

ChatGPT and library users: AI risks of hallucinations and misinformation

Bolaji Oladokun

Federal University of Technology, Ikot Abasi, Nigeria

Bolaji.oladokun@yahoo.com

Vivien Oluchi Emmanuel

Rivers State University, Port Harcourt, Nigeria

Emmanuel.vivien@ust.edu.ng

Onome Queen Osagie

University of South Africa, South Africa

Onomeosagie@gmail.com

Adefunke Olanike Alabi

University of Lagos, Lagos State, Nigeria

oladesh@yahoo.com

Research – Full text

Received: 15.03.2025

Accepted: 15.07.2025

Published: 31.08.2025

Copyright (c) 2025,
Bolaji Oladokun, Vivien
Oluchi Emmanuel,
Onome Queen Osagie,
Adefunke Olanike Alabi



This work is licensed
under a Creative
Commons Attribution
4.0 International
License.

Abstract

Purpose: The paper aims to explore the implications of using ChatGPT by library users, focusing on the potential risks of AI hallucinations, the reliability of AI-generated content for research, and strategies to mitigate these risks. The aim is to provide a comprehensive overview of ChatGPT's impact on research and content generation, highlighting the critical role of libraries in guiding users towards responsible AI usage.

Methodology/Design: A systematic review was employed to harvest relevant literature from Scopus, Web of Science, and Google Scholar. Sources between 2020 and 2024 were included in the study. This approach involved identifying, evaluating, and synthesizing research articles, reports, and studies related to ChatGPT, AI hallucinations, and their implications for library services.

Findings: The findings indicate that while ChatGPT offers significant advantages in terms of accessibility and efficiency, its reliance on research and content generation poses considerable risks. These include the dissemination of misinformation, erosion of critical thinking skills, and ethical concerns related to bias. The study highlights the need for improved training data, human oversight, and user education to mitigate these risks effectively.

Implications: The implications of this study are critical for libraries and their users. Libraries must implement comprehensive strategies to ensure the responsible use of AI tools like ChatGPT. This includes educating users about the limitations of AI, encouraging critical evaluation of AI-generated content, and promoting verification through trusted sources.

Originality: This essay provides a unique and thorough examination of the challenges and opportunities presented by ChatGPT in the context of library services. It combines insights from reviews with practical recommendations, and offers a balanced perspective on how to leverage AI technology while addressing its inherent risks.

Keywords

Artificial intelligence, AI, ChatGPT, hallucination, misinformation, library users.

Introduction

The advent of artificial intelligence (AI) has revolutionized the way we access and interact with information. AI-powered chatbots, such as ChatGPT, have emerged as increasingly popular tools for research and content generation, touting their ability to provide rapid and accurate responses to a wide range of queries (Bhattacharya et al., 2024). However, beneath their allure of convenience and efficiency lies a pressing concern: the risk of misinformation. ChatGPT, in particular, has been hailed as a breakthrough in natural language processing, capable of generating human-like responses to complex queries. However, recent studies have revealed a troubling trend: ChatGPT's propensity for hallucinations and reliance on inaccurate sources (Augenstein et al., 2023; Nahar et al., 2024; Oladokun et al., 2025). This phenomenon poses significant risks to library users who rely on ChatGPT for research and content generation, as inaccurate information can lead to flawed conclusions, perpetuate knowledge gaps, and undermine academic integrity.

Studies have revealed that ChatGPT's inaccuracies have the potential to compromise the quality of research and content generation in libraries (Koos & Wachsmann, 2023; Lakavath & Satish, 2023). Despite this, there is a lack of awareness among library users about the potential dangers of relying on ChatGPT for research and content generation. This knowledge gap necessitates a concerted effort to educate library users about the risks of ChatGPT's inaccuracies and to promote critical thinking and information literacy in evaluating AI-powered chatbot responses.

This study, therefore, explores the potential dangers of relying on ChatGPT for research and content generation, to educate library users about the risks of ChatGPT's inaccuracies, and to promote a culture of critical thinking and information literacy in libraries. The subsequent sections of this paper are organized as follows: an overview of ChatGPT, the Phenomenon of AI Hallucinations, perceived Risks Associated with the reliance on ChatGPT for research and content generation, implications of the use of ChatGPT by library users, and mitigating strategies for preventing the incidence of ChatGPT hallucinations. Finally, the paper presents the conclusions drawn from the study outcomes.

Methodology

This study employed a systematic review to investigate the potential risks and implications of using ChatGPT for library users. A systematic review was chosen as it allows for a

comprehensive and structured synthesis of existing research, ensuring a reliable and unbiased understanding of the subject matter. Initially, a total of 54 studies were identified through comprehensive searches of databases such as Scopus, Web of Science, and Google Scholar to ensure a wider capture of relevant studies. Following title screening, abstract review, and full-text analysis based on inclusion and exclusion criteria, 36 articles were deemed eligible and finally included in the study.

To ensure transparency and rigor, the study followed a systematic review protocol aligned with PRISMA (Preferred Reporting Items for Systematic Reviews and Meta-Analyses) guidelines. This protocol detailed the inclusion and exclusion criteria. The inclusion criteria were: (1) articles published between 2020 and 2024 to ensure contemporary relevance, (2) articles written in English to maintain uniformity in analysis, and (3) studies directly addressing AI, misinformation, ChatGPT, or library services. Exclusion criteria included articles not meeting the timeframe, studies not written in English, and those with ambiguous or limited relevance to the research objectives. The final selection was determined through an iterative process of title screening, abstract review, and full-text analysis.

The collected literature was analyzed thematically, focusing on emerging patterns, risks, and strategies related to ChatGPT usage in library services. Thematic analysis enabled the identification and categorization of recurring themes across studies, ensuring a robust synthesis of findings. Ethical considerations were observed by ensuring proper citation of all reviewed literature, upholding academic integrity, and avoiding data manipulation. Since this study relied on publicly available research, there were no direct interactions with participants, thereby eliminating privacy concerns.

Overview of ChatGPT

ChatGPT, developed by OpenAI, represents a significant advancement in artificial intelligence, specifically in the realm of natural language processing (NLP) (Abdullah et al., 2022). This paper provides an overview of ChatGPT, with a particular focus on its development and underlying technology, specifically the GPT-4 architecture, as well as its diverse use cases and growing popularity in various fields. ChatGPT is built on the GPT-4 architecture, a state-of-the-art language model that stands for Generative Pre-trained Transformer 4 (Nazir & Wang, 2023). This architecture represents the fourth iteration in a series of progressively more advanced models designed by OpenAI. Yenduri et al. (2024) state that the development of GPT-4 involved extensive pre-training on a vast and diverse corpus of text data from the internet, allowing it to learn patterns, structures, and nuances of human language. The model employs a transformer architecture, which consists of numerous layers of attention mechanisms (Niu et al., 2021). These mechanisms enable the model to weigh the importance of different words in a sentence, thereby understanding context more effectively than previous models.

One of the key innovations in GPT-4 is its ability to handle significantly larger datasets and more complex computations, thanks to advancements in both hardware and algorithmic efficiency (Chang, 2023). This allows GPT-4 to generate more coherent and contextually relevant

responses. The model's training process involves two main phases: pre-training, where the model learns to predict the next word in a sentence, and fine-tuning, where it is further refined using a narrower dataset with human feedback to enhance its accuracy and appropriateness (Naseem et al., 2021). The underlying technology of GPT-4 is designed to optimize both performance and scalability. This has been achieved through improved tokenization methods, better handling of long-range dependencies in text, and more sophisticated techniques for reducing biases and ensuring ethical AI use (Ray, 2023). These advancements have collectively contributed to GPT-4's ability to generate human-like text that is often indistinguishable from that written by humans.

The Phenomenon of AI Hallucinations

Artificial Intelligence (AI) has made remarkable strides in recent years, particularly in the domain of natural language processing (NLP). However, despite its advancements, AI systems like ChatGPT can exhibit a phenomenon known as hallucinations. According to Maleki et al. (2024), AI hallucinations refer to instances where AI generates information or responses that are plausible-sounding but factually incorrect or nonsensical. Hill (2023) states that AI hallucinations occur when an AI model, such as ChatGPT, produces text that appears coherent and reasonable but lacks a basis in reality or factual accuracy. These hallucinations can range from minor inaccuracies to entirely fabricated content. Unlike simple errors, hallucinations are characterized by their deceptive plausibility, making them particularly challenging to identify and correct. One prominent example of an AI hallucination involves generating false historical facts (Ji et al., 2023). For instance, if asked about a historical figure, an AI might confidently provide detailed but entirely fabricated information about that person's life and achievements. Another example is the creation of fictitious scientific data or medical advice, where the AI generates plausible-sounding but incorrect explanations or recommendations (Haupt & Marks, 2023). These hallucinations can mislead users, especially those relying on AI for accurate and reliable information.

The causes of AI hallucinations are multifaceted, stemming from the inherent limitations and design of AI systems. One primary cause is the nature of the training data. AI models like GPT-4 are trained on vast datasets harvested from the internet, which include both accurate and inaccurate information (Hadi et al., 20203). During training, the AI learns patterns and structures in the data, but it does not develop a true understanding of the content (Holzinger, 2018). This can lead to the generation of false or misleading information that reflects the imperfections of the training data. Another significant cause is the probabilistic nature of AI language models. These models generate responses based on the likelihood of word sequences, aiming to produce coherent and contextually appropriate text (Pandey et al., 2024). However, this probabilistic approach can result in the creation of plausible but incorrect information, especially when the model encounters gaps in its training data or ambiguous queries.

Moreover, the complexity of language itself contributes to AI hallucinations (Sovrano et al., 2023). Human language is rich with nuance, context, and implicit knowledge that AI models struggle to fully grasp (Sayers et al., 2021). Ambiguities and subtleties in language can lead AI to generate incorrect responses that seem reasonable but are factually inaccurate. This issue is exacerbated when AI attempts to answer questions outside its training data's scope, leading to confident yet erroneous statements. Furthermore, biases inherent in the training data also play a role in AI hallucinations (Rawte et al., 2023). If the data used to train the AI contains biases or misinformation, these can be propagated and amplified in the AI's responses. This can result in the generation of skewed or incorrect information that reflects the biases present in the original data.

Perceived Risks Associated with the Reliance on ChatGPT for Research and Content Generation

The advent of ChatGPT and similar AI-driven language models has revolutionized the fields of research and content generation. These tools offer remarkable capabilities, such as the ability to generate coherent and contextually relevant text, assist with data analysis, and streamline the creation of written materials (Bonner et al., 2023). However, the increasing reliance on ChatGPT for these purposes is accompanied by significant perceived risks. These risks span from the potential for misinformation to ethical concerns and the undermining of critical thinking skills. One of the most pressing risks associated with the reliance on ChatGPT for research and content generation is the potential for misinformation (Whalen & Mouza, 2023). Despite its advanced capabilities, ChatGPT can generate responses that are factually incorrect or misleading. This risk is particularly concerning in academic and professional research, where accuracy is paramount. Misinformation can propagate through scholarly articles, reports, and other content, leading to the dissemination of false information. This not only undermines the credibility of the content but also has broader implications for public knowledge and policy-making.

ChatGPT's responses are generated based on patterns and data within its training corpus, but the model lacks the ability to verify the authenticity and credibility of the sources it references (Ray, 2023; Sohail et al., 2023). Consequently, there is a significant risk that AI-generated content may rely on unreliable or biased information. Researchers and content creators who depend on ChatGPT without cross-verifying its outputs may inadvertently introduce errors into their work. This can be particularly problematic in fields where source integrity is critical, such as medicine, law, and science. Adding to the foregoing, the use of ChatGPT in research and content generation raises several ethical issues (Ray, 2023; Wu et al., 2024). One concern is the potential for AI-generated content to perpetuate biases present in the training data (Labajova, 2023). If the data used to train ChatGPT contains biases related to race, gender, or other social factors, the model may reproduce these biases in its outputs. This can reinforce stereotypes and contribute to inequality in various domains. Additionally, the use of AI for content creation raises questions about authorship and intellectual property.

Reliance on ChatGPT for generating research and content undermines critical thinking and analytical skills (Koos & Wachsmann, 2023). When researchers and writers depend heavily on AI to produce text, they may become less engaged in the critical processes of evaluation, synthesis, and original thought. This can lead to a reduction in the quality of scholarly work and a diminished capacity for independent analysis. Over time, the erosion of these skills could have profound implications for education and professional practice, where critical thinking is essential. In a different light, Al-Raimi et al. (2024) note that when users become accustomed to the convenience and efficiency of AI-generated content, they may become less inclined to engage in traditional research methods and writing practices. This dependency can reduce the incentive to develop and maintain essential skills, such as thorough literature review, data analysis, and original writing. In the long term, this shift could impact the depth and rigor of research and content across various fields.

Implications of the Use of ChatGPT by Library Users

Library users, including students, researchers, and scholars, rely on accurate and reliable information to support their academic and professional pursuits. However, the integration of ChatGPT, an advanced AI language model, into library services has sparked considerable debate. While the technology offers potential benefits, its use by library users also presents significant negative implications. Given that ChatGPT can generate responses that are not always accurate or reliable, its integration into library settings may adversely affect the quality of research and scholarly writing.

One of the primary concerns with the use of ChatGPT in libraries is the risk of misinformation (Whalen & Mouza, 2023). ChatGPT generates responses based on patterns learned from a vast corpus of text, which includes both accurate and inaccurate information (Ray, 2023). Consequently, users may receive information that sounds plausible but is factually incorrect. In the context of academic research, where precision and accuracy are crucial, this misinformation can lead to flawed research outcomes, propagation of false information, and a general decline in the quality of scholarly work. Rosyanafi et al. (2023) observe that the reliance on ChatGPT for research purposes undermines the integrity of the research process. Traditional research methods involve critical evaluation of sources, thorough literature reviews, and meticulous cross-referencing. However, ChatGPT's convenience leads to a shortcut mentality, where users accept AI-generated content at face value without rigorous verification. This shift can erode essential research skills, such as critical thinking and analytical reasoning, ultimately compromising the depth and reliability of academic research.

As noted by Maita et al. (2024), a significant negative implication of using ChatGPT is the potential erosion of critical thinking skills among library users. When users become dependent on AI-generated responses, they may be less inclined to engage in independent analysis and critical evaluation of information. Over time, this reliance can diminish their ability to scrutinize sources, question assumptions, and develop well-reasoned arguments. In academic and professional environments, the decline in critical thinking skills can result in superficial

understanding and poorly constructed research outputs. In addition, ChatGPT's responses are influenced by the data it was trained on, which can contain inherent biases (Hassani & Silva, 2023). When library users rely on ChatGPT, they risk encountering and propagating biased information. This can reinforce existing prejudices and contribute to unequal representation in scholarly work. Additionally, ethical concerns arise regarding the transparency and accountability of AI-generated content (Abdikhakimov, 2023). Users may not fully understand how ChatGPT generates responses, leading to unquestioning acceptance of potentially flawed or biased information.

The use of ChatGPT in libraries diminishes the traditional role of librarians as trusted guides in the research process (Oladokun et al., 2024). However, librarians play a crucial role in helping users navigate information sources, assess the credibility of materials, and develop effective research strategies (Roy & Chanada, 2024). It may be stated that as users increasingly turn to ChatGPT for immediate answers, the demand for librarian expertise may decline. This shift can reduce the personalized support and mentorship that librarians provide, leading to a less robust research environment. While ChatGPT offers quick answers, it lacks the nuanced understanding and contextual awareness that human experts provide (Gill & Kaur, 2023). This further implies that over-reliance on AI can create a dependency that limits users' ability to conduct independent research and problem-solving. This dependency is particularly concerning in situations where AI-generated responses are incorrect or misleading, leaving users without the skills to identify and correct these errors.

Mitigation Strategies for Preventing Incidence of ChatGPT Hallucination

The phenomenon of AI hallucinations poses significant challenges in various applications, including research and content generation. As pointed out, the mitigation strategies include improved training data, fine-tuning, fact-checking, and user feedback. First, one of the most effective strategies to reduce AI hallucinations is to improve the quality of the training data (Ji et al., 2024). This involves curating datasets from reputable sources and continuously updating them to reflect the latest, most accurate information. This is to say that the likelihood of generating hallucinations can be significantly reduced by training AI on higher-quality data. Additionally, implementing rigorous data validation processes to filter out incorrect or misleading information before it is used in training can further enhance data quality.

While adding to the strategies for mitigating hallucinations, Gunjai et al. (2024) propose the need to fine-tune AI models with domain-specific data. Fine-tuning involves taking a pre-trained model and further training it on a narrower, more specialized dataset. This process helps the model become more adept at handling specific types of queries and reduces the chance of producing inaccurate responses. For example, a model fine-tuned with medical literature will be more reliable when answering health-related questions. This specialization ensures that the AI provides more precise and contextually relevant information within a particular domain.

Also, incorporating robust fact-checking mechanisms into the AI response generation process is essential to mitigate hallucinations (McIntosh et al., 2023). These mechanisms involve cross-referencing AI-generated content with trusted databases and sources in real-time. When the AI generates a response, it can be automatically checked against a repository of verified information. If discrepancies are detected, the system flags the response for further review or correction. This approach ensures that users receive more accurate and trustworthy information, reducing the spread of misinformation.

In addition to the fact-checking strategy, Tonmoy et al. (2024) state that actively incorporating user feedback is another critical strategy for mitigating AI hallucinations. Users who encounter incorrect or misleading responses can provide feedback, which can then be analyzed and used to improve the model. In a similar light, other scholars stated that limiting the scope of AI applications and setting clear boundaries for its use can help prevent hallucinations (Ray, 2023; Magesh et al., 2024). This suggests that developers can reduce the likelihood of the model venturing into areas where it lacks sufficient training data or contextual understanding.

Conclusion

The reliance on ChatGPT for research and content generation by library users presents significant risks, primarily due to the potential for AI hallucinations. These hallucinations can result in the dissemination of misinformation, undermining the integrity of scholarly work and eroding critical thinking skills. Furthermore, the use of ChatGPT without proper oversight may perpetuate biases and ethical concerns, reducing the reliability of the information provided. As libraries increasingly integrate AI into their services, it is crucial to acknowledge these risks and address them proactively. Libraries and librarians must play a pivotal role in guiding users to exercise caution when utilizing AI tools like ChatGPT. This can be achieved by encouraging the critical evaluation of AI-generated content and emphasizing the importance of verifying information through trusted and reliable sources. While ChatGPT offers valuable capabilities, its use in libraries requires a balanced approach that prioritizes accuracy and ethical considerations. Libraries must implement mitigating strategies, including improved training data, human oversight, and user education, to ensure the responsible use of AI. By doing so, libraries can enhance the benefits of AI while safeguarding the quality and reliability of information, ultimately supporting the scholarly and research endeavors of their users.

Implications of the Study

The findings of this study have critical implications for both policy and practice in the field of librarianship. At the policy level, it indicates the need for libraries to establish guidelines for the ethical use of AI tools, such as ChatGPT. Policymakers must prioritize initiatives that integrate critical information literacy into library services, ensuring users can discern the limitations of AI-generated content. Additionally, there is a pressing need for frameworks that mandate fact-checking and verification of AI responses before dissemination to users. This aligns with promoting transparency, accuracy, and reliability in library services.

From a practical standpoint, the study highlights the necessity for libraries to take a proactive role in educating users about the risks of misinformation associated with AI tools. Librarians should be equipped with the skills to guide users in critically evaluating AI-generated responses and verifying them against credible sources. Furthermore, libraries should invest in staff training and workshops focused on AI literacy to ensure librarians remain effective mediators between users and technology. Implementing these strategies can strengthen the role of libraries in fostering informed and critical engagement with AI, ultimately enhancing the quality and reliability of research outcomes for library users.

Suggestions for Future Research

Future research should explore the long-term implications of integrating AI tools like ChatGPT into library services, particularly focusing on their impact on information literacy, critical thinking, and research quality among users. It is essential to investigate how reliance on AI-generated content affects academic performance and professional competence over time, as well as the potential risks of dependency. Additionally, studies could examine the effectiveness of various mitigation strategies, such as librarian-led AI literacy programs and real-time fact-checking mechanisms, in reducing misinformation and enhancing user confidence in navigating AI tools. Comparative studies analyzing user interactions with different AI systems across diverse academic disciplines and cultural contexts would provide deeper insights into the adaptability and limitations of such tools. Furthermore, interdisciplinary research is needed to address the ethical challenges posed by biases and inaccuracies in AI-generated content, developing robust frameworks to ensure fairness, accountability, and transparency. Finally, the evolving capabilities of AI language models warrant ongoing assessment, particularly as they intersect with emerging technologies like augmented reality and blockchain, to understand their implications for the future of libraries and information management.

Limitations of the Study

One limitation of this study is its reliance on published literature, which inherently restricts the scope of the information and perspectives available in existing research. This limits the ability to capture emerging trends and user experiences that may not yet be documented in academic sources. Additionally, the study exclusively focuses on articles published in English, which may exclude relevant research conducted in other languages or different geographical contexts, potentially narrowing the diversity of perspectives on the use of AI in libraries. Furthermore, while the systematic review methodology ensures a comprehensive analysis, it does not allow for direct empirical data collection or the inclusion of real-time user feedback, which could provide a more nuanced understanding of how library users interact with AI tools like ChatGPT. The study also did not explore the specific technical limitations or underlying algorithms of ChatGPT itself, which could provide valuable insights into its inherent inaccuracies and limitations.

Acknowledgment: The authors appreciate all scholars and authors whose work has been consulted during this paper.

References

- Abdullah, M., Madain, A., & Jararweh, Y. (2022, November). ChatGPT: Fundamentals, applications and social impacts. In *2022 Ninth International Conference on Social Networks Analysis, Management and Security (SNAMS)* (pp. 1-8). Ieee.
- Abdikhakimov, I. (2023, June). Unraveling the Copyright Conundrum: Exploring AI-Generated Content and its Implications for Intellectual Property Rights. In *International Conference on Legal Sciences* (Vol. 1, No. 5, pp. 18-32).
- Al-Raimi, M., Mudhsh, B. A., Al-Yafaei, Y., & Al-Maashani, S. (2024, March). Utilizing artificial intelligence tools for improving writing skills: Exploring Omani EFL learners' perspectives. In *Forum for Linguistic Studies* (Vol. 6, No. 2, pp. 1177-1177).
- Augenstein, I., Baldwin, T., Cha, M., Chakraborty, T., Ciampaglia, G. L., Corney, D., ... & Zagni, G. (2023). Factuality challenges in the era of large language models. *arXiv preprint arXiv:2310.05189*.
- Bhattacharya, P., Prasad, V. K., Verma, A., Gupta, D., Sapsomboon, A., Viriyasitavat, W., & Dhiman, G. (2024). Demystifying ChatGPT: An In-depth Survey of OpenAI's Robust Large Language
- Bonner, E., Lege, R., & Frazier, E. (2023). Large Language Model-Based Artificial Intelligence in the Language Classroom: Practical Ideas for Teaching. *Teaching English with Technology*, 23(1), 23-41. Models. *Archives of Computational Methods in Engineering*, 1-44.
- Chang, E. Y. (2023, December). Examining gpt-4: Capabilities, implications and future directions. In *The 10th International Conference on Computational Science and Computational Intelligence*.
- Gill, S. S., & Kaur, R. (2023). ChatGPT: Vision and challenges. *Internet of Things and Cyber-Physical Systems*, 3, 262-271.
- Gunjal, A., Yin, J., & Bas, E. (2024, March). Detecting and preventing hallucinations in large vision language models. In *Proceedings of the AAAI Conference on Artificial Intelligence* (Vol. 38, No. 16, pp. 18135-18143).
- Haupt, C. E., & Marks, M. (2023). AI-generated medical advice—GPT and beyond. *Jama*, 329(16), 1349-1350.
- Hadi, M. U., Qureshi, R., Shah, A., Irfan, M., Zafar, A., Shaikh, M. B., ... & Mirjalili, S. (2023). A survey on large language models: Applications, challenges, limitations, and practical usage. *Authorea Preprints*.

- Hassani, H., & Silva, E. S. (2023). The role of ChatGPT in data science: how ai-assisted conversational interfaces are revolutionizing the field. *Big data and cognitive computing*, 7(2), 62.
- Hill, M. (2023). *Hallucinating Machines: Exploring the ethical implications of generative language models* (Doctoral dissertation, Open Access Te Herenga Waka-Victoria University of Wellington).
- Holzinger, A. (2018, August). From machine learning to explainable AI. In *2018 world symposium on digital intelligence for systems and machines (DISA)* (pp. 55-66). IEEE.
- Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., ... & Fung, P. (2023). Survey of hallucination in natural language generation. *ACM Computing Surveys*, 55(12), 1-38.
- Koos, S., & Wachsmann, S. (2023). Navigating the Impact of ChatGPT/GPT4 on Legal Academic Examinations: Challenges, Opportunities and Recommendations. *Media Iuris*, 6(2).
- Magesh, V., Surani, F., Dahl, M., Suzgun, M., Manning, C. D., & Ho, D. E. (2024). Hallucination-Free? Assessing the Reliability of Leading AI Legal Research Tools. *arXiv preprint arXiv:2405.20362*.
- Maleki, N., Padmanabhan, B., & Dutta, K. (2024). AI Hallucinations: A Misnomer Worth Clarifying. *arXiv preprint arXiv:2401.06796*.
- Maita, I., Saide, S., Putri, A. M., & Muwardi, D. (2024). ProsCons of Artificial Intelligence-ChatGPT Adoption in Education Settings: A Literature Review and Future Research Agendas. *IEEE Engineering Management Review*.
- McIntosh, T. R., Liu, T., Susnjak, T., Watters, P., Ng, A., & Halgamuge, M. N. (2023). A culturally sensitive test to evaluate nuanced gpt hallucination. *IEEE Transactions on Artificial Intelligence*.
- Nahar, M., Seo, H., Lee, E. J., Xiong, A., & Lee, D. (2024). Fakes of Varying Shades: How Warning Affects Human Perception and Engagement Regarding LLM Hallucinations. *arXiv preprint arXiv:2404.03745*.
- Naseem, U., Razzak, I., Khan, S. K., & Prasad, M. (2021). A comprehensive survey on word representation models: From classical to state-of-the-art word representation language models. *Transactions on Asian and Low-Resource Language Information Processing*, 20(5), 1-35.
- Nazir, A., & Wang, Z. (2023). A comprehensive survey of ChatGPT: Advancements, applications, prospects, and challenges. *Meta-radiology*, 100022.
- Niu, Z., Zhong, G., & Yu, H. (2021). A review on the attention mechanism of deep learning. *Neurocomputing*, 452, 48-62.
- Labajová, L. (2023). The state of AI: Exploring the perceptions, credibility, and trustworthiness of the users towards AI-Generated Content.

- Lakavath, H., & Satish, C. (2023). Use of ChatGPT in Library Services: A Study. *Journal of Applied Research of Information Technology and Computing*, 14(1-3), 25-30.
- Oladokun, B.D, Yusuf, M., & Dogara, K. (2024). Students' Attitudes and Experiences with ChatGPT as a Reference Service Tool in a Nigerian University: A Comprehensive Analysis of User Perceptions. *Gamification and Augmented Reality*, 2, 36-36.
- Oladokun, B. D., Enakrire, R. T., Emmanuel, A. K., Ajani, Y. A., & Adetayo, A. J. (2025). Hallucination in Scientific Writing: Exploring Evidence from ChatGPT Versions 3.5 and 4o in Responses to Selected Questions in Librarianship. *Journal of Web Librarianship*, 19(1), 62-92.
- Pandey, R., Waghela, H., Rakshit, S., Rangari, A., Singh, A., Kumar, R., ... & Sen, J. (2024). Generative AI-Based Text Generation Methods Using Pre-Trained GPT-2 Model. *arXiv preprint arXiv:2404.01786*.
- Rawte, V., Sheth, A., & Das, A. (2023). A survey of hallucination in large foundation models. *arXiv preprint arXiv:2309.05922*.
- Ray, P. P. (2023). ChatGPT: A comprehensive review on background, applications, key challenges, bias, ethics, limitations and future scope. *Internet of Things and Cyber-Physical Systems*.
- Roy, R., & Chanda, A. (2024). Hello ChatGPT: Transforming Academia and Libraries through AI. *Tuijin Jishu/Journal of Propulsion Technology*, 45(2).
- Rosyanafi, R. J., Lestari, G. D., Susilo, H., Nusantara, W., & Nuraini, F. (2023). The dark side of innovation: Understanding research misconduct with chat gpt in nonformal education studies at universitas negeri surabaya. *Jurnal Review Pendidikan Dasar: Jurnal Kajian Pendidikan dan Hasil Penelitian*, 9(3), 220-228.
- Sayers, D., Sousa-Silva, R., Höhn, S., Ahmed, L., Allkivi-Metsoja, K., Anastasiou, D., ... & Yayilgan, S. Y. (2021). The Dawn of the Human-Machine Era: A forecast of new and emerging language technologies.
- Sohail, S. S., Farhat, F., Himeur, Y., Nadeem, M., Madsen, D. Ø., Singh, Y., ... & Mansoor, W. (2023). The future of gpt: A taxonomy of existing chatgpt research, current challenges, and possible future directions. *Current Challenges, and Possible Future Directions (April 8, 2023)*.
- Sovrano, F., Ashley, K., & Bacchelli, A. (2023, July). Toward eliminating hallucinations: Gpt-based explanatory ai for intelligent textbooks and documentation. In *CEUR Workshop Proceedings* (No. 3444, pp. 54-65). CEUR-WS.
- Tonmoy, S. M., Zaman, S. M., Jain, V., Rani, A., Rawte, V., Chadha, A., & Das, A. (2024). A comprehensive survey of hallucination mitigation techniques in large language models. *arXiv preprint arXiv:2401.01313*.
- Whalen, J., & Mouza, C. (2023). ChatGPT: Challenges, opportunities, and implications for teacher education. *Contemporary Issues in Technology and Teacher Education*, 23(1), 1-23.

Wu, X., Duan, R., & Ni, J. (2024). Unveiling security, privacy, and ethical concerns of ChatGPT. *Journal of Information and Intelligence*, 2(2), 102-115.

Yenduri, G., Ramalingam, M., Selvi, G. C., Supriya, Y., Srivastava, G., Maddikunta, P. K. R., ... & Gadekallu, T. R. (2024). Gpt (generative pre-trained transformer)—a comprehensive review on enabling technologies, potential applications, emerging challenges, and future directions. *IEEE Access*.