

Beyond Repair: An Exploratory Study of Non-Lexical Cues in Human–Robot Clarification Requests

Vanessa Vanzan¹, Talha Bedir¹, Bora Kara¹,

Vladislav Maraev¹, Erik Lagerstedt¹, Christine Howes¹

¹University of Gothenburg

Abstract

Clarification requests are interpreted in relation to their interactional context, yet little controlled work has examined how non-lexical cues shape responses to them. We report an exploratory Wizard-of-Oz study in which 18 participants discussed a moral dilemma with a Furhat robot that produced clarification requests in three conditions: baseline, filled pause, and laughter. Responses following laughter-prefaced clarification requests showed higher observed proportions of reformulation and sentential responses, whereas the filled-pause condition showed the lowest proportions. These preliminary observations motivate further investigation into how non-lexical cues contribute to the pragmatic interpretation of clarification requests in human–robot dialogue.

1 Introduction

1.1 Clarification requests and interactional cues

Clarification requests (CRs); also called next turn repair initiators in the Conversation Analysis literature (Schegloff et al., 1977), are a central resource for conversational repair, helping interlocutors resolve problems of speaking, hearing, and understanding (Purver, 2004; Ginzburg, 2012; Dingemanse and Enfield, 2024). CRs vary in how specifically they identify the source of trouble. In particular, restricted offers repeat or reformulate a candidate referent, as in “the doctor?”, thereby displaying partial understanding while inviting confirmation or elaboration (Dingemanse and Enfield, 2024). Recent work highlights that clarification requests should be understood as situated interactional actions whose interpretation depends on their placement within an ongoing interaction (Healey et al., 2025). Non-lexical cues may therefore affect otherwise how identical CRs are understood.

1.2 Laughter and filled pauses

Laughter is a ubiquitous social signal and an important interactional resource in conversation. It can perform interactional work beyond signalling humour (Glenn, 2003; Ginzburg et al., 2020; Maz-zocconi et al., 2020; Maraev et al., 2021); laughter-prefaced repetitions may also perform different actions from questioning repetitions with similar lexical content (Jefferson, 1972). Related work shows that laughter contributes to the organisation of repair sequences (Yamamoto, 2022). Filled pauses, meanwhile, are associated not only with planning and word retrieval but also with signalling uncertainty and coordinating interaction (Levelt, 1983; Bortfeld et al., 2001; Clark, 1996; Goodwin, 1979).

Despite these observations, little controlled experimental work has examined how such cues shape responses to CRs. We report an exploratory human–robot study in which a Furhat robot¹ produced restricted-offer CRs preceded by silence, a filled pause, or laughter. Using a robot allowed us to manipulate the preceding cue systematically while keeping the lexical content of the CR constant. Building on work suggesting that laughter influences the interpretation of repair sequences (Jefferson, 1972; Yamamoto, 2022), we expected laughter-prefaced clarification requests to encourage more elaborated responses than filled-pause-prefaced or baseline clarification requests.

2 Pilot study

Eighteen participants took part in a Wizard-of-Oz interaction with a Furhat robot while discussing the Balloon Task moral dilemma (Lavelle et al., 2012). The task elicits extended reasoning and stance-taking, creating natural opportunities for clarification requests to be inserted. Sessions lasted between 5 and 12 minutes, and all interactions were

¹www.furhatrobotics.com.

audio- and video-recorded for subsequent annotation. Participants were assigned to one of three between-subject conditions: Baseline Clarification Request (BCR; $n = 5$), Filled-Pause Clarification Request (FCR; $n = 7$), or Laughter Clarification Request (LCR; $n = 6$).

The robot produced restricted-offer CRs by repeating the character currently under discussion (e.g., “the doctor?”, “the pilot?”, “the teacher?”, “the prodigy?”). Restricted-offer CRs were chosen because they allow the preceding non-lexical cue to be manipulated while keeping the lexical content of the repair initiation constant. The lexical CR was identical across conditions; only the preceding cue varied, as shown in. All CRs were prerecorded and matched in duration. An experimenter triggered them at contextually appropriate points when participants were discussing or justifying their evaluation of a character.

Condition	Example
BCR	[silence] + “the doctor?”
FCR	“Hmm... the doctor?”
LCR	“Hahaha, the doctor?”

Table 1: Clarification request conditions

Across the study, 76 CRs were produced. Responses were manually annotated following [Healey et al. \(2011\)](#). The annotation categories were non-exclusive. Here we focus on reformulation, continuation of prior reasoning, and sentential responses, as these properties capture different forms of response elaboration.

3 Exploratory observations

Participants responded to 73 of the 76 clarification requests, suggesting that the robot’s repair initiations were generally treated as requiring a response. Table 2 summarises the observed response proportions.

Response type	BCR	FCR	LCR
Reformulation	7/23 (30%)	4/24 (17%)	15/29 (52%)
Continuation	11/23 (48%)	7/24 (29%)	18/29 (62%)
Sentential	12/23 (52%)	6/24 (25%)	18/29 (62%)

Table 2: Observed response proportions by condition. Categories were not mutually exclusive

The LCR condition showed the highest observed proportions for all three response properties, whereas the FCR condition showed the lowest. The contrast between LCR and BCR was smaller than the contrast involving FCR, particularly for sen-

tential responses. The descriptive pattern is therefore compatible both with laughter encouraging elaboration and with filled pauses making shorter responses more likely.

One illustrative example of a reformulation response from the LCR condition is shown below:

- (1) **User:** uh is perhaps William
Agent: Hahaha, the pilot?
User: I think it’s William

The participant reformulates the initial reference into a more explicit statement following the laughter-prefaced clarification request.

Another example from the FCR condition shows a contrasting response pattern:

- (2) **User:** the doctor
Agent: Hmm... the doctor?
User: yes

Unlike the LCR example, the participant responds with a minimal confirmation rather than reformulating or expanding the previous contribution.

4 Discussion

The preliminary observations are consistent with previous work arguing that CRs should be understood as situated interactional actions whose interpretation depends on their placement within an ongoing interaction ([Healey et al., 2025](#)). Although the lexical content of the CRs remained constant, different preceding non-lexical cues were associated with different response patterns.

The higher observed proportions of reformulation and sentential responses following LCR may be compatible with Jefferson’s (1972) observation that laughter-prefaced repetitions and questioning repetitions may perform different interactional work. The lower proportions of reformulation and sentential responses in the FCR condition are compatible with previous accounts describing filled pauses as signals associated with planning, uncertainty, and the coordination of interaction ([Clark, 1996](#); [Goodwin, 1979](#)). However, the present data do not establish how participants interpreted either cue. Different non-lexical cues may lead interlocutors to orient differently to identical CRs. The observations motivate further qualitative analysis and larger-scale studies.

Limitations

This study should be regarded as an exploratory pilot study. The sample size was small, with only 18 participants distributed across three between-subject conditions, and multiple clarification requests were produced by each participant. Consequently, the reported response patterns should be interpreted as descriptive rather than confirmatory. In addition, the data were annotated by a single annotator, and inter-annotator agreement was therefore not assessed. Future work will extend this study by collecting data from a substantially larger participant sample and by validating the annotation scheme through independent annotation by multiple annotators. These steps will enable participant-level analyses and a more robust evaluation of whether the observed response patterns generalise beyond the present dataset.

Acknowledgements

This research was supported by ERC Starting Grant DivCon: Divergence and convergence in dialogue: The dynamic management of mismatches (101077927). Howes was further supported by the Swedish Research Council grant (VR project 2014-39) for the establishment of the Centre for Linguistic Theory and Studies in Probability (CLASP) at the University of Gothenburg. We are also thankful to Simon Forsberg for his assistance with programming, to Viktoria Paraskevi Daniilidou for her help with Furhat remote API implementation, to Amelie Robrecht-Hilbig for valuable comments, and to the anonymous reviewers.

References

- Heather Bortfeld, Silvia D Leon, Jonathan E Bloom, Michael F Schober, and Susan E Brennan. 2001. Disfluency rates in conversation: Effects of age, relationship, topic, role, and gender. *Language and speech*, 44(2):123–147.
- Herbert H. Clark. 1996. *Using Language*. Cambridge University Press.
- Mark Dingemanse and N. J. Enfield. 2024. [Interactive repair and the foundations of language](#). *Trends in Cognitive Sciences*, 28(1):30–42.
- Jonathan Ginzburg. 2012. *The interactive stance: Meaning for conversation*. Oxford University Press, Oxford.
- Jonathan Ginzburg, Chiara Mazzocconi, and Ye Tian. 2020. Laughter as language. *Glossa: a journal of general linguistics*, 5(1):1–51.
- Phillip Glenn. 2003. *Laughter in interaction*. Cambridge University Press, Cambridge.
- Charles Goodwin. 1979. The interactive construction of a sentence in natural conversation. *Everyday language: Studies in ethnomethodology*, 97:101–121.
- Patrick G. T. Healey, Arash Eshghi, Christine Howes, and Matthew Purver. 2011. Making a contribution: Processing clarification requests in dialogue. In *Proceedings of the 21st Annual Meeting of the Society for Text and Discourse*, pages 11–13, Poitiers, France.
- Patrick GT Healey, Julian Hough, Dirk Vom Lehn, Elif Ecem Özkan, and Rose McCabe. 2025. Negotiating shared understanding: Coding repair in social interaction. *Research on Language and Social Interaction*, 58(3):281–302.
- Gail Jefferson. 1972. Side sequences. In David Sudnow, editor, *Studies in Social Interaction*, pages 294–338. Free Press, New York.
- Mary Lavelle, Patrick G. T. Healey, and Rosemarie McCabe. 2012. Is nonverbal communication disrupted in interactions involving patients with schizophrenia? In *Proceedings of the 34th Annual Conference of the Cognitive Science Society*, pages 1772–1777.
- Willem J. M. Levelt. 1983. Monitoring and self-repair in speech. *Cognition*, 14(1):41–104.
- Vladislav Maraev, Bill Noble, Chiara Mazzocconi, and Christine Howes. 2021. [Dialogue act classification is a laughing matter](#). In *Proceedings of the 25th Workshop on the Semantics and Pragmatics of Dialogue - Full Papers*, Potsdam, Germany. SEMDIAL. Part of PhD thesis, 12 pages.
- Chiara Mazzocconi, Ye Tian, and Jonathan Ginzburg. 2020. [What’s your laughter doing there? a taxonomy of the pragmatic functions of laughter](#). *IEEE Transactions on Affective Computing*, pages 1–1.
- Matthew Purver. 2004. [The theory and use of clarification requests in dialogue](#). Phd dissertation, University of London.
- Emanuel A. Schegloff, Gail Jefferson, and Harvey Sacks. 1977. The preference for self-correction in the organization of repair in conversation. *Language*, 53(2):361–382.
- Mai Yamamoto. 2022. *Delayed Laughter and Open Class Repair Initiators in English and Japanese*. Ph.D. thesis, Purdue University.