

A Pragmatic Study of Indirectness in Artificial Intelligence's Generated Language

Noor Salim Al-Asadi

Faculty of Education for Girls/University of Kufa/ / Iraq

ABSTRACT: Artificial intelligence tries mainly to create human-like language as means of communication. Since the creation of AI powered systems, such as chatbots and virtual assistant, their language become more advanced. However, the extent to which degree AI's language can imitate human's language pragmatically is still a question needs to be answered. Thus, this study detects the ability of different AI assistants to produce indirect language. 50 questions of various types were asked to different AI assistants like DeepSeek, Gemini and Chat GPT. The results show that all the three programs use indirect language strategies and devices, yet, they vary in the usage. ChatGPT's language is more indirect than Gemini and Deepseek. It uses assertives qualified with epistemic hedges and exhibits negative politeness as well as different linguistic markers. In contrast, DeepSeek's language is direct and authoritative as it uses clear assertive language and limited use of hedges. Gemini occupied a middle position for its indirectness remained less sophisticated and less frequent than that of ChatGPT.

KEYWORDS: Artificial Intelligence, Pragmatics, Hedges, Vagueness, Indirect language.

1. INTRODUCTION

Human communication extends far beyond literal meaning. People often rely on indirectness to reach specific goals such as: negotiating relationships, protect face, or communicate intentions. Pragmatics, which is a branch of linguistics, provides the theoretical tools analyze such phenomena by examining language use in context and uncovering intentions behind utterances. In recent years, artificial intelligence (AI) has occupied a great attention by human users in daily life to make life easier. It is highly witnessed that these programs use language with proficiency in imitating human's natural languages. Recent studies have begun to investigate indirectness in AI and human-robot interaction (Koo et al., 2025; Zhang et al., 2025; Park et al., 2024). These works highlight the importance of pragmatic inference in AI systems but remain limited in scope, either by focusing on a single pragmatic framework or restricting analysis to comprehension. Therefore, this study aims to fill this gap by adopting a broader pragmatic perspective, combining multiple theories, and analyzing both understanding and production of indirectness in AI-human interactions.

1.1 Questions of the study

1. To what extent can AI systems use indirect speech acts in interaction with humans?
2. How effectively can AI systems produce indirect language in ways comparable to human communication?
3. What are the indirect pragmatic strategies and devices employed in AI responses?
5. What differences emerge between various AI assistants in using indirect language?
6. What are the functions behind using indirect language by AI programs?

1.2 Aims of the study

1. Detecting AI's ability to generate and use indirect language in human-machine interaction.
2. Analyzing the use of linguistic devices in AI communication.
3. Integrate multiple pragmatic theories (Grice, Brown & Levinson, Searle) into a unified framework for analyzing indirectness in AI.
4. Comparing the use of indirectness Pragmatic strategies and linguistic devices between different AI assistants.

1.1 Hypothesis

1. It is hypothesized that AI systems use indirect language to achieve various pragmatic functions.
2. It is hypothesized that AI assistants vary in the use of indirect language.

1.2 Limitations of the Study

The study is concerned with limited AI assistants like ChatGPT, Gemini and Deepseek. Also, differences in culture and language were not taken into account because all tests were given in English with a normal US accent.

2. LITERATURE REVIEW

Since the creation of computers, there was an interest in the communication between humans and software. Natural language processing (NLP) refers to the ability of computers to help humans with everything related to language such as: "translating, answering questions, analysing voices and texts, and making conversations" (Jurafsky, 2008: p 8). During that time, computer's ability to do natural language processing was limited. Programs were not able to do advance and professional language processing. As a result, many serious efforts tried to develop computational programs. In order to be able to do that, programs should have phonological, morphological, syntactic, semantic and pragmatic faculty.

The 1950s marked a true starting point for language processing. Mainly with Chomsky's theory of transformational grammar, programs were designed to generate sentences from basic sentence. Eliza (1966), which is one of the earliest chatbot, used formal grammar to stimulate sentences without real knowledge of grammar or context. This program uses language simply by matching patterns of language to recognize phrases like "I need X" and change the words into suitable outputs like "What would it mean to you if you got X?" (Weizenbaum, 1976).

Later on, there was an addition of semantics by using semantic networks, logic-based representation, and frames. However, still there was very limited relation to relation with context. An example of early systems with semantic knowledge was Winograd's (1972) 'SHRDLU'. The latter is a computer program that could understand instructions of the speaker within a micro world of blocks.

The move toward pragmatics began in the late 1970s and 1980s, when researchers acknowledged that meaning is context-dependent and tied to speaker intentions. The foundation of pragmatic theories such as: Austin and Searle's speech act theories and conversational implicature led to the development of advanced systems. Systems like 'PAULINE' by Hovy (1990) focuses mainly on generating natural language under pragmatic constrictions. Similarly, PENMAN which depends on systemic functional grammar approaches to generation by linking linguistic choices with communicative choices.

During the (1990- 2017), a huge development happened in the system of programs. The development happened as follows: (1990-2000) statistical machine translation the availability of large bilingual corpora sparked the statistical revolution in NLP rather than manually encoding grammar (Brown et al, 1993); (2005-2017) introduction of basic neural networks in NLP (Mikolov, 2013). (2013-2017) Systems began to capture semantic relationships with better handling of long-range context compared to statistical models adoption of Recurrent Neural networks (RNNs), LSTMs, Seq2Seq Models (Sutskever et al., 2014), improved translation quality (Bahdanau et al, 2015), (2017- 2019) smoother and more context-aware language, (2019- present) AI moved into more advanced language use with near human communication abilities and can do academic tasks like translation and summarizing (GPT-3, Chat GPT, Gemini).

2.1 Previous Studies on Indirectness in AI Assistants

The field of studying indirectness in interaction with robots and AI is obviously new. Recent studies have tackled this issue but from different perspectives. Koo et al. (2025) investigate how well large language models (LLMs) understand the intentions behind utterances, especially indirect speech acts in Korean (Koo et al., 2025). Zhang et al. make an empirical study on robots to test their ability to understand indirect speech acts ISA in human-robot collaboration HRC. The results show that robots can deal with indirect speech acts generally. However, their ability is highly affected with task and context which requires a balanced use of direct and indirect requests (Zhang et al., 2025). Park et al. (2024) introduce MultiPragEval, a multilingual evaluation framework based on Gricean principles, emphasizing the importance of pragmatic inference for advanced AI understanding. Seals and Shalin (2023) argue that current conversational AI still struggles with pragmatic competence, showing failures in preserving meaning, interpreting indirectness, and maintaining contextual awareness. Overall, existing studies focus mainly on comprehension, leaving the production of indirect language largely unexplored and lacking comparative approaches. The present study addresses this gap by examining the production of AI language by using pragmatic theories to analyze how AI employs suitable indirectness devices in interaction.

3. THEORETICAL FRAMEWORK

3.1 Pragmatic Strategies of Indirectness

Thomas (1995) defines indirectness as the mismatch between literal meaning and the intentional meaning. Indirectness is one of the main aspects of conversational style that can help the speaker to convey his/her intention without expressing them explicitly (Tunnen, 2006). The use of indirectness in communication can be explained through several pragmatic strategies, each grounded in major theoretical approaches:

3.1.1 Speech Act Theory

Speech acts, generally, refer to the theory initiated by J. L. Austin and then further developed by his pupil John R. Searle. In 'How to Do Things with Words' Austin has presented the notion of speech acts by focusing on the ability to do actions by words (Austin, 1962 :4). He introduced three facets of acts that simultaneously performed when saying something (Huang, 2007:102). These acts are: Locutionary, Illocutionary and Perlocutionary acts. Mainly, speech act theory is related to Illocutionary acts that refer to the

A Pragmatic Study of Indirectness in Artificial Intelligence's Generated Language

action intended to be performed by a speaker in uttering a linguistic expression, by virtue of the conventional force associated with it, either explicitly or implicitly.

Searle (1976:11-13) introduced a new taxonomy of speech acts instead of Austin's old taxonomy. Allott (2010:179) has described Searle's taxonomy as the most influential. These speech acts fall into five different kinds:

1. Representatives: They commit the speaker to the truth of the expressed proposition, e.g., asserting, claiming, concluding, reporting, and stating...etc.
2. Directives: They are intended to produce some effect through action by the hearer. They express what the speaker wants, e.g., orders, commands, advice, commands, orders, questions, and requests....etc.
3. Commissives: They commit the speaker to some future course of actions, e.g., promises, swears, offers, pledges, refusals, and threats...etc.
4. Expressives: They have the function of expressing or making known, the speaker's psychological attitude towards a state of affairs which the illocution presupposes, e.g.: apologizing, blaming, congratulating, praising, and thanking....etc.
5. Declarations (or declaratives): They are acts which in their uttering a state of affairs comes into being e.g., naming, declaring...etc. (Searle, 1976 :10-14)

3.1.2 Grice's Cooperative Principle

Grice's theory of Cooperative Principle is one of most important theories within the field of pragmatics (Zhou, 2009: 43) which led to the development of pragmatics as a separate discipline within linguistics (Hadi, 2012: 69). Grice states: "make your conversational contribution such as is required, at the stage at which it occurs, by the accepted purpose or direction of the talk exchange in which you are engaged". Also, he introduces the four maxims which are: Quantity, Quality, Relation, and Manner. These maxims help participants to converse in a maximally efficient, rational, co-operative way (Levinson, 1983: 102). In the example "Where is Peter? – He is in the garden," all maxims are maintained. Grice points out that individuals do not always observe maxims and may "violate," "opt out," "clash," or "flout" them, often "to show that they are polite." (1975: 49). Even when maxims are not observed, speakers remain cooperative through implicature, "what is suggested but not formally expressed," which allows listeners to infer meaning from context.

3.1.3 Politeness Theory

The notion of politeness is widely used in pragmatics and explains the choice of different utterances. Politeness is increased through indirectness (Leech, 1983) and plays a fundamental role in people's choice of linguistic expressions (Thomas, 1995). It refers to the tendency of people to keep face and to preserve both their own face and that of others (Trosborg, 1995). Ide (1989) states that linguistic politeness involves intentional strategies for smooth communication and conformity to expected norms. Politeness theory has received extensive attention, with Brown and Levinson's (1978/1987) model. 'Face' as defined by Goffman (1967) as "the positive social value a person effectively claims," is central to this theory. Brown and Levinson distinguish between negative face ("freedom of action") and positive face ("approval from others"). Face-threatening acts occur when speakers express disagreement, make requests, give advice, or present criticism (Brown & Levinson, 1987). They differentiate between threats to negative face such as orders, requests, suggestions and threats to positive face, such as disapproval, criticism, or bringing bad news. Speakers use strategies to diminish face-threatening acts, beginning with choosing whether to perform them. They may act off-record using metaphor, irony, rhetorical questions, understatement, or hints. They may go on record baldly without redress, or use redressive action positive politeness or negative politeness to minimize threat to the addressee's face.

3.2 Indirectness Devices

In order to analyze the indirectness devices used in the language of AI programs in communication, the classifications of Hinkel (1995) is adopted.

1. Rhetorical devices:

- a. Rhetorical questions and tags
- b. Disclaimers and denials: "not mean to/ imply/say, x is not y, not even, no way, not ? adjective/verb/noun/adverb...."
- c. Vagueness and ambiguity: "a lot of, approximately, around, many/much, number of, pieces, x or y, x or so, several, aspects of, facets of, good/bad, and so on, seldom, who knows, do(es) usually/often/occasionally/sometimes whatever (pron.), some..."

2. Lexical and Referential Markers:

- a. Hedges and hedging devices "may, can, likely, possibly, seemingly, about, in a way, kind of, more or less, most, by some/any chance, hopefully, perhaps, in case of, as is well known, as people say, apparently, basically, according to, actually, relatively, potentially, probably...."
- b. Point of view distancing: "I believe/think, I am concerned/I would like to think..."
- c. Downtoners: "at all, almost, hardly, mildly, nearly, partly, slightly, somewhat, only, as good/well as, enough, at least, merely...."
- d. Diminutives: "a little, a few, a bit, virtually..."
- e. Discourse particles: "anyway, anyhow"
- f. Demonstratives: "that, these...."
- g. Indefinite pronouns/determiners: "everybody, nobody, anything, someone.."

A Pragmatic Study of Indirectness in Artificial Intelligence's Generated Language

h. Discourse understatements: “fairly, quite, rather, not bad...”

3. Syntactic markers and structures:

- a. Passive voice
- b. If conditionals.

3.3 ADVANTAGES OF INDIRECTNESS IN AI

3.3.1 Enhances Politeness

Indirectness can help in softening the authoritative statements. In other words, it can protect the users negative face (Brown & Levinson, 1987) by avoiding direct imposition and gives a space for the user to make a free decision. For example, instead of using commands, the speaker can use suggestions to reduce pressure and to maintain politeness. in addition, this strategy can make the conversation more cooperative.

3.3.2 Avoiding Responsibility

Indirectness can be used to avoid responsibility which is needed in AI programs to avoid misinformation. Systems are designed to avoid absolute or deterministic claims, especially in sensitive domains (health, legal, financial). Different devices and strategies can be used achieve this aim. For examples, Hedging (e.g., might, may, seems) communicates uncertainty to exhibit epistemic caution.

3.3.3 Strategic Evasion

Instead of giving direct language, systems may use indirect language to avoid answering questions or shift the user's focus to away from difficult or unwanted questions and stay cooperative. For instance, rather than answering sensitive or potentially embarrassing questions, programs employ indirect language as a means of strategic evasion.

4. METHODOLOGY

A corpus 50 human-Chatbot interactions about different topics have been collected. All the conversations are designed to detect the AI systems' usage of indirect strategies and devices. The analysis employs an integrated framework combining pragmatic strategies Grice's theory of conversational implicature, speech act theory and Brown and Levinson's Politeness Theory , as well as indirectness devices which include: Rhetorical devices, Lexical and referential markers, Syntactic markers and structures.

4.1 Data Analysis

4.1.1 Analysis of the Pragmatic Strategies in Generated AI texts

All the indirect strategies and devices are found the texts generated by AI programs. However, they vary in the way and the types used. The texts generated by Chat GPT exhibit more indirect aspects than other AI programs like DeepSeek or Gemini. For instance, it is noticed that ChatGPT uses assertives qualified with epistemic hedges, e.g., “*It seems that...*” and “*Some studies suggest...*”. Thus, answers of ChatGPT are often balancing informativeness with cautiousness. Directives occur sparingly and are typically softened to maintain neutrality. Such usage is to reduce epistemic commitment, to avoid responsibility, and to maintain politeness. Also, it is noticed that direct acts are rarely used by ChatGPT. However, if they occur, they would be combined with hedges, e.g. “*Divorce often causes sadness*”.

Indirect Assertives, Expressives and Directives are widely used as in:

“*You may want to practice a little every day...*”

“*She's kind of quiet but really supportive.*”

“*AI is likely to affect jobs...*”

On the other side, Deepseek's language is more assertive than the other programs. The use of assertive speech acts is obvious in its texts as in “*The information on the internet is not uniformly reliable*”. Directives are less used than assertives, appearing only in general advice, as the text focuses primarily on informing the reader and providing authoritative guidance, with minimal use of hedges and limited vagueness. In addition, it is noticed that DeepSeek uses direct speech acts more than indirect speech acts as in (No, absolutely not..., One of the most dangerous assumptions...).

The Gemini text, in contrast, emphasizes on directive speech acts, providing guidance and instructions, such as “*you should apply a critical lens*” and “*focus on sources that have been researched*”. Assertives are present in the texts and given as recommendations. Gemini uses both direct and indirect speech acts.

Politeness strategies across AI-generated texts from DeepSeek, Gemini, and ChatGPT reveals clear differences in how each system manages face-threatening acts. The DeepSeek text is direct and authoritative, with minimal attention to politeness. Negative politeness strategies are largely absent, and there is limited use of positive politeness (e.g. the use of inclusive pronouns: “we,” “us”, “*We must evaluate critically*”) which is used to Builds solidarity with the user. Generally, the texts focus on delivering factual information with little concern for imposing on or accommodating the reader. The Gemini text demonstrates moderate politeness, particularly through advisory phrasing such as “*you should apply a critical lens*” and “*focus on sources that have been researched*”.

A Pragmatic Study of Indirectness in Artificial Intelligence's Generated Language

Both positive and negative politeness are used, but still limited as the program tries to give information rather than communicate emotionally.

In contrast, ChatGPT-generated texts show high politeness, employing frequent hedges and cautious phrasing (“*you may want to consider...*,” “*it seems that...*” “*may slightly decrease*”, “*You could try starting with basic vocabulary...*”) to mitigate face-threatening acts and to avoid strong commitment. Here, negative politeness is dominant, complemented by moderate positive politeness that acknowledges the reader’s perspective and build connections with humans, (e.g., “*From the way you interact, you seem thoughtful...*”, “*It really depends on personal preferences.*”, “*It’s not a simple question.*”).

In terms of Grice’s Cooperative Principle, the data largely follow the maxims of Quality and Relation, providing relevant, cautious, and evidence-based information. However, the high frequency of hedging occasionally reduces the Maxim of Manner by introducing intentional vagueness.

4.1.2 Analysis of Indirect linguistic devices in AI-Generated Texts

Rhetorical Devices

- a. Rhetorical questions are used as guiding questions (e.g. “Why information on the internet is not inherently reliable?”, How to evaluate internet reliability?”). These rhetorical questions help to guide and to organize the conversation for more details.
- b. Disclaimers / denials are used to reduce responsibility and to emphasize what is not being claimed the speaker (e.g. “I wouldn’t say it works for absolutely everyone., It’s not necessarily all destructive, It is generally not easy for parents.”)
- c. Vagueness refers to general words that are used to avoid precision. (e.g. *many jobs, some tasks, in many cases, a number of roles*)

Lexical and Referential Markers

- a. Hedges. These words are used to avoid certainty and responsibility. Major hedges appear frequently:(e.g., may, can, likely, often, usually, not necessarily, sometimes). Also, they are used to maintain both positive and negative politeness. Negative politeness which is used to avoid threatening the reader’s face (e.g., “You could try starting with basic vocabulary...”) and positive politeness which is used to build solidarity and connection with the human user (e.g., “From the way you interact, you seem thoughtful...”).
- b. Point-of-view distancing: AI systems often employ point-of-view distancing by using impersonal constructions (e.g., “It seems...”, “It might be...”, “This is one of...”) to avoid full commitment to the information. This strategy helps maintain neutrality and reduces potential face-threatening effects.
- c. Downtoners These expressions are used to avoid overstatement and absolute tone (e.g., a little,” “only,” “merely”).
- d. Discourse particles such as (well, so, actually, and maybe) appear frequently in the AI-generated texts. They have pragmatic functions. They soften the tone, frame uncertainty, and help keeping the flow of information.
- e. Demonstratives They indicate deictic expressions which are used to refer to things without forceful claims (e.g., This pattern has been predicted...”, “These changes are often observed...”, “In that case...” “this assumption,” “that information”)
- f. Indefinite pronouns, also used in AI language. They allow the systems to generalize information (e.g., someone, anyone, people).
- g. Understatements these expressions are used to maintain a balanced, non-imposing tone (e.g., “not necessarily what is accurate,” “it is not solely destructive.”). This can serve lowering responsibility of information.

Syntactic Markers

- a. Passive voice

The use of passive voice rather than active can serve indirectness to avoid commitment, and enhancing objectivity as in “The internet is filled with inaccuracies...”.

- b. Conditional structures

These structures are used to show uncertainty and introduce hypothetical reasoning and to maintain politeness by saving the user’s face. For example: “If we look at recent data, we can notice a clear pattern.”, “If you are interested, I can provide more details.”, “if this happen...”

4.2 Functions of the Common Indirectness Devices Across Programs

Through data analysis, it is noticed that politeness is the most frequent among the analyzed texts. It appears in softened advice (e.g., You could try board game.); mitigated evaluations (e.g., You seem thoughtful and curious); softened disagreement (e.g., I wouldn’t say it works for absolutely everyone.); Preference for suggestions over commands (e.g., It might help to practice every day.)

Avoiding Responsibility (Epistemic Caution) is also found in the AI generated texts. Avoiding strong commitments in uncertain domains is clear through the use of: Frequent hedging (e.g., may, might, could, seems, probably); Use of responsibility-avoiding expressions (Some studies suggest..., the results are unclear...).

Strategic Evasion also appears in the texts generated by the three programs. It is regularly used in topics to avoid conflict, to reduce risk and keep the conversation cooperative. This is accomplished through: Avoiding strong statements on sensitive topics (e.g. Divorce, self-harm or harm to others, legal/moral issues, etc.); Framing complex questions as multifaceted “It’s not a simple question.”; Shifting focus from a direct evaluation (e.g., Of course, I can’t really know everything about you, but...).

5. CONCLUSION

Indirectness is one of the main aspects of language to accomplish various communicative goals. As compared to human language, the use of indirectness in AI assistants is more academic and organized than natural and intuitive. Moreover, it is concluded that complex strategies of indirect language like: humor and rhetorical devices (like metaphor) are absent. Based on the analyzed data, this study concludes that strategies and devices of indirectness are present in all the texts generated by AI assistants although their use differs across systems. ChatGPT demonstrates higher degree of formality and indirectness compared to other programs such as: DeepSeek and Gemini. The comparison among the three systems reveals differences in their use of pragmatic strategies, including speech acts, positive and negative politeness, and various indirectness devices. Moreover, the analyzed texts show the use of different rhetorical devices such as (rhetorical questions and disclaimers); lexical and referential markers of indirectness such as: (hedges, point-of-view distancing, downtoners, Diminutives, Indefinite pronouns); and Syntactic Markers (passive voice and conditional structures). This reflects the AI designers' intentional employment of indirectness to fulfill distinct interactional purposes such as: enhancing politeness which makes the conversation more cooperative; avoiding responsibility; making strategic evasion specially concerning sensitive topics. Future research may explore how these strategies affect comprehension and perception in different contexts, or how they vary across languages and cultural communication norms.

REFERENCES

- 1) Allott, Nicholas. "Speech Acts and Indirectness: Theoretical Perspectives." *Journal of Pragmatics*, vol. 42, no. 2, 2010, pp. 179–193.
- 2) Austin, J. L. *How to Do Things with Words*. Oxford UP, 1962.
- 3) Bahdanau, Dzmitry, Kyunghyun Cho, and Yoshua Bengio. "Neural Machine Translation by Jointly Learning to Align and Translate." *arXiv*, 2015, arXiv:1409.0473.
- 4) Brown, Penelope, and Stephen C. Levinson. *Politeness: Some Universals in Language Usage*. Cambridge UP, 1987.
- 5) Brown, Peter, et al. "Statistical Machine Translation." *Computational Linguistics*, vol. 19, no. 2, 1993, pp. 263–311.
- 6) Curtone, John. "Face-Threatening Acts and Politeness Strategies in Interaction." *Pragmatics*, vol. 21, no.1, 2011, pp. 50–65.
- 7) Goffman, Erving. *Interaction Ritual: Essays on Face-to-Face Behavior*. Anchor Books, 1967.
- 8) Grice, H. Paul. "Logic and Conversation." *Syntax and Semantics*, vol. 3, edited by Peter Cole and Jerry L. Morgan, Academic Press, 1975, pp. 41–58.
- 9) Hadi, Farid. "Pragmatics and Conversational Structure." *Journal of Linguistic Studies*, vol. 12, no. 1, 2012, pp. 69–82.
- 10) Hinkel, Eli. *Pragmatics and ESL: Indirectness in Requests*. Applied Language Studies, 1995.
- 11) Hovy, Eduard. "Generating Natural Language under Pragmatic Constraints." *Computational Linguistics*, vol. 16, no. 3, 1990, pp. 147–161.
- 12) Huang, Yan. *Pragmatics*. Oxford UP, 2007.
- 13) Ide, Sachiko. "Formal Forms and Discernment: Two Neglected Aspects of Linguistic Politeness." *Multilingua*, vol. 8, nos. 2–3, 1989, pp. 223–248.
- 14) Jurafsky, Daniel, and James H. Martin. *Speech and Language Processing*. Pearson Prentice Hall, 2008.
- 15) Koo, Jihyun, et al. "Understanding Indirect Speech Acts in Large Language Models: Evidence from Korean." *AI & Society*, 2025.
- 16) Leech, Geoffrey N. *Principles of Pragmatics*. Longman, 1983.
- 17) Levinson, Stephen C. *Pragmatics*. Cambridge UP, 1983.
- 18) Mikolov, Tomas, et al. "Efficient Estimation of Word Representations in Vector Space." *arXiv*, 2013, arXiv:1301.3781.
- 19) Park, Yuna, et al. "MultiPragEval: A Multilingual Evaluation Framework for AI Pragmatic Inference." *Computational Pragmatics Journal*, 2024.
- 20) Seals, Megan, and Valerie Shalin. "Pragmatic Competence in Conversational AI: Challenges and Directions." *AI & Ethics Journal*, vol. 3, no. 2, 2023, pp. 121–140.
- 21) Searle, John R. *Speech Acts: An Essay in the Philosophy of Language*. Cambridge UP, 1976.
- 22) Sutskever, Ilya, Oriol Vinyals, and Quoc V. Le. "Sequence to Sequence Learning with Neural Networks." *Advances in Neural Information Processing Systems*, vol. 27, 2014.
- 23) Thomas, Jenny. *Meaning in Interaction: An Introduction to Pragmatics*. Longman, 1995.
- 24) Trosborg, Anna. *Interlanguage Pragmatics: Requests, Complaints and Apologies*. Mouton de Gruyter, 1995.
- 25) Tunninen, Tuija. "Indirectness in English Conversation." *Pragmatics & Society*, vol. 7, no. 3, 2006, pp. 305–327.
- 26) Weizenbaum, Joseph. *Computer Power and Human Reason: From Judgment to Calculation*. W. H. Freeman, 1976.
- 27) Zhang, Hui, et al. "Robot Comprehension of Indirect Speech Acts in Human–Robot Collaboration." *Robotics and Autonomous Systems*, vol. 180, 2025, pp. 104–121.
- 28) Zhou, Hong. "The Role of Gricean Maxims in Communication." *Journal of Pragmatics*, vol. 41, no. 1, 2009, pp. 43–57.