



Institutionality of AI-Human Interaction: A case study of PDFgear Copilot using conversation analytic approach

Rehna Sotto

University of Jyväskylä, Finland

Corresponding Email: rehna.j.sotto@jyu.fi

Received: 13-07-2025

Reviewed: 055-08-2025

Accepted: 30-09-2025

Abstract

This study sought to advance our understanding of the anthropomorphic features of artificial intelligence (AI) and its usability in research and society, in general. Specifically, this study examined how ‘institutional’ is the interaction between an AI tool (PDFgear Copilot) and the human researcher during an AI-integrative process of a systematic literature review. In the final screening process, a total of 104 publications were found relevant and chatted with PDFgear Copilot for the summarization of the methods, sample, theoretical foundations, and findings. The conversations from these chats were analyzed using the features of institutional interaction by Drew & Heritage (1992). The results revealed that AI exhibits a ‘normative response’ in 96 research articles and 8 with a ‘disruptive response’ in the conversation sequence organization. When given follow-up questions, it was observed in 17 research articles that AI showed more anthropomorphic traits with similarity to ordinary conversation when AI expressed a degree of uncertainty and answer limitation. Overall, this study provides implications for information technology professionals in advancing AI’s human-like features and for researchers in further exploring the possibility of utilizing AI in research.

Keywords: AI, PDFgear Copilot, AI-human interaction, conversation analysis, anthropomorphism

Introduction

Methodological innovation in scientific research has been widespread over the years. This innovation encourages researchers to find ways to make research methods more efficient and useful to society (Jewitt et al., 2017). In the digitalized world, some of these innovations help researchers find answers to societal problems with the use of technology (Xiao & Su, 2022). The scientific breakthroughs with the use of artificial intelligence (AI), for example, created debate among scholars about its uses in enhancing methodological innovations in academic research (Ahmad et al., 2023) and organizational studies (Jöhnk et al., 2021). However, researchers from different fields offered different views on the application of AI, which continued the long-running debate as to whether it is ethical or not. Some research from

an ethical perspective question the morality of AI (Wilson et al., 2022) and the consequences of ethical scientific practice in research (Bouhouita-Guermech et al., 2023). On the other side of the societal spectrum, some researchers in education, business, and engineering offered some insights into the use of AI. This includes advancements of pedagogy in education (Sperling et al., 2024), in enhancing academic well-being among learners (Xiao et al., 2024), improving the healthcare system (Bajwa et al., 2021), and enhancing human work efficiency (Zirar et al., 2023).

In recent times, one of the sought-after discussions surrounds the Large Language Models (LLMs). These LLMs are AI-powered technology having anthropomorphic capabilities (Deshpande et al., 2023). These AI technologies have different uses for humans, such as the construction of essays, summarization of text, finding literature for research, and other features that apply to any scientific and societal fields. Technological advancement is within everybody's reach, and this societal development is inevitable. However, there is less research on how these LLMs could potentially aid in doing research and innovating an existing method to answer scientific inquiries. This paper expands research related to the use of an AI-powered tool to provide knowledge on the actual use of these technologies in the actual research process, which in turn helps create an informed and ethical digital society.

Literature review

AI tool: PDFgear Copilot

The topic of the use of artificial intelligence (AI) has been widely spread through business innovation, technological advancement in healthcare, and other research endeavors (Rashid & Kausik, 2024). One of the popular topics is the use of AI-powered LLMs in generating answers to a question through its chat features, finding suitable research journals for a specific topic, making summaries of reports, and more (Giannakopoulos et al., 2023; Johnson & Paulus, 2024). However, some of the available LLMs require a subscription process to use their services. Not known to many is that some AI-powered or powered by ChatGPT PDF assistants are free to download and use, and could be available to everyone. One of these is the PDFgear Copilot. This is an AI-powered technology (GPT 3.5) that uses natural language processing (NLP) and is believed to be useful in research (McLean, 2024). Below is a summary of PDFgear Copilot features (pdfgear.com).

Chat function. This feature means that a researcher, for example, can chat with the PDF and generate questions a researcher would wish AI to answer based on the text of the PDF. This AI tool allows the use of natural language to understand the question and give answers in a human-like way of interacting.

Summarization. This tool has the capability of summarizing lengthy files, like research paper information, into a brief and concise summary, which would help save time and effort in finding the important information and condensing the information.

Institutionality of AI-Human Interaction: A case study of PDFgear Copilot using conversation analytic approach

Extraction of key information. This feature allows the customization of the information a researcher, for instance, would want to extract. Using simple keywords and phrases, the prompt will generate the information that relates to the chosen words.

Secondary Confirmation Command. This feature allows clarification of interaction between the user and AI in case of misinterpretation or asking for clarity of the questions being posted. This feature further allows correction to the misunderstanding of concepts on the side of AI.

Institutionality of Interaction

The anthropomorphic nature of AI drives the transformation in which humans interact with machines (Alabed et al., 2022). One of the anthropomorphic features is its ability to be able to generate answers to human questions and act like a service provider of knowledge (Simas & Ulbricht, 2024). Due to this anthropomorphism of AI, humans became the clients in the interaction, relying upon the services of these AI-powered platforms. According to the concepts of conversation analysis, drawing specifically on the institutional interaction principles, an interaction becomes institutional when it is goal or task-oriented (Ruusuvuori, 2000). In this case, we can assume that in this service encounter, when utilizing AI-powered platforms and the like, there is an institutionality between AI-human interaction in the digital space.

The features of institutional interaction based on Drew & Heritage (1992) include interactional asymmetry, organization of interaction, sequence organization, turn-taking, and lexical choices. Interactional asymmetry means that in the interaction between a ‘professional’ and the ‘client’, the ‘professional’ is considered more knowledgeable on the subject matter. The ‘professional’ typically has control over the interaction. In the organization of the organization, this investigates the established activities that the participants have created during the conversation to meet the goal of the conversation. Sequence organization, on the other hand, describes how adjacency pairs like question-answer have been reduced, expanded, or structurally adapted. In terms of turn-taking, this could investigate the flow of conversation and knowing who was allowed to speak and if there is a delegation of the next speaker. Lastly, the lexical choice examines the choice of word usage that is specific to the goal-setting conversation.

With the application of these features, we could gain a deeper understanding of the institutionality of AI-human interaction. For instance, on how AI acts as the professional during the service interaction, in exploring the series or organization of activities during the service encounter (asking for questions or help from AI), in selecting the next writer, if AI can expand answers or will it be very limited, and on AI use of words to reflect the academic field. The research question, therefore, in this study is, **“How institutional is the interaction between AI and human during the goal-oriented conversation?”**.

Exploring these features could provide knowledge of how anthropomorphic AI interacts with humans, the institutionality of its interaction, and gain an understanding of the advantages and disadvantages of utilizing AI.

Research method

Data Source

Figure 1 below illustrates how I gathered the data for this study during the Spring of 2024. In the process, first, I conducted a search across engines about my research topic of interest, which generated 740 articles. After gathering them, I manually removed the duplicates using a spreadsheet. Then, I screened the research abstracts against the criteria. There were 104 research publications that were found to be relevant. Instead of the traditional systematic review method, where a researcher must read all 104 publications, I chose to innovate the method and deploy an AI-integrative process with the help of an AI tool, PDFgear Copilot. This was done to all 104 relevant research papers to assess them for eligibility and inclusion. I chatted with the human-like AI and was tasked to extract the key information. To be consistent with the process, I sent similar chat messages for every paper. This interaction was then copied and saved in Microsoft Word to be ready for analysis.

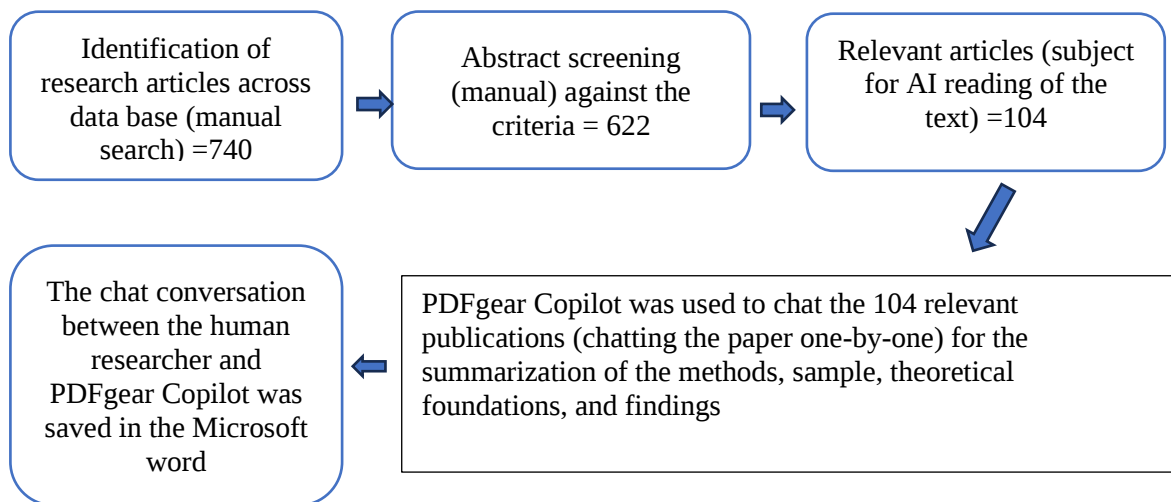


Figure 1. Data gathering process

Analysis

I analyzed the data gathered from my interaction with the AI tool, PDFgear Copilot, in all 104 papers that were screened for eligibility and inclusion. To explore the human-like nature of AI in interaction, I chose to employ conversation analysis to analyze the institutional features of the interaction between me, as the researcher (R), and the AI (PDFgear Copilot). In this study, I was considered the client who was asking questions or requesting services to the service provider AI. The analysis will be guided by the institutional features of interaction cited by Drew & Heritage (1992). Specifically, I looked at how ‘institutional’ is the AI-human interaction through the features of general organization, sequence organization, and turn-taking, and lexical choice during the goal-oriented conversation. This process also allowed me to point out how AI treated the questions as a ‘real question’ and elaborated on the similarities in an ordinary conversation.

Institutionality of AI-Human Interaction: A case study of PDFgear Copilot using conversation analytic approach

Results

Normative response

In the organization of interaction, of 104 publications investigated, 96 showed what can be classified as ‘normative response’ in answering the task question in the chat. The interaction follows the adjacency pairing of the question-answer sequence. The AI first threw questions to the researcher (R) on what to do with the paper, offering a service-like nature to its client, giving an option to either use the prompt or not, and allowing to chat about the research article in question. Below is the general organization of conversation,

- 1 R : ((opening the PDF file))
- 2 AI: ((sending the available prompts in the chat))
- 3 R : ((chat the first question))
- 4 AI: ((answer the first question))
- 5 -->AI does not select the next writer
- 6 R : ((self-select, ask the second question))
- 7 AI: ((answers only the second question))

The sequence continues to be the same until there are no more questions from ‘R’. From this kind of organization, it was noted that the sequence tends to be similar when the publications were empirical studies. It can also be noted that AI failed to select the next writer in the chat to continue the flow of conversation or use a follow-up form of questions, which may not imitate a human service professional. For instance, it lacks what can be commonly heard from human service professionals, “*Is there anything I could help you with?*” or “*Do you have other questions?*”. This result may reveal how a certain AI tool (with its programming) may not fully execute human-like features, but can serve humans in some ways, like answering a specific question and giving the key information that humans are searching for.

Disruptive response

Eight out of one hundred four publications exhibited what I classified as a ‘disruptive response’ in answering the first question in the chat box. Instead of the perceived normal question-answer adjacency pairs in the conversation, another question as a follow-up or clarification was sent back to me. The organization of interaction follows a question-question-answer sequence. This kind of sequence was observed when the publications investigated were literature review articles. The clarification on the chat messages is bound by its orientation to the lexical choices that were sent in the chat and the limitations of its built-in programming capacity. I will elaborate on these findings one by one in the subsections below.

Lexical-oriented

As mentioned, when the article in question is a literature article, the AI asked a follow-up question instead of answering the task question. When a certain word was changed and sent

to the chat box as a new question, the AI was able to detect and give the answer when it was satisfied with the word usage. Thus, it was bounded by the choice of words that made the AI respond and act in the right direction. Below is a representative excerpt from the chat.

1 R : summarize the **article** in terms of methods, sample,
2 theoretical foundations, findings
3 AI: I'm sorry, **but which article** are you referring to?
Can
4 you please provide more specific information so that I
5 can answer your question accurately?
6 -->*gave the number links*
7 R : summarize the **pdf** in terms of methods, sample,
8 theoretical foundations, findings
9 AI: I'm sorry, **but which PDF** are you referring to? Can
you
10 please provide the name or topic of the PDF so that I can
11 answer your question accurately?
12 -->*gave the same number links*
13 R : summarize the **current PDF file** in terms of methods,
14 sample, theoretical foundations, findings
15 AI: The **current PDF** file, titled... ((explained the
answer))

In another scenario, using the 'current PDF' lexical choice did not work. The AI still wanted some clarifications and the need for writing the article title was found necessary, like in the below representative excerpt:

1 R : summarize the **article** in terms of ...
2 AI: I'm sorry, but could you please provide **which**
article you
3 are referring to? There are several articles mentioned on
4 pages 12 and 13 of this PDF.
5 -->*gave the number links*
6 R : summarize the **current pdf** in terms of ...
7 AI: I'm sorry, but **there doesn't seem to be a single**
article
8 **in this PDF**. The pages contain references to several
9 articles and research studies. Could you please specify

Institutionality of AI-Human Interaction: A case study of PDFgear Copilot using conversation analytic approach

- 10 which of these you are referring to?
- 11 -->gave the same number links
- 12 R : summarize the **current pdf**, (**title of the article**), in terms...
- 13 AI: **Certainly!** (**title of article**), is an article..
- 14 ((explained the answer))

From these excerpts, it can be noted that AI showcased human-like features in conversing in the chat by uttering phrases like ‘I am sorry’, ‘can/could you please’, ‘certainly’, and the use of the personal pronoun ‘I’ in its answer. In terms of turn-taking, the AI was able to select the next writer in the chat box by asking a follow-up or clarification question, making it imitate the humanness of conversation, or in other words, showing aspects of ordinary conversations. These results revealed that human-like features of AI are more prevalent when it is not an empirical study article. This result further suggests that the built-in programming of this AI tool is made for a certain function, which researchers in the field need to be aware of.

AI's built-in capacity

Another observation where it had a question-question-answer sequence was when the built-in program capacity was reached. This means that AI would not be able to do the summarization tasks, like in the examples mentioned above, rather, it admitted its limitations in summarizing long pages of research articles. This can be shown in the below representative excerpt:

- 1 R : summarize the **article** in terms of ...
- 2 AI: I'm sorry, but I don't know **which article** you are
- 3 referring to as the given pages of the PDF contain
- 4 information 4about various articles. Could you please
- 5 **provide more information or specific page numbers?**
- 6 -->gave the number links
- 7 R : summarize the **current pdf** in terms of ...
- 8 AI: I apologize, but **I am not able to provide a**
- 9 **comprehensive**
- 10 **summary** of the entire PDF **as it consists of 37 pages**
- 11 covering various topics from different fields of study.
- 12 Could you **please specify a particular segment or page**
- 13 **range** that you would like me to summarize?
- 13 -->gave the same number links

From this kind of scenario, the AI could not provide any summary of the publication in question. This could mean that researchers need to read the full paper to extract the main points of the research. This result may suggest that AI summarization features in this specific kind of tool could potentially summarize the paper within its page range. The built-in programming needs improving to be able to extract the summary of long research articles to make it a more comprehensible and reliable aid for researchers.

AI's answer limitation

Aside from the noted human-like features of conversation which occurred when the research article was a literature review (as already mentioned above in section 3.1 and 3.2), observations from the other two follow-up questions (i.e., is the article related to shift work?, do the findings focus on team leader interaction?) that were thrown in the chat suggest the humanness of answers. This is when the AI is more likely to have not found the information from the research article at hand or if it did not know the answer to the questions. This kind of scenario occurred in 17 research articles, out of the total 104, being examined. Representative excerpts below show these instances.

Instance 1.

- 1 R : is the article related to shift work?
- 2 AI: **I don't know.** The information provided in the PDF
does
- 3 not mention anything about shift work.
- 4 -->*gave the number links*

Instance 2.

- 1 R : is the article related to shift work?
- 2 AI: **I don't know.**
- 3 -->*gave the number links*

Instance 3.

- 1 R : do the findings focus on team leader interaction?
- 2 AI: **I'm sorry,** but I couldn't find any information
regarding
- 3 team leader interaction in the given pages... **If you
have**
- 4 **if you have any other questions,**
- 5 **please don't hesitate to let me know.**
- 6 -->*gave the number links*

Institutionality of AI-Human Interaction: A case study of PDFgear Copilot using conversation analytic approach

The humanness of the conversation displayed the AI's limitations (I don't know, I'm sorry), uncertainty (I don't know), and service-oriented approach (please don't hesitate to let me know). This kind of situation does not typically exist when the information being asked of the AI can be found in the article. This result, however, encourages researchers to search for the information themselves and read the full article. Also, this further implies that an AI tool could not always augment a researcher's job and search for key information.

Pros and Cons

Based on the grounding of the above-mentioned (sections 3.2 and 3.3) observations, the AI-human interaction has its advantages and disadvantages during the AI-integrative process of systematic literature review. In terms of the positive outcome, the AI tool, specifically PDFgear Copilot, was able to deliver the needed information in all 96 empirical articles being examined. It provided the article summary where the methods, sample, theoretical foundation, and findings were presented, and it provided the number of links for me (the researcher is the only live human involved) to click, read, and confirm the appropriateness of the information in the chat message, and make own verdict of the information presented. It helped me to find the information in a quick manner and, the goal was met in a shorter period. The conversation did not possess a degree of AI humanness, however, it serves its duty in the summarization and extraction of vital information from the research article.

However, the disadvantages are in the limitation of answers, the degree of uncertainty, and the built-in program-related features in detecting the important information. It could not function properly when the article reached beyond the page limit, and if it is a literature review. The presence of humanness in the conversation could be present, but it is not efficient in doing its task, which requires an intervention of editing the chat messages or reading the full paper. The results, therefore, question the importance of the human-like feature of the AI tool for the necessity and correctness of extracting and giving the right information of the research paper to the human user. This further inquires into the importance of developing AI with anthropomorphic characteristics in interacting with humans, and to what degree is more important in the process of sharing the knowledge or information in connection with the AI tool's function and usability.

Discussion

The purpose of this study is to examine the institutionality of the interaction between AI and human in goal-oriented conversation and elaborate on the pros and cons of interacting with and utilizing AI. The results of the study found that AI exhibits limited human-like features of interaction when the articles being examined are empirical studies. The AI can directly answer the task or question, but has inadequate word prompts or phrases to encourage the continuous flow of conversation, which may lead to questioning the trustworthiness of AI. This observation can be related to previous studies questioning the anthropomorphic characteristics of AI in conversation with humans (Placani, 2024).

However, an opposite observation was noted when the articles being investigated were literature review articles or long research articles. The AI showcased more anthropomorphic features by using words or phrases for clarity of the questions being thrown in the chat box. Another scenario where AI showed human-like prompts was when the program had reached its limit and when it did not know the answer or could not find the information from the research paper. This raised questions about how human-like AI should be in the conversation to be enough to imitate humans, to increase the trustworthiness of the application. Trustworthiness of AI requires understanding how the machine works and is actually used (Lahusen et al., 2024), and is assessed by the human users (Ferrario, 2024; Schlicker et al., 2025). It was also observed that the interaction between the AI and the human (the researcher) is more researcher-driven to meet its goal in the conversation. This suggests that AI tools are still evolving and may have not reached their full anthropomorphism and autonomous potential. This situation requires understanding of the potential risk of the AI's autonomous feature and its anthropomorphic design, to evaluate its validity and reliability (Abramoff et al., 2020). In addition, this could also imply that, in this case, the AI tool still needs a human to configure and update its features to be able to become a more responsible and reliable aid in doing research.

Implications for Designing AI's anthropomorphic feature

This study provides implications for IT professionals in finding ways in which prompts and answers could become more service-oriented without compromising the trustworthiness and reliability of the AI tool. For instance, based on the results, AI only answers the question (question-answer sequence) and exhibits more of a bot than human-like when the research papers being asked were empirical studies. When it exhibited human-like features, when papers were literature reviews or beyond page range, the lexical choices were mostly limited to 'I am sorry' and 'I do not know', and the use of 'please'. An in-depth comparative study of human behavior and AI behavior from different applications (including other AI tools) may guide the development a more human-like features of AI. Programmers may benefit from collaborating with other experts studying naturally occurring data to compare, redesign, and improve AI's features to reach its goal of imitating human-human interaction and its autonomous potential.

Implications for AI's usability in research

Integrating AI in research is possible, but with caution. Researchers may need to be aware of the features of the specific AI tool they may be using to understand the limitations and the potential risks of these tools in augmenting researchers' jobs, like reading the full paper and extracting the needed information from the research paper at hand. For instance, in this case, the PDFgear Copilot may be useful to a wide range of users, like students and researchers. It is free to download, and one only needs an internet connection to chat with the research paper (in PDF format). This could provide a cheaper option for integrating AI in research-related tasks or information-seeking tasks, and in making a research synthesis grid. However, researchers would still need to make a confirmatory check of the correctness of information, trustworthiness, and reliability of the information that an AI tool has given in the chat in order to make ethical-sounding research.

Conclusion

This study expands research into the usability of large language models (LLMs) such as GPT-powered PDF chat assistants that use natural language processing (NLP) in academic research. This provides insights into the integration of AI in innovating a research method, the extent to which degree of anthropomorphic features are necessary to deliver its functionality in doing research, and the perceived advantages and disadvantages of utilizing AI. By far, as to my knowledge, this is the first to explore the use of a conversation analytic approach, drawing on the institutional interaction features by Drew & Heritage (1992), in investigating the institutionality of AI-human interaction. However, this study also has its own limitations. The study focused on one case, which is the case of AI tool named PDFgear Copilot. Even taking this into account, the results could not be undermined as it contributes to the understanding of how an AI tool is actually used and how AI interacts in the process. Future research could investigate and involve other AI tools to have an in-depth comparison of the results. With regards to methodology, future research could be replicated by gathering a conversation between AI and human and analyzing signals identifying the AI's humanness in the conversation, and examining what constitutes trust and reliability in the conversation. Future research could also explore other ways of integrating AI in reviews, like utilizing AI in the whole process rather than combining AI and the traditional way.

Declaration of conflicting interest

The author reports no conflicting interests.

References

- Abràmoff, M. D., Tobey, D., & Char, D. S. (2020). Lessons learned about autonomous AI: Finding a safe, efficacious, and ethical path through the development process. *American Journal of Ophthalmology*, 214, 134-142. <https://doi.org/10.1016/j.ajo.2020.02.022>
- Ahmad, S. F., Han, H., Alam, M. M., Rehmat, M. K., Irshad, M., Arraño-Muñoz, M., & Ariza-Montes, A. (2023). Impact of artificial intelligence on human loss in decision making, laziness and safety in education. *Humanities & Social Sciences Communications*, 10(1), 311-14. <https://doi.org/10.1057/s41599-023-01787-8>
- Alabed, A., Javornik, A., & Gregory-Smith, D. (2022). AI anthropomorphism and its effect on users' self-congruence and self-AI integration: A theoretical framework and research agenda. *Technological Forecasting & Social Change*, 182, 121786. <https://doi.org/10.1016/j.techfore.2022.121786>
- Bajwa, J., Munir, U., Nori, A., & Williams, B. (2021). Artificial intelligence in healthcare: Transforming the practice of medicine. *Future Healthcare Journal*, 8(2), e188-e194. <https://doi.org/10.7861/fhj.2021-0095>

- Bouhouita-Guermech, S., Gogognon, P., & Bélisle-Pipon, J. (2023). Specific challenges posed by artificial intelligence in research ethics. *Frontiers in Artificial Intelligence*, 6, 1149082. <https://doi.org/10.3389/frai.2023.1149082>
- Deshpande, A., Rajpurohit, T., Narasimhan, K., & Kalyan, A. (2023). Anthropomorphization of AI: Opportunities and risks. *arXiv.org*.
- Drew, P., Heritage, J. (1992). *Talk at Work. Interaction in institutional settings*, Cambridge: University Press
- Ferrario, A. (2024). Justifying our credences in the trustworthiness of AI systems: A reliabilistic approach: Justifying our credences in the trustworthiness. *Science and Engineering Ethics*, 30(6), . <https://doi.org/10.1007/s11948-024-00522-z>
- Giannakopoulos, K., Kavadella, A., Aaqel Salim, A., Stamatopoulos, V., & Kaklamanos, E. G. (2023). Evaluation of the performance of generative AI large language models ChatGPT, Google Bard, and Microsoft Bing Chat in supporting evidence-based dentistry: Comparative Mixed Methods Study. *Journal of Medical Internet Research*, 25(1), e51580. <https://doi.org/10.2196/51580>
- Jewitt, C., Xambo, A., & Price, S. (2017). Exploring methodological innovation in the social sciences: The body in digital environments and the arts. *International Journal of Social Research Methodology*, 20(1), 105-120. <https://doi.org/10.1080/13645579.2015.1129143>
- Johnson, C. W., & Paulus, T. M. (2024). Generating a reflexive AI-assisted workflow for academic writing. *Qualitative Report*, 29(10), 2772-2792. <https://doi.org/10.46743/2160-3715/2024.7634>
- Jöhnk, J., Weißert, M., & Wyrski, K. (2021-02-01). Ready or not, AI comes— An interview study of organizational AI readiness factors. *Business & Information Systems Engineering*, 63(1), 5-20. <https://doi.org/10.1007/s12599-020-00676-7>
- Lahusen, C., Maggetti, M., & Slavkovik, M. (2024). Trust, trustworthiness and AI governance. *Scientific Reports*, 14(1), 20752-10. <https://doi.org/10.1038/s41598-024-71761-0>
- McLean, D. (2024, April 24). *10 Best AI tools for Research in 2024 (Compared)*. Elegant Themes Blog. <https://www.elegantthemes.com/blog/business/best-ai-tools-for-research>
- PDFgear (n.d.). PDF task made easy. Retrieved from PDFgear website
- Placani, A. (2024). Anthropomorphism in AI: Hype and fallacy. *AI and Ethics (Online)*, 4(3), 691-698. <https://doi.org/10.1007/s43681-024-00419-4>
- Rashid, A. B., & Kausik, M. A. K. (2024). AI revolutionizing industries worldwide: A comprehensive overview of its diverse applications. *Hybrid Advances*, 7, 100277. <https://doi.org/10.1016/j.hybadv.2024.100277>
- Ruusuvuori, J. (2000). *Control in the medical consultation: Practices of giving and receiving the reason for the visit in primary health care*. Tampere University Press. <https://trepo.tuni.fi/handle/10024/67087>
- Schlicker, N., Baum, K., Uhde, A., Sterz, S., Hirsch, M. C., & Langer, M. (2025). How do we assess the trustworthiness of AI? Introducing the trustworthiness assessment model

Institutionality of AI-Human Interaction: A case study of PDFgear Copilot using conversation analytic approach

- (TrAM). *Computers in Human Behavior*, 170, 108671. <https://doi.org/10.1016/j.chb.2025.108671>
- Simas, G., & Ulbricht, V. R. (2024). Human-AI interaction: An Analysis of anthropomorphization and user engagement in conversational agents with a focus on ChatGPT. *Intelligent Human Systems Integration (IHSI 2024): Integrating People and Intelligent Systems*, 119(119). <https://doi.org/10.54941/ahfe1004510>
- Sperling, K., Stenberg, C., McGrath, C., Åkerfeldt, A., Heintz, F., & Stenliden, L. (2024). In search of artificial intelligence (AI) literacy in teacher education: A scoping review. *Computers and Education open*, 6, 100169. <https://doi.org/10.1016/j.caeo.2024.100169>
- Wilson, A., Stefanik, C., & Shank, D. B. (2022). How do people judge the immorality of artificial intelligence versus humans committing moral wrongs in real-world situations? *Computers in Human Behavior Reports*, 8, 100229. <https://doi.org/10.1016/j.chbr.2022.100229>
- Xiao, D., & Su, J. (2022). Role of technological innovation in achieving social and environmental sustainability: Mediating roles of organizational innovation and digital entrepreneurship. *Frontiers in Public Health*, 10, 850172. <https://doi.org/10.3389/fpubh.2022.850172>
- Xiao, J., Alibakhshi, G., Zamanpour, A., Zarei, M. A., Sherafat, S., & Behzadpoor, S. (2024). How AI literacy affects students' educational attainment in online learning: Testing a structural equation model in higher education context. *International Review of Research in Open and Distance Learning*, 25(3), 179-198. <https://doi.org/10.19173/irrodl.v25i3.7720>
- Zirar, A., Ali, S. I., & Islam, N. (2023). Worker and workplace artificial intelligence (AI) coexistence: Emerging themes and research agenda. *Technovation*, 124, 102747. <https://doi.org/10.1016/j.technovation.2023.102747>