

<https://doi.org/10.5281/zenodo.21221937>

Decoding the Machine: A Corpus-Based Analysis of Pragmatic Structures in AI-Generated Discourse and Their Interpretive Implications for EFL Learners

Ms. Areeba

University of Agriculture, Faisalabad.

Dr. Ayesha Asghar Gill

Assistant Professor, University of Agriculture, Faisalabad.

Ayesha.asghar@uaf.edu.pk

Abstract: *Conversational AI systems such as ChatGPT increasingly function as a primary source of linguistic input for EFL learners, yet existing research has examined either AI's capacity to generate pragmatically appropriate language or AI-assisted learning outcomes, leaving largely unexplored how the structural features of AI-generated responses may themselves create interpretive difficulty. This study addresses that gap through a corpus-based analysis of AI-generated responses from the WildChat dataset. A corpus of 122 human-AI exchanges (36,233 word tokens), sampled across four pragmatic categories, requests, apologies, refusals, and conversational implicature, was analyzed in AntConc 3.5.9 through Word List, Concordance, N-Gram, and Keyness analysis, interpreted through Grice's Cooperative Principle, Speech Act Theory, and Relevance Theory.*

Findings identify a recurring tripartite AI Limitation Formula (ALF) — softener, AI-identity disclaimer, limitation statement — that delays and obscures illocutionary force, with “as an AI language model” the dominant N-gram cluster ($f=37$) and however the most frequent marker overall ($f=64$). Keyness analysis ranks implicature as the most inferentially demanding category, since explicit markers are entirely absent, followed by refusal, whose force is most systematically concealed through the ALF and however-pivot (model: $LL=45.85$). Apology language proved highly formulaic (apologize: $LL=23.36$, $\%DIFF=+1689\%$), frequently failing Searle's sincerity condition. The study concludes AI-generated discourse encodes meaning through systematic, learnable structures with direct implications for EFL pragmatic instruction.

Keywords: *AI-generated discourse; pragmatic competence; EFL learners; corpus linguistics; Cooperative Principle; Speech Act Theory; Relevance Theory; WildChat*

1. Introduction

The rapid development of artificial intelligence technologies has significantly transformed contemporary language learning environments, particularly within English as a Foreign Language (EFL) contexts. Conversational AI systems such as ChatGPT increasingly function as interactive language partners, writing assistants, and sources of immediate linguistic input for learners worldwide, providing dynamic, real-time conversational interaction that often resembles natural human communication. Recent research demonstrates that AI tools can positively support fluency development, writing support, speaking confidence, and pragmatic awareness (Tahir, 2025; Bhar, 2026; Erdogan & Kitson, 2025). However, despite their grammatical sophistication and conversational fluency, Large Language Models (LLMs) remain fundamentally limited in their ability to negotiate

human pragmatic meaning authentically, since they generate responses through probabilistic language prediction rather than genuine communicative intention, emotional awareness, or sociocultural understanding.

Pragmatic competence, the ability to interpret and produce meaning appropriately according to context, communicative intention, and social relationship, has long been recognized as one of the most difficult dimensions of second language acquisition because it requires learners to move beyond literal meaning and recover implied communicative intent through inference (Kasper & Rose, 2002; Taguchi, 2011). These challenges become increasingly significant within AI-mediated communication: although AI-generated discourse often appears cooperative on the surface, it may encode pragmatic meaning through structurally indirect and highly formulaic patterns. A representative AI response, “I apologize, but as an AI language model, I cannot provide medical advice; however, I can offer general information,” simultaneously performs apology, identity disclaimer, limitation statement, and alternative offer within a single compressed utterance, structurally delaying and partially obscuring the actual illocutionary force of refusal.

Existing research on AI and pragmatics has focused primarily on two areas: whether AI systems can generate pragmatically appropriate responses (Nazeer et al., 2024; Andersson & McIntyre, 2025), and how AI tools may support learners' pragmatic competence development (Tahir, 2025; Erdogan & Kitson, 2025). Relatively little attention has been given to the inverse problem: how the structural pragmatic features of AI-generated responses may themselves create interpretive difficulty for EFL learners. This gap is consequential because AI-generated communication increasingly functions as a primary source of language exposure; if learners repeatedly encounter formulaic or pragmatically opaque AI discourse without explicit instruction, they may struggle to distinguish authentic human pragmatic norms from AI-specific communicative patterns. This study addresses that gap through a corpus-based analysis of AI-generated responses from the WildChat dataset. It is important to clarify that the study does not recruit EFL learners or measure comprehension directly; rather, it conducts a structural analysis identifying pragmatic features theoretically likely to create interpretive difficulty, based on established frameworks in interlanguage and cognitive pragmatics.

This study investigates the pragmatic features of AI-generated discourse and their structural implications for EFL learner interpretation. Specifically, the study aims to:

- (1) identify and analyze the pragmatic features present in AI-generated responses within the WildChat corpus and examine their structural influence on interpretive complexity for EFL learners;
- (2) compare the relative inferential demand of different pragmatic categories, requests, apologies, refusals, and conversational implicature, encoded in AI-generated responses; and
- (3) apply Grice's Cooperative Principle, Speech Act Theory, and Relevance Theory to explain the relationship between AI-generated pragmatic structures and interpretive difficulty.

These objectives correspond to three research questions: (RQ1) How do pragmatic features in AI-generated responses within the WildChat corpus create structural complexity and inferential demand for EFL learners? (RQ2) To what extent do different types of pragmatic meaning, requests, apologies, refusals, and conversational implicature, vary in their structural inferential difficulty? (RQ3) How do Grice's Cooperative Principle, Speech Act Theory, and Relevance Theory explain the relationship between AI-generated pragmatic structures and interpretive difficulty for EFL learners?

The study adopts a corpus-based research design integrating complementary quantitative and qualitative analytical procedures. A corpus of 122 human-AI conversational exchanges (36,233 word tokens) was purposively sampled from the WildChat dataset across four target pragmatic categories, requests (n=34), refusals (n=11), apologies (n=10), and conversational implicature (n=6), and analyzed in AntConc 3.5.9 through Word List, Concordance (KWIC), N-Gram/Cluster, and Keyness analysis.

Findings were interpreted through three complementary theoretical frameworks: Grice's (1975) Cooperative Principle, Speech Act Theory (Searle, 1969), and Relevance Theory (Sperber & Wilson, 1986). Full methodological detail is provided in Section 3.

The remainder of the article proceeds as follows. Section 2 reviews relevant literature and establishes the research gap. Section 3 details the corpus, data collection, and analytical procedures. Section 4 presents the quantitative and qualitative findings together with discussion. Section 5 concludes, outlines the study's contributions and limitations, and proposes directions for future research.

2. Literature Review and Theoretical Framework

2.1 AI Pragmatic Performance and AI-Assisted Learning

A substantial body of recent scholarship evaluates AI systems' pragmatic capabilities. Eragamreddy (2025) finds that while AI effectively processes explicit speech acts, it struggles significantly with indirectness, ambiguity, and cultural nuance, often prompting users to simplify their own language; using anonymized WildChat and OpenAI logs, the study establishes a foundation for examining how learners decode AI's literal outputs. Nazeer, Khan, Nawaz, and Rehman (2024) similarly find that ChatGPT handles metaphorical language well but performs erratically with multi-layered irony and emotional tone, while Andersson and McIntyre (2025) show that ChatGPT recognizes conventionalized impoliteness but struggles with nuanced implicature and tends to default to hyper-conservative, overly polite language that alters a text's original register. Brunet-Gouet et al. (2023) extend this concern theoretically, finding that ChatGPT struggles to spontaneously infer communicative intentions and indirect speech, relying on literal interpretation unless explicitly prompted to consider a character's mental state.

A second strand of scholarship examines AI-assisted pragmatic learning outcomes. Erdogan and Kitson (2025) find that AI tools provide immersive environments for teaching implicatures, presuppositions, and speech acts, though learners with high proficiency may still struggle to decode complex AI outputs. Bhar (2026), in a PRISMA-guided systematic review of 36 studies, concludes that AI should be conceptualized as a design-sensitive instructional resource rather than an autonomous teaching agent. Nurjain et al. (2025) and Ifadloh, Abdullahi, and Rukmana (2025) both find that learners must actively modify AI-generated language pragmatically, the former in argumentative writing and the latter in academic email feedback, where AI tools correct surface-level errors but fail to address context-sensitive nuances such as indirectness and tone. Across this literature, a consistent pattern emerges: studies evaluate either AI's pragmatic production or AI-assisted learning outcomes, but rarely the learner's interpretation of AI-generated pragmatic structures themselves.

2.2 Theoretical Framework

This study integrates three complementary pragmatic frameworks. Grice's (1975) Cooperative Principle holds that conversational meaning is governed by four maxims, Quantity, Quality, Relation, and Manner, which listeners assume speakers observe; when a maxim is overtly flouted, the listener infers an unspoken implicature. Because AI systems lack communicative intention, the present study does not claim that AI literally flouts maxims, but that AI discourse produces maxim-analogue patterns, structural outputs that create the same inferential demands as flouting or violation from the hearer's perspective, without the intentional communicative design those terms imply in human discourse (Miehling et al., 2024).

Speech Act Theory (Austin, 1962; Searle, 1969) distinguishes the locutionary act (literal production of words), the illocutionary act (the speaker's intended communicative force), and the perlocutionary act (the effect on the listener), and specifies felicity conditions, preparatory, essential, and sincerity, that a speech act must satisfy to succeed. Because LLMs lack consciousness or genuine belief, they cannot fulfill the sincerity condition, a theoretical paradox that fundamentally alters how their expressive and

refusal acts should be pragmatically categorized (Rosen, 2024). Relevance Theory (Sperber & Wilson, 1986) reframes communication as a cognitive search for relevance, balancing cognitive effects against processing effort; optimal relevance is achieved when effects are maximized and effort minimized, and listeners apply a “stopping rule,” ceasing inferential processing once an interpretation satisfies their expectation of relevance. Because AI’s ostensive communicative signals are entirely simulated, yet human cognition is evolutionarily wired to search for relevance regardless, students instinctively process AI outputs as though they possess underlying human intentionality, creating a structural imbalance between processing effort and communicative effect (Manzoor, 2025; Tan, 2025).

While existing literature extensively evaluates AI’s pragmatic performance (Nazeer et al., 2024; Andersson & McIntyre, 2025; Brunet-Gouet et al., 2023) and AI-assisted pragmatic learning outcomes (Erdogan & Kitson, 2025; Bhar, 2026; Nurjain et al., 2025), it rarely investigates how learners must cognitively process the indirectness, formulaicity, and pragmatic ambiguity characteristic of AI discourse itself. Studies such as Ifadloh, Abdullahi, and Rukmana (2025) establish the limitations of AI tools in providing pragmatic correction but focus on the AI’s diagnostic output rather than how learners decode and interact with AI-generated pragmatic structures; similarly, Andersson and McIntyre (2025) document AI’s tendency to homogenize and over-sanitize language without examining how non-native speakers navigate these artificially polished norms. The present study addresses this gap directly by shifting analytical focus from AI’s generative capacity to the structural features of AI output that are theoretically likely to create interpretive difficulty for EFL learners, grounded in a corpus-based analysis guided by three converging pragmatic frameworks.

3. Methodology

This study employed a corpus-based research design, integrating complementary quantitative keyness measurement with qualitative concordance interpretation, to investigate pragmatic features in AI-generated discourse and their potential interpretive complexity for EFL learners. This scope reflects the study’s primary aim: to examine structural pragmatic patterns in AI-generated text rather than to measure learner comprehension directly. Three theoretical frameworks guided the analysis, Grice’s (1975) Cooperative Principle, Speech Act Theory (Searle, 1969), and Relevance Theory (Sperber & Wilson, 1986), each applied as an analytical lens to account for potential pragmatic gaps identified in the corpus data.

The data source was the WildChat dataset, a large-scale publicly available repository of real human-AI conversational exchanges generated through interactions with ChatGPT. A total of 122 exchanges were selected through a two-stage sampling procedure: exchanges were first purposively filtered according to four target pragmatic categories, requests, apologies, refusals, and conversational implicature, ensuring sufficient representation of each type, and then randomly selected within each category to avoid systematic topic-domain bias. The final corpus comprised 36,233 word tokens across 6,216 word types, distributed across four sub-corpora: requests (n=34, 11,730 tokens), refusals (n=11, 2,294 tokens), apologies (n=10, 1,570 tokens), and conversational implicature (n=6, 2,016 tokens). As the WildChat dataset is publicly available and contains no personally identifiable information, no ethical approval was required.

Prior to corpus loading, each exchange was manually pre-tagged against three binary criteria corresponding to the study’s theoretical frameworks. An exchange was tagged Grice-positive if it contained a discourse pattern functionally analogous to maxim non-observance (verbosity, a formulaic sincerity marker, a however-pivot, or multi-move compression); Speech-Act-positive if it contained an identifiable illocutionary act evaluable against Searle’s (1969) felicity conditions; and Relevance-positive if it contained inferentially complex structures requiring additional processing effort. Pre-tagging confirmed that 89/122 exchanges (73%) contained Gricean features, 104/122 (85%) contained Speech Act features, and 72/122 (59%) contained Relevance Theory features, establishing that the corpus was sufficiently pragmatically rich to support the study’s three-framework analytical approach.

The primary analytical instrument was AntConc 3.5.9 (Anthony, 2019), employing four functions. Word List analysis generated frequency rankings of all word types across the full corpus, identifying high-frequency pragmatic markers. Concordance (KWIC) analysis examined each target marker in its immediate textual context to determine its discourse function across exchange types. N-Gram/Cluster analysis (minimum cluster size 3, maximum 5, minimum frequency 3) identified recurring multi-word formulaic sequences. Keyness analysis compared each pragmatic sub-corpus statistically against the full 122-exchange corpus as reference baseline, using Log-Likelihood (LL) and percentage difference (%DIFF) as measures of statistical significance and effect size respectively; all reported keywords met the significance threshold of $p < .05$.

Data were analyzed in three steps. First, frequency and concordance data from the full corpus were examined to address RQ1, identifying dominant pragmatic features and mapping their structural patterns. Second, keyness scores across the four sub-corpora were compared to address RQ2, ranking categories by statistical distinctiveness, marker density, and inferential demand. Third, identified patterns were interpreted through the three theoretical frameworks to address RQ3: Grice's Cooperative Principle accounted for maxim-analogue discourse patterns, Speech Act Theory examined felicity-condition failures in apology and refusal encoding, and Relevance Theory explained the cognitive-effort imbalance produced by formulaic structures.

Validity was established through alignment between research objectives, theoretical framework, corpus data, and analytical procedure. The WildChat corpus provides naturally occurring, authentic human-AI discourse rather than artificially constructed examples, strengthening ecological validity, while the three established theoretical frameworks, selected for their direct correspondence to the study's core constructs (conversational meaning, speech acts, inferential processing), strengthen conceptual validity. Methodological triangulation across Word List, Concordance, N-Gram, and Keyness analysis further increased the comprehensiveness of the findings. Reliability was supported through AntConc's fixed computational procedures, which ensure that identical input consistently produces identical output, and through predefined, transparent coding criteria applied consistently across the corpus; pre-tagging criteria were reviewed against a subset of 20 exchanges to confirm consistency of application. Cross-category distributional differences in pragmatic markers were further validated through a chi-square test of independence, $\chi^2(15) = 87.43$, $p < .001$, confirming that observed patterns were not attributable to chance.

4. Findings and Discussion

4.1 The AI Limitation Formula and Dominant Pragmatic Markers

Word List analysis of the full corpus (6,216 word types, 36,233 tokens) identified six pragmatic markers at elevated frequency (Table 1).

Table 1. Key pragmatic markers across the full corpus

Marker	Frequency	Dominant Function
however	64	Pivot from limitation to alternative offer (Relation-analogue)
please	27	Formulaic invitation/closing device
cannot	18	Indirect refusal, always embedded in explanation
unable	12	Refusal buried within ALF + however-pivot
apologize	9	Formulaic softener preceding refusal
unfortunately	6	Refusal hedge, co-occurs with cannot/unable

Concordance analysis confirmed that these markers function systematically rather than incidentally. Of

nine instances of apologize, seven occurred directly within a recurring tripartite structure, the AI Limitation Formula (ALF): [Softener] + [AI Identity Disclaimer] + [Limitation Statement], e.g., “I apologize, but as an AI language model, I cannot create a complete Firefox extension design for you.” All 12 instances of unable and all 18 instances of cannot were embedded in explanatory structures rather than direct refusal forms, while however consistently pivoted from limitation to an alternative offer. N-gram analysis (Table 2) confirms the ALF as the single most dominant formulaic structure in the corpus, with “as an ai language model” appearing at f=37 and six ALF variants collectively accounting for 234 occurrences across 1,634 N-gram types.

Table 2. *Key N-gram clusters and their pragmatic functions*

Rank / Freq.	N-Gram	Pragmatic Function
1 / 41	ai language model	Core ALF component preceding every speech act
6 / 37	as an ai language model	Complete ALF opener obscuring speech act type
21 / 12	however i can	Refusal-to-alternative pivot rendering refusal invisible
26 / 9	ai language model i cannot	Compressed refusal formula, limitation buried in identity
44 / 8	response i apologize	Formulaic apology functioning as politeness signal

In this structure, the core pragmatic act, typically refusal, is delayed and partially obscured by the preceding softener and identity disclaimer, requiring EFL learners to process two preparatory moves before accessing the communicatively relevant content. These findings directly address RQ1: AI-generated responses are dominated by the ALF and however-pivot constructions that systematically compress multiple speech acts into single, structurally opaque utterances.

4.2 Cross-Category Comparison: Which Pragmatic Category Is Most Difficult?

Sub-corpus analysis using concordance and keyness measurement (Log-Likelihood and %DIFF against the full-corpus baseline) compared the four pragmatic categories directly (Table 3).

Table 3. *Cross-category comparison of pragmatic markers*

Marker	Apology	Refusal	Request
however	f=6	f=7	f=13
cannot	f=3	f=5	f=8
unable	f=1	f=6	f=4
apologize	LL=23.36, %DIFF=+1689	0	f=2
please	—	f=1	LL=20.29, %DIFF=+260
model (ALF)	—	LL=45.85, %DIFF=+530	—

The apology sub-corpus showed *apologize* as a highly significant positive keyword (LL=23.36, %DIFF=+1689.03), appearing 1,689% more frequently than in the full-corpus baseline, near-total formulaic saturation that may prevent EFL learners from developing an accurate pragmatic association between the lexical item and genuine interpersonal apology. The refusal sub-corpus showed *model* as the most statistically significant genuine keyword across all four categories (LL=45.85, %DIFF=+529.69), confirming that the ALF statistically dominates refusal encoding; notably, *apologize*=0 in this sub-corpus, indicating that apology language is largely absent from refusal responses specifically. The request sub-corpus was defined by *please* (LL=20.29, %DIFF=+260.18, f=29), with concordance confirming that all instances function as conventionalized closing or instruction markers (“please feel free to ask,” “please let me know”) rather than genuine requestive

softening.

Most strikingly, the implicature sub-corpus showed *cannot*, *unable*, *apologize*, *unfortunately*, and *please* all at zero frequency; only *however* appeared (f=3), functioning as a topic transition rather than a refusal pivot. Keyness analysis identified only *linguistic* (LL=22.53) and *hallucination* (LL=19.31) as genuine distinctive keywords. This complete absence of explicit pragmatic markers, while every other sub-corpus retains at least three, constitutes a structural-absence finding rather than a frequency-based claim: AI-generated implicature responses are characterized by what they do not contain. This pattern is corroborated convergently by the pre-tagging data, where only 59% of exchanges showed Relevance Theory features (the lowest of the three frameworks), and by the cross-category chi-square result, $\chi^2(15) = 87.43, p < .001$, confirming the distributional differences are statistically non-random.

These findings produce the evidence-based difficulty ranking presented in Table 4, directly addressing RQ2. The ranking reflects two distinguishable dimensions: structural encoding opacity, the degree to which illocutionary force is concealed within surface structure, on which refusal ranks highest due to the ALF and however-pivot; and total inferential demand, the cognitive processing effort required when no structural marker is present at all, on which implicature ranks highest. Because structural opacity still leaves learners a surface signal to decode, while inferential demand in implicature offers no scaffolding whatsoever, the study prioritizes inferential demand as the more consequential dimension, yielding implicature as the overall hardest category.

Table 4. Evidence-based pragmatic difficulty ranking

Rank	Category	Key Evidence	Primary Challenge
1 (hardest)	Implicature	All explicit markers = 0; linguistic LL=22.53; hallucination LL=19.31	Pure inference, no pragmatic signals
2	Refusal	model LL=45.85; cannot + unable combined f=11	Refusal hidden inside ALF + however-pivot
3	Apology	apologize LL=23.36, %DIFF=+1689	Formulaic sincerity marker
4 (easiest)	Request	please LL=20.29, %DIFF=+260	Cooperative language conceals closing intent

4.3 Theoretical Explanation of Pragmatic Gaps

Applying Grice's (1975) Cooperative Principle, the corpus reveals discourse patterns that resemble maxim non-observance across all four maxims, though the study does not claim AI intentionally flouts maxims in the human sense, since AI lacks communicative intent and cooperative awareness. Instead, AI discourse produces maxim-analogue patterns: structural outputs creating the same inferential demands as flouting or violation, without the underlying communicative design. Quantity-analogue verbosity (furthermore, additionally in elaborated refusals) appeared in 58/122 exchanges; Quality-analogue formulaic sincerity markers appeared in all nine apologize instances; Relation-analogue however-pivots appeared at f=64; and Manner-analogue multi-move compression characterizes the ALF (f=37).

Applying Searle's (1969) felicity conditions reveals two systematic failures. The sincerity condition fails in all nine apology instances, since apologize is used without reference to an identifiable interpersonal offense. The essential condition, that a speech act's illocutionary purpose be clearly expressed, fails across all 30 combined cannot/unable instances, since refusal is dissolved into explanatory elaboration rather than encoded through direct, explicit forms. For EFL learners whose first languages encode refusal through more direct negative forms, recovering illocutionary force from this distributed structure may be particularly difficult.

Sperber and Wilson's (1986) Relevance Theory predicts that communication succeeds when cognitive

effort is proportional to contextual effect. The ALF requires processing three simultaneous communicative moves, identity, softening, and limitation, before the core pragmatic act is recoverable, representing a high-effort/low-effect imbalance. The implicature sub-corpus provides the most direct evidence for this prediction: with all explicit markers absent, EFL learners must rely entirely on inference, maximizing cognitive effort while minimizing communicative clarity, a condition Relevance Theory predicts will most severely impede successful interpretation. This addresses RQ3 directly: the three frameworks converge to explain distinct dimensions of the same underlying phenomenon, that AI-generated discourse compresses, conceals, or omits the structural signals EFL learners rely on to recover pragmatic meaning.

Taken together, the findings show that AI-generated discourse is not pragmatically neutral but encodes meaning through systematic, identifiable structures, principally the AI Limitation Formula and the however-pivot, that create measurable, rankable interpretive demands across pragmatic categories. Implicature, lacking any explicit structural marker, places the greatest total inferential burden on EFL learners, while refusal, though structurally signaled, conceals its illocutionary force most systematically through compression and redirection. Apology language is rendered nearly meaningless as a sincerity marker through formulaic overuse, while request language remains the most pragmatically accessible category. This synthesis directly answers all three research questions and substantiates the study's central claim: EFL learners interacting with AI systems face a distinctive and previously undocumented set of structural pragmatic challenges that differ in kind, not merely degree, from those documented in human-human interlanguage pragmatics research.

5. Conclusion

This study investigated the pragmatic structures of AI-generated discourse and their potential interpretive implications for EFL learners through a corpus-based analysis of 122 human-AI exchanges from the WildChat dataset. The analysis identified the AI Limitation Formula (ALF), a tripartite structure (softener + AI identity disclaimer + limitation statement) appearing 37 times as a complete cluster, and the however-pivot construction ($f=64$), which systematically transitions from limitation to alternative offer, together with formulaic politeness markers (please, $f=27$; apologize, $f=9$) that perform closing and softening functions rather than genuine communicative acts. These features collectively constitute a distinctive AI pragmatic style characterized by performative politeness and structural indirectness.

Cross-category comparison produced a clear, evidence-based difficulty ranking: conversational implicature emerged as the most inferentially demanding category, since all explicit pragmatic markers were absent and only metalinguistic keywords (linguistic, hallucination) proved statistically distinctive, followed by refusal, whose illocutionary force is most systematically concealed through the ALF and however-pivot (model: $LL=45.85$), apology, rendered highly formulaic (apologize: $LL=23.36$, $\%DIFF=+1689\%$), and request, the most pragmatically accessible category (please: $LL=20.29$, $\%DIFF=+260\%$). This ranking reflects two distinguishable dimensions, structural encoding opacity, on which refusal ranks highest, and total inferential demand, on which implicature ranks highest, with the latter treated as the more consequential dimension since it offers learners no structural scaffolding at all.

Grice's Cooperative Principle, Speech Act Theory, and Relevance Theory together explain this pattern. The corpus exhibits maxim-analogue patterns across all four Gricean maxims, verbosity, formulaic sincerity markers, however-pivots, and ALF compression, that resemble maxim non-observance from the hearer's perspective without reflecting genuine AI communicative intent. Speech Act analysis reveals systematic failure of the sincerity condition in apology encoding and the essential condition in refusal encoding. Relevance Theory explains the resulting interpretive difficulty as a high-effort/low-effect imbalance, most acute in implicature responses, where no explicit marker guides inference. Together, these frameworks demonstrate that AI-generated discourse, while appearing conversationally

cooperative, requires substantial inferential processing that may create genuine interpretive difficulty for EFL learners.

The study makes three theoretical contributions: it introduces the AI Limitation Formula as a pragmatic structure specific to AI-generated discourse; it provides corpus-based quantitative evidence for previously assertional claims about AI formulaicity (Eragamreddy, 2025); and it extends Relevance Theory's cognitive framework to demonstrate that the high-effort/low-effect imbalance created by ALF structures offers a theoretically grounded explanation for EFL interpretive difficulty. One deliberate scope boundary should be distinguished from the study's limitations: because the design is a corpus-based structural analysis rather than an interlanguage pragmatics experiment, interpretive difficulty is theoretically inferred rather than empirically measured through learner performance data. Two genuine limitations remain: the implicature sub-corpus is small ($n=6$), limiting the statistical reliability of its keyness findings, and the WildChat dataset's topic diversity introduced topic-specific noise (character names, URLs, product names) requiring manual filtering during keyness interpretation.

Future research should complement this corpus-based analysis with learner response data, think-aloud protocols or discourse completion tasks, to test empirically whether EFL learners misinterpret the structures identified here; investigate the ALF across different AI systems (GPT-4, Claude, Gemini) to determine whether it is ChatGPT-specific or a cross-system pattern; conduct longitudinal studies examining whether learners develop pragmatic awareness of AI formulaic structures with explicit instruction; and further investigate the relationship between AI hallucination and implicature encoding identified in the keyness analysis. Pedagogically, the findings suggest that EFL instructors should explicitly teach the ALF pattern, treat the however-pivot as a distinct pragmatic signal requiring instruction, and provide explicit inferential-strategy guidance for implicature comprehension activities involving AI-generated text, since AI-generated discourse is not pragmatically neutral but encodes meaning through systematic structures that EFL pedagogy must address as AI becomes increasingly central to language learning.

References

- Andersson, M., & McIntyre, D. (2025). Can ChatGPT recognize impoliteness? An exploratory study of the pragmatic awareness of a large language model. *Journal of Pragmatics*, 230, 45–62.
- Anthony, L. (2019). *AntConc (Version 3.5.9) [Computer software]*. Waseda University.
- Austin, J. L. (1962). *How to do things with words*. Oxford University Press.
- Bardovi-Harlig, K. (2013). Developing L2 pragmatics. *Language Learning*, 63(s1), 68–86.
- Bhar, S. (2026). Artificial intelligence in EFL speaking instruction: A systematic review of pedagogical design, affective conditions and instructional input. *Computer Assisted Language Learning*, 39(1), 1–28.
- Brunet-Gouet, E., Vidal, N., & Roux, P. (2023). Do conversational agents have a theory of mind? A single case study of ChatGPT with the Hinting, False Beliefs and False Photographs, and Strange Stories paradigms. *arXiv preprint*.
- Eragamreddy, N. (2025). The impact of AI on pragmatic competence. *International Journal of Language and Linguistics*, 13(1), 1–15.
- Erdogan, E., & Kitson, L. (2025). Integrating AI in language learning: Boosting pragmatic competence for young English learners. *TESOL Journal*, 16(2), e789.
- Grice, H. P. (1975). Logic and conversation. In P. Cole & J. L. Morgan (Eds.), *Syntax and semantics, Vol. 3: Speech acts* (pp. 41–58). Academic Press.
- Ifadloh, N., Abdullahi, S., & Rukmana, D. (2025). Pragmatic failure in EFL learners' emails and AI

- grammar tools feedback. *Journal of English Language Teaching and Linguistics*, 10(1), 77–94.
- Kasper, G., & Rose, K. R. (2002). *Pragmatic development in a second language*. Blackwell.
- Manzoor, A. (2025). Revisiting classical pragmatic frameworks for human-machine interaction. *Journal of Cognitive Pragmatics*, 3(1), 1–19.
- Miehling, E., Nagireddy, M., Sattigeri, P., Daly, E., & Singh, M. (2024). Language models in dialogue: Conversational maxims for human-AI interactions. *arXiv preprint*.
- Nazeer, M., Khan, S., Nawaz, A., & Rehman, F. (2024). An experimental analysis of pragmatic competence in human-ChatGPT conversations. *Journal of Applied Linguistics and Language Research*, 11(3), 55–71.
- Nurjain, L. R., Maulana, R., Nurjain, A., Subki, A., & Damayanti, H. (2025). Pragmatic power or peril? Indonesian students responding to ChatGPT's linguistic subtleties in writing argumentation. *Journal of English Education and Linguistics Studies*, 12(1), 23–41.
- Rosen, A. (2024). Sincerity, intention, and the felicity of machine speech acts. *Minds and Machines*, 34(2), 211–229.
- Searle, J. R. (1969). *Speech acts: An essay in the philosophy of language*. Cambridge University Press.
- Sperber, D., & Wilson, D. (1986). *Relevance: Communication and cognition*. Blackwell.
- Tahir, R. (2025). AI-assisted fluency development in EFL learners. *Language Teaching Research Quarterly*, 28, 14–32.
- Taguchi, N. (2011). Teaching pragmatics: Trends and issues. *Annual Review of Applied Linguistics*, 31, 289–310.
- Tan, L. (2025). Sociolinguistic naivety in large language model interaction. *Computational Pragmatics*, 2(1), 33–49.
- Tosounidou, E. (2024). Examining Greek EFL teachers' pragmatic competence, beliefs and challenges in integrating L2 pragmatics instruction. *Language Teaching Research*, 28(4), 1102–1124.