequently proposed as a central response to generative AI use in scholarship, but a generic disclosure statement may provide insufficient information about how AI contributed to a research project. A statement such as “AI was used during manuscript preparation” does not establish whether the system was used for grammar correction, translation, outlining, literature searching, reference suggestions, code generation, statistical interpretation, theoretical reasoning, or development of conclusions. Meaningful disclosure should therefore distinguish between different levels of assistance, including language-level support such as grammar and stylistic refinement; organizational support such as outlining and restructuring; knowledge-support activities such as literature searching and summarization; analytical assistance such as coding, classification, modelling, or statistical interpretation; and substantive intellectual assistance involving hypotheses, theoretical reasoning, interpretation, or conclusions.

The degree of disclosure and verification should increase as the substantive contribution of AI becomes greater. This approach allows readers, editors, reviewers, and institutions to understand the actual role played by generative AI rather than treating all forms of assistance as equivalent. Current publisher policies similarly distinguish routine language or spelling assistance from substantive generative-AI use that requires declaration. Nature Portfolio also emphasizes that authors retain responsibility for originality, accuracy, and research integrity when AI tools are used. Disclosure should therefore function as an accountability mechanism that enables meaningful evaluation of the research process rather than as a formal or ceremonial statement included only to satisfy publication requirements.

## 6.3 Institutional Responsibility Extends Beyond Individual Authors

Generative-AI-related research integrity should not be understood solely as a matter of individual researcher behaviour. Universities establish research policies and training requirements, supervisors influence research practices, journals define disclosure standards, publishers control editorial infrastructures, reviewers evaluate methodological quality, funders shape research incentives, and indexing and academic-evaluation systems can influence pressures to publish. Because these actors collectively shape the research environment, effective governance requires coordinated institutional responsibility.

Publishers and editorial offices likewise have an important role in developing balanced policies. Regulations that are excessively restrictive may discourage legitimate uses of AI for accessibility, language improvement, or technical assistance, whereas overly permissive policies may weaken accountability and make inappropriate AI use difficult to identify. Importantly, AI should not be viewed only as a threat to research integrity; it can also support integrity by assisting with the identification of statistical anomalies, problematic citations, manipulated images, and other potential weaknesses in scholarly work. The desirable future is therefore not scholarship without AI, but scholarship in which AI use remains traceable, verifiable, proportionate, transparently disclosed, and subordinate to accountable human judgment.

## 7. Conclusion

Generative artificial intelligence has become part of contemporary scholarly practice, and attempts to treat it as a temporary anomaly are increasingly unrealistic. Its capacity to support language, information



organization, coding, analysis, and manuscript development can improve research efficiency and widen participation in academic communication. These benefits, however, do not remove the foundations on which research credibility depends.

The researcher who signs the manuscript remains responsible for understanding the work, verifying the evidence, defending the methods, checking the references, and correcting the record when errors occur.

Future research should empirically test the proposed risk framework using multidisciplinary researcher surveys, journal-submission audits, reference-verification datasets, reviewer studies, and institutional case investigations. Comparative research should also examine whether disciplinary differences require different levels of disclosure and whether structured AI-use declarations improve reader trust and editorial decision-making.

Research integrity in the generative-AI era will ultimately depend less on whether scholarship contains AI-assisted language than on whether the scholarly community preserves verifiability, transparency, accountability, and human intellectual responsibility as non-negotiable conditions of credible research.

## References

1. Floridi, L., & Cowls, J. (2019). A unified framework of five principles for AI in society. *Harvard Data Science Review*, 1(1).
2. Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People—An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28, 689–707.
3. Mittelstadt, B. D., Allo, P., Taddeo, M., Wachter, S., & Floridi, L. (2016). The ethics of algorithms: Mapping the debate. *Big Data & Society*, 3(2), 1–21.
4. Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1, 389–399.
5. European Commission High-Level Expert Group on Artificial Intelligence. (2019). *Ethics Guidelines for Trustworthy AI*. European Commission.
6. IEEE. (2019). *Ethically Aligned Design: A Vision for Prioritizing Human Well-being with Autonomous and Intelligent Systems*. IEEE Standards Association.
7. Cath, C., Wachter, S., Mittelstadt, B., Taddeo, M., & Floridi, L. (2018). Artificial intelligence and the “Good Society”: The US, EU, and UK approach. *Science and Engineering Ethics*, 24, 505–528.
8. Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623.
9. Burrell, J. (2016). How the machine “thinks”: Understanding opacity in machine learning algorithms. *Big Data & Society*, 3(1), 1–12.
10. Kroll, J. A., Huey, J., Barocas, S., Felten, E. W., Reidenberg, J. R., Robinson, D. G., & Yu, H. (2017). Accountable algorithms. *University of Pennsylvania Law Review*, 165(3), 633–705.



11. Mittelstadt, B. D. (2019). Principles alone cannot guarantee ethical AI. *Nature Machine Intelligence*, 1, 501–507.
12. Vollmer, S., Mateen, B. A., Bohner, G., Király, F. J., Ghani, R., & Jonsson, P. (2020). Machine learning and artificial intelligence research for patient benefit: 20 critical questions on transparency, replicability, ethics, and effectiveness. *BMJ*, 368, l6927.
13. Horbach, S. P. J. M., & Halfman, W. (2020). Innovating editorial practices: Academic publishers at work. *Research Integrity and Peer Review*, 5, 11.
14. Haustein, S. (2016). Grand challenges in altmetrics: Heterogeneity, data quality and dependencies. *Scientometrics*, 108, 413–423.
15. Ioannidis, J. P. A. (2005). Why most published research findings are false. *PLoS Medicine*, 2(8), e124.
16. Munafò, M. R., Nosek, B. A., Bishop, D. V. M., Button, K. S., Chambers, C. D., Percie du Sert, N., Simonsohn, U., Wagenmakers, E.-J., Ware, J. J., & Ioannidis, J. P. A. (2017). A manifesto for reproducible science. *Nature Human Behaviour*, 1, 0021.
17. Nosek, B. A., Alter, G., Banks, G. C., Borsboom, D., Bowman, S. D., Breckler, S. J., Buck, S., Chambers, C. D., Chin, G., Christensen, G., Contestabile, M., Dafoe, A., Eich, E., Freese, J., Glennerster, R., Goroff, D., Green, D. P., Hesse, B., Humphreys, M., ... Yarkoni, T. (2015). Promoting an open research culture. *Science*, 348(6242), 1422–1425.
18. Resnik, D. B., Elliott, K. C., & Miller, A. K. (2015). A framework for addressing ethical issues in scientific research. *Ethics & Biology Engineering & Medicine*, 6(1–2), 1–14.
19. Committee on Publication Ethics. (2019). *COPE Guidelines: Core Practices*. Committee on Publication Ethics.