



Research And Academic Writing Using AI: Present Status and Future Perspectives

Akhil Sharma¹ and Prof. U.C. Sharma²

^{1,2}Department of Library and Information Science, Dr. Bhimrao Ambedkar University, Agra

sharmaakhil61996@gmail.com¹, scumesh51@yahoo.co.in²

KEYWORD

Artificial Intelligence (AI); Academic Writing (AW); Research Writing (RW); Large Language Models (LLM); ChatGPT; generative AI; research ethics; publication ethics;

ABSTRACT

Artificial Intelligence has entered everyday research work in a very short span of time. Since the public release of conversational language models in late 2022, students, research scholars, teachers and editors have started using such tools to search literature, plan drafts, improve language and manage references. This paper examines the present status of artificial intelligence in research and academic writing and considers the directions this practice may take in the coming years. The study follows a descriptive and analytical review design. It draws on peer reviewed articles, journal editorials, publisher policy statements and survey reports issued between 2022 and 2026, together with an examination of the tools themselves. The paper traces the growth of artificial intelligence from early theoretical work and rule-based chat programs to the transformer architecture that supports present day writing assistants. It then explains how these systems are used at each stage of the research cycle, from topic selection and literature search to drafting, language editing, data analysis and reference formatting. The benefits recorded include saving of time, better readability, fairer treatment of authors who do not write in English as a first language, and easier access to a wide body of literature. The problems recorded include invented references, weak reasoning, hidden bias, data privacy risk, unequal access and unreliable detection software. A table of twenty-three widely used tools, their area of strength and their web addresses is provided for working researchers. The paper closes with recommendations on disclosure, training, institutional policy and human verification.

1. Introduction

Writing is the final and most demanding stage of research. A study is judged not only by the quality of its data but also by the clarity with which its argument is placed before readers. For a long time this stage depended almost entirely on the skill and patience of the author. Word processors, spelling checkers and reference managers reduced some of the mechanical labour, but the sentences still had to be built one by one. That position began to change when LLM became available to ordinary users through simple chat windows.

The change was sudden. Within weeks of the release of ChatGPT in November 2022, scholars in library and information science were already asking how such a tool would affect teaching, reference services and the writing of scholarly work [1]. Commentators in the wider scientific press argued that conversational AI should be treated as a permanent feature of research practice rather than a passing novelty, and they set out early priorities for handling it responsibly [2]. A survey of more than sixteen hundred researchers carried out for Nature found that scientists were both hopeful and uneasy: many expected these tools to become important to their fields within a decade, while a large share worried about errors, bias and the spread of unreliable work [3].

Corresponding Author: Akhil Sharma, Research Scholar, Department of Library and Information Science, Dr. Bhimrao Ambedkar University, Agra

Email: sharmaakhil61996@gmail.com

Journals responded quickly. Editorial offices refused to accept AI systems as authors, on the plain ground that a program cannot take responsibility for a claim, cannot agree to publication and cannot answer a challenge after publication [4]. The Committee on Publication Ethics issued a position statement in February 2023 stating that AI tools cannot meet the requirements of authorship, and that authors who use such tools must say clearly which tool was used and for what purpose [5]. Assisted writing was therefore permitted, but hidden assisted writing was not. This dual position, permission joined with disclosure, is the central fact of present practice. It places a duty on the individual researcher that many are not yet trained to carry. A scholar who uses a language model to polish a paragraph is doing something that every publisher accepts. A scholar who allows the same model to supply references without checking them is doing something that can end a career. The distance between the two acts is small in effort and large in consequence.

The present paper was written against this background. Its purpose is to give a clear and usable account of where matters stand: what these tools are, how they came about, what they do well, where they fail, which of them are in common use, and what a research community should do next. The treatment is deliberately plain. RW on this subject is often addressed to specialists in computing; the aim here is to speak to the working scholar in any discipline who has to make daily decisions about these tools without any formal preparation for doing so

2. Objectives of the study

The study was carried out with the following objectives:

1. To trace the growth of AI and to identify the developments that made present day writing assistants possible.
2. To describe the present status of AI use at each stage of the research and academic writing process.
3. To identify the benefits reported by researchers, editors and publishers from the use of these tools.
4. To identify the weaknesses, risks and ethical problems attached to their use in AW.
5. To prepare a working list of widely used tools showing the special strength of each, its main use in RW and its web address.
6. To suggest measures for institutions, journals and individual researchers so that these tools may be used with honesty and care.

3. Scope of the study

The study covers the use of AI in the preparation of scholarly work, that is, in searching, reading, planning, drafting, editing, analysing, formatting and checking. It deals with generative systems built on LLM, and also with the search, mapping and reference tools that now carry machine learning inside them. The period examined runs from November 2022, when conversational models reached the general public, to the middle of 2026.

The study is not limited to any one discipline. Examples are drawn from medicine, computing, education and library and information science, because these fields produced the earliest published evidence, but the practices described apply across the humanities, social sciences and sciences alike. Both the author side and the publisher side of scholarly communication are considered, since editorial policy shapes what authors may do.

The sources examined are of four kinds: peer reviewed articles and letters on AI in scientific writing; editorials and formal policy statements issued by journals, publishers and ethics bodies; survey reports on researcher behaviour and attitude; and the published descriptions and interfaces of the tools themselves. The method is descriptive and analytical rather than experimental. No fresh survey of respondents was conducted, and no new dataset was gathered.

4. Limitations of the Study

Certain limits should be stated plainly so that the findings are read in the right spirit.

1. The subject moves faster than the literature about it. Models are revised every few months, and a description of tool behaviour written today may be partly outdated within a year.

2. The study rests on published sources and on the author's examination of the tools. It does not include a survey of researchers, so the attitudes reported here are taken from the surveys of others rather than gathered afresh.
3. Much of the earliest and strongest empirical work comes from medicine and computing. Fields with thinner published evidence are therefore represented mainly through general commentary.
4. The tools listed in Table 1 are those in common use at the time of writing. The list is selective, not exhaustive, and no endorsement of any product is intended. Features, pricing and web addresses change without notice.
5. Reported figures on how many researchers use AI come from self-report surveys with differing samples and question wording, so they should be read as indications of direction rather than exact measures.
6. English language sources dominate the available literature, which may understate practice in other language communities.

5. History of AI

The idea that a machine might handle language did not begin with the present decade. In 1950 Alan Turing asked whether machines could think and proposed a practical test built entirely around written conversation: if a person exchanging typed messages could not tell machine from human, the question of thinking had been answered in a workable way [6]. Language was thus placed at the centre of the subject from the very start. The term AI itself was adopted at a summer meeting held at Dartmouth College in 1956, which is generally taken as the formal beginning of the field.

The first working demonstration of machine conversation came a decade later. In 1966 Joseph Weizenbaum described ELIZA, a program that held a written exchange with a user by matching patterns in the input and turning them back as questions [7]. ELIZA understood nothing, and Weizenbaum said so clearly, yet users spoke to it as though it did. That early observation, that people readily attribute understanding to a system that produces fluent language, remains one of the most useful warnings in the whole subject.

Progress after this was uneven. The nineteen seventies and eighties were dominated by expert systems, which encoded human rules for narrow tasks such as medical diagnosis. These systems worked within their boundaries and failed outside them, and the effort of writing rules by hand proved too heavy. Interest and funding fell away in periods that later writers named the AI winters.

The recovery came from a change of method. Instead of writing rules, researchers allowed systems to learn patterns from data. Statistical methods improved machine translation and speech recognition through the nineteen nineties and early two thousands. From 2012 onwards, deep neural networks trained on large datasets and fast graphics processors produced sharp gains in image and speech tasks, and attention turned back to language.

The decisive step for writing came in 2017, when a team of researchers described the transformer architecture. This design allowed a model to weigh the relation between all words in a passage at once rather than reading strictly in sequence, and it could be trained efficiently on very large amounts of text [8]. Every widely used writing assistant today rests on this design. Successive generative pre-trained transformer models grew in size and ability through 2018, 2019 and 2020, but they remained mainly in the hands of developers.

The public turning point was the release of ChatGPT in November 2022. A plain chat window removed the technical barrier, and the tool reached a very large user base within weeks. The effect on scholarly publishing was immediate and disorderly. Within two months several papers and preprints had listed the chatbot itself as an author or co-author, which forced editors to state a position they had never needed to state before [9]. Between 2023 and 2026 the field settled into competition among several capable model families, with steady improvement in reasoning, in the handling of long documents, and in the ability to work from sources supplied by the user rather than from memory alone.

6. Use of AI in Research and AW

AI now appears at almost every stage of the research cycle. The uses described below are those reported in the literature and observable in the tools in current circulation.

6.1 Topic selection and idea generation

Researchers use language models to open up a subject at the earliest stage: to list possible angles on a broad theme, to sharpen a vague question into a researchable one, to suggest variables that might be examined, and to argue against a proposed design so that weaknesses appear early. The output at this stage is treated as raw material for thinking rather than as a source to be cited.

6.2 Literature search and screening

Search has changed more than any other stage. Tools that map citation networks show how an area developed and which papers connect to which. Evidence search tools return studies matched to a stated claim and extract population, method and outcome into a table for comparison. Citation context tools show whether later work supported or disputed a finding rather than merely counting citations. These functions shorten the screening stage of a review considerably, though the final judgement on relevance stays with the researcher.

6.3 Reading and summarising

Reading assistants attached to a PDF explain difficult passages, unpack equations and tables, and answer questions about the paper on screen. Systems that answer only from documents uploaded by the user have become popular in academic settings, because restricting the model to supplied sources reduces the risk of invented content. Researchers working outside their first discipline report that this use saves a great deal of time.

6.4 Planning and drafting

At the drafting stage the common practice is to ask for an outline, to convert rough notes into connected prose, to expand a compressed argument, or to produce an alternative version of a paragraph that has become tangled. An early discussion in critical care medicine described exactly this pattern of use, that is, help with organising material, with a first draft and with proofreading, and stated that such output must never replace human judgement and must always be reviewed by experts before use [10]. That description of good practice has held up well.

6.5 Language editing and translation

This is the most widely accepted use of all. Grammar checking, style improvement, correction of academic phrasing and translation are permitted by essentially every publisher, and they are the uses most often declared by authors. For scholars in India and in other countries where English is a second or third language, this

function has direct consequences for the chances of acceptance, since manuscripts have long been rejected on language grounds despite sound content.

6.6 Data handling and analysis

Language models write and explain code in R, Python and SPSS syntax, which allows researchers with limited programming background to carry out analyses that were previously out of reach. Conversational analysis tools accept a data file and return descriptive statistics, tests and charts. The risk here is different from the risk in writing: a wrong statistical procedure produces a confident and well formatted result that looks entirely correct, so verification of method matters more than verification of language.

6.7 References, formatting and submission

Reference managers now carry machine learning for metadata extraction and duplicate detection, and language models are used to convert a bibliography from one style to another, to draft cover letters, to check a manuscript against journal requirements and to prepare structured abstracts. This is clerical work, and handing it over releases time for thinking.

6.8 Use by journals and reviewers

Publishers use these tools on their own side, to screen submissions for text overlap and image manipulation, to check that references exist, to identify suitable reviewers and to sort incoming work. The use of language models by peer reviewers is more contested, since uploading an unpublished manuscript to a commercial service breaches the confidentiality of review, and several publishers have restricted the practice for that reason. Early guidance from Nature on ground rules for the use of such tools set the pattern that most journals later followed, that is, no authorship, full disclosure and complete author responsibility for content [12].

An examination of the potential and the threats of these tools in AW, published in 2023, argued that the technology would improve efficiency while raising serious questions of originality, accountability and training, and recommended that institutions prepare guidance rather than wait for problems to appear [11]. Three years later that recommendation is still only partly acted upon in many universities.

7. Advantages

The first advantage is time. Tasks that consumed long hours, screening search results, reformatting references, correcting grammar across a whole thesis, converting notes into readable prose, are completed in minutes. Researchers surveyed on the subject reported faster processing of information and saving of time as the clearest gains from AI in their work [3].

The second is language equity. Scholars who write in English as a second or third language have long carried a burden that native speakers do not. Assisted editing narrows that gap and allows work to be judged on its content. Given that a large part of world research output now comes from countries where English is not the first language, the effect on fairness is considerable.

The third is access to knowledge. Summarising and explanation tools let a researcher enter an unfamiliar area quickly, which supports interdisciplinary work. A social scientist can grasp the outline of a statistical method, and a scientist can follow an argument in policy studies, without a long period of preparatory reading.

The fourth is improved structure and readability. Models are good at organising material into an order that readers can follow, at maintaining a consistent register and at removing repetition. Reviewers spend less effort on presentation and more on substance.

The fifth is support for research design. Used as an opponent rather than a servant, a language model will list objections to a proposed method, point to variables that have not been controlled and name assumptions that have not been stated. This is a rehearsal of peer review before submission.

The sixth is help with routine analysis. Code generation opens methods to researchers who lack programming training, and the explanation of existing code helps them understand what a borrowed script actually does.

The seventh is accessibility. Text to speech, speech to text, summarising and simplification support researchers with visual difficulty, with reading difficulty and with limited typing ability, and thus widen participation in scholarship.

The eighth is a lowered barrier for beginners. A first year research scholar with no senior guidance available can get an immediate explanation of what a literature review should contain or how a methodology chapter is usually arranged. In institutions where supervision is thin, this is a real gain.

8. Limitations of Using AI in AW

The most serious weakness is fabrication. Language models produce text that is statistically likely rather than text that is true, and they will invent citations that carry plausible author names, journal titles and identification numbers. A study published in *Cureus* recorded exactly this behaviour, showing confident output supported by references that did not exist [13]. A separate investigation of references generated by the same class of tool found a large proportion to be false or unverifiable [14]. For AW this is the central danger, because a fabricated reference is a serious matter of integrity even when the author inserted it in good faith.

The second weakness is that fluent output is difficult to judge. When scientific abstracts generated by a language model were placed before blinded human reviewers along with genuine ones, reviewers correctly identified about two thirds of the generated abstracts and wrongly flagged a share of the real ones, and they described the task as surprisingly hard [15]. Text that reads well is therefore no evidence that the content behind it is sound.

The third weakness is the absence of accountability. Scholarship rests on a named person who can be questioned and who answers for what was published. A program cannot hold that position, and this is the reason journals refused authorship rather than any hostility to the technology. Editors set out the point directly when the first such claims appeared [16], and ethics bodies confirmed it in formal statements [5].

The fourth weakness is shallow reasoning. These systems arrange existing material well but do not generate genuine insight, do not know which finding matters and cannot form a research judgement. Work written mainly by such tools tends to be correct in form, safe in content and empty of contribution.

The fifth weakness is bias. Training data reflects what has already been published, which favours certain countries, languages, institutions and viewpoints. A model asked about a subject will reproduce the dominant position and may quietly omit regional scholarship, which is a real concern for research in Indian and other non-western contexts.

The sixth weakness is data privacy. Uploading unpublished data, participant information or a manuscript under review to a commercial service may breach ethics approval, funder conditions, institutional rules and the confidentiality of peer review.

The seventh weakness is unreliable detection. Software claiming to identify machine written text is not dependable. Work published in Patterns found that such detectors misclassified more than half of the writing samples produced by non-native English writers as machine generated, while classifying native samples almost perfectly [17]. Simpler vocabulary and shorter sentences make honest writing look artificial. Any disciplinary action based on a detector score alone is therefore unjust, and this is a matter of practical concern for Indian universities that have begun to use such scores.

The eighth weakness is uneven access. The strongest models sit behind subscriptions priced in foreign currency. Researchers at well-funded institutions obtain better assistance than those at smaller colleges, and a technology described as levelling the field may deepen the divide instead.

The ninth weakness is loss of skill. Writing is a way of thinking. A research scholar who never struggles through a difficult paragraph may not develop the ability to construct an argument at all, and the loss appears later, when the tool cannot supply what was never learned.

The tenth weakness is uncertainty about rights and ownership. The legal position of material generated by these systems, and of the copyrighted work used to train them, remains unsettled in most jurisdictions, and authors carry that uncertainty into their own publications.

9. AI Tools, Their Specialization and Uses

Table 1 lists tools in common academic use at the time of writing, arranged by function. Inclusion is descriptive and does not amount to a recommendation. Features and access conditions change frequently, and researchers should verify the current position, and their own institutional policy, before use.

Table 1: AI tools, specialisation, uses and links

AI Tool	Specialisation	Main use in research and AW	Link
ChatGPT (OpenAI)	General purpose large language model	Idea generation, outlining, drafting, rewriting, explaining concepts, basic code	https://chat.openai.com
Claude (Anthropic)	General purpose large language model with long document handling	Reading and summarising long papers, structured drafting, critique of arguments	https://claude.ai
Connected Papers	Visual citation mapping	Building a visual map of related work around a seed paper	https://www.connectedpapers.com
Consensus	Evidence search across scientific papers	Checking whether a claim is supported by published studies	https://consensus.app
DeepL	Machine translation	Translating source material and drafts between languages	https://www.deepl.com
Elicit	Literature discovery and evidence tables	Finding empirical papers, extracting method and outcome data into tables	https://elicit.com
Gemini (Google)	General purpose multimodal model	Drafting, summarising, working with text, images and tables together	https://gemini.google.com
Grammarly	Grammar and style checking	Correcting grammar, punctuation, tone and clarity in drafts	https://www.grammarly.com

AI Tool	Specialisation	Main use in research and AW	Link
HIX Bypass	Multilingual humanisation and broad compatibility with major LLMs like GPT-4 and Claude.	Helping international researchers smooth out and humanise AI-assisted translations of their research papers across 30+ languages.	hix.ai/bypass
Julius AI	Conversational data analysis	Running statistics and producing charts from uploaded data files	https://julius.ai
Litmaps	Citation network mapping	Tracing how a research area developed over time	https://www.litmaps.com
NotebookLM (Google)	Grounded notebook over uploaded sources	Question answering limited to the researcher's own uploaded documents	https://notebooklm.google.com
Paperpal	Editing built for manuscripts	Language editing plus submission readiness checks for journals	https://paperpal.com
Perplexity AI	Search based answering with citations	Quick fact finding with source links during background reading	https://www.perplexity.ai
QuillBot	Paraphrasing and summarising	Rewording the author's own sentences and shortening notes	https://quillbot.com
Research Rabbit	Literature mapping and alerts	Snowball searching and monitoring new work in a narrow area	https://www.researchrabb.it
Scholarcy	Automatic article summarising	Producing summary cards and highlights from uploaded papers	https://www.scholarcy.com
SciSpace	Paper reading assistant	Explaining difficult passages, tables and equations inside a PDF	https://scispace.com
Scite	Citation context analysis	Seeing whether later papers support or dispute a cited finding	https://scite.ai
Semantic Scholar	Academic search engine with AI ranking	Locating papers, tracking citations, reading auto generated summaries	https://www.semanticscholar.org
Trinka AI	Academic language editing	Subject aware grammar correction and technical style improvement	https://www.trinka.ai
Undetectable AI	Matches specific readability levels (e.g., University, Doctorate) with high readability.	Restructuring AI-drafted literature reviews and essays to mimic a natural academic cadence and bypass institutional detectors like Turnitin.	undetectable.ai
Writefull	Academic phrasing support	Suggesting standard academic phrasing, titles and abstracts	https://www.writefull.com

10. Results and Findings

The review produced the following findings.

Finding 1. Adoption is broad and normal. Assisted writing has moved from experiment to routine within about three years. Surveys of researchers report widespread use, with the largest share concentrated in language editing, proofreading, translation and literature searching rather than in the generation of original argument [3].

Finding 2. Editing dominates over authoring. Across the surveys examined, the tasks most often handed to these tools are clerical and linguistic. Estimates of the amount of published text substantially produced by a machine remain modest, and one systematic examination of published introductions concluded that undisclosed

machine writing affected only a small percentage of the papers tested [18]. The common pattern is a human argument dressed in machine assisted language, not a machine argument.

Finding 3. Policy has converged on two rules. Independent bodies arriving separately at the same conclusion produced a settled position: AI cannot be an author, and its use must be disclosed [4], [5], [12], [16]. Guidance from medical journal editors adds that authors must describe how the tool was used and remain fully responsible for accuracy, integrity and originality [19].

Finding 4. Fabrication is the leading technical risk. Invented references and confident false statements are recorded consistently in the literature and are not confined to any single model or version [13], [14]. This risk falls on the author, since it is the author who signs the paper.

Finding 5. Human judgement of machine text is weak. Trained reviewers identified generated abstracts at a rate well short of certainty and also misjudged genuine work [15]. Fluency cannot be used as a test of authenticity.

Finding 6. Detection software is not fit to be used as evidence. Detectors carry a measured bias against writers whose first language is not English [17]. Their output may inform a conversation but must not decide a case.

Finding 7. The equity effect runs in both directions. Assisted editing reduces the language disadvantage faced by non-native writers, while paid access and biased detection create fresh disadvantages for the same group. The net result depends on institutional decisions rather than on the technology.

Finding 8. Disclosure practice lags behind disclosure policy. Publishers require declarations, but surveys of author behaviour indicate that a substantial share of those who use these tools do not report the fact. The gap arises partly from fear of prejudice against declared use and partly from simple lack of awareness.

Finding 9. Training has not kept pace with adoption. Researchers are using tools whose failure modes they have not been taught to recognise. The literature repeatedly calls for structured preparation, and the calls have largely not been answered at institutional level [2], [11].

Finding 10. The direction of movement is towards grounded systems. Development is shifting from models that answer from memory to systems that answer from sources supplied by the user and show where each statement came from. For academic purposes this is the more important change, because it addresses fabrication at its origin.

11. Conclusion and Recommendations

AI has become part of the ordinary machinery of RW. It is used to search, to read, to plan, to draft, to correct and to format, and it is used by scholars at every level and in every discipline. The evidence assembled here does not support either of the two loud positions in the debate. The tools are not a threat that must be shut out of the academy, and they are not a substitute for scholarship. They are instruments, and like all instruments they extend the reach of a competent user and expose the weakness of an incompetent one.

What has settled is the rule of practice. Assistance is permitted and disclosure is required, and responsibility rests with the human author without any share passing to the tool. What has not settled is the

preparation of researchers to work within that rule. The distance between present policy and present practice is a training problem before it is an ethical one, and it can be closed. The following measures are proposed.

1. Disclose use in every submission. Name the tool, name the version and state exactly what it was used for, in the methods section or the acknowledgements as the journal requires. A short, accurate declaration protects the author far better than silence.

2. Verify every reference personally. Open each citation, confirm that it exists, confirm that the authors, year, volume and pages are correct, and confirm that the source actually supports the statement made. No exception should be allowed to this rule.

3. Keep the argument human. Use these tools for language, structure, search and clerical work. The research question, the interpretation of results and the conclusion should come from the researcher, because these are the parts that carry the contribution.

4. Protect confidential material. Do not upload unpublished data, participant information or manuscripts received for peer review to any commercial service without institutional clearance.

5. Frame institutional policy in writing. Universities should issue a plain statement covering permitted uses, required declarations and the procedure to be followed when misuse is suspected. Silence at institutional level pushes the whole burden onto individual students and supervisors.

6. Do not act on detector scores alone. Because these systems carry a measured bias against non-native English writers [17], a score must never by itself establish misconduct. Any inquiry should rest on drafts, notes, data files and a conversation with the author about the work.

7. Provide training rather than prohibition. Research methodology courses should include practical sessions on prompt construction, on verification of output, on privacy and on disclosure. Prohibition drives the practice underground and removes the chance to teach it well.

8. Address the access gap. Institutions, and particularly libraries, should negotiate group subscriptions to research tools in the same way they negotiate database access, so that the quality of assistance available to a scholar does not depend on personal ability to pay.

9. Prefer grounded tools for scholarly work. Where a system can be restricted to sources supplied by the user and can show the origin of each statement, it should be preferred to one that answers from memory.

10. Review guidance regularly. Any policy framed today will need revision within a year or two. Institutions should name a committee responsible for keeping the guidance current rather than treating it as settled once issued.

Looking ahead, three developments appear likely. Assistance will be built directly into the platforms researchers already use, so that the question will cease to be whether to use these tools and become how to use them well. Verification will become the central academic skill, more valuable than fluent composition. And scholarly value will rest more openly than before on originality of thought, since assistance with expression will be available to everybody. That last shift may prove healthy. If machines take over the labour of arranging words, what remains for the scholar is the question worth asking and the judgement of what the answer means, which is what research was always supposed to be about.

REFERENCES

1. Lund, B. D., & Wang, T. (2023). Chatting about ChatGPT: How may AI and GPT impact academia and libraries? *Library Hi Tech News*, 40(3), 26–29. <https://doi.org/10.1108/LHTN-01-2023-0009>
2. van Dis, E. A. M., Bollen, J., Zuidema, W., van Rooij, R., & Bockting, C. L. (2023). ChatGPT: Five priorities for research. *Nature*, 614(7947), 224–226. <https://doi.org/10.1038/d41586-023-00288-7>
3. Van Noorden, R., & Perkel, J. M. (2023). AI and science: What 1,600 researchers think. *Nature*, 621(7980), 672–675. <https://doi.org/10.1038/d41586-023-02980-0>
4. Thorp, H. H. (2023). ChatGPT is fun, but not an author. *Science*, 379(6630), 313. <https://doi.org/10.1126/science.adg7879>
5. Committee on Publication Ethics. (2023, 13 February). *Authorship and AI tools: COPE position statement*. <https://publicationethics.org/cope-position-statements/ai-author>
6. Turing, A. M. (1950). Computing machinery and intelligence. *Mind*, 59(236), 433–460.
7. Weizenbaum, J. (1966). ELIZA — a computer program for the study of natural language communication between man and machine. *Communications of the ACM*, 9(1), 36–45.
8. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998–6008.
9. Stokel-Walker, C. (2023). ChatGPT listed as author on research papers: many scientists disapprove. *Nature*, 613(7945), 620–621. <https://doi.org/10.1038/d41586-023-00107-z>
10. Salvagno, M., Taccone, F. S., & Gerli, A. G. (2023). Can artificial intelligence help for scientific writing? *Critical Care*, 27(1), 75. <https://doi.org/10.1186/s13054-023-04380-2>
11. Dergaa, I., Chamari, K., Zmijewski, P., & Ben Saad, H. (2023). From human writing to artificial intelligence generated text: examining the prospects and potential threats of ChatGPT in academic writing. *Biology of Sport*, 40(2), 615–622. <https://doi.org/10.5114/biolsport.2023.125623>
12. Nature. (2023). Tools such as ChatGPT threaten transparent science; here are our ground rules for their use. *Nature*, 613(7945), 612. <https://doi.org/10.1038/d41586-023-00191-1>
13. Alkaissi, H., & McFarlane, S. I. (2023). Artificial hallucinations in ChatGPT: implications in scientific writing. *Cureus*, 15(2), e35179. <https://doi.org/10.7759/cureus.35179>
14. Athaluri, S. A., Manthena, S. V., Kesapragada, V. S. R. K. M., Yarlagadda, V., Dave, T., & Duddumpudi, R. T. S. (2023). Exploring the boundaries of reality: investigating the phenomenon of artificial intelligence hallucination in scientific writing through ChatGPT references. *Cureus*, 15(4), e37432. <https://doi.org/10.7759/cureus.37432>
15. Gao, C. A., Howard, F. M., Markov, N. S., Dyer, E. C., Ramesh, S., Luo, Y., & Pearson, A. T. (2023). Comparing scientific abstracts generated by ChatGPT to real abstracts with detectors and blinded human reviewers. *npj Digital Medicine*, 6(1), 75. <https://doi.org/10.1038/s41746-023-00819-6>
16. Flanagan, A., Bibbins-Domingo, K., Berkwits, M., & Christiansen, S. L. (2023). Nonhuman “authors” and implications for the integrity of scientific publication and medical knowledge. *JAMA*, 329(8), 637–639. <https://doi.org/10.1001/jama.2023.1344>
17. Liang, W., Yuksekgonul, M., Mao, Y., Wu, E., & Zou, J. (2023). GPT detectors are biased against non-native English writers. *Patterns*, 4(7), 100779. <https://doi.org/10.1016/j.patter.2023.100779>
18. Desaire, H., Isom, M., & Hua, D. (2024). Almost nobody is using ChatGPT to write academic science papers (yet). *Big Data and Cognitive Computing*, 8(10), 133. <https://doi.org/10.3390/bdcc8100133>
19. International Committee of Medical Journal Editors. (2024). *Recommendations for the conduct, reporting, editing, and publication of scholarly work in medical journals*. <https://www.icmje.org/recommendations/>