

Improving Association Discovery through Multiview Analysis of Social Networks

Muhieddine Shebaro

Texas State University

Lia Nogueira de Moura

Texas State University

Jelena Tešić

jtesic@txstate.edu

Texas State University

Research Article

Keywords: multi-modal mining, fake news, COVID, social network analysis, network science

Posted Date: July 24th, 2023

DOI: <https://doi.org/10.21203/rs.3.rs-3179561/v1>

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

[Read Full License](#)

Additional Declarations: No competing interests reported.

Version of Record: A version of this preprint was published at Social Network Analysis and Mining on February 20th, 2024. See the published version at <https://doi.org/10.1007/s13278-023-01197-3>.

Improving Association Discovery through Multiview Analysis of Social Networks

Muhieddine Shebaro^{1†}, Lia Nogueira de Moura MSc^{1†},
Jelena Tešić PhD^{1*}

^{1*}Computer Science, Texas State University, San Marcos, 78666, TX,
USA.

*Corresponding author(s). E-mail(s): jtesic@txstate.edu;
Contributing authors: m.shebaro@txstate.edu; lia.lnm@gmail.com;

[†]These authors contributed equally to this work.

Abstract

The rise of social networks has brought about a transformative impact on communication and the dissemination of information. However, this paradigm shift has also introduced many challenges in discerning valuable conversation threads amidst fake news, malicious accounts, background noise, and trolling. In this study, we address these challenges by focusing on propagating fake news labels. We evaluate the efficacy of community-based modeling in effectively addressing these challenges within the context of social network discussions using the state-of-art benchmark. Through a comprehensive analysis of millions of users engaged in discussions on a specific topic, we unveil compelling evidence demonstrating that community-based modeling techniques yield precision, recall, and accuracy levels comparable to those achieved by lexical classifiers. Remarkably, these promising results are achieved even without considering the textual content of *tweets* beyond the information conveyed by hashtags. Moreover, we explore the effectiveness of fusion techniques in tweet classification and underscore the superiority of a combined community and lexical approach, which consistently delivers the most robust outcomes and exhibits the highest performance measures. We illustrate this capability with specific network graphs constructed based on Twitter interactions related to the COVID-19 pandemic, showcasing the practicality and relevance of our proposed methodology.

Keywords: multi-modal mining, fake news, COVID, social network analysis, network science

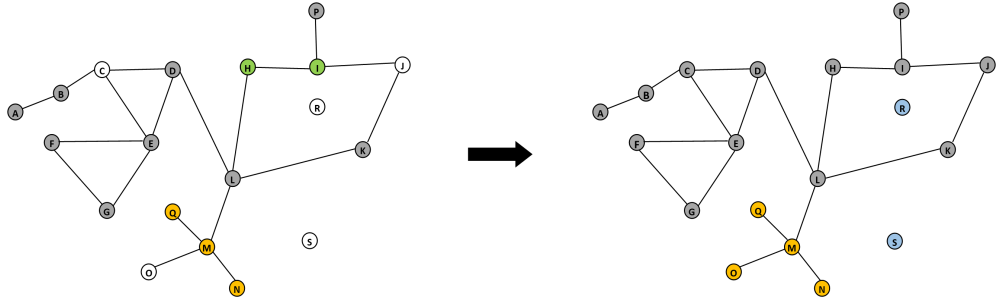


Fig. 1 Community attribute enrichment: analyze labeled data set in a network graph and extract community labels from the graph analysis of the network. Gray nodes are nodes with non-conspiracy content. Yellow nodes are promoting/discussing 5G conspiracy topics. White nodes are test nodes. Light blue are unknown nodes (indeterminate). Green nodes discuss other conspiracies.

1 Introduction

The advent of social networks has conferred significant importance on these platforms as principal channels for news consumption among a significant segment of the populace. The interconnectivity of online users within these networks facilitates the swift propagation of information, surpassing the conventional scope of traditional news media outlets like newspapers and television. Nevertheless, this inherent interconnectivity also amplifies the ease with which inaccurate and deceptive information can increase, particularly within the context of users’ social network connections. This study seeks to examine the potential utility of the structural characteristics of social network user connections in identifying and addressing false information, specifically within the domain of Twitter.

Can we classify the Tweet without knowing the content tweet? In this paper, we explore the social network context, Twitter’s rich network of interaction, i.e., connections, tags, retweets, and mentions, and how they influence the labeling of the content. We test the observation that people in the same social network group or discussion thread tend to quote and discuss similar resources and have shared topic items, shed new light on the challenges posed by social network dynamics, and offer an effective means of tackling them through community-based modeling. By revealing the comparable performance of community-based approaches to traditional lexical classifiers, we contribute to advancing tweet classification methodologies. Our research opens up exciting avenues for further exploration and application, paving the way for more sophisticated network selection and fusion methods that leverage both community attributes and lexical modeling to enhance the accuracy and effectiveness of tweet classification in the ever-evolving landscape of social networks. Our findings carry substantial implications for understanding the dynamics of social networks and advancing methodologies for tweet classification. By harnessing the power of community attributes and models, our research uncovers the invaluable contextual information embedded within social network interactions involving tweet authors and

objects. Furthermore, we present tangible evidence of our ability to capture comprehensive information by constructing network graphs that encapsulate crucial features such as retweets, mentions, replies, and quote networks.

To this end, we propose an enrichment of Tweet classification with a network-based analysis of the Twitter network, as illustrated in Figure 1. We relate the content of the *tweets* using multi-modal lexical analysis, employ community discovery by building a network of retweets, mentions, and hashtags, and employ network analysis on structural data mined from Twitter. To this end, if available, we have developed a robust lexical-based analysis for Tweet content that considers colloquialisms, abbreviations, and OCR text in images. We have also developed a scalable data science package that downloads, saves, and analyzes Twitter data at scale, providing robust content analysis of noisy communities on Twitter [1–3]. We evaluate the approach in the MediaEval 2020 Fake News task benchmark and COVID-19 (+) Twitter data set. We demonstrate the value of the author’s network in content classification on the MediaEval Fake News Detection Task 2020, which offers two Fake News Detection sub-tasks on COVID-19 and 5G conspiracy topics [4]. More specifically, they detect misinformation claims that the construction of the 5G network and the associated electromagnetic radiation triggered the SARS-CoV-2 virus. This benchmark challenge looked only at Tweet classification of COVID-19-related *tweets* in two ways: (1) multi-class labeling: 5G-Corona-Conspiracy, Other-Conspiracy, and Non-Conspiracy, and (2) binary labeling: Unknown-or-Non-Conspiracy and Any-Conspiracy. This paper shows that tweet classification on the author’s network only (without analyzing tweet content) performs similarly to tweet content classification.

Table 1 Tweet by a user with strong 5G Corona Conspiracy community ties. Community-based detection identified the group and augmented the lexical classification.

Content: <i>Does #5G cause #COVID2019 #coronavirus? No, of course not! Does non-ionizing #wireless radiation accelerate viral replication and contribute to #AntibioticResistance? Yes.</i>
Ground Truth: 5g_corona_conspiracy
Lexical model Prediction: non_conspiracy
Reply connection network majority prediction: 5g_corona_conspiracy
of edges in labeled 5g_corona_conspiracy set: 11
of edges in the other_conspiracy dataset: 0
of edges in the non_conspiracy conspiracy dataset: 0
% of <i>tweets</i> in the detected community that are from 5g_corona_conspiracy dataset: 100%
% of <i>tweets</i> in the detected community that are from other_conspiracy dataset 0%
% of <i>tweets</i> in the detected community that are from non_conspiracy dataset 0%

2 Related Work

This section reviews the related work on fake news detection on Twitter. The prevalence of “fake news” raises significant concerns. Recent research shows that fake news sharing is fueled by the same psychological motivations that drive other forms of partisan behavior, including sharing partisan news from traditional and credible news

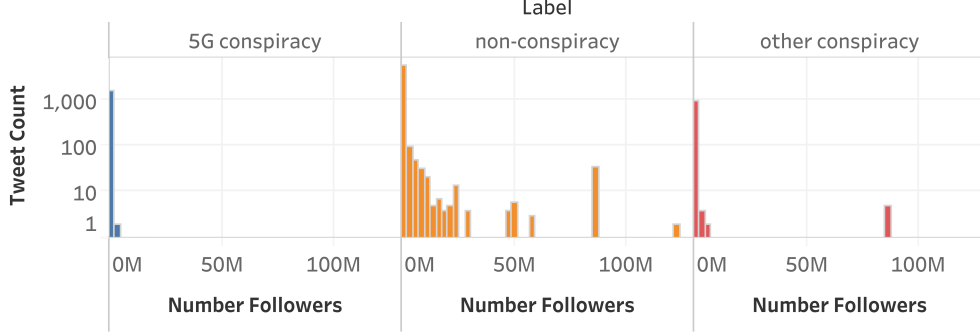


Fig. 2 Distribution of the feature *user_followers_count* for the different class labels (5G, non, other).

sources [5]. Given the widespread proliferation of misinformation online and the growing reliance on social media for news consumption, it is essential to comprehend how people evaluate and engage with posts of low credibility. This study examines users’ responses to fake news posts on their Facebook or Twitter feeds, seemingly originating from accounts they follow. To explore this phenomenon, we conducted semi-structured interviews with 25 participants who regularly employ social media for news consumption. Employing a browser extension unbeknownst to the participants, we temporarily introduced fake news into their feeds and observed their subsequent interactions. Participants provided insights into their browsing experiences and decision-making processes through this process. Our findings highlight various reasons individuals refrain from investigating posts of low credibility, including a tendency to accept content from trusted sources at face value and a reluctance to invest additional time in verification. Moreover, we outline the investigative techniques employed by participants to verify the trustworthiness of posts, encompassing both the functionalities provided by the platform and impromptu strategies. Building upon our empirical findings, we propose design guidelines to assist users in assessing the credibility of posts with low levels of credibility [6]. Twitter data has been used to understand the influence of fake news during the 2016 US presidential election [7]; it has also been used to analyze the COVID-19 and the 5G Conspiracy Theory [8] and the COVID-19 Twitter narrative among U.S. governors and cabinet executives [9]. Using logistic regression to classify Tweets based on topic [10] shows that the content of the Tweet dominates in correct Tweet classification. Writing style and frequency of word usage emerged as relevant features in the lexical analysis [11]. Two primary directions of leveraging community information are adapting deep learning techniques to learn the underlying characteristics of the Tweets in communities (e.g., [12]) or exploring the structural and sharing patterns of the topic (e.g., [13]).

2.1 Context Through Connections

Community-based modeling of social networks that leverages the spread of information in social media through retweets and comments has improved NLP-based modeling

[11]. Structural and sharing patterns in the Twitter verse are rich, and the definition of communities on Twitter is multi-dimensional. Users in the community can share geographic proximity and interconnections with mutual friends, groups, and topics of interest. Careful mapping of psychological profiles of over 2,300 American Twitter users linked to behavioral sharing data and sentiment analyses of more than 500,000 news story headlines finds that the individuals who report hating their political opponents are the most likely to share political fake news and selectively share content that is useful for derogating these opponents [5]. Factual News Graph (FANG) was proposed as a graphical social context representation and learning framework for fake news detection focusing on representation learning. FANG has captured. Social context to a degree if the topic is well represented and has generalized to related tasks, such as predicting the factuality of reporting of a news medium [14]. Similar unsupervised graph embedding methods on the graphs from the Twitter users' social network connections are used to find that the users engaged with fake news are more tightly clustered than users only engaged in factual news [15]. Graph-based approaches focus on bi-clique identification, graph-based feature vector learning, and label spreading on Twitter [16]. Still, they do not scale well to the number and heterogeneity of the topics examined. Schroeder and al. developed a framework for capturing and analyzing vast amounts of Twitter data. It consists of the primary data capturing component (Twitter API), the proxy, the storage, and experiment wrappers, which are connected to the storage and the proxy. The proxy provides quota leasing, an external API allowing users to execute calls with the same syntax, and request caching. The storage supports diverse types of databases and file storage. And experiment wrappers constitute a setup for analytical tasks and collecting data. Example experiments include follower analysis of an account for fake followers detection and network analysis of a user for determining the position of an account concerning its surrounding network [17].

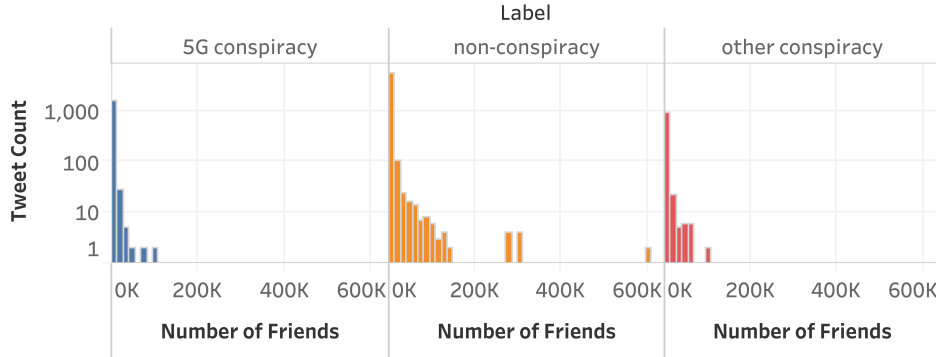


Fig. 3 Distribution of the feature *user_friends_count* for the different class labels (5G, non, other).

2.2 Lexical Aspects

Beyond using lexical and community features, other novel avenues of harnessing a tweet have been explored in various tasks. The prediction of lexical aspects of tweets with the #MeToo hashtag, a movement that has recently emerged against sexual assault and advocating women’s rights, has been performed by capitalizing on both textual and visual modalities. Still, contextual embeddings and transformer language models weren’t employed because they were computationally expensive [18]. Many similar works dealing with these same type of modalities, however, has put the pre-served version of BERT and a generic Deep Neural Network (DNN) to use for feature extraction, significantly boosting generalization performance such as [19] for ultimately developing a profiling system to identify anonymous and potentially nefarious users’ genders and [20] for finding disaster tweets. The concept of multi-modal tweet fusion is even introduced in the context of geosciences [21], where the authors proposed incorporating contextual hydrology information to classify flood-related tweets effectively. This proposal has yielded promising success along several metrics. It sheds light on the importance of not restricting models, regardless of their type, in feeding on lexical data and not neglecting other discriminate information. Another pivotal and salient modality is the location features of geo-tagged (longitude and latitude) tweets for sentiment analysis [22]. The tweets’ word embeddings were obtained and merged with the vectorized location features to create a set of hybrid representations. These representations enhance accuracy in classifying sentiment compared to the baseline GloVe model using a convolutional neural network (CNN) and a bi-directional long short-term memory recurrent neural network (LSTM). Graph Neural networks per-

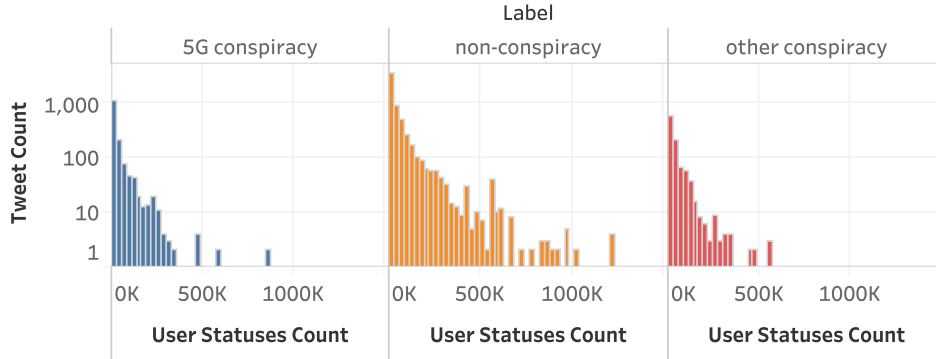


Fig. 4 Distribution of the feature *user_statuses_count* for the different class labels (5G, non, other).

form well in multi-modal contexts. For instance, Gao et al. presented MM-GNN, a novel framework that addresses inquiries by providing information from images. MM-GNN incorporates visual, semantic, and numeric modalities to represent an image

as a graph. Next, the node features are refined by leveraging contextual information from these modalities (using message passing), which improves performance in question-answering tasks [23].

There is also a multitude of state-of-the-art Graph neural network (GNN) variants that have been developed to resolve current issues of vanilla baseline GNNs. For instance, SelfSAGCN was created to alleviate over-smoothing and when labeled data are severely scarce using "Identity Aggregation" and "Semantic Alignment" techniques [24]. Due to the limited memory resources when loading the entire attributed graph into the network for processing in current GNNs, Bi-GCN was designed in which it binarizes both the network parameters and input node features and produces comparable results as baseline vanilla models such as GraphSage and GCN [25]. In addition, sparsely and noisily labeled graphs have been dealt with via the novel NRGNN variant [26]. Another notable GNN framework, Tail-GNN, is based on the concept of neighborhood translation in the structurally rich head nodes to be transferred to the structurally limited tail nodes to enhance their representations and uncover missing neighborhood nodes [27].

Unsupervised graph clustering approaches have considered merely structural information, and in recent years attributed graph clustering has gained substantial attention: it integrates additional attribute data about vertices into the clustering task to enhance its result. State-of-the-art graph neural networks suffer from training data bias and vertex feature dependency [28].

Social media platforms have become a vital source of information during the outbreak of the pandemic (COVID-19). The phenomenon of fake information or news spread through social media has become increasingly prevalent and a powerful tool for information proliferation. Detecting fake news is crucial for the betterment of society. Existing fake news detection models focus on increasing the performance, which leads to overfitting and lag generalizability. Hence, these models require training for various datasets of the same domain with significant variations in the distribution. In our work, we have addressed this overfitting issue by designing a robust distribution generalization of transformers-based generative adversarial network (RDGT-GAN) architecture, which can generalize the model for COVID-19 fake news datasets with different distributions without retraining. Based on our experimental findings, it is evident that the proposed model outperforms the current state-of-the-art (SOTA) models in terms of performance.

Social media provides a rapid, simple, and accessible platform for people to communicate and share news online. However, the information published on this platform is not always trustworthy. As a result, malicious actors often use social media to disseminate fake news or mislead news readers, such as with personal or political attacks that could spark protests or riots. In this paper, we propose a learning technique for detecting fake news sources (i.e., fake users) on the Twitter platform. Three features—tweet content, published time, and social graph—have been defined and extracted from Twitter to create a deep neural network (DNN) as a predictive model. We conducted experiments on PolitiFact, a standard FakeNewsNet dataset. The results show that the proposed approach outperforms traditional baselines with 98.7% accuracy [29].

Table 2 Tweet content has all the words, and the lexical approach misclassified it. The community approach provided enough attributes for the fusion run to identify it correctly.

Content: <i>Explaining why beneficial effects from cannabis on intestine inflammation conditions like ulcerative colitis and Crohn’s disease have been reported often. If the endocannabinoid isn’t present, inflammation isn’t balanced; the body’s immune cells attack the intestinal lining.</i>
Ground Truth: non_conspiracy
Lexical model Prediction: 5g_corona_conspiracy
All connections network majority prediction: non_conspiracy
of connections in the 5g_corona_conspiracy dataset: 0
of connections in the other_conspiracy dataset: 129
of connections in the non_conspiracy conspiracy dataset: 185
% of <i>tweets</i> in the community that are from 5g_corona_conspiracy dataset: 10%
% of <i>tweets</i> in the community that are from other_conspiracy dataset: 25%
% of <i>tweets</i> in the community that are from non_conspiracy dataset: 65%

3 Methodology

This paper uses a scalable approach to gather, discover, analyze, and summarize joint sentiments of Twitter communities, extract community and network features, and improve the lexical-based baseline for Tweet classification using community information [2].

3.1 Content Analysis, Transformation, and Feature Selection

The tweets we analyzed had a content capacity of 280 characters. That limit tends to produce a writing style that differs from most corpora. To achieve brevity, users employ a lexicon that includes abbreviations, colloquialisms, *hashtags*, and *emojis*, and *tweets* may contain frequent misspellings. The context of a Tweet is also more affluent, as it resides in a rich network of retweets and replies. To this end, we employ lexical-based analysis and community analysis for Tweet content and context. The **Lexical Analysis Pipeline** implements the transformation of Twitter content, feature extraction, and modeling to make predictions for the NLP-based task [30].

In the *transformation* step, we tested several pre-processing, tokenization, and normalization techniques. We measured the influence of each transformation approach to predict performance on the part of the development set, turning off the feature and comparing the performance using 5-fold measures. Removing punctuation, preserving URLs, and normalizing several specific terms (e.g., ‘U.K.’ to ‘UK’) in the Tweet contributed to better content classification, as expected for the short tweet content. Stemming did not influence the classification recall on this small development set, nor did lemmatization. We speculate that the Tweet content was too short and the data was too small to derive any meaningful conclusion, and therefore we did not apply either. *Feature extraction* from Tweet content was implemented in two ways: encoding terms as vectors representing either the occurrence of terms in the text (*Bag-Of-Words*) or the impact of terms on a document in a corpus (TF-IDF). We extended the feature set in the *tweets* using *Optical Character Recognition* (OCR) of embedded images.

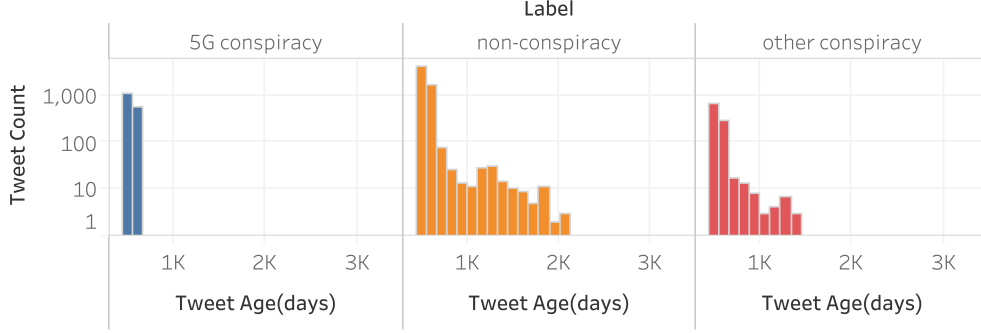


Fig. 5 Distribution of the feature *tweet_age* for the different class labels (5G, non, other).

3.2 Rich Graph Network Analysis

We apply the **Community Analysis Pipeline** for community discovery in networks created from user and hashtag connections to construct seven different networks from the raw Twitter data: *All Users Connections*, a network created from the labeled data set, with each vertex in the network being a user and each edge of the network being the connection between two users by either a retweet, quote, reply, mention, or friendship; *Retweet Connections*, which is similar to *All Users Connections*, but with each edge being the connection between two users by retweets only; *Mention Connections* which is similar to *All Users Connections*, but with each edge being the connection between two users by mentions only; *Reply Connections*, which is similar to *All Users Connections*, but with each edge being the connection between two users by replies only; *Quote Connections*, which is similar to *All Users Connections*, but with each edge being the connection between two users by quotes only; *Friends Connections*, which is similar to *All Users Connections*, but with each edge being the connection between two users by friendship only; and *Hashtag Connections* is a network created from the labeled data set, with each vertex in the network being a hashtag and each edge of the network being the connection between two hashtags used together in the same Tweet. We have developed an in-house scalable package *pytwanalysis* [1–3] to collect and save information-rich Twitter data, create networks, and discover communities in the data.

3.2.1 Community Labeling

We utilized all networks to learn the user attributes and *tweets* relevant to the community and topic. First, we found communities using an adapted Louvain method [1, 31]. We labeled each community with one of the three conspiracy categories (5G, non, other) based on the majority of the *tweets* for that community. If we found a community with more *tweets* with the *5G* label as opposed to *non* or *other*, we assigned the *5G* label to unlabeled *tweets* in that community. Figure 1 demonstrates a simplification of this method. We applied the method to all seven networks for community discovery and assigned seven community labels (from seven networks) to each Tweet, listed as

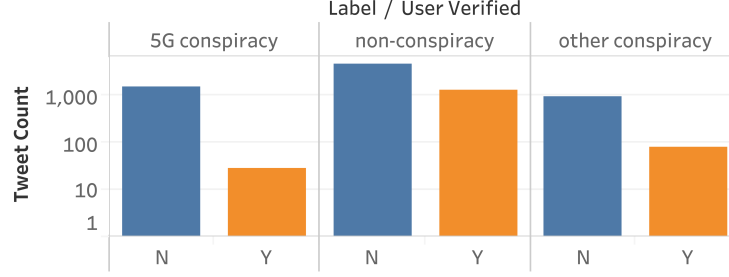


Fig. 6 Distribution of the feature *user_verified* for the different class labels (5G, non, other)

features 1 through 7 on Table 3. For the *Hashtag Connections* network, because one Tweet can have multiple hashtags, then one Tweet could belong to multiple hashtag communities. In that case, the majority logic selects the most common community found for that Tweet. The remaining *tweets* that did not belong to any community or that belonged to a community with *tweets* strictly originating from the test data set were assigned as *Unknown*. Many *Unknowns* were found because many *tweets* did not have any connections with other users in the labeled data sets (i.e., no retweets, replies, quotes, mentions, friends, or hashtags). An additional combined label was created with a combination of the other seven labels, listed as feature eight on Table 3. The combined label first uses the label from the quote network; if the quote network has an unknown value, it uses the value from the reply network, followed by the mention, all user connections, retweets, friends, and hashtag networks. The order of use for each network in the combined label was decided based on the evaluation metrics for the predictions coming from each network (Table 9). The community discovery approach can be helpful for data sets in which users are well-connected to each other.

User connectivity was also extracted from the graphs created from the development data sets. *User connectivity* is a feature that shows the degree of connectivity between each user in the *All Users Connections* network for each of the provided classification labels, driven by the observation that if vertices are well-connected, their content is similar. See features 9 through 12 on Table 3.

3.2.2 Attribute Labeling

User Attributes in the *tweets* are also extracted from the Twitter data. The produced networks can contain several disconnected *tweets*, so we expand the suite of network features and extract four additional user attributes and one Tweet attribute as follows: 1. *user_followers_count* (Fig. 2); 2. *user_friends_count* (Fig. 3); 3. *user_statuses_count* (Fig. 4); 4. *user_verified* (Fig. 6); 5. *tweet_age (days since creation)* (Fig. 5). Since the community majority selection predictions generated many unknown assignments, we used an additional classifier to help predict labels for *tweets* that were disconnected from the network. Since we have different types of features, we used the versatile Random Forest classifier that can work well with a mixture of categorical and numerical features. Community features 1 through 12 from Table 3 and user features

Table 3 Community attributes as explained in 3.2.1.

#	Community Feature
1	lv_comty_usr_all(majory_label)
2	lv_comty_usr_rt(majory_label)
3	lv_comty_usr_mention(majory_label)
4	lv_comty_usr_reply(majory_label)
5	lv_comty_usr_quote(majory_label)
6	lv_comty_usr_friend(majory_label)
7	lv_comty_usr_ht(majory_label)
8	lv_comty(majory_label)_combined
9	usr_degree_in_5g_corona_conspiracy
10	usr_degree_in_non_conspiracy
11	usr_degree_in_other_conspiracy
12	usr_degree_combined

1. to 5. listed above are used as input to the Random Forest classifier. The distribution of data for the features in the labeled data is shown in Figure 2, Figure 3, Figure 4, Figure 5, and Figure 6.

Community features 8 through 20 from Table 3 and user features from 1 through 5 are input to the multi-label (5G, non, other) Random Forest classifier. Because of the number of unknown predictions from the community assignments, this additional classifier helps predict labels for *tweets* that were disconnected from the network. Since we have different types of features, we used the versatile Random Forest classifier that can work well with a mixture of categorical and numerical features.

First, we create three different networks from the raw data: *User Connections* from provided data: vertex is a user, and each edge is the connection between two users by either a retweet, quote, reply, or mention; *Hashtag Connections* from provided data: vertex in the network is a hashtag, and edge exists between two hashtags if they were used together in the same tweet; and *User Connections 8M*: a network created from provided data and the auxiliary dataset of over 8M *tweets*, where vertices and edges of the network created the same way as the *User Connections* network. Next, we extract the degree of connectivity for each of the provided conspiracy labels (5G, non, and other) driven by the observation that if vertices are well connected, their content is similar. We employ the *Louvain Community* discovery method to discover communities in all three networks and apply to specific *tweets* information from each network analyzed [3]. We labeled each community with one of the three conspiracy categories (5G, non, other) based on the majority of the labels for that community associated with the tweet label. If we find a community where 5G labels are more significant than non other, we will use the 5G label to assign the label to unlabeled *tweets* in that community. These assignments were done based on the combination of communities in all three networks. *tweets* that did not belong to any community, or belonged to a community with *tweets* strictly originating from the test dataset, were assigned based on their degree of connectivity, and the remaining were assigned as *Unknown*. Many unknowns were found because many *tweets* did not have any connections with other users in the given datasets (no retweets, replies, quotes, mentions, or hashtags).

3.3 Modality Overlap Analysis

In this subsection, we aim to explore and determine whether the communities derived from different modalities exhibit low overlap, signifying complementary information, or if there is a considerable amount of overlap, suggesting redundancy or similar underlying structures. Quantifying this measure may help identify the modalities that contribute the unique information and design fusion methods accordingly. For example, it can allow us to identify which modalities should be assigned more weight to get the best performance in classification tasks.

3.3.1 Network Construction

After undergoing multiple pre-processing steps, a network has been constructed from the COVID-19 (+) data set, which consists of 8 million *tweets*. First, replies, quotes, and retweets are the selected connection modes of the network. Unlike in the case of quotes and retweets, we have found that there is no elaborate information present (full_text, media_url...etc.) replied by *tweets* in COVID-19 (+). Hence, we removed any edges constructed in the replies connection mode, where the target node is not found within the 8 million *tweets* due to the inability to extract textual and visual features from it. To reduce sparsity in the network, every target node should be connected to at least ten nodes. Otherwise, the isolated nodes or the nodes' connections falling under this threshold are pruned. Moreover, isolated nodes and duplicate edges were eliminated, and the first occurrence of any duplicate was kept. As a result, the total number of nodes and edges dropped to 3,407,903 and 3,316,523, respectively. For simplicity, every node ID, designated by its tweet ID, was mapped to values ranging from 0 to 3,407,902.

3.3.2 Visual and Textual Feature Extraction

We find that 154,923 *tweets* had images in COVID-19 (+). Some of the *tweets* were suspended, impeding some of the retrieval of the images. We also assigned the name of each image to its corresponding tweet ID, preserving the link between the tweet and the image. *VGG16* model pre-trained on *ImageNet* was employed as a feature extractor for all the images. On the other hand, textual embeddings were produced by a trained adapted version of BERT for COVID *tweets* called *BERTweet* by VinAIRsearch [32]. We utilized the baseline normalizations as elaborated below in subsection 3.1 but with a few alterations that include removing usernames, all special characters, hash-tags, contractions, non-English *tweets* if present, links (which not only incorporates "https://t.co/," but also "http" and "www"), and emojis. These additional textual normalizations were applied, and BERTweet features were subsequently extracted.

3.3.3 Augmented Network Construction

We seek to obtain an infused network that is comprised of the network above as well as a visual similarity graph. The latter is built by computing the cosine similarity between each node's image DNN features in the preprocessed network. The edges are hence, formed between each node and its five most visually similar nodes. The number

of edges bumped up to 4,091,138 in our processed COVID (+) network. The motive behind this is that the GNN will aggregate features from the neighboring nodes of not only those from replies, quotes, and retweets but also from the nodes with an image that’s visually like it.

3.3.4 Graph Neural Network Training

To leverage all modalities and aggregate features from neighborhood nodes, the adjacency, and the feature matrices are fed to an unsupervised GNN framework. The selected model for training the graph neural network is GraphSage [33], which produces an embedding output of size 50 dimensions. The hyperparameters are epoch = 1, batch size = 50, layer size = 50, and learning rate = 0.001, with Adam as an optimizer. The choice of this variant of GNN is ascribed to the fact that GraphSage utilizes the neighborhood sampling concept, which it renders scalable. GraphSage GNN has been trained separately on both the constructed and visually infused networks with the same textual feature matrix representing the nodes’ features.

3.3.5 Clustering

Both networks have been clustered using the Louvain Algorithm [34]. However, the rest has been clustered using HDBSCAN (Hierarchical DBSCAN) [35]. It is faster than regular DBSCAN. The minimum cluster size has been set to 10. Due to the memory constraints associated with clustering high dimensional textual embeddings and extensive data, the number of dimensions of the text has been reduced to 10 using the PCA method. However, the dimensions are intact when generating GNN embeddings.

4 Experimental Setup

4.1 Data Sets

The task at hand deals with highly imbalanced datasets as outlined in Table 4 for details). Generating fake *tweets* using the most predictive or most common terms for each class led to the over-fitting of most classifiers. We took a different route and adjusted class weights to account for imbalanced data when possible. The MediaEval Fake News Detection Task 2020 looks into *tweets* for misinformation claims that the construction of the 5G network and the associated electromagnetic radiation triggered the SARS-CoV-2 virus. We have received a labeled data set of approximately 6,000 *tweets* related to COVID-19, 5G, and their corresponding metadata; see details in Table 4). Note that all of our training was done using the development set, which contains 1,120 *tweets* labeled for 5G-COVID conspiracy, 688 *tweets* for another conspiracy, and 4,138 for non-conspiracy *tweets*, as shown in Table 4. This data set is small and very imbalanced. Thus, we extended the labeled data set with a new COVID-19 (+) data set that contains *tweets* related to #Coronavirus, #Covid19, and #Covid-19, collected from March through September 2020, with over 3.2 million users and 8 million *tweets* [3]. From the 8 million *tweets*, we filtered only the *tweets* that can make a connection in the existing networks created from the labeled data. After applying

Table 4 MediaEval 2020, COVID-19 (+), and friendship data sets. For MediaEval 2020, note that the number of users in each set does not add up to the total number of users, as the same user can have *tweets* in different data sets.

Dataset	Tweet Count	User Count
1. Fake News [4]	8,854	7,475
Development Labels	Tweet Count	User Count
5g_corona_conspiracy	1,120	1,053
other_conspiracy	688	638
non_conspiracy	4,138	3,643
Total	5,946	5,197
Test Labels	Tweet Count	User Count
5g_corona_conspiracy	532	512
other_conspiracy	346	334
non_conspiracy	2,030	1,832
Total	2,908	2,639
2. Friends of Fake News [4]		3,385,981
3. COVID-19 (+) [3]	771,203	657,785

the filter, we ended with 771,203 COVID-19 Tweets. The COVID-19 (+) data set was used to augment the feature space for classification. We also extended knowledge about user relationships by using the Twitter API to retrieve a list of friends for each user in the labeled data set. A total of 3,385,981 users were retrieved, but that number does not include 100% of the users in the friendship list, as some of the previously existing users are not accessible anymore (e.g., the account is suspended).

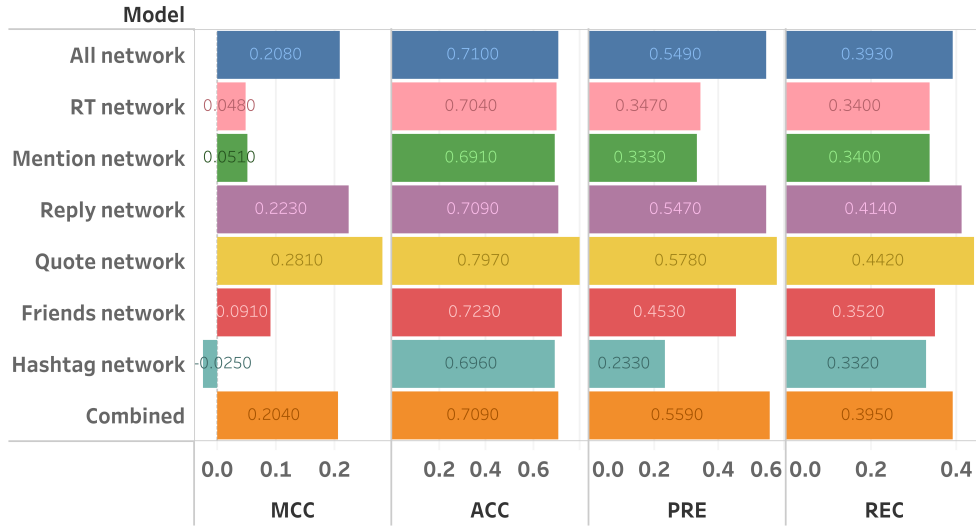


Fig. 7 Comparison of the multi-class community majority assignment excluding the unknowns for the different types of networks, as detailed in section *Multi-class without Unknowns* in Table 9

4.2 Measures

We measured the performance of the proposed methods on a tiny labeled subset of test data in Table 4. MediaEval officially reported that the metric used for evaluating the multi-class classification performance was the multi-class generalization of the *Matthews correlation coefficient* (MCC) [4, 36, 37]. MCC has advantages in bioinformatics over F1 and accuracy, as it considers the balance ratios of the four confusion matrix categories (true positives, true negatives, false positives, and false negatives). In a social network analysis, we are more interested in missed *tweets* (false negatives) and true positives. For this reason, we discuss our results from the perspective of precision, recall, and accuracy. We employ the adjusted Rand index (ARI) metric to measure the overlap between modalities and compare the partitions. We have already tested the lexical *classification* pipeline incorporating a variety of classifiers: Naive Bayes, Support Vector Machine, Random Forest, Multilayer Perceptron, Stochastic Gradient Descent, and a Logistic Regression classifier, and ended up using Logistic Regression, which has been shown to perform best for the content-based classification in [30]. We compared the performance of the classifiers on validation sets, both for the multi-class and binary classification subtasks.

5 Results and Analysis

5.1 Lexical Analysis Pipeline

Table 5 Logistic regression (LR) and logistic regression with OCR (LR-OCR) modeling scores for Multi-class and binary labeling of MediaEval 2020 test set.

Labeling	Multi-class				Binary			
Model	MCC	ACC	PRE	REC	MCC	ACC	PRE	REC
LR	0.435	0.749	0.597	0.569	0.492	0.789	0.749	0.743
LR-OCR	0.379	0.706	0.459	0.384	0.492	0.789	0.749	0.742

While the TF-IDF vectorizer captures the importance of terms well, we found better results using a *Bag-Of-Words* model in Section 5, likely due to the high occurrence and variety of colloquialisms and abbreviations. Table 5 shows the metrics for the multi-class and binary predictions using the Logistic Regression classifier [30]. The lexical analysis pipeline’s baseline results in this paper improve upon Data Lab’s best multi-class logistical regression (LR) model MediaEval 2020 submission [30] using cross-validation and regularization. The new best MCC result for the LR used in this paper is **0.435** for multi-class and **0.492** for binary classification.

5.2 Community Analysis Pipeline

Table 9 shows the metrics for the multi-class and binary predictions using the Louvain community majority assignment for each type of network with and without the COVID-19 (+) data set. Results are intuitive, as community majority assignments using the combined connections network with the COVID-19 (+) data set perform the

Table 6 Ternary (runs 001 - 004) and binary (runs 011 - 014) labeling scores returned by benchmark engine (MCC), and our analysis on development set (MediaEval 2020) released ground-truth (MCC, Precision, Recall, Acc). Model abbreviations: LR for logistic regression; LR-OCR for logistic regression w OCR; CL for community labeling; LR-CL for fusion run. The team has places second in the competition.

Evaluation Set		Test	Development			
Ternary	Model	MCC	MCC	Prec	Recall	Acc
001	LR	0.431	0.431	0.624	0.510	0.766
002	LR-OCR	0.363	0.465	0.599	0.565	0.767
003	CL	0.081	0.170	0.388	0.229	0.281
004	LR-CL	0.363	0.442	0.462	0.430	0.725
Binary	Model	MCC	MCC	Prec	Recall	Acc
011	LR	0.437	0.487	0.770	0.720	0.856
012	LR-OCR	0.428	0.516	0.780	0.737	0.862
013	CL	0.091	0.219	0.604	0.615	0.748
014	LR-CL	0.091	0.244	0.613	0.631	0.743

best over the range of measures. The table also shows the number of *tweets* that were classified as unknown when they did not belong to any community. The additional results for the Random Forest classifier are included in the table for comparison. Note that the total for each model is always 2,908, which is the number of labeled *tweets* in the test set.

The *Community Contribution Analysis* MediaEval 2020 development set is small and only captures fragments of the community. The number of unknown community assignments is large. It skews the use of community attributes, as shown by the low performance in section *Multi-class with Unknowns* in Table 9. Thus, we separate the evaluation in the multi-class community majority assignment into evaluation including the unknowns and evaluation excluding the unknowns. The metrics without the unknowns were calculated separately so that we could evaluate how well we could classify the *tweets* that did belong to a community, as shown in section *Multi-class without Unknowns* in Table 9 and Figure 7. Results calculated without the unknowns show comparative performance with the lexical pipeline.

The results in Table 9 show that the performance of community modeling is **comparable** to the lexical model if unknown assignments are excluded, and the quality of the predictions in different types of networks are broken down. Networks created from *quotes* and *replies* seem to yield the best results. Our initial premise is that similar topics and news are shared with the people who quote each other or participate in the same discussion thread, so this finding confirms the value of that correlation. On the other hand, the hashtag network’s predictions do not provide excellent results, as many of the same hashtags are used in both conspiracy and non-conspiracy-labeled data.

Labeling Considerations: The main challenge of the community approach is scale; the annotations and the topic should be prevalent in the data set to benefit from the community-based analysis truly. The COVID-19 (+) data set was obtained by finding an **intersection** of our originally mined data set of 8 million Tweets; see Section 4.1.

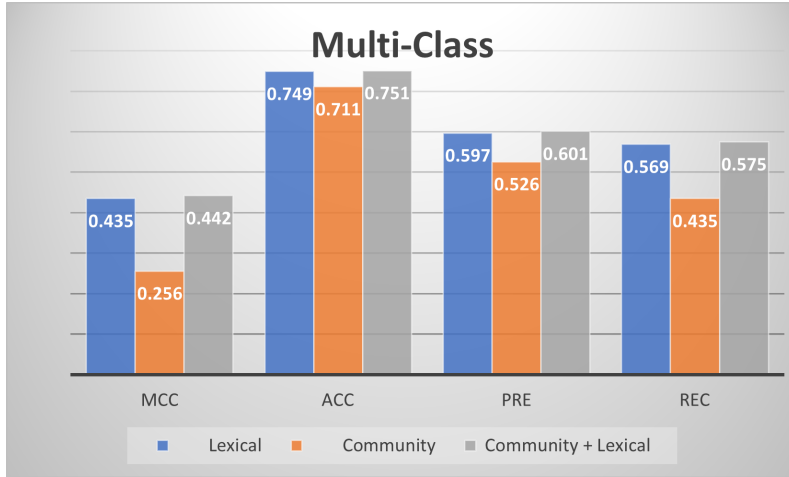


Fig. 8 Modeling comparisons on multi-class for the test set for Multi-Class classification. Community-only classification offers comparable precision and accuracy without even considering tweet text. Fusion of the lexical and community method offers the best performance across the board.

		Lexical Model	Community Network							Random Forest
			All	Retweet	Mention	Reply	Quote	Friends	Hashtag	
Community Network	Lexical Model	100%	70%	33%	46%	20%	17%	65%	22%	72%
	All	70%	100%	41%	57%	27%	22%	82%	28%	85%
	Retweet	33%	41%	100%	80%	68%	69%	41%	56%	37%
	Mention	46%	57%	80%	100%	62%	54%	54%	49%	52%
	Reply	20%	27%	68%	62%	100%	81%	28%	61%	22%
	Quote	17%	22%	69%	54%	81%	100%	27%	67%	19%
	Friends	65%	82%	41%	54%	28%	27%	100%	34%	77%
	Hashtag	22%	28%	56%	49%	61%	67%	34%	100%	25%
Random Forest		72%	85%	37%	52%	22%	19%	77%	25%	100%

Table 7 Overlap in the community multi-class predictions by the method: the percentage shows the overlap between the predictions of two methods out of the 2908 test records.

Community-based analysis with the auxiliary data brought the value of community connections to this analysis; compare model and model+ in Table 9. The COVID-19 (+) data set improved the connectivity in the network, which consequently improved the number of *tweets* that were able to be classified. The number of unknowns from the all connection network (All) decreased from 198 (All) to 108 (All+) when an analysis of the same labeled data was done within the more extensive network, and the MCC score jumped from 0.089 to 0.180. Using the Random Forest classifier over community and attribute labels improves the overall performance of the classification; see Table 9. The classifier can assign values for *tweets* that could not be classified with the community majority assignments since it uses additional features apart from the community features; see Section 3.2.2.

Table 10 summarizes the correct classification results that the network modeling produces that the lexical one does not. The community predictions perform comparably for cases where the Tweet was not isolated from the network. Figure 7 illustrates the overall multi-class detection overlap by the method. The highest overlap occurs between the *all connections* network predictions and the Random Forest model, which is expected since the network predictions were used as features for the Random Forest model. The lexical model overlaps most with the *all connections* network predictions and Random Forest. Other methods that have high overlap in their predictions are the *all connections* network with the *friends* network, the *retweet* network with the *mention* network, and the *quote* network with the *reply* network.

Table 8 Modeling comparisons on multi-class and binary results for the test set of MediaEval 2020

Labels	Multi-class				Binary			
	MCC	ACC	PRE	REC	MCC	ACC	PRE	REC
Lexical-(LogisticRegression)	0.435	0.749	0.597	0.569	0.492	0.789	0.749	0.743
Community-(RandomForest)	0.256	0.711	0.526	0.435	0.368	0.751	0.704	0.666
Community + Lexical	0.442	0.751	0.601	0.575	0.493	0.789	0.750	0.743

5.3 Combining Community and Lexical Attributes

In this experiment, we combine the logic of the lexical pipeline, as described in Section 3.1, and the community pipeline, as described in Section 3.2. We use the prediction of the lexical pipeline as a new input feature for the community pipeline that uses the Random Forest classifier. The combination of features that provided the best results was the following: lexical_prediction, user_followers_count, user_friends_count, user_statuses_count, user_verified, tweet_age, lv_comty_usr_all(majory_dataset), and lv_comty(majory_dataset)-combined.

Community modeling does not consider the tweet’s content beyond hashtags: it models the interactions with the tweet (mentions, quotes, retweets, replies), and with the author (friends). The model trained on community-based and lexical-based features achieved the highest MCC score on the test set, as shown in Table 8. Binary lexical and community classifications (non-conspiracy vs. conspiracy) perform better than the lexical multi-class baseline. Recent work has shown different dispersion patterns regardless of the conspiracy topic [38], and our community and lexical binary capture this observation well, as it outperforms across four different measures of classification efficiency; see Table 8 for details.

5.4 Quantifying Modality Overlap

Table 11 shows that multiple modalities seem to capture specific information, and it is not relevant for community discovery at a global scale due to the negligible overlap between the modalities. However, communities produced by each modality might have value for specific discovery and mining tasks. The low overlap provides insights into

Table 9 Predictions for the community labeling using MediaEval development data and Auxiliary COVID-19 (+) data set. Performance measures (MCC, Precision, Recall, Accuracy) were computed for every type of network for multi-class classification, including the unknown predictions, for multi-class classification, excluding the unknown predictions, and for binary classification.

			Multi (Unknowns)				Multi (No Unknowns)				Binary predictions			
Community Predictions - Majority Selection														
Description	Total	Unknowns	MCC	ACC	PRE	REC	MCC	ACC	PRE	REC	MCC	ACC	PRE	REC
All network	2908	198	0.089	0.664	0.425	0.249	0.101	0.713	0.566	0.352	0.276	0.733	0.694	0.598
RT network	2908	2908									0.000	0.698	0.349	0.500
Mention network	2908	2095	0.027	0.192	0.386	0.084	0.204	0.686	0.514	0.403	0.123	0.703	0.632	0.529
Reply network	2908	2474	0.036	0.098	0.361	0.051	0.234	0.654	0.481	0.448	0.137	0.706	0.644	0.533
Quotes network	2908	2659	0.064	0.067	0.457	0.035	0.461	0.783	0.609	0.597	0.110	0.704	0.663	0.518
Friends network	2908	390	0.091	0.627	0.405	0.232	0.074	0.724	0.540	0.346	0.231	0.722	0.680	0.574
Hashtag network	2908	2158	-0.002	0.174	0.326	0.065	0.070	0.675	0.434	0.345	0.058	0.699	0.636	0.506
Combined	2908	154	0.142	0.675	0.391	0.270	0.161	0.713	0.522	0.377				
Community Predictions - Majority Selection - COVID-19 (+) Dataset														
Description	Total	Unknowns	MCC	ACC	PRE	REC	MCC	ACC	PRE	REC	MCC	ACC	PRE	REC
All network +	2908	108	0.180	0.683	0.412	0.283	0.208	0.710	0.549	0.393	0.345	0.743	0.692	0.655
RT network +	2908	1636	0.012	0.308	0.261	0.112	0.048	0.704	0.347	0.340	0.231	0.724	0.700	0.567
Mention network +	2908	1107	0.006	0.428	0.250	0.157	0.051	0.691	0.333	0.340	0.209	0.716	0.661	0.568
Reply network+	2908	2107	0.040	0.195	0.410	0.085	0.223	0.709	0.547	0.414	0.134	0.704	0.632	0.534
Quote network +	2908	2296	0.075	0.168	0.433	0.070	0.281	0.797	0.578	0.442	0.138	0.707	0.668	0.528
Friends network +	2908	392	0.101	0.625	0.340	0.235	0.091	0.723	0.453	0.352	0.243	0.725	0.682	0.581
Hashtag network +	2908	2076	-0.001	0.199	0.174	0.071	-0.025	0.696	0.233	0.332	-0.017	0.697	0.349	0.500
Combined +	2908	80	0.180	0.689	0.419	0.288	0.204	0.709	0.559	0.395				
ML Classifier														
Description	Total	Unknowns	MCC	ACC	PRE	REC	MCC	ACC	PRE	REC	MCC	ACC	PRE	REC
Random Forest	2908	0	0.256	0.711	0.526	0.435					0.368	0.751	0.704	0.666

Table 10 Comparison of the predictions between the community and lexical models. The test data set has 2,908 labeled Tweets. *Equal to lexical* is the number of predictions for that model that were classified the same as the lexical model. *Unique* is the number of predictions that the model predicted differently than the lexical model.

Lexical Model vs Community Predictions				
Lexical Model Multi-class : correct 2,177; incorrect 731				
	Equal to Lexical		Unique	
Model	Correct	Incorrect	Correct	Incorrect
All network	1726	470	261	451
RT network	799	635	96	1378
Mention network	1106	592	139	1071
Reply network	499	662	69	1678
Quote network	443	686	45	1734
Friends network	1604	517	214	573
Hashtag network	523	671	60	1654
Random Forest	1772	434	297	405
Lexical Model Binary : correct 2,293; incorrect 615				
	Equal to Lexical		Unique	
Model	Correct	Incorrect	Correct	Incorrect
All network	1810	265	350	483
RT network	1783	292	323	510
Mention network	1767	299	316	526
Reply network	1737	305	310	556
Quote network	1746	304	311	547
Friends network	1788	295	320	505
Hashtag network	1705	319	296	588
RandomForest	1855	286	329	438

Table 11 ARI & number of communities between five multi-modal modes for COVID-19 (+). 1: Network, 2: Text Embeddings, 3: Graph Neural Network (GNN) embeddings, 4: Augmented network with visual edges, 5: GNN embeddings produced by training on augmented network with visual edges and text embeddings, 6: Number of communities.

ARI	COVID-19 (+)				
	1	2	3	4	5
1	1.0	0.084	0.0002	0.124	0.001
2	0.084	1.0	0.0004	0.053	0.0265
3	0.0002	0.0004	1.0	0.0001	-0.001
4	0.124	0.053	0.0001	1.0	0.0138
5	0.001	0.0265	-0.001	0.0138	1.0
6	91,380	81,252	30,995	67,146	87,505

the effectiveness of different modalities in capturing the underlying patterns within multi-modal tweet data and how much they complement each other.

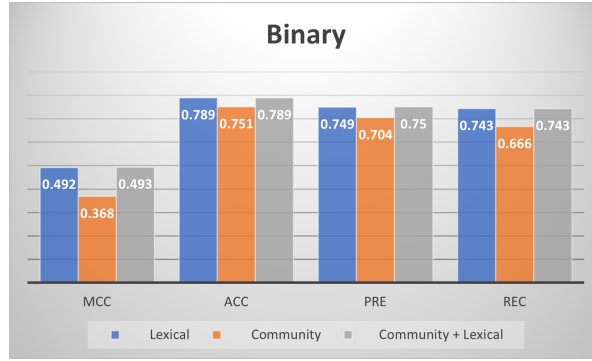


Fig. 9 Modeling comparisons on binary results for the test set for Binary classification. Community-only classification offers comparable precision and accuracy without even considering tweet text. Fusion of the lexical and community method offers the best performance across the board.

6 Discussion and Outlook

In conclusion, this research highlights the significant influence of community behavior in tweet classification, suggesting that it carries a comparable weight to tweet content. By introducing a community-based approach to tweet classification, we successfully utilized six distinct community network knowledge graphs to classify tweet content accurately. Our findings demonstrate the advantages of incorporating community attributes and models into the lexical baseline for tweet classification.

Notably, community networks offer valuable contextual information for understanding tweet communication, and our study reveals that community-only modeling is as informative as content modeling, as it encompasses crucial details regarding social network interactions with the tweet object. Remarkably, our community modeling techniques, implemented on a large-scale real network, achieved comparable precision, recall, and accuracy to a lexical classifier, even without considering tweet content beyond hashtags. Furthermore, we have shown that essential fusion techniques outperform lexical and network baselines. In contrast, the combination of community and lexical approaches produces the most robust outcomes and superior performance measures, as evidenced by the MediaEval Fake News task results. The complex knowledge graph depicted in Figure 7, which encompasses retweet, mentions, reply, and quote networks, illustrates our ability to capture and incorporate comprehensive network information. Moving forward, we plan to explore enhanced network selection and fusion methods in conjunction with Lexical Modeling and Friends Network, aiming to improve tweet classification’s effectiveness and accuracy.

References

- [1] Nogueira, L.: pytwanalysis Package. <https://pypi.org/project/pytwanalysis/>
- [2] Nogueira, L., Tešić, J.: pytwanalysis: Twitter data management and analysis at scale. In: International Conference on Social Network Analysis Management and Security (SNAMS2021) (2021). <https://emergingtechnet.org/SNAMS2021/>

- [3] Nogueira, L.: Social network analysis at scale: Graph-based analysis of Twitter trends and communities. Master's thesis, Texas State University (Dec 2020). <https://digital.library.txstate.edu/handle/10877/12933>
- [4] Pogorelov, K., Schroeder, D.T., Burchard, L., Moe, J., Brenner, S., Filkukova, P., Langguth, J.: Fake News: Coronavirus and 5g conspiracy task at MediaEval 2020. In: Working Notes Proceedings of the MediaEval 2020 Workshop. MediaEval, ??? (2020). <http://ceur-ws.org/Vol-2882/>
- [5] Osmundsen, M., Bor, A., Vahlstrup, P.B., Benchmann, A., Petersen, M.B.: Partisan polarization is the primary psychological motivation behind political fake news sharing on Twitter. *American Political Science Review* **115**(3), 999–1015 (2021) <https://doi.org/10.1017/S0003055421000290>
- [6] Geeng, C., Yee, S., Roesner, F.: Fake news on Facebook and twitter: Investigating how people (don't) investigate. In: Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems. CHI '20, pp. 1–14. Association for Computing Machinery, New York, NY, USA (2020). <https://doi.org/10.1145/3313831.3376784> . <https://doi.org/10.1145/3313831.3376784>
- [7] Bovet, A., Makse, H.A.: Influence of fake news in Twitter during the 2016 us presidential election. *Nature communications* **10**(1), 1–14 (2019)
- [8] Ahmed, W., Vidal-Alaball, J., Downing, J., Seguí, F.L.: Covid-19 and the 5g conspiracy theory: a social network analysis of twitter data. *Journal of Medical Internet Research* **22**(5), 19458 (2020)
- [9] Sha, H., Hasan, M.A., Mohler, G., Brantingham, P.J.: Dynamic topic modeling of the covid-19 Twitter narrative among us governors and cabinet executives. arXiv preprint arXiv:2004.11692 (2020)
- [10] al., I.: Using logistic regression method to classify tweets into the selected topics. In: Intl. Conf. on Advanced Computer Science and Information Systems (ICACSIS), pp. 385–390. IEEE, NY (2016)
- [11] Zhou, Zafarani: Fake news detection: An interdisciplinary research. In: WWW Proceedings, p. 1292. ACM, NY (2019)
- [12] al., M.: Fake News Detection on Social Media Using Geometric Deep Learning (2019)
- [13] al., K.: An anatomical comparison of fake news and trusted-news sharing patterns on Twitter. *Computational and Mathematical Organization Theory* (2020)
- [14] Nguyen, V.-H., Sugiyama, K., Nakov, P., Kan, M.-Y.: Fang: Leveraging social context for fake news detection using graph representation. In: Proceedings of the 29th ACM International Conference on Information and Knowledge Management.

- CIKM '20, pp. 1165–1174. Association for Computing Machinery, New York, NY, USA (2020). <https://doi.org/10.1145/3340531.3412046> . <https://doi.org/10.1145/3340531.3412046>
- [15] Su, T.: Automatic fake news detection on Twitter. PhD thesis, University of Glasgow (2022)
 - [16] Gangireddy, S.C.R., P, D., Long, C., Chakraborty, T.: Unsupervised fake news detection: A graph-based approach. In: Proceedings of the 31st ACM Conference on Hypertext and Social Media. HT '20, pp. 75–83. Association for Computing Machinery, New York, NY, USA (2020). <https://doi.org/10.1145/3372923.3404783> . <https://doi.org/10.1145/3372923.3404783>
 - [17] Schroeder, D.T., Pogorelov, K., Langguth, J.: Fact: a framework for analysis and capture of Twitter graphs. In: 2019 Sixth International Conference on Social Networks Analysis, Management and Security (SNAMS), pp. 134–141 (2019). <https://doi.org/10.1109/SNAMS.2019.8931870>
 - [18] Bansal, S.: A mutli-task mutlimodal framework for tweet classification based on cnn (grand challenge). 2020 IEEE Sixth International Conference on Multimedia Big Data (BigMM), Multimedia Big Data (BigMM), 2020 IEEE Sixth International Conference on, 456–460 (2020)
 - [19] Suman, C., Naman, A., Saha, S., Bhattacharyya, P.: A multimodal author profiling system for tweets. IEEE Transactions on Computational Social Systems, Computational Social Systems, IEEE Transactions on, IEEE Trans. Comput. Soc. Syst **8**(6), 1407–1416 (2021)
 - [20] Gao, W., Li, L., Zhu, X., Wang, Y.: Detecting disaster-related tweets via multimodal adversarial neural network. IEEE MultiMedia, MultiMedia, IEEE **27**(4), 28–37 (2020)
 - [21] Bruijn, J.A., Moel, H., Weerts, A.H., Ruiter, M.C., Basar, E., Eilander, D., Aerts, J.C.J.H.: Improving the classification of flood tweets with contextual hydrological information in a multimodal neural network. Computers and Geosciences **140** (2020)
 - [22] Lim, W.L., Ho, C.C., Ting, C.: Sentiment analysis by fusing text and location features of geo-tagged tweets. IEEE Access, Access, IEEE **8**, 181014–181027 (2020)
 - [23] Gao, D., Li, K., Wang, R., Shan, S., Chen, X.: Multi-modal graph neural network for joint reasoning on vision and scene text. 2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), Computer Vision and Pattern Recognition (CVPR), 2020 IEEE/CVF Conference on, CVPR, 12743–12753 (2020)

- [24] Yang, X., Deng, C., Dang, Z., Wei, K., Yan, J.: Selsagcn: Self-supervised semantic alignment for graph convolution network. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), pp. 16775–16784 (2021)
- [25] Wang, J., Wang, Y., Yang, Z., Yang, L., Guo, Y.: Bi-gcn: Binary graph convolutional network. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), pp. 1561–1570 (2021)
- [26] Dai, E., Aggarwal, C., Wang, S.: Nrgnn : Learning a label noise resistant graph neural network on sparsely and noisily labeled graphs. Proceedings of the 27th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, 227–236 (2021)
- [27] Liu, Z., Nguyen, T.-K., Fang, Y.: Tail-gnn : Tail-node graph neural networks. Proceedings of the 27th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, 1109–1119 (2021)
- [28] Dai, E., Wang, S.: Say no to the discrimination: Learning fair graph neural networks with limited sensitive attribute information. In: Proceedings of the 14th ACM International Conference on Web Search and Data Mining, pp. 680–688. Association for Computing Machinery, New York, NY, USA (2021). <https://doi.org/10.1145/3437963.3441752>
- [29] Bhatia, T., Manaskasemsak, B., Rungsawang, A.: Detecting fake news sources on twitter using deep neural network. In: 2023 11th International Conference on Information and Education Technology (ICIET), pp. 508–512 (2023). <https://doi.org/10.1109/ICIET56899.2023.10111446>
- [30] Magill, A., Tomasso, M.: Fake News Twitter Data Analysis. <https://github.com/DataLab12/fakenews> (2020)
- [31] Aynaud, T.: python-louvain 0.14: Louvain algorithm for community detection. <https://github.com/taynaud/python-louvain> (2020)
- [32] Nguyen, D.Q., Vu, T., Nguyen, A.T.: BERTweet: A pre-trained language model for English Tweets. In: Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations, pp. 9–14 (2020)
- [33] Hamilton, W., Ying, Z., Leskovec, J.: Inductive representation learning on large graphs. In: Guyon, I., Luxburg, U.V., Bengio, S., Wallach, H., Fergus, R., Vishwanathan, S., Garnett, R. (eds.) Advances in Neural Information Processing Systems, vol. 30. Curran Associates, Inc., ??? (2017)
- [34] Blondel, V.D., Guillaume, J.-L., Lambiotte, R., Lefebvre, E.: Fast unfolding of communities in large networks. Journal of Statistical Mechanics: Theory and Experiment **2008**(10), 10008 (2008) <https://doi.org/10.1088/1742-5468/2008/>

- [35] Campello, R.J.G.B., Moulavi, D., Sander, J.: Density-based clustering based on hierarchical density estimates. In: Pei, J., Tseng, V.S., Cao, L., Motoda, H., Xu, G. (eds.) *Advances in Knowledge Discovery and Data Mining*, pp. 160–172. Springer, Berlin, Heidelberg (2013)
- [36] Chicco, D., Jurman, G.: The advantages of the Matthews correlation coefficient (MCC) over f1 score and accuracy in binary classification evaluation. *BMC genomics* **21**(1), 1–13 (2020)
- [37] Baldi, P., Brunak, S., Chauvin, Y., Andersen, C.A., Nielsen, H.: Assessing the accuracy of prediction algorithms for classification: an overview. *Bioinformatics* **16**(5), 412–424 (2000)
- [38] al., V.: The spread of true and false news online. *Science* **359**(6380), 1146–1151 (2018)