

Does Geo-co-location Matter? A Case Study of Public Health Conversations during COVID-19

Paiheng Xu, Louiqa Raschid, Vanessa Frias-Martinez

University of Maryland, College Park
paiheng@cs.umd.edu, lraschid@umd.edu, vfrias@umd.edu

Abstract

Social media platforms like Twitter (now X) have been pivotal in information dissemination and public engagement, especially during COVID-19. A key goal for public health experts was to encourage prosocial behavior that could impact local outcomes such as masking and social distancing. Given the importance of local news and guidance during COVID-19, the objective of our research is to analyze the effect of localized engagement, on social media conversations. This study examines the impact of geographic co-location, as a proxy for localized engagement between public health experts (PHEs) and the public, on social media. We analyze a Twitter conversation dataset from January 2020 to November 2021, comprising over 19 K tweets from nearly five hundred PHEs, along with approximately 800 K replies from 350 K participants. Our findings reveal that geo-co-location is associated with higher engagement rates, especially in conversations on topics including masking, lockdowns, and education, and in conversations with academic and medical professionals. Lexical features associated with emotion and personal experiences were more common in geo-co-located contexts. This research provides insights into how geographic co-location influences social media engagement and can inform strategies to improve public health messaging.

1 Introduction

Social media platforms such as Twitter (now X) often play a crucial role in information dissemination and public engagement. The COVID-19 pandemic presented a complex social media communication challenge along many dimensions: public health experts had to provide vital guidance to the public, in a scenario where the science was constantly being updated, and in the midst of a social media deluge of mis- and dis-information.

A key goal for public health experts was to encourage prosocial behavior that could impact In Real Life (IRL) outcomes (Miles et al. 2022). Pro-social behavior including masking and social distancing was critical during the initial stages, in the absence of vaccines and treatments. Given the differences in COVID-19 effects across regions, and the lack of unifying federal guidance, states and local government played an important role in developing regulations. The public had to rely on news and guidance, at the local level, to receive timely and relevant information (Fischer 2020).

Given the importance of local news and guidance during COVID-19, the objective of our research is to analyze the effect of localized engagement, on social media conversations. We use geographic co-location (which we formally define in the paper) between public health experts and the public as a proxy for localized engagement. Prior research has shown that the physical distance between social media users can impact their online interactions; for example, the probability of online interaction between users in a social network lessens exponentially with an increase in the physical distance between these users (Liben-Nowell et al. 2005; Leskovec and Horvitz 2008; Newman, Barabási, and Watts 2011). It has also been shown that users on Twitter consume and retweet more Tweets from users who are co-located in their own country (Kulshrestha et al. 2012; Cuevas et al. 2014). A study on Reddit (Bozarth et al. 2023) found that news sharing and user interactions occur more often within a state, when compared to across states. While there has been some recent research on social media engagement during COVID-19 (Rao et al. 2023; Bojja et al. 2020; Gallagher et al. 2021), there is very limited prior research on conversations involving public health experts and their peers (Rao et al. 2024).

In this paper, we hypothesize that geographic co-location between public health experts and the public may promote greater engagement on social media. Our novel research advances the state of the art by analyzing the impact of geographic co-location between public health experts and the public on the following: (1) The engagement rate, as measured by the count of participants and replies in a conversation comprising an original tweet and replies. (2) The lexical features that are expressed in the conversation. We also evaluate the effect of geographic co-location on (3) the COVID-19 topics being discussed and on (4) the profile (profession) of the public health expert. We note that our research is the first to understand the role of geographic co-location on engagement. The study also contributes to the literature on the effectiveness of public health messaging during COVID-19.

Our research is based on an X (Twitter) dataset spanning January 2020 through November 2021. Starting with a seed set of thirty public health experts (PHEs), we identified a larger group of almost five hundred PHEs or their expert peers in adjacent domains, e.g., medical communications and outreach. We then curated a dataset of conversations,

comprising over 19 K posts (tweets) from the PHEs, and approximately 800 K replies from almost 350 K participants in the conversations.

We used Carmen (Dredze et al. 2013; Zhang, DeLucia, and Dredze 2022) to determine the geo-location (US state) of the PHEs and the participants. We then identified geo-co-located pairs - where the PHE and participant are co-located in the same US state, and non-geo-co-located pairs, respectively. Linguistic Inquiry and Word Count (LIWC) (Pennebaker et al. 2015) was used to identify lexical features that capture levels of self-disclosure and psychological processes. We used machine learning models from (Rao et al. 2023) to identify COVID-19 topics in the original post (tweet), and Large Language Models (LLMs) to determine the profile (professions) of the PHE authors.

Our analysis reveals the following key findings:

1. Statistical tests show that the engagement rates when the PHEs and the public are geo-co-located are significantly higher than the engagement rates when they are not geo-co-located.
2. Conversations discussing COVID-19 issues such as masking, lockdown, and education have higher geo-co-located engagement rates. In contrast, vaccine-related conversations show an opposite trend, eliciting more engagement from non-geo-co-located participants.
3. Based on lexical analysis, the presence of first-person pronouns and emotional expressions, e.g., anxiety, in a tweet, are associated with a larger gap between the geo-co-located and non-geo-co-located engagement rates. This indicates that conversations that share personal experiences or emotions are associated with higher geo-co-located engagement rates.
4. Replies from participants who are geo-co-located with PHEs are generally more positive in sentiment, and express more emotions and positive personal issues when the PHEs share personal experiences / feelings, in comparison to non-geo-co-located replies.
5. Geo-co-location has a greater impact on engagement for conversations originated by PHEs representing academic and medical professions, and a somewhat lesser impact for PHEs in media and policy / political professions.

This research provides valuable novel insights into understanding the impact of localization or geographic co-location, on social media engagement. It has the potential to inform how public health messaging strategies can harness geo-co-location, to enhance engagement, and to maintain a positive discourse around public health guidance in online communities.

2 Related Work

Geo-co-location in social networks Many studies have shown that users on social media platforms tend to connect or interact preferentially with users who are geographically closer over those who are farther away (Leskovec and Horvitz 2008; Crandall et al. 2010; Scellato et al. 2011; Laniado et al. 2018). Similar findings have also been replicated for users on cellphone social networks (Hong, Frias-Martinez, and Frias-Martinez 2016; Frias-Martinez et al.

2012b). Our study focuses on message-exchange interactions (i.e., replying to an original post) on platforms such as Twitter, where interactions typically reflect a shared interest in specific topics (Kwak et al. 2010). We specifically investigate the shared geo-location between users, which differs from prior research focusing on the correlation between users' geo-location and content of discussion (Hu, Farnham, and Talamadupula 2015; Pavalanathan and Eisenstein 2015; Cheke et al. 2020). While many researchers have used Twitter interactions in the form of replies and mentions as predictors of geographically closer connections (Sadilek, Kautz, and Bigham 2012; McGee, Caverlee, and Cheng 2013; Jurgens et al. 2015), few have considered the opposite direction, i.e., how geo-co-location influences users' interaction behaviors. Kulshrestha et al. (2012) analyze the information flow between different countries by counting the number of information producers and consumers in each country and they consider all followers of a user as the consumers of the tweets shared by this user. Similarly, Cuevas et al. (2014) count the number of retweets and compare how tweets are shared across countries; while Bozarth et al. (2023) finds geographical diffusion of news articles on Reddit is rare among states in the U.S. However, all these studies are either limited to a country-level analysis or fail to investigate the content involved in the interactions. With Twitter releasing its API v2 in late 2020, which facilitates the collection of all conversations sparked by an original post, we can more closely study how geo-co-location affects the tweets and their replies.

Geo-location data on Twitter Twitter provides geo-location information in the metadata of tweets and users, including coordinates and the Place object from tweets, as well as location information in user profiles. The geo-location of a tweet can also be inferred from its content (Mahmud, Nichols, and Drews 2012; Izbicki, Papalexakis, and Tsotras 2019). Given that our study focuses on the geo-location of users, we rely on tools that extract geo-location from user-provided metadata in tweets and user profiles (Dredze et al. 2013; Zhang, DeLucia, and Dredze 2022). We acknowledge that Twitter users (Wojcik and Hughes 2019) and those who use Twitter's geo-location services (Sloan and Morgan 2015; Pavalanathan and Eisenstein 2015; Karami et al. 2021) may be a biased sample of the population. Despite that limitation, Twitter's geo-location data has been shown to be valuable in many applications such as disaster response (Crooks et al. 2013; Ghahremanlou, Sherchan, and Thom 2015; Hong and Frias-Martinez 2020), urban planning (Frias-Martinez et al. 2012a; Milusheva et al. 2021) and public health surveillance (Jordan et al. 2018; Xu, Dredze, and Broniatowski 2020; Sigalo, Frias-Martinez et al. 2023; Xu, Broniatowski, and Dredze 2024; Sigalo et al. 2023).

Content understanding for interactions on social media Numerous text analysis tools are employed to comprehend user interactions on social media, such as topic analysis (Cheke et al. 2020; Li et al. 2024), linguistic analysis (Pavalanathan and Eisenstein 2015; Choi et al. 2021; Wood-Doughty et al. 2021), and text classification models for constructs like sentiment (Lwin et al. 2020) and emo-

tion (Wheaton, Prikhidko, and Messner 2021). With the recent advance in Large Language Models (LLMs), many works have explored the possibilities of using LLMs to understand social questions (Ziems et al. 2024; Zhou et al. 2024). In this work, we utilize a wide range of text analysis tools, including topic modeling (Eisenstein, Ahmed, and Xing 2011), LIWC for linguistic analysis (Pennebaker et al. 2015), sentiment classification (Loureiro et al. 2022), and LLMs, to answer research questions related to geo-co-location engagement.

3 Data

The Twitter conversation dataset used in this study is centered around COVID-19 discussions. The creation of the dataset commenced with a curated seed user list of 30 Public Health Experts (PHEs) (see Appendix A), both academics and medical professionals, who were active in COVID-19 discussions. These 30 experts were handpicked by academic colleagues in public health. A network of retweets of these 30 seed experts was then constructed using a dataset of over 1 billion publicly available COVID-19 Tweets (Chen et al. 2020). All users in the retweet network were ranked using eigenvector centrality (Ghosh and Lerman 2010) and the Top 500 ranked users were identified. A small count of users were manually identified as organizations or bots and were removed. The result was an expanded set of 489 PHEs or PHE adjacent influential users. An example of a PHE adjacent influential user may be a medical reporter or a policy expert or an appointee in a medical or health agency who actively retweets COVID-19 tweets from the seed PHEs.

The next step was to collect all of the original posts (tweets) of PHEs, and the resulting conversations, comprised of replies¹. The dataset covered a period from the early days of the pandemic, starting in January 21, 2020, through November 4, 2021. We only consider the original posts (tweets) of the 489 PHEs, i.e., we exclude their own retweets, replies, and quoted tweets, giving a dataset of 144 K tweets. The final step was to collect all of the replies for the original tweets. Unfortunately, restrictions and rate limits imposed by Twitter on the academic API limited our ability to rehydrate all of the replies (conversations). We were able to obtain all reply tweets for 19.5 K original tweets, from 462 PHEs, that had at least one reply. The final dataset included approximately 786 K reply tweets from 345 K unique participants in conversations with the 462 PHEs. Each tweet received an average of 40.24 replies, with a median of 5 replies per original PHE tweet, and approximately 70% of the tweets did not have replies.

In the rest of the paper, we refer to an original PHE tweet as a *tweet* and a reply tweet as a *reply*. We refer to the PHEs as either *authors* or *PHE authors*. All users who responded with a reply tweet to an original tweet are referred to as *participants* in the conversation. Finally, we refer to the union of participants who reply to any of the original tweets of a PHE author as the *audience* of the PHE author.

¹<https://developer.twitter.com/en/docs/twitter-api/conversation-id>

4 Geo-Co-Location and Engagement

We report on the activity statistics of PHE authors and participants, comparing those who share geo-location and those who do not. We then define geo-co-located engagement. Prior research has shown that social media users who share their location information are more active (Rzeszewski and Beluch 2017).

4.1 Geo-Located Activity Statistics

We use the Carmen library (Dredze et al. 2013; Zhang, DeLucia, and Dredze 2022) to infer the geo-location of the PHE authors and the participants. Carmen uses the following features: location information and self-description from the user profile, coordinates, and the Place object from the meta-data associated with tweets. The most frequently mentioned (inferred) state is selected as the home state for the user. We note that Carmen could infer the geo-location (state) of over 2/3 of the PHEs (Table 1) and over 1/3 of the participants (Table 2).

We report on activity statistics comparing PHE authors, and participants, who share or do not share geo-location, in Tables 1 and 2, respectively. The mean and standard deviation (Std) of the count of replies per participant, for tweets from PHE authors with geo-location, is 2.29 ± 6.55 , in comparison to 1.91 ± 7.82 for tweets from PHE authors without geo-location information. We used the two-sample Kolmogorov-Smirnov (KS) test (Hodges Jr 1958), a non-parametric statistical test, to assess the statistical significance of the difference between the samples. It rejected the null hypothesis ($p < 0.001$), pointing to a statistically significant increase in engagement from participants for PHE authors with geo-location information. The mean and Std of the count of replies per tweet, for PHE authors with geo-location is 23.08 ± 53.93 , in comparison to 25.98 ± 60.79 for PHE authors without geo-location. However that difference is not statistically significant (the null hypothesis was not rejected with $p\text{-value} > 0.01$).

We conclude that geo-location information for PHE authors does appear to impact engagement at the participant level, with higher engagement for PHEs that have geo-location information; however, that impact is not significant when measuring the intensity of responses.

Next, we consider the statistics for the *participants* in Table 2. We observe that participants with geo-location information post significantly more replies, with a mean and standard deviation of 2.45 ± 8.33 , than participants without geo-location, with mean and standard deviations of 2.38 ± 7.43 . This difference is statistically significant (two-sample KS test rejected the null hypothesis with $p < 0.05$).

To summarize, these activity statistics reflect that sharing geo-location information is an important signal for engagement in social media conversations, for both PHE authors and participants. This finding is also consistent with prior results that social media users who share their location information are more active (Rzeszewski and Beluch 2017). This further motivates our study to understand the impact of geographic co-location on social media engagement. Next, we formalize our definition of geo-co-located engagement and then we present our results in Section 5.

	# Authors	# Tweets	# Replies	# Participants	Mean (Std) followers	Mean (Std) replies per tweet	Mean (Std) replies per participant
With state	342	16,465	380k	166k	157k (654K)	23.08 (53.93)	2.29 (6.55)
No state	120	4,432	115k	60k	160k (438K)	25.98 (60.79)	1.91 (7.82)

Table 1: Statistics for PHEs with and without geo-location in the COVID-19 Conversation Dataset. We **bold** the statistic comparison between the PHEs with and without state that is significant from a two-sample KS test ($p < 0.001$).

	# Participants	Mean (Std) replies
With state	73,912	2.45 (8.33)
No state	131,760	2.38 (7.43)

Table 2: Statistics for participants with and without geo-location in the COVID-19 Conversation Dataset. We **bold** the statistic comparison between participants with and without state that is significant from a two-sample KS test ($p < 0.05$).

4.2 Geo-co-located Engagement

Geo-co-location (*gcl*) For each (tweet, reply) pair in the dataset, we refer to the participant being geo-co-located (*gcl*) with the author of the tweet, if they are both located in the same state. Conversely, the participant and the author are not geo-co-located (*non-gcl*) if the participant and the author are located in different states. Participants with no geo-location, i.e., state, information are excluded from the *gcl* and *non-gcl* groups for further analysis.

Engagement Rate For a tweet T from PHE author A , we define the *gcl audience* as the set of all *gcl participants* across all the tweets created by A in the dataset. The *gcl engagement rate* for a tweet T from author A is then defined as the fraction of A 's *gcl audience* that replied to their tweet T i.e.,

$$gcl\ engagement\ rate = \frac{\text{count of } gcl\ participants\ for\ T}{\text{count of } gcl\ audience\ for\ A},$$

The *non-gcl engagement rate* for tweet T of author A is similarly defined using *non-gcl participants* and *non-gcl audience*.

Justification for Engagement Rate Using the total audience size of an author as the normalization term for the proposed engagement rate enables direct comparisons between *gcl* and *non-gcl* engagement rates. On average, the *gcl* replies take up 14.2% of the replies for a Tweet while *non-gcl* replies take up 20.4%. With the proposed measurement, we can compare the distributions of *gcl* and *non-gcl* engagement rates with appropriate statistical tests.

5 Understanding the Effect of Geo-co-location on Engagement

Our research on analyzing the effect of *gcl* on engagement is organized and presented as follows: First, we test our hypothesis that *gcl* promotes greater engagement, i.e., attracts more replies. Next, we explore this hypothesis disaggregated by COVID-19 topics i.e., whether greater *gcl* engagement is dependent on the type of issue being discussed. We then

drill down into the content, to identify lexical features that may reveal additional insights into the emotions and concerns that are expressed in *gcl* versus *non-gcl* conversations. Finally, we determine if the PHE author profile, and in particular, the professional, has an impact on attracting greater rates of *gcl* engagement.

5.1 Does *gcl* promote more engagement?

We first investigate whether geo-co-location (*gcl*) promotes greater engagement by comparing the distribution of *gcl* and *non-gcl* engagement rates. Since computing *gcl* and *non-gcl* engagement rates requires authors to have both a *gcl* and *non-gcl* audience, we filter the dataset to only include authors that meet this constraint. This results in a subset of 256 authors and 15,926 original tweets. The 256 authors represent approximately 75% of the 342 authors that have state information. However, we note that the subset of 15,926 tweets includes approximately 97% of the tweets from the 342 authors, and it covers $\approx 76\%$ of the entire tweet dataset.² Thus, our analysis focuses on the more active authors who share their location, and on a large, representative fraction of the dataset.

Table 3 displays the engagement rates (ER) for *gcl* and *non-gcl* audiences, as well as for audiences where the location of the participant could not be inferred (labeled as *no-loc*). The table shows that a higher average engagement rate is observed for *gcl* audiences, in comparison to *non-gcl* audiences (mean values: 1.89% vs 1.76%). A two-sample KS test confirmed the significance of this difference ($n = 15,926$, $p < 0.001$). In addition, the *no-loc* engagement rate is very similar to the *non-gcl* rate (mean values: 1.76% vs 1.78%). This validates that our focus on comparing *gcl* versus *non-gcl* subsets is a reasonable experimental filter.

Finally, Figure 1 provides a histogram of the *gcl* (red hatched bars) and *non-gcl* (blue bars) engagement rates. The significant rightward skew of the *gcl* histogram clearly illustrates its dominance in the region with a higher engagement rate region; for example, we see the histogram extend to a 100% engagement rate.

In summary, the statistical significance of the *gcl* versus *non-gcl* distributions, and the insights from the histograms, confirm our hypothesis that geographical co-location between PHE authors and participants promotes greater engagement on social media.

5.2 What Tweets attract more *gcl* engagement?

Impact of COVID-19 Concerns We first examine whether the *gcl* engagement rates vary across Tweets with

² $78.8\% \times 97\% \approx 76\%$, where 78.8% is the percentage of original posts from the 342 authors with state, calculated from Table 1.

	Mean (Std) participants per tweet	Mean (Std) ER
<i>gcl</i>	0.99 (3.32)	1.89% (8.05%)
<i>non-gcl</i>	5.16 (13.09)	1.76% (6.10%)
<i>no-loc</i>	11.18 (26.56)	1.78% (5.53%)

Table 3: Engagement rate (ER) for different types of participants. *no-loc* means there is no location information for the participants. We **bold** the statistic comparison between *gcl* and *non-gcl* groups that is significant from a two-sample KS test ($p < 0.05$).

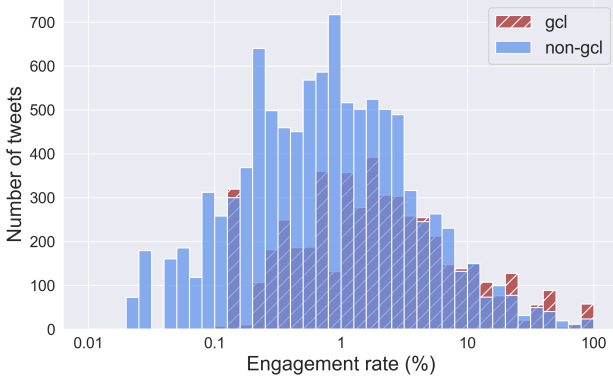


Figure 1: Distribution of *gcl* and *non-gcl* engagement rate. Tweets with 0% engagement rate are not shown in the plot.

different topics, and how these variations relate to *gcl* versus *non-gcl* engagement. We adopt a method from Rao et al. (2023) to identify tweets that discuss different COVID-19 issues. Their method extracts issue-relevant keywords from Wikipedia articles by identifying the most distinctive keywords when comparing issue-specific articles with general articles on the pandemic and politics in the U.S. using SAGE (Eisenstein, Ahmed, and Xing 2011). These manually verified keywords are then used to label tweets through keyword matching. Rao et al. (2023) validate their method by evaluating its results on a test sample with manual annotations. We replicate this method and retain issues that have F1 scores higher than 0.7; the issues are *masking*, *lockdowns*, *education*, and *vaccines*. Details can be found in Appendix B, where Table 10 shows sample tweets about each issue.

We repeat the engagement rate analysis in Section 5.1 on each of the subsets of tweets identified representing different COVID-19 issues namely, masking, lockdowns, education, and vaccines. Table 4 shows that masking-related tweets have the highest engagement rates for both *gcl* and *non-gcl* audiences and present the largest gap between these two types of engagement rates (two-sample KS test was significant at $p < 0.001$). This is consistent with Rao et al. (2023) that finds that masking-related tweets gained more replies. Similarly to masking issues, lockdown- and education-related Tweets also show higher engagement rates for *gcl* audiences when compared to *non-gcl*, and these differences were all statistically significant. On the other hand, **vaccine-related Tweets attract more engagement from**

	Masking	Lockdowns	Education	Vaccines
<i>gcl</i> ER	2.32%	1.49%	1.94%	1.52%
<i>non-gcl</i> ER	1.88%	1.39%	1.80%	1.73%
Diff	0.44%*	0.10%*	0.14%*	-0.21%*
# authors	162	107	168	202
# Tweets	929	326	812	2240

Table 4: Difference between *gcl* and *non-gcl* engagement rate (ER) for Tweets discussing different COVID-19 issues. * indicates that KS test results for *gcl* ER and *non-gcl* ER are significant at $p < .001$.

non-gcl audiences.

LIWC Analysis LIWC has been extensively used to study psychological patterns and emotional states in social media texts, providing insights into user behavior and social dynamics (Schwartz et al. 2013; Jiang and Wilson 2018; Choi et al. 2021). To explore whether psychological and emotional states - measured via LIWC categories - might affect the *gcl* engagement rates, when compared to *non-gcl* audiences, we propose the following approach. We count the occurrences of the LIWC categories in each original Tweet and normalize them by the number of tokens in the Tweet. We use the TweetTokenizer from NLTK (Bird, Klein, and Loper 2009) for tokenization. Next, we run coefficient regression analyses with the occurrences of LIWC categories in the original Tweets as independent variables, and with the engagement rates as dependent variables, to assess the effect of LIWC categories on engagement rates. We include the following LIWC categories: first person singular pronouns (*i*), first person plural pronouns (*we*), and second person pronouns (*you*) because the usage of pronouns reveal levels of self-disclosure (De Choudhury and De 2014; Wang, Burke, and Kraut 2016) and they are also some of the most distinctive features in the tweets with high *gcl* engagement rate based on our preliminary analysis (see Appendix C). Additionally, we choose subcategories in psychological processes potentially relevant to COVID-19 discussions. These include: (1) all subcategories in the affective process, i.e., positive emotion (*posemo*), negative emotion (*negemo*) and its subcategories (*anxiety* (*anx*), *anger* and *sadness* (*sad*)); (2) selected subcategories in the social processes, i.e., *family* and *friends*; (3) selected subcategories in the biological processes, i.e., *body* and *health*; (4) all subcategories in the personal concerns, i.e., *work*, *leisure*, *home*, *money*, *religion* (*relig*), and *death*; and (5) swear words (*swear*) from informal language, resulting a total of 19 variables. We define the linear regression as follows:

$$\text{Engagement Rate} = \beta_0 + \sum_{L_i \in \text{LIWC}} (\beta_i L_i) + \epsilon,$$

where L_i and β_i are selected LIWC categories and their coefficients, i indexes the categories and starts with 1, β_0 and ϵ are the constant and error term respectively. We repeat the analysis for three types of engagement rate as the dependent variable, i.e., *gcl*, *non-gcl* and the difference between *gcl* and *non-gcl* (denoted as *Diff*).

Additionally, we repeat the coefficient regression analyses while controlling for number of followers (# followers), as this is recognized as one of the driving factors for engagement on social media (Cha et al. 2010; Bakshy et al. 2011). We also control for the normalization terms of the proposed engagement rates, i.e., the sizes of *gcl* and *non-gcl* audiences (*gcl* audience and *non-gcl* audience). We min-max normalize these controlling factors. Moreover, we include four categorical variables to indicate whether the Tweet belongs to each of the four COVID-19 issues.

Table 5 show the coefficient regression analysis for the model with both LIWC categories and control variables. Similar results were observed without the controls. We discuss the results using Diff as the dependent variable, since it cancels out the effect of LIWC categories that are significant for both *gcl* and *non-gcl* engagement rates, such as *posemo*, *family*, *gcl* audience and *non-gcl* audience, and allows us to focus on the LIWC categories that have a significant effect on increasing or decreasing the *gcl* engagement rate. We can observe that the coefficients for *i* and *anx* are significant and positive, indicating that a 1% increase of word usage in these two categories would increase the difference between *gcl* and *non-gcl* engagement rates by 0.1% and 0.13% respectively (*gcl* ER is higher). This suggests that **PHEs sharing personal experiences and feelings (especially anxiety) attract higher engagement from *gcl* audiences**. On the other hand, the table also shows a significant negative coefficient for Tweets discussing vaccines (*vaccines_True*), suggesting an increase in *non-gcl* engagement rates with respect to *gcl* for vaccine-related Tweets. This is consistent with our results in Table 4. We also experimented with a mixed-effect model (Lindstrom and Bates 1988) grouped by authors and found similar results.

Qualitative Analysis Table 6 presents sample tweets and their *gcl* and *non-gcl* replies. The author of the first tweet is located in California yet the tweet is about Michigan, resulting in a higher *non-gcl* engagement rate. The other two examples demonstrate that sharing personal experiences (with first-person pronouns) attracts higher *gcl* and *non-gcl* engagement rates. However, their *gcl* engagement rates are higher than *non-gcl* ones, which is aligned with the regression results in Table 5, given that the occurrence of *i* is associated with sharing personal experiences or feelings.

5.3 How are *gcl* replies different than *non-gcl* ones?

Our linguistic analysis has focused on the original Tweets. Now, we shift to the replies and attempt to understand the differences in language usage between *gcl* replies and *non-gcl* replies. We first do a sentiment analysis and continue with a coefficient regression analysis to understand the effect of LIWC features on *gcl* and *non-gcl* engagement.

Sentiment We start by looking into the sentiment expressed in the *gcl* and *non-gcl* replies. We use a RoBERTa-base model (Loureiro et al. 2022) trained on 124M tweets from January 2018 to December 2021, and fine-tuned for

	Diff	<i>gcl</i>	<i>non-gcl</i>
Selected LIWC categories			
i	0.102*** (0.038)	0.420*** (0.037)	0.318*** (0.028)
<i>posemo</i>	0.038 (0.024)	0.087*** (0.023)	0.049*** (0.017)
<i>negemo</i>	-0.070 (0.046)	0.012 (0.044)	0.082** (0.034)
anx	0.128* (0.069)	-0.052 (0.067)	-0.179*** (0.050)
<i>family</i>	0.104 (0.086)	0.236*** (0.083)	0.132** (0.063)
<i>death</i>	0.020 (0.041)	-0.041 (0.040)	-0.061** (0.030)
Control variables			
# followers	-0.001 (0.008)	0.011 (0.008)	0.012* (0.006)
<i>gcl</i> audience	0.000 (0.004)	-0.016*** (0.003)	-0.016*** (0.003)
<i>non-gcl</i> audience	-0.002 (0.003)	-0.025*** (0.003)	-0.023*** (0.002)
<i>masking_True</i>	0.004 (0.003)	0.004 (0.003)	0.000 (0.002)
<i>education_True</i>	-0.001 (0.003)	-0.001 (0.003)	-0.000 (0.002)
<i>lockdowns_True</i>	-0.003 (0.005)	-0.003 (0.004)	-0.000 (0.003)
<i>vaccines_True</i>	-0.004** (0.002)	-0.005*** (0.002)	-0.001 (0.001)
R-squared	0.002	0.024	0.031
R-squared Adj.	0.000	0.023	0.029

Table 5: LIWC regression results of LIWC categories’ impact on different types of engagement rates. We only show LIWC categories that are significant in at least one of the regression models (columns). Standard errors in parentheses. * $p < .1$, ** $p < .05$, *** $p < .01$. Full table in Appendix D.

sentiment analysis with the TweetEval benchmark (Rosenthal, Farra, and Nakov 2017). We classify each reply as positive, neutral, or negative. For each original Tweet T , we compute the positive rate for its *gcl* replies and *non-gcl* replies, respectively. We find that *gcl* replies are slightly more positive and less negative than *non-gcl* replies. The mean percentage of *gcl* replies labeled as positive is 14% and the standard deviation is 29%, while the mean percentage for *non-gcl* replies labeled as positive is $13\% \pm 22\%$. The difference between the two distributions was statistically significant (KS test with $p < 0.001$ and $n = 3,521$). The mean percentage of *gcl* and *non-gcl* replies labeled as negative were also significantly different ($p < 0.001$) with mean and standard deviation values of $43\% \pm 39\%$ and $45\% \pm 29\%$, respectively.

LIWC analysis To examine differences in language usage between *gcl* and *non-gcl* replies, we propose to run a logistic regression where the dependent variable indicates whether a reply is *gcl* or *non-gcl* and the independent variables are the LIWC categories described in Section 5.2. A coefficient

Tweet - <i>Author state</i>	<i>gcl</i> ER	<i>non-gcl</i> ER	<i>gcl</i> reply - <i>participant state</i>	<i>non-gcl</i> reply - <i>participant state</i>
2 emerging US surges: Michigan (B.1.1.7) and Northeast (B.1.1.7+B.1.526). Both variants are vaccine responsive. They need very aggressive vaccination efforts + avoid relaxing mitigation. – CA	0.7%	1.8%	The little inflection that is not receiving adequate attention [URL] – CA	Michigan has a comprehensive VAX plan in place. [url] – MI
Anybody have ideas for pandemic-era holiday or birthday gifts? I’m thinking books, obviously, and magazine subscriptions. Gift certificates to nearby restaurants with delivery or safe take-out? Fun masks – to grim? Outdoor gear? [url] – DC	42.9%	5.7%	Gift certificates to nearby local movie theatres that are streaming like @[user] – DC	Flowbee, Netflix subscription, Zoom subscription.... Oh wait, do you mean for me or other people? – WA
Had another one of my “I’m in a crowded indoor place where no one is masked and then I suddenly remember Covid” dreams last night. How many of you have those too? – NY	41.7%	38.3%	All the time. But I genuinely wonder in the dream, “wait. Is ‘it’ over so it’s ok no one around me is masked??” – NY	I’ve had about 15 ... no joke! Horror films in my head running for exits through unmasked hoards. – CA

Table 6: Sampled tweets and replies with *gcl* and *non-gcl* engagement rate (ER). First-personal pronouns in the tweets that indicates self-disclosure are colored in red.

analysis will allow us to identify linguistic features that are uniquely associated with *gcl* or *non-gcl* replies. As the Table 7 shows, negative emotions (particularly anger) are more prevalent in *non-gcl* replies; while the use of pronouns such as *we* and *you* is more common in *gcl* replies, which might suggest more direct communication with the *gcl* PHE authors. Interestingly, *gcl* and *non-gcl* replies are associated with different types of personal concerns. For example, *gcl* replies involve topics related to *work*, *home*, and *money* more frequently, while *non-gcl* replies are more likely to discuss *relig* and death and biological processes (i.e., *body* and *health*).

From Original Tweets to Replies We extend our analysis by investigating how specific language (LIWC categories) in the original Tweets might elicit similar language (LIWC categories) in their corresponding *gcl* and *non-gcl* replies. To carry out this analysis, we run logistic regressions with the independent variables being all the LIWC categories in the original Tweet T , and the dependent variable being one binarized LIWC category present in the replies for that tweet T . Formally, the model is defined as:

$$L_i^{\text{gcl_replies}} = \beta_0 + \sum_{L_i^T \in \text{LIWC}} (\beta_i L_i^T) + \epsilon,$$

where $L_i^{\text{gcl_replies}}$ is a binary value of 1 if the i -th LIWC category is present in at least one of the *gcl* replies to T , and L_i^T is the i -th LIWC score of T . We run this model for each binarized LIWC category (dependent variable), and for *gcl* and *non-gcl* replies, separately. A comparison of the regression coefficients for the *gcl* and *non-gcl* groups (see Figure 2) reveals a strong diagonal effect for both groups, with slightly larger diagonal coefficients in *gcl* replies for most personal concerns (*work*, *home*, *relig*, and *death*) and swear words. This suggests that *gcl* replies may be more

likely to discuss personal concerns compared to *non-gcl* replies, when these concerns have been raised in the original Tweets. However, there is a stronger agreement in *non-gcl* replies for LIWC categories associated with affective processes, including *posemo* and *negemo* and all its subcategories, *anx*, *anger*, and *sad*. This suggests that, while *gcl* and *non-gcl* replies discuss similar types of personal concerns when they are mentioned in the original Tweets, *non-gcl* replies are more likely to share similar emotions to the original Tweets. We hypothesize that this may be due to *non-gcl* participants having weaker social connections with the authors, leading them to share similar emotions and potentially avoiding conflict, while *gcl* participants express more freely a diverse set of emotions independently of the ones that are raised in the original Tweets from local authors.

More interestingly, when observing the heatmaps horizontally, the coefficients for i are significant and positive in almost all regressions for *non-gcl* replies, while only significant and positive in *posemo*, social processes (*family* and *friend*), and some personal concerns (*leisure*, *home*, *relig*) for *gcl* replies. This pattern indicates that **authors sharing personal experiences tend to elicit all kinds of psychological and emotional states from *non-gcl* replies but mostly positive *gcl* replies. This finding reveals that geo-co-located messaging could be used as a strategy for enhancing engagement and maintaining a positive discourse in online communities.**

5.4 What types of PHEs have more *gcl* engagement?

As understanding who disseminates information is as crucial as the content itself on social media, we shift our focus toward the tweet authors. In the public health context, we want to know what kinds of PHEs attract more *gcl* engagement.

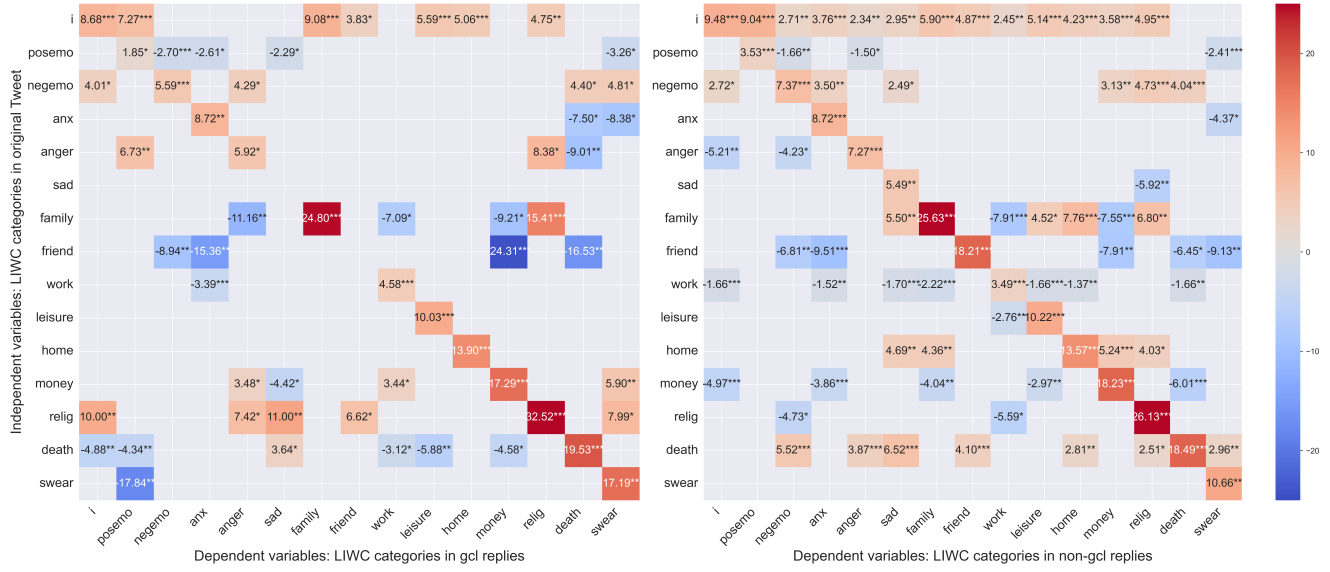


Figure 2: How *gcl* replies and *non-gcl* with specific LIWC categories correlate with LIWC features in the original Tweet. Each column is a logistic regression. Each cell is the coefficient. $p < 0.1$ is marked with *, $p < 0.05$ with **, and $p < 0.01$ with ***.

Variable	Coefficient	Odds ratio	95% CI
i	-0.46	0.63	(0.39, 1.03)
we	1.45	4.25***	(2.15, 8.40)
you	2.05	7.73***	(5.12, 11.68)
posemo	0.04	1.04	(0.78, 1.40)
negemo	-0.95	0.39**	(0.20, 0.76)
anx	-0.51	0.60	(0.16, 2.20)
anger	-1.43	0.24**	(0.08, 0.69)
sad	0.70	2.01	(0.73, 5.54)
family	-1.39	0.25	(0.06, 1.06)
friend	-1.33	0.26	(0.05, 1.30)
body	-1.98	0.14***	(0.05, 0.39)
health	-1.31	0.27***	(0.14, 0.52)
work	0.73	2.08***	(1.38, 3.13)
leisure	-0.15	0.86	(0.37, 1.99)
home	3.80	44.48***	(12.12, 163.25)
money	1.35	3.85***	(1.78, 8.31)
relig	-1.22	0.29*	(0.11, 0.82)
death	-1.77	0.17*	(0.06, 0.50)
swear	0.66	1.93	(0.74, 4.98)

Table 7: Logistic regression results on LIWC categories for *gcl* vs. *non-gcl* replies. Positive coefficients and odds ratios larger than 1 indicate the LIWC categories are more often in *gcl* replies. 95% CI is the confidence interval for the odds ratio. $p < 0.05$ is marked with *, $p < 0.01$ with **, and $p < 0.001$ with ***.

followers Our analysis begins with a key characteristic of Twitter users—the number of followers. We split the 346 PHEs with available state information into quartiles based on their follower counts. We then compare the distributions between *gcl* and *non-gcl* engagement rates of tweets whose authors are in each quartile group, using the KS test. Table 8

	# followers & quartiles	(0, 4105) < 25%	[4105, 15703] 25-50%	[15703, 67164] 50-75%	[67164, 10.4M] ≥ 75%
<i>gcl</i> ER		12.16%	3.25%	2.37%	1.11%
<i>non-gcl</i> ER		12.16%	2.91%	2.16%	1.06%
Diff		0.00%	0.34%*	0.21%*	0.05%*
# authors		25	62	79	90
# Tweets		222	2291	4007	9406

Table 8: Difference between *gcl* and *non-gcl* engagement rate for Tweets whose authors have different numbers of followers. * indicates that KS test results for the *gcl* ER and *non-gcl* ER comparison is significant at $p < .001$.

shows that the engagement rate difference between *gcl* and *non-gcl* groups is statistically significant for authors with a volume of followers larger than the 25% quartile *i.e.*, authors with a volume of followers that is above the lower quartile of the follower distribution have higher *gcl* engagement rates when compared to *non-gcl*. **This finding points to authors with larger volumes of followers and with geo-location information being able to engage more with *gcl* audiences when compared to *non-gcl* ones.**

Professions Given that the authors in the dataset are very active Twitter users and have rich information in their user descriptions, it is feasible to identify the professions of these authors. With the advancement of Large Language Models (LLMs) in understanding social context (Ziems et al. 2024), we use *gpt-4* to annotate these authors’ professions based on their descriptions and usernames. After preliminary analysis, we categorize the users into four profession types namely, media, academic, medical, and political-related. For each profession type, we include a few professions in the prompt. For example, journalists and reporters for media-related professions. We instruct the model to gen-

erate a binary prediction for each profession type, as well as the rationale for the prediction, to increase robustness and accuracy. We show the complete prompt, as well as sample data and prediction in Appendix F. A member of the research team manually reviewed the classification results for 50 authors and found around or above 90% accuracy for all four profession types.

As Table 9 shows, the differences between *gcl* and *non-gcl* engagement rates are all significant ($p < 0.001$) across the four profession types, with *gcl* engagement rates being always significantly higher than *non-gcl*. **Academic and medical-related professions attract the highest *gcl* engagement rates (2.09% and 1.94%, respectively), as well as the largest differences with respect to their *non-gcl* counterparts with changes in engagement of 0.2% and 0.21%.**

Professions	Media	Academic	Medical	Political
<i>gcl</i> ER	1.63%	2.09%	1.95%	1.72%
<i>non-gcl</i> ER	1.55%	1.89%	1.73%	1.65%
Diff	0.08%*	0.20%*	0.23%*	0.07%*
# authors	131	96	105	20
# Tweets	9107	5610	6734	1394
Test Acc.	88%	98%	94%	98%

Table 9: Differences between *gcl* and *non-gcl* engagement rates for Tweets whose authors have different professions. Test Acc is gpt4’s accuracy on a sample of 50 authors. * indicates that KS test results for the *gcl* ER and *non-gcl* ER comparison are significant at $p < .001$.

6 Limitations and Future Work

We enumerate several limitations and areas of future work. Our study is based on a dataset collected on Twitter, where users who use location services may be more active than those who do not (Rzeszewski and Beluch 2017). Additionally, Twitter users may provide inaccurate location information (Hecht et al. 2011) and the geo-location inference tool we adopt is imperfect, leading to potential inaccuracies in our analysis. Future work could benefit from using more advanced geo-inference tools, which leverage deep learning models (Zhang et al. 2023) or require richer user data including tweet content and users’ following network (Tian et al. 2020), to improve the accuracy and coverage of geo-location information. Moreover, it would also be interesting to investigate how our findings generalize to discussions of other topics beyond public health messaging and on other platforms such as Reddit.

Although we employed a carefully validated COVID-19 issue detection model (Rao et al. 2023) and sentiment classification model (Loureiro et al. 2022), as well as the widely adopted LIWC for linguistic analysis, the models are not perfect and LIWC partially captures the linguistic properties as it focuses on word-level features. Future work should consider these contexts when interpreting the conclusions from this study.

7 Conclusion

This study highlights the critical role of geographic co-location in influencing social media engagement. By analyzing a dataset of Twitter conversations during COVID-19, we found that users who are *gcl* with PHEs exhibit higher engagement rates, particularly in discussions about masking, lockdowns, and education. Conversely, vaccine-related discussions tend to attract higher engagement from *non-gcl* participants. Our findings indicate that emotional language and personal experiences significantly enhance engagement from *gcl* participants who reply more positively in sentiment. Particularly, *gcl* participants express more positive emotions and personal issues in response to tweets that share personal experiences or emotions. Furthermore, The impact of *gcl* is greater for PHEs in academic and medical fields compared to media and political professions. These findings highlight the need to consider geographic co-location in public health messaging to improve engagement and maintain a positive discourse around public health guidance in online communities.

References

- Bakshy, E.; Hofman, J. M.; Mason, W. A.; and Watts, D. J. 2011. Everyone’s an influencer: quantifying influence on twitter. In *Proceedings of the fourth ACM international conference on Web search and data mining*, 65–74.
- Bird, S.; Klein, E.; and Loper, E. 2009. *Natural language processing with Python: analyzing text with the natural language toolkit*. ” O’Reilly Media, Inc.”.
- Bojja, G. R.; Ofori, M.; Liu, J.; Ambati, L. S.; et al. 2020. Early Public Outlook on the Coronavirus Disease (COVID-19): A Social Media Study. *Amcis 2020 Proceedings*.
- Bozarth, L.; Quercia, D.; Capra, L.; and Šćepanović, S. 2023. The role of the big geographic sort in online news circulation among US Reddit users. *Scientific Reports*, 13(1): 6711.
- Cha, M.; Haddadi, H.; Benevenuto, F.; and Gummadi, K. 2010. Measuring user influence in twitter: The million follower fallacy. In *Proceedings of the international AAAI conference on web and social media*, volume 4, 10–17.
- Cheke, N.; Chandra, J.; Dandapat, S. K.; et al. 2020. Understanding the impact of geographical distance on online discussions. *IEEE Transactions on Computational Social Systems*, 7(4): 858–872.
- Chen, E.; Lerman, K.; Ferrara, E.; et al. 2020. Tracking social media discourse about the covid-19 pandemic: Development of a public coronavirus twitter data set. *JMIR public health and surveillance*, 6(2): e19273.
- Choi, M.; Budak, C.; Romero, D. M.; and Jurgens, D. 2021. More than meets the tie: Examining the role of interpersonal relationships in social networks. In *Proceedings of the International AAAI Conference on Web and Social Media*, volume 15, 105–116.
- Crandall, D. J.; Backstrom, L.; Cosley, D.; Suri, S.; Huttenlocher, D.; and Kleinberg, J. 2010. Inferring social ties from geographic coincidences. *Proceedings of the National Academy of Sciences*, 107(52): 22436–22441.

- Crooks, A.; Croitoru, A.; Stefanidis, A.; and Radzikowski, J. 2013. # Earthquake: Twitter as a distributed sensor system. *Transactions in GIS*, 17(1): 124–147.
- Cuevas, R.; Gonzalez, R.; Cuevas, A.; and Guerrero, C. 2014. Understanding the locality effect in Twitter: measurement and analysis. *Personal and Ubiquitous Computing*, 18: 397–411.
- De Choudhury, M.; and De, S. 2014. Mental health discourse on reddit: Self-disclosure, social support, and anonymity. In *Proceedings of the international AAAI conference on web and social media*, volume 8, 71–80.
- Dredze, M.; Paul, M. J.; Bergsma, S.; and Tran, H. 2013. Carmen: A twitter geolocation system with applications to public health. In *Workshops at the twenty-seventh AAAI conference on artificial intelligence*.
- Eisenstein, J.; Ahmed, A.; and Xing, E. P. 2011. Sparse additive generative models of text. In *Proceedings of the 28th international conference on machine learning (ICML-11)*, 1041–1048.
- Fischer, S. 2020. Local news availability does not increase pro-social pandemic response.
- Frias-Martinez, V.; Soto, V.; Hohwald, H.; and Frias-Martinez, E. 2012a. Characterizing urban landscapes using geolocated tweets. In *2012 International conference on privacy, security, risk and trust and 2012 international conference on social computing*, 239–248. IEEE.
- Frias-Martinez, V.; Soto, V.; Virseda, J.; and Frias-Martinez, E. 2012b. Computing cost-effective census maps from cell phone traces. In *Workshop on pervasive urban applications*.
- Gallagher, R. J.; Doroshenko, L.; Shugars, S.; Lazer, D.; and Foucault Welles, B. 2021. Sustained online amplification of COVID-19 elites in the United States. *Social Media+ Society*, 7(2): 205630512111024957.
- Ghahremanlou, L.; Sherchan, W.; and Thom, J. A. 2015. Geotagging twitter messages in crisis management. *The Computer Journal*, 58(9): 1937–1954.
- Ghosh, R.; and Lerman, K. 2010. Predicting influential users in online social networks. *arXiv preprint arXiv:1005.4882*.
- Hecht, B.; Hong, L.; Suh, B.; and Chi, E. H. 2011. Tweets from Justin Bieber's heart: the dynamics of the location field in user profiles. In *Proceedings of the SIGCHI conference on human factors in computing systems*, 237–246.
- Hodges Jr, J. 1958. The significance probability of the Smirnov two-sample test. *Arkiv för matematik*, 3(5): 469–486.
- Hong, L.; Frias-Martinez, E.; and Frias-Martinez, V. 2016. Topic models to infer socio-economic maps. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 30.
- Hong, L.; and Frias-Martinez, V. 2020. Modeling and predicting evacuation flows during hurricane Irma. *EPJ Data Science*, 9(1): 29.
- Hu, Y.; Farnham, S.; and Talamadupula, K. 2015. Predicting user engagement on twitter with real-world events. In *Proceedings of the International AAAI Conference on Web and Social Media*, volume 9, 168–177.
- Izbicki, M.; Papalexakis, V.; and Tsotras, V. 2019. Geolocating Tweets in any Language at any Location. In *Proceedings of the 28th ACM International Conference on Information and Knowledge Management*, 89–98.
- Jiang, S.; and Wilson, C. 2018. Linguistic signals under misinformation and fact-checking: Evidence from user comments on social media. *Proceedings of the ACM on Human-Computer Interaction*, 2(CSCW): 1–23.
- Jordan, S. E.; Hovet, S. E.; Fung, I. C.-H.; Liang, H.; Fu, K.-W.; and Tse, Z. T. H. 2018. Using Twitter for public health surveillance from monitoring and prediction to public response. *Data*, 4(1): 6.
- Jurgens, D.; Finethy, T.; McCorriston, J.; Xu, Y.; and Ruths, D. 2015. Geolocation prediction in twitter using social networks: A critical analysis and review of current practice. In *Proceedings of the international AAAI conference on web and social media*, volume 9, 188–197.
- Karami, A.; Kadari, R. R.; Panati, L.; Nooli, S. P.; Bheemreddy, H.; and Bozorgi, P. 2021. Analysis of geotagging behavior: Do geotagged users represent the Twitter population? *ISPRS International Journal of Geo-Information*, 10(6): 373.
- Kulshrestha, J.; Kooti, F.; Nikraves, A.; and Gummadi, K. 2012. Geographic dissection of the Twitter network. In *Proceedings of the International AAAI Conference on Web and Social Media*, volume 6, 202–209.
- Kwak, H.; Lee, C.; Park, H.; and Moon, S. 2010. What is Twitter, a social network or a news media? In *Proceedings of the 19th international conference on World wide web*, 591–600.
- Laniado, D.; Volkovich, Y.; Scellato, S.; Mascolo, C.; and Kaltenbrunner, A. 2018. The impact of geographic distance on online social interactions. *Information Systems Frontiers*, 20(6): 1203–1218.
- Leskovec, J.; and Horvitz, E. 2008. Planetary-scale views on a large instant-messaging network. In *Proceedings of the 17th international conference on World Wide Web*, 915–924.
- Li, Z.; Mao, A.; Stephens, D.; Goel, P.; Walpole, E.; Dima, A.; Fung, J.; and Boyd-Graber, J. 2024. Improving the TENOR of Labeling: Re-evaluating Topic Models for Content Analysis. In *Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics (Volume 1: Long Papers)*, 840–859.
- Liben-Nowell, D.; Novak, J.; Kumar, R.; Raghavan, P.; and Tomkins, A. 2005. Geographic routing in social networks. *Proceedings of the National Academy of Sciences*, 102(33): 11623–11628.
- Lindstrom, M. J.; and Bates, D. M. 1988. Newton—Raphson and EM algorithms for linear mixed-effects models for repeated-measures data. *Journal of the American Statistical Association*, 83(404): 1014–1022.
- Loureiro, D.; Barbieri, F.; Neves, L.; Espinosa Anke, L.; and Camacho-collados, J. 2022. TimeLMs: Diachronic Language Models from Twitter. In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics: System Demonstrations*, 251–260. Dublin, Ireland: Association for Computational Linguistics.

- Lwin, M. O.; Lu, J.; Sheldenkar, A.; Schulz, P. J.; Shin, W.; Gupta, R.; and Yang, Y. 2020. Global sentiments surrounding the COVID-19 pandemic on Twitter: analysis of Twitter trends. *JMIR public health and surveillance*, 6(2): e19447.
- Mahmud, J.; Nichols, J.; and Drews, C. 2012. Where is this tweet from? inferring home locations of twitter users. In *Proceedings of the International AAAI Conference on Web and Social Media*, volume 6, 511–514.
- McGee, J.; Caverlee, J.; and Cheng, Z. 2013. Location prediction in social media based on tie strength. In *Proceedings of the 22nd ACM international conference on Information & Knowledge Management*, 459–468.
- Miles, A.; Andiappan, M.; Upenieks, L.; and Orfanidis, C. 2022. Using prosocial behavior to safeguard mental health and foster emotional well-being during the COVID-19 pandemic: A registered report of a randomized trial. *PloS one*, 17(7): e0272152.
- Milusheva, S.; Marty, R.; Bedoya, G.; Williams, S.; Ressor, E.; and Legovini, A. 2021. Applying machine learning and geolocation techniques to social media data (Twitter) to develop a resource for urban planning. *PloS one*, 16(2): e0244317.
- Monroe, B. L.; Colaresi, M. P.; and Quinn, K. M. 2008. Fightin' words: Lexical feature selection and evaluation for identifying the content of political conflict. *Political Analysis*, 16(4): 372–403.
- Newman, M.; Barabási, A.-L.; and Watts, D. J. 2011. *The structure and dynamics of networks*. Princeton university press.
- Pavalanathan, U.; and Eisenstein, J. 2015. Confounds and Consequences in Geotagged Twitter Data. In *Proceedings of the 2015 Conference on Empirical Methods in Natural Language Processing*, 2138–2148.
- Pennebaker, J. W.; Boyd, R. L.; Jordan, K.; and Blackburn, K. 2015. The development and psychometric properties of LIWC2015.
- Rao, A.; Guo, S.; Wang, S. Y. N.; Morstatter, F.; and Lerman, K. 2023. Pandemic Culture Wars: Partisan Differences in the Moral Language of COVID-19 Discussions. In *2023 IEEE International Conference on Big Data (BigData)*, 413–422. IEEE.
- Rao, A.; Sabri, N.; Guo, S.; Raschid, L.; and Lerman, K. 2024. #EPITWITTER: Public Health Messaging during the COVID-19 Pandemic. *Under review*.
- Rosenthal, S.; Farra, N.; and Nakov, P. 2017. SemEval-2017 task 4: Sentiment analysis in Twitter. In *Proceedings of the 11th international workshop on semantic evaluation (SemEval-2017)*, 502–518.
- Rzeszewski, M.; and Beluch, L. 2017. Spatial characteristics of twitter users—Toward the understanding of geosocial media production. *ISPRS International Journal of Geo-Information*, 6(8): 236.
- Sadilek, A.; Kautz, H.; and Bigham, J. P. 2012. Finding your friends and following them to where you are. In *Proceedings of the fifth ACM international conference on Web search and data mining*, 723–732.
- Scellato, S.; Noulas, A.; Lambiotte, R.; and Mascolo, C. 2011. Socio-spatial properties of online location-based social networks. In *Proceedings of the international AAAI conference on web and social media*, volume 5, 329–336.
- Schwartz, H. A.; Eichstaedt, J. C.; Kern, M. L.; Dziurzynski, L.; Ramones, S. M.; Agrawal, M.; Shah, A.; Kosinski, M.; Stillwell, D.; Seligman, M. E.; et al. 2013. Personality, gender, and age in the language of social media: The open-vocabulary approach. *PloS one*, 8(9): e73791.
- Sigalo, N.; Awasthi, N.; Abrar, S. M.; Frias-Martinez, V.; et al. 2023. Using COVID-19 vaccine attitudes on Twitter to improve vaccine uptake forecast models in the United States: infodemiology study of Tweets. *JMIR infodemiology*, 3(1): e43703.
- Sigalo, N.; Frias-Martinez, V.; et al. 2023. Using COVID-19 Vaccine Attitudes Found in Tweets to Predict Vaccine Perceptions in Traditional Surveys: Infodemiology Study. *JMIR infodemiology*, 3(1): e43700.
- Sloan, L.; and Morgan, J. 2015. Who tweets with their location? Understanding the relationship between demographic characteristics and the use of geoservices and geotagging on Twitter. *PloS one*, 10(11): e0142209.
- Tian, H.; Zhang, M.; Luo, X.; Liu, F.; and Qiao, Y. 2020. Twitter user location inference based on representation learning and label propagation. In *Proceedings of The Web Conference 2020*, 2648–2654.
- Wang, Y.-C.; Burke, M.; and Kraut, R. 2016. Modeling self-disclosure in social networking sites. In *Proceedings of the 19th ACM conference on computer-supported cooperative work & social computing*, 74–85.
- Wheaton, M. G.; Prikhidko, A.; and Messner, G. R. 2021. Is fear of COVID-19 contagious? The effects of emotion contagion and social media use on anxiety in response to the coronavirus pandemic. *Frontiers in psychology*, 11: 567379.
- Wojcik, S.; and Hughes, A. 2019. Sizing up Twitter users. <https://www.pewresearch.org/internet/2019/04/24/sizing-up-twitter-users/>. [Online; accessed 9-January-2024].
- Wood-Doughty, Z.; Xu, P.; Liu, X.; and Dredze, M. 2021. Using Noisy Self-Reports to Predict Twitter User Demographics. In *Proceedings of the Ninth International Workshop on Natural Language Processing for Social Media*, 123–137.
- Xu, P.; Broniatowski, D. A.; and Dredze, M. 2024. Twitter social mobility data reveal demographic variations in social distancing practices during the COVID-19 pandemic. *Scientific reports*, 14(1): 1165.
- Xu, P.; Dredze, M.; and Broniatowski, D. A. 2020. The twitter social mobility index: Measuring social distancing practices with geolocated tweets. *Journal of medical Internet research*, 22(12): e21499.
- Zhang, J.; DeLucia, A.; and Dredze, M. 2022. Changes in tweet geolocation over time: A study with carmen 2.0. In *Proceedings of the Eighth Workshop on Noisy User-generated Text (W-NUT 2022)*, 1–14.

- Zhang, J.; DeLucia, A.; Zhang, C.; and Dredze, M. 2023. Geo-seq2seq: Twitter user geolocation on noisy data through sequence to sequence learning. In *Findings of the Association for Computational Linguistics: ACL 2023*, 4778–4794.
- Zhou, Y.; Xu, P.; Wang, X.; Lu, X.; Gao, G.; and Ai, W. 2024. Emojis decoded: Leveraging chatgpt for enhanced understanding in social media communications. *arXiv preprint arXiv:2402.01681*.
- Ziems, C.; Held, W.; Shaikh, O.; Chen, J.; Zhang, Z.; and Yang, D. 2024. Can large language models transform computational social science? *Computational Linguistics*, 1–55.

A Seed PHEs

The usernames for the 30 seed PHEs on Twitter are: EricTopol, PeterHotez, ashishkjha, trvr, EpiEllie, JuliaRaifman, devisridhar, meganranney, luckytran, asosin, DrLeanaWen, dremilyporterm, DrJaime-Friedman, davidwdowdy, BhramarBioStat, geochurch, DrEricDing, michaelmina_lab, Bob_Wachter, Jennifer-Nuzzo, mtosterholm, MonicaGandhi9, cmyeaton, nataliexdean, angie_rasmussen, ProfEmilyOster, mlipsitch, drlucymcbride, ScottGottliebMD, CDCDirector, and Surgeon_Genera

B COVID-19 Issues

The lockdown issue comprises content pertaining to early state and federal mitigation efforts, such as quarantines, stay-at-home orders, business closures, reopening, and calls for social distancing. The masking issue is defined by discussions on the use of face coverings, mask mandates, mask shortages, and anti-mask sentiment. Education-related content involves tweets about school closures, reopening of educational institutions, homeschooling, and online learning during the pandemic. The vaccine issue pertains to discussions about COVID-19 vaccines, vaccine mandates, anti-vaccine sentiment, and vaccine hesitancy in the U.S.

Issue	Sample Tweets
Lockdowns	This is a GREAT idea. We're all in this together. Take care of each other. #Stay-Home #TakeItSeriously #FlattenTheCurve #COVID19
Masking	We're in the middle of a pandemic and y'all are still coughing and sneezing without covering your mouths ? Come on now.
Education	More glimmers of hope as we "safely" move forward and open up Texas AM University while containing #COVID19 .
Vaccines	You are joking right? Zero sympathy for anti-vaxxers who quit their jobs rather than get vaccinated . They put us all at risk and make the pandemic prolonged for the world.

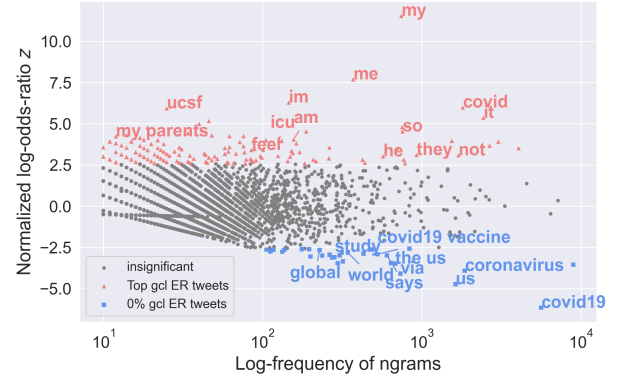
Table 10: Sample Tweets for each COVID-19 issue. Issue-relevant keywords are in red. Table from Rao et al. (2023).

C Lexical Analysis for Geo-co-located Tweets

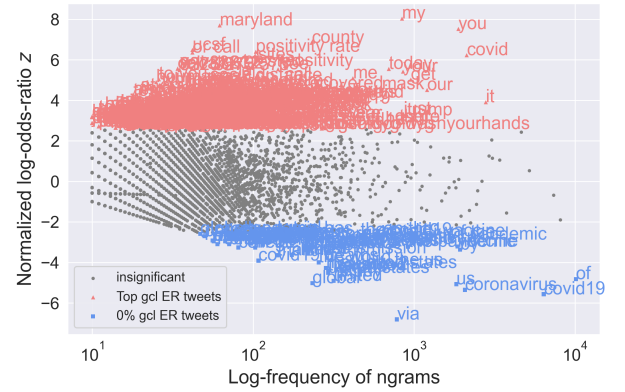
To gain insights into what lexical features are most distinctive for *gcl* tweets, we compare Tweets that have 0% *gcl* engagement rate ($N = 10,854$, bottom 68% of the tweets) with the top 10% percentile of tweets (3.4% *gcl* engagement rate with $N = 1,576$). We also compare 0% *gcl* ER tweets with the top 20% percentile of tweets (1.0% *gcl* engagement rate with $N = 3,190$).

We follow Monroe, Colaresi, and Quinn (2008) and use normalized log-odds-ratio z to find out the n -grams ($n \in \{1, 2, 3\}$) more associated with Tweets with higher and lower *gcl* engagement rates. We find that the top group

shares more personal experiences/feelings ("my", "me"), while the bottom group talks about the national or global situation of COVID-19 with keywords such as "covid19", "us", "united states", "study", "the world" (Figure 3).



(a) Top 10% percentile of tweets



(b) Top 20% percentile of tweets

Figure 3: Most distinctive n -grams ($n \in \{1, 2, 3\}$): tweets with 0% *gcl* engagement rate (bottom 68% of the tweets) vs. (a) the top 10% percentile of tweets (3.4% *gcl* engagement rate) and (b) the top 20% percentile of tweets (1.0% *gcl* engagement rate). Only n -grams with 99.5% confidence level are colored ($z > 2.57$)

D Full Regression Results for Section 5.2

The complete results of the LIWC regression analysis in Section 5.2 is shown in Table 11.

E Full Heatmap for Section 5.3

The heatmaps with all LIWC categories from Section 5.3 are shown in Figure 4.

F Prompt for Profession Classification

The prompt we used for profession classification in Section 5.4 is as follows:

For the following Twitter user profile description and username, infer whether the user's profession is

	Diff	<i>gcl</i>	<i>non-gcl</i>
const	0.002 (0.002)	0.023*** (0.002)	0.021*** (0.001)
i	0.102*** (0.038)	0.420*** (0.037)	0.318*** (0.028)
we	-0.012 (0.030)	-0.005 (0.029)	0.007 (0.022)
you	-0.023 (0.039)	-0.008 (0.038)	0.014 (0.028)
posemo	0.038 (0.024)	0.087*** (0.023)	0.049*** (0.017)
negemo	-0.070 (0.046)	0.012 (0.044)	0.082** (0.034)
anx	0.128* (0.069)	-0.052 (0.067)	-0.179*** (0.050)
anger	0.105 (0.073)	0.028 (0.071)	-0.077 (0.053)
sad	0.057 (0.074)	-0.017 (0.072)	-0.074 (0.054)
family	0.104 (0.086)	0.236*** (0.083)	0.132** (0.063)
friend	-0.021 (0.093)	0.051 (0.090)	0.072 (0.068)
body	-0.065 (0.055)	-0.061 (0.053)	0.004 (0.040)
health	-0.017 (0.026)	-0.017 (0.025)	-0.000 (0.019)
work	0.007 (0.018)	-0.011 (0.018)	-0.018 (0.013)
leisure	-0.059 (0.039)	-0.034 (0.038)	0.025 (0.028)
home	0.019 (0.060)	-0.010 (0.058)	-0.029 (0.043)
money	-0.035 (0.042)	-0.055 (0.040)	-0.020 (0.030)
relig	-0.009 (0.081)	-0.079 (0.079)	-0.070 (0.059)
death	0.020 (0.041)	-0.041 (0.040)	-0.061** (0.030)
swear	-0.041 (0.157)	-0.056 (0.152)	-0.015 (0.114)
# followers	-0.001 (0.008)	0.011 (0.008)	0.012* (0.006)
gcl audience	0.000 (0.004)	-0.016*** (0.003)	-0.016*** (0.003)
non-gcl audience	-0.002 (0.003)	-0.025*** (0.003)	-0.023*** (0.002)
masking.True	0.004 (0.003)	0.004 (0.003)	0.000 (0.002)
education.True	-0.001 (0.003)	-0.001 (0.003)	-0.000 (0.002)
lockdowns.True	-0.003 (0.005)	-0.003 (0.004)	-0.000 (0.003)
vaccines.True	-0.004** (0.002)	-0.005*** (0.002)	-0.001 (0.001)
R-squared	0.002	0.024	0.031
R-squared Adj.	0.000	0.023	0.029

Table 11: LIWC regression results of LIWC categories' impact on different types of engagement rates. Standard errors in parentheses. * $p < .1$, ** $p < .05$, *** $p < .01$

related to each of the professions as follows: media professions (e.g., journalist, reporter, writer, editor, and author), academic (e.g., professor, researcher, and Scientist), medical professions (e.g., physician, doctor, epidemiologist), and politicians. For each Twitter user, output the results for each of the four professions and provide a brief explanation in the json format of '{"screen_name": USERNAME, "MediaProfessions": true/false, "Academic": true/false, "MedicalProfessions": true/false, "Politicians": true/false, "Explanation": "[Explanation]"}'. Only output true for the best-matched professionals if possible. All outputs should be in the json format as described above.

We then concatenate it with usernames and user descriptions of five authors for each query in the following format:

```
username: {screen_name}
description: {description}
```

Note that {screen_name} and {description} are replaced with actual usernames and descriptions. A sample output from ChatGPT with an anonymized user name is:

```
{"screen_name": "ANONYMIZED",
"MediaProfessions": false, "Academic":
true, "MedicalProfessions": true,
"Politicians": false, "Explanation":
"The user is a Clinical Psychology PhD
Student and does research in Global
Mental Health, which indicates that
he is involved in academic and medical
professions."}
```

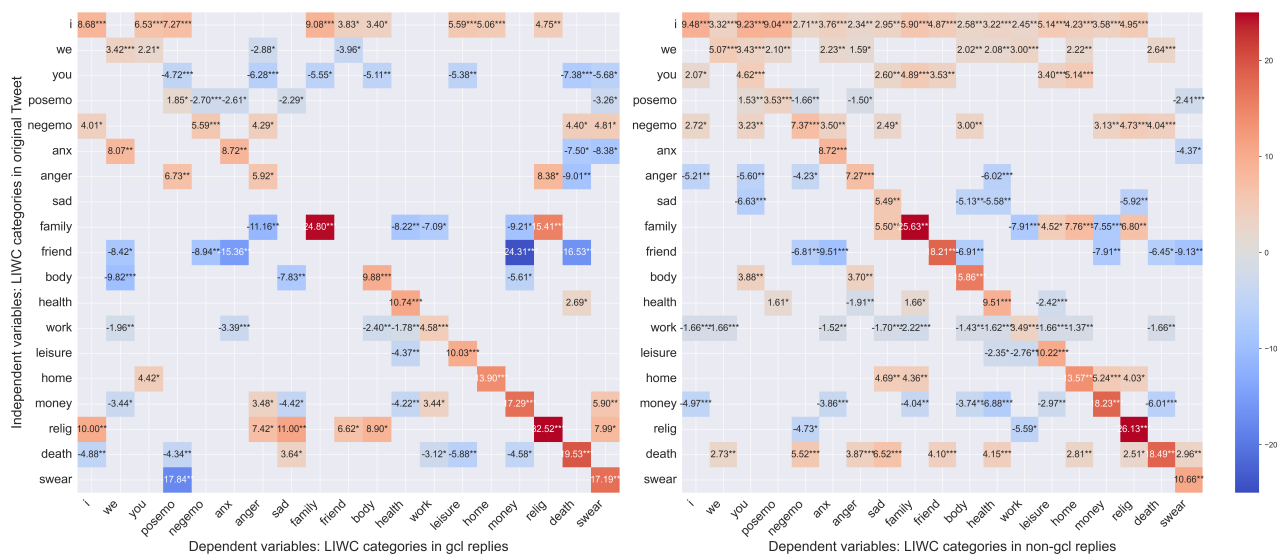


Figure 4: How *gcl* replies and *non-gcl* with specific LIWC categories correlate with LIWC features in the original Tweet. Each column is a logistic regression. Each cell is the coefficient. $p < 0.1$ is marked with *, $p < 0.05$ with **, and $p < 0.01$ with ***. Only selected LIWC categories are shown for demonstration purposes.