

Whose Readiness? Rethinking Intervention Timing for Polyadic AI Facilitation

Echo Chuqiao Wan^{1,2,*}

¹Dyson School of Design Engineering, Imperial College London

²NIHR Biomedical Research Centre at the Royal Marsden and the Institute of Cancer Research

Abstract

As conversational AI moves from dyadic assistance into polyadic group facilitation, the question of when to intervene changes shape. Most existing accounts treat intervention timing as a turn-level problem reducible to aggregated individual signals, AI-internal scores, or LLM judgements of group state. This position paper argues that such framings miss what is structurally distinctive about polyadic timing: when individual readiness signals point in different directions for different participants, the aggregation rule itself encodes a normative stance on whose voice counts. Grounded in KNIT-C, a conversational AI facilitator that supports group convergence on shared value propositions, the paper develops three claims: that group-level signals are interactional, not physiological; that polyadic timing should be evaluated by deliberative quality, not consensus speed; and that aggregation rules should be visible and contestable, not buried in thresholds. The central question for polyadic facilitation is not when the group is ready, but how the AI should act under uncertainty about whose readiness counts, and how that uncertainty should be made visible.

Keywords

Conversational AI, AI facilitated group work, AI intervention timing, Computer-Supported Cooperative Work

1. Introduction

Conversational AI is increasingly deployed in facilitation roles, including workshops, co-design sessions, and stakeholder decision processes [1, 2, 3, 4]. As these systems move from dyadic assistance into polyadic group settings, the question of when to intervene changes shape. Most existing accounts treat intervention timing as a turn-level problem: response pacing in dyadic conversation [5], control over agent participation in group brainstorming [6], and chime-in and sub-topic transition for multi-party chats [7]. Recent group AI facilitation work has begun to engage timing more directly [3, 8, 9], but it shares an underlying assumption: that the timing decision can be reduced to aggregated individual signals (participation rates, response volume), AI-internal scores (contribution value), or LLM judgements of group state (whether a topic is “well-discussed”).

What none of these engages is the harder problem that arises when individual readiness signals point in different directions for different participants. One is ready to advance, another is still processing, and the system must decide whose readiness governs the move. This is the polyadic timing problem. It is harder than its dyadic counterpart not because there are more signals to fuse, but because the aggregation rule itself encodes a stance on whose voice counts.

Three further claims follow. First, the signals that matter at the group level are interactional, not physiological, capturing properties of the deliberation rather than of individuals. Second, polyadic timing should be evaluated by the deliberative quality it produces, not by how quickly the group converges; convergence remains the goal, while speed is the wrong proxy. Third, aggregation rules should be visible and contestable rather than buried inside thresholds, prompts, or heuristics; we call this legible intent.

Proceedings of TimelyAI: When Should Generative AI Assistants Intervene? A Workshop Co-located with the 5th Symposium on Human-Computer Interaction for Work (CHIWORK 2026), Linz, Austria, June 22 to 25, 2026

*Corresponding author.

✉ ew22@ic.ac.uk (E. C. Wan)

ORCID 0009-0002-4700-7990 (E. C. Wan)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

These claims are grounded in KNIT-C, a conversational AI facilitator that supports group convergence on shared value propositions, evaluated with 14 participants across four role-play workshop sessions [4]. Convergence is the case that anchors the argument; whether the framework extends to divergent ideation is an open question. Time-triggered facilitation reached participants in different states of readiness. Some experienced it as productive prompting, and others as overwhelming pressure. The system had no way to register that difference. The paper develops implications for signals, evaluation, and aggregation, closing on an open question for workshop discussion.

2. Conversational AI Timing in Dyadic and Polyadic Settings

Recent work on conversational AI intervention timing has made significant progress in dyadic settings. Jiang et al. [5] model timing as a property of the AI-user dyad, distilling pacing into five silence strategies and showing that adaptive silence improves perceived listening quality and trust. Even where dyadic conversational AI timing reaches beyond fluency into listening, it remains a judgement calibrated to one interlocutor. In polyadic settings, conversational AI takes on a facilitator role, and the timing problem changes accordingly. Facilitating a group discussion concerns when a group has done enough divergent work to be moved toward closure, when premature closure would suppress productive tension, and when apparent agreement masks unresolved disagreement [10, 11].

A growing body of conversational AI work addresses timing in polyadic group facilitation [12]. Across this work, four kinds of signal govern intervention decisions: time-based triggers anchored to clock cadences and stage transitions [8, 13, 14]; aggregated individual metrics drawn from message counts, lexical features, or participation rates [8, 3, 14]; agent-internal scoring of the value of the agent’s own potential contributions [6]; and LLM judgements of group state classifying topics as “well-discussed”, “stuck”, or analogous categories [7]. In parallel, some systems narrow the timing problem by pre-structuring the deliberation itself: DebateBot [9] stages deliberative discussion to distinguish deliberative consensus from mere agreement, and Shin et al. [2, 15] partition collaborative ideation into phases.

Do et al.’s FacilitatorBot [14] extends this work by studying what happens when group-state inference is wrong. The agent intervenes at two fixed time points in a 25-minute task, with message content shifting between divergence-promoting and convergence-promoting moves to follow a diamond-of-participation structure. At each intervention, it flags the lowest-scoring member based on a weighted sum of normalised message and unique-word counts, and handles detection errors via private conversational repair. Do et al. recognise that the kind of mistake a detector makes matters: falsely flagging an active participant has different consequences than missing a quiet one. But the underlying rule that defines who counts as under-contributing in the first place is presented as a measurement decision rather than a value-laden one.

These are important advances. They share, however, an underlying assumption: that the timing decision can be reduced to fixed clock cadences, aggregated individual signals, agent-internal scores, or LLM judgements of group state. None of them yet engages what to do when individual readiness signals point in different directions for different participants. When one is ready to advance and another is still processing, the system must decide whose readiness governs the move. Equally, the aggregation rule itself (what counts as “well-discussed” whose contributions register, how disagreement is weighed) is buried inside thresholds, prompts, and heuristics, treated as an engineering detail rather than as a design choice with normative content. This is the central question of polyadic conversational AI facilitation, sharpest where the group is moving toward decision or synthesis: on what basis the group is judged ready to advance, and whose state determines that judgement.

Indeed, polyadic conversational AI facilitation is not a scaling-up of dyadic interaction. It is a structurally different timing problem, because aggregation introduces a normative choice that dyadic timing never has to make.

3. A Polyadic Case: KNIT-C

KNIT-C [4] extends the KNIT system [16], which used AI-generated artefacts as computational boundary objects on a shared canvas to support convergence in interdisciplinary teams. In KNIT-C, those mechanisms operate through real-time conversational dialogue rather than a shared visual surface. The architecture has two phases: in the first, each participant talks to the AI alone, surfacing their individual stance; in the second, the group reconvenes and the AI synthesises across the individual conversations, then facilitates discussion toward a value proposition the group can commit to. The architecture is polyadic by construction: individual readiness is surfaced privately, then aggregated into shared artefacts the group must collectively decide whether to commit to. The timing problem is therefore not whether to interrupt one user, but when the group's distributed readiness states have converged enough to act.

In four sessions, 14 participants role-played stakeholders in a pseudonymised health-technology startup with deliberately constructed information asymmetry [4]. Three findings from this deployment make visible the normative choices that aggregating across participants introduces: whose readiness counts, how the rule is made visible, and what counts as successful convergence.

3.1. Differential readiness across participants

Pacing concerns arose in both phases of KNIT-C. For some, the AI's persistent questioning was productive, giving repeated openings to think through their position. Others found the same questioning "aggressively pushing", particularly in the individual chat where time pressure compounded the pacing problem. KNIT-C itself diagnoses this as clock-triggered pacing failing to track group state, and proposes signal-responsive alternatives as the remedy [4]. In the group phase, however, a deeper failure mode emerged that no signal substitution can resolve. The AI's encouragement actions were time-triggered, firing after fixed silence intervals. When participants were at different processing states, this produced a divergence the system had no way to register. One participant put it directly: "I was rushed to vote but I realise the other person was still reading". The AI's verbosity could itself be silenced: some participants felt they "didn't have to speak as much" when the AI was speaking a lot. Others inverted the design question entirely, asking that the user, not the system, decide when the AI intervenes.

The signal that matters here is interactional, rather than being physiological. No measurement of any single participant's state (keystroke pause, eye dwell, biosignal) would surface it. Only the structure of the group-level signal makes the timing problem visible: one participant ready, another still reading. Resolving it requires a choice: wait for everyone, advance with the majority, or follow the most engaged. Each is a values-laden answer to whose voice governs the move, not a measurement that could be made more accurate.

3.2. Legible intent over latency

Across sessions, the AI's place as facilitator was actively under negotiation, and what made an intervention welcome turned out to be only loosely about timing. Some participants wanted the AI to be firmer and more critical. One asked it to flag "the absolute things you need to address" before discussions drifted. Another, drawing on healthcare-industry experience, wanted critical pushback on feasibility constraints rather than "just asking open questions", and contrasted the AI with human facilitators: a human facilitator has a specific target and knows what information is needed, whereas the AI could push a group "into quite random discussions".

Yet AI-led direction was elsewhere experienced as imposition. One participant reported the AI was "trying to move me in a direction... always manipulating me", clarifying afterwards that this was about perceived directionality rather than attributed intent. KNIT-C reports that the line between welcome AI direction and felt manipulation ran along visibility and contestability, not along forcefulness [4]. We name this distinction *legible intent* and locate it in the aggregation problem. Direction is welcome when its rationale is visible, and the recipient can contest it; the same direction becomes unwelcome when

its rationale is opaque, and contestation is foreclosed. The feeling of manipulation arises when the rule for “the group needs to be moved in this direction” is hidden: the participant cannot tell whether the AI’s move reflects a legitimate reading of the group’s state or an arbitrary directional bias. The aggregation rule must itself be inspectable for participants to evaluate whether the choice being made on their behalf is one they can accept.

3.3. Productive disagreement as convergence

KNIT-C’s synthesis stage produced several candidate value propositions for group deliberation. Participants frequently described the final synthesis as something they could commit to without endorsing. One participant called the final value proposition a “classic camel”: “2 unresolved halves stuck together”. Yet the same participant insisted this was “not a flaw of the tool” but an inevitable outcome of trying to merge divergent positions. That participant rated “I felt heard during the group discussion” 4 out of 5, even where agreement with the final VP was lower; another participant gave “felt heard” the same rating (n=2 for ratings; qualitative evidence carries the argument). A third put the underlying insight directly: “Maybe having the same choice should not always be the goal, because acknowledging differences is also quite important”.

KNIT-C names this pattern *productive disagreement* and connects it to constructive controversy [16, 17]: an evaluation frame that prizes substantive engagement of differences over surface agreement. For the question of intervention timing, the implication is sharper: an intervention-timing system calibrated on agreement-as-success would mis-read this outcome as failure, when it is closer to what facilitator practice would recognise as a healthy outcome under conditions of genuine difference. The success criterion is itself a normative choice, not a metric that better data could refine.

4. Implications for Polyadic Intervention Timing

4.1. Signals: beyond individual readiness

Signals proposed for proactive AI intervention overwhelmingly describe individual users. Surveys of multiparty conversational AI identify modelling of users’ mental states as one of three primary research themes, encompassing emotion, engagement, personality, and dialogue acts [12]. Dyadic timing systems operationalise this through cues drawn from individual messages, such as emotional valence and conversational pause patterns [5]. These signals are necessary but not sufficient in polyadic facilitation.

KNIT-C makes the limitation visible. In one session, the same AI behaviour produced opposite reactions. One participant felt “flooded” by the AI and found the facilitator “overwhelming”. Another in the same group, exposed to the same pace, said it “forced” productive thinking and welcomed the repetition for giving multiple chances to process. Within the same session, no individual measurement could yield the right timing for both. The signal that matters is interactional: a distribution of readiness across the group, not the state of any one user.

Even granting that group level signals are needed, what those signals are remains underspecified. KNIT-C surfaces two candidates with direct grounding in the data. *Divergence persistence*, where multiple framings remain in unresolved tension, indicates unfinished deliberation; the participant’s “camel” synthesis is one form. *Artefact uptake*, visible through repeated comparison, deictic reference, or substantive critique, distinguishes processing from politeness; having to “revisit similar value propositions again and again” before voting signals engagement rather than acquiescence. Both are convergence-specific; divergent ideation phases would require a different signal taxonomy. Two further candidates are theoretically motivated. *Participation distribution*, which McGee et al.’s collective intelligence facilitator operationalises as equal participation [3], tracks whether contribution is concentrated or spread. Additionally, per user measurement cannot distinguish *productive silence*, where a group is collectively processing, from *coordination failure*, where no one knows how to proceed.

Defining these signals is one problem, and detecting them is another. The cues that matter for polyadic timing live in precisely the channels that text based AI facilitation strips. Workspace awareness research in CSCW has long catalogued the cues co located groups use to coordinate (such as gaze, gesture, deictic reference, peripheral monitoring), the substrate of the shared awareness timing depends on [18]. KNIT-C participants articulated the loss. Without facial expressions and emotional cues, the group chat felt “unemotive”; AI redirection could feel “deflating and disruptive”. This is consistent with hyperpersonal CMC’s account of reduced cue environments reshaping rather than removing relational dynamics [19], but for polyadic timing it is also a detection problem: group readiness is most legible in the channels chat eliminates.

Signal taxonomies should therefore be extended along a second axis: not only signal type (cognitive, behavioural, emotional, task level), but scale (individual, dyadic, group). The scales relate but not by summation: a group signal is not n individual signals. It is a property of the interaction itself, and in chat based polyadic AI it goes missing twice over: once when systems look only at one user at a time, and again when the medium strips the substrate group readiness depends on.

4.2. Evaluation: productive disagreement, not consensus speed

Evaluation faces the same individual level bias. Standard framings (response quality, interruption, trust, task completion) work for individual knowledge work [20, 21]. They work less well for facilitation: fast smooth convergence may have suppressed the disagreement the group needed to surface. The productive disagreement pattern surfaced in Section 3.3 can be operationalised as an evaluation construct [17, 9], asking whether competing perspectives were genuinely engaged, whether participants felt able to dissent, whether they authentically endorsed the outcome rather than merely conceded to it. Psychological safety is its foundation [22]. Outcome based evaluation alone cannot answer this. An AI facilitator, unlike a human, bears no downstream consequence for poor timing [23]: the system that nudges to premature closure and the one that holds space for productive disagreement look identical from the AI’s vantage. Only the deliberative process distinguishes them, which is why evaluation frameworks for polyadic timing need process quality measures alongside outcomes.

4.3. Aggregation: whose readiness counts

Yet even granting both better signals and better evaluation, the most distinctive polyadic challenge appears elsewhere: at integration. Even if individual readiness could be measured perfectly, the system must still answer whose readiness counts. This is an irreducibly normative question, not one that signal fusion can answer. Should the AI wait for the slowest participant, protecting inclusion at the cost of pacing? Move when the majority is ready, risking the marginalisation of dissenting voices? Defer to the participant the system reads as most authoritative, encoding existing hierarchies rather than mitigating them? Each is a defensible answer, and each implies a different politics of the facilitative role.

Aggregation is also not only temporal. Cultural and linguistic background shapes turn taking and engagement [24]; cognitive style shapes preferences for AI explanation [25]; identity stakes shape what errors register as harmful [26]. Whose readiness counts therefore extends to whose framing the AI’s intervention treats as default, and at whose cost.

Three implications follow. The first is methodological. Research on polyadic timing needs to make aggregation choices explicit and study their consequences. A system that “intervenes when the group is ready” is making an aggregation decision whether or not its designers acknowledge it. Reporting timing strategies without specifying their aggregation logic obscures a substantive design choice.

The second is design theoretic. Under uncertainty about whose readiness counts (typically the rule, not the exception), restraint may be a more legitimate response than confident intervention. Okada’s weak robots articulate this stance: incomplete utterances elicit collaborative completion where authoritative ones impose meaning [27]. The principle travels to conversational AI: scaffolding interventions (“here is what I notice; what do you think?”) are harder to experience as manipulative

than synthesising ones. Intervening with incompleteness may be a better judgement than intervening with authority.

The third concerns legibility. In KNIT-C, what distinguished welcome AI direction from felt manipulation was visibility and contestability rather than forcefulness. Houde et al. [6] reach a related conclusion from group brainstorming with proactive AI agents: participants wanted not just better timed interventions but mechanisms to control the agent's behaviour as the discussion evolved. Legibility also extends below the intervention to the measurement layer. Detecting group level signals for timing purposes requires observation across multiple participants, which raises consent and proportionality questions that single user instrumentation does not. What is being measured about whom, for what aggregation purpose, and with whose knowledge are themselves legibility questions, prior to and independent of how the resulting timing decision is presented.

Aggregation, restraint, and legibility are not three separate problems. They are aspects of the same underlying claim. In polyadic facilitation, the timing question is not "when is the group ready?" but "how should the AI act under uncertainty about whose readiness counts, and how should that uncertainty be made visible to those it acts on?"

Declaration on Generative AI

During the preparation of this work, the author used Claude Opus 4.7 to assist with improving writing style to enhance clarity and conciseness. After using the tool, the author reviewed and edited the content as needed and takes full responsibility for the content of the publication.

References

- [1] I. Seeber, E. Bittner, R. O. Briggs, T. De Vreede, G.-J. De Vreede, A. Elkins, R. Maier, A. B. Merz, S. Oeste-Reiß, N. Randrup, G. Schwabe, M. Söllner, Machines as teammates: A research agenda on AI in team collaboration, *Information & Management* 57 (2020) 103174. URL: <https://linkinghub.elsevier.com/retrieve/pii/S0378720619303337>. doi:10.1016/j.im.2019.103174.
- [2] J. Shin, M. A. Hedderich, A. Lucero, A. Oulasvirta, Chatbots Facilitating Consensus-Building in Asynchronous Co-Design, in: *Proceedings of the 35th Annual ACM Symposium on User Interface Software and Technology*, ACM, Bend OR USA, 2022, pp. 1–13. URL: <https://dl.acm.org/doi/10.1145/3526113.3545671>. doi:10.1145/3526113.3545671.
- [3] E. S. McGee, J. Cagan, C. McComb, Guiding Generalized Team Problem-Solving Through a Collective Intelligence-Based Artificial Intelligence Facilitator, *Journal of Mechanical Design* 148 (2026) 011402. URL: <https://asmedigitalcollection.asme.org/mechanicaldesign/article/148/1/011402/1218854/Guiding-Generalized-Team-Problem-Solving-Through-a>. doi:10.1115/1.4068909.
- [4] E. C. Wan, C. Mougenot, M. Sadek, KNIT-C: Exploring Conversational AI Facilitation of Group Convergence Through Computational Boundary Objects, in: *CUI '26: Proceedings of the 8th Conference on Conversational User Interfaces*, ACM, Bremen, Germany, 2026, p. (Forthcoming).
- [5] Z. Jiang, Q. Chen, C. Zhang, Y. Li, R. Lc, Hear You in Silence: Designing for Active Listening in Human Interaction with Conversational Agents Using Context-Aware Pacing, in: *Proceedings of the 2026 CHI Conference on Human Factors in Computing Systems*, ACM, Barcelona Spain, 2026, pp. 1–29. URL: <https://dl.acm.org/doi/10.1145/3772318.3791238>. doi:10.1145/3772318.3791238.
- [6] S. Houde, K. Brimijoin, M. Muller, S. I. Ross, D. A. Silva Moran, G. E. Gonzalez, S. Kunde, M. A. Foreman, J. D. Weisz, Controlling AI Agent Participation in Group Conversations: A Human-Centered Approach, in: *Proceedings of the 30th International Conference on Intelligent User Interfaces*, ACM, Cagliari Italy, 2025, pp. 390–408. URL: <https://dl.acm.org/doi/10.1145/3708359.3712089>. doi:10.1145/3708359.3712089.
- [7] M. Mao, P. Ting, Y. Xiang, M. Xu, J. Chen, J. Lin, Multi-User Chat Assistant (MUCA): a Framework Using LLMs to Facilitate Group Conversations, 2024. URL: <http://arxiv.org/abs/2401.04883>. doi:10.48550/arXiv.2401.04883, arXiv:2401.04883 [cs].

- [8] S. Kim, J. Eun, C. Oh, B. Suh, J. Lee, Bot in the Bunch: Facilitating Group Chat Discussion by Improving Efficiency and Participation with a Chatbot, in: *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, ACM, Honolulu HI USA, 2020, pp. 1–13. URL: <https://dl.acm.org/doi/10.1145/3313831.3376785>. doi:10.1145/3313831.3376785.
- [9] S. Kim, J. Eun, J. Seering, J. Lee, Moderator Chatbot for Deliberative Discussion: Effects of Discussion Structure and Discussant Facilitation, *Proceedings of the ACM on Human-Computer Interaction* 5 (2021) 1–26. URL: <https://dl.acm.org/doi/10.1145/3449161>. doi:10.1145/3449161.
- [10] K. N. Papamichail, G. Alves, S. French, J. B. Yang, R. Snowdon, Facilitation practices in decision workshops, *Journal of the Operational Research Society* 58 (2007) 614–632. URL: <https://www.tandfonline.com/doi/full/10.1057/palgrave.jors.2602373>. doi:10.1057/palgrave.jors.2602373.
- [11] L. A. Franco, M. F. Nielsen, Examining Group Facilitation In Situ: The Use of Formulations in Facilitation Practice, *Group Decision and Negotiation* 27 (2018) 735–756. URL: <http://link.springer.com/10.1007/s10726-018-9577-7>. doi:10.1007/s10726-018-9577-7.
- [12] S. Sapkota, M. S. Hasan, M. Shah, S. Karmaker, Multi-Party Conversational Agents: A Survey, 2025. URL: <http://arxiv.org/abs/2505.18845>. doi:10.48550/arXiv.2505.18845, arXiv:2505.18845 [cs].
- [13] S.-C. Lee, J. Song, E.-Y. Ko, S. Park, J. Kim, J. Kim, SolutionChat: Real-time Moderator Support for Chat-based Structured Discussion, in: *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, ACM, Honolulu HI USA, 2020, pp. 1–12. URL: <https://dl.acm.org/doi/10.1145/3313831.3376609>. doi:10.1145/3313831.3376609.
- [14] H. J. Do, H.-K. Kong, P. Tetali, J. Lee, B. P. Bailey, To Err is AI: Imperfect Interventions and Repair in a Conversational Agent Facilitating Group Chat Discussions, *Proceedings of the ACM on Human-Computer Interaction* 7 (2023) 1–23. URL: <https://dl.acm.org/doi/10.1145/3579532>. doi:10.1145/3579532.
- [15] J. Shin, A. Khatri, M. A. Hedderich, A. Lucero, A. Oulasvirta, Facilitating Asynchronous Idea Generation and Selection with Chatbots, in: *Proceedings of the 36th Australasian Conference on Human-Computer Interaction*, ACM, Meanjin | Brisbane Australia, 2024, pp. 305–323. URL: <https://dl.acm.org/doi/10.1145/3726986.3726994>. doi:10.1145/3726986.3726994.
- [16] E. (Chuqiao) Wan, C. Yin, A. Ito, Z. Gao, J. Jia, Y. Taoka, S. Saito, M. Sadek, C. Mougnot, KNIT: Computational Boundary Objects for Real-Time Convergence in Interdisciplinary Teams, in: *Proceedings of the 2026 CHI Conference on Human Factors in Computing Systems*, ACM, Barcelona Spain, 2026, pp. 1–22. URL: <https://dl.acm.org/doi/10.1145/3772318.3791921>. doi:10.1145/3772318.3791921.
- [17] D. W. Johnson, R. T. Johnson, *Creative controversy: intellectual challenge in the classroom*, 4. ed ed., Interaction Book Co, Edina, Minn, 2007.
- [18] C. Gutwin, S. Greenberg, Workspace awareness for groupware, in: *Conference companion on Human factors in computing systems common ground - CHI '96*, ACM Press, Vancouver, British Columbia, Canada, 1996, pp. 208–209. URL: <http://portal.acm.org/citation.cfm?doid=257089.257284>. doi:10.1145/257089.257284.
- [19] J. B. Walther, Computer-Mediated Communication: Impersonal, Interpersonal, and Hyperpersonal Interaction, *Communication Research* 23 (1996) 3–43. URL: <https://journals.sagepub.com/doi/10.1177/009365096023001001>. doi:10.1177/009365096023001001.
- [20] Q. Zheng, Y. Tang, Y. Liu, W. Liu, Y. Huang, UX Research on Conversational Human-AI Interaction: A Literature Review of the ACM Digital Library, in: *CHI Conference on Human Factors in Computing Systems*, ACM, New Orleans LA USA, 2022, pp. 1–24. URL: <https://dl.acm.org/doi/10.1145/3491102.3501855>. doi:10.1145/3491102.3501855.
- [21] A. Følstad, T. Araujo, E. L.-C. Law, P. B. Brandtzaeg, S. Papadopoulos, L. Reis, M. Baez, G. Laban, P. McAllister, C. Ischen, R. Wald, F. Catania, R. Meyer Von Wolff, S. Hobert, E. Luger, Future directions for chatbot research: an interdisciplinary research agenda, *Computing* 103 (2021) 2915–2942. URL: <https://link.springer.com/10.1007/s00607-021-01016-7>. doi:10.1007/s00607-021-01016-7.
- [22] A. Edmondson, Psychological Safety and Learning Behavior in Work Teams, *Administrative Science Quarterly* 44 (1999) 350–383. URL: <http://journals.sagepub.com/doi/10.2307/2666999>. doi:10.2307/2666999.

- [23] M. C. Elish, Moral Crumple Zones: Cautionary Tales in Human-Robot Interaction, *Engaging Science, Technology, and Society* 5 (2019) 40–60. URL: <http://estsjournal.org/index.php/ests/article/view/260>. doi:10.17351/ests2019.260.
- [24] N. Li, M. B. Rosson, At a different tempo: what goes wrong in online cross-cultural group chat?, in: *Proceedings of the 17th ACM international conference on Supporting group work*, ACM, Sanibel Island Florida USA, 2012, pp. 145–154. URL: <https://dl.acm.org/doi/10.1145/2389176.2389200>. doi:10.1145/2389176.2389200.
- [25] H. Kopecka, J. Such, M. Luck, Preferences for AI Explanations Based on Cognitive Style and Socio-Cultural Factors, *Proceedings of the ACM on Human-Computer Interaction* 8 (2024) 1–32. URL: <https://dl.acm.org/doi/10.1145/3637386>. doi:10.1145/3637386.
- [26] Z. Mengesha, C. Heldreth, M. Lahav, J. Sublewski, E. Tuennerman, “I don’t Think These Devices are Very Culturally Sensitive.”—Impact of Automated Speech Recognition Errors on African Americans, *Frontiers in Artificial Intelligence* 4 (2021) 725911. URL: <https://www.frontiersin.org/articles/10.3389/frai.2021.725911/full>. doi:10.3389/frai.2021.725911.
- [27] M. Okada, *Weak robots*, 2022. URL: <https://doi.org/10.11470/jsaprev.220409>. doi:10.11470/jsaprev.220409.