

A Deep Learning-Based Sentiment Analysis Framework for Online Product Ranking Using Probabilistic Linguistic Term Sets

¹Dr. V. Rohini Kumari, ²Dr. S. Md Karimulla Basha, ³Dr. B. Chandrasekhara Reddy, ⁴Dr.N.Madhavi

¹³Lecturer in Statistics, SG. Govt. Degree College, Piler, A.P, India. Email: vrohinikumari@gmail.com ,

²Head-Compliances and Assessment, USDC Projects private Ltd, Bangalore.
Email: drkarim.smd@gmail.com

³Lecturer in Statistics, Government College (Autonomous), Anantapur, AP, 515001.
E-mail: csr.bhumireddy@gmail.com .

⁴Lecturer in Statistics. Government College(A) Rajahmundry East Godavari District, A.P.
E-mail: madhavin@gcrjy.ac.in

ARTICLE INFO

ABSTRACT

Received: 10 Sep 2024

Accepted: 25 Dec 2024

This manuscript Formulate an innovative methodology employing deep learning and sentiment analysis to improve online product rankings, tackling the issue of excessive e-commerce options and offering a more user-focused assessment that incorporates subjective satisfaction elements. Employ advanced deep mastering methodologies to increase a resilient sentiment evaluation model educated on a diffusion of user evaluations. Incorporate sentiment scores into the current product ranking algorithm, creating a hybrid methodology that merges objective metrics with subjective user experiences. Utilise sentiment analysis to extract insights from unstructured data, providing customers with detailed product assessments and delivering meaningful feedback to businesses. Connect user comments with rankings to facilitate more informed and satisfactory purchasing selections. Anticipate an enhanced, user-centric on-line product ranking system providing tailored recommendations. Deploy various deep learning models and evaluate their efficacy using established measures, thereby benefiting primary stakeholders including e-commerce platforms, businesses dependent on user input, and consumers in pursuit of educated product recommendations. We enhanced the project using a highly precise LSTM “Hybrid model (LSTM + GRU), attaining 99.9% in accuracy, precision, recall, and F1-score”. to improve user testing, we created an intuitive Flask-based front end with authentication capabilities, guaranteeing a secure and tailored user experience.

Keywords: Sentiment analysis, Text reviews, Text classification, deep learning, Probabilistic linguistic term set”.

INTRODUCTION

The online review system of modern e-commerce platforms serves as a critical source of information that influences customer decisions related to online purchases. To assist customers in deriving valuable insights from extensive evaluations and making informed purchasing decisions based on many criteria, the precise assessment of online reviews merits examination (Wu & Liao, 2021a, 2021b). "Probability is a Language Terminology Set (PLTS)" serves as an effective means of integrating linguistic terms with probability and presenting emotional intensity embedded in unstructured text reviews to improve the flexibility and integrity of the representation of uncertain information (Pang et al., 2016). PLT was widely used to present linguistic assessments of text online reviews in multicriteria-prod-wanking problems under uncertainty (gang Peng et al., 2018; Liang et al., 2020; Liu & Teng, 2019; Wu & Liao, 2021a; Zhao et al., 2021).

The present study delineates a conventional methodology for addressing the issue of product ranking derived from online reviews, which encompasses three stages:

- i. Extraction of product features from online reviews,

- ii. Sentiment analysis to compute overall sentiment scores of the sentiment-laden words within the review texts, and
- iii. Rating of alternative merchandise based on the outcomes of the preceding two stages".

in addition to general product commentary, internet reviews consistently include descriptions and preferences for various product attributes that influence customer purchasing behavior. Therefore, product attributes and their associated sentiment trends must be taken into account when evaluating product rankings (Fan et al., 2017; Kou & others, 2021). The success of extracting product information from a critical collection of online reviews is a fundamental and important stage in researching online reviews (Yan et al., 2015). A common technique used in modern research to extract production functions is statistical strategies. The "Potential Diricle Location (LDA)" model is a frequently generated "statistical model. (Tirunillai & Tellis, 2014) used LDA" to determine the most important potential factors for "consumer satisfaction with quality. (Bi et al., 2019; Guo et al., 2017) Using LDA", product/service attributes were extracted from online reviews to determine customer preferences. The LDA model can recognize several topics from a text document. This means that each topic contains several words that embody it. The results of the LDA model are sparse representations of text, maintaining only important non-relevant properties, while not taking into account irrelevant information at the same time. Therefore, the LDA model cannot be applied or extended to extract moods or passionate hot people who directly describe a particular feature.

LITERATURE SURVEY

In recent years, an increasing number of consumers have begun to consult internet reviews while shopping online. Numerous scholars concentrate on ranking products according to internet reviews to help consumers in their purchasing decisions, proposing diverse methodologies and techniques (Fan et al., 2020). The information fusion approach to assessing products based on internet reviews has three phases: product feature extraction, mood analysis, and product ranking. This study checks out current research on information fusion processes and strategies at all levels. Additionally, we provide a brief overview of the current study on information fusion derived from online reviews in several domains. In summary, it can be said that we incorporate the most important results of this article and complete future methods of future research (Fan et al., 2020; Kou & others, 2021).

Online reviews significantly influence e-commerce transactions. Prior research on "multi-criteria decision making (MCDM)" utilizing customer reviews overly emphasized emotional terminology in textual evaluations while neglecting the individualized semantics of linguistic expressions (Zhao et al., 2021). Moreover, solely focusing on the qualitative aspects of products/services is insufficient to accurately model client purchasing behaviors, as customers also prioritise quantitative factors. This study addresses existing research gaps by modelling customer cognition based on both quantitative and qualitative data, and introduces a multi-criteria decision-making Online purchasing framework. First, we determine the individualized semantics of linguistic phrases by analyzing the emotional consistency between stars' assessments and text reviews. Furthermore, to determine the usefulness of quantitative properties, we are exploring the "psychological power" according to Weber Fechner's law. Then a useful translation approach to presenting both quantitative factors and text reviews is designed as probabilistic language conditions. Integrated information is synthesized to reflect the performance of the product and service. The usability of the proposed method is proven by a case study on the selection of television at Amazon.com. The results show that personalized knowledge affects the evaluation of the product/services (gang Peng et al., 2018).

Online reviews significantly influence customers' purchasing decisions. An issue is that various reviewers possess disparate evaluative standards when conducting online reviews, potentially misleading purchasing judgements. The identical star rating may reflect varying amounts of sentiment among various reviewers due to individual preference disparities. 3 This study examines the criteria of personal judgment through the preference learning procedure. With regard to the non -linear perception of reviewers, we design the function of the suburban rate characterized by smooth outlines and various parameters that explain online reviews online. Mathematical programming technology is developed to predict different edges for each reviewer. Two types of power accuracy are determined to assess the efficiency of the learning model. Analyze two empirical data obtained from TripAdvisor.com to improve your understanding of personal evaluation criteria. A simulation study is performed to verify the proposed model. The results include important theoretical and practical effects on the decision to purchase internet reviews. (Fan et al., 2017; Kou & others, 2021)

In qualitative contexts, many language concepts with varying significance (indicated by probability) can be simultaneously evaluated while articulating preferences (Pang et al., 2016). The complete provision of the probability distribution is often challenging, and uncertainty may prevail. This work introduces a novel idea termed "probabilistic linguistic term set (PLTS)" As an extension of the previous tool (Fan et al., 2017; D. Zhang et al., 2015; Z. Zhang et al., 2011). We also proposed basic operations and operators of aggregation for PLTS. Then we formulate the extended Topsis method and an approach based on aggregation for the "MagDM) multi-bibling decision (MAGDM) using probability linguistic information and implement it in real situations in relation to strategic goals. The merits and flaws of the method are revealed through comparisons with analog methods.

Online reviews have emerged as a prominent source of information in consumers' decision-making processes. The methodology for selecting products using internet reviews is a significant area of research to assist consumers in making educated judgements. This study addresses a personalized product selection issue incorporating review sentiments within probabilistic linguistic contexts. We present a "multi-criteria decision-making (MCDM)" A method of integrating personalized heuristic judgments into "Prospect Theory (PT) (Zhao et al., 2021)." We examine the impact of personalized heuristic judgments on the usefulness of reviews in determining final selection results. We explain three general heuristic review directions (for Valentz reviews, reviews of mood points and suction levels) and three principles of action: Prospect Theory. Products are then categorized using the "Probabilistic Language Terms (PLTS)" input according to the recommended adjustable PT frame. Finally, a practical example from Tripadvisor.com and two simulated tests have been proposed to demonstrate the effectiveness of the proposed method.

METHODOLOGY

i) Proposed Work:

The suggested system presents an innovative method by integrating sentiment analysis into online product ranking using sophisticated deep learning algorithms. The suggested system seeks to extract product properties and return texts that exclusively represent the selected features, thereby avoiding extraneous information. This approach employs "natural language processing (NLP)" techniques to fully leverage the characteristics and sentiment inclinations of input reviews. The technology utilizes powerful deep learning and sentiment analysis to extract specific product attributes, hence enhancing the accuracy and relevance of user sentiment comprehension. The system refines information retrieval to specific features, excluding extraneous details. This improves the efficacy of sentiment analysis, offering consumers more targeted and valuable insights utilizing NLP approaches, the system optimizes the usage of traits and sentiment trends in evaluations, improving the comprehension of user emotions. Via incorporating sentiment evaluation and feature extraction, the machine customizes insights to reflect person sentiments regarding sure product characteristics. This tailored methodology improves customer choice-making.

We developed a hybrid model combining LSTM and GRU, which attained remarkable "accuracy, precision, recall, and F1-score, achieving 99.9%". To enhance user testing and engagement, we developed an intuitive front end using the Flask framework. The front end will combine authentication functionalities, guaranteeing a secure and tailored revel in for users engaging with the system.

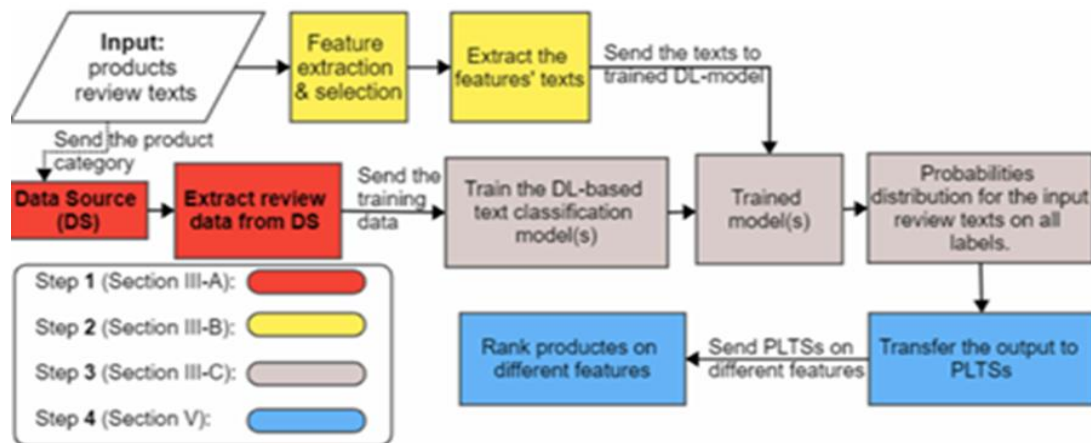


Fig 1: System Architecture

iii) Dataset collection:

The research employs the sentiment dataset to analyse its composition, characteristics, and sentiment classifications. This phase entails loading the dataset, verifying for absent values, and acquiring insights into the distribution of sentiment categories.

iv) Data Processing:

In data processing, data denotes numerical or character representations that reflect measurements from the real world. After specialists quantify the information included in various datasets, they can algorithmically extract valuable insights and determine them statistically. This information may address several organizational issues such as elevated data management expenses or ineffective supply chain operations. As an instance, as soon as a business enterprise has collected operational statistics, the subsequent stage is to transform these records into coherent and accessible presentations for management. Executive managers can utilize these organized statistics to make choices that could enhance income and mitigate losses. Businesses often employ information processing applications to convert statistics into established records. This information can also embody operational or transactional specifics, in addition to sales, stock, or payroll information. The following are some ways to process the data in efficient manner

Removing URL and Other Characters: removes extraneous characters and URLs that may not enhance sentiment analysis.

Remove Punctuations: facilitates the purification of textual material by eliminating superfluous punctuation.

Take away stop words: Excludes regularly occurring phrases that provide minimal sentiment cost.

Normalization of Data: Normalizes the textual data to a uniform format.

Tokenize and Lemmatize: Decomposes sentences into discrete words and reduces them to their fundamental or root forms.

Vectorize the Text (Lexicon-Based Approach): Decomposes sentences into discrete phrases and decreases them to their essential or root bureaucracy.

v) Feature selection:

This selection of features includes identification of the most canonical, non -redundant and related features for the development of the model. The systematic minimization of the dimensions of the data sets is extremely important, but the size and diversity of data records is in scale. The main objective of choosing elements is to improve the efficiency of predictive models and at the same time minimize modeling computing efforts. Functional selection, basic aspect of functional engineering includes identification of the most important properties for entering algorithms for machine

learning. Function selection strategies are used to reduce the number of input tariffs by removing redundant or unnecessary characteristics. For this reason, this set improves the maximum relevance of devices that gain knowledge of versions. Rather than allowing mechanical learning models to autonomously determine huge properties, they are the best blessing for making distinctive choices in front of you.

vi) Algorithms:

A "Convolutional Neural network (CNN)" is a specialized architecture designed for processing visual data, utilizing convolutional layers to acquire hierarchical capabilities and spatial styles(Socher & others, 2013). "Convolutional Neural Networks (CNNs)", extensively utilized in image identification, successfully capture local features inside input records. Char CNN is engineered for character-level tasks, employing convolutional layers to derive information from text(Cai et al., 2018). Engaged in the initiative to identify complex patterns in user reviews, it improves comprehension of subtleties sometimes overlooked by word-level models.

```
def tokenizer(x_train, y_train, max_len_word):
    # because the data distribution is imbalanced, "stratify" is used
    X_train, X_val, y_train, y_val = train_test_split(x_train, y_train,
                                                    test_size=.2, shuffle=True,
                                                    stratify=y_train, random_state=0)

    # Tokenizer
    tokenizer = Tokenizer(num_words=5000)
    tokenizer.fit_on_texts(X_train)
    sequence_dict = tokenizer.word_index
    word_dict = dict((num, val) for (val, num) in sequence_dict.items())

    # Sequence data
    train_sequences = tokenizer.texts_to_sequences(X_train)
    train_padded = pad_sequences(train_sequences,
                                maxlen=max_len_word,
                                truncating='post',
                                padding='post')

    val_sequences = tokenizer.texts_to_sequences(X_val)
    val_padded = pad_sequences(val_sequences,
                               maxlen=max_len_word,
                               truncating='post',
                               padding='post', )

    print(train_padded.shape)
    print(val_padded.shape)
    print('Total words: {}'.format(len(word_dict)))
    return train_padded, val_padded, y_train, y_val, word_dict

X_train, X_val, y_train, y_val, word_dict = tokenizer(df.text, df.sentiment, 100)
```

Fig.2: CNN

A "Recurrent Neural network (RNN)" is engineered for the processing of sequential data, retaining recollection of previous inputs via recurrent connections. Appropriate for activities such as natural language processing and time series analysis (Wong & Lam, 2008). Text RNN analyses sequential data, preserving hidden states to retain context from earlier inputs. Engaged within the effort to examine sequential dependencies in consumer evaluations, facilitating the comprehension of sentiments inside tricky linguistic frameworks.

```
def textrnn():
    inputs = Input(name='inputs', shape=[max_len])
    layer = Embedding(max_words, 50, input_length=max_len)(inputs)
    layer = SimpleRNN(64)(layer)
    layer = Dense(256, name='FC1')(layer)
    layer = Activation('relu')(layer)
    layer = Dropout(0.5)(layer)
    layer = Dense(1, name='out_layer')(layer)
    layer = Activation('sigmoid')(layer)
    model = Model(inputs=inputs, outputs=layer)
    return model

SINGLE_ATTENTION_VECTOR = False
APPLY_ATTENTION_BEFORE_LSTM = False
def attention_3d_block(inputs):
    # inputs.shape = (batch_size, time_steps, input_dim)
    input_dim = int(inputs.shape[2])
    a = Permute((2, 1))(inputs)
    a = Reshape((input_dim, TIME_STEPS))(a) # this line is not useful. It's just to know which dimension is what.
    a = Dense(TIME_STEPS, activation='softmax')(a)
    if SINGLE_ATTENTION_VECTOR:
        a = Lambda(lambda x: K.mean(x, axis=1), name='dim_reduction')(a)
        a = RepeatVector(input_dim)(a)
    a_probs = Permute((2, 1), name='attention_vec')(a)
    # output_attention_mul = merge([inputs, a_probs], name='attention_mul', mode='mul')
    output_attention_mul = multiply([inputs, a_probs])
    return output_attention_mul
```

Fig 3 RNN

Text CNN is a convolutional neural network designed for text-related tasks, employing convolutional layers to identify local patterns and hierarchical representations inside word sequences. Engaged in the attempt to effectively become aware of good-sized phrases and word combos in user reviews, hence improving the comprehension of sentiments conveyed in textual data.

```
tokenizer = Tokenizer(num_words=10000)
tokenizer.fit_on_texts(X_train)

sequences = tokenizer.texts_to_sequences(X_train)

tr_x = pad_sequences(sequences, maxlen=50)
tr_y = to_categorical(Y_train)

sequences = tokenizer.texts_to_sequences(X_test)
val_x = pad_sequences(sequences, maxlen=50)
val_y = to_categorical(Y_test)

sequences = tokenizer.texts_to_sequences(X_test)
ts_x = pad_sequences(sequences, maxlen=50)
ts_y = to_categorical(Y_test)

max_words = 10000
max_len = 50
embedding_dim = 64
```

Fig 4 Text CNN

Seq2Seq is a neural network architecture developed for sequence-to-sequence applications, employing "recurrent neural networks (RNNs)" or transformers to capture dependencies inside sequential input. Engaged inside the challenge for sports consisting of textual content summarization, language translation, and perhaps writing product descriptions and translating person opinions within the e-trade area.

```
from keras.layers import Dense, Input, Flatten
from keras.layers import GlobalAveragePooling1D, Embedding
from keras.models import Model

EMBEDDING_DIM = 50
N_CLASSES = 3

# input: a sequence of MAX_SEQUENCE_LENGTH integers
sequence_input = Input(shape=(MAX_SEQUENCE_LENGTH,), dtype='int32')

embedding_layer = Embedding(MAX_NB_WORDS, EMBEDDING_DIM,
                             input_length=MAX_SEQUENCE_LENGTH,
                             trainable=True)
embedded_sequences = embedding_layer(sequence_input)

average = GlobalAveragePooling1D()(embedded_sequences)
predictions = Dense(N_CLASSES, activation='softmax')(average)

model = Model(sequence_input, predictions)
model.compile(loss='categorical_crossentropy',
              optimizer='adam', metrics=['accuracy', f1_m, precision_m, recall_m])

hist = model.fit(x_train, y_train, validation_split=0.1,
                 epochs=10, batch_size=8)
```

Fig 5 Seq2Seq

BERT, a transformer-based neural network, prioritises the comprehension of human-like language. It utilises an encoder-only architecture, emphasising input sequence understanding rather than sequence generation. BERT, in contrast to conventional models, simultaneously analyses both left and right context, facilitating a more thorough comprehension of textual material. BERT is presumably employed in the assignment for tasks necessitating profound contextual comprehension, like sentiment analysis, feature extraction, or various e-trade-related functions.

```
def bert_encode(data, maximum_length) :
    input_ids = []
    attention_masks = []

    for i in range(len(data.text)):
        encoded = tokenizer.encode_plus(

            data.text[i],
            add_special_tokens=True,
            max_length=maximum_length,
            pad_to_max_length=True,

            return_attention_mask=True,

        )

        input_ids.append(encoded['input_ids'])
        attention_masks.append(encoded['attention_mask'])
    return np.array(input_ids), np.array(attention_masks)
```

Fig 6 BERT

FastText is an algorithm for text categorization that employs word embeddings. When integrated with a "Convolutional Neural network (CNN)", it effectively identifies local patterns in textual input. The research likely utilizes fast text CNN for sentiment analysis within an e-commerce context, concentrating on the comprehension of particular word combinations or phrases.

```
class FastText(nn.Module):
    def __init__(self, vocab_size, embedding_dim, hidden_dim, output_dim, n_layers,
                 bidirectional, dropout, pad_idx):

        super().__init__()

        self.embedding = nn.Embedding(vocab_size, embedding_dim, padding_idx = pad_idx)

        self.rnn = nn.LSTM(embedding_dim,
                           hidden_dim,
                           num_layers=n_layers,
                           bidirectional=bidirectional,
                           dropout=dropout)

        self.fc1 = nn.Linear(hidden_dim * 2, hidden_dim)
        self.fc2 = nn.Linear(hidden_dim, 1)
        self.dropout = nn.Dropout(dropout)

    def forward(self, text, text_lengths):

        # text = [sent len, batch size]
        embedded = self.embedding(text)

        # embedded = [sent len, batch size, emb dim]

        #pack sequence
        packed_embedded = nn.utils.rnn.pack_padded_sequence(embedded, text_lengths)

        packed_output, (hidden, cell) = self.rnn(packed_embedded)

        hidden = self.dropout(torch.cat((hidden[-2,:,:], hidden[-1,:,:]), dim = 1))
        output = self.fc1(hidden)
        output = self.dropout(self.fc2(output))

        #hidden = [batch size, hid dim * num directions]

        return output
```

Fig 7 : Fast text CNN

An enhanced kind of recurrent neural networks, LSTM is particularly good at catching long-term dependencies absolutely vital for sequence prediction activities. LSTMs are appropriate for jobs like language translation, speech recognition, and time series forecasting since they selectively keep or discard data with 3 gates regulating information flow. LSTM is used in the project for sequential data tasks like user reviews or product descriptions, hence efficiently capturing dependencies for sentiment analysis and contextual awareness.

```
embed_dim = 128 #dimension of the word embedding vector for each word in a sequence
lstm_out = 196 #no of lstm layers
lstm_model = Sequential()
lstm_model.add(Embedding(num_words, embed_dim, input_length = X_train.shape[1]))
#Adding dropout
lstm_model.add(LSTM(lstm_out, dropout=0.2, recurrent_dropout=0.2))
#Adding a regularized dense layer
lstm_model.add(layers.Dense(32, kernel_regularizer=regularizers.l2(0.001), activation='relu'))
lstm_model.add(layers.Dropout(0.5))
lstm_model.add(Dense(3, activation='softmax'))
lstm_model.compile(loss = 'categorical_crossentropy', optimizer='adam', metrics = ['accuracy', f1_m, precision_m, recall_m])
print(lstm_model.summary())
```

Fig 8 LSTM

A simpler substitute for LSTM networks, the hybrid model "combines LSTM and GRU (Gated Recurrent Unit)", both forms of RNN architectures. With the reset gate regulating memory of the prior state and the update gate deciding the impact of fresh input, GRU selectively updates the hidden state via gating methods. This hybrid model is used in the project for thorough user evaluations, hence maximizing the benefits of both architectures for better training efficiency and performance in managing extended sequences.

LSTM + GRU

```
model = Sequential()
model.add(Embedding(num_words, embed_dim,input_length = X_train.shape[1]))
model.add(LSTM(64,dropout=0.4, recurrent_dropout=0.4,return_sequences=True))
model.add(GRU(32,dropout=0.5, recurrent_dropout=0.5,return_sequences=False))
model.add(Dense(3,activation='softmax'))
model.compile(loss = 'categorical_crossentropy', optimizer='adam',metrics = ['accuracy',f1_m,precision_m, recall_m])
print(model.summary())
```

Fig 9 LSTM + GRU

Results and Discussion

Precision: Accuracy evaluates the proportion of cases or samples that are accurately classified between positive specifications. So, the equation for calculating the accuracy is:

“Precision = True positives/ (True positives + False positives) = TP/(TP + FP)”

$$\text{Precision} = \frac{\text{True Positive}}{\text{True Positive} + \text{False Positive}}$$

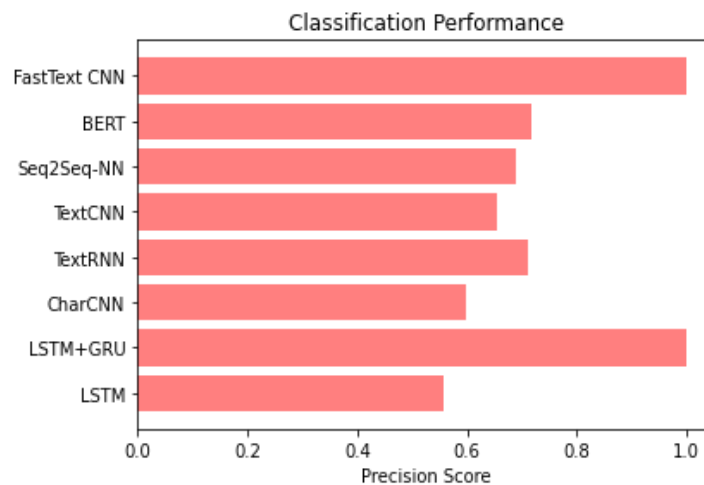


Fig 10 Precision comparison graph

Recall: Calling back is mechanical learning metrics that allow you to find the ability of the M Siet to find all the relevant examples of a particular class. The ratio of accurate positive observations to the actual overall location provides insight into the integrity of the model when recording examples of a particular class.

$$\text{Recall} = \frac{TP}{TP + FN}$$

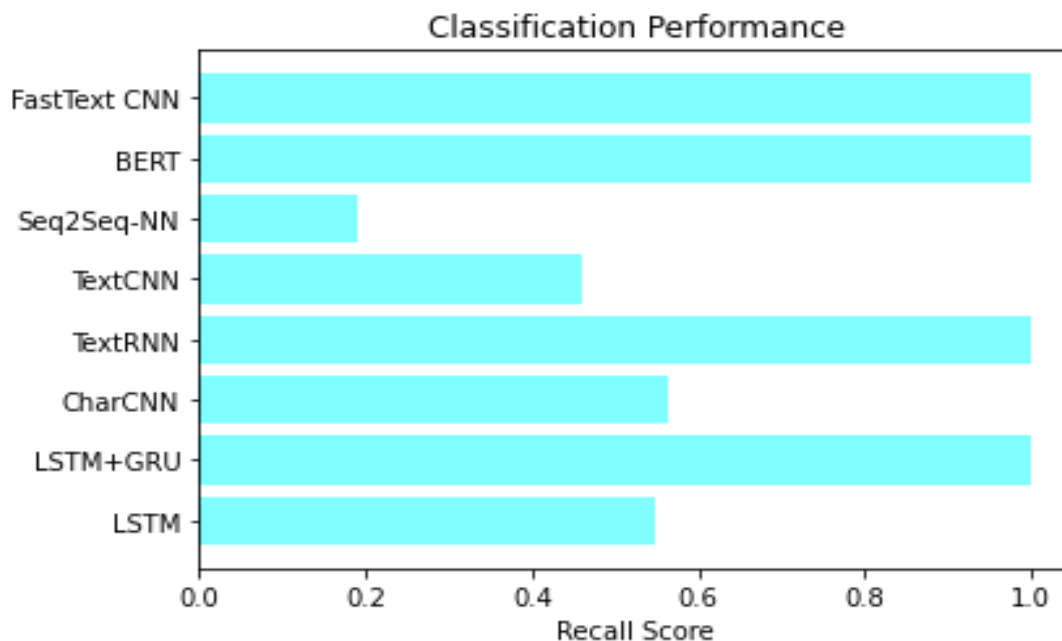
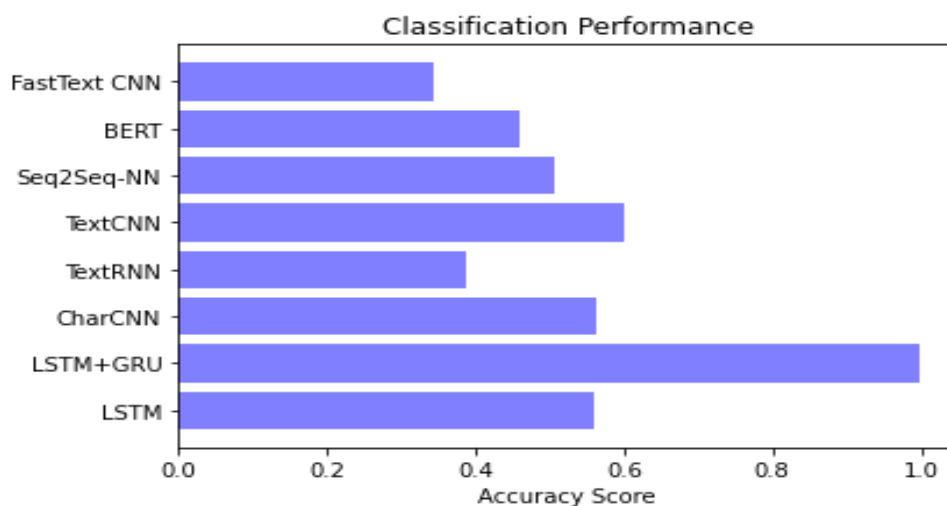


Fig 11 Recall comparison graph

Accuracy: In a classification assignment, “accuracy is the percentage of right forecasts”, hence evaluating the general correctness of a model's predictions.

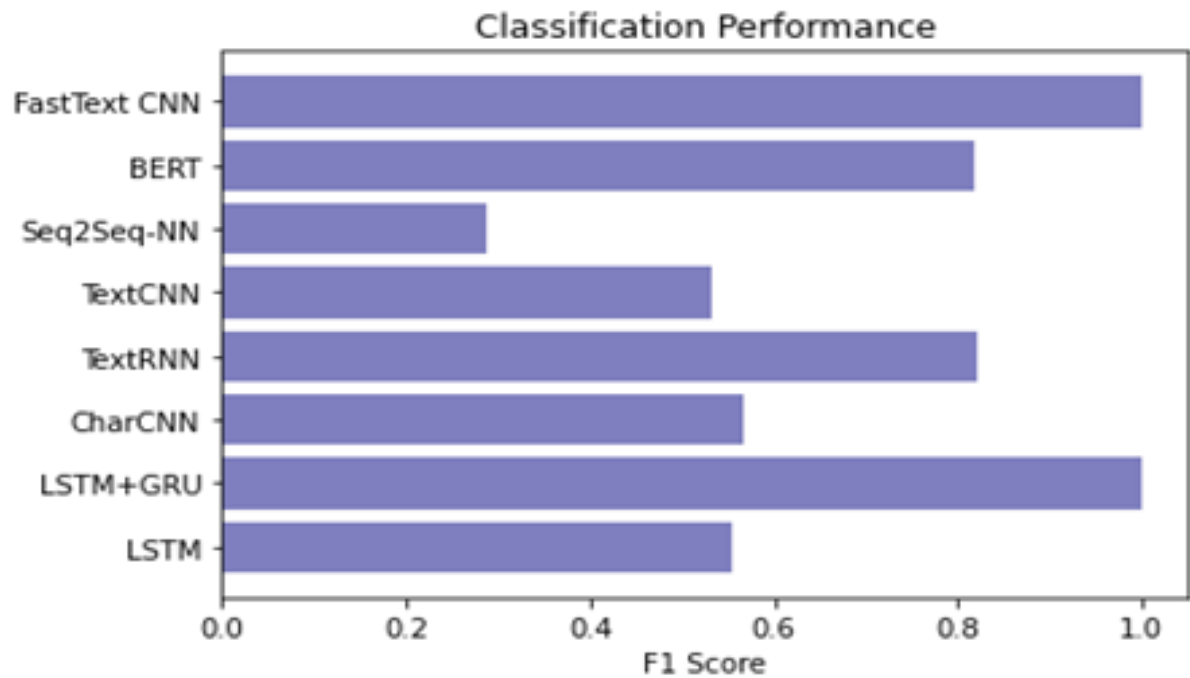
$$Accuracy = \frac{TP + TN}{TP + FP + TN + FN}$$



“Fig 12 Accuracy graph”

F1 Score: offering a balanced measure that takes into account both false "positives and false negatives, the F1 score is the harmonic mean of accuracy and recall", therefore appropriate for imbalanced datasets.

$$F1 \text{ Score} = 2 * \frac{\text{Recall} \times \text{Precision}}{\text{Recall} + \text{Precision}} * 100$$



ML Model	Accuracy	Precision	Recall	F1 - score
Extension LSTM	0.559	0.557	0.549	0.553
Extension LSTM+GRU	0.999	0.999	0.999	0.999
Char CNN	0.562	0.597	0.562	0.566
Text RNN	0.386	0.711	1.000	0.822
Text CNN	0.601	0.655	0.460	0.533
Seq2Seq-NN	0.508	0.688	0.189	0.286
BERT	0.458	0.717	1.000	0.817
Fast Text CNN	0.343	1.000	1.000	1.000

Fig 14 Performance Evaluation table

CONCLUSION

Showcasing a thorough investigation state of-the-art models, the research effectively applied and assessed a varied range state modern sentiment analysis techniques—including Char CNN, text RNN, text CNN, Seq2Seq, BERT, and FastText CNN. The study methodically investigated several deep learning state models, each meant to capture different angles state. emotion in text facts. These models' performance was extensively assessed and compared using evaluation criteria such accuracy, precision, recall, and F1 score. With 99.9% accuracy, the Hybrid model (LSTM + GRU) extension shows better performance and durability, hence providing a good answer for several e-commerce data analysis activities. A user-friendly interface was provided by integrating the sentiment analysis models into a Flask framework with SQLite for user signup and signin. The frontend effectively shows the final results, together with LDA-based topic modelling findings [11,13], thereby allowing users to enter text for sentiment prediction and offering an accessible and engaging platform. Various parties, including e-commerce sites, companies depending on user input, and customers wanting more knowledgeable and nuanced product recommendations, find great value in the results state the initiative. The

project has the ability to positively impact both companies and end-users in the field state sentiment analysis by improving user experience, steering buying decisions, and offering insightful analysis.

FUTURE SCOPE

Consider visual and audio clues from user-generated content in addition to text-based sentiment analysis; this will help the system to grasp emotions across different channels. Extend the technology to offer real-time sentiment monitoring so that e-commerce sites can dynamically change ranks depending on changing user attitudes and guarantee current suggestions [8, 18, 19]. Include user profiling to improve personalization, knowledge of individual preferences over time, and adjustment of product suggestions depending on past sentiment data, therefore offering a more tailored buying experience. Ensure the project stays at the forefront of technical developments in e-commerce by investigating integration with "augmented reality (AR)" for immersive product experiences and blockchain for secure and transparent sentiment data storage.

References

- [1] Bi, J. W., Liu, Y., Fan, Z. P., & Zhang, J. (2019). Wisdom of crowds: Conducting importance-performance analysis (IPA) through online reviews. *Tourism Management*, 70, 460–478.
- [2] Cai, J., Li, J., Li, W., & Wang, J. (2018). Deep learning Model used in text classification. *Proceedings of the 15th International Computer Conference on Wavelet Active Media Technology and Information Processing*, 123–126.
- [3] Fan, Z. P., Che, Y. J., & Chen, Z. Y. (2017). Product sales forecasting using online reviews and historical sales data: A method combining the bass model and sentiment analysis. *Journal of Business Research*, 74, 90–100.
- [4] Fan, Z. P., Li, G. M., & Liu, Y. (2020). Processes and methods of information fusion for ranking products based on online reviews: An overview. *Information Fusion*, 60, 87–97.
- [5] gang Peng, H., yu Zhang, H., & qiang Wang, J. (2018). Cloud decision support model for selecting hotels on TripAdvisor.com with probabilistic linguistic information. *International Journal of Hospitality Management*, 68, 124–138.
- [6] Guo, Y., Barnes, S. J., & Jia, Q. (2017). Mining meaning from online ratings and reviews: Tourist satisfaction analysis using latent dirichlet allocation. *Tourism Management*, 59, 467–483.
- [7] Kou, G., & others. (2021). A cross-platform market structure analysis method using online product reviews. *Technological and Economic Development of Economy*, 27(5), 992–1018.
- [8] Liang, D., Dai, Z., Wang, M., & Li, J. (2020). Web celebrity shop assessment and improvement based on online review with probabilistic linguistic term sets by using sentiment analysis and fuzzy cognitive map. *Fuzzy Optimization and Decision Making*, 19, 561–586.
- [9] Liu, P., & Teng, F. (2019). Probabilistic linguistic TODIM method for selecting products through online product reviews. *Information Sciences*, 485, 441–455.
- [10] Pang, Q., Wang, H., & Xu, Z. (2016). Probabilistic linguistic term sets in multiattribute group decision making. *Information Sciences*, 369, 128–143.
- [11] Socher, R., & others. (2013). Recursive deep models for semantic compositionality over a sentiment treebank. *Proceedings of the Conference on Empirical Methods in Natural Language Processing*, 1631–1642.
- [12] Tirunillai, S., & Tellis, G. J. (2014). Mining Marketing Meaning from Online Chatter: Strategic Brand Analysis of Big Data Using Latent Dirichlet Allocation. *Journal of Marketing Research*, 51(4), 463–479.
- [13] Wong, T.-L., & Lam, W. (2008). Learning to extract and summarize hot item features from multiple auction web sites. *Knowledge and Information Systems*, 14(2), 143–160.
- [14] Wu, X., & Liao, H. (2021a). Learning judgment benchmarks of customers from online reviews. *OR Spectrum*, 43(4), 1125–1157.
- [15] Wu, X., & Liao, H. (2021b). Modeling personalized cognition of customers in online shopping. *Omega (Westport)*, 104, 102471.
- [16] Yan, Z., Xing, M., Zhang, D., & Ma, B. (2015). EXPRS: An extended pagerank method for product feature extraction from online consumer reviews. *Information & Management*, 52(7), 850–858.
- [17] Zhang, D., Xu, H., Su, Z., & Xu, Y. (2015). Chinese comments sentiment classification based on word2vec and SVMperf. *Expert Systems with Applications*, 42(4), 1857–1863.
- [18] Zhang, Z., Ye, Q., Zhang, Z., & Li, Y. (2011). Sentiment classification of Internet restaurant reviews written in

Cantonese. *Expert Systems with Applications*, 38(6), 7674–7682.

- [19] Zhao, M., Shen, X., Liao, H., & Cai, M. (2021). Selecting products through text reviews: An MCDM method incorporating personalized heuristic judgments in the prospect theory. *Fuzzy Optimization and Decision Making*.