

GENERATIVE ARTIFICIAL INTELLIGENCE IN SCIENTIFIC RESEARCH: A QUANTITATIVE SOCIAL SCIENCE PERSPECTIVE

Žarko Radenović¹

¹PhD, Innovation center University of Niš, z_radjenovic89@outlook.com

Abstract

Generative artificial intelligence is increasingly transforming scientific research by supporting idea generation, methodological design, literature search, data interpretation, and academic writing. This study examines the use, perceived benefits, and limitations of generative AI tools in scientific research from a quantitative social science perspective. Based on Innovation Center University of Niš survey data collected from 100 respondents, the findings indicate that a substantial majority of participants have already used generative AI tools in their research activities. Specifically, 82% of respondents reported previous use of generative AI, while 18% had not used such tools. Among users, the most frequently reported reasons for using generative AI included searching the literature with similar methodological frameworks, generating initial research ideas, identifying possible shortcomings and corrections, and obtaining consultative support in methodological decision-making. Respondents who had not used generative AI most often cited insufficient reliability, ethical concerns, limited competence in using AI tools, and the absence of perceived need at the current stage of research. The results further suggest that researchers generally perceive generative AI as a tool that can contribute to the efficiency of scientific work. On a five-point scale, the majority of respondents expressed moderate to high agreement with the statement that generative AI can improve research efficiency. In terms of the methodological framework, the greatest perceived benefits were identified in selecting specific methods and instruments for research implementation, generating research writing ideas, defining research objectives, and formulating hypotheses. These findings indicate that generative AI is not primarily viewed as a replacement for researchers, but as an auxiliary tool that can enhance productivity, support methodological reasoning, and improve the organization of research tasks. The study contributes to the emerging discussion on the responsible integration of generative AI in scientific research, particularly within the social sciences. It highlights both the practical potential of AI-assisted research and the need for methodological literacy, ethical awareness, and critical evaluation of AI-generated outputs. The paper concludes that generative AI can significantly support scientific research when used transparently, critically, and in accordance with established academic standards.

Keywords: generative artificial intelligence; scientific research; research methodology; research ethics; academic integrity; survey research; quantitative social science

INTRODUCTION

Generative artificial intelligence systems (GenAI), especially large language model interfaces, have become embedded in everyday knowledge work. In research contexts, they are used for idea generation, literature exploration, coding support, drafting, translation, summarization, and methodological reflection. The shift is not merely technical. It changes how researchers search, formulate, justify, and communicate knowledge. Foundation models are powerful because they can be adapted across many tasks, but this generality also creates downstream risks when errors, biases, or opaque assumptions travel into specialized scientific workflows (Bommasani et al., 2021). The scholarly debate has therefore moved beyond the question of whether researchers use GenAI. The more relevant question is how they use it, where they perceive value, and which research phases require the strongest safeguards. Prior work has emphasized both opportunities and risks. GenAI can accelerate drafting and provide structured suggestions, but it can also produce plausible falsehoods, invent citations, reproduce bias, and blur responsibility for scholarly claims (Bender et al., 2021; Salvagno et al., 2023; Thorp, 2023). In this sense, GenAI represents a socio-technical condition of research rather than a neutral productivity tool. The present paper contributes a descriptive quantitative perspective by analyzing survey findings from an author-provided presentation on GenAI in scientific research. Its focus is deliberately methodological: rather than treating AI use only as a writing or publication issue, the analysis examines where researchers locate benefits and risks across key stages of research design. This emphasis is

important because methodological decisions - such as defining objectives, formulating hypotheses, selecting instruments, and sampling, shape the validity and credibility of later results. The central research problem is the tension between usefulness and trust. Researchers may value GenAI precisely in complex areas where uncertainty is high, but those are also the areas in which unsupported suggestions may mislead a study design. This paper addresses the following research questions: (1) How widespread is reported GenAI use among surveyed researchers? (2) What motivations and tools dominate current usage? (3) Which methodological benefits and ethical risks are most salient? (4) Where do benefit and risk perceptions overlap, and what does this imply for responsible research practice?

LITERATURE BACKGROUND

The rapid uptake of generative artificial intelligence (GenAI) in scientific work has been accompanied by growing calls for institutional, ethical, and methodological governance. Although GenAI systems are increasingly used to support literature exploration, academic writing, coding, data interpretation, and research design, their integration into scholarly practice remains contested. The central issue is not only whether researchers use these tools, but how they use them, under what conditions, and with what forms of accountability. Recent debates suggest that GenAI should not be treated as a neutral technical instrument. Rather, it is a socio-technical system that reshapes research routines, redistributes epistemic authority, and introduces new responsibilities for researchers, institutions, journals, and peer reviewers. The broader digital transformation literature provides an important foundation for understanding GenAI as part of a wider movement toward data-driven, platform-based, and automated decision support. Studies on digitalization in tourism and e-booking systems show that digital technologies increasingly shape how organizations collect information, communicate with users, organize services, and support decision-making processes. Simović et al. (2020) examine tourism in the digital age from an e-booking perspective, while Vlahović et al. (2024) emphasize the role of e-booking systems in tourism-related digital transformation. Similarly, Marjanović et al. (2019) apply multi-criteria ranking to selected travel-industry indicators among EU members, showing how quantitative and digital tools can support comparative evaluation and strategic decision-making. These studies are relevant because GenAI can be understood as a more advanced form of digital support. It does not only process or present information, but also generates explanations, recommendations, summaries, and methodological suggestions. Research on digital indicators and performance measurement also supports this interpretation. Rađenović et al. (2025) analyze digital mapping of business performance indicators of agricultural holdings in Serbia, demonstrating how digital tools can assist in organizing, visualizing, and interpreting complex performance data. Vranjanac and Rađenović (2022) examine EU countries through hierarchical clustering based on circular economy performance indicators, illustrating the value of quantitative methods for identifying patterns across multidimensional datasets. These studies show that contemporary research increasingly depends on digital tools for classification, ranking, mapping, and interpretation. However, GenAI differs from conventional digital or statistical tools because it produces natural-language outputs that may appear authoritative even when they require additional verification. The question of data integrity is also central to the responsible use of digital technologies. Rađenović (2020) discusses blockchain technology as a data integrity tool in an e-health scenario, emphasizing the importance of trust, traceability, and secure information management in digital systems. Although blockchain and GenAI are different technologies, both raise important questions about reliability, verification, and accountability. In scientific research, this means that GenAI should not be evaluated only through its productivity benefits, but also through its ability to support transparent, traceable, and methodologically responsible research practice. van Dis et al. (2023) argue that conversational AI requires research communities to revise priorities around accountability, transparency, and independent verification. Their argument is particularly relevant because scientific knowledge depends not only on producing plausible explanations, but also on maintaining traceable procedures through which claims can be checked, challenged, and replicated. In this context, GenAI creates a tension between efficiency and epistemic control. On the one hand, it can accelerate preliminary writing, idea generation, literature mapping, and methodological comparison. On the other hand, it can also obscure the origin of ideas, introduce fabricated or weakly supported claims, and make it more difficult to distinguish between human judgement and machine-generated suggestion.

Publication ethics organizations have responded to this tension by emphasizing responsibility and disclosure. The Committee on Publication Ethics (COPE, 2023) states that AI tools cannot be listed as authors because they cannot take responsibility for the integrity, originality, and accuracy of scholarly work. However, COPE also recognizes that AI tools may be used in the research and writing process, provided that such use is transparently disclosed when it materially contributes to the manuscript. This position is important because it separates the use of GenAI from authorship status. The researcher may use AI assistance, but the researcher remains responsible for the final argument, evidence, interpretation, and ethical compliance of the work.

Three conceptual issues are especially relevant for the present study. The first concerns the epistemic status of GenAI outputs. GenAI outputs are probabilistic rather than evidentiary. Large language models generate text by predicting likely sequences of language based on patterns in training data; they do not verify truth, evaluate evidence, or understand methodological validity in the way scientific researchers are expected to do. As a result, they can produce fluent and persuasive responses that are factually inaccurate, methodologically inappropriate, or unsupported by valid sources. Bender et al. (2021) characterize this problem through the metaphor of “stochastic parrots,” warning that language models may generate convincing linguistic forms without grounded understanding. Thorp (2023) and Salvagno et al. (2023) similarly highlight the danger of plausible but incorrect outputs in scientific writing, particularly when researchers rely on AI-generated explanations, summaries, or citations without independent verification. This issue is especially significant in methodological work. In scientific research, methods are not merely procedural choices; they structure what can be known, how data are collected, and how conclusions are justified. If GenAI suggests an inappropriate sampling strategy, measurement instrument, statistical procedure, or interpretation, the problem may not be immediately visible at the level of wording. A text can appear academically polished while concealing weak design logic. Therefore, the methodological use of GenAI requires a higher standard of scrutiny than simple language editing. Researchers must evaluate whether AI-generated suggestions are theoretically coherent, empirically appropriate, ethically acceptable, and compatible with the research question.

The second conceptual issue concerns the breadth of ethical risk. Ethical concerns about GenAI are often discussed through familiar categories such as plagiarism, fabrication, and academic dishonesty. However, the risks are broader and more structurally embedded. Weidinger et al. (2022) identify a wide taxonomy of risks associated with language models, including discrimination, exclusion, privacy violations, misinformation, malicious use, and overreliance. In research contexts, these risks can appear in subtle ways. Bias may enter through suggested literature, categories, concepts, or assumptions. Privacy risks may arise when researchers upload sensitive data, interview transcripts, unpublished manuscripts, or institutional documents into external AI systems. Authorship ambiguity may occur when AI-generated wording or conceptual framing becomes difficult to separate from the researcher’s own contribution. Overreliance may weaken methodological literacy by encouraging researchers to accept generated recommendations without sufficient critical evaluation. These risks are not evenly distributed across all phases of research. Some stages, such as grammar correction or formatting assistance, may involve relatively limited epistemic risk. Other stages, such as hypothesis formulation, instrument construction, sampling design, interpretation of results, and selection of analytical methods, involve substantially higher stakes. In these phases, GenAI does not merely support communication; it may influence the structure of the research itself. This distinction is central to the present paper because the survey findings show that respondents perceive both high benefits and high risks in the same methodological areas. Such overlap suggests that the most useful applications of GenAI may also be the areas that require the strongest safeguards. The third conceptual issue is human accountability. UNESCO’s guidance on generative AI in education and research places human-centered governance at the center of responsible use. Miao and Holmes (2023) emphasize transparency, data protection, institutional policy, capacity building, and the preservation of human agency. This perspective is useful because it avoids both technological optimism and technological rejection. GenAI is not framed as inherently beneficial or inherently harmful. Instead, its value depends on the governance conditions under which it is used, the competencies of users, and the safeguards applied to its outputs. For scientific research, this means that GenAI should be integrated into workflows in ways that strengthen, rather than replace, critical reasoning and methodological responsibility.

A related concern is the concentration of authority in foundation models. Bommasani et al. (2021) argue that foundation models create both opportunities and risks because their general-purpose architecture allows them to be adapted across many domains. In research practice, this generality can be attractive, the same tool may assist with literature review, translation, coding, questionnaire drafting, statistical explanation, and manuscript editing. However, generality can also create misplaced confidence. A tool that performs well in one task may be unreliable in another, particularly when the task requires domain-specific knowledge, ethical sensitivity, or methodological precision. Therefore, the use of GenAI in research should be task-specific, critically supervised, and clearly documented. The literature also points to a distinction between productivity gains and knowledge quality. GenAI may reduce the time needed for routine academic tasks, help researchers overcome writer's block, produce alternative formulations, or identify gaps in an argument. Salvagno et al. (2023) note that AI may assist scientific writing, but they also stress that its outputs require expert review. This distinction is crucial for evaluating GenAI in scientific contexts. Efficiency alone is not a sufficient criterion of responsible use. A faster research process is not necessarily a better one if it weakens verification, transparency, or intellectual ownership. Responsible integration therefore requires researchers to ask not only whether GenAI saves time, but whether it improves or compromises the quality of reasoning, evidence, and interpretation. Taken together, the literature suggests that GenAI should be conceptualized as a conditional methodological assistant. It may help researchers explore alternatives, compare methodological options, draft preliminary wording, identify possible weaknesses, summarize debates, or generate questions for further investigation. However, it should not replace domain expertise, ethics review, source verification, instrument validation, statistical reasoning, or theoretical judgement. The appropriate role of GenAI is therefore consultative rather than authoritative. It can support the research process, but it cannot legitimate the research process. This paper uses the distinction between assistant and substitute as the interpretive lens for the survey findings. The survey results are analyzed not simply as indicators of adoption, but as evidence of how researchers position GenAI within their own methodological practice. Particular attention is given to the overlap between perceived benefits and perceived ethical risks. If researchers value GenAI most in the same phases where they also identify the greatest ethical danger, then responsible-use frameworks must focus less on general acceptance or rejection and more on differentiated governance across specific research tasks. In this sense, the literature background provides the conceptual foundation for interpreting GenAI as a useful but limited tool whose scientific legitimacy depends on transparency, validation, and continued human accountability.

METHODOLOGY

The study uses a descriptive quantitative survey design. The objective is not causal inference but empirical mapping. The paper summarizes reported patterns of GenAI use, perceived benefits, and ethical concerns among academic respondents. The source data are an author-provided survey presentation titled "Generative Artificial Intelligence in Scientific Research: A Quantitative Social Science Perspective." The presentation reports $N = 100$ unless otherwise noted, with the motivations question restricted to respondents who had used AI tools ($N = 84$). The sample includes a broad range of academic and teaching ranks. Research Associate/Assistant respondents formed the largest group (28%), followed by Research Associate/Docent (21%), Senior Research Associate/Associate Professor (17%), Research Adviser/Full Professor (17%), Research Assistant Trainee (16%), and Assistant with PhD (1%). This composition indicates that the survey captured both early-career and senior academic perspectives, which is important because exposure to AI tools and responsibility for methodological decisions may vary by rank. The analysis covers six measurement domains: academic profile, AI adoption, motivations for AI use, perceived methodological benefits, perceived ethical risks, perceived methodological locations of ethical problems, and tools known or used. Percentages are interpreted as descriptive proportions of respondents. For multiple-response items, percentages should be read as endorsement rates rather than mutually exclusive categories. The analysis is based on descriptive statistics: frequencies, percentages, ranking of response categories, and comparison of benefit and risk distributions. Special attention is given to benefit-risk overlap, defined as a methodological segment that respondents rate highly in both the benefit and ethical-risk areas. This approach is appropriate because the available data are aggregated survey results rather than individual-level microdata. Inferential statistics, subgroup tests, and multivariate modeling were therefore not conducted. The paper

analyzes aggregated, de-identified survey results from a presentation. No individual participant data are included. Because the topic concerns responsible AI use, the interpretation follows a conservative principle that GenAI is treated as a tool requiring disclosure, verification, and human accountability, consistent with contemporary publication-ethics and research-ethics guidance (COPE, 2023; Resnik & Hosseini, 2025).

RESULTS AND DISCUSSION

Figure 1 shows that the sample was not concentrated in a single academic rank. The modal group was Research Associate/Assistant (28%), but senior ranks were strongly represented: Senior Research Associate/Associate Professor and Research Adviser/Full Professor each accounted for 17% of the sample. This distribution matters because respondents were likely to have practical experience with research design, supervision, publication, or methodological decision-making.

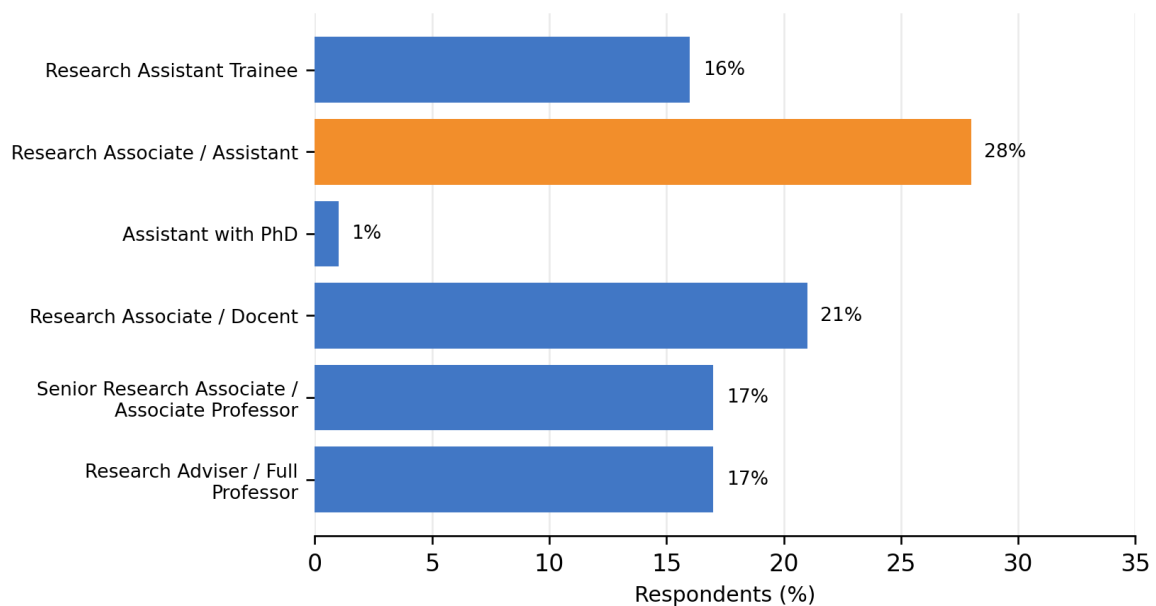


Figure 1. Academic/teaching rank of respondents. Source: author-provided survey presentation, N = 100.
Source: Author's research

AI adoption was high: 82% of respondents reported having used generative AI tools in scientific research, while 18% had not. The finding suggests that GenAI is already a routine element of the research environment for a large majority of respondents. The non-user group remains analytically important, however, because it may reflect ethical caution, skill barriers, institutional restrictions, disciplinary differences, or lack of perceived need.

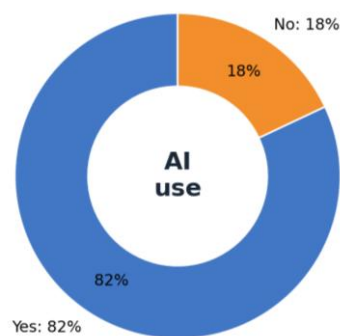


Figure 2. Share of respondents who had used generative AI tools in scientific research. Source: author-provided survey presentation, N = 100.
Source: Author's research

Among respondents who had used AI tools (N = 84), the most frequent motivation was searching literature with similar methodological frameworks (39.3%). This is a noteworthy result because it positions GenAI as a methodological orientation tool rather than only a writing assistant. The next tier included generating initial research ideas and identifying shortcomings or corrections, each endorsed by 14.3% of AI users. Consultative support in methodological decision-making was endorsed by 13.1%. More precise definition of methodological elements and productivity/time pressure were both low at 3.6%, suggesting that respondents did not primarily describe GenAI use as simple time-saving. Instead, the dominant pattern was exploratory and methodological.

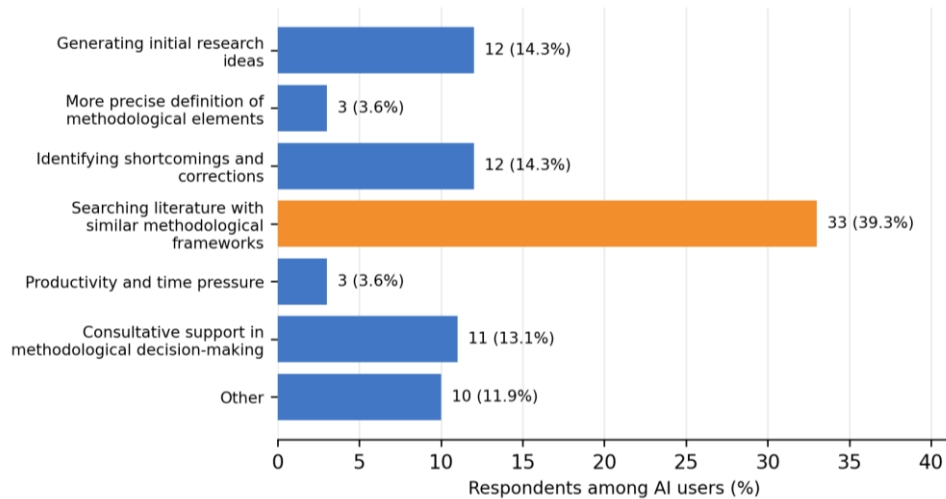


Figure 3. Main motivations for GenAI use among respondents who had used AI tools. Source: author-provided survey presentation, N = 84.
Source: Author's research

Perceived benefits were concentrated in research design and writing support. Selecting specific methods and instruments was the highest-ranked benefit (43%), followed very closely by generating research-writing ideas (42%). Defining research objectives (22%), formulating research hypotheses (21%), and determining the research subject (19%) formed a middle cluster. Sampling decisions had the lowest endorsement (13%), indicating that respondents may view sampling as requiring more context-specific, expert, or discipline-bound judgment. The benefit profile therefore suggests selective trust: researchers appear willing to use GenAI for structuring and exploring methodological options, but less willing to rely on it for sampling decisions.

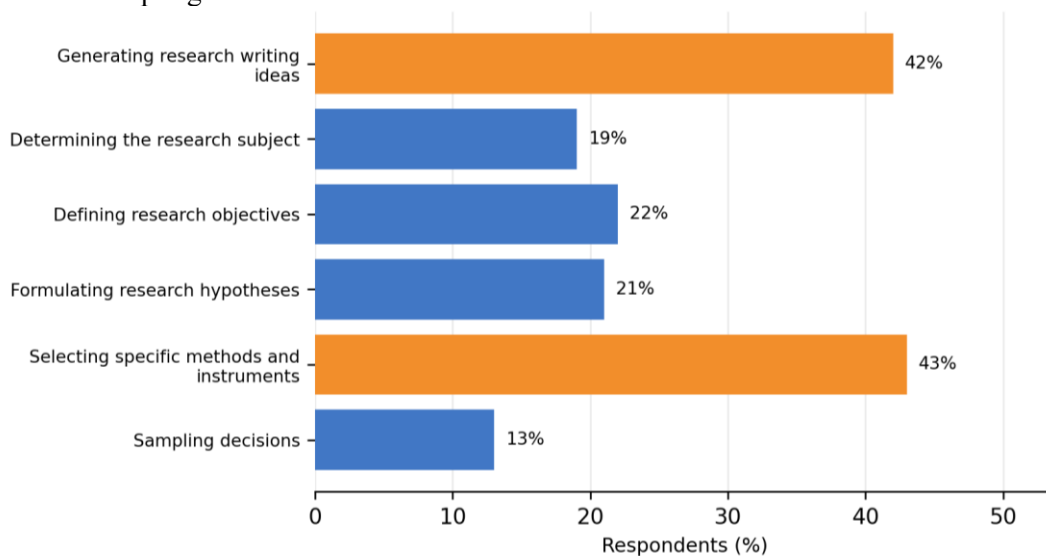


Figure 4. Perceived methodological benefits of GenAI. Source: author-provided survey presentation, N = 100.
Source: Author's research

The leading ethical concern was unreliability of generated content (49%). This result aligns with the broader literature on hallucination, factual error, and unsupported claims in large language model outputs. Authorship responsibility was the second strongest concern (39%), followed by threats to academic integrity (34%). Algorithmic bias (29%) and privacy/data protection (27%) also attracted substantial concern. The distribution implies that respondents understand risk not only as technical inaccuracy but as a governance problem: unreliable content, unclear responsibility, and integrity threats all affect the credibility of the scientific record.

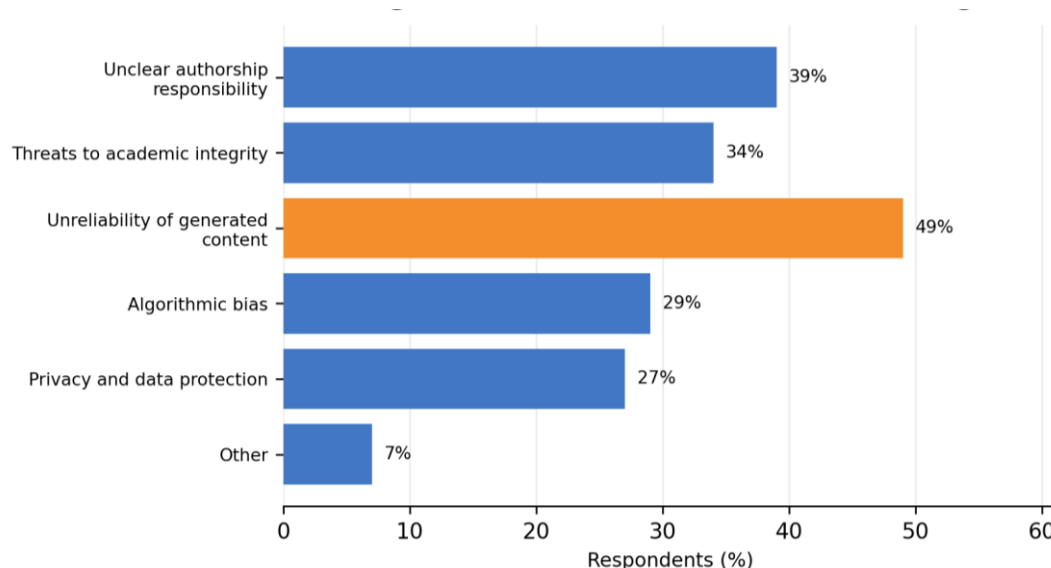


Figure 5. Most significant ethical risks associated with GenAI in research. Source: author-provided survey presentation, N = 100.

Source: Author's research

Respondents located ethical problems most strongly in selecting specific methods and instruments (41%) and formulating research hypotheses (38%). Defining research objectives (31%) and generating writing ideas (30%) also had high endorsement. Determining the research subject and sampling decisions were each endorsed by 20%. The pattern indicates that respondents perceive ethical vulnerability most acutely in parts of research design where AI suggestions may steer later decisions. This finding supports the view that AI governance cannot be restricted to final manuscript writing; it must also address earlier stages of question formation, design selection, and instrument choice.

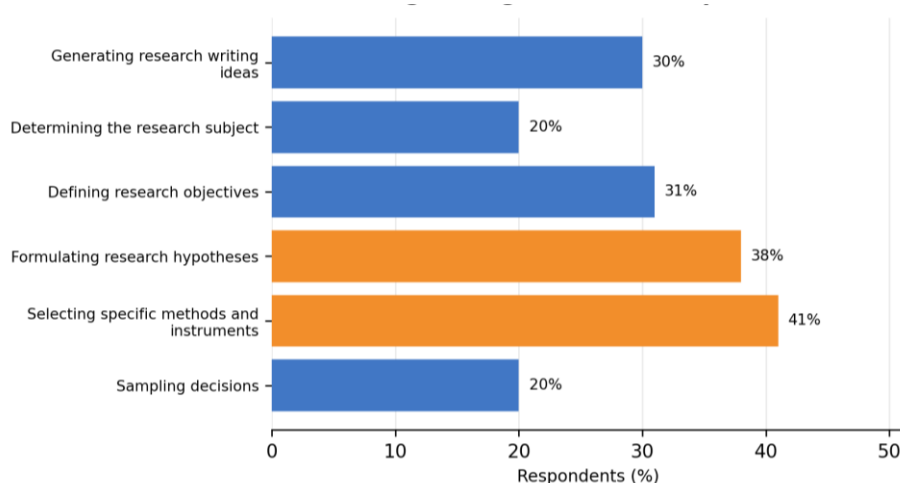


Figure 6. Methodological segments most exposed to ethical problems. Source: author-provided survey presentation, N = 100.

Source: Author's research

The tool landscape was strongly dominated by general-purpose AI assistants. ChatGPT was known or used by 84% of respondents, while Google Gemini was endorsed by 46%. Claude (22%) and Microsoft Copilot (18%) followed, whereas specialized academic tools such as Perplexity AI (10%), Scite (10%), Elicit (5%), and Consensus (4%) were much less common. This imbalance matters because general-purpose systems may be convenient but are not necessarily optimized for transparent scholarly retrieval, citation verification, or domain-specific methodological reasoning. The results therefore support the need for tool literacy: researchers should understand not only what a tool can produce, but also what it cannot verify.

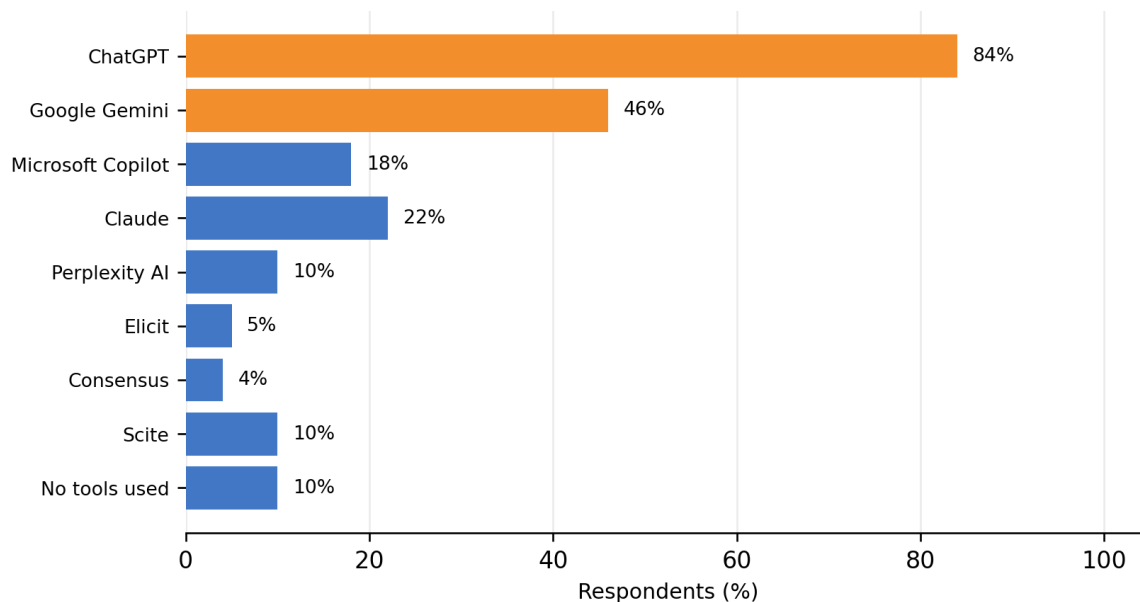


Figure 7. Generative AI tools known and used in research work. Source: author-provided survey presentation, N = 100.
Source: Author's research

The most important quantitative pattern is the overlap between benefit and risk in the same methodological phase. Selecting specific methods and instruments was the highest perceived benefit (43%) and also the highest perceived location of ethical risk (41%). This nearly symmetrical pattern suggests that researchers are not simply enthusiastic or fearful. Instead, they recognize a double condition: GenAI is useful where methodological design is complex, but precisely because those decisions are consequential, AI-generated suggestions require strong safeguards. This is the central empirical insight of the survey.

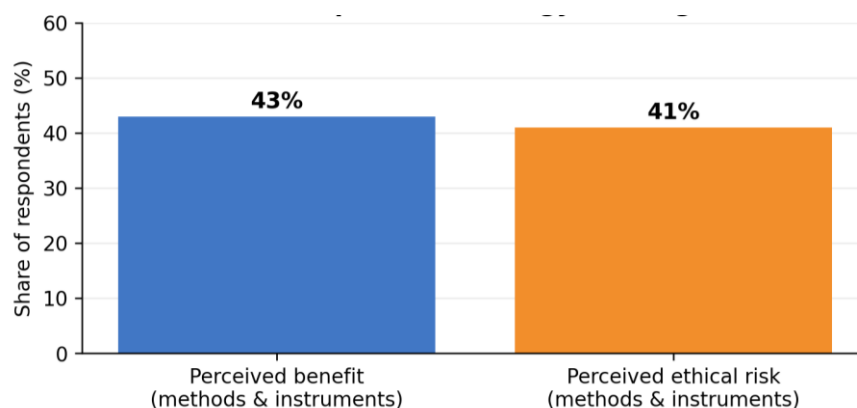


Figure 8. Benefit-risk overlap in the selection of methods and instruments. Source: author-provided survey presentation, N = 100.
Source: Author's research

The findings support three main interpretations. First, GenAI adoption is already widespread among surveyed researchers. The 82% adoption rate means that discussions of research integrity can no longer assume AI use is marginal. Institutions and conferences should therefore move from prohibition-centered assumptions to disclosure, literacy, and verification frameworks. This does not mean unrestricted use; rather, it means that responsible-use rules need to reflect actual practice. Second, respondents primarily associate GenAI with methodological exploration. The strongest motivation among users was searching literature with similar methodological frameworks, and the strongest perceived benefit was selecting methods and instruments. This contrasts with a narrow view of GenAI as only a drafting tool. Researchers are using AI to think with, compare options, and identify potential corrections. This creates value but also increases the consequences of error. When AI is used near the beginning of a research process, a misleading suggestion can shape the entire design. Third, the benefit-risk overlap indicates conditional trust rather than uncritical adoption. Respondents value AI most in methods and instruments, but they also identify this area as the greatest ethical-risk location. This pattern is consistent with the literature warning that GenAI can generate confident but unreliable claims (Thorp, 2023), reproduce bias (Bender et al., 2021), and create responsibility gaps. The practical implication is that GenAI should be positioned as a structured assistant: useful for brainstorming alternatives and prompting reflection, but never sufficient for final methodological decisions without human review. The dominance of general-purpose tools adds another layer to the discussion. ChatGPT and Gemini were far more visible than specialized academic tools. General-purpose tools may lower access barriers, but they also create risks when users treat fluent responses as validated evidence. Responsible methodology therefore requires researchers to verify references in bibliographic databases, compare AI suggestions against methodological literature, document AI use in protocols or methods sections when relevant, and avoid entering sensitive or confidential data into systems without approved safeguards. For conference and journal communities, these findings point to a balanced governance model. Manuscript policies should require disclosure of substantive AI assistance, especially in literature review, data analysis, coding, or methodology formulation. Peer review guidelines should include awareness of AI-generated citation errors and unsupported claims. Research training should include practical exercises in prompt evaluation, source checking, bias identification, and documentation of AI-assisted workflows. Such steps preserve the human accountability that scientific authorship requires while acknowledging that AI assistance is now common.

CONCLUSION

This paper analyzed survey results on generative AI in scientific research using an IMRAD framework. The findings show high adoption, strong methodological use, and clear ethical caution. The central result is not simply that researchers use AI, nor that they fear it. Rather, the key finding is that researchers identify the same methodological phase, selecting methods and instruments, as both the greatest benefit and the greatest ethical risk. This overlap captures the contemporary challenge of GenAI in science. It is most useful where judgment is most consequential. The appropriate response is neither rejection nor uncritical acceptance. GenAI should be integrated as a transparent, verified, and human-accountable research assistant. Its outputs can support reflection, comparison, and early-stage drafting, but claims, citations, methodological choices, and interpretations must remain the responsibility of human researchers. Scientific value emerges when AI support is combined with methodological literacy, source verification, data protection, and ethical accountability. The study has several limitations. First, the analysis is based on aggregated survey results from a presentation rather than the original respondent-level dataset. As a result, subgroup comparisons by rank, discipline, gender, institution type, or prior AI experience could not be conducted. Second, several variables appear to allow multiple responses, which means that percentages should be interpreted as endorsement rates rather than mutually exclusive distributions. Third, the survey captures self-reported perceptions rather than observed AI behavior. Respondents may overestimate or underestimate both their use and their verification practices. Fourth, the sample size is modest, so the results should be treated as exploratory and context-specific rather than nationally or globally representative. Future research should collect respondent-level data, include disciplinary comparisons, distinguish between general-purpose and academic-specific AI tools, and measure actual verification practices. Longitudinal work would also be valuable because GenAI tools, institutional policies, and scholarly norms are changing rapidly.

References

1. Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (pp. 610–623). Association for Computing Machinery. <https://doi.org/10.1145/3442188.3445922>
2. Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., Bernstein, M. S., Bohg, J., Bosselut, A., Brunskill, E., Brynjolfsson, E., Buch, S., Card, D., Castellon, R., Chatterji, N., Chen, A., Creel, K., Davis, J. Q., Demszky, D., ... Liang, P. (2021). *On the opportunities and risks of foundation models*. arXiv. <https://arxiv.org/abs/2108.07258>
3. Marjanović, I., Rađenović, Ž., & Marković, M. (2019). EU members multi-criteria ranking according to selected travel industry indicators. In V. Bevanda & S. Štetić (Eds.), *4th international thematic monograph: Modern management tools and economy of tourism sector in present era* (pp. 639–654). <https://doi.org/10.31410/tmt.2019.639>
4. Miao, F., & Holmes, W. (2023). *Guidance for generative AI in education and research*. UNESCO. <https://unesdoc.unesco.org/ark:/48223/pf0000386693>
5. Radjenovic, Z. (2020). The cost-saving role of blockchain technology as a data integrity tool: E-health scenario. *KnE Social Sciences*, 4(1), 339–352. <https://doi.org/10.18502/kss.v4i1.5995>
6. Rađenović, Ž., Janjić, I., Talić, M., Đokić, M., Rađenović, T., Boskov, T., & Krstić, B. (2025). Digital mapping of business performance indicators of agricultural holdings in Serbia. *Economics of Agriculture*, 72(1), 225–240.
7. Resnik, D. B., & Hosseini, M. (2025). The ethics of using artificial intelligence in scientific research. *AI and Ethics*. <https://doi.org/10.1007/s43681-024-00493-8>
8. Salvagno, M., Taccone, F. S., & Gerli, A. G. (2023). Can artificial intelligence help for scientific writing? *Critical Care*, 27, Article 75. <https://doi.org/10.1186/s13054-023-04380-2>
9. Simović, O., Rađenović, Ž., Perović, D., & Vujačić, V. (2020). Tourism in the Digital Age: E-booking Perspective. *ENTRENOVA-ENTERPRISE RESEARCH INNOVATION*, 6(1), 616–627.
10. Thorp, H. H. (2023). ChatGPT is fun, but not an author. *Science*, 379(6630), 313. <https://doi.org/10.1126/science.adg7879>
11. 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>
12. Vlahović, O., Rađenović, Ž., Perović, D., Vujačić, V., & Davidović, K. (2024). Digital transformation in tourism: The role of e-booking systems. *Croatian Regional Development Journal*, 5(2), 129–145.
13. Vranjanac, Ž., & Rađenović, Ž. (2022). EU countries hierarchical clustering towards circular economy performance indicators. *Facta Universitatis, Series: Working and Living Environmental Protection*.
14. Weidinger, L., Mellor, J., Rauh, M., Griffin, C., Uesato, J., Huang, P.-S., Cheng, M., Glaese, M., Balle, B., Kasirzadeh, A., Kenton, Z., Brown, S., Hawkins, W., Stepleton, T., Biles, C., Birhane, A., Haas, J., Rimell, L., Hendricks, L. A., ... Gabriel, I. (2022). Taxonomy of risks posed by language models. In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency* (pp. 214–229). Association for Computing Machinery. <https://doi.org/10.1145/3531146.3533088>