

Narrative detection in online patient communities

Anne Dirkson

Suzan Verberne

Wessel Kraaij

Leiden Institute of Advanced Computer Science, Leiden University
Niels Bohrweg 1, 2333 CA Leiden, the Netherlands
`{a.r.dirkson, s.verberne, w.kraaij}@liacs.leidenuniv.nl`

Abstract

Although narratives on patient forums are a valuable source of medical information, their systematic detection and analysis has so far been limited to a single study. In this study, we examine whether psycho-linguistic features or document embeddings can aid identification of narratives. We also investigate which features distinguish narratives from other social media posts. This study is the first to automatically identify the topics discussed in narratives on a patient forum. Our results show that for classifying narratives, character 3-grams outperform psycho-linguistic features and document embeddings. We found that narratives are characterized by the use of past tense, health-related words and first-person pronouns, whereas non-narrative text is associated with the future tense, emotional support words and second-person pronouns. Topic analysis of the patient narratives uncovered fourteen different medical topics, ranging from tumor surgery to side effects. Future work will use these methods to extract experiential patient knowledge from social media.

1 Introduction

Nowadays, online patient forums are the main medium by which patients exchange their narratives. These narratives mainly recount their own experiences with their condition. As such, they contain experiential knowledge [Bor76], defined as the knowledge that patients gain from their own experiences. In recent years, such experiential knowledge has increasingly been recognized as valuable and complementary to empirical knowledge [CBC⁺13]. Consequently, more health-related applications are making use of patient forum data, for instance to track public health trends [SOG⁺16] and to detect adverse drug responses [SGN⁺15]. Experiential knowledge is also valuable for patients themselves: patients indicate that they strongly rely on experiences and information provided on patient forums [SHBL16]. This is especially true for patients with a rare disease, for which medical professionals often lack expertise and the number of studies is limited [AKG08].

To understand the experiential knowledge on patient forums, forum posts that contain narratives must first be identified. As of yet, research into systematically distinguishing patient narratives on patient forums is limited to a single study on Dutch forum data [VBSEng], which uses words as only features. We expand upon this work using a different data set by examining whether document embeddings and psycho-linguistic features can improve the identification of patient narratives. We expect so, because these aggregated features are less dependent on

Copyright © 2019 for the individual papers by the paper’s authors. Copying permitted for private and academic purposes. This volume is published and copyrighted by its editors.

In: A. Jorge, R. Campos, A. Jatowt, S. Bhatia (eds.): Proceedings of the Text2StoryIR’19 Workshop, Cologne, Germany, 14-April-2019, published at <http://ceur-ws.org>

individual terms, which may overlap significantly between narratives and factual statements about the same topic. Secondly, we explore how narratives differ from other types of posts by studying which features are influential in identifying narratives and which posts are classified incorrectly. Thirdly, we analyze how prevalent narratives are on a cancer patient forum and which topics these narratives discuss.

2 Related Work

Narratives on patient forums have mainly been studied qualitatively (e.g. [vUKDT⁺09]). The automatic identification of narratives on a patient forum is limited to the study by Verberne *et al.* [VBSEng] on a Dutch cancer forum. They identified narratives with a F_1 of 0.911 using only the lower-cased words of the posts as features. They also found that various linguistic factors (1st person singular, 3rd person and negations) and psychological processes (social processes and religion) were correlated with the presence of narratives. These psycho-linguistic features were measured using the Linguistic Inquiry and Word Count (LIWC) method [TP10].

Additionally, research into self-reported adverse drug responses (ADRs) has led to the development of classifiers for differentiating between factual statements of ADRs and personal experiences of ADRs on social media [BY12, NSO⁺15, SG15]. However, these classifiers are highly specific and thus not suitable for identifying patient narratives in general.

Another closely related field is the classification of personal health mentions on social media i.e. posts that mention a person who is affected as well as their specific condition, such as: ‘my granddad has Alzheimer’s’. Presently, only two studies have investigated this task. The first by Lamb *et al.* [LPD13] focused on separating flu awareness from actual flu reports on social media. More recently, Karisani *et al.* [KA18] introduced WESPAD, a classifier for personal health mentions, which attains state-of-the-art performance for seven different health domains including stroke, depression and flu infection. Nonetheless, a personal health mention alone is not sufficient to consider the post a narrative, and thus these classifiers are also inadequate for our purpose.

3 Methods

3.1 Data

Our data consists of an open, international Facebook forum for patients with Gastrointestinal Stromal Tumor (GIST)¹. It is moderated by GIST Support International and consists of 36,722 posts with a median length of 20 tokens.

3.2 Preprocessing

The data was lowercased and tokenized with NLTK. Due to the noisy nature of user-generated content, especially in the spelling of medical terms, we applied a tailored preprocessing pipeline² to our data. Firstly, an existing normalization pipeline for social media [Sar17]³ was used to normalize tokens to American English and to expand generic abbreviations used on social media. Hereafter, domain-specific abbreviations were expanded with a lexicon of 42 non-ambiguous abbreviations, generated based on 1000 posts and annotated by a domain expert and the first author. Spelling mistakes were detected using a combination of relative frequency and edit distance to possible candidates and corrected using weighted Levenshtein distance. Correction candidates were derived from the corpus itself. Drug names were normalized using the RxNorm database [Nat]. Non-English posts were removed using *langid* [LB12]. Punctuation was removed, but stop words were not, as we expect function words to play a role in the expression of narratives.

3.3 Supervised classification

3.3.1 Manual annotation of example data

We randomly selected 1050 posts for annotation. The annotators were asked to indicate per message whether it contains a personal experience. They were not provided with its context. Personal experiences did not need to be about the author but could be about someone else. This definition was based on earlier work by Verberne *et al.* [VBSEng] and van Uden-Kraan *et al.* [vUKDT⁺08]. The first 50 posts were annotated individually by the

¹<https://www.facebook.com/groups/gistsupport/>

²The preprocessing scripts can be found at: <https://github.com/AnneDirkson/LexNorm>

³<https://bitbucket.org/asarker/simplenormalizerscripts>

first author and another PhD student to improve the annotation guidelines.⁴ The remaining 1000 posts were divided equally into six sets of 200 posts, with 40 posts (20%) overlapping between all sets. The overlap was used to calculate the pairwise Cohen’s kappa. There were seven annotators in total: six PhD students and one GIST patient. Each sample was assigned to an annotator, apart from one sample which was divided between two PhD students. To be able to include the overlapping sample in the classification, we opted to use the annotations of the GIST patient for these 40 posts.⁵

3.3.2 Feature sets

Four feature sets were derived from the text data: word unigrams, character n-grams (using the CountVectorizer function in sklearn), psycho-linguistic features, and document embeddings. For both word unigrams and character n-grams, we investigated whether TF-IDF weighting would improve performance compared to raw counts. Additionally, we explored whether stemming or lemmatising the data prior to extracting the unigrams could improve performance. Psycho-linguistic features were based on the LIWC 2015 [TP10]. Punctuation categories were discarded, resulting in 82 LIWC features in total. LIWC is a well-known method for investigating psychological processes in text and includes both linguistic (e.g. first-person pronouns) and psychological categories (e.g. positive emotions). The last feature set consisted of document embeddings: a doc2vec model [LM14] was trained on the labeled training data for each fold in the cross-validation. We combine a distributed memory model with a distributed bag of words model, as recommended by Le and Mikolov [LM14]. We also attempted to train document embeddings first on the unsupervised data and then re-train on the supervised data, but this led to nonsensical classification features.

3.3.3 Supervised classification algorithms

Classifiers were evaluated separately for each feature set. We ignored all posts that had been left empty by the annotator (the annotator chose neither yes nor no): three posts were ignored for this reason. For word unigrams, character n-grams and psycho-linguistic features, we compared four sklearn classification algorithms: Multinomial Naive Bayes (MNB), linear Support Vector Classification (LinearSVC), Stochastic Gradient Descent (SGD) with log loss, and K Nearest Neighbours (KNN). These were chosen according to the following criteria: (1) known to perform well on text data, (2) recommended for small data sets and (3) able to calculate probabilistic outcomes. The latter enabled us to use probabilistic ensembles. The doc2vec representations combined with Logistic Regression were used as classifier in itself: the document representations were tagged with the labels of the training data. This model was then used to derive vector representations for new documents. To test if a combination of feature types could improve performance, we evaluated soft voting (argmax of the sums of the predicted probabilities) of the best individual classifiers for the best performing variants of each feature set. Significance testing was done with pair-wise t-tests.

To evaluate the performance, the average F_1 score of a 10-fold cross validation was used. For each run, hyper-parameters were tuned for that specific training set using a 10-fold grid search on the training data. The tuning grids were based on sklearn documentation: C from 10^{-3} to 10^3 (steps of x10) for LinearSVC and Logistic Regression; number of neighbors from 3 to 11 (steps of 2) for KNN; and max iterations from 2 to 2048 (steps of x2) and alpha from 10^{-8} to 10^{-2} (steps of x10) for SGD. The dimensionality of the document vectors was tuned with a grid of 100 to 400 (steps of 100).

3.4 Topic modelling of the whole data set

To label the remaining data, the best performing classifier was used with the hyper-parameter settings that were optimal in the majority of the training sets. To investigate which topics are discussed in the patient narratives, we used topic modelling with non-Negative Matrix Factorization of the TF-IDF weighted tokens without stopwords. Topic coherence, measured using TC-W2V [OGCC15], was used to select the number of topics. Topic labels were assigned manually by exploring the words with the highest weights and the top-ranked (i.e. most relevant) messages per topic.

⁴The annotation guidelines can be found at: <https://github.com/AnneDirkson/NarrativeFilter>

⁵The annotated data is available upon request in order to protect the privacy of the patients

4 Results

4.1 Annotated data

The data was slightly imbalanced, with 37.7% of the posts containing a narrative, resulting in a majority baseline of roughly 0.62. The inter-annotator agreement was substantial ($\kappa = 0.69$).

4.2 Classifier evaluation

A Linear SVC on character 3-grams achieves the highest F_1 score (Table 1), although character 4-grams ($p = 0.526$), stemmed unigrams ($p = 0.930$) and lemmatised unigrams ($p = 0.587$) do not perform significantly worse. Character 5- and 6-grams also do not perform worse overall ($p = 0.122$ and $p = 0.169$), but their recall is significantly lower ($p = 0.023$ and $p = 0.029$). The classifiers for the best performing document embeddings (DBOW+DM) and psycho-linguistic features, however, are significantly worse overall than character 3-grams ($p = 0.0055$ and $p = 0.026$ respectively). Employing TF-IDF weighting does not aid any of the unigram or character n-gram features. Additionally, neither feature selection ($F_1=0.761$) nor word boundaries ($F_1=0.796$) improve the performance of character 3-grams. Using a range of character n-grams, namely 3-to-4 ($F_1=0.814$), 3-to-5 ($F_1=0.814$), or 3-to-6 ($F_1=0.812$), also does not boost performance.

Ensemble classification did not perform better than character 3-grams alone (see Table 2). Nevertheless, an ensemble of all four feature types is significantly more precise than all other classifiers ($p = 0.0048$ compared to the second best). To further explore why ensemble classification does not manage to improve overall performance, we investigated the predictions of individual classifiers. As can be seen in Table 3, there is a high degree of overlap between the predictions based on character 3-grams and the other feature sets (88.3%, 83.8% and 84.4% respectively). Consequently, the vast majority of the predictions cannot be improved by complementing character 3-grams with these feature sets. Interestingly, 4.7% of the posts are misclassified by all feature sets. Considering the non-overlapping predictions, the percentage of correct predictions was higher for character 3-grams than for either document embeddings or psycho-linguistic features in a pair-wise comparison. Thus, it appears that adding these features would be more detrimental than beneficial to narrative classification.

Table 1: Mean Test Score (10-fold CV) For Best Classifiers Per Feature Set

Feature set		Size	Classifier	F_1 (SD)	Recall (SD)	Precision (SD)
Unigrams	Original	4,078	SGD	0.795 (0.025)	0.788 (0.074)	0.811 (0.055)
	Stemmed	3,205	SGD	0.814 (0.031)	0.793 (0.047)	0.840 (0.049)
	Lemmatised	3,777	SGD	0.808 (0.039)	0.810 (0.059)	0.813 (0.070)
Character n-grams	3-grams	5,086	SVC	0.815 (0.035)	0.844 (0.047)	0.793 (0.058)
	4-grams	16,496	SVC	0.811 (0.027)	0.827 (0.068)	0.844 (0.029)
	5-grams	36,349	SGD/SVC	0.796 (0.023)	0.784 (0.059)	0.817 (0.069)
	6-grams	60,443	SGD	0.793 (0.040)	0.797 (0.042)	0.795 (0.079)
LIWC		82	SVC	0.773 (0.031)	0.805 (0.044)	0.752 (0.077)
Doc2vec	DBOW	400	LogReg	0.737 (0.029)	0.751 (0.056)	0.735 (0.066)
	DM	400	LogReg	0.762 (0.039)	0.749 (0.062)	0.785 (0.070)
	DM+DBOW	800	LogReg	0.772 (0.037)	0.803 (0.064)	0.749 (0.055)

Table 2: Mean Test Score (10-fold CV) For Ensemble Classification. * DM+DBOW variant.

Feature sets	F_1 (SD)	Recall (SD)	Precision (SD)
3-grams + LIWC + Doc2vec* + Stemmed Unigrams	0.770 (0.029)	0.703 (0.065)	0.859 (0.053)
3-grams + LIWC + Doc2vec*	0.795 (0.037)	0.772 (0.072)	0.829 (0.065)
3-grams + LIWC	0.706 (0.032)	0.624 (0.059)	0.828 (0.073)
3-grams + Doc2vec*	0.755 (0.048)	0.735 (0.089)	0.786 (0.040)

Table 3: Comparison of Predictions of Classifiers for Different Feature Sets. * DM+DBOW variant.

Compared to		Both		Difference	
		Correct(%)	Incorrect(%)	In Favour of 3-grams(%)	In Favour of Other Method(%)
Character 3-grams	LIWC	75.0	8.8	8.4	7.7
	Doc2Vec*	74.8	9.6	8.6	6.9

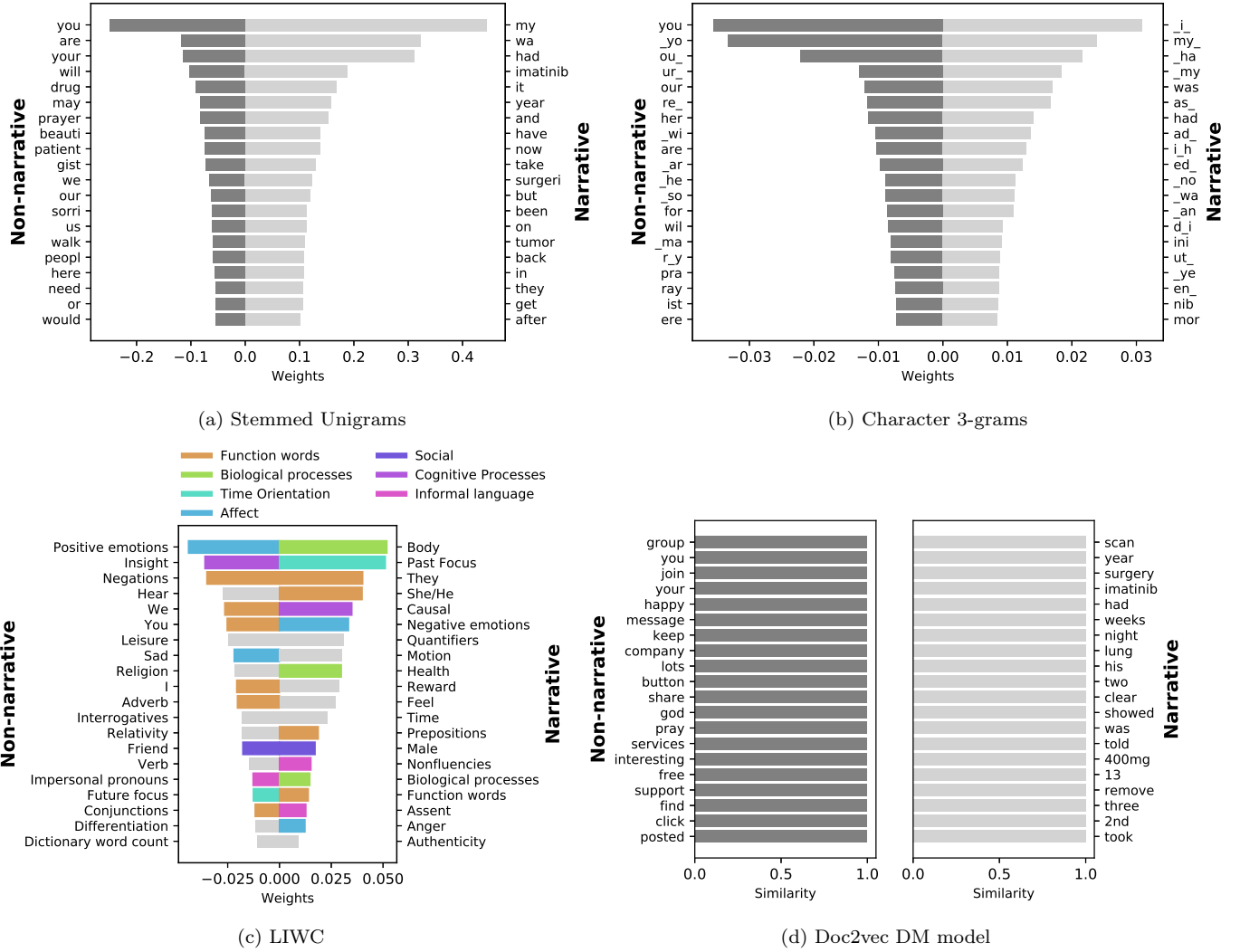


Figure 1: The 20 Most Influential Features In Individual Classifiers. In (b) underscores represent spaces.

4.3 Influential features

Narratives are typically distinguished by terms relating to the past tense (*was*, *had*, *years*), health (*imatinib*, *tumor*, *surgeri*) and first-person narrative (*my*, *i*) (see Figure 1). This is corroborated by the character 3-grams, psycho-linguistic features and document embeddings. Some of the important terms for non-narrative texts are also health-related (*patients*, *gist*) and first-person narrative (*we*, *us*), which showcases the difficulty of the task at hand. In general, non-narrative texts seem to focus more on emotional support (*prayer*, *share*, *may*), second-person narrative (*you*, *your*) and the future (*may*, *will*). The psycho-linguistic features additionally reveal that narratives contain more mentions of causality and negative emotions. In contrast, non-narrative texts seem to contain more positive emotions. Lastly, as predicted, function words appear important for classifying narratives in social media, and it is thus advisable to not remove stopwords.

4.4 Error analysis for the best performing classifier

Error analysis reveals that a significant proportion of the errors is due to incorrect annotation: 36.9% of the false positives and 36.2% of the false negatives were labelled incorrectly (see Table 4). Specifically, annotators have difficulty correctly labelling discussions about personal medical facts or side effects as narratives (e.g. *i have been on imatinib 5 months and lost 1/3 of my hair*). Conversely, annotators may incorrectly judge posts that give emotional support, external information or advice to be narratives while they are not (e.g. *i may be wrong but*

Table 4: Error Analysis for best classifier (Character 3-gram Classification of Narratives)

False positives		False negatives	
Reasons for misclassification	Frequency	Reasons for misclassification	Frequency
Mislabelling	24	Mislabelling	17
Emotional support/thanks	15	Unknown	12
Information/advice	13	Lack of context	7
Lack of context	7	Question	5
Question	4	Non-medical narratives	3
Unknown	1	Hypothetical	1
Empty post	1	Empty post	2
TOTAL	65	TOTAL	47

total gastrectomy sounds very extreme for two small gist’).

The incorrect labelling may have impacted the automated classification such that these categories are also more difficult for the computer to distinguish. The classifier does, however, appear to outperform human judgement and to some extent ‘correct’ their mistakes. In fact, its performance may be underestimated by the metrics based on these incorrect labels. Other types of posts that appears challenging for the computer are posts that lack context or contain questions. The former are often answers to unknown questions posed earlier in the thread.

4.5 Frequency and content of patient narratives

4.5.1 Automated narrative detection in unsupervised data

The percentage of narratives in the unlabelled data is 37.0 %, which is comparable to the annotated sample. This results in a total of 13.436 posts for topic modelling.⁶

4.5.2 Topic modelling

The TC-W2V metric [OGCC15] identifies the optimal number of topics to be fourteen. The resulting topics relate to different aspects of the medical process for GIST patients (see Table 5). Note that imatinib is the most commonly used medication.

5 Discussion

The detection of narratives was most optimal when using character 3-grams. Their strength is in their ability to cluster relevant word types based on suffixes and prefixes. This is especially relevant in the medical domain e.g. all cancer medication for GIST ends in ‘*nib*’. In contrast, psycho-linguistic features appear to suffer from oversimplification, because they aggregate words that define *different* classes into one category e.g. *we* and *my* into the umbrella category of first person pronouns (see Figure 1). The use of document embeddings may have been hampered by the small size of the data. An alternative explanation could be that incorrect labelling impacts these features more strongly than word-based features.

Narratives could be differentiated most strongly by their use of past tense, first-person narrative and health-related words. The first two are in line with linguistic definition of a narrative. The stronger focus on health, however, may indicate that patients prefer to share their own health experiences than health information from external sources.

Annotating narratives appears a challenging task, despite providing annotators with a guideline based on previous work [VBSEng] and validated through initial annotation by two annotators. This is underscored by our inter-annotator agreement ($\kappa = 0.69$) which was comparable to that of Verberne *et al.* [VBSEng] ($\kappa = 0.71$). Our classifier performed less well than their system ($F_1 = 0.91$), which may be explained by their larger sample of annotated data (2.051 posts).

Inevitably, our results depend on the choice of what constitutes a narrative and how annotators interpret this definition. It appears that especially the line between a medical fact about oneself and a medical experience is fuzzy for annotators. Future studies could perhaps use this knowledge to develop clearer guidelines.

⁶The code for unsupervised narrative filtering is shared at: <https://github.com/AnneDirkson/NarrativeFilter>

Table 5: Most Important Topics Discussed In Patient Forum Narratives. Topic labels were assigned manually.
 * Cancer medication

Topic labels	Top 10 words	Top-ranked post for the topic
Tumor location	tumor stomach removed liver small cm mitotic metastases rate intestine	‘i only had one tumor on my stomach’
(Emotional) Coping	take get time doctor like also know imatinib* day would	‘i completely understand i started 400 imatinib after surgery in and have lots of bad days [...]
Duration of Treatment	years imatinib* almost ago 10 taking two still 11 12	‘about 1 and 1/2 years’
Types of Scans	scan ct pet results next today last showed week cat	‘oops one is a ct scan and one is a pet scan’
Diagnosis of GIST	gist diagnosed cancer specialist oncologist husband anyone ago surgeon found	‘that was my gist’
Other Medication	sunitinib* regorafenib* sorafenib* imatinib* working 37 exon nilotinib* trial stopped drug	‘i have this on sunitinib’
Side Effects	side effects imatinib* effect different fatigue eyes bad 400mg time	‘and no side-effects’
Tumor Surgery	surgery remove since weeks first post surgeon second shrink done	‘just had surgery’
Absence of Tumor Recurrence	disease evidence still years today post since resection year far	‘no evidence of disease no evidence of disease’
Recurrence of Work, Medication or Tumor	back came come hair go went weeks took coming lost	‘i started imatinib after i went back to work’
Emotional support	good luck news best far hope bad goes well keep pretty	‘all my best and good luck’
Dosage of Medication	mg 400 800 imatinib* 600 take day taking since started	‘11 years of imatinib since 2003 at 600 mg and since november 2009 at 800 mg [...]
Timing of Scans	months every scans three ct six year two first month	‘my doctor said 3 years’
Ingesting imatinib	one year last took imatinib* day another old got time	‘take imatinib’

6 Conclusion

For the detection of patient narratives on social media, psycho-linguistic features and document embeddings are outperformed by character 3-grams. These narratives are associated with the past tense, health and first-person pronouns, whereas non-narrative text is associated with the future tense, emotional support and second-person pronouns. The patient narratives could be subdivided into discussions of fourteen different medical topics, ranging from surgery to side effects. Future work will develop automated methods for the extraction of patient knowledge from the narratives.

7 Acknowledgements

This work was financed by the SIDN fonds. The authors also thank H. Vos, G. Wiggers, W. Verschoof, A. Brandsen, D. Gawehns, P. Dhar, M. Vinkenoog and G. van Oortmerssen of Leiden University for annotating the data.

References

- [AKG08] Ségolène Aymé, Anna Kole, and Stephen Groft. Empowerment of patients: lessons from the rare diseases community. *The Lancet*, 371(9629):2048–2051, 2008.
- [Bor76] Thomasina Borkman. Experiential Knowledge: A New Concept for the Analysis of Self-Help Groups. *Social Service Review*, 50(3):445–456, 1976.
- [BY12] Jiang Bian and Fan Yu. Towards Large-scale Twitter Mining for Drug-related Adverse Events. In *SHB12*, pages 25–32, 2012.

- [CBC⁺13] Pam Carter, Roger Beech, Domenica Coxon, Martin J. Thomas, and Clare Jinks. Mobilising the experiential knowledge of clinicians, patients and carers for applied health-care research. *Contemporary Social Science*, 8(3):307–320, 2013.
- [KA18] Payam Karisani and Eugene Agichtein. Did you really just have a heart attack? *Proceedings of the 2018 World Wide Web Conference on World Wide Web - WWW 18*, 2018.
- [LB12] Marco Lui and Timothy Baldwin. langid.py: An Off-the-shelf Language Identification Tool. In *Proceedings of the 50th annual meeting of the association of computational linguistics*, pages 25–30, 2012.
- [LM14] Quoc Le and Tomas Mikolov. Distributed Representations of Sentences and Documents. In *Proceedings of the 31st international conference on machine learning*, 2014.
- [LPD13] Alex Lamb, Michael J Paul, and Mark Dredze. Separating Fact from Fear: Tracking Flu Infections on Twitter. In *Proceedings of NAACL-HLT*, pages 789–795, 2013.
- [Nat] National Library of Medicine (US). Rxnorm.
- [NSO⁺15] Azadeh Nikfarjam, Abeed Sarker, Karen O’Connor, Rachel Ginn, and Graciela Gonzalez. Pharmacovigilance from social media: mining adverse drug reaction mentions using sequence labeling with word embedding cluster features. *Journal of the American Medical Informatics Association: JAMIA*, 22(3):671–81, 2015.
- [OGCC15] Derek O’Callaghan, Derek Greene, Joe Carthy, and Pádraig Cunningham. An analysis of the coherence of descriptors in topic modeling. *Expert Systems with Applications*, 42(13):5645–5657, 2015.
- [Sar17] Abeed Sarker. A customizable pipeline for social media text normalization. *Social Network Analysis and Mining*, 7(45), 2017.
- [SG15] Abeed Sarker and Graciela Gonzalez. Portable automatic text classification for adverse drug reaction detection via multi-corpus training. *Journal of Biomedical Informatics*, 53:196–207, 2015.
- [SGN⁺15] Abeed Sarker, Rachel Ginn, Azadeh Nikfarjam, Karen O’Connor, Karen Smith, Swetha Jayaraman, Tejaswi Upadhaya, and Graciela Gonzalez. Utilizing social media data for pharmacovigilance: A review. *Journal of Biomedical Informatics*, 54:202–212, 2015.
- [SHBL16] Edin Smailhodzic, Wyanda Hooijsma, Albert Boonstra, and David J. Langley. Social media use in healthcare: A systematic review of effects on patients and on their relationship with healthcare professionals. *BMC Health Services Research*, 16(1):442, 2016.
- [SOG⁺16] Abeed Sarker, Karen O’Connor, Rachel Ginn, Matthew Scotch, Karen Smith, Dan Malone, and Graciela Gonzalez. Social Media Mining for Toxicovigilance: Automatic Monitoring of Prescription Medication Abuse from Twitter. *Drug Safety*, 39(3):231–240, 2016.
- [TP10] Yla R. Tausczik and James W. Pennebaker. The psychological meaning of words: LIWC and computerized text analysis methods. *Journal of Language and Social Psychology*, 29(1):24–54, 2010.
- [VBSEng] Suzan Verberne, Anika Batenburg, Remco Sanders, and Mies Van Eenbergen. Social processes of online empowerment on a cancer patient discussion form: using text mining to analyze linguistic patterns of empowerment processes. *JMIR Cancer*, Forthcoming.
- [vUKDT⁺08] Cornelia F. van Uden-Kraan, Constance H. Drossaert, Erik Taal, Bret R Shaw, Erwin R. Seydel, and Mart A. F. J. van de Laar. Empowering processes and outcomes of participation in online support groups for patients with breast cancer, arthritis, or fibromyalgia. *Qualitative Health Research*, 18(3):405–417, 2008.
- [vUKDT⁺09] Cornelia F. van Uden-Kraan, Constance H.C. Drossaert, Erik Taal, Erwin R. Seydel, and Mart A. F. J. van de Laar. Participation in online patient support groups endorses patients’ empowerment. *Patient Education and Counseling*, 74(1):61–69, 2009.