

Detecting Human Emotion by Text Classification

U Brunda¹, Palakuru Akhilesh², Dr.K.Kalaiselvi³

Final year, SRM Institute Of Science and Technology, Kattankulathur-603203, Chennai, Tamil Nadu, India
Assistant Professor, SRM Institute Of Science and Technology, Kattankulathur-603203, Chennai, Tamil Nadu, India
Email:bu8486@srmist.edu.in, pa4393@srmist.edu.in, mkkalai1981@gmail.com

Abstract— Nowadays, it's fairly usual to share moments on social media. By communicating thoughts, ideas, and enjoyable experiences over text, we can express our feelings without needing a lot of words. To investigate people's opinions, sentiments, and emotions, for instance, businesses may target YouTube as an abundant source of data. A greater comprehension of an author's emotions is often possible through emotion analysis. Analyzing expressions as positive, negative, or neutral has been the focus of almost all projects evaluating Telugu social media. We'll categorize the expressions in this essay into groups based on the emotions of happiness, fury, fear, disgust, and melancholy. Different approaches have been used in the case of other languages to automatically recognize textual emotions, however few of them were based on deep learning. Now let's talk about the system we utilized to classify the feelings stated in Telugu YouTube comments. For tasks requiring phrase classification, our model includes an XLM-RoBERTa and Multilingual BERT that was specifically trained on our dataset using trained word vectors. We contrasted the outcomes of our method with those of other machine learning techniques. The architecture of our deep learning technique is a word-based, end-to-end network, phrase, and document vectorization procedures. The proposed deep learning strategy was tested using the Telugu YouTube comments dataset, and the results were promising when compared to more traditional machine learning methods.

I. INTRODUCTION

As social media has become more popular, internet users can now voice their opinions on a wide range of subjects. Social networking sites are increasingly being used for a variety of activities, such as the advertising of products, the sharing of news, and the recognition of achievements.

Emotion analysis, often known as opinion mining, is the study of how to infer from textual data how individuals feel about a particular thing, person, or organization.

Market analysis, e-commerce, social media monitoring, and many more areas are examples of contemporary applications for emotion analysis. Telugu is the fifteenth most frequently spoken language in the world, with more than 75 million native speakers. The creation of a technique for Telugu text emotion analysis will benefit several people and organizations.

Everyday life brings us into contact with a variety of events, which leads to the formation of opinions regarding those occurrences. A person's emotions are strong feelings they have in reaction to their circumstances or interpersonal relationships. It has a big impact on consumer decision-making in a lot of different areas, such e-commerce, restaurants, movies, interests, and satisfaction with a service or a product. Additionally, it affects our health! Users can now voice their opinions about a comment, picture, or event using Facebook's replies feature, which has just undergone some changes. These reactions include angry, happy, love, and surprise.

Examples for emotion analysis as given below:

1. నేనుపుట్టినసంవత్సరాలకేమానాన్నచనిపోయాడు...ఎలాఉంటాడోకూడాతెలీదు
Nenu puttina sanvatsaralake ma nanna chanipoyadu.. Ela untadho kuda teledu – SAD
2. స్టోరీచండాలంగాడంది
Story chandalanga undi - ANGRY
3. ఈచిత్రం 200 గ్రాస్నుదాటవచ్చు
Ee chitram 200 gross nu datavacchu – TRUST
4. సినిమా విఎఫ్ఎక్స్ చూసినేనుఆశ్చర్యపోయాను
Cinema vfx chusi nenu ascharyapoya – SURPRISE
5. చాలాట్రీప్లీలుఉన్నందుననేనుసినిమానిఆత్రంగాచూడాలనుకుంటున్నాను
Chala twist lu unnanduna nenu cinemani aatranga chudalanukuntunnnanu – ANXIETY

In academic circles, emotional analysis is seen as a kind of higher, more developed version of sentiment analysis. Sentiment analysis is used to classify texts (posts, words, or documents) as neutral, positive, or negative. Emotional analysis, on the other hand, is a more extensive and in-depth investigation of user emotions with the goal of examining the psychology of various user behaviors and illuminating deeper human emotional meanings including anger, disgust, trust, grief, delight, and surprise.

The English language has a good reputation in the field of emotion detection, including the accessibility of datasets and dictionaries, in contrast to Telugu, which has a dearth of resources.

In this study, We look into automatic emotion recognition for Telugu language using MULTI LINGUAL BERT using four steps: word, sentence, document vectorization, and classification. Displaying the performance and precision that deep learning has so far attained, we also compared this methodology to other machine learning techniques. We applied our techniques to analyze user sentiment in the YouTube comments dataset.

II. RELATED WORK

A. NAILA ASLAM(2022) *Sentiment Analysis and Emotion Detection on Cryptocurrency Related Tweets Using Ensemble LSTM-GRU Model*

The methodological approach that they took. Random Forest, Decision Trees, KNN, and SVM, as well as the Metrics That They Employed in Order to Obtain the Output Precision, Accuracy, and the F1 score Remember, one of the problems with the model is that the balancing of the dataset through the use of the random under sampling shows that performance is diminished due to there being fewer training data.

B. A Majeed (2022)*Emotion Detection in Roman Urdu Text Using Machine Learning*

This research develops detection of human for Roman Urdu sentences with a specific size of dataset and they are mapped to six different classes of emotions. They used methods like KNN, Decision tree, SVM, and Random Forest. The final result shows KNN as the best model with better F-measure score compared to the other approaches.

C. T. Balomenos(2022) *Emotion Analysis in Man-Machine Interaction Systems*

This paper is developed for extracting the emotions from related image sequences. This uses an advanced intelligent rule-based system. It helps the MMI to deal with specific emotion states such as frustration and anger.

D. Abdullah (2022) *Multimodal Emotion Recognition using Deep Learning Smsa*

This paper is a review of emotional recognition of multimodal signals and unimodal solutions as they have higher accuracy. This improves better understanding of physiological signals and emotional awareness.

E. Omkar Gokhale, Shantanu Patanka, Onkar Litake, Aditya Mandke, Dipali Kadam(2022)

Emotion analysis in Tamil

It is the overview of shared task of emotion analysis in Tamil. The task is split into two in which one of the task includes social media Tamil comments that are annotated with specific suited emotions and in the other task fine-grained emotions are annotated for the social media comments in Tamil. The metrics used for evaluating models are Precision, Recall and Micro average.

F. LINJIAN LI(2021) A Novel Emotion Lexicon for Chinese Emotional Expression Analysis on Weibo: Using Grounded Theory and Semi-Automatic Methods

The downsides for the model include the methodology that was utilized (ALO and SC-LIWC), the metrics that were used to generate the output (Precision, recall, and F1), and the methodology that was used (ALO). Only users from China's Weibo were surveyed for this data set.

The strength of the relationship that each word has with the corresponding emotion category was not included in the lexicon.

G. Chang Liu, Taiao Liu, Shuojue Yang, And Yajun Du(2021)

Individual Emoticon Recognition Approach Combined Gated Recurrent Unit with Emotion Distribution Model.

This paper proposes a model called semantic emoticon emotion recognition (SEER). First the input text is divided into four categories with emotion dictionary and emoticons. Then it is combined by a bidirectional gated recurrent unit (Bi-GRU) a network with an emotion-vector-capturing attention mechanism. Lastly, a emoticon distribution model is constructed to obtain emotion vectors from various social network data. Fourth, we combine the emoticon emotion characteristics in text with the texts semantic emotional components using various fusion weights based on the various types of input short messages. Depending on the resulting emotion vector, we finally divide the short text emotions into six categories.

H. BHARATHI RAJA CHAKRAVARTHI (2021)Dataset for identification of homophobia and transophobia in Multilingual youtube comments.

This paper describes the process of building the dataset, qualitative analysis of data, and inter-annotator agreement. In addition, we create baseline models for the dataset.

I. FERDOUS AHMED(2020) Emotion Recognition from Body Movement

The methodologies that were applied were SVM,LDA,GNV,DT, and KNN. The metrics that were applied for the purpose of obtaining the output were f-Score,p-Score, and Accuracy. The limitations of the model are as follows: Observed a marginal drop in performance across the board in action-independent cases

J. ZISHAN AHMAD, RAGHAV JINDAL, ASIF EKBAL and PUSHPAK BHATTACHHARYY (2020)

Borrow from rich cousin: transfer learning for emotion detection using cross lingual embedding. Expert Systems with Applications.

This paper is mapped with the emotions of disaster domain sentences in Hindi. Dataset is created for the disaster domain sentences. The models used here are CNN and Bi-LSTM (Bi-Directional Long Short Term Memory). For Hindi emotion categorization, the neural networks are trained on the available datasets, and then the weights are adjusted using one of four transfer learning techniques.

K. Zhenzhong Lan, Mingda Chen, Piyush Sharma, and Rady Soricut (2019)

ALBERT: A Lite BERT for Self-supervised learning of language representations.

For BERT to use less memory and train more quickly, we provide two parameter-reduction strategies. Detailed empirical data demonstrates that our suggested methods produce models that scale far better than the original BERT. We also employ a self-supervised loss that emphasises modelling inter-sentence coherence, and we demonstrate that it consistently facilitates tasks that require multi-sentence inputs.

L. Alexis Connaeu, Kartikay Khandelwal, Naman Goyal,Vishrav Chaudhary (2019)

Unsupervised cross-lingual representation learning at scale.

Using more than two terabytes of filtered CommonCrawl data, we train a Transformer-based masked language model on 100 different languages. On a number of cross-lingual benchmarks, including +14.6% average accuracy on XNLI, +13% average F1 score on MLQA, and +2.4% F1 score on NER, our model, called XLM-R, greatly surpasses multilingual BERT (mBERT). Low-resource languages are where XLM-R excels, with Swahili's XNLI accuracy increasing by 15.7% and Urdu's by 11.4% over earlier XLM models. The trade-offs between (1) positive transfer and capacity dilution and (2) the performance of high and low resource languages at scale are among the important aspects that must be considered in order to accomplish these advantages, and we also give a thorough empirical study of these factors.

M. Stephen Merity, Nitish Shirish Keskar, And Richard Socher. (2019)

An analysis of neural language modeling at multiple scales.

We provide a model architecture and training method that, when applied to the WikiText-103 data set, achieves state-of-the-art performance while being significantly quicker than an NVIDIA cuDNN LSTM-based model by a factor of two. using the Quasi-Recurrent Neural Network (QRNN), an Longer sequences within batches and softmax with weight tying.

N. Jeremyhoward And Sebastian Ruder (2019)

Universal language model fine-tuning for text classification.

We describe strategies that are essential for fine-tuning a language model and propose Universal Language Model Fine-tuning (ULMFiT), a powerful transfer learning method that may be used for any NLP application. On six text classification tasks, our approach greatly exceeds the state-of-the-art, lowering the error on most datasets by 18–24%. Furthermore, it matches the performance of training from scratch on 100x more data with only 100 labeled instances. Our pre-trained models and code.

O. VINAY KUMAR JAIN, SHISHIR KUMAR, and STEVEN LAWRENCE FERNANDES (2019)

Extraction of emotions from multilingual text using intelligent text processing and computational linguistics.

Every word of emotion in a tweet is significant in decision-making, hence an efficient pre-processing technique has been utilized to maintain the significance of multilingual emotional words. The Naive Bayes algorithm and Support Vector Machine (SVM) are used to classify tweets' sentiments in exquisite detail.

III. PROPOSED METHODOLOGY

Using a masked language modeling (MLM) aim, a model was trained on the top 104 languages with the most content on Wikipedia. Maximum accuracy was thus possible. To make sure the outcomes were trustworthy, this was done. The inaugural release took place in the repository you are currently seeing, and this page acted as its introduction. This model can distinguish between distinct varieties of English because it is sensitive to capitalization, a feature that both varieties of English share.

The team in charge of Hugging Face has been in charge of writing this model card because the team in responsibility of releasing BERT was not accountable for doing so. This is so that Hugging Face can be released, which is the responsibility of the team in charge of it.

The XLM-RoBERTa of BERT model is a transformers model that has been pre trained on a big corpus of data that includes information in a variety of languages using an unsupervised learning approach. The BERT model has been trained using this data. The model in question has been designated as the XLM-RoBERTa model. This suggests that an automatic method was used to produce inputs and labels from those texts, and that it was pre trained exclusively on the raw texts, with no human tagging of them in any way (which explains why it can use a significant amount of data that is publicly available). This also suggests that the raw texts were not in any manner labeled by humans before it was pre-trained. Additionally, this implies that there were no human annotations of any kind on any of the raw texts. This also shows that the raw texts were not in any manner, shape, or form labeled by humans. This can be inferred from the absence of human labeling. To be more precise, it underwent pre-training with the idea of achieving the following objectives:

The first step in masked language modeling (MLM) is to take a sentence as input, randomly mask 15% of the words, then run the entire masked sentence through the model and ask it to predict the words that have been hidden. This procedure is carried out again and again until the model correctly predicts the words that have been obscured. This process is repeated multiple times until the model can correctly anticipate the words that are hidden. Multiple iterations of this approach are carried out until the model is able to correctly anticipate the words that have been hidden. Unlike traditional recurrent neural networks (RNNs), which normally view the words one after another, and autoregressive models like GPT, which internally mask the future tokens, this approach does not view the words in a sequential sequence. This approach instead presents the words in the order that a reader would use if they were reading them out to themselves. This approach examines the words in the same sequence as one would read them aloud to themselves. In other words, it adheres to the speech's organic flow. On the other hand, this approach takes into account the words in the presentation's order of arrangement. As a result, the model is able to obtain a complete representation of the text that takes into consideration both orientations.

While the model is being trained, the Next Sentence Prediction (NSP) method combines two masked sentences into a single input. This improves the model's capacity for learning. As a result, the learning efficiency of the model is increased. This is carried out in order to enable the model to produce predictions that are more accurate.

Although it is not always the case, it is possible that they will correlate to sentences that were written next to one another in the original text. Although not always the case. Do not, however, take this too literally because it is not always the case. But on the other side, they can decide not to follow through in the end. The model must then decide whether or not the two sentences were located in the text directly after one another. If you have access to a dataset with labeled sentences, for example, you may use the characteristics the BERT model produced as inputs to train a typical classifier.

The standard classifier will be able to learn from the tagged sentences as a result. As a result, the model is able to develop an internal representation of the languages that are included in the training set. This will give the standard classifier the chance to learn about the classification of sentences. Then, using this representation of the languages, features that are useful for later tasks in the process can be extracted.

A. Algorithm: Xlm-Roberta

1. Importing XLMRobertaTokenizer and XLMRobertaForSequenceClassification from transformers.
2. Model is named as xlm-roberta-base.
3. Input that is the tokenizer is given as XLMRobertaTokenizer.from_pretrained(MODEL_TYPE).
4. The module is downloaded of 100 percent.
5. Checking the size of the vocab.
6. Verifying whether the special tokens are present or not.
7. Model inputs are given such as
input_ids (type: torch tensor)
attention_mask (type: torch tensor)
labels (type: torch tensor)
8. The very first input is the 'input_ids'. These represents the sentences which also represent tokens.
9. The second is the 'token_type_ids'.
10. Third is the 'attention_mask'. It has the same length as of 'input_ids' and it also tells the model which tokens in the 'input_ids' are working and which are padding.
11. To indicate token or a special word '1' is used and for padding '0' is used.
12. Third input also consists of 'labels'.
13. A tokenizer is used to create XLM-RoBERTa input for both one and two input sentences.
14. The sequence of tokens are decoded.
15. The truncated tokens will return in a list called overflowing_tokens.
16. Data is loaded.
17. Creates folds according to the requirement for training and testing.
18. Displays the required folds.

IV. IMPLEMENTATION

We'll outline the data in this part that was used, as well as our methodologies, in order to recognise emotions in Telugu YouTube videos using a deep learning approach XLM-RoBERTa.

The three modules of the project for implementation are:

- a. Dataset creation
- b. Training dataset
- c. Testing dataset

A. Dataset Creation

The dataset of Telugu YouTube comments was provided and used for the training of the model, which is an ordinal classification task based on the intensity of feelings: You must classify a comment into one of five ordinal classes of intensity for the emotion represented by the letter E if it is offered together with an emotion that begins with that letter. One comment has been added to the dataset for each of the following emotions, for a total of one thousand comments includes rage, fear, joy, disgust, and sadness.

Our dataset was split into two sets: 500 comments made up the training set, and 100 comments made up the testing set. 90% of the dataset had to be used for training, while just 10% was necessary for testing. The test dataset was only used to test the created model and offer an indicator of how well the trained model is working. To train the model, the training datasets were classifier and to optimise the parameters. The model was not given access to the test dataset.



Fig2: Created dataset

Data Preprocessing is done. Because our dataset was in Telugu, we had to perform some specialized pre-processing in order to identify the most effective pattern i.e, training the dataset. The following are the steps that we followed:

- In addition, we have the option of including a step that gets rid of stop words in the input text. Stop words include things like prepositions, conjunctions, and other similar words.

The dataset is divided in two types as data used for training for 80 percent and 20 percent for testing the data. The trained dataset is tested with different algorithms like XLM-RoBERTa and Multilingual BERT. To test the dataset necessary python libraries for Colab code execution dataset as pandas data frame are imported. We have used seaborn's count plot to count various emotions. The task is to find the best machine learning algorithm with good accuracy.

897

```
[ ] import pandas as pd
from sklearn.model_selection import train_test_split

df = pd.read_excel('sample_data (1) 2.xlsx', names=['text', 'label'])

df.head()

[ ] import numpy as np
import seaborn as sns
import matplotlib.pyplot as plt
from collections import defaultdict

import torch
import torch.nn.functional as F
from torch import nn, optim
from torch.utils.data import Dataset, DataLoader
from sklearn.metrics import confusion_matrix, classification_report
```

V. RESULTS & DISCUSSION

This model is trained using XLM-RoBERTa algorithm and Multilingual BERT algorithm with around 600 Telugu sentences mapped with the emotions happy, neutral, disgust and anxiety. This model gives the accuracy of 77 percentage for XLM-RoBERTa algorithm and for the algorithm Multilingual BERT it gives 53 percentage.

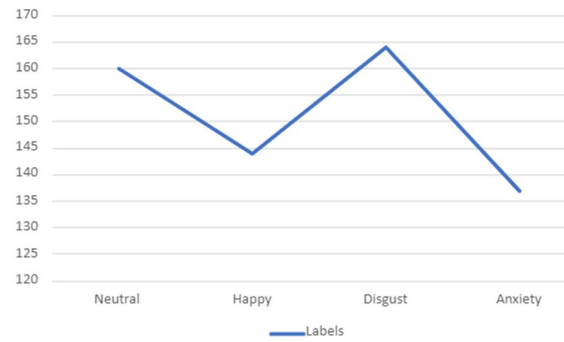


Fig 3: Graphical representation count of mapped dataset

	precision	recall	f1-score	support
anxiety	0.6500	0.7000	0.6700	70
disgust	0.7581	0.7059	0.7934	74
happy	0.7217	0.7636	0.8286	73
neutral	0.6886	0.6250	0.7085	64
accuracy	0.7789			161
macro avg	0.6896	0.6986	0.7801	161
weighted avg	0.6957	0.6959	0.7822	161

Fig 4: Accuracy of 77 Percentage using XLM-roberta algorithm

```
print(classification_report(y_test, y_pred, target_names=class_names))
```

	precision	recall	f1-score	support
anxiety	0.42	0.47	0.34	70
disgust	0.38	0.59	0.41	74
happy	0.32	0.38	0.39	73
neutral	0.56	0.76	0.66	64
accuracy	0.53			161
macro avg	0.42	0.55	0.45	161
weighted avg	0.45	0.52	0.46	161

Fig 5: Model with 53 percent accuracy for multilingual BERT algorithm

VI. CONCLUSION

This study attempted to classify comments made on social media. We applied the XLM-RoBERTa and Multilingual BERT strategies. With a macro-averaged f1 score of 0.77 for XLM-RoBERTa and for Multilingual BERT f1 score is of 0.53. XLM-RoBERTa method outscored all other models. Overall, the models are seen to identify emotions like anxiety, happiness, neutral, and disgust. The models are far less accurate in classifying more complex emotions like fear, rage, and melancholy. To enhance the performance of the models, alternative strategies, such as genetic algorithm-based ensembling, can be tested in the future.

REFERENCES

- [1] Emotion Analysis in Man-Machine Interaction Systems – T.Balomenos 2022 https://link.springer.com/chapter/10.1007/978-3-540-30568-2_27
- [2] Multimodal emotion recognition using deep learning SMSA -Abdullah 2022 https://scholar.google.com/scholar?cluster=11062434886599925582&hl=en&as_sdt=0,5
- [3] Emotion Analysis in Tamil - Omkar Gokhale, Shantanu Patanka, Onkar Litake, Aditya Mandke, Dipali Kadam 2022 https://scholar.google.com/scholar?cluster=11062434886599925582&hl=en&as_sdt=0,5
- [4] Emotion detection in roman Urdu text using machine learning - A Majeed 2022 https://dl.acm.org/doi/abs/10.1145/3417113.3423375?casa_token=fS0ijtmLfAIAAAAAA:a0hgZLiFwIAYnp3E5x5fvxZ9TAXGgYBBZ_XmDBI0xiY0NU1nfJvK5xXkwfMGPTguPBNBwrfb4GmjBQ
- [5] Sentiment Analysis and Emotion Detection on Cryptocurrency Related Tweets Using Ensemble LSTM-GRU Model - NAILA ASLAM, FURQAN RUSTAM, ERNESTO LEE, PATRICK BERNARD WASHINGTON, AND IMRAN ASHRAF 2022 <https://ieeexplore.ieee.org/abstract/document/9751065>
- [6] A Novel Emotion Lexicon for Chinese Emotional Expression Analysis on Weibo: Using Grounded Theory and Semi-Automatic Methods - LIANG XU, LINJIAN LI, ZEHUA JIANG, ZAOYI SUN, XIN WEN, JIAMING SHI, RUI SUN, AND XIUYING QIAN 2021 <https://ieeexplore.ieee.org/abstract/document/9139939>
- [7] Emotion Recognition from Body Movement - FERDOUS AHMED, A. S. M. HOSSAIN BARI, AND MARINA L. GAVRILOVA 2020 <https://ieeexplore.ieee.org/abstract/document/8945309>
- [8] Individual Emotion Recognition Approach Combined Gated Recurrent Unit with Emoticon Distribution Model - CHANG LIU,TAIAO LIU, SHUOJUE YANG,AND YAJUN DU 2021 <https://ieeexplore.ieee.org/abstract/document/9597507>
- [9] Borrow from rich cousin: transfer learning for emotion detection using cross lingual embedding. Expert Systems with Applications - Zishan Ahmad, Raghav Jindal, Asif Ekbal, and Pushpak Bhattacharyya. 2020. https://www.sciencedirect.com/science/article/pii/S0957417419305536?casa_token=kEwcg0YM4DIAAAAAA:Gxkl7Wj7_hbtk7vdEPDEwzd7eqgnW_-4xRCL5c8PxV0GRuLYhpHcieOkW895-482sC5rtYWEyiOO
- [10] Dataset for identification of homophobia and transphobia in Multilingual youtube comments. - Bharathi Raja Chakravarthi, Ruba Priyadharshini, Rahul Ponnusamy, Prasanna Kumar Kumaresan, Kayalvizhi Sampath, Durairaj Thenmozhi, Sathiyaraj Thangasamy, Rajendran Nallathambi, and John Phillip McCrae 2021 <https://arxiv.org/abs/2109.00227>
- [11] ALBERT: A Lite BERT for Self-supervised learning of language representations. - Zhenzhong Lan, Mingda Chen, Sebastian Goodman, Kevin Gimpel, Piyush Sharma, and Radu Soricut 2020 <https://arxiv.org/abs/2003.10592v1>
- [12] Unsupervised cross-lingual representation learning at scale. Alexis Conneau, Kartikay Khandelwal, Naman Goyal, Vishrav Chaudhary, Guillaume Wenzek, Francisco Guzmán, Edouard Grave, Myle Ott, Luke Zettlemoyer, and Veselin Stoyanov. 2019. <https://arxiv.org/abs/1911.02116>
- [13] An analysis of neural language modeling at multiple scales. - Stephen Merity, Nitish Shirish Keskar, and Richard Socher. 2019. <https://arxiv.org/abs/1803.08240>
- [14] Universal language model fine-tuning for text classification. - Jeremy Howard and Sebastian Ruder. 2019. <https://arxiv.org/abs/1801.06146>
- [15] Extraction of emotions from multilingual text using intelligent text processing and computational linguistics. - Vinay Kumar Jain, Shishir Kumar, and Steven Lawrence Fernandes. 2019. https://www.sciencedirect.com/science/article/pii/S1877750317301035?casa_token=xGYROGWQ9aIAAAAAA:O-oKcOAp1QWNbdRDt6pGb-XK8bus-y1sWsk2SCdQGMdYxUs6ycXsePgqUht0qlwpwtDgwxuUzJZ