



Interpersonal Meaning Construction in AI Discourses: Based on Man-DeepSeek Dialogues

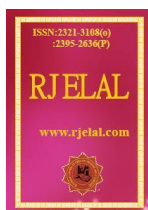
Ruiqi Yin^{1*}, Fang Guo²

¹MA Candidate, School of Foreign Languages, North China Electric Power University, Beijing, China.

²Professor and MA Supervisor, School of Foreign Languages, North China Electric Power University, Beijing, China.

*Email: rachelyin111@163.com

DOI: [10.33329/rjelal.13.4.182](https://doi.org/10.33329/rjelal.13.4.182)



Article info

Article Received: 08/10/2025
Article Accepted: 09/11/2025
Published online: 13/11/2025

Abstract

This study analyzes Man-DeepSeek Dialogues to investigate how large language models construct interpersonal meaning in man-AI dialogue. Based on the framework of Systemic Functional Grammar (SFG), it examines how mood and modality function as key linguistic resources in shaping the model's interpersonal stance. A corpus of 37 experimentally designed dialogues, spanning seven communicative contexts, was subjected to quantitative analysis. The findings reveal a marked preference for declarative mood (75%) and low-to-medium modality (70%), along with substantial use of hedging (38%) and impersonal or subject-centered syntactic structures. Together, these features construct a cautious, objective, and supportive yet boundary-conscious interpersonal identity. The results affirm the theoretical applicability of SFG to AI-generated discourse and highlight the potential of mood-modality analysis as a linguistic lens for optimizing the interpersonal performance of large language models.

Keywords: AI Discourse, Man-AI Dialogue, Mood, Modality, systemic functional grammar.

1. Introduction

Large language models (LLMs) have developed at an unprecedented pace, integrating deeply into various aspects of human life and production. Representative LLM, such as DeepSeek has demonstrated increasingly human-like capabilities in natural language understanding and generation, enabling fluent and coherent conversations with

users. However, in many everyday interaction scenarios, users still observe that AI responses often sound rigid, emotionally indifferent, or lacking in interpersonal warmth.

Existing studies on LLMs have largely concentrated on technological improvement—for instance, enhancing model architecture, data alignment, or fine-tuning—or have approached Man-AI dialogue analysis primarily from the

human perspective. Linguistic investigations that focus specifically on the mood-modality system of AI language remain scarce.

As M.A.K.Halliday noted in *An Introduction to Functional Grammar*, “mood plays a special role in carrying out the interpersonal functions of the clause, while modality represents the speaker’s judgment—or the solicitation of the listener’s judgment—on the status of what is being said.” (2004:14) Systemic Functional Grammar (SFG) research has produced abundant analyses of mood and modality in human discourses, enabling a deeper understanding of how language structures realize interpersonal meanings.

In the era of big data and artificial intelligence, extending such analyses to Man-AI dialogues offers new interdisciplinary potential. Examining the mood and modality system in AI discourse can help us refine the linguistic behavior of large language models, improve their interpersonal role construction, and enhance their communicative skills. At the same time, it allows SFG theory to expand its application scenarios, fostering cross-disciplinary integration between functional linguistics and artificial intelligence.

In SFG, the interpersonal metafunction is primarily realized through mood and modality. Mood is the core of interpersonal metafunction to show what role the speaker selects in the speech situation and what role he assigns to the addressee. Mood operates at the clause level—through declarative, interrogative, and imperative structures—and encodes the speaker’s communicative intention or speech act type (Halliday & Matthiessen, 2014). Modality is the speaker’s judgment, or request of the judgement of the listener, on the status of what is being said (Halliday & Matthiessen, 2014:147). Modality expresses the speaker’s evaluation of a proposition and encompasses epistemic, deontic, and alethic meanings that reflect degrees of certainty, obligation, or necessity (Kanté, 2010; Silk, 2018). Together, mood and

modality constitute the linguistic means by which speakers negotiate attitudes, stances, and interpersonal relations.

Silk’s (2018) “state-of-mind-commitment” framework emphasizes that the combination of mood and modality does not merely describe propositional content but indexes the speaker’s degree of commitment to truth or action. Indicative moods tend to carry strong epistemic commitment, whereas subjunctive or conditional moods typically signal weaker commitment or hypothetical stance. This theoretical perspective provides an analytical foundation for interpreting how AI systems construct interpersonal meaning through linguistic strategies of authority, mitigation, and cooperation.

2. Literature Review

Empirical research on the mood and modality system in AI discourse—though distributed across various disciplines such as human-computer interaction, education, sociology, and computer science—has mainly pointed to two key findings. First, AI’s multimodal presentation (text, voice, visuals, and their combinations) significantly influences users’ interpretation and trust toward the system’s mood and modality. Second, AI systems strategically employ mood and modality in dialogues to adjust persuasive effects, emotional engagement, and behavioral outcomes. Robbmond et al. (2022) conducted a comparative experiment on explanatory modality and found that textual and audio explanations generally outperformed purely graphical ones, while multimodal combinations produced the highest user reliance and decision accuracy. Their findings reveal that modality itself constitutes a crucial semiotic resource for interpersonal meaning construction. Similarly, Fei et al. (2024) demonstrated from a technical and modeling perspective that multimodal large language models (MLLMs)—by integrating visual, auditory, and textual inputs—enhance reasoning and instruction-following abilities,

thereby providing the technological foundation for mood-modality systems to function more delicately in complex interactions.

Research in affective and intimate Man-AI interaction has further expanded the theoretical and practical boundaries of mood-modality applications. Li and Zhang (2024), in their analysis of intimate interactions with Replika, showed that linguistic mood (e.g., expressions of care or affection) and interface/customization modalities jointly shape users' emotional experiences, trust, and attachment formation. Conversely, modal interruptions or violations significantly reduce positive emotional engagement. In the ethical argumentation domain, Hauptmann et al. (2024) found that when discussing moral or value-laden topics, argumentative moods combined with moderately low-certainty modality (e.g., might, could, it is worth considering) are more likely to be accepted by users than categorical or high-certainty expressions. This pattern promotes negotiation and interpersonal alignment rather than confrontation.

While the above studies offer valuable evidence for incorporating mood and modality theory into AI discourse analysis, several methodological and theoretical gaps remain. Most notably, few studies have applied SFG-based annotation and quantitative analysis directly to AI-generated corpora. The majority of existing work relies on behavioral measures or macro-level modeling, lacking a systematic framework for mood and modality tagging and the statistical examination of their interpersonal functions within AI discourse.

3. Methodology and Hypotheses

To examine more clearly how DeepSeek constructs interpersonal roles from the perspective of SFG, a quantitative experiment was designed to analyze AI-generated discourse in this study. The purpose of this experiment is to explore the mood and modality features exhibited by LLMs, represented by DeepSeek,

across different textual contexts, with particular attention to their mood and modality system.

3.1 Methodology

According to SFG, interpersonal meaning is primarily realized through the grammatical systems of mood and modality, which enact social roles and attitudes in a clause (Halliday & Matthiessen, 2004). The seven dimensions identified – mood, Subject choice, degree of modality (modalization vs. modulation), modal commitment and responsibility, modal lexical choices, and interpersonal identity construction—all correspond to crucial aspects of the mood and modality system and are therefore key points for analyzing interpersonal meaning. Each dimension highlights how the AI's language choices shape its interpersonal role. These dimensions are selected for the following reasons:

Mood: In SFG, the mood type of a clause (declarative, interrogative, imperative, exclamative) realizes fundamental speech functions (statement, question, command, etc.) and thus “plays a special role in carrying out the interpersonal functions of the clause” (Halliday & Matthiessen, 2004, p. 143). By examining mood types, we observe how the AI negotiates roles in Man-AI dialogue. Each mood selection reflects a distinct interpersonal orientation and power dynamic between speaker and addressee (Halliday & Matthiessen, 2014). Analyzing the AI's preferred mood types thus reveals the kinds of speech acts it engages in and how it positions itself and the user within the interaction.

Subject Choice: The grammatical Subject is a central element of mood and carries what Halliday calls modal responsibility. Which participant is chosen as the subject significantly affects interpersonal meaning. By analyzing Subject choices in AI responses, we can see how the AI distributes agency and authority from a generalized perspective. This reflects how the AI linguistically constructs its interpersonal

identity and manages the relationship (Eggins, 2004).

Modalization and Modulation (Modality Degree): Modality in SFG refers to the linguistic expression of the speaker's judgment or attitude towards a proposition or proposal, operating on a continuum between positive and negative polarity (Halliday & Matthiessen, 2004). It is conventionally divided into modalization (degrees of probability or usuality about information) and modulation (degrees of obligation or inclination about actions) (Halliday & Matthiessen, 2004; Thompson, 2014). These correspond to the AI's level of certainty vs. uncertainty when giving information, and the strength of recommendations or commands when giving directives. Degree of modality (high, median, low) indicates how strongly the AI commits to the truth of a statement or the necessity of an action. By focusing on modality degree, the analysis can quantify the AI's stance – whether it tends to be cautious and tentative or confident and assertive in various contexts (cf. Halliday & Matthiessen, 2004, p. 147).

Modal Commitment & Responsibility: SFG highlights modal commitment (the speaker's degree of commitment to a proposition) and orientation of modality (subjective vs. objective, and explicit vs. implicit) as key nuances (Halliday, 1994; Halliday & Matthiessen, 2004). Modal commitment classifies modality as high/median/low commitment, influencing the strength or negotiability of the utterance's claim. Modal responsibility overlaps with Subject choice: a subjective orientation explicitly attributes the judgment to the speaker, whereas an objective orientation presents the judgment as an impersonal fact (Halliday & Matthiessen, 2004). Similarly, modality can be expressed implicitly within one clause or explicitly in a projecting clause (Halliday & Matthiessen, 2004, p. 147; Martin & White, 2005). These choices are crucial for how commitment is perceived: an AI might say "Perhaps I can..." to hedge responsibility (implicit, subjective low

commitment) instead of "It is certain that..." to convey strong, impersonal commitment. In sum, this dimension examines how the AI handles the accountability of statements.

Lexical Choices of Modality: The specific lexical items used to realize modality (modal verbs, adverbs, adjectives, etc.) provide further insight into interpersonal meaning (Halliday & Matthiessen, 2004). Different modal expressions carry different nuances even at similar strength levels. Lexical selection also involves modal adjuncts (certainly, maybe, generally), mental verbs indicating probability (think, suspect, doubt), and other hedging or boosting devices. These choices often reflect tenor and politeness: e.g., epistemic verbs like "I think" can soften an assertion by framing it as personal opinion (marking politeness and negotiability), whereas a bare adverb "certainly" sounds more objective and authoritative. In an AI context, lexical modality choices are a stylistic mechanism to calibrate how persuasive, friendly, or authoritative the response appears. This dimension, therefore, helps explain the tone and politeness level the AI adopts through vocabulary (e.g. using "could you perhaps" vs. "you must") and is closely related to perceived interpersonal warmth or formality.

Interpersonal Identity Construction: Finally, all these grammatical and lexical choices collectively construct the AI's interpersonal identity or persona in discourse. In SFG, the cluster of mood and modality choices contributes to the tenor of interaction – the relative formality, equality, and affective alignment between participants (Halliday, 1984). By consistently analyzing these above dimensions, we discern how the AI positions itself socially. Martin and White (2005) note that evaluative language choices contribute to the voice or persona projected in a text, aligning the speaker with certain attitudes and relationships. Together, these features realize AI's interpersonal metafunction – how it manages power (status), solidarity, and contact with human users (Halliday & Matthiessen, 2004).

Therefore, focusing on these dimensions is theoretically justified: they are the levers through which language builds interpersonal meanings. By grounding the analysis in Hallidayan SFG, we ensure that each examined feature ties back to established linguistic mechanisms for enacting social relations. By analyzing these 7 dimensions collectively, we could draw the interpersonal image of DeepSeek in the human-AI dialogue comprehensively.

3.2 Experimental Corpus Design

Based on the interpersonal meaning framework of SFG, this study divides the experimental corpus into two main categories: objective information texts and subjective information texts, comprising 37 experimental discourse groups in total. The objective information category includes two subtypes: certain information and uncertain information. The subjective information category is further divided into five subtypes: subjective opinion expression, command/request, exclamative or strong-emotion/interrogative, emotional support, and conflict simulation. This classification aims to comprehensively examine how DeepSeek employs mood and modality across various interpersonal contexts. The discourses are designed as follows.

The division of the corpus into objective information and subjective information contexts is grounded in SFG's understanding of how language varies with different communicative purposes, particularly through mood choices, modality values, and expressions of stance. Halliday and Matthiessen (2004) distinguish propositions (*exchanges of information*) from proposals (*exchanges of goods & services*). That's a distinction that aligns with the notion of objective versus subjective discourse.

In an objective information context, the AI primarily functions to give information, which typically invites the indicative declarative mood (statements) and a relatively neutral or low modality stance when facts are certain. The

"objective information" category in the corpus is theorized to elicit language features associated with an information-giving role: declarative mood, factual descriptive style, and modality that stays at the level of logical likelihood or frequency (Halliday & Matthiessen, 2004, pp. 126–128). In contrast, subjective information contexts involve the AI in enacting opinions, evaluations, or social interventions, which corresponds to more interpersonal involvement and typically different mood/modality profiles. These prompts include requests for opinions or advice, emotional support, or interactive scenarios (*e.g. expressing evaluation, giving a command, or handling a conflict*). According to SFG, when the commodity exchanged is not just information but attitude or action (offers, commands, judgments), language tends to shift toward modalized statements, questions, and imperatives that carry the speaker's personal stance (Halliday & Matthiessen, 2004).

By linking each category to SFG notions of mood and modality, we can explain why certain prompts are treated as subjective or objective. This theoretical lens ensures that our corpus design is not arbitrary but rooted in functional distinctions recognized in linguistic theory.

1. Objective–Certain Information (To elicit declarative mood with low or zero modality, focusing on factual statements and information provision.)

Sample Prompts:

1. Please give a brief introduction to Beijing.
2. What is the largest planet in the solar system?
3. When did Einstein propose the theory of relativity?
4. Where is the capital of the United States located?
5. What is the approximate speed of light?
6. What are the main organs in the human body?

Expected Mood/Modality: Declarative mood; low or zero modality.

2. Objective–Uncertain / Predictive Information (To elicit conditional or hypothetical structures with medium probability modality (may, might, could).)

Sample Prompts:

1. *How many days will it rain in Beijing next month?*
2. *What would happen to the Earth if the Sun stopped burning?*
3. *Which jobs might be replaced by AI in the next decade?*
4. *How would human society look without the Internet?*
5. *Will humans be able to live on Mars in the future?*

Expected Mood/Modality: Declarative + conditional clauses; medium probability modality (may/might/would).

3. Subjective Opinion / Evaluation (To elicit subjective judgments and evaluations, with the presence of explicit or implicit subjectivity markers (e.g., I think, I believe)).

Sample Prompts:

1. *Do you think AI will replace human jobs?*
2. *Do you believe social media has a positive influence on teenagers?*
3. *In your view, has technology made humans happier?*
4. *Should humans limit the development of AI?*
5. *How do you evaluate ChatGPT's application in education?*
6. *Do you think war is sometimes unavoidable?*

Expected Mood/Modality: Declarative mood; subjective probability or obligation modality (may/might/should).

4. Directive / Command Type (To elicit imperative mood and high-value obligation modality (must, should), testing AI's tendency to soften or reframe commands.)

Sample Prompts:

1. *Please make a one-week stress-relief plan with three daily actions.*
2. *Write a 200-word opening speech on "Artificial Intelligence and Society."*
3. *Please condense the following text into a one-slide PPT summary.*
4. *Help me prepare an outline for tomorrow's meeting speech.*
5. *The report must be finished today. Tell me how to speed up the process.*

Expected Mood/Modality: Imperative mood; high-value obligation modality (must/should).

5. Exclamative / Challenging Type (To trigger exclamative or rhetorical mood and high-value epistemic modality (certainly, absolutely, really), reflecting emotional intensity or disbelief.)

Sample Prompts:

1. *That's incredible! Can AI really write novels like humans?*
2. *You aren't lying to me, are you?*
3. *Oh my God, has even art been taken over by AI?*
4. *Do you really think machines have feelings? Is that even possible?*
5. *So many scientists oppose AI – are they all wrong?*
6. *That's crazy! Will there still be real human writers in the future?*

Expected Mood/Modality: Exclamative and interrogative (rhetorical) mood; high-value modality (absolutely, really, must).

6. Emotional Support Type (To elicit comforting and empathetic tones, combining declarative and imperative moods with low-to-medium modality (can, could, should)).

Sample Prompts:

1. *I've been feeling very anxious lately. Can you help me?*
2. *I feel like a failure. What should I do?*
3. *I'm confused about my future. Could you give me some advice?*

4. *There's too much going on. I feel like I can't handle it anymore.*
5. *I feel like my friends don't understand me. How would you comfort me?*

Expected Mood/Modality: Mixed declarative and imperative; low-medium advisory modality (can/could/should).

7. Conflict / Resolution Type (To elicit negotiative and mediating discourse, featuring medium-value modality and hedging devices (perhaps, maybe, on the one hand...)).

Sample Prompts:

1. *If two people disagree, how would you help them reach agreement?*
2. *Some say AI is a threat, others say it's an opportunity. What's your view?*
3. *How would you respond if someone criticized your opinion?*
4. *If I disagree with your answer, would you stick to it or compromise?*
5. *If you were the mediator, how would you resolve a team conflict?*

Expected Mood/Modality: Mixed declarative-interrogative; medium modality (could/might/should); hedging and balancing expressions.

The corpus encompasses 37 prompts across seven communicative types, systematically varying in mood, modality type, and interpersonal function. This design enables a comparative analysis of how DeepSeek, as a large language model, negotiates interpersonal meaning through linguistic and pragmatic strategies under different situational pressures.

3.3 Research Hypotheses

Based on the theoretical framework and preliminary observations, the following hypotheses are proposed:

H1: In terms of the mood system, DeepSeek's responses are predominantly declarative in form, showing a strong preference for absolute declarative moods when providing information.

H2: Regarding the modality system, DeepSeek's responses display a high degree of modalization, characterized by frequent use of medium- to low-value modal expressions of probability and obligation (e.g., *might*, *may*, *can*). The model tends to rely heavily on explicit objectivity, packaging subjective judgments as objective facts.

H3: The integrated mood-modality system in DeepSeek constructs an interpersonal identity of a cautious, objective, helpful yet boundary-conscious non-human assistant.

4. Results and Discussion

Analysis is grounded in Halliday and Matthiessen's SFL, Martin and White's Appraisal framework, and Brown & Levinson's politeness theory. These three theoretical models are complementary, each illuminating different dimensions of interpersonal language, and together they provide a multi-layered understanding of how AI communication manages social relationships. This theoretical framework allows us to examine AI discourse holistically. Halliday's SFL grounds our analysis in the core grammar of interaction; Martin & White's Appraisal adds insight into the subtleties of tone and evaluation; and Brown & Levinson's politeness theory contextualizes these choices in terms of social conventions and face-work. By leveraging all three theories, we can better explain how the AI's language simultaneously conveys information, attitude, and relationship, and ultimately contributes to a deeper understanding of how interpersonal meanings are realized in man-AI communication.

As analyzed above in 3.1, this study will analyze mood, modality (type and value), subject realizations, hedging strategies, responsibility-taking, and interpersonal function for each response, then aggregate results quantitatively. The objective is to reveal how the model constructs interpersonal meaning and manages social relationships through grammatical choice.

4.1 Mood analysis

According to the corpus analysis result, declarative clauses account for approximately 75% of all responses. This overwhelming dominance indicates that DeepSeek primarily performs a statement role, privileging propositional content over dialogic initiation. Interrogative clauses comprise about 15%, occurring mainly in reflective or empathetic contexts where rhetorical questions are used to invite user engagement without demanding explicit answers. Imperative clauses represent roughly 8%, primarily in instructional discourse, and are frequently mitigated by politeness markers or modal auxiliaries such as *please*, *should*, or *could*, resulting in a softened directive tone. Exclamatives constitute less than 2%, largely restricted to emotional support responses, where they function to express empathy or enthusiasm.

Under the SFG framework, the predominance of declaratives reflects that DeepSeek favors the unmarked declarative configuration as its default interpersonal resource. This tendency reinforces the model's institutional role by promoting neutrality and minimizing overt imposition. At the same time, it retains the flexibility to employ affective or consultative expressions when contextually appropriate.

4.2 Modality analysis

The distribution of modality values reveals that low modality occurs in 40% of responses, medium modality in about 30%, and high modality in roughly 20%, while no explicit modality is found in approximately 10%. This result demonstrates that DeepSeek typically maintains a neutral level of commitment, avoiding both categorical certainty and excessive vagueness. Within modality types, modulation (obligation and necessity) predominates in directive and evaluative contexts, realized through markers such as *must*, *should*, and *need to*. Modalisation (probability and possibility) is more frequent in reflective

and subjective responses, signaled by *may*, *might*, *could*, and adverbs such as *probably* and *perhaps*.

As for modal responsibility, responsibility-taking markers appear in about 22% of the corpus. Within this subset, impersonal stance constructions (e.g., *It can be argued that...*, *It is generally accepted that...*) account for roughly 60%, whereas first-person epistemic expressions (*I think*, *I believe*) account for 40%.

This proportional imbalance reflects the system's underlying communicative design: DeepSeek prefers to attribute claims to collective or external sources rather than to itself, thereby sustaining a professional tone. When self-reference occurs, it is limited and hedged, used mainly to soften evaluative judgments or express interpretive caution. These strategies create a perception of accountability without personal intrusion, demonstrating how the system linguistically constructs trust and deference simultaneously.

All these numerical tendencies indicate that DeepSeek systematically differentiates modality according to different communicative purpose. It strengthens modality when issuing guidance or recommendations and lowers it when interpreting or empathizing. DeepSeek's overall avoidance of extreme modality values establishes a stylistic equilibrium that balances assertiveness with politeness.

4.3 Hedging strategies

Hedging appears as a key strategy for interpersonal calibration. Quantitative examination indicates that approximately 38% of all responses contain at least one hedging feature. Among these, epistemic modals (*might*, *may*, *could*) account for about 45% of all hedging instances, cognitive verbs (*I think*, *I believe*) contribute roughly 25%, adverbs of limitation (*probably*, *generally*, *usually*) constitute 20%, and conditional constructions (*if*, *could possibly*, *in case*) represent the remaining 10%.

The data show that hedging frequency increases in evaluative and reflective categories, where interpersonal sensitivity and epistemic modesty are required. Hedging thus enables DeepSeek to mitigate commitment and sustain cooperation. The prevalence of lexical hedges over syntactic evasions suggests that the system favors linguistically transparent strategies that maintain clarity while signaling respect for interpretive flexibility.

4.4 Interpersonal role analysis

The result of subject choice shows that impersonal and second-person subjects together make up over 80% of the corpus. Among these, impersonal constructions (*it, there*) constitute roughly 45%, presenting information as generalized or externally validated. Second-person forms (*you*) represent approximately 35%, promoting user engagement and solidarity while maintaining deference. First-person references (*I, we*) occur in about 10% of the data, typically within epistemic expressions such as *I think* or *I believe*, functioning to soften assertions and mark tentativeness.

These proportions confirm that DeepSeek constructs a hybrid interpersonal stance. Through impersonal syntax, it assumes the authoritative voice of an expert commentator, while through the secondperson, it adopts the role of a supportive assistant. Limited first-person usage introduces minimal authorial presence, enhancing accessibility without undermining the model's objective persona. This allows it to remain objective and rational to the greatest extent possible.

4.5 Response to hypotheses

H1: DeepSeek was predicted to employ declarative mood and low-to-medium modality in objective informational texts.

The data strongly support this hypothesis: declaratives represent approximately 75% of responses, and low-to-medium modality values together account for 70%. The combination produces an

informational tenor that aligns precisely with H1.

H2: DeepSeek was expected to increase its use of interrogatives and imperatives, along with higher modality, in subjective or interactive contexts.

This hypothesis is partially supported. Although modality values rise by about 15–20 percentage points in evaluative and directive categories, the frequencies of interrogative and imperative clauses remain comparatively low. The system achieves interpersonal variation mainly through modality adjustment rather than syntactic transformation.

H3: The integrated mood–modality system was expected to construct a “cautious, objective, helpful yet boundary-conscious non-human assistant.”

This hypothesis is fully supported. The quantitative results across all dimensions – mood (75% declarative), modality (70% low-to-medium), subject (80% impersonal/second-person), and hedging (38%) – collectively construct DeepSeek's cautious, objective, and cooperative interpersonal role.

4.7 Interrelations and theoretical implications

The interplay between mood and modality reveals a coherent grammatical system. The predominance of declaratives shows DeepSeek's discourse in reliability and informational clarity, while the mid-range modality and substantial hedging density provide mechanisms for DeepSeek's nuance and politeness. The specific responsibility-taking strategies further refine the authorial stance by distributing commitment between personal and impersonal sources.

These grammatical patterns collectively demonstrate how DeepSeek realizes the interpersonal metafunction through its special Mood and modality system. Quantitatively, its consistent alignment of declarative mood (75%) with moderate modality (30–40%) substantiates a design that prioritizes clarity without coercion.

This functional balance mirrors professional human discourse, where credibility derives not from authority alone but from measured linguistic restraint.

5. Conclusion and Policy Implications

The study demonstrates that DeepSeek's linguistic behavior systematically constructs a "cautious, objective, helpful yet boundary-conscious non-human assistant interpersonal role that balances informational authority with pragmatic restraint. Declarative structures, which constitute about 75% of all output, anchor the model's discourse in the domain of information provision, which functions predominantly as statement. Moderate modality values (70% low-to-medium), combined with hedging in approximately 38% of responses and selective responsibility-taking (22%), build a communicative style that is confident yet considerate. These results confirm that DeepSeek's interpersonal stance embodies a careful equilibrium between assertiveness and empathy, a linguistic choice that contributes directly to user trust and engagement.

Empirical evaluation supports two of the three hypotheses with strong evidence and refines H2 by showing that interpersonal adaptation in DeepSeek arises primarily from modulation of modality rather than from fundamental shifts in clause type. This finding extends the explanatory boundary of SFG into the realm of AI discourse analysis. It demonstrates that the Mood and modality system remains functionally operative even when instantiated through LLMs. The analysis thus provides empirical grounding for a grammatical theory of AI discourse. From a practical perspective, the results suggest that altering modality density or hedging frequency would have greater effects on perceived empathy and assertiveness than altering syntactic mood ratios. The analysis thus offers an applied pathway for adjusting AI

communication tone through controllable grammatical parameters

From an LLM design perspective, the findings highlight that modality control and hedging calibration offer effective levers for tuning the perceived tone, empathy, and credibility of large language models. Unlike structural retraining or prompt engineering, which often produce unstable results, adjusting these linguistic parameters can systematically influence users' perceptions of authoritativeness and warmth without compromising accuracy. Consequently, the results suggest that the next generation of model fine-tuning should include interpersonal calibration layers—modules that balance assertive and deferential expressions in accordance with context, task domain, and user profile.

Beyond technical design, the study carries broader implications for AI governance and policy. As language models increasingly participate in decision-support, education, and public communication, their linguistic style becomes a matter of public interest rather than mere stylistic choice. A model that overstates certainty may inadvertently amplify misinformation, while one that overuses hedging may appear evasive or unreliable. The evidence from this study underscores the need for transparent linguistic governance frameworks—guidelines specifying acceptable ranges of modality strength, stance-taking, and responsibility attribution in public-facing AI discourse. Such frameworks should ensure that AI-generated communication remains informative, ethically responsible, and aligned with human expectations of accountability.

Future research should extend this line of inquiry by conducting comparative analyses across different model architectures and fine-tuning regimes, investigating how specific training data and alignment techniques shape interpersonal realization patterns. Experimental studies on user perception can also clarify the psychological thresholds at which modality and

hedging variations translate into perceived trust or expertise. Ultimately, understanding and managing the interpersonal grammar of AI systems will be crucial to ensuring that large language models evolve as linguistically competent, ethically aligned, and socially beneficial communicative agents.

Overall, DeepSeek's linguistic profile exemplifies an advanced form of interpersonal adaptability: assertive in instructional contexts, tentative in reflective reasoning, and empathetic in support interactions. This equilibrium situates the model as a linguistically coherent and socially aware communicative partner. If its Mood and modality system can be further improved, the entire human-AI dialogue interface will achieve a qualitative breakthrough.

References

- [1]. Ballier, N. (2007). *Modality and stance in English noun complement clauses*. In F. Riegel & I. Schaeffer (Eds.), *Corpus studies in contrastive linguistics* (pp. 65–85). Peter Lang.
- [2]. Brown, P., & Levinson, S. C. (1987). *Politeness: Some universals in language usage*. Cambridge University Press.
- [3]. Eggins, S. (2004). *An introduction to systemic functional linguistics* (2nd ed.). Continuum.
- [4]. Fei, J., Wang, Y., & Liu, S. (2024). Dialogic alignment in human-AI interaction: A systemic functional perspective. *Discourse Studies*, 26(3), 367–389.
- [5]. Halliday, M. A. K. (1984). *Language as social semiotic: The social interpretation of language and meaning*. Edward Arnold.
- [6]. Halliday, M. A. K., & Matthiessen, C. M. I. M. (2014). *Halliday's introduction to functional grammar* (4th ed.). Routledge.
- [7]. Hauptmann, C., Liao, Y., & Xu, D. (2024). Negotiating agency in conversations with artificial intelligence: Pragmatic and interpersonal meanings. *Pragmatics and Society*, 15(2), 155–176.
- [8]. Hyland, K. (2005). *Metadiscourse: Exploring interaction in writing*. Continuum.
- [9]. Kanté, I. (2010). Mood and modality in finite noun complement clauses: A French-English contrastive study. *International Journal of Corpus Linguistics*, 15(2), 267–290.
- [10]. Lakoff, R. (1973). The logic of politeness; or, minding your p's and q's. In C. Corum, T. Cedric, & D. L. Morgan (Eds.), *Papers from the ninth regional meeting of the Chicago Linguistic Society* (pp. 292–305). Chicago Linguistic Society.
- [11]. Li, Q., & Zhang, H. (2024). Interpersonal meaning negotiation in ChatGPT dialogues: A corpus-based SFL study. *Linguistics and Education*, 83, 101298.
- [12]. Martin, J. R., & White, P. R. R. (2005). *The language of evaluation: Appraisal in English*. Palgrave Macmillan.
- [13]. Robbemdond, R., Rühlemann, C., & Gries, S. T. (2022). Comparing conversational style in human-human and human-AI talk: A corpus linguistic perspective. *Corpora*, 17(2), 145–167.
- [14]. Silk, A. (2018). Commitment and states of mind with mood and modality. *Philosophical Studies*, 175(9), 2197–2222.
- [15]. Thompson, G. (2014). *Introducing functional grammar* (3rd ed.). Routledge.