



Does Fake News in Different Languages Tell the Same Story? An Analysis of Multi-level Thematic and Emotional Characteristics of News about COVID-19

Lina Zhou¹ · Jie Tao² · Dongsong Zhang¹

Accepted: 12 August 2022 / Published online: 26 September 2022

© The Author(s), under exclusive licence to Springer Science+Business Media, LLC, part of Springer Nature 2022

Abstract

Fake news is being generated in different languages, yet existing studies are dominated by English news. The analysis of fake news content has focused on lexical and stylometric features, giving little attention to semantic features. A few studies involving semantic features have either used them as the inputs to classifiers with no interpretations, or treated them in isolation. This research aims to investigate both thematic and emotional characteristics of fake news at different levels and compare them between different languages for the first time. It extends a state-of-the-art topic modeling technique to extract news topics and introduces a divergence measure to assess the importance of thematic characteristics for identifying fake news. We further examine associations of the thematic and emotional characteristics of fake news. The empirical findings have implications for developing both general and language-specific countermeasures for fake news.

Keywords Fake news · Cross-lingual analysis · Topic modeling · Theme · Emotion · COVID-19

1 Introduction

Fake news is being generated and disseminated at a staggering rate. Social media enables rapid news dissemination and lowers barriers to reach a broad audience (Abonizio et al., 2020), yet it also contributes to the proliferation of fake news. Fake news has the potential to influence political outcomes, lure consumers into deceptive marketing schemes, defame business firms or celebrities, and mislead the public into making wrong decisions (Akhter et al., 2021). Moreover, fake news spreads across the globe in different languages, magnifying its impact. The COVID-19 pandemic has led to devastating consequences, uncertainty, and impact on every aspect of humanity at a globe scale, feeding an

enormous amount of fake news disseminated broadly across different languages on various social media platforms. Thus, the pandemic provides an opportunity for researchers to understand the characteristics of fake news in different languages to better inform the research and practice on fake news detection and intervention.

Previous studies on fake news detection have predominantly focused on English news (Davoudi et al., 2022; Pérez-Rosas et al., 2018; Shu et al., 2019). Studies on the characteristics of fake news in other languages remain scarce. Given that culture, language, political views, and religion may influence the way that news is generated, perceived, and disseminated, it is important to understand the characteristics of fake news in different languages. Despite that a few studies have explored the detection of fake news in multiple languages (Abonizio et al., 2020; Faustini & Covões, 2020), they have focused on developing language-independent detection models rather than understanding the characteristics of fake news in a multi-lingual setting. In addition, current development of fake news detection models has primarily drawn on lexical and stylometric features from news content while giving little attention to semantic features. Some studies have explored extracting topics using topic modeling techniques to understand fake news (Paixão et al., 2020; Sabeeh et al., 2021). However, Paixão et al.

✉ Lina Zhou
lzhou8@uncc.edu

Jie Tao
jtao@fairfield.edu

Dongsong Zhang
dzhang15@uncc.edu

¹ University of North Carolina at Charlotte, Charlotte, NC 28233, USA

² Fairfield University, Fairfield, CT 06824, USA

(2020) extract topics from fake and real news separately and merely list sample topics without providing any statistical evidence for a comparison between the two types of news or their effects on fake news detection. Sabeeh et al. (2021) utilize an LDA (Latent Dirichlet Allocation) based method for computing topic match scores between news headlines and bodies to transform multi-class data in fake news detection. Moreover, semantic features may be extracted at different levels of granularity. More specific features such as individual topics may be less generalizable for fake news detection. Word embedding models have the potential to capture the meaning of words for fake news detection (e.g., Faustini & Covões, 2020; Paixão et al., 2020)), yet it remains difficult to interpret those representations. Deep neural network models that incorporate word embeddings for fake news detection have centered on improving the detection performance (e.g., Sabeeh et al., 2021) rather than explaining the characteristics of fake news. Without being transformed into human understandable knowledge, those trained models and outputs would have limited impacts on humans' battling against fake news and on developing targeted countermeasures and mitigation strategies.

To address the above-mentioned research gaps, this study investigates the semantic characteristics of fake news in theme and emotion aspects across different languages. Specifically, it aims to answer three research questions. First, *how to characterize the themes and emotions expressed in news contents?* A related question is how to extract themes (used interchangeably with topics hereafter) from news texts effectively. Second, *how do the thematic characteristics of fake news that distinguish itself from real news vary across different languages?* Third, *how do the emotional characteristics of fake news that distinguish itself from real news vary across different languages?*

We address those research questions by analyzing COVID-19 related news. Since the start of the pandemic, health agencies and policy makers have been battling the proliferation of fake news while working to curb the spread of COVID-19. However, it is proven to be difficult for them to keep track of the growth of false information and even harder to address the real concerns of the public (Nwankwo et al., 2020). The global issue is not alleviated albeit the efforts from different social media platforms. In this study, we choose English and Chinese news because, according to Statista,¹ they are the top-2 most common languages used on the Internet. For either language, we first collect fake news datasets in relation to COVID-19 and extract themes from the news by developing a transformer-based topic modeling framework. Then, we design semantic features at different

levels of granularity to characterize themes and emotions expressed in the news. Next, we identify and compare news themes and emotional characteristics that can help distinguish fake news from real news in either language. Finally, we examine the effect of language on thematic and emotional characteristics of fake news.

The rest of the paper is organized as follows. We first review related work in Section 2 and then introduce our research method in Section 3. Subsequently, we report the analysis results in Section 4 and finally conclude the paper with Section 4.

2 Related Work

Since news contents are the primary source of semantic information, such as themes and emotions, we first review the literature on textual features of fake news, followed by an introduction to the topic modeling techniques and their application in fake news detection given that topic modeling remains the dominant method for extracting themes or topics from text. Next, we review literature on fake news detection in multiple languages. Finally, based on our review of the related work and their limitations, we propose research questions and hypotheses.

2.1 Textual Features of Fake News

Automatic detection of online fake news mainly draws on features from news contents and social context (Shu et al., 2017). News contents consist of news title, text body, and images and videos embedded in news. Accordingly, textual and visual features can be extracted from the news contents to support fake news detection. Moreover, textual features can be represented at different levels of granularity such as word, sentence, and article levels. Sample visual features of news include clarity score, coherence score, and so on. Social context features can be derived from users' social engagements during news consumption on social media platforms, which are further divided into user-based (e.g., source credibility and number of followers/followees), post-based (e.g., number of responses received), and network-based features (e.g., centrality measures) (Shu et al., 2019). Unlike content-based features, social context features may not be readily available on some social media platforms. Among different types of content features, textual features are the most commonly used (Zhang & Ghorbani, 2020), and thus are utilized in this study.

A variety of textual features have been used to detect fake news. Based on a classification framework for text-based cues to online deception (Zhou et al., 2004), which has been widely used to guide feature engineering for fake news detection, we classify textual features into the following

¹ www.statista.com/statistics/262946/share-of-the-most-common-languages-on-the-internet/

categories: lexical, morphological, syntactic, semantic, and discourse features. Lexical features are based on individual words or terms in text. Morphological features are based on the analysis of structure and parts of words such as parts-of-speech. Syntactic features are about linguistic constituents such as phrases through analyzing the structure of sentences in news text by following certain syntactic rules. Semantic features are related to the meaning of news text; and discourse features focuses on the use, purpose, or functions of text by considering its context.

Fake news detection research has used *lexical features*, such as news length, subjectivity (Ozbay & Alatas, 2020; Reis et al., 2019), percentage of uppercase characters/exclamation marks/questions marks, the number of unique words, spelling errors (Faustini & Covões, 2020), word diversity, readability (e.g., Flesch Reading Ease and the SMOG Index) (Choudhary & Arora, 2021), sentiment (e.g., sentiment polarity of news), psycholinguistic features (e.g., LIWC features) (Paixão et al., 2020; Reis et al., 2019), and n-grams (Bakir & McStay, 2018); *morphological features* such as proportion of adjectives/adverbs/nouns (Faustini & Covões, 2020), *syntactic features*, such as noun phrases (Zhou et al., 2004) and CFG-based features (Pérez-Rosas et al., 2018); and *semantic features*, such as word embeddings (Faustini & Covões, 2020; Paixão et al., 2020) and latent topics (Paixão et al., 2020; Sabeeh et al., 2021). Existing studies have primarily focused on lexical features while giving little attention to semantic features. Word embeddings enable the words with the same meaning to have similar representations, but they are difficult to explain. Their extraction of latent topics has mainly relied on LDA (Blei et al., 2003) while overlooking more recent development in topic modeling techniques (see Section. 2.2 for discussion). More importantly, prior studies either list sample extracted topics to gain a qualitative understanding only (Paixão et al., 2020) or use extracted topics as the input features to classification models without providing any interpretations (Sabeeh et al., 2021). Furthermore, those studies treat each topic in isolation without considering the relationships between topics and the overall topic distributions.

2.2 Topic Modeling Techniques and Applications in Fake News Research

An important step in combating online misinformation is to understand its common themes proliferating through social media, so that corresponding factual information can be provided (Nwankwo et al., 2020). Topic modeling refers to using unsupervised learning techniques for text analysis to determine clusters of terms that represent the topics of the documents. Topic modeling involves representing words and grouping similar word representations to infer topics from text documents. The objectives of topic modeling include

discovering latent topics present in a textual corpus, annotating documents based on topic loadings, and using the identified topics and terms associated with those topics to organize, search, understand, and summarize text. For example, if a topic model indicates that many social media posts are discussing face masks, then governments can provide information about different types of face masks, correct wearing methods, and their effectiveness in providing protection against COVID-19.

There are a number of traditional topic modeling techniques, including LDA (Blei et al., 2003), Non Negative Matrix Factorization (NMF) (Dhillon & Sra, 2005), Latent Semantic Analysis (LSA) (Dumais, 2004), and Pachinko Allocation Model (PAM) (Li & McCallum, 2006). Among them, LDA is the most commonly deployed method, which describes each document by the probabilistic distribution of topics and describes each topic by the probabilistic distribution of the co-occurrence of words. One add-on method to LDA is NMF, which decomposes (or factorizes) a high-dimensional vector into two matrices, namely a term-topic matrix and a topic-document matrix, in which the coefficients (e.g., weights for the topics) are non-negative. LSA extracts the relationships among different words in a document corpus by determining the optimal number of topics through an iterative process. PAM is a variation of LDA. Unlike LDA, however, PAM models correlations among the generated topics.

Since transformer-based models gained significant attention from the Natural Language Processing (NLP) community, they have also been applied to the topic modeling task. Compared to the traditional topic modeling techniques, which mainly rely on the co-occurrence of words, transformer-based models utilize the semantic information captured via text embeddings. Transformer-based models such as BERT (Bidirectional Encoder Representations from the Transformers) (Devlin et al., 2019) combined with the class-based TF-IDF (c-TF-IDF) metric (Grootendorst, 2020) can create easy-to-interpret topics via dense clusters (Grootendorst, 2020). Transformer-based models usually yield superior results compared to traditional machine learning and deep learning models with much less engineering time (i.e., fine tuning versus training) because they exclusively use the multi-head attention mechanism and are trained on massive generic corpora. Additionally, transformer-based models produce better text representations than discrete text representation techniques (e.g., word2vec (Mikolov et al., 2013)) because transformer-based embeddings are context dependent (i.e., the same word has different embeddings in different contexts), position-aware (i.e., taking the positions of words into consideration in terms of representation), able to generate representations beyond a word level (e.g., sentence representation rather than the average of word representations

using word2vec), and better at handling out-of-vocabulary words. Applying transformer-based models to topic modeling typically goes through three main steps: 1) learning real-valued vector representations of text documents using selected transformer models; 2) performing dimension reduction (e.g., using Uniform Manifold Approximation and Projection UMAP (McInnes & Healy, 2018)) on the learned document representations, followed by clustering the dimension-reduced representations based on their semantic similarities in the embedding space; and 3) extracting topics from the clusters using the c-TF-IDF metric (Grootendorst, 2020), and representing each of the topics using a set of terms. Similar to the classic tf-idf metric, the final c-TF-IDF value for any term t is derived by multiplying its tf and idf values.

Most of the research on fake news detection relies on supervised learning algorithms (Sabeeh et al., 2021), which is heavily dependent on the quality of labels. A few recent studies have also leveraged topic modeling techniques. Sabeeh et al. (2021) built a two-step fake news detection model by using a pre-trained BERT model to assist the first-step classification and discovering the topics of the headline and body of news articles to support the second-step classification. However, the BERT model was used for the extraction of word embeddings, and the extracted topics mainly served as the input features to support fake news detection without any explicit interpretation or illustration. Another study (Gupta et al., 2022) identified topics and key themes emerging in COVID-19 fake and real news to understand the strategies that fake news writers used to lure people to read and spread fake news about COVID-19. They extracted five topics from fake and real news separately. A comparison between the two sets of topics reveals common themes shared between fake and real news, such as health hazards, spread statistics, and counter measures, as well as some major differences. The study did not provide any rationale for choosing five topics, which can have an impact on topic quality. Given that labeled fake news remains scarce and news in the real-world consists of a mixture of real and fake news, it would make a more ecological sense to combine fake and real news for topic extraction. Ito et al. (2015) determined whether a user was a domain expert or biased by analyzing and comparing the topical divergence of their tweets with those of other users. The findings of their study suggest that using topical features is an effective way of assessing user credibility on Twitter. However, their study used topics as signals to enhance user credibility classification rather than interpret those topics for fake news detection. Importantly, all of the above studies employed LDA for topic modeling, which suffers from the limitations discussed earlier. Moreover, none of them analyzed fake news in more than one language. Furthermore, they analyzed topics in isolation

without looking into their relationships and patterns among individual news.

2.3 Fake News in Different Languages

Social media platforms augment the speed at which social media content reaches a broad audience (Blanco-Herrero & Calderón, 2019). During the COVID-19 pandemic, countries have endorsed the cross-regional statement on “infodemic”, and the spread of fake news is considered “as dangerous to human health and security as the pandemic itself.” However, compared with research on English fake news, there is a scarce of studies on fake news in other languages, partly because of the lack of the related datasets and the challenges in analyzing news in multiple languages.

With an increasing recognition of the importance of studying fake news in non-English languages, scholars have made some efforts to collect fake news datasets in various languages (e.g., Abonizio et al., 2020; Kishore Shahi & Nandini, 2020; Posadas-Durán et al., 2019; Yang et al., 2021). For instance, FIRE (Forum for Information Retrieval Evaluation) hosted the first shared task focusing on fake news detection in the Urdu language in 2020 (Amjad et al., 2020). The availability of such datasets enable researchers to explore the detection of fake news in other languages. However, existing research mainly focuses on developing either models for non-English languages only, or language-independent models for fake news detection. For instance, Al-Ash et al. (2019) deployed ensemble learning methods for Indonesian fake news detection. Du et al. (2021) detected COVID-19 misinformation in Chinese using a deep learning framework and a curated Chinese real and fake news dataset according to existing fact-checked news in English. Kar et al. (2021) proposed a BERT-based model augmented by additional relevant features extracted from tweets for multiple Indic-languages (e.g., Hindi and Bengali) besides English.

Another related research stream is focused on developing language-independent models that rely on general text features for fake news detection. For example, Faustini and Covões (2020) developed a generic approach to detecting fake news in three different languages: English, Portuguese, and Bulgarian. They used a number of input features, including frequency counts of text features (e.g., proportion of uppercase characters, number of sentences, and number of words per sentence), word2vec representations, and bag-of-words with tf-idf. Their results show that text length and word2vec are the most important features across different languages. Abonizio et al. (2020) evaluated language-independent textual features, such as complexity, stylometric, and psychological features, for detecting fake news in American English, Brazilian Portuguese, and Spanish using traditional machine learning techniques, such as Support Vector

Machines, Random Forest, and eXtreme Gradient Boosting. However, both studies were focused on lexical features without providing interpretable semantic features that signal fake news and comparing such characteristics of fake news across different languages. Dementieva and Panchenko (2020) proposed an approach to detecting fake news using multilingual evidence. Instead of using the original fake news in different languages, they matched English news titles to non-English ones using an online translator.

Although fake news has penetrated into every social media platform, the thematic characteristics of fake news that distinguish itself from real news remain largely under explored when they are compared across different languages and/or at different levels of granularity. Studies on the emotional characteristics of fake news also suffer from similar limitations despite that emotional appeals play a significant part in producing, proliferating, and promoting fake news (Bakir & McStay, 2018; Horner et al., 2021; Martel et al., 2020; Paschen, 2020). This research is aimed to address the above-mentioned limitations.

2.4 Research Questions and Hypothesis Development

Our literature review shows that, while topics have been employed as predictive features of fake news detection models, they have rarely been used to understand the unique characteristics of fake news. Based on an analysis of a dataset collected between March and May of 2020, Gupta et al. (2022) suggest that fake and real news have some major differences in their themes. They found that compared with real news, fake news had less coverage on the global impact of COVID-19 (e.g., global economy and oil price), and had exclusive coverage on the themes such as virus origin. Traditional news media, which are common sources for collecting real news, have preferences in topic coverage. For example, both BuzzFeed and The New York Times publish stories about governments or politics most frequently, followed by crime or terrorism, science, and health or technology (Tandoc, 2017). On the other hand, a systematic analysis of sample misinformation from a corpus of fact-checked claims about COVID-19 found that the topics of false claims included conspiracy theories, virus transmission, virus origins, public preparedness, and vaccine development (Brennen et al., 2020). To combat misinformation, Center for Disease Control and Prevention (CDC) has shared a number of posts containing misinformation and facts about COVID-19 vaccines on its website. There are also warnings from the U.S. Food and Drug Administration and the Federal Trade Commission that companies selling fraudulent products claim to be able to treat or prevent COVID-19 (FDA, 2020). Therefore, this study aims to answer the first question as follows:

RQ1. What are the thematic characteristics of COVID-19 related fake news that distinguish itself from real news?

In this study, we use two variables to characterize the overall topic distribution within a news article. One is topic concentration defined as the focus of topic coverage in a piece of news, and the other is topic uncertainty defined as the lack of clarity about the overall theme of a piece of news. Authoring fake news is “a highly structured process designed to play on human tendencies” (George et al., 2021, page 1071). Earlier research on online deception behavior suggests that deceivers may be motivated to establish trust relationships with their remote partners to make up for the deceivers’ lack of actual memory (Zhou et al., 2004). In addition, like fabricating online fake reviews (Zhang et al., 2016), one of the strategies for crafting fake news is to imitate the look and feel of real news. The analysis of 225 pieces of COVID-19 misinformation (Brennen et al., 2020) show that most of them (59%) involved various forms of content manipulation by twisting, recontextualizing, and reworking existing true information. In comparison, less misinformation (38%) is completely fabricated (i.e., using various actors, methods, and tactics to create fake news). The reconfigured misinformation accounts for 87% of social media interactions in the sample, which is in contrast with about 12% of fabricated content (Brennen et al., 2020). Dogo et al. (2020) suggest that the opening sentences of fake articles are topically deviated more from the rest of the articles as compared to real news. Thus, we predict that in general, fake news is likely to show greater thematic variation and uncertainty by covering a wide range of topics than real news and propose the first two hypotheses as the following:

- H1. Fake news has lower topic concentration than real news.
- H2. Fake news has higher topic uncertainty than real news.

Despite fake news being created to look like real news, one indicator of fake news lies in the presence of its author’s personal opinions or lack of objectivity (Tandoc Jr. et al., 2021). Emotions have consequences for public behaviors and opinions (Brader et al., 2011). Fake news is more opinion-based than real news (Gupta et al., 2022). Studies have shown that emotional processing may play a role in susceptibility to fake news. For instance, inducing reliance on emotion would result in greater belief in fake news stories compared to no emotion or inducing reliance on reasoning (Martel et al., 2020). Activating emotional reactions to either spread or suppress fake news might help fake news creators to achieve their business or political objectives (Horner et al., 2021). Therefore, we propose the following hypotheses:

H3. Fake news expresses a higher level of overall emotion than real news.

Emotion expressions range from negative to positive polar. In this study, we refer to emotional polarity as a spectrum of emotion state ranging from the negative extreme to the positive extreme. Recent studies have suggested that fake news articles are designed to induce inflammatory emotions in readers (Bakir & McStay, 2018). Paschen (2020) and Gupta et al. (2022) found that fake news was more negative than real news. Given that fake news intends to mislead others, authoring fake news can be considered as a type of deception. Deceiving induces a negative experience, in which deceivers may inadvertently leak their intention in their messages. Studies on online deception behavior (Zhou et al., 2004) suggest that deceivers might take a low-key approach due to the possible arousal of guilt by deception, demonstrating negative affect to disassociate themselves from their messages. Thus, we propose that.

H4. Fake news expresses a lower level of emotional polarity than real news.

H5. Fake news expresses a higher level of negative emotion than real news.

Negative emotion can be manifested in various discrete emotion states such as anger, sadness, and anxiety. Studies (Gupta et al., 2022; Paschen, 2020) have shown that fake news displays specific negative emotions such as anger more than real news. Yet the expression of anger in online news decreases its credibility because readers perceive such an expression in the headlines as a lack of cognitive effort of the author when writing the news (Deng & Chau, 2021). It is important to recognize that the expression of emotion is context-dependent. In case of the pandemic, anger could be triggered by various factors, such as the action or no action on implementing certain intervention policies.

A comparison of four emotions in COVID-19 related fake and real news found that sadness was among the most dominant types of emotion in both types of news (Gupta et al., 2022). The reporting of COVID-19 cases and related deaths is expected to evoke sadness. Moreover, an analysis of tweets shows that real tweets contain more sadness than fake stories (Vosoughi et al., 2018). In addition, the pandemic hangover and uncertainty with the evolving pandemic situation also arouses anxiety. Thus, we propose the following hypotheses about discrete emotions:

H6. Fake news expresses a higher level of anger than real news.

H7. Fake new expresses a lower level of a) sadness and b) anxiety than real news.

The Internet is multilingual, so is fake news. Interpersonal deception theory (Buller & Burgoon, 1996) suggests that deceivers are engaged in information, behavior, and image management. Building on the extensive research on deception behavior in face-to-face communication (DePaulo et al., 2003), online deception research (e.g., Zhou, 2005; Zhou et al., 2004) has witnessed significant progress over the past two decades. Moreover, deception behaviors are situated in relational and interactional context. Language is related to culture, which is a key context of communication. However, the number of studies on deception behavior in other languages is much fewer than that in English. A study on Chinese online deception behavior found that deceivers exhibited a tendency to use less complex and diversified texts in their messages (Zhou & Sung, 2008). Another study of Spanish speakers suggested that linguistic and psychological processes would be most effective for differentiating deceptive from true statements (Almela et al., 2012). Thus, we expect the thematic and emotional characteristics of fake news to differ between different languages.

On the other hand, it is possible to leverage information from one language, either through translation or equivalent semantic categories, to build deception classifiers for different languages (Pérez-Rosas & Mihalcea, 2014). According to a few recent studies on building fake news detection models across multiple languages (Abonizio et al., 2020; Faustini & Covões, 2020), similar features are shared among news in different languages. However, to the best of our knowledge, none of those studies has attempted to characterize the thematic and emotional features of fake news across different languages.

According to the economics of emotion theory (Bakir & McStay, 2018), fake news is created to evoke emotional responses in readers that will help fake news creators achieve business or other objectives by gaining readers' attention. In addition, the results of a user study show that "participants who reported high levels of emotions were more likely to take actions that would spread or suppress the fake news, participants who reported low levels of emotions were more likely to ignore or disengage from the spread of false news" (Horner et al., 2021, p. 1039). This is because emotion can reflect the strategies that fake news creators or algo-journalism designers choose to influence readers by either promoting or inhibiting the associated topics (Paschen, 2020). On the other hand, the emotion expressions can vary with the topics under discussion. The literature still lacks for such kind of understanding, not to mention a comparative analyses of news in different languages. Therefore, we propose the following set of research questions:

RQ2. How do the thematic characteristics of fake news in different languages compare?

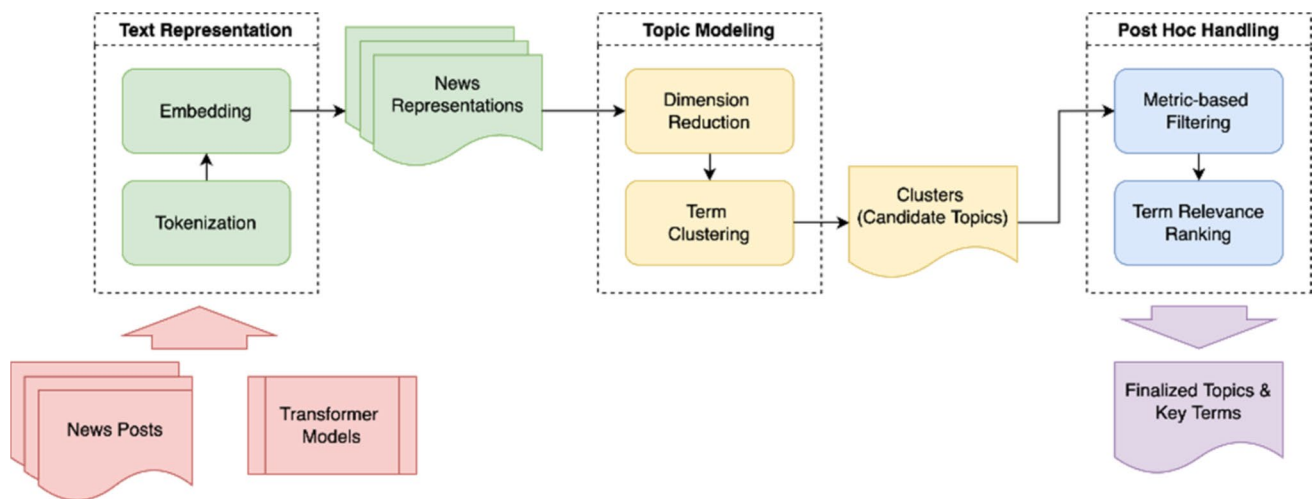


Fig. 1 Transformer-based topic modeling (TM²)

RQ3. How do the emotional characteristics of fake news in different languages compare?

RQ4. How do the associations of thematic and emotional characteristics of fake news in different languages compare?

3 Method

3.1 Datasets

Given the widespread and lasting impact of COVID-19, we chose it as the context of news. We collected English data across various sources, including COVID-19 Fake News Dataset (Patwa et al., 2021), NewsGuard Coronavirus Misinformation Tracking Center (NewsGuard, 2021), CODA-19 (Banik, 2020), Poynter CoronaVirusFacts (Poynter, 2021), CovidLies (Hossain et al., 2020), and AVax tweets dataset (Muric et al., 2021). Most of these datasets contain both real news and fake news. Additionally, we also collected English real news on COVID-19 from reputable sources, such as WHO and CDC websites and Twitter accounts of WHO, CDC, and the CDC director. We examined the selected news articles by verifying their relevance to COVID-19 and label accuracy. For the relevance to COVID-19, we first created embeddings of the news content using a transformer that was fine-tuned on COVID-19 texts, and then clustered the embeddings. The texts whose embeddings could not be clustered were from the datasets. Finally, we obtained an English dataset of 26,064 news, consisting of 13,778 (52.86%) fake news articles. Compared with English, Chinese fake news datasets on COVID-19 are rarer. We were able to identify a few quality Chinese datasets, such as CHECKED (Yang et al., 2021), Infodemic (Luo et al., 2021), and CrossFake (Du et al.,

2021). The final merged Chinese dataset contains 4,224 news articles, containing 3,512 (83.14%) fake news.

3.2 TransforMer-based Topic Modeling (TM²)

Topic extraction is accomplished through building topic models. Traditional topic modeling techniques (e.g., LDA, NMF) are limited in two main aspects. First, these techniques rely on the patterns extracted from the language space built with the specific text data used in the analysis, which may lead to poor generalizability and coverage. Second, the traditional techniques capture context as simple term co-occurrences (e.g., bag-of-words), which are difficult to capture the sentence-level context.

To address the aforementioned limitations, we proposed a topic modeling method (TM²) by extending a state-of-the-art transformer-based model, which is pre-trained on a large amount of generic text data, and uses the self-attention mechanism to capture the contextual information embedded in the text data. In addition, we also considered different transformer models. Figure 1 depicts the architecture of TM², which consists of three main components: text representation, topic modeling, and post-hoc handling.

- Text representation.** We selected Sentence-BERT (SBERT) as the embedding model (Reimers & Gurevych, 2020) in this study to represent news articles at the sentence/document level (rather than at the word level as word2vec does). We made the selection for two reasons. First, some SBERT models are pre-trained for natural language inference tasks, which are suitable for downstream tasks (e.g., topic modeling). Second, SBERT provides several native functions to compute semantic similarities, which makes our analysis more efficient.

To use transformer-based models like SBERT for text representation, the news articles went through preprocessing steps such as tokenization. It is worth noting that we represented each news article as an embedding vector in this study. The out-of-the-box SBERT tokenizers worked well for English texts, but not so well for Chinese texts. Thus, we employed a third-party tokenizer (jieba) for the Chinese news articles. After tokenization, news representation was learned via a chosen SBERT model. Although it is well recognized that transformer models in general need fine-tuning (retrained using the problem-specific data) to reach optimal performance, we found it not the case for topic modeling because the topics tend to contain too many non-meaningful words (e.g., stop words) after fine-tuning. Hence, we used the stock version of the models. In view that multiple SBERT models were available for English and Chinese texts, we designed several evaluation metrics, such as coherence, diversity, and average weighted F1 scores (See Section 3.3 for details) to select the most appropriate model for news representations.

- **Topic Modeling.** Topic modeling comprises dimension reduction and term clustering. Using the SBERT models, news articles were represented in a high-dimensional space (e.g., 768 dimensions). Given that most clustering algorithms do not work well with high-dimensional data, *dimension reduction* is deemed necessary. To this end, we selected the UMAP algorithm (McInnes & Healy, 2018) because it preserves the global structure of original new features and does not require extensive running time and computational restrictions. *Term clustering* groups news representations (i.e., vectors in a language space) into different clusters, or candidate topics. Since the topics in a language space can vary in terms of their levels of density, we selected HDBSCAN (Campello et al., 2013), a density-based clustering method, to identify candidate topics. Moreover, the technique supports the identification of the most important candidate topics with interpretable representations.
- **Post-hoc Handling.** This component mainly consists of two strategies for filtering candidate topics to select a final set of topics and their associated key terms: *metric-based filtering* and *term relevance ranking*. In metric-based filtering, we first measured the similarities among clusters (i.e., candidate topics). To this end, we developed topic-based tf*idf (T_tf*idf) by adapting c-tf*idf to topics, as shown in Eq. (1).

$$T_tf * idf = \frac{t_i}{w_i} \times \log \frac{m}{\sum_j^n t_j} \quad (1)$$

where t_i is the frequency of term t in topic i ; w_i is the total number of terms in i ; m is the average number of terms per topic, and $\sum_j^n t_j$ is the sum of frequency counts of t

across all n topics. Compared to the tf*idf metric, T_tf*idf is able to better measure the importance of each term to its associated topics collectively. Despite the similarity in terms of obtaining topic loadings between T_tf*idf and word-topic matrices in LDA, they have two key differences: 1) T_tf*idf measures the importance of a term to a certain topic rather than to the entire dataset, and 2) the calculation of the metric uses document embedding, which captures more contextual information in text compared to word2vec. *Term relevance ranking* aims to ensure that topic terms are coherent within a topic yet diverse across different topics. To maintain the coherence of all the terms belonging to the same topic, we introduced the topic optimization step by employing Maximal Marginal Relevance (MMR) (Carbonell & Goldstein, 1998), as shown in Eq. (2):

MMR =

$$\text{Arg max}_{D_i \in T} \left[\lambda \left(\text{sim}_1(w_i, T) - (1 - \lambda) \max_{D_j \notin T} (\text{sim}_2(w_i, w_j)) \right) \right] \quad (2)$$

where w_i and w_j denote a term in document D_i and D_j , respectively, and D_i contains topic T , but D_j does not. sim_1 measures the maximal pairwise similarity between w_i and all terms in topic T , and sim_2 measures the similarity between w_i and w_j . $\lambda \in [0, 1]$ is a constant. MMR (Xia et al., 2015) ranks query search results based on their relevance, which has been widely used to extract key phrases from text to improve the representativeness of the phrases for a document. In this study, we extended MMR to topic modeling to improve the representativeness of the selected terms for each topic, while promoting the diversity among different terms. For each topic, we selected the top- k terms ranked in a descending order of their MMR scores, where k was determined heuristically (Aletas & Stevenson, 2013). Compared to the coherence scores, which are widely used in LDA to measure the co-occurrence of topical words (i.e., words appearing in the same document), MMR reduces redundancy and therefore improves the diversity of terms under each topic. The final outputs of the post-hoc handling component contain final sets of topics and representative terms for each topic, as well as the probabilities of the topics for each news article in terms of topic loadings.

3.3 Evaluation Settings and Metrics for Topic Modeling Methods

We performed both intrinsic and extrinsic evaluations of the TM^2 model. The intrinsic evaluation metrics include coherence score (O'Callaghan et al., 2015) and diversity. We introduced diversity to measure intra-topic coverage (i.e., term coverage within a topic), which is different from other inter-topic diversity metrics (e.g., (Tran et al., 2013)). In view that topics have

been used as the input features for fake news detection (e.g., Xu et al., 2020), we operationalized the extrinsic metric as the performance of fake news detection using the extracted topics, specifically weighted average F1 score.

- **Coherence.** Traditional coherence-based metrics are calculated based on the co-occurrences of the top- l topic terms in the same document (e.g., news) (Mimno et al., 2011). Given that we used transformer models to represent news, it was a natural choice to measure the similarity between topic terms within the representation language space. We operationalized coherence as the average pairwise cosine similarity among the top- l terms of each topic. Specifically, we first derived the average pairwise cosine similarities by topic, and then averaged the similarity values across all the selected topics.
- **Diversity.** Diversity is operationalized as the average ratio of unique topic terms across all the topics. We define a unique term as one whose Jaccard distance to all other top- l terms of the same topic in the embedding space is larger, compared to the average pairwise Jaccard distance of all terms in the news datasets. A higher diversity value indicates a better topic coverage.
- **Average weighted F1 scores (average F1).** F1 is a harmonic mean of precision and recall of fake news detection models using the extracted topics as input features, where precision is defined as the percentage of detected fake (or real) news that is actually fake (or real), and recall as percentage of actual fake (or real) news that is correctly detected. The average F1 scores were yielded via fivefold cross validations, and weighted to avoid any bias from the imbalanced news classes. We adopted Tree-based Pipeline Optimization Tool (TPOT) to compute average F1. TPOT relies on the genetic algorithm to optimize various machine learning models (e.g., support vector machine and XGBoost). It considers different models in each generation, and selects the best models for the subsequent generation. TPOT optimizes the models via a sequence of tasks, including pre-processing, feature engineering, and hyperparameter optimization.

We applied the elbow method to the above three metrics to determine an optimal number of topics to represent the English and Chinese news articles separately. In addition, we compared different types of transformer models to select the best one for fake news detection. We also considered XGBoost as a baseline classifier because it demonstrated a superior performance to other models in a previous study (Lin et al., 2019). For all the models, we first extracted bi- and tri-grams from the news and represented them with their T_{tf*idf} scores, then performed hyperparameter tuning, and finally determined the number of topics and reported the performances of the selected models.

3.4 Variables and Measurements

We measured each of the top- k topics extracted from news article n with its topic loading, which was defined as the probability of n discussing the corresponding topic. Additionally, we also selected variables to measure the news thematic characteristics at the topic pair and news levels.

To support the extraction of topic pairs, we first selected the top-5 topics of news article n based on a descending order of their ranking within n . Incorporating the ranking enables the analysis to focus on the most dominant topics in the news. For two news articles that contain the same set of topics, their rankings could be different. Additionally, the number of topics was limited to 5 because news articles in our datasets are relatively short. Then, we used a moving window of size 2 to extract topic pairs from n . Inspired by an early study (Ito et al., 2015), we employed Kullback–Leibler divergence (AlSumait et al., 2009) (referred to as KLD hereafter) as a measure for topic pairs, $KLD(R_p \parallel F_p)$, where R_p and F_p denote the distributions of topic pair p in real news and fake news, respectively. The scores of KLD range from 0 to infinitely large, with 0 indicating that the two distributions are exactly the same, and a higher value indicating a greater distance between the two distributions of the topic pair. In the context of this research, the topic pairs with higher KLD scores are expected to have stronger discriminatory power for news credibility. In light of the unbalanced sample sizes between real and fake news, we performed repeated random sampling on the majority class for N times to derive the KLD score for each topic pair, where N was set to 1,000.

At the news level, we introduced two variables: topic concentration and topic uncertainty. We operationalized *topic concentration* with Herfindahl–Hirschman Index (HHI) (Rhoades, 1995), as shown in Eq. (3). The coefficient of HHI reaches the maximum value when a news article is concentrated on a single topic, or the minimum value when the topic loadings are perfectly evenly distributed across different topics. We operationalized *topic uncertainty* with Shannon’s entropy (Shannon, 1948), as shown in Eq. (4). Equation (5) shows the calculation of the probability of topic i in news article n . Generally, a uniform probability distribution of different topics yields the maximum entropy, whereas a skewed probability distribution toward a single topic yields the minimum entropy.

$$HHI(n) = - \sum_{i=1}^m (p(t_{in}))^2 \quad (3)$$

$$Entropy(n) = - \sum_{i=1}^m p(t_{in}) \log_2 p(t_{in}) \quad (4)$$

$$p(t_{in}) = \frac{t_{in}}{\sum_{i=1}^k t_{in}} \quad (5)$$

Table 1 Performances of different embedding and classification models

Embedding models	Diversity	Average F1	
		Proposed	Baseline
(a) English news			
all-MiniLM-L12-v2	.7834	.9064	.8641
all-distilroberta-v1	.7767	.8864	.8214
all-mpnet-base-v2	.7653	.8812	.8199
(b) Chinese news			
paraphrase-multilingual-MiniLM-L12-v2	.7962	.8824	.7625
distiluse-base-multilingual-cased-v1	.7321	.8801	.7832
distiluse-base-multilingual-cased-v2	.7415	.8695	.7346

Values in bold denote the results of the best performing models.

where t_{in} denotes the loading of topic i for news article n ; m is the total number of selected topics from the news dataset, and k is the number of selected topics for each news article.

We measured emotion at two different levels: overall emotion and discrete emotions. The variables for measuring the overall emotion include *overall emotion*, *negative emotion*, and *emotional polarity*. The measures for the discrete emotions consist of three variables—sadness, anxiety, and anger. We used TextBlob (Loria, 2020) to measure the emotional polarity of English news, and *cn-sentiment-measures*, specifically absolute proportional difference (Chen, 2021), to measure Chinese news. Both of their values are in the range of $[-1, 1]$, with 1 being extremely positive and -1 extremely negative. The values of other variables were derived from the outputs of LIWC (Linguistic Inquiry and Word Count) 2015 (Pennebaker et al., 2007). The tool has been widely used to analyze social media data and has been extended from English to many other languages such as Chinese. The measures of these variables were normalized by the total word count in the news.

We performed independent-samples t-tests to test the research hypotheses and answer the proposed research questions about the thematic and emotional characteristics of fake news. We also performed Pearson's correlation analysis between topic loadings and emotional polarity to understand the associations between the thematic and emotional characteristics of fake news.

4 Results

We first report the results of topic model comparisons, and present the topics extracted from English and Chinese news separately. Then, we report the statistical analysis results to answer the research questions and test the hypotheses.

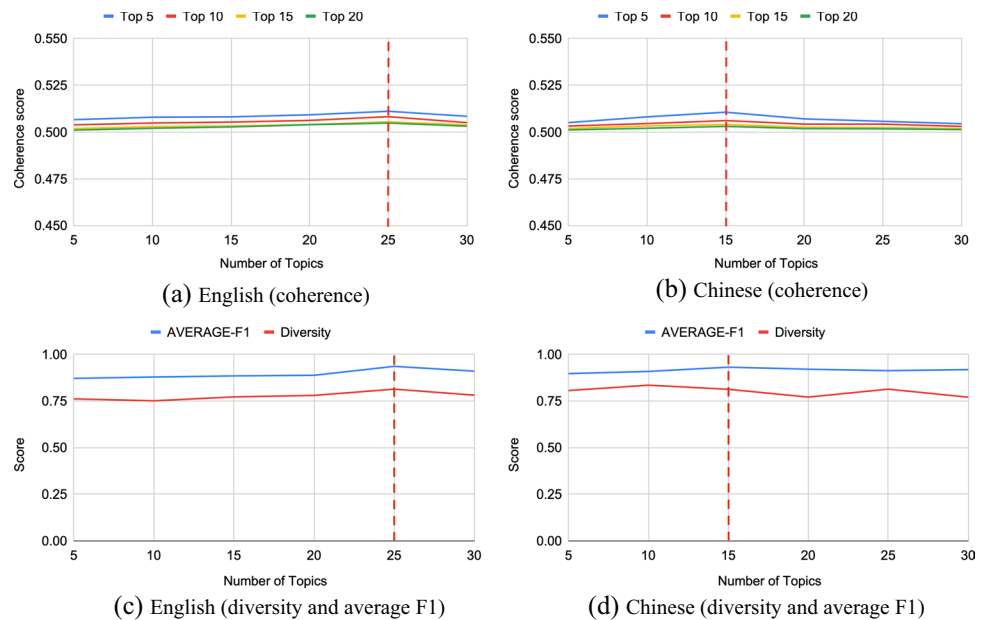
4.1 Comparison of Different Topic Modeling Methods

We implemented the proposed TM² framework using a Python package called BERTopic (Grootendorst, 2020), which relies on the embedding models from another study (Reimers & Gurevych, 2020). In search for the best embedding model, we selected several SBERT models (pretrained for the semantic search task), including *all-mpnet-base-v2*, *all-distilroberta-v1*, and *all-MiniLM-L12-v2* for the English news, and *paraphrase-multilingual-MiniLM-L12-v2*, *distiluse-base-multilingual-cased-v1*, and *distiluse-base-multilingual-cased-v2* for the Chinese news. The diversity and average-f1 measures of different embedding models are reported in Table 1. The best performances for each metric are highlighted in bold in the table, with the average F1 being over 90.6% for the English news and over 88.2% for the Chinese news. The embedding model that produced the best performances in terms of both average F1 and diversity was *all-MiniLM-L12-v2* for the English news, and the embedding model that produced the best performances was *paraphrase-multilingual-MiniLM-L12-v2* for the Chinese news. Our proposed model outperforms the baseline model in terms of average F1 by 4.23% for the English news and outperforms the baseline model by approximately 12% for the Chinese news, respectively. These results demonstrate the effectiveness of TM².

The results of applying the elbow method for selecting the optimal number of topics for the news datasets are shown in Fig. 2. Based on the results, we selected 25 topics for the English news and 15 topics for the Chinese news. The extracted topics from the entire news datasets in each language are listed in Table 2.

To illustrate the extracted topics and demonstrate their quality, we list three sample topics of English and Chinese fake news, respectively, along with their top 10-terms in Table 3.

Fig. 2 Elbow method plot for determining an optimal number of topics



4.2 Comparison of Individual Topics between Fake and Real News

As a part of the effort to answer RQ1, we compared the differences between the topics of fake and real news. The descriptive statistics of topic loadings of the extracted topics, and the results of independent-samples t-tests are reported in Table 2.

Table 2(a) shows that the fake and real English news differ across all the topics ($p < 0.001$) except for *travel*. Specifically, some topics have more coverage in fake news than real news, such as *getting vaccinated*, *depopulation*, *unvaccinated people*, *virus origin*, *pandemic in Italy*, *US President*, *back to school*, *chicken contamination*, *India lockdown*, and *Bill Gates*. In contrast, some other topics have more coverage in the real news than in the fake news, such as *pandemic*, *vaccinequity*, *face mask*, *deaths*, *cases*, *UK lockdown*, *India fighting COVID-19*, *COVID-19 vaccine*, *flu vaccination*, *protective measures*, *cases in Nigeria*, *covidview report*, *case updates*, and *testing*.

Table 2(b) shows that most of the topics differ significantly between the fake and real Chinese news. Specifically, some topics have a higher level of presence in the fake than real news ($p < 0.001$), such as *viral infection*, *protective measures*, *pandemic in US*, and *preventative measures*. Some other topics have a higher level of presence in the real news than in the fake news, such as *deaths*, *confirmed cases*, *fighting COVID-19*, *PCR testing*, and *pandemic in specific countries* (all at $p < 0.001$), as well as *prevention and control* ($p < 0.01$) and *waste water testing* ($p < 0.05$). However, no difference was detected in some topics, such as *Shenshan Hospital* and *vaccine* ($p > 0.05$), between fake and real Chinese news.

4.3 Comparison of Topic Pairs between Fake and Real News

As another part of the effort to answer RQ1, we extracted topic pairs and compared them between fake and real news. Figure 3 lists top-25 topic pairs from the English and Chinese news ranked in a descending order of their KLD scores. They are sensical to human interpretations.

Figure 3a shows that topic pairs, such as (*pandemic*, *deaths*), (*pandemic*, *COVID-19 vaccine*), (*US President*, *face mask*) can best discriminate between fake and real English news. In addition, combining *pandemic* (a real news topic) with another topic, such as *deaths*, *COVID-19 vaccine*, *pandemic in Italy*, *UK lockdown*, *cases*, and *depopulation* can help discriminate fake from real news in English. Similar findings are also observed for *COVID-19 vaccine*, another real news topic. It is interesting to observe from the selected topic pairs that they are a combination of fake news topics and real news topics (those with a significantly higher coverage in real news than fake news) with only a few exceptions.

Figure 3b shows that (*confirmed cases*, *viral infection*), (*Shenshan hospital*, *viral infection*), (*deaths*, *confirmed cases*) are the top-3 topic pairs that help discriminate between fake and real news in Chinese. In addition, *viral infection* — a fake news topic, accounts for a significant percentage of the selected topic pairs when combined with another topic, such as *confirmed cases*, *Shenshan hospital*, *deaths*, *prevention & control*, *pandemic in Japan*, *pandemic in Spain*, *COVID-19 vaccine*, and *protective measures*. Similar findings hold for *confirmed cases* (a real news topic). Furthermore, we observe from the selected topic pairs that a combination of fake and real news topics and a combination of two real news topics are the most common types of topic pairs.

Table 2 Descriptive statistics (mean [std]) and t-test results of the extracted topics

ID	Topic	Fake news	Real news	<i>t</i> -value (F-R)	<i>p</i> -value
(a) English					
E1	getting vaccinated	0.043 [0.140]	0.005 [0.022]	31.852	< .001***
E2	pandemic	0.006 [0.031]	0.022 [0.104]	−16.120	< .001***
E3	Vaccinequity	0.015 [0.023]	0.039 [0.126]	−21.220	< .001***
E4	face mask	0.010 [0.050]	0.018 [0.092]	−8.808	< .001***
E5	depopulation	0.016 [0.083]	0.003 [0.003]	19.082	< .001***
E6	deaths	0.015 [0.043]	0.036 [0.121]	−18.508	< .001***
E7	unvaccinated people	0.038 [0.106]	0.011 [0.019]	29.878	< .001***
E8	cases	0.016 [0.027]	0.048 [0.147]	−23.524	< .001***
E9	virus origin	0.025 [0.106]	0.007 [0.022]	18.820	< .001***
E10	pandemic in Italy	0.028 [0.104]	0.009 [0.019]	20.111	< .001***
E11	U.S. President	0.021 [0.087]	0.010 [0.024]	14.279	< .001***
E12	UK lockdown	0.004 [0.033]	0.012 [0.087]	−9.244	< .001***
E13	India fighting COVID-19	0.009 [0.010]	0.024 [0.107]	−15.314	< .001***
E14	COVID-19 vaccine	0.017 [0.050]	0.025 [0.092]	−7.975	< .001***
E15	flu vaccine	0.010 [0.045]	0.016 [0.075]	−7.674	< .001***
E16	protective measures	0.010 [0.028]	0.024 [0.098]	−15.621	< .001***
E17	cases in Nigeria	0.005 [0.004]	0.018 [0.102]	−14.349	< .001***
E18	Covidview report	0.011 [0.042]	0.020 [0.081]	−10.565	< .001***
E19	back to school	0.009 [0.054]	0.005 [0.024]	7.948	< .001***
E20	chicken contamination	0.008 [0.054]	0.004 [0.018]	8.313	< .001***
E21	India lockdown	0.017 [0.081]	0.006 [0.015]	15.216	< .001***
E22	Bill Gates	0.009 [0.063]	0.001 [0.003]	14.158	< .001***
E23	case updates	0.006 [0.029]	0.014 [0.093]	−9.531	< .001***
E24	testing	0.005 [0.031]	0.007 [0.048]	−3.999	< .001***
E25	travel	0.010 [0.045]	0.009 [0.056]	1.534	0.125
(b) Chinese					
C1	Deaths	0.033 [0.065]	0.058 [0.171]	−6.615	< .001***
C2	Shenshan hospital	0.070 [0.175]	0.068 [0.160]	0.350	0.727
C3	prevention & control	0.017 [0.054]	0.025 [0.095]	−3.299	0.001**
C4	confirmed cases	0.025 [0.030]	0.057 [0.159]	−11.105	< .001***
C5	viral infection	0.039 [0.114]	0.020 [0.102]	4.117	< .001***
C6	protective measures	0.048 [0.106]	0.030 [0.123]	4.067	< .001***
C7	pandemic in U.S	0.049 [0.161]	0.019 [0.068]	4.930	< .001***
C8	fighting COVID-19	0.017 [0.035]	0.028 [0.105]	−4.850	< .001***
C9	COVID-19 vaccine	0.026 [0.067]	0.022 [0.102]	1.325	0.186
C10	pandemic in Japan	0.011 [0.008]	0.030 [0.126]	−8.869	< .001***
C11	British Prime Minister	0.016 [0.016]	0.029 [0.123]	−6.373	< .001***
C12	preventative measures	0.055 [0.147]	0.007 [0.038]	8.658	< .001***
C13	wastewater surveillance	0.019 [0.097]	0.030 [0.135]	−2.538	0.011
C14	PCR testing	0.006 [0.006]	0.014 [0.072]	−6.463	< .001***
C15	pandemic in Spain	0.002 [0.002]	0.020 [0.118]	−8.768	< .001***

***: $p < .001$; **: $p < .01$

Topics in bold denote overlapping topics between news in the two languages.

4.4 Comparison of News-Level Topic Characteristics between Fake and Real News

To test hypotheses H1 and H2, we analyzed news-level topic characteristics in terms of topic concentration and

uncertainty between fake and real news. The descriptive statistics and t-test results of the overall topic characteristics are reported in Table 4.

The table shows that English fake news has a lower level topic concentration ($p < 0.001$) yet a higher level

Table 3 Sample topics and their top-10 terms

ID	Top-10 terms	Topic
(a) English news		
E22	gates, bill, bill gates, depopulation, vaccine, billgatesbioterrorist, his, he, foundation, gates foundation	Bill Gates
E14	vaccine, vaccines, covid, covid vaccine, more, covid vaccines, mrna, dose, vaccinated, get	COVID-19 vaccine
E9	China, Chinese, coronavirus, Wuhan, in Wuhan, in China, the coronavirus, virus, video, bats	Virus origin
(b) Chinese news		
C2	神山(Shenshan), 医院(Hospital), 一线(frontline), 神山 医院 (Shenshan hospital), 武汉(Wuhan), 医疗 (medical), 英雄 (hero), 医疗队(medical team), 军舰 (naval ship), 湖北 (Hubei)	Shenshan hospital
C6	口罩(face mask), 消毒剂(sanitizer), 消毒(sanitizing), 洗手 (handwashing), 佩戴(wear), 清洗(wash), 医用 口罩 (medical mask), 接触(contact), 防护(protection), 酒精 (alcohol)	Protective measures
C9	疫苗(vaccine), 病毒(coronavirus), 研究(research), 抗体 (antibody), 蛋白(protein), 突变(mutation), 研发(R&D), 临床试验 (clinical trial), 灭活疫苗 (inactivated vaccine), 检测 (testing)	COVID-19 vaccine

of topic uncertainty than its real-news counterpart ($p < 0.001$). The same findings also hold for the Chinese news. Thus, hypotheses H1 and H2 are supported.

4.5 Comparison of Emotional Characteristics between Fake and Real News

The descriptive statistics of the overall and discrete emotion features of fake and real news are reported in Table 5. The results of independent-samples t-tests are reported in Table 6.

The test results show that the overall emotion ($p < 0.001$) and negative emotion ($p < 0.001$) in fake news are higher than those in real news, and the emotional polarity in fake news is lower than that in real news for both English and Chinese ($p < 0.001$). In addition, our examination of the three discrete emotions shows that there is a higher level of anger ($p < 0.001$), yet a lower level of sadness ($p < 0.001$) and anxiety ($p < 0.01$), in English fake news than in real news. In addition, the analysis of Chinese news yields a higher level of anger ($p < 0.05$) and sadness ($p < 0.05$) in fake news than in real news, yet no difference in anxiety ($p > 0.05$) between the two types of news. Thus, hypotheses H3 ~ H6 are supported, and H7 is partly supported.

4.6 Associations of Fake News Topics with Emotion

To answer RQ4, we analyzed the correlations between the loadings of fake news topics and their emotional polarity. Among the English fake news topics (those having a significantly higher coverage in fake news than in real ones), some have positive associations with emotional polarity, such as *getting vaccinated*, *depopulation*, *back to school*, and *India lockdown* ($p < 0.001$), while some others have negative associations with emotional polarity, such as *Bill Gates* ($p < 0.001$), *unvaccinated people*, and *pandemic in Italy* ($p < 0.05$). Among the Chinese fake news topics, such

as *viral infection* and *preventative measures*, they all have negative correlations with emotional polarity.

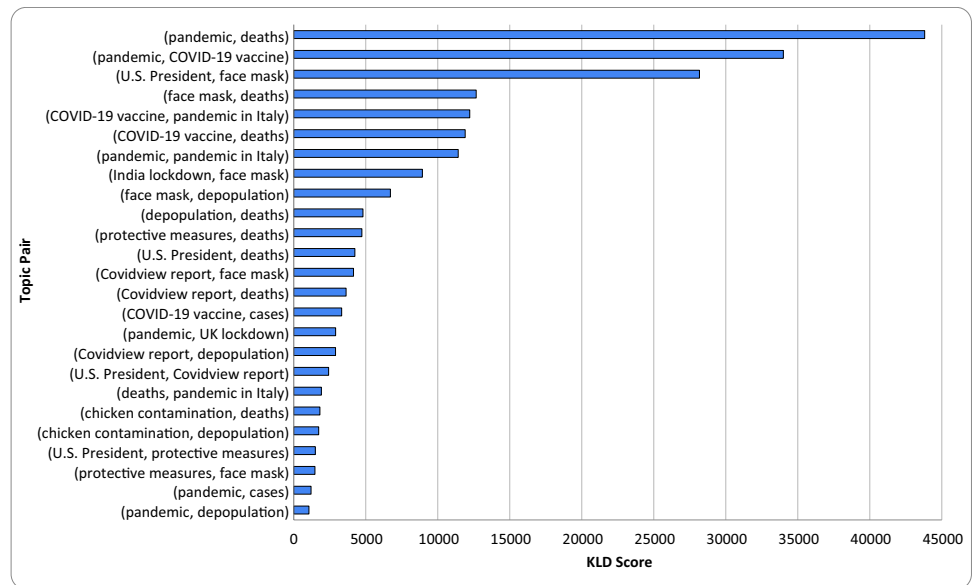
We further analyzed the correlations of the fake news topics for fake and real news separately and plotted their correlation coefficients in Fig. 4. In the English fake news, *US President* is positively, and *Bill Gates* is negatively, associated with emotional polarity. However, *virus origin* and *chicken contamination* does not show any correlation with polarity. The results on Chinese news show that *pandemic in US* is positively associated with emotional polarity ($p < 0.05$) only in the fake news, and *viral infection* and *preventative measures* are negatively associated with polarity ($p < 0.001$) only in real news. However, there is no correlation between *protective measures* and polarity ($p > 0.05$). Interestingly, the correlations of some topics with emotional polarity are in opposite directions between fake and real news. For instance, *cases in Nigeria* has a negative correlation with emotional polarity in English fake news but a positive correlation in English real news.

4.7 Comparison between English and Chinese News

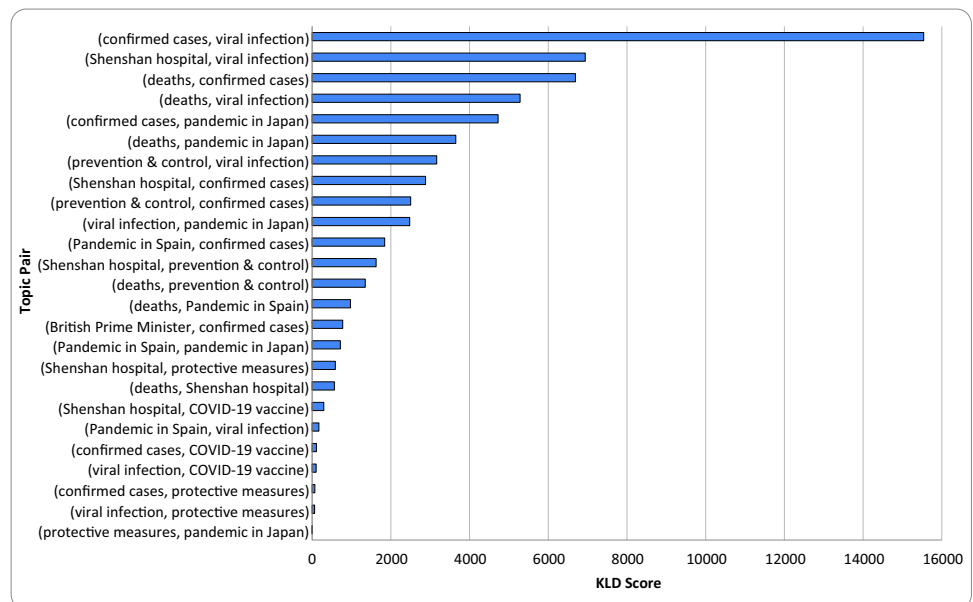
Based on the separate results for English and Chinese news as reported earlier, we draw a comparison between the two languages to answer research questions RQ2 ~ RQ4.

Cross-referencing the list of extracted topics between English and Chinese news in Table 2a and b reveals a number of overlapping topics between the two languages, such as *deaths*, *protective measures* and *COVID-19 vaccine*, which are highlighted in bold. In addition, some other topics covered by the news in one language are very similar to those in the other language, despite that the topics have different levels of specificity in the news of different languages. For instance, English news covers topics of *cases*, *testing*, *India fighting COVID-19*, and the corresponding topics in Chinese news are *confirmed cases*, *PCR testing*, and *fighting COVID-19*. Among them, the

Fig. 3 KLD scores of topic pairs



(a) English news



(b) Chinese news

effects of some topics, such as *deaths*, (*confirmed*) *cases*, (*PCR*) *testing*, and (*India*) *fight COVID-19*, are consistent between the languages. However, the effects of other topics are inconsistent and even show opposite directions between the two languages. For instance, *protective measures* has a higher coverage in English real news than in fake news ($p < 0.001$), yet has a higher coverage in Chinese fake news than the real news ($p < 0.001$). *COVID-19 vaccine* occurred more frequently in English real news than in fake news ($p < 0.001$), yet it does not show any difference between Chinese real and fake news ($p > 0.05$). In addition, it is worth noting that none of the overlapping topics

has a higher coverage in fake news than in real news. In other words, news in the two languages is unlikely to share the same fake news stories.

In view that individual fake news topics are divergent between English and Chinese news, we compared the topic pairs in terms of the overall distributions of their KLD scores instead of the topic pairs themselves between English and Chinese news. To this end, we performed log-transformations of the KLD scores and sorted the topic pairs in a descending order of their $\text{Log}(\text{KLD})$ values. The plot of $\text{Log}(\text{KLD})$ -order of topic pairs is depicted in Fig. 5. We performed Kolmogorov–Smirnov test to compare the distribution of KLD scores

Table 4 Descriptive statistics (mean [std]) and T-test results of news-level topic characteristics

Language	Variable	Fake news	Real news	<i>t</i> -value (F-R)	<i>p</i> -value
English	topic uncertainty	1.377 [0.898]	1.276 [0.887]	9.112	< .001***
	topic concentration	0.195 [0.262]	0.252 [0.316]	−15.763	< .001***
Chinese	topic uncertainty	1.335 [0.859]	1.069 [0.759]	7.693	< .001***
	topic concentration	0.273 [0.264]	0.370 [0.325]	−8.557	< .001***

Table 5 Descriptive statistics of emotion features

Variable	English		Chinese	
	Fake news	Real news	Fake news	Real news
overall emotion	3.540 [4.548]	2.924 [3.481]	6.707 [7.141]	4.703 [4.052]
emotional polarity	0.039 [0.217]	0.127 [0.198]	0.265 [0.348]	0.366 [0.239]
negative emotion	1.744 [3.293]	1.158 [2.321]	2.736 [4.222]	2.013 [2.550]
anger	0.701 [2.283]	0.156 [0.796]	0.484 [1.832]	0.222 [0.918]
sad	0.242 [1.159]	0.302 [1.067]	0.250 [1.446]	0.135 [0.722]
anxiety	0.310 [1.271]	0.361 [1.201]	0.402 [1.893]	0.288 [0.902]

Table 6 T-test results of emotion features between fake and real news

Variable	English		Chinese	
	<i>t</i> value (F-R)	<i>p</i> value	<i>t</i> value (F-R)	<i>p</i> value
overall emotion	12.335	< .001***	7.255	< .001***
emotional polarity	−34.391	< .001***	−8.670	< .001***
negative emotion	16.737	< .001***	4.412	< .001***
anger	26.276	< .001***	3.719	< .001***
sad	−4.350	< .001***	2.064	0.039*
anxiety	−3.316	0.001**	1.573	0.116

***: $p < .001$; **: $p < .01$; *: $p < .05$

of topic pairs between English and Chinese news. The analysis did not yield a significant results ($p > 0.05$), suggesting that the two distributions are identical.

Our analyses of the news-level thematic features (see Section 4.4 for test results) between English and Chinese news suggest that they are identical in terms of topic concentration and uncertainty.

The comparison of overall emotion features reveals identical patterns with respect to fake news between the two languages. When drilling down to discrete emotions; however, the comparisons between English and Chinese news yield mixed findings. Like English, Chinese fake news shows a higher level of anger than real news. Unlike the English news, however, the Chinese fake news shows a higher level of sadness than real news, but no difference in anxiety between fake and real news.

Our comparisons of the emotion polarity associated with fake news topics suggest that polarity of English fake news

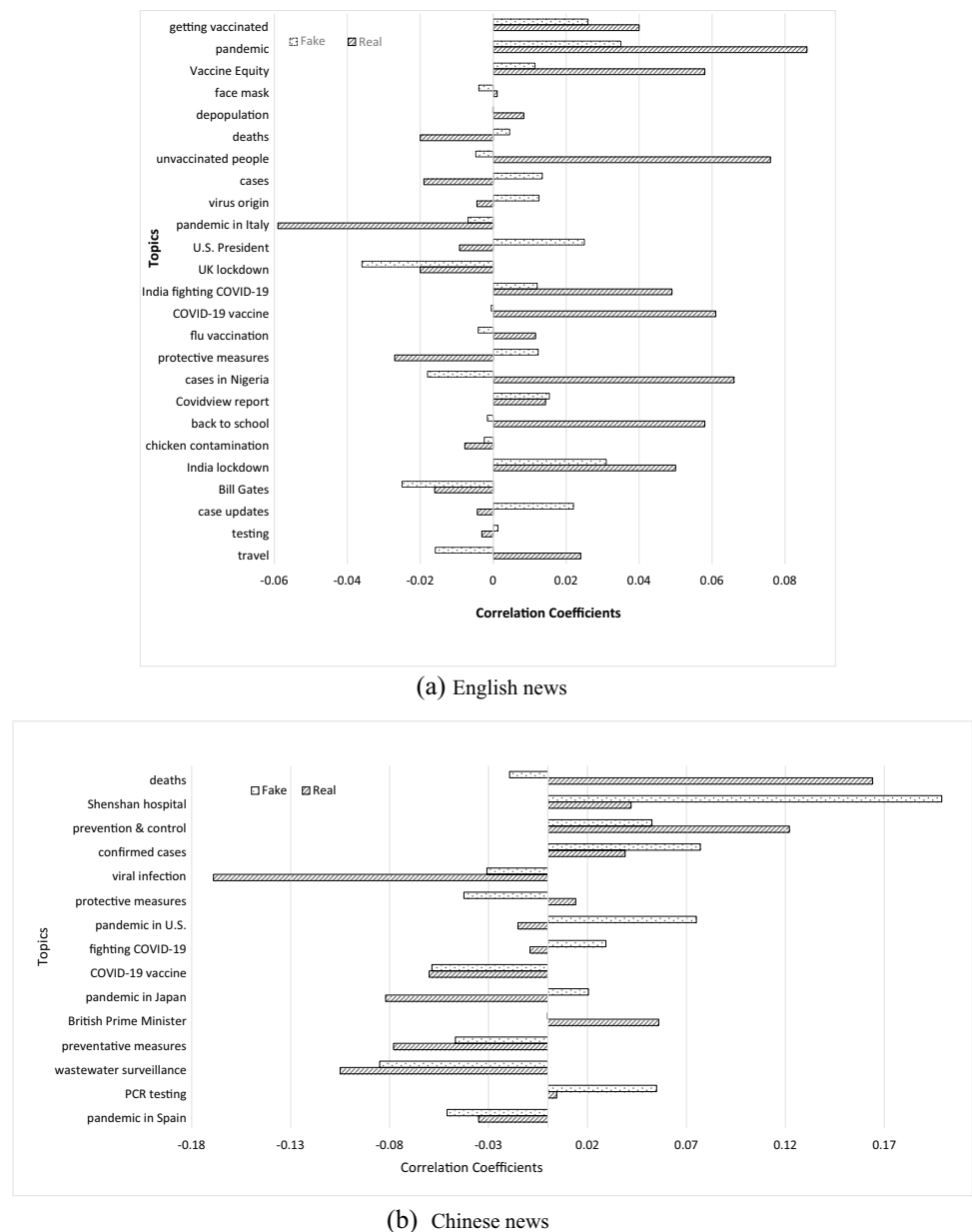
topics in the entire news datasets is generally divided, while that of Chinese fake news topics tends to have a negative orientation. If we limit the analysis to fake news only, the same pattern repeats for the English news, whereas a positive orientation starts to emerge for the Chinese news. They suggest that English fake news may adopt either a promoting or demoting framing strategies, while Chinese fake news mainly focuses on the promoting strategy to influence others' opinions.

5 Discussion

To explore the thematic and emotion patterns of fake news across different languages, we analyzed and compared English and Chinese fake news in the context of COVID-19 pandemic by answering four research questions and testing seven research hypotheses. Based on our empirical results, we obtain the following main findings. First, fake and real news differs in thematic characteristics. We identify a number of topics that have higher coverage in fake news than real news or vice versa in individual languages. Second, fake news has lower topic concentration but higher topic uncertainty than real news in both languages. Third, the overall emotion, negative emotion, and anger is higher in fake news than their real counterparts, and emotional polarity is lower in fake news than real news, irrespective of language. Fourth, there are cross-language differences in fake news topics, a few discrete emotions of fake news, and the associations between fake news topics and emotional polarity.

We make the following observations from the different thematic characteristics of fake news between the selected

Fig. 4 Correlation coefficients between topic loadings and emotional polarity



languages. English fake news involves a number of conspiracy theories such as *depopulation*, *virus origin*, *Bill Gates* (e.g., microchip the world through vaccination), and *chicken contamination* (leading to COVID-19 related deaths). Chinese fake news, in contrast, tends to focus on virus transmission and measures for protection and prevention of COVID-19. In addition, some of the topics reflect region- or country-specific themes, such as the topic of *US President* in English fake news and *Shenshan Hospital* in Chinese news. Further, news about *getting vaccinated* and *back to school* have been popular in countries like the U.S., which might have led to multiple topics related to vaccination in English fake news. However, they did not appear to be controversial topics in China.

The findings on sadness and anxiety of the Chinese fake news are unexpected. Contrary to the prediction, Chinese fake news shows a higher level of sadness than real news. This could be attributable to the governments' responses to COVID-19 and implementation of strict prevention and control policies since the pandemic outbreak, which quickly kept the number of cases under control (Sun et al., 2020). As a result, sadness was not the overtone of real news in China but was exploited by fake news to mislead the public. Along a similar vein, despite the multiple waves of COVID-19 variants up to the time of our data collection, it did not arouse as much anxiety in China as in some other countries.

This study makes multifold research contributions. First, this research identifies the thematic patterns that

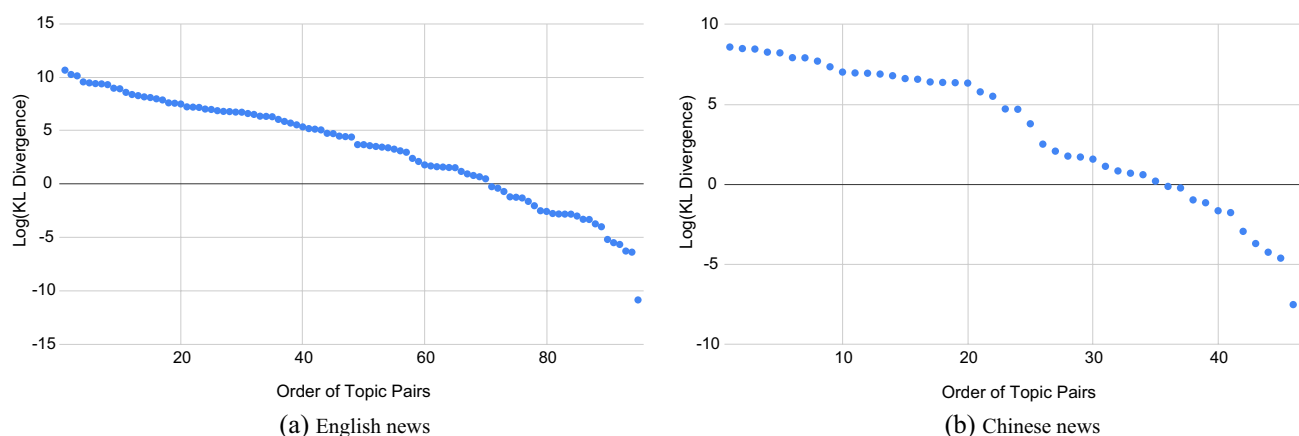


Fig. 5 Log(KLD) distribution plot

help distinguish fake news from real news in two different languages for the first time. Second, this is the first study that compares the emotional characteristics of fake news in a multi-lingual setting. Third, it characterizes the themes of fake news at multiple levels, including individual topics, topic pairs, and news level, and characterizes the emotions expressed in fake news at overall, polarity, and discrete emotion levels. Fourth, it extends a transformer-based topic modeling method to extract topics from news and designs extrinsic evaluation metrics for the performance of topic models. The empirical results demonstrate a superior performance of our proposed topic modeling method. Last but not the least, we analyze the associations between the thematic and emotional characteristics of fake news, which have implications for understanding fake news strategies.

The research findings have significant implications for fake news detection and multilingual analysis of social media data. On the one hand, the overall thematic and emotional characteristics, such as topic concentration and uncertainty and overall and negation emotions, appear to be generalizable across different languages. Thus, it is possible to leverage the common characteristics of fake news to build a cross-lingual component of an automated fake news detection model. These patterns may be considered as design guidelines for cross-lingual fake news detection, as well as countermeasures for enhancing the performance of existing detection models. On the other hand, specific themes and emotions indicative of fake news may not transfer across different languages, suggesting that fake news detection models should also incorporate targeted and language-dependent characteristics of fake news to improve their effectiveness for individual languages.

Our proposed framework for extracting topics from multilingual news has been applied to both English and Chinese news in this study. Our comparisons of different transformer models for topic modeling suggest the most effective

models. The framework is generalizable, and the selected models can be used to extract topics from text in other languages. In addition, our proposed multi-level metrics for thematic and emotional characteristics of fake news can be used to support analyzing other types of text.

We acknowledge several limitations of the study, which may present future research opportunities. First, since there are many English-speaking countries, a finer-grained analysis of fake news (e.g., at the country level with the help of geo-tagged data) may provide better explanations for the observed thematic and emotion characteristics. In addition, it would be interesting to test the generality of our findings about the characteristics of fake news to other languages such as Spanish. Second, in view of the complementary nature of the multi-level semantic characteristics of fake news to the lexical and stylometric features of fake news commonly used in the literature, combining them is expected to boost the performance of fake news detection. Adopting data resampling strategies to balance the fake and real news might contribute to a further improvement of model performances. Third, given that fake news datasets are more readily available in English, it would be promising to leverage fake news datasets and knowledge in a high-resource language to facilitate developing models for fake news detection in another low-resource language using transfer learning. Fourth, the fake news topics identified in this study are limited to the timeframe of our data collection. In view that fine-grained fake news thematic characteristics such as individual fake news topics will most likely evolve over time, it would be interesting to identify the trends of fake news themes through analyzing news content along the temporal dimension. Last but not least, we did not differentiate the emotions expressed across different topics within a single news in this study. Future studies may consider topic-based sentiment analysis to measure the emotions associated with individual topics more precisely.

6 Conclusion

This study characterizes, extracts, and compares the themes and emotions of fake news between English and Chinese news in the context of COVID-19. Moreover, it examines the thematic and emotional characteristics of fake news at multiple levels. Our empirical results reveal that the coarse-grained fake news characteristics are consistent across different languages while the fine-grained fake news characteristics differ significantly between the two different languages. The findings have implications both for enhancing the performance of general fake news detection models and countermeasures and for developing cross-lingual fake news detection models. In addition, the findings of this study contribute to gaining a deeper understanding of the strategies for creating fake news. Furthermore, our proposed topic modeling method and variables for measuring thematic and emotional characteristics of fake news can be extended to support other text analytics tasks.

Funding This work was partly supported by the National Science Foundation [Award #s: CNS 1917537 and SES 1912898] and the School of Data Science and Belk College of Business at UNC Charlotte. Any opinions, findings, and conclusions, or recommendations expressed in this paper are those of the authors and do not necessarily reflect the views of the above funding agency.

Data Availability The data that support the findings of this study are partly available from the third party as listed in Section 3.1 of the manuscript. The remaining part of the datasets are not publicly available due to the privacy policies of the data sources. Information on how to obtain the data is available from the authors on request.

Declarations

Competing Interests Lina Zhou is on the Executive Group of Information Systems Frontier and receives no compensation as a member of the group.

References

- Abonizio, H. Q., de Moraes, J. I., Tavares, G. M., & Barbon Junior, S. (2020). Language-independent fake news detection: English, Portuguese, and Spanish mutual features. *Future Internet*, 12(5), 87. <https://doi.org/10.3390/fi12050087>
- Akhter, M. P., Zheng, J., Afzal, F., Lin, H., Riaz, S., & Mehmood, A. (2021). Supervised ensemble learning methods towards automatically filtering Urdu fake news within social media. *PeerJ Computer Science*, 7, e425. <https://doi.org/10.7717/peerj-cs.425>
- Al-Ash, H. S., Putri, M. F., Mursanto, P., & Bustamam, A. (2019). Ensemble learning approach on Indonesian fake news classification. *3rd International Conference on Informatics and Computational Sciences (ICICoS)* (pp. 1–6). <https://doi.org/10.1109/ICICoS48119.2019.8982409>
- Aletras, N., & Stevenson, M. (2013). Evaluating topic coherence using distributional semantics. *Proceedings of the 10th international conference on computational semantics (IWCS)* (pp. 13–22). Potsdam, Germany.
- Almela, Á., Valencia-García, R., & Cantos, P. (2012, Apr). Seeing through deception: A computational approach to deceit detection in written communication. *Proceedings of the workshop on computational approaches to deception detection*, Avignon, France (pp. 15–22).
- AlSumait, L., Barbará, D., Gentle, J., & Domeniconi, C. (2009). *Topic Significance Ranking of LDA Generative Models*. In W. Buntine, M. Grobelnik, D. Mladenić, & J. Shawe-Taylor (Eds.), *Machine learning and knowledge discovery in databases*. ECML PKDD 2009. Lecture Notes in Computer Science (vol. 5781, pp. 67–82). Springer. https://doi.org/10.1007/978-3-642-04180-8_22
- Amjad, M., Sidorov, G., Zhila, A., Gelbukh, A., & Rosso, P. (2020). *Urdu-Fake@FIRE2020: shared track on fake news identification in Urdu*. Forum for information retrieval evaluation, Hyderabad, India (pp. 37–40). <https://doi.org/10.1145/3441501.3441541>
- Bakir, V., & McStay, A. (2018). Fake News and the Economy of Emotions Problems, causes, solutions. *Digital Journalism*, 6(2), 154–175. <https://doi.org/10.1080/21670811.2017.1345645>
- Banik, S. (2020). COVID fake news dataset. *Zenodo*. <https://doi.org/10.5281/zenodo.4282522>
- Blanco-Herrero, D., & Calderón, C. A. (2019). Spread and reception of fake news promoting hate speech against migrants and refugees in social media: Research plan for the doctoral programme education in the knowledge society. *Proceedings of the seventh international conference on technological ecosystems for enhancing multiculturalism, León, Spain* (pp. 949–955). <https://doi.org/10.1145/3362789.3362842>
- Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003). Latent dirichlet allocation. *Journal of Machine Learning Research*, 3, 993–1022.
- Brader, T., Marcus, G., & Miller, K. L. (2011). Emotion and public opinion. *The Oxford handbook of American public opinion and the media* (pp. 384–401). <https://doi.org/10.1093/oxfordhb/9780199545636.003.0024>
- Brennen, J. S., Simon, F., Howard, P. N., & Nielsen, R. K. (2020). Types, sources, and claims of COVID-19 misinformation. RISJ Factsheet. *Reuters institute for the study of journalism*. <https://reutersinstitute.politics.ox.ac.uk/types-sources-and-claims-covid-19-misinformation>
- Buller, D. B., & Burgoon, J. K. (1996). Interpersonal deception theory. *Communication Theory*, 6(3), 203–242. <https://doi.org/10.1111/j.1468-2885.1996.tb00127.x>
- Campello, R. J. G. B., Moulavi, D., & Sander, J. (2013). *Density-Based Clustering Based on Hierarchical Density Estimates*. Advances in Knowledge Discovery and Data Mining, Berlin, Heidelberg.
- Carbonell, J., & Goldstein, J. (1998). The use of MMR, diversity-based reranking for reordering documents and producing summaries. *Proceedings of the 21st annual international ACM SIGIR conference on research and development in information retrieval, Melbourne, Australia* (pp. 335–336). <https://doi.org/10.1145/290941.291025>
- Chen, D. (2021). *Chinese sentiment measures*. <https://github.com/dhchen/cn-sentiment-measures>
- Choudhary, A., & Arora, A. (2021). Linguistic feature based learning model for fake news detection and classification. *Expert Systems with Applications*, 169, 114171. <https://doi.org/10.1016/j.eswa.2020.114171>
- Davoudi, M., Moosavi, M. R., & Sadreddini, M. H. (2022). DSS: A hybrid deep model for fake news detection using propagation tree and stance network. *Expert Systems with Applications*, 198, 116635. <https://doi.org/10.1016/j.eswa.2022.116635>
- Dementieva, D., & Panchenko, A. (2020). Fake news detection using multilingual evidence. *IEEE 7th international conference on data science and advanced analytics (DSAA)* (pp. 775–776). <https://doi.org/10.1109/DSAA49011.2020.00111>
- Deng, B., & Chau, M. (2021). The Effect of the Expressed Anger and Sadness on Online News Believability. *Journal of Management*

- Information Systems, 38(4), 959–988. <https://doi.org/10.1080/07421222.2021.1990607>
- DePaulo, B. M., Lindsay, J. J., Malone, B. E., Muhlenbruck, L., Charlton, K., & Cooper, H. (2003). Cues to deception. *Psychological Bulletin*, 129(1), 74–112. <https://doi.org/10.1037/0033-2909.129.1.74>
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019, June 2–June 7). *BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding*. Proceedings of NAACL-HLT 2019, Minneapolis, Minnesota. 4171–4186. <https://doi.org/10.18653/v1/n19-1423>
- Dhillon, I. S., & Sra, S. (2005). Generalized nonnegative matrix approximations with Bregman divergences. *Proceedings of the 18th international conference on neural information processing systems, Vancouver, British Columbia, Canada* (pp. 283–290).
- Dogo, M. S., Deepak, P., & Jurek-Loughrey, A. (2020). Exploring thematic coherence in fake news. *ECML PKDD 2020 workshops. Communications in Computer and Information Science* (vol. 1323, pp. 571–580). Springer. https://doi.org/10.1007/978-3-030-65965-3_40
- Du, J., Dou, Y., Xia, C., Cui, L., Ma, J., & Yu, P. S. (2021). Cross-lingual COVID-19 fake news detection. *2021 International Conference on Data Mining Workshops (ICDMW), Auckland, New Zealand* (pp. 859–862). <https://doi.org/10.1109/ICDMW53433.2021.00110>
- Dumais, S. T. (2004). Latent semantic analysis. *Annual Review of Information Science and Technology*, 38(1), 188–230. <https://doi.org/10.1002/aris.1440380105>
- Faustini, P. H. A., & Covões, T. F. (2020). Fake news detection in multiple platforms and languages. *Expert Systems with Applications*, 158, 113503. <https://doi.org/10.1016/j.eswa.2020.113503>
- George, J., Gerhart, N., & Torres, R. (2021). Uncovering the truth about fake news: A research model grounded in multi-disciplinary literature. *Journal of Management Information Systems*, 38(4), 1067–1094. <https://doi.org/10.1080/07421222.2021.1990608>
- Grootendorst, M. (2022). BERTopic: Neural topic modeling with a class-based TF-IDF procedure. *arXiv preprint arXiv:2203.05794*.
- Gupta, A., Li, H., Farnoush, A., & Jiang, W. (2022). Understanding patterns of COVID infodemic: A systematic and pragmatic approach to curb fake news. *Journal of Business Research*, 140, 670–683. <https://doi.org/10.1016/j.jbusres.2021.11.032>
- Horner, C. G., Galletta, D., Crawford, J., & Shirsat, A. (2021). Emotions: The unexplored fuel of fake news on social media. *Journal of Management Information Systems*, 38(4), 1039–1066. <https://doi.org/10.1080/07421222.2021.1990610>
- Hossain, T., Logan IV, R. L., Ugarte, A., Matsubara, Y., Young, S. D., & Singh, S. (2020). COVIDLies: Detecting COVID-19 misinformation on social media. *Proceedings of the 1st Workshop on NLP for COVID-19 (Part 2) at EMNLP 2020, Online. ACL*. <https://doi.org/10.18653/v1/2020.nlpCOVID19-2.11>
- Ito, J., Song, J., Toda, H., Koike, Y., & Oyama, S. (2015). Assessment of tweet credibility with LDA features. *Proceedings of the 24th international conference on world wide web, Florence, Italy* (pp. 953–958). <https://doi.org/10.1145/2740908.2742569>
- Kar, D., Bhardwaj, M., Samanta, S., & Azad, A. (2021). No rumours please! A multi-Indic-Lingual approach for COVID fake-tweet detection. *2021 Grace Hopper Celebration India (GHCI)*, 1–5. <https://doi.org/10.1109/ghci50508.2021.9514012>
- Kishore Shahi, G., & Nandini, D. (2020). FakeCovid - A multilingual cross-domain fact check news dataset for COVID-19. *International workshop on cyber social threats*. <https://doi.org/10.5281/zenodo.3965870>
- Li, W., & McCallum, A. (2006). Pachinko allocation: DAG-structured mixture models of topic correlations. *Proceedings of the 23rd international conference on machine learning, Pittsburgh, Pennsylvania, USA* (pp. 577–584). <https://doi.org/10.1145/1143844.1143917>
- Lin, J., Tremblay-Taylor, G., Mou, G., You, D., & Lee, K. (2019). Detecting fake news articles. *Proceedings of the 2019 IEEE international conference on big data*. Los Angeles, CA (pp. 3021–3025). <https://doi.org/10.1109/BigData47090.2019.9005980>
- Loria, S. (2020). *Textblob documentation* (Release 0.16.0). <https://buildmedia.readthedocs.org/media/pdf/textblob/latest/textblob.pdf>
- Luo, J., Xue, R., Hu, J., & El Baz, D. (2021). Combating the infodemic: A Chinese infodemic dataset for misinformation identification. *Healthcare*, 9(9), 1094. <https://doi.org/10.3390/healthcare9091094>
- Martel, C., Pennycook, G., & Rand, D. G. (2020). Reliance on emotion promotes belief in fake news. *Cognitive Research: Principles and Implications*, 5(1), 47. <https://doi.org/10.1186/s41235-020-00252-3>
- McInnes, L., & Healy, J. (2018). UMAP: Uniform manifold approximation and projection for dimension reduction. *Journal of Open Source Software*, 3(29), 861. <https://doi.org/10.21105/joss.00861>
- Mikolov, T., Chen, K., Conrado, G., & Dean, J. (2013). Efficient estimation of word representations in vector space. *Proceedings of the workshop at ICLR, Scottsdale* (pp. 1–12).
- Mimno, D., Wallach, H. M., Talley, E., Leenders, M., & McCallum, A. (2011). Optimizing semantic coherence in topic models. *Proceedings of the conference on empirical methods in natural language processing, Edinburgh, United Kingdom* (pp. 262–272). <https://aclanthology.org/D11-1024.pdf>
- Muric, G., Wu, Y., & Ferrara, E. (2021). COVID-19 vaccine hesitancy on social media: Building a public twitter data set of antivaccine content, vaccine misinformation, and conspiracies. *JMIR Public Health and Surveillance*, 7(11), e30642. <https://doi.org/10.2196/30642>
- NewsGuard. (2021). *Coronavirus misinformation tracking dataset*. <https://www.newsguardtech.com/coronavirusmisinformation-tracking-center/>. Accessed 20 March 2021
- Nwankwo, E., Okolo, C., & Habonimana, C. (2020). Topic modeling approaches for understanding COVID-19 misinformation spread in Sub-Saharan Africa. *AI for social good workshop*. https://csrcs.seas.harvard.edu/files/csrcs/files/ai4sg_2020_paper_70.pdf
- O’Callaghan, D., Greene, D., Carthy, J., & Cunningham, P. (2015). An analysis of the coherence of descriptors in topic modeling. *Expert Systems with Applications*, 42(13), 5645–5657. <https://doi.org/10.1016/j.eswa.2015.02.055>
- Ozbay, F. A., & Alatas, B. (2020). Fake news detection within online social media using supervised artificial intelligence algorithms. *Physica A: Statistical Mechanics and Its Applications*, 540, 123174. <https://doi.org/10.1016/j.physa.2019.123174>
- Paixão, M., Lima, R., & Espinas, B. (2020). Fake news classification and topic modeling in Brazilian Portuguese. *2020 IEEE/WIC/ACM international joint conference on Web Intelligence and Intelligent Agent Technology (WI-IAT)*. Melbourne, Australia (pp. 427–432). <https://doi.org/10.1109/WIIAT50758.2020.00063>
- Paschen, J. (2020). Investigating the emotional appeal of fake news using artificial intelligence and human contributions. *Journal of Product & Brand Management*, 29(2), 223–233. <https://doi.org/10.1108/JPB-12-2018-2179>
- Patwa, P., Sharma, S., Pykl, S., Guptha, V., Kumari, G., Akhtar, M. S., ..., Chakraborty, T. (2021). Fighting an infodemic: COVID-19 fake news dataset. In T. Chakraborty, K. Shu, H.R. Bernard, H. Liu, & M.S. Akhtar (Eds.), *Combating online hostile posts in regional languages during emergency situation*. CONSTRAINT 2021. Communications in Computer and Information Science (vol. 1402, pp. 21–29). Springer. https://doi.org/10.1007/978-3-030-73696-5_3
- Pennebaker, J. W., Chung, C. K., Ireland, M., Gonzales, A., & Booth, R. J. (2007). *The development and psychometric properties of LIWC2007*. <http://www.liwc.net/LIWC2007/LanguageManual.pdf>
- Pérez-Rosas, V., & Mihalcea, R. (2014). Cross-cultural deception detection. *Proceedings of the 52nd annual meeting of the association for computational linguistics* (Vol. 2, pp. 440–445) Baltimore, Maryland.
- Pérez-Rosas, V., Kleinberg, B., Lefevre, A., & Mihalcea, R. (2018). Automatic detection of fake news. *Proceedings of the 27th international conference on computational linguistics* Santa Fe, New Mexico, USA.
- Posadas-Durán, J., Gómez-Adorno, H., Sidorov, G., & Escobar, J. J. M. (2019). Detection of fake news in a new corpus for the Spanish

- language. *Journal of Intelligent & Fuzzy Systems*, 36, 4869–4876. <https://doi.org/10.3233/JIFS-179034>
- Poynter. (2021). *The CoronaVirus facts database*. <https://www.poynter.org/coronavirusfactsalliance/>. Accessed 20 March 2021
- Reimers, N., & Gurevych, I. (2020, Nov). Making monolingual sentence embeddings multilingual using knowledge distillation. *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*. Online (pp. 4512–4525). <https://doi.org/10.18653/v1/2020.emnlp-main.365>
- Reis, J. C. S., Correia, A., Murai, F., Veloso, A., & Benevenuto, F. (2019). Supervised Learning for Fake News Detection. *IEEE Intelligent Systems*, 34(2), 76–81. <https://doi.org/10.1109/MIS.2019.2899143>
- Rhoades, S. A. (1995). Market share inequality, the HHI, and other measures of the firm-composition of a market. *Review of Industrial Organization*, 10(6), 657–674. <http://www.jstor.org/stable/41798607>. Accessed 22 Aug 2021
- Sabeeh, V., Zohdy, M., & Al Bashairah, R. (2021). Fake news detection through topic modeling and optimized deep learning with multi-domain knowledge sources. In R. Stahlbock, G. M. Weiss, M. Abou-Nasr, C. Y. Yang, H. R. Arabnia, & L. Deligiannidis (Eds.), *Advances in data science and information engineering*. Transactions on Computational Science and Computational Intelligence (pp. 895–907). Springer.
- Shannon, C. E. (1948). A mathematical theory of communication. *The Bell System Technical Journal*, 27(3), 379–423. <https://doi.org/10.1002/j.1538-7305.1948.tb01338.x>
- Shu, K., Sliva, A., Wang, S., Tang, J., & Liu, H. (2017). Fake News detection on social media: A data mining perspective. *SIGKDD Explorations Newsletter*, 19(1), 22–36. <https://doi.org/10.1145/3137597.3137600>
- Shu, K., Wang, S., & Liu, H. (2019). Beyond news contents: The role of social context for fake news detection. *Proceedings of the Twelfth ACM international conference on web search and data mining, Melbourne VIC, Australia* (pp. 312–320). <https://doi.org/10.1145/3289600.3290994>
- Sun, J., Chen, X., Zhang, Z., Lai, S., Zhao, B., Liu, H., Wang, S., Huan W., Zhao, R., Ng, M.T.A., & Zheng, Y. (2020). Forecasting the long-term trend of COVID-19 epidemic using a dynamic model. *Scientific Reports*, 10(1), 21122. <https://doi.org/10.1038/s41598-020-78084-w>
- Tandoc, E. C. (2017). Five ways BuzzFeed is preserving (or transforming) the journalistic field. *Journalism*, 19(2), 200–216. <https://doi.org/10.1177/1464884917691785>
- Tandoc Jr., E., Thomas, R., & Bishop, L. (2021). What Is (Fake) News? Analyzing News Values (and More) in Fake Stories. *Media and Communication*, 9(1), 110–119. <https://doi.org/10.17645/mac.v9i1.3331>
- Tran, N. K., Zerr, S., Bischoff, K., Niederee, C., & Krestel, R. (2013). *Topic Cropping: Leveraging Latent Topics for the Analysis of Small Corpora*. In T. Aalberg, C. Papatheodorou, M. Dobrev, G. Tsakonas, & C. J. Farrugia (Eds.), *Research and advanced technology for digital libraries*. TPD 2013. Lecture Notes in Computer Science (vol. 8092, pp. 297–308). Springer.
- U.S. Food and Drug Administration (FDA). (2020). Coronavirus update: FDA and FTC warn seven companies selling fraudulent products that claim to treat or prevent COVID-19. *FDA news release*. <https://www.fda.gov/news-events/press-announcements/coronavirus-update-fda-and-ftc-warn-seven-companies-selling-fraudulent-products-claim-treat-or>. Accessed 22 Aug 2021
- Vosoughi, S., Roy, D., & Aral, S. (2018). The spread of true and false news online. *Science*, 359(6380), 1146–1151. <https://doi.org/10.1126/science.aap9559>
- Xia, L., Xu, J., Lan, Y., Guo, J., & Cheng, X. (2015). Learning maximal marginal relevance model via directly optimizing diversity evaluation measures. *Proceedings of the 38th international ACM SIGIR conference on research and development in information retrieval, Santiago, Chile* (pp. 113–122). <https://doi.org/10.1145/2766462.2767710>
- Xu, K., Wang, F., Wang, H., & Yang, B. (2020). Detecting fake news over online social media via domain reputations and content understanding. *Tsinghua Science and Technology*, 25(1), 20–27. <https://doi.org/10.26599/TST.2018.9010139>
- Yang, C., Zhou, X., & Zafarani, R. (2021). CHECKED: Chinese COVID-19 fake news dataset. *Social Network Analysis and Mining*, 11(1), 58. <https://doi.org/10.1007/s13278-021-00766-8>
- Zhang, X., & Ghorbani, A. A. (2020). An overview of online fake news: Characterization, detection, and discussion. *Information Processing & Management*, 57(2), 102025. <https://doi.org/10.1016/j.ipm.2019.03.004>
- Zhang, D., Zhou, L., Kehoe, J. L., & Kilic, I. Y. (2016). What online reviewer behaviors really matter? Effects of verbal and nonverbal behaviors on detection of fake online reviews. *Journal of Management Information Systems*, 33(2), 456–481. <https://doi.org/10.1080/07421222.2016.1205907>
- Zhou, L. (2005). An empirical investigation of deception behavior in Instant Messaging. *IEEE Transactions on Professional Communication*, 48(2), 147–160. <https://doi.org/10.1109/tpc.2005.849652>
- Zhou, L., Burgoon, J. K., Nunamaker, J. F., & Twitchell, D. (2004). Automated linguistics based cues for detecting deception in text-based asynchronous computer-mediated communication: An empirical investigation. *Group Decision & Negotiation*, 13(1), 81–106. <https://doi.org/10.1023/b:grup.0000011944.62889.6f>
- Zhou, L., & Sung, Y. (2008). Cues to deception in online Chinese groups. *Proceedings of Hawaii International Conference on System Sciences (HICSS-41), Big Island, HI, USA*. <https://doi.org/10.1109/hicss.2008.109>

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Springer Nature or its licensor holds exclusive rights to this article under a publishing agreement with the author(s) or other rightsholder(s); author self-archiving of the accepted manuscript version of this article is solely governed by the terms of such publishing agreement and applicable law.

Lina Zhou is a Professor in the Department of Business Information Systems and Operations Management at the Belk College of Business at UNC Charlotte. Her research interests include deception detection, social media analytics, business analytics, and so on. She has (co-)authored numerous articles published in various journals and conference proceedings. Her work on fake news detection was partly supported by the National Science Foundation (NSF) and School of Data Science and Belk College of Business at UNC Charlotte.

Jie Tao is an Associate Professor of Information Systems and Operations Management in the Dolan School of Business at Fairfield University. Dr. Tao received a Doctoral degree in Information Systems from Dakota State University. His recent research interests majorly include Machine Learning and Natural Language Processing. He published several papers in prestigious Information Systems journals. He also presented research articles at several academic conferences. Dr. Tao received the AIS ATLAS award for his services for AIS in the year 2013, and the Best Paper award at HICSS 2020. He is also an Nvidia Deep Learning Institute certified instructor. He served as the faculty advisor for one of the top 10 student teams in IBM WAGC 2016.

Dongsong Zhang is a Belk Endowed Chair Professor in Business Analytics at the Belk College of Business, UNC Charlotte. His current research interests include social computing, health IT, mobile HCI, business intelligence, and online communities. He has published approximately 150 research articles and received a dozen research grants and awards from the NSF, National Institutes of Health (NIH), U.S. Department of Education, and Google Inc., etc.