



Gratification and Utility as Determinants of AI Adoption among Professionals: A Qualitative Cross-Industry Analysis

Yuansheng Liu¹ and Xinjue Jiang^{2*}

¹ Department of Communication, Faculty of Modern Languages and Communication, Universiti Putra Malaysia, Seri Kembangan, 43400, Malaysia

² Department of Language and Literacy Education, Faculty of Education, Universiti Malaya, 50603 Kuala Lumpur, Wilayah Persekutuan Kuala Lumpur, Malaysia

*xinjuejiang123@gmail.com

Abstract. Rapid deployment of Artificial Intelligence (AI) technology in formal and informal sectors of the economy in the last five years or so has made it important to understand how professionals in various sectors perceive this technology and what determines their perception. This paper purported to examine the factors causing the deification of AI among professionals in China who are considerably familiar with AI and adept at using it, from a Uses and Gratification Theory and Utility Theory perspective. Using focus group discussions to collect primary qualitative data and thematic analysis, it is found that knowledge of AI limitations does not deter its adoption. Users have taken a mitigation or early-adopter mindset to manually curtail drawbacks, while considering the flaws of AI as pains of using a technology in its early days, and having faith in its prospects as well as its accuracy and efficiency.

Keywords: AI, hallucination, uses and gratifications, utility theory, LLM, confidence in AI.

1 Introduction

1.1 Background

Over a very short time, in less than five years between when AI first became a cause for global uproar and now, the rapid deployment of this technology across all sectors has boosted the necessity for understanding how professionals perceive its advantages and its limitations. The concern for this paper is that, for some professionals, in particular contexts, in perception, mostly favorable to AI, raises to the level of ‘deification,’ which is the phenomenon that elevates an entity to a god-like status, characterized by undue trust or faith and deliberate ignorance of its shortcomings. As per Gillespie et al. (2025), 58% of people considered AI systems to be trustworthy, with 65% having faith in their accuracy or reliability, while 52% showed some concern over ethics or safety questions. Conversely, the faith in the value of AI or its prospects does not go down

despite admission of its flaws. Finland, being a good example of this, shows that people consider AI helpful, despite not having high trust in it (Gillespie et al., 2025). In the same vein, attitudes towards AI seem to vary across national boundaries. Power (2025) reports that in China, 70% people believe AI can resolve serious socioeconomic problems, with the number being very high among 18 to 34-year-olds compared to Western nations. Notably, Chong et al. (2023) showed that in the context of human designers, confidence in AI declines with exposure to poor performance of this technology but does not get any expected boost from subsequent good performance of AI tools. These findings are contrary to the findings of Chong et al. (2022), which showed, using the example of chess puzzles, that confidence in AI and in oneself both are proportional to the quality of performance. This shows that the patterns of AI adoption and confidence in this technology are not uniform across sectors, and people in different professional fields value AI differently.

It is important to note here that trust and optimism regarding AI exists despite the drawbacks and limitations being public knowledge, qualifying this confidence as a form of deification, characterized by selective blindness towards particular information. Hoffman (2025) reports that AI hallucination, whereby even advanced Large Language Models (LLMs) fabricate information confidently, and a lack of “truly novel” ideas are the key limitations of AI at present. Hallucinations are a critical issue, as they can have severe consequences unless the AI-generated results are thoroughly scrutinized. Examples include Med-Gemini of Google hallucinating a non-existent part of the brain and ChatGPT of OpenAI falsely accusing a man of murder (Hoffman, 2025). Conversely, since AI depends on being trained by the existing corpus of human knowledge, its ideas, some degree of error and misinformation should be expected. Rather, it seems, the obligation to re-view the AI-generated responses can be a major inconvenience and a point of inefficiency. Neuberger (2025) also claims that AI results are mere autocompletion as LLMs predict, based on known patterns, how a conversation progresses and how to respond to certain questions. This implies that AI does not have any critical thinking skills, which makes AI responses, particularly on complex and nuanced matters, unreliable. Jeung and Huang (2023) reveal that in their study, the participants showed a mismatch between expected reactions to mistakes made by AI and actual reactions, which were curiosity rather than frustration or annoyance. Specifically, when the AI gives a mistaken response, users become interested in understanding why the AI gave that response, instead of criticizing the LLMs. However, some participants did express some degree of disappointment due to AI mistakes (Jeung & Huang, 2023). This paper explores why this deification of AI persists despite cases of AI making mistakes or LLMs behind the AI having fatal flaws, which are widely reported. It examines how the perception of AI is formed across key sectors, using Uses and Gratification Theory and Utility Theory.

1.2 Research Aim

The aim of this study is to understand the role of user gratification in the deification of AI and assess how this tendency affects the perception of AI limitations.

2 Theoretical Framework

2.1 Uses and Gratification Theory (UGT)

According to UGT, media consumption is driven by purpose despite appearing otherwise. In this theory, consumers are not manipulated by media; they are aware of its uses, implying that they make a conscious decision to choose one media outlet over another (Ungvarsky, 2020). As such, media houses compete for greater viewership among themselves, and one of the ways they do this is by curating content to cater to consumer interests. Falgoust et al. (2022) argue, quoting Shao, that users engage with media in three ways: consuming, participating, and producing, and also assert that different uses of media are driven by different needs. According to UGT, there are five key consumer needs, namely cognitive, integrative, affective, social integrative and tension-release, while four bases to gratify consumers include modality, agency, interactivity, and navigability (Vinney, 2025). This paper examines what role gratification plays in driving positive interpretation of the capabilities of AI.

2.2 Utility Theory (UT)

As per UT, the choice of consumable materials is freely made by consumers so long as there is no external constraint on doing so. Using 20th-century theorists like Neumann and Morgenstern, this theory proposes that economic actions of humans are determined by what utility they ascribe to an object (Hou, 2025). While it was previously believed that these choices are strictly mathematical, modern understanding dictates that humans only organize their choices by order of preference. This theory is applied here to assess how the perceived utility of certain AI tools and services influences the perception of AI among professionals.

2.3 Literature Review

While the pattern of AI adoption and its degree vary from one sector to another, the confidence in AI among professionals in that field showcases some commonalities. Painter et al. (2023) argue, in the context of healthcare, that trustworthiness, seamless “interoperability with existing systems” are determiners of confidence in AI. Evidence shows that human professionals accept AI decisions when they are confident in themselves and reject suggestions when they are not (Chong et al., 2022). As such, when human self-confidence is impacted by AI performance, better AI performance should result in greater confidence in technology, while poor performance should shatter confidence, with the self-confidence of the users playing a mediatory role in this human-machine engagement. Chong et al. (2023) also identify the minimal impact of AI-suggested actions as a factor in losing confidence in it. Hence, the question of why LLM power-users continue to use AI, showing a high degree of confidence and trust, despite exposure to AI mistakes, becomes pertinent. On a slightly different note, Übellacker (2025) identifies that professionals are bound to refine “their interpretations of how AI

fits into their work” upon encountering the various limitations of AI, including AI hallucinations, biases, token-length constraints, and so on. Consequently, this paper is expected to show the qualitative drivers of confidence in AI in the context of Chinese LLM users.

3 Methodology

3.1 Sampling and Data Collection

This study follows a primary qualitative research design, where data were collected using focus group discussions of 16 participants, divided into 4 groups as per sectors of their professions. Discussed is semi-structured as the moderator is present and initiates the conversation, but does not intervene unless it is to redirect the conversation back to the topic at hand. It used a snowball sampling method, which is to say, each participant was to recruit more participants to the study. This method is useful when the people conducting the study do not have ready access to the target group (Johnson, 2014). The final participants were selected according to the selection criteria. This means the participants must be 21 years or older, employed in the private sector, have experience using AI tools or LLMs, live and work in China, and consider themselves “tech-savvy” and among the “power-users.” Finally, the four sectors are Information Technology (IT), healthcare, education or academia, and software development.

3.2 Data Analysis

To analyze qualitative data, a thematic data analysis method was used in this paper. As per Clarke and Braun (2017), thematic data analysis means identifying patterns in the data and codifying the common elements to derive themes. Codes include semantic and latent codes, which respectively mean literal and underlying or implied meaning. Themes assessed and discussed in light of theories and established knowledge to develop a complete understanding.

4 Findings

4.1 Summary of Coding

Table 1. Small Coding Table for Thematic Analysis.

Theme	Representative Excerpt	Semantic Code	Latent Code
Mitigation strategies	"...compulsively verify AI outputs..."	Manual review of output	Awareness of flaws creates new "molded" workflows.
Perceived Superiority	"...it is quicker than any human..."	Efficiency in manual labor	Ignoring limitations in favor of speed.
Market Pressure	"AI is the future, isn't it?"	Influence of slogans	Fear of missing out; following trends without foresight.

Gratification	"...extra time to reduce workload and stress..."	Relief of task loads	Satiating "tension-release" needs via UGT.
Rationalization	"Those cases [hallucinations] are outliers."	Dismissal of errors	Selective blindness to maintain the "deification".
External Influence	"AI is the future, isn't it?"	A catchy slogan to get the point across	Influence of popular narratives
Doubt	"...very early stages of adopting AI..."	Much work to be done to complete AI integration	Cautious and careful approach to AI adoption.

Note. This table is a shortened but representative version of the complete details of extracted data and its codification.

4.2 Summary of Findings and Discussion

A key finding of the study is that awareness of flaws of AI does not correlate with willingness to adopt it, which is to say, knowledge of how AI can and has made mistakes does not imply that the participants are reluctant to adopt it. A representative case is a participant recalling that they faced trouble due to mistakes by AI and chose to verify AI responses thereafter (Table 1). Thus, a workaround or mitigation mentality prevails, whereby professionals are willing to adopt AI, and when faced with challenges, they will add a new step to their workflow, like manual oversight of AI, instead of foregoing AI completely. Notably, most participants seem to downplay the flaws of AI as outliers and not enough to deter them from using this technology. As Sundar and Kim (2019) pointed out, the performance of a machine is often associated with "precision and low error rates." This confirmation bias works in favor of AI, as professionals, despite encountering troubles with this technology, are likely to conceive of them as 'one-off' incidents and not a systemic issue. Hangl et al. (2023) also identify that despite the difficulty in integrating with existing workflows and processes being a major barrier in AI adoption, improved decision-making, and efficiency, as well as increased accuracy and productivity, remain among the major drivers of the same. Moreover, the belief that AI is more accurate and faster than human workers is almost universal. Some participants showcase an "early adopter" mindset that enduring some setbacks in the early days of a technology will hasten future improvement. Among the gratification and utility drivers behind these attitudes, the convenience of quickly accessing information in a succinct way seems highly relevant. Gillespie et al. (2023) found that market trends pressurize individual professionals as well as organizations to adopt AI-based technology, and management staff are increasingly proficient with it. As such, using AI is quickly becoming the default in different sectors, and companies are also pushing AI aggressively on employees to keep up with the market demand.

Awareness of AI flaws translates not into actions to improve the AI models but molding the human part of the workflow to better suit the necessity of using AI "properly" (Table 1). This generally means redundant processes to quality check AI outputs, even when the AI response is accurate. This is the reason why in the data, admissions of flaws in AI are often tacit or implied but not clearly asserted. Notably, even the few participants who argued to point out AI limitations are also willing to accept AI (Table 1). However, since the sample for this study has been limited by geographical region, these results may not be replicated in regions with distinct society and

culture. This geo-graphical variation is supported by research, as Li et al. (2025) find through their study of cross-country assessment of levels of trust in AI that countries like France and Australia demonstrate a degree of skepticism towards AI, while nations such as China and India showcase the most trust. Burke and Akhtar (2023) report that data bias is a major limitation and point of concern regarding AI adoption, as present LLMs depend heavily on already available information to trace patterns and build models, which can be hindered due to the data itself being biased or not being impartially fed to the LLM. As such, data usage-related questions or doubts about AI also include the security risk of possessing data, accessibility, and ethical dilemmas about data ownership. In a technology-driven market, showcasing adeptness in using AI maintains the status and credibility of professionals, while the quickness of AI responses relieves some work pressure. Peer expectations are also a determinant of behavior, as in healthcare, where margin for error is almost zero, AI adoption is mostly in assistive roles, while in education, it is often rejected for task completion due to the high value of self-learning and skill development.

5 Conclusion

The overall findings indicate that the deification of AI is not a result of a lack of knowledge of its shortcomings, but rather the very highly perceived utility of this new technology, so that occasional setbacks are not enough to deter people from it. While encountering a challenge with AI, professionals are more likely to think about workarounds rather than discarding it altogether. They are willing to reshape their workflows to accommodate AI, while managing its flaws manually, rather than forswearing its benefits. It is evident here that manual supervision of AI outputs is necessary, and not optional, to avoid the pitfalls of AI adoption. Another major aspect is that users perceive AI as a superior technology compared to existing tools at the disposal of the average person. AI is thought to be faster, easier to use, and more accurate compared to humans, hence more reliable. Notably, even when the users acknowledge the drawbacks of AI, they showcase great optimism that those issues will be fixed soon. This also highlights the need for critical knowledge of how present models of LLMs and AI work. AI adoption for professionals can be a purely pragmatic commercial decision, with market trends, demand, and popular narratives pushing for this technology aggressively, and competitor companies adopting it rapidly. However, these findings have some key limitations: Since qualitative data was collected using self-reporting (participants expressing their own opinions, assuming they are speaking truthfully), the scope for generalizing the results was limited. Conversely, since the data and findings are qualitative, the study lacks empirical evidence in the findings dimension, and future research will benefit from a mixed-method approach or a quantitative study. Future studies may also broaden the focus from China to other countries to improve the scope for results being generalized.

Acknowledgments

A third-level heading in 9-point font size at the end of the paper is used for general acknowledgements, for example: This study was funded by X (grant number Y).

Disclosure of Interests

It is now necessary to declare any competing interests or to specifically state that the authors have no competing interests. Please place the statement with a third-level heading in 9-point font size beneath the (optional) acknowledgements, for example: The authors have no competing interests to declare that are relevant to the content of this article. Or: Author A has received research grants from Company W. Author B has received a speaker honorarium from Company X and owns stock in Company Y. Author C is a member of committee Z.

References

1. Burke, S. A., & Akhtar, A. (2023). The shortcomings of artificial intelligence: A comprehensive study. *International Journal of Library and Information Science*, 15(2), 8-13. <https://doi.org/10.5897/IJLIS2023.1068>
2. Chong, L., Raina, A., Goucher-Lambert, K., Kotovsky, K., & Cagan, J. (2023). The evolution and impact of human confidence in artificial intelligence and in themselves on AI-assisted decision-making in design. *Journal of mechanical design*, 145(3), 1-6. <https://doi.org/10.1115/1.4055123>
3. Chong, L., Zhang, G., Goucher-Lambert, K., Kotovsky, K., & Cagan, J. (2022). Human confidence in artificial intelligence and in themselves: The evolution and impact of confidence on adoption of AI advice. *Computers in Human Behavior*, 127(2022), 1-10. <https://doi.org/10.1016/j.chb.2021.107018>
4. Falgoust, G., Winterlind, E., Moon, P., Parker, A., Zinzow, H., & Madathil, K. C. (2022). Applying the uses and gratifications theory to identify motivational factors behind young adults' participation in viral social media challenges on TikTok. *Human Factors in Healthcare*, 2(2022), 1-14. <https://doi.org/10.1016/j.hfh.2022.100014>
5. Gillespie, N., Lockey, S., Curtis, C., Pool, J., & Akbari, A. (2023). *Trust in artificial intelligence: A global study*. The University of Queensland and KPMG Australia. <https://aiunplugged.io/wp-content/uploads/2023/10/Trust-in-Artificial-Intelligence.pdf>
6. Gillespie, N., Lockey, S., Ward, T., Macdade, A., & Hassed, G. (2025). *Trust, attitudes, and use of artificial intelligence: A global study 2025*. The University of Melbourne and KPMG. <https://doi.org/10.26188/28822919>
7. Hangl, J., Behrens, V. J., & Krause, S. (2022). Barriers, Drivers, and Social Considerations for AI Adoption in Supply Chain Management: A Tertiary Study. *Logistics*, 6(3), 2-22. <https://doi.org/10.3390/logistics6030063>
8. Hoffman, S. (2025). *Understanding the Limitations of AI*. AlphaSense. <https://www.alpha-sense.com/resources/research-articles/limitations-of-ai/>
9. Hou, H. (2025). Utility Theory Application in Decision-Making Behavior for Energy Use and Management: A Systematic Review. *Energies*, 18(8), 1-19. <https://doi.org/10.3390/en18082125>

10. Jeung, J. L., & Huang, J. Y. C. (2023, October). Correct me if I am wrong: exploring how AI outputs affect user perception and trust. In *Companion publication of the 2023 conference on computer-supported cooperative work and social computing* (pp. 323-327). <https://doi.org/10.1145/3584931.3606997>
11. Johnson, T. P. (2014). Snowball sampling: Introduction. In N. Balakrishnan, T. Colton, B. Everitt, W. Piegorsch, F. Ruggeri, & J. L. Teugels (Eds.), *Wiley StatsRef: Statistics Reference Online*. John Wiley & Sons. <https://doi.org/10.1002/9781118445112.stat05720>
12. Li, Y., Wu, B., Huang, Y., Liu, J., Wu, J., & Luan, S. (2025). Warmth, competence, and the determinants of trust in artificial intelligence: A cross-sectional survey from China. *International Journal of Human-Computer Interaction*, 41(8), 5024-5038. <https://doi.org/10.1080/10447318.2024.2356909>
13. Neuberger, C. (2025, November 11). *Beyond the Buzz: Understanding the Real Limitations of AI*. Invenity. <https://www.invenity.com/2025/11/11/the-real-limitations-of-ai/>
14. Power, J. (2025, November 19). *Trust in AI far higher in China than West, poll shows*. Al Jazeera. <https://www.aljazeera.com/economy/2025/11/19/trust-in-ai-far-higher-in-china-than-west-poll-shows><https://www.aljazeera.com/economy/2025/11/19/trust-in-ai-far-higher-in-china-than-west-poll-shows>
15. Sundar, S. S., & Kim, J. (2019, May). Machine heuristic: When we trust computers more than humans with our personal information. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems* (pp. 1-9). <https://dl.acm.org/doi/10.1145/3290605.3300768>
16. Übellacker, T. (2025). Making sense of AI limitations: how individual perceptions shape organizational readiness for AI adoption. *ArXiv*. <https://doi.org/10.48550/arXiv.2502.15870>
17. Ungvarsky, J. (2020). *Uses and gratifications theory*. EBSCO. <https://www.ebsco.com/research-starters/communication-and-mass-media/uses-and-gratifications-theory>
18. Vinney, C. (2025, November 13). What Is Uses and Gratifications Theory in Media Psychology?. VeryWell Mind.

Open Access This chapter is licensed under the terms of the Creative Commons Attribution-NonCommercial 4.0 International License (<http://creativecommons.org/licenses/by-nc/4.0/>), which permits any noncommercial use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.

