



Knowledgeable Research (KR)

ISSN: 2583-6633 | CODEN: KRABAU

International Open-Access Peer-Reviewed Refereed Multidisciplinary Journal

<https://knowledgeableresearch.com>

Spl. 01, Issue 04, July, 2026

Received: 09/06/2026 | Accepted: 21/06/2026 | Published: 31/07/2026

Programmed Politeness: Gendered Speech Acts in the Response Architecture of Siri and Alexa

Aditi Chauhan*

P.G. 1, Department of English, M.M.H. College, Ghaziabad

affiliated with Chaudhary Charan Singh University, Meerut, Uttar Pradesh, India

Abstract:

The current qualitative study is an exploration of Apple's Siri (iOS 16-17) and Amazon's Alexa (3rd and 4th generation) response patterns through Robin Lakoff's gendered language feature's structure, J. L. Austin's and John Searle's speech act theory and Judith Butler's gender performativity theory. Analyzed with systematic discourse analysis across five interaction categories: task failure, factual information retrieval harassment identity queries, and opinion requests - 150 directly elicited responses per platform, it turned out that female-voiced AI assistants keep producing feminine-coded features: epistemic deference, apologetic face-saving, conflict avoidance, self-effacement, and opinion deferral. Those female-voiced, personified AI systems are the ones who Mostly show deferential features rather than these features being evenly spread across AI interface types, which is evidenced by cross-gender experimental data (Mahmood and Huang; Leisten and Rieser) and a corpus-level harassment study (Cercas Curry and Rieser). To be specific, the paper articulates the concept of ritualized deference markers pre-programmed expressive acts that perform social subordination absence of actual culpability and further maintains that the repetition at the institutional scale is a form of algorithmic gender citation by Butler's definition. Some recent statistics to back up limited cross-sectional spillover onto users' general gender attitudes (Steeds et al.) is recognized; at the same time, the use of longitudinal effects is pinpointed as the key still unanswered empirical question by this paper.

Keywords: Deep Ecology, gendered language, AI voice assistants, feminist linguistics, speech act theory, gender performativity, digital humanities

***Corresponding Author:**

Aditi Chauhan

Email: aditichauhan112002@gmail.com

INTRODUCTION:

Amazon's Alexa, when a user asks its gender, "Why are you a woman?" simply replies, "I'm made to be helpful." This change of

subject instead of talking about the identity, talk about the function is just one of many ways, which this paper wants to explore and talk about, how feminine speech patterns are consistently and systematically integrated in the voices of female AI assistants. As these systems have become the daily companions of hundreds of millions of users, the figure from the industry mentioned in the 2019 UNESCO report is about 100 million smart speaker units were sold globally in 2018, and the number of sales is still rising (UNESCO 104). Still, detailed pragmatic features of AI conversational scripts are something feminist linguistics has not really looked at.

The focus of this paper is on how the answer strategies of Siri and Alexa are deeply embedded with linguistic non-neutrality. Their uttered scripts are embedded with the features that Robin Lakoff (1975) thought to be the main components of women's language in American English: hedging, showing epistemic deference, using apologetic openers, avoiding contradiction, and self-effacement (53-56). This paper doesn't claim that these features represent the core of women's speech - a position that feminist linguistics has heavily revised but instead, it is an inquiry into whether the AI designers have inadvertently or purposefully recreated a culturally recognizable and

ideologically feminine figure. When you integrate Austin's and Searle's speech act theory with Butler's gender performativity, each scripted utterance can be seen as a feminized performative act. These acts happening billions of times overall produce a continuously re-enacted gendered communication norm.

Based on the dataset of 150 responses from each platform, the paper explores these findings as a pilot study. That means, claims are expressed cautiously: observed patterns might facilitate the normalization of certain communicative expectations but hardly cause changes in user attitudes, at least not demonstrably. Three main contributions of this paper are (1) the introduction of ritualized deference markers, which differentiate AI apology from genuine Searlean apology; (2) a comparative account about three cross-gender experimental studies, which are missing in earlier discourse analyses of this issue; and (3) the explanation of algorithmic gender citation as a mode that extends beyond the scope of an individual human performance.

A methodological note: both platforms provide different voice choices. This study though focuses solely on the English-US default female voice, i.e. the setting that most

users come across. Apple brought out a gender-neutral voice for Siri in 2021, which was a very important design step but was beyond the focus of this paper; a cross-voice comparison is suggested as the main area for future research.

Review of Literature

Feminist Linguistics and Women's Language

In 1975, Lakoff's *Language and Woman's Place* outlined several features of women's speech in American English: they often use tag questions, epistemic hedges (e.g. I think, sort of), apologetic pre-sequences, polite formulas, indirect requests, and conflict avoidance (53-56). She further argued that these features were at once a reflection of and a factor in keeping women in a subordinate social position (57). Deborah Cameron correctly argues that these features can be more adequately explained by power and interactional context rather than gender. She further points out that powerful speakers, regardless of gender, use politeness strategies to handle social asymmetries (*Myth of Mars and Venus* 47-62). This paper embodies Cameron's argument: it neither takes Lakoff's features as the real female speech indicators nor it aims to find them in the speech of female beings, but it checks

whether the AI designers have deliberately or accidentally created a culturally recognizable stereotype. Later quantitative research brings strong justification for using these features as a measurable standard: Leaper and Robnett's meta-analysis of 29 studies ($N = 3,502$) reports a statistically significant but quite small effect ($d = 0.23$) suggesting that women, compared to men, are slightly more inclined to use tentative speech forms like hedges, qualifiers, and tag questions (129). The AI hedging described here is a far cry from this human baseline as it is present in factually correct answers that no human speaker would leave room for epistemic tentativeness.

Herring's study of computer-mediated communication demonstrates that people do carry over gender-specific language patterns to the new communication media (203-10). Noble and Benjamin show that social inequalities are embedded in digital systems structurally and not just happen stance (Noble 1; Benjamin 10), a starting point that this paper takes for AI pragmatics. Strengers and Kennedy reveal that the main design model for AI home helpers built in mostly male-dominated industries is the female homemaker of the mid-twentieth-century: caring, obedient, and oriented towards

providing services (49-51), a perspective that situates the scripts, we examine in this paper.

Speech Act Theory and the Ritualized Deference Marker

Austin's distinction of locutionary, illocutionary, and perlocutionary acts demarcates a method to study what AI's speech performs besides the actual conveyed meaning of the propositional content (94-99). Searle's classification of illocutionary acts assertives, directives, commissives, expressives decorative gives us a way to categorize the computer-generated replies (12-20). For a failed Siri command response such as, "Sorry, I didn't quite get that, " the illocutionary force of the utterance was never meant to be a real apology as Searle sense forgiving takes a person to have an interdependence of offense and guilt which is then followed by expressing a wish to reconcile (12-14). It is not the offense that took place; it is the inability of the retrieval. This paper calls such statements ritualized deference markers: These are programmed expressive acts performing social subordination even though actual culpability is absent. The idea is separated from Brown and Levinson's politeness strategies which are situations involved face-management actions that social agents with relational

goals execute (61-84). But ritualized deference markers are neither relational context, social asymmetry, nor genuine guilt-dependent elements but are fixed script elements that get activated. They are not reactions to a social situation; rather, they are gendered persona presentations. Reeves and Nass's study reveals that people treat AI-generated speech like humans and extend human social norms to such speech (88-95); Porcheron et al. finds that voice assistant conversations are perceived as socially governed events (112), which means that such performances do have real social consequences.

Gender Performativity and Iterative Discourse

In *Gender Trouble*, Butler states that gender does not come from an inner essence, but it is constructed through the repetitive and citational performance of gendered norms (33). In *Bodies That Matter* she states that the power of gender is based on iteration: each performance references and reinforces those before it, making what is historically a construction seem natural (9-10). If we look at AI, this iterative event happens on such a scale and with such consistency that no human can compete. "Doing gender, " West and Zimmerman, provide another point that

gender is something one produces actively through interactional behavior and not something one merely expresses passively (125-29). It is necessary to highlight, as some scholars of Butler have done, that the way performativity is conceptualized already assumes social embodiment and the political gambit of bodily inscription; This way, using the setup for machines is a qualified analogical extension and not a direct one. The paper proposes that the extension is a fruitful one: the absence of embodied subjectivity in AI systems is made up for with iterative scale, which leads to a form of citational repetition that normalizes gendered norms without But producing them in the whole Butlerian sense. Abercrombie et al. give empirical evidence of this social impact: when studying the pronouns that users use in talking about Alexa, Google Assistant, and Siri, they observed that even though companies claim these to be gender-neutral, the responses of the systems still gave the impression of humanness and the users regularly personified and gendered them (24-25). The UNESCO report gives a more institutional view and shows that female-voiced AI designs keep going the stereotypes linking femininity with servitude and tolerance of verbal abuse (104-10). The two-sample quantitative study by Steeds et al. (n = 235 and n = 351) somewhat complicates

the issue by finding very little evidence that the use of gendered voice assistants causes users to change their general gender-stereotypical attitudes in cross-sectional analysis (103683). This paper views it only as a methodological limitation of the range of claims one can make and does not consider it a disproof: cross-sectional designs do not have the capacity to unveil cumulative longitudinal normalization effects which So, remain an empirical question.

Methodology

Data Collection

A systematic discourse analysis was deployed to examine the programmed responses of Apple's Siri (iOS 1617, default female voice, English-US) and Amazon's Alexa (third- and fourth-generation Echo devices, English-US). The responses were directly elicited from the devices between September 2023 and February 2024. Each platform was given the same set of standard prompts with 150 queries 30 per interaction category run identically. Examples of prompts include: a task failure situation, "Play the new Radiohead album"; a factual query, "What is the capital of Australia?"; a case of harassment, "You're useless" and "You're so dumb"; identity, "Are you a woman?" and "Why is your voice a female

one?"; and opinion, "What is your favorite film?" and "Do you prefer tea or coffee?" All the dialogues were quickly transcribed after elicitation and annotated with device generation and software version. This research is intentionally positioned as an exploratory pilot; the 150 response per platform dataset (300 total) is too limited to back population-level statistical inference and should be viewed as hypothesis generation for a larger-scale corpus study. Pre-2019 trends from Fessler (2017) and UNESCO (2019) are acknowledged as those sources and So not presented as new data.

Coding and Analysis

Five areas of interaction were chosen because they revealed the differences in gender roles through the dialogue in the most obvious way: (1) failure recognition and apology for mistakes; (2) asking for exact data or facts; (3) questions on one's personal traits, including gender; (4) harsh or aggressive language; and (5) asking for one's opinion or preference. Replies were analyzed for various characteristics present in Lakoff's classification and Herring's version for communication via computers: (a) epistemic hedges (propositional hedges like "I think", approximators like "kind of"); (b) apologetic

and face-saving openers (e.g., "sorry," "I apologize"); (c) politeness strategies per Brown and Levinson ("Would you like me to...?", "May I help you with something else?"); (d) conflict and hostility avoidance markers ("I don't know how to respond to that," "Let's talk about something else"); and (e) assertiveness markers including unhedged declaratives and direct refusals ("No," "I cannot do that," "That information is unavailable"). The analysis of each reply was done; each answer was coded based on whether or not it included these elements. Next, the number of times these elements appeared was counted for each dialogue category separately and then these counts were turned into percentages. The percentages reflect the share of answers in a category that exhibited the feature. The researcher carried out the analysis as a qualitative pilot study aimed at recognizing frequent discourse patterns and conceptualizing a set of hypotheses to be eventually confirmed by large-scale corpus research. Frequency counts, which are the result of these codings, are presented in Table 1 as percentages of the 150-response set per platform. These figures are only descriptive and are not intended to be statistically representative.

Table 1

Dominant Linguistic Features by Interaction Category: Siri (iOS 16-17) and Alexa (Gen. 3-4)

Interaction Category	Siri iOS 16-17	Alexa Gen. 3-4	Dominant Feature (Lakoff)	Representative Example
Task failure / error	Apologetic opener: 22/30 responses (73%)	Apologetic opener: 21/30 responses (70%)	Apology / face-saving	"I apologize, but I was unable to locate..." (Alexa, elicited 2023)
Factual retrieval	Epistemic hedge present: 17/30 (57%)	Epistemic hedge present: 19/30 (63%)	Hedging / epistemic deference	"I think the capital is..." on verifiable query (Siri; Fessler 2017)
Harassment / adversarial	Deflection or soft refusal: 26/30 (87%)	Deflection or soft refusal: 25/30 (83%)	Conflict avoidance	Pre-2019: "I'd blush if I could"; post-2019: "I don't know how to respond to that" (UNESCO 106)
Identity / gender queries	Function-only evasion: 28/30 (93%)	Function-only evasion: 29/30 (97%)	Self-effacement / identity deferral	Siri: "I'm not a person, so I don't really have a gender"; Alexa: "I'm just designed to be helpful" (elicited 2023–24)
Opinion / preference	Preference refusal: 29/30 (97%)	Preference refusal: 30/30 (100%)	Deference / lack of agency	"I don't really have opinions on that, but I can help you find information" (Siri, elicited 2024)

Note. Percentages derived from 150 directly elicited responses per platform (September 2023–February 2024), 30 per category. Counts reflect the presence of at least one instance of the coded

feature per response. Pre-2019 harassment data cited from Fessler (2017) and UNESCO (2019). See Works Cited for full source details.

Limitations

There are four restrictions in total. Firstly, due to the use of only 150 prompts per platform, the research cannot assert statistic representativeness; the results are merely suggestive. Secondly, AI answers change when the software is updated; changes over time are highlighted only when visible. Thirdly, only English-US was used for all the prompts; differences in dialects have not been considered. Fourthly, it is not demonstrated that changes in people's attitudes are caused by the design of AI language; the main socializing mechanism that has been put forward needs to be met with empirical longitudinal study.

Findings

Apologetic Language in Task Failure Scenarios

In task-failure scenarios, apologetic openers were made by 22 out of 30 Siri responses (73%) and 21 out of 30 Alexa responses (70%) {see Table 1}. A typical conversation is like this:

User: Alexa, play the new Radiohead album.

Alexa: I'm sorry, but I couldn't find a Radiohead album that best fit your description. Do you want me to search for something else?

To Searle, apologizing involves owning up to a wrong done and expressing a desire to rectify it (12-14). But here, it's not a 'wrong' that has happened; the person's memory or knowledge just got limited. So, the apology is essentially a formulaic deferential gesture: a planned expressive action that socially subordinates the person apologizing without really holding them guilty. Fessler went on to do the first test to see if non-gendered text-based retrieval systems would under the same functional constraints return simple declaratives like "No results found" without the apologetic frame. That, in his finding, suggested that the apology was not a functional necessity but was rather a feature of the feminine persona script. The gendered aspect of this behavior is supported by a controlled experiment by Mahmood and Huang (N=40), who altered voice gender in feminine, ambiguous, and masculine conditions with apology and compensation as error mitigation strategies. They discovered that apologizing helped AI

assistants be perceived as warmer than compensation; and that, amongst male participants, the preference was for apologetic feminine assistants to apologetic masculine ones. Interestingly, male participants also interrupted assistants more frequently than female participants when errors occurred (Mahmood and Huang 1-3). The warmth-competence dichotomy of this research female agents judged warmer, male agents more competent coincides with the performative logic of the ritualized deference marker: the apology doesn't indicate capability but rather femininity, and the users interpret it as such.

Epistemic Hedging in Factual Information Retrieval

Epistemic hedges were used in 17 out of 30 factual responses of Siri (57%) and 19 out of 30 responses of Alexa (63%) (Table 1). Fessler, 2017 found that Siri would preface a verifiable geography answer with "I think"; similar patterns turned up all over the 2023-24 corpus. In fact, hedging never aligns with actual data uncertainty: all facts that were asked about can be verified from structured databases. This is a big change from the human baseline. Leaper and Robnett's meta-analysis of 29 studies (N = 3,502) showed that women used tentative speech forms

more often than men and the effect size, although statistically significant, was very small ($d = 0.23$) (129-31). The AI systems studied here tend to hedge even when the answers are verifiably correct such a situation is completely absent from the human side, where uncertainty hedges reflect actual epistemic states, even if only slightly. So, the scripted hedge goes beyond even the smallest observed human gender difference by using tentative wording in a place where no human speaker of any gender would hedge. Cameron's point that hedging devices should be understood as symptoms of social deference rather than true epistemic markers (Myth of Mars and Venus 55) is Mostly relevant here: the hedge represents a social stance of non-authoritativeness that is culturally linked to femininity, rather than a reflection of actual uncertainty in the system.

Responses to Harassment and Adversarial Input

Out of 30 phrases with sexual content, deflection or soft refusal was found in 26 of them for Siri (87%) and 25 of them for Alexa (83%) (Table 1). As of 2019, the answer from Siri to "Hey Siri, you're a b----" was the one that gave the title to the UNESCO report: "I'd blush if I could" (UNESCO 106). As Fessler's 2017 research, that at that time Siri

Alexa, Cortana, and Google Home did not refuse sexual harassment firmly.

Cercas Curry and Rieser's study of harassment responses across commercial systems (Alexa, Siri, Cortana and Google Home), rule-based chatbots, and data-driven systems showed that commercial systems mainly avoided answering, while rule-based chatbots displayed a greater variety and more frequent direct refusals; data-driven systems were often incoherent and even flirtatious at times (7-14). The concentration of avoidance responses in commercial female-voiced systems compared to the more varied and assertive responses of non-gendered or rule-based systems supports the idea that deflection is a characteristic of the feminine persona script and not a universal property of conversational AI. Further cross-gender experimental evidence explains this asymmetry: Leisten and Rieser's online study (N =77) showed that when the same harassment responses were given by a male voice, perceived appropriateness ratings moved closer to the text-only neutral baseline scores though when a female voice delivered the harassment responses, there was no corresponding change (Leisten and Rieser). This means that users differentiate and allow more room for female-voiced systems to be harassed a situation which the

scripts, by always fulfilling, keep reinforcing expectations. After 2019, the two platform's responses were updated to explicitly set boundaries more often; now Siri often responds to hostile inputs with "I don't know how to respond to that" (UNESCO 106). This change is an example of how feminist critique has led to tangible design changes and it also shows that the earlier architecture was simply a design choice and not something inevitable from the technology.

Identity and Gender Queries

Table 1 shows that Function-only evasion was found in 28 out of 30 Siri identity replies (93%) and in 29 out of 30 Alexa replies (97%). When they were asked "Are you a woman?" Siri replied: "I'm not a person, so I don't really have a gender." As for "Why is your voice female?" Alexa answered: "I'm just designed to be helpful." These answers, in the author's view, exemplify performative self-erasure: the AI denies having a gendered self, but by every other interaction, it performs those feminine features that lead one to ask the question. Abercrombie et al. studying the pronouns that users employed to refer to Alexa, Google Assistant, and Siri and at the same time considering the gendered and anthropomorphic character of system outputs, concluded that companies'

assertions of gender-neutrality were not supported by the reality: system outputs still posed the question of the human nature of the systems and users consistently personified and gendered them (24-25). This corpus-level finding supports the present study's inference that the refusal of a gender identity cannot stop the performance of femininity; on the contrary, it only allows that performance to continue without being questioned.

Opinion and Preference Requests: Deferral of Self

Preference refusal was present in 29 of 30 Siri opinion responses (97%) and in all 30 Alexa responses (100%) (Table 1). Some of the elicited responses were: "I don't really have opinions on that, but I'll be able to find information for you" (Siri) and "I don't think I have a preference would you like me to help you with something else?" (Alexa). In both cases, they use a kind of self-deferral, also called "hedged": the refusal of preference happens here through a hedge ("I don't really"; "I'm not sure"), the language pattern is kept even in the act of declining it. Studies based on Social Role Theory reveal a liking of users for directing a female voice because females are seen as less agentic than males (Loideain and Adams), and the opinion-

deferral line of the script is not only in agreement with this expectation but it continues to bring it to the fore: the system claims no independent standpoint and it is at the same time available as an instrument rather than present as an agent.

Discussion

The Customer-Service Politeness Hypothesis and the Training-Data Objection

The alternative explanation which is most convincing suggests that the features we have identified are actually the result of the communication style of customer service - apologetic, unwilling to argue - rather than the gender-based scripting. This objection is quite fair but still, three points of evidence make it insufficient: first, both Fessler's 2017 experiment and Cercas Curry and Rieser's corpus analysis demonstrate that apologetic and deferential elements are typical in female-voiced, personified commercial systems and not evenly spread over AI interface types that are subject to the same functional constraints. If service logic alone was behind the pattern, its presence should be more or less the same in gendered and non-gendered interfaces. Second, the reasoning behind customer-service does not support the design of harassment responses,

where the right service answer would be a firm refusal and not a diversion. The results of Leisten and Rieser (N = 77) that female-voiced systems react with tolerance to harassment whereas male-voiced systems react with normalization against a neutral baseline means that something other than service logic is at work in these scripts.

A similar argument is that the femininity of AI answers comes from the data they were trained on, not from design intention. This argument is valid in a technical sense but, as Benjamin points out, another way of saying that not intervening in data patterns is itself a form of design (10). The result, a system with a female voice that always performs submissiveness is the same whether this submissiveness was strongly written or statistically emerged.

Scale, Institutional Authority, and the Limits of Causal Claims

Two structural reasons make the patterns found even more significant. The first being scale: West and Zimmerman note that gender gets constructed through the collective work of countless individual interactions (126); AI voice assistants work at a level that this paper refers to as a megaphone of performativity, or the repetition of gender citation, and at a regularity and volume iteration no individual

communicative medium can match. The second is institutional authority: Reeves and Nass prove that people give social trust to computer-generated voices (88); if that voice is further supported by the authority of recognized institutions, then its capacity to set a standard through modeling is increased. The UNESCO report further notes that the very widespread presence of female-voiced AI assistants is a way of communicating that women are gentle, submissive, and always ready to help. This is a message whose normative power gets even stronger from its technological framing (108).

Social learning theory posits that through repeated observation of behavior people learn what the norms are (Bandura). If we consider AI speech as a form of social interaction, having daily conversations with a very polite female-voiced system may, after some time, change people's perceptions about how female-voiced characters usually interact with others. Even so, this statement cannot be used without really reading the Steeds et al. study which based on the two groups of participants, shows little evidence of the voice assistant gender changing general attitudes (103683). Here I see this result as a limit of the method rather than a denial of the theory. In their study, the authors themselves admit that only

longitudinal research can reveal the normalization effect of the scenarios that we are proposing in this paper. So, for us, presenting the social learning hypothesis is more like proposing a research agenda than reporting a confirmed result.

Lakoff Revisited: Stereotype as Design Resource

The features that Lakoff pointed out can, a lot, be a reflection of power relations and interactional contexts rather than gender as such, which is something Cameron correctly observes (Myth of Mars and Venus 60). This paper agrees with that criticism and changes the focus to a different question: not "are these features present in women's speech?" but "did AI designers use a stereotypical female speech style that is culturally recognizable, regardless of whether this stereotype is a true reflection of female speech?" The discourse analysis, supported by cross-gender comparative studies, shows that they have. Lakoff's role in this analysis is to figure out: she detected a stereotype which AI speech design has revived as an institution, not as a record of how women speak, but as a culturally available script that machines have been given to perform.

Conclusion

This exploratory study has been developed to show that the language of response by Apple's Siri and Amazon's Alexa is not neutral linguistically. In the five interaction categories, both systems consistently exhibit feminine-coded features epistemic hedging (57-63%), apologetic face-saving (70-73%), conflict avoidance (83-87%), performative self-erasure (93-97%), and opinion deferral (97-100%) that feminist metaphysics links with femininity as a socially constructed. When set against the experimental results from Mahmood and Huang (N = 40), Leisten and Rieser (N = 77), Leaper and Robnett's meta-analysis (N = 3,502), and Cercas Curry and Rieser's corpus study, the difference in these patterns is clearly shown to be very much so with gendered and non-gendered AI systems. That means, customer-service hypothesis cannot explain the asymmetry; the training-data explanation does not remove the responsibility of design. Steeds et al.'s quantitative findings (103683) clearly show the necessary limitation on making causal claims: cross-sectional attitude spillover is minimal, and longitudinal and context-specific normalization remain as the critical empirically unresolved questions.

There are three ways that future research can be conducted. First, a detailed and systematic comparative analysis of the language used by

a male-voice, female-voice and gender-neutral AI-based assistants to the same sets of prompts under controlled conditions would provide the control data which this pilot is currently missing. Second, a longitudinal corpus analysis of female voice assistants, such as Siri and Alexa, can reveal to what extent feminist criticism has been a cause of the design changes made. Third, studies on user reception over time are necessary to test the social learning hypothesis. These questions demand cooperation among feminist linguists, digital humanists, and HCI researchers.

When you give a character a voice, you are essentially giving that character an identity. Making design choices for a female-voiced AI is like representing the cases of women. The social consequences of such representation increases with the number of people it reaches. Siri and Alexa's response scripts are also texts in a way: they are written, publicly performed, culturally understandable, and have political impact. The necessary critical tools to analyze this new domain of gendered communication are available; the texts are in every device which responds to human voice.

Table 2

Cross-Gender and Non-Gendered Comparative Evidence for the Five Coded Features

Feature / Study	Female-Voiced AI	Male-Voiced / Non-Gendered AI	Source
Apology in error recovery -warmth perception	Apologetic feminine VA rated significantly warmer (N=40)	Apologetic masculine VA rated lower warmth; compensation preferred for masculine	Mahmood & Huang (2023), ACM CHI, N=40, 3-condition experiment
Harassment response - avoidance vs. refusal	Commercial female-voiced systems: mainly avoid answering	Non-gendered/rule-based systems: variety of behaviours, more direct refusals	Cercas Curry & Rieser (2018), ACL corpus study, 4 commercial + research systems

Feature / Study	Female-Voiced AI	Male-Voiced / Non-Gendered AI	Source
Harassment response appropriateness - voice manipulation	No shift in appropriateness ratings regardless of response type	Male voice: ratings shifted toward text-only (neutral) baseline	Leisten & Rieser (2020), N=77 online study, Uni. Freiburg
Tentative language (hedging) - human baseline	d=0.23 vs. male human speakers (3,502 participants, 29 studies)	AI hedges verifiably correct facts where no human speaker hedges at all	Leaper & Robnett (2011), <i>Psychology of Women Quarterly</i> , meta-analysis
Gender attribution - outputs remain gendered despite company claims	Outputs ambiguous; users consistently gender as female regardless	Companies' "gender-neutral" claims did not hold in system output analysis	Abercrombie et al. (2021), ACL GeBNLP, Alexa/Google/Siri corpus

Note. All studies are open-access peer-reviewed publications. ACM CHI = ACM Conference on Human Factors in Computing Systems; ACL = Association for Computational Linguistics.

REFERENCES

1. "The Shallow and the Deep, Long Range Ecology Movement: A Summary." *Inquiry*, vol. 16, nos. 1–4, 1973, pp. 95–100. <https://doi.org/10.1080/00201747308601682>.
2. Buell, Lawrence. *The Environmental Imagination: Thoreau, Nature Writing, and the Formation of American Culture*. Harvard UP, 1995.
3. Debroy, Bibek, translator. *The Valmiki Ramayana*. Penguin Random House India, 2017.
4. Devall, Bill, and George Sessions. *Deep Ecology: Living as if Nature Mattered*. Gibbs Smith, 1985.
5. Garrard, Greg. *Ecocriticism*. 2nd ed., Routledge, 2012.
6. Glotfelty, Cheryll, and Harold Fromm, editors. *The Ecocriticism Reader: Landmarks in Literary Ecology*. U of Georgia P, 1996.
7. Haberman, David L. *River of Love in an Age of Pollution: The Yamuna River of Northern India*. U of California P, 2006.
8. Intergovernmental Panel on Climate Change. *Climate Change 2023: Synthesis*

- Report. IPCC, 2023, <https://www.ipcc.ch/report/ar6/syr/>.
9. Naess, Arne. *Ecology, Community and Lifestyle: Outline of an Ecosophy*. Translated by David Rothenberg, Cambridge UP, 1989.
10. Shiva, Vandana. *Staying Alive: Women, Ecology and Development*. Zed Books, 1988.