

SENTIMENTAL REFLECTION OF GLOBAL CRISES: CZECH AND UKRAINIAN VIEWS ON POPULAR EVENTS THROUGH THE PRISM OF INTERNET COMMENTARY¹

KATERYNA HORDIIENKO – ZDENĚK JOUKL

Faculty of Arts, Palacký University in Olomouc, Czech Republic

HORDIIENKO, Kateryna – JOUKL, Zdeněk: Sentimental reflection of global crises: Czech and Ukrainian views on popular events through the prism of internet commentary. *Jazykovedný časopis (Journal of Linguistics)*, 2024, Vol. 75, No. 1, pp. 43–61.

Abstract: Social media have become a part of our lives, and their use helps us learn about events and comment on them with certain emotions. The purpose of our study was to determine the most frequent tone (positive, negative, neutral) of comments on impactful emergency and crisis news in the Czech Republic and Ukraine on a specific topic (pandemics, war, natural disaster etc.) using the sentiment analysis method. The methods of the study included a theoretical analysis of literature, social media (Twitter, Telegram), a Python program using: large language models GPT-3.5-Turbo and Twitter-XLM-RoBERTa, processing and interpretation of results (psycholinguistic).

Keywords: sentiment analysis, emotions, social media, news, posts, tweets, artificial intelligence, GPT-3.5-Turbo, and Twitter-XLM-RoBERTa

1. INTRODUCTION

Social networks have become an important part of our daily lives due to the widespread use of the Internet. The information and communication community, which is developing in a controversial and multidirectional manner, implies the emergence of fundamentally new and interdisciplinary areas of scientific research and development (Nemesh 2017). The dynamics of virtual discourse influence personality development and its consciousness levels, and unveil the characteristic behaviours of various communities.

The growing popularity of the Internet has elevated it to the rank of the main source of universal information and the basis for the formation of social opinion (Wankhade et al. 2022). Digital technologies inadvertently fuel the need to be constantly informed, and this thirst for information leads to a faster spread of unverified news, as anyone can share, comment on, and access information for free (Bonet-Jover et al. 2023). Accordingly, users use various online resources to express

¹ This publication was made possible thanks to targeted funding provided by the Czech Ministry of Education, Youth and Sports for specific research, granted in 2023 to Palacký University Olomouc (IGA_FF_2023_029).

their views, opinions, and emotions about events, products, people, and things around them (Vidhya et al. 2021; Wankhade et al. 2022).

A special place in the commenting activity is occupied by news that evokes an emotional response from commenters and makes them react more actively, i.e., they become popular (Hordiienko – Joukl 2023a). For example, the presence of negative words in a headline increases the click-through rate (CTR) for this headline and thus the response, i.e., the commenting of the audience (Robertson et al. 2023).

Accordingly, researchers often use sentiment analysis to analyse and predict readers' opinions and reactions (and/or measure audience fragmentation; cf. Yang et al. 2020). Also, they automate the process of interpreting comments in terms of polarity and extract important information from large volumes of data (Veselovská 2015). Therefore, our study focuses on the statistical comparison of quantitative results of sentiment analysis of news items consumed by Czechs and Ukrainians daily and their relationship to the sentiment of comments (positive, negative, neutral), which are the product of the emotional state of commenters from different samples and can be compared for the reliability of the results. This data aids in understanding emotional reactions and public opinions, providing valuable support to organizations, analysts, government officials, corporations, suppliers, and others in their decision-making strategies for information consumers (Vidhya et al. 2021; Shamantha et al. 2019; Yang et al. 2020).

The article is organized as follows: Section 2 – Related Work; Sec. 3 – Objectives of the Study; Sec. 4 – Methodology includes the basis of sentiment analysis and its summary, Research Challenges and Contribution of the Paper; Sec. 5 – a description of the Methods and Sources used in the research, Sec. 6 – stages of research Implementation, Sec. 7 – the obtained Results, Sec. 8 – Discussion of the obtained results, determination of the potential of the results for future research, Sec. 9 – Conclusion of our research work; and Sec. 10 – Acknowledgments.

2. RELATED WORK

Social media are online applications that allow users to share their moods, status, and opinions with their virtual social circle (Wan – Gau 2015; Rosenthal et al. 2015). On social media, users post their statuses or thoughts to share with the world (Vidhya et al. 2021). Among Ukrainian users, the most popular platforms for commenting on events are TikTok, Telegram, Instagram, and Facebook. Czech users use Facebook more and Telegram less, but we were unable to use Facebook data in our study due to the social network's policies. Therefore, we chose a less popular platform with a more open policy, namely Twitter and Telegram (Hordiienko – Joukl 2023c).

Sentiment analysis is a very popular type of content analysis method in computational linguistics because a large number of comments, reviews, tweets,

feedback, and other reactions are generated on websites such as e-commerce sites and social networks, and these are used to study the author's emotional state and assess sentiment. In addition, an important context is the general mood of the news we consume daily and its relationship to the mood of comments, which are a product of the emotional state of content consumers (Robertson et al. 2023). Emotions have polarities: joy, surprise, and love are considered positive emotions, while sadness, anger, and disgust are considered negative emotions (Hung – Alias 2023). The emotions studied in this article can be defined as positive, negative, and neutral (cf. Vidhya et al. 2021; Liu 2012). Classifying the emotions of a text according to a specific popular news item can increase the accuracy of sentiment analysis, create a better summary of social opinion, identify the interests of a person or group, and predict behaviour, which can lead to cyberbullying prevention (Pomytkina – Podkopaieva – Hordiienko 2021; Bahan – Navalna – Istomina 2022).

The main approaches to using the method in research include machine learning, hybrid, and lexical (Ravi – Ravi 2015), where the first and second are dominant (Araque et al. 2017), but the latter is usually incorporated into a machine learning approach to improve results (Sánchez-Rada – Iglesias 2019). We will use LLMs (large language models) to see if a sentiment analyser based on at the moment widely used models such as GPT or BERT can be relied upon to understand the current context and current sentiment about an event.

3. OBJECTIVES OF THE STUDY

1. Justify the importance of public opinion and sentiment on popular news topics and the collection of this opinion online using Twitter (or X) and Telegram among Ukrainian and Czech online community users.
2. Categorize comments by news topics (e.g., sports, disasters, diseases, politics and society, war, holidays, etc.) and keywords.
3. Identify the most prevalent tone of comments across a selection of posts and tweets in the Czech Republic and Ukraine on a specific topic.
4. Identify research challenges related to solving the problem of sentiment classification.

4. METHODOLOGY

The search was conducted for posts from the time period of 2022-2023 on Twitter (or X) and Telegram in the main popular channels and publications where readers have the opportunity to comment and rate them, e.g.: BBC Ukraine, Unian.ua, Espresso.tv, Prm.ua for the Ukrainian language and the following for the Czech language: Idnes.cz, Hospodářské noviny, accounts of politicians, accounts of popular people, using the most resonant events that aroused wide public interest

and emotional response of the audience by topics (e.g., “natural disasters”, “catastrophes”, “diseases”, “economy”, “politics”, “changes”, “war”, “violence”, “death”, etc.) and keywords. The news that received a large number of comments, i.e., became popular, was included in the study. We analysed 43 Czech tweets from Twitter and 38 Ukrainian posts from Telegram, and processed 3000 Czech and Ukrainian comments on them using Python script to use the Twitter-XLM-RoBERTa model and GPT-3.5-Turbo-0125 model. The use of Ukrainian and Czech content is in line with the opportunity to determine the general mood of popular news and its impact on the emotions of commentators without the socio-cultural context and with anonymity. Classification of sentiment in sentiment analysis was carried out with the following tasks: data collection, data cleaning, classification of comments by news topics and keywords, launching a pipeline for AI models (tokenization, encoding/decoding, where features are identified, comments classified by tone). These tasks were completed with all the collected comments. These stages of the study are described in the “Implementation” section, and a number of research problems that were identified based on the study are listed in the “Research Challenges” section.

4.1 Research Challenges

1. Text classification analysis is about determining which features or markers can help to accurately identify different classes. For sentiment analysis, these would be positive, neutral, or negative classes. A common approach to text mining is to treat each word in a piece of text as a feature or individual data. However, a single word does not reflect the meaning of the text and may have a different meaning when used in a different context. Therefore, it is better to include other features of the text that provide context and accuracy (Hung – Alias 2023).

2. There are other problems associated with sentiment analysis, such as text context, ambiguity, individual informal writing style, negation words, sarcasm, and irony. Although the recognition of sarcasm and irony in text has improved significantly and is being actively researched, the lack of sufficient tools for working with different languages, and detecting implicit meanings, and ambiguity, needs to be more actively addressed (Wankhade et al. 2022). For example, “You know, we have a topic for discussion. That was the plan.” This review comment features positive words in a negative context.

3. Different models may have different accuracy with the subdivision of sentiment categories into binary and ternary classes (Kumar – Prabhu 2018). Singh et al. 2022 have different sentiment categories, i.e., sadness, joy, fear, anger.

4. Most studies in sentiment analysis aim to identify the polarity of sentiment hidden in a text written in one language. Sentiment research is limited in its use of languages. Therefore, our study takes on the challenge of experimenting with data-poor (although still well-covered) languages (Kumar – Prabhu 2018).

5. The difference between formal written text and text on social media is that the latter does not necessarily follow standard language patterns and is considered informal text but allows for the free expression of feelings and thoughts through different means (Hung – Alias 2023).

6. The presence of irrelevant content, such as ads or gifs, had to be removed manually, while the rest was easily removed using regular expressions according to the structure of the tweets (author name, account name, interpunct character, date, comment text).

7. The original design and research plan for processing CommonCrawl's WET files did not work because it is not always possible to process 21 TB of data, and it did not reflect user comments appropriately because today most user comments are stored dynamically using JavaScript and similar tools. Automatically downloading such a data format would be too complicated and beyond the scope of this article. While it may be easy enough to find relevant comments due to the size of the database, they are very sparsely distributed, making processing very inefficient and causing a waste of resources.

8. The procedure of sentiment analysis is more complicated for use in social media due to the presence of user reactions, abbreviations, jargon, humour, different concepts, and relationships between users of different nationalities (Yuna et al. 2022). Interpreting the news and comments of the majority without taking into account cultural peculiarities can cause false sentiments and stimulate inadequate activity among people of other cultures. In addition, the analysis may be based on the GPT-3.5-Turbo, which has a political and cultural perspective, but at the time of writing this publication was based only on data up to 2021. These problems create obstacles to interpreting sentiments and determining the respective polarity of sentiments of commentators of different nationalities (Hordiienko – Joukl 2023b). The identified problems will be explored in a future publication that will be a continuation of this study. The limitation is that we will not research what the drivers of the sentiments are, this is usually addressed by aspect-based sentiment analysis (Mehra 2023; Tamchyna – Fiala – Veselovská 2015).

4.2 Contribution of the Paper

In terms of theoretical contribution, this study is a novel contribution to the existing literature that has presented the importance of investigating the sentiment of informal comments in language-constrained settings with hidden sentiment polarity in the text. Although the process of applying sentiment analysis is well-researched and often used to determine public opinion, it remains important to study the reliability of public sentiment on popular news topics among users in different countries (Ukraine, Czech Republic) and the general mood of the news consumed daily. This study determines the existence of a link between the mood of popular news and the mood of comments and identifies the most important topics for society. The results will be the impetus for the next study of the cultural and temporal aspects of commentary interpretation.

5. METHODS AND SOURCES

This section contains information about the research methods used and the research design, which includes the stages used in the study, such as Twitter (or X) and Telegram datasets, data cleaning, thematic sorting, and classification algorithms used in the study. The flowchart of the research stages is shown in Fig. 1. You can find a description of the stages of the empirical research in Chapter 6. Implementation.

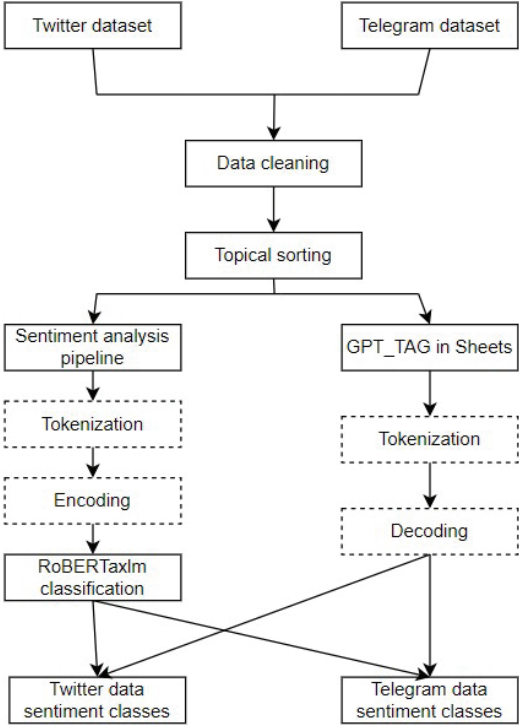


Fig. 1. Flowchart of the research stages

In our study, sentiment analysis is based on the Twitter-XLM-RoBERTa and GPT-3.5-Turbo models because they meet our research interest in exploring the capabilities of modern technologies. This type of models is called large language models (LLMs). LLMs are built on transformer models, trained on vast and diverse dataset and parameters, which allows them to process extensive input data, find long-distance dependencies, and learn complex patterns. Large language models are widely used for sentiment analysis purposes because they are currently the most effective, adapted to specific industries like finance (Li et al. 2023), healthcare (Başarslan – Kayaalp 2020; Li et al. 2023), or commerce, and allow for more

accurate classification of sentiment in specific contexts (for instance, Jain et al. 2023 crafted a version of BERT adapted to evaluate airline reviews, and they also made the model examine their pre-defined cultural aspects). We were interested in the context of comments and news within X, or Twitter, and Telegram. For example, Tan, Lee and Lim (2023) achieved an accuracy of 91.52% for Twitter with a modified version of RoBERTa (different adapted versions are used).

Twitter-RoBERTaXLM stands for Twitter-Robustly Optimized Bidirectional Encoder Representations from Transformers, multilingual, which means that this is an LLM that overcomes the limitations of unidirectional methods, and compared to BERT, RoBERTa has dynamic masking, i.e., more attention masks are tested for the result (Tan et al. 2023). This model (Barbieri et al. 2022) is specifically trained for sentiment analysis on multilingual Twitter data.

GPT-based approaches overcome the performance of others, and GPT-3.5-Turbo is trained on many different kinds of data, including social media data. In the experiment by Kheiri and Karimi 2023, GPT-3.5-Turbo outperformed RoBERTa in the task of sentiment analysis, showing a better understanding of cultural context, emoji, slang, sarcasm, or abbreviations.

Microblogging service X, formerly Twitter, is a social networking platform that serves to share short thoughts and allows interaction through replies to comments. Telegram is a cloud-based instant messaging service that provides a higher level of security and privacy. Both services are characterized by a high level of bot usage, and a certain percentage of comments or posts on both X and Telegram may be created by bots. However, the main objective of our study was to determine the general sentiment in comments and news; we did not set ourselves the goal of detecting bots.

For this study, 3000 comments in Czech from 43 news tweets and 3000 comments in Ukrainian from 38 news posts were collected. The data was cleaned and anonymized using the regular expressions (e.g., author names were deleted) and processed using the Python program. The comments were assigned to the tweets and posts to which they responded to allow for an overall analysis of the comments on individual news items. That is why the tone of each comment was evaluated using GPT-3.5-Turbo and Twitter-XLM-RoBERTa, and the number of negative, neutral, and positive comments was highlighted. The use of 2 models and the opinions of users from different countries make it possible to compare the results with each other and get a more robust result.

6. IMPLEMENTATION

Data collection

Data can be collected from the Internet through web crawling, social media, news channels, e-commerce websites, forums, web blogs, and some other websites. First, we conducted data collection for our sentiment analysis (Wankhade et al.

2022). We collected comments manually using social media, most of which were written during 2022 and 2023. In total, we collected 43 Czech tweets from Twitter and 38 Ukrainian posts from Telegram, with 3000 Czech and Ukrainian comments each. The search for posts was based on the largest number of comments that aroused wide public interest and an emotional response (positive, neutral, or negative) among a large number of users of popular online newsgroups, motivating them to leave a comment under the news. Another important factor was to search for the main popular newsgroups and publications that readers follow, trust, and have the opportunity to comment on and rate.

Data cleaning

During the data cleaning phase, irrelevant content had to be deleted. Some content, like advertisements, required manual deletion, while it was easy to delete the rest by employing regular expressions. Tweets and posts have a regular structure composed of the author's name, account name, interpunct character (U+00B7 in Unicode), date, and finally the text of the comment. Sometimes, there were empty comments because they contained only a gif image. We anonymized and extracted the comments by using the following regular expressions and substitutions, which extract the row with the author's name, the row with the account name, and the row with the interpunct character and leave the date to keep the comment borders identifiable (Definition 1):

$$\text{Regex: } \backslash N^* \backslash n @ \backslash w + \backslash n \cdot \backslash n \quad (1)$$

where “N” – everything that is not a newline; “*” – any number of times (including zero); “\n” – newline; “@” – an at because the comments were always associated with a username; “\w” – any word character; “+” – at least one, then we have a newline again, an interpunct character, and a newline. Substitution with nothing follows, which leaves us with comments separated by dates. Now we can easily remove newlines (\n) and replace the dates with newlines separating individual comments, thereby obtaining a table with comments, each on a separate row.

The regex for Telegram was simpler in Definition 2:

$$\backslash N^* [\backslash N^*] \backslash n \quad (2)$$

where “\N” is any character that is not a newline, “*” means any number of times, “[\N*]” matches square brackets with anything inside except for newlines, and we conclude by a newline (\n). Substitution with nothing follows. This way, we obtain a table with comments, each on a separate row.

Topical sorting

Subsequently, to systematize and analyse the selected comments, we created a file of comments in Ukrainian and Czech, which contains information about the time

of publication, source, news topic, search keywords, tweet/post title, the comment itself, and the sentiment according to GPT-3.5-Turbo and Twitter-XLM-RoBERTa.

The table with the comments had to be organized clearly, with the original post being identifiable. This enabled us to compare the overall sentiments of individual posts. Thus, comments were categorized by topics and keywords. We identified the key search words in Czech, Ukrainian, and English to organize tweets and articles into main topics, and based on them, we processed 42 topics of news articles.

Classification by topic was done qualitatively: given the smaller number of news headlines (i.e., short texts), a manual keyword-based thematic analysis, enriched by observation of other emerging thematic clusters, was more appropriate. For example, a headline might contain the keyword 'war' but report a famous person's opinion (and other headlines also referred to some famous person), then the headline fell under two different topics (war and famous people).

GPT-3.5-Turbo implementation

Given that we have the data in the form of a table, it is convenient to utilize the GPT-3.5-Turbo model directly in the file through the tool GPT for Sheets (this requires paid API access). GPT for Sheets introduced several commands designed for GPT specifically, like GPT_SUMMARIZE, GPT_EDIT, GPT_CLASSIFY, or GPT_TAG. For our purpose, we will use the GPT_TAG function. In the following example, we will explain its format and usage (Definition 3):

=GPT_TAG(E2;"negativní, neutrální, pozitivní";\$J\$5:\$K\$19;1;0;5;"gpt-3.5-turbo") (3)

where, "GPT_TAG" – function for classifying the text in cell E2 (and gradually in the following cells); "negativní, neutrální, pozitivní" – output variants negative, neutral, positive; "\$J\$5:\$K\$19" – examples of classification in cells "from - to"; 1 – top_k parameter, i.e., the maximal number of outputs; 0 – temperature parameter, or, creativity, with 0 meaning the most probable output and having the lowest creativity; 5 – the number of tokens, where 1 should suffice, but especially with Ukrainian, with a lower number we kept getting shortened answers; "gpt-3.5-turbo" – the used model.

In definition, the first parameter means the cell with the comment; the second one is the list of possible results (i.e., sentiments) to be obtained; then a list of examples follows; and finally, the employed model. We gave the function four examples (4-shot learning) for each of the three possible sentiment outcomes, e.g., "Це фантастична новина! Нарешті якийсь прогрес." to be classified as positive in Ukrainian or "Je to naprostá katastrofa. Zdá se, že nic nefunguje dobře." to be labelled negative in Czech.

The GPT-3.5-Turbo model, called via an API in Google Sheets, now tokenizes the data, runs the decoding process to generate the most probable results based on the inputs, and displays the sentiment categories in the Sheets.

Twitter-XLM-RoBERTa model pipeline

Twitter-XLM-RoBERTa pipeline for text classification specifically (pipelines are highly abstracted functions that represent a series of operations and thus simplify programming) includes embeddings or bidirectional encoding and all the necessary configurations and task-specific components required for the model to return results.

To make the Twitter-RoBERTaXLM analysis possible, we transmitted our comment data into an Office Open XML Workbook file on a local computer. Each row contained one user comment, and there was no header. Vladimír Matlach (2023) provided a Python script for us, which employs a multilingual Twitter-XLM-roBERTa-base model from Barbieri, Anke & Camacho-Collados (2022). It uses a pipeline and tokenizer of the model and writes the evaluations into a file.

7. RESULTS

In the study, 3000 Czech comments were used. Both GPT-3.5-Turbo and Twitter-XLM-RoBERTa evaluated the majority of them as negative; for GPT-3.5-Turbo, it was 69.93%, i.e., 2098 comments, and for Twitter-XLM-RoBERTa, 57.76%, i.e. 1739 comments were negative. Only 280 comments (9.33%) of the comments were positive for GPT-3.5-Turbo, and Twitter-XLM-RoBERTa outputted a more positive figure of 594 comments (19.8%). Besides, 622 comments (20.73%) were neutral for GPT-3.5-Turbo, but 667 comments (22.23%) were for Twitter-XLM-RoBERTa.

The figures are similar in Ukrainian: 2068 comments (68.93%) were negative for GPT-3.5-Turbo, and 1738 comments (57.93%) were negative for Twitter-XLM-RoBERTa. Only 227 comments (7.56%) were positive for GPT-3.5-Turbo, but 370 comments (12.33%) for Twitter-XLM-RoBERTa. Moreover, of 705 comments (23.5%) were neutral for GPT-3.5-Turbo, and 892 (29.73%) for Twitter-XLM-RoBERTa. The results of the analysis of Ukrainian and Czech comments are presented in Table 1.

Table 1. The sentiment of comments on popular posts and tweets among Ukrainian and Czech readers of news on social media

Czech Comments					Ukrainian Comments				
Model	GPT-3.5-Turbo		Twitter-XLM-RoBERTa		Model	GPT-3.5-Turbo		Twitter-XLM-RoBERTa	
	com-ments	%	com-ments	%		com-ments	%	com-ments	%
Sentiment					Sentiment				
negative	2098	69.93	1739	57.76	negative	2068	68.93	1738	57.93
neutral	622	20.73	667	22.23	neutral	705	23.5	892	29.73
positive	280	9.33	594	9.8	positive	227	7.56	370	12.33

The results show that the majority of comments in both Czech and Ukrainian using both models are negative, which is indicative of the general mood of readers, the use of profanity in online content to ventilate emotions, reduce tension as a defence against situations that create negativity in real life (Pomytkina et al. 2021). In both cases, there were fewer positive comments because the reader must be able to defend their boundaries and not violate the boundaries of others, understand and manage their own emotions (self-regulation). For example, in a study by Yakovytska et al. (2022), 55% of respondents have a low level of control over their own emotions. That is, when we try to calm down, it is much more difficult to accept and experience emotions in an environmentally friendly way for others, showing respect. Our results are contrary to the results from Singh et al. 2022, where the majority of tweets were positive, while negative tweets accounted for the smallest percentage.

Illustrative data obtained as a result of the theoretical and empirical study are presented in Figures 2 and 3.

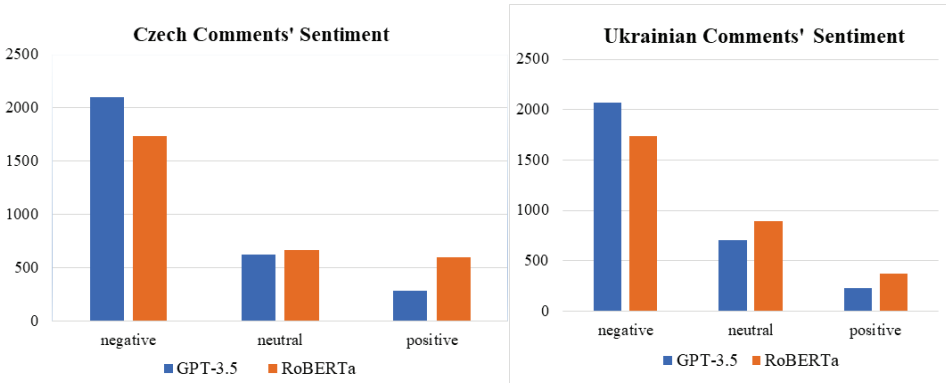


Fig. 2 and Fig. 3. Sentiments of Czech and Ukrainian comments on popular news

Thus, we can conclude that GPT-3.5-Turbo and Twitter-XLM-RoBERTa demonstrate similar results with an average difference of up to 6%, while the average difference between Ukrainian and Czech comments is 3.6%.

It is important to note that, according to our research methodology, the general mood of the news we consume daily affects the mood of the comments. Therefore, we analysed the news by sentiment and compared it to each other. The stages of the work correspond to the study of comments. We selected 43 news tweets in Czech and 38 news posts in Ukrainian. The results of the analysis of news in the form of posts and tweets are presented in Figures 4 and 5.

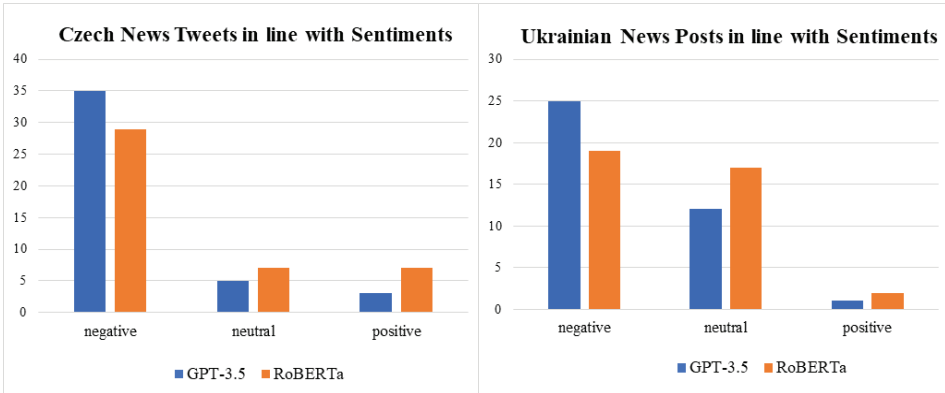


Fig. 4 and Fig. 5. Sentiments of Czech tweets and Ukrainian posts

The illustrative data demonstrates the dominance of negative news posts and tweets. The identified popular events with a negative tone can be analysed as follows: 35 Czech tweets (81.4%) and 25 Ukrainian posts (65.8%) were negative according to GPT-3.5-Turbo, 6 news items (14%) fewer according to Twitter-XLM-RoBERTa’s assessment for Czech tweets, and 6 news items (15.8%) fewer for Ukrainian posts. Instead, positive news accounts for the smallest percentage. Only 3 Czech tweets (7%) were positive according to GPT-3.5-Turbo, while Twitter-XLM-RoBERTa had 7 (16.3%). For Ukrainian posts, GPT-3.5-Turbo identified 1 (2.6%), and Twitter-XLM-RoBERTa identified 2 (5.3%). Interestingly, Twitter-XLM-RoBERTa identified the same number of positive and neutral tweets in Czech tweets, which draws our attention to the ambiguity of identification and confirms the need for an additional model for verification. In summary, GPT-3.5-Turbo and Twitter-XLM-RoBERTa demonstrate similar analysis results with an average difference of up to 10%, with the average difference between Ukrainian and Czech comments being 3.6%. The results of the study of Ukrainian and Czech news are presented in Table 2.

Table 2. News sentiment among Ukrainian posts and Czech tweets

Czech Comments					Ukrainian Comments				
Model	GPT-3.5-Turbo		Twitter-XLMRoBERTa		Model	GPT-3.5-Turbo		Twitter-XLM-RoBERTa	
	com-ments	%	com-ments	%		com-ments	%	com-ments	%
Sentiment					Sentiment				
negative	35	81.4	29	67.4	negative	25	65.8	19	50
neutral	5	11.6	7	16.3	neutral	12	31.6	17	44.7
positive	3	7	7	16.3	positive	1	2.6	2	5.3

A qualitative analysis of the tabular data showed that people’s tendency to pay attention to negative news is a natural stimulus that automatically activates the threat response, increases physiological activation, and focuses on it. This, in turn, helps to form some knowledge on the topic and avoid potentially harmful or painful experiences (Robertson et al. 2023). To help the psyche “survive”, i.e., to relieve tension when reading news with a negative context, the user can respond with a correspondingly negative comment. Therefore, most of the news stories and comments we studied have a negative sentiment. Also, these results can be explained by a well-known phenomenon, the negativity bias. It is a tendency of the news (or even a journalistic principle) to focus on the negative because of the hard-wired instinct of humans to pay more attention to what can be harmful to them. Thus, in their reactions, people will be focusing again on the negative, through which they can release their emotions. In addition, negative comments may result from the popularity of negative tone in social media news.

According to topical sorting, the processed comments on news posts/tweets were categorized by topics, into positive, neutral, and negative, and compared by percentage. The number of comments is the sum of the sentiment using the GPT-3.5-Turbo and Twitter-XLM-RoBERTa models in the amount of 12000 comments, 6000 comments each for Czech and Ukrainian. The Czech comments were categorized into 23 groups according to the following topics, namely: natural hazards, war, assassination attempt, future, covid, democracy, hate, history, LGBT, currency, migration, sport, market, water, elections, agriculture, animals, army, popular people, gas, president, Russia, politics and society. The results are presented in the diagram in Figure 6 (inspired by Rocha 2022).

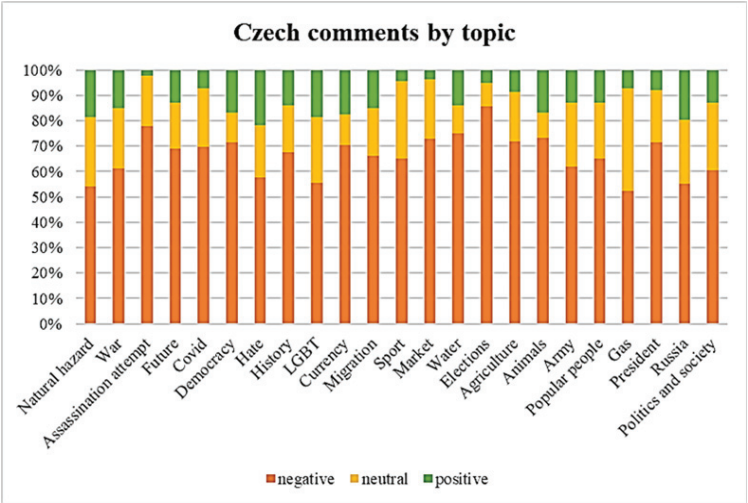


Fig. 6. Sentiments of Czech comments by topic of news posts

In terms of percentage, the most negative sentiment of Czech comments is reflected in the tweet topics “assassination attempt” (78%) and “elections” (85%), neutral – “sport” (30.6%) and “gas” (40.4%), positive – “hate” (21.9%), “Russia” (19.6%), “LGBT” (18.5%), “natural hazard” (18.4%). The results indicate the topics that are of the greatest concern to society and evoke the appropriate sentiment. We would like to justify these results by the positive sentiment: the topic “hate” was referred to in the context that the person who commits violence must be held accountable for it; the topic “Russia” – sanctions and requests for the names of those who took bribes for spreading Russian propaganda; the topic “LGBT” – everyone’s personal choice; the topic “natural hazard” – sympathy and support for those who suffered. It is worth noting that the most popular topics in terms of the number of comments were “war” (3532 comments) and “politics and society” (1350 comments). This means that despite the pronounced sentiment in terms of percentage by topic, the engagement of readers-commenters may be higher for other news tweets.

Ukrainian comments were analysed in the same way and categorized into 18 main topics, namely: economics, healthcare, international relations, natural disaster, natural spectacle, politics, and society, president, social problems, sport, war, animals, celebration, elections, harassment, hate, LGBT, prices, water. The results are shown in the diagram in Fig. 7.

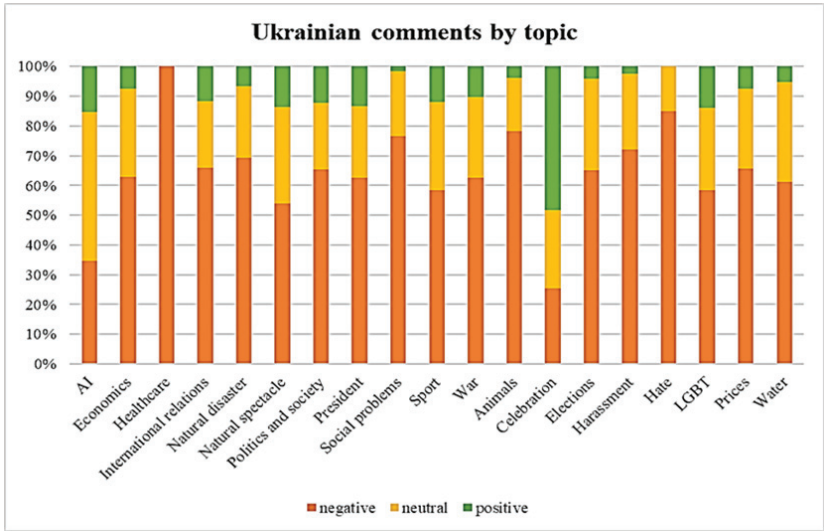


Fig. 7. Sentiments of Ukrainian comments by topic of news posts

Thus, the most negative sentiment of Ukrainian comments is seen in the tweet topics of “healthcare” (100%) and “hate” (85%); neutral – “AI” (50%), “water” (33.4%) and natural spectacle (32.3%); positive – “celebration” (48%). Ukrainian

comments were analysed in the same way and categorized into 19 main topics, namely “war” (1839 comments) and “international relations” (2049 comments).

Having described the detailed results, it is apt to include inter-rater agreement measures. For our situation of two rating LLMs, Cohen’s Unweighted Kappa is used, the value of which ranges from -1 to 1, where 1 is perfect agreement, 0 no agreement beyond chance and negative values agreement less than chance. For the calculation, JASP software was used (JASP Team 2020). Confidence interval of 95% provides a range within which the true kappa value is expected to lie with 95% confidence; this interval is defined by its upper and lower bounds

Table 3. Cohen’s Unweighted Kappa on Czech and Ukrainian comments

Cohen’s Unweighted Kappa			Confidence interval (CI)	
Ratings	Unweighted kappa	Asymptotic Standard Error (SE)	Lower CI 95%	Upper CI 95%
RoBERTa – GPT-3.5 (CZ)	0.467	0.014	0.439	0.495
RoBERTa – GPT-3.5 (UKR)	0.452	0.015	0.423	0.482

The values of Cohen’s Unweighted Kappa were the following: 0.467 for Czech and 0.452 for Ukrainian, indicating some (moderate) agreement between the raters. The variability (standard error is 0.014 for Czech and 0.015 for Ukrainian. Confidence interval lies between 0.423 and 0.482. The difference is calculated for each case, i.e., not only positive & negative difference is taken as a difference, but also neutral & positive or neutral & negative is taken as a difference.

Thus, the dominant sentiment is negative, and the least common one is positive. The most popular topics of news posts and tweets are about war, international relations, politics and society, and events related to these topics can be considered popular in the virtual space.

8. DISCUSSION

Our work was inspired by studies on the use of sentiment analysis [9; 10; 11; 12; 13; 21-23; 28], verbal expression of emotions and their classification [1; 8; 14; 15], the use of various models and APIs for sentiment analysis [2; 11; 19; 14; 24; 29; 30], news and social media research [3; 4; 18; 26] and the psychological aspect of news popularity (perception, emotions, self-regulation on the Internet) [16; 17; 18; 25]. The current gaps in research are that interpreting news and commentary without taking into account cultural characteristics can provide a general misunderstanding of sentiment for those who focus on social opinions in various fields of business, science, and government [7; 20; 27]. In

addition, the GPT3.5-Turbo-based analysis, which has a temporal aspect of training on data up to 2021 and an error in measurements with the Twitter-XLM-RoBERTa model (a problem with the interpretation of irony, humour, and profanity), gives impetus to verify the results using an additional model. That is, the temporal, socio-cultural aspects and the reliability of the obtained sentiments are the prospects of our future research.

9. CONCLUSION

Public opinion and sentiment of popular topics of news tweets and posts show identical indicators of different groups of online community users. The results of the sentiment analysis of comments show that the majority of comments, both Czech and Ukrainian, using different models, are negative, which indicates the general mood of readers. We perceive this as a way of venting emotions, reducing the tension from the negative in real life. The dominance of negative sentiment was also found in the sentiment analysis of Czech news tweets (according to GPT-3.5-Turbo – 81.4%, Twitter-XLM-RoBERTa – 67.4%) and Ukrainian posts (according to GPT-3.5-Turbo – 65.8%; Twitter-XLM-RoBERTa – 50%). These results can be explained by the negativity bias, i.e., people’s focus on negative news is a manifestation of an instinctive response to a natural stimulus that automatically activates the threat response, increases physiological activation, focuses on it, and makes them pay more attention to what may be harmful to reduce emotional stress. In addition, negative comments may be the result of the popularity of the negative tone in social media news.

The processing and classification of comments on 42 news topics revealed that the most sentimental among Czechs and Ukrainians are negative comments on the topics “assassination attempt” (78%), “elections” (85%), “healthcare” (100%), “hate” (85%), while the most popular topics are “war” (3532 Czech comments) and “politics and society” (1350 Czech comments), “war” (1839 Ukrainian comments) and “international relations” (2049 Ukrainian comments).

Thus, in accordance with the aim and objectives of the study, the most frequent tone of comments on popular news in the Czech Republic and Ukraine on a specific topic was determined using the method of sentiment analysis.

ACKNOWLEDGMENTS

The authors would like to thank the IGA project mentor Mgr. Kateřina Lesch, Ph.D., Department of General Linguistics, Palacký University in Olomouc, and the Czech Ministry of Education, Youth and Sports for supporting this work.

References

BAHAN, Myroslava – NAVALNA, Maryna – ISTOMINA Alla (2022): Individual Verbal Codes of Spontaneous Emotional Psychoregulation of Modern Ukrainian Youth. In: *Psycholinguistics*, Vol. 31, No. 2, pp. 6–32.

BARBIERI, Francesco – ANKE, Luis Espinosa – CAMACHO-COLLADOS, Jose (2022): XLM-T: Multilingual Language Models in Twitter for Sentiment Analysis and Beyond. In: *2022 Language Resources and Evaluation Conference, LREC 2022*, pp. 258–266.

BASARSLAN, Muhammet Sinan – KAYAALP, Fatih (2020): Sentiment Analysis with Machine Learning Methods on Social Media. In: *Advances in Distributed Computing and Artificial Intelligence Journal (ADCAIJ)*, Vol. 9, No. 3, pp. 5–15. DOI 10.14201/ADCAIJ202093515

BONET-JOVER, Alba – SEPÚLVEDA-TORRES, Robiert – SAQUETE, Estela – MARTÍNEZ-BARCO, Patricio – NIETO-PÉREZ, Mario (2023): RUN-AS: a novel approach to annotate news reliability for disinformation detection. In: *Language Resources and Evaluation*. DOI 10.1007/s10579-023-09678-9

HORDIIENKO, Kateryna – JOUKL, Zdeněk (2023a): Psychological Basis of the Criminogenic Tone Formation of the Text and Its Automatic Determination. In: *IV International Scientific and Practical Conference “Individuality in the Psychological Dimensions of Communities and Professions”*, pp. 17–20.

HORDIIENKO, Kateryna – JOUKL, Zdeněk (2023b): Sentiment Analysis of Criminogenic Comments on the Internet: Psychological Approach. In: *POLIT. Modern Problems of Science. Humanities: Theses of Reports XXIII International Science and Practice Conf. of Higher Education Graduates and Young Scientists*, pp. 98–99.

HORDIIENKO, Kateryna – JOUKL, Zdeněk (2023c): Sentiment Analysis of Different Nationalities' Internet Comments on Extraordinary News: Cultural Aspect. In: *Summer School of Linguistics (SSoL 2023)*.

HUNG, Lai Po – ALIAS, Suraya (2023): Beyond Sentiment Analysis: A Review of Recent Trends in Text-Based Sentiment Analysis and Emotion Detection. In: *Journal of Advanced Computational Intelligence and Intelligent Informatics (JACIII)*, Vol. 27, No. 1, pp. 84–95. DOI 10.20965/jaciii.2023.p0084

KHEIRI, Kiana – HAMID, Karimi (2023): SentimentGPT: Exploiting GPT for Advanced Sentiment Analysis and its Departure from Current Machine Learning. In: *arXiv*, 2307.10234v2.

KUMAR, M. R. Pawan – PRABHU, Jayagopal (2018): Role of Sentiment Classification in Sentiment Analysis: A Survey. In: *Annals of Library and Information Studies*, Vol. 65, pp. 196–209.

JAIN, Praphula Kumar – QUAMER, Waris – PAMULA, Rajendra – SARAVANAN, Vijayalakshmi (2023): SpSAN: Sparse self-attentive network-based aspect-aware model for sentiment analysis. In: *Journal of Ambient Intelligence and Humanized Computing*, Vol. 14, pp. 3091–3108. DOI 10.1007/s12652-021-03436-x.

LIU, Bing (2012): Sentiment analysis and opinion mining. In: *Synthesis Lectures on Human Language Technologies*, Vol. 5, No. 1, pp. 1–167. DOI 10.2200/S00416ED1V01Y201204HLT016

LI, Xianzhi – CHAN, Samuel – ZHU, Xiaodan – PEI, Yulong – MA, Zhiqiang – LIU, Xiaomo – SHAH, Shah (2023): Are ChatGPT and GPT-4 General-Purpose Solvers for Financial Text Analytics? A Study on Several Typical Tasks. In: *arXiv* 2305.05862v2.

LI, Shanghao – XIE, Zerong – CHIU, Dickson K. W. – HO Kevin K. W. (2023): Sentiment Analysis and Topic Modeling Regarding Online Classes on the Reddit Platform: Educators versus Learners. In: *Applied Sciences*, Vol. 13, No. 4, p. 2250. DOI 10.3390/app13042250.

MEHRA, Payal (2023): Unexpected Surprise: Emotion Analysis and Aspect Based Sentiment Analysis (ABSA) of User Generated Comments to Study Behavioral Intentions of Tourists. In: *Tourism Management Perspectives*, Vol. 45, Article 101063. DOI 10.1016/j.tmp.2022.101063

NEMESH, Olena (2017): *Virtual Activity of Personality: Structure and Dynamics of Psychological Content*. Kyiv: Slovo. 221 p.

POMYTKINA, Liubov – PODKOPAIEVA, Yuliia – HORDIIENKO, Kateryna (2021): Peculiarities of Manifestation of Student Youth' Roles and Positions in the Cyberbullying Process. In: *International Journal of Modern Education and Computer Science*, Vol. 13, No. 6., pp. 1–10. DOI 10.5815/ijmecs.2021.06.01

ROBERTSON, Claire E. – PRÖLLOCHS, Nicolas – SCHWARZENEGGER, Kaoru – PÄRNAMETS, Philip – VAN BAVEL, Jay J. – FEUERRIEGEL, Stefan (2023): Negativity drives online news consumption. In: *Nature Human Behaviour*, Vol. 7, No. 5, pp. 812–822. DOI 10.1038/s41562-023-01538-4

ROCHA, Bruno (2022, June 12). Sentiment Analysis with NSTagger: Ranking popular subreddits by the negativity/hostility of its comments. In: *SwiftRocks*. <https://swiftrocks.com/sentiment-analysis-reddit-negativity>

SÁNCHEZ-RADA, J. Fernando – IGLESIAS, Carlos A. (2019): Social Context in Sentiment Analysis: Formal Definition, Overview of Current Trends and Framework for Comparison. In: *Information Fusion*, Vol. 52, pp. 344–356. DOI 10.1016/j.inffus.2019.05.003

SHAMANTHA, Rai B. – SHETTY, Sweekriti M. – RAI, Prakhyath (2019): Sentiment Analysis Using Machine Learning Classifiers: Evaluation of Performance. In: *IEEE 4th International Conference on Computer and Communication Systems (ICCCS)*, pp. 21–25. DOI 10.1109/CCOMS.2019.8821650

SINGH, Chetanpal – IMAM, Tasadduq – WIBOWO, Santoso – GRANDHI, Srimannarayana (2022): A Deep Learning Approach for Sentiment Analysis of COVID-19 Reviews. In: *Appl. Sci.*, Vol. 12, No. 8, p. 3709. DOI 10.3390/app12083709

TAMCHYNA, Aleš – FIALA, Ondřej – VESELOVSKÁ, Kateřina (2015): Czech aspect-based sentiment analysis: A new dataset and preliminary results. In *CEUR Workshop Proceedings*, Vol. 1422, pp. 95–99.

TAN, Long Tan – LEE, Chin Poo – LIM, Kian Ming (2023): RoBERTa-GRU: A Hybrid Deep Learning Model for Enhanced Sentiment Analysis. In: *Applied Sciences (Switzerland)*, Vol. 13, No. 6. DOI 10.3390/app13063915

YAKOVYTSKA, Lada – POMYTKINA, Liubov – SYNYSHYNA, Viktoriia ICHANSKA, Olena – HORDIIENKO Kateryna (2022): Computer addiction as a new way of personal self-realization of student youth. In: *Journal of Physics: Conference Series, XIV International Conference on Mathematics, Science and Technology Education (ICon-MaSTEd)*, Vol. 2288 012040. DOI 10.1088/1742-6596/2288/1/012040

YANG, Tian – MAJÓ-VÁZQUEZ, Sílvia – NIELSEN, Rasmus K. – GONZÁLEZ-BAILÓN, Sandra (2020): Exposure to news grows less fragmented with an increase in mobile

access. In: *Proceedings of the National Academy of Sciences of the United States of America*, Vol. 117, No. 46, pp. 28678–28683. DOI 10.1073/pnas.2006089117

YUNA, Di – XIAOKUN, Liu – JIANING, Li – LU, Han (2022): Cross-Cultural Communication on Social Media: Review From the Perspective of Cultural Psychology and Neuroscience. In: *Front. Psychol., Sec. Cultural Psychology*, Vol. 13, Article 858900. DOI 10.3389/fpsyg.2022.858900

VESELOVSKÁ, Kateřina (2017): *Sentiment analysis in Czech*. Prague: Institute of Formal and Applied Linguistics as the 16th publication in the series Studies in Computational and Theoretical Linguistics. 171 p.

VIDHYA, Ravi – GOPALAKRISHNAN, Pavithra – VALLAMKONDU, Nanda Kishore (2021): Sentiment Analysis Using Machine Learning Classifiers: Evaluation of Performance. In: *Proceedings of the First International Conference on Computing, Communication and Control System (I3CAC 2021)*. DOI 10.4108/eai.7-6-2021.2308565

MAYUR, Wankhade – RAO, Annavarapu Chandra Sekhara – KULKARNI, Chaitanya (2022): A Survey on Sentiment Analysis Methods, Applications, and Challenges. In: *Artificial Intelligence Review*, Vol. 55, pp. 5731–5780. DOI 10.1007/s10462-022-10144-1

Resumé

SENTIMENTÁLNY ODRAZ GLOBÁLNYCH KRÍZ: ČESKÝ A UKRAJINSKÝ POHĽAD NA POPULÁRNE UDALOSTI CEZ PRIZMU INTERNETOVÝCH KOMENTÁROV

Článok predstavuje analýzu sentimentu v neformálnych komentároch na sociálnych sieťach v Českej republike a na Ukrajine k mimoriadnym udalostiam, a to pomocou modelov umelej inteligencie. Cieľom štúdie je určiť prevládajúci emocionálny tón (pozitívny, negatívny, neutrálny) komentárov k závažným krízovým príspevkom na vybrané témy (pandémia, vojna, prírodné katastrofy a pod.) s využitím metódy analýzy sentimentu. Výskum zahŕňal teoretické spracovanie literatúry, analýzu sociálnych médií (Twitter, Telegram), program Python využívajúci modely strojového učenia GPT-3.5-Turbo a Twitter-XLM-RoBERTa, ako aj spracovanie a psycholingvistickú interpretáciu výsledkov. Celkom bolo zozbieraných 43 českých tweetov a 38 ukrajinských príspevkov z Telegramu, ktoré obsahovali 3000 českých a 3000 ukrajinských komentárov. Väčšina z nich bola napísaná v priebehu rokov 2022 a 2023. Zistilo sa, že väčšina komentárov, českých aj ukrajinských, je negatívna, čo odkazuje na všeobecnú náladu čitateľov. Štúdia tiež poukazuje na problematickosť analýzy sentimentu v prípade príspevkov na sociálnych sieťach v dôsledku prítomnosti reakcií používateľov, skratiek, sarkazmu, humoru alebo tiež rôznych národností používateľov. Článok poskytuje všeobecný pohľad na verejnú mienku a postoje ku globálnym krízam a ich reflexiu prostredníctvom internetových komentárov.