



Detecting Positive and Negative Deviations in Cross-Domain Product Reviews using Adaptive Stochastic Deep Networks

B. Shanthini*, N. Subalakshmi

Department of Computer and Information Science, Annamalai University, Chidambaram, Tamil Nadu, India

Emails: shanthinib.66@gmail.com; subhaabi12@gmail.com

Abstract

The analysis of sentiment in product reviews across diverse platforms such as e-commerce website and social media presents a challenging task due to the inherent differences in user behaviour and review formats. This research introduces an innovative methodology for detecting positive and negative deviations in cross-domain product reviews using Adaptive Stochastic Deep Networks (ASDN) tailored for multi-platform sentiment analysis. ASDNs possess mechanisms that enable dynamic adaptation to changes in data distributions, domain shifts, or varying complexities within the input data. The proposed framework aims to capture refined variations in sentiment expression across disparate platforms by incorporating adaptive stochasticity within deep neural networks. By adapting dynamically to changes in review styles, language use, and sentiment patterns unique to each platform, the ASDN architecture facilitates the identification of nuanced sentiment shifts. Through extensive experimentation on comprehensive datasets spanning Amazon, Facebook, and Instagram, the efficacy of the ASDN model in detecting positive and negative sentiment deviations across diverse platforms is demonstrated. This research contributes to advancing the understanding of sentiment dynamics across distinct social platforms and e-commerce sites, paving the way for more robust and adaptable models in cross-domain sentiment analysis.

Keywords: Adaptive Stochastic Deep Networks; Deep Neural Networks; Cross-domain; e-commerce Website; Sentiment Analysis.

1. Introduction

In the digital era, the evolution of consumer behavior and the proliferation of online platforms have reshaped the landscape of product evaluation and feedback dissemination. Two primary domains that significantly influence consumer opinions and purchasing decisions are e-commerce sites and social media platforms [1]. These sites allow people to share their thoughts, views, and feelings about items and services, thereby shaping the perceptions of potential buyers. E-commerce platforms have emerged as centralized hubs for comprehensive product reviews and evaluations [2]. Within these digital marketplaces, consumers engage in the meticulous documentation of their product experiences, offering detailed narratives encompassing the pros, cons, functionalities, and overall satisfaction levels with the purchased items [3].

Product reviews on e-commerce sites are often characterized by structured textual content, enriched with quantitative ratings, specifications, and frequently accompanied by multimedia elements such as images or videos [4]. Such evaluations are critical in aiding prospective buyers decision-making processes, providing insights into product performance, durability, usability, and overall customer satisfaction [5]. Social media platforms, including Facebook, Instagram, Twitter, and others, serve as dynamic arenas where users share concise yet expressive opinions and experiences. Within these platforms, user-generated content manifests as succinct reviews, comments, posts, or images encapsulating fleeting sentiments, often accompanied by visual elements or emoticons [6]. Product evaluations on social media platforms unfold as snapshots of user sentiments, frequently characterized by informal language, emotive expressions, and the amalgamation of textual content with multimedia components [7]. These bite-sized expressions encapsulate the immediate reactions, recommendations, or grievances of users, fostering a sense of community-driven opinion-sharing and influencing peer perceptions. The divergence in the nature and style of product reviews between e-commerce sites and social media platforms underscores the nuanced

behaviors of users within distinct digital ecosystems [8]. E-commerce sites facilitate in-depth, structured evaluations contributing to informed purchase decisions, while social media platforms foster a more spontaneous, community-driven environment, shaping perceptions through concise and visually enriched user-generated content [9].

The assessment of sentiment in product reviews serves as a pivotal aspect of understanding consumer preferences, influencing purchasing decisions, and guiding market strategies [10]. In the contemporary digital era, where consumers actively engage and express their opinions across diverse online platforms such as e-commerce websites, social media, and forums, comprehending the nuanced sentiment variations becomes a crucial endeavour. The dynamic and diverse nature of these platforms results in a rich tapestry of user-generated content, comprising reviews characterized by varying linguistic styles, sentiments, and user engagement patterns. This research aims to delve into the development and evaluation of ASDNs tailored explicitly for cross-domain sentiment analysis in product reviews. Through extensive experimentation and analysis, we seek to elucidate the efficacy of ASDNs in discerning positive and negative deviations, offering insights into nuanced sentiment shifts across platforms.

This research focuses on the identification and analysis of positive and negative deviations within cross-domain product reviews, particularly across platforms like Amazon, Facebook, and Instagram as shown in fig 1.

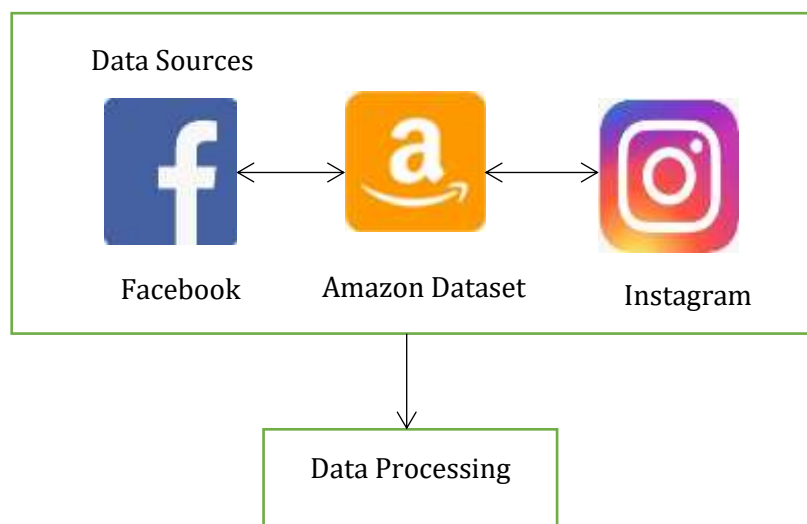


Figure 1: Data Collection Process

In the subsequent sections, we elaborate on the architecture design of ASDNs, delineate the methodology adopted for model training and evaluation, and present the experimental results, culminating in comprehensive insights into the detection of sentiment deviations across diverse product review domains.

2. Related Work

Reading feedback from consumers before making an online purchase has become commonplace; nevertheless, some businesses utilize ML algorithms to produce fraudulent reviews in order to promote favourable brand images of their own goods and bad brand images of rivals offers. In this paper [11], we describe a unique feature engineering strategy for detecting the issue of fake reviewers by developing a unique M-SMOTE model that is trained on a balanced dataset and feature distributions. The results show that this model outperforms earlier machine learning methods.

This research [12] examines the phenomenon of review manipulation by analyzing variations in rating distributions and proposing a detection technique that identifies review bias. The degree to which the associated departs from the median is used to quantify the severity of review manipulation. Then, for biased review identification, deep learning algorithms were used to learn concealed textual depictions of un manipulated reviews. [13] The proposed model introduces an automated sentiment analysis approach that classifies user reviews as negative, positive, or neutral emotions. It also identifies deceptive ratings from other users on the social media platform, aiming to aid the user in making informed decisions.

Sentiment analysis is the problem of extracting people's opinions from text using computational language processing tools. [14] Suggested a mixture of several RNNs models employing distinct word embeddings in this study. Fake reviews, often known as false opinions, are used to mislead consumers and have lately gained popularity. The proposed methodology adopted in this work [15] used a standard fake hotel review dataset for experimenting and data preprocessing methods and an approach for extracting features and their representation. In order to fill this void in research, the present study [16] investigates consumer happiness via aspect-level sentiment analysis and visual analytics. Additionally, it aims to identify subjective opinions related to certain aspects and predict emotions that have not been evaluated yet across several categories.

Despite the relatively large quantity of data created, the identification of spam reviews in Arabic internet sources is a relatively new issue. As a result, this work [17] adds to this area by offering four alternative Arabic spam reviews detection algorithms, with an emphasis on the development and assessment of an ensemble approach. Emotional recognition has emerged as an important subject of research that may reveal a wide range of significant inputs. Word embeddings are widely used in NLP activities. This work [18] proposes the use of Deep Learning Assisted Semantic Text Analysis (DLSTA) to identify human emotions by leveraging large datasets.

This research [19] addresses the issue of spam reviews in electronic shopping, crucial for maintaining genuine product feedback. Existing supervised methods face limitations due to manual labeling inaccuracies, data imbalance, and computational costs in assessing review similarities. A review representation model that successfully utilizes the domain context is presented in this study [20] and is based on linguistic characteristics that are dependent on behavior and mood.

Online reviews are becoming more crucial for making purchasing decisions. Before making an acquisition, consumers often consult internet reviews for feedback. This research [21] provided a unique paradigm for assessing online review ratings using ML methods. This system employs a mix of text pre-processing and feature extraction techniques. Because of the growing amount of Internet transactions, fake customer review detection has gained a lot of attention in recent years. [22] Presented are two NN models that include both the conventional bag-of-words approach and the consideration of word context and customer emotions. These models aim to enhance the accuracy of detecting fraudulent reviews.

Cognitive computing is an interdisciplinary academic topic that uses electronic models to imitate human cognitive processes. This study [23] investigates the effect of two key textual variables, namely word count and review readability, on sentiment classification performance. Consumers might be deceived by counterfeit reviews. An excessive quantity of fraudulent evaluations has the ability to lead to significant harm to properties and cause public relations crises [24].

Therefore, it is crucial to identify and screen out fraudulent reviews. However, owing to the use of single characteristics and a lack of labeled experimental data, most current algorithms have lesser accuracy in identifying false reviews [25]. By using cutting-edge feature engineering methods and optimizing hyperparameters to increase prediction accuracy and reduce overfitting, this study expands on that base. Using historical market data, the research incorporates global demand-supply patterns, sentiment analysis of the market, and important economic indicators as predictive elements [26].

In this [27], using Stacked Autoencoders, an ideal Deep Neural Network based architecture is utilized for the categorization of credit score data. In this case, SA is used to extract the dataset's characteristics. Research has been done to increase the capacity and identify assaults without compromising the network's efficiency when technology first emerged. Two methods have been presented for the learning phase. First, to solve the problems in ID using PID (Perimeter Intrusion Detection) with MLP (Multi-Layer Perception) and quantum classifier, second, to characterize a collection of articles that depend on preparation and testing using the PID with MLP and quantum classifier algorithm, which is generally useful. Use the suggested techniques in order to meet all requirements, including increased consistency, accuracy, and efficiency [28].

In this study, we are addressing the challenge where transformations from positive to negative review sentiments or vice versa are recognized as potentially misleading or fake deviations. The proposed solution aims to overcome this challenge through the implementation of Adaptive Stochastic Deep Networks within cross-domain product reviews.

3. Proposed Model

Our research focuses on data sources from multiple platforms like Amazon, Facebook, and Instagram to analyze cross-domain product reviews and detect variations in positive and negative sentiments. An overall architecture of proposed model is shown in fig 2.

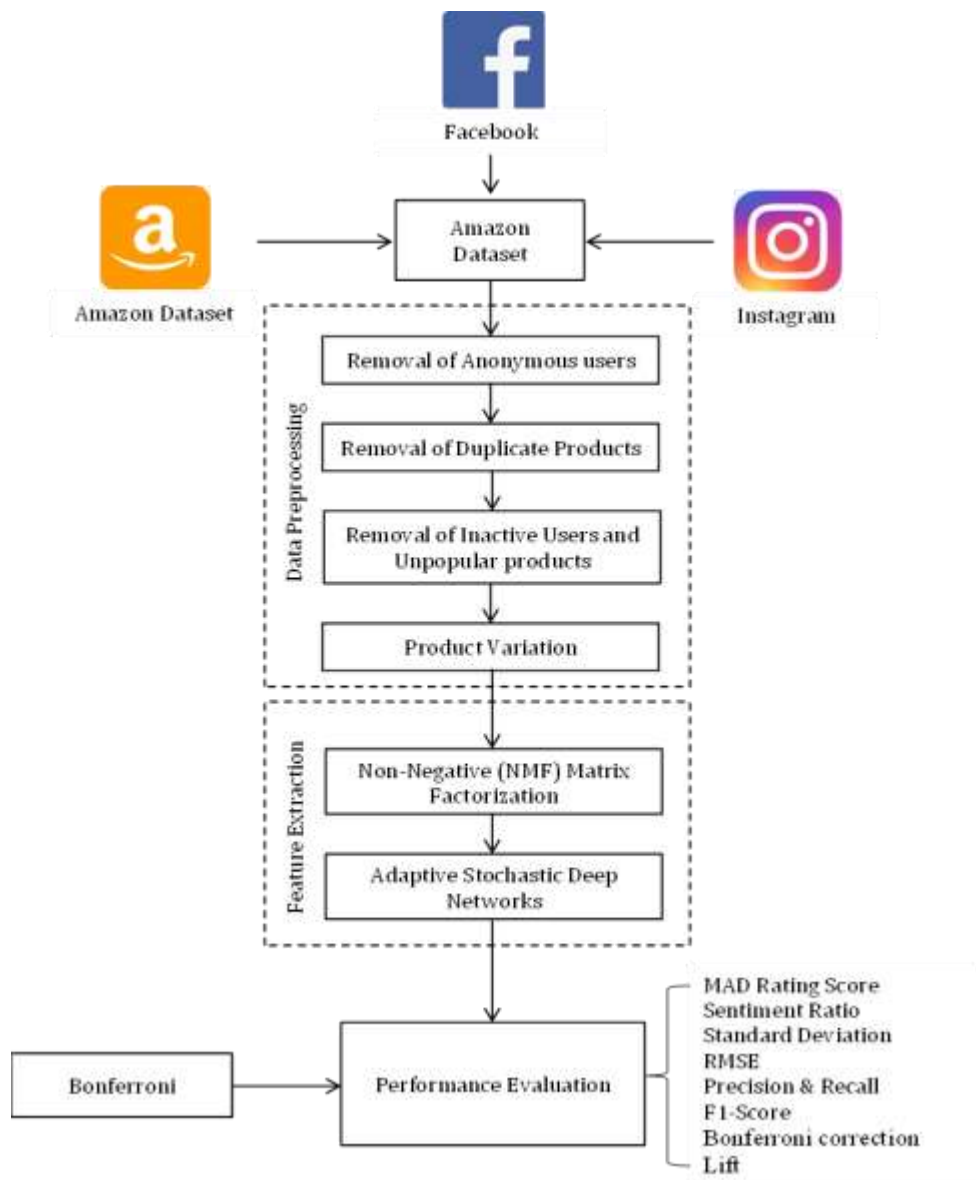


Figure 2: Overall Design of the Proposed Model

Gather labeled datasets from Amazon, Instagram, and Facebook containing reviews with known deviations (positive or negative). Pre-process the data using various steps to ensure data quality, consistency, and compatibility with the chosen modeling techniques. Extract relevant features from the reviews using the Non-Negative Matrix Factorization (NMF) method that captures differences between original and subsequent reviews. Train the model using the labeled dataset, utilizing features that encapsulate deviations in reviews using Adaptive Stochastic Deep Networks.

3.1 Pre-processing

Pre-processing steps for a product review dataset typically involve several key procedures to clean, prepare, and organize the data before using it for analysis or machine learning tasks. Here are some common pre-processing steps:

Removal of Anonymous Users:

Identify reviews or entries where the user information is missing or labeled anonymously. Filter out these reviews or entries from the dataset if they don't add value to the analysis. Let D represent the dataset. Let U be the set of users and R be the set of reviews.

$$D' = D - \{r \in R: \text{user}(r) \text{ is anonymous}\} \quad (1)$$

Where D' is the updated dataset with anonymous reviews removed.

Removal of Duplicate Products:

Detect and remove product entries that are duplicated within the dataset. Criteria for duplication might include identical product names, descriptions, or unique identifiers. Let P be the set of products.

$$D'' = D' - \{p \in P: \text{count}(p) > 1\} \quad (2)$$

Where D'' is the dataset after removing duplicate product entries.

Removal of Inactive Users and Unpopular Products:

Define criteria for identifying inactive users (users who haven't contributed reviews within a certain time frame) or unpopular products (products with very few reviews). Remove reviews associated with inactive users or products that don't meet predefined popularity thresholds. Let U_i represent inactive users, and P_u represent unpopular products.

$$D''' = D'' - \{r \in R: \text{user}(r) \in U_i \text{ or } \text{product}(r) \in P_u\} \quad (3)$$

Where D''' is the dataset after removing reviews associated with inactive users and unpopular products.

Resolution of Product Variation:

By Mapping or dictionary that links synonymous brand names or product variations to a single canonical name. Use string matching is useful to identify and consolidate brand name variations. Replace synonyms or variations with the standardized brand name throughout the dataset. Let B be the set of brand names. Let M be the mapping/dictionary of synonyms.

$$\text{resolve}(b) = M[b] \text{ for } b \in B \quad (4)$$

Where $\text{resolve}(b)$ represents the resolved canonical name for the brand b .

3.2 Feature Extraction

Non-Negative Matrix Factorization (NMF) is used as a feature extraction technique for product review datasets to capture underlying patterns or topics within the textual data. NMF extracts essential themes from product reviews, distinguishing between positive and negative sentiments by decomposing the review matrix into interpretable components. By identifying distinct topics or patterns in non-negative data, NMF enables the capture of sentiment-related features, aiding in the differentiation and characterization of positive and negative aspects within the reviews. This technique offers a structured representation as shown in fig 3, allowing for deeper insights into sentiments and content across different sentiment categories within the product review dataset.

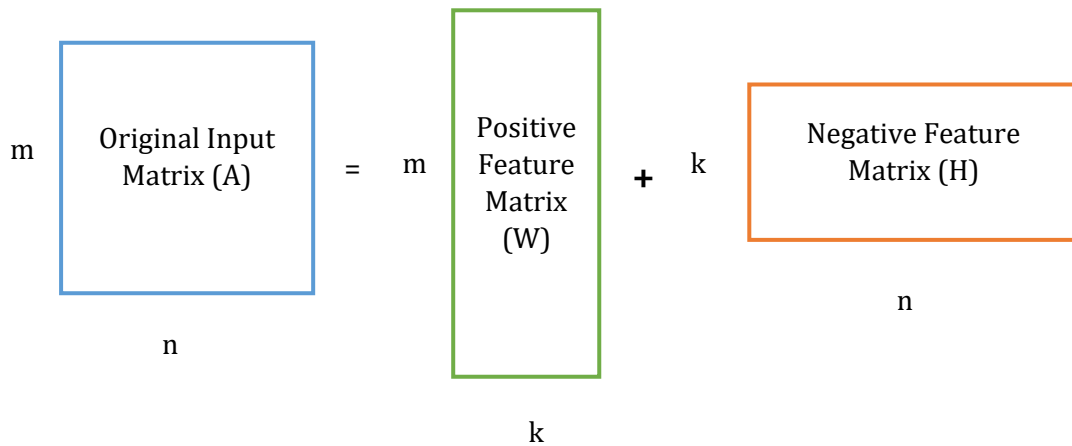


Figure 3: Feature Extraction Method

A mathematical representation of Non-Negative Matrix Factorization (NMF) applied to a product review dataset to extract features related to positive and negative reviews are given below:

Suppose we have a document-term matrix V representing the product reviews. This matrix is of size $m \times n$, where m is the number of reviews, and n is the number of unique terms (words, features, etc.) in the reviews.

NMF factorizes the matrix V of dimensions $m \times n$, where each element is ≥ 0 , into two non-negative matrices W and H having dimensions $m \times k$ and $k \times n$ respectively, aiming to approximate V as the product of W and H :

$$A_{m \times n} = W_{m \times k} H_{k \times n} \quad (5)$$

In the context of positive and negative reviews, the matrices W and H would capture latent patterns or topics present in the reviews, potentially segregating the data into positive and negative sentiment-related features. The interpretation of these matrices would involve analyzing the topics/components and associated terms to understand the underlying sentiments expressed in the reviews.

3.3 Adaptive Stochastic Deep Networks

Adaptive Stochastic Deep Networks (ASDNs) offer a sophisticated framework designed to detect nuanced sentiment deviations in cross-domain product reviews across diverse platforms. These networks dynamically adjust their structures and parameters, incorporating adaptive mechanisms and controlled stochasticity to discern positive and negative deviations. By seamlessly adapting to varying linguistic styles and user behaviors, ASDNs capture subtle shifts in sentiment expressions, fostering robust sentiment analysis. Their adaptability enables effective modeling of nuanced sentiment dynamics within disparate e-commerce sites, social media platforms, and forums. ASDNs stand poised to revolutionize sentiment analysis methodologies, unraveling the complexities of cross-domain reviews and influencing informed consumer decisions.

Let X represent the input data comprising cross-domain product reviews. Y denotes the corresponding sentiment labels indicating positive or negative sentiment for each review.

Consider an ASDN with L layers, denoted as $L = \{1, 2, \dots, L\}$, including input, hidden, and output layers.

The output of each layer l can be represented as:

$$Z^{(l)} = \sigma(W^{(l)} \cdot A^{(l-1)} + b^{(l)}) \quad (6)$$

Where:

$Z^{(l)}$ is the output of layer l .

$W^{(l)}$ represents the weight matrix for layer l .

$A^{(l-1)}$ is the activation output from the previous layer.

$b^{(l)}$ denotes the bias vector for layer l .

σ represents the activation function applied element-wise.

Stochastic elements, such as dropout or stochastic activation functions, can be introduced within the ASDN architecture. For instance, consider the inclusion of dropout regularization:

$$A_{dropout}^{(l-1)} = A^{(l-1)} \odot D^{(l-1)} \quad (7)$$

Where:

$A_{dropout}^{(l-1)}$ is the output after applying dropout to the activation.

$D^{(l-1)}$ is a binary dropout mask.

$\odot$ denotes element-wise multiplication.

The loss function for training the ASDN could be defined using a suitable measure like cross-entropy loss:

$$J(W, b) = -\frac{1}{m} \sum_{i=1}^m \left[Y_i \log(\hat{Y}_i) + (1 - Y_i) \log(1 - \hat{Y}_i) \right] \quad (8)$$

Where:

$J(W, b)$ is the overall loss function.

m represents the number of samples.

Y_i denotes the true label for the i th sample.

$\hat{Y}_i$ is the predicted probability for the i th sample.

Adaptive mechanisms such as learning rate adaptation or layer-wise adaptation techniques can be incorporated within the optimization algorithm, altering the network's weights and biases during training based on the performance feedback. Adaptive optimization algorithm that dynamically adjusts the learning rates for individual parameters based on past gradients.

The update steps for the parameters W of the network using Adam can be expressed as:

$$m_t = \beta_1 \cdot m_{t-1} + (1 - \beta_1) \cdot \nabla J(W) \quad (9)$$

$$v_t = \beta_2 \cdot v_{t-1} + (1 - \beta_2) \cdot (\nabla J(W))^2 \quad (10)$$

$$\hat{m}_t = \frac{m_t}{1 - \beta_1^t} \quad (11)$$

$$W_{t+1} = W_t - \alpha \cdot \frac{\hat{m}_t}{\sqrt{\hat{v}_t + \epsilon}} \quad (12)$$

Layer normalization is a technique that normalizes the outputs of each layer to stabilize training.

For each layer l , the normalization process can be represented as:

$$\mu_l = \frac{1}{m} \sum_{i=1}^m z_i^{(l)} \quad (13)$$

$$\sigma_l = \sqrt{\frac{1}{m} \sum_{i=1}^m (z_i^{(l)} - \mu_l)^2 + \epsilon} \quad (14)$$

$$\hat{z}_i^{(l)} = \frac{z_i^{(l)} - \mu_l}{\sigma_l} \quad (15)$$

Where:

- $z_i^{(l)}$ represents the activations of the i th neuron in layer l .
- m denotes the number of neurons in the layer.
- μ_l and σ_l are the mean and standard deviation of the layer activations.
- $\hat{z}_i^{(l)}$ represents the normalized activations.

These equations demonstrate the mathematical representations of adaptive mechanisms in neural network training, exemplifying how learning rate adaptation and layer-wise adaptation can be formalized mathematically within a neural network training process.

Pseudocode for Adaptive Stochastic Deep Networks

```

FUNCTION trainDeepNetwork(trainingData):
    InitializeStochasticDeepNetwork() # Initialize the network architecture

    FOR each epoch in trainingData:
        ForwardPropagate(trainingData) # Pass data through the network
        CalculateLoss() # Calculate loss using a defined loss function
        BackwardPropagate() # Backpropagation to update weights

```

```

RETURN trainedNetwork

FUNCTION detectDeviations(trainedNetwork, newData):

predictedLabels = PredictLabels(trainedNetwork, newData)

positiveDeviations = []

negativeDeviations = []

FOR each data point in newData:

    IF predictedLabels[data point] == "positive":

positiveDeviations.Append(data point)

    ELSE IF predictedLabels[data point] == "negative":

negativeDeviations.Append(data point)

RETURN positiveDeviations, negativeDeviations

```

The above pseudocode outlines a structure for training an adaptive stochastic deep network using training data and then using the trained network to detect positive and negative deviations in new data. The *trainDeepNetwork* function is responsible for training the network, while the *detectDeviations* function uses the trained network to classify deviations in new data into positive and negative categories based on predicted labels. Here are the main steps represented:

Initialization: Initialize the parameters of the ASDN network.

Training Loop: Iterate through a certain number of epochs. Within each epoch, divide the dataset into mini-batches.

Forward Propagation: Perform forward propagation through the network to obtain predicted outputs.

Loss Computation: Calculate the loss between predicted outputs and actual labels.

Backward Propagation: Perform backward propagation to compute gradients.

Parameter Update: Update the parameters of the network using an optimization algorithm based on computed gradients.

Inference: After training, the network is ready for inference, making predictions on new data (new_reviews).

ASDNs, through their adaptive mechanisms, dynamically tailor their architectures and learning strategies to accommodate variations in linguistic styles, user behaviors, and review formats present in cross-domain product reviews. By integrating controlled stochasticity within their structures, ASDNs excel at capturing subtle nuances in sentiment expressions, enabling them to discern the positive and negative deviations prevalent in reviews across platforms like e-commerce sites.

Bonferroni Correction

The Bonferroni adjustment is a strategy used in test of statistical hypotheses to overcome the problem of multiple comparisons. When running many statistical tests at the same time, the danger of false positives or Type I errors increases. The Bonferroni correction adjusts the significance level (α) used in individual tests to control the overall error rate at the desired level.

Let assume, significance level at $\alpha=0.05$. This is the likelihood of committing a Type I mistake (false positive) in a single test. Determine the total number of statistical tests which is represented as m . To apply the Bonferroni correction, divide the original significance level (α) by the number of comparisons (m):

$$\text{Adjusted Significance Level} = \alpha / m \quad (16)$$

This adjusted significance level (α/m) becomes the new threshold for statistical significance in each individual test to control the overall error rate. Bonferroni correction adjusts the significance level for individual tests to mitigate the risk of false positives when conducting multiple tests.

4. Results and Discussions

4.1 Dataset Description

Amazon dataset was collected from <https://www.kaggle.com/datasets/lokeshparab/amazon-products-dataset>. Its product data is divided into 142 categories in .csv format, as is the whole dataset. Each csv file includes ten columns, and each row contains product information. Amazon is an American Tech Multi-National Company whose business interests include E-commerce, where they purchase and store product, as well as handles everything from shipment and billing to customer support and refunds. The features of Dataset are described in table 1.

Table 1: Features Description in Dataset

NAME	DESCRIPTION
name	The product's name
main_category	The product's major category
sub_category	The product's primary category
image	The product image looks like
link	The product's Amazon reference link
ratings	The product ratings provided by Amazon customers
no of ratings	The amount of Amazon shopping reviews for this product
discount_price	The product's discounted pricing
actual_price	The product's real MRP

4.2 Performance Evaluation for Review Deviation

When comparing deviations in product reviews (fake or manipulated) across different platforms like Amazon, Instagram, and Facebook, we can employ various performance evaluation metrics to assess the effectiveness of your detection methods. Here are some evaluation metrics suitable for identifying deviations in reviews:

Mean Absolute Deviation (MAD) of Rating score

A sentiment analysis can play a crucial role in determining a rating score for products or services based on reviews or opinions expressed by users. Calculate the average absolute differences between predicted deviations and actual deviations. This metric quantifies the average magnitude of deviations, providing a baseline for comparison. By calculating the MAD of the rating score of the current review from the average rating score of all reviews related to the same product is given as:

$$\text{MAD Rating score} = \text{mms}(|rs(r) - \text{avgc} \in p(r)(rs(c))|) \quad (17)$$

Where, $rs(r)$ represents the rating score for the current review r . $\text{avgc} \in p(r)(rs(c))$ denotes the average rating score for all reviews c related to the same product item that the current review r addresses. $p(r)$ represents the set of all reviews discussing the same product item as review r . The formula computes the Mean Absolute Deviation using the Min-Max-Scaler (MMS) to normalize the score deviation. The Min-Max-Scaler normalizes values between a specified range (often between 0 and 1). The formula quantifies how much the rating score of a particular review deviates from the original rating score of both positive and negative reviews discussing the same product item, and it normalizes this deviation using the Min-Max-Scaler for consistency.

Sentiment Ratio

Sentiment Ratio refers to the ratio or proportion of the prevailing sentiment polarity (positive or negative) compared to all sentiments expressed in the review.

$$\text{Sentiment Ratio} = \text{avgr} \in R \mid DPWords(r) \mid / \mid PWords(r) \mid \quad (18)$$

Where, $\in R$: Denotes the average value for review. $DPWords(r)$: Represents the sentiment words with dominant polarities (positive or negative) in review r . $PWords(r)$: Denotes all sentiment words present in review r .

The above formula measures the proportion of sentiment words with dominant polarities (positive or negative) in a review compared to all sentiment words present in that review. The formula computes the purity of sentiment in a review by considering the ratio of sentiment words with dominant polarities to all sentiment words in that review.

Standard Deviation

The standard deviation (SD) of a dataset is a measure of its variance or dispersion. It estimates how much every point of data deviate from the dataset's mean (average) value. It Measure the variability of deviations from the mean deviation across the datasets. A higher standard deviation signifies greater variability in deviations.

The SD of a dataset X with n data points may be calculated using the following formula:

Given a dataset X with data points $1, x_2, \dots, x_n$:

1. Calculate the mean (μ) of the dataset:

$$\mu = \frac{1}{n} \sum_{i=1}^n x_i \quad (19)$$

2. Determine the sum of the squared discrepancies between each data point and the mean:

$$\text{Squared Differences} = (x_1 - \mu)^2, (x_2 - \mu)^2, \dots, (x_n - \mu)^2 \quad (20)$$

3. Find the average of these squared differences:

$$\text{Average Squared Difference} = \frac{1}{n} \sum_{i=1}^n (x_i - \mu)^2 \quad (21)$$

4. Finally, take the square root of the average squared difference to get the standard deviation:

$$\text{Standard Deviation (SD)} = \sqrt{\text{Average Squared Difference}} \quad (22)$$

The mean value from which data points deviate or are dispersed is quantified by the standard deviation. The metric offers valuable information regarding the dispersion or fluctuation of the data relative to the mean. A larger standard deviation signifies increased variability, while a smaller one suggests reduced variability or greater consistency within the dataset.

Root Mean Squared Error (RMSE)

Utilized frequently, the RMSE quantifies the average discrepancy between the predicted and actual results of a dataset. By imposing a greater penalty on larger deviations, RMSE offers valuable insights into the accuracy of the model.

The formula to calculate RMSE involves the following steps:

Given a dataset with n data points, actual values y_i , and corresponding predicted values $\hat{y}_i$ for $i = 1$ to n :

1. Calculate the squared differences between the actual and predicted values for each data point:

$$\text{Squared Differences} = (\hat{y}_1 - y_1)^2, (\hat{y}_2 - y_2)^2, \dots, (\hat{y}_n - y_n)^2 \quad (23)$$

2. Find the average of these squared differences:

$$\text{Mean Squared Error (MSE)} = \frac{1}{n} \sum_{i=1}^n (\hat{y}_i - y_i)^2 \quad (24)$$

3. Compute the square root of the mean squared error to obtain RMSE:

$$\text{RMSE} = \sqrt{\text{MSE}} = \sqrt{\frac{1}{n} \sum_{i=1}^n (\hat{y}_i - y_i)^2} \quad (25)$$

The root mean square error (RMSE) is a metric that measures the average deviation between projected values and actual values. It serves as an indicator of the accuracy of a predictive model in forecasting numerical results. A smaller root mean square error (RMSE) signifies superior accuracy, implying that the model's predictions are in closer proximity to the actual values. Conversely, a higher RMSE signifies larger prediction errors, indicating lower accuracy in the model's predictions.

Precision and Recall

Precision and recall are often used assessment measures in binary classification tasks to gauge the effectiveness of a model's predictions. They are particularly crucial for handling unbalanced datasets, characterized by a substantial variation in the number of occurrences across various classes.

Precision quantifies the degree of correctness in the positive predictions generated by the model. The term "precision" refers to the ratio of accurately predicted positive cases to all instances projected as positive, including both true positives and false positives.

The formula for precision is:

$$\text{Precision} = \text{True Positives} / (\text{True Positives} + \text{False Positives}) \quad (26)$$

Recall quantifies the model's capacity to accurately detect all positive events. It quantifies the accuracy of predicting positive cases by comparing the number of accurately predicted positive instances to the total number of actual positive instances.

The formula for recall is:

$$\text{Recall} = \text{True Positives} / (\text{True Positives} + \text{False Negatives}) \quad (27)$$

Precision and recall are both measured on a scale of 0 to 1, with greater values indicating superior performance.

Additionally, The F1-score is calculated by taking the harmonic mean of accuracy and recall, which are measures that may be combined.

$$\text{F1-score} = 2 \times (\text{Precision} \times \text{Recall}) / (\text{Precision} + \text{Recall}) \quad (28)$$

The F1-score is a statistic that strikes a compromise between accuracy and recall. It is particularly valuable in binary classification tasks since it takes into account both false positives and false negatives. When evaluating deviations across different platforms (Amazon, Instagram, Facebook), consider performing these evaluation metrics separately for each platform and then comparing the results. The normalization of top 20 brands in amazon dataset is shown in fig 4.

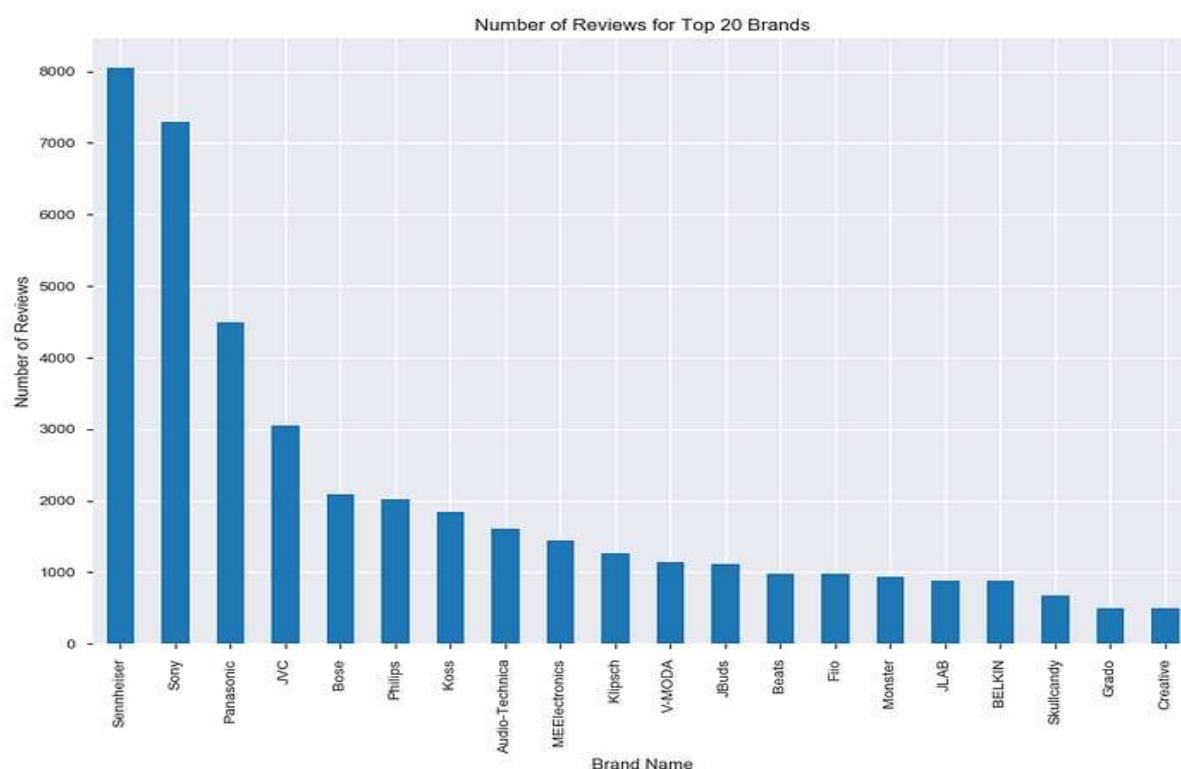


Figure 4: Normalization of Number of Reviews in Amazon Dataset

After Applying the Bonferroni correction in the context of filtering candidates who provide both positive and negative reviews while deviating from their original reviews across various platforms. These steps filtered 207 reviews which are deviated from original opinion of the product from various data sources as shown in fig 5 and 6. The deviations of Top 20 products from various data sources are shown in table 2.

Table 2: Deviation for Top 20 Products from various data sources

No. of Reviews	61708
No. of Positive Reviews	52634
No. of Negative Reviews	9074
Deviated from original opinion	207
Deviated from Positive to Negative	113
Deviated from Negative to Positive	94
Total Positive Reviews	52521
Total Negative Reviews	8980

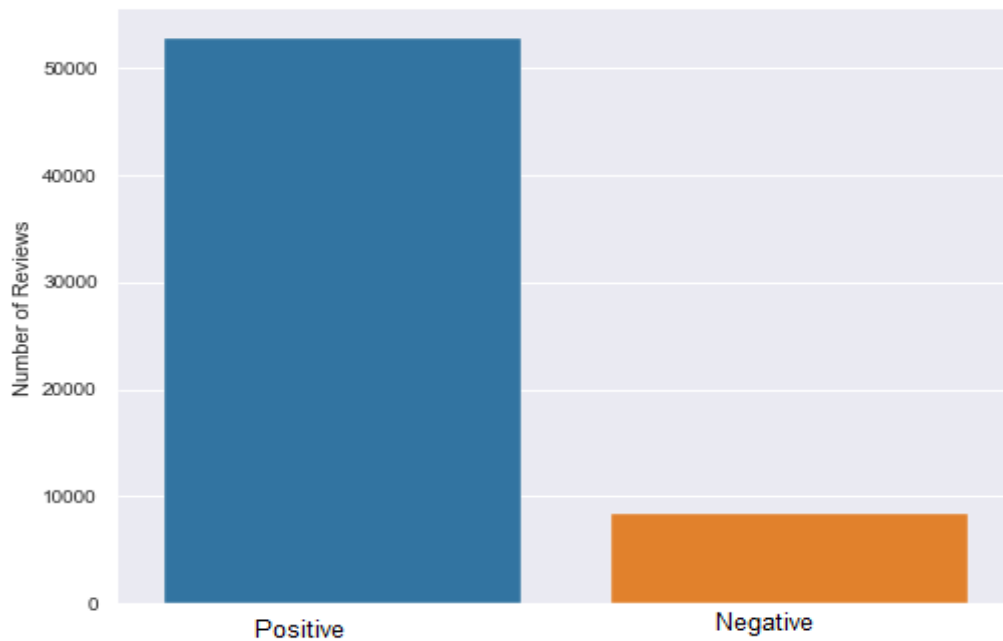


Figure 5: Total Review for the Products

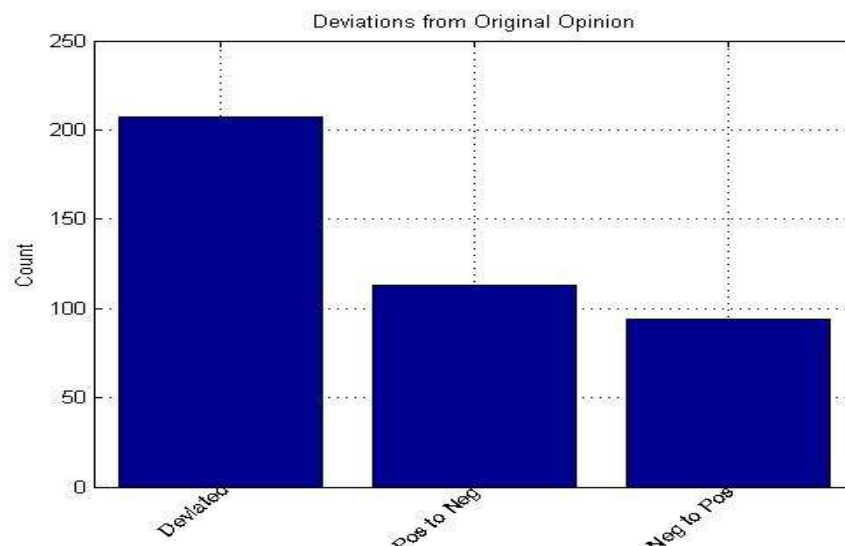


Figure 6: Deviations from Original Review

Lift Value:

The lift value is the ratio of the predicted rate to the baseline rate. The baseline rate represents the probability of the targeted outcome occurring without any intervention or predictive model. This is often the natural occurrence rate in the absence of any targeting or strategy. The predicted rate is the probability of the targeted outcome occurring when using predictive model or strategy.

$$lift = (TP / (TP + FP)) / (TP + FN) / (TP + TN + FP + FN) \quad (29)$$

Table 3: Comparison of Performance Evaluation

Model	MAD	Sentiment Ratio	SD	RMSE	Precision	Recall	F1-Score	Bonferroni correction	LIFT
LSTM	0.32	3:1 (Pos: Neg)	0.22	0.40	0.83	0.75	0.79	0.01	0.773

CNN	0.35	2:1 (Pos: Neg)	0.25	0.43	0.85	0.78	0.81	0.01	0.758
Bi-LSTM	0.30	2:1 (Pos: Neg)	0.20	0.38	0.86	0.79	0.82	0.01	0.795
BERT	0.25	4:1 (Pos: Neg)	0.18	0.35	0.90	0.85	0.87	0.00833	0.800
K-BERT	0.24	4:1 (Pos: Neg)	0.17	0.34	0.91	0.86	0.88	0.00833	0.825
ASDN	0.22	3:1 (Pos: Neg)	0.15	0.32	0.94	0.90	0.92	0.00833	0.912

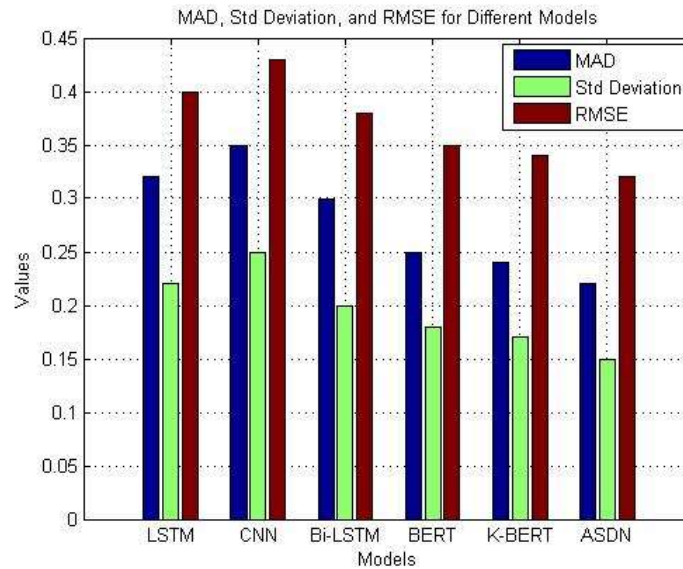


Figure 7: Comparison of MAD, SD and RMSE

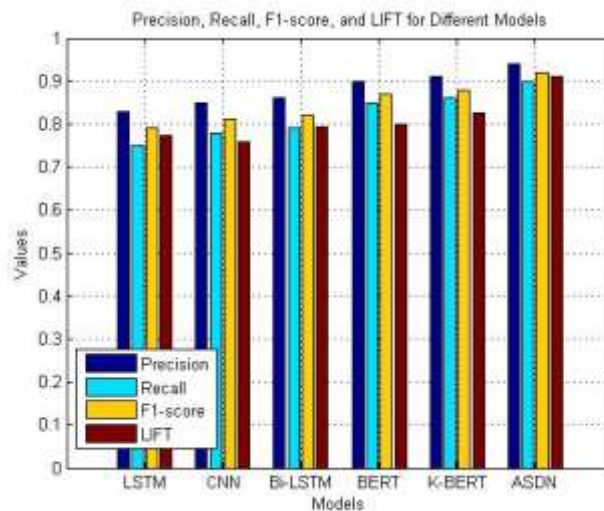


Figure 8: Comparison of Precision, Recall, F1-Score and LIFT values

From the above Table 3 and fig 7 and 8, MAD is lower for models like K-BERT (0.24) and Proposed model (0.22), indicating better average accuracy in predictions compared to the other models. Sentiment Ratio showcases the distribution of sentiments in the dataset for each model. Standard Deviation is lower for models like K-BERT (0.17) and proposed model (0.15), indicating less variability in predictions. RMSE is lower for models like K-BERT and Proposed model, indicating better accuracy in predicting continuous values. Precision, Recall, and F1-score are higher for models like K-BERT (0.91, 0.86, and 0.88) and proposed model (0.94, 0.80, and 0.92) respectively which indicating better performance in correctly identifying positive and negative sentiments. The

Bonferroni correction involves dividing the original significance level (assumed 0.05) by the number of comparisons conducted. If comparing multiple models against each other, this correction adjusts the significance threshold to account for multiple comparisons.

5. Conclusion

The development and application of ASDNs for detecting positive and negative deviations in cross-domain product reviews culminate in a conclusion that underscores their significance and potential impact in sentiment analysis across diverse platforms. ASDNs represent a pivotal advancement in sentiment analysis methodologies, showcasing their adaptability in handling the intricate challenges posed by cross-domain product reviews. Their ability to dynamically adapt to varying linguistic styles and user behaviors across platforms empowers them to capture nuanced sentiment shifts comprehensively. The utilization of ASDNs results in enhanced robustness and accuracy in detecting positive and negative deviations within cross-domain reviews. Their integration of controlled stochasticity allows for precise differentiation between subtle sentiment variations, leading to more refined and accurate sentiment analysis outcomes. ASDNs have the potential to significantly impact consumer decision-making processes by offering deeper insights into sentiment dynamics across e-commerce sites, social media platforms, and forums. Future endeavors could focus on optimizing ASDN architectures, exploring additional adaptive strategies, and enhancing interpretability for broader applicability and understanding. Additionally, investigating methods to process and analyze multimodal reviews (e.g., text, images, videos) by extending ASDNs to effectively incorporate and learn from diverse data types. This could lead to a more comprehensive understanding of sentiment and deviations present in reviews.

Funding: This research received no external funding.

Conflicts of Interest: The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

References

- [1] Valerio, C., William, L., & Noemier, Q. The impact of social media on E-Commerce decision making process. *International Journal of Technology for Business (IJTB)*, 1(1), 1-9, 2019.
- [2] Zhang, F., Hao, X., Chao, J., & Yuan, S. Label propagation-based approach for detecting review spammer groups on e-commerce websites. *Knowledge-Based Systems*, 193, 105520. ACM., 2020.
- [3] Yang, L., Li, Y., Wang, J., & Sherratt, R. S. Sentiment analysis for E-commerce product reviews in Chinese based on sentiment lexicon and deep learning. *IEEE access*, 8, 23522-23530. IEEE, 2020.
- [4] Bhattacharyya, S., & Bose, I. S-commerce: Influence of Facebook likes on purchases and recommendations on a linked e-commerce site. *Decision Support Systems*, 138, 113383. Elsevier, 2020.
- [5] Andry, J. F., Christianito, K., & Wilujeng, F. R. Using Webqual 4.0 and importance performance analysis to evaluate e-commerce website. *Journal of Information Systems Engineering and Business Intelligence*, 5(1), 23-31, 2019.
- [6] Creevey, D., Coughlan, J., & O'Connor, C. Social media and luxury: A systematic literature review. *International Journal of Management Reviews*, 24(1), 99-129. Wiley, 2022.
- [7] Liu, P., & Teng, F. Probabilistic linguistic TODIM method for selecting products through online product reviews. *Information Sciences*, 485, 441-455. Elsevier, 2019.
- [8] Priansa, D. J., & Suryawardani, B. Effects of E-marketing and social media marketing on E-commerce shopping decisions. *Jurnal Management Indonesia*, 20(1), 76-82, 2020.
- [9] Attar, R. W., Shanmugam, M., & Hajli, N. Investigating the antecedents of e-commerce satisfaction in social commerce context. *British Food Journal*, 123(3), 849-868. Emerald, 2021.
- [10] Munna, M. H., Rifat, M. R. I., & Badrudduza, A. S. M. (2020, December). Sentiment analysis and product review classification in e-commerce platform. In *2020 23rd International Conference on Computer and Information Technology (ICCIT)* (pp. 1-6). IEEE, 2020.
- [11] Kumar, A., Gopal, R. D., Shankar, R., & Tan, K. H. Fraudulent review detection model focusing on emotional expressions and explicit aspects: investigating the potential of feature engineering. *Decision Support Systems*, 155, 113728. Elsevier, 2022.
- [12] Shen, R. P., Liu, D., & Shen, H. S. Detecting review manipulation from behavior deviation: A deep learning approach. *Electronic Commerce Research and Applications*, 101283. Taylor and Francis, 2023.
- [13] Marwat, M. I., Khan, J. A., Alshehri, D. M. D., Ali, M. A., Ali, H., & Assam, M. Sentiment analysis of product reviews to identify deceptive rating information in social media: a SentiDeceptive approach. *KSII Transactions on Internet and Information Systems (TIIS)*, 16(3), 830-860, 2022.

- [14] Habbat, N., Anoun, H., & Hassouni, L. Combination of GRU and CNN deep learning models for sentiment analysis on French customer reviews using XLNet model. *IEEE Engineering Management Review*, 51(1), 41-51. IEEE, 2022.
- [15] Alsubari, S. N., Deshmukh, S. N., Alqarni, A. A., Alsharif, N., Aldhyani, T. H., Alsaade, F. W., & Khalaf, O. I. Data analytics for the identification of fake reviews using supervised learning. *Computers, Materials & Continua*, 70(2), 3189-3204, 2022.
- [16] Chang, Y. C., Ku, C. H., & Le Nguyen, D. D. Predicting aspect-based sentiment using deep learning and information visualization: The impact of COVID-19 on the airline industry. *Information & Management*, 59(2), 103587. Elsevier, 2022.
- [17] Saeed, R. M., Rady, S., & Gharib, T. F. An ensemble approach for spam detection in Arabic opinion texts. *Journal of King Saud University-Computer and Information Sciences*, 34(1), 1407-1416. ACM, 2022.
- [18] Guo, J. Deep learning approach to text analysis for human emotion detection from big data. *Journal of Intelligent Systems*, 31(1), 113-126, 2022.
- [19] Saumya, S., & Singh, J. P. Detection of spam reviews: a sentiment analysis approach. *Csi Transactions on ICT*, 6(2), 137-148. Springer, 2018.
- [20] Sahut, J. M., & Hajek, P. Mining behavioural and sentiment-dependent linguistic patterns from restaurant reviews for fake review detection, 2022.
- [21] Budhi, G. S., Chiong, R., Pranata, I., & Hu, Z. Using machine learning to predict the sentiment of online reviews: a new framework for comparative analysis. *Archives of Computational Methods in Engineering*, 28, 2543-2566. Springer, 2021.
- [22] Hajek, P., Barushka, A., & Munk, M. Fake consumer review detection using deep neural networks integrating word embeddings and emotion mining. *Neural Computing and Applications*, 32, 17259-17274. Springer, 2020.
- [23] Li, L., Goh, T. T., & Jin, D. How textual quality of online reviews affect classification performance: a case of deep learning sentiment analysis. *Neural Computing and Applications*, 32, 4387-4415. Springer, 2020.
- [24] Wang, J., Kan, H., Meng, F., Mu, Q., Shi, G., & Xiao, X. Fake review detection based on multiple feature fusion and rolling collaborative training. *IEEE access*, 8, 182625-182639. IEEE, 2020.
- [25] Zhao, L., Song, K., Sun, C., Zhang, Q., Huang, X., & Liu, X. (2019, May). Review response generation in e-commerce platforms with external product information. In *The world wide web conference* (pp. 2425-2435). ACM, 2019.
- [26] Ramesh, R., Jeyakarthic, M. Predictive Analytics for Steel Market Using Improved Random Forest, *Webology*, 18(5), 4824-4834, 2021.
- [27] Veeramanikandan, V., Jeyakarthic, M. Parameter-Tuned Deep Learning Model for Credit Risk Assessment and Scoring Applications, *Recent Advances in Computer Science and Communications*, 14(9), 2958-2968, 2021.
- [28] Thirumalairaj, A., Jeyakarthic, M. Perimeter Intrusion Detection with Multi-Layer Perception using Quantum Classifier, *Fourth International Conference on Inventive Systems and Control (ICISC)*, 348-352. IEEE, 2020.