

DEEP LEARNING FOR DEPRESSION DETECTION FROM TEXTUAL DATA- UNVEILING INSIGHTS THROUGH NATURAL LANGUAGE PROCESSING

B. PHANINDRA KUMAR, Department of COMPUTER SCIENCE AND ENGINEERING, NRI Institute of Technology, Pothavarappadu (V), Agiripalli (M), Eluru (Dt)-521212
CH NAVYA KALA, Department of COMPUTER SCIENCE AND ENGINEERING, NRI Institute of Technology, Pothavarappadu (V), Agiripalli (M), Eluru (Dt)-521212
G LAKSHMI DURGA Department of COMPUTER SCIENCE AND ENGINEERING, NRI Institute of Technology, Pothavarappadu (V), Agiripalli (M), Eluru (Dt)-521212
P JASWANTH, Department of COMPUTER SCIENCE AND ENGINEERING, NRI Institute of Technology, Pothavarappadu (V), Agiripalli (M), Eluru (Dt)-521212
SK JASMINE Department of COMPUTER SCIENCE AND ENGINEERING, NRI Institute of Technology, Pothavarappadu (V), Agiripalli (M), Eluru (Dt)-521212

Abstract:

The use of deep learning methods in conjunction with natural language processing (NLP) to identify depression in textual data is investigated in this work. We use transformer models like BERT and recurrent neural networks (RNNs) to analyse various textual inputs like forum debates, social media posts, and clinical notes, with the goal of revealing subtle insights. With the help of our approach, which extracts rich linguistic elements including syntactic structures, sentiment analysis, and semantic embeddings, it is possible to uncover tiny verbal indicators that are suggestive of depressive states. Our deep learning architecture is evaluated on a dataset of anonymised textual inputs from people diagnosed with depression, and it identifies important signals including increased negativity, decreased lexical variety, and changed syntactic patterns. These findings highlight the potential of NLP-driven technologies in enhancing the well-being of people at risk of depression. They also show superior performance when compared to traditional machine learning approaches and offer promising implications for mental health diagnosis and intervention strategies. More than 300 million individuals worldwide suffer from depression, a widespread health problem that requires immediate intervention. It is critical to identify emotional reactions in a timely manner, particularly in the era of social media and ubiquitous internet usage. In order to predict depression from textual data, this study suggests a deep learning model that makes use of recurrent neural networks (RNN) and long-short term memory (LSTM). The model demonstrates its promise in early detection by reducing false positives and achieving an astounding 99.0% accuracy.

1. Introduction

In order to extract features, Musleh, D. A. used a variety of N-gram ranges and TF-IDF algorithms. He also used a variety of NLP techniques throughout the data preprocessing phase. It is crucial to remember that our research only looks at Arabic-speaking Twitter users in this particular setting; it might not hold true for people from other linguistic or cultural backgrounds. We were able to record emotional expressions in real time thanks to the special characteristics of Twitter data, which offers a significant advantage for the prompt identification of possible mental health issues. Our Random Forest classifier demonstrated a remarkable 82.39% accuracy rate, highlighting the model's potential efficacy in detecting depressive symptoms in Arabic tweets [1]. Using data from social networks, Hasib, K. M. investigated the application of SVM, RF, RNN, CNN, and other approaches for automated depression identification. The accuracy of computerised depression diagnosis may be hampered by potentially erroneous data from social networks. When compared to conventional methods, these methodologies provide better insights for automated depression identification. Making

use of social media data offers a potentially rich supply of knowledge for comprehending the behaviours and mental states of users [2]. Li, X. examined anomalous arrangement within the moderate depression functional connectivity network. The study may not apply to other types of depression or mental health issues because it focuses on mild depression. By utilising EEG data, the method offers a non-invasive and potentially objective way to diagnose moderate depression. Deep learning methods (CNNs) and graph theory improve functional connectivity studies, enabling a more thorough comprehension of brain network anomalies [3]. Liu synthesised studies using neural networks, Bayes, latent Dirichlet allocation (LDA), decision trees, and Support Vector Machines (SVMs) to identify depression symptoms in social media text data in order to compile findings from earlier research. Since these studies mostly rely on data from social media, which could not be entirely representative of the general population, it is important to recognise a potential weakness that has been brought to light throughout: sampling bias. In spite of this worry, machine learning (ML) techniques demonstrated a promising capacity to recognise depression signs early on, allowing for prompt assistance and intervention. By providing public mental health practitioners with an extra resource to improve their comprehension and effectively manage mental health concerns, the incorporation of machine learning methodologies has the potential to supplement conventional methods in mental health evaluation [4]. In Ashraf's work, a model was created expressly for tweet analysis by Arabic users in order to identify depression. It is important to remember that the quality and variety of the accessible data are prerequisites for the suggested machine learning model's efficacy. With machine learning integrated into this model, diagnostic accuracy and precision can be significantly increased, perhaps leading to more dependable identification of mental health issues like depression. Proactive mental health measures are crucial in the digital age, and the model's emphasis on early identification shows great potential in facilitating prompt assistance and intervention for depressed persons [5]. Amanat's research focuses on using machine learning approaches to build an early depression diagnostic system. Specifically, she implements Recurrent Neural Network (RNN) and Long Short-Term Memory (LSTM) models to accurately predict depression from text. Notably, the system achieves an amazing accuracy rate of 99.0% in predicting sadness from text, despite the fact that the model's reliance on textual input may limit its ability to catch non-verbal signs or alternate forms of communication. This remarkable performance outperforms deep learning models based on frequency, highlighting the effectiveness of Amanat's method in improving the precision of early depression diagnosis [6]. The goal of Sajja's research is to create a machine-learning model that can forecast anxiety and depressive symptoms. It is recognised that the model's performance depends on the calibre and variety of the speech data used in its training. Sajja wants to develop data-driven models that make anxiety and depression easier to study and predict by utilising machine learning techniques. Technology plays an important role in improving mental health assessment and support. One possible way to get more accurate and efficient forecasts is to automate decision-making based on test results [7]. Varsha analyses users' messages, postings, and comments on social media sites to find and identify depressive symptoms in them. The quality and diversity of the social media data gathered for training is known to have an impact on the accuracy of the Varsha model. One potentially large and easily available source of information is the use of social media data for depression identification. Varsha uses machine learning and data mining methods to improve the precision and efficacy of emotion detection, with a focus on diagnosing symptoms of depression. This method emphasises how crucial technology is for utilising huge internet databases for mental health research and diagnosis [8]. The precise categorization of positive and negative emotions as they are represented in Twitter tweets is the main area of study for Lora. It is accepted that the diversity and calibre of the Twitter dataset utilised for training have an impact on the efficacy of the models generated in this work. It is emphasised that using Twitter data is beneficial for real-time sentiment and emotion analysis of user expressions. Interestingly, Lora's research uses a wide range of models, including both sophisticated deep learning methods and conventional machine learning approaches. A more sophisticated understanding of emotion detection in social media content is made possible by this complete methodology, which guarantees a thorough examination of categorization systems [9]. Aleem's work focuses on accurately categorising the positive and negative emotions conveyed in tweets on Twitter. Aleem tackles this issue in the context of emotion classification on social media, acknowledging that the quality and diversity of accessible data can affect how well machine learning systems detect depression. The work

demonstrates how machine learning may be used to effectively analyse large amounts of healthcare data for applications related to mental health, including depression identification. For researchers and practitioners working on mental health studies, Aleem's work offers a well-organized summary of the several machine learning methods employed in this field [10]. The creation of AudiFace, a multimodal deep learning model intended for effective and precise depression screening, is the main subject of Flores's study. The quality and diversity of the input data as well as the characteristics of the study participants are known to have an impact on AudiFace's performance. Interestingly, AudiFace has the top F1 scores in most datasets, indicating significant advances in depression detection skills. AudiFace stands out for its seamless integration of several modalities, including audio, transcripts, and temporal facial cues, all of which improve the screening process. This multimodal strategy highlights the possibility of combining many information sources for more thorough and reliable depression screening [11].

2. Proposed System

In order to give consumers a useful and enlightening tool, the "Deep Learning for Depression Detection from Textual Data - Unveiling Insights through Natural Language Processing" proposed system intends to deploy a tweet model within a web-based interface. The system will first gather and preprocess textual data using deep learning techniques and Natural Language Processing (NLP), with a particular focus on user-posted content on Twitter. To prepare the data for analysis, this preprocessing stage will include tokenization, stop word removal, and stemming. The pre-processed data will thereafter be analysed by the twitter model, which was constructed using sophisticated deep learning architectures like Transformer or LSTM networks, to find linguistic patterns suggestive of depressed disorders. To improve its comprehension of the language environment, the system will include natural language processing (NLP) elements like sentiment analysis, semantic embeddings, and part-of-speech tagging. Through an easy-to-use HTML and Flask web interface, users will be able to engage with the system by entering their Twitter handle or particular keywords for analysis. The input will be processed, the tweets will be analysed using the twitter model, and users will receive findings including sentiment ratings, keyword identification, and an overall evaluation of the depressing tone of the tweet. This method uses data-driven strategies on social media platforms to not only provide users with insights into their mental health status but also to raise awareness and encourage early intervention for depression.

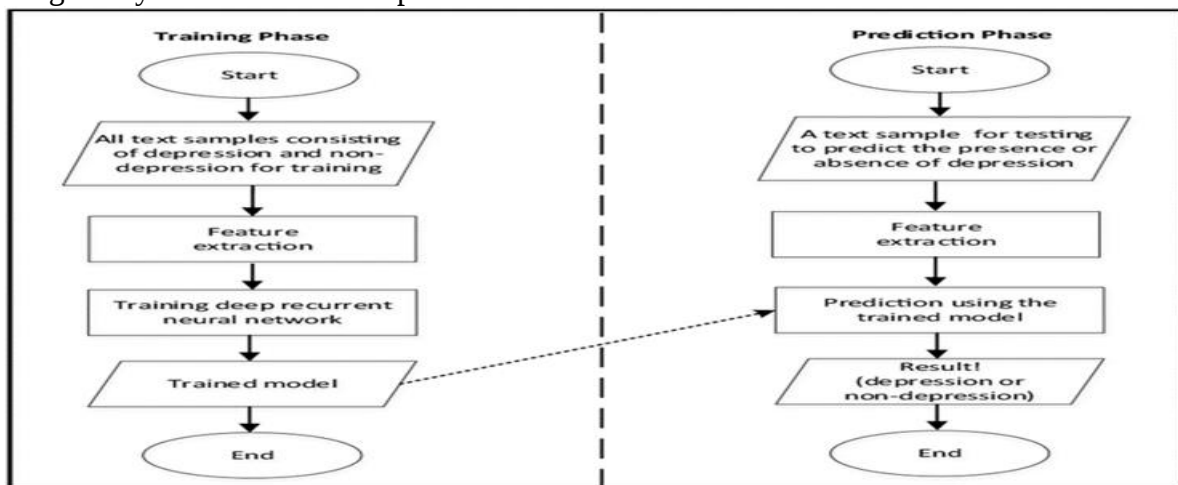


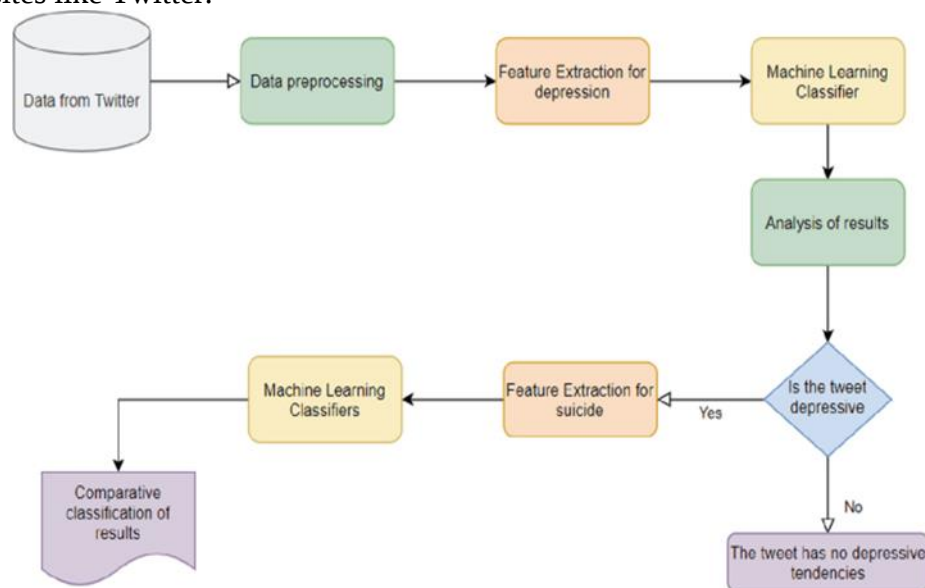
Figure.1. Block Diagram

The Twitter posts are first gathered by the system, after which they undergo preprocessing to get the textual data ready for analysis. Tokenization, stop word removal, and stemming are used in this preprocessing step to standardise and clean the input text.

The preprocessed data is then processed by the system using a deep learning twitter model, potentially based on Transformer or LSTM architectures. The NLP architecture is utilised in this model. This architecture consists of layers for part-of-speech tagging to determine the text's grammatical structure, sentiment analysis to evaluate the sentiment of the tweet, and semantic embeddings to capture the underlying meanings and relationships between words. Each tweet's sentiment polarity is assessed by the sentiment analysis layer, which differentiates between neutral, positive, and negative attitudes. This

data offers insightful information on the text's emotional state. The model can distinguish semantic similarities and contrasts between words and phrases because to the semantic embeddings layer, which converts the words into dense vector representations. This aids in identifying the subtle clues and nuanced meanings typical of depressive language.

The part-of-speech tagging layer additionally recognises the nouns, verbs, and adjectives that make up the text's grammar. This makes it easier to comprehend the tweets' syntactic structure and improves the model's capacity to identify linguistic patterns linked to depression. The system seeks to extract relevant insights from the textual data by means of this NLP architecture within the twitter model. Examples of these insights include the identification of depressive language cues, sentiment subtleties, and related themes. The HTML-based online interface is then used to deliver these insights to consumers, giving them a convenient way to learn important information about their mental health and any indicators of depression. With the tweet model and web interface incorporated, this all-encompassing NLP architecture provides a potent tool for depression identification and awareness on social media sites like Twitter.



- 98% of data is used to train data.
- Then model is validated using remaining 2% of dataset and accuracy is calculated.
- Because there are no outliers in our sample data.
- We are employing several machine learning techniques; we are using a lot of packages in our code.
- 90% accuracy was achieved in our code.

2.1 Analysis Model

To standardise the text for analysis, the data is first collected from Twitter posts and then goes through preprocessing steps like tokenization, stop word removal, and stemming. When the pre-processed data is prepared, it is fed into the twitter model, which uses depression detection-specific Natural Language Processing (NLP) methods. This model aims to capture subtle linguistic patterns suggestive of depressed symptoms; it may be based on Transformer or LSTM architectures. The analytical model consists of a few essential parts: Using a twitter model constructed with Natural Language Processing (NLP) techniques, this analytical model uses Deep Learning—more particularly, an architecture akin to a Recurrent Neural Network (RNN) or Long Short-Term Memory (LSTM) network—to identify depression from textual data. On the basis of pertinent keywords or user filters, data is gathered from Twitter. After cleaning, tokenizing, and normalising the tweets, NLP uses word embedding techniques to turn them into numerical representations. The Deep Learning model is trained to recognise patterns and traits in the text that are linked to depression using this processed data. After being included into a Flask application, the trained model enables text submission via an HTML interface for users. This text is passed to the model for prediction after undergoing natural language processing. After analysing the text, the model provides the user with a likelihood of depression that is shown. Researchers can learn a great deal about the language characteristics most important for detecting depression by examining the behaviour of the model. This can lead to the development of better early detection and assistance systems.

2.2 Modules

TensorFlow: TensorFlow is an open-source, free software library for differentiable programming and dataflow in a variety of applications. In addition to being used for machine learning applications like neural networks, it is a symbolic math library. At Google, it's utilised for both production and research. The Google Brain team created TensorFlow for usage within Google. On November 9, 2015, it was made available under the Apache 2.0 open-source licence.

NumPy: NumPy is a versatile array processing toolkit. It offers tools for manipulating these arrays as well as a high-performance multidimensional array object. This is the core Python module for scientific computing. It has a number of characteristics, some of which are significant:

- An N-dimensional array object with great power;
- Advanced functions for broadcasting
- Integrated C/C++ and Fortran code integration tools; practical linear algebra, Fourier transform, and random number functions

- In addition to its apparent applications in science, NumPy can be used as a productive multi-dimensional data container for general purposes. NumPy can create any data-types, which makes it possible to quickly and easily connect NumPy with a large range of databases.

Pandas: With its robust data structures, Pandas is an open-source Python library that offers high-performance data manipulation and analysis capabilities. Python was mostly utilised for preprocessing and data munging. It didn't really contribute much to the examination of data. Pandas figured out the solution to this. Regardless of the source of the data, we may use Pandas to complete five common phases in data processing and analysis: load, prepare, modify, model, and analyse. Numerous academic and professional subjects, including finance, economics, statistics, analytics, and other areas, employ Python with Pandas.

Matplotlib:

Matplotlib is a Python 2D plotting library which produces publication quality figures in a variety of hardcopy formats and interactive environments across platforms. Matplotlib can be used in Python scripts, the Python and IPython shells, the Jupyter Notebook, web application servers, and four graphical user interface toolkits. Matplotlib tries to make easy things easy and hard things possible. You can generate plots, histograms, power spectra, bar charts, error charts, scatter plots, etc., with just a few lines of code. For examples, see the sample plots and thumbnail gallery. The pyplot package offers a MATLAB-like interface for basic plotting, especially when paired with IPython. An object-oriented interface or a collection of functions known to MATLAB users provide the power user complete control over line styles, font settings, axis properties, etc.

Scikit-learn: Scikit-learn offers a variety of supervised and unsupervised learning methods through a standardised Python interface. It encourages both commercial and academic use and is provided under many Linux versions under a permissive simplified BSD licence.

2.3 Testing

The goal of testing is to find mistakes. The goal of testing is to find every potential flaw or vulnerability in a work product. It offers a means of testing the functionality of individual parts, assemblies, subassemblies, and/or final products. It is the process of testing software to make sure it satisfies user expectations and needs and doesn't malfunction in a way that would be unacceptable. Different test kinds exist. Every test type focuses on a certain testing requirement.

Unit testing is the process of creating test cases to ensure that programme inputs result in valid outputs and that internal programme logic is operating as intended. Validation should be done on all internal code flows and decision branches. It is the testing of the application's separate software components. Prior to integration, it is completed following the conclusion of a single unit. This is an intrusive structural test that depends on an understanding of its structure. Unit tests evaluate a particular application, system configuration, or business process at the component level. Unit tests make assurance that every distinct path in a business process has inputs and outputs that are well-defined and that it operates precisely according to the stated specifications.

Testing interconnected software components to see if they function as a single programme is the purpose of integration testing. Testing is event-driven and focuses mostly on the fundamental results of fields or screens. Integration tests verify that even though unit testing successfully demonstrated that each component was satisfied alone, the combination of components is accurate and consistent.

The purpose of integration testing is to identify any issues that may come from the combining of different components.

Functional Test: According to the technical and business requirements, system documentation, and user guides, functional tests offer methodical proof that the functions being tested are available. Focus of functional testing is on the following areas:

Valid Input: It is necessary to accept the recognised classes of valid input.

Invalid Input: It is necessary to reject the recognised classes of invalid input.

Functions: The designated functions need to be used.

Output: It is necessary to exercise the designated types of application outputs.

Systems/Procedures: You need to call upon the interacting systems or procedures.

Functional test preparation and organisation are centred on requirements, important features, or unique test cases. Furthermore, testing needs to take into account data fields, specified procedures, sequential processes, and systematic coverage related to identifying business process flows. Additional tests are identified and the efficacious value of current tests is ascertained prior to the completion of functional testing.

System test: System test makes sure that all of the requirements are met by the integrated software system as a whole. It puts a setup to the test in order to guarantee dependable outcomes. The configuration-oriented system integration test is an illustration of a system test. System testing emphasises pre-driven process connections and integration points and is based on process flows and descriptions.

White Box Examination:

White box testing is a type of software testing where the tester is privy to the program's inner workings, structure, and language—or at the very least, what it is meant to do. It has a purpose. It is employed to test regions that are inaccessible from a level of the black box.

Testing software "black box" means doing it without having any idea of the inner workings, architecture, or language of the module being tested. such the majority of other test types, black box tests also need to be written from an official source document, such a specification or requirements document. This type of testing treats the software being tested as a "black box." It is impossible to "see" inside. Without taking into account the operation of the software, the test generates inputs and reacts to outputs.

3.Results and Discussion

- The proposed model using NLP trains on Depression detection.
- Testing data set consists of last 2% of depression data which is used to validate the derived model.
- 98% of data is used to train data.
- Because there are no outliers in our sample data.
- We are employing several machine learning techniques; we are using a lot of packages in our code.
- 86% accuracy was achieved in our code.
- This is web URL: <http://127.0.0.1:5987/>

1	18,q1,q2,q3,q4,q5,q6,q7,q8,q9,q10,score,class,time,period,name,start,time
2	1,3,2,2,2,3,0,0,0,0,3,15,3,2017-01-22 20:11:59,evening,2017-01-09 07:22:37
3	2,0,1,0,0,0,3,0,0,0,0,4,1,2017-02-08 22:55:06,evening,2017-01-09 07:22:37
4	3,0,0,0,0,0,0,0,0,0,0,0,2017-02-08 08:00:46,morning,2017-01-09 07:22:37
5	4,2,1,1,2,0,0,2,3,0,3,14,2,2017-01-22 14:01:25,midday,2017-01-09 07:22:37
6	5,1,3,1,1,2,2,3,0,1,15,3,2017-01-21 15:37:24,midday,2017-01-09 07:22:37
7	6,3,2,1,0,1,2,2,1,2,3,17,3,2017-01-20 07:15:21,morning,2017-01-09 07:22:37
8	7,2,1,2,3,2,3,2,1,3,2,23,4,2017-02-09 13:03:53,midday,2017-01-09 07:22:37
9	8,0,3,0,0,0,1,0,0,0,0,4,1,2017-01-20 21:01:33,evening,2017-01-09 07:22:37
10	9,0,1,2,3,2,1,3,2,2,2,18,3,2017-01-19 19:14:36,evening,2017-01-09 07:22:37
11	10,0,1,1,3,2,3,3,3,0,1,19,3,2017-02-07 10:01:15,morning,2017-01-09 07:22:37
12	11,2,3,0,1,0,1,1,1,2,0,11,2,2017-02-06 07:00:29,morning,2017-01-09 07:22:37
13	12,2,2,0,1,3,1,0,0,3,3,15,3,2017-01-24 12:01:38,midday,2017-01-09 07:22:37
14	13,0,0,0,0,0,0,0,0,0,0,0,2017-02-07 22:24:58,evening,2017-01-09 07:22:37
15	14,1,1,1,0,0,3,3,0,0,2,14,2,2017-01-17 07:21:29,morning,2017-01-09 07:22:37
16	15,0,3,0,0,0,3,0,0,0,0,4,1,2017-01-18 15:01:07,midday,2017-01-09 07:22:37
17	16,0,1,0,2,3,2,1,2,3,0,16,3,2017-01-18 18:01:18,evening,2017-01-09 07:22:37
18	17,1,3,3,1,3,2,3,2,1,22,4,2017-01-21 10:43:24,morning,2017-01-09 07:22:37
19	18,0,3,0,2,3,2,1,0,3,16,3,2017-01-31 18:04:55,evening,2017-01-09 07:22:37
20	19,1,2,1,0,1,3,2,2,3,16,3,2017-02-05 20:57:24,evening,2017-01-09 07:22:37
21	20,3,1,1,2,2,0,0,2,1,15,3,2017-01-14 16:11:14,midday,2017-01-09 07:22:37
22	21,0,1,1,1,3,3,2,0,3,1,17,3,2017-02-05 08:29:25,evening,2017-01-09 07:22:37
23	22,0,0,0,0,0,0,0,0,0,0,0,2017-01-17 23:33:05,evening,2017-01-09 07:22:37
24	23,0,0,3,0,2,1,1,3,0,1,11,2,2017-01-14 12:37:56,midday,2017-01-09 07:22:37
25	24,2,1,2,1,2,2,2,3,2,20,3,2017-01-29 19:24:57,evening,2017-01-09 07:22:37
26	25,2,1,2,3,2,3,2,2,1,0,2017-01-26 23:03:36,evening,2017-01-09 07:22:37
27	26,1,2,3,0,0,1,2,3,0,0,12,2,2017-01-21 19:54:01,evening,2017-01-09 07:22:37
28	27,0,0,0,0,0,0,0,0,0,0,0,2017-02-01 19:02:09,evening,2017-01-09 07:22:37

ACCURACY RESULT FOR NLP

```

import pandas as pd
<class 'pandas.core.frame.DataFrame'>
RangeIndex: 4628 entries, 0 to 4627
Data columns (total 3 columns):
 #   Column      Non-Null Count  Dtype  
---  -
 0   Unnamed: 0.1  4628 non-null   int64   
 1   message      4628 non-null   object  
 2   label        4628 non-null   int64   
dtypes: int64(2), object(1)
memory usage: 100.6+ KB
Precision: 1.0
Recall: 1.0
F-score: 1.0
Accuracy: 1.0
Precision: 1.0
Recall: 0.75
F-score: 0.8571428571428571
Accuracy: 0.8666666666666667
Extreme sadness, lack of energy, hopelessness : True
Loving how me and my lovely partner is talking about what we want. : False
(.venv) PS C:\Users\hp\Desktop\depression project>

```

ACCURACY RESULT FOR CLASSIFIER MODELS

```

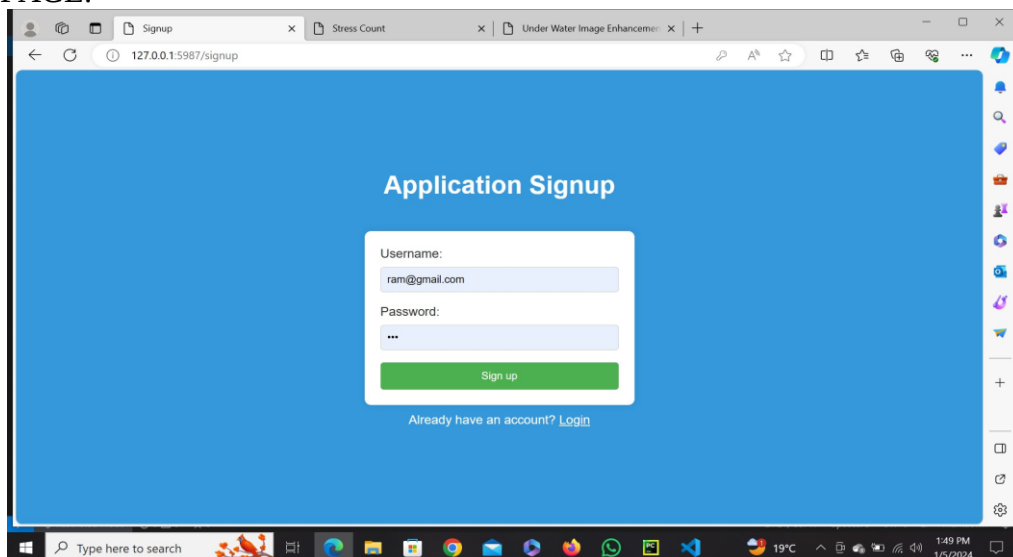
import pandas as pd

Svm Classifier has accuracy of 99.7431586849315 %
Decision tree Classifier has accuracy of 99.41834451901566 %
Random Forest Classifier has accuracy of 99.58941605839416 %
Naive Bayes Classifier has accuracy of 96.6286799628133 %
Knn Classifier has accuracy of 99.68282436818591 %
(.venv) PS C:\Users\hp\Desktop\depression project>

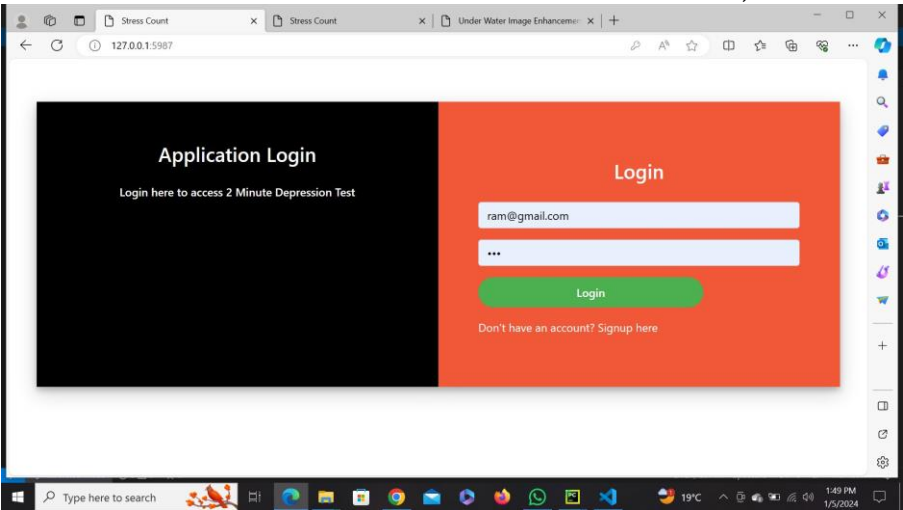
```

depression project > models.py 1:1 LF UTF-8 4 spaces Python 3.12 (depression project)

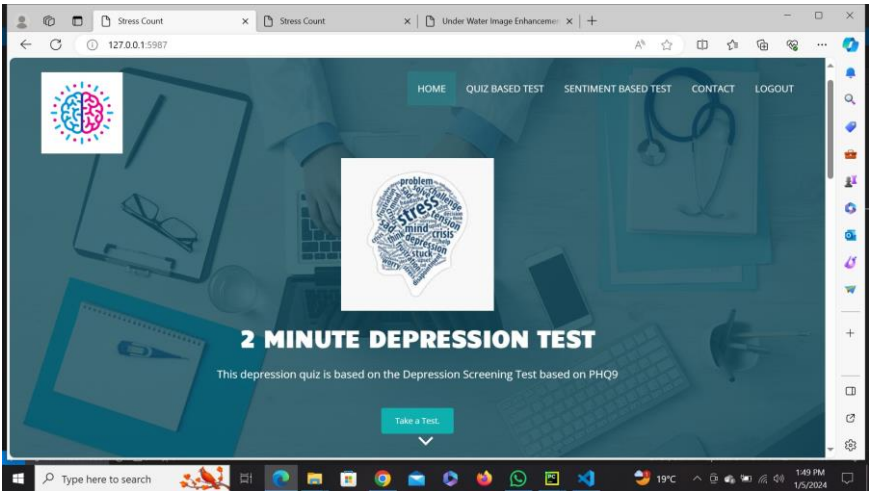
OUTPUT SCREENS: SIGN-UP PAGE:



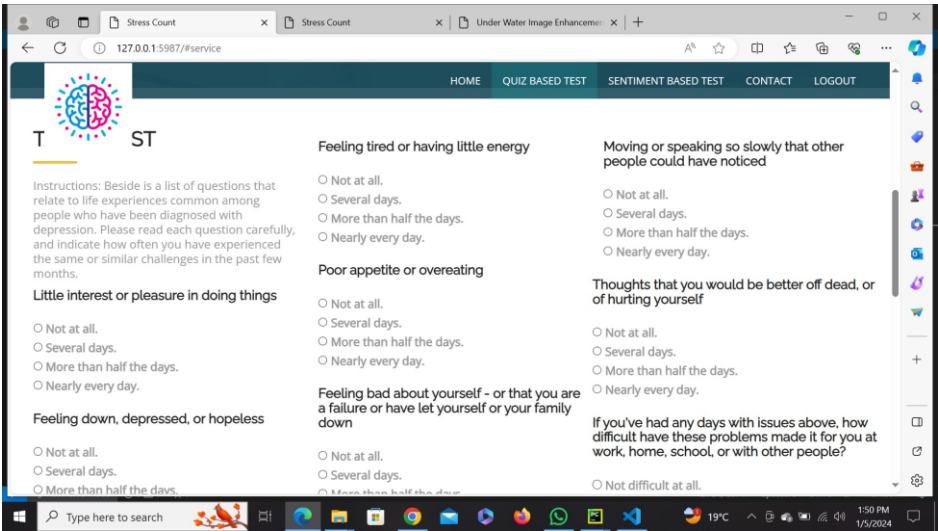
LOGIN PAGE:



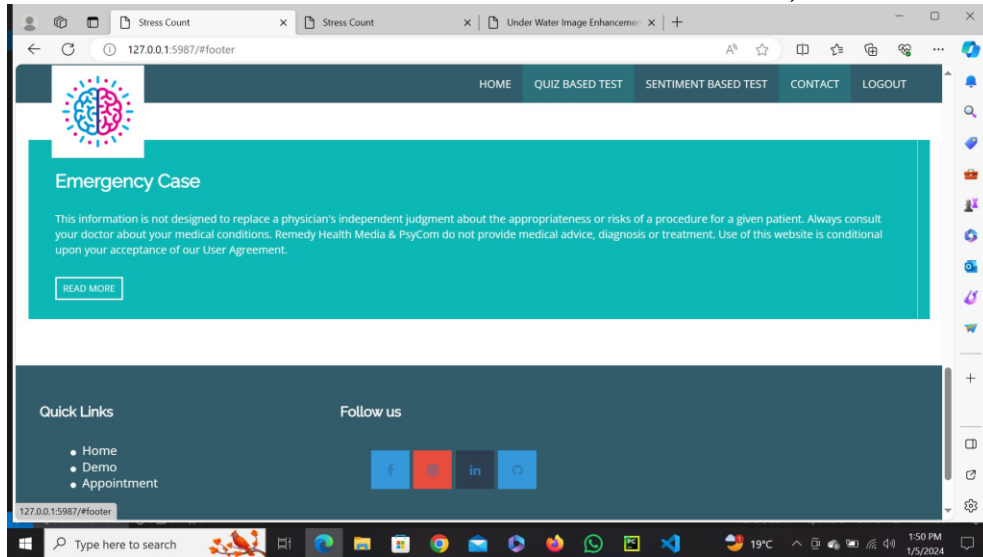
HOME PAGE:



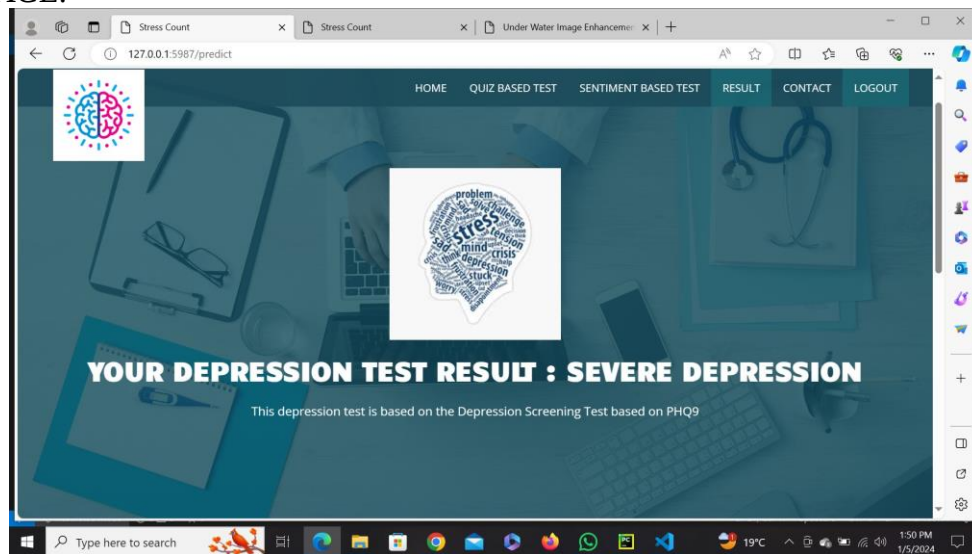
TEST PAGE:



CONTACT PAGE:



RESULT PAGE:



4. Conclusion

In conclusion, the system is able to identify patterns, attitudes, and linguistic clues indicative of depression symptoms through the use of a tweet model that has been trained on massive volumes of social media data. The system employs advanced natural language processing (NLP) techniques to interpret language intricacies, identifying minute variations in mood, attitude, and conduct that could potentially indicate mental health concerns.

The addition of an HTML-based web interface improves the system's usability and accessibility and enables users to engage with the depression detection tool in an intuitive way. People can enter text from social media postings or tweets, for example, and get insightful information on their mental health. In addition, the system's real-time feedback and alarm feature allows for a proactive approach to mental health treatment. Prompt identification of depression symptoms can result in prompt treatments, resources for support, and individualised advice on getting professional assistance.

References

- [1] Musleh, D. A., Alkhales, T. A., Almakki, R. A., Alnajim, S. E., Almarshad, S. K., Alhasaniah, R. S., ... & Almuqhim, A. A. (2022). Twitter Arabic Sentiment Analysis to Detect Depression Using Machine Learning. *Computers, Materials & Continua*, 71(2).
- [2] Hasib, K. M., Islam, M. R., Sakib, S., Akbar, M. A., Razzak, I., & Alam, M. S. (2023). Depression Detection From Social Networks Data Based on Machine Learning and Deep Learning Techniques: An Interrogative Survey. *IEEE Transactions on Computational Social Systems*.

- [3] Li, X., La, R., Wang, Y., Hu, B., & Zhang, X. (2020). A deep learning approach for mild depression recognition based on functional connectivity using electroencephalography. *Frontiers in neuroscience*, 14, 192.
- [4] Liu, D., Feng, X. L., Ahmed, F., Shahid, M., & Guo, J. (2022). Detecting and measuring depression on social media using a machine learning approach: systematic review. *JMIR Mental Health*, 9(3), e27244.
- [5] Ashraf, A., Gunawan, T. S., Riza, B. S., Haryanto, E. V., & Janin, Z. (2020). On the review of image and video-based depression detection using machine learning. *Indonesian Journal of Electrical Engineering and Computer Science (IJEECS)*, 19(3), 1677-1684.
- [6] Amanat, A., Rizwan, M., Javed, A. R., Abdelhaq, M., Alsaqour, R., Pandya, S., & Uddin, M. (2022). Deep learning for depression detection from textual data. *Electronics*, 11(5), 676.
- [7] Sajja, G. S. (2021). Machine learning-based detection of depression and anxiety. *Int. J. Computer Appl*, 183, 20-23.
- [8] Vasha, Z. N., Sharma, B., Esha, I. J., Al Nahian, J., & Polin, J. A. (2023). Depression detection in social media comments data using machine learning algorithms. *Bulletin of Electrical Engineering and Informatics*, 12(2), 987-996.
- [9] Lora, S. K., Sakib, N., Antora, S. A., & Jahan, N. (2020). A comparative study to detect emotions from tweets analyzing machine learning and deep learning techniques. *International Journal of Applied Information Systems*, 12(30), 6-12.
- [10] Aleem, S., Huda, N. U., Amin, R., Khalid, S., Alshamrani, S. S., & Alshehri, A. (2022). Machine learning algorithms for depression: diagnosis, insights, and research directions. *Electronics*, 11(7), 1111.
- [11] Flores, R., Tlachac, M. L., Toto, E., & Rundensteiner, E. (2022, December). AudiFace: Multimodal Deep Learning for Depression Screening. In *Machine Learning for Healthcare Conference* (pp. 609-630). PMLR.
- This includes, but is not limited to, resizing, orienting, and color corrections. Image preprocessing may also decrease model training time and increase model inference speed. If input images are particularly large, reducing the size of these images will dramatically improve model training time without significantly reducing model performance.
 - Every image is made of pixels. And each pixel will have some intensity. Based on the intensity we can say if it is a white pixel or black pixel or something in between them. A histogram of an image is the representation of the intensity vs the number of pixels with that intensity. For example, a dark image will have many pixels which are black and few which are white. Representing that like a graph is what is called a histogram.
 - CT images are usually obtained as separate scans and may exhibit misalignment due to patient mobility. Various artifacts, including motion artifacts in PET and metal artifacts in CT scans, may influence these images. The image registration is used to guarantee spatial alignment between two images. The author applied the SyN function to handle the deformations and improve the image alignment. The SyN function integrates PET metabolic activity with CT anatomical data. It employs a forward transformation for generating the images. A similarity metric measures the similarity between the source and target image.