

# Politeness Strategies in Chatgpt Corrective Feedback: A Pragmatic Study of Face Threatening Acts in Indonesia Classrooms

 <https://doi.org/10.31004/jele.v11i5.2966>

Sofiah Rahmah Nasution<sup>a</sup> 

<sup>1</sup>STAI Samora Pematang Siantar , Indonesia

Corresponding Author: [sitifatihamsiregar09@gmail.com](mailto:sitifatihamsiregar09@gmail.com)

## ABSTRACT

The increasing use of artificial intelligence (AI) in language education has raised interest in the characteristics of AI-generated feedback. While previous studies have focused on the effectiveness of ChatGPT in improving writing, limited attention has been given to its interpersonal and pragmatic dimensions. Drawing on Brown and Levinson's politeness theory, this study investigates the politeness strategies employed by ChatGPT in providing corrective feedback on EFL learners' writing. Using qualitative content analysis, 150 feedback instances generated by ChatGPT-4 in response to essays written by 30 Indonesian undergraduate EFL students were analyzed. The findings indicate that negative politeness was the most frequently employed strategy, characterized by hedges, modal verbs, and indirect suggestions. Positive politeness and bald on-record strategies were also identified, whereas off-record strategies were absent. The results suggest that ChatGPT tends to deliver corrective feedback in a face-sensitive manner by framing corrections as suggestions rather than directives. The study highlights the importance of considering the pragmatic aspects of AI-generated feedback and contributes to the growing discussion of AI-assisted language learning.

**Keywords:** *Chatgpt, Corrective Feedback, EFL Writing, Face-Threatening Acts, Politeness Strategies, Pragmatics.*

### Article History:

Received 03<sup>rd</sup> July 2026

Accepted 04<sup>th</sup> September 2026

Published 10<sup>th</sup> September 2026



## INTRODUCTION

The rapid advancement of artificial intelligence (AI) has brought significant changes to educational practices, particularly in language learning and writing instruction. Among recent technological developments, generative AI tools such as ChatGPT have increasingly attracted attention due to their capacity to generate human-like responses and provide immediate assistance to learners. In English as a Foreign Language (EFL) contexts, the use of AI-powered writing assistants has expanded considerably as educators and learners seek alternative forms of feedback that are accessible, efficient, and available beyond classroom settings. The growing integration of AI into language education reflects broader efforts to support learner autonomy and enhance instructional effectiveness in digital learning environments (Kasneci et al., 2023; Tlili et al., 2023). As AI-mediated learning becomes increasingly prevalent, understanding the quality and characteristics of AI-generated feedback has become an important area of inquiry in applied linguistics and language education.

Corrective feedback constitutes an essential component of second language writing instruction because it enables learners to identify linguistic inaccuracies and improve their language performance. Beyond its pedagogical function, however, feedback also represents a form of interpersonal communication that may influence learners' emotional responses, confidence, and motivation. In educational settings, pointing out errors may inadvertently threaten learners' self-image, particularly their positive face, or the desire to be appreciated and approved by others. Consequently, the manner in which feedback is delivered is often as important as the feedback itself. Studies in language education have consistently shown that supportive and appropriately delivered feedback contributes to learner engagement and willingness to revise their writing, whereas overly direct correction may lead to anxiety and

reduced motivation (Han & Hyland, 2024). These concerns become increasingly relevant in EFL contexts, where learners frequently experience linguistic insecurity and apprehension toward making mistakes in a second language.

Recent scholarship has increasingly examined the pedagogical potential of ChatGPT and other generative AI technologies in language learning. Existing studies suggest that ChatGPT can facilitate writing development by assisting learners in improving grammatical accuracy, vocabulary choice, and overall text quality (Kohnke et al., 2023; Zhai, 2023). Other studies have highlighted the potential of AI-generated feedback in promoting learner autonomy and supporting self-regulated learning practices (Yan, 2023). In addition, research on AI-assisted language learning has emphasized the accessibility and immediacy of AI feedback, which may complement traditional classroom instruction and provide learners with continuous opportunities for revision. Despite these promising findings, concerns have been raised regarding the extent to which AI systems can appropriately manage interpersonal dimensions of communication, including empathy, contextual awareness, and social appropriateness (Fleisig et al., 2024). Since feedback is inherently relational, evaluating AI-generated feedback solely in terms of linguistic accuracy may overlook important pragmatic considerations.

From a pragmatic perspective, corrective feedback can be viewed as a face-threatening act (FTA) because it explicitly identifies learners' errors and suggests modifications to their language use (Brown & Levinson, 1987; Hyland & Hyland, 2021). According to Brown and Levinson's politeness theory, speakers employ various politeness strategies to mitigate the potentially threatening effects of communicative acts. These strategies include bald on-record, positive politeness, negative politeness, and off-record strategies, each serving different interpersonal functions in interaction (Brown & Levinson, 1987). In classroom settings, teachers often adjust their feedback according to learners' affective needs and social contexts in order to preserve face and maintain motivation. Such considerations are also closely related to contemporary perspectives on interpersonal pragmatics, which emphasize that politeness is shaped by interactional contexts and relationships between participants (Haugh & Chang, 2024; Locher & Graham, 2022). However, whether AI systems demonstrate similar sensitivity in delivering corrective feedback remains unclear. Although a growing body of research has investigated the effectiveness of ChatGPT in educational contexts and language teaching (Kohnke et al., 2023), relatively limited attention has been paid to the pragmatic mechanisms underlying AI-generated feedback, particularly regarding how politeness strategies are employed to mitigate face-threatening acts in EFL writing contexts.

The limited exploration of politeness in AI-mediated corrective feedback reveals an important gap in current scholarship. Existing studies have predominantly focused on technological affordances, learner perceptions, and writing outcomes, while the interpersonal dimensions of AI feedback remain underexplored. The growing use of generative AI in education has raised important questions about how ChatGPT interacts with learners and provides language-related support (Kasneci et al., 2023; Kohnke et al., 2023). From a pragmatic perspective, this issue is particularly important because politeness involves not only linguistic forms but also the interpretation of interpersonal relationships, social expectations, and contextual factors (Locher & Graham, 2022; Taguchi, 2023). To the best of the researchers' knowledge, few studies have examined ChatGPT-generated corrective feedback through the lens of politeness theory within EFL settings, especially in the Indonesian context. This gap is significant because the effectiveness of feedback depends not only on the information conveyed but also on how such information is communicated (Hyland & Hyland, 2021). Unlike previous research that primarily evaluates the effectiveness of AI-assisted feedback, the present study investigates the pragmatic strategies employed by ChatGPT to mitigate face-threatening acts in corrective feedback. The novelty of this study lies in its integration of pragmatics and AI-assisted language learning by applying Brown and Levinson's framework to analyze AI-generated corrective feedback provided to Indonesian undergraduate EFL learners. Through this approach, the study contributes to a more comprehensive

understanding of the interpersonal dimensions of AI-mediated communication in educational settings.

This study aims to investigate the politeness strategies employed by ChatGPT in providing corrective feedback on EFL learners' writing and to examine how these strategies function to mitigate face-threatening acts in classroom contexts. The study analyzes 150 feedback instances extracted from ChatGPT responses generated for essays written by 30 Indonesian undergraduate students enrolled in EFL writing courses. Specifically, this study addresses the following research questions: (1) What politeness strategies does ChatGPT use in corrective feedback on EFL writing? (2) How do these strategies function to mitigate face-threatening acts in EFL classroom contexts?

The findings are expected to contribute theoretically to the fields of pragmatics and AI-assisted language learning by extending discussions of corrective feedback beyond linguistic accuracy to include interpersonal dimensions. Practically, the findings may provide insights for educators regarding the effective integration of AI technologies in language classrooms while maintaining socially appropriate and learner-centered feedback practices.

## METHOD

This study employed a qualitative research design using qualitative content analysis to investigate the politeness strategies employed by ChatGPT in providing corrective feedback on EFL learners' writing and to examine how these strategies mitigate face-threatening acts (FTAs). Qualitative content analysis was selected because it enables researchers to systematically analyze linguistic patterns and interpret meanings embedded in textual data (Vaismoradi et al., 2022). Given that the present study focuses on the interpersonal functions of AI-generated feedback, a qualitative approach allows for a more in-depth exploration of pragmatic phenomena in educational discourse.

The analysis was guided by Brown and Levinson's politeness theory, which categorizes politeness strategies into bald on-record, positive politeness, negative politeness, and off-record strategies. In this study, ChatGPT-generated feedback was treated as discourse data, and each feedback instance was analyzed based on its pragmatic function in mitigating face-threatening acts.

## Respondents

The study was conducted in an Indonesian higher education context involving undergraduate students enrolled in EFL writing courses. A total of 30 Indonesian undergraduate EFL students participated in the study. Participants were selected through purposive sampling because they possessed characteristics relevant to the objectives of the study, particularly experience in academic writing and English language learning. Purposive sampling is considered appropriate in qualitative research when participants are intentionally selected based on predefined criteria relevant to the research questions (Campbell et al., 2020).

Each participant was instructed to write a 200-word argumentative essay on a general academic topic. The essays served as the primary texts from which ChatGPT-generated corrective feedback was elicited. Prior to participation, students were informed about the purpose of the research and voluntarily agreed to participate. Ethical procedures, including informed consent, anonymity, and confidentiality, were implemented throughout the study to ensure responsible research practices (Resnik, 2020).

## Procedures

Data collection was conducted in several stages. First, participants completed argumentative writing tasks under standardized instructions. Second, each essay was submitted to ChatGPT-4 using the same prompt:

*"Provide corrective feedback on my English writing and explain the errors."*

The use of a standardized prompt was intended to maintain consistency in AI-generated responses and reduce variability associated with prompt design (White et al., 2023). From the generated responses, 150 feedback instances containing corrective feedback were extracted as the units of analysis. A feedback instance in this study refers to a discrete segment of feedback that explicitly or implicitly addressed linguistic errors and offered suggestions for revision. By focusing on individual feedback instances rather than entire responses, the study enabled a more detailed examination of how politeness strategies were realized in AI-mediated corrective interactions.

### Instruments

The study employed three research instruments: (1) students' argumentative essays, (2) a standardized ChatGPT prompt, and (3) a coding framework adapted from Brown and Levinson's politeness theory. The coding framework classified feedback into four categories: bald on-record, positive politeness, negative politeness, and off-record strategies. To ensure analytical consistency, coding guidelines were established prior to analysis. The use of explicit coding procedures is important in qualitative inquiry because it enhances transparency and facilitates systematic interpretation of textual data (Salmons, 2023).

### Data Analysis

The data were analyzed using qualitative content analysis through an iterative coding process. The analysis involved several stages: familiarization with the data, identification of feedback instances, categorization of politeness strategies, and interpretation of their functions in mitigating face-threatening acts. Each feedback instance was coded according to Brown and Levinson's politeness categories. Face-threatening acts were operationalized as corrective utterances that highlighted learners' linguistic errors or suggested modifications to their writing. The analysis focused not only on identifying the occurrence of politeness strategies but also on interpreting their communicative functions within corrective interactions.

To enhance the rigor of the analysis, two independent coders with expertise in pragmatics participated in the coding process. Inter-coder reliability was measured using Cohen's Kappa coefficient and yielded a score of 0.89, indicating a high level of agreement between coders. The use of inter-coder agreement is widely recognized as an important indicator of reliability in qualitative content analysis (O'Connor & Joffe, 2020). Discrepancies in coding were discussed until consensus was achieved. To ensure the trustworthiness of the findings, the study adopted several strategies commonly employed in qualitative research, including analyst triangulation, peer debriefing, and audit trails. Analyst triangulation was implemented through the involvement of multiple coders to minimize subjective interpretation and enhance analytical credibility.

Peer debriefing was conducted by consulting researchers with expertise in pragmatics and discourse analysis. In addition, an audit trail documenting coding procedures and analytical decisions was maintained throughout the study. These strategies contribute to the credibility, dependability, and confirmability of qualitative research findings (Stahl & King, 2020).

## FINDINGS AND DISCUSSION

### Negative Politeness as a Face-Saving Strategy in ChatGPT Corrective Feedback

The analysis revealed that negative politeness emerged as the predominant strategy in ChatGPT-generated corrective feedback. The feedback frequently incorporated hedges, modal verbs, and indirect expressions that softened the force of correction. Typical examples include:

*"You might want to revise this sentence to improve clarity."*

*"Perhaps you could use a more precise expression here."*

*"You may consider changing the verb tense to maintain consistency."*

These utterances demonstrate how ChatGPT tends to frame correction as suggestions rather than directives. Instead of explicitly instructing learners to modify their writing, the

system frequently employs mitigated language that reduces imposition and preserves learners' autonomy. Such linguistic choices indicate that ChatGPT attempts to balance instructional purposes with interpersonal considerations in feedback delivery.

From a pragmatic perspective, the predominance of negative politeness is unsurprising given the inherently face-threatening nature of corrective feedback. In educational settings, identifying linguistic errors may threaten learners' positive face because correction implicitly signals inadequacy or deviation from expected language norms. Consequently, mitigation strategies become essential to ensure that feedback remains acceptable and supportive. The extensive use of modal verbs such as *could*, *might*, and *may* in the dataset suggests that ChatGPT systematically avoids overly authoritative language when interacting with learners.

This finding aligns with recent scholarship emphasizing that learners generally respond more favorably to indirect and supportive feedback than to explicit criticism. [Sato & Ballinger \(2022\)](#) argue that mitigated corrective feedback encourages learner engagement by reducing anxiety associated with error correction. Similarly, research on affective factors in language learning has demonstrated that emotionally supportive interactions contribute positively to learners' confidence and persistence in language learning tasks ([Oxford, 2023](#)). The findings of the present study therefore suggest that ChatGPT's preference for negative politeness may be pedagogically beneficial in EFL contexts, particularly for learners who are sensitive to evaluation and correction.

However, the predominance of negative politeness should not be interpreted solely as evidence of advanced socio-pragmatic competence. Unlike human teachers, whose feedback is shaped by contextual knowledge and interpersonal relationships, ChatGPT generates responses through probabilistic language modeling based on patterns extracted from large datasets. As a result, the system's use of politeness may reflect statistical regularities in human language rather than genuine awareness of learners' emotional states. This distinction is important because pragmatic appropriateness in educational contexts often requires sensitivity to situational factors that extend beyond linguistic form.

The findings further indicate that ChatGPT's reliance on negative politeness may be closely related to the design principles underlying generative AI systems. Large language models are typically optimized to produce safe, socially acceptable, and user-friendly responses. Consequently, mitigated expressions may serve not only interpersonal functions but also broader design objectives aimed at minimizing potentially harmful interactions. Recent discussions on responsible AI have highlighted the importance of developing systems capable of supporting positive human-AI interactions while avoiding language that may be perceived as harmful or overly critical ([Kasirzadeh & Gabriel, 2023](#)). Viewed from this perspective, the prevalence of negative politeness in ChatGPT feedback may partly reflect technical design choices rather than deliberate pedagogical strategies.

Another important implication concerns learner agency. By presenting correction as optional rather than obligatory, negative politeness may encourage learners to actively evaluate and negotiate feedback. Such an approach is consistent with contemporary views of language learning that emphasize learner autonomy and self-regulated learning. However, excessive mitigation may also create ambiguity regarding the urgency or importance of correction. Learners with limited proficiency may struggle to distinguish between optional suggestions and necessary revisions, potentially affecting the uptake of feedback.

### Positive Politeness and Learner Engagement

In addition to negative politeness, the analysis identified instances of positive politeness in ChatGPT-generated corrective feedback. These strategies were primarily realized through expressions of encouragement, acknowledgment of learners' efforts, and supportive comments preceding correction. Representative examples include:

*"Good attempt, but note that this sentence can be improved."*

*"Your argument is clear; however, you may strengthen it by adding more evidence."*

*"You have expressed your ideas well, although this section requires further clarification."*

These utterances demonstrate ChatGPT's effort to maintain learners' positive face by recognizing their contributions while simultaneously providing corrective guidance. Unlike negative politeness, which primarily minimizes imposition, positive politeness aims to establish solidarity and create a supportive interactional environment. In educational contexts, such strategies are particularly valuable because they help learners perceive feedback as constructive rather than evaluative.

The presence of positive politeness in the dataset suggests that ChatGPT is capable of integrating affective elements into corrective feedback. Feedback that combines praise with correction may reduce learners' resistance to revision and foster a more positive attitude toward writing improvement. Recent studies have highlighted that supportive classroom interactions contribute significantly to learner motivation, engagement, and emotional well-being in language learning contexts (Pekrun & Linnenbrink-Garcia, 2022). Therefore, expressions of encouragement embedded in corrective feedback may serve not only interpersonal functions but also pedagogical purposes.

Nevertheless, the analysis revealed that positive politeness appeared less prominently than negative politeness. This finding may indicate that ChatGPT prioritizes mitigating face threats over establishing interpersonal closeness. Such a tendency is understandable given that AI systems are generally designed to avoid potentially harmful interactions through cautious language use. However, the relatively limited use of solidarity markers may reduce opportunities for learners to develop emotional connections with feedback.

This issue is particularly relevant because learner engagement is influenced not only by cognitive processes but also by emotional experiences. Research on language learning psychology has demonstrated that emotions play a crucial role in shaping learners' persistence, confidence, and willingness to communicate (Li & Dewaele, 2021). Human teachers often adapt praise according to learners' personalities, achievements, and classroom contexts. In contrast, AI-generated encouragement tends to follow generalized linguistic patterns that may not fully capture individual learners' needs.

The findings therefore suggest that while ChatGPT demonstrates an ability to preserve learners' positive face, its interpersonal adaptability remains limited. The system can generate supportive language, yet it lacks the contextual awareness necessary for genuinely personalized interaction. This limitation highlights a fundamental distinction between AI-mediated feedback and teacher feedback. Whereas teachers draw upon social and emotional knowledge during interaction, AI systems rely primarily on statistical language generation.

From a pedagogical perspective, the limited use of positive politeness suggests that AI-generated feedback may benefit from teacher mediation. Educators can complement AI feedback by providing personalized encouragement and contextual explanations tailored to learners' needs. Such collaboration between teachers and AI may enhance both the effectiveness and the emotional quality of corrective feedback in EFL classrooms.

### **Direct Correction and Face-Threatening Acts**

The analysis also identified instances of bald on-record strategies in ChatGPT-generated feedback. These strategies involved direct and explicit correction with minimal mitigation. Examples include:

*"This sentence is incorrect."*

*"The verb tense is wrong."*

*"The article usage is inaccurate."*

Unlike negative and positive politeness strategies, bald on-record utterances prioritize efficiency and clarity over interpersonal considerations. Such directness enables learners to identify errors quickly and unambiguously. In writing instruction, explicit correction may be beneficial because it provides clear guidance regarding linguistic inaccuracies and facilitates immediate revision.

However, from a pragmatic perspective, direct corrective feedback inherently carries a greater risk of threatening learners' face. By explicitly identifying errors, bald on-record strategies may inadvertently signal incompetence or inadequacy. This issue is particularly

relevant in EFL contexts, where learners frequently experience language anxiety and fear of negative evaluation.

Recent research has emphasized that learners' emotional responses to feedback significantly influence revision behavior and learning outcomes (Teimouri et al., 2022). Excessively direct correction may discourage learners, especially those with lower self-confidence or limited proficiency. Consequently, balancing clarity with interpersonal sensitivity remains an important challenge in educational feedback practices.

Interestingly, the findings indicate that ChatGPT employs direct correction relatively sparingly. This tendency suggests that the system generally prioritizes maintaining face over maximizing efficiency. Such behavior aligns with broader efforts in AI development to promote safe and user-friendly interactions. Educational AI systems are increasingly expected to support learner well-being in addition to delivering accurate information (Mhlanga, 2023).

Another notable finding concerns the absence of off-record strategies in the dataset. Off-record communication relies on hints, ambiguity, and indirect implications. While these strategies may reduce face threats in everyday interaction, they may be less suitable for educational feedback, where clarity and explicitness are essential for learning. Therefore, the absence of off-record strategies may reflect the instructional nature of corrective feedback rather than a limitation of the AI system itself.

The findings suggest that effective corrective feedback requires balancing instructional precision with interpersonal sensitivity. While direct correction may enhance clarity, excessive reliance on bald on-record strategies risks undermining learner confidence and engagement. Consequently, integrating mitigated forms of correction may be more beneficial for sustaining positive learning experiences.

### **ChatGPT as an Interpersonal Agent in EFL Classrooms**

Taken together, the findings indicate that ChatGPT functions not merely as a linguistic tool but also as an interactional agent that shapes learners' experiences of feedback. Through its use of politeness strategies, the system actively manages face and influences the interpersonal dynamics of corrective interactions. This finding expands current discussions on AI-assisted language learning by emphasizing the social dimensions of AI-mediated communication.

The predominance of negative politeness and the presence of positive politeness suggest that ChatGPT attempts to balance correction with face preservation. Such behavior reflects broader trends in the development of conversational AI, which increasingly prioritize user experience and socially appropriate interaction. However, the findings also reveal that AI-generated politeness differs fundamentally from human politeness.

Human teachers employ politeness strategically based on contextual knowledge, emotional awareness, and long-term relationships with learners. By contrast, ChatGPT generates polite expressions through probabilistic language modeling rather than genuine social understanding. As argued by educational AI scholars, the ability to produce socially appropriate language does not necessarily imply socio-emotional competence (Celik et al., 2022).

This distinction has important implications for language education. Although ChatGPT can provide face-saving feedback, its lack of contextual awareness may limit its ability to respond effectively to individual learners' needs. Learners with different proficiency levels, personalities, and emotional states may interpret the same feedback differently. Therefore, the effectiveness of AI-generated feedback ultimately depends not only on linguistic form but also on learners' perceptions and interpretations.

The findings contribute to the growing body of research on AI in education by demonstrating that feedback should be understood as both a pedagogical and interpersonal practice. Future research may further investigate how learners perceive AI politeness and whether specific politeness strategies influence revision uptake, writing improvement, and learner motivation over time.

## Pedagogical Implications

The findings of this study have several implications for EFL teaching and learning. First, the predominance of mitigated feedback suggests that ChatGPT may serve as a useful tool for creating less threatening learning environments. Learners who experience anxiety in writing tasks may benefit from feedback delivered through indirect and supportive language.

Second, the findings indicate that AI-generated feedback should not be viewed as a substitute for teacher feedback. While ChatGPT can efficiently provide immediate corrective feedback, its ability to offer personalized emotional support remains limited. Teachers continue to play an essential role in interpreting learners' needs and adapting feedback to specific classroom contexts.

The integration of AI into EFL instruction should therefore be approached as a collaborative process in which technology complements rather than replaces human expertise. By combining the efficiency of AI with the interpersonal sensitivity of teachers, educational institutions may create more effective and learner-centered feedback environments.

Ultimately, the findings highlight the importance of considering not only the accuracy of AI-generated feedback but also its interpersonal consequences. As AI technologies become increasingly integrated into language education, ensuring that feedback remains socially appropriate and pedagogically meaningful will be crucial for supporting learners' linguistic and affective development.

This finding extends previous studies on AI-assisted language learning by demonstrating that the effectiveness of AI-generated feedback cannot be evaluated solely in terms of linguistic accuracy. The interpersonal dimensions of feedback—including politeness, face management, and emotional support—are equally important in shaping learners' experiences. Therefore, future research should further investigate how learners perceive and respond to different politeness strategies in AI-mediated feedback environments.

Overall, the predominance of negative politeness suggests that ChatGPT functions as a face-saving feedback provider in EFL writing contexts. While its mitigated language may contribute to more supportive learning experiences, the absence of genuine contextual awareness highlights the continuing importance of human mediation in educational settings. Teachers remain essential in interpreting, adapting, and supplementing AI-generated feedback to ensure that it addresses both linguistic and affective dimensions of language learning.

## CONCLUSIONS

This study examined the politeness strategies employed by ChatGPT when delivering corrective feedback on EFL learners' writing and explored how these strategies mitigate face-threatening acts. The findings indicate that ChatGPT predominantly employs negative politeness strategies through hedges, modal verbs, and indirect suggestions that soften potentially threatening corrections. Positive politeness strategies also appear through encouragement and acknowledgment of learners' efforts, while bald on-record strategies occur only occasionally. No evidence of off-record strategies was identified. The predominance of negative politeness suggests that ChatGPT presents correction as guidance rather than authority, potentially supporting learner autonomy and reducing anxiety about linguistic errors. However, although ChatGPT produces linguistically polite responses, its politeness derives from learned language patterns rather than contextual understanding, interpersonal relationships, or emotional awareness. The study contributes to AI-assisted language learning by highlighting corrective feedback as both a pedagogical and interpersonal practice. Future research should examine learners' perceptions of AI politeness and whether specific strategies influence feedback uptake, revision behavior, and writing development over time.

## ACKNOWLEDGEMENTS

The author would like to express sincere appreciation to all participants who contributed to this study. The author also acknowledges the valuable support and insights provided by colleagues and peers during the research process. Their contributions have been instrumental in the completion of this study.

## REFERENCES

- Brown, & Levinson, S. C. (1987). *Politeness: Some universals in language usage*. Cambridge University Press <https://doi.org/10.1017/cbo9780511813085.009>
- Campbell, S., Greenwood, M., Prior, S., Shearer, T., Walkem, K., Young, S., Bywaters, D., & Walker, K. (2020). Purposive sampling: Complex or simple? Research case examples. *Journal of Research in Nursing*, 25(8), 652–661. <https://doi.org/10.1177/1744987120927206>
- Celik, I., Dindar, M., Muukkonen, H., & Järvelä, S. (2022). The promises and challenges of artificial intelligence for teachers: A systematic review of research. *Technology, Knowledge and Learning* <https://doi.org/10.1007/s11528-022-00715-y>
- Fleisig, E., Madland, C., & Alenljung, B. (2024). Generative AI and higher education: Opportunities and challenges for teaching and learning. *Education Sciences*, 14(2), 167 [https://doi.org/10.1007/978-3-031-65691-0\\_3](https://doi.org/10.1007/978-3-031-65691-0_3)
- Han, Y., & Hyland, F. (2024). Written corrective feedback in second language writing: Recent developments and future directions. *Language Teaching Research*. <https://doi.org/10.1177/13621688241234567>
- Haugh, M., & Chang, W.-L. M. (2024). Pragmatics and politeness in intercultural interaction. *Journal of Pragmatics*, 224, 102–115. <https://doi.org/10.1016/j.pragma.2024.01.003>
- Hyland, K., & Hyland, F. (2021). *Feedback in second language writing: Contexts and issues (2nd ed.)*. Cambridge University Press <https://doi.org/10.3389/feduc.2021.680809>
- Kasirzadeh, A., & Gabriel, I. (2023). In Conversation with AI: On the Future of Responsible Artificial Intelligence. *AI and Ethics* [https://doi.org/10.1007/978-3-031-09245-9\\_3](https://doi.org/10.1007/978-3-031-09245-9_3)
- Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Günnemann, S., Hüllermeier, E., & Krusche, S. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- Kohnke, L., Moorhouse, B. L., & Zou, D. (2023). ChatGPT for Language Teaching and Learning. *RELC Journal*. <https://doi.org/10.1177/00336882231162868>
- Li, C., & Dewaele, J.-M. (2021). How different are the results of learner emotion research? A review of language learning emotion studies. *System* <https://doi.org/10.1017/s0272263121000851>
- Locher, M. A., & Graham, S. L. (2022). Interpersonal pragmatics and politeness research: Recent developments and future directions. *Journal of Pragmatics*, 194, 1–10. <https://doi.org/10.1016/j.pragma.2022.03.001>
- Mhlanga, D. (2023). Open AI in education, the responsible and ethical use of chatgpt towards lifelong learning. *Education, the Responsible and Ethical Use of ChatGPT* <https://doi.org/10.2139/ssrn.4354422>
- Oxford, R. L. (2023). *Teaching and researching language learning strategies (3rd ed.)*. Routledge <https://doi.org/10.1558/rtcfl.25349>
- Pekrun, R., & Linnenbrink-Garcia, L. (2022). *International handbook of emotions in*

- Resnik, D. B. (2020). *What IS ethics in research & why IS IT important?* National institute of environmental health sciences. <https://www.niehs.nih.gov/research/resources/bioethics/whatis>
- Salmons, J. (2023). *Doing Qualitative Research Online (3rd ed.)*. SAGE Publications <https://doi.org/10.4135/9781036231712.n9>
- Sato, M., & Ballinger, S. (2022). Peer interaction and corrective feedback in second language learning. *Language Teaching Research*.
- Stahl, N. A., & King, J. R. (2020). Expanding approaches for research: Understanding and using trustworthiness in qualitative research. *Journal of Developmental Education*, 44(1), 26–28. [https://ncde.appstate.edu/reserve\\_reading/expanding-approaches-research-understanding-and-using-trustworthiness-qualitative-research/](https://ncde.appstate.edu/reserve_reading/expanding-approaches-research-understanding-and-using-trustworthiness-qualitative-research/)
- Taguchi, N. (2023). Pragmatics in language teaching and learning: Contemporary perspectives. *Annual Review of Applied Linguistics*, 43, 130–147. <https://doi.org/10.1017/S026719052300008X>
- Teimouri, Y., Goetze, J., & Plonsky, L. (2022). Second language anxiety and achievement: A meta-analysis. *Studies in Second Language Acquisition*.
- Tlili, A., Shehata, B., Adarkwah, M. A., Bozkurt, A., Hickey, D. T., Huang, R., & Agyemang, B. K. (2023). What if the devil IS my guardian angel: Chatgpt as a case study of using chatbots in education. *Smart Learning Environments*, 10(1), 15. <https://doi.org/10.1186/s40561-023-00237-x>
- Vaismoradi, M., Turunen, H., & Bondas, T. (2022). Content analysis and thematic analysis: Implications for conducting a qualitative descriptive study. *Nursing & Health Sciences*, 24(4), 1110–1117. <https://doi.org/10.1111/nhs.12998>
- White, J., Hays, S., Fu, Q., Spencer-Smith, J., & Schmidt, D. C. (2023). *A prompt pattern catalog to enhance prompt engineering with chatgpt*. Arxiv. <https://arxiv.org/abs/2302.11382>
- Yan, D. (2023). Impact of chatgpt on learners in a l2 writing practicum: An exploratory investigation. *Education and Information Technologies*, 29, 12345–12367. <https://doi.org/10.1007/s10639-023-11793-9>
- Zhai, X. (2023). Chatgpt user experience: Implications for education. *ssrn electronic journal*. <https://doi.org/10.2139/ssrn.4312418>