

From discourses to emotions: A Twitter data-driven
study of environmental worry using deep neural
networks.

Candidate Number : XFVN7

Supervisor: Doctor Yi Gong

Pages : 61

Wordcount: 8436

Department of Arts and Sciences

University College London

BASC0024: Final Year Dissertation

BASc Arts and Sciences

Submitted in May 2023

From discourses to emotions: A Twitter data-driven study of environmental worry using deep neural networks.

Abstract

Sustainability and climate change have become increasingly prominent issues in contemporary society. According to a recent survey conducted by the World Economic Forum, almost half of all young people ranked climate change and environmental degradation as the most serious issues facing our world, surpassing concerns over war or inequality (Ipsos 2022). This generation frequently uses social media platforms, such as Twitter, to voice their opinions, share their feelings, and engage in dynamic and participatory conversations about environmental issues.

This dissertation aims to investigate the discourses surrounding the climate crisis and its emotions, adopting an interdisciplinary approach that integrates sentiment analysis to offer deeper insights into climate psychology. A Twitter dataset containing over 40,000 English-written tweets, collected over four months, was pre-processed for this study. Employing the Long Short-Term Memory model, this study analysed interpersonal expressions used about climate change, with a particular focus on negative emotive content.

The results of this investigation indicate that environmental worry is the most prevalent form of emotion expressed on Twitter about climate change, and it encompasses a range of multifaceted expressions. Additionally, the vocabulary utilised in tweets employing negative or positive emotional content also differed significantly. Consequently, this research offers valuable insights into developing effective online communication strategies to inspire pro-climate action.

Preface

My interest in emotions was sparked when I wrote my final year French Baccalaureate paper on how our intentions can impact our actions. In addition, I gained experience in Machine Learning (ML) when I worked as a junior data scientist in the Data Analytics team of a large insurance company during the previous summer. Initially, I had planned to merge these two areas into an original theoretical work, but upon further research, I realised that this connection had already been established and was a developing field of study. This discovery only furthered my interest in exploring how affections and ML could intersect.

This existing field of study can be referred to as sentiment analysis. Sentiment analysis is commonly seen as a subarea of Natural Language Processing (NLP) in Artificial Intelligence (AI), which identify and evaluate opinions expressed in text using automated methods. Though sentiment analysis originated from computer science, it has spread to management sciences and social sciences in the last years from predicting the stock market using opinions from board posts (Mokhtari, Yen, and J. Liu 2018), interpreting public opinions in the context of electoral policies (Nooralahzadeh, Arunachalam, and Chiru 2013), or even predicting movie success and box-office revenue detecting emotions from online reviews (Ha et al. 2019). It is characterised by its interdisciplinary nature, drawing upon psychology, linguistics, and computer science knowledge to build its predictive models. With the arrival of deep learning techniques such as deep neural networks, the performance of emotion detection achieved relatively high performance in classifying sentences into emotion classes.

This research project aims to analyse how people express their eco-anxiety on Twitter using emotion analysis. Eco-anxiety refers to the anxiety caused by the threat of climate change. To conduct this analysis, the project utilises a Long Short-Term Memory (LSTM) model that is coded in Python. Additionally, basic linear algebra is employed to provide understanding a deeper of the model architecture and to evaluate it.

Our study builds upon previous research and combines methodologies from the fields of psychology and computer science to examine the range of negative emotions experienced in response to the ecological crisis. Through this interdisciplinary approach, we aim to develop new insights and draw original conclusions and narratives. By exploring eco-anxiety and its various manifestations online, our project seeks to provide insight into the keywords used to detect negative narratives in order to improve feelings about the more informative climate change, ultimately leading to a more informative and constructive discourse about climate change.

Acknowledgments

I would like to express my gratitude to Dr. Yi Gong, my supervisor, for their feedback during the research of this paper. My family and friends deserve special thanks for their invaluable support during this project. I would also like to extend gratitude to Mr. Elon Musk for his contribution in making research on this topic more challenging through his decision to restrict access to Twitter’s free API. Additionally, I am also grateful to the Arts and Sciences Department at UCL for providing me with a rich and fulfilling education over the course of three years.

Contents

1	Introduction	1
2	Literature review	5
2.1	Literature review on emotions and climate change.	5
2.2	Literature review on emotion analysis	10
3	Proposed Methodology	12
3.1	Methodology	12
3.2	Model: Long Short-Term Memory (LSTM)	12
3.2.1	Word Embeddings	16
3.3	Evaluation Metric	17
3.4	Datasets	18
3.4.1	Training Dataset	18
3.4.2	Twitter Dataset	19
4	Results	24
4.1	Results of the exploratory data analysis	24
4.2	Comparison of LSTM and LSTM with Glove Performance	27
4.3	In-depth analysis of negative sentiments using LSTM-GloVe	29
5	Discussion, Limitations and Suggestions	33
5.1	Discussion	33
5.2	Limitations	34
5.3	Suggestions	37
6	Conclusion	39
A	Appendix	42
A.1	Appendix : Code	43
A.2	A Comprehensive Reference Guide: Key Terminology for Acronyms and Python Libraries	51

List of Figures

3.1	Proposed methodology architecture diagram.	13
3.2	A single memory block of LSTM RNN architecture. Source: (Rahuljha 2020)	14
3.3	Architecture Comparison of LMST Models with Embeddings: Left-hand Side depicts Keras Embeddings, while Right-hand Side depicts Glove Embeddings in Word2Vec Format.	15
3.4	TSNE (t-Distributed Stochastic Neighbor Embedding) plot representing the top 100-words embeddings of our dataset extracted from the ‘Glove’ library.	17
3.5	Initial distribution of labels in the training dataset is represented in graph and table form (in percent).	18
3.6	Time Distribution of Retrieved Tweets – excluding retweets.	21
3.7	Keywords Distribution of Retrieved Tweets.	21
3.8	Features of the Twitter dataset retrieved used for modelling.	21
3.9	An Example of tweet pre-processed. This is a fictional tweet.	22
4.1	Geographical Distribution of Retrieved Tweets from the column “Location”, which contain 1899 non-null values.	24
4.2	Tweet word count analysis: Frequency of tweets with 20 words or less (Left) and Boxplot of overall word count distribution (Right).	25
4.3	Tree representing the distribution of the most common words in the dataset. The size of the box represents its frequency.	25
4.4	Comparison of Word Frequency in the Dataset: The table on the right displays the most frequently occurring words, while the table on the left depicts some infrequent words in the dataset.	26
4.5	Word cloud representing the topics identified by the LDA model trained on one of the dictionaries of the library ‘gensim’. The code for this figure can be find in Appendix A.4. Its uses LDA model.	26
4.6	Pie chart and histogram illustrating the distribution of the sentiments using the ‘VaderSentiment’ library.	27
4.7	Plot of the training and validation curves (accuracy vs.epoch) for the LSTM model with Keras embeddings (Right) and LSTM model with GloVe Embeddings (Left).	28
4.8	Report Tables comparing Precision, Recall, and F1-score Values of LSTM Models with Keras (Right) and GloVe Embeddings (Left)	29
4.9	Pie plots of the majority label predicted, and the minority label predicted from LSTM datasets.	30

4.10	Frequently Occurring Words in Positive and Negative Labeled Tweets, with Some Words Common in Both Categories.	31
4.11	Knowledge Graph representing information as a network of interconnected nodes using the top 100 most frequent words labelled as 'Negative' in the Twitter dataset.	32
5.1	Confusion Matrix comparing the actual and predicted labels for the LSTM model with Glove embeddings.	35
A.1	Code 1 : Retrieving twitter dataset through hashtag filters	43
A.2	Code 2 : Geolocating tweets using user profile location.	44
A.3	Code 3: Generating a map showing the geographical distribution of tweets by country	44
A.4	Code 4 : Generating the Topics' Wordcloud	45
A.5	Code 5 : Function generating the TSNE Plot of Word embedded using Glove	46
A.6	Code 6 : Creating a pie chart to show the distribution of Positive, Negative, and Neutral sentiments from VaderSentiment library	46
A.7	Code 7: Creating a Knowledge Graph	47
A.8	Code 8: Wordcloud for both positive and negative sentiment words combined	48
A.9	Code 9: Training of Model 1 - LSTM	49
A.10	Code 10: Training of Model 2 - LSTM with GloVe	49
A.11	Code 11: Plotting the Results confusion matrix	49
A.12	Code 12: Example 1 of a tweet and its labeled emotions	50
A.13	Code 13: Example 2 of a tweet and its labeled emotions	50

List of Tables

2.1	A glossary outlining the key definitions of the keywords used in the paper. .	6
-----	---	---

Chapter 1

Introduction

The climate and the biodiversity crisis are probably the greatest challenges of our time (Innocenti et al. 2021). The prevailing scientific agreement affirms the occurrence and persistence of global warming, which will have significant impacts on both public and environmental health (Zimbra et al. 2018). The Anthropogenic causes of climate change can no longer be denied, and its systematic consequences involve the increase in the prevalence and intensity of extreme weather events, temperatures, gradual climate changes (e.g., rising sea levels, loss of biodiversity, ...) and increased risks of global pandemics (S. Taylor 2020).

As public awareness of its irreversible impact increases, individuals are increasingly experiencing distressing mental health symptoms. Indeed, becoming aware of climate change as a global hazard and its associated risks can pose challenges to an individual's emotional and social well-being (Alexander 2021). Distress can result in negative emotions, including fear, sadness, anger, anxiety, frustration, and hopelessness. In 2021, a research study found that more than three-quarters of young people in ten surveyed countries thought their future is frightening, with more than half of them thinking that humanity is doomed (Cauberghe et al. 2021). For every person physically affected by a climate disaster, 40 are affected psychologically, says the report from the Grantham Institute at Imperial College London demonstrating that climate change is as much a psychological problem as

an environmental issue (Lawrance et al. 2021).

Although various authors have demonstrated that climate change affects mental health by causing emotional distress, eco-anxiety defined by the American Psychiatric Association as “a chronic fear of environmental doom”(Soutar and Wand 2022), has been primarily researched when analysing the negative effect of the threat on people’s mental health. This has resulted in underestimating or neglecting other negative emotions related to climate change (Stanley et al. 2021).

To bridge this gap, this study aims to delve deeper into the variety of emotions regarding climate change, focusing on its negative expression. While positive emotions can enhance our overall well-being, it is also crucial to acknowledge that negative emotions also play an important role in maintaining good mental health. Differentiating negative emotions is necessary as “while negative emotions are unpleasant, their degree of activation [can] differ” (Stanley et al. 2021,p.4). Indeed, negative emotions have different purposes in signalling an oncoming threat and consequently different degrees of activation to inhibit action or drive behaviour change. It has been established that individuals who experience eco-anger are more likely to engage in pro-climate activism, whereas those experiencing eco-depression exhibit less adaptive behaviour(Stanley et al. 2021)(Alm, Roth, and Sproat n.d.). Enhancing comprehension of the extent and function of emotions with climate change may enhance involvement by tailoring communication to the specific negative emotion and aiding individuals in managing climate-induced stressors (León, Negrodo, and Erviti 2022).

Traditional research methods like interviews and questionnaires have been widely used for analysing emotions (Peters 2022)(Bouman et al. 2020)(Pihkala 2020)(Jabreel and Moreno 2019). However, these methods fail to capture the dynamic nature and the mass reach of such issues like climate change, whose knowledge and implications are continually shifting. To address this, we propose a new approach that integrates social media data analysis into climate psychology. Our methodology employs a deep neural networks model that can be compared to a ”microscope” making invisible phenomena in large amounts

of unstructured text visible (McGillivray and Tóth 2020). From this novel viewpoint, we focus on interpersonal conversations on Twitter about climate change.

Based on the social cognitive theory, we made the argument that expression on climate change could be understood in terms of expressed emotions (Budziszewska and Jonsson 2022). We classify under the umbrella term “environmental worry” all the traditional negative emotions such as anger, sadness, frustration, etc.

This study aims to study the discourse on climate expressions on Twitter to get a clearer picture of how people emotionally react to the climatic crisis, especially looking closer at negative emotions. Beyond “fact-checking”, this study aims at a better understanding of the circulation of different narratives related to climate change and those related to eco-anxiety to improve the understanding of the issue.

To guide the scope of our research, we formulated the following research question:

Using data gathered from Twitter over four months, how do English-speaking individuals communicate their negative emotions related to climate change?

To answer this question, we pose two sub-research questions (RQs):

- **RQ1:** How can we effectively translate emotions from textual social media content using computational methods?
- **RQ2:** What could be the keywords to use in social media to generate an appropriate anger propitious to generate actions ?

Through the exploration of these research questions, the following objectives are presented:

1. Examine social media narratives and conversations on climate change, especially focusing on its negative emotive content.
2. Gain knowledge from tweets to guide strategies to address public concerns on climate change and reduce anxiety behaviours in the future.

Plan of the dissertation:

The dissertation will follow three main steps.

First, we will conduct a literature review to position our research within the context of climate psychology and sentiment analysis. Second, we will develop an interdisciplinary methodology by explaining the underlying assumptions of our models using basic algebra and describing the dataset used. Finally, we will present and discuss our results, answering the research question, and highlighting the limitations of our study. We will conclude by providing suggestions for further work in the field of eco-anxiety.

Chapter 2

Literature review

Psychology has only recently started studying the effects of climate change on mental health. This delay in research is significant (Peters 2022). Sentiment analysis and climate psychology are relatively new fields of study (Pihkala 2020), with the latter still in the process of establishing fundamental definitions, while the former continues to develop its computational methods. Although there has been exponential growth in our understanding of sentiment analysis, research in climate psychology has yet to be extensively explored.

We will look at how these fields intersect with social media studies through a literature review.

2.1 Literature review on emotions and climate change.

The research established that exposure to the climatic events of climate change negatively impacts physical health, mental health, and social relationships (Corral-Verdugo 2021). The phenomenology of climate change's consequences varies widely; they might be immediate or gradual, short-term, or long-term. Panic attacks, a decrease in appetite, irritability, weakness, and lack of sleep are a few of the symptoms that can result from climate anxiety (Wang et al. 2014). Though climate change and mental health are now coming forth in literature, the exact definitions of the keywords and key concepts re-

lated to this issue remain unclear due to their novelty and complexity (Pihkala 2020). It is worth underlining that in some papers the connection between climatic events and their emotions was described through the introduction of new terms such as eco-guilt (Jabreel and Moreno 2019), ecological grief, eco-anger (Roxburgh et al. 2019), or even solastalgia (Pihkala 2020). Additionally, discussions about “Ecological Anxiety disorder” in Cultural Geography or “Anthropocene disorders” in Environmental Humanities do not mean actual disorders of a pathological nature (Pihkala 2020). Clarifications of what we mean by “environmental worry” are thus needed. This lack of agreement presents a significant challenge for researchers seeking to investigate the relationship between climate change and mental health, as it makes it difficult to establish common ground and develop a shared understanding of the issue. We provided below a Glossary defining the keywords used in this paper.

Keyword	Definition
”Negative emotion	An unpleasant emotional reaction that expresses negative affect and often disrupts progress towards one’s goals. Examples include anger, envy, sadness, and fear” (Association n.d.).
Affect	The positive or negative feelings that an individual has towards an event or object, which can inform judgments and decisions by providing a quick evaluation heuristic (Brosch 2021).
Eco-anxiety	A type of anxiety characterized by non-specific worry about our relationship to support environments, often involving a chronic fear of environmental doom (Soutar and Wand 2022).
Environmental worry	General concern or anxiety about the state of the environment, including its impact on human health, wildlife, and ecosystems.
Climate anxiety	Anxiety that is specifically related to Anthropogenic climate change, caused by human activities such as burning fossil fuels and deforestation (Brosch 2021).

Table 2.1: A glossary outlining the key definitions of the keywords used in the paper.

The dearth of scholarly investigations on the adverse emotional experiences associated with climate change has prompted the inception of this interdisciplinary research study. Despite the rising prevalence of eco-anxiety, investigating this phenomenon poses substantial challenges, as Pihkala has acknowledged, “the situation is further complicated by the fact that there is not yet a common interdisciplinary field of studies on the subject”

(2018,p.546) when writing on eco-anxiety.

Climate psychology scholarship is the predominant emerging field that aims to assist human beings in mitigating and adapting to climate change consequences. The most common publication formats were review papers, publications using a qualitative methodology, and theoretical papers (Peters 2022). Researchers in this field tend to use interviews and thematic analysis as their primary research methodologies. However, their ontological and epistemological positions can influence the validity and reliability of their findings, limiting the scope and potential impact of their research if they are too rigid in their beliefs. In their systematic review, researchers argued that though the field is only a recent one, there appears to be a predominance of studies situated within a qualitative approach (Tam, A. K. .-. Leung, and Clayton 2021).

Additionally, "literature on climate change tends to focus on a deficit discourse by presenting individuals as unwilling to change or having insufficient information" (Peters, p. 514). In addition, despite the increased apparent demand for therapeutic support, very little existing literature seems to address relevant and specific therapeutic interventions. Therefore, there is an urgent need to look beyond the disciplinary field of climate psychology scholarship and explore new methods to determine mitigating and adaptive behaviour in response to climate change (Peters 2022).

It is crucial to comprehend these negative emotions and sentiments because they play a significant role in determining people's responses and reactions to the challenges posed by the climate crisis. Climate psychologists, using the social cognitive theory developed by Albert Bandura in the 1970s (Bandura 1986 - 1986), suggest that individuals' self-efficacy, or belief in their ability to achieve goals, influences their behaviour. Furthermore, accumulating research in the affective sciences has demonstrated that emotions and affect have a substantial influence on human information processing, decision-making, and behaviour. Climate change perceptions and actions are predominantly driven by affect and emotions, as research in the field has shown (Brosch 2021)(Bouman et al. 2020)(Stanley et al. 2021). In social psychology research, anger predicts collective action, indicating that eco-anger

could be more strongly related to collective action on climate change. Researchers used online surveys to establish a relationship between eco-anxiety, eco-depression, and actions and found that experiencing eco-anger predicted better mental health outcomes, greater engagement in pro-climate activism, and personal behaviours than experiencing eco-anxiety (Fernandez et al. 2016)(Roxburgh et al. 2019). Furthermore, experimental studies that investigated the effects of inducing collective guilt for human-caused environmental damages, such as signing an environmental petition, and using positive and negative environmental messages found that pride increased intentions to invest in environmental protection, guilt increased willingness to repair environmental damages, and anger increased tendencies to punish others for negative environmental actions (Brosch 2021).

Framing literature points out the insufficient scholarly attention paid to frames occurring from casual conversations despite their effects on public perception (Jang and Hart 2015). Methodological challenges have led to a scarcity of empirical research on framing interpersonal communications during the climate crisis. Responding to this gap in the literature, we consider the space of Twitter conversations to be the natural field where interpersonal everyday conversations a societal topic is unobtrusively measured and analysed.

Indeed, social media has become a rich repository for researchers to help in understanding the public's opinions and feelings about current events. As a source of news or a place of debate, social media platforms have an important role in communications for climate change. As of 2021, Twitter has 330 million active users (Sirisha and Chandana 2022). A 2016 Pew Research Centre survey found that "62% of American adults now get news on social media sites, with 18% doing so regularly" (Roxburgh et al. 2019,p.51). The growth of social media as a source of news means platforms like Twitter are joining legacy media as important mediators of discourse on climate change, making it a valuable and unique data source for sentiment analysis and opinion mining. Surveys and experimental evidence shows that concern about climate change increases with current media coverage, sometimes coupled with direct experience such as fluctuations in local weather conditions.

Finally, social media is driven by interactions and creates a new participatory culture, which transforms passive individuals into active producers who produce and share contexts. People unable to express themselves on such a distressing theme such as the climate emergency might be more likely to express themselves on Twitter due to its engaging platforms and ‘anonymity’ (J. Leung et al. 2021). The literature review on climate change on Twitter did not delve into the characteristics of the climate change conversation, nor did it conduct a detailed analysis of incidents. Instead, the review centered on the quantity of Twitter posts as an indicator of the level of interest in the topic (e.g., such as the pike number of tweets around COP22). Especially in the climate change context, recent evidence shows that selective media use and climate change perceptions mutually reinforce each other, leading to opinion polarisation (León, Negredo, and Erviti 2022). We thus expect Twitter users to rely on issues frames to help reinforce their political positions. In addition, social media can play a decisive role in bringing closer to people the topic, facilitating public awareness, and fostering action to tackle it. The popularity and influence of negative tweets could potentially lead to a “snowball effect” on social media, as regular social media users pick up sentimental feelings and opinions from tweets by others (J. Leung et al. 2021). However, research on the role of social media posts and climate awareness has barely scratched the surface.

In conclusion, the relationship between climate change and negative emotions is gaining recognition in the literature. However, there is a lack of clarity in the definitions of the keywords and concepts related to this issue, presenting a significant challenge for researchers. Researchers demonstrated that the emotional experiences associated with climate have a substantial influence on human decision-making and behaviour. As a result, there is an urgent need to expand beyond the disciplinary field of climate psychology study to investigate new approaches for determining mitigating and adapting behaviour in response to climate change.

2.2 Literature review on emotion analysis

With the availability of state-of-the-art ML and NLP algorithms, sentiment analysis has become a key tool in understanding human behaviour and offers numerous challenging and fascinating research problems (O. Bruna 2016). Since 2000, sentiment analysis, commonly seen as a sub-area of NLP, has grown to be one of the most active field research fields in AI (B. Liu 2015). Its theoretical framework touches every core of NLP, such as lexical semantics, word sense, discourse analysis as well as information extraction (Colneric and Demsar 2020). It is defined as “an [...] area of research motivated to improve the automated recognition of sentiment expressed in the text” (Zimbra et al. 2018,p.2). In other words, through its models, sentiment analysis aims to assign to textual data a value of its emotional weight or content. With the rise of AI, two branches have been developed within its field. On the one hand, traditional sentiment analysis is done with the following three labels outputs such as “Positive”, “Negative” and “Neutral” (Antonakaki, Fragopoulou, and Ioannidis 2021)(Markowitz and Guckian 2018). On the other hand, the recent development of state-of-the-art NLP algorithms allowed us to analyse the emotions of users in a more nuanced and detailed way classifying emotions from frustration to anger or happiness (Markowitz and Guckian 2018) (Jain and Sandhu 2015).

Because of its multidisciplinary applications, research in the field does not only advances the state of research in NLP but also management science, linguistics, political science, or even economy, as these fields are concerned with customer or public opinion and discourses (B. Liu 2015).

In this brief survey of the literature, we will describe the various attributes of sentiment models by examining methods such as supervised and unsupervised learning, lexicon-based techniques, and their application in opinion mining on social media platforms.

Recognising users’ emotions is a major challenge for both humans and machines. ML approaches are used to find different valuable patterns within huge and complex data. Most current ML algorithms solve the current task as a text classification problem. To assign “emotional content” value to an input, the machine needs to have an accurate ground

truth for emotion. While supervised learning approaches rely on labelled training data with sentiment labels to calibrate the model’s parameters, unsupervised learning identifies patterns in the data directly, often using probabilistic models without prior knowledge of the feature class distributions (Jabreel and Moreno 2019)(Mao et al. 2019)(Gautam et al. 2022).

The literature has shown that supervised learning approaches are more commonly used in emotion detection, as they typically yield better results than unsupervised learning. However, the quality of the training dataset is critical to the performance of the model (Colneric and Demsar 2020). One approach to obtaining knowledge of the polarity and intensity of emotional words is using public sentiment lexicons (Jabreel and Moreno 2019). A sentiment analysis approach using pre-existing dictionaries of words labelled as positive or negative can be employed. However, it is important to note that words such as ”Happy” and ”Sad” may appear in similar emotional contexts but represent different emotions. Therefore, relying solely on word co-occurrence may not be enough to capture the emotional information of these words (Mao et al. 2019).

As phrases play a key role in determining the most appropriate set of emotions that must be assigned to a tweet, we have prioritized the use of a human-annotated dataset of tweets with their labelled emotions to train our classifier. By using this dataset, our ML model can learn to recognize patterns in the training data and generalize them to make predictions on new, unseen data.

Additionally, word embeddings such as Global Vectors for Word Representation (Glove) (Chatterjee et al. 2019)(Mao et al. 2019) or Generative Pre-trained Transformer (GPT) are NLP that assign a vector of real numbers, typically with several hundred dimensions, to a word. These embeddings capture the meaning or semantic relationship of the word to other words in its context. This allows the classifier to fully comprehend the content of a word based on its surrounding context (Chatterjee et al. 2019)(McGillivray and Tóth 2020).

Chapter 3

Proposed Methodology

3.1 Methodology

“One of today’s cardinal tasks is to marry the philosopher’s literate ethics with the scientist’s commitment to numerate analysis. Words are important, but they often require a numerate cast”. (Hardin 1998,p.2) Bringing concepts between humanities and language technology, we translated a qualitative research problem into quantitative research goals. Indeed, analyses of social media information often adopt qualitative research methodologies by applying a grounded theory or phenomenological approach (Innocenti et al. 2021)(Budziszewska and Jonsson 2022). By integrating the tools of deep learning into a topic traditionally researched using qualitative methods such as in-depth interviews, we aim to provide a new outlook on eco-anxiety and advance a fundamental understanding of climate psychology. Integrating results from various paradigmatic approaches is difficult and requires precise methodology justified by the literature review.

3.2 Model: Long Short-Term Memory (LSTM)

In the below **Methodology** diagram 3.1, we present an overview of the many processes we elaborated to select our model.

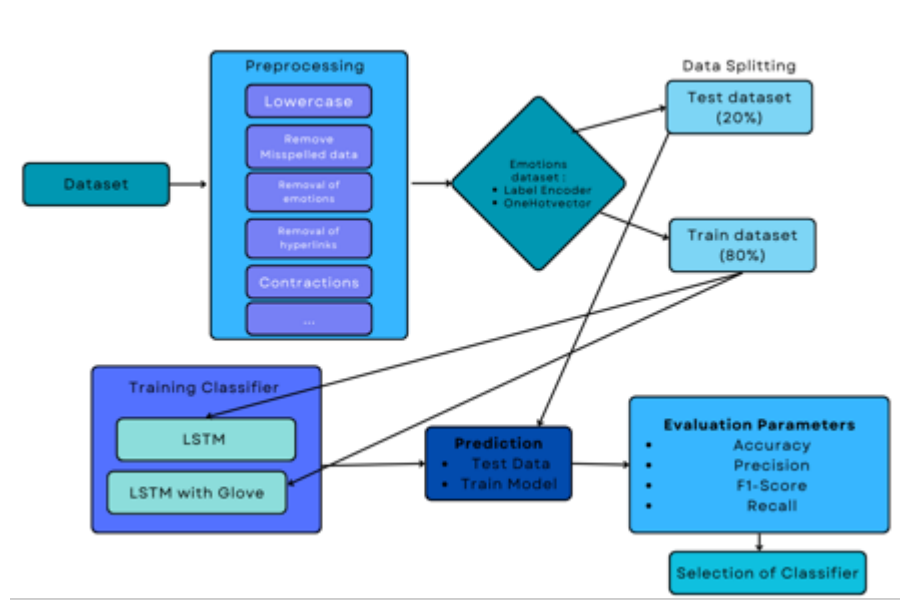


Figure 3.1: Proposed methodology architecture diagram.

NLP is a subfield of AI that deals with understanding and deriving insights from human languages such as text and speech. LSTM is an artificial neural network well suited for text analysis as its network architecture can quickly learn grammatical dependencies. To understand the LSTM concept, it is essential to understand RNN architecture as it is built on its past work. RNN is a type of supervised deep learning that uses a feedforward artificial neural network which can handle variable-length sequence input. Unlike traditional feedforward neural networks, RNNs use feedback loops to process sequences to maintain memory over time. In the traditional RNN algorithm, recurrent units have very simple structures that have no memory units and additional gates (Chatterjee et al. 2019). There is only a simple multiplication of inputs and previous outputs, which is passed through the corresponding activation function. Nonetheless, an LSTM recurrent unit contains gates, which are used to maintain memory for long periods (Sirisha and Chandana 2022). Figure 3.2 highlights the main components that interact with each other to process sequential data:

1. **The input gate:** This gate controls the flow of new input data into the cell. It takes the input data and decides how much of it to let in.

2. **The forget gate:** This gate controls the retention of information from the previous time step. It decides how much of the previous cell state to forget.
3. **The cell state:** This is the "memory" of the cell, which stores information over time. It is updated based on the input gate and forget gate.
4. **The output gate:** This gate controls the output of the cell. It produces the final output of the LSTM cell.

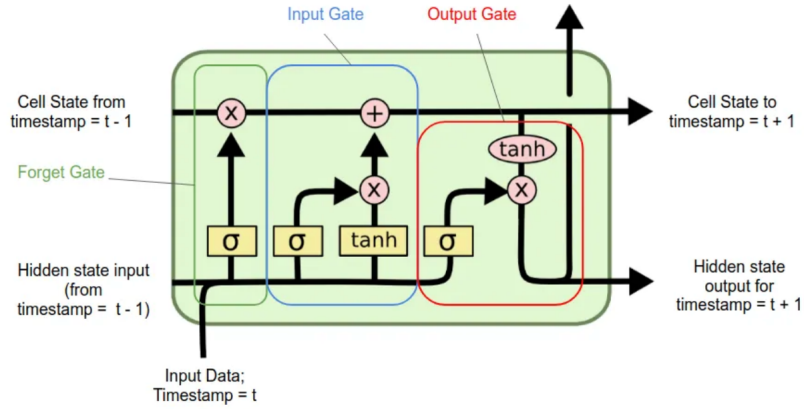


Figure 3.2: A single memory block of LSTM RNN architecture. Source: (Rahuljha 2020)

LSTM recurrent unit:

Our model is mainly composed of three layers which are embedding, RNN algorithm using LSTM recurrent units, and a fully connected sequential dense neural network (Saxena 2023). We created the embedding layer with a vector space of a dense dimension of 16 and an input sequence length depending on the length of the training dataset. Then for LSTM, we used many of 250 several the last layer of our RNN network, we used a fully connected neural network with a hidden layer containing 128 neurons and an output layer composed of 6 output neurons. The detail of our model architecture is provided in more detail by 3.3 .

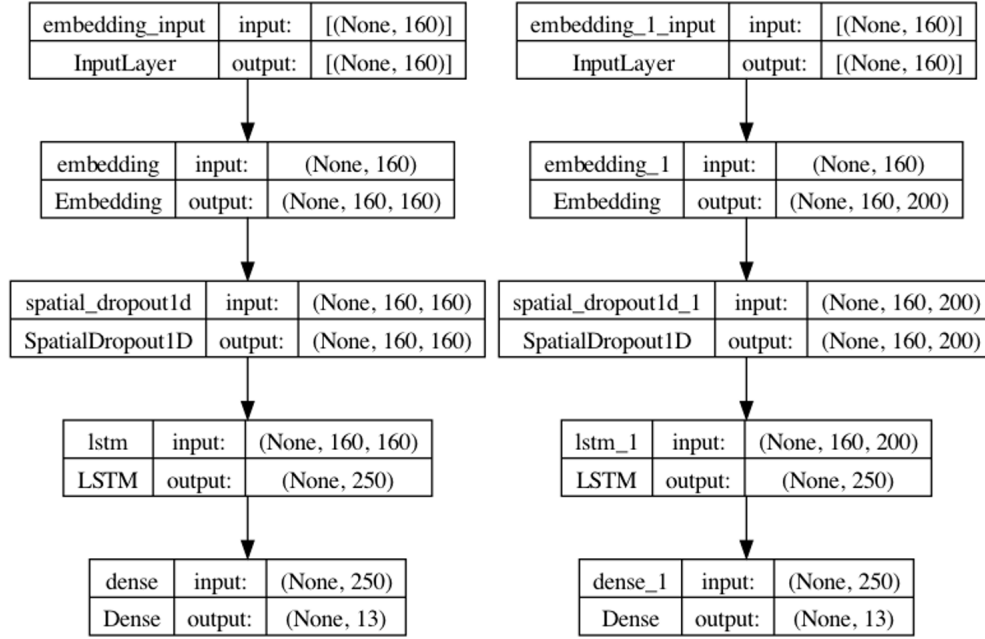


Figure 3.3: Architecture Comparison of LMST Models with Embeddings: Left-hand Side depicts Keras Embeddings, while Right-hand Side depicts Glove Embeddings in Word2Vec Format.

The code of the models can be found for Figure A.8 and Figure A.8 in **Appendix**

Activation function:

To prevent overfitting, we set a dropout rate of 20% and employed SoftMax as the activation function. The SoftMax activation function is defined as:

$$\text{softmax}(z_i) = \frac{e^{z_i}}{\sum_{j=1}^n e^{z_j}} \quad (3.1)$$

for $j = 1$ to n , where z_i is the i -th element of the input vector, and n is the number of elements in the input vector. It is used as the final activation function in neural networks for multi-class classification tasks and takes a vector of real numbers as input and outputs another vector of the same length, with each element representing the probability of the input belonging to the corresponding class (Saxena 2023).

Loss function:

Additionally, we used Adam and Categorical CrossEntropy functions for the loss and optimizer parameters of the network, respectively. Our findings indicate that this RNN architecture with LSTM recurrent units and fully connected sequential dense neural network produced excellent results in handling variable-length sequence inputs. The Categorical CrossEntropy loss function is defined as:

$$loss = - \sum y_{true} \log y_{pred} \quad (3.2)$$

, where y_{true} and y_{pred} are the true and predicted output vectors, respectively. The goal of training the model is to minimize the loss function by adjusting the weights and biases of the model using an optimizer such as the Adam optimizer.

3.2.1 Word Embeddings

The processing of words by ML models requires a numerical representation. Word embeddings convert words into vectors that represent several facets of their meaning and their semantic connections to other words (for instance, the vector for "Queen" provides information such as gender, status etc.) (Chatterjee et al. 2019). **GloVe** word embeddings approach to speech is "You shall know a word by the company it keeps" (Mao et al. 2019,p.2). Indeed, by leveraging statistical information, these numerical representations allow for the recognition of words with comparable meanings by identifying similarities or dissimilarities between words based on their vector representation. First-words word embedding techniques, such as GloVe, were created utilizing deep learning and trained on massive, unlabelled corpora using various architectures such as CNN, RNN and Transformers.

Before applying the LSTM model with GloVe, we converted both inputs of datasets into GloVe embeddings in the **Word2vec** format (Gb and Jacob 2022). Figure ??represents how the input is spatially computed for word embeddings.

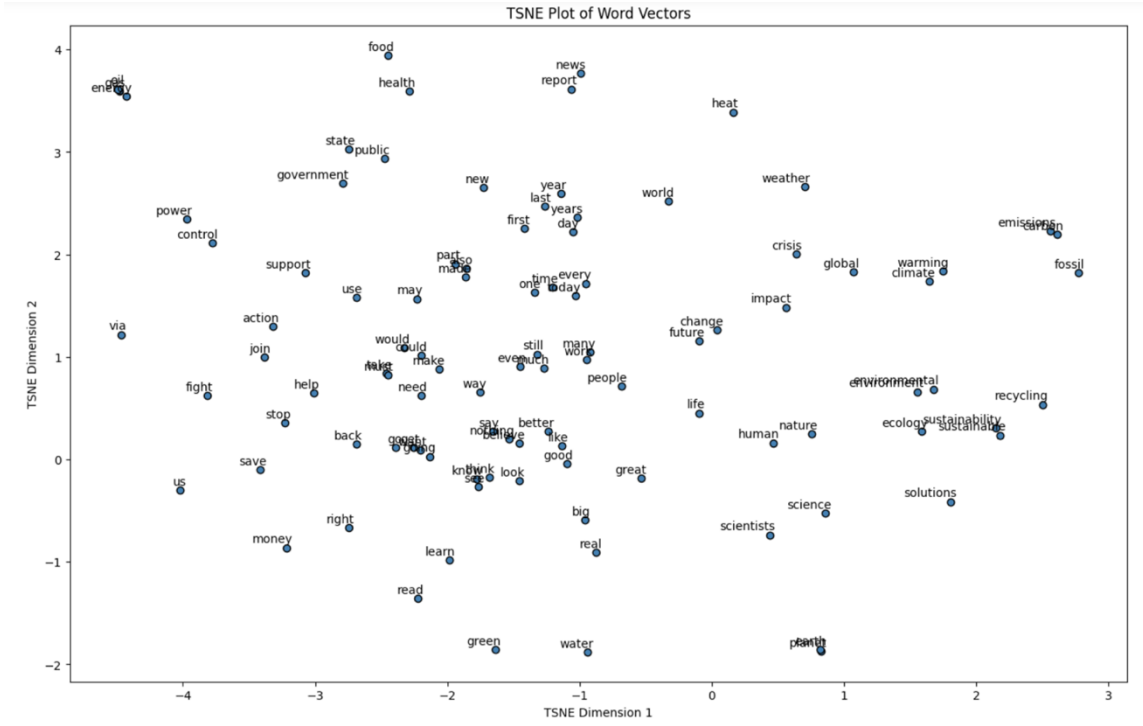


Figure 3.4: TSNE (t-Distributed Stochastic Neighbor Embedding) plot representing the top 100-words embeddings of our dataset extracted from the ‘Glove’ library.

3.3 Evaluation Metric

When it comes to categorising emotions, the commonly used measures for evaluation are precision and recall rate. The Macro F1-score, which is the unweighted mean of all the per-class F1 scores, was also selected to measure the accuracy when measuring the performance of both models. The formula to F1 score the F1 score for a particular class is as follows: $F_1 = 2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}$

Where : $\text{Precision} = \frac{\text{True Positives}}{\text{True Positives} + \text{False Positives}}$ and $\text{Recall} = \frac{\text{True Positives}}{\text{True Positives} + \text{False Negatives}}$

Precision pertains to the ratio of true positives to positive predictions, while recall pertains to the ratio of correctly predicted positive cases to all positive cases. The optimal F1 score is 1, while the lowest score is 0.

When comparing two predictive models, it is essential to consider the requirements of the problem and select evaluation metrics accordingly. The Macro average F1 Score was

deemed a suitable evaluation metric primarily due to its capacity to assess the performance of classification models in the presence of class imbalance. In our training dataset, class imbalance is important as the largest category (neutral) is over twenty times larger than the three smallest classes (enthusiasm, anger, and boredom). Although accuracy may be a more straightforward metric to comprehend, class imbalance may result in deceptive accuracies, hence our decision to avoid it.

3.4 Datasets

3.4.1 Training Dataset

Dataset

The dataset used for training was obtained from Kaggle.com, a platform for data science and machine learning competitions that provides various open-source datasets and notebooks (Benedicto 2020). This dataset comprises 33,871 tweets and is utilised for binary classification, with each tweet labelled with one emotion. We selected this dataset for its focus on tweets and the diverse range of emotions/labels it presents.

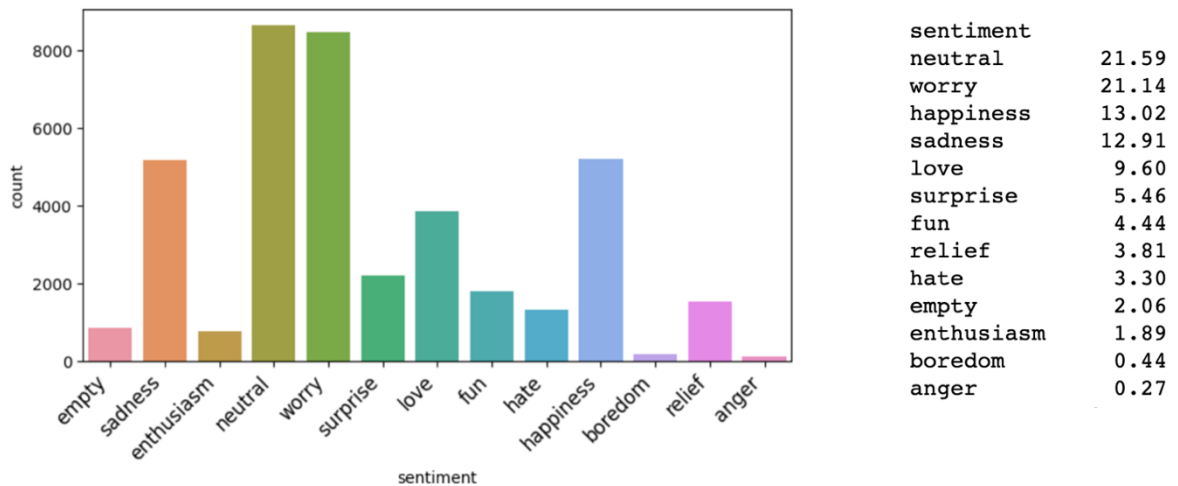


Figure 3.5: Initial distribution of labels in the training dataset is represented in graph and table form (in percent).

Pre-processing

The sentiment column of this dataset is the target variable of the model. To optimise the learning of our classifier, we encoded each category into numerical representations using the functions Label Encoding and One-Hot Encoding. These steps considerably improve the model performance, and the interpretability of the results as well as reduce memory usage. The data were divided into training and testing datasets through a random split using the `train_test_split()` function from the **scikit-learn** library. The proportion of data reserved for testing was set to 20%. This approach enables the model to be trained on the training dataset and assessed on the testing dataset, providing an estimate of its performance on new and unseen data. This approach mitigates the issue of overfitting and offers a more reliable indication of how the model will perform on unseen data.

3.4.2 Twitter Dataset

Ethical consideration

To begin with, ethical and legal are high-priority concerns when retrieving social media data for research purposes (Townsend and Wallace 2017)(Wainman 2018). Though the interest in the field of social media is increasing within the researchers' communities, no consistent approach to ethics has been provided to researchers in this sphere. For this research project, the Twitter academic Application Programming Interface (API) has been used to retrieve all English-language tweets posted from the UK that mention climate change in 3 months from January 2023 to March 2023 using the **Tweepy** library. We informed Twitter of our research intentions, and our research project has been accepted as being consistent with the conditions of the platform. By the Twitter agreement, "when their data [the data of Tweeters] has public disposition it may be viewed and used for research" (Gold 2020, p.6) so no ethics review was necessary. However, though the data being used has a current public disposition, the unclear boundaries between "public" and "private" spaces may not always be easily deduced from posting and there is still a current debate within academia in classifying Twitter data usage as falling between primary and

secondary analysis. We could not obtain implicit and informed consent for data use from the research participants as it is unrealistic for large quantitative research to obtain consent between researcher and research participants. By complying with the Twitter agreement, “implicit consent for data use is only available as an ethical defence when compared to the current state of the dynamic Twitter dataset”. This has implications for research design since data retrieval must account for this, hence the dataset has been synchronized to the state of the online Twitter data set. Consequently, no data was retained. Moreover, anonymity is a key consideration in research ethics (J. Taylor and Pagliari 2017) . The username of the tweeter wasn’t gathered, and no direct quote from a tweet is included in this dissertation. Only results derived from aggregated data will be released.

Pre-processing

To analyse emotions through Twitter, it is necessary to organise the collected data into a corpus, which is a collection of texts in linguistic studies that can be analysed for a particular purpose. Creating a corpus involves three main stages: collecting raw data, annotating it, and analysing it.

Collecting raw data

Using the Tweepy library, we were able to retrieve a dataset of over 40,000 tweets. All tweets were collected from public accounts, and private accounts were excluded. We aimed to mine tweets related to the topic of climate change, and we achieved this by filtering through various hashtags that were identified in the top trends of Twitter search (#climate, #climatechange, #ecology, #climateaction, etc.). To focus our analysis on distinct expressions about climate change, we removed all duplicate tweets utilising the `drop_duplicates()` method in the **Pandas** library. The unique tweets obtained during the 4 months are outlined in Figure 3.6, along with the distribution of keywords in our dataset as illustrated in Figure 3.7. As indicated in Figure 3.8 , our resulting dataset contained precisely 10,032 unique tweets. The code in Appendix A.1 extracts user and tweet information, including description, location, following, fol-

lowers, total tweets, retweets, hashtags, and date of the tweet. Columns 'Language', 'extracted_user_location', 'country', 'location', 'coordinates', 'state', 'clean_msg' and 'message_len' were added during the pre-processing step.

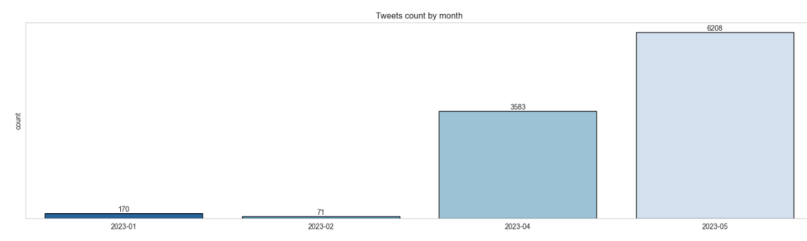


Figure 3.6: Time Distribution of Retrieved Tweets – excluding retweets.

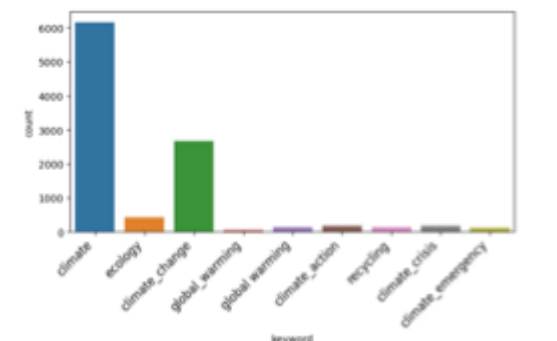


Figure 3.7: Keywords Distribution of Retrieved Tweets.

```
<class 'pandas.core.frame.DataFrame'>
RangeIndex: 10032 entries, 0 to 10031
Data columns (total 17 columns):
#   Column                                Non-Null Count  Dtype
---  -
0   Description                           8777 non-null   object
1   Location                             6790 non-null   object
2   Following                            10032 non-null  float64
3   Follower                             10032 non-null  float64
4   Tweets                              10032 non-null  float64
5   Retweet                             10032 non-null  float64
6   original                             10032 non-null  object
7   Date                                10032 non-null  object
8   keyword                             10032 non-null  object
9   Language                            10032 non-null  object
10  extracted_user_location              1899 non-null   object
11  country                             9673 non-null   object
12  location                             1899 non-null   float64
13  coordinates                         1899 non-null   float64
14  state                              1899 non-null   float64
15  clean_msg                           10032 non-null  object
16  message_len                         10032 non-null  int64
dtypes: float64(7), int64(1), object(9)
memory usage: 1.3+ MB
```

Figure 3.8: Features of the Twitter dataset retrieved used for modelling.

Pre-processing

Tweets are short and often informal (at most 280 characters). A tweet may contain images, videos, URLs, emoticons, and hashtags. For this analysis, we focused only on tweets written in English, for this we cleaned using the library ‘`langdetect()`’ from Google’s language-detection. A dataset cleaning has been done before commencing the process of modelling. This cleaning process was conducted manually, involving the removal of mentions, contractions, hyperlinks, emoticons, some punctuations, and white spaces. We also corrected common misspelled data.

```
In [20]: 1 clean_text("@ttt hasn't the sun become warmer? https://thisanexample ☀")
Out[20]: 'has not the sun become warmer? sun with face'
```

Figure 3.9: An Example of tweet pre-processed. This is a fictional tweet.

Data visualisation of the Twitter Dataset

Before modelling, we performed an exploratory analysis on the dataset to better understand the cleaned dataset and to identify any potential issues with data quality. Furthermore, we gathered additional data, such as location and the average number of words and characters, to account for outliers. We found 1899 non-null locations, which represents almost 20% of the dataset. To extract the location of the tweets, we utilised the **Nominatim** library by using the user’s description location and plotted the geographical distribution of tweets through a **Matplotlib** map. To further examine the distribution of the most common words, we created a Treemap and compiled a table of the most frequent and rare words. Additionally, we conducted a small textual analysis using open-source libraries. Initially, we created a thematic analysis using the open-source library, **Genism**, which grouped the data into four sub-themes using the generative probabilistic model known as Latent Dirichlet Allocation (LDA). The equation of the LDA model is : where :

$$p(z_{i,j}|z_{-i,j}, w) \propto \frac{(n_{i,\cdot}^{(k)} + \alpha)(n_{\cdot,j}^{(k)} + \beta)}{\sum_{k'} (n_{i,\cdot}^{(k')} + \alpha)(n_{\cdot,j}^{(k')} + \beta)} \quad (3.3)$$

where:

- N is the number of words in the document
- M is the number of documents in the corpus
- K is the number of topics
- α is the hyperparameter for the per-document topic distribution
- β is the hyperparameter for the per-topic word distribution
- z is the topic assignment for each word in the document
- w is the specific word

After that, we ran an overview of the Sentiment Score of the Twitter dataset using the **VaderSentiment** library. Valence Aware Dictionary for Sentiment Reasoning (VADER) is a text sentiment analysis model that is sensitive to both polarity (Positive/Negative) and intensity (Strength) of sentiment.

Chapter 4

Results

4.1 Results of the exploratory data analysis

This preliminary analysis yielded several key insights into the dataset. Firstly, Figure 4.1 shows the geographical distribution of the tweets we retrieved, with missing values shaded in. It revealed that half of the tweets with location data were from the United Kingdom and the United States, while India, Kenya, and Canada were also important sources.

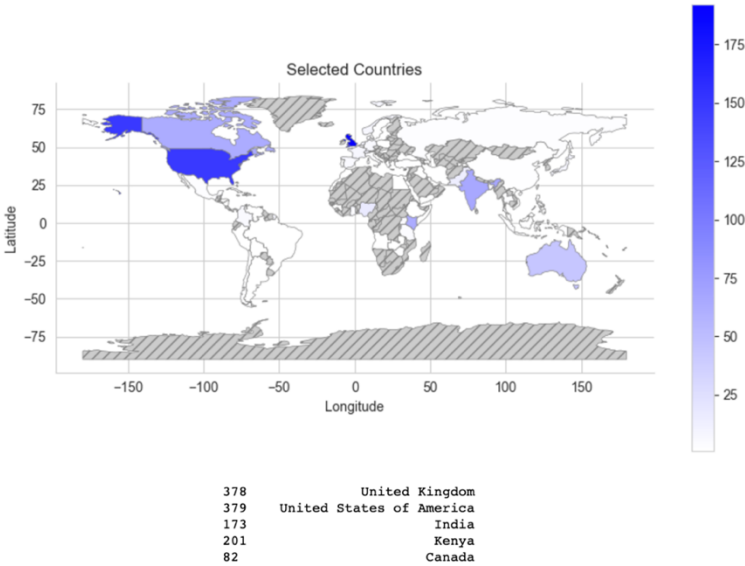


Figure 4.1: Geographical Distribution of Retrieved Tweets from the column “Location”, which contain 1899 non-null values.

Secondly, we found that the average tweet length was 28 words, with 178 characters on average. The Boxplot in Figure 4.2 shows the distribution of the word count, with a few outliers that we carefully analysed before removing some of them to improve the accuracy of the model.

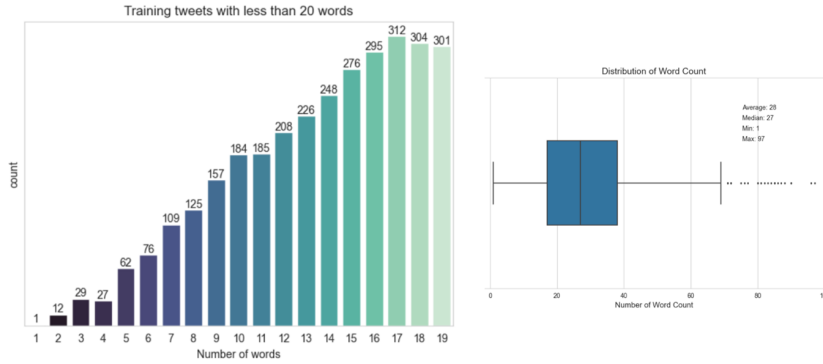


Figure 4.2: Tweet word count analysis: Frequency of tweets with 20 words or less (Left) and Boxplot of overall word count distribution (Right).

Thirdly, the Tree Map in Figure 4.3 revealed that the word "climate" appeared in more than 60% of the tweets, which was expected given that it was one of our main keywords. Words like "us" and "people" highlighted the sense of responsibility felt by users, while the term "global" emphasized the worldwide impact of the climate crisis. However, it is important to be cautious when interpreting these words, as their meaning can vary depending on context and cultural perspectives.

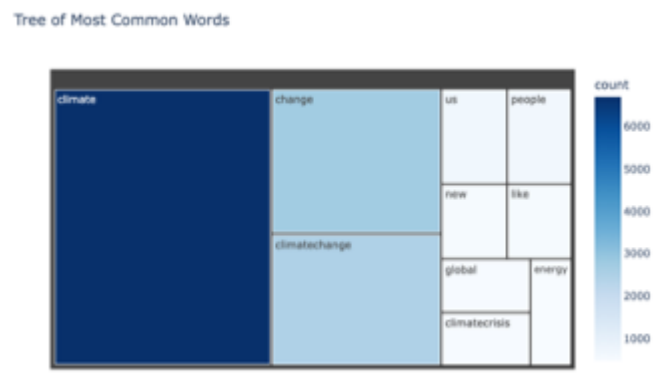


Figure 4.3: Tree representing the distribution of the most common words in the dataset. The size of the box represents its frequency.

climate	6694	climateregulations	1
change	2736	revising	1
climatechange	2484	gorgonproject	1
us	698	boundarydamproject	1
people	680	petranovaproject	1
new	552	sleipnerproject	1
like	539	ghgsemissions	1
global	528	carbonstorage	1
climatecrisis	526	eradicating	1
energy	485	transportforum	1

Figure 4.4: Comparison of Word Frequency in the Dataset: The table on the right displays the most frequently occurring words, while the table on the left depicts some infrequent words in the dataset.

Fourthly, Figure 4.5 displays the initial word cloud which showcases the themes recognized by the LDA model. Through analysing the co-occurrence of words in the dictionary, the model has identified the four key topics. The first topic centres around nature and urgency, while the second topic emphasizes the significance of individuals. The third topic seems to highlight the fundamental principles that drive change, while the fourth topic appears to be more closely related to weather and climate conditions. As for the final topic, it remains challenging to analyse. Finally, using the Vader Sentiment library, we found that 44.6% of tweets were labelled as positive, 35.6% as negative, and 19.8% as neutral. The histogram of the Sentiment Distribution has a peak value around the Sentiment Score of 0,1 indicating that most Positive tweets almost belong to the ‘Neutral’ category.



Figure 4.5: Word cloud representing the topics identified by the LDA model trained on one of the dictionaries of the library ‘gensim’. The code for this figure can be find in Appendix A.4. Its uses LDA model.

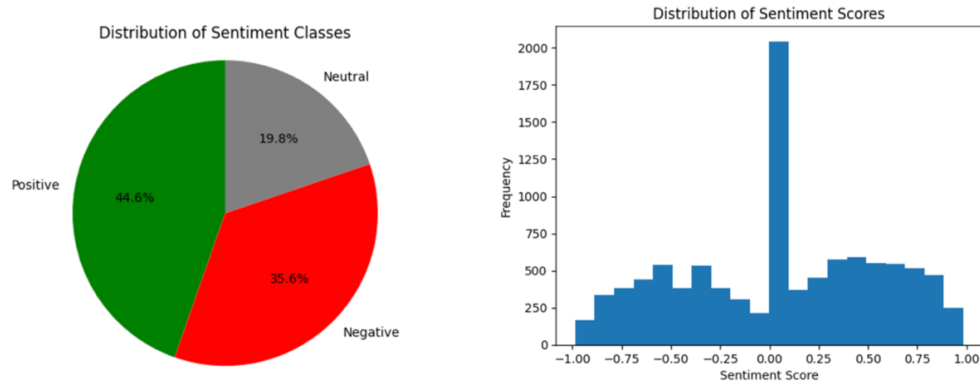


Figure 4.6: Pie chart and histogram illustrating the distribution of the sentiments using the 'VaderSentiment' library.

Overall, this preliminary analysis provided a glimpse into the opinions and attitudes towards climate change expressed on Twitter, which could be further explored in a more detailed analysis of emotions.

4.2 Comparison of LSTM and LSTM with Glove Performance

Training process.

Experiments were carried out to assess the effectiveness of supervised learning in categorising emotions using two models, each consisting of 5 cycles (Epoch). Each Epoch comprised 32 batches, with a portion of the dataset used to train the neural network. The training and validation loss curves (depicted in Figure 4.7) were used to evaluate the learning process of both classifiers. These curves allowed us to assess if the model was overfitting - meaning it was learning from noisy data. The graph on the right indicates that while the accuracy of the training set increased with each Epoch, it eventually plateaued around 0.36 before dropping around Epoch 2. This suggests that the model was memorizing the training data and not adapting well to new data, indicating an overfitting issue. Conversely, the graph on the left demonstrates that the training has converged,

and the increasing validation and training accuracies show good generalization ability for unseen data. In the case of the LSTM with GloVe model, it is learning the underlying patterns in the dataset without overfitting or underfitting.

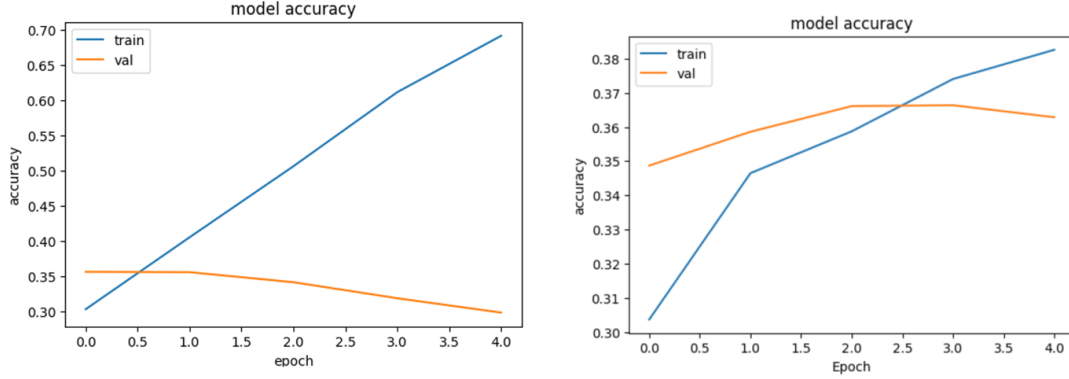


Figure 4.7: Plot of the training and validation curves (accuracy vs.epoch) for the LSTM model with Keras embeddings (Right) and LSTM model with GloVe Embeddings (Left).

Performance

When comparing the overall performance of models, we picked the weighted F1 score as it considers the number of samples in each class, while the macro F1 score treats all classes equally, regardless of their size. This choice was justified in the **Methodology** section. The report tables below provide a summary of the results for each model. On the one hand, for the LSTM model, the overall macro weighted average F1 score was equal to 0.28 and the accuracy was equal to 0.30. On the other hand, for the LSTM model with GloVe embeddings, the weighted F1 score reached 0.33 and the accuracy was 0.36. From the table 4.8, we can conclude that the LSTM with Glove embeddings model is a better classifier for the Twitter dataset. The highest F1 score is obtained for the “Neutral” class, which is 0.46.

Classification Report for LSTM With Keras Embeddings:					Classification Report for LSTM With Glove Embeddings:				
	precision	recall	f1-score	support		precision	recall	f1-score	support
Empty	0.07	0.03	0.04	160	Empty	0.00	0.00	0.00	160
Sadness	0.27	0.36	0.31	1033	Sadness	0.35	0.29	0.32	1033
Enthusiasm	0.00	0.00	0.00	168	Enthusiasm	0.00	0.00	0.00	168
Neutral	0.34	0.47	0.39	1739	Neutral	0.37	0.61	0.46	1739
Worry	0.34	0.26	0.29	1673	Worry	0.36	0.39	0.37	1673
Surprise	0.10	0.07	0.08	425	Surprise	0.19	0.01	0.02	425
Love	0.37	0.38	0.38	791	Love	0.51	0.41	0.45	791
Fun	0.10	0.07	0.08	362	Fun	0.18	0.03	0.05	362
Hate	0.20	0.17	0.18	244	Hate	0.30	0.24	0.26	244
Happiness	0.31	0.32	0.32	1044	Happiness	0.33	0.46	0.38	1044
Boredom	0.00	0.00	0.00	34	Boredom	0.14	0.03	0.05	34
Relief	0.12	0.07	0.09	284	Relief	0.17	0.01	0.03	284
Anger	0.00	0.00	0.00	27	Anger	0.00	0.00	0.00	27
accuracy			0.30	7984	accuracy			0.36	7984
macro avg	0.17	0.17	0.17	7984	macro avg	0.22	0.19	0.18	7984
weighted avg	0.28	0.30	0.28	7984	weighted avg	0.33	0.36	0.33	7984

Figure 4.8: Report Tables comparing Precision, Recall, and F1-score Values of LSTM Models with Keras (Right) and GloVe Embeddings (Left)

In conclusion, the model LSTM with GloVe is a better classifier for the Twitter dataset than the LSTM with the pre-build embeddings from the Keras library. Consequently, it is the one we will use for the in-depth analysis.

4.3 In-depth analysis of negative sentiments using LSTM-GloVe

Multi-class labelled emotions.

Upon training our classifier, the retrieved dataset was processed using LSTM with GloVe. Each tweet was given a percentage for each of the 13 emotions. The resulting pie chart was generated by determining the majority label and majority labels of each tweet. Additionally, in Appendix A.11, we included two unique and fictional tweets to examine this classification in greater detail. Using LSTM-GloVe, 46.8% of the tweets were categorized as "Neutral," while 41.2% were labelled as "Worry." Among the emotions identified in the tweets, "Happiness" was represented in 5.8% and "Sadness" in 3.4% of them. In most tweets, "Boredom" was the least frequent emotion assigned, while "Anger," "Love," and "Enthusiasm" were also not commonly detected.

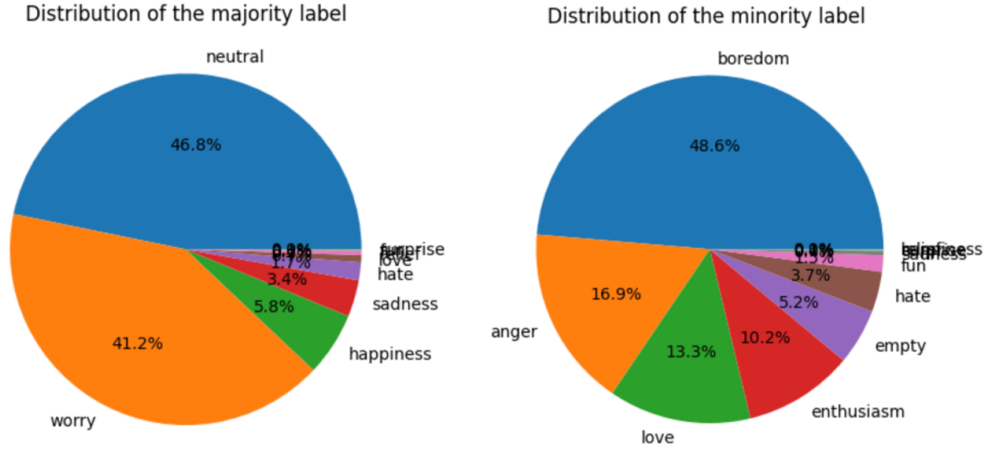


Figure 4.9: Pie plots of the majority label predicted, and the minority label predicted from LSTM datasets.

Focus on negative emotions. We decided to group the eight different negative emotions in our labels under a new corpus to enhance clarity. These negative emotions include "Empty", "Sadness", "Hate", "Neutral", "Worry", "Boredom", "Relief", and "Anger". Additionally, we grouped the remaining five emotions under the Positive dataset, which includes tweets labelled as "Enthusiasm", "Surprise", "Love", "Fun", and "Happiness". By creating a word cloud in Figure 4.10, we identified the 100 most frequent words in both corpora. The red words indicate the most frequent word in the 'Positive' corpus, whereas the black labels show the most frequent words in the negative corpus. When a word is present in both corpora, it appears in blue. It can be observed that the term "research" is commonly associated with Positive discourse, while "leaders" is often linked to negative messages. The word "scientists" is a recurring term in both Positive and Negative discourses.

Chapter 5

Discussion, Limitations and Suggestions

5.1 Discussion

This research project represents the first attempt to analyse online "emotional" personal discourses on climate change using automated methods. We built two supervised machine learning models based on emotion detection and compared their effectiveness and strength through experimental comparisons of two different neural network embedding architectures. Our results confirm the importance of considering word embedding structure for text classification tasks, demonstrating that the choice of neural network model should not be the only focus of sentiment analysis. Other parameters, such as the training dataset and model hyperparameters, play a key role in model accuracy and should be evaluated carefully. The low F1 score obtained underscores the subjectivity of emotional expression in brief messages and the challenge of identifying emotions for algorithms.

Our analysis yielded contradictory results for the representation of "Negative" and "Positive" sentiment when using the open-sourced library **VaderSentiment** and our trained classifier. While Vader Sentiment is recommended for social media analyses, it is not specifically trained on tweets, which have a unique language due to the platform's char-

acter limit and the prevalence of slang and acronyms. Thus, we suggest using a model trained on a platform’s corpus when analysing social media data. Furthermore, Vader Sentiment’s Sentiment Score is not particularly useful for categorizing sentiment into categories, as it is often neutral, but the category of Neutral values is relatively low.

Our analysis confirms that worry is prevalent among those discussing the climate crisis, especially among young people. Although most tweets analysed were labelled as neutral, we collected a large database of expressions on the matter that challenges studies that suggest a lack of engagement and discussions around climate change. We found that negative emotions expressed towards climate change are multi-faceted, with most tweets analysed expressing more than three emotions. This supports the method of multi-labelling to capture the range of emotions in Twitter users’ discourse. Additionally, knowing the difference between anxiety and eco-anger, we were unable to comprehend because of the sample of labelled tweets as ”Anger”. It is essential to distinguish between these two negative emotions in order to more accurately target communications to elicit action (anger) rather than inaction (anxiety). Thus, we encourage additional research on this topic. Finally, we suggest that identifying relevant vocabulary for distressing posts could improve communication about climate change.

5.2 Limitations

Limitations of the data analysis

The issue of the class imbalance is the data analysis’s most obvious weakness. The confusion matrices were observed, and one is displayed in Fig. X. These observations showed that the LSTM worried,correctly,tweets, models are biased towards categories with a lot of training instances (such as concern and neutral), whereas they perform badly for categories with very few examples (such as boredom). We can see that most tweets when they are not predicted correctly are classified as neutral, worry or happy. Before training the deep neural network, one might include morphological changes, synonyms, or derivatives of words to help with this. The performance of the models would have improved by

reducing the number of target classes by combining certain labels.

LSTM Sentiment Analysis
Confusion Matrix

	empty	sadness	enthusiasm	neutral	worry	surprise	love	fun	hate	happiness	boredom	relief	anger
empty	0	11	0	87	36	0	2	1	3	16	3	1	0
sadness	0	298	0	230	359	1	23	3	40	77	1	1	0
enthusiasm	0	11	0	77	36	1	9	0	1	33	0	0	0
neutral	0	84	1	1057	299	5	46	10	23	209	1	4	0
worry	1	257	0	516	653	4	53	8	45	130	0	6	0
surprise	0	42	0	171	95	4	21	2	3	87	0	0	0
love	0	41	0	123	54	0	326	10	8	226	0	3	0
fun	0	17	0	134	49	3	15	11	2	131	0	0	0
hate	0	34	0	53	83	0	2	1	58	12	1	0	0
happiness	0	36	0	292	82	3	122	13	6	485	0	5	0
boredom	0	10	0	7	10	0	1	0	2	3	1	0	0
relief	0	14	0	114	46	0	21	2	2	81	0	4	0
anger	0	1	0	12	8	0	1	0	3	2	0	0	0
	empty	sadness	enthusiasm	neutral	worry	surprise	love	fun	hate	happiness	boredom	relief	anger

Predicted Label

Figure 5.1: Confusion Matrix comparing the actual and predicted labels for the LSTM model with Glove embeddings.

The F1-score per class can be computed using the F1-score formula described in the **Methodology**. Using the "neutral" class as an example, we can calculate the true positive (TP), false positive (FP), and false negative (FN) values as follows: $TP = 1057$ (the (10, 5) element in the matrix), $FP = 682$ (the first row without the neutral element), and $FN = 1816$ (the first column without the neutral element). We can then compute the precision and recall for the "neutral" class as follows: $Precision = TP / (TP + FP)$ and $Recall = TP / (TP + FN)$. Substituting the values we obtained, we have $Precision = 1057 / (1057 + 682)$ and $Recall = 1057 / (1057 + 1816)$. We have $F1\ score \approx 0.458$. This result is verified by the Table of Report in the **Results** section.

Additionally, regarding noise, since no humans were involved in the classification, we could not properly address humorous, sarcastic, or ironic tweets and did not portray any emotional content. Expressions and opinions about social and political issues on social media are much harder because they often need analysis at the pragmatic level, which can be difficult to teach without background knowledge to the classifier.

Moreover, as underlined in the literature review, climate change is categorised online by its polarisation, often providing fertile ground for misinformation to spread. From our initial dataset, we were only able to use 25% of the values questioning the data quality of the retrieved dataset. Due to the importance of opinions on social media, the prevalence of bot intervention has increased over the past decade, and detecting spam is a challenging task for our model if we do not provide the recognition pattern of spam tweets. Filtering out these posts or users is a challenging task that relies heavily on qualitative human-crafted datasets of sentiment vocabulary and pre-classified “ground truths”. An essential restriction to note is that not all content can be labelled as emotional content, and some content, particularly news-related content, can be informative. Finally, sentiment and lexicons play an essential role in developing efficient sentiment analysis. Different people can interpret different emotions in various ways. It is possible that the training dataset contains incorrect labeling, thus impacting the ability of the LSTM network to learn from its pattern representations. To sum up, the difficulty of understanding context, sarcasm, an imbalance in class size, ambiguity in natural language, and rapidly expanding Internet slang and acronyms for tweet sentiment analysis further complicate the task of understanding emotions.

Limitations of Twitter as a reliable research source.

When studying emotions and opinion, researchers face a trade-off between data quality and data newness. As researchers have more control of which participants to recruit and which questions are asked, data quality is more ensured in surveys and experiments. Yet, it is often costly data, especially if a large sample of panel data is needed. In their article, Chen et al. highlight that “newness” is a strength of social media data and is especially

useful for studying emerging life sciences issues where the right questions to ask are elusive (2023).

Other sources of bias may also include the quality of search terms for data retrieval as well as the quality of Twitter API and third platforms in returning representative results. Indeed, the free API does not return the same results as the paid options access. The discrepancies between free and paid options in return results have alarmed researchers as it is high enough to produce different results in content analysis and hiding the full descriptive statics of issue frames represented in the Twitter verse (Chen, Duan, and Yang 2023). Despite its popularity as a source of accessible digital trace data, the sample of the general Twitter data is still debated and appears highly contingent on the research question. We must express caution about extrapolating Twitter users' opinions, such as the level of worry towards the environmental crisis Twitter users are younger and "only 10% of users created 80%" (Chen, Duan, and Yang 2023,p.122) thus the validity of generalisations needs to be assessed along with the biases within political discussions in Twitter. Additionally, it is important to note that social media behaviour is not the same as behaviour in the physical world.

5.3 Suggestions

Given the scope of this dissertation, it was necessary to make certain assumptions and simplifications to develop a relevant and engaging research topic. Unfortunately, my laptop's limited memory capacity prevented me from utilising more advanced computational models, like Bert or Roberta, for analysing textual data (Sirisha and Chandana 2022) . Nonetheless, exploring the potential of training different models for future research may lead to improved accuracy in emotion classification tasks. Additionally, a promising avenue for future research involves expanding the model to include other climate-relevant positive emotions, such as hope (Bouman et al. 2020). As stated in the **Introduction**, concern may be less associated with personal actions and may be less experimental and personal. It would also be fascinating to see how theories of behaviour change can be

translated into computational methods, like classifying tweets based on users' behavioural stages using the 5 Doors Theory (Fernandez et al. 2016) (e.g, understanding, conviction).

Chapter 6

Conclusion

Despite the urgent warnings being disseminated by the scientific community, climate change does not appear to be a priority for a significant portion of the global population. This discrepancy between scientific understanding and public sentiment engagement has led to this research project. The literature on the increasing disconnect between individuals and climate change posits the primary hypothesis of eco-anxiety from the standpoint of climate psychology. As a result, scholarly investigations in this domain have predominantly centred on the manifestation of anxiousness pertaining to climate change, frequently extending this ecological anxiety to negative emotions.

This paper conducted an analysis of the various forms of discourse surrounding climate change through sentiment analysis. This was achieved by implementing an LSTM classifier to scrutinise tweets. Notably, this marks the first instance in which such an approach has been employed. The findings indicated that while worry was the most commonly recognised emotion, other emotional states such as anger, sadness, or frustration were also widely observed in the discourse of users pertaining to the subject matter.

The objective of this study was to evaluate the range of perspectives regarding climate change through the development of a sentiment classifier using sentiment analysis techniques applied to tweets. This marks the initial application of sentiment analysis in this context. Our findings indicate that while worry was the most reported emotion, users also

frequently expressed feelings of sadness, anger, and frustration in their discussions on the subject matter. Utilising a multi-classifier labelling approach was essential in achieving a comprehensive comprehension of the emotional spectrum in this analysis.

To assess the efficacy of our model, we developed and conducted an analysis on two distinct supervised LSTM models trained on a Twitter corpus. The utilisation of a pre-labelled training dataset tailored to the lexicon of tweets was essential in enhancing the precision of our analysis, considering the shortness in characters of tweets and the incorporation of slang language. Furthermore, our study highlighted the importance of meticulously selecting the input embedding for NLP classification performance, as evidenced by the differentiation of our models based solely on their embeddings, with one utilising Glove embedding.

Using data visualisation tools, we proceeded to conduct additional analysis on the particular lexicon and themes that were identified as conveying negative emotional content. Furthermore, we are furnishing the source code for our analysis and urging readers to develop a multi-label classifier pertaining to a selected topic on Twitter.

While our approach has some general and sometimes unavoidable limitations, the most significant one is the low F1-score, which reveals the strong link between the classification of our emotion and the distribution of the training dataset. In addition, the subjective nature of emotion poses a significant challenge for both human annotators and machines in determining the accurate distribution of the "ground" truth. To enhance the validity of our findings, it would be advantageous to train our classifier on a larger dataset and employ NLP models such as BERT or GPT that require significant computational resources. The present study offers a preliminary structure for examining the interplay between emotions and climate change discourse on social media. We believe that researching negative emotions is one way to channel them effectively. Empirical data indicates that the dissemination of resilient models or innovative solutions may facilitate proactive measures. It is anticipated that the identification of keywords and themes in this analysis will mitigate the incapacitating effects of negative emotions through the promotion of more innovative

communication tactics that are less fixated on fear and exigency. The findings of this study hold significance for scholars within the discipline and have the potential to stimulate cross-disciplinary partnerships. Additionally, governmental, and non-governmental organisations may benefit from these results by improving their communication strategies and facilitating the shift towards a sustainable future.

Word count : 8436 words.

Appendix A

Appendix

The following section contains the code appendix, which provides a comprehensive listing of the key programming code used in the project as well as a comprehensive reference guide on key terminology.

A.1 Appendix : Code

```

1 def scrape(words, numtweet, scraped_data):
2     tweets_per_request = 100 # Tweepy allows up to 100 tweets per request
3
4     tweets = tweepy.Cursor(api.search_tweets, q=words, lang="en",
5                             tweet_mode='extended').items(numtweet)
6
7     list_tweets = [tweet for tweet in tweets]
8
9     i = 1
10
11     # we will iterate over each tweet in the list for extracting information about each tweet
12     for tweet in list_tweets:
13         description = tweet.user.description
14         location = tweet.user.location
15         following = tweet.user.friends_count
16         followers = tweet.user.followers_count
17         totaltweets = tweet.user.statuses_count
18         retweetcount = tweet.retweet_count
19         hashtags = tweet.entities['hashtags']
20         date = tweet.created_at
21
22         try:
23             text = tweet.retweeted_status.full_text
24         except AttributeError:
25             text = tweet.full_text
26         hashtext = list()
27         for j in range(0, len(hashtags)):
28             hashtext.append(hashtags[j]['text'])
29
30         ith_tweet = [description, location, following,
31                     followers, totaltweets, retweetcount, text, hashtext, date]
32
33         printtweetdata(i, ith_tweet, scraped_data)
34         i = i+1

```

```

1 print("Enter Twitter HashTag to search for")
2 words = input()
3
4 print("Enter Number of tweets to be scraped")
5 numtweet = int(input())
6 scraped_data = []
7 print("Fetching tweets...")
8 scrape(words, numtweet, scraped_data)
9 print('Scraping has completed!')

```

```

Enter Twitter HashTag to search for
climate
Enter Number of tweets to be scraped
2
Fetching tweets...

```

Figure A.1: Code 1 : Retrieving twitter dataset through hashtag filters

```

import spacy
from geopy.geocoders import Nominatim
from geopy.extra.rate_limiter import RateLimiter
from tqdm.notebook import tqdm

nlp = spacy.load("en_core_web_sm")
geolocator = Nominatim(user_agent="my-custom-user-agent")
geocode = RateLimiter(geolocator.geocode, min_delay_seconds=5, max_retries=5)

# update the extracted_user_location column directly when geocoding is successful or not
for i, location in tqdm(df3['extracted_user_location'].items()):
    if location is not np.nan:
        try:
            geocoded_location = geocode(location, language="en")
            if geocoded_location is not None:
                df3.at[i, 'location'] = geocoded_location
                df3.at[i, 'coordinates'] = tuple(geocoded_location.point)
                df3.at[i, 'state'] = geocoded_location.address.split(',')[0]
                df3.at[i, 'country'] = geocoded_location.address.split(',')[1]
            else:
                # if location not found, add a zero value to the list
                df3.at[i, 'location'] = 0
        except:
            # if there is any exception, add a zero value to the list
            df3.at[i, 'location'] = 0

# save the dataframe to CSV after each iteration
df3.to_csv("updated_data_2.csv", index=False)

```

Figure A.2: Code 2 : Geolocating tweets using user profile location.

```

import pandas as pd
import geopandas as gpd
import matplotlib.pyplot as plt
from matplotlib.colors import LinearSegmentedColormap

# Load world countries shapefile data using geopandas
world = gpd.read_file(gpd.datasets.get_path('naturalearth_lowres'))

# Count the number of times each country appears in the data column
counts = df['country'].value_counts()
counts = counts.rename_axis('name').reset_index(name='count')

# Merge the counts with the world shapefile data based on country name
merged = world.merge(counts, on='name', how='left')

colors = []
for i in range(10):
    color = (i / 9, i / 9, 1.0)
    colors.append(color)
cmap = LinearSegmentedColormap.from_list("", list(reversed(colors)))

# Count the number of times each country appears in the data column
counts = df['country'].value_counts()
counts = counts.rename_axis('name').reset_index(name='count')

# Merge the counts with the world shapefile data based on country name
# Set the color for missing features (i.e. countries without a value in the count column)
missing_color = '#cccccc'

# Plot a choropleth map of selected countries based on the count column
ax = merged.plot(column='count', cmap=cmap, legend=True, edgecolor='gray', linewidth=0.5, figsize=(10, 6), missing=missing_color)

# Set plot title and axis labels
ax.set_title('Selected Countries')
ax.set_xlabel('Longitude')
ax.set_ylabel('Latitude')

# Remove the side and top spines from the plot
ax.spines['right'].set_visible(False)
ax.spines['top'].set_visible(False)

# Adjust the legend font size and position
leg = ax.get_legend()
if leg is not None:
    leg.set_bbox_to_anchor((1.2, 1))
    for txt in leg.get_texts():
        txt.set_fontsize('large')

# Show the map
plt.show()

```

Figure A.3: Code 3: Generating a map showing the geographical distribution of tweets by country

```

import pandas as pd
import nltk
from nltk.tokenize import word_tokenize
from nltk.corpus import stopwords
from gensim import corpora, models
import matplotlib.pyplot as plt
from wordcloud import WordCloud

# preprocess the text in the 'clean_text' column
nltk.download('stopwords')
nltk.download('punkt')
stop_words = set(stopwords.words('english'))

def preprocess_text(text):
    tokens = word_tokenize(text)
    filtered_tokens = [word for word in tokens if not word.lower() in stop_words and word.isalpha()]
    return filtered_tokens

df3['clean_tokens'] = df3['clean_msg'].apply(preprocess_text)

# identify themes in each row of the 'clean_tokens' column
dictionary = corpora.Dictionary(df3['clean_tokens'])
corpus = [dictionary.doc2bow(tokens) for tokens in df3['clean_tokens']]

lda_model = models.LdaModel(corpus, num_topics=5, id2word=dictionary, passes=10)

df3['topics'] = lda_model[corpus]

# get the top words for each topic
top_words_per_topic = []
for i, topic in lda_model.show_topics(num_topics=-1, formatted=False):
    topic_words = [word for word, _ in topic]
    top_words_per_topic.append(topic_words)

# plot word clouds for each topic
num_topics = len(top_words_per_topic)
fig, axes = plt.subplots(nrows=1, ncols=num_topics, figsize=(20, 10), sharex=True, sharey=True)
for i in range(num_topics):
    wordcloud = WordCloud(width=800, height=400, background_color='white', random_state=42).generate(' '.join(top_words_per_topic[i]))
    axes[i].imshow(wordcloud, interpolation='bilinear')
    axes[i].set_title(f'Topic {i+1}', fontsize=20)
    axes[i].axis('off')

plt.tight_layout()
plt.show()

```

Figure A.4: Code 4 : Generating the Topics' Wordcloud

```

import matplotlib.pyplot as plt
from sklearn.manifold import TSNE
import numpy as np
from collections import Counter

def tsne_plot(model, sentences, top_n=50):
    """
    Function to create and visualize TSNE plot of word vectors.
    Args:
        model: gensim word2vec or GloVe model
        sentences: list of sentences to plot
        top_n: number of top words to select based on frequency
    """

    # Create a frequency distribution of all words in the sentences
    all_words = [word for sentence in sentences for word in sentence if word in model.key_to_index]
    word_freq = Counter(all_words)

    # Select the top n words based on frequency
    top_words = [word for word, freq in word_freq.most_common(top_n)]

    # Extract vectors and labels for the top words
    tokens = []
    labels = []
    for word in top_words:
        tokens.append(model[word])
        labels.append(word)

    # Convert tokens to a 2D array
    tokens = np.array(tokens)

    # Create a TSNE model with 2 dimensions
    tsne_model = TSNE(perplexity=40, n_components=2, init='pca', n_iter=2500, random_state=23)

    # Fit and transform the vectors using TSNE
    tsne_vectors = tsne_model.fit_transform(tokens)

    # Create a scatter plot of the TSNE vectors
    plt.figure(figsize=(16, 10))
    plt.scatter(tsne_vectors[:, 0], tsne_vectors[:, 1], c='steelblue', edgecolors='k')
    for label, x, y in zip(labels, tsne_vectors[:, 0], tsne_vectors[:, 1]):
        plt.annotate(label, xy=(x, y), xytext=(5, 2), textcoords='offset points', ha='right', va='bottom')

    plt.title('TSNE Plot of Word Vectors')
    plt.xlabel('TSNE Dimension 1')
    plt.ylabel('TSNE Dimension 2')
    plt.show()

```

Figure A.5: Code 5 : Function generating the TSNE Plot of Word embedded using Glove

```

from vaderSentiment.vaderSentiment import SentimentIntensityAnalyzer
import pandas as pd
import matplotlib.pyplot as plt

# Initialize the sentiment analyzer
analyzer = SentimentIntensityAnalyzer()

# Define a function to classify the sentiment of each message
def get_sentiment_class(msg):
    score = analyzer.polarity_scores(msg)['compound']
    if score > 0.05:
        return 'Positive'
    elif score < -0.05:
        return 'Negative'
    else:
        return 'Neutral'

# Apply the function to the clean_msg column of the DataFrame to get the sentiment classes
sentiment_classes = df['clean_msg'].apply(get_sentiment_class)

# Count the number of messages in each sentiment class
sentiment_counts = sentiment_classes.value_counts()

# Create a pie chart to show the distribution of positive, negative, and neutral sentiments
labels = sentiment_counts.index
sizes = sentiment_counts.values
colors = ['green', 'red', 'grey']
plt.pie(sizes, labels=labels, colors=colors, autopct='%1.1f%%', startangle=90)
plt.axis('equal')
plt.title('Distribution of Sentiment Classes')
plt.show()

```

Figure A.6: Code 6 : Creating a pie chart to show the distribution of Positive, Negative, and Neutral sentiments from VaderSentiment library

```

# extract subject
source = [i[0] for i in entity_pairs]

# extract object
target = [i[1] for i in entity_pairs]

kg_df = pd.DataFrame({'source':source, 'target':target, 'edge':relations})

# count the number of occurrences of each edge
edge_counts = kg_df.groupby(['source', 'target']).size().reset_index(name='count')

# keep only the edges with the highest counts
edge_counts = edge_counts.sort_values('count', ascending=False).head(40)

# create a directed-graph from a dataframe
G = nx.from_pandas_edgelist(edge_counts, "source", "target", edge_attr=True, create_using=nx.MultiDiGraph())

# plot the directed graph
plt.figure(figsize=(12, 8))
pos = nx.spring_layout(G, k=0.5)
nx.draw_networkx_nodes(G, pos, node_color='lightblue', node_size=500)
nx.draw_networkx_edges(G, pos, edge_color='grey', arrowsize=10, width=1)
nx.draw_networkx_labels(G, pos, font_size=10, font_family='sans-serif')
plt.axis('off')
plt.show()

```

Figure A.7: Code 7: Creating a Knowledge Graph

```

from wordcloud import WordCloud, STOPWORDS
import matplotlib.pyplot as plt

# Filter the dataframe to only include rows with negative emotions
df_neg = df2[(df2.max_label == "empty") | (df2.max_label == "sadness") | (df2.max_label == "hate") |
             (df2.max_label == "neutral") | (df2.max_label == "worry") | (df2.max_label == "boredom") |
             (df2.max_label == "relief") | (df2.max_label == "anger")]

# Join the text from the filtered dataframe into a single string
text_neg = " ".join(df_neg["text"])

# Filter the dataframe to only include rows with positive emotions
df_pos = df2[(df2.max_label == "enthusiasm") | (df2.max_label == "surprise") | (df2.max_label == "love") |
             (df2.max_label == "fun") | (df2.max_label == "happiness")]

# Join the text from the filtered dataframe into a single string
text_pos = " ".join(df_pos["text"])

# import necessary packages
import matplotlib.pyplot as plt
from wordcloud import WordCloud, STOPWORDS

# import necessary packages
import matplotlib.pyplot as plt
from wordcloud import WordCloud, STOPWORDS

# define stopwords to remove from the text
stopwords = set(STOPWORDS)
stopwords.update({"amp", "im", "u", "will"})

# create a WordCloud object for positive and negative sentiment words combined
wordcloud = WordCloud(width = 800, height = 800,
                      background_color = 'white',
                      stopwords = stopwords,
                      min_font_size = 10)

# generate wordcloud for positive sentiment words
wordcloud.generate_from_text(text_pos)
positive_words = set(list(wordcloud.words_.keys())[:100])

# generate wordcloud for negative sentiment words
wordcloud.generate_from_text(text_neg)
negative_words = set(list(wordcloud.words_.keys())[:100])

# create a dictionary of colors for positive, negative, and overlapping words
color_dict = {}
for word in positive_words - negative_words:
    color_dict[word] = "red"
for word in negative_words - positive_words:
    color_dict[word] = "black"
for word in positive_words & negative_words:
    color_dict[word] = "blue"

# create a WordCloud object for both positive and negative sentiment words combined
wordcloud = WordCloud(width = 800, height = 800,
                      background_color = 'white',
                      stopwords = stopwords,
                      min_font_size = 10,
                      color_func=lambda *args, **kwargs: color_dict.get(args[0], "black")).generate(text_pos + " " + text_neg)

# plot the WordCloud for both positive and negative sentiment words combined
plt.figure(figsize = (8, 8), facecolor = None)
plt.imshow(wordcloud)
plt.axis("off")
plt.tight_layout(pad = 0)
plt.title("Frequently Occurring Words in Positive and Negative Labeled Tweets, with Some Words Common in Both Categories")
plt.show()

```

Figure A.8: Code 8: Wordcloud for both positive and negative sentiment words combined

```

1 embed_dim = 200
2 lstm_out = 250
3
4 model_lstm_gve = Sequential()
5 model_lstm_gve.add(Embedding(len(w_idx) + 1, embed_dim, input_length = X_test_pad.shape[1], weights=[embedding_matrix]))
6 model_lstm_gve.add(SpatialDropout1D(0.2))
7 model_lstm_gve.add(LSTM(lstm_out, dropout=0.2, recurrent_dropout=0.2))
8 model_lstm_gve.add(keras.layers.Dense(13, activation='softmax'))
9 #adam rmsprop
10 model_lstm_gve.compile(loss = "categorical_crossentropy", optimizer='adam', metrics = ['accuracy'])
11 print(model_lstm_gve.summary())

```

Model: "sequential_1"

Layer (type)	Output Shape	Param #
embedding_1 (Embedding)	(None, 160, 200)	6044000
spatial_dropout1d_1 (SpatialDropout1D)	(None, 160, 200)	0
lstm_1 (LSTM)	(None, 250)	451000
dense_1 (Dense)	(None, 13)	3263

=====
Total params: 6,498,263
Trainable params: 454,263
Non-trainable params: 6,044,000
None

Figure A.9: Code 9: Training of Model 1 - LSTM

```

1 embed_dim = 160
2 lstm_out = 250
3
4 model = Sequential()
5 model.add(Embedding(len(w_idx) + 1, embed_dim, input_length = X_test_pad.shape[1]))
6 model.add(SpatialDropout1D(0.2))
7 model.add(LSTM(lstm_out, dropout=0.2, recurrent_dropout=0.2))
8 model.add(keras.layers.Dense(13, activation='softmax'))
9 #adam rmsprop
10 model.compile(loss = "categorical_crossentropy", optimizer='adam', metrics = ['accuracy'])
11 print(model.summary())

```

Model: "sequential"

Layer (type)	Output Shape	Param #
embedding (Embedding)	(None, 160, 160)	4035200
spatial_dropout1d (SpatialDropout1D)	(None, 160, 160)	0
lstm (LSTM)	(None, 250)	411000
dense (Dense)	(None, 13)	3263

=====
Total params: 5,249,463
Trainable params: 5,249,463
Non-trainable params: 0
None

Figure A.10: Code 10: Training of Model 2 - LSTM with GloVe

```

def conf_matrix(y, y_pred, title):
    fig, ax = plt.subplots(figsize=(17,17))
    labels = ["empty", "sadness", "enthusiasm", "neutral", "worry", "surprise", "love", "fun", "hate", "happiness"]
    ax = sns.heatmap(confusion_matrix(y, y_pred), annot=True, cmap="Blues", fmt='g', cbar=False, annot_kws={"size":12})
    plt.title(title, fontsize=20)
    ax.xaxis.set_ticklabels(labels, fontsize=14)
    ax.yaxis.set_ticklabels(labels, fontsize=14)
    ax.set_ylabel('True Label', fontsize=18)
    ax.set_xlabel('Predicted Label', fontsize=18)
    plt.show()

```

```

fig = conf_matrix(y_test.argmax(1), y_pred_lstm.argmax(1), 'LSTM Sentiment Analysis\nConfusion Matrix')

```

Figure A.11: Code 11: Plotting the Results confusion matrix

56 2021 with climate change we only have 11 years to reign in emissions to do that we need to tax the out of you c
limateemergency climatecult greentaxes
Name: text, dtype: object

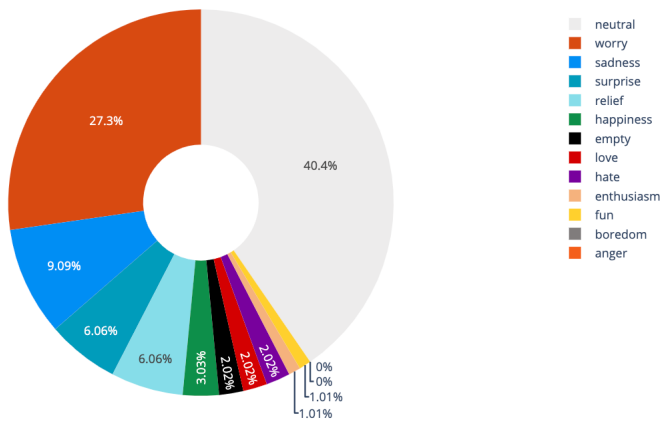


Figure A.12: Code 12: Example 1 of a tweet and its labeled emotions

567 americas largest banks are failing on climate change tell them all to stop funding fossil fuels now
Name: text, dtype: object

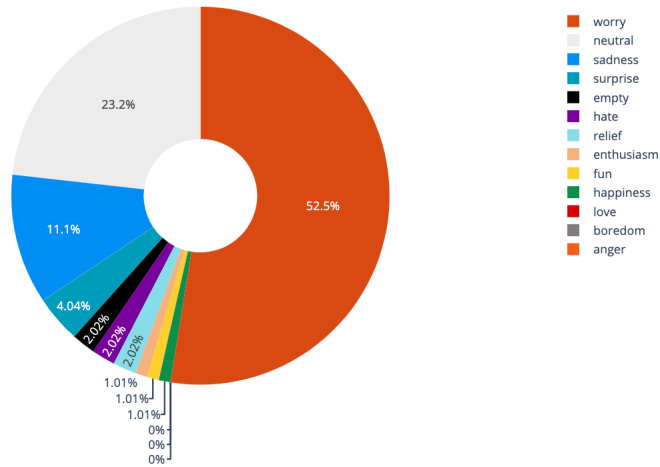


Figure A.13: Code 13: Example 2 of a tweet and its labeled emotions

A.2 A Comprehensive Reference Guide: Key Terminology for Acronyms and Python Libraries

Acronyms:

AI Artificial Intelligence

GloVe Global Vectors for Word Representation

GPT Generative Pre-trained Transformer

LDA Latent Dirichlet Allocation

LSTM Long Short-Term Memory

ML Machine Learning

NLP Natural Language Processing

RNN Recurrent Neural Networks

VADER Valence Aware Dictionary for Sentiment Reasoning

List of Python libraries mentioned:

Gensim

Keras

Matplotlib

Nominatim

Pandas

Scikit-learn

Tweepy

VaderSentiment

Word2vec

Bibliography

- Alexander, Marcalee (2021). "Climate change and health: The day for tomorrow". In: *The Journal of Climate Change and Health* 4. ISSN: 26672782. DOI: 10.1016/j.joclim.2021.100062.
- Alm, Cecilia Ovesdotter, Dan Roth, and Richard Sproat (n.d.). "Emotions from text: machine learning for text-based emotion prediction". In: *Proceedings of Human Language Technology Conference and Conference on Empirical Methods in Natural Language Processing*, pp. 579–586.
- Antonakaki, Despoina, Paraskevi Fragopoulou, and Sotiris Ioannidis (2021). "A survey of Twitter research: Data model, graph structure, sentiment analysis and attacks". In: *Expert Systems with Applications* 164. ISSN: 09574174. DOI: 10.1016/j.eswa.2020.114006.
- Association, American Psychological (n.d.). *APA Dictionary of Psychology*. Web Page. URL: <https://dictionary.apa.org/negative-emotion>.
- Bandura, Albert (1986 - 1986). *Social foundations of thought and action : a social cognitive theory / Albert Bandura*. eng. Prentice-Hall series in social learning theory. Englewood Cliffs, N.J: Prentice-Hall. ISBN: 013815614X.
- Benedicto, Manuel (2020). *Figure Eight Labelled Textual Dataset- Tweets containing a descriptor of the emotion they evoke*. Dataset. URL: <https://www.kaggle.com/datasets/manuelbenedicto/figure-eight-labelled-textual-dataset>.
- Bouman, Thijs et al. (2020). "When worry about climate change leads to climate action: How values, worry and personal responsibility relate to various climate actions". In: *Global Environmental Change* 62. ISSN: 09593780. DOI: 10.1016/j.gloenvcha.2020.102061.
- Brosch, Tobias (2021). "Affect and emotions as drivers of climate change perception and action: a review". In: *Current Opinion in Behavioral Sciences* 42, pp. 15–21. ISSN: 23521546. DOI: 10.1016/j.cobeha.2021.02.001.
- Budziszewska, M. and S. E. Jonsson (2022). "Talking about climate change and eco-anxiety in psychotherapy: A qualitative analysis of patients' experiences". In: *Psychotherapy* 59.4. Budziszewska, Magdalena Jonsson, Sofia Elisabet eng University of Warsaw; Faculty of Psychology/ Ministry of Science and Higher Education/ 2022/06/28 Psychotherapy (Chic). 2022 Dec;59(4):606-615. doi: 10.1037/pst0000449. Epub 2022 Jun 27., pp. 606–615. ISSN: 1939-1536 (Electronic) 0033-3204 (Linking). DOI: 10.1037/pst0000449. URL: <https://www.ncbi.nlm.nih.gov/pubmed/35758982>.
- Cauberghe, V. et al. (2021). "How Adolescents Use Social Media to Cope with Feelings of Loneliness and Anxiety During COVID-19 Lockdown". In: *CYBERPSYCHOLOGY, BEHAVIOR, AND SOCIAL NETWORKING* 24.4. Cauberghe, Verolien Van Wesen-

- beeck, Ini De Jans, Steffi Hudders, Liselot Ponnet, Koen eng 2020/11/14 Cyberpsychol Behav Soc Netw. 2021 Apr;24(4):250-257. doi: 10.1089/cyber.2020.0478. Epub 2020 Oct 20., pp. 250–257. ISSN: 2152-2723 (Electronic) 2152-2715 (Linking). DOI: 10.1089/cyber.2020.0478. URL: <https://www.ncbi.nlm.nih.gov/pubmed/33185488>.
- Chatterjee, Ankush et al. (2019). “Understanding Emotions in Text Using Deep Learning and Big Data”. In: *Computers in Human Behavior* 93, pp. 309–317. ISSN: 07475632. DOI: 10.1016/j.chb.2018.12.029.
- Chen, K., Z. Duan, and S. Yang (2023). “Twitter as research data -Tools, costs, skill sets, and lessons learned”. In: *Politics Life Sciences* 41.1. Chen, Kaiping Duan, Zening Yang, Sijia eng 2023/03/07 Politics Life Sci. 2023 Mar;41(1):114-130. doi: 10.1017/pls.2021.19., pp. 114–130. ISSN: 1471-5457 (Electronic) 0730-9384 (Linking). DOI: 10.1017/pls.2021.19. URL: <https://www.ncbi.nlm.nih.gov/pubmed/36877114>.
- Colneric, Niko and Janez Demsar (2020). “Emotion Recognition on Twitter: Comparative Study and Training a Unison Model”. In: *IEEE Transactions on Affective Computing* 11.3, pp. 433–446. ISSN: 1949-3045 2371-9850. DOI: 10.1109/taffc.2018.2807817.
- Corral-Verdugo, Victor (2021). “Psychology of climate change (Psicología del cambio climático)”. In: *PsyEcology* 12.2, pp. 254–282. ISSN: 2171-1976 1989-9386. DOI: 10.1080/21711976.2021.1901188.
- Fernandez, Miriam et al. (2016). *Talking climate change via social media*. Conference Paper. DOI: 10.1145/2908131.2908167.
- Gautam, Madhu et al. (2022). “Sentiment Analysis about COVID-19 vaccine on Twitter data: Understanding Public Opinion”. eng. In: *2022 6th International Conference on Intelligent Computing and Control Systems (ICICCS)*. IEEE, pp. 1487–1493. ISBN: 9781665410359.
- Gb, Shushma and I. Jeena Jacob (2022). *A Semantic Approach for Computing Speech Emotion Text Classification Using Machine Learning Algorithms*. Conference Paper. DOI: 10.1109/iceeict53079.2022.9768465.
- Gold, Nicolas (2020). *Using Twitter Data in Research Guidance for Researchers and Ethics Reviewers*. Web Page.
- Ha, Hyoji et al. (2019). “An Improved Study of Multilevel Semantic Network Visualization for Analyzing Sentiment Word of Movie Review Data”. In: *Applied Sciences* 9.12. ISSN: 2076-3417. DOI: 10.3390/app9122419.
- Hardin, Garrett (1998). “Extensions of ”The Tragedy of the Commons””. eng. In: *Science (American Association for the Advancement of Science)* 280.5364, pp. 682–683. ISSN: 0036-8075.
- Innocenti, Matteo et al. (2021). “Psychometric properties of the Italian version of the Climate Change Anxiety Scale”. In: *The Journal of Climate Change and Health* 3. ISSN: 26672782. DOI: 10.1016/j.joclim.2021.100080.
- Ipsos (2022). *Majority across 34 countries describe effects of climate change in their community as severe*. Newspaper Article. URL: <https://www.ipsos.com/en/climate-change-effects-displacements-global-survey-2022>.
- Jabreel, Mohammed and Antonio Moreno (2019). “A Deep Learning-Based Approach for Multi-Label Emotion Classification in Tweets”. In: *Applied Sciences* 9.6. ISSN: 2076-3417. DOI: 10.3390/app9061123.

- Jain, Ubeeka and Amandeep Sandhu (2015). "A Review on the Emotion Detection from Text using Machine Learning Techniques". In: *International Journal of Current Engineering and Technology* 5.4.
- Jang, S. Mo and P. Sol Hart (2015). "Polarized frames on "climate change" and "global warming" across countries and states: Evidence from Twitter big data". In: *Global Environmental Change* 32, pp. 11–17. ISSN: 09593780. DOI: 10.1016/j.gloenvcha.2015.02.010.
- Lawrance, E et al. (2021). *The impact of climate change on mental health and emotional wellbeing: current evidence and implications for policy and practice*. Report. Grantham Institute- Climate Change and the Environment- Imperial College London. DOI: 10.25561/88568.
- León, Bienvenido, Samuel Negredo, and María Carmen Erviti (2022). "Social Engagement with climate change: principles for effective visual representation on social media". In: *Climate Policy* 22.8, pp. 976–992. ISSN: 1469-3062 1752-7457. DOI: 10.1080/14693062.2022.2077292.
- Leung, J. et al. (2021). "Anxiety and Panic Buying Behaviour during COVID-19 Pandemic- A Qualitative Analysis of Toilet Paper Hoarding Contents on Twitter". In: *Int J Environ Res Public Health* 18.3. Leung, Janni Chung, Jack Yiu Chak Tisdale, Calvert Chiu, Vivian Lim, Carmen C W Chan, Gary eng Switzerland 2021/01/31 Int J Environ Res Public Health. 2021 Jan 27;18(3):1127. doi: 10.3390/ijerph18031127. ISSN: 1660-4601 (Electronic) 1661-7827 (Print) 1660-4601 (Linking). DOI: 10.3390/ijerph18031127. URL: <https://www.ncbi.nlm.nih.gov/pubmed/33514049>.
- Liu, Bing (2015). *Sentiment Analysis: Mining Opinions, Sentiments, and Emotions*. eng. Cambridge University Press. ISBN: 9781107017894.
- Mao, Xingliang et al. (2019). "Sentiment-Aware Word Embedding for Emotion Classification". In: *Applied Sciences* 9.7. ISSN: 2076-3417. DOI: 10.3390/app9071334.
- Markowitz, Ezra M. and Meghan L. Guckian (2018). "Climate change communication". In: *Psychology and Climate Change*, pp. 35–63. ISBN: 9780128131305. DOI: 10.1016/b978-0-12-813130-5.00003-5.
- McGillivray, Barbara and Gábor Mihály Tóth (2020). *Applying Language Technology in Humanities Research Design, Application, and the Underlying Logic*. Palgrave Macmillan. ISBN: 978-3-030-46492-9. URL: <https://doi.org/10.1007/978-3-030-46493-6>.
- Mokhtari, Sohrab, Kang K Yen, and Jin Liu (2018). "Effectiveness of Artificial Intelligence in Stock Market Prediction Based on Machine Learning". In: *International Journal of Computer Applications*, pp. 1–10.
- Nooralahzadeh, Farhad, Viswanathan Arumachalam, and Costin-Gabriel Chiru (2013). *2012 Presidential Elections on Twitter – An Analysis of How the US and French Election were Reflected in Tweets*. Conference Paper. DOI: 10.1109/cscs.2013.72.
- O. Bruna H. Avetisyan, J. Holub (2016). "Emotion models for textual emotion classification". In: *Journal of Physics* 772.Conference Series, pp. 1–6. DOI: 10.1088/1742-6596/772/1/012063.
- Peters, Fatima (2022). "The role of critical methodologies in climate psychology scholarship: themes, gaps, and futures". In: *South African Journal of Psychology* 52.4, pp. 510–521. ISSN: 0081-2463. DOI: 10.1177/00812463221130158.
- Pihkala, Panu (2020). "Anxiety and the Ecological Crisis: An Analysis of Eco-Anxiety and Climate Anxiety". In: *Sustainability* 12.19. ISSN: 2071-1050. DOI: 10.3390/su12197836.

- Rahuljha (2020). *LSTM gradients-Detailed mathematical derivation of gradients for LSTM cells*. Web Page. URL: <https://towardsdatascience.com/lstm-gradients-b3996e6a0296>.
- Roxburgh, Nicholas et al. (2019). "Characterising climate change discourse on social media during extreme weather events". In: *Global Environmental Change* 54, pp. 50–60. ISSN: 09593780. DOI: 10.1016/j.gloenvcha.2018.11.004.
- Saxena, Shipra (2023). *Learn About Long Short-Term Memory (LSTM) Algorithms*. Web Page. URL: <https://www.analyticsvidhya.com/blog/2021/03/introduction-to-long-short-term-memory-lstm/>.
- Sirisha, Uddagiri and Boleem Sai Chandana (2022). "Aspect based Sentiment and Emotion Analysis with ROBERTa, LSTM". In: *International Journal of Advanced Computer Science and Applications* 13.11, pp. 766–774. URL: www.ijacsa.thesai.org.
- Soutar, C. and A. P. F. Wand (2022). "Understanding the Spectrum of Anxiety Responses to Climate Change: A Systematic Review of the Qualitative Literature". In: *Int J Environ Res Public Health* 19.2. Soutar, Catriona Wand, Anne P F eng Review Systematic Review Switzerland 2022/01/22 Int J Environ Res Public Health. 2022 Jan 16;19(2):990. doi: 10.3390/ijerph19020990. ISSN: 1660-4601 (Electronic) 1661-7827 (Print) 1660-4601 (Linking). DOI: 10.3390/ijerph19020990. URL: <https://www.ncbi.nlm.nih.gov/pubmed/35055813>.
- Stanley, Samantha K. et al. (2021). "From anger to action: Differential impacts of eco-anxiety, eco-depression, and eco-anger on climate action and wellbeing". In: *The Journal of Climate Change and Health* 1. ISSN: 26672782. DOI: 10.1016/j.joclim.2021.100003.
- Tam, Pong, Angela K. -y Leung, and Susan Clayton (2021). "Research on climate change in social psychology publications: A systematic review". In: *Asian Journal of Social Psychology* 24.2, pp. 117–143. ISSN: 1367-2223 1467-839X. DOI: 10.1111/ajsp.12477.
- Taylor, Joanna and Claudia Pagliari (2017). "Mining social media data: How are research sponsors and researchers addressing the ethical challenges?" In: *Research Ethics* 14.2, pp. 1–39. ISSN: 1747-0161 2047-6094. DOI: 10.1177/1747016117738559.
- Taylor, S. (2020). "Anxiety disorders, climate change, and the challenges ahead: Introduction to the special issue". In: *Journal of Anxiety Disorders* 76. Taylor, Steven eng Netherlands 2020/09/30 J Anxiety Disord. 2020 Dec;76:102313. doi: 10.1016/j.janxdis.2020.102313. Epub 2020 Sep 22., p. 102313. ISSN: 1873-7897 (Electronic) 0887-6185 (Print) 0887-6185 (Linking). DOI: 10.1016/j.janxdis.2020.102313. URL: <https://www.ncbi.nlm.nih.gov/pubmed/32992267>.
- Townsend, Leanne and Claire Wallace (2017). "Chapter 8: The Ethics of Using Social Media Data in Research: A New Framework". In: *The Ethics of Online Research*. Advances in Research Ethics and Integrity, pp. 189–207. ISBN: 978-1-78714-486-6 978-1-78714-485-9. DOI: 10.1108/s2398-601820180000002008.
- Wainman, Ruth (2018). *How should I use social media as a research source?* Web Page. URL: <https://blogs.ucl.ac.uk/rdm/2018/12/how-should-i-use-social-media-as-a-research-source/>.
- Wang, Yuan et al. (2014). *Hashtag Graph Based Topic Model for Tweet Mining*. Conference Paper. DOI: 10.1109/icdm.2014.60.
- Zimbra, David et al. (2018). "The State-of-the-Art in Twitter Sentiment Analysis". In: *ACM Transactions on Management Information Systems* 9.2, pp. 1–29. ISSN: 2158-656X 2158-6578. DOI: 10.1145/3185045.