



To Reply or to Quote: Comparing Conversational Framing Strategies on Twitter

HIMANSHU ZADE, SPENCER WILLIAMS, THERESA T. TRAN, and
 CHRISTINA SMITH, Human Centered Design & Engineering, University of Washington
 SUKRIT VENKATAGIRI, Department of Computer Science, Swarthmore College
 GARY HSIEH and KATE STARBIRD, Human Centered Design & Engineering and Center for an
 Informed Public, University of Washington

Social media platform affordances allow users to interact with content and with each other in diverse ways. For example, on Twitter,¹ users can like, reply, retweet, or quote another tweet. Though it's clear that these different features allow various types of interactions, open questions remain about how these different affordances shape the conversations. We examine how two similar, but distinct conversational features on Twitter — specifically reply vs. quote — are used differently. Focusing on the polarized discourse around Robert Mueller's congressional testimony in July 2019, we look at how these features are employed in conversations between politically aligned and opposed accounts. We use a mixed methods approach, employing grounded qualitative analysis to identify the different conversational and framing strategies salient in that discourse and then quantitatively analyzing how those techniques differed across the different features and political alignments. Our research (1) demonstrates that the quote feature is more often used to broadcast and reply is more often used to reframe the conversation; (2) identifies the different framing strategies that emerge through the use of these features when engaging with politically aligned vs. opposed accounts; (3) discusses how reply and quote features may be re-designed to reduce the adversarial tone of polarized conversations on Twitter-like platforms.

CCS Concepts: • **Human-centered computing** → **Social network analysis**; *Empirical studies in HCI*;

Additional Key Words and Phrases: Twitter, social network analysis, microblogging platforms, social media platforms, quote tweets, reply tweets, conversational framing

¹On July 24, 2023, Twitter was rebranded to “X”. In this article, we refer to the platform that we studied as Twitter.

This work was completed while the author was at the Center for an Informed Public, University of Washington. This work was supported by National Science Foundation (NSF) awards IIS-1749815 and IIS-1715078. Sukrit Venkatagiri was additionally supported by Craig Newmark Philanthropies and the John S. and James L. Knight Foundation at the CIP. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the authors and do not necessarily reflect the views of the above supporting organizations or the NSF.
 Authors' addresses: H. Zade, S. Williams, T. T. Tran, and C. Smith, Human Centered Design and Engineering, University of Washington, Seattle 98105, Washington; e-mails: {himanz, sw1918, thettran, csmith79}@uw.edu; S. Venkatagiri, Department of Computer Science, Swarthmore College, Swarthmore 19081, Pennsylvania; e-mail: sukritv@uw.edu; G. Hsieh and K. Starbird, Human Centered Design and Engineering and Center for an Informed Public, University of Washington, Seattle 98105, Washington; e-mails: {garyhs, kstarbi}@uw.edu.



This work is licensed under a Creative Commons Attribution-NoDerivs International 4.0 License.

© 2024 Copyright held by the owner/author(s).

2834-5533/2024/1-ART9

<https://doi.org/10.1145/3625680>

ACM Reference format:

Himanshu Zade, Spencer Williams, Theresa T. Tran, Christina Smith, Sukrit Venkatagiri, Gary Hsieh, and Kate Starbird. 2024. To Reply or to Quote: Comparing Conversational Framing Strategies on Twitter. *ACM J. Comput. Sustain. Soc.* 2, 1, Article 9 (January 2024), 27 pages.
<https://doi.org/10.1145/3625680>

1 INTRODUCTION

Social media platforms offer users a range of features to engage with the content and with each other. For example, Reddit users can upvote and downvote content, Facebook users can reshare, comment, or add an emotional response (e.g., like, love, care, haha, wow, sad, angry) to a post. These different features — and especially how users employ these different features — shape the kinds of interactions that take place on the platform. At scale, these interactions structure the broader networks, communities, and overall discourse that occurs on the platform [38, 73], and beyond.

It is possible that some of these features exacerbate problematic behavior on social media platforms. For example, researchers have repeatedly documented political polarization within social media platforms [4, 43, 52, 83]. Conover et al. found that politically motivated individuals used mention networks on Twitter to provoke polarized conversations by injecting partisan content into information streams that involve ideologically-opposed users [21]. Journalists are increasingly highlighting how online platforms can potentially radicalize users within groups formed around conspiracy theories [19, 20, 72, 88].

To alleviate some of the toxicities and promote a healthy discourse, online platforms continue to experiment with their features. For example, Twitter has experimented with their conversational affordances by promoting users to quote a tweet rather than simply retweet it to make users consider why they want to amplify the tweet [36]. They discontinued the change after the 2020 US Presidential election upon learning that it did not promote thoughtful amplification as anticipated [37]; about 45% of the additional quote tweets comprised a single word, and 70% contained less than 25 characters. This natural experiment suggests that platforms are aware that small affordances can impact the health of the discourse they facilitate, even if they do not know how to improve things. More recently, there is discussion about how certain features catalyze negative user behavior [5], especially as the emerging microblogging platform *Mastodon* debates if they should introduce the quote feature on its platform in a way that preserves its “antiviral” design [81].

This research examines how two similar, but distinct features for engagement on Twitter — the reply tweet and the quote tweet — afford different types of interactions within and across groups of users with different political alignments. Twitter introduced the quote feature — referred to as *retweet with comment* — in April 2015 in addition to the already offered reply feature. Figure 1 illustrates the difference between a response posted toward an original tweet as a “reply” and a response posted as a “quote”. In the remainder of this article, we refer to the original tweet as the *root tweet*, and the reply tweet and the quote tweet collectively as the *response tweet*.

Garimella et al. examined the role of introducing the quote feature on Twitter discourse in 2016 [38] — by looking for the presence of disagreement and insult — and found that in the early days, the quoting mechanism led to more civil discourse as compared to replies. We extend their line of research and examine how the features of reply and quote (before the changes made in August 2020) are employed by looking into broader conversational and framing techniques. We examine each pair of a root-tweet and a response-tweet and ask three related **research questions**:

- RQ1. How do conversational and framing strategies differ, if at all, between the two response types: reply and quote?



Fig. 1. The left image refers to the Reply tweet, while the right image refers to the Quote tweet.

- RQ2. How do conversational and framing strategies vary when there is political (mis)alignment between the author of the root tweet and the author of the response tweet?
- RQ3: How do conversational and framing strategies vary between reply and quote tweets when accounting for their authors' political (mis)alignment (i.e., interaction effect)?

To answer these questions, we focused on the Twitter discourse around Robert Mueller's testimony (July 24, 2019) about the "Mueller Report" [60] for curating a collection of tweet responses that are both politically charged and bring politically-diverse perspectives. We selected two subsets of 170 root tweet-reply tweet pairs and 170 root tweet-quote tweet pairs where half of the tweets (85 from each set) came from politically-aligned Twitter accounts, and the other half came from politically-opposed Twitter accounts. We developed a codebook using a grounded theory-lite approach and by adopting Entman's *framing* lens [30]. We then coded the responses (reply tweet or quote tweet) based on how they extended the root tweet — with a focus on identifying different conversational and framing techniques.

With respect to our first research question, we found that users employed reply and quote tweet affordances to target different audiences and resulting in differences in framing. Reply tweets were more often directed toward the author of the root tweet (73.5% vs. 43.3%), while quote tweets served to "broadcast" the root tweet's message to a different audience (56.7% vs. 26.5%). Reply tweets also more frequently reframed the topic of the root tweet compared to quote tweets (46% vs. 34%). These results were statistically significant. Concerning our second research question, the distribution of codes confirmed a highly polarized discourse, e.g., providing supporting evidence when politically aligned or being condescending when politically opposed. With respect to our third research question, we found asymmetric adoption of the replying or quoting mechanism when the response was politically aligned or opposed. For instance, when quoting a root tweet that they opposed, users were more likely to include evidence for refuting the root tweet. However, when quoting a root tweet they supported, users included evidence less often and instead focused on amplifying the message. We found no such variation in the use of evidence when replying to politically aligned or opposed users.

In summary, our work makes three contributions:

- (1) We develop a coding framework — guided by our interpretation of data and by previous work on framing [30] — that identifies conversational and framing strategies adopted by the response tweets toward the root tweet. Our approach expands upon previous work by looking at conversational strategies that surpass disagreement and use of insult [38] and facilitate a broader analysis of Twitter discourse.
- (2) Using our coding framework, we facilitate an understanding of how users employ the reply and quote features of Twitter to build, refine, revise, and amplify conversations that they prefer *and* challenge or counter conversations that they oppose.

- (3) We offer insights into the implications of our findings, focusing on how quotes could serve as a mechanism to promote polarization by building a refutation toward what one dislikes and merely amplifying what one likes on online platforms like Twitter and *maybe* Mastodon in the near future.

2 RELATED WORK

Conversations have been studied through discourse analysis by many researchers. We start by understanding the role of the platform features in influencing these conversations. Next, we discuss the use of framing and linguistic techniques to influence online conversations. Later, we look into the contentious process of framing and counter-framing adopted toward influencing online conversations.

2.1 Role of Platform Features in Shaping Online Conversations

The design of social media platforms includes different features to allow users to engage with the content and to interact with each other, and these features shape the conversations that take place through them. For example, Twitter affords the formation of communities ephemerally around the life of a topic through the use of hashtags [13], while Reddit fosters a more lasting sense of community through the “subreddit” feature over shared interests [11]. Similar features on different platforms — or even on the same platform but at different periods of their evolution — can lead to different perceptions of utility and patterns of use. For example, Instagram users (during the first few years of that platform’s lifetime) posted pictures on the platform as a way of archiving these images [51], while on Snapchat the inability to keep an image within the conversation for an extended period of time made its users perceive and invoke those images less as photographic objects and more as a form of conversation [64]. In this article, we examine how two similar, but distinct conversational features on Twitter — specifically reply vs. quote — are used in different ways to engage in and shape the discourse there.

The reply feature on Twitter was initially a user-driven convention. The platform began to provide functional support for the feature in 2007 [80] and continues to iterate on its design [86]. Twitter introduced the quote feature — or “retweet with comment” as it was called back then — in April 2015 [77]. One of the early studies by Garimella et al. investigated the impact of introducing this feature on the online discourse [38]. They examined the direct network connections between the root tweeter and the quote tweeter and found that the quote feature in 2016 afforded pushing the political discourse away from the root user’s network toward new audiences. They also found that the quote feature leads to more civil responses than the reply feature by examining the responses for the presence of disagreement and use of insults. We extend their line of research and compare how quote tweets and reply tweets engage with the subject topic of the root tweets, focusing on *the linguistic, conversational, and framing strategies employed in the response tweet*. Our analysis thus looks into several dimensions beyond use of disagreement and insult, e.g., use of sarcasm, use of evidence, and so on.

Comparative analysis of the affordances of Facebook and YouTube platform around political expression suggests that salient social identification on Facebook led to a more civil discourse there than on Youtube as seen through threaded conversations witnessed on these two platforms in 2010 [44]. Does shifting attention from the root tweet author’s audience to the response tweet author’s audience — as afforded by quotes — correlate with a user responding in a more civil manner [38], or has this changed with the times? Here, we characterize how users engage in conversations differently through both the reply and the quote features.

2.2 Influencing Online Conversations through Framing and Linguistic Techniques

Frames, as defined by Goffman [39], are interpretive lenses through which we, as individuals and groups, classify information and make sense of experiences. The concept of framing has been invoked as a device for making sense of the political discourse around us [31, 35, 62, 66].

The concept of frames — to highlight and make salient certain information, while obscuring other information — has been studied either by treating it as a goal in itself, or as a process to achieve a goal. The former involves focusing on the *frames* that emerged in a given discourse. For example, Egbunike et al. identified different frames to discuss the unfolding of the #OccupyNigeria protests as it was reported in newspapers and social media [29]. The second approach of studying this construct is to treat *framing* as a process through which frames are created, refined, and reshaped. Entman describes framing as “selecting some aspects of a perceived reality and making them more salient in a communicating text, in such a way as to promote a particular problem definition, causal interpretation, moral evaluation, and/or treatment recommendation” [30]. Framing, in this view, can be a strategy of affecting how others interpret the world. For example, emphasizing different aspects in the title of similar policy initiatives can impact political decision-making and shape how people vote [14]. Our research focuses on the conceptualization of framing that takes place in online environments through a collective process [59, 68, 71, 78]. For example, Meraz et al. described how crowds consisting of both regular and “elite” or influential Twitter accounts came together to co-construct, revise, and distribute frames across the platform. [58]

Most closely related to our work, researchers have investigated online discourse through linguistics techniques [26, 76]. Conversation-based linguistic characteristics have been effectively used toward identifying coordination in a disaster-related discourse [70]. Ferguson et al. successfully used a coding-framework modeled after linguistic techniques to identify exploratory dialogue that occurred in an online conference [34]. Variations of Ferguson’s framework have been found useful to compare the interaction patterns across subreddits that facilitate learning, e.g., *r/AskScience*, *r/AskAcademia*, and so on. [45]. Here, we developed a codebook to identify different conversation and framing techniques that emerged through response tweets as users tend to revise and reshape the discourse across different interactive features.

2.3 Influencing Conversations through Framing Contests

Responses to a post on social media platforms can range from a simple agreement indicating an endorsement [10, 84] to a strong disagreement indicating disapproval [8, 47]. This article captures how both the elements of agreeing and disagreeing could be used together in a nuanced manner to shape the conversation differently and to steer the narrative in a desired direction.

The process of achieving a shared sense of reality is contentious, i.e., it offers individuals involved in the sensemaking an opportunity to create and challenge different considerations of the relevant events [7]. The interpretive and contentious process of framing happens both within a group [49] and across different opposed groups [78]. While such a framing contest can be productive to register the dissenting perspectives [22, 49], it can be detrimental when used exclusively to counter-attack as a mechanism for undermining the opposition [9, 18, 85]. Given the agency of participating groups in this contestation, framing contests can be studied best by accounting for the structure of the discourse in which the interpretive sensemaking occurs [7, 35]. Our approach accounts for the structure of the Twitter discourse by identifying the two most prominent perspectives about Mueller’s testimony.

Stewart et al. studied counter-framing practices on Twitter by mapping the structure of the discourse and examined how two opposed groups negotiate the context of police-related shootings [78]. We adopt this approach and identify two distinct groups within the Mueller discourse on

Twitter that are ideologically opposed. To identify group affinity, we use the feature of *retweeting network interactions* that are known to roughly predict someone's attitude on Twitter [23, 24, 54]. By identifying the two most prominent opposed groups that engage in the political discourse around Mueller's testimony, we then study how reply and quote features represent different framing practices within politically aligned groups and across politically opposed groups.

3 DATA COLLECTION AND FILTERING

3.1 Collecting the Twitter Discourse about Robert Mueller's Politically Charged Testimony

To answer our two questions, we needed to identify an online conversation that contained: (a) diverse usage of the reply and quote features on tweeter, and (b) involved politically charged conversations. Thus, we focused our study on the online discourse about the polarizing testimony of Special Counsel Robert Mueller on July 24, 2019 to the **United States (U.S.)** Congress about his investigation into Russian interference in the 2016 U.S. presidential election [60]. This discourse coalesced into two primary positions, one that sought to interpret the report as implicating U.S. President Trump in crimes related to election interference in 2016 and supported Mueller, and another that sought to interpret the report as exonerating U.S. President Trump in those crimes [48]. While one side responded in support of the charges implied by Mueller in the report, the other side cited Mueller's refusal to state those charges outright as rationale for opposition toward the claims in the report [69]. This diversity of public opinion about Mueller's testimony reflected in the online discourse about it and allowed us to curate a rich data set with enough conversations facilitated by the reply and quote mechanisms between politically aligned and opposed accounts.

We used the Twitter streaming API to collect 4.2 million tweets containing the search term "Mueller" or "mueller" during a period of five days (from July 22, 2019 to July 26, 2019) corresponding to Mueller's testimony and two days before and after that testimony.

3.2 Identifying the Two Most Prominent Perspectives in the Discourse

We then mapped out the underlying structure of the conversation by generating a network graph to identify information sharing trajectories and infer the political affinity of different accounts in the data. To infer that structure, researchers [28, 43, 61, 65, 78] have frequently used modularity analysis techniques[6]; higher modularity implies a stronger connection and ideological similarity between the accounts of that module. To aid modularity analysis, we only included accounts that were connected to other accounts and indicated ideological similarity to other accounts. Therefore, we selected accounts that were retweeted three times, or were retweeted twice by the same user, or contributed at least three unique tweets. The filtered data set consisted of 261,199 active unique Twitter accounts that produced 368,773 tweets.

Using the modularity analysis algorithm in the Gephi network visualization software [6], we generated a retweet network graph. This algorithm clusters together accounts that retweet common account(s) (indicating ideological proximity to each other) while pushing apart accounts that lack any common connections (indicating ideological dissimilarity). Consistent with prior research [21, 78], we focused on two primary clusters that represent the dominant polarity within the discourse as shown in Figure 2.

To validate that the two primary clusters aligned with the two politically-opposed narratives around Mueller's testimony, we examined the frequent bi-grams and tri-grams and found that the phrases popular in each cluster were unique. Phrases like "obstruction of justice", "election security bills", "Russian hack", "read the Mueller Report", and so on. occurred most frequently in the green cluster (right side in Figure 2). These phrases delineate the polarizing (left-leaning) narratives



Fig. 2. The retweet graph for tweet-data about Mueller’s testimony. The “green” cluster includes Twitter accounts extending support to Mueller and the “pink” cluster includes Twitter accounts that were opposed to Mueller in the context of Mueller’s testimony. We refer to these clusters as “left-leaning” and “right-leaning” clusters respectively from here on.

that supported Mueller’s testimony. Phrases like “the steele dossier”, “highly conflicted Robert”, “he doesn’t know”, “retweet realdonaldtrump”, and so on. dominated the pink cluster (left side in Figure 2). These phrases delineated the (right-leaning) counter-narrative that discredited the testimony and opposed Mueller’s testimony. We also inspected the top ten most retweeted accounts, their profile descriptions, and their recent tweets from both the clusters to determine whether they were politically left-leaning or right-leaning, and then mapped these political affiliations onto all the accounts in that cluster.

3.3 Filtering the Collection for Qualitative Coding

Next, we identified a smaller sample of root tweet-response tweet pairs for qualitative analysis. Instead of focusing on the most influential accounts (the ones that received massive engagement), we focused our analysis on accounts that were quoted fewer than 1,000 times (corresponding to the 99th percentile in our data set). In order to ensure that we can retrieve a sufficient number of replies for each of these accounts (explained below), we also excluded accounts that got less attention and were quoted tweeted fewer than two times by accounts in both the left- and right-leaning clusters (indicating low reach, from 0 to the 63rd percentile).

This left us with 3,372 right-leaning and 4,787 left-leaning Twitter accounts. To reduce the scope for our qualitative analysis, we randomly selected 100 accounts — 50 from each of the two clusters. For each of these accounts, we randomly choose one Twitter account within the same cluster that quoted it, and another from a politically-opposed cluster that quoted it. We refer to these responses as *quote tweets*. To balance the number of quote tweets for both clusters, we coded 170 out of the 200 quote tweets and reply tweets (as explained below).

Retrieving the replies. Twitter’s Streaming API returned different metadata for quotes and replies at the time of collecting this data in 2019 (Twitter only introduced the ability to retrieve all replies in a conversation in it’s API v2 in mid-2020 [15]). For a tweet that contains a tracked keyword, the API returned both original tweets and any quote tweet (even if the quote tweet itself did not contain the keyword). However, the API did not return replies to a tweet that contained a keyword, unless the reply itself had the keyword. Thus, to accurately capture reply tweets that match the same inclusion criteria as for quote tweets, we retrieved all the mentions of the selected 100 Twitter

Table 1. Step-by-step Description of How we Filtered a Total of 340 Responses — 170 Replies and 170 Quotes; 170 Responses Politically Aligned to the Root Tweet Account and 170 Opposed to the Root Tweet — from a Large Pool of about 4.2 Million Tweets from the Data Set

Step#	Filtering Procedure	Resultant data sample
Step 1	We used the Twitter streaming API to collect tweets containing the term “Mueller” or “mueller” between July 22 and July 26, 2019.	4.2 million tweets and 659,316 unique Twitter accounts
Step 2	To aid clustering accounts as per their inferred political affinity (i.e., modularity analysis), we filtered accounts that were retweeted three times, or were retweeted twice by same user, or contributed at least three unique tweets.	368,773 tweets produced by 261,199 unique Twitter accounts
Step 3	To identify the two most polarizing perspectives in the discourse, we limited to accounts that belonged to the two dominant account-clusters as per the modularity analysis.	347,797 tweets produced by 111,659 unique Twitter accounts
Step 4	To skip the most influential accounts, we limited to accounts that were quoted less than 1000 times (99th percentile); to ensure some online attention, we also skipped accounts that were only quoted once or twice.	247,602 tweets produced by 3,372 right-leaning and 4,787 left-leaning unique Twitter accounts
Step 5	To facilitate qualitative analysis, we randomly selected a small number of accounts.	50 right-leaning and 50 left-leaning unique Twitter accounts
Step 6	To accurately capture reply tweets that match our inclusion criteria of quote tweets, we retrieved all the mentions of the selected 100 Twitter accounts from Twitter and identified their reply tweets posted by accounts within either of the two clusters. .	170 root tweet-reply tweet pairs (85 politically aligned, 85 opposed) and 170 root tweet-quote tweet pairs (85 politically- aligned, 85 opposed)

accounts from Twitter and identified their reply tweets that were posted by accounts within either of the two clusters in Figure 2. After accounting for the lack of replies for some accounts, our overall sample of responses consisted of 85 root tweets and four response tweets for every root tweet: (1) a quote tweet from a politically-aligned account, (2) a quote tweet from a politically-opposed account, (3) a reply tweet from a politically-aligned account, and (4) a reply tweet from a politically-opposed account. Table 1 depicts the overall process of how we filtered a small set of 170 root tweet-reply tweet pairs and 170 root tweet/quote tweet pairs from the larger pool in our data-set. The limitation of 170 replies and 170 quotes comes from the need to balance the number of replies and quotes across politically-aligned and opposed accounts and the additional step needed to collect the replies as the Twitter API in 2019 did not provide direct access to threaded replies.

Confirming the sufficiency of the sample-size. To confirm that our data sample filtered from the larger set for manual coding is large enough to allow the use of inferential statistics, we conducted a power analysis using the G-power tool [33]. Given that the assigned labels from codebook will serve as the outcome variables, we chose a two-tailed a priori analysis for the z-test family suitable for logistic regression. The power analysis determined that we needed a sample size of 308 — smaller than our sample of 340 — assuming an effect size corresponding to odds ratio of 1.5 with 80 percent power.

4 METHODS

4.1 Methodological Approach

We adopted a mixed-methods analysis to answer our research question. In the first part of the analysis, we adopted a qualitative grounded theory-lite approach[16, 46] to identify common thematic characterizations that emerged through the data as illustrated in Table 2. This led us to identifying the different conversational framing strategies prevalent in the Mueller online discourse.

Table 2. Definitions of Codes from the Codebook when Categorizing Conversational Framing Strategies in Root Tweet-response Tweet Pairs

Category (Codes)	Definition
Framing technique: Engage or reframe (Engage, Reframe)	This refers to how the response tweet extends the root tweet. It could simply add some <i>engage</i> with the message in the same frame, or it could try to change the frame of the topic and thus <i>reframe</i> .
Framing technique: Twisting of words (Yes, No) – coded if <i>Engage</i> or <i>reframe</i> is “Reframe”)	This is a hierarchical code within refocus. It captures instances where the response tweet borrows words from the root tweet and pivots the context around those words in an attempt to refocus the conversation.
Framing technique: Articulate a frame* (Explain, Describe, Project)	This captures how the response tweet extends the root tweet. If it adds details to the situation (relative to the context of the event as described in the root tweet), we code it as <i>describe</i> . If it diagnoses the cause for the situation (so that it leads to the event as described in the root tweet), we code it as <i>explain</i> . If it suggests a potential event that might result out of the event in the root tweet, we code it as <i>project</i> .
Personal Evaluation* (Character, Credibility, Competence)	We use the code <i>character</i> , <i>credibility</i> and/or <i>competence</i> if the response tweet judges a person in the root tweet based on their sense of morality, trustworthiness, and/or skills/abilities respectively.
Level of verifiability (Expression, Claim, Evidence)	<i>Claim</i> refers to a statement of truth that is yet to be verified but can be verified. <i>expression</i> is a statement that cannot be verified, whereas <i>evidence</i> is a statement of truth that has been verified (as used by the responder).
Target audience (Root tweeter, Broader audience, Both)	This refers to whom the response targets. If directed (explicit and clear) toward the root tweeter, we code it as <i>root tweeter</i> . If directed more toward a larger audience, we code it as <i>broader audience</i> ; <i>both</i> otherwise.
Directed valence (Support, Oppose, Neutral)	It captures the general sentiment of <i>support</i> or <i>oppose</i> as the response tweet expresses toward the root tweet.
Language: Tone (Positive, Negative, Neutral)	This refers to the use of <i>positive</i> language, <i>negative</i> language or <i>neutral</i> language in the response.
Language: Attitude (Reverent, Condescending, None)	We code a response tweet as <i>reverent</i> or <i>condescending</i> depending on whether the response tweet treats the root tweeter with respect OR if it attempts to establish a sense of superiority or mockery toward the root tweeter respectively.

Categories marked with an asterisk were not mutually exclusive, i.e., coders could assign any one or all of the codes from these categories to one root-response tweet pair. We ascertained a shared understanding of the different codes amongst the coders using an inter-coder reliability metric as measured by Cohen’s Kappa (≥ 0.6) across all categories [56].

In the second part, we adopted a more quantitative approach to understand which of the conversational framing strategies result from the distinct conversational features (reply and quote) on Twitter. Table 3 illustrates the distribution of codes from our coding scheme across the response type and across the type of accounts involved in that conversational exchange.

Table 3. Frequency Distribution of Codes from the Coding Scheme as Analyzed Across Both the Factors: (1) Type of Feature: Reply and Quote (2) Type of Discourse: Politically Aligned and Politically Opposed

Type of response and political alignment (Columns)	Reply		Quote		Interaction	Main	Main
Categories and Codes (Rows)	Politically aligned	Politically opposed	Politically aligned	Politically opposed	Reply/Quote X Politically aligned/opposed	Reply/Quote	Politically aligned/opposed
Target audience:							
Root tweeter	23	38	11	28			
Broader audience	12	10	33	18	No	Yes	Yes
Both	50	37	41	39			
Response frame:							
Engage	55	37	67	44	No	Yes	Yes
Reframe	30	48	18	40			
Twisting of words:							
Yes	3	3	3	10	No	Yes	No
No	27	45	15	30			
Level of verifiability:							
Evidence	11	15	4	20			
No evidence	74	70	79	64	Yes	-	-
Personal evaluation:							
Character	23	22	17	29			
No Character	62	63	68	56	Yes	-	-
Articulate a frame:							
Projection	15	19	32	16			
No projection	70	66	53	69	Yes	-	-
Tone:							
Positive	19	2	21	2			
Negative	36	54	38	50	No	No	Yes
Neutral	30	29	26	33			
Directed valence:							
Support	70	13	67	22			
Oppose	13	72	13	58	No	No	Yes
Neutral	2	0	5	5			
Directed attitude:							
Reverent	12	0	10	1			
Condescending	26	40	25	42	No	No	Yes
None	46	45	50	42			

Last 3 columns indicate the presence of any interaction effects across the two factors, main effects otherwise. Table 6 describes the regression analysis for the interaction effect, while Table 4 and Table 5 describe the regression analysis for the main effects. *Note: *Twisting of words* is a sub-category within *Response frame* and was only coded if the Response frame was coded as *Reframe*.

4.2 Approach to Coding: Identifying Conversational Framing Strategies in Root Tweet-Response Tweet Pairs

We employed Henwood’s grounded theory-lite approach [46] (derived from Charmaz’s description of grounded theory [16]) to interpret and derive the different conversational strategies and framing techniques as expressed in the response tweet (reply or quote) toward the root tweet. The generative part of the coding process involved a group of six researchers working in pairs and assigned a code toward each root tweet-response tweet pair, capturing how the response tweet interacts with the root tweet and root tweeter. The pairwise coding approach is motivated by previous research

analyzing conversational threads on microblogging sites like Twitter [82, 89]. We answered several questions of the response specific to the subject topic as described in the tweet like “To whom does the response tweeter target in the conversation?”, “Does the response tweet make any personal judgment toward a person mentioned in the root tweet (including the tweeter)?” and so on. and generated the codes.

To develop the initial coding scheme, we first analyzed 60 (out of the 340) root tweet-response tweet pairs comprising 15 unique root tweets and 4 unique response tweets for each root tweet and varying in the political affinity of the response tweeter (politically aligned or opposed) and the response feature (reply or quote). Researchers were not aware of the political affinity (inferred from the network graph) of the root tweet or whether the response was a reply or a quote to mitigate any bias in the coding process.

Second, we augmented our coding scheme to include additional categories relevant to framing and highlighted by Entman [30] that were missing in our initial, inductively-generated scheme. In particular, we added a new category *Articulate a frame* with the codes *Description*, *Explanation*, and *Projection* modeled after the Entman’s framing concepts of problem definition, causal interpretation, and treatment recommendation respectively. We also refined the *Personal evaluation* category to capture evaluations not only based on one’s *Character*, but also based on one’s *Credibility* and *Competence*. Thus, the final codebook emerged through dual processes — both inductive and deductive.

Third, using codes from the revised codebook, researchers individually coded 120 more root tweet-response tweet pairs (out of the 340) in two successive rounds—coding 60 pairs in each round. At the end of each round, researchers collectively discussed their assigned codes, disambiguated their understanding of the codes, and refined the codes and their definitions in these two iterations for achieving a shared understanding of the different codes (inter-coder reliability as measured by Cohen’s Kappa ≥ 0.6 across all categories indicating moderate agreement between coders [56]). With a refined understanding of the codebook (summary in Table 2), researchers then used consensus-coding to code the remaining and re-code the already inspected root tweet-response tweet pairs. Eventually, we coded all 340 root tweet-responses.

4.3 Approach to Analysis: Identifying the Distribution of Strategies Across Reply/Quote and Political Alignment

We analyzed the coded 340 root-response tweet pairs to identify which conversational framing strategies result from the use of distinct conversational features (reply and quote) on Twitter. We used a multinomial regression to model the conversational and framing strategies as captured by the codes (treated as nominal variables) by considering main effects of the type of response (reply or quote) and political alignment (politically-aligned or politically-opposite), and their interaction effects. In cases where the interaction effect was not significant, we drop the interaction terms from our model and focus on the main effects as depicted in Table 3. Since every root tweet occurred four times in our data (paired with two unique reply tweets and two unique quote tweets), we included the root tweet as a random effect to control for its repetition.

In the findings below, we first report on the qualitatively identified conversational and framing strategies. Next, we report on the quantitative findings to introduce any interesting and unique phenomena across replies and quotes as we observed in the data.

5 FINDINGS: A CODEBOOK TO CATEGORIZE CONVERSATIONAL AND FRAMING STRATEGIES IN ROOT-RESPONSE TWEET PAIRS

To answer our research question about how two distinct features of reply and quote could afford different types of interactions within and across groups of users with different political alignments,

we developed a codebook. Informed both by our grounded, interpretive process to surface salient themes and linguistic techniques — as well as by previous work on framing [30] — in the tweets, our codebook captures several aspects of conversational framing. Salient features include:

- We introduce five framing techniques (across two dimensions) that a tweet employed relative to the root tweet. The first dimension categorizes responses based on whether they *engage* with the topic of conversation or if they *reframe* it by bringing attention to new topics; the second dimension offers different types of framing articulations that are directly borrowed from Entman’s conceptualization: *describe*, *explain*, and *project*. We also included a third dimension that is only applied when the response *reframes* the content of the root-tweet or not.
- Framing strategies embedded in moral foundations are known to influence people’s political attitudes [25]. To accommodate and benefit from this understanding, we refined Entman’s conceptualization of personal evaluation into three distinct codes so we can differentiate between evaluations based on personal attributes like moral *character*, *competence*, and/or *credibility* that are particularly important in a political discourse.
- For identifying the extent to which conversations in political discourse are framed around evidence, our framework also categorizes a statement as an *expression*, a *claim*, or an *evidence*. We found this dimension useful to differentiate responses that were aimed at intensifying a message from those that offered a robust refutation toward a message.

In addition, we also identified several categories as summarized in Table 2. We now describe all of these in detail.

5.1 Framing Techniques

Our research identified five different framing techniques (across two dimensions) that a tweet employed relative to the root tweet.

Engage or reframe (Engage, Reframe). First, we coded each tweet as to whether it *engaged* with the current frame or *reframed* the conversation. These two framing categories were mutually exclusive.

The first type of response tweets directly engage with the frame presented in the root tweet, either by intensifying it (for tweets that also had a directed valence of *support*) or by directly challenging it (for tweets that had a directed valence of *oppose*).

Root tweet: “Mueller’s testimony underlined what was already clear: the President of the United States broke the law and would be under criminal indictment if he did not hold that office. He is not above the law. Congress should begin impeachment proceedings.”

Response tweet 1: “About time. Thank you.”

Response tweet 2: “NOTHING THERE”

In the example above, the first tweet is the root tweet, which initiated this conversation. The second tweet is a response tweet that *engaged* with the current frame in a supportive way. The third tweet is a response tweet that engaged with the original frame by directly challenging it.

The second type of response tweets, those in the *reframe* category, functioned to shift the frame in a meaningful way. For example:

Root tweet: Searing, sad picture of a man who gave his all to his country during an arduous half century, from Vietnam, through leading the FBI and his final turn as a witness, explaining his probe of @realDonaldTrump. RobertMueller is everything Trump is not.

Response tweet: Mueller met with his attorneys all week. This was all an act to make himself look like a fool. All fake to get out of being prosecuted.

Reframing often presented as drastic shifts in meaning like the example above, which indirectly challenged the original frame (of Mueller as a heroic figure) by presenting a new, counter-frame (of Mueller as a criminal). Response tweets, such as this one, that had a directed valence of *oppose* can be considered as a method of contesting an existing frame [49]. Reframing also took more subtle forms, even occurring between two aligned accounts as the response tweet shifted the frame to highlight different elements of the unfolding event.

The *reframe* technique occurred most often through alterations of the problem definition, though some response tweets included new content along other framing dimensions, e.g., an added moral evaluation and/or treatment recommendation.

Twisting of words (Yes, No). One salient type of reframing involved what we called *twisting of words* where the response tweet accepted and anchored on a small number of words in the root tweet and shifted the frame around those words. For example:

Root tweet: Every Trump supporter who calls Mueller confused and senile seems to willfully ignore Trump’s ignorance of basic math and grammar

Response tweet: His basic math and grammar only made him like A billion dollars richer than you and the person who’s over you

The response tweet here shifts the frame of conversation by anchoring around the original words and re-interpreting them. We treated *Twisting of Words* as its own code category, a sub-category of *Reframe*.

Articulate a Frame (Explain, Describe, Project). A second dimension of framing was articulating a frame. In developing this dimension, we relied upon Entman’s four-part definition of framing as “select(ing) some aspects of a perceived reality... to promote a particular problem definition, causal interpretation, moral evaluation and/or treatment recommendation” ([30]: page 52). Our analysis surfaced three different types of frame articulations that aligned with Entman’s conceptualization: *describe*, *explain*, and *project*.

Response tweets that provided a description of the underlying situation were coded as *describe*. For example:

Root tweet: Cut through all the noise. This is the 35 seconds that you really need to watch. #MuellerHearings [VIDEO] of Mueller testimony

Response tweet: Mueller looks like he is going to cry. He knows he’s screwed no matter what, and is being used as a pawn on live television.

Tweets with *describe* code function to frame the conversation by, as Entman says, promote a problem definition or understanding of the underlying situation. In this case, the response tweet adds a frame to the original tweet, suggesting that Mueller is uncomfortable and caught up in the strategies and manipulations of others.

Response tweets were coded as *explain* if they attempted to explain — i.e., provide a “causal interpretation” for — some aspect of the situation. In the example below, the response tweet adds an explanation for Muller presenting a surprise witness, claiming that he’s incapable of testifying by himself.

Root tweet: Top R on House Judiciary Rep Collins slams Chairman Nadler for potentially allowing a surprise witness at Mueller hearing tomorrow – Aaron Zebley, counsel for Mueller [image of text]

Response tweet: the great and powerful Mueller is a shaky-voiced clueless old man who is incapable of testifying by himself.

Finally, response tweets that presented a projected outcome or treatment recommendation were coded as *project*. For example:

Root tweet: .@HouseDemocrats suffered a huge setback with #Mueller’s testimony yesterday. Now, grasping at straws, @OversightDems retaliate by issuing a subpoena for all of Jared & Ivanka’s personal emails & texts. The relentless, baseless attacks against @POTUS @realDonaldTrump must stop!

Response tweet: No, Trump will face 10 counts of obstruction of justice charges in 2020 when both of you are looking for new jobs. Shame on you for letting Putin pick our President

Here, the response tweet contests the original frame and presents an alternative frame, one that projects that Donald Trump will be indicted for obstruction of justice.

Though certainly an underlying current in much of the conversation, the moral evaluation dimension did not present as a salient category in our analysis. Few response tweets explicitly invoked a moral evaluation. However, we did surface categories related to personal evaluations that included a moral judgment (which we discuss next).

5.2 Personal Evaluation (Character, Credibility, Competence)

To defend their own perspective in the political back and forth, many responses discussed the merit of a tweet’s content by evaluating the root tweeter or another person of interest in the root tweet. We found these evaluations to be based on three distinct characteristics of the people involved in the tweet. The first was a person’s moral code of conduct. We coded such responses as *character*, e.g., “Because crimes against humanity are wrong and people care”.

Often the intent of evaluating a person of interest in the tweet was primarily to discredit their authority by calling them “fake” or invoking reasons why they cannot be trusted. This led us to the second code of *credibility*. The third characteristic that was used as a basis of judgment toward people involved in the tweet was that of their skill and ability. For example,

Root tweet: “Mueller sounds like he’s drugged. Is this somnolent incoherence what the Democrats expected from him?”

Response tweet 1: “Mueller cannot answer questions & he wants the questions re-played”

Response tweet 2: “It’s called thinking”

The root tweet in this example sets the topic as Mueller’s dissatisfying conduct in the testimony, which the first response supports by judging Mueller’s ability to answer the questions. The second response challenges the root tweet by referring to Mueller as someone who is capable of thinking. We coded both these responses as personal evaluations based on *competence*.

5.3 Level of Verifiability (Expression, Claim, Evidence)

We also coded each response tweet for level of verifiability: expression, claim, or evidence. Codes in this category were mutually exclusive.

Perhaps due to the political nature of the data, we found that a large part of the responses were statements with an expression (often of emotion) that could not be verified. For example, comments like “This didn’t age well” and “PERFECT!” are *expressions* that cannot be verified.

Our second category, *claim* refers to a statement that is yet to be verified but can be verified. This code was applied to responses like “this has been debunked time and time again” with the

understanding that either the content of the root tweet has been debunked several times or it has not, i.e., the content has a concrete potential to be verified.

We mapped responses offering statements that have been verified as *evidence*. For example, “According to American law, a prosecutor is not supposed to exonerate a person,” cites a source (American law) for its claim (that a prosecutor is not supposed to exonerate a person). We therefore coded this tweet as providing evidence. The *evidence* category includes widely accepted statements and sentiments, such as, “... justice and “innocent until proven guilty” are important values of US justice system.”

5.4 Target Audience (Root Tweeter, Broader Audience, Both)

Social media users are known to have “imagined audiences” for their content [53, 75], and these mental models — or folk theories [32] — of who sees their content shape their decisions about how and what to share [3, 55]. In our data, we noted that different posts seemed to engage with different target audiences.

For example, a response tweet that begins by stating “I believe you are against Voter ID and FOR illegal aliens” attempts to directly engage with the author of the root tweet. On the other hand, responses like, “These people might end up saving our democracy” suggest a broader target audience, where the response tweeter is not calling out or engaging directly with the root tweeter. Often, the response tweet speaks both to the root tweeter and to the broader audience, e.g., “Keep lying @<root tweeter>! Trumpism corrupts good people. <Downward-pointing finger emoji> Perfect example. All for clicks!”

Drawing from these three categories, we coded each tweet for its target audience, i.e., whether it attempted to engage directly with the *root tweeter*, was intended for the *broader audience*, or attempted to engage at *both* levels.

5.5 Directed Valence (Support, Oppose, Neutral)

For each root-response tweet pair, we captured the valence of the response tweet toward the root tweet using three mutually exclusive codes: *support*, *oppose*, and *neutral*. Not surprisingly, response tweets from politically-aligned accounts were more likely to support the root tweet, while those from politically-misaligned accounts were more likely to oppose the root tweet — though the trend was clearer for reply tweets than for quote tweets.

5.6 Language

Varying the linguistic techniques used within a sentence is known to invoke different attitudes and behavior. For example, positive or negative wordings can influence the subsequent online expressions [50]. Further, the adoption of such linguistic techniques can vary across different groups. For example, sarcasm is found to occur more often in a right-leaning attack-discourse [2]. To benefit from this understanding, we coded our data for two categories about the *tone* and *sincerity* of the language used in the responses.

Tone (Positive/Negative). We coded each response tweet depending on the language within the response as either *positive* — e.g., “You did a great job saving our Republic today!” — or *negative* — e.g., “This didn’t age well” — , or *neutral* — e.g., “There’s a list, and we need to start somewhere”. Since codes in this category were not relative to the root tweet, either tone could be used to support or to challenge the root tweet.

Attitude (Reverent/Condescending). This category differentiates the response tweets that try to establish a sense of superiority and/or mockery toward the root tweet(er) from those that try to treat the root tweet(er) with respect and/or honor. Examples of *condescending* include responses

Table 4. Odds Ratios for *Target Audience*, *Response Frame*, and *Twisting of Words* Calculated using Logistic Regression Modeled After the Type of Feature (Reply and Quote) and Political Alignment (Politically Aligned and Politically Opposed) and Controlled for Random Effects

	Odds ratio	CI [95%]	p
Target audience (<i>Root tweeter, Broader public, Both; reference: Broader public</i>)			
Root tweeter/(Intercept)	1.15	[0.88, 1.49]	.296
Root tweeter/Political alignment (Politically aligned)	0.30*	[0.16, 0.58]	<.001
Root tweeter/Type of feature (Reply)	3.81*	[1.98, 7.34]	<.001
Both/(Intercept)	1.37*	[1.08, 1.75]	.01
Both/Political alignment (Politically aligned)	0.72	[0.40, 1.27]	.254
Both/Type of feature (Reply)	2.55*	[1.42, 4.59]	.002
Response frame (<i>Engage, Reframe; reference: Engage</i>)			
Reframe/(Intercept)	0.92	[0.75, 1.11]	.369
Reframe/Political alignment (Politically aligned)	0.36*	[0.23, 0.56]	<.001
Reframe/Type of feature (Reply)	1.67*	[1.06, 2.62]	.027
Twisting of words (<i>Yes, No; reference: No</i>)			
Yes/(Intercept)	0.33*	[0.15, 0.66]	<.003
Yes/Political alignment (Politically aligned)	0.60	[0.12, 2.31]	.484
Yes/Type of feature (Reply)	0.20*	[0.04, 0.71]	.021

like “Sure. Just like there was concrete evidence of Russian collusion. Democrats who keep falling for this are #NotVeryBright”. On the other hand, responses that illustrate gratitude toward the root tweet — e.g., “Thank you sir. Justice, “innocent until proven guilty” and “equal under the law” are important values of the American justice system” — were labeled with the *reverent* code.

Out of all the categories, only *Articulate a frame* and *Personal evaluation* were not mutually exclusive, i.e., coders could use any one or all of the codes when annotating the data. All other categories were mutually exclusive. Thus, coders could assign exactly one of the *expression*, *claim*, and *evidence* codes from the category of “level of verifiability” to a root-response tweet pair.

6 FINDINGS: REPLY/QUOTE AND POLITICAL ALIGNMENT AFFORD UNIQUE CONVERSATIONAL AND FRAMING STRATEGIES

In this part of the findings, we focus on how the strategies that we discussed above differed between the reply and quote features. First, we report on the primary differences in conversational and framing strategies between reply and quote, and then discuss additional differences after accounting for political alignment of the authors of those tweets.

6.1 RQ1: Conversational Strategies Vary Between Reply vs. Quote

The primary focus of this research was to identify how the mechanisms of replying and quoting are used differently by Twitter users to develop, shape, refine and promote preferred conversational frames. At the same time, we wanted to examine how these mechanisms are employed to challenge, appropriate, and counter opposing conversational frames — particularly in a politically charged conversation. Across the different dimensions that we identified in our codebook, we found three dimensions where the strategies were employed differently between reply and quote tweets (Table 4).

We found that the *Target audience* differed between replies and quotes, $X^2(df = 2, N = 340) = 16.65, p = .00024$. Specifically, a higher percentage of replies targeted the root tweeter (73.5% vs. 43.3% for quotes), while a higher percentage of quotes targeted the broader audience (56.7% vs. 26.5% for replies). This difference was significant in our logistic regression analysis ($p < 0.001$), indicating that the odds of replies engaging directly with the root tweeter (over the Broader audience)

were 3.8 times higher than that of quotes (refer Table 4). Such a framing to shift the conversation away from the root tweeter toward a broader audience — in accordance with findings from prior research [38] — imparts a sense of indirectness in the quoting mechanism and influences its usage differently than that of replies in different scenarios discussed below.

Another key difference between reply and quote was that *Reframing* was more common in reply tweets than in quote tweets. We found that 46% of the reply tweets reframed the topic of conversation set in the root tweet as compared to only 34% quote tweets (refer Table 3). The odds of reframing the conversation were 1.67 times higher than simply engaging with the original frame in responses that made use of the replying mechanism as compared to the quoting mechanism (Table 4). This result was surprising given the potential of the quoting mechanism to reinterpret another tweet by adding to the conversation.

We also found that quotes were more likely used to *Twist words* than replies. This phenomena is a special case of reframing that occurred in 8.3% of the replies and 28.9% of the quotes. There was a five times higher chance of twisting of words to occur (relative to no twisting of words) in responses expressed through quotes as compared with replies (refer Table 4). Quote tweets often used the root tweet as “evidence”, anchored around the original words, and re-interpreted the meaning. Interestingly, though quotes were less likely overall (than replies) to reframe a conversation, they were much more likely to “twist words” when they do.

6.2 RQ2: Conversational Strategies Vary When Response is Politically Aligned vs. Opposed

We found three differences (refer Table 5) in conversational and framing strategies that depended on who you talk to during a political exchange online, i.e., people who share similar opinions or dissimilar opinions.

Given the political nature of the data, we found more conversations that happened between two accounts sharing similar opinions to be largely supportive of each other, while those between dissimilar opinions expressed opposition, $X^2(df = 2, N = 340) = 130.155, p < 0.00001$. The odds of having support (over opposition) expressed through a response between aligned accounts were about 20 times as compared to a response between dissonant accounts. Similarly, we found that the tone was generally more positive when engaging with an account with similar political alignment, $X^2(df = 2, N = 222) = 34.816, p < 0.00001$. When we controlled for the type of response, our model indicated with a statistical significance that positive tone was about 14 times more likely to occur between politically-aligned accounts than politically-opposed accounts. This theme continued as we found responses that were structured to express respect (category: *attitude*) to be 35 times more likely to emerge through engagements happening on the same side of the political-aisle than those happening across the political-aisle, $X^2(df = 2, N = 340) = 26.839, p < 0.00001$. These findings can be explained by the presence of polarization on a Twitter platform — particularly as seen in the case of politically charged topics like Brett Kavanaugh’s nomination to serve as a Justice on the US Supreme Court in 2018 [23].

We also found that accounts engaged in the same topic as set in the root tweet with an impassioned behavior of (mostly) supporting the topic in the root tweet when the account happens to be politically-aligned, but reframed the topic of discussion otherwise, $X^2(df = 1, N = 299) = 18.169, p = 0.0002$. When we controlled for the type of response, the odds of a response trying to refocus were 2.8 times higher if the response were toward a politically opposed account than if it were an aligned account. This aligns with previous findings that people promote a discourse if it aligns with their own beliefs, but challenge it otherwise. For example, politically motivated individuals use mention networks on Twitter to interact with ideologically-opposed users by injecting partisan content into information streams [21].

Table 5. Odds Ratios for *Directed Valence*, *Directed Attitude*, and *Tone of Language* Calculated using Logistic Regression Modeled after the Type of Feature (Reply and Quote) and Political Alignment (Politically Aligned and Politically Opposed) and Controlled for Random Effects

	Odds ratio	CI [95%]	p
Directed valence (<i>Support, Oppose, Neutral; reference: Oppose</i>)			
Support/(Intercept)	0.57*	[0.45, 0.72]	<.001
Support/Political alignment (Politically aligned)	20.17*	[11.42, 35.60]	<.001
Support/Type of feature (Reply)	0.65	[0.37, 1.14]	.134
Neutral/(Intercept)	0.26*	[0.17, 0.42]	<.001
Neutral/Political alignment (Politically aligned)	7.84*	[2.25, 27.32]	<.001
Neutral/Type of feature (Reply)	0.15*	[0.03, 0.71]	.017
Attitude (<i>Reverent, Condescending, None; reference: Condescending</i>)			
Reverent/(Intercept)	0.11*	[0.04, 0.29]	<.001
Reverent/Political alignment (Politically aligned)	35.21*	[4.63, 268.05]	<.001
Reverent/Type of feature (Reply)	1.12	[0.45, 2.83]	.803
None/(Intercept)	1.02	[0.85, 1.24]	.773
None/Political alignment (Politically aligned)	1.77*	[1.13, 2.79]	.013
None/Type of feature (Reply)	1.01	[0.64, 1.58]	.972
Tone (<i>Positive, Negative, Neutral; reference: Negative</i>)			
Positive/(Intercept)	0.20*	[0.12, 0.34]	<.001
Positive/Political alignment (Politically aligned)	14.06	[4.82, 41.01]	<.001
Positive/Type of feature (Reply)	0.88	[0.44, 1.77]	.722
Neutral/(Intercept)	0.78*	[0.64, 0.94]	<.011
Neutral/Political alignment (Politically aligned)	1.27	[0.79, 2.03]	.318
Neutral/Type of feature (Reply)	0.98	[0.61, 1.55]	.919

Our coding indicated that about 70% of the responses involving politically opposed accounts attempted to engage with the root tweeter and about 57% of the responses involving politically aligned accounts reached out to a broader audience, $X^2(df = 2, N = 340) = 15.55, p < .001$. When we controlled for the type of response, i.e., quote or reply, we found that an account is about three times more likely to focus on the individual root tweeter if the two accounts have an opposing outlook as compared to sharing a similar political outlook. The cognizance of having a disagreement with the root tweeter about the topic under discussion can promote such a framing technique — to move away from the topic of discussion, to focus on whom you disagree with, and consequently direct the disagreement toward them.

6.3 RQ3: Conversational Strategies Vary Across Both Response Type and Political Alignment

In the sections above, we presented the three conversational strategies that differed between quotes and replies. Here, we present the other conversational strategies that though did not differ between quotes and replies directly, did differ when we factor in political alignment (refer Table 6).

We found that although character is similarly mentioned between replies and quotes, there are more nuanced differences when we factor in political alignment. Specifically, although replies had similar frequencies in mentions of character between political alignment, the use of character in quotes exhibited a different pattern as illustrated in Figure 3. Quote tweets referenced an individual's character more often (29 times) when engaging with politically opposed accounts than aligned accounts (17 times). The odds for referring to one's character in a reply tweet were 0.67 times that of a quote tweet when the accounts were politically opposed, but the odds increased to 1.48 times when the accounts were politically aligned.

Table 6. Odds Ratios for *Character*, *Evidence*, and *Projection* Calculated using Logistic Regression Modeled after the Type of Feature (Reply and Quote) and Political Alignment (Politically Aligned and Politically Opposed) and Controlled for Random Effects

	Odds ratio	CI [95%]	p
Character (Yes, No; reference: No)			
Yes/(Intercept)	0.51*	[0.33, 0.80]	.004
Yes/Political alignment (Politically aligned)	0.48*	[0.24, 0.96]	<.04
Yes/Type of feature (Reply)	0.67	[0.34, 1.30]	.242
Yes/Political alignment (Politically aligned) X Type of feature (Reply)	2.2*	[0.83, 5.87]	.004
Evidence (Yes, No; reference: No)			
Yes/(Intercept)	0.31*	[0.18, 0.49]	<.001
Yes/Political alignment (Politically aligned)	0.16*	[0.04, 0.45]	.001
Yes/Type of feature (Reply)	0.69	[0.32, 1.47]	.344
Yes/Political alignment (Politically aligned) X Type of feature (Reply)	4.32*	[1.11, 19.37]	<.001
Projection (Yes, No; reference: No)			
Yes/(Intercept)	0.23*	[0.13, 0.39]	<.001
Yes/Political alignment (Politically aligned)	2.6*	[1.31, 5.34]	.007
Yes/Type of feature (Reply)	1.24	[0.59, 2.64]	.569
Yes/Political alignment (Politically aligned) X Type of feature (Reply)	0.29*	[0.10, 0.79]	<.001

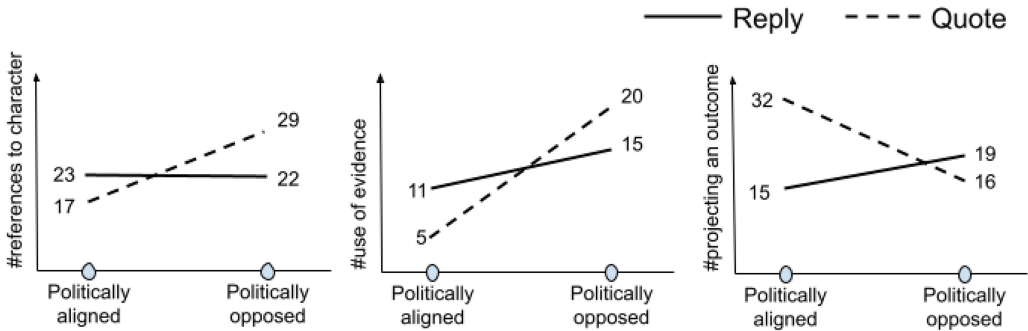


Fig. 3. Interaction effect for (left) referring to one's character, (center) use of evidence, and (right) making a projection. Please note that the Y axis for all the 3 graphs differ in range.

We found the tendency of adopting *Evidence* within one's response — to support the response tweeter's own take toward the root tweet — was similar to referencing an individual's *character* as seen above. Evidence-based responses emerged 35 times when engaging with a politically opposed account as compared to 16 times when engaging with a politically aligned account. As Figure 3 illustrates, this varied with the choice of a reply tweet or quote tweet. When responding across the ideological divide, the odds for use of evidence in replies were 0.69 times as compared to the odds for use of evidence in quotes; with accounts sharing similar opinions, these odds increased to about 3 times. We believe this pattern for the use of evidence can be explained by the disagreements — and the perceived need for making a convincing refutation — across the ideological divide.

When we coded to *articulate the frame* within the tweet, 15 replies to politically similar accounts projected outcomes as opposed to 19 replies to politically opposed accounts. This trend changed for quote tweets, with 32 politically aligned quote tweets making some projection as compared to only 16 politically opposed quote tweets (Figure 3). The odds for a reply to project an outcome or suggest a potential solution toward the event set in the root tweet were 1.24 times higher than

that of a quote when the tweet pairs were politically opposed; the odds ratio however reduced to 0.354 when the tweet pairs were politically aligned. The need to *describe* in details and/or *explain* a disagreement toward a politically opposed root tweet(er) — implying a lesser focus on *projecting* an outcome — diminishes when a response engages with politically aligned account. This accompanied by the preference of quotes to *engage* with the same topic as set by the root tweet can explain the higher number of politically-aligned projections. Examples include *expressions* like “#IMPEACH”, “#SpeakerCortez won’t be happy” that often urge an action and seemed to merely amplify the message of the root tweeter. We now discuss the implications of these findings.

7 DISCUSSION

In order to discuss the implications of our findings, we first address the limitations of our research. We then summarize the key contributions of our coding framework and then focus on what differentiates the quoting mechanism from the replying mechanism, how the antagonistic nature of quoting can lead to contestation, polarization, and “dunking”, and finally recommend design implications.

7.1 Limitations

Our research focused on the interactions between root tweets and response tweets. One limitation of this work is that at the time of data collection, there was no easy way to collect the reply threads for tweets in a Streaming API collection — even as quotes were picked up with the relevant root tweets. We attempted to overcome this limitation through a follow-up data collection, but there are inconsistencies between those data samples, potentially limiting their equivalence for an effective comparison.

The findings presented in this work highly depend on the nature of the data collected back in 2019. For example, later versions of the Twitter API allowed easier retrieval of the threaded responses and would have facilitated collecting a larger sample size, replies in particular. Though relatively limited in size, we have ascertained the statistical sufficiency of our sample size to observe an effect size corresponding to odds ratio of 1.5 with 80 percent power.

Additionally, the findings are also influenced by the state of Twitter algorithms at play behind the reply and quote features and also people’s adoption of these features unique to the platform, as seen around July 2019. With the recent changes made to the platform, e.g., introducing the “For you” and “Following” dual-timelines [57], it remains to be seen how these affordances get appropriated by users as they get more experience with them and how they cope as the platform has and will continue to evolve over the coming years. Regardless, this research adds to the scholarship around how microblogging platform affordances impact online discourses in significant ways discussed below.

7.2 Generalizable Codebook to Identify Conversational and Framing Strategies on Microblogging Platforms

This research contributes a theoretically grounded coding framework that identifies conversational and framing strategies adopted by the response tweets toward the root tweet. We expand upon previous work by looking at conversational strategies that surpass disagreement and the use of insult [38] and facilitate a broader analysis of Twitter discourse. We demonstrated the use of our framework to categorize the root-response tweet pairs to help us compare and contrast the use of reply and quote features on the Twitter platform in political discourse.

Though grounded in data specific to a polarized online discourse, we designed our coding framework with conversational attributes in mind — e.g., changing the focal topic, targeting a specific audience, and so on. — that are broadly applicable to any topic. Beyond this research, our

proposed framework could therefore be useful to understand how responses are used in an online discourse to build, refine, revise, and amplify conversational frames that one prefers and challenge or counter-frames that they oppose regardless of any topical limitations. For example, the presented framework could be easily employed to study how recent online conversations around vaccinations or student loan forgiveness were shaped differently using different conversational features.

There are several open questions across research domains where our codebook could be applied on platforms beyond Twitter. For example, researchers and similar professionals often have trouble navigating scientific discourse in online spaces [12], so quantitative work examining the relationship between different discursive and presentation styles and the resulting discourse could suggest best practices for professionals discussing controversial topics. In addition, with Mastodon [79] and other Twitter alternatives like Threads [27] gaining prominence, our codebook could aid researchers in understanding how the different affordances across these platforms affect the conversations that occur there.

7.3 Comparing Reply and Quote Features

This research examined how Twitter users adopt the reply and quote features in online political discourse. We found three key differences between reply tweets and quote tweets.

First, we found that reply tweets often communicated directly with the root tweeter, while quote tweets were used to broadcast a message to the quote tweeter's following on Twitter; bloggers [1] and researchers [40] both acknowledge how broadcasting helps increasing one's audience. This difference in *target audience* suggests that the quoting mechanism is more often used to shift the conversation away from the root tweeter toward a broader audience — in accordance with findings from prior research [38]) — and imparts a sense of indirectness to the response.

Second, quote tweets continued to *engage* with the topic of conversation as set by the root tweet, while reply tweets more often *reframed* the conversation, shifting attention to different dimensions of the debate that were not salient in the original tweet. This result was surprising, given the potential of quote tweets to add a comment and reinterpret the message of the root tweet. We suspect that the motivation to broadcast that root tweet to their own followers on Twitter might limit them from altering the message, but instead sharing it with them by adding some signal of support or opposition.

Third, replies and quotes differed in how they are used when responding to politically aligned or opposed accounts. Responses that came in through the quoting mechanism more frequently used *evidence* to support their contestation and present a robust refutation when the root tweeter was politically opposed; quotes exhibited a similar pattern when it came to questioning the root tweeter's *character*. However, the use of evidence or questioning one's character through replies was less partial to the political affiliation of the root tweeter. The dedicated attention of quote tweets to accounts that are opposed to making an evidence-based refutation explains how the quoting mechanism can provoke contestation.

We now discuss how these key differences influence the use of quotes vs. replies in an online discourse — with a focus on the quoting mechanism — toward provoking contestation, polarization, and *dunking*. It is important to note that our data included replies and quotes authored by unique Twitter accounts, thereby reducing the chance that the variations in conversational framing and strategies that we observed resulted from individual differences.

7.3.1 Quotes amplify the “like” and contest the “dislike”. We found that quotes tweets were characterized by broadcasting to a wider audience, diverting a response to the quote tweeter's own following, and *twisting words* to relay a sense of antagonism to the politically opposed. This resulted

in the quoting mechanism promoting a more passive-aggressive tone when conversing with the root tweet. In other words, quoting allowed a user to be aggressive (just like a reply), but the aggression was diverted (unlike a reply).

This passive-aggressiveness afforded by quotes can explain why quote tweets have been found to challenge the narrative that is prominent amongst the ideologically opposed accounts in an on-line discourse on Twitter [78]. The dedicated use of evidence-based responses to counter-argue and refute messages from accounts that are ideologically opposed can also corroborate the tendency of quotes to provoke contestation [42].

The conversational strategy of reframing to alter the focus of the conversation allows a reply to *describe* in detail and/or *explain* their stance toward the message of the root tweet(er). This makes replies more suitable to challenge the message of not only politically opposed but also politically aligned accounts. The (relatively) lesser focus on reframing and more on engaging within the topic of conversation — accompanied by the potential to reach out to a broader audience — may lead to the quoting mechanism being used to intensify the message of politically aligned accounts.

Quotes thus seem to either broadcast their agreement and amplify their approval toward the politically aligned, or provoke a contestation and intensify the extent of dislike toward the politically opposed — potentially feeding into the affective polarization [21, 83] in an already polarized Twitter discourse.

7.3.2 Quotes lead to “dunking”. Our research presents an interesting opportunity to consider why designers of emerging platforms like Mastodon have early on opposed introducing quoting mechanism to discourage dunking-like behaviors [81]. An early investigation into the usage of the quote feature — referred to as *retweet with comment* back in 2016 — suggested that quoting might help construct a more civil online discourse than that facilitated by replying [38]. However, subsequent research has found that quoting can promote contestation by drawing attention to disagreements [78] and/or by interpreting the root tweet out of (temporal) context to invoke humor and critique [42].

One possible explanation for this change comes back to a specific design choice in how the quote feature was implemented — in that, it enabled users to shift the audience away from one that is primarily followers of the root tweeter (as it is in the case of a Reply) to an audience that is primarily followers of the quote tweeter. This design choice — which has been updated slightly since this data was collected [67] — is likely implicated in what we perceived as a *diversion* strategy and has contributed to the phenomenon of “dunking.” To dunk implies that one is able to both easily and dramatically score political points. Dunking often takes place through a quote tweet intended to refute the root tweet while also mocking or otherwise denigrating the root tweeter to get attention [74]. Let us consider an example of dunking from our data:

Root tweet: “So...ummm...Bob Mueller is old. And this hearing is just painful to watch.”

Quote tweet: “Is that all you’ve got <root tweeter’s name>?”

In this example, the primary motivation of the quote tweeter is to ridicule the politically-opposed root tweeter and to suggest that they could have done better when expressing their opposition toward the testimony. While a reply tweet could fulfill the intent of ridiculing, it would primarily get attention from the root tweeter and the root tweeter’s following, who will likely support and resist the quote tweeter’s refutation of the root tweet. The use of quoting mechanism here facilitates diverting the ridicule away from their own following on Twitter. This broader audience, which is likely politically aligned with the quote tweeter and hence opposed to the root tweeter, can now witness the ridicule and also participate in intensifying the mockery. Thus, quotes can increase the chance of further ridiculing the root tweeter and might avoid any direct confrontation from them.

7.4 Design Implications

As we discussed above, the adoption of the quoting mechanism in a polarized discourse differs from that of the replying mechanism due to the potential for using quote tweets to divert attention away from the root tweeter to fetch attention from a broader audience comprising the quote tweeter's following on Twitter, and to indirectly contest the frames of politically opposed accounts.

One potential recommendation here is to give users more control over who can quote their online posts. This could reduce the potential for indirect contestation and related behaviors like dunking. At the time of this research, Twitter allowed its users some control over who can reply to their tweets, e.g., "everyone," "people you follow," or "people you mention." Similar choices to decide who participates through quotes could assist the root tweeters to better keep track of and moderate the emerging discourse. We are glad to report that Twitter introduced changes aligning with these suggestions in late 2020 [87]. We suspect that such a design choice could also negatively affect the ability of Twitter users to challenge falsehoods and hold high visibility accounts — who are often the most quoted — accountable for problematic tweets.

Another strategy for alleviating some of the toxicities that emerge through quotes would be to increase their visibility to the root tweeter and their audience. Twitter took a step in this direction — implemented in 2020 [67] — by making it easier to view the quotes for a specific tweet that appears in your feed. Though this might result in more divisive discourse in the short term, as a root tweeter's audience is now more easily able to mobilize in defense of their tweets, over time, it could reduce the tendency for people to use the quote feature for passive aggressive attacks and dunks — as their imagined audience changes and they begin to make different decisions on how and what to post. This is certainly something to be investigated in future research.

Some of the lessons about platform design choices from this research have already been incorporated into Twitter, and will continue to impact the design of more microblogging platforms in future. These lessons become particularly important as the designers of an emerging platform like Mastodon, who were previously opposed to the quoting mechanism to preserve its "antiviral" design and dissuade dunking-like behaviors [81], negotiate the possibility of introducing it in spirit of public interest. This research and the coding framework will serve a useful guide and evaluate future design choices toward fostering healthy online discourses on emerging microblogging platforms like Mastodon, Threads, and so on.

8 CONCLUSION

The affordances of social media platforms are known to offer users different capacities toward political participation. Some of these affordances can trigger unwanted engagements compromising the health of some online environments. For example, reply is more likely to be used for trolling on Breitbart and IGN, but is more likely to be used on CNN for starting new discussions [17]. Reacting to perceptions of how some affordances have negative effects on the health of discourse, many platforms have changed how they facilitate and make visible certain kinds of engagements — such as replies and comments [41, 63]. The presented research can help us to understand how distinct features like reply and quote encourage conversations by promoting different framing strategies and hence shape resulting behaviors as we engage with people from all sides on online platforms.

ACKNOWLEDGMENTS

We thank University of Washington's Department of **Human Centered Design and Engineering (HCDE)** and the **Center For an Informed Public (CIP)** for providing supportive communities.

REFERENCES

- [1] Alex. 2020. How to use quote retweets to grow your audience (with examples). Retrieved from <https://tweethunter.io/blog/how-to-use-quote-retweets-to-grow-your-audience-with-examples>. Accessed on July 30, 2023.
- [2] Ashley A. Anderson and Heidi E. Huntington. 2017. Social media, science, and attack discourse: How twitter discussions of climate change use sarcasm and incivility. *Science Communication* 39, 5 (2017), 598–620.
- [3] Ahmer Arif, John J. Robinson, Stephanie A. Stanek, Elodie S. Fichet, Paul Townsend, Zena Worku, and Kate Starbird. 2017. A closer look at the self-correcting crowd: Examining corrections in online rumors. In *Proceedings of the 2017 ACM Conference on Computer Supported Cooperative Work and Social Computing*. 155–168.
- [4] Christopher A. Bail, Lisa P. Argyle, Taylor W. Brown, John P. Bumpus, Haoan Chen, M. B. Fallin Hunzaker, Jaemin Lee, Marcus Mann, Friedolin Merhout, and Alexander Volfovsky. 2018. Exposure to opposing views on social media can increase political polarization. *Proceedings of the National Academy of Sciences* 115, 37 (2018), 9216–9221.
- [5] Hilda Bastian. 2022. Reflecting on twitter, white flight, & “quote tweet” tensions at mastodon. Retrieved from <https://absolutelymaybe.plos.org/2022/12/01/reflecting-on-twitter-white-flight-quote-tweet-tensions-at-mastodon/>. Accessed on July 30, 2023.
- [6] Mathieu Bastian, Sebastien Heymann, and Mathieu Jacomy. 2009. Gephi: An open source software for exploring and manipulating networks. *Icswm* 8, 2009 (2009), 361–362.
- [7] Robert D. Benford and David A. Snow. 2000. Framing processes and social movements: An overview and assessment. *Annual Review of Sociology* 26, 1 (2000), 611–639.
- [8] Leticia Bode, Alexander Hanna, Junghwan Yang, and Dhavan V. Shah. 2015. Candidate networks, citizen clusters, and political expression: Strategic hashtag use in the 2010 midterms. *The ANNALS of the American Academy of Political and Social Science* 659, 1 (2015), 149–165.
- [9] Arjen Boin, Paul ‘t Hart, and Allan McConnell. 2009. Crisis exploitation: Political and policy impacts of framing contests. *Journal of European Public Policy* 16, 1 (2009), 81–106.
- [10] Danah Boyd, Scott Golder, and Gilad Lotan. 2010. Tweet, tweet, retweet: Conversational aspects of retweeting on twitter. In *Proceedings of the 2010 43rd Hawaii International Conference on System Sciences*. IEEE, 1–10.
- [11] Brian C. Britt, Rebecca K. Britt, Jameson L. Hayes, and Jeyoung Oh. 2020. Continuing a community of practice beyond the death of its domain: Examining the tales of link subreddit. *Behaviour & Information Technology* 41, 1 (2020), 1–22.
- [12] Michael Brüggemann, Ines Lörcher, and Stefanie Walter. 2020. Post-normal science communication: Exploring the blurring boundaries of science and journalism. *Journal of Science Communication* 19, 3 (2020), A02. DOI: <https://doi.org/10.22323/2.19030202>
- [13] Axel Bruns and Jean E. Burgess. 2011. The use of twitter hashtags in the formation of ad hoc publics. In *Proceedings of the ECPR General Conference*.
- [14] Monika Butler and Michel André Maréchal. 2007. Framing effects in political decision making: Evidence from a natural voting experiment. *University of St. Gallen Economics Discussion Paper* 2007-04 (2007).
- [15] Ian Cairns and Priyanka Shetty. 2020. Introducing a new and improved twitter API. Retrieved from https://blog.twitter.com/developer/en_us/topics/tools/2020/introducing_new_twitter_api.html. Accessed on July 30, 2023.
- [16] Kathy Charmaz. 2006. *Constructing Grounded Theory: A Practical Guide Through Qualitative Analysis*. Sage.
- [17] Justin Cheng, Christian Danescu-Niculescu-Mizil, and Jure Leskovec. 2015. Antisocial behavior in online discussion communities. In *Proceedings of the International AAAI Conference on Web and Social Media* 9, 1 (2015), 61–70.
- [18] Dennis Chong and James N. Druckman. 2013. Counterframing effects. *The Journal of Politics* 75, 1 (2013), 1–16.
- [19] Ben Collins and Brandy Zadrozny. 2020. Twitter bans 7,000 QAnon accounts, limits 150,000 others as part of broad crackdown. *NBC News* July 21 (2020). <https://www.nbcnews.com/tech/tech-news/twitter-bans-7-000-qanon-accounts-limits-150-000-others-n1234541>. Accessed on July 30, 2023.
- [20] Zadrozny Collins and Ben Collins. 2019. Trump, QAnon and an impending judgment day: Behind the facebook-fueled rise of the epoch times. *NBC News* (2019).
- [21] Michael Conover, Jacob Ratkiewicz, Matthew Francisco, Bruno Gonçalves, Filippo Menczer, and Alessandro Flammini. 2011. Political polarization on twitter. In *Proceedings of the International AAAI Conference on Web and Social Media*.
- [22] Patrick G. Coy and Lynne M. Woehrl. 1996. Constructing identity and oppositional knowledge: The framing practices of peace movement organizations during the persian gulf war. *Sociological Spectrum* 16, 3 (1996), 287–327.
- [23] Kareem Darwish. 2019. Quantifying polarization on twitter: The kavanaugh nomination. In *Proceedings of the International Conference on Social Informatics*. Springer, 188–201.
- [24] Kareem Darwish, Peter Stefanov, Michaël Aupetit, and Preslav Nakov. 2020. Unsupervised user stance detection on twitter. In *Proceedings of the International AAAI Conference on Web and Social Media*. 141–152.
- [25] Martin V. Day, Susan T. Fiske, Emily L. Downing, and Thomas E. Trail. 2014. Shifting liberal and conservative attitudes using moral foundations theory. *Personality and Social Psychology Bulletin* 40, 12 (2014), 1559–1573.
- [26] Munmun De Choudhury and Emre Kiciman. 2017. The language of social support in social media and its effect on suicidal ideation risk. In *Proceedings of the International AAAI Conference on Web and Social Media*. 32–41.

- [27] Chase DiBenedetto. 2023. Threads already has a hate speech problem, civil rights groups warn. Retrieved from <https://mashable.com/article/threads-hate-speech-disinformation>. Accessed on July 30, 2023.
- [28] Haifeng Du, Marcus W Feldman, Shuzhuo Li, and Xiaoyi Jin. 2007. An algorithm for detecting community structure of social networks based on prior knowledge and modularity. *Complexity* 12, 3 (2007), 53–60.
- [29] Nwachukwu Egbunike. 2015. Framing the# occupy nigeria protests in newspapers and social media. *Open Access Library Journal* 2, 05 (2015), 1.
- [30] Robert M. Entman. 1993. Framing: Toward clarification of a fractured paradigm. *Journal of Communication* 43, 4 (1993), 51–58.
- [31] Robert M. Entman. 2010. Media framing biases and political power: Explaining slant in news of campaign 2008. *Journalism* 11, 4 (2010), 389–408.
- [32] Motahhare Eslami, Karrie Karahalios, Christian Sandvig, Kristen Vaccaro, Aimee Rickman, Kevin Hamilton, and Alex Kirlik. 2016. First I "like" it, then I hide it: Folk theories of social feeds. In *Proceedings of the 2016 CHI Conference on Human Factors in Computing Systems*. 2371–2382.
- [33] Franz Faul, Edgar Erdfelder, Albert-Georg Lang, and Axel Buchner. 2007. G* power 3: A flexible statistical power analysis program for the social, behavioral, and biomedical sciences. *Behavior Research Methods* 39, 2 (2007), 175–191.
- [34] Rebecca Ferguson, Zhongyu Wei, Yulan He, and Simon Buckingham Shum. 2013. An evaluation of learning analytics to identify exploratory dialogue in online discussions. In *Proceedings of the 3rd International Conference on Learning Analytics and Knowledge*. 85–93.
- [35] Peer C. Fiss and Paul M. Hirsch. 2005. The discourse of globalization: Framing and sensemaking of an emerging concept. *American Sociological Review* 70, 1 (2005), 29–52.
- [36] Vijaya Gadde and Kayvon Beykpour. 2020. Additional steps we're taking ahead of the 2020 US election. Retrieved from https://blog.twitter.com/en_us/topics/company/2020/2020-election-changes.html. Accessed on July 30, 2023.
- [37] Vijaya Gadde and Kayvon Beykpour. 2020. An update on our work around the 2020 US elections. Retrieved from https://blog.twitter.com/en_us/topics/company/2020/2020-election-update.html. Accessed on July 30, 2023.
- [38] Kiran Garimella, Ingmar Weber, and Munmun De Choudhury. 2016. Quote rts on twitter: usage of the new feature for political discourse. In *Proceedings of the 8th ACM Conference on Web Science*. ACM, 200–204.
- [39] Erving Goffman. 1974. *Frame Analysis: An Essay on the Organization of Experience*. Harvard University Press.
- [40] Xuanjun Gong, Richard Huskey, Haoning Xue, Cuihua Shen, and Seth Frey. 2023. Broadcast information diffusion processes on social media networks: Exogenous events lead to more integrated public discourse. *Journal of Communication* (2023), jqad014.
- [41] Doug Gross. 2014. Online comments are being phased out. Retrieved from <https://www.cnn.com/2014/11/21/tech/web/online-comment-sections/index.html>. Accessed on July 30, 2023.
- [42] Pedro Guerra, Roberto Nalon, Renato Assunção, and Wagner Meira Jr. 2017. Antagonism also flows through retweets: The impact of out-of-context quotes in opinion polarization analysis. In *Proceedings of the International AAAI Conference on Web and Social Media*.
- [43] Pedro Henrique Calais Guerra, Wagner Meira Jr, Claire Cardie, and Robert Kleinberg. 2013. A measure of polarization on social media networks based on community boundaries.. In *Proceedings of the International AAAI Conference on Web and Social Media*.
- [44] Daniel Halpern and Jennifer Gibbs. 2013. Social media as a catalyst for online deliberation? Exploring the affordances of Facebook and YouTube for political expression. *Computers in Human Behavior* 29, 3 (2013), 1159–1168.
- [45] Caroline Haythornthwaite, Priya Kumar, Anatoliy Gruzd, Sarah Gilbert, Marc Esteve del Valle, and Drew Paulin. 2018. Learning in the wild: Coding for learning and practice on reddit. *Learning, Media and Technology* 43, 3 (2018), 219–235.
- [46] Karen Henwood and Nick Pidgeon. 2003. Grounded theory in psychological research. In *Qualitative Research in Psychology: Expanding Perspectives in Methodology and Design*, P. M. Camic, J. E. Rhodes, and L. Yardley (Eds.). American Psychological Association, Washington, 131–155.
- [47] Sarah J. Jackson and Brooke Foucault Welles. 2015. Hijacking# myNYPD: Social media dissent and networked counterpublics. *Journal of Communication* 65, 6 (2015), 932–952.
- [48] Elaine Kamarck. 2019. The mueller testimony: Two narratives. Retrieved from <https://www.brookings.edu/blog/fixgov/2019/07/25/the-mueller-testimony-two-narratives/>. Accessed on July 30, 2023.
- [49] Sarah Kaplan. 2008. Framing contests: Strategy making under uncertainty. *Organization Science* 19, 5 (2008), 729–752.
- [50] Adam DI Kramer, Jamie E. Guillory, and Jeffrey T. Hancock. 2014. Experimental evidence of massive-scale emotional contagion through social networks. *Proceedings of the National Academy of Sciences* 111, 24 (2014), 8788–8790.
- [51] Eunji Lee, Jung-Ah Lee, Jang Ho Moon, and Yongjun Sung. 2015. Pictures speak louder than words: Motivations for using instagram. *Cyberpsychology, Behavior, and Social Networking* 18, 9 (2015), 552–556.
- [52] Jae Kook Lee, Jihyang Choi, Cheonsoo Kim, and Yonghwan Kim. 2014. Social media, network heterogeneity, and opinion polarization. *Journal of Communication* 64, 4 (2014), 702–722.
- [53] Eden Litt. 2012. Knock, knock. who's there? The imagined audience. *Journal of Broadcasting & Electronic Media* 56, 3 (2012), 330–345.

- [54] Walid Magdy, Kareem Darwish, Norah Abokhodair, Afshin Rahimi, and Timothy Baldwin. 2016. #isisisnotislam or# deportallmuslims? Predicting unspoken views. In *Proceedings of the 8th ACM Conference on Web Science*. 95–106.
- [55] Alice E. Marwick and Danah Boyd. 2011. I tweet honestly, I tweet passionately: Twitter users, context collapse, and the imagined audience. *New Media and Society* 13, 1 (2011), 114–133.
- [56] Mary L. McHugh. 2012. Interrater reliability: The kappa statistic. *Biochemia Medica: Biochemia Medica* 22, 3 (2012), 276–282.
- [57] Ivan Mehta. 2023. Twitter brings its “for you” and “following” dual-timeline view to the web. Retrieved from <https://techcrunch.com/2023/01/13/twitter-brings-its-for-you-and-following-dual-timeline-view-to-the-web/>. Accessed on July 30, 2023.
- [58] Sharon Meraz and Zizi Papacharissi. 2013. Networked gatekeeping and networked framing on# Egypt. *The International Journal of Press/Politics* 18, 2 (2013), 138–166.
- [59] Sharon Meraz and Zizi Papacharissi. 2016. Networked framing and gatekeeping. *The Sage Handbook of Digital Journalism*. London: Sage (2016), 95–112.
- [60] Robert S. Mueller. 2019. *The Mueller Report: Report on the Investigation Into Russian Interference in the 2016 Presidential Election*. WSBLD.
- [61] Mark EJ Newman. 2006. Modularity and community structure in networks. *Proceedings of the National Academy of Sciences* 103, 23 (2006), 8577–8582.
- [62] Zhongdang Pan and Gerald M. Kosicki. 1993. Framing analysis: An approach to news discourse. *Political Communication* 10, 1 (1993), 55–75.
- [63] Jay Peters. 2020. Twitter’s new reply-limiting feature is already changing how we talk on the platform. Retrieved from <https://www.theverge.com/2020/5/23/21266969/twitter-new-reply-limiting-feature-how-using-changing-talk>. Accessed on July 30, 2023.
- [64] James Pierce and Eric Paulos. 2014. Counterfunctional things: Exploring possibilities in designing digital limitations. In *Proceedings of the 2014 Conference on Designing Interactive Systems*. ACM, 375–384.
- [65] Barbara Poblete, Ruth Garcia, Marcelo Mendoza, and Alejandro Jaimes. 2011. Do all birds tweet the same? Characterizing Twitter around the world. In *Proceedings of the 20th ACM International Conference on Information and Knowledge Management*. 1025–1030.
- [66] Philip Pond and Jeff Lewis. 2019. Riots and Twitter: Connective politics, social media and framing discourses in the digital public sphere. *Information, Communication and Society* 22, 2 (2019), 213–231.
- [67] Jon Porter. 2020. Twitter quote tweets are now easier to find. Retrieved from <https://www.theverge.com/2020/9/1/21409925/twitter-quote-tweets-counter-retweet-with-comments-interface-ratiod>. Accessed on July 30, 2023.
- [68] Liza Potts, Joyce Seitzinger, Dave Jones, and Angela Harrison. 2011. Tweeting disaster: Hashtag constructions and collisions. In *Proceedings of the 29th ACM International Conference on Design of Communication*. 235–240.
- [69] Lily Puckett. 2019. Mueller hearings: How the left and right handled the ‘brilliant’ and ‘painful’ testimonies. Retrieved from <https://www.independent.co.uk/news/world/americas/mueller-hearings-testimony-republican-democrat-gop-trump-reaction-a9019666.html>. Accessed on July 30, 2023.
- [70] Hemant Purohit, Andrew Hampton, Valerie L. Shalin, Amit P. Sheth, John Flach, and Shreyansh Bhatt. 2013. What kind of# conversation is Twitter? Mining# psycholinguistic cues for emergency coordination. *Computers in Human Behavior* 29, 6 (2013), 2438–2447.
- [71] Daniel M. Romero, Brendan Meeder, and Jon Kleinberg. 2011. Differences in the mechanics of information diffusion across topics: Idioms, political hashtags, and complex contagion on twitter. In *Proceedings of the 20th International Conference on World Wide Web*. 695–704.
- [72] Kevin Roose. 2020. You’ve avoided the internet fever swamps, until now: What is QAnon, the viral and sometimes violent pro-Trump conspiracy theory? *New York Times* (2020).
- [73] Mattia Samory, Vincenzo-Maria Cappelleri, and Enoch Peserico. 2017. Quotes reveal community structure and interaction dynamics. In *Proceedings of the 2017 ACM Conference on Computer Supported Cooperative Work and Social Computing*. 322–335.
- [74] Heather Schwedel. 2017. “Dunking” is delicious sport. Retrieved from <https://slate.com/technology/2017/12/dunking-is-delicious-and-also-probably-making-twitter-terrible.html>. Accessed on July 30, 2023.
- [75] Bryan Semaan, Heather Faucett, Scott Robertson, Misa Maruyama, and Sara Douglas. 2015. Navigating imagined audiences: Motivations for participating in the online public sphere. In *Proceedings of the 18th ACM Conference on Computer Supported Cooperative Work and Social Computing*. 1158–1169.
- [76] Eva Sharma, Koustuv Saha, Sindhu Kiranmai Ernala, Sucheta Ghoshal, and Munmun De Choudhury. 2017. Analyzing ideological discourse on social media: A case study of the abortion debate. In *Proceedings of the 2017 International Conference of the Computational Social Science Society of the Americas*. 1–8.
- [77] Catherine Shu. 2015. Twitter officially launches its “retweet with comment” feature. Retrieved from <https://techcrunch.com/2015/04/06/retweetception/>. Accessed on July 30, 2023.

- [78] Leo Graiden Stewart, Ahmer Arif, A. Conrad Nied, Emma S. Spiro, and Kate Starbird. 2017. Drawing the lines of contention: Networked frame contests within# BlackLivesMatter discourse. *Proceedings of the ACM on Human-Computer Interaction* 1, CSCW (2017), 1–23.
- [79] Chris Stokel-Walker. 2022. Should I join mastodon? A scientists’ guide to Twitter’s rival. *Nature* (2022). DOI :<https://doi.org/10.1038/d41586-022-03668-7>. Accessed on July 30, 2023.
- [80] Biz Stone. 2007. Are you twittering @ me? Retrieved from https://blog.twitter.com/official/en_us/a/2007/are-you-tweeting-me.html. *Twitter Blog* (May 2007).
- [81] Clive Thompson. 2022. Twitter alternative: How mastodon is designed to be “antiviral”. Retrieved from <https://uxdesign.cc/mastodon-is-antiviral-design-42f090ab8d51#>. Accessed on July 30, 2023.
- [82] Peter Tolmie, Rob Procter, Mark Rouncefield, Maria Liakata, and Arkaitz Zubiaga. 2018. Microblog analysis as a program of work. *ACM Transactions on Social Computing* 1, 1 (2018), 1–40.
- [83] Joshua A. Tucker, Andrew Guess, Pablo Barberá, Cristian Vaccari, Alexandra Siegel, Sergey Sanovich, Denis Stukal, and Brendan Nyhan. 2018. Social media, political polarization, and political disinformation: A review of the scientific literature. *Political Polarization, and Political Disinformation: A Review of the Scientific Literature* (2018).
- [84] Nikki Usher, Jesse Holcomb, and Justin Littman. 2018. Twitter makes it worse: Political journalists, gendered echo chambers, and the amplification of gender bias. *The International Journal of Press/Politics* 23, 3 (2018), 324–344.
- [85] Quintan Wiktorowicz. 2004. Framing jihad: Intramovement framing contests and al-qaeda’s struggle for sacred authority. *International Review of Social History* 49, S12 (2004), 159–177.
- [86] Suzanne Xie. 2020. New conversation settings, coming to a tweet near you. Retrieved from https://blog.twitter.com/en_us/topics/product/2020/new-conversation-settings-coming-to-a-tweet-near-you.html. Accessed on July 30, 2023.
- [87] Suzanne Xie. 2020. Twitter brings its “for you” and “following” dual-timeline view to the web. Retrieved from https://blog.twitter.com/en_us/topics/product/2020/new-conversation-settings-coming-to-a-tweet-near-you. Accessed on July 30, 2023.
- [88] Sanam Yar and Ian Prasad Philbrick. 2020. The rise of QAnon. *New York Times* (2020). <https://www.nytimes.com/2020/08/13/briefing/qanon-kamala-harris-coronavirus-your-thursday-briefing.html>. Accessed on July 30, 2023.
- [89] Arkaitz Zubiaga, Maria Liakata, Rob Procter, Geraldine Wong Sak Hoi, and Peter Tolmie. 2016. Analysing how people orient to and spread rumours in social media by looking at conversational threads. *PLoS One* 11, 3 (2016), e0150989.

Received 17 April 2023; revised 15 July 2023; accepted 30 July 2023