



LUND UNIVERSITY

Detection of Stance-Related Characteristics in Social Media Text

Simaki, Vasiliki; Simakis, Panagiotis; Paradis, Carita; Kerren, Andreas

Published in:

SETN '18 Proceedings of the 10th Hellenic Conference on Artificial Intelligence

DOI:

[10.1145/3200947.3201017](https://doi.org/10.1145/3200947.3201017)

2018

Document Version:

Peer reviewed version (aka post-print)

[Link to publication](#)

Citation for published version (APA):

Simaki, V., Simakis, P., Paradis, C., & Kerren, A. (2018). Detection of Stance-Related Characteristics in Social Media Text. In *SETN '18 Proceedings of the 10th Hellenic Conference on Artificial Intelligence Association for Computing Machinery (ACM)*. <https://doi.org/10.1145/3200947.3201017>

Total number of authors:

4

General rights

Unless other specific re-use rights are stated the following general rights apply:

Copyright and moral rights for the publications made accessible in the public portal are retained by the authors and/or other copyright owners and it is a condition of accessing publications that users recognise and abide by the legal requirements associated with these rights.

- Users may download and print one copy of any publication from the public portal for the purpose of private study or research.
- You may not further distribute the material or use it for any profit-making activity or commercial gain
- You may freely distribute the URL identifying the publication in the public portal

Read more about Creative commons licenses: <https://creativecommons.org/licenses/>

Take down policy

If you believe that this document breaches copyright please contact us providing details, and we will remove access to the work immediately and investigate your claim.

LUND UNIVERSITY

PO Box 117
221 00 Lund
+46 46-222 00 00

See discussions, stats, and author profiles for this publication at: <https://www.researchgate.net/publication/324756885>

Detection of Stance-Related Characteristics in Social Media Text

Chapter · July 2018

DOI: 10.1145/3200947.3201017

CITATIONS

0

4 authors, including:



Vasiliki Simaki

Lancaster University

23 PUBLICATIONS 27 CITATIONS

[SEE PROFILE](#)



Carita Paradis

Lund University

108 PUBLICATIONS 1,111 CITATIONS

[SEE PROFILE](#)



Andreas Kerren

Linnaeus University

156 PUBLICATIONS 1,053 CITATIONS

[SEE PROFILE](#)

Some of the authors of this publication are also working on these related projects:



The London-Lund Corpus 2 of spoken British English (LLC 2) [View project](#)



How the human mind makes use of contraries in everyday life. A new multidimensional approach to contraries in perception, language, reasoning and emotions. [View project](#)

All content following this page was uploaded by [Vasiliki Simaki](#) on 25 April 2018.

The user has requested enhancement of the downloaded file.

Detection of Stance-Related Characteristics in Social Media Text

Vasiliki Simaki

Centre for Corpus Approaches to Social Science,
Lancaster University, LA1 4YW, United Kingdom
Centre for Languages and Literature,
Lund University, Sweden
v.simaki@lancaster.ac.uk

Carita Paradis

Centre for Languages and Literature, Lund University
221 00 Lund, Sweden
carita.paradis@englund.lu.se

Panagiotis Simakis

XPLAIN
153 44 Athens, Greece
simakis@xplain.co

Andreas Kerren

Department of Computer Science and Media Technology,
Linnaeus University
351 95 Växjö, Sweden
andreas.kerren@lnu.se

ABSTRACT

In this paper, we present a study for the identification of stance-related features in text data from social media. Based on our previous work on stance and our findings on stance patterns, we detected stance-related characteristics in a data set from Twitter and Facebook. We extracted various corpus-, quantitative- and computational-based features that proved to be significant for six stance categories (CONTRARIETY, HYPOTHETICALITY, NECESSITY, PREDICTION, SOURCE OF KNOWLEDGE, and UNCERTAINTY), and we tested them in our data set. The results of a preliminary clustering method are presented and discussed as a starting point for future contributions in the field. The results of our experiments showed a strong correlation between different characteristics and stance constructions, which can lead us to a methodology for automatic stance annotation of these data.

CCS CONCEPTS

• **Computing methodologies** → **Natural language processing**; **Lexical semantics**;

KEYWORDS

stance-taking, text clustering, feature extraction, social media text

ACM Reference Format:

Vasiliki Simaki, Panagiotis Simakis, Carita Paradis, and Andreas Kerren. 2018. Detection of Stance-Related Characteristics in Social Media Text. In *SETN '18: 10th Hellenic Conference on Artificial Intelligence, July 9–15, 2018, Rio Patras, Greece*. ACM, New York, NY, USA, 7 pages. <https://doi.org/10.1145/3200947.3201017>

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

SETN '18, July 9–15, 2018, Rio Patras, Greece

© 2018 Copyright held by the owner/author(s). Publication rights licensed to the Association for Computing Machinery.

ACM ISBN 978-1-4503-6433-1/18/07...\$15.00

<https://doi.org/10.1145/3200947.3201017>

1 INTRODUCTION

The detection and analysis of speaker stance, or stance-taking, is a topic in semantics and discourse analysis where real communication between people and their interaction is in focus. Researchers in this field observe and identify the verbal ways that speakers use to make assessments and/or position themselves towards a topic, an event, and/or an idea. The literature in this field covers a wide range of various concepts related to stance such as modality [9, 17], evidentiality [8, 12, 29], evaluation/appraisal [25, 31], subjectivity [6, 11, 41, 42], and sentiment [20, 21, 39, 40, 44]. In our previous work on stance, we adopted Du Bois' stance triangle [7], in order to provide a definition on stance-taking as "the performance by humans in communication, actions taken by speakers to express their beliefs, evaluations and attitudes toward (i) objects, scenes and events, and (ii) toward propositions, and speakers' viewpoints on what is talked about", and to propose an original functional-cognitive framework based on notional stance categories [35].

The detection of stance in discourse has recently been addressed from a text mining perspective. The identification of text characteristics related to stance, and the classification of texts according to the stance expressed has been an important methodology in most works on this topic. In an effort to go beyond sentiment analysis tasks, and enlighten the readers about the speakers' positioning towards a topic/event/idea, researchers proposed various methods and tools to capture stance in discourse automatically, namely, whether the speaker is in favor of or against the given topic/event/idea. The automatic identification of speaker stance is an important means for the monitoring of opinions in wide audiences, and it has already an increasing use in texts extracted from social media.

We focus on discourse extracted from social media, since this is the most popular channel that people of various backgrounds use on a daily basis to express their opinions about topics such as politics, music, lifestyle, environment, or personal matters. This activity produces a massive number of sound/video data, images, and texts everyday that needs to be further analysed and grouped according to different criteria that are set in each case. Text data from social media provide important information about social media users, their preferences, habits and the trends. Monitoring the social media users' positioning towards different topics (e.g., product reviews, news, social issues) is an interesting task with potential

applications in other disciplines such as sociology, political science, management and marketing studies. Besides that, there is still an active interest from a linguistics perspective, as new stance-related linguistic characteristics can be observed from a text type that is relatively new. The plethora of social media data, and the availability of tools for their analysis support linguistic studies, in which new insights about the language of positioning can be discovered.

In our previous work on stance, we proposed an original framework, based on notional stance categories [35]. We compiled a manually annotated corpus according to this framework, and performed various qualitative, quantitative and computational tasks in order to identify patterns in language that are associated to each stance. We highlighted six out of ten stance categories as the most frequent ones in our corpus: CONTRARIETY, HYPOTHETICALITY, NECESSITY, PREDICTION, SOURCE OF KNOWLEDGE and UNCERTAINTY, and showed that stances such as CONTRARIETY and NECESSITY, are more distinctive than other [33]. In this paper, we applied the corpus-, quantitative-, and computational-based features that were highlighted as important in our previous studies in a data set extracted from social media. In this preliminary study, we managed to detect these features in non-annotated data from social media on the basis of which we performed clustering experiments. We grouped our data into six clusters in an effort to associate these groups to the corresponding stance categories. Our results pointed to the challenge of this task, and the difficulty of applying the features and method in the specific set of data. We analysed the results, and compared them to the previous work, debating whether a methodology for the automatic stance annotation of social media data according to our framework is possible.

The remainder of this paper is organised as follows. Section 2 describes the background work of this study. In Section 3, the methodology of this study and the data used are described. Section 4 presents and discusses our experimental work and findings. Finally, Section 5 concludes this paper.

2 BACKGROUND WORK

Previous studies in automatic identification of stance focus their interest on whether a speaker supports a fact/event/idea or not, and researchers mostly perform classification experiments of texts into pro or con stance classes. Stance classification is connected to the fields of Subjective Language Identification [44], Opinion Mining, and Sentiment Analysis [28].

In one of the early studies in the field, Whitelaw et al. [43] combined functional taxonomies with APPRAISAL, and performed sentiment analysis tasks. They discovered that the semantic information that the APPRAISAL categories carry improved the sentiment classification. The fact that sentiment and stance are closely related concepts in text mining is also supported by the findings of a recent study by Mohammed et al. [27], which showed that the sentiment features improved the stance classification results, achieving an accuracy up to 69%. Stance identification methods have been applied to data extracted from ideological online debates [1, 13–16, 37, 38], in data from online sources towards political topics [2, 3, 26], in student essays [10], and in microblogging such as Twitter posts [30].

The majority of these studies addressed stance-taking as a binary issue, however, recent studies used a wider spectrum of stance concepts to support or deny a claim/rumour [4, 5, 45]. The identification of stance has also been addressed from an Information Visualisation perspective namely the uVSAT tool for visual stance analysis created by Kucher et al. [23]. This tool contains multiple approaches for analysing temporal and textual data as well as exporting stance markers in order to prepare a stance-oriented training data set.

Within the StaViCTA project¹ project, we addressed speaker stance from an interdisciplinary perspective combining knowledge and tools from the fields of semantics, corpus linguistics, computational linguistics, and visualisation. As a first step, we proposed an original stance framework consisting of ten notional stance categories: AGREEMENT/DISAGREEMENT, CERTAINTY, CONTRARIETY, HYPOTHETICALITY, NECESSITY, PREDICTION, SOURCE OF KNOWLEDGE, TACT/ RUDENESS, UNCERTAINTY, and VOLITION. These stance categories were identified and attributed to sentences extracted from blog sources thematically related to the 2016 UK referendum. Two experts manually annotated these sentences with semantic criteria using the ALVA annotation tool [22], and the final output of this procedure resulted to the Brexit Blog Corpus (BBC)², a gold standard resource of 1,682 sentences (35,492 words in total). BBC was evaluated, the inter-coder reliability was calculated, and the stance categories and their co-occurrences were discussed and analysed. Since the beginning, the compilation of BBC was aiming to be evaluated statistically and computationally, in order (i) to test the proposed framework's efficiency, and (ii) to provide new insights and linguistic patterns for the identification of stance in discourse.

In the quantitative analysis of the Brexit Blog Corpus that followed [33], we aimed to identify features that determine the formal profiles of six stance categories (CONTRARIETY, HYPOTHETICALITY, NECESSITY, PREDICTION, SOURCE OF KNOWLEDGE, and UNCERTAINTY) in a subset of the BBC. The study had two parts: firstly, it examined a large number of formal linguistic features such as punctuation, words and grammatical categories occurred in the sentences in order to describe the specific characteristics of each category, and secondly, it compared characteristics in the entire data set in order to determine stance similarities in the data set. We showed that among the six stance categories in the corpus, CONTRARIETY and NECESSITY were the most discriminative ones, with the former using longer sentences, more conjunctions, more repetitions and shorter forms than the sentences expressing other stances. NECESSITY had longer lexical forms but shorter sentences, which were syntactically more complex.

In our most recent study [34], we used the same BBC subset of the annotated data as a springboard to approach stance identification from the opposite point of view, namely from how stance is realised in text. Our aim was to identify specific constructions that are related to the six stance categories, and to this end, we followed a two-step experimental procedure. Firstly, we performed a quantitative analysis of the annotated corpus data in order to identify significant lexical forms that are stance-specific for each category. Secondly, we performed a meta-annotation procedure of the data.

¹StaViCTA project: <http://cs.lnu.se/stavicta/>

²Publicly available here: <https://snd.gu.se/sv/catalogue/study/snd1037>

One of the BBC annotators was asked to single out the constructions that triggered his annotation decision in the previous study where functional-semantic criteria were only used. We compared the results of the two techniques, and proposed a list of constructions of stanced discourse as particularly salient expressions of each stance type.

Apart from the quantitative and analytical studies presented above, we also performed a text classification task in order to show that these categories can be automatically detected. In this study [32], we proposed a large set of lexical and syntactic linguistic features that were tested, and classification experiments were implemented using different algorithms. We achieved accuracy of up to 30% for the six-class experiments, which was not fully satisfactory. As a second step, we calculated the pair-wise combinations of the stance categories. Confirming our quantitative study [33], the CONTRARIETY and NECESSITY binary classification achieved the best results with up to 71% accuracy. This result was encouraging and highlighted the fact that each stance category has a different level of distinctness, with CONTRARIETY and NECESSITY as the most discriminative ones.

3 METHODOLOGY AND DATA

In this section, we describe the methodology that was followed throughout this study, and the data that we used to perform our experiments.

3.1 Methodology

In Section 2, we described our previous work on the identification of stance in texts, and the linguistic resource that was created and annotated for that purpose. BBC is a gold standard corpus with highly informative and accurately annotated content, but it faces some limitations in terms of generalisation and applicability. There are two major issues we realised that need to be dealt with in our work on stance after BBC. Firstly, the stance-related characteristics from our previous studies need to be evaluated against a different set of data. The stance features that are corpus-, quantitative-, and computational-based were detected and identified as significant entities of stanced sentences extracted from a specific text type (blog posts and comments) towards a specific thematic area (2016 UK referendum). It is an important task to confirm these findings using a different data set consisting of text chunks covering a wider thematic orientation. Secondly, it is very important to confirm or not the efficiency of the proposed stance framework in order to use it for the annotation of other text data too. The efficiency of our stance framework is two-pronged. It aims (i) to examine whether our framework covers to a sufficient extent the large spectrum of the different stances that people take when positioning towards a topic/event/idea, and (ii) to estimate the frequency of these stances in discourse, and the linguistic patterns used to express each stance.

In our first study on this topic [35], we already observed the paucity of some of the proposed stances in the BBC (AGREEMENT/DISAGREEMENT, CERTAINTY, TACT/RUDENESS and VOLITION), and, in the following studies, we continued with six stances (the six most frequent stances in the BBC). In this study, we followed the same principle, aiming to bring together all the stance-related characteristics in order to evaluate them in a data set from social media.

The characteristics that we focused in our previous tasks were at character-, word-, syntactic-level, extracted manually, automatically, or statistically from our data set depending the respective methodology that was followed. In Table 1, we present all our statistical features for the identification of each stance category.

As reported in Section 2, the six-class classification accuracy of the BBC subset did not prove to be satisfactory (30%). In order to shed some light on this problem, and test the suitability of the features used in the given task, we performed feature selection experiments, using the ReliefF algorithm proposed by Koronenko et al. [19]. This algorithm improves the reliability of the probability approximation since it is robust to incomplete data and generalised to multiclass problems. ReliefF was implemented using the WEKA toolkit³, and the nine highest-ranked features are shared with features of the statistical analysis: *word length*, *character number/sentence*, *word number/sentence*, *commas frequency*, *short words frequency*, *punctuations frequency*, *hapax legomena frequency*, *different forms frequency*, and *digits frequency*. The fact that these features are also statistically significant confirmed the existence of linguistic patterns in stanced sentences. In our list of features, we also added the stance constructions identified in our latest study [34]. In this work, we derived corpus-based constructions that were stance-related by following a quantitative and a qualitative analysis (as described in Section 2). The output of these two methods is presented in Table 2. With bold, we present the stance constructions that were confirmed by both methods, and the rest of the entries in the table are the constructions that were highlighted by only one method. In order to distinguish the two construction classes, we attributed a different coefficient to each of them: we weighted with 1 the constructions confirmed by two methods, and with 0.5 the constructions confirmed by one method.

The purpose of this study was to detect these features in a different data set. To this end, we used a subset of our social media corpus, as described in Section 3.2. In these data, that are not annotated according to our stance framework, we estimated the stance-related features as described above. Then, we implemented clustering experiments in order to split the data set into six clusters of texts showing similar characteristics. This method may lead to an automatic way of using these characteristics for the attribution of social media texts to distinctive classes associated to our stance categories, and therefore, to an automatic annotation process. Our methodology aims to contribute not only to the text mining community, but also to corpus linguistics, computational linguistics and semantics since this is the first study using a framework based on notional stance categories.

3.2 Data Description

Here, we used data from our social media text corpus [36]. This data set consists of 712,033 posts (13,424,523 words and 89,347,103 characters in total). The posts were extracted from the official Facebook and Twitter profiles of public figures like actors, authors, singers, athletes, politicians, and they were annotated with the author's sociodemographic information. To extract the data, we used the *Facepager* software [18]. The average size of the corpus posts is 125 characters per post. The topics discussed vary from personal

³WEKA toolkit: <https://www.cs.waikato.ac.nz/ml/weka/>

Table 1: The features derived from the Brexit Blog Corpus, for each stance category. The + symbol means that this feature was significant and had the *highest* distribution/frequency in the data of the specific stance; while the - symbol shows that the feature was significant but had the *lowest* distribution in the specific data.

CONTRARIETY	HYPOTHETICALITY	NECESSITY	PREDICTION	SOURCE OF KNOWLEDGE	UNCERTAINTY
+ word number/sent. + commas freq. + adverbs freq. + conjunctions freq. + hapax dislegomena freq. - fullstops freq. - hapax legomena freq. - different forms freq.	+ spaces freq. + punctuation freq. + short word freq. - adjectives freq.	+ verbs freq. + word length + different forms freq. + hapax legomena freq. + fullstops freq. + pronouns freq. - word number/sent. - character number/sent. -hapax legomena freq. - commas freq. - adverbs freq.	- verbs freq.	+ digits freq. + nouns freq. + prepositions freq. - conjunctions freq. - short words freq. - pronouns freq.	+ adjectives freq. - prepositions freq.

Table 2: The stance constructions of the BBC subset. The constructions confirmed by two methods are highlighted in bold.

CONTRARIETY	HYPOTHETICALITY	NECESSITY	PREDICTION	SOURCE OF KNOWLEDGE	UNCERTAINTY
but not while are than and	if would a/an we could in will	must need/needs should to about let who have	be may will to going back might next think it is not the	as has that his said was/were the by he I it show to	could I maybe may might probably think

branding to opinions about social and political matters, nature, etc. The corpus was compiled from September to December 2015, and data from 838 different users (535 male and 302 female users) were manually annotated with information about the author’s gender, age, professional activity, national variety of the English, and any other additional information available such as his/her educational background or professional details. The annotation labels were given according to the information that the users provide about themselves in their social media accounts, and in some cases according to the information that Wikipedia⁴ entries or other internet sources provide (as most authors are well-known personalities).

For this study, we extracted a subset of data that are close to the BBC sentences in terms of the average sentence length. The mean length of the BBC sentences is 21 words per sentence. The data set we extracted consists of 156,156 sentences (3,190,973 words in total) with an average length of 20.33 words per sentence.

4 EXPERIMENTS

The first step of this study was to estimate the features presented in Section 1 in the data set. We used the BBC features that showed to be highly ranked according to the feature selection process. The feature extraction process was implemented in Python, and we used spacy⁵ for the sentence parsing and the POS-tagging. Our final feature set consists of 21 features, presented in Table 3 with a short description.

After the feature extraction process, we normalised our features, and we used the scikit-learn library⁶, and more specifically the k-means algorithm for the data clustering into 6 groups. The k-means method [24] is one of the methods that use the euclidean distance, which minimizes the sum of the squared euclidean distance between the data points and their corresponding cluster centers. In order to visualise the results of the clustering process, we performed a Primary Component Analysis (PCA) dimensionality reduction

⁴https://en.wikipedia.org/wiki/Main_Page

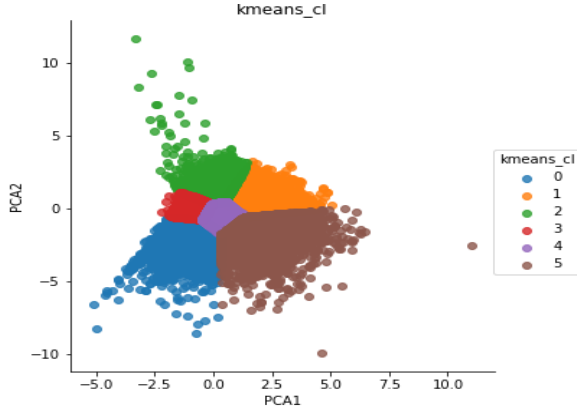
⁵spacy: <https://spacy.io/>

⁶scikit-learn library: <http://scikit-learn.org/stable/>

Table 3: The feature set of this study in an alphabetical order.

Feature name	Description
adj	Adjectives frequency
adv	Adverbs frequency
average_sent_len_chars	Character number/sentence
average_sent_len_words	Word number/sentence
average_word_length	Average word length/sentence
comma_freq	Commas frequency
conj	Conjunctions frequency
contrariety	CONTRARIETY constructions
digit_freq	Digits frequency
fullstop_freq	Fullstops frequency
hapax	Hapax legomena frequency
hypotheticality	HYPOTHETICALITY constructions
n_diff_words	Different forms frequency
necessity	NECESSITY constructions
noun	Nouns frequency
prediction	PREDICTION constructions
pron	Pronouns frequency
source_of_knowledge	SOURCE OF KNOWLEDGE constructions
spaces_freq	Spaces frequency
uncertainty	UNCERTAINTY constructions
verb	Verbs frequency

method (in two dimensions, PCA1 and PCA2), and the results are presented in Figure 1.

**Figure 1: The plot of the six clusters in two-dimensional space.**

As we see in Figure 1, the text clusters are very close to each other, which means that the distance between the texts clustered in the six groups are close in terms of the features that were extracted. Since our data are not annotated, it is hard to decide which cluster is associated to which stance category.

In Figure 2, we present the correlation matrix for our features. We observe a strong correlation between features that are dependent from each other, such as the different metrics for the sentence length (average_sent_len_chars and average_sent_len_words), or the dependency between the frequency of verbs and pronouns. The

higher saturated blue shows the strong correlation between the feature couples. An interesting finding derived from this matrix, is the dependency between the stance construction features of the different stance categories. This fact can be explained in terms of the shared constructions among the six features. What turns out to be problematic in this case is that the results of the previous study on this topic (the one that proposed the list of stance constructions) highlighted stop-words among these constructions. Instead of looking for more complex stance expressions that can be significant for the identification of the corresponding stance, we searched for the forms that were detected as important, but it turned out that for the current task these forms were not effective. A possible solution to this may be the extraction of more complex features.

The clustering results highlighted that our data needs to be filtered in a more refined way, and that more features need to be derived and added. One possible way to deal with the short distance between the six clusters is to identify the texts that do not express any stance, by distinguishing in a first step the neutral from the stanced texts. After a closer look in the clusters' content, we realised that apart from texts where stance can be identified, there are plenty stance-free texts, where the content of the post can be characterised as neutral. In many of these cases, the text follows a picture upload, as in the example *"hey look! this stunning #portrait of me by @rosswatsonart is being used to promote a #nudeart?"*. Also, many of the texts may not express any stance, but they are highly polarised in terms of sentiment, like in the example *"happy anniversary to the man who changed my life and made every day a dream come true! love you @andywgrant #fb"*. Another category of texts that can be characterised as stance-free are promotional posts with self-branding content, as the example *"looking for some new music for the weekend? check my #liveinthefuture top 10 chart at beatport <http://t.co/2cpl73ixiy>"*. This is a category of texts in our data set that are expected considering their source (public figures, celebrities, etc.). Therefore, a first grouping of the data into neutral and stanced can be a future step that can improve our results.

The challenge of this task is also located in the different type of the BBC sentences and the social media text data. Both data sources are social media networks, but differences can be identified, especially in the cases where the data are extracted from Twitter, in which the non-linguistic characteristics (e.g., the hashtags and the mentions) are very salient. Additionally, the features extracted from BBC may be too fitting to the specific resource, and unefficient or unsuitable to be used in different data.

5 CONCLUSIONS

In this study, we tested stance features from our previous studies in a data set from social media texts. We grouped these data into six clusters in order to identify our stance categories, and associate them to the corresponding clusters. We analysed our results, and the correlation of the clustering features highlighted that the features we used for the specific task need to be refined and enriched. This investigation is the basis for the task we aim namely, the automatic annotation of text data in terms of notional stance categories. We derived important conclusions about the nature of the task, about the features, and the method. As highlighted in all our previous studies within this field, stance identification is a challenging task,

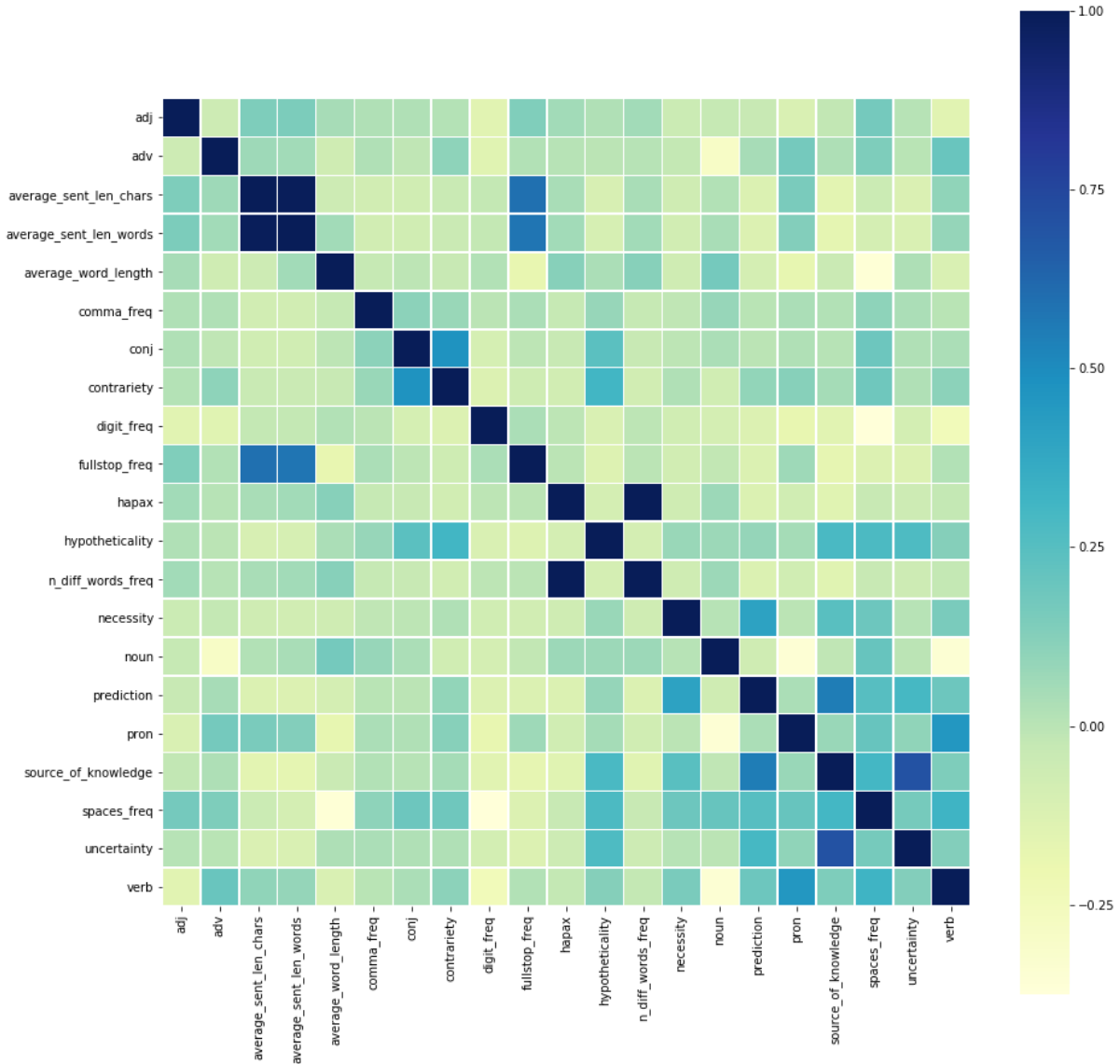


Figure 2: Correlation matrix of the extracted features.

especially when more refined stance concepts are employed without the important information that a manually annotated resource can provide.

ACKNOWLEDGMENTS

This research is funded by the StaViCTA project, supported by the Swedish Research Council (framework grant the Digitized Society Past, Present, and Future, No. 2012-5659).

REFERENCES

- [1] Pranav Anand, Marilyn Walker, Rob Abbott, Jean E Fox Tree, Robeson Bowmani, and Michael Minor. 2011. Cats rule and dogs drool!: Classifying stance in on-line debate. In *Proceedings of the 2nd workshop on computational approaches to subjectivity and sentiment analysis*. Association for Computational Linguistics, 1–9.
- [2] Isabelle Augenstein, Tim Rocktäschel, Andreas Vlachos, and Kalina Bontcheva. 2016. Stance detection with bidirectional conditional encoding. *arXiv preprint arXiv:1606.05464* (2016).
- [3] Isabelle Augenstein, Andreas Vlachos, and Kalina Bontcheva. 2016. Usfd at semeval-2016 task 6: Any-target stance detection on twitter with autoencoders. In *Proceedings of the 10th International Workshop on Semantic Evaluation (SemEval-2016)*. 389–393.
- [4] Roy Bar-Haim, Indrajit Bhattacharya, Francesco Dinuzzo, Amrita Saha, and Noam Slonim. 2017. Stance Classification of Context-Dependent Claims. In *Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics: Volume 1, Long Papers*, Vol. 1. 251–261.
- [5] Roy Bar-Haim, Lilach Edelstein, Charles Jochim, and Noam Slonim. 2017. Improving claim stance classification with lexical knowledge expansion and context utilization. In *Proceedings of the 4th Workshop on Argument Mining*. 32–38.
- [6] Emile Benveniste. 1971. Subjectivity in language. *Problems in general linguistics* 1 (1971), 223–30.

- [7] John W Du Bois. 2007. The stance triangle. In *Stancetaking in discourse: Subjectivity, evaluation, interaction*, Robert Englebretson (Ed.). Amsterdam: John Benjamins, 139–182.
- [8] Lena Ekberg and Carita Paradis. 2009. Evidentiality in language and cognition. *Functions of Language* 16, 1 (2009), 5–7.
- [9] Roberta Faccinetti, Frank Palmer, and Manfred Krug. 2003. *Modality in contemporary English*. Vol. 44. Walter de Gruyter.
- [10] Adam Faulkner. 2014. Automated classification of stance in student essays: An approach using stance target information and the Wikipedia link-based measure. *Science* 376, 12 (2014), 86.
- [11] Dylan Glynn and Mette Sjölin. 2015. Subjectivity and Epistemicity. In *Subjectivity and Epistemicity: Stance strategies in discourse and narration*, Dylan Glynn and Mette Sjölin (Eds.). Lund University Press.
- [12] Xiang Gu. 2015. Evidentiality, subjectivity and ideology in the Japanese history textbook. *Discourse & Society* 26, 1 (2015), 29–51.
- [13] Kazi Saidul Hasan and Vincent Ng. 2013. Extra-linguistic constraints on stance recognition in ideological debates. In *Proceedings of the 51st Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, Vol. 2. 816–821.
- [14] Kazi Saidul Hasan and Vincent Ng. 2013. Frame semantics for stance classification. In *Proceedings of the Seventeenth Conference on Computational Natural Language Learning*. 124–132.
- [15] Kazi Saidul Hasan and Vincent Ng. 2013. Stance classification of ideological debates: Data, models, features, and constraints. In *Proceedings of the Sixth International Joint Conference on Natural Language Processing*. 1348–1356.
- [16] Kazi Saidul Hasan and Vincent Ng. 2014. Why are you taking this stance? identifying and classifying reasons in ideological debates. In *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*. 751–762.
- [17] Issa Kanté. 2010. Mood and modality in finite noun complement clauses: A French-English contrastive study. *International journal of corpus linguistics* 15, 2 (2010), 267–290.
- [18] Till Keyling and Jakob Jünger. 2013. Facepager (Version, fe 3.3). *An application for generic data retrieval through APIs* (2013).
- [19] Igor Kononenko, Edvard Šimec, and Marko Robnik-Šikonja. 1997. Overcoming the myopia of inductive learning algorithms with RELIEFF. *Applied Intelligence* 7, 1 (1997), 39–55.
- [20] Athanasia Koumpouri, Iosif Mporas, and Vasileios Megalooikonomou. 2015. Evaluation of Four Approaches for Sentiment Analysis on Movie Reviews: The Kaggle Competition. In *Proceedings of the 16th International Conference on Engineering Applications of Neural Networks (INNS)*. ACM, 23.
- [21] Athanasia Koumpouri, Iosif Mporas, and Vasileios Megalooikonomou. 2015. Opinion Recognition on Movie Reviews by Combining Classifiers. In *International Conference on Speech and Computer*. Springer, 309–316.
- [22] Kostiantyn Kucher, Carita Paradis, Magnus Sahlgren, and Andreas Kerren. 2017. Active learning and visual analytics for stance classification with ALVA. *ACM Transactions on Interactive Intelligent Systems (TiIS)* 7, 3 (2017), 14.
- [23] Kostiantyn Kucher, Teri Schamp-Bjerede, Andreas Kerren, Carita Paradis, and Magnus Sahlgren. 2016. Visual analysis of online social media to open up the investigation of stance phenomena. *Information Visualization* 15, 2 (2016), 93–116.
- [24] James MacQueen et al. 1967. Some methods for classification and analysis of multivariate observations. In *Proceedings of the fifth Berkeley symposium on mathematical statistics and probability*, Vol. 1. Oakland, CA, USA, 281–297.
- [25] James R Martin and Peter R White. 2003. *The language of evaluation*. Vol. 2. Springer.
- [26] Saif Mohammad, Svetlana Kiritchenko, Parinaz Sobhani, Xiao-Dan Zhu, and Colin Cherry. 2016. A Dataset for Detecting Stance in Tweets.. In *LREC*.
- [27] Saif M Mohammad, Parinaz Sobhani, and Svetlana Kiritchenko. 2017. Stance and sentiment in tweets. *ACM Transactions on Internet Technology (TOIT)* 17, 3 (2017), 26.
- [28] Bo Pang, Lillian Lee, et al. 2008. Opinion mining and sentiment analysis. *Foundations and Trends® in Information Retrieval* 2, 1–2 (2008), 1–135.
- [29] Kristen Precht. 2003. Stance moods in spoken English: Evidentiality and affect in British and American conversation. *TEXT-THE HAGUE THEN AMSTERDAM THEN BERLIN- 23*, 2 (2003), 239–258.
- [30] Ashwin Rajadesingan and Huan Liu. 2014. Identifying users with opposing opinions in Twitter debates. In *International conference on social computing, behavioral-cultural modeling, and prediction*. Springer, 153–160.
- [31] Jonathon Read and John Carroll. 2012. Annotating expressions of appraisal in English. *Language resources and evaluation* 46, 3 (2012), 421–447.
- [32] Vasiliki Simaki, Carita Paradis, and Andreas Kerren. 2017. Stance classification in texts from blogs on the 2016 British referendum. In *International Conference on Speech and Computer*. Springer, 700–709.
- [33] Vasiliki Simaki, Carita Paradis, and Andreas Kerren. 2018. Evaluating stance-annotated sentences from the Brexit Blog Corpus: A quantitative linguistic analysis. *ICAME Journal* 42 (2018).
- [34] Vasiliki Simaki, Carita Paradis, and Andreas Kerren. 2018. A two-step procedure to identify stance constructions in discourse from political blogs. *Forthcoming* (2018).
- [35] Vasiliki Simaki, Carita Paradis, Maria Skeppstedt, Magnus Sahlgren, Kostiantyn Kucher, and Andreas Kerren. 2017. Annotating speaker stance in discourse: the Brexit Blog Corpus. *Corpus Linguistics and Linguistic Theory* (2017).
- [36] Vasiliki Simaki, Panagiotis Simakis, Carita Paradis, and Andreas Kerren. 2017. Identifying the Authors' National Variety of English in Social Media Texts. In *Proceedings of the International Conference Recent Advances in Natural Language Processing, RANLP 2017*. 671–678.
- [37] Swapna Somasundaran and Janyce Wiebe. 2010. Recognizing stances in ideological on-line debates. In *Proceedings of the NAACL HLT 2010 Workshop on Computational Approaches to Analysis and Generation of Emotion in Text*. Association for Computational Linguistics, 116–124.
- [38] Dhanya Sridhar, Lise Getoor, and Marilyn Walker. 2014. Collective stance classification of posts in online debate forums. In *Proceedings of the Joint Workshop on Social Dynamics and Personal Attributes in Social Media*. 109–117.
- [39] Maite Taboada. 2016. Sentiment analysis: an overview from linguistics. *Annual Review of Linguistics* 2 (2016), 325–347.
- [40] Marjan Van de Kauter, Bart Desmet, and Véronique Hoste. 2015. The good, the bad and the implicit: a comprehensive approach to annotating explicit and implicit sentiment. *Language resources and evaluation* 49, 3 (2015), 685–720.
- [41] Arie Verhagen. 2005. *Constructions of intersubjectivity: Discourse, syntax, and cognition*. Oxford University Press on Demand.
- [42] Peter RR White. 2003. Beyond modality and hedging: A dialogic view of the language of intersubjective stance. *TEXT-THE HAGUE THEN AMSTERDAM THEN BERLIN- 23*, 2 (2003), 259–284.
- [43] Casey Whitelaw, Navendu Garg, and Shlomo Argamon. 2005. Using appraisal groups for sentiment analysis. In *Proceedings of the 14th ACM international conference on Information and knowledge management*. ACM, 625–631.
- [44] Janyce Wiebe, Theresa Wilson, Rebecca Bruce, Matthew Bell, and Melanie Martin. 2004. Learning subjective language. *Computational linguistics* 30, 3 (2004), 277–308.
- [45] Arkaitz Zubiaga, Elena Kochkina, Maria Liakata, Rob Procter, Michal Lukasik, Kalina Bontcheva, Trevor Cohn, and Isabelle Augenstein. 2018. Discourse-aware rumour stance classification in social media using sequential classifiers. *Information Processing & Management* 54, 2 (2018), 273–290.