

Dynamic Label Adjusted and Key Term Based Product Review Analysis Framework for Weakly Labelled Data

V.Uma Devi¹ and Dr. Vallinayagi V²

¹Research Scholar, Register Number: 182212621620010, ²Associate Professor & Head,

¹Sri Sarada College for Women, Affiliated to Manonmaniam Sundaranar University, Abishekapatti, Tirunelveli.

²Department of Computer Science, Sri Sarada College for Women, Affiliated to Manonmaniam Sundaranar University, Abishekapatti, Tirunelveli.

¹Professoruma@gmail.com, ²vallinayagimahesh@gmail.com

ABSTRACT

Sentiment analysis is the rising and the important field carried out on social network. Product reviews are very much important and one of the major factor which is considered by the new buyers. For upcoming buyers, these reviews help in making decisions. The key challenge lies in considering the weakly labelled data. The main objective of this work is to make an effective prediction of the product review sentiment in weakly labelled data. This system consists of two kinds of model to strengthen the review sentiments process with the help of RNN and CNN. First model (KT_CNN) is constructed with the help of Yake based Key Terms learning through Convolutional Network added with Blended Sentiment Polarity. The second model (DLA_RNN) is designed with polarity distance based dynamic label adjusted data with Long Short Term Memory network. The deep features extracted from the both models are concatenated to represent the sentiment of the product review. The proposed (KT_CNN_DLA_RNN) method archives the F1-Score up to 88.71% for the Amazon Data Products dataset.

1. Introduction

Nowadays, the booming of buying or shopping of goods or services over the internet has been increased. After purchasing, buyers show special interest in writing or reviewing comments about their online shopping in social networks. The comments are very much valuable for other buyers. With the help of these resources, decision making process is done by an outlook customer. It is tremendously essential to have a justifiable commencement of how customers notice the products. Customers communicate their view through the comments and the products. Customers' positive opinions direct to the success of a business and negative opinions probably foremost to its downfall [14]. Social media websites, namely LinkedIn, Instagram, Twitter, YouTube, Facebook and e-commerce websites, namely Amazon, Flipkart, Snapdeal, Myntra, and Jabong etc. are being commonly utilized to converse perspectives efficiently. Mainly, e-commerce companies have offered people with numerous opportunities to explore online content [22]. Handing over a positive or a negative sentiment to the reviews can assist companies recognize their peoples and also facilitate peoples to make superior decisions [20]. In the mean time, the number of reviews also develop speedily. At the same time, severe information overload problem also occurs.

Sentiment analysis is considered as a protracted standing research matter. Earlier, sentiment classification methods usually divided into two categories, namely lexicon-based methods and machine learning methods [3]. In Lexicon-based methods [1, 4], first a sentiment lexicon is constructed with help of opinion words. Then, classification rules are designed based on emerging opinion words and prior syntactic knowledge. Here, substantial efforts are required in lexicon construction and rule design. The demerit is, it can compact only with implicit opinions in an ad-hoc way and not with weakly labelled data [5]. Then the machine learning method came into emerged [6]. Following machine learning feature extraction also came into emerged [7]. Currently, deep learning has also emerged as an effective way for solving sentiment classification problems [3]. The purpose of sentiment analysis is not narrow down to specific purpose such as product, sociology, marketing, advertising and movie reviews [24]. It also have utilized in other areas, namely news, sport, election, politics, etc. For example, sentiment analysis is also utilized to categorize people's opinions about a definite nominee or political party in online political debates [23].

The key challenge is how to accurately sentiment analysis the weakly labelled data. In this work, a novel framework has been proposed for review sentences sentiment classification. Document level and sentence level sentiments cannot present adequate data that is vital for decision making [21]. The framework treats review ratings as weak labels. It is concluded by considering sentiment on the aspect, whether it is positive, negative or neutral or has a rating usually ranges from 1 to 5. For example, with ratings 5-stars scale, the user has commented the product as poor, ratings below 3-stars as the good product. Hence, here the contradiction takes place that rating is 5, but the product is poor. These contradictions are called weakly labelled data.

2. Related Works

A lot of research works are going on product review sentiment analysis. Product review sentiment classification is performed using the following methods, namely syntax-based techniques[27], dictionary-based techniques [28], phrase-level techniques [29], frequency-based techniques [30], supervised machine learning techniques [31], hybrid techniques [32] and unsupervised machine learning techniques [33].

Ding et al [1] have proposed an opinion mining technique which is the holistic lexicon-based approach for classifying reviews. Zhu et al [2] have applied an Aspect-based opinion polling technique in customer reviews for sentiment analysis. Pang et al [6] have performed the sentiment classification work by applying NaiveBayes technique. Naive Bayes is one of the most popular machine learning algorithms. Dave et al [8] have used n-grams for Opinion extraction and semantic classification of product reviews. Mullen et al [9] have applied Part-of-speech (POS) information and syntactic relations for extracting features. Sentiment classification is performed by applying Support Vector Machine (SVM). Ziyu Guan et al [11] have considered ratings as weak supervision signals. Here, high level representation is collected from the general sentiment distribution of sentences through rating information. On top of the embedding layer, a sentiment layer is added. Saiful [13] has applied Recurrent Neural Network(RNN) for sentiment classification of Bengali text. BiLSTM technique is also combined with RNN to improve accuracy. Chen Long et al [15] have proposed weakly-supervised deep learning framework to deal with the sentiment classification. Chenghua Lin et al [16] have proposed a new probabilistic modelling method named Joint Sentiment Topic (JST). This model identifies sentiment and topic concurrently from text which depends on Latent Dirichlet Allocation (LDA). This technique also provides indeed coherent and informative results.

Quanzeng You et al [17] have applied Convolutional Neural Network (CNN) for image sentiment classification. Here, a paragraph vector is trained for analysing the textual sentiment classification. Zhang Min [18] has applied a mechanism called Weakly Supervised Mechanism (WSM) to pre-train the parameters of the proposed model. In the weakly labelled data, labelled data is used for the fine-tune initialised parameters. The effect of noise data is decreased on the initial consequences. Here, CNN and Bi-directional Long Short-term Memory (Bi-LSTM) is concatenated with WSM to form WSM-CNN-LSTM for the sentiment classification. Akbar et al [14] have proposed Aspect-Based Sentiment Analysis (ABSA) technique for the market products to analyse customer opinions. It has two main tasks, namely Aspect Sentiment Classification (ASC) and Aspect Extraction (AE). Hence, analyzing customer opinions in a review is a difficult task. Here a deep language model called BERT is also used. Deep language model is supported by two simple modules called Parallel Aggregation and Hierarchical Aggregation. Then, Conditional Random Fields(CRFs) is also used for the sequence labelling task process. Rakesh Kumar and Shashi Shekhar [19] have applied high level representation learning in sentiment classification. This high level representation captures sentiment distribution from sentences. Sentiment classification is performed using efficient deep learning method. In deep learning framework, classification layer is added on top of embedding layer. Here, labelled sentences are utilized for supervised fine tuning.

Waqar Muhammad et al [20] have compared machine learning methods with lexicon based methods. Here, five machine learning techniques is compared, namely Logistic Regression (LR), Naive Bayes (NB), Random Forests (RF), Support Vector Machines (SVM) and K-Nearest Neighbor (K-NN). Then three lexicon based techniques is compared, namely SenticNet, Happiness and SentiStrength. In absence of reviews of a particular website, labelled product reviews from other websites can be efficiently utilized. The supervised techniques are trained to accomplish an equivalent performance to the unsupervised lexicon based techniques. This approach also benefits by enveloping all of the product reviews which the lexicon based approaches be unsuccessful to do so. The results depict that reviews from Amazon are the easiest to classify, pursued by reviews from Reevo, and Facebook reviews are the hardest to classify. Ameen Banjar et al [22] have proposed a method called Aspect Based Sentiment Analysis - Polarity Estimation of customer Reviews (ABSA-PER). This proposed method has three steps, namely data pre-processing, Aspect Co-occurrence Calculation (CAC) and polarity estimation.

Kia Dashtipour et al [23] have classified the subject's sentiment as, namely positive, negative, and neutral. This work contains deep-learning-driven with context-aware and Persian sentiment analysis method. Here, CNN and LSTM are applied which are the two efficient deep learning methods.

3. The Proposed Work

In the proposed architecture we used in the input of the model is product review data with rating. This proposed approach combines the strength of two models with different feature extraction on review statements with key terms. Both two kinds of models are presented in two sections; the overall block diagram of the proposed system is shown in the figure 1.

Section 31. Clearly describe the feature extraction using Key Term based CNN (KT_CNN) model and subsequent section scribes the Dynamic Label Adjusted RNN (DLA_RNN) model.

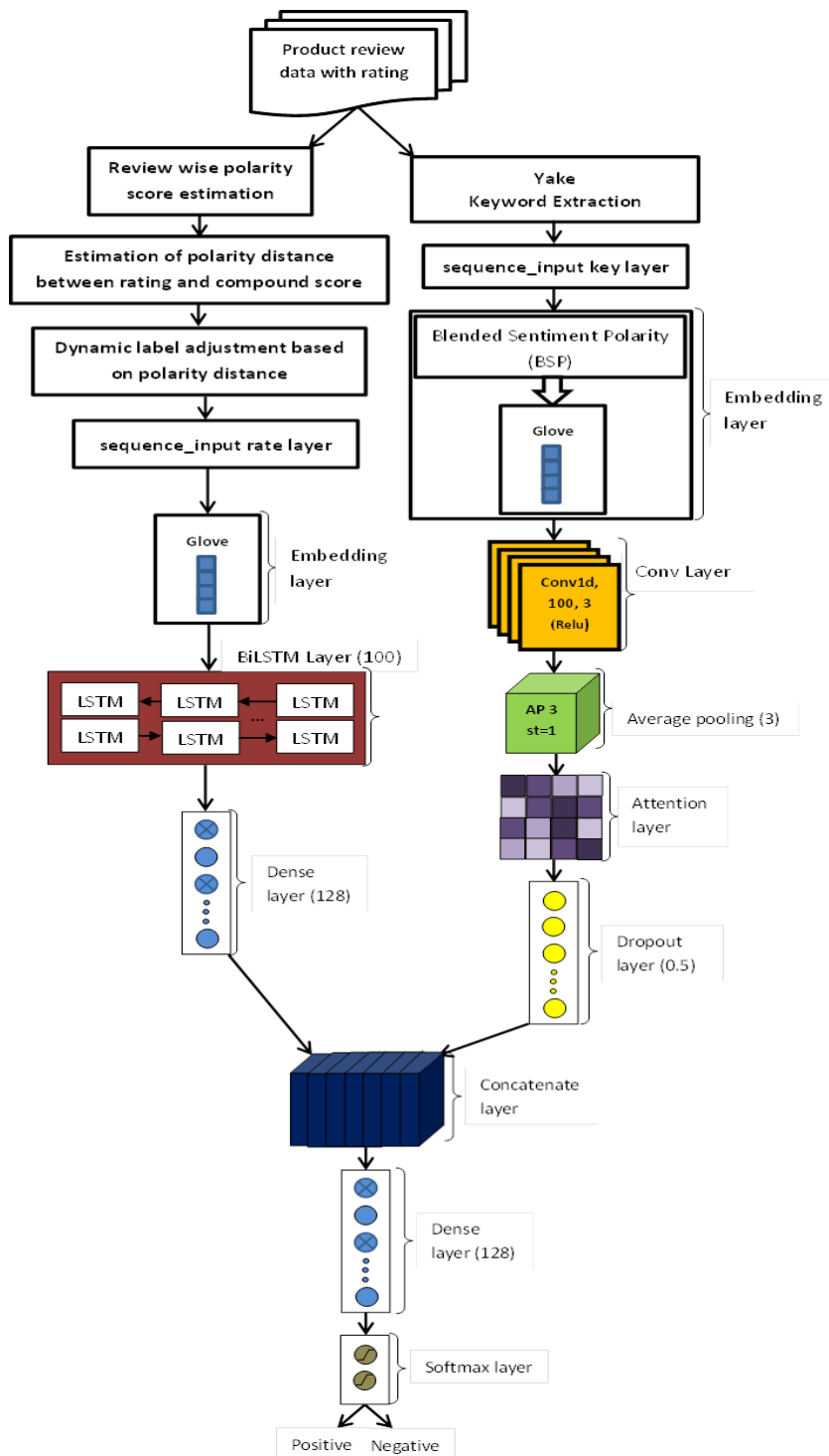


Figure 1. Overall Architecture of the Proposed Dynamic Label Adjusted and Key Term Based Product Review Analysis

3.1. Key Term Based CNN Model (KT_CNN)

This branch initially extracts the key terms from the review with the help of Yake model and those extracted key words are arranged in the order of occurrence in the sentence and given as input of the CNN model with embedding.

3.1.1. Yet Another Keyword Extractor (Yake)

YAKE is called a light-weight unsupervised automatic keyword extraction technique. This technique respites on text statistical features extracted from single documents. The most important keywords of a text are also chosen. YAKE is also called in terms of Term Frequency Independent (TFI). TFI means a probable keyword with no conditions are assigned in the sentence frequency and in the minimum frequency.

YAKE [34] has five main procedures to be followed. First text pre-processing and candidate term identification, second feature extraction, third computing term score, fourth n-gram generation and computing candidate keyword score. Then, finally data deduplication and ranking. In the first stage, pre-processes are performed in the document into a machine-readable format in order to recognize potential candidate terms. Here, first sentences are divided into chunks when punctuation is available. Then, each chunk is divided into tokens. Each token is then changed to lowercase and tagged with proper delimiters. This stage is a significant and vital step to recognize better candidate terms.

The second stage acquires a list of individual terms as input and represents them by a set of five statistical features [36]. They are Y_{LR} that estimates the number of various terms that occur to the left side and right side of the candidate word. Y_{Pos} estimates the values on those words which occur at the start of the sentence itself. Y_{case} replicate the casing aspect of a word. $Y_{difsent}$ quantifies how frequently a candidate word emerges within various sentences. Y_{freq} how frequently the word occurs. Then, all these five statistical features are applied for the estimation of the $Y(w)$ score for each term as shown in the eq.(1).

$$Y(w) = \frac{Y_{LR} \times Y_{Pos}}{Y_{case} + \frac{Y_{freq}}{Y_{LR}} + \frac{Y_{difsent}}{Y_{LR}}} \quad (1)$$

In the third stage, the obtained features are heuristically combined into a single score as represented in the eq.(1). In the fourth stage, the candidate keywords are generated by assigning scores depending on their importance [35]. The sliding window concept of 3-grams is applied for generating a contiguous sequence of 1, 2 and 3-gram candidate keywords. Similar keywords are combined by applying deduplication distance similarity measure in the fifth stage. The list of final keywords is then arranged by their relevance scores. Every candidate keyword will then be allocated with a final $Y(kw)$, which has the smaller score with more meaningful keyword.

$$Y(kw) = \frac{\nabla_{w \in kw} Y(w)}{TF(kw) \times (1 + \sum_{w \in kw} Y(w))} \quad (2)$$

Where $Y(kw)$ is the score of the candidate keyword. At last the system will yield a list of important keywords, formed by 1, 2, 3-grams. Yake series-independent computerized keyword extractor is provided keywords to the input layer [39].

3.1.2. Blended Polarity Bending (BSP)

The output of the YAKE technique is given as input to the embedding layer. The embedding layer is combined with the Glove and the Blended polarity bending concept proposed in the paper [37].

VADER [12] is a rule-based sentiment analysis tool and lexicon analyzer. It is particularly accustomed to sentiments uttered in public media. VADER also utilizes a collection of catalogue of lexical features from a sentiment lexicon. VADER sentimental classification depends on a dictionary which maps lexical characters to emotion intensities called as sentiment scores. The sentiment score of a text is calculated by summing up the intensity of every word in the text available.

VADER is very much intellectual sufficient to recognize the fundamental situation. This also identifies with the emphasis of capitalization and punctuation. According to semantic point of reference, sentiments are commonly labelled as neither positive nor negative. VADER informs on the subject of the positivity and the negativity score. Along with explains positive or negative a sentiment appears. VADER's sentiment lexicon is available [25].

The TextBlob consist of a library for handing out the textual data. Its important feature is sentiment analysis. Based on the polarity of the sentiments, the values positive, negative and neutral are estimated. In the TextBlob, the output will have only one value, positive, negative and neutral [40].

As described in [37] the score for each polarity from Vader and TextBlob is combine with proper weight and added with the Glove Embedding coefficients.

3.1.3. GloVe Embeddings

GloVe Embeddings are a kind of phrase embedding that encode the co-incidence chance ratio among phrases as vector differences. GloVe makes use of a weighted least squares goal J that minimizes the distinction among the dot manufactured from the vectors of phrases and the logarithm in their range of co-occurrences:

$$J = \sum_{i,j=1}^V f(X_{ij})(w_i^T \tilde{w}_j + b_i + \tilde{b}_j - \log X_{ij})^2 \quad (4)$$

where w_i and b_i are the word vector and bias respectively of word i , $\tilde{w}_j$ and $\tilde{b}_j$ are the context word vector and bias respectively of word k , X_{ij} is the number of times word i occurs in the context of word j , and f is a weighting function that assigns lower weights to rare and frequent co-occurrences [47].

3.1.4. Convolutional Neural Network (CNN)

Convolutional Neural Network (CNN) is applied with Natural Language Processing (NLP) to extract features in embedding vectors of a specific sentence [38]. CNN is applied to extract valuable features, namely relationships and phrases between words that are nearer together in a specific sentence. Hence CNN is used in NLP text classification. CNN usually consider a sentence of word vectors. CNN generates a phrase vector for all sub phrases and also grammatically correct phrases.

This NLP CNN is made up of 1D convolutional and pooling layer. This pooling layer is provided as input to a series of dense layers. This method also facilitates by reducing the dimensionality of the text and performs as a summary of sorts which is then provided.

Let us consider two 1D vectors, namely X and Y . X represents the primary vector. Y represents the corresponding filter applied. The convolution between X and Y is estimated at a value n is shown using Eq.(5).

$$(X * Y)[n] = \sum_{m=-M}^M X[n - m]Y[m] \quad (5)$$

Let us consider word-vectors as $w_i \in V^k$. Then, the concatenated word vectors of a n -word sentence is $w_{i:n} = w_1 + w_2 + \dots + w_n$. Here, the filter v is a vector and it is represented as shown in Eq. (6).

$$Z_i = f(v^T w_{i:i/h-1} + b) \quad (6)$$

Where $Z = [Z_1, Z_2, \dots, Z_{n-h+1}] \in \mathbb{R}^{n-h+1}$.

The output of CNN is provided as input to the max-pooling layer. The output of the max-pooling layer represented as shown in Eq. (7).

$$Z = \max \{Z\} \quad (7)$$

3.1.5. Average Pooling

The concept of common or suggest for pooling and extracting the features, this is the primary convolution- based deep neural network. As shown in Fig. 3, an average pooling layer plays down-sampling through dividing the input into rectangular pooling areas and computing the common values of every region [40].

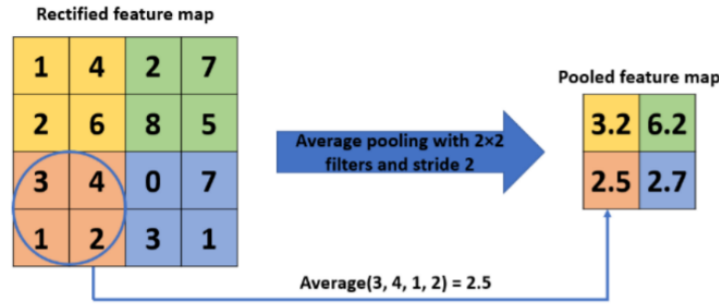


Figure 3Average Pooling Layer Process

3.1.6. Attention:

a) Word Attention

Word Attention Not all words contribute equally to the representation of the sentence meaning. The attention mechanism is used to extract such words that are important to the meaning of the sentence and aggregate the representation of those informative words to form a sentence vector. Specifically,

$$u_{it} = \tanh (W_w h_{it} + b_w)$$

$$\alpha_{it} = \frac{\exp \left(u_{it}^T u_w\right)}{\sum_t \exp \left(u_{it}^T u_w\right)}$$

$$s_i = \sum_t \alpha_{it} h_{it} \quad (8)$$

ument classification:

$$p = \operatorname{softmax}\left(W_c^v + b_c\right)$$

We use the negative log likelihood of the correct labels as training loss:

$$L = -\sum_d \log p_{dj},$$

where j is the label of document d

b) Sentence Attention

Sentence Attention to reward sentences that are clues to correctly classify a review, we again use attention mechanism and introduce a sentence level context vector u_s and use the vector to measure the importance of the sentences.

This yields

$$u_i = \tanh \left(W_s h_i + b_s\right)$$

$$\alpha_i = \frac{\exp \left(u_i^T u_s\right)}{\sum_i \exp \left(u_i^T u_s\right)}$$

$$v = \sum_i \alpha_i h_i \quad (9)$$

where v is the review vector that summarizes all the information of sentences in a review.

3.1.7 Dropout Layer

Dropout layer is used to eliminate the memorization of the network. The working principle of this method is performed by random deletion of some nodes of the network. The dropout rate is 0.5.

3.2. Polarity based Dynamic Label Adjusted Model

3.2.1. Review Polarity Score Estimation

This model focused on the prediction process of the weakly labelled data in an effective form. This module finds the polarity score of each review with the help of VADER. The fetched polarity score falls in eight kinds of category or in eight range as shown in the figure 2. In this work those values are considered in five modes such as most negative, negative, neutral, positive, most positive. Such as -4 to -2 as most negative, -2 to 0 as negative, 0 to 1 as neutral, 1 to 2 as positive and 2 to 4 as most positive.

3.2.2. Polarity Distance based Dynamic Label Adjustment

In the case of weekly labelled data, there is an occurrence of contradiction between rating and review statement, such as the rating may be 5 star but the review is in the form of negative meaning. This issue is handled by this dynamic label adjustment based on polarity distance. This section first assign a value for the review based on the polarity score as described in 3.2.1. Where it assigns most negative as 1, negative as 2, neutral as 3, negative as 4 and positive as 5. Then it finds the absolute distance between that polarity label and the rating of the review with the help of the following equation (10).

$$\text{Polarity}_{\text{dist}} = || \text{Polarity_Score_Value}(\text{review}) - \text{Rating}(\text{review}) || \quad (10)$$

The estimated polarity distance is less than or equal to three means the label is adjusted based on the polarity score as shown in the equation 11.

$$\text{if } \text{Polarity}_{\text{dist}} \leq 3 \begin{cases} \text{label} \leftarrow \text{Positive} & \text{if polarity score} == 4 \text{ or } 5 \\ \text{label} \leftarrow \text{Negative} & \text{if polarity score} == 1 \text{ or } 2 \\ \text{label} \leftarrow \text{Neutral} & \text{if polarity score} == 3 \end{cases} \quad (11)$$

Based on the above process the label is dynamically adjusted based on the value of polarity distance. For example if the polarity score of the review is 1 and the rating is 4 means there is a contradiction between the rating and the actual representation of the review, hence this approach change the label of the review from the positive to negative if it is originally stated as positive.

3.2.3. Bilstm

In Recurrent Neural Networks (RNNs), BiLSTM is a popular method [19]. Hidden states are reorganized gradually by BiLSTM, by utilizing earlier hidden state of final stage and present contribution. This method learns gating system to become skilled at dependencies of long term which is most generally applied. Memory cell is also separately preserved. It is constant equivalence of memory circuit. An interior state of memory cell controls the reset, the write and the read operations by forget, input and output gates.

3.2.4. Dense

A dense layer contains neurons of network layer. [ML repository (2019)]. Dense layer receives the input from all layers using the neurons. The layer has a weight matrix w , a bias vector b , and the activations of previous layer a . The dense layer is 128. The dense layer is defined by,

$$y = (a(x.w) + b) \quad (12)$$

where a is the element-wise argument, w is a weights matrix and bias is a bias vector created by the layer.

The dense layer and the dropout layer concatenate using concatenate layer. Next concatenate layer input to dense layer. The dense layer is 128. And final used softmax layer are used to detect the class labels of the input images

4. Experiment Result

In this section, we present the empirical evaluation of WDE on reviews collected from Amazon.com.

Table 1: Statistics of the labeled dataset.

	Positive	Negative	Total
Subjective	3750	2024	5774
Objective	1860	4120	5980
Total	5610	6144	11754

4.1. Dataset description

Weakly labelled data

We collected Amazon customer reviews of 3 domains: digital cameras, cell phones and laptops. All unlabeled reviews were extracted from the Amazon data product dataset [McAuley et al., 2015][41]. The labeled dataset was randomly split into training set (80%), and test set (20%) and we maintain the proportion as shown in Table 1.

4.2. Evaluation metrics

The performance metrics used to evaluate the classification results are Precision, Recall, F-measure and Accuracy.

Those metrics are computed based on the values of true positive (TP), false positive (FP), true negative (TN) and false negative (FN) assigned classes [75].

i. Accuracy

Accuracy is the most intuitive performance measure and it is simply a ratio of correctly predicted observation to the total observations. It is given by:

$$Accuracy = \frac{TP+TN}{TP+FP+FN+TN} \quad (13)$$

ii. Precision

Precision is the ratio of correctly predicted positive observations to the total predicted positive observations. It is given by:

$$Precision = \frac{TP}{TP+FP} \quad (14)$$

iii. Recall (Sensitivity)

Recall is the ratio of correctly predicted positive observations to the all observations in actual class. It is given by:

$$Recall = \frac{TP}{TP+FN} \quad (15)$$

iv. F1 score

F1 Score is the weighted average of Precision and Recall. Therefore, this score takes both false positives and false negatives into account. It is given by:

$$F1Score = \frac{2*(Recall*Precision)}{(Recall+Precision)} \quad (16)$$

Results and Discussion

Table 1: Comparison of Accuracy and F1-Score

Method	Accuracy	Macro-F1
Lexicon	0.722	0.721
SVM	0.818	0.818
NBSVM	0.826	0.825
SSWE	0.835	0.834
SentiWV	0.808	0.807
MemNet	0.839	0.838
CNN-rand	0.847	0.847
CNN-rand11m	0.849	0.848
CNN-weak	0.771	0.771
LSTM-rand	0.845	0.845
LSTM-rand11m	0.85	0.849
WDE-CNN	0.877	0.876

WDE-LSTM	0.879	0.879
CNN	80.98	79.815
LSTM	85.51	81.8
LSTM with VADER	90.88	84.85
BSP	91.65	85.95
KT_CNN	92.81	86.97
DLA_RNN	93.58	87.92
KT_CNN_DLA_RNN	94.47	88.71

Table 2: Comparison of Accuracy, Precision, Recall and F1-Score

	CNN	LSTM	LSTM with VADER	BSP	KT_CNN	DLA_RNN	KT_CNN_DLA_RNN
Accuracy	80.98	85.51	90.88	91.65	92.28	93.78	94.57
Precision	80.325	81.116	85.12	85.49	86.76	87.58	88.27
Recall	79.312	82.496	84.59	86.43	87.19	88.27	89.17
F1-score	79.81	81.80	84.85	85.95	86.97	87.92	88.71

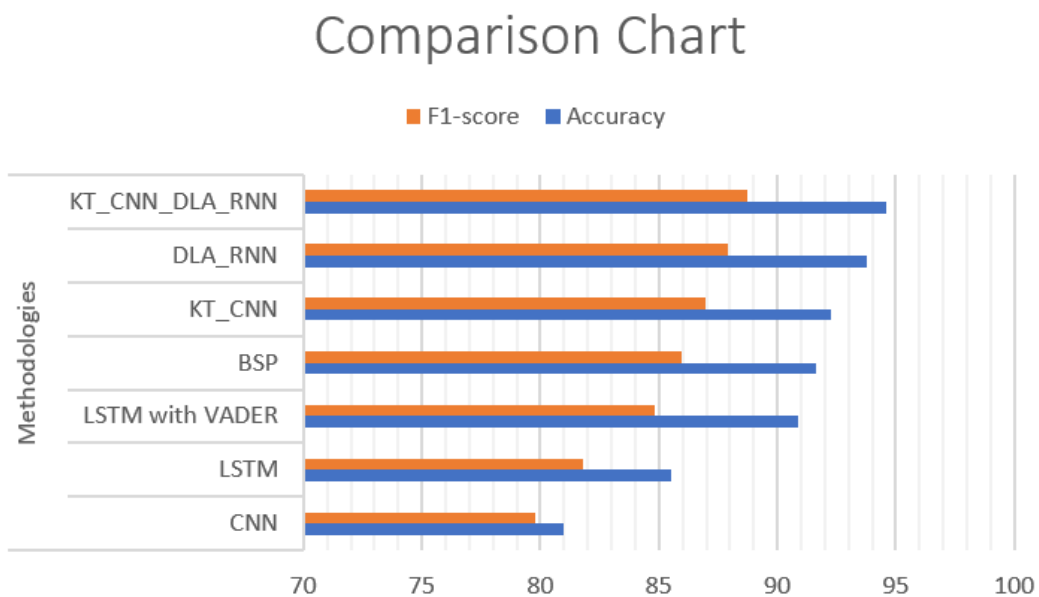


Fig. 2: Comparison of accuracy and F1-Score between various methods

Table 2 illustrates the accuracy comparison. In this table, proposed3 method provide 2.92 higher accuracy than the proposed method. Proposed3 method provide 3.69 better accuracy than LSTM with VADER method. Proposed3 method provide 9.06 higher accuracy than LSTM Proposed3 method provide 13.59 better accuracy than CNN method. Fig. 2 indicates the bar diagram of various methods of Accuracy. According to the diagram Proposed3 method provide accuracy comparing to other methods.

Table 2 illustrates the F1-Score comparison. In this table, proposed3 method provide 2.76 higher F1-Score than the proposed method. Proposed3 method provide 3.86 better F1-Score than LSTM with VADER method. Proposed3 method provide 6.91 higher F1-Score than LSTM Proposed3 method provide 8.902 better F1-Score than CNN method. Fig. 5 indicates the bar diagram of various methods of F1-Score. According to the diagram Proposed3 method provide F1-Score comparing to other methods.

5. Conclusion

In this paper, we presented an effective framework especially designed for sentiment analysis on product review with weakly labelled data. One of key challenge in product review sentiment analysis is there a contradiction between the review statement and its corresponding rating. This issue is handled by this paper with the help of two different models namely Key Terms based CNN (KT_CNN) and Dynamic Label Adjusted criteria with LSTM (DLA_RNN). From the experimental result and analysis on the Amazon Data Product dataset, it clearly shows that the proposed model attains significant results compared to state of art works. The dynamic label adjustment using polarity score deviation helps the system to improve the performance along with role of key terms in review. The proposed (KT_CNN_DLA_RNN) method archives the F1-Score up to 88.71% for the Amazon Data Products dataset. The proposed approach achieves 6.91 higher F1-Score than LSTM and 8.902 than the traditional CNN method.

References

1. X. Ding, B. Liu, and P. S. Yu, 'A holistic lexicon-based approach to opinion mining', In WSDM, pp. 231–240, 2008.
2. J. Zhu, H. Wang, M. Zhu, B. K. Tsou, and M. Ma, 'Aspect-based opinion polling from customer reviews', IEEE Transactions on Affective Computing, Vol. 2, No.1, pp. 37–49, 2011.
3. Wei Zhao, Ziyu Guan, Long Chen, Xiaofei He, Deng Cai, Beidou Wang and Quan Wang, 'Weakly-supervised Deep Embedding for Product Review Sentiment Analysis', IEEE Transactions on Knowledge and Data Engineering, 2017.
4. P. D. Turney, 'Thumbs up or thumbs down?: semantic orientation applied to unsupervised classification of reviews', In ACL, pages 417–424, 2002.
5. L. Zhang and B. Liu, 'Identifying noun product features that imply opinions', In ACL, pp. 575–580, 2011.
6. B. Pang, L. Lee, and S. Vaithyanathan, 'Thumbs up?: sentiment classification using machine learning techniques. In EMNLP, pp 79–86, 2002.
7. B. Pang and L. Lee, 'Opinion mining and sentiment analysis. Foundations and trends in information retrieval', Vol. 2, No. 1-2, pp. 1–135, 2008.
8. K. Dave, S. Lawrence, and D. M. Pennock, 'Mining the peanut gallery: Opinion extraction and semantic classification of product reviews', In WWW, pp. 519–528, 2003.
9. T. Mullen and N. Collier, 'Sentiment analysis using support vector machines with diverse information sources. In EMNLP, Vol. 4, pp. 412–418, 2004.
10. P.D. Turney and M.L. Littman, 'Unsupervised learning of semantic orientation from a hundred-billion-word corpus', Technical Report ERB-1094, National Research Council Canada, Institute for Information Technology, 2002.
11. Ziyu Guan, Long Chen, Wei Zhao, Yi Zheng, Shulong Tan and Deng Cai 'Weakly-Supervised Deep Learning for Customer Review Sentiment Classification', Proceedings of the Twenty-Fifth International Joint Conference on Artificial Intelligence (IJCAI-16), 2016.

12. C.J. Hutto and Eric Gilbert, 'VADER: A Parsimonious Rule-based Model for Sentiment Analysis of Social Media Text', Association for the Advancement of Artificial Intelligence, (www.aaai.org), 2014.
13. Md Saiful Islam, 'A Deep Recurrent Neural Network with BiLSTM model for Sentiment Classification', IEEE, International Conference on Bangla Speech and Language Processing(ICBSLP), 2018.
14. Akbar Karimi, Leonardo Rossi and Andrea Prati, 'Improving best performance for aspect-based sentiment analysis', arXiv:2010.11731v2, 2021.
15. Chen Long, Guan Ziyu, He Jinhong and Peng Jinye, 'A Survey on Sentiment Classification', Journal of Computer Research and Development, Vol. 24, No. 56, pp. 1150-1170, 2017.
16. Chenghua Lin; Yulan He; Richard Everson and Stefan Ruger, 'Weakly Supervised Joint Sentiment-Topic Detection from Text', IEEE Transactions on Knowledge and Data Engineering, Vol. 24, No. 6, 2012.
17. Quanzeng You, Jiebo Luo, Hailin Jin and Jianchao Yang, 'Joint Visual-Textual Sentiment Analysis with Deep Neural Networks', Proceedings of the 23rd ACM international conference on Multimedia, pp. 1071–1074, 2015.
18. Zhang Min, 'Drugs Reviews Sentiment Analysis using Weakly Supervised Model', IEEE International Conference on Artificial Intelligence and Computer Applications (ICAICA), 2019.
19. Rakesh Kumar and Shashi Shekhar, 'Product Review Sentiment Analysis Using Weakly-supervised Deep Embedding Using CNN and MLSTM Techniques', International Journal of Control and Automation, Vol. 12, No. 6, 2019, pp. 108-116.
20. Waqar Muhammad, Maria Mushtaq, Khurum Nazir Junejo and Muhammad Yaseen Khan, 'Sentiment Analysis of Product Reviews in the Absence of Labelled Data using Supervised Learning Approaches', Malaysian Journal of Computer Science, Vol. 33, No. 2, pp. 118-132, 2020.
21. Zohreh madhoushi, Abdul razak hamdan and Suhaila zainudin, 'Aspect-Based Sentiment Analysis Methods in Recent Years', Asia-Pacific Journal of Information Technology and Multimedia, Vol. 7 No. 2, pp. 79 -96, 2019.
22. Ameen Banjar, Zohair Ahmed, Ali Daud, Rabeeh Ayaz Abbasi and Hussain Dawood, 'Aspect-Based Sentiment Analysis for Polarity Estimation of Customer Reviews on Twitter', Computers, Materials & Continua, Vol.67, No.2, pp. 2203- 2225, 2021.
23. Kia Dashtipour, Mandar Gogate, Ahsan Adeel, Hadi Larijani and Amir Hussain, 'Sentiment Analysis of Persian Movie Reviews Using Deep Learning', Entropy, 2021, 23, 596.
24. C.J. Hutto and Eric Gilbert, 'VADER: A Parsimonious Rule-based Model for Sentiment Analysis of Social Media Text', Conference: Proceedings of the Eighth International AAAI Conference on Weblogs and Social Media At: Ann Arbor, MI, 2015.
25. <http://comp.social.gatech.edu/papers/>
26. Zeinab Rajabi, Mohammad Reza Valavi and Maryam Hourali 'A Context-Based Disambiguation Model for Sentiment Concepts Using a Bag-of-Concepts Approach', Cognitive Computation, Vol. 12, pp.1299–1312, 2020.
27. K. Ravi and V. Ravi, A survey on opinion mining and sentiment analysis: Tasks, approaches and applications', Knowledge Based System, Vol. 89, pp. 14–46, 2015.
28. C. M. Faisal, A. Daud, F. Imran and S. Rho, 'A novel framework for social web forums thread ranking based on semantics and post quality features', The Journal of Supercomputing, Vol. 72, no. 11, pp. 4276–4295, 2016.
29. H. U. Khan, A. Daud, U. Ishfaq, T. Amjad, N. Aljohani et al., 'Modelling to identify inertial bloggers in the blogosphere: A survey', Computers in Human Behavior, vol. 68, pp. 64–82, 2017.
30. B. Kama, M. Ozturk, P. Karagoz, I. H. Toroslu and O. Ozay, 'A web search enhanced feature extraction method for aspect-based sentiment analysis for Turkish informal texts', Big Data Analytics and Knowledge Discovery, LNCS, vol. 9829, pp. 225–238, 2016.

31. H. H. Do, P. W. C. Prasad, A. Maag and A. Alsadoon, 'Deep learning for aspect-based sentiment analysis: A comparative review', *Expert Systems with Applications*, vol. 118, pp. 272–299, 2019.
32. P. Takala, P. Malo, A. Sinha and O. Ahlgren, 'Gold-standard for topic-specific sentiment analysis of economic texts', in *Proceedings of the Ninth Int. Conf. on Language Resources and Evaluation*, Reykjavik, Iceland, pp. 2152–2157, 2014.
33. K. Schouten and F. Frasincar, 'Survey on aspect-level sentiment analysis', *IEEE Transactions on Knowledge and Data Engineering*, Vol. 28, no. 3, pp. 813–830, 2016.
34. Ricardo Campos, Vítor Mangaravite, Arian Pasquali, Alípio Jorge, Célia Nunes and Adam Jatowt, 'YAKE! Keyword extraction from single documents using multiple local features', *Elsevier, Information Sciences*, Vol. 509, pp. 257–289, 2020.
35. Ricardo Campos, Vítor Mangaravite, Arian Pasquali, Alípio Mário Jorge, Célia Nunes and Adam Jatowt, 'YAKE! Collection-Independent Automatic Keyword Extractor', *Springer, ECIR 2018, LNCS 10772*, pp. 806–810, 2018.
36. Eirini Papagiannopoulou and Grigorios Tsoumakas, 'A Review of Keyphrase Extraction', *arXiv:1905.05044v2*, 2019.
37. V. Uma Devi and Dr. Vallinayagi V, 'Blended Sentiment Polarity (BSP) for Product Review Analysis', *International Journal of Advanced Science and Technology*, Vol. 29, No. 05, pp. 4456 - 4466.
38. Wei Wang and Jianxun Gang, "Application of Convolutional Neural Network in Natural Language Processing", 2018 *International Conference on Information Systems and Computer Aided Education (ICISCAE)*.
39. Campos R., Mangaravite V., Pasquali A., Jorge A.M., Nunes C., Jatowt A. (2018) YAKE! Collection-Independent Automatic Keyword Extractor. In: Pasi G., Piwowarski B., Azzopardi L., Hanbury A. (eds) *Advances in Information Retrieval. ECIR 2018. Lecture Notes in Computer Science*, vol 10772. Springer, Cham. https://doi.org/10.1007/978-3-319-76941-7_80
40. Gholamalinezhad, H., & Khosravi, H. (2020). Pooling Methods in Deep Neural Networks, a Review. *arXiv preprint arXiv:2009.07485*.
41. [McAuley et al., 2015] Julian McAuley, Rahul Pandey, and Jure Leskovec. Inferring networks of substitutable and complementary products. In *SIGKDD*, pages 785–794, 2015.