



Word-level neutrosophic sentiment similarity

Florentin Smarandache^a, Mihaela Colhon^{b,*}, Ștefan Vlăduțescu^c, Xenia Negrea^c

^a Department of Mathematics, University of New Mexico, USA

^b Department of Computer Science, University of Craiova, Romania

^c Faculty of Letters, University of Craiova, Romania

HIGHLIGHTS

- A new word-level similarity measure defined by means of the words' sentiment scores.
- The similarity measure is defined without considering the words' lexical category.
- The resulted scores correctly distance the neutral words from the sentiment words.

ARTICLE INFO

Article history:

Received 16 July 2018

Received in revised form 5 February 2019

Accepted 17 March 2019

Available online 20 March 2019

MSC:

03B52

68T50

Keywords:

Word-level similarity

Neutrosophic sets

Sentiwordnet

Sentiment relatedness

ABSTRACT

In the specialized literature, there are many approaches developed for capturing textual measures: textual similarity, textual readability and textual sentiment. This paper proposes a new sentiment similarity measures between pairs of words using a fuzzy-based approach in which words are considered *single-valued neutrosophic sets*. We build our study with the aid of the lexical resource *SentiWordNet 3.0* as our intended scope is to design a new word-level similarity measure calculated by means of the sentiment scores of the involved words. Our study pays attention to the polysemous words because these words are a real challenge for any application that processes natural language data. After our knowledge, this approach is quite new in the literature and the obtained results give us hope for further investigations.

© 2019 Elsevier B.V. All rights reserved.

1. Introduction

Semantic textual similarity is a measure of the degree of semantic equivalence between some pieces of texts [1]. This measure is exploited in many natural language processing (NLP) tasks, very actual at the present moment, such as paraphrase recognition [2], tweets search [3], image retrieval by caption [4,5], query reformulation [6] or automatic machine translation evaluation [7]. In information retrieval (IR) the user's query is usually expressed by means of a short sequence of words based on which the most similar documents related to the query must be returned to the user.

On the other hand, textual sentiment analysis consists of measuring the attitude or emotional affect of the text. Using this kind of data very actual research fields such as affective computing or sentiment analysis can understand and predict human emotions [8] as their basic tasks are emotion recognition [9,10]

and polarity detection [11–14]. Emotion recognition means to find a set of emotion triggers while polarity detection is usually designed as a binary classifier with “positive” and “negative” outputs [15,16].

In a world full of indeterminacy [17] the reality cannot be drawn only using two colours: “white” and “black” or “positive” and “negative” or “true” and “false” because uncertainty plays a determinant role. Fuzzy set theory has been used in many studies where uncertainty plays a determinant role. Natural language texts contain large amount of uncertain information [18] mainly caused by: 1. the polysemy of same words (for example, the English word “line” has more than 20 distinct senses); 2. the fact that different words can have the same meaning (for example “stomach pain” and “belly ache”); 3. the ambiguities of natural language construction which can happen at many levels of analysis, both syntactic and semantic, which imply different interpretations for the same words or phrases. If we consider also the natural diversity in subjectivity of any natural language utterance, we can conclude that this domain can be regarded as uncertain one.

* Corresponding author.

E-mail addresses: smarand@unm.edu (F. Smarandache), colhon.mihaela@ucv.ro (M. Colhon), vladutescu.stefan@ucv.ro (Ș. Vlăduțescu), xenia.negrea@ucv.ro (X. Negrea).

To deal with large amount of uncertain knowledge, many fuzzy based systems have been developed, but they still remained weak explored in the domain of identifying the sentiment orientation of sentences. The detection of the polarity or subjectivity predictors in written text usually implies to compute the terms grade membership in various pre-defined or computed categories [19,20]. These studies usually require a pre-defined sentiment lexicon for detecting the sentiment words. If this step ends successfully, they have to compute the distance between the identified words and the class centroid in order to measure the fuzzy membership [21–23]. Each membership function is interpreted as the appurtenance degree of the analysed piece of text to a certain sentiment class [24].

These systems could benefit from on a robust word-level similarity component. Most of the existing approaches for determining the semantic similarity between words do not incorporate the words' sentiment information. The present study focuses on the task of measuring the sentiment similarity at a word-level.

Sentiment similarity indicates the similarity of word pairs from their underlying sentiments. In the linguistic literature, sentiment similarity has not received enough attention. In fact, the majority of previous works employed semantic similarity as a measure to also compute the sentiment similarity of word pairs [25,26]. Nevertheless, some works stated that sentiment similarity can reflect better the similarity between sentiment words than semantic similarity measures [27].

Following [28] we consider that the sentiment information is crucial in finding the similarity between two concepts, in particular, between two words. In this assumption, in this study we propose a new sentiment similarity measure between pairs of words using a neutrosophic approach [29–33] and with the aid of the SentiWordNet 3.0 [34] lexical resource. Our intended scope is to suggest a new measure for the sentiment similarity degree of two words which takes into account not only the “positive” and “negative” sentiment labels but also their more refined derivatives such as: “objective”, “weak positive”, “weak negative”, “strong positive” and “strong negative”.

1.1. Justification

An important number of word-level similarity measures were defined using lexico-semantic information. Based on the syntactic category of the involved words we can have a *similarity measures* or a *relatedness measures*. Most similarity measures are computed for words within the same category, usually for nouns and verbs. Still, many similarity approaches consider the semantics and not the lexical category in the process of similarity findings as in the case when the verb “mary” should be found semantically equivalent with nouns such as “wife” or “husband” [1] and not necessarily with another verb.

Corresponding, the relatedness measures are used to compute the similarity degree between words with different categories, e.g. between a noun and a verb such as “tears” and “to cry” [35]. Nevertheless, this restriction is not always obey, as many word similarity measures are developed without paying attention to the syntactic category of the involved words [36]. When defining our proposal we do not differentiate words upon their part of speech as we consider the *sentiment similarity* just as the inverse difference value between the sentiment polarities of two words. Thus, in what follows, the terms similarity and relatedness will be considered equivalent.

There is another important aspect of the proposed measure: it has a symmetric dimension, following thus the key assumption of the most similarity models even if this idea is not universally true, especially when it comes to model human similarity judgments [37]. “Asymmetrical similarity occurs when an object

with many features is judged as less similar to a sparser object than vice versa” [38] such as, for example, when comparing a very frequent word with an infrequent word as “boat” with “dinghy” [37].

The reason we choose a *symmetric* measure to model the proposed word-level similarity measure is determined by two aspects of the study:

1. it treats the words as independent entities, defined only by their SentiWordNet scores and therefore, additional information such as word frequency are not considered
2. by following a neutrosophic approach, the proposed method aggregates all the scores corresponding to all the senses a word can have in a *single-valued neutrosophic set* representation and thus, information about a particular sense are not computed and the words are treated as entities with a single facet

1.2. WordNet

WordNet thesaurus is a collection of nouns, verbs, adjectives and adverbs, being a graph-formed dictionary with a unique organization based on word sense and synonyms [39]. Graph-based structures are widely used in natural language processing applications such as [40,41]. In WordNet structure there are two main forms of word representations: lemma and synset [42]. The synsets are considered “logical groups of cognitive synonyms” or “logical groups of word forms” which are inter-connected by “semantic pointers” with the purpose of describing the semantic relatedness between the connected synsets. These relations were used to find similarity measures between word senses based on the lengths of the relationships between them.

The “net” structure of the WordNet is constructed by means of the lexical or conceptual links differentiated upon the part of speech of the words from the connected synsets. The noun synsets are connected through the “hyperonymy” (and its inverse, “hyponymy”) and the “meronymy” (and its inverse, “holonymy”) relations. The verbs are linked through the “troponym”, “hyponymy” and “entailment” relations. Adjectives point to their antonyms or to the related nouns while adverbs are linked to adjectives through the “pertainym” relation.

1.3. SentiWordNet as a sentiment lexicon

SentiWordNet extends the usability of WordNet to another dimension, by mapping a large number of WordNet synsets to sentiment scores indicating their “positivity”, “negativity” and “objectivity” [42]. Always, the sum of these three values is 1.0.

Because SentiWordNet is built upon the WordNet data, the common problem that is observed at WordNet appears also at SentiWordNet senses: the too fine-grained synsets make hard the distinguishing between the senses of a word. As a direct consequence, the scoring of synsets are even more difficult to predict. The main problem is how much the related synsets and glosses or even the terms of the same synset share or not the same sentiment.

Table 1 presents some sentiment scores examples of the most positive and the most negative words' senses in SentiWordNet [43]. It is important to mention that all the SentiWordNet scores were obtained after weighting 8 classifiers and averaging their classifications [44].

With the construction of this lexical resource, a wide category of tasks, usually in the domain of Opinion Mining (or Sentiment Analysis) started to take shape. Here are three categories of tasks that can be implemented by making usage of the synsets sentiment scores [44]:

Table 1

Example of scores in SentiWordNet [43].

Synsets & sentiment score	Positive score	Negative score	Neutral score
good#1 (0.75, 0, 0.25)	0.75	0	0.25
superb#1 (0.875, 0, 0.125)	0.875	0	0.125
abject#1 (0, 1, 0)	0	1	0
bad#1 (0, 0.625, 0.325)	0	0.625	0.325
unfortunate#1 (0, 0.125, 0.875)	0	0.125	0.875

- *subjectivity–objectivity polarity*: its scope is to determine whether the given text is subjective or objective [11,45];
- *positivity–negativity polarity*: its scope is to determine whether the text is positive or negative on its subject matter [11,46];
- *strength of the positivity–negativity polarity*: its scope is to determine how positive or negative the given text is. More precisely, these tasks have to decide if the opinion expressed by a text is weakly or strongly positive/negative [12,29];
- *extracting opinions from a text*, which firstly implies to determine if the given text includes an opinion or not, and (if it is the case) to determine the author of the opinion, the opinion subject and/or the opinion type [26].

Sentiment analysis was defined for textual content analysis but recent studies perform this kind of analysis on visual content such as images and videos [4]. Performing sentiment analysis on visual content implies to identify the “visual concepts that are strongly related to sentiments” and to label these concepts with few lexical terms (for example, in [4] the authors propose a visual labelling mechanism by means of adjective-noun pairs as usually opinion detection is based on the examination of adjectives in sentences [19]).

This paper is dedicated to the problem of sentiment similarity between pairs of words using a neutrosophic approach in which a word is interpreted as a *single-valued neutrosophic set* [47,48]. At our knowledge, this is the second study that addresses the problem of words sentiment data using neutrosophic concepts. With the intended scope of filling the gap concerning the objectivity aspect of some words, the previous study [49] addresses the problem of the so-called “neutral words” with the aid of neutrosophic measures applied on the words’ sentiment scores.

The study presented in this paper includes and extends the work initiated in [49] as it addresses all types of words, whether sentiment words or objective words. The proposed formalism can be used in any sentiment analysis task as it determines the sentiment polarity of a word by computing its similarity with some seed words (words whose sentiment labels are known or provided). The considered similarity measures can be of great help also for the text similarity techniques that pair the words of the involved texts in order to quantify the degree to which the analysed texts are semantically related [1,50]. In these techniques, pairs of text sequences are aligned based on the similarity measures of their component words.

The remainder of the paper is organized as follows: in the following section we summarize the most recent studies in the domain of similarity measures with focus on the investigated neutrosophic concepts. Section 3 describes the method we designed for constructing a new word-level similarity measure using the sentiment scores of the involved words and applying the neutrosophic theory. In Section 4 the evaluation results are given. The final section sketches the conclusions and the future plan directions.

2. Similarity measures. Related works

There is an important number of works concerning the semantic similarity with different levels of granularity starting from the word-to-word similarity to the document-to-document similarity (important issue for any search engine) [1,35].

Many approaches have been proposed with the intended scope of capturing the semantic similarity between words: Latent Semantic Analysis (LSA) [51], Point-wise Mutual Information (PMI) [52] (for estimate the sentiment orientation) or numerous WordNet based similarity measures. Much attention has recently been given to calculating the similarity of word senses, in support of various natural language learning and processing tasks. One can use the shortest path or the Least Common Subsumer (LCS) depth length algorithm to calculate the distance between the nodes (words) as a measure of similarity between word senses [36,42]. One difficulty here is that some words have different meanings (senses) in different contexts, and thus different scores for each sense.

Such techniques can be applied within a semantic hierarchy, or ontology, such as WordNet. WordNet acts as a thesaurus, in that it groups words together based on their meanings. The semantic distance between words can be estimated as the number of vertices that connect the two words. Another approach makes usage of a large corpus (e.g. Wikipedia) to count the terms that appear close to the words being analysed in order to construct two vectors and compute a distance (e.g. cosine). In this method, the similarity degree between the two entities is given by the cosine value of the angle determined by their vectors representation [53].

The similarity problems are also modelled using concepts from fuzzy set theory and it is our belief (which will be further proved) that neutrosophic theory, that was defined in order to generalize the concepts of classic set and fuzzy set, offers more appropriate tools. Indeed, in a *Neutrosophic Set* the indeterminacy, which is so often encountered in real-life problems such as decision support [54], is quantified explicitly [30,31] as it will be shown in what follows.

2.1. Fuzzy and neutrosophic sets

A *fuzzy set* is built from a reference set called universe of discourse which is never fuzzy. Let us consider U - the universe of discourse. A fuzzy set A over U is defined as:

$$A = \{(x_i, \mu_A(x_i)) \mid x_i \in U\}$$

where $\mu_A(x_i) \in [0, 1]$ represents the membership degree of the element $x_i \in U$ in the set A [55,56].

Now, if we take A be a *intuitionistic fuzzy set* (IFS) in the universe of discourse U , then the set A is defined as [57]:

$$A = \{(x, \mu_A(x), \nu_A(x)) \mid x \in U\}$$

where $\mu_A(x) : U \rightarrow [0, 1]$ is the membership degree and $\nu_A(x) : U \rightarrow [0, 1]$ represents the non-membership degree of the element $x \in U$ in A , with $0 \leq \mu_A(x) + \nu_A(x) \leq 1$.

The concept of *neutrosophic set* A in the universe of discourse U is defined as an object having the form [47]:

$$A = \{ \langle x : t_A(x), i_A(x), f_A(x) \rangle, x \in U \}$$

where the functions $t_A(x), i_A(x), f_A(x) : U \rightarrow [0, 1]$ define respectively the degree of membership, the degree of indeterminacy, and the degree of non-membership of a generic element $x \in U$ to the set A .

If on a neutrosophic set A we impose the following condition on the membership functions $t_A, i_A, f_A : U \rightarrow [0, 1]$:

$$0 \leq t_A + i_A + f_A \leq 3, x \in A$$

then the resulted set $A \subset U$ is called a *single-valued neutrosophic set* [58]. We can also write $x(t_A, i_A, f_A) \in A$.

Corresponding to the notions of neutrosophic set and single-valued neutrosophic set, similar works have been done on graph-theory resulting the notions of neutrosophic graphs [59] and single-valued neutrosophic graphs [60] and on number-theory resulting the concept of neutrosophic numbers and single valued trapezoidal neutrosophic number [61,62].

2.2. Neutrosophic similarity measures

Neutrosophic distance and similarity measures were applied in many scientific fields such as decision making [63,64], pattern recognition [65,66], medical diagnosis [67,68] or market prediction [69].

In this section we enumerate the similarity measures together with their complements – the distance measures, that are applied and then compared in the proposed neutrosophic method for words similarity (see Section 3).

Intuitionistic fuzzy similarity measure between two IFSs A and B satisfies the following properties [70]:

- (1) $0 \leq S(A, B) \leq 1$
- (2) $S(A, B) = 1$ if $A = B$
- (3) $S(A, B) = S(B, A)$
- (4) $S(A, C) \leq S(A, B)$ and $S(A, C) \leq S(B, C)$ if $A \subseteq B \subseteq C$ for any A, B, C - intuitionistic fuzzy sets.

We have that similarity and distance (dissimilarity) measures are complementary, which implies $S(A, B) = 1 - d(A, B)$.

Let $A = \{(x, \mu_A(x), \nu_A(x)) \mid x \in U\}$, $B = \{(x, \mu_B(x), \nu_B(x)) \mid x \in U\}$ be two IFSs in the universe $U = \{x_1, \dots, x_n\}$. Several distance measures between A and B were proposed in the literature, from which we consider here only the *Normalized Euclidean distance* for two IFSs [71]:

$$d_{IE}(A, B) = \sqrt{\frac{1}{2n} \sum_{i=1}^n ((\mu_A(x_i) - \mu_B(x_i))^2 + (\nu_A(x_i) - \nu_B(x_i))^2)} \quad (1)$$

which will be called in what follows as *Intuitionistic Euclidean distance measure*.

In general a *similarity measure* between two single-value neutrosophic sets A and B is a function defined as [33,72,73]:

$$S : NS(X)^2 \rightarrow [0, 1]$$

where NS denotes the *Neutrosophic Set* concept.

The *Euclidean distance* or the *Euclidean dissimilarity measure* between two single-value neutrosophic elements $x_1(t_A^1, i_A^1, f_A^1)$, $x_2(t_A^2, i_A^2, f_A^2) \in A$ is defined as [72,73]:

$$d_E(x_1, x_2) = \sqrt{\frac{1}{3} [(t_A^1 - t_A^2)^2 + (i_A^1 - i_A^2)^2 + (f_A^1 - f_A^2)^2]} \quad (2)$$

Properties of the Euclidean distance. If x_1 and x_2 are two neutrosophic elements and $d_E(x_1, x_2)$ denotes the *Euclidean distance* as in definition (2), then the following properties are fulfilled:

1. $d_E(x_1, x_2) \in [0, 1]$
2. $d_E(x_1, x_2) = 0$ if and only if $x_1 = x_2$ (or $t_A^1 = t_A^2$, $i_A^1 = i_A^2$ and $f_A^1 = f_A^2$)
3. $d_E(x_1, x_2) = 1$ if and only if $|t_A^1 - t_A^2| = |i_A^1 - i_A^2| = |f_A^1 - f_A^2| = 1$
For examples: $x_1(1, 1, 1)$ and $x_2(0, 0, 0)$; or $x_1(1, 0, 0)$ and $x_2(0, 1, 1)$; or $x_1(0, 1, 0)$ and $x_2(1, 0, 1)$, etc.

The *Euclidean similarity measure* or the complement of the *Euclidean distance* between two neutrosophic elements $x_1(t_A^1, i_A^1, f_A^1)$, $x_2(t_A^2, i_A^2, f_A^2) \in A$ is defined as [72,73]:

$$s_E(x_1, x_2) = 1 - d_E(x_1, x_2)$$

$$= 1 - \sqrt{\frac{1}{3} [(t_A^1 - t_A^2)^2 + (i_A^1 - i_A^2)^2 + (f_A^1 - f_A^2)^2]} \quad (3)$$

Properties of the Euclidean similarity measure. If x_1 and x_2 are two neutrosophic elements and $s_E(x_1, x_2)$ denotes the *Euclidean similarity measure* as in definition (3), then the following properties are fulfilled:

1. $s_E(x_1, x_2) \in [0, 1]$
2. $s_E(x_1, x_2) = 0$ if and only if $x_1 = x_2$ (or $t_A^1 = t_A^2$, $i_A^1 = i_A^2$ and $f_A^1 = f_A^2$)
3. $s_E(x_1, x_2) = 1$ if and only if $|t_A^1 - t_A^2| = |i_A^1 - i_A^2| = |f_A^1 - f_A^2| = 1$
For examples: $x_1(1, 1, 1)$ and $x_2(0, 0, 0)$; or $x_1(1, 0, 0)$ and $x_2(0, 1, 1)$; or $x_1(0, 1, 0)$ and $x_2(1, 0, 1)$, etc.

The *Euclidean distance* between two neutrosophic elements can be extended to the *Normalized Euclidean distance* or *Normalized Euclidean dissimilarity measure* as follows.

Let A and B be two *single-valued neutrosophic sets* from the universe of discourse U ,

$A = \{x_i \in U, \text{ where } t_A(x_i), i_A(x_i), f_A(x_i) \in [0, 1], \text{ for } 1 \leq i \leq n \text{ and } n \geq 1\}$,

and

$B = \{x_i \in U, \text{ where } t_B(x_i), i_B(x_i), f_B(x_i) \in [0, 1], \text{ for } 1 \leq i \leq n \text{ and } n \geq 1\}$

The *Normalized Euclidean distance* between the two single-valued neutrosophic sets A and B is defined as [72–75]:

$$d_{nE}(A, B) = \left\{ \frac{1}{3n} \sum_{i=1}^n (t_A(x_i) - t_B(x_i))^2 + (i_A(x_i) - i_B(x_i))^2 + (f_A(x_i) - f_B(x_i))^2 \right\}^{\frac{1}{2}} \quad (4)$$

Properties of the Normalized Euclidean distance between two Neutrosophic Sets. If A and B are two single-valued neutrosophic sets then the *Normalized Euclidean distance* between A and B follows the distance measures properties:

1. $d_{nE}(A, B) \in [0, 1]$
2. $d_{nE}(A, B) = 0$ if and only if $A = B$ or for all $i \in \{1, 2, \dots, n\}$, $t_A(x_i) = t_B(x_i)$, $i_A(x_i) = i_B(x_i)$ and $f_A(x_i) = f_B(x_i)$
3. $d_{nE}(A, B) = 1$ if and only if for all $i \in \{1, 2, \dots, n\}$, $|t_A(x_i) - t_B(x_i)| = |i_A(x_i) - i_B(x_i)| = |f_A(x_i) - f_B(x_i)| = 1$

The *Normalized Euclidean similarity measure* or the complement of the *Normalized Euclidean distance* between two single-valued neutrosophic sets A and B is defined as [30,72–75]:

$$s_{nE}(A, B) = 1 - d_{nE}(A, B) \quad (5)$$

which implies

$$s_{nE}(A, B) = 1 - \left\{ \frac{1}{3n} \sum_{i=1}^n (t_A(x_i) - t_B(x_i))^2 + (i_A(x_i) - i_B(x_i))^2 + (f_A(x_i) - f_B(x_i))^2 \right\}^{\frac{1}{2}} \quad (6)$$

Properties of the Normalized Euclidean Similarity Measure between two Neutrosophic Sets If A and B are two single-valued neutrosophic sets then the *Normalized Euclidean Similarity Measure* between A and B follows the similarity measures properties:

1. $s_{nE}(A, B) \in [0, 1]$
2. $s_{nE}(A, B) = 0$ if and only if $A = B$ or for all $i \in \{1, 2, \dots, n\}$, $t_A(x_i) = t_B(x_i)$, $i_A(x_i) = i_B(x_i)$ and $f_A(x_i) = f_B(x_i)$
3. $s_{nE}(A, B) = 1$ if and only if for all $i \in \{1, 2, \dots, n\}$, $|t_A(x_i) - t_B(x_i)| = |i_A(x_i) - i_B(x_i)| = |f_A(x_i) - f_B(x_i)| = 1$

Another commonly used distance measure for two single-valued neutrosophic sets A and B is *Normalized Hamming distance measure* defined as [76]:

$$d_{nH}(A, B) = \frac{1}{3n} \sum_{i=1}^n (|t_A(x_i) - t_B(x_i)| + |i_A(x_i) - i_B(x_i)| + |f_A(x_i) - f_B(x_i)|) \quad (7)$$

3. Proposed approach

In this section we present a method designed for determining the semantic distance between pairs of words using a neutrosophic approach in which a word is interpreted as a *single-valued neutrosophic set* [47,48]. The semantic distances are determined without taking into account the part of speech data of the involved words. In our approach, the words are internally represented as vectors of three values, their corresponding SentiWordNet scores (shortly, SWN scores). Thus, any lexical and syntactical information about words is discarded.

In what follows we describe all the involved data, the theoretical concepts and the representations used in the implementation of the proposed similarity method.

3.1. Word-level neutrosophic sentiment similarity

In this study we address the problem of sentiment similarity between pairs of words by following the neutrosophic approach firstly proposed in [49] in which a word w is interpreted as a *single-valued neutrosophic set* [47,48] having the representation:

$$w = (\mu_{\text{truth}}(w), \mu_{\text{indeterminacy}}(w), \mu_{\text{false}}(w)) \quad (8)$$

where $\mu_{\text{truth}}(w)$ denotes the *truth membership degree* of w , $\mu_{\text{indeterminacy}}(w)$ represents the *indeterminacy membership degree* of w and $\mu_{\text{false}}(w)$ represents the *false membership degree* of the word w , with $\mu_{\text{truth}}(w), \mu_{\text{indeterminacy}}(w), \mu_{\text{false}}(w) \in [0, 1]$.

Similar with [49] we use the SentiWordNet lexical resource (shortly, SWN) in order to fuel the proposed approach with data. More precisely, the three membership degrees of the words representation (see Eq. (8)) are the positive, neutral and, correspondingly, the negative scores provided by SentiWordNet.

Problem definition. We propose and evaluate a method for the problem of determining the sentiment class of a word w by measuring its distance from several *seed words*, one seed word for each sentiment class. In this assumption, we propose the usage of three semantic distances: *Intuitionistic Euclidean distance*, *Euclidean distance* and *Hamming distance*. We work with 7 seed words, each seed word being a representative sentiment word for each of the seventh sentiment degrees: *strong positive*, *positive*, *weak positive*, *neutral*, *weak negative*, *negative* and *strong negative*. We prove that all the considered theoretical concepts work very well as we apply and evaluate them on all the SentiWordNet words (that is, 155 287 words).

If w_1 and w_2 are highly similar, we expect the semantic distance value to be closer to 0, otherwise semantic relatedness value should be closer to 1. We consider SentiWordNet sentiment scores as the only features of the words.

As we have already pointed out, in this approach, a word internal representation consists of its SWN scores. In this assumption, a word w can be considered a *single-valued neutrosophic set* and thus, all the properties involving this concept can be used and applied.

In order to exemplify this assumption, let us consider the verb “scam”. In the SWN dataset this word has a single entry, that is it has a single SWN score triplet:

$$\text{scam} = (0, 0.125, 0.875)$$

By following the neutrosophic assumption in which a word is considered a single-value neutrosophic set, the representation of the word w becomes:

$$w(t_w, i_w, f_w)$$

where:

- the degree of membership, t_w , is the word positive score,
- the degree of indeterminate-membership, i_w , is the word neutral score,
- the degree of non-membership, f_w , is the word negative score.

Obviously the conditions imposed on these degree values are preserved: $t_w, i_w, f_w \in [0, 1]$ and $0 \leq t_w + i_w + f_w = 1 \leq 3$.

For the considered example we have: $t_{\text{scam}} = 0$, $i_{\text{scam}} = 0.125$ and $f_{\text{scam}} = 0.875$, which implies $\text{scam}(0, 0.125, 0.875)$.

Let us now consider the general case in which a word w can appear in more than one synset in the SentiWordNet lexicon, meaning that the word has more than one sense. In this case we have n SWN score triplets for a single word w , with $n \geq 1$.

In order to construct the neutrosophic word representation, a single scores triplet must be provided. For this reason, for every word w with n senses, $n \geq 1$, we implemented the *weighted average formula* (after [77]) over all its positive, negative and, respectively, neutral scores obtaining in this manner three sentiment scores for all the three facets of a word sentiment polarity:

- the overall positive score of the word w :

$$t_w = \frac{t_{w^1} + \frac{1}{2}t_{w^2} + \dots + \frac{1}{n}t_{w^n}}{1 + \frac{1}{2} + \dots + \frac{1}{n}} \quad (9)$$

- the overall neutral score of the word w :

$$i_w = \frac{i_{w^1} + \frac{1}{2}i_{w^2} + \dots + \frac{1}{n}i_{w^n}}{1 + \frac{1}{2} + \dots + \frac{1}{n}} \quad (10)$$

- the overall negative score of the word w :

$$f_w = \frac{f_{w^1} + \frac{1}{2}f_{w^2} + \dots + \frac{1}{n}f_{w^n}}{1 + \frac{1}{2} + \dots + \frac{1}{n}} \quad (11)$$

where w^1 denotes the first sense of the word w , w^2 represents the second sense of the word w , etc.

In order to calculate the overall scores of a word w we use the weighted average formula because it considers frequencies of the words' senses: the score of the first sense (which is the most frequent) is preserved entirely, while the rest of the scores, which correspond to the less used senses, appear divided accordingly (by $1/2, 1/3$, etc.).

The sentiment class of a word is determined by computing a single score upon these overall scores. This unique score will represent the average of the differences between the positivity and negativity scores calculated per each sense.

More precisely, for a word w with n senses, the single sentiment score is determined by following the already defined mechanism for words' scores calculus based on SentiWordNet triplets (see [42]) which implies to determine the average weighted difference between their positive and negative scores such as:

$$\frac{1}{n} \sum_{i=1}^n \omega_i (\text{pos}_i - \text{neg}_i)$$

where the weights ω_i are chosen taking into account several word characteristics which can carry different levels of importance in conveying the described sentiment [42] (such as part of speech) and n represents the number of synsets in which the word w

```

function sent_class(score)
  sent_class <- "neutral"
  IF (score > 0.5) THEN sent_class <- "strong positive" ELSE
  IF (0.25 < score <= 0.5) THEN sent_class <- "positive" ELSE
  IF (0 < score <= 0.25) THEN sent_class <- "weak positive" ELSE
  IF (-0.25 <= score < 0) THEN sent_class <- "weak negative" ELSE
  IF (-0.5 <= score < -0.25) THEN sent_class <- "negative" ELSE
  IF (score < -0.5) THEN sent_class <- "strong negative"

  return sent_class
endfunction

```

Fig. 1. The *sent_class* function.

```

function distance(dist, sent_class_w1, sent_class_w2)
  IF (dist is between Table2(sent_class_w1, sent_class_w2))
    return true
  return false
endfunction

```

Fig. 2. The *evaluate* function.

appears, that is the number of its senses. The average is used in order to ensure that the resulted scores are ranging between -1 and 1 [42].

Let us consider a word w with n senses, $w_1, w_2, \dots, w_n$. In this study the overall score of the word w is determined using the formula [42,77]:

$$\text{score} = \frac{(t_{w_1} - f_{w_1}) + \frac{1}{2}(t_{w_2} - f_{w_2}) + \dots + \frac{1}{n}(t_{w_n} - f_{w_n})}{1 + \frac{1}{2} + \dots + \frac{1}{n}} \quad (12)$$

As we have already pointed out, the values of *score* vary between -1 (meaning that the word w is a “strong negative” word) and 1 (the word w is a “strong positive” word).

Usually sentiment analysis applications deal with binary (positive vs. negative) or ternary (positive vs. negative vs. objective) classifications which normally leads to very good state-of-the-art accuracy (more than 70%) [42]. In this study, using the sentiment scores defined for the SentiWordNet synsets, we consider all the degrees of sentiments referred in the literature:

- *strong positive/negative word*: great difference between the positive/ negative scores and the negative/positive scores of the word (usually, above 0.5)
- *positive/negative word*: the positive/negative scores are greater than the negative/positive ones (the difference is smaller than 0.5 but greater than 0.25)
- *weak positive/negative word*: small difference between the positive/ negative scores and the negative/positive ones
- *neutral word*: the neutral scores subsume the positive and negative scores.

We defined a set of rules in order to uniquely map the general score of a word to one of the following sentiment classes: “strong positive”, “positive”, “weak positive”, “neutral”, “weak negative”, “negative”, “strong negative”. The rules are given in an algorithmic form under the *sent_class* function in Fig. 1.

If w_1 and w_2 are two words: $w_1(t_{w_1}, i_{w_1}, f_{w_1})$, $w_2(t_{w_2}, i_{w_2}, f_{w_2})$, the distance measures between w_1 and w_2 are as follows:

1. Intuitionistic Euclidean distance:

$$d_{IE}(w_1, w_2) = \sqrt{\frac{1}{2}[(t_{w_1} - t_{w_2})^2 + (f_{w_1} - f_{w_2})^2]} \quad (13)$$

2. Euclidean distance:

$$d_E(w_1, w_2)$$

$$= \sqrt{\frac{1}{3}[(t_{w_1} - t_{w_2})^2 + (i_{w_1} - i_{w_2})^2 + (f_{w_1} - f_{w_2})^2]} \quad (14)$$

3. Hamming distance:

$$d_H(w_1, w_2) = \frac{1}{3}[|t_{w_1} - t_{w_2}| + |i_{w_1} - i_{w_2}| + |f_{w_1} - f_{w_2}|] \quad (15)$$

4. Experimental setup

We evaluate the accuracy of the considered mechanism by implementing the Normalized Euclidean and, in order to give terms of comparison, we also evaluate the Normalized Hamming distance and Intuitionistic Euclidean distance in the same scenario.

In Table 2 we give the values we impose on the distance measures with respect to the sentiment classes of the involved two words. The values of Table 2 are symmetrical and for this reason only the values under the main diagonal are given.

Obviously, we considered the smallest distance values in cases of words having the same sentiment class (these cases are given on the diagonal). A strong value for distance value means that the two words are completely dissimilar from the sentiment polarity point of view. For example, a word having “negative” sentiment class (or shortly, a negative word) and a word with “positive” sentiment class (a positive word) must have the distance value d bigger than 0.65, where d cannot be greater than 1.

Based on Table 2 values, the evaluation of the distance values with respect to the sentiment classes of the involved words is depicted in Fig. 2.

For the evaluation scenario we chose seven “seed words”, one for each sentiment class and we iterate through the lexical resource and calculate the distance measures between each of the seven seed words and all the words that appear in SentiWordNet (155287 words in total).

Resuming, the algorithmic form of the evaluation scenario for the proposed word-level sentiment similarity method is given in Fig. 3.

4.1. Evaluation scores

In Table 3 we present the selected seed words together with the results obtained by implementing and evaluating all the three distance measures proposed for this study: Normalized Euclidean

```

foreach seed_w <- seed word from sentiment classes: {strong positive, positive,
                                                    weak positive, neutral,
                                                    weak negative, negative,
                                                    strong negative}

  seed_sent_class <- sentiment class of seed_w

  foreach w <- word from SentiWordNet
    score <- overall score of w
    w_sent_class <- sentiment_class(score)

    foreach d <- distance from {dIF, dE, dH}
      evaluate(d, w_sent_class, seed_sent_class)
    endfor
  endfor
endfor

```

Fig. 3. The evaluation scenario.

Table 2

The used distance measure values with respect to the words sentiment classes.

Strong positive	[0, 0.2)						
Positive	[0, 0.3)	[0, 0.2)					
Weak positive	[0.25, 0.5)	[0, 0.3)	[0, 0.2)				
Neutral	[0.3, 0.65)	[0.3, 0.65)	[0, 0.3)	[0, 0.2)			
Weak negative	(0.65, 1]	(0.65, 1]	[0.25, 0.5)	[0, 0.3)	[0, 0.2)		
Negative	(0.65, 1]	(0.65, 1]	(0.65, 1]	[0.3, 0.65)	[0, 0.3)	[0, 0.2)	
Strong negative	(0.65, 1]	(0.65, 1]	(0.65, 1]	[0.3, 0.65)	[0.25, 0.5)	[0, 0.3)	[0, 0.2)
Sent. classes	Strong positive	Positive	Weak positive	Neutral	Weak negative	Negative	Strong negative

Table 3

Evaluation scores.

Seed word	Similarity distance precision		
	Euclidean distance	Hamming distance	Intuitionistic Euclidean distance
Sent. class: Strong positive Word: singable#a Overall scores: (0.75, 0.0, 0.25)	0.8411	0.8580	0.8808
Sent. class: Positive Word: spunky#a Overall scores: (0.5416, 0.2083, 0.25)	0.7714	0.7725	0.8059
Sent. class: Weak positive Word: immunized#a Overall scores: (0.5, 0.375, 0.125)	0.0392	0.0608	0.1219
Sent. class: Neutral Word: hydrostatic#a Overall scores: (0.0, 0.0, 1.0)	0.9676	0.9489	0.9570
Sent. class: Weak negative Word: misguided#a Overall scores: (0.25, 0.4583, 0.2916)	0.0973	0.1070	0.1279
Sent. class: Negative Word: reformable#a Overall scores: (0.125, 0.5, 0.375)	0.8259	0.8260	0.8573
Sent. class: Strong negative Word: unworkmanlike#a Overall scores: (0.0, 0.75, 0.25)	0.8542	0.8764	0.8875

distance, Normalized Hamming distance and Intuitionistic Euclidean distance measure.

The obtained accuracy results are mainly influenced by the way in which the considered seed words can be distinguished from the most preponderant words of this lexical resource, that is from the *neutral words* as they are the most frequent words of the SentiWordNet resource.

As it can be seen in Table 3 and Fig. 4 the considered distance measures have a similar behaviour: all the distance measures have more than 77% precision for the most of the considered seed words, which is above the average precision (70%) recognized in the specialized literature for the sentiment classifiers accuracy.

The highest precision (more than 74%) is achieved by applying the distance measures between the *neutral seed word* and all the SentiWordNet's words. Also very good scores (more than 82%) were achieved by applying the distances between the *negative seed word* and SentiWordNet words, then we have the scores corresponding to the *strong positive seed word* (more than 0.84 as precision) and finally the scores corresponding to the *positive seed word* (more than 77% precision).

But these very good results were not achieved for the *weak positive seed word* and *weak negative seed word* where the precision is almost zero. This failure can be caused by the fact that these particular sentiment words cannot be distinguished very well from the most preponderant words of SentiWordNet, that is from the *neutral words*.

We can therefore conclude that all the considered distance measures can distinguish very well the words of the most important sentiment classes from the point of view of a sentiment classifier: the (strong) positive or negative words and the neutral words. Still, the proposed measures are not capable for measuring the similarity of *weak sentiment words* with the rest of the sentiment words.

The most important conclusion that comes from the performed experiment is that the behaviour of all the considered distance measures is very similar – almost identical (see Fig. 4). We interpret this result as a proof for the robustness of the considered theory.

5. Conclusions and future work

In the latest years there has been developed a relatively large number of word-to-word similarity studies that can be grouped in two main categories: distance-oriented measures applied on structured representations and metrics based on distributional similarity learned from large text collections [50].

In this paper we propose a sentiment similarity method that fits in the first category of similarity studies and which takes into account only the sentiment aspects of the words and not their lexical category. We follow here recent text similarity approaches such as [1,28] defined around the same hypothesis which postulates that knowing the sentiment is beneficial in measuring the similarity.

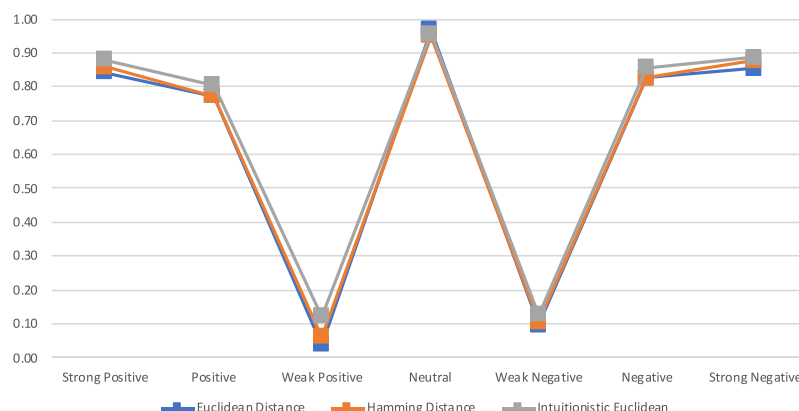


Fig. 4. The graphical visualization of the similarity distances precision.

Our proposal is formalized in a domain that was never used before for this kind of task – the neutrosophic theory, as it uses neutrosophic sets for representing the sentiment aspects of the words. The neutrosophic set is a generalization of the intuitionistic fuzzy set concept, and thus our proposal is in line with the recent fuzzy based studies that started to emerge for text processing tasks [20,78,79]. Indeed, fuzzy logic is capable of dealing with linguistic uncertainty as it considers the classification problem to be a “degree of grey” problem rather than a “black and white” problem [20] (the last one is the most used approach in sentiment analysis tasks).

For this first approach we obtained very promising results. Indeed, by applying distance measures on the neutrosophic words representations we shown that we can thus obtain a similarity method as we manage very clear to distinguish the words of the most important sentiment classes from the rest of the considered words: the SentiWordNet entries, that is, 155 287 words of all possible sentiment classes.

We also plan to extend our study to sequences of words with the intended scope of designing a method that can be applied for measuring documents similarity.

Conflict of interest

The authors declare that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

References

- [1] A. Kashyap, L. Han, R. Yus, J. Sleeman, T. Satyapanich, S. Gandhi, T. Finin, Robust semantic text similarity using LSA, machine learning, and linguistic resources, *Lang. Resour. Eval.* 50 (1) (2016) 125–161, <http://dx.doi.org/10.1007/s10579-015-9319-2>.
- [2] B. Dolan, C. Quirk, C. Brockett, Unsupervised construction of large paraphrase corpora: exploiting massively parallel news sources, in: *Proceedings of the 20th International Conference on Computational Linguistics*, in: COLING '04, Stroudsburg, PA, USA, 2004, <http://dx.doi.org/10.3115/1220355.1220406>, Article 350.
- [3] B. Sriram, D. Fuhr, E. Demir, H. Ferhatosmanoglu, M. Demirbas, Short text classification in twitter to improve information filtering, in: *Proceedings of the 33rd International ACM SIGIR conference on Research and Development in Information Retrieval (SIGIR '10)*, New York, USA, 2010, pp. 841–842, <http://dx.doi.org/10.1145/1835449.1835643>.
- [4] D. Borth, J. Rongrong, T. Chen, T. Breuel, S.F. Chang, Large-scale visual sentiment ontology and detectors using adjective noun pairs, in: *Proceedings of the 21st ACM International Conference on Multimedia (MM '13)*, New York, USA, 2013, pp. 223–232, <http://dx.doi.org/10.1145/2502081.2502282>.
- [5] T.A.S. Coelho, P.P. Calado, L.V. Souza, B. Ribeiro-Neto, R. Muntz, Image retrieval using multiple evidence ranking, *IEEE Trans. Knowl. Data Eng.* 16 (4) (2004) 408–417, <http://dx.doi.org/10.1109/TKDE.2004.1269666>.
- [6] D. Metzler, S. Dumais, C. Meek, Similarity measures for short segments of text, in: G. Romano G. Amati (Ed.), *Proceedings of the 29th European conference on IR research (ECIR'07)*, Springer-Verlag, Berlin, Heidelberg, 2007, pp. 16–27.
- [7] D. Kauchak, R. Barzilay, Paraphrasing for automatic evaluation, in: *Proceedings of the Main Conference on Human Language Technology Conference of the North American Chapter of the Association of Computational Linguistics (HLT-NAACL '06)*, Stroudsburg, PA, USA, 2006, pp. 455–462, <http://dx.doi.org/10.3115/1220835.1220893>.
- [8] E. Cambria, B. Schuller, Y. Xia, C. Havasi, New avenues in opinion mining and sentiment analysis, *IEEE Intell. Syst.* 28 (2) (2013) 15–21, <http://dx.doi.org/10.1109/MIS.2013.30>.
- [9] R.A. Calvo, S. D'Mello, Affect detection: an interdisciplinary review of models methods and their applications, *IEEE Trans. Affect. Comput.* 1 (1) (2010) 18–37, <http://dx.doi.org/10.1109/T-AFFC.2010.1>.
- [10] H. Gunes, B. Schuller, Categorical and dimensional affect analysis in continuous input: current trends and future directions, *Image Vis. Comput.* 31 (2) (2013) 120–136, <http://dx.doi.org/10.1016/j.imavis.2012.06.016>.
- [11] B. Pang, L. Lee, A sentimental education: sentiment analysis using subjectivity summarization based on minimum cuts, in: *Proceedings of the 42nd Meeting of the Association for Computational Linguistics (ACL'04)*, 2004, pp. 271–278, <http://dx.doi.org/10.3115/1218955.1218990>.
- [12] B. Pang, L. Lee, Seeing stars: exploiting class relationships for sentiment categorization with respect to rating scales, in: *Proceedings of the 43rd Meeting of the Association for Computational Linguistics (ACL'05)*, 2005, pp. 115–124, <http://dx.doi.org/10.3115/1219840.1219855>.
- [13] B. Pang, L. Lee, Opinion mining and sentiment analysis, *Found. Trends Inf. Retr.* 2 (1–2) (2008) 1–135, <http://dx.doi.org/10.1561/15000000011>.
- [14] B. Liu, *Sentiment Analysis and Opinion Mining*, Morgan & Claypool Publishers, 2012.
- [15] E. Cambria, Affective computing and sentiment analysis, *IEEE Intell. Syst.* 31 (2) (2016) 102–107, <http://dx.doi.org/10.1109/MIS.2016.31>.
- [16] L. Augustyniak, P. Szymanski, T. Kajdanowicz, W. Tulgłowicz, Comprehensive study on lexicon-based ensemble classification sentiment analysis, *Entropy* 18 (1) (2016) 4, <http://dx.doi.org/10.3390/e18010004>.
- [17] S. Fathi, H. ElGhawalby, A.A. Salama, A neutrosophic graph similarity measures, in: F. Smarandache, S. Pramanik (Eds.), *New Trends in Neutrosophic Theory and Applications*, Pons Editions, 2016, pp. 291–301.
- [18] A.F. Smeaton, Progress in the application of natural language processing to information retrieval tasks, *Comput. J.* 35 (3) (1992) 268–278, <http://dx.doi.org/10.1093/comjnl/35.3.268>.
- [19] A. Papadopoulos, J. Desbiens, Method and system for analysing sentiments, Google Patents, US Patent App. 14/432, 436, 2015. <https://encrypted.google.com/patents/US20150286953>.
- [20] C. Jefferson, H. Liu, M. Cocea, Fuzzy approach for sentiment analysis, in: *Proceedings of International Conference on Fuzzy Systems (FUZZ-IEEE 2017)*, Naples, Italy, 2017, pp. 1–6, <http://dx.doi.org/10.1109/FUZZ-IEEE.2017.8015577>.
- [21] R. Batuwita, V. Palade, FSVM-CIL: fuzzy support vector machines for class imbalance learning, *IEEE Trans. Fuzzy. Syst.* 18 (3) (2010) 558–571, <http://dx.doi.org/10.1109/TFUZZ.2010.2042721>.
- [22] M. Dragoni, G. Petrucci, A fuzzy-based strategy for multi-domain sentiment analysis, *Internat. J. Approx. Reason.* 93 (2018) 59–73, <http://dx.doi.org/10.1016/j.ijar.2017.10.021>.
- [23] C. Zhao, S. Wang, D. Li, Determining fuzzy membership for sentiment classification: a three-layer sentiment propagation model, *PLoS ONE* 11 (11) (2016) e0165560, <http://dx.doi.org/10.1371/journal.pone.0165560>.
- [24] K. Howells, A. Ertugan, Applying fuzzy logic for sentiment analysis of social media network data in marketing, *Procedia Comput. Sci.* 120 (2017) 664–670, <http://dx.doi.org/10.1016/j.procs.2017.11.293>.

- [25] P. Turney, M. Littman, Measuring praise and criticism: inference of semantic orientation from association, *ACM Trans. Inform. Syst.* 21 (4) (2003) 315–346, <http://dx.doi.org/10.1145/944012.944013>.
- [26] S.M. Kim, E. Hovy, Extracting opinions, opinion holders, and topics expressed in online news media text, in: *Proceedings of the Workshop on Sentiment and Subjectivity in Text (ACL'06)*, 2006, pp. 1–8.
- [27] M. Mohtarami, H. Amiri, L. Man, P.T. Thanh, L.T. Chew, Sense sentiment similarity: an analysis, in: *Proceedings of the Twenty-Sixth Conference on Artificial Intelligence (AAAI'12)*, 2012, pp. 1706–1712.
- [28] R. Balamurali, S. Mukherjee, A. Malu, P. Bhattacharyya, Leveraging sentiment to compute word similarity, in: *Proceedings of the 6th International Global Wordnet Conference (GWC'12)*, Matsue, Japan, 2012, pp. 10–17.
- [29] T. Wilson, J. Wiebe, R. Hwa, Just how mad are you? finding strong and weak opinion clauses, in: *Proceedings of the 21st Conference of the American Association for Artificial Intelligence (AAAI'04)*, San Jose, US, 2004, pp. 761–769.
- [30] J. Ye, Similarity measures between interval neutrosophic sets and their applications in multicriteria decision-making, *J. Intell. Fuzzy Syst.* 26 (1) (2014) 165–172, <http://dx.doi.org/10.3233/IFS-120724>.
- [31] J. Ye, Vector similarity measures of simplified neutrosophic sets and their application in multi-criteria decision-making, *Int. J. Fuzzy Syst.* 16 (2) (2014) 204–215.
- [32] J. Ye, Multiple attribute group decision-making method with completely unknown weights based on similarity measures under single valued neutrosophic environment, *J. Intell. Fuzzy Syst.* 27 (6) (2014) 2927–2935, <http://dx.doi.org/10.3233/IFS-141252>.
- [33] J. Ye, Q.S. Zhang, Single valued neutrosophic similarity measures for multiple attribute decision-making, *Neutrosophic Sets Syst.* 2 (2014) 48–54.
- [34] S. Baccianella, A. Esuli, F. Sebastiani, SentiWordNet 3.0: an enhanced lexical resource for sentiment analysis and opinion mining, in: *Proceedings of the Seventh International Conference on Language Resources and Evaluation (LREC 2010)*, 2010, pp. 2200–2204.
- [35] V. Rus, N. Niraula, R. Banjade, Similarity measures based on latent dirichlet allocation, in: A. Gelbukh (Ed.), *Computational Linguistics and Intelligent Text Processing (CICLing 2013)*, in: *Lecture Notes in Computer Science*, vol. 7816, Springer, Berlin, Heidelberg, 2013, pp. 459–470, http://dx.doi.org/10.1007/978-3-642-37247-6_37.
- [36] I. Atoum, C.H. Bong, Joint distance and information content word similarity measure, in: *Soft Computing Applications and Intelligent Systems. Communications in Computer and Information Science*, Vol. 378, Springer, Berlin, Heidelberg, 2013, pp. 257–267, http://dx.doi.org/10.1007/978-3-642-40567-9_22.
- [37] J.M. Gawron, Improving sparse word similarity models with asymmetric measures, in: *Proceedings of the 52nd Annual Meeting of the Association for Computational Linguistics (ACL'14)*, 2014, pp. 296–301.
- [38] R.L. Goldstone, in: R.A. Wilson, F.C. Keil (Eds.), *The MIT encyclopedia of the cognitive sciences*, Cambridge, UK, 1999, pp. 757–759.
- [39] C. Fellbaum (Ed.), *WordNet: An Electronic Lexical Database*, MIT Press, Cambridge, MA, 1998.
- [40] V. Negru, G. Grigoraș, D. Dănculescu, Natural language agreement in the generation mechanism based on stratified graphs, in: *Proceedings of the 7th Balkan Conference on Informatics (BCI 2015)*, Craiova, Romania, 2015, p. 36, <http://dx.doi.org/10.1145/2801081.2801121>.
- [41] F. Hristea, On a dependency-based semantic space for unsupervised noun sense disambiguation with an underlying Naïve Bayes model, in: *Proceedings of the Joint Symposium on Semantic Processing. Textual Inference and Structures in Corpora*, Trento, Italy, 2013, pp. 90–94, <http://aclweb.org/anthology/W13-3825>.
- [42] E. Russell, *Real-Time Topic and Sentiment Analysis in Human-Robot Conversation* (Master's Theses of Marquette University), 2015, p. 338.
- [43] J. Kreutzer, N. White, *Opinion Mining using SentiWordNet*, Uppsala University, 2013.
- [44] A. Esuli, F. Sebastiani, SENTIWORDNET: a publicly available lexical resource for opinion mining, in: *Proceedings of the 5th Conference on Language Resources and Evaluation (LREC'06)*, 2006, pp. 417–422.
- [45] H. Yu, V. Hatzivassiloglou, Towards answering opinion questions: separating facts from opinions and identifying the polarity of opinion sentences, in: *Proceedings of the 8th Conference on Empirical Methods in Natural Language Processing (EMNLP'03)*, 2003, pp. 129–136, <http://dx.doi.org/10.3115/1119355.1119372>.
- [46] P. Turney, Thumbs up or thumbs down? semantic orientation applied to unsupervised classification of reviews, in: *Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics (ACL'02)*, 2002, pp. 417–424, <http://dx.doi.org/10.3115/1073083.1073153>.
- [47] F. Smarandache, *Neutrosophy: Neutrosophic Probability, Set, and Logic*, American Research Press, Rehoboth, USA, 1998.
- [48] F. Smarandache, *Symbolic Neutrosophic Theory*, Europa Nova, Brussels, 2015.
- [49] M. Colhon, Ș. Vlăduțescu, X. Negrea, How objective a neutral word is? a neutrosophic approach for the objectivity degrees of neutral words, *Symmetry-Basel* 9 (11) (2017) 280, <http://dx.doi.org/10.3390/sym9110280>.
- [50] R. Mihalcea, C. Corley, C. Strapparava, Corpus-based and knowledge-based measures of text semantic similarity, in: A. Cohn (Ed.), *Proceedings of the 21st National Conference on Artificial Intelligence*, in: *AAAI'06*, vol. 1, AAAI Press, Boston, Massachusetts, 2006, pp. 775–780.
- [51] T.K. Landauer, P.W. Foltz, D. Laham, Introduction to latent semantic analysis, *Discourse Processes* 25 (1998) 259–284, <http://dx.doi.org/10.1080/01638539809545028>.
- [52] P. Isola, D. Zoran, D. Krishnan, E.H. Adelson, Crisp boundary detection using pointwise mutual information, in: D. Fleet, T. Pajdla, B. Schiele, T. Tuytelaars (Eds.), *Computer Vision – ECCV 2014*, in: *Lecture Notes in Computer Science*, vol. 8691, Springer, 2014, pp. 799–814, http://dx.doi.org/10.1007/978-3-319-10578-9_52.
- [53] B. Hajian, T. White, Measuring semantic similarity using a multi-tree model, in: *Proceedings of the 9th Workshop on Intelligent Techniques for Web Personalization and Recommender Systems (ITWP 2011)*, 2011, p. 1.
- [54] M. Khan, L.H. Son, M. Ali, H.T.M. Chau, N.T.N. Na, F. Smarandache, Systematic review of decision making algorithms in extended neutrosophic sets, *Symmetry* 10 (8) (2018) 314, <http://dx.doi.org/10.3390/sym10080314>.
- [55] N. Werro, *Fuzzy Classification of Online Customers. Fuzzy Management Methods*, Springer International Publishing, Springer International Publishing, Switzerland, 2015.
- [56] F. Ali, D. Kwak, P. Khan, S.R. Islam, K.H. Kim, K.S. Kwak, Fuzzy ontology-based sentiment analysis of transportation and city feature reviews for safe traveling, *Transp. Res. C* 77 (2017) 33–48, <http://dx.doi.org/10.1016/j.trc.2017.01.014>.
- [57] K.T. Atanassov, Intuitionistic fuzzy sets, *Fuzzy Sets Syst.* 20 (1) (1986) 87–96, http://dx.doi.org/10.1007/978-3-7908-1870-3_1.
- [58] H. Wang, F. Smarandache, Y. Zhang, R. Sunderraman, Single valued neutrosophic sets, in: *Technical Sciences and Applied Mathematics, Infinite Study*, 2012, pp. 10–14.
- [59] H.M. Malik, M. Akram, F. Smarandache, Soft rough neutrosophic influence graphs with application, *Mathematics* 6 (7) (2018) 125, <http://dx.doi.org/10.3390/math6070125>.
- [60] S. Broumi, M. Talea, F. Smarandache, A. Bakali, Single valued neutrosophic graphs: degree, order and size, in: *Proceedings of IEEE International Conference on Fuzzy Systems (FUZZ-IEEE)*, 2016, pp. 2444–2451, <http://dx.doi.org/10.1109/FUZZ-IEEE.2016.7738000>.
- [61] S. Broumi, A. Bakali, M. Talea, F. Smarandache, L. Vladareanu, Computation of shortest path problem in a network with sv-trapezoidal neutrosophic numbers, in: *Proceedings of International Conference on Advanced Mechatronic Systems (ICAMechS)*, 2016, pp. 417–422, <http://dx.doi.org/10.1109/ICAMechS.2016.7813484>.
- [62] S. Broumi, A. Bakali, M. Talea, F. Smarandache, L. Vladareanu, Applying dijkstra algorithm for solving neutrosophic shortest path problem, in: *Proceedings of International Conference on Advanced Mechatronic Systems (ICAMechS)*, 2016, pp. 412–416, <http://dx.doi.org/10.1109/ICAMechS.2016.7813483>.
- [63] M. Akram, N. Ishfaq, S. Sayed, F. Smarandache, Decision-making approach based on neutrosophic rough information, *Algorithms* 11 (5) (2018) 59, <http://dx.doi.org/10.3390/a11050059>.
- [64] M. Akram, F. Shumaiza Smarandache, Decision-making with bipolar neutrosophic TOPSIS and bipolar neutrosophic ELECTRE-I, *Axioms* 7 (33) (2018) 33, <http://dx.doi.org/10.3390/axioms7020033>.
- [65] P.T.M. Phuong, P.H. Thong, L.H. Son, Theoretical analysis of picture fuzzy clustering: convergence and property, *J. Comput. Sci. Cybern.* 34 (1) (2018) 17–32, <http://dx.doi.org/10.15625/1813-9663/34/1/12725>.
- [66] P.H. Thong, L.H. Son, Online picture fuzzy clustering: a new approach for real-time fuzzy clustering on picture fuzzy sets, in: *Proceedings of International Conference on Advanced Technologies for Communications*, 2018, pp. 193–197.
- [67] R.T. Ngan, B.C. Cuong, T.M. Tuan, L.H. Son, Medical diagnosis from images with intuitionistic fuzzy distance measures, in: *Proceedings of International Joint Conference, IJCRS 2018*, Quy Nhon, Vietnam, 2018, pp. 479–490, http://dx.doi.org/10.1007/978-3-319-99368-3_37.
- [68] G. Shahzadi, M. Akram, A.B. Saeid, An application of single-valued neutrosophic sets in medical diagnosis, *Neutrosophic Sets Syst.* 18 (2017) 80–88, <http://dx.doi.org/10.5281/zenodo.1175619>.
- [69] S. Jha, R. Kumar, L.H. Son, J.M. Chatterjee, M. Khari, N. Yadav, Neutrosophic soft set decision making for stock trending analysis, *Evolv. Syst.* (2018) 1–7, <http://dx.doi.org/10.1007/s12530-018-9247-7>.
- [70] S.M. Chen, C.H. Chang, A novel similarity between Atanassov's intuitionistic fuzzy sets based on transformation techniques with applications to pattern recognition, *Inform. Sci.* 291 (2015) 96–114, <http://dx.doi.org/10.1016/j.ins.2014.07.033>.
- [71] K.T. Atanassov, *Intuitionistic Fuzzy Sets. Theory and Applications*, Physica-Verlag, Wyrzburg, 1999.

- [72] J. Ye, Single valued neutrosophic minimum spanning tree and its clustering method, *Int. J. Intell. Syst.* 23 (3) (2014) 311–324, <http://dx.doi.org/10.1515/jisys-2013-0075>.
- [73] J. Ye, Clustering methods using distance-based similarity measures of single-valued neutrosophic sets, *Int. J. Intell. Syst.* 23 (4) (2014) 379–389, <http://dx.doi.org/10.1515/jisys-2013-0091>.
- [74] S. Broumi, F. Smarandache, Several similarity measures of neutrosophic sets, *Neutrosoph. Sets Syst.* 1 (2013) 54–62, <http://dx.doi.org/10.5281/zenodo.30312>.
- [75] S. Broumi, F. Smarandache, New distance and similarity measures of interval neutrosophic sets, in: *Proceedings of 17th International Conference on Information Fusion (FUSION'14)*, 2014, pp. 1–7.
- [76] P. Majumdar, S.K. Samanta, On similarity and entropy of neutrosophic sets, *J. Intell. Fuzzy Syst.* 26 (3) (2014) 1245–1252, <http://dx.doi.org/10.3233/IFS-130810>.
- [77] P. Tönberg, The Demo Java Class for the SentiWordNet Website, <http://sentiwordnet.isti.cnr.it/code/SentiWordNetDemoCode.java> (accessed 2018).
- [78] H. Liu, M. Cocea, Fuzzy rule based systems for interpretable sentiment analysis, in: *Proceedings of Ninth International Conference on Advanced Computational Intelligence (ICACI'17)*, 2017, pp. 129–136, <https://doi.org/10.1109/ICACI.2017.7974497>.
- [79] D. Chandran, K.A. Crockett, D. Mclean, A. Crispin, An automatic corpus based method for a building multiple Fuzzy word dataset, in: *Proceedings of IEEE International Conference on Fuzzy Systems (FUZZ-IEEE)*, Istanbul, Turkey, 2015, pp. 1–8, <https://doi.org/10.1109/FUZZ-IEEE.2015.7337877>.