

UDC 81'33+81'27:004.8

DOI <https://doi.org/10.24919/2308-4863/102-29>**Olesia STOIKA,***orcid.org/0000-0002-7695-6100**Doctor of Pedagogical Sciences, Professor,
Professor at the Department of Applied Linguistics
Uzhhorod National University
(Uzhhorod, Ukraine),**Chief Research Fellow at the Department of Comparative Pedagogy
Institute of Pedagogy of the National Academy of Educational Sciences of Ukraine
(Kyiv, Ukraine) olesya.stoyka@uzhnu.edu.ua***Bohdana SAVKA,***orcid.org/0009-0002-3673-7699**2nd year master's student at the Faculty of International Economic Relations
Uzhhorod National University
(Uzhhorod, Ukraine) savka.bohdana@student.uzhnu.edu.ua*

COOPERATION WITHOUT A PARTNER: THE MIGRATION OF PRAGMATIC LABOUR IN USER-AI DISCOURSE

This article describes the communicative situation created by the spread of conversational interfaces built on large language models, a situation for which linguistic theory was not designed: users mitigate directives, thank the system and apologise for their own imprecision, although a language model holds no beliefs, assumes no obligations and possesses no face requiring protection. Pragmatic description of this behaviour is entirely successful, yet the frameworks themselves presuppose communicative intention on both sides.

The article examines what becomes of pragmatic competence once that presupposition is withdrawn. It argues that established frameworks retain their descriptive adequacy while surrendering their reciprocity: they no longer model an exchange between interlocutors but the sustained pragmatic work of a single participant, directed at an addressee that can neither endorse nor return it.

The theoretical groundwork is laid by J. Austin, J. Searle, H. P. Grice, S. Levinson, D. Sperber and D. Wilson, G. Leech, P. Brown and J. Habermas; the survival of social norms in machine-directed speech is explained by B. Reeves, C. Nass and Y. Moon, A. Gambino, J. Fox and R. Ratan; empirical refinement comes from L. Panfili, R. Bowman, Z. Yin and E. Miehling; the applied dimension from J. White and J. D. Zamfirescu-Pereira, the metapragmatic from C. Caffi and J. Verschuere, the sociolinguistic from P. Joshi, A. Radford, L. Bilaniuk and S. Sokolova.

The study employs critical analysis and synthesis of sources, conceptual analysis of pragmatic categories with respect to their reciprocity requirement, comparison of the status of the maxims in human dialogue and in user-AI discourse, and the modelling of a two-level account of the artificial addressee.

Established frameworks are shown to retain their descriptive adequacy while surrendering their reciprocity: interpretive labour migrates to the user, and the Gricean maxims are re-encoded from expectations into design specifications within the prompt. The contradiction between automatic and strategic politeness is resolved through a distinction between an agent model and a resource model of the system, and the term operational politeness is proposed. Assistant-directed code-switching is reclassified as compensation for asymmetric channel availability, so that speakers of thinly resourced languages carry an inferential and a channel burden at once. Prospects for further study lie in eliciting a corpus of machine-directed speech from Ukrainian users, testing the distinction between inferential and channel burden against languages of comparable resourcing status, and reconceiving prompt design within AI-literacy curricula as a linguistic rather than a technical skill.

Key words: *pragmatics, human-AI interaction, Gricean maxims, metapragmatic competence, operational politeness, low-resource languages, Ukrainian discourse.*

Олеся СТОЙКА,

orcid.org/0000-0002-7695-6100

*доктор педагогічних наук, професор,
професор кафедри прикладної лінгвістики
Ужгородського національного університету
(Ужгород, Україна),*

*головний науковий співробітник відділу порівняльної педагогіки
Інституту педагогіки Національної академії педагогічних наук України
(Київ, Україна) olesya.stoyka@uzhnu.edu.ua*

Богдана САВКА,

orcid.org/0009-0002-3673-7699

*студентка II курсу магістратури факультету міжнародних економічних відносин
Ужгородського національного університету
(Ужгород, Україна) savka.bohdana@student.uzhnu.edu.ua*

КООПЕРАЦІЯ БЕЗ ПАРТНЕРА: ЗМІЩЕННЯ ПРАГМАТИЧНОГО НАВАНТАЖЕННЯ В ДИСКУРСІ «КОРИСТУВАЧ – ШТУЧНИЙ ІНТЕЛЕКТ»

У статті охарактеризовано комунікативну ситуацію, яку створило поширення діалогових інтерфейсів на основі великих мовних моделей і для якої лінгвістичну теорію не було розроблено: користувачі пом'якшують спонукання, дякують системі та вибачаються за власну неточність, хоча мовна модель не має ані переконань, ані зобов'язань, ані «обличчя», яке потребувало б захисту. Прагматичний опис такої поведінки цілком успішний, проте самі теорії передбачають наявність комунікативної інтенції в обох учасників.

У статті досліджено, що відбувається з прагматичною компетенцією, коли це припущення знімається. Обґрунтовано, що усталені теоретичні моделі зберігають описову спроможність, але втрачають взаємність: вони описують уже не обмін між співрозмовниками, а безперервну прагматичну працю одного учасника, спрямовану на адресата, який не здатен ані підтвердити її, ані відповісти на неї.

Теоретичним підґрунтям слугують праці Дж. Остіна, Дж. Серля, Г. П. Грайса, С. Левінсона, Д. Спербера і Д. Вілсона, Дж. Ліча, П. Браун та Ю. Габермаса; перенесення соціальних норм на машину пояснили Б. Рівз, К. Насс і Ю. Мун, А. Гамбіно, Дж. Фокс і Р. Ратан; емпіричне уточнення – Л. Панфілі, Р. Боумен, З. Їнь та Е. Мілінг; прикладний вимір – Дж. Вайт і Дж. Д. Замфіреску-Перейра, метапрагматичний – К. Каффі та Й. Вершуерен, соціолінгвістичний – П. Држоші, А. Редфорд, Л. Біланюк і С. Соколова.

Застосовано критичний аналіз і синтез джерел, концептуальний аналіз прагматичних категорій на предмет вимоги взаємності, зіставлення статусу максим у людському діалозі та в дискурсі «користувач – ШІ», а також моделювання дворівневого опису штучного адресата.

Обґрунтовано, що усталені прагматичні моделі зберігають описову спроможність, але втрачають взаємність: інтерпретаційне навантаження переміщується на користувача, а максими Грайса перекодовуються з очікувань у проєктні вимоги запиту. Суперечність між автоматичною та стратегічною ввічливістю знято через розрізнення агентної та ресурсної моделей системи; запропоновано поняття операційної ввічливості. Перехід на іншу мову перекваліфіковано як компенсацію нерівної доступності каналу, через що носії мов з обмеженим ресурсом несуть подвійне навантаження – інференційне та каналне. Перспективи подальших досліджень пов'язані зі збиранням корпусу машиноспрямованого мовлення українських користувачів, перевіркою розрізнення інференційного та каналного навантаження на матеріалі мов зі зіставним ресурсним статусом і переосмисленням проєктування запиту в програмах ШІ-грамотності як лінгвістичної, а не суто технічної навички.

Ключові слова: прагматика, взаємодія «людина – ШІ», максими Грайса, метапрагматична компетенція, операційна ввічливість, мови з обмеженими ресурсами, український дискурс.

Problem statement. The spread of natural-language interfaces has produced a communicative situation for which linguistic theory was not designed. Users address language models and voice assistants in ordinary speech: they mitigate directives, thank the system for its answers and apologise for their own imprecision. Yet a language model holds no beliefs to be sincere about, assumes no obligations, possesses no face requiring protection and occupies no position from which uptake could be granted or withheld. Pragmatic description of this behaviour is nevertheless entirely

successful – implicature, deixis, indirectness and face-work are all observably present in machine-directed speech. The difficulty is that the frameworks producing that description assume both participants possess communicative intention. The unresolved issue is therefore not whether users extend pragmatic norms to machines, which empirical work has settled, but what becomes of pragmatic competence once the reciprocity in which it was theorised is withdrawn – and whether the competences it now demands are equally available to speakers of every language.

Analysis of recent research. Scholarship bearing on this problem has developed along four largely independent lines. The foundational pragmatic tradition supplies the conceptual apparatus: J. Austin (1962) established that utterances constitute actions assessed as felicitous rather than true; J. Searle (1976) systematised illocutionary force into five classes defined by illocutionary point, direction of fit and expressed psychological state; H. P. Grice (1975) formulated the Cooperative Principle and derived conversational implicature from the exploitation of its maxims; S. Levinson (1983) demarcated pragmatics as the domain of meaning remaining once truth conditions have been assigned; D. Sperber and D. Wilson (1995) replaced the code model of communication with an inferential account grounded in mutually manifest assumptions; G. Leech (1983) and P. Brown and S. Levinson (1987) accounted for the relational dimension through the Politeness Principle and face theory respectively. A second line, developed within human–computer interaction research, explains why these norms survive contact with machines. B. Reeves and C. Nass (1996) demonstrated that users respond to media much as they respond to social life; C. Nass and Y. Moon (2000) attributed this to ethopoeia and mindlessness, automatic social scripts running despite conscious awareness of their inapplicability; A. Gambino, J. Fox and R. Ratan (2020) have since argued that prolonged exposure to conversational systems produces distinct human–media scripts rather than borrowed human ones. A third line supplies empirical refinement: L. Panfili and colleagues (2021) documented asymmetric maxim application in voice-assistant use and proposed an additional Priority Maxim specific to human–computer exchange; R. Bowman and colleagues (2024) identified a non-linear relationship between mimicked politeness and user trust; Z. Yin and colleagues (2024) established that impolite prompting measurably degrades model output; E. Miehl and colleagues (2024) proposed a systematic reinterpretation of the maxims for human–AI communication. A fourth line, oriented toward practice, treats the user’s own linguistic choices: J. White and colleagues (2023) catalogued reusable prompt patterns, while J. D. Zamfirescu-Pereira and colleagues (2023) showed that non-expert users design prompts opportunistically and misjudge how far a formulation generalises. Finally, work on linguistic resourcing – P. Joshi and colleagues (2020), A. Radford and colleagues (2022) – together with L. Bilaniuk’s (2005) account of Ukrainian language practice and S. Sokolova’s (2022) analysis of language choice under conditions of unequal availability, supplies the material for the sociolinguistic dimension of the problem.

The aim of the article. The aim of this article is to determine what becomes of pragmatic competence

when it is exercised toward an addressee incapable of the intentions it presupposes. Four objectives follow: to identify where classical pragmatic frameworks require reciprocity and what survives its loss; to reconceptualise prompt design as metapragmatic competence; to resolve the contradiction between automatic and strategic politeness through a distinction between two models of the artificial addressee; and to establish that the resulting competence is unevenly available across languages.

Presentation of the main material. Pragmatics established itself as an autonomous discipline by treating meaning as something achieved between participants rather than encoded within sentences. That achievement is theorised, throughout the tradition, as bilateral. Austin assesses performatives not as true or false but as felicitous or defective, and the conditions of felicity include a sincerity condition binding the speaker to a psychological state together with an uptake requirement that places part of the act’s success beyond the speaker’s control (Austin, 1962: 55–73). Searle’s taxonomy rests on the same foundation: each of the five classes is defined by an illocutionary point, a direction of fit and a psychological state the speaker necessarily expresses (Searle, 1976: 344–369). Grice states the requirement most explicitly – a speaker means something only by intending to produce an effect through the audience’s recognition of that very intention, a recognition that must function as the hearer’s reason for responding rather than as its cause (Grice, 1975: 41–58).

Sperber and Wilson establish the same requirement from the opposite direction. Linguistic decoding, on their account, yields only a schematic representation serving as evidence about the speaker’s intention; the hearer must infer the remainder. Communication succeeds not through mutual knowledge, which would demand an infinite regress of recursive checks, but through mutual manifestness: a cognitive environment of assumptions perceptible or inferable to both parties (Sperber, Wilson, 1995: 39–43). This is the most demanding of the conditions surveyed here and the one a language model most conspicuously fails, since manifestness is a property of a mind for which something can be evident, and a system with no cognitive environment cannot share one. Habermas states the requirement socially rather than cognitively, grounding communicative action oriented toward mutual understanding in a lifeworld of implicit shared knowledge from which participants define their situation (Habermas, 1984: 285–286, 320). Whether stated as sincerity, intention recognition, manifestness or lifeworld, every framework surveyed here requires two minds.

Language models satisfy none of these conditions, and the point is architectural rather than empirical. A system manipulating symbols according to their formal relations possesses no semantics of its own; whatever intentionality it appears to display is contributed by those who design and interpret it (Searle, 1980: 417–424). Gubelmann extends this to contemporary systems: large models develop emergent internal representations that may reasonably be called cognitively autonomous, but lack the material and moral autonomy that agency requires (Gubelmann, 2024: 9–20). An AI system therefore does not perform speech acts. It produces sequences that users interpret as speech acts.

This yields the article’s starting position. Pragmatic frameworks applied to human–AI interaction do not forfeit their descriptive adequacy but their reciprocity. They continue to identify implicature, indirectness, face-work and deixis precisely, yet what they describe is no longer an exchange between interlocutors. It is the sustained pragmatic activity of one participant, maintained toward an addressee that can neither endorse it nor return it.

If the addressee registers nothing, the persistence of cooperative behaviour becomes the phenomenon requiring explanation, and two incompatible explanations are available. The first is cognitive. Reeves and Nass established that individuals respond to mediated entities much as to social ones, automatically rather than through deliberation (Reeves, Nass, 1996: 10–12, 22–25); Nass and Moon named the mechanism *ethopoeia*, since minimal human-like cues suffice to activate interpersonal scripts and no effective mechanism exists for suppressing them once triggered (Nass, Moon, 2000: 81–103). Epistemic vigilance sharpens the account: comprehension proceeds from a provisional stance of trust withdrawn only when active monitoring supplies grounds for doubt (Sperber et al., 2010: 367–369), and in machine-directed interaction that monitoring is frequently not engaged (Panfili et al., 2021).

The second explanation is instrumental. Yin and colleagues demonstrated across three languages that impolite prompting produces measurably worse output: length inflates, stereotypical bias increases, refusal rates rise (Yin et al., 2024: 9–15). Politeness before a language model is therefore not merely tolerated but rewarded mechanically, by a system trained on data in which cooperative framing correlates with cooperative response.

The two accounts sit uneasily together. The first holds that machine-directed politeness is automatic, unsuppressible and below awareness; the second that it is calculable and motivated by output quality. A single behaviour cannot easily be both a reflex and

a tactic. The literature has accumulated strong evidence for each without confronting the tension, which cannot be resolved at the level at which it arises. Its resolution requires a distinction developed below.

The first consequence of pragmatic asymmetry concerns the distribution of interpretive work. In human dialogue that work is shared: the speaker formulates efficiently in the expectation that the hearer will enrich, disambiguate and repair, and the hearer does so, returning signals that permit adjustment. When the addressee can do none of this, the whole of that work reverts to the single participant capable of performing it. We describe this as the migration of pragmatic labour.

Its signature is best captured through metapragmatics. Caffi defines metapragmatic competence as the region of speaker competence permitting control, planning and feedback upon an ongoing interaction by relating language to context (Caffi, 2009: 625–626); Verschueren describes the same capacity as continuous choice-making governed by variability, negotiability and adaptability (Verschueren, 1999: 55–60). In ordinary conversation this operates below awareness. Prompt design renders it fully explicit: the user knows the formulation selected will determine what returns, evaluates the return against the intention, and reformulates accordingly.

The clearest expression is what happens to the Gricean maxims. In human dialogue they are expectations the speaker relies upon and exploits through flouting. Before a system incapable of holding expectations they cannot function so, surviving instead as instructions the user must encode. Miehling and colleagues reach the same conclusion from the design side, arguing that Quality must be extended to require explicit reasoning and Relation refined to track a user’s evolving goals rather than the immediate query (Miehling et al., 2024: 14420–14422) – in both cases describing work that no longer occurs by default and must be stipulated. The maxims are not abandoned; they change grammatical mood, from what one may assume to what one must specify (see Table 1).

The technical literature has independently catalogued the strategies through which this encoding is performed without recognising them as pragmatic. White and colleagues describe reusable prompt patterns: assigning the system a persona, requiring it to propose improved formulations of the user’s question, obliging it to append its reasoning (White et al., 2023). Each is a codified metapragmatic operation – role framing calibrates register, constraint setting functions as metadiscursive regulation, appended reasoning institutes feedback the system cannot otherwise supply.

Users register the migration. Panfili and colleagues observed that when a voice assistant failed, users attributed the failure to their own insufficient specificity rather than to the technology (Panfili et al., 2021: 296) – self-attribution of communicative failure being a defining marker of metapragmatic awareness. Zamfirescu-Pereira and colleagues found conversely that non-experts design prompts opportunistically, expecting the system to recover implicit cues as a human would (Zamfirescu-Pereira et al., 2023). The competence is therefore not an extension of ordinary communicative ability but a distinct skill, unevenly held, that must be deliberately acquired.

This supplies the distinction required to resolve the difficulty deferred above. Users behave toward these systems as though they were partners and simultaneously design prompts as though they were instruments. Some research treats the second as a correction of the first, as though experienced users abandon social response. The evidence does not sup-

port that reading: politeness persists among expert users, and Gambino, Fox and Ratan report specialised human-media scripts developing rather than social scripts being extinguished (Gambino, Fox, Ratan, 2020: 71–85). We therefore propose that users hold two models of the artificial addressee concurrently, at different levels of the interaction (see Table 2).

The models do not conflict because they answer different questions. The agent model governs how the user speaks; the resource model governs what the user decides to do. A user may thank a system spontaneously and in the same session calculate that a particular formulation will yield a better result. The automatic and instrumental accounts of machine-directed politeness are therefore both correct and are not descriptions of the same thing: ethopoeia supplies the politeness, the resource model supplies the motive for not suppressing it.

This yields a further consequence for pragmatic theory. Leech separated the Politeness Principle from the

Table 1

Conversational maxims in human dialogue and in user–AI discourse

Maxim	Status in human dialogue	Status in user–AI discourse
Quantity	An expectation the hearer is assumed to observe; under-specification is repaired inferentially	A quantity specification the user must state in advance; nothing repairs an under-specified prompt but the user
Relation	A shared orientation to the accepted purpose of the exchange, continuously negotiated	A thematic boundary the user must delimit and re-delimit, as the system holds no model of the evolving purpose
Manner	A presumption of orderliness against which obscurity signals implicature	A formal requirement encoded in the prompt; obscurity signals nothing and yields degraded output
Quality	A sincerity commitment the speaker undertakes and the hearer provisionally trusts	An instruction the user issues – acknowledge uncertainty, supply grounds – since the system holds no beliefs to be sincere about

Source: compiled by the authors on the basis of Grice (1975) and Miehling et al. (2024).

Table 2

The dual model of the artificial addressee

Aspect	Agent model	Resource model
Level of operation	Spontaneous, automatic, below awareness	Deliberate, strategic, consciously monitored
Governing mechanism	Ethopoeia; inferential comprehension; social scripts	Metapragmatic planning; evaluation of output against intention
Observable behaviour	Politeness markers, indirect requests, gratitude, apology, hedging	Prompt structuring, persona assignment, constraint setting, reformulation
Treats the system as	A social entity capable of registering the exchange	A processing channel with determinate properties and limits
Attribution of failure	To the system, as uncooperative or evasive	To the user's own formulation, as insufficiently specified
Governs	The manner of the exchange	The strategy of the exchange

Source: developed by the authors.

Cooperative Principle on the grounds that the former maintains social relations at some cost to informational efficiency and the latter the reverse (Leech, 1983: 7–10). Under the resource model that separation does not collapse so much as become empty: the Politeness Principle has no work to perform before an addressee possessing no face to protect, and what survives is cooperative behaviour retaining politeness morphology while losing the motivation that gave the morphology its point. The term operational politeness is proposed for this configuration – face-work whose function is neither relational nor ritual but processual, directed at a system’s output behaviour rather than an interlocutor’s standing. The two principles are therefore not independent axes of communicative motivation in general; their independence holds only where the addressee’s response is socially rather than formally determined.

The dual model also accounts for a finding otherwise anomalous. Bowman and colleagues found that systems mimicking positive politeness were perceived as caring until users recognised the capacity for care was absent, whereupon trust declined; excessive negative politeness was experienced as condescending by users who had already elected to interact (Bowman et al., 2024). This is not a paradox but a boundary being drawn: the agent model responds to the politeness, the resource model evaluates it against what the system is known to be and withdraws the response when the two diverge too far.

The resource model has one further and more consequential application. If the interpretive burden rests on the user, the resources available to that user determine how far it can be carried – and those resources are distributed structurally rather than accidentally. Only a small proportion of the world’s languages is represented in current language technologies, and systems described as language-agnostic are typically built on a narrow, typologically related set; Ukrainian occupies an intermediate position, its substantial informal web presence permitting gains from unsupervised pre-training while the annotated collections supervised speech systems require remain thin (Joshi et al., 2020: 6282–6284). The consequence is measurable: in one widely used recognition system Ukrainian was allocated fewer than seven hundred training hours against more than four hundred thousand for English, with word error rates rising accordingly and the gap widening in the smaller models cheap enough for on-device deployment (Radford et al., 2022). This imbalance is not linguistically neutral – the volume of Russian-language speech data available reflects an imposed historical language policy rather than any property of the languages concerned (Bilaniuk, 2005: 13–15).

Nor is the difficulty confined to accuracy. Voice interfaces are monolingual by architecture: language identification occurs early and commits the system to a single model, so items embedded from another language cannot be decoded. Yet the everyday speech such systems are meant to receive is, in Ukraine, routinely mixed. Bilaniuk documents switching as a contextualising device and a group style rather than a symptom of deficient competence, and describes bilingual exchanges in which each participant maintains a different language without breakdown of understanding (Bilaniuk, 2005: 121–123, 158, 173–176). A monolingual interface therefore does not merely inconvenience such speakers; it withdraws a resource they use as a matter of course.

How should the resulting shift of language be classified? The obvious candidate is accommodation: speakers converge toward an interlocutor’s way of speaking, and such convergence characteristically flows toward whoever holds power. We argue that this does not survive contact with the machine case, and that its failure is informative. Accommodation requires an interlocutor possessing a language of their own, a partner whose preference one moves toward. A recognition system holds no preference: it does not favour English, it processes English better. The asymmetry is one of capability, not orientation. This is precisely the distinction the dual model makes available – accommodation belongs to the agent model, which governs manner, whereas selecting a code the system can receive at all belongs to the resource model, which governs strategy. What the machine case shares is not accommodation but asymmetric availability. Sokolova describes the abandonment of one’s habitual language because the information required exists only in another, reporting Ukrainian-language internet use falling well below Ukrainian use in face-to-face settings precisely because the Ukrainian-language segment is thin (Sokolova, 2022). The structure is identical. We therefore reclassify assistant-directed code-switching as compensation for an unevenly resourced channel rather than accommodation to an addressee.

Two consequences follow. The first is that each such switch is subtractive. A bilingual speaker’s code choice is constrained not by what they can produce but by what the addressee can receive; where the addressee cannot receive Ukrainian, Ukrainian falls outside the repertoire the exchange can draw on. The cost is not only quantitative: moving into an in-group language claims closeness and shared ground, whereas moving toward an external, higher-status code carries distance (Brown, Levinson, 1987: 110–111), so an interface that will not accept Ukrainian sets the exchange at a default remove.

The second is where the article's claims converge. Pragmatic labour migrates to the user, but it does not migrate equally. A user operating in a well-resourced language bears an inferential burden: the whole work of specification, disambiguation and repair the system cannot perform. A user operating in a thinly resourced language bears that burden and, in addition, a channel burden – the prior work of securing a code the system can receive at all. The two compound. Pragmatic labour in human–AI discourse is therefore not merely displaced onto one participant but distributed unevenly according to a hierarchy of linguistic resourcing that is itself the sediment of historical language policy – a finding invisible to any research programme conducted exclusively in English.

Conclusions. The analysis presented here supports a single organising claim: pragmatic frameworks applied to human–AI interaction retain their descriptive power while surrendering their reciprocity. Every framework surveyed above – from Austin's felicity conditions to Habermas's lifeworld – presupposes a partner capable of holding beliefs, assuming obligations and supplying uptake. Language models supply none of these, yet they activate the entire apparatus those constructs describe. What the frameworks now model is not an exchange but the sustained pragmatic work of one participant.

Three findings follow. Interpretive labour migrates wholly to the user, converting conversational maxims from expectations into design specifications and automatic competence into deliberate metapragmatic work. The contradiction between automatic and strategic politeness dissolves once users are recognised to

sustain an agent model governing manner alongside a resource model governing strategy, under which the Politeness Principle loses its object. And the resource model permits assistant-directed code-switching to be reclassified as compensation for asymmetric channel availability rather than accommodation to an addressee, revealing that speakers of thinly resourced languages carry an inferential and a channel burden at once.

The account defines its own principal limitation. It anticipates that politeness, indirectness, anthropomorphising address and code choice will vary with the user's working model of the system, but specifies neither the direction nor the magnitude of that variation, and no usable corpus of Ukrainian machine-directed interaction exists against which the prediction might be tested.

Further research should accordingly proceed in three directions. Empirically, controlled elicitation of machine-directed speech from Ukrainian users would establish how far the two models govern observable behaviour. Theoretically, the distinction between inferential and channel burden should be tested against other languages occupying comparable positions in the resourcing hierarchy. In applied terms, the recognition of prompt design as metapragmatic competence bears directly on AI literacy curricula, which currently treat it as a technical rather than a linguistic skill. That conversational systems will not receive mixed input is, in the end, a design decision rather than a technical limit. The pragmatic asymmetry described here is structural; the inequality layered upon it is not.

BIBLIOGRAPHY

1. Austin J. L. *How to do things with words*. Oxford : Oxford University Press, 1962. URL: <https://silverbronzo.wordpress.com/wp-content/uploads/2017/10/austin-how-to-do-things-with-words-1962.pdf>.
2. Bilaniuk L. *Contested tongues: Language politics and cultural correction in Ukraine*. Ithaca : Cornell University Press, 2005. URL: <https://diasporiana.org.ua/wp-content/uploads/books/27428/file.pdf>.
3. Bowman R., Cooney O., Newbold J. W., Thieme A., Clark L., Doherty G., Cowan B. Exploring how politeness impacts the user experience of chatbots for mental health support. *International Journal of Human-Computer Studies*. 2024. Vol. 184. Article 103181. DOI: <https://doi.org/10.1016/j.ijhcs.2023.103181>.
4. Brown P., Levinson S. C. *Politeness: Some universals in language usage*. Cambridge : Cambridge University Press, 1987. URL: https://www.academia.edu/26395652/Politeness_Some_universals_in_language_usage.
5. Caffi C. *Metapragmatics*. *Concise encyclopedia of pragmatics* / ed. by J. L. Mey. 2nd ed. Amsterdam : Elsevier, 2009. P. 625–626.
6. Gambino A., Fox J., Ratan R. A. Building a stronger CASA: Extending the computers are social actors paradigm. *Human-Machine Communication*. 2020. Vol. 1. P. 71–85. DOI: <https://doi.org/10.30658/hmc.1.5>.
7. Grice H. P. *Logic and conversation. Syntax and semantics*. Vol. 3: *Speech acts* / ed. by P. Cole, J. L. Morgan. New York : Academic Press, 1975. P. 41–58. URL: <https://thegricean.github.io/xprag-ling247/readings/grice1975.pdf>.
8. Gubelmann R. Large language models, agency, and why speech acts are beyond them (for now): A Kantian-cum-pragmatist case. *Philosophy & Technology*. 2024. P. 9–20. DOI: <https://doi.org/10.1007/s13347-024-00696-1>.
9. Habermas J. *The theory of communicative action*. Vol. 1: *Reason and the rationalization of society* / transl. by T. McCarthy. Boston : Beacon Press, 1984. URL: <https://teddykw2.wordpress.com/wp-content/uploads/2012/07/jurgen-habermas-theory-of-communicative-action-volume-1.pdf>.
10. Joshi P., Santy S., Budhiraja A., Bali K., Choudhury M. The state and fate of linguistic diversity and inclusion in the NLP world. *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*. 2020. P. 6282–6284. DOI: <https://doi.org/10.18653/v1/2020.acl-main.560>.

11. Leech G. N. *Principles of pragmatics*. London : Longman, 1983. DOI: <https://doi.org/10.4324/9781315835976>.
12. Levinson S. C. *Pragmatics*. Cambridge : Cambridge University Press, 1983. DOI: <https://doi.org/10.1017/CBO9780511813313>.
13. Miehling E., Nagireddy M., Sattigeri P., Daly E., Piorkowski D., Richards J. Language models in dialogue: Conversational maxims for human-AI interactions. Findings of the Association for Computational Linguistics: EMNLP 2024. 2024. P. 14420–14422. DOI: <https://doi.org/10.18653/v1/2024.findings-emnlp.843>.
14. Nass C., Moon Y. Machines and mindlessness: Social responses to computers. *Journal of Social Issues*. 2000. Vol. 56, № 1. P. 81–103. DOI: <https://doi.org/10.1111/0022-4537.00153>.
15. Panfili L., Duman S., Nave A., Ridgeway K. P., Eversole N., Sarikaya R. Human–AI interactions through a Gricean lens. *Proceedings of the Linguistic Society of America*. 2021. Vol. 6, № 1. P. 288–302. DOI: <https://doi.org/10.3765/plsa.v6i1.4971>.
16. Radford A., Kim J. W., Xu T., Brockman G., McLeavey C., Sutskever I. Robust speech recognition via large-scale weak supervision. *arXiv*. 2022. DOI: <https://doi.org/10.48550/arXiv.2212.04356>.
17. Reeves B., Nass C. *The media equation: How people treat computers, television, and new media like real people and places*. Stanford : CSLI Publications, 1996. URL: <https://www.researchgate.net/publication/37705092>.
18. Searle J. R. *A taxonomy of illocutionary acts. Language, mind, and knowledge* / ed. by K. Gunderson. Minneapolis : University of Minnesota Press, 1976. P. 344–369. URL: <https://hdl.handle.net/11299/185220>.
19. Searle J. R. Minds, brains, and programs. *Behavioral and Brain Sciences*. 1980. Vol. 3, № 3. P. 417–424. DOI: <https://doi.org/10.1017/S0140525X00005756>.
20. Sokolova S. The problem of choosing the language of communication. *Cognitive Studies / Études cognitives*. 2022. № 22. Article 2649. DOI: <https://doi.org/10.11649/cs.2649>.
21. Sperber D., Clément F., Heintz C., Mascaro O., Mercier H., Origgi G., Wilson D. Epistemic vigilance. *Mind & Language*. 2010. Vol. 25, № 4. P. 359–393. DOI: <https://doi.org/10.1111/j.1468-0017.2010.01394.x>.
22. Sperber D., Wilson D. *Relevance: Communication and cognition*. 2nd ed. Oxford : Blackwell, 1995. URL: https://monoskop.org/images/e/e6/Sperber_Dan_Wilson_Deirdre_Relevance_Communication_and_Cognition_2nd_edition_1996.pdf.
23. Verschueren J. *Understanding pragmatics*. London : Edward Arnold, 1999. URL: https://www.researchgate.net/publication/282016316_Understanding_Pragmatics.
24. White J., Fu Q., Hays S., Sandborn M., Olea C., Gilbert H., Elnashar A., Spencer-Smith J., Schmidt D. C. A prompt pattern catalog to enhance prompt engineering with ChatGPT. *arXiv*. 2023. DOI: <https://doi.org/10.48550/arXiv.2302.11382>.
25. Yin Z., Wang H., Horio K., Kawahara D., Sekine S. Should we respect LLMs? A cross-lingual study on the influence of prompt politeness on LLM performance. *Proceedings of the Second Workshop on Social Influence in Conversations (SICON 2024)*. 2024. P. 9–35. DOI: <https://doi.org/10.18653/v1/2024.sicon-1.2>.
26. Zamfirescu-Pereira J. D., Wong R. Y., Hartmann B., Yang Q. Why Johnny can't prompt: How non-AI experts try (and fail) to design LLM prompts. *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*. New York : Association for Computing Machinery, 2023. P. 1–21. DOI: <https://doi.org/10.1145/3544548.3581388>.

REFERENCES

1. Austin, J. L. (1962). *How to do things with words* (pp. 55–73). Oxford University Press. <https://silverbronzo.wordpress.com/wp-content/uploads/2017/10/austin-how-to-do-things-with-words-1962.pdf>
2. Bilaniuk, L. (2005). *Contested tongues: Language politics and cultural correction in Ukraine* (pp. 13–15, 121–123, 158, 173–176). Cornell University Press. <https://diasporiana.org.ua/wp-content/uploads/books/27428/file.pdf>
3. Bowman, R., Cooney, O., Newbold, J. W., Thieme, A., Clark, L., Doherty, G., & Cowan, B. (2024). Exploring how politeness impacts the user experience of chatbots for mental health support. *International Journal of Human-Computer Studies*, 184, Article 103181. <https://doi.org/10.1016/j.ijhcs.2023.103181>
4. Brown, P., & Levinson, S. C. (1987). *Politeness: Some universals in language usage* (pp. 110–111). Cambridge University Press. https://www.academia.edu/26395652/Politeness_Some_universals_in_language_usage
5. Caffi, C. (2009). Metapragmatics. In J. L. Mey (Ed.), *Concise encyclopedia of pragmatics* (2nd ed., pp. 625–626). Elsevier.
6. Gambino, A., Fox, J., & Ratan, R. A. (2020). Building a stronger CASA: Extending the computers are social actors paradigm. *Human-Machine Communication*, 1, 71–85. <https://doi.org/10.30658/hmc.1.5>
7. 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. <https://thegricean.github.io/xprag-ling247/readings/grice1975.pdf>
8. Gubelmann, R. (2024). Large language models, agency, and why speech acts are beyond them (for now): A Kantian-cum-pragmatist case. *Philosophy & Technology*, 9–20. <https://doi.org/10.1007/s13347-024-00696-1>
9. Habermas, J. (1984). *The theory of communicative action: Vol. 1. Reason and the rationalization of society* (T. McCarthy, Trans.; pp. 285–286, 320). Beacon Press. <https://teddykw2.wordpress.com/wp-content/uploads/2012/07/jurgen-habermas-theory-of-communicative-action-volume-1.pdf>
10. Joshi, P., Santy, S., Budhiraja, A., Bali, K., & Choudhury, M. (2020). The state and fate of linguistic diversity and inclusion in the NLP world. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics* (pp. 6282–6284). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2020.acl-main.560>
11. Leech, G. N. (1983). *Principles of pragmatics* (pp. x, 7–10). Longman. <https://doi.org/10.4324/9781315835976>
12. Levinson, S. C. (1983). *Pragmatics* (pp. 1–18, 61–63). Cambridge University Press. <https://doi.org/10.1017/CBO9780511813313>

13. Miehl, E., Nagireddy, M., Sattigeri, P., Daly, E., Piorkowski, D., & Richards, J. (2024). Language models in dialogue: Conversational maxims for human-AI interactions. In *Findings of the Association for Computational Linguistics: EMNLP 2024* (pp. 14420–14422). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2024.findings-emnlp.843>
14. Nass, C., & Moon, Y. (2000). Machines and mindlessness: Social responses to computers. *Journal of Social Issues*, 56(1), 81–103. <https://doi.org/10.1111/0022-4537.00153>
15. Panfili, L., Duman, S., Nave, A., Ridgeway, K. P., Eversole, N., & Sarikaya, R. (2021). Human–AI interactions through a Gricean lens. *Proceedings of the Linguistic Society of America*, 6(1), 288–302. <https://doi.org/10.3765/plsa.v6i1.4971>
16. Radford, A., Kim, J. W., Xu, T., Brockman, G., McLeavey, C., & Sutskever, I. (2022). *Robust speech recognition via large-scale weak supervision*. arXiv. <https://doi.org/10.48550/arXiv.2212.04356>
17. Reeves, B., & Nass, C. (1996). *The media equation: How people treat computers, television, and new media like real people and places* (pp. 10–12, 22–25). CSLI Publications. <https://www.researchgate.net/publication/37705092>
18. Searle, J. R. (1976). A taxonomy of illocutionary acts. In K. Gunderson (Ed.), *Language, mind, and knowledge* (pp. 344–369). University of Minnesota Press. <https://hdl.handle.net/11299/185220>
19. Searle, J. R. (1980). Minds, brains, and programs. *Behavioral and Brain Sciences*, 3(3), 417–424. <https://doi.org/10.1017/S0140525X00005756>
20. Sokolova, S. (2022). The problem of choosing the language of communication. *Cognitive Studies / Études cognitives*, 22, Article 2649. <https://doi.org/10.11649/cs.2649>
21. Sperber, D., Clément, F., Heintz, C., Mascaro, O., Mercier, H., Origgi, G., & Wilson, D. (2010). Epistemic vigilance. *Mind & Language*, 25(4), 359–393. <https://doi.org/10.1111/j.1468-0017.2010.01394.x>
22. Sperber, D., & Wilson, D. (1995). *Relevance: Communication and cognition* (2nd ed., pp. 39–43, 244–247). Blackwell. https://monoskop.org/images/e/e6/Sperber_Dan_Wilson_Deirdre_Relevance_Communication_and_Cognition_2nd_edition_1996.pdf
23. Verschueren, J. (1999). *Understanding pragmatics* (pp. 55–60). Edward Arnold. https://www.researchgate.net/publication/282016316_Understanding_Pragmatics
24. White, J., Fu, Q., Hays, S., Sandborn, M., Olea, C., Gilbert, H., Elnashar, A., Spencer-Smith, J., & Schmidt, D. C. (2023). *A prompt pattern catalog to enhance prompt engineering with ChatGPT*. arXiv. <https://doi.org/10.48550/arXiv.2302.11382>
25. Yin, Z., Wang, H., Horio, K., Kawahara, D., & Sekine, S. (2024). Should we respect LLMs? A cross-lingual study on the influence of prompt politeness on LLM performance. In *Proceedings of the Second Workshop on Social Influence in Conversations (SICoN 2024)* (pp. 9–35). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2024.sicon-1.2>
26. Zamfirescu-Pereira, J. D., Wong, R. Y., Hartmann, B., & Yang, Q. (2023). Why Johnny can't prompt: How non-AI experts try (and fail) to design LLM prompts. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems* (pp. 1–21). Association for Computing Machinery. <https://doi.org/10.1145/3544548.3581388>

Дата першого надходження статті до видання: 14.08.2026

Дата прийняття статті до друку після рецензування: 07.09.2026

Дата публікації (оприлюднення) статті: 25.09.2026

Стаття поширюється на умовах
ліцензії відкритого доступу (CC BY 4.0)

