



Drugs Review Sentiment Analysis

Akshay K. Patil¹, Manoj V. Nikum^{*2}

Student Of MCA, SJRIT Dondaicha, KBC NMU Jalgaon, Maharashtra, India¹

Assistant Professor & HOD, MCA Department, SJRIT Dondaicha, KBC NMU JALGAON, Maharashtra, India^{*2}

Abstract: Since coronavirus has shown up, inaccessibility of legitimate clinical resources is at its peak, like the shortage of specialists and healthcare workers, lack of proper equipment and medicines etc. The entire medical fraternity is in distress, which results in numerous individual's demise. Due to unavailability, individuals started taking medication independently without appropriate consultation, making the health condition worse than usual. As of late, machine learning has been valuable in numerous applications, and there is an increase in innovative work for automation. This paper intends to present a drug recommender system that can drastically reduce specialists heap. In this research, we build a medicine recommendation system that uses patient reviews to predict the sentiment using various vectorization processes like Bow, TF-IDF, Word2Vec, and Manual Feature Analysis, which can help recommend the top drug for a given disease by different classification algorithms. The predicted sentiments were evaluated by precision, recall, f1score, accuracy, and AUC score. The results show that classifier LinearSVC using TF-IDF vectorization outperforms all other models with 93% accuracy.

Keywords: Index Terms—Drug, Recommender System, Machine Learning, NLP, Smote, Bow, TF-IDF, Word2Vec, Sentiment analysis

I. INTRODUCTION

With the number of coronavirus cases growing exponentially, the nations are facing a shortage of doctors, particularly in rural areas where the quantity of specialists is less compared to urban areas. A doctor takes roughly 6 to 12 years to procure the necessary qualifications. Thus, the number of doctors can't be expanded quickly in a short time frame. A Telemedicine framework ought to be energized as far as possible in this difficult time [1].

Clinical blunders are very regular nowadays. Over 200 thousand individuals in China and 100 thousand in the USA are affected every year because of prescription mistakes. Over 40% medicine, specialists make mistakes while prescribing since specialists compose the solution as referenced by their knowledge, which is very restricted [2][3]. Choosing the top-level medication is significant for patients who need specialists that know wide-based information about microscopic organisms, antibacterial medications, and patients [6]. Every day a new study comes up with accompanying more drugs, tests, accessible for clinical staff every day. Accordingly, it turns out to be progressively challenging for doctors to choose which treatment or medications to give to a patient based on indications, past clinical history.

With the exponential development of the web and the web-based business industry, item reviews have become an imperative and integral factor for acquiring items worldwide. Individuals worldwide become adjusted to analyze reviews and websites first before settling on a choice to buy a thing. While most of past exploration zeroed in on rating expectation and proposals on the E-Commerce field, the territory of medical care or clinical therapies has been infrequently taken care of. There has been an expansion in the number of individuals worried about their well-being and finding a diagnosis online. As demonstrated in a Pew American Research center survey directed in 2013 [5], roughly 60% of grown-ups searched online for health-related subjects, and around 35% of users looked for diagnosing health conditions on the web. A medication recommender framework is truly vital with the goal that it can assist specialists and help patients to build their knowledge of drugs on specific health conditions.

A recommender framework is a customary system that proposes an item to the user, dependent on their advantage and necessity. These frameworks employ the customers' surveys to break down their sentiment and suggest a recommendation for their exact need. In the drug recommender system, medicine is offered on a specific condition dependent on patient reviews using sentiment analysis and feature engineering. Sentiment analysis is a progression of strategies, methods, and tools for distinguishing and extracting emotional data, such as opinion and attitudes, from language [7]. On the other hand, Feature engineering is the process of making more features from the existing ones; it improves the performance of models.



This examination work separated into five segments: Intro- duction area which provides a short insight concerning the need of this research, Related works segment gives a concise insight regarding the previous examinations on this area of study, Methodology part includes the methods adopted in this research, The Result segment evaluates applied model results using various metrics, the Discussion section contains limita- tions of the framework, and lastly, the conclusion section.

II. LITERATURE SURVEY

With a sharp increment in AI advancement, there has been an exertion in applying machine learning and deep learning strategies to recommender frameworks. These days, recom- mender frameworks are very regular in the travel industry, e-commerce, restaurant, and so forth. Unfortunately, there are a limited number of studies available in the field of drug proposal framework utilizing sentiment analysis on the grounds that the medication reviews are substantially more intricate to analyze as it incorporates clinical wordings like infection names, reactions, a synthetic names that used in the production of the drug [8].

The study [9] presents GalenOWL, a semantic-empowered online framework, to help specialists discover details on the medications. The paper depicts a framework that suggests drugs for a patient based on the patient's infection, sensitivi- ties, and drug interactions. For empowering GalenOWL, clin- ical data and terminology first converted to ontological terms utilizing worldwide standards, such as ICD-10 and UNII, and then correctly combined with the clinical information.

Leilei Sun [10] examined large scale treatment records to locate the best treatment prescription for patients. The idea was to use an efficient semantic clustering algorithm estimating the similarities between treatment records. Likewise, the author created a framework to assess the adequacy of the suggested treatment. This structure can prescribe the best treatment regimens to new patients as per their demographic locations and medical complications. An Electronic Medical Record (EMR) of patients gathered from numerous clinics for testing. The result shows that this framework improves the cure rate. In this research [11], multilingual sentiment analysis was performed using Naive Bayes and Recurrent Neural Network (RNN). Google translator API was used to convert multilingual tweets into the English language. The results exhibit that RNN with 95.34% outperformed Naive Bayes, 77.21%.

The study [12] is based on the fact that the recommended drug should depend upon the patient's capacity. For example, if the patient's immunity is low, at that point, reliable medicines ought to be recommended. Proposed a risk level classification method to identify the patient's immunity. For example, in excess of 60 risk factors, hypertension, liquor addiction, and so forth have been adopted, which decide the patient's capacity to shield himself from infection. A web-based prototype system was also created, which uses a decision support system that helps doctors select first-line drugs. Xiaohong Jiang et al. [13] examined three distinct al- gorithms, decision tree algorithm, support vector machine (SVM), and backpropagation neural network on treatment data. SVM was picked for the medication proposal module as it performed truly well in each of the three unique bound- aries - model exactness, model proficiency, model versatility. Additionally, proposed the mistake check system to ensure analysis, precision and administration quality.

Mohammad Mehedi Hassan et al. [14] developed a cloud- assisted drug proposal (CADRE). As per patients' side effects, CADRE can suggest drugs with top-N related prescriptions. This proposed framework was initially founded on collabora- tive filtering techniques in which the medications are initially bunched into clusters as indicated by the functional description data. However, after considering its weaknesses like com- putationally costly, cold start, and information sparsity, the model is shifted to a cloud-helped approach using tensor decomposition for advancing the quality of experience of medication suggestion.

Considering the significance of hashtags in sentiment anal- ysis, Jiugang Li et al. [15] constructed a hashtag recommender framework that utilizes the skip-gram model and applied con- volutional neural networks (CNN) to learn semantic sentence vectors. These vectors use the features to classify hashtags using LSTM RNN. Results depict that this model beats the conventional models like SVM, Standard RNN. This explo- ration depends on the fact that it was undergoing regular AI methods like SVM and collaborative filtering techniques; the semantic features get lost, which has a vital influence in getting a decent expectation.

III. METHODOLOGIES

The dataset used in this research is Drug Review Dataset (Drugs.com) taken from the UCI ML repository [4]. This dataset contains six attributes, name of drug used (text), review (text) of a patient, condition (text) of a patient, useful



count (numerical) which suggest the number of individuals who found the review helpful, date (date) of review entry, and a 10-star patient rating (numerical) determining overall patient contentment. It contains a total of 215063 instances. Fig. 1 shows the proposed model used to build a medicine recommender system. It contains four stages, specifically, Data preparation, classification, evaluation, and Recommendation.

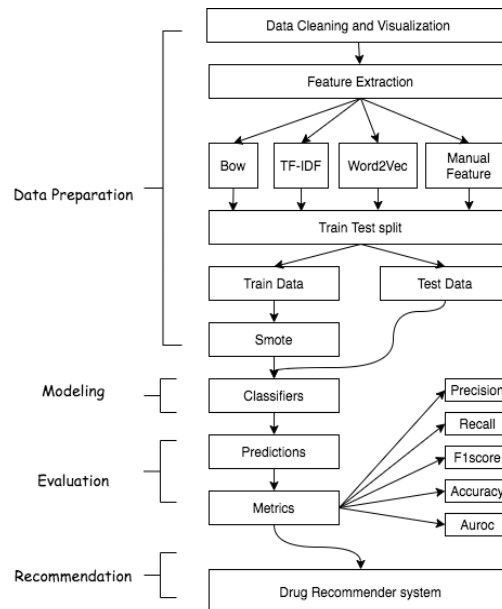


Fig. 1. Flowchart of the proposed model

A. Data Cleaning and Visualisation

Applied standard Data preparation techniques like checking null values, duplicate rows, removing unnecessary values, and text from rows in this research. Subsequently, removed all 1200 null values rows in the conditions column, as shown in Fig. 2. We make sure that a unique id should be unique to remove duplicacy.

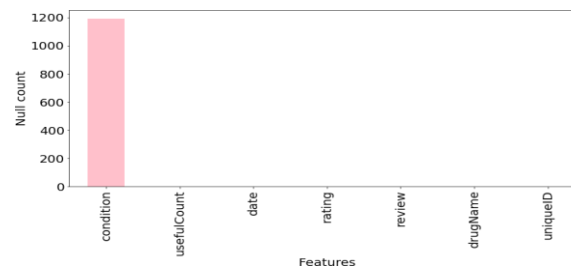


Fig. 2. Bar plot of the number of null values versus attributes

Fig. 3 shows the top 20 conditions that have a maximum number of drugs available. One thing to notice in this figure is that there are two green-colored columns, which shows the conditions that have no meaning. The removal of all these sorts of conditions from final dataset makes the total row count equals to 212141.

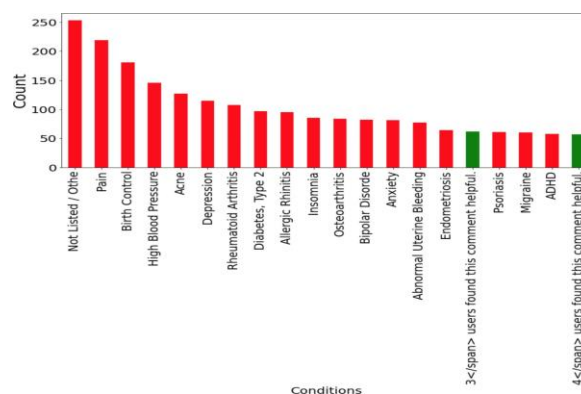


Fig. 3. Bar plot of Top 20 conditions that has a maximum number of drugs available



Fig. 4 shows the visualization of value counts of the 10-star rating system. The rating beneath or equivalent to five featured with cyan tone otherwise blue tone. The vast majority pick four qualities; 10, 9, 1, 8, and 10 are more than twice the same number. It shows that the positive level is higher than the negative, and people's responses are polar. The condition and drug column were joined with review text because the condition and medication words also have predictive power. Before proceeding to the feature extraction part, it is critical to clean up the review text before vectorization. This process is also known as text preprocessing. We first cleaned the reviews after removing HTML tags, punctuations, quotes, URLs, etc. The cleaned reviews were lowercased to avoid duplication, and tokenization was performed for converting the texts into small pieces called tokens. Additionally, stopwords,

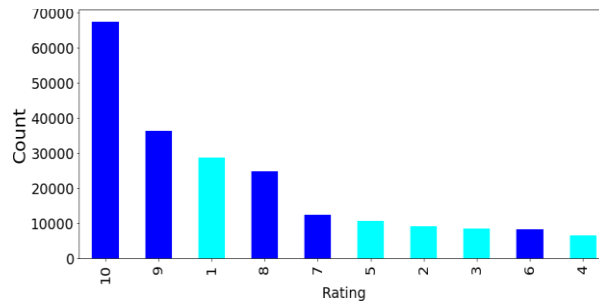


Fig. 4. Bar plot of count of rating values versus 10 rating number

for example, “a, to, all, we, with, etc.,” were removed from the corpus. The tokens were gotten back to their foundations by performing lemmatization on all tokens. For sentiment analysis, labeled every single review as positive and negative based on its user rating. If the user rating range between 6 to 10, then the review is positive else negative.

B. Feature Extraction

After text preprocessing, a proper set up of the data required to build classifiers for sentiment analysis. Machine learning algorithms can't work with text straightforwardly; it should be changed over into numerical format. In particular, vectors of numbers. A well known and straightforward strategy for feature extraction with text information used in this research is the bag of words (Bow) [16], TF-IDF [17], Word2Vec [18]. Also used some feature engineering techniques to extract features manually from the review column to create another model called manual feature aside from Bow, TF-IDF, and Word2Vec.

1) *Bow*: Bag of words [16] is an algorithm used in natural language processing responsible for counting the number of times of all the tokens in review or document. A term or token can be called one word (unigram), or any subjective number of words, n-grams. In this study, (1,2) n-gram range is chosen. Fig. 5 outlines how unigrams, digrams, and trigrams framed from a sentence. The Bow model experience a significant drawback, as it considers all the terms without contemplating how a few terms are exceptionally successive in the corpus, which in turn build a large matrix that is computationally expensive to train.

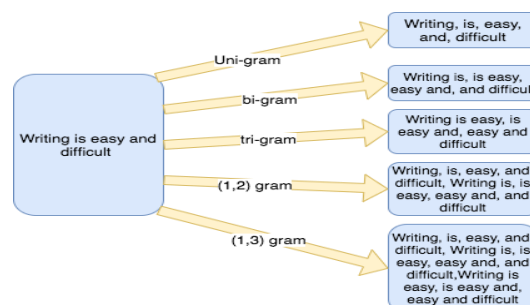


Fig. 5. Comparison of various types of grams framed from a sentence

2) *TF-IDF*: TF-IDF [17] is a popular weighting strategy in which words are offered with weight not count. The principle was to give low importance to the terms that often appear in the dataset, which implies TF-IDF estimates relevance, not a recurrence. Term frequency (TF) can be called the likelihood of locating a word in a document.

$$tf(t, d) = \log(1 + freq(t, d)) \quad (1)$$

Inverse document frequency (IDF) is the opposite of the number of times a specific term showed up in the whole corpus. It catches how a specific term is document specific.



C. Train Test Split

We created four datasets using Bow, TF-IDF, Word2Vec, and manual features. These four datasets were split into 75% of training and 25% of testing. While splitting the data, we set an equal random state to ensure the same set of random numbers generated for the train test split of all four generated datasets.

TABLE I
LIST OF FEATURES EXTRACTED MANUALLY FROM USER REVIEWS

Feature	Description
Punctuation	Counts the number of punctuation
Word	Counts the number of words
Stopwords	Counts the number of stopwords
Letter	Counts the number of letters
Unique	Counts the number of unique words
Average	Counts the mean length of words
Upper	Counts the uppercase words
Title	Counts the words present in title

IV. RESULTS

In this work, each review was classified as positive or negative, depending on the user's star rating. Ratings above five are classified as positive, while negative ratings are from one to five-star ratings. Initially, the number of positive ratings and negative ratings in training data were 111583 and 47522, respectively. After applying smote, we increased the minority class to have 70 percent of the majority class examples to curb the imbalances. The updated training data contains 111583 positive classes and 78108 negative classes. Four different text representation methods, namely Bow, TF-IDF, Word2Vec,

TABLE III
BAG-OF-WORDS

Model	Class	Prec	Rec	F1	Acc.	AUC
LogisticRegression	negative	0.85	0.87	0.86	0.91	0.90
	positive	0.94	0.93	0.94		
Perceptron	negative	0.87	0.85	0.86	0.92	0.898
	positive	0.94	0.94	0.94		
RidgeClassifier	negative	0.80	0.87	0.84	0.90	0.892
	positive	0.94	0.91	0.93		
MultinomialNB	negative	0.81	0.85	0.83	0.89	0.881
	positive	0.93	0.92	0.92		
SGDClassifier	negative	0.80	0.85	0.82	0.89	0.878
	positive	0.93	0.91	0.92		
LinearSVC	negative	0.84	0.87	0.86	0.91	0.90
	positive	0.94	0.93	0.94		

Table IV manifests the metrics on the TF-IDF vectorization method. LinearSVC increased the TF-IDF vectorization method performance to 93%, which is more noteworthy than the accuracy achieved by perceptron (91%) using bag of



words model. There was a close competition between LinearSVC, perceptron, and ridge classifier, with only a 1% difference. However, LinearSVC was picked as the best algorithm since the AUC score is 90.7%, which is greater than all other algorithms.

TABLE IV TF-IDF

Model	Class	Prec	Rec	F1	Acc.	AUC
LogisticRegression	negative	0.79	0.74	0.76	0.86	0.826
	positive	0.89	0.92	0.90		
Perceptron	negative	0.89	0.83	0.86	0.92	0.895
	positive	0.93	0.96	0.94		
RidgeClassifier	negative	0.89	0.84	0.86	0.92	0.897
	positive	0.93	0.95	0.95		
MultinomialNB	negative	0.85	0.83	0.84	0.90	0.883
	positive	0.93	0.94	0.93		
SGDClassifier	negative	0.76	0.57	0.65	0.82	0.745
	positive	0.83	0.92	0.88		
LinearSVC	negative	0.89	0.86	0.87	0.93	0.907
	positive	0.94	0.96	0.95		

The performance metrics of various classification methods on Word2Vec can be analyzed using Table V. The best accuracy is 91% by the LGBM model. Random forest and catboost classifier provide comparable sort of results whereas decision tree classifier performed poorly. Analyzing the region operating curve score, we can easily manifest that the LGBM has the highest AUC score of 88.3%.

TABLE V WORD2VEC

Model	Class	Prec	Rec	F1	Acc.	AUC
DecisionTree Classifier	negative	0.61	0.69	0.65	0.78	0.751
	positive	0.86	0.81	0.84		
RandomForest Classifier	negative	0.86	0.77	0.81	0.89	0.858
	positive	0.91	0.95	0.93		
LGBM Classifier	negative	0.86	0.82	0.84	0.91	0.883
	positive	0.93	0.94	0.93		
CatBoost Classifier	negative	0.81	0.79	0.80	0.88	0.855
	positive	0.91	0.92	0.92		

Table VI displays the performance metrics of four different classification algorithms on manually created features on user reviews. Compared to all other text classification methods, the results are not pretty impressive. However, the random forest achieved a good accuracy score of 88%.

TABLE VI
MANUAL FEATURE SELECTION

Model	Class	Prec	Rec	F1	Acc.	AUC
DecisionTree Classifier	negative	0.65	0.75	0.69	0.80	0.816
	positive	0.88	0.83	0.85		
RandomForest Classifier	negative	0.79	0.81	0.80	0.88	0.857
	positive	0.92	0.91	0.91		
LGBM Classifier	negative	0.74	0.74	0.74	0.85	0.787
	positive	0.89	0.89	0.89		
CatBoost Classifier	negative	0.72	0.73	0.73	0.84	0.804
	positive	0.88	0.88	0.88		

After evaluating all the models, the prediction results of Perceptron (Bow), LinearSVC (TF-IDF), LGBM (Word2Vec), and RandomForest (Manual Features) was added to give combined model predictions. The main intention is to make sure that the recommended top drugs should be classified correctly by all four models. If one model predicts it wrong,



then the drug's overall score will go down. These combined predictions were then multiplied with normalized useful count to get an overall score of each drug. This was done to check that enough people reviewed that drug. The overall score is divided by the total number of drugs per condition to get a mean score, which is the final score. Fig. 8 shows the top four drugs recommended by our model on top five conditions namely, Acne, Birth Control, High Blood Pressure, Pain and Depression.

condition	drugName	Score
Acne	Retin-A	0.069334
Acne	Atralin	0.088545
Acne	Magnesium hydroxide	0.088545
Acne	Retin A Micro	0.097399
Birth Control	Mono-Linyah	0.005448
Birth Control	Gildess Fe 1.5 / 30	0.005987
Birth Control	Ortho Miconor	0.006149
Birth Control	Lybrel	0.027766
High Blood Pressure	Adalat CC	0.303191
High Blood Pressure	Zestril	0.305851
High Blood Pressure	Toprol-XL	0.362589
High Blood Pressure	Labetalol	0.367021
Pain	Neurontin	0.158466
Pain	Nortriptyline	0.171771
Pain	Pamelor	0.231829
Pain	Elavil	0.304513
Depression	Remeron	0.124601
Depression	Sinequan	0.146486
Depression	Provigil	0.240185
Depression	Methylin ER	0.328604

Fig. 8. Recommendation of top four drugs on top five conditions

V. DISCUSSION

The results procured from each of the four methods are good, yet that doesn't show that the recommender framework is ready for real-life applications. It still need improvements. Predicted results show that the difference between the positive and negative class metrics indicates that the training data should be appropriately balanced using algorithms like Smote, Adasyn [24], SmoteTomek [25], etc. Proper hyperparameter optimization is also required for classification algorithms to improve the accuracy of the model. In the recommendation framework, we simply just added the best-predicted result of each method. For better results and understanding, require a proper ensembling of different predicted results. This paper intends to show only the methodology that one can use to extract sentiment from the data and perform classification to build a recommender system.

VI. CONCLUSION

Reviews are becoming an integral part of our daily lives; whether go for shopping, purchase something online or go to some restaurant, we first check the reviews to make the right decisions. Motivated by this, in this research sentiment analysis of drug reviews was studied to build a recommender system using different types of machine learning classifiers, such as Logistic Regression, Perceptron, Multinomial Naive Bayes, Ridge classifier, Stochastic gradient descent, LinearSVC, applied on Bow, TF-IDF, and classifiers such as Decision Tree, Random Forest, Lgbm, and Catboost were applied on Word2Vec and Manual features method. We evaluated them using five different metrics, precision, recall, f1score, accuracy, and AUC score, which reveal that the Linear SVC on TF-IDF outperforms all other models with 93% accuracy. On the other hand, the Decision tree classifier on Word2Vec showed the worst performance by achieving only 78% accuracy. We added best-predicted emotion values from each method, Perceptron on Bow (91%), LinearSVC on TF-IDF (93%), LGBM on Word2Vec (91%), Random Forest on manual features (88%), and multiply them by the normalized usefulCount to get the overall score of the drug by condition to build a recommender system. Future work involves comparison of different over-sampling techniques, using different values of n-grams, and optimization of algorithms to improve the performance of the recommender system.

REFERENCES

- [1]. Telemedicine, <https://www.mohfw.gov.in/pdf/Telemedicine.pdf>
- [2]. Wittich CM, Burkle CM, Lanier WL. Medication errors: an overview for clinicians. Mayo Clin Proc. 2014 Aug;89(8):1116-25.
- [3]. CHEN, M. R., & WANG, H. F. (2013). The reason and prevention of hospital medication errors. Practical Journal of Clinical Medicine, 4.



- [4]. Drug Review Dataset, <https://archive.ics.uci.edu/ml/datasets/Drug%2BReview%2BDataset%2B%2528Drugs.com%2529#>
- [5]. Fox, Susannah, and Maeve Duggan. "Health online 2013. 2013." URL: <http://pewinternet.org/Reports/2013/Health-online.aspx>
- [6]. Bartlett JG, Dowell SF, Mandell LA, File TM Jr, Musher DM, Fine MJ. Practice guidelines for the management of community-acquired pneumonia in adults. Infectious Diseases Society of America. Clin Infect Dis. 2000 Aug;31(2):347-82. doi: 10.1086/313954. Epub 2000 Sep 7. PMID: 10987697; PMCID: PMC7109923.
- [7]. Fox, Susannah & Duggan, Maeve. (2012). Health Online 2013. Pew Research Internet Project Report.
- [8]. T. N. Tekade and M. Emmanuel, "Probabilistic aspect mining approach for interpretation and evaluation of drug reviews," 2016 International Conference on Signal Processing, Communication, Power and Embedded System (SCOPES), Paralakhemundi, 2016, pp. 1471-1476, doi: 10.1109/SCOPES.2016.7955684.
- [9]. Doulaverakis, C., Nikolaidis, G., Kleontas, A. et al. GalenOWL: Ontology-based drug recommendations discovery. J Biomed Semant 3, 14 (2012). <https://doi.org/10.1186/2041-1480-3-14>
- [10]. Leilei Sun, Chuanren Liu, Chonghui Guo, Hui Xiong, and Yanming Xie. 2016. Data-driven Automatic Treatment Regimen Development and Recommendation. In Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD '16). Association for Computing Machinery, New York, NY, USA, 1865–1874. DOI:<https://doi.org/10.1145/2939672.2939866>
- [11]. V. Goel, A. K. Gupta and N. Kumar, "Sentiment Analysis of Multilingual Twitter Data using Natural Language Processing," 2018 8th International Conference on Communication Systems and Network Technologies (CSNT), Bhopal, India, 2018, pp. 208-212, doi: 10.1109/CSNT.2018.8820254.
- [12]. Shimada K, Takada H, Mitsuyama S, et al. Drug-recommendation system for patients with infectious diseases. AMIA Annu Symp Proc. 2005;2005:1112 Y. Bao and X. Jiang, "An intelligent medicine recommender system framework," 2016 IEEE 11th Conference on Industrial Electronics and Applications (ICIEA), Hefei, 2016, pp. 1383-1388, doi: 10.1109/ICIEA.2016.7603801.
- [13]. Zhang, Yin & Zhang, Dafang & Hassan, Mohammad & Alamri, Atif & Peng, Limei. (2014). CADRE: Cloud-Assisted Drug Recommendation Service for Online Pharmacies. Mobile Networks and Applications. 20. 348-355. 10.1007/s11036-014-0537-4.
- [14]. J. Li, H. Xu, X. He, J. Deng and X. Sun, "Tweet modeling with LSTM recurrent neural networks for hashtag recommendation," 2016 International Joint Conference on Neural Networks (IJCNN), Vancouver, BC, 2016, pp. 1570-1577, doi: 10.1109/IJCNN.2016.7727385.
- [15]. Zhang, Yin & Jin, Rong & Zhou, Zhi-Hua. (2010). Understanding bag-of-words model: A statistical framework. International Journal of Machine Learning and Cybernetics. 1. 43-52. 10.1007/s13042-010-0001-0.
- [16]. J. Ramos et al., "Using tf-idf to determine word relevance in document queries," in Proceedings of the first instructional conference on machinelearning, vol. 242, pp. 133–142, Piscataway, NJ, 2003.
- [17]. Yoav Goldberg and Omer Levy. word2vec Explained: deriving Mikolov et al.'s negative-sampling word-embedding method, 2014; arXiv:1402.3722.
- [18]. Danushka Bollegala, Takanori Maehara and Kenichi Kawarabayashi. Unsupervised Cross-Domain Word Representation Learning, 2015; arXiv:1505.07184.
- [19]. Textblob, <https://textblob.readthedocs.io/en/dev/>.
- [20]. van der Maaten, Laurens & Hinton, Geoffrey. (2008). Visualizing data using t-SNE. Journal of Machine Learning Research. 9. 2579-2605.
- [21]. N. V. Chawla, K. W. Bowyer, L. O. Hall and W. P. Kegelmeyer. SMOTE: Synthetic Minority Over-sampling Technique, 2011, Journal Of Artificial Intelligence Research, Volume 16, pages 321-357, 2002; arXiv:1106.1813. DOI: 10.1613/jair.953.
- [22]. Powers, David & Ailab,. (2011). Evaluation: From precision, recall and F-measure to ROC, informedness, markedness & correlation. J. Mach.Learn. Technol. 2. 2229-3981. 10.9735/2229-3981
- [23]. Haibo He, Yang Bai, E. A. Garcia and Shutao Li, "ADASYN: Adaptive synthetic sampling approach for imbalanced learning," 2008 IEEE International Joint Conference on Neural Networks (IEEE World Congress on Computational Intelligence), Hong Kong, 2008, pp. 1322-1328, doi: 10.1109/IJCNN.2008.4633969.
- [24]. Z. Wang, C. Wu, K. Zheng, X. Niu and X. Wang, "SMOTETomek-Based Resampling for Personality Recognition," in IEEE Access, vol. 7, pp. 129678-129689, 2019, doi: 10.1109/ACCESS.2019.2940061.