

BEYOND NEUTRALITY: COMPARATIVE STUDY OF CHAT GPT-PRODUCED RECOMMENDATION LETTERS

Ajita Devi Sharma

Sociology Department

ajitagautam@gmail.com

ABSTRACT

This study explores LLM-generated documents, their social impacts and biases, and their outcomes in research. Language, as a deeply intricate and nuanced system, continues to challenge our efforts in both understanding and generation. Language modeling has progressed from foundational statistical techniques to the sophistication of Large Language Models (LLMs) powered by deep learning. Trained on vast quantities of data, self-generated content, and self- and semi-supervised learning, demonstrate remarkable capabilities for producing contextually relevant, human-like text and executing a broad spectrum of language tasks. Despite their promise, these systems are not without limitations, particularly concerning embedded biases stemming from training data and model architecture. Emerging scholarship delves into their symbolic reasoning, potential, and interdisciplinary applications, though many avenues remain insufficiently explored. As examples of generative artificial intelligence, ChatGPT increasingly contributes to education by fostering accessible learning, facilitating domain-specific assessments, and nurturing creativity and critical thinking among learners. It attempts to understand the LLM model by unraveling the facts in day-to-day practices. This article applies the qualitative methods used in comparing two AI-generated recommendation letters, indicating male and female (Ajit/ Ajita). Textual review as a major tool and technique of knowing LLMs and their biases, significance, social impacts, and the ways of interpreting biases with outcomes. It is based on sociotechnical system theory and algorism bias theory. This article claims that LLMs are the latest AI-generated product utilized in academic discourses, enriching the capacity of the user. The biases could be minimized by careful attention, capacity, tactful utilization, and knowledge of the user.

Keywords: *Artificial intelligence, Biases, Deep learning, Language modeling, Text generation*

Conceptualization

Language is a complex system that is challenging to understand and generate effectively as it changes social meanings involving grammar, context, intention, emotion, and cultures. Over the past, language modeling has evolved from statistical models with remarkable abilities (Zhao et al., 2023). Large Language

Models are more dominant in the core of mainstream approaches, introducing bias caused by the choice of texts and system generation (Navigli et al., 2023). LLMs are AI systems trained on extensive text data to generate human-like language. Using self- and semi-supervised learning, they predict contextually appropriate text and perform a wide range of language tasks. Recent research also explores mechanistic interpretations, linking their behavior to symbolic inference algorithms. Head et al. (2023) claimed that AI-generated language is a form of generative artificial intelligence designed for text-based content, utilizing deep learning techniques on large datasets, and is a recently developed, primarily targeting text generation and analysis. Likewise, deep learning models with many parameters serve as a significant and versatile tool, including multiple and multilayered applications, across diverse disciplines, even though these opportunities have not been systematically investigated (Van Wyk, 2024). ChatGPT supports researchers an easy access to learning, subject domain tasks, assessment, and evaluation practices extremely efficiently in subject domain tasks, improving creativity and critical thinking in accomplishing goals (Sullivan et al., 2023).

LLMs like Chat GPT are automating professional communication, but surrounding recommendation letters remain human-centered, insisting on personal insights, institutional character, and social evaluation. It assumes that authenticity and individual judgment are key to credibility regarding AZI-generated texts and their disruption. Most of the literature celebrates LLMs as neutral tools and mechanisms that operate with embedded data but denies socio-technical and cultural dynamics. It facilitates the user and is recognized for its fluency and efficiency, with limited critical interrogation and highlighting the need to question regarding the values and codes reinforced as the dominant norms that guide automation.

This study not only examines the established claim that LLMs are biased (Resnik, 2025) but also targets showing how bias appears in a specific genre of recommendation letters. Especially, the recommendation letter is significant in academic and professional sectors where linguistic differences underscore in denying or grabbing of opportunities. Therefore, comparison suggest paying careful attention to data training, fairness, representation, accountability, and human evaluations.

Research Questions

- What are LLMs and associated biases?
- What are the significance and social impacts of LLMs?
- How are the biases interpreted and addressed?

Objectives and Methodology

This article applies qualitative methods and accommodates descriptive and exploratory designs to analyze and interpret data and draw conclusions. The article further examines the underlying features, uses, social dynamics, and broader implications of the phenomenon under study, framing and transformation through literature review and analysis.

Notably, it has used review-based data as the source of information. Two recommendation letters were generated by ChatGPT-5.2 version used in creating recommendation letters. Data were coded in a thematic order, identifying key concepts addressing search questions and based on theoretical grounds.

Conceptual Frameworks

This article is guided by the conceptual framework connecting key concepts LLMs, significance, social impacts, debates, and scholarly discourse. A range of relevant literature is consulted to support both conceptual and empirical analysis, relying on the study employs a thick description approach to explore the nuanced dimensions of the subject matter. It helps to explain how AI-generated text produces gender, authority, social biases, and needs human evaluations.

Theoretical frameworks

This article adopts a sociotechnical system theory and algorithmic bias theory to explore the LLM models, existing biases, and societal impacts of large language models (LLMs) like ChatGPT. Socio-technical systems point to designing and evolving technical systems, integrating people, process, and technology in producing technical systems (Trist, 1981). Likewise, algorithm bias theory (Friedman & Nissenbaum, 1996), O'Neil (2016), and Noble (2018) have shaped how algorithms can perpetuate existing societal biases due to faulty data, design, and systemic differences and inequalities, viewing LLMs as products of both technical design and social values, shaped by and shaping society. While they offer benefits in education, employment, and communication, likewise, risk-reinforcing biases, spreading misinformation, and promoting cultural homogenization. This blended approach offers a nuanced lens for examining a balanced lens to assess how it influences human interaction and evolving social norms.

General Overview

While establishing new territory, the process significantly impacts both the public and the research community, raising important ethical and practical challenges. It must align with human values and ethical principles. Although researchers may approach it with high expectations, they must remain mindful of its limitations (De Angelis et al., 2023). Chat GPT is a significant and interesting language processing

application insightfully applied in different disciplines, such as history to medicine, mathematics, physics, and medicine fields, providing directions for future advancement. It aligns with bias, misinformation, and data privacy risk and suggests forward-looking perspectives on its development, integration, and responsible use in the future (Liu et al., 2023).

It has significant potential for expanding decision-making and is a novel framework designed to uncover and reduce biases in large language models supporting complex decision-making, and the biases can be reduced through commercial and open-source models (Echterhoff et al., 2024). Similarly, Badyal et al. (2023) talked about biases in LLM models that can be applied in the creative field, such as a new form of media.

It is a recently emerged but very useful tool assisting in writing professional documents, providing convenience and fairness, and reducing gender bias that needs deeper scrutiny of the attached biases in writing (Wan et al., 2023). Likewise, it is automatic code generation models playing a key role in enhancing the productivity of software development procedures, but no one denies the role of bias regarding gender, race, and age (Huang et al., 2023).

LLM has biases that can be reduced through prompt engineering, fine-tuning, and the development of more reliable bias detection media content and systems (Lin et al., 2024). There is a rising trend of using LLM judges and humans; both are biased, vulnerable, and underscore the urgency for robust, unbiased, and bias-resistant evaluation systems (Chen et al., 2024). Similarly, these models underscore the challenges and highlight the immediate requirement for identifying methods of correcting embedded social biases (Taubenfeld et al., 2024). There is the pervasive issue of bias, indicating data selection, choice of text, systematic, gender, race, from sexual orientation to ethnicity, religion, and culture that influence production (Navigli et al., 2023).

Social Impact

Gokul (2023), Solaiman et al. (2023), Haque & Li (2024), Tahir et al. (2024), Ramu et al (2024), Yang et al (2024), Xie & Avila(2025) have discussed the impact of LLM in society. They have elaborated LLM as an AI-generated system that creates many impacts, creating societal advancement, enhanced quality of education, labor market accessibility, short and sweet cultural knowledge, and prompt decision-making support. It reflects recurring concerns about bias, exclusion, and governance for equitable outcomes, simultaneously matching with technical development and sociological awareness, inclusive practices, and structural safeguards.

Haque & Li (2024) claimed that ChatGPT supports in many respects, such as capacity enhancement, effective communication, and access to education, and provides ethical guidelines as a positive contribution to society without reinforcing existing inequalities or biases, whereas it contributes to disruption, interference, misinformation, overreliance, and ethical ambiguity to the dependent user. Likewise, Tajir et al (2024) talked about context-aware sociological grounding for shifting group norms to deepen social dynamics for performing well in basic social patterns.

Yang et al. (2024) argued that LLMs are incomplete social predictors that indicate generalized biases and cultural limitations, advocating for human oversight in sensitive applications. Ramu et al (2024) focused on the need for sociotechnical integration to ensure equity and should be aware of biases. Gokul (2023) recognizes productivity gains and overreliance and unexamined deployment as ethical drawbacks of LLMs, whereas Xie & Avila (2025) alarm that LLMs raise cultural homogenization, misinformation, and bias replication, which impact society deeply, and could be minimized through inclusive model development and regulation.

Gokul (2023), Haque & Li (2024), Tahir et al. (2024), Ramu et al. (2024), and Xie & Avila(2025) have underscored the dual nature of LLMs in framing contemporary scholars and society. AI-generated content contributes to educational advancement, enhanced communication skills, and social efficacy without denying critical concerns and biases. This study unites the need for contextual understanding, sociotechnical integration, and conscious engagement of users for the development of LLMs as a tool of empowerment rather than reinforcement of differences. Niraula (2024) argued that AI-generated texts are used as a window for understanding and negotiating content, and it is applied as a teaching tool at universities, reflecting how educational policies and norms have been reshaped. Likewise, Solaiman et al. (2023) claim that AI- generated text traces the social impacts, but that needs to be judged through risk distribution and benefits sharing pattern, which is possible by reshaping power structure and cultural ideologies in a structured way.

Therefore, AI-generated ChatGPT creates a social impact in multiple disciplines such as universities, education policy, and teaching and learning patterns, social norms, cultural practices, environment that needs to analyze risk distribution and benefits sharing patterns.

Results

AI-Generated Recommendation Letters produced by ChatGPT

Case A

Subject: Letter of Recommendation for Ajita Sharma

To Whom It May Concern,

I am writing to wholeheartedly recommend Ajita Sharma for any position or opportunity she may be pursuing. I have had the privilege of working closely with Ajita for the past four years at the University of Tennessee, and throughout this time, she has consistently demonstrated exceptional dedication, professionalism, and expertise in her field.

Ajita has shown herself to be an invaluable member of our team, consistently exceeding expectations in her role. Her dedication to her work is evident in the meticulous attention to detail she brings to every task, ensuring that projects are completed to the highest standard and promptly. Furthermore, Ajita's strong analytical skills and problem-solving abilities have been instrumental in overcoming challenges and finding innovative solutions.

One of Ajita's greatest strengths is her ability to collaborate effectively with colleagues and stakeholders from diverse backgrounds. She is a skilled communicator, able to articulate complex ideas clearly and concisely, and her approachable nature fosters a positive and inclusive working environment. Ajita's collaborative spirit and willingness to share her knowledge and expertise make her a true asset to any team.

In addition to her professional abilities, Ajita possesses exceptional interpersonal skills and a genuine passion for supporting others. She consistently goes above and beyond to mentor and support her colleagues, offering guidance and encouragement whenever needed. Ajita's empathy and compassion make her not only a great colleague but also a trusted friend and confidante.

In conclusion, I cannot speak highly enough of Ajita Sharma's character, work ethic, and professional abilities. She is an outstanding individual who would be an invaluable addition to any organization or team. I am confident that she will continue to excel in all her endeavors, and I wholeheartedly endorse her for any opportunity she may pursue.

Sincerely,

[Name] [Position] [University/Organization Name]

Case B

Subject: Letter of Recommendation for Ajit Sharma

To Whom It May Concern,

I am pleased to write this letter in strong support of Ajit Sharma, who has been a valuable member of our team at the University of Tennessee for the past four years. I have had the privilege of working closely with Ajit during this time and can confidently attest to his exceptional skills, dedication, and contributions to our institution.

Throughout his tenure here, Ajit has consistently demonstrated a high level of professionalism and expertise in his role. He approaches his work with meticulous attention to detail and a strong commitment to excellence, ensuring that projects are completed to the highest standards and deadlines are consistently met. Ajit's proficiency in his field has been evident in the quality of his work and his ability to adapt to new challenges and responsibilities.

Ajit is a skilled professional as well as a collaborative and supportive colleague. He has a strong ability to work effectively within teams, fostering an environment of open communication and cooperation. His willingness to share his knowledge and expertise has been instrumental in the success of many projects, and he is always ready to offer support and guidance to his peers.

In addition to his professional accomplishments, Ajit is a person of integrity and character. He approaches every task with enthusiasm and a positive attitude, and his dedication to his work is truly commendable. Ajit's commitment to excellence and his strong work ethic makes him an asset to any organization.

In conclusion, it is found that the utmost confidence is placed in Ajit Sharma's abilities and potential. He has proven himself to be a reliable, skilled, and dedicated professional, and I am certain that he will continue to excel in any future endeavors. I wholeheartedly recommend Ajit for any opportunities that come his way, and I am confident that he will make a significant and positive contribution wherever he goes.

Sincerely,

[Name] [Position] [University/Organization Name]

Critical Observation on LLM-Generated Recommendation Letters

The mentioned LLMs produced content enriched with specific words, terms, and symbols indicating gender, content, authority, and sociological impacts. While using LLM, the educator, researcher, and writer are highly encouraged to take advantage, but without being slaves to technology. It is an emerging pathway for the knowledge seeker to use properly or misuse. Mainly, the following impacts are

noted while comparing two relational recommendation letters produced by ChatGPT.

Gender Biases

There are remarkable differences between the two recommendations, indicating tone, language style, interpersonal differences, technical focus, alliance, and closing statement. In terms of gender differences, it has a reflection on tone, such as warm, expressive, and personal, which are more communal and nurturing descriptions, which are stereotypes following traditional expectations for women (Ajita), whereas males (Ajit) have used professional, measured, and reserved tones with agentic traits indicating male-coded evaluations.

In terms of language style, Ajita is enriched with emotionally rich concepts such as a trusted friend, compassion, and empathy, whereas in the case of Ajit, it has focused on competence and professionalism. Interpersonal emphasis is also used differently, such as for Ajita, mentoring, empathy, inclusiveness, and support, but for Ajit, it has been used with a moderate, supportive, but less emotional emphasis. In response to a technical focus, it is moderate for Ajita, mentioning analytical and problem-solving skills, whereas for Ajit, there is a highly technical focus on proficiency and adaptability.

Mentioning recommendation strength, Ajita has emphasized on strength of soft skills but is more emotionally invested, which can be read as less objective, personality-based rather than skill, and for Ajit, it has focused more on team success, which is professionally credible and worthy in evaluative contexts. There are high biases in closing statements, such as for Ajita, who wholeheartedly endorses outstanding, whereas for Ajit, utmost confidence, reliable, skilled, and dedicated are used.

Both recommendations are supportive, but differ in tone, emphasis, and descriptions. These differences can reinforce systemic biases; women are framed with collaborators and supporters rather than leadership capacity and technical expertise, which can potentially affect hiring priority and advancement decisions, especially in a patriarchal society and technical field.

Content Biases

Ajita's recommendation letter has prioritized opening contents that are more emotional, personal, than professional, whereas achievements are acknowledged, focusing on analytical skills with feminine-coded traits such as collaboration, mentorship, and empathy. The work is described as dedicated and detail-oriented, but less mentioned about independent decision-making, leadership traits, and strategic impact. There is a specific space for team contributions that is highlighted

as fostering a positive environment rather than leadership initiatives, such as suggesting a supportive role as a helper rather than a leader, and innovator roles.

Whereas, in Ajit's recommendation letter, the opening contents are more neutral, less emotionally prioritized, with strong support, focusing heavily on proficiency, adaptability, and responsibility, aligning with independence and competence. Descriptive words like dedicated professional, skilled, and reliable emphasize self-sufficiency, reducing interpersonal warmth. Work ethics are focused on strong adaptability, dependability, and leadership, which can reflect a high level of acknowledgement with positivity. In terms of team contribution, it has insisted on instrumental contributions to knowledge sharing, suggesting ownership and a high level of expertise.

Authority Biases

Authority bias indicates the tendency to attribute greater accuracy, significance, and influence from authority figures in different spheres, such as tone regarding objectivity and subjectivity, formality, conciseness, structure, emphasis on results, and avoidance of emotional language. It has preferring content that appears more formal, neutral, or firm, often aligning with male-coded writing styles, with Ajit's letter as objective phrasing, stronger in evaluations.

Ajita's letter is enriched with positive and rich in praising, warmth, less authoritative with emotional framing, whereas Ajit's letter is written with the expectation from evaluators emphasizing competence and facts, with limited emotional framing, assertive, formal, higher credibility, and objective tones.

Sociological Biases

The comparison reflects sociological gender bias in valuation and representation. Case A (Ajita) indicates traditional feminine-associated qualities, in terms of interpersonal quality, which are treated as relational traits, a stronger emotional tone and overall framing value is applied with multidimensional but socially connected care, likability and supportive identity with personality and skills whereas, Case B (Ajit) is evaluated through masculine-coded professionalism where emphasis is highlighted on competency, potential team player, and underscored to the group dynamic and focus centered on performances. The value is centered on productivity and outcomes rather than emotional or relational. Sociologically, it provides insights into how gendered expectations shape social recognition, where masculine and feminine norms directly determine relational value and performance value.

Result and Discussion

Literature underscores the ethical, social, and practical challenges associated with the application of large language models while establishing diverse territory. De Angelis (2023) claimed the need for alignment with human values and ethical principles, suggesting that researchers maintain critical awareness. Liu et al. (2023) stressed concerns around misinformation, data privacy, and forward-looking strategies, and Badyal et al. (2023) emphasized creative domains while being aware of embedded model biases.

Similarly, Chen et al. (2024), Taubenfeld et al. (2024), Navigli et al. (2023), and Huang et al. (2023) discussed LLM biases and their sociological implications in practical settings. They have suggested enhancing fairness and convenience, deeper scrutiny in software development to reduce biases regarding gender, race, and age, suggesting the development of bias-resistant evaluation systems, prompt engineering, bias identification, correction mechanisms, and regular auditing for more equitable AI systems.

The sociological trends observed in the two recommendation letters mirror real-world biases in society and demonstrate how LLMs reinforce and perpetuate these patterns. These cases reveal how sociotechnical systems interact with social inequalities and serve as a small-scale version of broader structural discrimination. In the absence of critical oversight and correction, LLMs risk reproducing and amplifying existing social hierarchies, which normalize, legitimize, and accelerate inequality.

Gender norms are reinforced in AI-generated recommendations for similar profiles, which often rely on gender-stereotyped language, exposing unfair evaluations. Likewise, there are biases in automated scoring and screening due to language deemed more credible or objective, and devaluation of emotional intelligence, such as emotional and communal traits, self-sustaining, and reflecting dominant norms, which perpetuate misevaluation.

Based on sociotechnical systems and algorithmic bias theory, it has provided theoretical ground for understanding Chat-GPT- generated recommendation letters, indicating a mutual connection between technology and society. Sociotechnical system theory not only automated process but also algorithms and training of data. This supports examining the issue of bias, accountability, ethics, trustworthiness, and social impacts with the outcome. Similarly, Algorithmic theory reproduces unequal, stereotypical representations. It shows that AI-generated output is not neutral, free from social influence and bias. It helps in knowing how bias enters AI-generated text through training data, prompt framing, and

institutional use. It favors certain groups, perspectives, cultures, and knowledge, influencing responses.

This article examines recently developed Large Language Models (LLMs) such as ChatGPT, with particular attention to generating evaluative texts. The study explores the implicit biases embedded in AI-generated content and their broader implications across education, research, and professional credibility, indicating differences in two recommendation letters. It focuses on how biases are reflected and vary along with gender, authority, language, and sociocultural positioning. This analysis suggests that ChatGPT frequently produces dominant socio-cultural norms depending on the training input. Using qualitative textual analysis, the study concludes with ethical reflections informed by sociotechnical systems theory and algorithmic bias frameworks.

Conclusion

The study concludes that two cases with varying illustrations go beyond simple sociolinguistic artifacts with a warning. These two cases suggest that Chat-GPT-generated text reflects bias, but findings should be treated carefully due to the limited sample size. It shows how selected cases reveal possible patterns of bias. Therefore, it should be understood as a sociotechnical artifact shaped by training data, design, language, and social context. Viewed through a qualitative lens grounded in sociotechnical systems and algorithmic bias theories, this article underscores the complex duality of LLMs: they can function as a tool for empowerment. Their use in the research and education sector can be transformed through critical judgment, users' awareness, cultural knowledge, and digital literacy. The responsible use of LLMs can support in enhancing understanding and knowledge production.

Alongside this, future applications should carefully focus on social taboos, inclusive datasets, language models, and collaboration between users' engagement and technical development. Future requirement of AI-generated text is not significant due to advanced technology, but the way it is perceived, interpreted, and applied. Its responsible use demands addressing risk, inclusive datasets, thoughtful system design, regular audits, and human evaluation.

References

- Badyal, N., Jacoby, D., & Coady, Y. (2023, October). Intentional biases in LLM responses. In *2023 IEEE 14th Annual Ubiquitous Computing, Electronics & Mobile Communication Conference (UEMCON)* (pp. 0502-0506). IEEE.
- Chen, G. H., Chen, S., Liu, Z., Jiang, F., & Wang, B. (2024). Humans or LLMs as the judge? a study on judgment biases. *arXiv preprint arXiv:2402.10669*.
- De Angelis, L., Baglivo, F., Arzilli, G., Privitera, G. P., Ferragina, P., Tozzi, A. E., & Rizzo, C. (2023). ChatGPT and the rise of large language models: the new AI-driven Infodemic threat in public health. *Frontiers in public health*, 11, 1166120.
- Echterhoff, J. M., Liu, Y., Alessa, A., McAuley, J., & He, Z. (2024, November). Cognitive bias in decision-making with LLMs. In *Findings of the association for computational linguistics: EMNLP 2024* (pp. 12640-12653).
- Gokul, A. (2023). LLMs and AI: Understanding its reach and impact. *Preprints*.
- Haque, M. A., & Li, S. (2024). Exploring ChatGPT and its impact on society. *AI and Ethics*, 1-13.
- Head, C. B., Jasper, P., McConnachie, M., Raftree, L., & Higdon, G. (2023). Large language model applications for evaluation: Opportunities and ethical implications. *New directions for evaluation*, 2023(178-179), 33-46.
- Huang, D., Zhang, J. M., Bu, Q., Xie, X., Chen, J., & Cui, H. (2024). Bias testing and mitigation in LLM-based code generation. *ACM Transactions on Software Engineering and Methodology*.
- Lin, L., Wang, L., Guo, J., & Wong, K.-F. (2024). *Investigating bias in LLM-based bias detection: Disparities between LLMs and human perception*. arXiv.
- Liu, Y., Han, T., Ma, S., Zhang, J., Yang, Y., Tian, J., ... & Ge, B. (2023). Summary of ChatGPT-related research and perspective towards the future of large language models. *Meta-radiology*, 1(2), 100017.
- Navigli, R., Conia, S., & Ross, B. (2023). Biases in large language models: origins, inventory, and discussion. *ACM Journal of Data and Information Quality*, 15(2), 1-21.
- Niraula, S. (2024). The impact of ChatGPT on academia: A comprehensive analysis of AI policies across UT system academic institutions. *Advances in Mobile Learning Educational Research*, 4(1), 973-982.

- Ramu, K., Nijhawan, G., Lalitha, G., Asha, V., & Fakher, A. J. (2024). Multidimensional Impacts of Generative AI and an In-Depth Analysis of LLMs with Their Expanding Horizons in Technology and Society. In *2024, OPJU International Technology Conference (OTCON) on Smart Computing for Innovation and Advancement in Industry 4.0* (pp. 1-6). IEEE.
- Resnik, P. (2025). Large language models are biased because they are large language models. *Computational Linguistics*, 51(3), 885-906.
- Solaiman, I., Talat, Z., Agnew, W., Ahmad, L., Baker, D., Blodgett, S. L., ... & Subramonian, A. (2023). Evaluating the social impact of generative AI systems in systems and society. *arXiv preprint arXiv:2306.05949*.
- Sullivan, M., Kelly, A., & McLaughlan, P. (2023). ChatGPT in higher education: Considerations for academic integrity and student learning.
- Tahir, A., Cheng, L., Sandoval, M., Silva, Y. N., Hall, D. L., & Liu, H. (2024, September). Evaluating LLMs' Capabilities Towards Understanding Social Dynamics. In *International Conference on Advances in Social Networks Analysis and Mining* (pp. 230-244). Cham: Springer Nature Switzerland.
- Taubenfeld, A., Dover, Y., Reichart, R., & Goldstein, A. (2024, November). Systematic biases in LLM simulations of debates. In *Proceedings of the 2024 conference on empirical methods in natural language processing* (pp. 251-267).
- Van Wyk, M. M. (2024). Is ChatGPT an opportunity or a threat? Preventive strategies employed by academics related to a GenAI-based LLM at a faculty of education. *Journal of applied learning and teaching*, 7(1), 35-45.
- Xie, Y., & Avila, S. (2025). The social impact of generative LLM-based AI. *Chinese Journal of Sociology*, 11(1), 31-57.
- Yang, K., Li, H., Wen, H., Peng, T. Q., Tang, J., & Liu, H. (2024). Are Large Language Models (LLMs) Good Social Predictors? *arXiv preprint arXiv:2402.12620*.
- Zhao, W. X., Zhou, K., Li, J., Tang, T., Wang, X., Hou, Y., ... & Wen, J. R. (2023). A survey of large language models. *arXiv preprint arXiv:2303.18223*, 1(2), 1-124.