

Leverage Social Media Analysis and NLP to Study the Effectiveness of Antidepressants

Shreya Mandot, Birmingham City University, Birmingham, UK

Amna Dridi, Birmingham City University, Birmingham, UK

Zebedee Wanyonyi, Birmingham City University, Birmingham, UK

Edlira Vakaj, Birmingham City University, Birmingham, UK

ABSTRACT

Scientists now monitor population health by leveraging social media, where adults discuss adverse drug effects. By analysing social media conversations, scientists gain valuable insights into the public's perceptions with antidepressants, allowing for a better understanding of their impact on mental health. Antidepressants' efficacy is controversial, causing diverging opinions among the population, shaping attitudes towards treatment.

This study uses Twitter posts in the UK to understand antidepressants side effects through various topic modelling techniques (LDA, BERTopic, Top2Vec). To extract tweets related to antidepressants, we used a list of antidepressants approved by the Food and Drug Administrator (FDA) as keywords. The dataset contains 9,500 English tweets. BERTopic performed the best (coherence score: 0.70%), evaluated based on human judgment and topic's coherence score. Using FDA's side effect list as benchmark, this study found that most of the detected side effects align with the ones in the benchmark, indicating trustworthy insights from social media.

INTRODUCTION

Depression is a serious condition that affects adults of all ages all over the world and its prevalence is rising. Depression affects about 300 million individuals worldwide and is one of the leading causes of disability in high-income countries, accounting for more disability-adjusted life years than all other mental disorders. Many mental health illnesses, including mood, psychosis, and anxiety disorders, need the use of antidepressants Temmingh and Stein (2015). According to estimates, depression costs the economy 83 billion dollars a year in lost productivity, increased healthcare costs, and suicide. Depression is one of the world's most common psychiatric illnesses with symptoms such as unhappiness, loss of interest or pleasure, feelings of guilt and/or poor self-esteem, sleep and/or food abnormalities, exhaustion, and lack of focus. Mental health issues continue to rise as an illness and are likely to increase more in the current and near future by the ongoing COVID-19 pandemic and after the COVID-19 outbreak, researchers foresee a "tsunami of Depression" Tandon (2020). The mental health consequences because of big events like the COVID-19 pandemic can be long-lasting Galea, Merchant and Lurie (2020).

Major depressive disorder can lead to suicidal thoughts, inhibit behavioral function, and raise the likelihood of chronic conditions such as heart disease and obesity. General practitioners are well-positioned to identify and treat such illnesses, allowing for better mental health outcomes. Psychiatric medications are key to treating many mental health conditions, including mood, psychotic, and anxiety disorders. In such cases, antidepressant medications are frequently prescribed to treat the major depressive disorder, and many nations have noted notable increases in prescribing rates in recent decades. Antidepressants are one of the most prescribed types of medication in young and middle-aged adults, according to reports from the United States, Canada, and the United Kingdom, with 7% of adults aged 18-39 and 14% of adults aged 40-59 in use the antidepressants on regular basis Coupland et al. (2018).

The use of antidepressants is associated with several side effects and can cause medication nonadherence and withdrawal by interfering with optimal treatment dosage Lingam and Scott (2002). The common antidepressant side effects include gastrointestinal issues, weight gain or appetite, dry mouth, sleep problems, and sexual issues, among others. These antidepressants are used by one in every six British people, and they account for five of the top 40 drugs sold in the UK Min. As per NHS different antidepressants can cause a variety of negative effects. Therefore, there is a pressing need to get a deeper comprehension of the effect that antidepressants have on those suffering from mental illnesses. Antidepressants are the basis of major depressive disorder treatment Gartlehner et al. (2011). These are widely used and available worldwide and are organized into numerous medication categories with different action modes. Side effects of antidepressants should be known before they are used regularly, clinical trials for mental health do not adequately assess side effects, especially regarding long-term consequences. The antidepressants work well but they all have different risks, side effects, and different benefits. The antidepressant groups show common side effects such as insomnia, food disorders, pain, sexual issues, weight gain, dizziness, etc Jackson and St. Onge (2003). Because various side effects and frustration caused by it can also lead to abrupt withdrawal of these medicines Bull et al. (2002). Every person has a different reaction to the given medication, on other hand,

as per the traditional investigation patient's personal belief has an important influence on medicine. Currently, the side effect of these antidepressants is measured using pattern analysis and timing of responses of the patients from database systems that keep a record of events or incidents de Anta et al. (2022). These measures are usually biased Cipriani et al. (2018), and antidepressants' side effect depends on the Individual. The experience of every individual must be considered, but it is complex due to clinical variations. Traditional research like this is time-consuming and needs to be updated from time to time.

With the rapid growth in the use of antidepressants, there is a constant emphasis on precision psychiatry to understand the side effects and impact of these antidepressants. Multiple studies are being conducted to understand the adverse drug reactions, but it takes an expert reader and a lot of time to screen the side effects by looking at the relevant literature without a machine reading it. Initially, side effect detection was not well-structured and was obtained solely through communication between health professionals and patients or published case reports available in MEDLINE, PubMed, or other publicly available datasets Rison (2013). As a result, society requires an alternative method for detecting the side effects of clinical medications. Social media has the potential to generate novel and reliable data sources for drug side effects. As multiple factors affect the study of antidepressant side effects there is growing interest in understanding the potential impact of social media for this study. Individual organizations and users have shifted away from this traditional method and believe more in meeting online and building communities where they can share information related to antidepressant side effects and seek advice from other users. Nowadays people use social media to express their beliefs and emotions and concerns. Social Media allows researchers, scientists, and experts to obtain public health information in large amounts directly from people with relevant topics. As per the study conducted by Pew Research Centre Greenwood, Perrin and Duggan (2016), two third of adults worldwide use social media and more than 90% of citizens from the United Kingdom use social media. In recent years the study by collecting social media posts has gained a lot of popularity which gives a clear insight into patients' and non-patients beliefs towards medical treatment. It was noted by an expert Branley and Covey (2017) that people express more on social media which is instant and informal than in professional appointments with the doctor or in therapy sessions Alvarez-Mon et al. (2021). As this information is spontaneous it is more likely to demonstrate correct beliefs.

Various studies have taken place in health-related research using social media which focuses on getting information such as adverse drug reactions, medication abuse, sentiments analysis on such data Kaplan and Haenlein (2010) Korkontzelos et al. (2016). Extracting information from users' posts on social media has always been a challenge in the domain of natural language processing as this information might be noisy and inaccurate. Researchers tried to overcome the problem with the help of Twitter by using the total number of medications mentioned in the entire data. This research has limitations such as it is not possible to assess the side effect of medication on an individual level. Using social media to study the effects of medications on specific cohorts would necessitate the development of systems that can automatically differentiate between posts that express the use of the same medication or their side effects.

In this research, we want to look at a broader range of social media exposures and see if there was a link between social media and antidepressant side effects. To understand the side effect of the antidepressants we are going to leverage topic modelling techniques such as LDA, BERTopic, and Top2Vec using multiple feature extraction techniques. We are going to do a deep analysis in LDA using multiple feature extraction like Bag of Word (BOW), Term Frequency Inverse Document Frequency (TF-IDF), unigram bigram and at last FastText. For the topic modelling approach selecting the number of topics has always been a debate. Topic modelling is a method for determining which group of words (i.e., topic) from a collection of documents best represents the information in the collection. In our research, we aim to understand the result with the total number of medicines and with the total number of side effects. We will be using K=65 (total number of medicines provided) and K=35 (Number of side effects). This process will be repeated for each topic modelling algorithm and, we will check how BERTopic and Top2Vec define their number of topics. We will compare the results achieved using topic modelling techniques. Also, we will validate the results of the antidepressant side effects derived from social media to list of side effects for the antidepressant medicine provided by FDA. The essential contribution of this research is:

- Create the Drug-Reaction dataset in the English language using Twitter for tasks.
- related to the side effect of an antidepressant with the help of different antidepressant drug names.
- Implement Topic Modelling techniques like LDA, BERTopic, and Top2Vec which model themes present in a collection of posts, allowing for the identification of side effects of antidepressants.
 - Compare the results for a different number of topics and compare the results with the provided benchmark.

The rest of the project is organized as follows: In section 1.1, we elaborate on the Background and Rationale. In section 1.2, we specify the Aim of the project. In section 1.3, we declare the Objectives of the research. In section 1.4, we generate research questions. In chapter 2, we review the previous works related to this research. In chapter 3, we describe the research methodology. In chapter 4, we describe the dataset, data collection, and data pre-processing. In chapter 5, we evaluate the results. Eventually, in chapter 6, we conclude the research and describe the future work.

LITERATURE REVIEW

PSYCHIATRY DOMAIN RESEARCH

Depression is a disease of our modern era that poses a significant challenge to family physicians as well as mental health professionals Ramic et al. (2020). Antidepressants are the first primary care for patients dealing with depression. Antidepressants' impact on various neurotransmitter systems is linked to a significant number of their side effects Gilbody et al. (2003). The lack of research in this area and doctors' ignorance of this issue may be contributing factors to the negative side effects of antidepressants and their irrational use Turner et al. (2008). To limit the side effects of these antidepressants new antidepressant groups like Tetracyclic Antidepressants (TCA), Serotonin Norepinephrine Reuptake Inhibitors (SNRI) and Selective Serotonin Reuptake Inhibitors (SSRI) were discovered with the idea of increasing the selectivity of the target of action

to individual neurotransmitters López-Muñoz and Alamo (2009) Beck et al. (1961). Current antidepressant medicines are impactful and usually well accepted, but non-compliance is concerning. Around 70% of people taking antidepressants are non-compliant which means that they don't comply with rules and precautions while taking these drugs Lin et al. (1995). This led to patients missing the dose of medicines or abruptly stop these medications. As per the study suggested by Lin et al. (1995) 28% of patients stop taking antidepressants in first month of the treatment while 40-45% people stop before completing the requested therapy. As per author Katon et al. (1992) 60% of primary care patients stops taking these medications within 2 months of the prescribed treatment. Various studies have been conducted to understand the discontinuations of these medicines among which the side effects of these psychiatric drugs have been the most known reason. Antidepressants have been available for over 50 years, with iproniazid being the first drug to be used in the clinical diagnosis of depression. Tricyclic antidepressants are the first-generation antidepressant De Ridder (1982). The study in the field of antidepressant increased rapidly after the Imipramine, and other Tricyclic Antidepressants found to be effective in patients. Even though these antidepressants are effective it is no longer a drug of choice for the treatment of depression due to its severe side effects Gartlehner et al. (2011). These antidepressants were replaced by second generation antidepressants such as Selective Serotonin Reuptake Inhibitors (SSRIs) and, more recently, Selective Serotonin and Noradrenaline Reuptake Inhibitors (SNRIs). The tricyclic gives comparable effectiveness to that of second generation, however, have more severe side effects because of their pharmacologic behaviour and lower overdose threshold Bitter, Filipovits and Czobor (2011).

Various studies have been conducted to understand the side effects of psychiatric drugs. As per author Hirschfeld (2000) the main mechanism of psychiatric drugs and their baseline is not fully understood but the monoamine hypothesis shows that these drugs target the neurotransmitters of serotonin, norepinephrine, and dopamine, which are associated with feelings of well-being, alertness, and pleasure individual Gilbody et al. (2003). The most common side effects of tricyclic antidepressants are constipation, dizziness, and xerostomia Trindade et al. (1998). Because these drugs block central nervous system receptors, they can cause blurred vision, constipation, mouth sores, confusion, urinary retention, and tachycardia Güloğlu et al. (2011). It can cause orthostatic hypotension and dizziness because it blocks alpha-1 adrenergic receptors. TCAs can also lead to an increase the suicidal thoughts in individuals aged 24 or below. Cascade, Kalali and Kennedy (2009) surveyed on patients taking one of the following selective serotonin reuptake inhibitor antidepressants such as citalopram, escitalopram, fluoxetine, paroxetine, and sertraline. And reported side effects of antidepressants. Scientists noticed that 38% of 700 people have the most common side effects like insomnia, weight gain, and dysfunction. The other side effects included drowsiness, dry mouth, sleepiness, fatigue, sexual dysfunction, etc. To dig into this in detail research was conducted by David A Frommer et al. (1987) which demonstrated that the overdose of these antidepressants has unwanted side effects like the toxicity of the liver, sleep disorder, and hormonal imbalances.

SOCIAL MEDIA AND MACHINE LEARNING RESEARCH

Depression and mental disorder studies came to market much earlier than the internet and social media platform. Social media has become an integral part of the lives of many people dealing with a mental disorder. Social media refers to web and mobile platforms that enable people to connect with others inside a virtualized environment (such as Facebook, Twitter, Instagram, Snapchat, or LinkedIn) and communicate, co-create, or exchange various types of digital content, such as information, messages, photos, or videos Ahmed et al. (2019). Studies have also shown that people with a mental disorders, illnesses, and depression use social media more than the general population. As per the study done by Bucci, Schwannauer and Berry (2019) people with mental disorders use social media platforms to seek information and guidance about their mental health. Numerous studies have investigated the ability of postings on the Daily Power internet group to detect adverse drug reactions. In 2010, researcher Leaman et al. (2010) introduced a lexicon-based approach to detect the side effects of antidepressants. Experts performed various analyses using surveys or based on a user study. When people started to explain their views on social media author Xu et al Xu did a study on how people use social media platforms to discuss depression and antidepressants. An inventory for measuring depression suggested by Beck Beck et al. (1961) used 21 questionnaires to analyze the mental and physical well-being of individuals. Even after society's dominance for a post on social media, the data posted is unstructured and often heavy, and it gets difficult to perform data analysis.

With the realization that healthcare organizations and clinicians need to be more engaged with patients, over the past few years, healthcare organizations, clinicians, and the public have started using social media. Social media has great potential for improving interaction and patient involvement Househ, Borycki and Kushniruk (2014). According to research by Horvath et al. (2012) and Taggart et al. (2015) HIV patients most frequently cite the exchange of information and socializing with others as qualities of a perfect social network. Using a topic model, Wang, Huang and Gan (2016) demonstrated how users of a discussion board for pregnant women discussed topics that interested them and shared their experiences, worries, and concerns about medications. Breast cancer patients support one another and offer practical solutions to deal with drug side effects, according to research by Mao et al. (2013). Large sets of social media messages may contain hidden semantic structures that can be found using topic models. They might allow for a more thorough investigation of nonadherence behaviors. This investigation is based on patient testimonies of their own real-world drug decisions.

Experts applied the machine learning algorithm to the tweet datasets by doing manual annotations. Rachel Ginn and other partners Ginn et al. (2014) applied a Naïve Bayes and Support Vector Machine classifier on manually annotated adverse drug reaction twitter data. The classifier showed a significant performance and was used as a baseline for future development. In adjustment to the above research to classify the tweet comments author Akhtyamova, Alexandrov and Cardiff (2017) also applied Neural Networks (CNN) model with word2vec embedding. The research proposed by Zahra Rezaei and team Rezaei et al. (2020) used the "ask a patient" dataset in combination with the Twitter dataset and put forward a semi-supervised CNN deep learning techniques to classify the adverse drug effects in Twitter. In the end, Rezaei suggested that deep learning networks like FastText, CNN, and HAN were much faster and more accurate when compared to others. For predicting the treatment outcomes of Antidepressants experts used the Deep Neural Network of Deep Learning Tsai, Chang and Chen (2022).

For the study, the author applied the Univariate feature selection on the dataset and models were compared with and without feature selection and the results illustrated that the combination of feature extraction resulted in a good performance for outcome prediction.

In a variety of fields, including politics, topic modelling techniques have been used to examine the conceptual structure of textual data extracted from social media Yıldırım, Üsküdarlı and Özgür (2016). There has been various research applied on Tweet content using LDA to extract health-related keywords in tobacco use Prier et al. (2011), seasonal allergies, and childhood obesity Ghosh and Guha (2013). Social media has created an entirely new avenue for social science research, particularly when it comes to the intersection of human relations and technology. Although the informal language and sparse text make retrieving the underlying topics in tweets difficult, expert Weng et al. (2010) previously found that Latent Dirichlet Allocation (LDA) using the Bag of Word feature extraction technique produced decent results on tweets. The advancement of visual programming led researchers to apply topic modelling. Philip Resnik Resnik et al. (2015) did a study on language differences between depressive and non-depressive people where he uses the Twitter dataset created by Coppersmith. He explored topic modelling techniques like latent Dirichlet allocation (LDA), Supervised Topic Modelling (SLDA), for the automatic identification of depression. The SLDA algorithm performed poorly compared to vanilla LDA model based on n-grams. In a few research, experts have confirmed that LDA models, can uncover meaningful and potentially useful latent structures compared to other topic modelling algorithms. A study to predict the stability of antidepressant treatment routine with the major depressive disorder was performed on the electronic health record (EHR) or administrative data sets Hughes et al. (2020). Experts used supervised topic modelling called prediction-constrained (PC) topic modelling algorithm for the prediction purpose and demonstrated that the dataset information may help predict general treatment response but not specific medication response. A study to detect messages describing patients' noncompliant behaviors associated with an antidepressant drug of interest was done using the topic modelling algorithm like LDA where recall of 98.5% and precision of 32.6% are the respective values for the topic model approach's detection of instances of noncompliance behaviours 45.5% Abdellaoui et al. (2018). Utilizing LDA, patient forums have also been investigated. To find drug side effects, Yang, Kiang and Shang (2015) examined 1500 messages from patient forums.

Researchers examined the topics discussed in the unfavourable tweets for vaccine discourse and how they evolved by extracting them using the BERTopic model Lindelöf, Aledavood and Keller (2022). A study to understand the main subjects of conversation on Twitter about telemedicine for mental health and/or drug abuse was carried out with natural language processing using BERTweet and BERTopic model Baird, Xia and Cheng (2022). The BERTopic proved to be the best approach for determining the keywords related to telehealth. The researcher performed a study to determine suicidal behavior patterns in Tweet dataset using LDA and CNN Non-negative Matrix Factorization (NMF) Luo et al. (2020). The research showed NMF gives better accuracy to identify patterns for short texts. The theory of word embedding started a trend for word representation in vector form, which results in a more constructive method of representing the words. A comparison between two topic modelling techniques LDA and Top2vec was done in healthcare to understand covid-19 vaccine hesitancy where Top2Vec was able to reveal more determinant topics compared to LDA Ma, Zeng-Treitler and Nelson (2021). Numerous studies have been performed to compare the suitable topic modelling technique for short documents such as tweets. Likhitha, Harish and Kumar (2019) Performed LDA and Word2Vec text mining on tweets where word embedding technique gave better H-score than another model.

Experts started doing research on newly emerging topic modelling algorithms like BERTopic, Top2Vec. A Topic modelling Comparison of LDA, NMF, Top2Vec, and BERTopic to help understand Twitter posts was done by researchers where BERTopic gave exceptional results compared to other models Egger and Yu (2022). We have been seeing BERTopic as an emerging technique in the study of medicine. An Effective approach based on the BERT language model for Twitter sentiment analysis was done by experts Pota et al. (2020). Even after the popularity, reliability, and validity of LDA, BERTopic, and SLDA, there was always a concern in social science about the application of these traditional algorithms Egger and Yu (2022). In this study, we take steps toward using more complex models like using multiple feature extraction techniques with LDA such as Bag of Words based LDA, trying N-gram approach, TF-IDF based LDA, and FastText is shown to be providing better accuracy with a short text. We will be comparing the results for different LDA models. The Top2Vec model uses Distributed Bag of Words (DBOW) in doc2vec and a universal sentence encoder to produce jointly embedded document and word vectors as an alternative to the LDA approach Angelov (2020). In the same way, BERTopic uses continuous TF-IDF to analyse linguistic signals for anti-depression side effects detection and show encouraging results. Our research aims to compare these algorithms for better coherence score which is nothing but the degree of similarity measures between words in topics generated by a topic model and evaluating results for most deterministic antidepressant side effects.

METHODOLOGY

This section will detail the methodology adopted in this research work.

FEATURE EXTRACTION TECHNIQUES

Text modelling is difficult because it is messy, and techniques such as machine learning algorithms prefer well-defined fixed-length inputs and outputs. Feature extraction is an important feature when we deal with Natural Language Processing (NLP) as the data we have collected is in an unreadable format by computer. To pass the data to any algorithm we need the data in a readable format by computer. So, we need a feature extraction technique that can convert text into the matrix of vectors. There are many features extraction techniques like Bag of Words (BOW), Term Frequency Inverse Document Frequency (TF-IDF), n-gram, word embedding, NLP-based features like word count, noun count, Word2Vec, etc. Before applying the feature extraction technique, we will be doing text pre-processing so that we can get rid of the not used part of the data, noise, stop

words, punctuation, etc. In our research, we are going to use the below feature extraction techniques and compare them for topic modelling:

- Bag of words
- Term Frequency – Inverse Document Frequency
- N-gram
- FastText

TOPIC MODELLING

Topic modelling is a technique for automatically discovering the underlying topics that are present in a collection of documents. It is often used in natural language processing and information retrieval applications. Statistical techniques called topic modelling algorithms examine the words in documents to identify the themes that run throughout a large body of documents. The information collected from Twitter is in human language. It is usually unlabeled and unstructured and can have long text lengths. To effectively extract features from big corpus or paragraphs various text mining approaches were used among which topic modelling was the most used. It is used to identify hidden patterns within the text. Topic modelling is a useful technique for a wide range of applications, including information retrieval, text classification, and text summarization. It can help to identify the main themes in a large dataset of documents and can be used to organize and classify the documents based on their content. There are various well-known topic modelling algorithms such as LDA, latent semantic analysis LSA, Probabilistic LSA, and newly evolved algorithms such as NMF, Top2Vec, and BERTopic which are gaining more attention from experts. As per the study, LDA and BERTopic and Top2vec have gained the most accuracy while dealing with tweet analysis. We will be using these algorithms for our research and digging deeper to achieve the best results. We aim to implement the above topic modelling algorithm and compare them for the results achieved.

LATENT DIRICHLET ALLOCATION

In 2003, David Blei, Andrew Ng, and Michael I. Jordan proposed LDA as a topic discovery graphical model Blei, Ng and Jordan (2003). The fundamental tenet of topic modelling is that each latent topic in a document is expressed by a distribution of words. In the field of text mining, Latent Dirichlet Allocation (LDA) is the most widely used topic modelling technique. LDA is an unsupervised topic modelling technique. It is one of the commonly known topic modelling techniques to extract topics from a corpus. Latent means hidden. When we collect any unstructured document, it is made up of various words, and each topic is also made up of various words. The main objective of LDA is to find which topic a document belongs to, based on words in it. The technique converts the received document word into different matrices mainly document term matrix and topic word matrix. LDA is a probabilistic model, which means that it defines a probability distribution over the possible combinations of topics and words that could have generated the observed data. Given a set of documents and a specified number of topics, LDA estimates the parameters of the model (i.e., the distribution of topics over documents and the distribution of words over topics) that are most likely to have generated the data. In this study, we proposed to apply LDA for clustering analysis on our given dataset with multiple feature extraction techniques. We will compare the feature extraction techniques like Bag of Words, Unigram Bigram, TF-IDF, and FastText to derive the best results.

DATASET

As of January 1, 2021, Twitter, one of the biggest social media platforms, had over 1445,000,000 users and is growing by an estimated 935,000 users daily, producing about 15,100 tweets every second. It is a considered as potential gold mine of information for research. Twitter's application programming interface (API) makes part of its data publicly available and easily accessible. A Survey was done by researchers (Parker et al., 2013) shows that 42% of online users use social media to post or seek information about health conditions, and 26% of them discussed personal health issues. As Twitter has gained popularity in public health research in the last decade and it is easy to use with public APIs, we will be using a Twitter data source for our study. The first goal is to find the dataset suitable for our problem domain. For any machine learning technique gathering suitable data is the most important part.

DATA COLLECTION

The below section gives a list of requirements for data collection:

- Select keywords to find antidepressant side effects.
- Find the list of antidepressants.
- Extract the tweets for a specific timeline.
- Specify the language of tweets to be extracted.
- Find tweets for the world as people can talk about antidepressants from anywhere and are not limited to any country.
- To find the side effects in the United Kingdom, we need to filter the tweets specific to the UK using the location of tweets.

We limit our research to a set of FDA-approved antidepressants and antidepressant enhancement medications. For creating the dataset of medicines, we are going to use the list of antidepressant drugs used by experts Saha et al. (2022) in JMIR mental health journal. This list was developed with the approval of the Food and Drug Administration—in consultation with the psychiatrist co-author (JT). The list contains 297 brand names that are matched to 49 generic names of four different types of

antidepressants. These classes include selective serotonin reuptake inhibitors (SSRIs), serotonin-norepinephrine reuptake inhibitors (SNRIs), tricyclic antidepressants, and tetracyclic antidepressants Saha et al. (2022).

To build the corpus we queried around 9,800 Tweets from the Twitter API in English language mentioning the names of 65 medicines listed in Figure 4.1 . These medicines were from the group of antidepressants of various brands and generic names. antidepressants' medicine names as a keyword to extract tweets from Twitter. The next purpose is to find the tweets related to the keywords. For extracting the tweets, we will be using Twitter Application Processing Interface (API) with the permission of Twitter. We will be using the keywords like "SSRI", "SNRIs", "Tetracyclic", "Tricyclic" and more such keywords from Figure 1 is

agomelatine, amineptine, amitriptyline, amoxapine, bupropion, butriptyline, citalopram, clomipramine, desipramine, desvenlafaxine, dibenzepin, dosulepin, doxepin, duloxetine, escitalopram, etoperidone, fluoxetine, fluvoxamine, hydroxynefazodone, imipramine, iprindole, levomilnacipran, lofepramine, maprotiline, mazindol, meta-chlorophenylpiperazine, mianserin, mirtazapine, nefazodone, nisoxetine, nomifensine, norclomipramine, northiaden, nortriptyline, opipramol, oxaprotiline, paroxetine, protriptyline, reboxetine, sertraline, setiptiline, trazodone, triazoledione, trimipramine, venlafaxine, vilazodone, viloxazine, vortioxetine, zimelidine

Figure 1: List of antidepressants considered for this study.

given below. We search for public English posts containing specific keywords and gather the tweets post covid for the timeline of the year 2022.

DATASET FORMAT

The dataset will be collected in the following format. As shown in the table below we will be collecting unique tweets id and tweet text related to antidepressants which will be used for analyzing the side effects of antidepressants with the location of users and the date on which tweet is posted.

Table 1: Tweets Dataset

Tweet-Id	TweetText	Location	date
112223456	@dxldlez @Noahpinion Welbutrin is the one antidepressant that increases nausea https://t.co/zyvPQo7xlo	Wales,UK	2022-11-07
265496977	@postliberal I'm on this And Lexapro and gabapentin Only side effect I've seen is like crazy weight gain :)	United States	2022-11-14

Table 2 below shows the samples of tweets collected regarding the medicines. The tweet shows various kinds of side effects such as headache, sleep deprivation, insomnia, problems in staying focused and so on. We aim to use the topic modelling algorithm to extract more such side effects and check it against the benchmark created by us. As above- collected tweets show the wild character and are not cleaned, we will be performing data pre-processing steps on it in the section below.

Table 2: Tweets Gathered Related to Antidepressant

Tweet-Id
I caught cold having toothache and headache on my period and not to mention that I actually might have PMDD cuz I have the history as Escitalopram consumer *crying snuggled up in blanket on my bed
starting lexapro soonðŸ˜˜~(Donâ€™t tell me ur bad experience with it because i have severe hypochondria and will shut and cry and puke),
@RunInWithATree I take Lexapro and being without it really messes you up. I felt a lot of dizziness when I was getting used to it. Are you out of refills? Might be making the sleep deprivation worse.
@Poly40399120 I miss taking Sertraline, since they moved me onto Escitalopram my libido has been almost non-existent and I have a hard time trying to stay focused...
My mirtazapine dosage has been 15mg for almost 4 weeks and my sleep is still awful so I'll

hopefully be having a GP telephone appointment today to see if I can lower it anymore (I'm hoping I can go down to 7.5mg but idek if that's possible)

DATA PRE-PROCESSING

Data Pre-Processing and transforming is an important step in Machine learning. Without pre-processing, raw tweets are highly unstructured and contain redundant information. To address these issues, tweet preprocessing is carried out in a series of steps.

- We remove the duplicate tweets to remove redundancies.
- Lowercase transformation: This procedure is used to change the text's format to lowercase when parsing is being done. You can lowercase every word to account for variations in capitalization. It is also employed to prevent word repetition.
- Hyperlink Removal: Hyperlink are removed from the social media post using regular expression in python.
- Wild Character Removal: In the tweet dataset we have Emojis, Smileys, special character, Mentions etc. We need to remove it to make sure data is in readable format. To remove this, we use regular expressions.
- Numbers and repeating character removal: When we perform text processing, we eliminate numbers from 0-9 as they do not impact on finding the important words for side effect. Also, there will be few repeating words like 'ii', 'aa' which we will be removed to make the data clean and concise.
- Tokenisation: In plain English, tokenization is the process of dividing a phrase, sentence, paragraph, one or more text documents into smaller units.
- Stopword Removal: In any language, stop words are a group of frequently used words. Stop words in English include "the" "is" and "and" for instance. Stop words are used in NLP and text mining applications to filter out unnecessary words so that the important words can be the focus. To remove stop words, we use Natural Language Toolkit (NLTK) library.
- Lemmatization: In Natural Language Processing (NLP), a text normalization technique called lemmatization is used to convert any kind of word to its base root mode. Lemmatization is responsible for combining various word inflections with the same meaning into their root forms. For example, lemmatization will convert leaf's, leaves to leaf.
- Stemming: Stemming, which affixes to suffixes, prefixes, or the roots of words known as "lemmas," is the process of reducing a word to its stem. Natural language processing (NLP) and natural language understanding (NLU) both benefit from stemming.

EVALUATION

IMPLEMENTATION

In our experiment, we extracted the features using the python Scikit-Learn library. As mentioned in the above sections we used different feature extractions for topic modelling algorithms like LDA, BERTopic, and Top2Vec. For LDA we implemented feature extraction techniques like Bag of Words, Term Frequency Inverse Distribution, Unigram bigram, and fastText. For BERTopic we used continuous Bag of Words and for Top2vec we experimented using the universal sentence encoder embedding method. Regardless of the variety of topic modelling algorithms proposed, selecting an appropriate number of topics for a given corpus is a common challenge in successfully applying these techniques. Choosing too few topics results in overly broad results, while choosing too many results in "over- clustering" of a corpus into many small, highly similar topics.

Topic modelling techniques require a researcher to specify a k value or the number of topics in the corpus Gillies et al.(2022). In practice, this is a difficult and significant decision. There is no such standard rule for defining the number of topics and this is crucial as the validity of the results is dependent on the fitting model process. There are multiple ways to tackle this issue, but a solution based on a Bayesian approach that involves computing the log-likelihood for all possible models in each interval is most appreciated for its simplicity Griffiths and Steyvers (2004). Though this method has been quite popular there are a few disadvantages to using the method for the selection of the topics. The large amount of data required to determine all the conditional probabilities required in the rigorous application of the formula is one of the most commonly recognized problems with using a Bayesian approach Griffiths and Steyvers (2004). To determine an appropriate number of topics, compare the goodness-of-fit of models fitted with varying numbers of topics. The perplexity of a held-out set of documents can be used to assess the goodness-of-fit of a model. The perplexity of the model indicates how well it describes a set of documents. To check the goodness of fit of the model we have implemented the topic modelling techniques with the number of topics (n) as the total number of medicines that is, 65 and total number of side effects of the drugs as 36.

As previously stated, to select the LDA model that best fits the data, the number (n) of topics was varied between 36 and 65 topics, and the models' performance was evaluated in terms of coherence and perplexity. BERTopic modelling is an unsupervised topic modelling algorithm, and it automatically identifies the number of topics present in the textual data. Sometimes BERTopic generated too many topics or too few topics, there it gives a way to control this behavior where the number of topics to be created can be specified. In our experiment, we controlled the number of topics with n=36, 65 and gave a chance to BERTopic to identify its own topic where we compared the coherence score for each

of the methods. Top2Vec is a pre-trained model generating its number of topics. In our study, we haven't controlled the number of topics for Top2Vec and evaluated the results for several topics created by default.

EVALUATION OF TOPIC MODELLING

Topic modelling is a subcategory of natural language processing that is used to investigate text data. Topic models are frequently used to analyze unstructured text data, they do not offer any advice on the caliber of the topics generated. The secret to comprehending topic models is evaluation. It works by recognizing key patterns or topics in the document based on words or phrases with similar meanings.

Topic modelling evaluation is a step for knowing how well it fits your data to determine the patterns or trends in document. Here, we have created these topics to cluster our data with same meaning. There are 2 ways to evaluate topic models.

Human judgement: It can depend on observation as checking top 'n' words in the topic, or it be based on interpretation. For example, 'topic intrusion' for determining the topics that do not belong to the document. Even though human judgement results give a good result, it is expensive and time-consuming which makes us calculate quantitative metrics.

Quantitative Matrix: In the Quantitative matrix, the model will calculate the perplexity score and coherence score.

Perplexity Score: It is a score to measure how well your trained model predicts new data. A low perplexity score implies a good model which means it is good at predicting the words appearing in the document. Coherence: The degree of similarity measures between words in topics generated by a topic model is determined by coherence. As we get more alike words the more will be the coherence score.

RESULTS AND DISCUSSION

In this research, we have examined social media content to understand the side effects of antidepressants with the help of a topic-modelling approach. We performed data pre-processing steps and feature extraction techniques on given data to perform topic modelling algorithms like LDA, BERTopic, and Top2vec.

All the topic modelling techniques could find keywords related to the side effects of antidepressants. In the initial step, we performed an LDA model considering a feature extraction technique as Bag-Of Words. Researchers need to mention the number of topics or clusters to be created as a parameter while implementing this algorithm. In this study, we have created 35 and 65 individual topics to observe the side effects of the given drugs.

The figures 5.1 and 5.2 overview critical topics, respective percentage contributions in each document, and their associated related keywords.

Now, at first, we implement this LDA model with a number of topics to create as 36, fetching the side effects related keywords such as weight gain, attack, panic, and anxiety. As given in table 5.1 topics 1, 35, and 7 talks about the side effects, whereas topic one specifically implies the long-term side effects of antidepressants. The coherence score achieved by the Bag of Word-based LDA model is 0.473. Topic modelling uses this coherence score to determine how human-interpretative the topics are. We performed the same analysis again with a change in the number of topics to 65. With the increased number of topics, we achieved a coherence score of approximately 0.52. It showed the side effects like sleep, nausea, weight gain, migraine, sweating, anxiety, death (suicidal thoughts), depression, etc., shown in table 5.2. The word cloud plotted in figure 5.1 shows the adverse drug effects like pain, depression, nausea, weight gain, appetite, and panic attack. Topic 61 gives us insights into side effects such as stress and chronic effects of the drugs Paxil and Wellbutrin. Topic 29 gives an idea of side effects in women, such as pain and nausea. Topic 47 shows sleep and weight gain.

Table 3: Keyword collected for 36 Topics with BOW

Topic Number	Keywords
35	make, hear, weight, gain, alone, sleep, amount, regular, difference
13	food, fuck, probably, wellbutrin, dog, day, med, take, already, abuse
10	bupropion, antidepressant, die, doctor, mean, switch, know, give, state, trazodone
8	therapy, psychedelic, morning, trial, daily, wellbutrin, post, day, take, let
7	trazodone, anxiety, help, attack, panic, hard, big, community
1	effect, side, long, term, due, fact, bad, sleep, depression
31	health, put, mental, call, take, prescribe, time, day, love, combine



Figure 2: Word cloud of keyword collected for LDA with BoW



Figure 4: Word cloud of keyword collected for 65 Topics with BERTopic



Figure 3: Word cloud of keyword collected with Top2Vec

Our research compared multiple topic modelling techniques such as LDA, BERTopic, and Top2Vec. The below table shows the comparison of coherence scores achieved for each number of topics created. We compared our implemented models against our topic modelling evaluation technique, such as coherence score and human judgment. For LDA, we can see that even though the coherence score for BOW-based LDA and unigram bigram-based LDA for 65 topics are close to each other, as per human judgment, we have achieved more dominant side effects with the unigram bigram-based LDA approach. FastText has been proven to be the most suitable feature extraction technique in the past for short text, and our coherence score says so. We could only fetch a few side effects with this technique. We tried our FastText analysis with 36 and 65 topics, and we could not interpret any dominant side effects with both analyses. When the total number of topics is 36, there is little significant difference in the results produced by BOW-based LDA, Unigram Bigram-based LDA, and TF-IDF-based LDA. With Top2Vec total number of the topic created is 61 a coherence score of 0.46.

Further, the BERTopic modeling technique yields the best results when the model creates the topics. With a total number of topics created, 119, we achieved a coherence score of 0.71, which is low compared to the FastText model. Even with a low coherence score compared to the FastText model, we derived more meaningful side effects of antidepressants as per the drug names. On reducing the topic to 36 and 65, the coherence score reduced to 0.563 and 0.63, respectively, as compared to topics created by BERTopic on its own. We noticed that when we reduced the number of topics to 36 and 65 with BERTopic, it still gave us more side effects of antidepressants compared to others. After performing all the analysis, we could see that BERTopic has outperformed in all the scenarios and resulted in giving prominent side effects of the antidepressants.

Table 4: Topic Modelling Accuracy

Topic Modelling	Feature Extraction	Coherence Score with 35 Topics	Coherence Score with 65 Topics	Topics Created on Own
LDA	BOW	0.47	0.52	NA
LDA	Unigram Bigram	0.48	0.56	NA
LDA	TF-IDF	0.49	0.51	NA
LDA	FastText	0.60	0.74	NA
BERTopic	BERT, cTF-IDF	0.56	0.63	0.71
Top2Vec	Universal-Sentence-Encoder	NA	NA	0.47

To validate our results, we compared the results with the side effects of antidepressants from the food-drug-approved list of medicines. TheFDAlist shows the side effects like nausea, nervousness, problems sleeping, sexual problems, sweating, dizziness, weight gain, restlessness, drowsiness, focus (difficulty with attention, judgment, and thinking), spinning sensation, anxiety, suicidal thoughts, and behaviours, headache, stroke. Our topic modelling analysis reveals many keywords potentially related to prominent side effects of antidepressants. Such as, with LDA with unigram bigram, we could see the side effects like depression, panic attacks, sleep, pain, major depressive disorder, killing, death, feeling dizzy, insomnia, weight gain, periods miss, and ill. According to BERTopic, mirtazapine has side effects on sleep and appetite, trazodone has sleepiness, amitriptyline has nerve and migraine-related pain, Paxil has attacked, and escitalopram has suicidal thoughts.

There is a lack of simple and systematic methods for understanding antidepressant side effects Lexchin et al.(2003). Many side effects are known to be under-reported in clinical care, and this study provides a new way to assess the actual burden better and lived experience of patients. This study can also be used to inform patients about the full range of side effects and as a tool for more legal decision-making regarding medication selection. Furthermore, this research has implications for drug repurposing and developing new drugs that target treatments with fewer side effects

CONCLUSIONS AND FUTURE WORK

In conclusion, the findings of our linguistic study reveal a number of keywords that have the potential to be associated with the well-known adverse effects of antidepressants. By utilizing topic modelling techniques such as BERTopic, LDA, and Top2Vec, we were able to retrieve the adverse effects associated with antidepressants. The coherence score and human assessment are utilized in evaluating these topic models. We could see that among the several implementations of LDA with different kinds of feature extraction, unigram and bigram generated the most potential side effects from drugs. Even though the FastText result provided us with the highest level of coherence compared to all other topic modelling techniques, we achieved negligible adverse effects with this method. BERTopic has outperformed all other topic modelling methodologies regarding the number of side effects produced. Despite having a lower coherence score for the BERTopic, our model extracted more severe antidepressant side effects from medication names than the FastText model. We also discovered that more excellent coherence scores do not necessarily reflect a more accurate representation of distributional information concerning human evaluations. This study illustrates how the use of social media can contribute to a better knowledge of the adverse effects of medication. We want to keep the clinical trial separate from our study because there are various qualities that social media platforms do not have but that clinical trials might have. It is possible to think of it as an additional resource for learning about the adverse effects of medicine.

Since Twitter only allows access to the most recent seven days of data, we were only able to analyze a tiny dataset for our study. Understanding the adverse effects of antidepressants is a continuing concern, so the future dataset can be expanded. Our study, which can be expanded to include Tweets in other languages including Spanish, French, and Italian, is only limited to English-language Tweets. Since Twitter's data may be accessed by the public through the Twitter API, we have selected Twitter as the medium for our research. It may have been the ideal platform for exploratory study due to its rigorous rules governing the maximum number of characters that can be used in a tweet. The methodology employed in this study should, however, also be applicable to other platforms because social media posts are frequently succinct and unstructured. Nevertheless, it is crucial to understand that social media differs in terms of user demographics, text display, and other aspects. Future research should investigate the efficacy of topic modelling algorithms on various social media sites. Although BERTopic is recommended in this study for the analysis of short texts, new topic modelling techniques are also being developed. Researchers recommend trying and evaluating various recently established algorithms to ascertain the side effects of antidepressants to maximize the usage of topic modelling methodologies.

REFERENCES

Abdellaoui, R., Foulquié, P., Texier, N., Faviez, C., Burgun, A., Schück, S. et al., 2018. Detection of cases of noncompliance to drug treatment in patient forum posts: topic model approach. *Journal of medical Internet research*, 20(3), p.e9222.

Ahmed, Y.A., Ahmad, M.N., Ahmad, N. and Zakaria, N.H., 2019. Social media for knowledge-sharing: A systematic literature review. *Telematics and informatics*, 37, pp.72–112.

Akhtyamova, L., Alexandrov, M. and Cardiff, J., 2017. Adverse drug extraction in twitter data using convolutional neural network. 2017 28th international workshop on database and expert systems applications (dexa). IEEE, pp.88–92.

Alvarez-Mon, M.A., Anta, L. de, Llaverio-Valero, M., Lahera, G., Ortega, M.A., Soutullo, C., Quintero, J., Barco, A. Asunsolo del and Alvarez-Mon, M., 2021. Areas of interest and attitudes towards the pharmacological treatment of attention deficit hyperactivity disorder: thematic and quantitative analysis using twitter. *Journal of Clinical Medicine*, 10(12), p.2668.

Angelov, D., 2020. Top2vec: Distributed representations of topics. arXiv preprint arXiv:2008.09470.

Anta, L. de, Alvarez-Mon, M.A., Ortega, M.A., Salazar, C., Donat-Vargas, C., Santoma-Vilaclara, J., Martin-Martinez, M., Lahera, G., Gutierrez-Rojas, L., Rodriguez-Jimenez, R. et al., 2022. Areas of interest and social consideration of antidepressants on english tweets: a natural language processing classification study. *Journal of personalized medicine*, 12(2), p.155.

Babu, N.V. and Kanaga, E., 2022. Sentiment analysis in social media data for depression detection using artificial intelligence: A review. *SN Computer Science*, 3(1), pp.1–20.

Baird, A., Xia, Y. and Cheng, Y., 2022. Consumer perceptions of telehealth for mental health or substance abuse: a twitter-based topic modeling analysis. *JAMIA open*, 5(2).

Beck, A.T., Ward, C.H., Mendelson, M., Mock, J. and Erbaugh, J., 1961. An inventory for measuring depression. *Archives of general psychiatry*, 4(6), pp.561–571.

Bitter, I., Filipovits, D. and Czobor, P., 2011. Adverse reactions to duloxetine in depression. *Expert opinion on drug safety*, 10(6), pp.839–850.

Blei, D.M., Ng, A.Y. and Jordan, M.I., 2003. Latent dirichlet allocation. *Journal of machine Learning research*, 3(Jan), pp.993–1022.

Branley, D.B. and Covey, J., 2017. Pro-ana versus pro-recovery: A content analytic comparison of social media users' communication about eating disorders on twitter and tumblr. *Frontiers in psychology*, 8, p.1356.

Bucci, S., Schwannauer, M. and Berry, N., 2019. The digital revolution and its impact on mental health care. *Psychology and Psychotherapy: Theory, Research and Practice*, 92(2), pp.277–297.

Bull, S.A., Hu, X.H., Hunkeler, E.M., Lee, J.Y., Ming, E.E., Markson, L.E. and Fireman, B., 2002. Discontinuation of use and switching of antidepressants: influence of patient-physician communication. *Jama*, 288(11), pp.1403–1409.

Cascade, E., Kalali, A.H. and Kennedy, S.H., 2009. Real-world data on ssri antidepressant side effects *Psychiatry (Edgmont)*, 6(2), p.16.

Cipriani, A., Furukawa, T.A., Salanti, G., Chaimani, A., Atkinson, L.Z., Ogawa, Y., Leucht, S., Ruhe, H.G., Turner, E.H., Higgins, J.P. et al., 2018. Comparative efficacy and acceptability of 21 antidepressant drugs for the acute treatment of adults with major depressive disorder: a systematic review and network meta-analysis. *Focus*, 16(4), pp.420–429.

Coppersmith, G., Dredze, M. and Harman, C., 2014. Quantifying mental health signals in twitter. *Proceedings of the workshop on computational linguistics and clinical psychology: From linguistic signal to clinical reality*. pp.51–60.

Coupland, C., Hill, T., Morriss, R., Moore, M., Arthur, A. and Hippisley-Cox, J., 2018. Antidepressant use and risk of adverse outcomes in people aged 20–64 years: cohort study using a primary care database. *BMC medicine*, 16(1), pp.1–24.

De Ridder, J., 1982. Mianserin: result of a decade of antidepressant research. *Pharmaceutisch Weekblad*, 4(5), pp.139–145.

Egger, R. and Yu, J., 2022. A topic modeling comparison between lda, nmf, top2vec, and bertopic to demystify twitter posts. *Frontiers in Sociology*, 7.

Frommer, D.A., Kulig, K.W., Marx, J.A. and Rumack, B., 1987. Tricyclic antidepressant overdose: a review. *Jama*, 257(4), pp.521–526.

Galea, S., Merchant, R.M. and Lurie, N., 2020. The mental health consequences of covid-19 and physical distancing: the need for prevention and early intervention. *JAMA internal medicine*, 180(6), pp.817–818.

Gartlehner, G., Hansen, R.A., Morgan, L.C., Thaler, K., Lux, L., Van Noord, M., Mager, U., Thieda, P., Gaynes, B.N., Wilkins, T. et al., 2011. Comparative benefits and harms of second-generation antidepressants for treating major depressive disorder: an updated meta-analysis. *Annals of internal medicine*, 155(11), pp.772–785.

Ghosh, D. and Guha, R., 2013. What are we 'tweeting' about obesity? mapping tweets with topic modeling and geographic information system. *Cartography and geographic information science*, 40(2), pp.90–102.

Gilbody, S., Whitty, P., Grimshaw, J. and Thomas, R., 2003. Educational and organizational interventions to improve the management of depression in primary care: a systematic review. *Jama*, 289(23), pp.3145–3151.

Gillies, M., Murthy, D., Brenton, H. and Olaniyan, R., 2022. Theme and topic: How qualitative research and topic modeling can be brought together. arXiv preprint arXiv:2210.00707.

Ginn, R., Pimpalkhute, P., Nikfarjam, A., Patki, A., O'Connor, K., Sarker, A., Smith, K. and Gonzalez, G., 2014. Mining twitter for adverse drug reaction mentions: a corpus and classification benchmark. *Proceedings of the fourth workshop on building and evaluating resources for health and biomedical text processing*. pp.1–8.

Greenwood, S., Perrin, A. and Duggan, M., 2016. Social media update 2016. *Pew Research Center*, 11(2), pp.1–18.

Griffiths, T.L. and Steyvers, M., 2004. Finding scientific topics. *Proceedings of the National academy of Sciences*, 101(suppl_1), pp.5228–5235.

Grootendorst, M., 2022. Bertopic: Neural topic modeling with a class-based tf-idf procedure. arXiv preprint arXiv:2203.05794.

Güloğlu, C., Orak, M., Üstündağ, M. and Altuncı, Y.A., 2011. Analysis of amitriptyline overdose in emergency medicine. *Emergency Medicine Journal*, 28(4), pp.296–299.

Heald, A.H., Stedman, M., Davies, M., Farman, S., Taylor, D., Bailey, S. and Gadsby, R., 2020. Quantifying the impact of patient-practice relationship quality on the levels of the average annual antidepressant practice prescribing rate in primary care in England. *The Primary Care Companion for CNS Disorders*, 22(1), p.27082.

Hirschfeld, R.M., 2000. History and evolution of the monoamine hypothesis of depression. *Journal of clinical psychiatry*, 61(6), pp.4–6.

Horvath, K.J., Danilenko, G.P., Williams, M.L., Simoni, J., Amico, K.R., Oakes, J.M. and Simon Rosser, B., 2012. Technology use and reasons to participate in social networking health websites among people living with hiv in the us. *AIDS and Behavior*, 16(4), pp.900–910.

Househ, M., Borycki, E. and Kushniruk, A., 2014. Empowering patients through social media: the benefits and challenges. *Health informatics journal*, 20(1), pp.50–58.

Huang, J., Peng, M., Wang, H., Cao, J., Gao, W. and Zhang, X., 2017. A probabilistic method for emerging topic tracking in microblog stream. *World Wide Web*, 20(2), pp.325–350.

- Hughes, M.C., Pradier, M.F., Ross, A.S., McCoy, T.H., Perlis, R.H. and Doshi-Velez, F., 2020. Assessment of a prediction model for antidepressant treatment stability using supervised topic models. *JAMA network open*, 3(5), pp. e205308–e205308.
- Islam, M., Kabir, M.A., Ahmed, A., Kamal, A.R.M., Wang, H., Ulhaq, A. et al., 2018. Depression detection from social network data using machine learning techniques. *Health information science and systems*, 6(1), pp.1–12.
- Jackson, K.C. and St. Onge, E.L., 2003. Antidepressant pharmacotherapy: considerations for the pain clinician. *Pain Practice*, 3(2), pp.135–143.
- Kaplan, A.M. and Haenlein, M., 2010. Users of the world, unite! the challenges and opportunities of social media. *Business horizons*, 53(1), pp.59–68.
- Katon, W., Von Korff, M., Lin, E., Bush, T. and Ormel, J., 1992. Adequacy and duration of antidepressant treatment in primary care. *Medical care*, pp.67–76.
- Korkontzelos, I., Nikfarjam, A., Shardlow, M., Sarker, A., Ananiadou, S. and Gonzalez, G.H., 2016. Analysis of the effect of sentiment analysis on extracting adverse drug reactions from tweets and forum posts. *Journal of biomedical informatics*, 62, pp.148–158.
- Leaman, R., Wojtulewicz, L., Sullivan, R., Skariah, A., Yang, J. and Gonzalez, G., 2010. Towards internet- age pharmacovigilance: extracting adverse drug reactions from user posts to health-related social networks. *Proceedings of the 2010 workshop on biomedical natural language processing*. pp.117–125.
- Lexchin, J., Bero, L.A., Djulbegovic, B. and Clark, O., 2003. Pharmaceutical industry sponsorship and research outcome and quality: systematic review. *BMJ*, 326(7400), pp.1167–1170.
- Likhitha, S., Harish, B. and Kumar, H.K., 2019. A detailed survey on topic modeling for document and short text data. *International Journal of Computer Applications*, 178(39), pp.1–9.
- Lin, E.H., Von Korff, M., Katon, W., Bush, T., Simon, G.E., Walker, E. and Robinson, P., 1995. The role of the primary care physician in patients' adherence to antidepressant therapy. *Medical care*, pp.67–74.
- Lindelöf, G., Aledavood, T. and Keller, B., 2022. Vaccine discourse on twitter during the covid-19 pandemic. *arXiv preprint arXiv:2207.11521*.
- Lingam, R. and Scott, J., 2002. Treatment non-adherence in affective disorders. *Acta Psychiatrica Scandinavica*, 105(3), pp.164–172.
- López-Muñoz, F. and Alamo, C., 2009. Monoaminergic neurotransmission: the history of the discovery of antidepressants from 1950s until today. *Current pharmaceutical design*, 15(14), pp.1563–1586.
- Luo, J., Du, J., Tao, C., Xu, H. and Zhang, Y., 2020. Exploring temporal suicidal behavior patterns on social media: Insight from twitter analytics. *Health informatics journal*, 26(2), pp.738–752.
- Ma, P., Zeng-Treitler, Q. and Nelson, S.J., 2021. Use of two topic modeling methods to investigate covid vaccine hesitancy. *Int. conf. on human beings*. vol. 384, pp.221–226.

CONTACT INFORMATION

Your comments and questions are valued and encouraged. Contact the author at:

Amna Dridi

Birmingham City University

Millennium Point, 1 Curzon Street,

Birmingham, B4 7XG

Email: amna.dridi@bcu.ac.uk

