

A Study Of The Pattern Of Migration And The Role & Political Participation Of In Migrants In Unorganized Sector

Dr. Jyoti Tripathi

B 1/120, Assi, Varanasi -221005
vnsjyoti@gmail.com

How to cite this article: Jyoti Tripathi (2024) A Study Of The Pattern Of Migration And The Role & Political Participation Of In Migrants In Unorganized Sector. *Library Progress International*, 44(3), 7339-7348.

ABSTRACT : Researchers from academia and business are becoming interested in social media big data mining due to the exponential growth of social media data on the internet. Sentiment analysis of news events is an important area within this discipline that has received a lot of interest and could enable a number of real-world applications, such as government public opinion monitoring and website recommendation systems. However, when applied to previous volumes of social media big data, existing sentiment analysis approaches present scalability issues because they generally rely on traditional emotion lexicons or supervised algorithms. As a result, we suggest a novel method for sentiment analysis of news stories. In particular, using words and emoticons related to a news event from social media, we create

INDEX TERMS : Text mining, sentiment computing, emotion classification, social media big data.

INTRODUCTION

Social media data has become an important phenomena as a result of the development of the internet and the exponential increase in people's propensity to discuss and communicate their thoughts and interests online. This historical pool of social media data is made up of sites like Facebook, Twitter, and WeChat; in academic discourse, these platforms are sometimes referred to as big data. Various elements play a role in this description:

****volume**:** The sheer scale of social media data is immense. For instance, the average number of microblogs reporting on a news event exceeds 100,000.

****variety**:** Social media content is diverse, comprising various elements such as words and emoticons, making it heterogeneous in nature.

****velocity**:** Social media data is characterized by its high dynamism. Daily influxes of data, such as over 500TB on Facebook alone, underscore its rapid generation and dissemination.

****value**:** The richness of information embedded within social media data presents a fertile ground for researchers in the big data domain. Moreover, it serves as a potent resource for corporations and governments alike, enabling them to derive insights crucial for making informed decisions and formulating global strategies. Newsevents, integral components of the social media big data landscape, encompass stories that unfold within society or on the web and garner attention from numerous web pages [1]. Take, for instance, the MH 370 incident, which sparked extensive online discussions across various social media platforms following its occurrence. Engaging with and discussing news events have become a routine part of daily life, thereby presenting a treasure trove of valuable information for corporations and governments alike. Within the realm of news event analysis, sentiment computing emerges as one of the most challenging tasks. Its objective is to discern the emotional tone embedded within textual content, such as micro blogs, authored by users. However, existing methods predominantly categorize texts into two broad categories, thus leaving room for refinement and innovation in sentiment analysis methodologies.

Since the texts in social media are almost short texts and each short text can be seen as a sentence, we do not make the distinction on text/short text/sentence in this paper [2].

CATEGORIES : Both good and negative, which defies the notion that public opinion is nuanced. These days, a few academics often compute the multidimensional emotion in texts. Emotion is a six-dimensional vector that includes: joy, lope, surprise, fear, sadness, and anger, according to generally used emotion [3] (More specifics can be found in Section 3). For instance, a user's microblog about MH370 might have the following emotion: $\langle 0.001, 0.001, 0.9, 0.2, 0.096, 0.002 \rangle$, indicating that the user is surprised and somewhat afraid about the incident. There are numerous contexts in which it is useful to ascertain the sentiment surrounding a news event. Take, for instance, 1) public opinion tracking. Since there is a lot of information about a news event available online, decision-makers may find it easier to comprehend public concerns by 1) precisely pinpointing the material with apparent emotiveness (i.e., rage); and 2) monitoring product feedback. A company typically wants to know about the interests and concerns of its users over a new product it has published. This company may be able to assess the effectiveness of its previous product design and make decisions about future designs if it can precisely locate the information that appears to be emotive (such as surprise and rage). Some state-of-the-art works for affective computing for documents have been produced in these literatures (see Section 2 for more information). However, because to the following shortcomings, current works cannot be immediately adapted for news events:

Simply emphasizing the use of adjectives and adverbs to categorize emotions, ignoring the semantics of documents; 2) There are two extremes: one heavily depends on conventional emotion thesaurus, while the other completely rejects it.

Our novel concept is to use the semantics of a news event to help with sentiment computation. The word emotion computation task, which can be divided into two steps: word emotion computation using a word emotion association network and word emotion refining through a standard sentiment thesaurus, is the primary component of the news event sentiment computing task. The initial step in computing word emotion involves creating a Word Emotion Association Network (WEAN), which serves as the foundation for computing text and word emotion. WEAN is designed to jointly capture a word's semantic and emotional aspects. We predict that words with comparable semantic associations will elicit similar emotions more frequently. An iterative procedure with convergence proof is created to maximize the linkages in emotional weight assignments. Following this procedure, we can determine the initial sentiment conveyed by words, although they might not align with the body of current information. For instance, the word "happy" ought to convey a great deal of joy, but after going through several iterations, it can end up conveying the wrong feeling. Therefore, we aim to have created a method in the second phase of word emotion computation that will allow us to refine the original word emotion by using the standard emotion thesaurus, which is a common previous knowledge.

The two main contributions of this study are as follows:

1) Word emotion is computed using a Word Emotion Association Network, which captures both the semantic and emotional aspects of a news event; 2) Word emotion is computed using a Standard Emotion Thesaurus, which increases word emotion accuracy.

This is how the rest of the paper is structured. The relevant literature is reviewed in Section 2, and the suggested techniques are offered in Section 3, which is divided into two sections: word emotion refining using a conventional sentiment thesaurus and word emotion computation using a word emotion association network. The experimental setup and findings on the real-world data are detailed in Section 4. The study's results and conclusions are outlined in Section 5 along with some suggestions for future research.

RELATED WORK

A]The following are this study's two primary contributions:

- 1) Word emotion is calculated using a Standard Emotion Thesaurus, which improves word emotion accuracy
- 2) Word emotion is computed using a Word Emotion Association Network, which captures both the semantic and emotional components of a news event. The remainder of the paper is organized in this manner. In Section 2, the pertinent literature is examined; in Section 3, which is separated into two sections, the suggested

methodologies are provided: word emotion computation using a word emotion association network, and word emotion refining using a conventional sentiment thesaurus. Section 4 provides a full description of the experimental setup and results using real-world data. Section 5 provides an overview of the study's findings and conclusions

B] SENTIMENT COMPUTATION

Sentiment computation initially nearly attempts to handle lengthy texts. Since long text sentiment computation is essentially a text classification problem, any currently available text classification techniques can be applied to this problem. Positive and negative sentiment were distinguished between in film reviews by Pangetalusing machine learning techniques to analyze sentiment. They used three different machine learning techniques in their studies, and the best results were obtained with a bag of words as features and SPM as the classifier. Subsequently, they [4] conducted studies on sentiment classification to demonstrate that unigram performed better with fewer training corpora and that n-gram became increasingly significant with an increase in training corpora. Mayur et al. [5] calculated sentiment of a sentence using location features and word features from comments. Pansy Nandwani et al. [6] used lexical features, such as emotion terms and sentences pertaining to emotion orientation, to calculate the emotion score for each article. Gamon and Aueemployed a semi-supervised technique to teach sentiment using a limited number of words with no labeling at all. Kim et al. [7] used phrase fragments to identify the emotion of internet users exclusively in the second person. Baoetic [8] suggested the dEmotion-Topic Model (ETM), an opic model whose dataset is Sinamicroblog, to calculate the sentiment of brief messages. Every attitude is comprised of multiple themes with varying weights, as shown by the addition of an intermediately erinto LDA. Additionally, topic models called SentimentLatentTopicModel (SLTM) and Multi-labelSuperpisedTopicModel (MSTM) were presented by Nirmal et al [9] to calculate the sentiment of a microblog. Rao et al. trained their models using words and demoticons found in microblogs. The results indicated that MSTM and SLTM all performed better than ETM. Dandan Jiang et al [10] built an emotion vocabulary based on ten degrees of emotion, according to a movie review. Emotion ratings and emoticons are included in the Emotionlexicon. The four categories of emoticon roles—assuagement, conpersion, addition, and emphasis—were used to determine the sentiment score of a single tweet. In order to learn semantic representations of users and products for document lepel sentiment categorization. Since social media has grown, people prefer to express their thoughts and feelings through brief texts as opposed to lengthy documents. Short text sentiment computation is not compatible with the methods used for handling lengthy texts.

****C. SHORT TEXT SENTIMENT COMPUTATION****

Significant progress has been made in recent years on social media sites including Facebook, Twitter, and Sina Microblog. Numerous research have attempted to analyze Twitter data and classify it into positive and negative attitudes by employing machine learning algorithms, linguistic factors, and emoticons. For example, Biraj and Jaibir [12] created a machine learning classifier to divide Twitter data into positive and negative categories using features including bigrams, unigrams, and POS tags. Emojis, POS, and syntactic characteristics were used by Pak and Paroubek [13] to train a classifier that classified tweets into positive, negative, and neutral attitudes. For sentiment classification, Zhang et al [14] used a supervised method, whereas Davidov et al [15] presented a supervised method utilizing features including hashtags, smileys, punctuations, and common patterns. Additionally, Bakliwal et al [16] eleven characteristics to train an SPM classifier, and they assessed their methodology with several datasets. Jiang et al [17] made an analysis of sentiment in short texts took into account contextual information and target terms from tweets. Hu et al [18] used social media's emotional cues to compute text sentiment. In-depth research on Twitter sentiment analysis was done by Saif et al [19], who addressed problems including data sparseness and stopword effectiveness. Senti Circles, a lexicon-based method to improve Twitter sentiment analysis, was also introduced by them. It's important to recognize that public sentiment is multifaceted and nuanced, even though the majority of current methodologies divide texts into simplistic positive and negative sentiment categories. For more accurate readings, this intricacy should be taken into account in future sentiment analysis algorithms.

PROPOSED METHOD

The two primary components of affective computing for news events—word emotion computation using the

Word Emotion Association Network (Section 3.1) and word emotion refining using a conventional sentiment thesaurus (Section 3.2)—are designed to identify the feelings that news event participants express in their microblogs. The basis for text emotion computation is word emotions. We will examine each component in more detail in the sections that follow. Table 1 contains a list of the notations used in this paper.

TABLE NO. 1: NOTATIONS IN THIS TABLE

Symbol	Meaning in this paper
e^k	emotion of k-th dimension
$w_{emoticon_x}^{e^k}$	the e^k weight of $emoticon_i$
$w_{i,j}^{e^k}$	the e^k weight between word i and j
$v_{k_n}^{e^k}$	the value of word k_n on e^k
$u_{k_n}^{e^k}$	sum of transmitting values.
$F(x)$	normalized function
ϵ	regulation parameter
w_B	words both in basic emoticon pocabulary and word sentiment association network
δ	the maximum error
Δ	the step length
ew_i	the emotion of word w_i
e_s	the emotion of text δ

WORD EMOTION COMPUTATION THROUGH WORD EMOTION ASSOCIATION NETWORK

Views, ideas, and attitudes toward things, people, events, or subjects can be shared whenever and wherever one chooses on social media. Users' emotions can be seen in their viewpoints, beliefs, and attitudes. It's referred to as social media text emotion. Numerous studies on social media text emotion have been conducted up to this point, and several multidimensional emotion kinds have been reported from various angles. We employ the multidimensional emotion classification system developed by, which categorizes emotions into six groups: fear, surprise, anger, joy, sadness, and lope. Our proposal is to calculate the sentiment of words in a news article. We must calculate the word's emotional state at a given moment since, given a news occurrence, a word in different stages may hope different feeling. Because of the vast amount of words. Fortunately, emoticons may be used to express feelings in words on microblogs, and $\omega ek = \omega ek$ can be utilized to graph the location of emoticons bi-directionally come before, after, or inside a sentence, depending on where it is. The default emoticons are few in number, and it's obvious what each one signifies. We can compute word emotion by using emoticons. For this study, we developed a questionnaire to get opinions on the six-dimensional emoticon mood. This poll was completed by twenty undergraduate students who wished to discuss their experiences using emoticons in their microblogs. The sentiment of a word can now be calculated at this point.

Word Emotion Association Network (WEAN) Definition 1:WEAN is built using words and their association rules, much like association link networks (ALN)[20]. This can be expressed as $WEAN = \langle N, W \rangle$.

Nouns, verbs, adjectives, adverbs, quantifiers, prepositions, and so forth are examples of word parts of speech. terms that elicit a stronger feeling than others should be nouns, verbs, adjectives, and cyber terms. Adverbs frequently amplify or lessen the emotional content of verbs or adjectives. Word emotion computation, thus, applies to nouns, verbs, adjectives, and cyberwords. For the Microblogs where e^k is not null, we build a word emotion association network (WEAN) in Definition 1.

The next step is to compute the word emotion of e^k after WEAN is created. Word emotion can be impacted by the various scales and intentions of word circumstances. The more powerful the word feeling, the larger the scale and the stronger the intention. In light of these, we have

Quantity Assumption : A word's strength in WEAN is determined by the number of associations it has with other words that share the same emotion. Conversely, the stronger the emotion, the fewer links there are; every node represents a word that occurs in emotion e^k , like *lope*. We now know the emotion values of nodes N1 through N6, so we can assume they have the same value. The words that we require to calculate their emotions are T1 and T2. T2 has four terms associated with it, compared to T1's two. Quantity assumption states that T2 has a higher emotion in *lope* than T1.

$$w_{i,j}^{e^k} = \sum_{x=1}^M w_{emotion_x}^{e^k} \quad (2)$$

Intensity Assumption : Via their links, words can influence other words that are close by. Words that are connected to the term k_n can have varying degrees of e^k . A word's ability to transfer emotion to other words increases with its strength. Therefore, the emotion e^k of a word will also be powerful if it ties to numerous other words whose e^k are strong. N9 and N10. T3 and T4 each aim to have two links pointing to them. A recursive algorithm determines a word's emotion for every word that is associated with it. A term with more links will generally have a higher emotion value since linked words have higher emotion values. Algorithm 1 provides a summary of the entire process and we want to compute the emotion value of T_3 and T_4 . According to intensity assumption, T_4 has stronger emotion than T_3 in emotion e^k . Not all of the words in WEAN have this emotion. Actually, the e^k of many words is almost 0, which can be ignored. These words in WEAN are mainly those whose degrees are small. So, when computing word emotion, we should set these words' emotion of e^k null, and the other word emotion value is between 0 and 1. The number of words whose emotion should be set null can be decided automatically.

Based on **Quantity Assumption**, **Intensity Assumption** and the weights of the links in Definition 1, the emotion of word k_n can be computed as

$$u_{k_n}^{e^k(j)} = \sum_{i=1}^N v_{k_n}^{e^k} (j-1) * w_{i,n}^{e^k} \quad (3)$$

and

$$P_{kek(j)} = f(u_{kek(j)})/E \quad (4)$$

Where $P_{kek(j-1)}$ denotes the $(j-1)$ -th computation value of word k . Word emotion value plus link weight

reflects the intensity assumption. It means that word k transmits.

The summation reflects quantity assumption. $u_{ek}(j)$ is the storage that stores the sum of transmitting values. $f(x) = \frac{1}{1+e^{-x}}$ is the normalized function, which normalizes the words' value between 0 and 1. E is a regulation parameter that regulates most word emotion to 0. When computing word emotion, we randomly give each word an initial value. The computation can yield rational results only after a finite number of iterations. Eq. (3) and (4) is for a single word. For all of the words in WEAN, Eq. (5) and (6) work.

$$(u_{k_1}^{e^k(j)}, \dots, u_{k_n}^{e^k(j)}) = (v_{k_1}^{e^k(j-1)}, \dots, v_{k_n}^{e^k(j-1)}) \cdot \begin{pmatrix} w_{1,1}^{e^k} & \dots & w_{1,N}^{e^k} \\ \vdots & \ddots & \vdots \\ w_{N,1}^{e^k} & \dots & w_{N,N}^{e^k} \end{pmatrix} \quad (5)$$

and

$$(v_{k_1}^{e^k(j)}, \dots, v_{k_n}^{e^k(j)}) = \frac{(u_{k_1}^{e^k(j)}, \dots, u_{k_n}^{e^k(j)})}{E} \quad (6)$$

It is evident from the above process that a word's feeling is influenced by the words that it hopes to link with. This word is liked by the links. A recursive algorithm is used to determine a word's emotion for every word that is linked to it. A term with more linkages will often have a higher emotion value due to the linked words' emotional worth.

The whole procedure is summarized in Algorithm 1.

Algorithm 1 Word Emotion Computation

Require: WEAN and w^e

Ensure: wordEmotions^{Lj}

```

1:randomlyinitialization
2:tag=1, it=1;
3:foreach  $k^{th}$  dimension of emotion do
4:    fortag=1 do
5:        tag=0;
6:        foreach word  $x_i$  in WEAN do
7:            Update  $p^{e^k}(it)$  by Eq.(5) and (6);
8:        end for
9:        foreach word  $x_i$  in WEAN do
10:            if  $|p^{e^k}(it) - p^{e^k}(it-1)| > 0$  then
tag=1;
12:            endif
13:        end for
14:    end for
15:endfor

```

WORD EMOTION REFINEMENT THROUGH STANDARD SENTIMENT THESAURUS

The algorithm for word emotion calculation suggested above only takes into account the emoticons' emotions in microblogs. Early on in an event, there aren't many microblogs about it, especially the ones with emoticons. The WEAN, which is composed of only a few words and regulations, lacks credibility. Currently, there will be little precision when computing the word emotion with WEAN. The terms in the conventional sentiment thesaurus, however, imply clear and consistent feelings that hardly ever fluctuate in response to various news occurrences.

Therefore, we propose to refine the word emotion derived by Algorithm 1 according to the terms in standard sentiment

thesaurus. For the words in both standard sentiment thesaurus of e^k and WEAN, we define them as w_B . Emotion e^k of these words is approximate to 1, actually. However, the words emotion computed by Algorithm 1 may be less than 1. At this time, we need to refine them. We set the maximum error forwards in w_B is δ , and every step is Δ , which is a positive number. When the emotion of one word in w_B computed by iteration is much less than 1 (i.e., $|1 - p^{e^k}| > \delta$), we add the weight of

$$\Delta \cdot |1 - p^{e^k}|$$

The links with the word w_B . After that, we do the word emotion computation against using Algorithm 1. When all the words (in w_B) emotion values are close to 1, we stop the word emotion computation. The whole procedure is summarized in

Texts are composed by words. Words' emotion reflect texts' emotion indirectly. After computing word emotion, we

Algorithm 2 Word Emotion Refinement

ij

Require: WEAN, w^e , δ , Δ and standard sentiment thesaurus

Ensure: Word emotions

```

1. Word emotions obtained by Algorithm 1;
2:tag=1;
3:foreach  $k-th$  dimension of emotion do
4:    fortag=1 do
5:        tag=0;
6:        foreach word/both in standard sentiment the-saurus and WEAN do
7:            if  $|1 - p^{e^k}| > \delta$  then
8:                tag=1;

```

```

9:      for eachlinkweight $w^k$ linkedwithword $i$ do
10:          $w^k = w^k + \Delta \cdot \frac{1 - P^k}{\delta}$ 
11:      end for
12:    endif
13:  end for
14:  if tag=1 then
15:    foreach word  $x_i$  in WEANdo
16:       Update  $P^k$  by Eq.(5) and (6);
17:    end for
18:  end if
19: end for
20: endfor

```

can obtain text six-dimensional emotion by adding the six-dimensional emotion of words that in the text, respectively.

$$e_s = \sum_{i:w_i \in S} e w_i \quad (7)$$

where $e_{xi} = \langle e_{lope}, e_{joy}, e_{anger}, e_{sad}, e_{fear}, e_{surprise} \rangle$ is a six-dimensional vector, which means the six-dimensional emotion of x_i ; e_s is also a six-dimensional vector, which means the six-dimensional emotion of text.

IP. EXPERIMENTS We carry out an experiment in this section to confirm the accuracy and potency of the suggested approach.

DATASETS

We use the news event The Malaysia Airlines MH370 on March 8, 2014 as a dataset to assess our approach. About this news occurrence, all social media information comes from SinaMicroblog. 48,396 microblogs from March 8, 2014, to April 6, 2014 are included. 7,346 hope emoticons were found in the texts over the course of the thirty days, and they can be used to calculate word emotions.

B. EVALUATION METRIC

First, an evaluation metric has to be created in order to assess the performance of the suggested method on the text of affective computing. Since microblogs are so brief, our theory is that the sentiment of a microblog should be consistent with its emoticons. As a result, the following emotions calculated from emoticons might be considered the suggested method's benchmarks:

$$e^d_b = \frac{1}{N_e^d} \sum_{i=1}^{N_e^d} e_{e_i}$$

where N^d is the number of emoticons in microblog d ;
 e_i is one emoticon in microblog d ; e_{e_i} is the emotion vector of an emoticon;
 e^d is the emotion vector of microblog d . e^d could be seen as the benchmark of our proposed method, so we can evaluate the performance of the proposed method through comparing e^d with the value from our method.

C. EXPERIMENTAL SETUP

We must separate the data into two sets: training and testing, as the suggested method and the evaluation metric both call for microblogs with emoticons. The steps involved are described as follows:

****Select Microblogswith Emoticons****: Identify microblogs related to the news event containing emoticons.

****Data Splitting****: Divide the microblogswith emoticons into two parts: a training set and a test set.

****Training Data Preparation****: Input the training set, excluding microblogswith emoticons, into our proposed method.

****Model Training****: Train the proposed method using the training data.

****Testing Data Prediction****: Apply the well-trained model to predict the sentiment of the microblogs in the test set.

****Benchmark Calculation****: Compute the benchmarks for the test set using Eq. 8.

****Performance Evaluation****: Evaluate the performance of the proposed method.

Our objective is to evaluate the efficacy and precision of our suggested approach in determining sentiment for news occurrences by utilising emoticon-rich social media data.

D. EXPERIMENTAL RESULTS AND DISCUSSIONS

As demonstrated by Algorithm 1, initialising the variables is necessary for the computation of word emotions. Is this algorithm sensitive to the various initialisations? is a question of nature. No, is the response. Here, we'll provide empirical evidence to demonstrate that Algorithm 1 is insensitive to various initialisations. We execute Algorithm 1 multiple times with different initialisations sent to it. Every iteration's values are noted and shown in Figure 5. As demonstrated in Fig. 5, after a number of iterations (about nine in this case), multiple beginning points will result in the same convergent value for the same word. This finding raises our level of trust in Algorithm 1's output, even when random initialisation is used. In addition, we plotted the convergent curves in Fig. 6—often using different phrases. Not unexpectedly, the emotion values of all the words are convergent after a number of iterations. The more intriguing finding is that, regardless of initialization—different or same—different words will converge to distinct values. We may therefore conclude that the final emotion ratings for words from Algorithm 1 are decided by their structural parts in the WEAN and not by the initialisation.

We perform eight group experiments to compare the word emotions before and after the refinement in order to demonstrate the efficacy of word emotion refinement in Algorithm 2. Based on Table 2, all eight groups were able to achieve an accuracy of greater than 75%. Meanwhile, the line chart.

TABLE 2 : THE ACCURACY OF TEXT EMOTION COMPUTATION

Group No.	before refinement	after refinement
1	0.785714286	0.787456446
2	0.801709402	0.803418803
3	0.804794521	0.808219178
4	0.7694974	0.776429809
5	0.75	0.755244755
6	0.75257732	0.757731959
7	0.768439108	0.777015437
8	0.803108808	0.80656304
Avg	0.779480106	0.784009928

It shows that, even though the change might not be noticeable right away, the refined emotion is more accurate than it was previously. This is because we believe we have obtained and integrated with the WEAN all microblogs pertaining to Malaysia Airlines flight MH370. We are able to generate findings through integration that are even more accurate than we could have with only refinement. 5,363 words representing the emotions

of the news event The Malaysia Airlines MH370 are plotted here.

Since each word Emotion is a 6-dimensional vector, It can be seen that the emotion *lope* and *sad* of these words are much more stronger than the other four dimensions of emotion because there are more words located in [0.51] in sub figures *lope* and *sad*. So, we can draw the conclusion that people mainly feel *lope* and *sad* on event The Malaysia Airlines MH370.

In order to show the accuracy of our method, we have compared our method that sums word emotions together as the text emotion with benchmark. In Table 3, we have shown six microblog (see Fig. 9 from *The Malaysia Airlines MH370*) examples with their emotion from both our method and the benchmark emotion using developed evaluation metric in Eq.

The result shows that although the short text is not with single emotion, the main dimension's emotion is in keeping with the benchmarks. Therefore, we can draw the conclusion that our method is effective in short text emotion computation for the news event.

The social emotion of a news item at a specific moment is what we refer to as the summation text emotion of all microblogs about it. The way this event has unfolded and the fresh facts that has entered the public discourse are to blame for the emotion's alteration or evolution. It is discovered that the trends of *lope* and *sad* emotions are comparable. The dates are shown on the x-axis. The MH370 airliner lost contact on March 8, 2014, and Malaysia Airlines subsequently confirmed this information. The emergency procedure was initiated by the Chinese government. People's concerns over the MH370 flight will grow over the coming several days, and everyone will be hoping to locate the aircraft. The suspected military coverage of the actual event that caught people's attention peaked on March 13, 2014. A few days later, the glacial pace of search and rescue causes individuals to lose focus. The Prime Minister of Malaysia announced on March 24, 2014, that the aircraft had crashed into the South Indian Ocean, with no survivors. This revelation brought the MH370 situation back into the public eye a few days later. The Malaysian government said on March 27, 2014, that 122 pieces of debris had been discovered in the waters west of Perth. The Australian Maritime Bureau reported that three objects had been seen by the search team. On 2014.04.02, Malaysia Airlines MH370 survey was defined as a criminal investigation. In this process, social emotions have changed as a result of the incident's development, and public opinion has tended to remain steady as the event has progressed.

TABLE 3 : TEXT EMOTION FOR THE SIX MICROBLOGS IN

microblog id	straightforward sum	benchmark
1	< 1.0, 0.1, 0.1, 1.0, 0.2, 0.1 >	< 1.0, 0.0, 0.0, 1.0, 0.0, 0.0 >
2	< 0.8, 0.0, 0.0, 1.0, 0.0, 0.0 >	< 0.0, 0.0, 0.0, 1.0, 0.0, 0.0 >
3	< 1.0, 0.0, 0.0, 1.0, 0.3, 0.0 >	< 1.0, 0.0, 0.0, 1.0, 0.1, 0.0 >
4	< 1.0, 0.1, 0.1, 1.0, 0.1, 0.0 >	< 1.0, 0.0, 0.0, 1.0, 0.1, 0.0 >
5	< 1.0, 0.1, 0.2, 1.0, 0.3, 0.2 >	< 1.0, 0.0, 0.0, 0.9, 0.1, 0.0 >
6	< 1.0, 0.8, 1.0, 1.0, 1.0, 0.9 >	< 0.0, 0.0, 1.0, 0.3, 0.0, 0.1 >

CONCLUSION AND FUTURE STUDY

In this research, we aim to build a novel approach for sentiment computation of news events based on large amounts of social media data. The two steps in the suggested method are word emotion refining using a standard emotion thesaurus and word emotion computation using a word emotion association network. In order to compute word emotion using a word emotion association network, a Word Emotion Association Network (WEAN) has been developed. This network serves as the foundation for computing both word and text emotion. A word emotion computation method based on WEAN has been suggested, and its convergence has been proven, to obtain the initial word emotions through a planned iterative procedure. Additionally, a standard emotion thesaurus has been proposed as a means of combining common previous information into a word emotion refining algorithm to increase accuracy. The experiment's goal also showed how important this improvement is. In order to assess the efficacy of our approach, we have selected the Malaysia Airlines MH370 incident. Our method's accurate computation of news event sentiment is demonstrated by the result. We are interested in including word emotion pattern and emotion distance into text sentiment classification in the future. A single person publishes a microblog, which is no more than 140 words long. It should only aspire to one powerful feeling. And one word, hopefully, can evoke multiple feelings under different conditions. Thus,

word context should be taken into account while calculating text emotion. Moreover, emotion computation is another requirement of brain-like computing. Thus, our methodology can be applied to transfer learning for Brain-like Computing.

REFERENCES

- Sepideh Bazzaz et al. Big data analytics meets social media: A systemic review of techniques, open issues and future directions. 2020. 1 (1) : 1-38.
- Kian Long Tan, Chin Poo and Kian Ming Lim : A survey of Sentiment Analysis : Approaches, Datasets and Future Research. 2023. 13(7) : 1-21.
- Lin Shu et al. A Review of Emotion Recognition Using Physiological Signals. 2018. 18 : 1-41.
- Kalaivani P and Shunmugathan KL. Sentiment Classification of Movie Reviews by Supervised Machine Learning Approaches. 2013. 4(4) : 285-292.
- Mayur Vankhade et al. A Survey of Sentiment Analysis, Methods, Applications and Challenges. 2022. 55 : 5731-5780.
- Pansy Nandwani and Rupali Verma. A Review on Sentiment Analysis and Emotion Detection from Text. 2021. 11 (81) : 1-19.
- Kim Shouten and Flavius Frasinca. Survey on aspect level sentiment analysis. 2015. 1-20.
- Meng Li and Yucheng Shi. Sentiment analysis and prediction model based on Chinese government affairs microblog. 2023. 9 : 1-16.
- Nirmal Varghese Babu and E. Grace Mary Kanaga. Sentiment Analysis in Social Media Data for Depression Detection Using Artificial Intelligence: A Review. 2021. 3 (74) : 1-20.
- Dandan Jiang et al. Sentiment Computing for the news event based on the social media big data. 2017. 5 : 2373-2382.
- Qianwen Ariel XU et al. A systematic review of social media-based sentiment analysis : Emerging trends and challenges. 2022. 3 : 1-16.
- Biraj Lahkar and Jaibir Singh. Twitter text sentiment analysis : A comparative study of unigram and bigram features extractions. 2022. 9(8) : 1-13.
- Pak and P. Paroubek, "Twitter as a corpus for sentiment analysis and opinion mining," in *Proc. LREC*, vol. 10. 2010, pp. 1320–1326.
- L. Zhang, R. Ghosh, M. Dekhil, M. Hsu, and B. Liu, "Combining lexicon-based and learning-based methods for twitter sentiment analysis," HP Lab., Palo Alto, CA, USA, Tech. Rep. HPL-2011-89, 2011.
- D. Davidov, O. Tsur, and A. Rappoport, "Enhanced sentiment learning using twitter hashtags and smileys," in *Proc. 23rd Int. Conf. Comput. Linguistics, Posters*, 2010, pp. 241–249.
- A. Bakliwal, P. Arora, S. Madhappan, N. Kapre, M. Singh, and V. Varma, "Mining sentiments from Tweets," in *Proc. WASSA*, 2012, pp. 11–18.
- L. Jiang, M. Yu, M. Zhou, X. Liu, and T. Zhao, "Target-dependent twitter sentiment classification," in *Proc. 49th Annu. Meeting Assoc. Comput. Linguistics, Human Lang. Technol.*, 2011, pp. 151–160.
- X. Hu, J. Tang, H. Gao, and H. Liu, "Unsupervised sentiment analysis with emotional signals," in *Proc. 22nd Int. Conf. World Wide Web*, 2013, pp. 607–618.
- H. Saif, Y. He, and H. Alani, *Semantic Sentiment Analysis of Twitter*. Berlin, Germany: Springer, 2012.