

Evaluation and Facilitation of Online Discussions in the LLM Era: A Survey

Anonymous ACL submission

Abstract

We present a survey of methods for assessing and enhancing the quality of online discussions, focusing on the potential of Large Language Models (LLMs). While online discourses aim, at least in theory, to foster mutual understanding, they often devolve into harmful exchanges, such as hate speech, threatening social cohesion and democratic values. Recent advancements in LLMs enable artificial facilitation agents to not only moderate content, but also actively improve the quality of interactions. Our survey synthesizes ideas from Natural Language Processing (NLP) and Social Sciences to provide (a) a new taxonomy on discussion quality evaluation, (b) an overview of intervention and facilitation strategies, (c) along with a new taxonomy of conversation facilitation datasets, (d) an LLM-oriented roadmap of good practices and future research directions, from technological and societal perspectives.

1 Introduction

Discussions, especially of complex or controversial topics, are a cornerstone of collective decision-making (Burton et al., 2024). In contrast to initial hopes of promoting mutual understanding (Rheingold, 2000), online discussions (especially in social media) often degenerate into hate speech, personal attacks, promoting conspiracy theories or propaganda – to the extent that they can even be considered a threat to social cohesion and democracy (Tucker et al., 2018; Mathew et al., 2019).

Natural Language Processing (NLP) and Machine Learning (ML) can potentially help improve the quality of online discussions. For example, automatic classifiers (Bang et al., 2023; Molina and Sundar, 2022) are already being used to help or even replace human moderators, by flagging posts that violate the law or policies of online discussion fora (Saeidi et al., 2021).

Social Science provides theories and applications for the facilitation of a discussion, but in

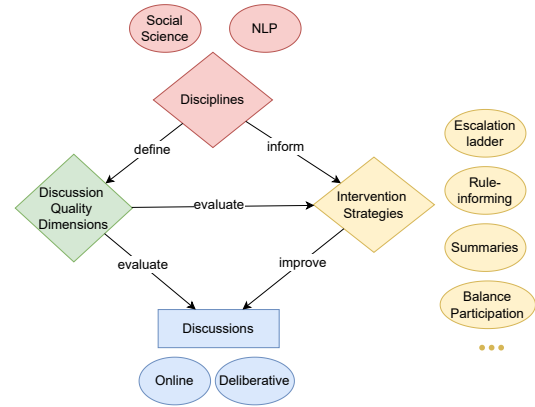


Figure 1: A conceptualization of this survey. We explore approaches from different disciplines, which recommend their own ways of evaluating and improving discussions.

specific contexts, such as teaching/learning (Mansour, 2024) or clinical discussions (Gelula, 1997), without much research devoted to online discussions, such as in social media. While prior NLP studies have explored LLM-facilitated discussions (Burton et al., 2024; Aher et al., 2023; Beck et al., 2024; Schroeder et al., 2024; Small et al., 2023; Cho et al., 2024), rarely does Social Science work examine how facilitation can be automated (Gimpel et al., 2024).

In this survey, we combine LLM-based methods, with ideas from Social Science (e.g., Deliberative Theory) when discussing how to evaluate online discussions, and when exploring intervention strategies. Figure 1 provides a high-level conceptualization of our work.

The main research question of this survey is *can LLMs be used effectively as facilitators in online discussions?* To explore this question, we focus on three key areas: (1) methods (potentially also LLM-based) for evaluating aspects of online discussions, (2) intervention strategies for facilitation,

and (3) available data resources relevant to facilitation. Specifically, we survey discussion evaluation aspects and introduce a new taxonomy (§4). We map tasks suited for ML models, LLMs, and humans, aggregate multidimensional insights on facilitation strategies (§5), and outline future possibilities for LLMs (§6). Additionally, we compare major datasets, dividing them into categories per task (§7). Our work focuses mostly on written thread-like discussions (§2).

Our findings show that (a) many discussion evaluation dimensions coexist in the literature; (b) LLM advancements show significant promise in improving the quality and timeliness of facilitation methods; (c) while surveying the existing datasets, we notice a scarcity of datasets for studying facilitation. We posit that LLM-generated discussions, could become an asset to develop and test automatic facilitation strategies in diverse artificial discussions, before testing the strategies and the LLM-based facilitator agents in more costly experiments with human participants.

2 Terminology

Given the numerous aspects to consider regarding discussion quality and facilitation, we clarify the terminology we use. We highly recommend consulting the Terminology Section of Appendix C and, especially, Table 3, where we explain our findings with regard to the terms used in the literature.

Facilitation vs. Moderation The term ‘moderation’ is more commonly used in NLP (Argyle et al., 2023), typically referring to the flagging and/or removal of unwanted content (‘content moderation’), while ‘facilitation’ is more prevalent in the Social Sciences, where it encompasses a broader scope, including active interventions (Vecchi et al., 2021; Kaner et al., 2007; Trenel, 2009). Given the limited attention to facilitation in NLP and the survey’s grounding in Social Science, we distinguish between the terms, even though they are sometimes used interchangeably in the literature.

Ex-Post moderation This survey mainly focuses on ‘Real-Time, Ex-Post-moderation’, i.e., moderation happening just after the user has posted some content. This is different from pre-moderation approaches, such as nudging users before they post harmful content (Argyle et al., 2023), or delaying the posting of user content until a moderator has had the chance to check it.

Discussion, Deliberation, Dialogue, Debate

The definitions of these terms often vary across literature (Russmann and Lane, 2016; Goñi, 2024). We focus on **discussions**, a general term for verbal/written exchanges (Russmann and Lane, 2016), and **deliberations**, a term for structured discussions focusing on opinion sharing (Degeling et al., 2015; Lo and McAvoy, 2023). This is in contrast to the (at least in theory) collaborative nature of **dialogues** (Rose-Redwood et al., 2018; Bawden, 2021; Goñi, 2024) and the competitive and organized nature of **debates** (Lo and McAvoy, 2023).

Tree-style discussions (or “threads”) are discussions which start from an Original Post (OP) with subsequent comments replying to either the OP or to other comments (Seering, 2020).

3 Comparison to Other Surveys

Only two studies have surveyed the field of NLP while also considering ideas from Social Science. However, they focus mainly on Argument Mining (AM). These are the studies of Wachsmuth et al. (2024) and Vecchi et al. (2021). Wachsmuth et al. (2024) focus primarily on discussion *evaluation* disregarding its relation to facilitation, which is one of the main goals of our survey. The survey of Vecchi et al. (2021) argues that advancing AM for social good requires a collaborative effort between AM and Social Science. They point out that traditional AM has prioritized the logical structure and soundness of arguments, while overlooking other important dimensions, such as civility, respectfulness, inclusiveness, originality, and the broader impacts of discussions—such as encouraging mutual understanding and problem-solving. Building on these notions, we incorporate ideas from Social Science into NLP-based approaches, discussing both discussion evaluation and facilitation, both with a focus on the potential of LLMs.

4 Discussion Quality Evaluation

Improving online discussions presupposes being able to define and measure *discussion quality*. While there have been attempts to provide frameworks for discussion quality evaluation (Kies, 2022; Gerber et al., 2018), none of them is directed towards facilitation. Crucially, most existing frameworks ultimately rely on human judgments as their reference point, yet human evaluation is expensive, slow, and shows low inter-rater agreement on dimensions that involve subjective interpretation,

such as pragmatic cues (Smith et al., 2022; Yeh et al., 2021; Khalid and Lee, 2022). This evaluation bottleneck motivates a taxonomy of evaluation methods that is both comprehensive and amenable to scalable automatic measurement.

In this work, we draw from the works of Bächtiger et al. (2022, 2010); Steenbergen et al. (2003); Falk and Lapesa (2023) and Kies (2022) to define a new social-science-informed taxonomy for discussion quality dimensions. While we present a structured taxonomy, it is important to note that the categories are not mutually exclusive. Rather, elements within the taxonomy may coexist within evaluation dimensions, complement one another, or serve as explanatory mechanisms for other dimensions. An example of the dimension interaction can be found in Table F in the Appendix. The grouped dimensions along with the NLP approaches are shown in the Appendix in Table 4.

4.1 Structure and Logic

Argument Structure and Analysis Argument Quality (AQ) is a multidimensional concept assessed through logical, rhetorical, and dialectical dimensions (Wachsmuth et al., 2017). The logical dimension focuses on the coherence and structure of the argument. The rhetorical dimension assesses persuasiveness, focusing on the argument’s style and emotional appeal. The dialectical dimension assesses the constructiveness of the argument. Empirically, threads with well-formed claim-evidence chains exhibit higher coherence and lower odds of devolving into ad-hominem attacks, making AQ scores, as a discussion quality dimension, an *early-warning* indicator of derailment (Chang and Danescu-Niculescu-Mizil, 2019). All the above dimensions of automatic argument-structure analysis can be used by a facilitator to keep the discussion fact-centered, inclusive, and on track (Falk et al., 2021; Falk and Lapesa, 2023).

Coherence and Flow ‘Coherence’, as described above, evaluates logical consistency, while ‘flow’ assesses smooth progression in discussions (Li et al., 2021). Both are essential tools for facilitators in their effort to redirect off-topic comments and guide transitions between topics during a discussion (Lambert et al., 2024; Park et al., 2012; Falk et al., 2024). A sudden drop in how well responses match the topic or question often comes before personal attacks or off-topic turns (Chang and Danescu-Niculescu-Mizil, 2019; Zhang et al.,

2018), making coherence and flow indicators of argument structure and a valuable early signal for facilitators.

Turn-taking How speakers alternate, the frequency of their turns, and the participants they address can serve as a diagnostic of conversational health. Balanced exchanges enhance coherence (Cervone and Riccardi, 2020), predict constructiveness (§4.3) (Niculae and Danescu-Niculescu-Mizil, 2016), and provide facilitators with actionable cues (Schroeder et al., 2024). To gauge speaking time, turn count, and word usage, researchers have applied metrics such as entropy (Niculae and Danescu-Niculescu-Mizil, 2016) and Gini coefficients (Schroeder et al., 2024).

Linguistic Markers Linguistic markers have been used to help model content and expression in online discussions (Wilson et al., 1984). Early methods used lexicons for sentiment, toxicity, politeness (§4.2 and 4.3) and collaboration evaluation (Lawrence et al., 2017; Avelle et al., 2024). For example, spikes in hedges (e.g., ‘maybe’, ‘I guess’) invite clarification requests by facilitators, while bursts of second-person pronouns, similarly to turn-taking, often foreshadow personal attacks and can prompt a civility nudge (Niculae and Danescu-Niculescu-Mizil, 2016).

Speech and Dialogue Acts Rooted in Speech Act Theory (Austin, 1975; Searle, 1969), dialogue acts have been employed to assess deliberative quality and analyze facilitation strategies (Fournier-Tombs and MacKenzie, 2021; Chen et al., 2024a). They characterize dialogue turns (e.g., interruption) to analyze interaction dynamics (Ferschke et al., 2012; Stolcke et al., 2000; Zhang et al., 2017; Al-Khatib et al., 2018). Positive (e.g., causal reasoning) or negative (e.g., disrespect) dialogue acts can be scored to reflect discussion quality and low scores may indicate the need for interventions (Ziems et al., 2024; Cimino et al., 2024; Martinenghi et al., 2024; Schroeder et al., 2024).

Pragmatic Comprehension Pragmatic comprehension—how context shapes meaning—is crucial to facilitation, as intended meanings often diverge from literal expressions (i.e., implicature). Humans resolve such ambiguity using social and commonsense knowledge. Grice’s maxims (Grice, 1975), a central pragmatic concept, can help explain this process by outlining the conversational principles people rely on to infer meaning, while

they have already been used to assess discussion quality (Jwalapuram, 2017; Langevin et al., 2021; Ngai et al., 2021; Nam et al., 2023).

4.2 Social Dynamics

Politeness Politeness serves as a cornerstone of prosocial behavior, an attribute that facilitators desire to foster in online discussion forums (Lambert et al., 2024). In the context of facilitation, it has mainly been studied in relation to conversational derailment (§7) (Zhang et al., 2018) and constructiveness (§4.3) (De Kock and Vlachos, 2021; Zhou et al., 2024).

Power and Status Power and status influence conversational dynamics, affecting language use and turn-taking (§4.1). Higher status speakers can control the flow of discussions and foster social inequalities. Interestingly, low-status individuals tend to mimic the linguistic styles of high-status speakers more than the opposite (Danescu-Niculescu-Mizil et al., 2012), and this can be used as a signal that there is high/low-status disparity in a discussion. Facilitators may intervene, then, to ensure that the right to speak is evenly distributed among participants, preventing projection of social biases and stereotypes.

Disagreement Disagreements, when constructive, improve discussions by fostering deeper understanding (Friess, 2018; De Kock and Vlachos, 2021). Assessing disagreement, however, is complex. The hierarchy of Graham, 2008 considers disagreement tactics ranging from name calling to refuting the central point. Along with other work on dispute tactics (Walker et al., 2012; Benesch et al., 2016; De Kock et al., 2022), it can be used to examine types of disagreements in a discussion.

4.3 Emotion and Behavior

Empathy Empathy is the ability to understand others' perspectives and emotions and respond correspondingly (Lipman, 2003; Xu and Jiang, 2024). Facilitators desire to foster empathy in online discussions, since it encourages prosocial behavior and boosts engagement (Xu and Jiang, 2024; Concannon and Tomalin, 2024; Lambert et al., 2024). To do so, they encourage users to share personal stories and experiences (Schroeder et al., 2024). Various coding schemes (Macagno et al., 2022), psychological indicators (e.g., the emotion-laden words of Furniturewala and Jaidka, 2024), and dimensions (e.g., perceived engagement such in Xu

and Jiang, 2024) have been used to detect both expressed and perceived empathetic traits.

Toxicity Toxicity in online discussions refers to harmful or disrespectful language that hinders productive discourse and can derail meaningful discussions (Avalle et al., 2024). Facilitation is key to maintaining healthy communication, requiring both early detection of toxicity and (in the case of more active facilitation) proactive de-escalation strategies, such as conversation redirection or positive engagement (§5). In the case of conventional moderation that only aims to flag or remove toxic content, debate persists over what content warrants removal (Warner et al., 2025; Habibi et al., 2024; Pradel et al., 2024).

Sentiment Sentiment analysis helps identify whether discussions are positive, negative, or neutral. In the context of facilitation, sentiment analysis gauges the tone of discussions, which influences the quality of interactions (De Kock and Vlachos, 2021). Positive sentiment contributions in online discussion forums usually signal prosocial behavior and hence are highly encouraged by facilitators (Lambert et al., 2024), while negative sentiments among discussants contribute to conversation toxicity (Avalle et al., 2024).

Controversy Controversy arises from divergent viewpoints, leading to polarized exchanges that can escalate to toxicity and derail online discussions (Avalle et al., 2024). Controversial comments have been shown to contribute to a decline in positive emotions and a sustained rise in anger (Hessel and Lee, 2019; Chen et al., 2024b). The spread of political leanings among discussants and sentiment distribution analysis are common approaches to measure controversy (Avalle et al., 2024).

Constructiveness Constructiveness fosters meaningful dialogue, especially in online discussions, by promoting resolution and cooperation (Shahid et al., 2024). It is often signalled by linguistic markers (§4.1) (De Kock et al., 2022; Falk et al., 2024). A facilitator can exploit a constructiveness score; threads trending upward are worth highlighting or summarizing, whereas a downward drift may trigger facilitation tactics such as slower, structured turn-taking or clarification prompts (De Kock and Vlachos, 2021).

4.4 Engagement and Impact

Engagement Engagement is desirable in online discussion platforms as it combines interest and participation (Lambert et al., 2024; Park et al., 2012). It is proxied by measures like reciprocity (Graham and Witschge, 2003; Stromer-Galley, 2007; Zhang et al., 2018), number of comments posted by each user (Avalle et al., 2024), discussion length (Adomavicius, 2021; Avalle et al., 2024), while Ferron et al. (2023) define subdimensions such as response diversity, interestingness, and specificity.

Persuasion Empirical literature has primarily examined factors influencing persuasion that align with other categories in our taxonomy, such as linguistic markers (§4.1) and turn-taking (§4.2)(Tan et al., 2016). Considering this connection, persuasion is not only an indicator of argument quality, but may also serve as a proxy for identifying additional markers signaling whether facilitator intervention is needed.

Diversity and Informativeness Diversity in online discussions refers to the presence of varied perspectives, backgrounds, and experiences, which can enrich conversations by fostering constructive exchanges (Irani et al., 2024; Zhang et al., 2024). To prevent echo chambers and promote inclusivity, facilitators can use diversity measures to encourage opinion diversity (Anastasiou et al., 2023), encouraging users to explore a broad range of perspectives on a given issue (Kim et al., 2021). Informativeness refers to the relevance and value of information shared in a discussion and is considered a building stock of prosociality, an attribute that facilitation tries to foster in online discussion platforms (Lambert et al., 2024).

4.5 LLM Approaches to Discussion Quality

LLMs can significantly aid in evaluating discussion quality, performing on par with humans in annotating argument structure, coherence, and flow across tasks like argument mining and synthesis (Mirzakhmedova et al., 2024; Rescala et al., 2024; Wang et al., 2023; Chen et al., 2024a; Irani et al., 2024; Anastasiou and De Liddo, 2024; Zhang et al., 2024; Mendonca et al., 2024; Zhang et al., 2023). LLMs can also reliably label dialogue acts (Ziems et al., 2024; Cimino et al., 2024; Martinenghi et al., 2024; Schroeder et al., 2024), as well politeness, power, disagreement, and toxicity (Zhou et al.,

2024; Ziems et al., 2024). However, they struggle with tasks involving social norms (such as irony and humor) and show limited accuracy in sentiment and engagement detection (Hu et al., 2023; Sravanthi et al., 2024; Furniturewala and Jaidka, 2024; Xu and Jiang, 2024). While LLMs show promise in measuring controversy and persuasion, performance drops at the discussion level, especially when aiming to measure diversity, informativeness, and generally when social and pragmatic understanding is necessary (Ziems et al., 2024; Avalle et al., 2024; Lawrence and Reed, 2020).

5 Intervention Strategies

5.1 When to Intervene

Picking the right moment to intervene is a crucial part of effective facilitation strategies. If a facilitator does not intervene when they should have, there is a risk of significant escalation, while intervening when unnecessary can increase toxicity (Schaffner et al., 2024; Trujillo and Cresci, 2022; Schluger et al., 2022; Cresci et al., 2022). It is imperative then (also considering the evaluation dimensions discussed in §4), for a facilitator to be able to recognize subtle cues hinting towards escalation, in order to defuse the situation, something that even experienced human facilitators are not confident to reliably do (Schluger et al., 2022). The NLP task of ‘Conversational Forecasting’ may contribute towards this direction. Given a conversation up to a point, a model attempts to predict if an event will occur in the future in that conversation. In our case, this is where a facilitator would intervene (Schluger et al., 2022). Traditional ML models can perform well on this task, although their performance varies (Falk et al., 2021; Park et al., 2012; Falk et al., 2024; Schluger et al., 2022).

5.2 How to Intervene

There is currently no agreed-upon taxonomy for facilitator interventions. Lim et al. (2011) propose a taxonomy that focuses on discussion facilitation, excluding, however, disciplinary or administrative actions, which are common in online discussions. Park et al. (2012) propose another taxonomy consisting of seven moderator functions, ranging from policing the discussion to solving technical issues. These functions roughly correlate with the volunteer moderator roles, as described by Seering (2020). More practical approaches can be found in facilitator manuals (eRulemaking Initia-

tive, 2017; MIT Center for Constructive Communication, 2024) and books (White et al., 2024).

Facilitators often have to decide what form of coercive measure to take to make sure the conversation remains healthy, without having to intervene repeatedly. Human interventions typically use an unofficial ‘escalation ladder’ (Figure 1), where the facilitator will progressively move from milder facilitation tactics to threatening, and finally disciplinary action (Seering, 2020). ‘Conversational moderation’ (Cho et al., 2024), where a facilitator first converses with the offender, has proven effective and is actively encouraged in some facilitator guidelines (The Commons, 2025). This is probably why disciplinary action is typically not the first choice of a facilitator (Schluger et al., 2022) and why it should reasonably be used as a last resort. Softer kinds of interventions that facilitators frequently use first include: setting and informing users about rules (Schluger et al., 2022; Seering, 2020), welcoming new users (Schluger et al., 2022), summarizing key points (Small et al., 2023; Falk et al., 2024), balancing participation (Kim et al., 2021; Fishkin et al., 2018), and aiding users improve their points (Tsai et al., 2024; Falk et al., 2024).

5.3 Personalized Interventions

It is worth stressing that intervention strategies should not be applied en masse, without considering the characteristics of each individual. Traditionally, massive application (or threatening) of disciplinary action has led to adverse effects community- and platform-wide (Trujillo and Cresci, 2022; Falk et al., 2021) and the creation of echo-chambers (Cho et al., 2024). There are also calls for research to move away from one-size-fits-all approaches and instead move towards personalized interventions (Cresci et al., 2022). Human facilitators are often able to personalize interventions per individual (Schluger et al., 2022), and we hypothesize that LLMs can also do so to some extent.

6 Towards LLM-based facilitation

Until recently, ML models used as facilitation agents were confined to either performing menial tasks, such as pasting automated messages (Seering, 2020; Schluger et al., 2022), suggesting facilitation actions (e.g., rejecting posts), possibly via human-in-the-loop frameworks (Fishkin et al.,



Figure 2: Capabilities of simpler ML, LLM, and human facilitation. Task complexity and cost increase from left to right. Intermediate tasks are handled suboptimally by the preceding method.

2018; Gelauff et al., 2023), identifying possibly escalatory comments (Schluger et al., 2022), or employing pre-programmed facilitative tactics, as in the work of Kim et al. (2021), where the model produces automated messages encouraging participation. However, older ML-based and rule-based facilitation are not effective enough to meet the high demands of most platforms (Seering, 2020; Schaffner et al., 2024).

Advances in LLMs enable the development of *facilitation agents* that more actively engage in discussions. These agents can warn users about policy violations (Kumar et al., 2024), suggest rephrasings to improve tone or persuasiveness (Bose et al., 2023), monitor turn-taking (Schroeder et al., 2024), and summarize or visualize key discussion points (Small et al., 2023). They can also assist in drafting group statements that reflect diverse viewpoints (Tessler et al., 2024). A brief, non-exhaustive summary of the capabilities of simpler ML models, LLMs, and humans can be found in Figure 2.

6.1 Administrating the Discussion

LLMs are able to tackle a variety of ‘administrative’ facilitation tasks that help structure discussions. For example, facilitators often summarize the views of the participants, seek confirmation of understanding, and share perspectives. This iterative summarization is a task LLMs may handle effectively (Small et al., 2023; Burton et al., 2024). However, Feng and Qin (2022) point out some challenges such as discussions with multiple participants, topic drifts, multiple co-references, diverse interactive signals, and diverse domain terminologies. Still, according to Jin et al. (2024), LLMs bring significant advantages over conventional ML methods, “notably in the quality and flexibility of

the generated texts and the prompting paradigm to alleviate the cost of training deep models”.

In some deliberative contexts, facilitators are also encouraged to begin a discussion with their own opinion (Small et al., 2023), although others disagree (MIT Center for Constructive Communication, 2024). This is a task LLMs can also handle, albeit less convincingly than Information Retrieval (IR) approaches (Karadzhov et al., 2021).

Finally, LLMs can help marginalized groups in discussions by offering translations of the discussions in their native languages, and by helping them phrase their opinions with proper grammar and syntax (Tsai et al., 2024; Burton et al., 2024). This can directly improve discussions by increasing their diversity (Section 4.4).

6.2 Evolving Traditional Automation Models

LLMs have been proven to be adept at NLP tasks such as the detection of hate speech (Shi et al., 2024), toxicity (Kang and Qian, 2024; Wang and Chang, 2022), and misinformation (Kang and Qian, 2024; Wang and Chang, 2022). These abilities make LLMs usable as drop-in replacements for traditional ML models for these tasks, suggesting that conversational LLM facilitation agents may be able to identify, and dynamically adapt to such phenomena properly. We note however that LLMs are much more expensive and less scalable than their simpler ML counterparts. Furthermore, LLM annotation has its own challenges: LLM survey responses (Jansen et al., 2023; Bisbee et al., 2024; Neumann et al., 2025) and annotations (Gligorić et al., 2024) are generally unreliable and surface-level. Non-deterministic behavior is also common in LLMs (Atil et al., 2025), but also particularly in closed-source models (Bisbee et al., 2024) on which a lot of research on LLM annotation hinges.

6.3 Fully Automated LLM-based Facilitation

There are indications that LLMs can be used as facilitators in the fullest capacity of the role. LLMs are able to predict optimal facilitation tactics (Schroeder et al., 2024), like traditional ML models (Al-Khatib et al., 2018). Furthermore, they have proven capable of developing and executing social strategies in other tasks, e.g., negotiation games, LLM interactions (Abdelnabi et al., 2024; Cheng et al., 2024a; Martinenghi et al., 2024). Given that relatively simple ML chatbots, which do not leverage generative text capabilities, have been reported to improve discussions (Kim et al., 2021), many

expect LLM-based facilitation to be a promising solution to the well-known bottleneck of human facilitation (Small et al., 2023; Seering, 2020; Burton et al., 2024; Schroeder et al., 2024). Notably, Cho et al. (2024) successfully use LLM facilitators with prompts based on Cognitive Behavioral Therapy to moderate a live discussion with human participants. Their work shows that LLM facilitators can adapt in their instructions to users, although they cannot by themselves affect the discussion with regard to cooperation and mutual respect between the participants.

Nevertheless, LLMs have inherent limitations that make them worse than humans in most social tasks (Figure 2) (Rossi et al., 2024). While human facilitators are encouraged to be neutral (White et al., 2024; eRulemaking Initiative, 2017), numerous studies point to biases in sociodemographic, statistical, and political terms in LLMs (R.Anthis et al., 2025; Hewitt et al., 2024; Rossi et al., 2024), which can be exacerbated during the course of a discussion (Taubenfeld et al., 2024).

7 Facilitation Datasets

In this section, we provide an overview of the most prominent datasets for online facilitation, considering their sizes and their relevance to core facilitation tasks. We propose the following new taxonomy of facilitation datasets: **Conversation Derailment** datasets, where the task is to predict when a conversation escalates, therefore requiring facilitator intervention; and **Facilitator Interventions** datasets, which include comments by facilitators in active discussions, sometimes annotated with the tactics employed. Some datasets contain information that can be used in multiple tasks. An overview of the surveyed datasets and their categories in our taxonomy can be found in Table 1.

8 LLM Discussion Facilitation Roadmap

Evaluation LLMs can serve as automated discussion quality annotators (§4). Are these annotators infallible? Not yet. Certain dimensions, especially those that are highly subjective (e.g., pragmatic understanding), remain challenging for LLMs to annotate accurately. But we must take into account that even human annotations tend to be polarized for such subjective quality dimensions (Argyle et al., 2023), largely due to sociodemographic background effects and personal biases (Beck et al., 2024; Sap et al., 2020).

Name	Task		Size	Content
Wikipedia Disputes (De Kock and Vlachos, 2021)	Conversation Derailment		7, 425 D	Includes annotations for several ‘dispute tactics.’
WikiConv (Hua et al., 2018)	Facilitator Interventions		91, 000, 000 D	Includes moderation meta-data such as comment edits and deletions.
Conversations Gone Awry (Zhang et al., 2018)	Conversation Derailment		4, 188 D	Predicts derailment by analyzing rhetorical tactics, human-annotated.
Chang and Danescu-Niculescu-Mizil (2019) (1)	Conversation Derailment		4, 188 D	Extends the ‘Conversations Gone Awry’ dataset.
Chang and Danescu-Niculescu-Mizil (2019) (2)	Conversation Derailment		6, 842 D	Based on the r/ChangeMyView subreddit.
Park et al. (2012)	Conversation Derailment	Facilitator Interventions	1, 678 C	Comprised of 4 datasets. Includes 19 intervention types belonging to 7 moderator roles.
RegulationRoom (Falk et al., 2021)	Conversation Derailment		3, 000 C	Extends the dataset of Park et al. (2012).
DeliData (Karadzhov et al., 2021)	Facilitator Interventions		500 D	Group discussions, includes task-oriented quality measure which may be used to approximate discussion quality.
Wiki-Tactics (De Kock et al., 2022)	Facilitator Interventions		213 D	Based on Wikipedia Disputes, includes moderation action metadata such as comment edits and deletions.
UMOD (Falk et al., 2024)	Facilitator Interventions		2, 000 C	Based on the r/ChangeMyView subreddit, annotated for facilitation tactics and AQ.
Fora (Schroeder et al., 2024)	Facilitator Interventions		262 D	Original dataset revolving around experience-sharing, annotated for facilitation tactics.

Table 1: Overview of reviewed datasets. Unnamed datasets are referred to by the names of the authors only. The size reflects the number of annotated conversations, disregarding unlabeled data. **D** indicates the number of discussions. **C** indicates the number of individual comments or dialogue turns.

On the other hand, prompted LLMs offer a more scalable and cost-effective alternative for annotating discussion quality compared to human annotation and traditional (or self-) supervised training on large annotated datasets. Using LLMs for annotation, however, requires careful model selection considering whether models are open or closed source, model size, model alignment, as well as prompt selection, and (if applicable) fine-tuning requirements. These choices should be tailored to the specific quality dimension being evaluated.

Facilitation Intervention types should be adapted to the different legal frameworks, rules, and social norms of each community/platform. While there are exhaustive surveys on intervention types and policies, such as that of Schaffner et al. (2024), there is yet no methodology to train human or artificial facilitators according to these factors. We posit that experiments using exclusively LLM user/facilitator-agents are necessary to sustainably test new facilitation strategies and interventions per community and platform, as in other NLP tasks that involve LLM-generated conversation (Ulmer et al., 2024; Cheng et al., 2024b; Park et al., 2022, 2023), before testing the resulting facilitators in costly experiments with human participants. Finally, the datasets presented in

Table 1 can be used to train and assess LLM facilitators in the future, as well as to generate additional data—similar to the existing ones, but with controlled modifications—to stress-test various facilitators in particular settings (e.g., predicting or recovering from a conversation derailment).

9 Conclusions

This survey examined online discussion evaluation and facilitation by bridging insights from Social Science and NLP, with a focus on the growing role of LLMs. We introduced a new discussion evaluation taxonomy, with categories that should remain flexible depending on the evaluation task and the characteristics of the discussion. In terms of intervention strategies, both human- and machine-driven advancements show significant promise in improving the quality of interventions, helping online discussions remain constructive, and resistant to derailment. Most facilitation datasets still originate from human online conversations, with research yet to fully explore the capabilities of LLMs. Taking the above into account, we believe that now is the time to embrace LLMs for facilitation to foster healthier and more constructive conversations.

10 Limitations

This survey is not without its limitations. While we have attempted to present a comprehensive overview of facilitation methods, certain techniques, such as summarization, could be explored in greater depth. Since summarization is a vast subfield of NLP, it was only briefly mentioned.

Moreover, it is important to highlight that most research on facilitation has been conducted solely in English-speaking online spaces. The inherent limitations of LLMs in handling other languages and cultural contexts must be considered. As a result, these findings may not be easily applicable to other regions of the world.

Finally, the majority of real-world online discussions and deliberations happen in the context of communities, where group dynamics (social behaviors, power structures, norms, and interactions) apply. Thus, a fuller review of facilitation would have to account for the internal dynamics of such communities, as well as the wider role of the facilitator as a figure that not only helps in the conversation but has a social status in the group as well.

11 Ethical Considerations

Although AI and LLMs in particular can be effectively used as discussion facilitators, offering dynamic, responsive discussion support, their deployment must meet strict transparency, safety, and accountability standards, especially for high-risk applications, as stated in the EU AI Act.¹ For example, a person or minority group may have been unfairly disadvantaged in an AI-enhanced deliberation. It is also necessary for the users to be aware that they are interacting with AI facilitators. Ideally the consent of the users should be sought before using any sort of AI-enhanced discussion platform.

Even if LLMs facilitators eventually achieve a high level of autonomy, it is advisable to maintain human oversight. Keeping a human-in-the-loop ensures greater transparency and enables effective error prevention, detection, and correction.

References

S. Abdelnabi, A. Gomaa, S. Sivaprasad, L. Schönherr, and M. Fritz. 2024. Cooperation, competition, and

¹<https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>

- maliciousness: LLM-stakeholders interactive negotiation. In *The Thirty-eight Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, pages 96–106, Vancouver, Canada.
- S. Adomavicius. 2021. [Putting the social in social media: How human connection triggers engagement](#). In *Proceedings of the New York State Communication Association*, volume 2017.
- G. Aher, R.I. Arriaga., and A.T. Kalai. 2023. Using large language models to simulate multiple humans and replicate human subject studies. In *Proceedings of the 40th International Conference on Machine Learning*, pages 337 – 371, Hawaii, USA.
- K. Al-Khatib, H. Wachsmuth, K. Lang, J. Herpel, M. Hagen, and B. Stein. 2018. [Modeling deliberative argumentation strategies on Wikipedia](#). In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 2545–2555, Melbourne, Australia.
- L. Anastasiou and A. De Liddo. 2024. A hybrid human-AI approach for argument map creation from transcripts. In *Proceedings of the First Workshop on Language-driven Deliberation Technology (DELITE)@ LREC-COLING 2024*, pages 45–51, Turin, Italy.
- L. Anastasiou, A. De Moor, B. Brayshay, and A. De Liddo. 2023. [A tale of struggles: an evaluation framework for transitioning from individually usable to community-useful online deliberation tools](#). In *Proceedings of the 11th International Conference on Communities and Technologies, C&T '23*, page 144–155, New York, NY, USA. Association for Computing Machinery.
- L.P. Argyle, C.A. Bail, E.C. Busby, J.R. Gubler, T. Howe, C. Rytting, T. Sorensen, and D. Wingate. 2023. Leveraging AI for democratic discourse: Chat interventions can improve online political conversations at scale. *Proceedings of the National Academy of Sciences*, 120(41):1–8.
- C. S. C. Asterhan and B. B. Schwarz. 2010. [Online moderation of synchronous e-argumentation](#). *International Journal of Computer-Supported Collaborative Learning*, 5:259–282.
- B. Atil, S. Aykent, A. Chittams, L. Fu, R. J. Passonneau, E. Radcliffe, G. R. Rajagopal, A. Sloan, T. Tudrej, F. Ture, Z. Wu, L. Xu, and B. Baldwin. 2025. [Non-determinism of "deterministic" llm settings](#). *Preprint*, arXiv:2408.04667.
- J. L. Austin. 1975. *How to Do Things with Words*. Oxford University Press.
- M. Avale, N. Di Marco, G. Etta, E. Sangiorio, S. Alipour, A. Bonetti, L. Alvisi, A. Scala, A. Baronchelli, M. Cinelli, et al. 2024. Persistent interaction patterns across social media platforms and over time. *Nature*, 628(8008):582–589.

793	Y. Bang, T. Yu, A. Madotto, Z. Lin, M. Diab, and P. Fung. 2023. Enabling classifiers to make judgments explicitly aligned with human values. In <i>Proceedings of the 3rd Workshop on Trustworthy NLP</i> , pages 311–325, Toronto, Canada.	846
794		847
795		848
796		849
797		850
		851
798	R. Bawden. 2021. Understanding dialogue: Language use and social interaction. <i>Computational Linguistics</i> , 47(3):703–705.	852
799		853
800		854
801	T. Beck, H. Schuff, A. Lauscher, and I. Gurevych. 2024. Sensitivity, performance, robustness: Deconstructing the effect of sociodemographic prompting. In <i>Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 2589–2615, Malta.	855
802		856
803		857
804		858
805		859
806		860
807		861
808	S. Benesch, D. Ruths, K. P. Dillon, H. M. Saleem, and L. Wright. 2016. Counterspeech on twitter: A field study. <i>dangerous speech project</i> .	862
809		
810		
811	J. Bisbee, J. D. Clinton, C. Dorff, B. Kenkel, and J. M. Larson. 2024. Synthetic replacements for human survey data? the perils of large language models . <i>Political Analysis</i> , 32(4):401–416.	863
812		864
813		865
814		866
		867
815	R. Bose, I. Perera, and B. Dorr. 2023. Detoxifying online discourse: A guided response generation approach for reducing toxicity in user-generated text. In <i>Proceedings of the First Workshop on Social Influence in Conversations</i> , pages 9–14, Toronto, Canada.	868
816		869
817		870
818		
819		
820	J. W. Burton, E. Lopez-Lopez, S. Hechtlinger, et al. 2024. How large language models can reshape collective intelligence. <i>Nature Human Behaviour</i> , 8:1643–1655.	871
821		872
822		873
823		874
824	A. Bächtiger, M. Gerber, and E. Fournier-Tombs. 2022. 83discourse quality index . In <i>Research Methods in Deliberative Democracy</i> . Oxford University Press.	875
825		876
826		
827	A. Bächtiger, S. Niemeyer, M. Neblo, M. R. Steenbergen, and J. Steiner. 2010. Disentangling diversity in deliberative democracy: Competing theories, their blind spots and complementarities . <i>Journal of Political Philosophy</i> , 18(1):32–63.	877
828		878
829		879
830		
831		
832	A. Cervone and G. Riccardi. 2020. Is this dialogue coherent? learning from dialogue acts and entities . In <i>Proceedings of the 21th Annual Meeting of the Special Interest Group on Discourse and Dialogue</i> , pages 162–174, online. Association for Computational Linguistics.	880
833		881
834		882
835		883
836		
837		
838	J. P. Chang and C. Danescu-Niculescu-Mizil. 2019. Trouble on the horizon: Forecasting the derailment of online conversations as they develop . In <i>Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)</i> , pages 4743–4754, Hong Kong, China. Association for Computational Linguistics.	884
839		885
840		886
841		887
842		888
843		889
844		900
845		901
	G. Chen, L. Cheng, L. A. Tuan, and L. Bing. 2024a. Exploring the potential of large language models in computational argumentation. In <i>Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 2309–2330.	
	M. Chen, L. Frermann, and J. H. Lau. 2024b. WHoW: A cross-domain approach for analysing conversation moderation . <i>Preprint</i> , arXiv:2410.15551.	
	P. Cheng, T. Hu, H. Xu, Z. Zhang, Y. Dai, L. Han, and N. Du. 2024a. Self-playing adversarial language game enhances llm reasoning . <i>ArXiv</i> , abs/2404.10642.	
	P. Cheng, T. Hu, H. Xu, Z. Zhang, Y. Dai, L. Han, and N. Du. 2024b. Self-playing adversarial language game enhances llm reasoning . <i>ArXiv</i> , abs/2404.10642.	
	H. Cho, S. Liu, T. Shi, D. Jain, B. Rizk, Y. Huang, Z. Lu, N. Wen, J. Gratch, E. Ferrara, and J. May. 2024. Can language model moderators improve the health of online discourse? In <i>Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)</i> , pages 7478–7496, Mexico City, Mexico.	
	G. Cimino, C. Li, G. Carenini, and V. Deufemia. 2024. Coherence-based dialogue discourse structure extraction using open-source large language models . In <i>Proceedings of the 25th Annual Meeting of the Special Interest Group on Discourse and Dialogue</i> , pages 297–316, Kyoto, Japan.	
	S. Concannon and M. Tomalin. 2024. Measuring perceived empathy in dialogue systems . <i>AI & Society</i> , 39:2233–2247.	
	S. Cresci, A. Trujillo, and T. Fagni. 2022. Personalized interventions for online moderation . In <i>Proceedings of the 33rd ACM Conference on Hypertext and Social Media</i> , page 248–251, New York, NY, USA.	
	C. Danescu-Niculescu-Mizil, L. Lee, B. Pang, and J. Kleinberg. 2012. Echoes of power: language effects and power differences in social interaction . In <i>Proceedings of the 21st International Conference on World Wide Web</i> , page 699–708, New York, NY, USA.	
	C. De Kock, T. Stafford, and A. Vlachos. 2022. How to disagree well: Investigating the dispute tactics used on Wikipedia. In <i>Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing</i> , pages 3824–3837, Abu Dhabi, United Arab Emirates.	
	C. De Kock and A. Vlachos. 2021. I beg to differ: A study of constructive disagreement in online conversations . In <i>Proceedings of the 16th Conference of the European Chapter of the Association for Computational Linguistics: Main Volume</i> , pages 2017–2027, Online.	

- C. Degeling, S. M. Carter, and L. Rychetnik. 2015. Which public and why deliberate? – a scoping review of public deliberation in public health and health policy research. *Social Science & Medicine*, 131:114–121.
- Cornell eRulemaking Initiative. 2017. [Ceri \(cornell e-rulemaking\) moderator protocol](#). Cornell e-Rulemaking Initiative Publications, 21.
- N. Falk, I. Jundi, E. M. Vecchi, and G. Lapesa. 2021. Predicting moderation of deliberative arguments: Is argument quality the key? In *Proceedings of the 8th Workshop on Argument Mining*, pages 133–141, Punta Cana, Dominican Republic.
- N. Falk and G. Lapesa. 2023. Bridging argument quality and deliberative quality annotations with adapters. In *Findings of the Association for Computational Linguistics: EACL 2023*, pages 2469–2488.
- N. Falk, E. Vecchi, I. Jundi, and G. Lapesa. 2024. Moderation in the wild: Investigating user-driven moderation in online discussions. In *Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 992–1013, St. Julian’s, Malta.
- X. Feng and B. Qin. 2022. A survey on dialogue summarization: Recent advances and new frontiers. In *Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence, IJCAI-22*, pages 5453–5460. Survey Track.
- A. Ferron, A. Shore, E. Mitra, and A. Agrawal. 2023. MEEP: Is this engaging? prompting large language models for dialogue evaluation in multilingual settings. In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pages 2078–2100, Singapore. Association for Computational Linguistics.
- O. Ferschke, I. Gurevych, and Y. Chebotar. 2012. Behind the article: Recognizing dialog acts in Wikipedia talk pages. In *Proceedings of the 13th Conference of the European Chapter of the Association for Computational Linguistics*, pages 777–786, Avignon, France.
- J. Fishkin, N. Garg, L. Gelauff, A. Goel, K. Munagala, S. Sakshuwong, A. Siu, and S. Yandamuri. 2018. Deliberative democracy with the online deliberation platform. In *The 7th AAAI Conference on Human Computation and Crowdsourcing (HCOMP 2019)*. HCOMP.
- E. Fournier-Tombs and M. K. MacKenzie. 2021. Big data and democratic speech: Predicting deliberative quality using machine learning techniques. *Methodological Innovations*, 14(2):20597991211010416.
- D. M. Friess. 2018. Letting the faculty deliberate: Analyzing online deliberation in academia using a comprehensive approach. *Journal of Information Technology & Politics*, 15(2):155–177.
- S. Furniturewala and K. Jaidka. 2024. Empaths at WASSA 2024 empathy and personality shared task: Turn-level empathy prediction using psychological indicators. In *Proceedings of the 14th Workshop on Computational Approaches to Subjectivity, Sentiment, & Social Media Analysis*, pages 404–411, Bangkok, Thailand. Association for Computational Linguistics.
- L. Gelauff, L. Nikolenko, S. Sakshuwong, J. Fishkin, A. Goel, K. Munagala, and A. Siu. 2023. *Achieving parity with human moderators*, pages 202–221. Routledge.
- M. H. Gelula. 1997. Clinical discussion sessions and small groups. *Surgical Neurology*, 47(4):399–402.
- M. Gerber, A. Bächtiger, S. Shikano, S. Reber, and S. Rohr. 2018. Deliberative abilities and influence in a transnational deliberative poll (europolis). *British Journal of Political Science*, 48(4):1093–1118.
- H. Gimpel, S.n Lahmer, M. Wöhl, et al. 2024. Digital facilitation of group work to gain predictable performance. *Group Decision and Negotiation*, 33:113–145.
- K. Gligori’c, T. Zrnica, C. Lee, E. J. Candès, and D. Jurafsky. 2024. Can unconfident llm annotations be used for confident conclusions? *ArXiv*, abs/2408.15204.
- J.I. Goñi. 2024. What is “dialogue” in public engagement with science and technology? bridging sts and deliberative democracy. *Minerva*.
- P. Graham. 2008. How to disagree. Accessed: 2024-06-24.
- T. Graham and T. Witschge. 2003. In search of online deliberation: Towards a new method for examining the quality of online discussions. *Communications*, 28(2):173–204.
- HP Grice. 1975. Logic and conversation. *Syntax and semantics*, 3.
- M. Habibi, D. Hovy, and C. Schwarz. 2024. The content moderator’s dilemma: Removal of toxic content and distortions to online discourse. *Preprint*, arXiv:2412.16114.
- J. Hessel and L. Lee. 2019. Something’s brewing! early prediction of controversy-causing posts from discussion features. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 1648–1659, Minneapolis, Minnesota. Association for Computational Linguistics.
- L. Hewitt, A. Ashokkumar, I. Ghezae, and R. Willer. 2024. Predicting results of social science experiments using large language models. Equal contribution, order randomized.

1009	J. Hu, S. Floyd, O. Jouravlev, E. Fedorenko, and E. Gibson. 2023. A fine-grained comparison of pragmatic language understanding in humans and language models . In <i>Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 4194–4213, Toronto, Canada.	1062
1010		1063
1011		1064
1012		1065
1013		
1014		1066
1015		1067
1016	Y. Hua, C. Danescu-Niculescu-Mizil, D. Taraborelli, N. Thain, J. Sorensen, and L. Dixon. 2018. WikiConv: A corpus of the complete conversational history of a large online collaborative community . In <i>Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing</i> , pages 2818–2823, Brussels, Belgium.	1068
1017		1069
1018		1070
1019		
1020		1071
1021		1072
1022		1073
1023		1074
1024	A. Irani, M. Faloutsos, and K. Esterling. 2024. Argusense: Argument-centric analysis of online discourse . In <i>Proceedings of the International AAAI Conference on Web and Social Media</i> , volume 18, pages 663–675.	1075
1025		
1026		1076
1027		1077
1028		1078
1029	B. J. Jansen, S. Jung, and J. Salminen. 2023. Employing large language models in survey research . <i>Natural Language Processing Journal</i> , 4:100020.	1079
1030		1080
1031		
1032	H. Jin, Y. Zhang, D. Meng, J. Wang, and J. Tan. 2024. A comprehensive survey on process-oriented automatic text summarization with exploration of llm-based methods . <i>arXiv preprint</i> .	1081
1033		1082
1034		1083
1035		1084
1036	P. Jwalapuram. 2017. Evaluating dialogs based on Grice’s maxims . In <i>Proceedings of the Student Research Workshop Associated with RANLP 2017</i> , pages 17–24, Varna.	1085
1037		
1038		1086
1039	S. Kaner, Le. Lind, C. Toldi, S. Fisk, and D. Berger. 2007. <i>Facilitator’s Guide to Participatory Decision-Making</i> . John Wiley & Sons/Jossey-Bass, San Francisco.	1087
1040		
1041		1088
1042		1089
1043	H. Kang and T. Qian. 2024. Implanting LLM’s knowledge via reading comprehension tree for toxicity detection . In <i>Findings of the Association for Computational Linguistics ACL 2024</i> , pages 947–962, Bangkok, Thailand and virtual meeting.	1090
1044		1091
1045		1092
1046		1093
1047		1094
1048	G. Karadzhov, T. Stafford, and A. Vlachos. 2021. Delidata: A dataset for deliberation in multi-party problem solving . <i>Proceedings of the ACM on Human-Computer Interaction</i> , 7:1 – 25.	1095
1049		
1050		1096
1051		1097
1052	B. Khalid and S. Lee. 2022. Explaining dialogue evaluation metrics using adversarial behavioral analysis . In <i>Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies</i> , pages 5871–5883, Seattle, United States. Association for Computational Linguistics.	1098
1053		1099
1054		
1055		1100
1056		1101
1057		
1058		1102
1059	R. Kies. 2022. Online deliberative matrix. In <i>Research Methods in Deliberative Democracy</i> , pages 148–162. Oxford University Press.	1103
1060		1104
1061		1105
		1106
		1107
	S. Kim, J. Eun, J. Seering, and J. Lee. 2021. Moderator chatbot for deliberative discussion: Effects of discussion structure and discussant facilitation . <i>Proc. ACM Hum.-Comput. Interact.</i> , 5(CSCW1).	1108
		1109
	D. Kumar, Y. A. AbuHashem, and Z. Durumeric. 2024. Watch your language: Investigating content moderation with large language models. <i>Proceedings of the International AAAI Conference on Web and Social Media</i> , 18(1):865–878.	1110
		1111
	C Lambert, Fk Choi, and E Chandrasekharan. 2024. "positive reinforcement helps breed positive behavior": Moderator perspectives on encouraging desirable behavior . <i>Proc. ACM Hum.-Comput. Interact.</i> , 8(CSCW2).	1112
		1113
		1114
	R. Langevin, R. J Lordon, T. Avrahami, B. R. Cowan, T. Hirsch, and G. Hsieh. 2021. Heuristic evaluation of conversational agents . In <i>Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems</i> , New York, NY, USA.	
	J. Lawrence, J. Park, K. Budzynska, C. Cardie, B. Konat, and C. Reed. 2017. Using argumentative structure to interpret debates in online deliberative democracy and erulemaking . <i>ACM Trans. Internet Technol.</i> , 17(3).	
	J. Lawrence and C. Reed. 2020. Argument mining: A survey . <i>Computational Linguistics</i> , 45(4):765–818.	
	Z. Li, J. Zhang, Z. Fei, Y. Feng, and J. Zhou. 2021. Conversations are not flat: Modeling the dynamic information flow across dialogue utterances . In <i>Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)</i> , pages 128–138, Online.	
	S.C.R. Lim, W. Cheung, and K. Hew. 2011. Critical thinking in asynchronous online discussion: An investigation of student facilitation techniques. <i>New Horizons in Education</i> , 59:52–65.	
	M. Lipman. 2003. <i>Thinking in Education</i> , 2 edition. Cambridge University Press.	
	Y. Liu, S. Ultes, W. Minker, and W. Maier. 2023. Unified conversational models with system-initiated transitions between chit-chat and task-oriented dialogues . In <i>Proceedings of the 5th International Conference on Conversational User Interfaces, CUI ’23</i> , New York, NY, USA.	
	J. Lo and P. McAvoy. 2023. <i>Debate and Deliberation in Democratic Education</i> , page 298–310. Cambridge Handbooks in Education. Cambridge University Press.	
	F. Macagno, C. Rapanta, E. Mayweg-Paus, and M. Garcia-Milà. 2022. Coding empathy in dialogue . <i>Journal of Pragmatics</i> , 192:116–132.	

1115	N. Mansour. 2024. Students' and facilitators' experiences with synchronous and asynchronous online dialogic discussions and e-facilitation in understanding the nature of science . <i>Education and Information Technologies</i> , 29:15965–15997.	1170
1116		1171
1117		1172
1118		1173
1119		1174
1120	A. Martinenghi, G. Donabauer, S. Amenta, S. Bursic, M. Giudici, U. Kruschwitz, F. Garzotto, and D. Og-nibene. 2024. LLMs of catan: Exploring pragmatic capabilities of generative chatbots through prediction and classification of dialogue acts in boardgames' multi-party dialogues . In <i>Proceedings of the 10th Workshop on Games and Natural Language Processing @ LREC-COLING 2024</i> , pages 107–118, Torino, Italia.	1175
1121		
1122		
1123		
1124		
1125		
1126		
1127		
1128		
1129	B. Mathew, R. Dutt, P. Goyal, and A. Mukherjee. 2019. Spread of hate speech in online social media. In <i>Proceedings of the 10th ACM Conference on Web Science</i> , page 173–182, New York, NY, USA.	1176
1130		1177
1131		1178
1132		1179
1133	J. Mendonca, I. Trancoso, and A. Lavie. 2024. ECoh: Turn-level coherence evaluation for multilingual dia-logues . In <i>Proceedings of the 25th Annual Meeting of the Special Interest Group on Discourse and Dia-logue</i> , pages 516–532, Kyoto, Japan.	1180
1134		
1135		
1136		
1137		
1138	N. Mirzakhmedova, M. Gohsen, C. H. Chang, and B. Stein. 2024. Are large language models reliable argument quality annotators? In <i>Conference on Ad-vances in Robust Argumentation Machines</i> , pages 129–146. Springer.	1181
1139		1182
1140		1183
1141		1184
1142		1185
1143	MIT Center for Constructive Communication. 2024. Unpublished training materials developed by the mit center for constructive communication. Guide given to human facilitators.	1186
1144		
1145		
1146		
1147	M.D. Molina and S.S. Sundar. 2022. When AI moder-ates online content: effects of human collaboration and interactive transparency on user trust. <i>Journal of Computer-Mediated Communication</i> , 27(4).	1187
1148		1188
1149		1189
1150		1190
1151	Y. Nam, H. Chung, and U. Hong. 2023. Language artificial intelligences' communicative performance quantified through the gricean conversation theory . <i>Cyberpsychology, Behavior, and Social Networking</i> , 26(12):919–923. PMID: 37976199.	1191
1152		1192
1153		1193
1154		1194
1155		1195
1156	T. Neumann, M. De-Arteaga, and S. Fazelpour. 2025. Should you use llms to simulate opinions? qual-ity checks for early-stage deliberation . <i>Preprint</i> , arXiv:2504.08954.	1196
1157		
1158		
1159		
1160	E.W.T. Ngai, M.C.M. Lee, M. Luo, P.S.L. Chan, and T. Liang. 2021. An intelligent knowledge-based chat-bot for customer service . <i>Electronic Commerce Re-search and Applications</i> , 50:101098.	1197
1161		1198
1162		1199
1163		1200
1164	V. Niculae and C. Danescu-Niculescu-Mizil. 2016. Conversational markers of constructive discussions . In <i>Proceedings of the 2016 Conference of the North American Chapter of the Association for Computa-tional Linguistics: Human Language Technologies</i> , pages 568–578, San Diego, California.	1201
1165		1202
1166		1203
1167		1204
1168		1205
1169		
	J. Park, S. Klingel, C. Cardie, M. Newhart, C. Farina, and J.J. Vallbé. 2012. Facilitative moderation for online participation in erulemaking . In <i>Proceedings of the 13th Annual International Conference on Digi-tal Government Research</i> , page 173–182, New York, NY, USA.	1206
		1207
	J.S. Park, J.C. O'Brien, C.J. Cai, M.R. Morris, P. Liang, and M.S. Bernstein. 2023. Generative agents: Inter-active simulacra of human behavior . <i>Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology</i> .	1208
		1209
		1210
	J.S. Park, L. Popowski, C.J. Cai, M.R. Morris, P. Liang, and M.S. Bernstein. 2022. Social simulacra: Creat-ing populated prototypes for social computing sys-tems . In <i>Proceedings of the 35th Annual ACM Sym-posium on User Interface Software and Technology</i> , UIST '22, New York, NY, USA.	1211
		1212
		1213
		1214
	F. Pradel, J. Zilinsky, S. Kosmidis, and Y. Theocharis. 2024. Toxic speech and limited demand for con-tent moderation on social media . <i>American Political Science Review</i> , 118(4):1895–1912.	1215
		1216
		1217
		1218
	N. Raj Prabhu, C. Raman, and H. Hung. 2021. Defining and quantifying conversation quality in spontaneous interactions . In <i>Companion Publication of the 2020 International Conference on Multimodal Interaction</i> , ICMI '20 Companion, page 196–205, New York, NY, USA.	1219
		1220
		1221
		1222
		1223
		1224
	J. R.Anthis, R. L., S. M. Richardson, A. C. Kozlowski, B. Koch, J. Evans, E. Brynjolfsson, and M. Bern-stein. 2025. Llm social simulations are a promising research method . <i>Preprint</i> , arXiv:2504.02234.	
	P. Rescala, M.H. Ribeiro, T. Hu, and R. West. 2024. Can language models recognize convincing argu-ments? In <i>Findings of the Association for Computa-tional Linguistics: EMNLP 2024</i> , pages 8826–8837, Miami, Florida, USA.	
	H. Rheingold. 2000. <i>The Virtual Community: Home-steading on the Electronic Frontier</i> . The MIT Press.	
	R. Rose-Redwood, R. Kitchin, L. Rickards, U. Rossi, A. Datta, and J. Crampton. 2018. The possibilities and limits to dialogue . <i>Dialogues in Human Geogra-phy</i> , 8(2):109–123.	
	L. Rossi, K. Harrison, and I. Shklovski. 2024. The prob-lems of llm-generated data in social science research . <i>Sociologica</i> , 18(2):145–168.	
	U. Russmann and A. Lane. 2016. Discussion, dialogue, and discoursel doing the talk: Discussion, dialogue, and discourse in action — introduction . <i>Interna-tional Journal of Communication</i> , 10.	
	M. Saeidi, M. Yazdani, and A. Vlachos. 2021. Cross-policy compliance detection via question answering. In <i>Proceedings of the 2021 Conference on Empiri-cal Methods in Natural Language Processing</i> , pages 8622–8632, Online and Punta Cana, Dominican Re-public.	

1335	L.L. Tsai, A. Pentland, A. Braley, N. Chen,	K. White, N. Hunter, and K. Greaves. 2024. <i>facilitating</i>	1391
1336	J.R. Enríquez, and A. Reuel. 2024. Generative	<i>deliberation - a practical guide</i> . Mosaic Lab.	1392
1337	AI for Pro-Democracy Platforms. <i>An</i>		
1338	<i>MIT Exploration of Generative AI</i> . https://mit-	T. P. Wilson, J. M. Wiemann, and D. H. Zimmerman.	1393
1339	genai.pubpub.org/pub/mn45hexw .	1984. <i>Models of turn taking in conversational inter-</i>	1394
		<i>action</i> . <i>Journal of Language and Social Psychology</i> ,	1395
1340	J.A. Tucker, A. Guess, P. Barberá, C. Vaccari, A. Siegel,	3(3):159–183.	1396
1341	S. Sanovich, D. Stukal, and B. Nyhan. 2018. Social		
1342	media, political polarization, and political disinforma-	Z. Xu and J. Jiang. 2024. <i>Multi-dimensional evaluation</i>	1397
1343	tion: A review of the scientific literature. <i>SSRN</i>	<i>of empathetic dialogue responses</i> . In <i>Findings of the</i>	1398
1344	<i>Electronic Journal</i> .	<i>Association for Computational Linguistics: EMNLP</i>	1399
		2024, pages 2066–2087, Miami, Florida, USA. As-	1400
1345	D. Ulmer, E. Mansimov, K. Lin, L. Sun, X. Gao, and	sociation for Computational Linguistics.	1401
1346	Y. Zhang. 2024. <i>Bootstrapping LLM-based task-</i>		
1347	<i>oriented dialogue agents via self-talk</i> . In <i>Findings of</i>	Y. Yeh, M. Eskenazi, and S. Mehri. 2021. <i>A comprehen-</i>	1402
1348	<i>the Association for Computational Linguistics: ACL</i>	<i>sive assessment of dialog evaluation metrics</i> . In <i>The</i>	1403
1349	2024, pages 9500–9522, Bangkok, Thailand.	<i>First Workshop on Evaluations and Assessments of</i>	1404
		<i>Neural Conversation Systems</i> , pages 15–33, Online.	1405
1350	E.M. Vecchi, N. Falk, I. Jundi, and G. Lapesa. 2021.	Association for Computational Linguistics.	1406
1351	Towards argument mining for social good: A survey.		
1352	In <i>Proceedings of the 59th Annual Meeting of ACL</i>	A. Zhang, B. Culbertson, and P. Paritosh. 2017. <i>Char-</i>	1407
1353	<i>and 11th International Joint Conference on NLP</i> ,	<i>acterizing online discussion using coarse discourse</i>	1408
1354	pages 1338–1352, Online.	<i>sequences</i> . <i>Proceedings of the International AAAI</i>	1409
		<i>Conference on Web and Social Media</i> , 11(1):357–	1410
1355	A. Veglis. 2014. Moderation techniques for social me-	366.	1411
1356	dia content. In <i>Social Computing and Social Media</i> ,		
1357	pages 137–148, Cham. Springer International Pub-	C. Zhang, L. D’Haro, C. Tang, K. Shi, G. Tang, and	1412
1358	lishing.	H. Li. 2023. <i>xDial-eval: A multilingual open-</i>	1413
		<i>domain dialogue evaluation benchmark</i> . In <i>Findings</i>	1414
1359	H. Wachsmuth, G. Lapesa, E. Cabrio, A. Lauscher,	<i>of the Association for Computational Linguistics:</i>	1415
1360	J. Park, E.M. Vecchi, S. Villata, and T. Ziegenbein.	<i>EMNLP 2023</i> , pages 5579–5601, Singapore.	1416
1361	2024. Argument quality assessment in the age of		
1362	instruction-following large language models. In <i>Pro-</i>	C. Zhang, L. F. D’Haro, Y. Chen, M. Zhang, and H. Li.	1417
1363	<i>ceedings of the 2024 Joint International Conference</i>	2024. A comprehensive analysis of the effective-	1418
1364	<i>on Computational Linguistics, Language Resources</i>	ness of large language models as automatic dialogue	1419
1365	<i>and Evaluation</i> , pages 1519–1538, Torino, Italia.	evaluators. In <i>Proceedings of the AAAI Conference</i>	1420
		<i>on Artificial Intelligence</i> , volume 38, pages 19515–	1421
1366	H. Wachsmuth, N. Naderi, Y. Hou, Y. Bilu, V. Prab-	19524.	1422
1367	hakaran, T.A. Thijm, G. Hirst, and B. Stein. 2017.		
1368	Computational argumentation quality assessment in	J. Zhang, J. Chang, C. Danescu-Niculescu-Mizil,	1423
1369	natural language. In <i>Proceedings of the 15th Confer-</i>	L. Dixon, Y. Hua, D. Taraborelli, and N. Thain. 2018.	1424
1370	<i>ence of the European Chapter of the Association for</i>	<i>Conversations gone awry: Detecting early signs of</i>	1425
1371	<i>Computational Linguistics: Volume 1, Long Papers</i> ,	<i>conversational failure</i> . In <i>Proceedings of the 56th An-</i>	1426
1372	pages 176–187.	<i>annual Meeting of the Association for Computational</i>	1427
		<i>Linguistics (Volume 1: Long Papers)</i> , pages 1350–	1428
1373	M. Walker, J. F. Tree, P. Anand, R. Abbott, and J. King.	1361, Melbourne, Australia.	1429
1374	2012. A corpus for research on deliberation and		
1375	debate. In <i>Proceedings of the Eighth International</i>	L. Zhou, Y. Farag, and A. Vlachos. 2024. <i>An LLM</i>	1430
1376	<i>Conference on Language Resources and Evaluation</i>	<i>feature-based framework for dialogue constructive-</i>	1431
1377	<i>(LREC’12)</i> , pages 812–817, Istanbul, Turkey.	<i>ness assessment</i> . In <i>Proceedings of the 2024 Con-</i>	1432
		<i>ference on Empirical Methods in Natural Language</i>	1433
1378	Y. Wang, X. Chen, B. He, and L. Sun. 2023. Context-	<i>Processing</i> , pages 5389–5409, Miami, Florida, USA.	1434
1379	tual interaction for argument post quality assessment.	Association for Computational Linguistics.	1435
1380	In <i>Proceedings of the 2023 Conference on Empiri-</i>		
1381	<i>cal Methods in Natural Language Processing</i> , pages	C. Ziems, W. Held, O. Shaikh, J. Chen, Z. Zhang, and	1436
1382	10420–10432.	D. Yang. 2024. <i>Can large language models trans-</i>	1437
		<i>form computational social science? Computational</i>	1438
1383	Ya. Wang and Y.T. Chang. 2022. <i>Toxicity detection</i>	<i>Linguistics</i> , 50(1):237–291.	1439
1384	<i>with generative prompt-based inference</i> . <i>ArXiv</i> ,		
1385	abs/2205.12390 .		
		A Acronyms	1440
1386	M. Warner, A. Strohmayr, M. Higgs, and L. Coventry.		
1387	2025. <i>A critical reflection on the use of toxicity</i>	NLP Natural Language Processing	1441
1388	<i>detection algorithms in proactive content moderation</i>		
1389	<i>systems</i> . <i>International Journal of Human-Computer</i>	ML Machine Learning	1442
1390	<i>Studies</i> , 198:103468.		

LLM Large Language Model

AM Argument Mining

ML Machine Learning

IR Information Retrieval

AQ Argument Quality

B Keywords for Literature Query

Keyword Selection

online discussions, deliberation, dialogue, discussion evaluation, discussion metrics, dialogue, deliberation, NLP, AI, discussion quality, argument mining, survey, LLM, conversation, moderation, facilitation, communication, democracy AI dialogue systems, group dynamics

Table 2: Keywords for search engine queries

C Terminology Background

Here we explain how the surveyed articles were selected. We also explain our reasoning for choosing and disambiguating certain terms (see §2). The definitions of the terms can be found in Table 3.

Facilitation vs. Moderation “Moderation”, as a term, is more common in Computer Science and NLP, while facilitation is prevalent in Social Sciences (Vecchi et al., 2021; Kaner et al., 2007; Trenel, 2009). Moderators enforce rules and ensure orderly interactions, usually with the threat of disciplinary action, though they can also act as community leaders (Falk et al., 2024; Seering, 2020; eRulemaking Initiative, 2017). Facilitators, on the other hand, guide discussions, promote participation, and structure dialogue, particularly in online deliberation and education platforms (Asterhan and Schwarz, 2010). Despite these distinctions, the terms are sometimes used interchangeably (Cho et al., 2024; Park et al., 2012; Kim et al., 2021), while it is also common for moderators to use facilitation tactics (eRulemaking Initiative, 2017; Park et al., 2012; Kim et al., 2021; Cho et al., 2024; Schluger et al., 2022).

Pre-moderation and Post-moderation Multiple taxonomies have been proposed to describe the temporal dimension of moderation; that is, when moderator action is applied in relation to when the content is visible to the users (Veglis, 2014; Schluger et al., 2022). These taxonomies are very

similar to each other, and usually boil down to the following distinctions:

- *Pre-moderation*: The user is dissuaded, or prevented from, posting harmful content. Pre-moderation techniques can include nudges at the writing stage (Argyle et al., 2023), reminders about platform rules (Schluger et al., 2022), or even a moderation queue where posts have to be approved before being visible to others (Schluger et al., 2022).
- *Real-Time*: The moderator is part of the discussion and intervenes like a referee would during a match.
- *Ex-post*: The moderator is called after a possible incident has been flagged and makes the final call.

Discussion, Deliberation, Dialogue, Debate

There is little to no consensus on how to properly define terms such as “discussion” and “dialogue” (Russmann and Lane, 2016; Goñi, 2024). In this section, we attempt to disambiguate the use of such terms for the purposes of our survey and based on the existing related work. First, our study focuses on **discussions**, a broader term encompassing various informal and formal exchanges, including online discussions in fora (Russmann and Lane, 2016), with which we are mainly concerned. In contrast, **dialogue** refers to collaborative interactions in which participants work toward a shared understanding and alignment (Rose-Redwood et al., 2018; Bawden, 2021; Goñi, 2024). Studies on dialogue emphasize its cooperative nature, aiming for mutual insight rather than competition (Bawden, 2021). Dialogue can also refer to dialogue systems, a major NLP sub-area, traditionally including both task-oriented dialogues and casual conversation (Eliza-like)² “chatbots” (Liu et al., 2023; Sun et al., 2021).

A more specific concept is **deliberation**, which involves structured discussions aimed at informed decision-making, often prioritizing reasoned argumentation and the consideration of diverse perspectives (Degeling et al., 2015; Lo and McAvoy, 2023). Meanwhile, **debate** is typically adversarial, where participants focus on persuading others or defending their positions. Unlike dialogue or deliberation, debate centers more on winning or

²<http://web.njit.edu/~ronkowitz/eliza.html>

Concept	Definition and Characteristics
Discussion	Broad term encompassing informal and formal exchanges, including online discussions in fora. Can involve elements of debate, dialogue, and deliberation.
Dialogue	Collaborative interaction aimed at shared understanding and alignment. Emphasizes cooperation rather than competition. Also refers to dialogue systems in NLP (task-oriented or chatbot conversations).
Deliberation	Structured discussion focusing on informed decision-making with reasoned argumentation and diverse perspectives. Less about persuasion, more about collective reasoning.
Debate	Adversarial interaction where participants aim to persuade or defend positions rather than achieve mutual understanding. Focused on rhetorical effectiveness.
Thread-style Discussions	Online discussions structured in tree/thread formats (e.g., Reddit). Can incorporate elements of all rhetorical styles (debate, dialogue, deliberation).
Discussion Quality	Subjective measure influenced by cultural background, engagement, and type of discussion. Defined by socio-dimensional aspects of participant experiences.
Moderation	Ensures orderly interactions by enforcing guidelines. Moderators can be volunteers or employees, often associated with disciplinary actions.
Facilitation	Encourages equal participation and organizes discussion flow. More common in deliberative and educational contexts, though often used interchangeably with moderation.

Table 3: Definition of terms used in this survey.

convincing, making it less about collective reasoning and more about rhetorical effectiveness (Lo and McAvoy, 2023). Debates also typically have much stricter (and enforced) rules than other discussions.

For this study, we specifically focus on online written discussions, particularly those occurring in **thread- or tree-style** formats (Seering, 2020). A thread is a collection of messages or posts grouped together in an online forum, discussion board, or messaging platform (such as Reddit). It begins with an initial post (often called the original post, or OP), and subsequent replies are ordered either chronologically or by relevance. Threads usually address a specific topic or question and allow users to engage in discussions about that subject. A thread may grow as users contribute more responses. It must be noted, however, that this type of discussion can contain elements from all the other discussion styles. For example, the adversarial element of the debates, or the argumentative element that can be found both in dialogues and deliberations.

Discussion Quality The success of a discussion is often subjective, influenced by a variety of factors such as the cultural background and linguistic proficiency of the participants (Zhang et al., 2018), as well as their level of engagement (See et al., 2019). It also depends on the type of the discussion, since some types of discussions, such as deliberations or debates, may not aim at con-

sensus. Given these complexities, we adopt the definition proposed by Raj Prabhu et al. (2021), which views the perceived *discussion quality* as a measurement that attempts to quantify interactions by taking into account multiple socio-dimensional aspects of individual experiences and abilities.

D Methodology

The search and article selection of this survey was conducted using specific keywords in academic search engines (e.g., Google Scholar, Semantic Scholar, Scopus), digital libraries and repositories (e.g., ACL Anthology, ACM Digital Library, IEEE Xplore, JSTOR). We focused on peer-reviewed publications written in English between 2014 and 2024, granting exceptions only for established works predating this period. Additionally, we reviewed other cited papers that appeared highly relevant, provided they were peer-reviewed and cited by more than 20 citations of other researchers, unless the topic was very niche, in which case we judged by its content. The search strategy incorporated keywords and phrases related to LLMs, discussion facilitation, and discussion evaluation. The list of keywords used is provided in Table 2. The search was further informed by existing survey articles, such as those by Vecchi et al. (2021) and Wachsmuth et al. (2024), which served as starting points both for identifying relevant literature and for specifying the vocabulary used in the keyword search.

E Discussion Quality Taxonomy

In this part of the Appendix, we present a table summarizing the discussion evaluation taxonomy (§4). The dimensions are outlined alongside both pre-LLM and LLM-based approaches, while also highlighting their respective contributions to facilitation. The dimensions are color-coded for clarity, with orange indicating associated dimensions that could serve as early signs of potential derailment, green marking signs of constructive growth—i.e., conversations going well or worth participating in—and pink denoting interaction dynamics.

F Online Discussion Example with Color-coded Politeness Markers

This table highlights key politeness-related linguistic features such as hedging, personal references, sentiment, and direct questions. These features are essential in the context of facilitation, where the goal is to guide conversations constructively, maintain safety, and foster mutual understanding. By identifying these elements, the facilitator (human or automatic) can better interpret the tone, intent, and emotional weight of each utterance. For example, detecting hedging or positive sentiment can guide the model to adopt a more collaborative tone, while recognizing negative sentiment or accusatory second-person references may prompt it to de-escalate tension and encourage constructive dialogue.

Dimension	Facilitation Use	Pre-LLM Approaches	LLM Approaches
Structure & Logic			
Argument structure & analysis	Spot claim-evidence chains; raise early-warning flags; keep debate fact-centred	Argument-mining pipelines; claim/premise detection; AQ scoring; graph & neural models	Zero/few-shot AQ labelling; argument-structure parsing; on-the-fly argument-map summaries
Coherence & flow	Detect topic drift; redirect or bridge gaps	Entity-grid & sequential coherence models; topic modelling; dialogue state tracking	Prompted coherence scoring; chain-of-thought flow checks; off-topic suggestions
Turn-taking	Monitor balance (entropy/Gini); nudge silent voices; avoid dominance	Turn-entropy / Gini metrics; rule-based alarms	Context-window turn counts; balanced-participation prompts
Language features	Track hedges, 2nd-person spikes, jargon; trigger clarification or civility nudges	Lexicon features; n-gram-based hedging detectors	Style-transfer rephrasers; embedding hedge detection; tone-repair suggestions
Speech & dialogue acts	Identify interruptions, proposals, question types; score deliberative quality	Dialogue-act tagging with ISO/DAMSL labels	Few-shot Dialogue Act tagging; tactic selection based on Dialogue Act patterns
Pragmatic comprehension	Resolve implicatures & sarcasm; surface hidden misunderstandings	Commonsense reasoning (Knowledge Base + neural); limited coverage	In-context reasoning; auto clarifying questions
Social Dynamics			
Politeness	Forecast derailment; issue civility nudges or positive reinforcement	Politeness lexicons; domain-independent classifiers	Annotation & polite rewrites; policy-violation explanations
Power & status	Detect dominance; invite low-status voices; rebalance floor	Style-matching, pronoun analysis; social-role features	Power imbalance estimation; moderator suggestions
Disagreement	Distinguish constructive vs destructive dissent; de-escalate	Graham-hierarchy / stance detection	Few-shot labelling; automatic reframing prompts
Emotion & Behavior			
Empathy	Encourage empathic turns; highlight emotional cues	Lexicon/coding empathy classifiers; affective features	Perceived-empathy scoring; supportive paraphrases
Toxicity	Flag harmful language; decide moderation step	BERT/toxicity classifiers; detox lexicons	Detection + rewrite suggestions; policy chat
Sentiment	Track emotional climate; intervene at negativity spikes	Lexicon & neural sentiment analysis	Prompt-based labelling; tone-shift detection
Controversy	Sense polarization; invite balancing views	Topic-polarity metrics; ideology models	Ideology tagging; polarity-aware summaries
Constructiveness	Stream score; escalate or summarize based on trend	Feature-based classifiers (linguistic, discourse)	Constructive-rewrite coaching
Engagement & Impact			
Engagement	Detect lulls or dominance; prompt interaction	Turn/word counts; reply-time gaps	Auto-recaps; invite quiet users
Persuasion	Spotlight evidence-based arguments; dampen manipulation	Lexical overlap; ethos/pathos/logos; persuasion prediction	Outcome prediction; neutral framing suggestions
Diversity & Informativeness	Monitor viewpoint spread & info density	Topic-diversity indices; IR-based scoring	Simulate perspectives; propose links

Table 4: Summary of discussion quality dimensions and corresponding pre-LLM and LLM-based facilitation strategies.

Turn	Utterance
0	Why should we help people based on race, and say “ we ’ll help everyone who’s black, because they could be poor” instead of just “ we ’ll help everyone who’s poor, in which black people make up a proportionally larger amount”?
1	That study is worse than useless unless it also distinguishes between “black sounding” names that are associated with wealth and poverty.
2	That wouldn’t discount it, that would just add another intersectional axis to investigate. >which I know without looking that it didn’t. How rational .
3	It’s certainly more rational than unquestioningly swallowing everything I read, as some people do. Did this study of yours also test difficult to pronounce Polish names, or Russian names? Or would that have interfered too much with the foregone conclusion they were attempting to reach?
4	Are you implying that’s what I have done? You may be the only one making assumptions here.

Table 5: Dissuccssion example from the Reddit Change My View dataset (Chang and Danescu-Niculescu-Mizil, 2019). Color indicates politeness-related features: **hedging**, **1st person reference**, **2nd person reference**, **direct questions**, **negative sentiment** and **positive sentiment**. The annotation was produced with a soon-to-be-released annotation toolkit for discussion evaluation.