

Identifying Idealised Vectors for Emotion Detection Using CMA-ES

Mohammed Alshahrani
School of Computer Science and
Electronic Engineering
University of Essex
Colchester, Essex
ma16562@essex.ac.uk

Spyridon Samothrakis
Institute for Analytics and Data
Science
University of Essex
Colchester, Essex
ssamot@essex.ac.uk

Maria Fasli
Institute for Analytics and Data
Science
University of Essex
Colchester, Essex
mfasli@essex.ac.uk

ABSTRACT

Detecting the emotional content of text is one of the most popular NLP tasks. In this paper, we propose a new methodology based on identifying “idealised” words that capture the essence of an emotion; we define these words as having the minimal distance (using some metric function) between a word and the text that contains the relevant emotion (e.g. a tweet, a sentence). We look for these words through searching the space of word embeddings using CMA-ES. Our method produces state of the art results, surpassing classic supervised learning methods.

CCS CONCEPTS

• **Computing methodologies** → **Supervised learning**;

KEYWORDS

Evolution Strategy, Emotion Detection

ACM Reference Format:

Mohammed Alshahrani, Spyridon Samothrakis, and Maria Fasli. 2019. Identifying Idealised Vectors for Emotion Detection Using CMA-ES. In *Genetic and Evolutionary Computation Conference Companion (GECCO '19 Companion)*, July 13–17, 2019, Prague, Czech Republic. ACM, New York, NY, USA, 2 pages. <https://doi.org/10.1145/3319619.3322057>

1 INTRODUCTION

Written text is used widely as a means of communication; it is often the case that textual communications online involve short informal messages. One of the most common Natural Language Processing (tasks) is trying to identify the emotions evoked by a sentence. The NLP methods that attack this problem fall broadly under the umbrella of Sentiment Analysis (SA) and emotion detection. Whereas SA is used to identify positive, neutral or negative polarity, emotion detection techniques are used to classify more specific emotions like anger, fear, sadness and love.

In start contract to the dearth of unsupervised learning methods, emotion detection has been performed extensively on tweets [1, 4, 7] using supervised learning methods. In this paper, we propose a

new method for performing supervised learning. We search the word vector space for an “idealised” emotional vector. In order to identify emotion, we adopt the Euclidean Distance function to calculate the distance between the word embedding of tweets and the “idealised” emotional vector. To the best of our knowledge, this work is the first to explore detecting emotions from tweets by adopting an evolutionary strategy with Word2Vec through distance functions. We show state of the art results in the twitter dataset of Mohammad [6].

The rest of the paper is organised as follows. The following section presents the dataset used and benchmark approach. Section 3 will present the main methods and models this paper is using. In Section 4, we describe the experimental work. Section 5 discusses the results obtained. The paper closes with the conclusions drawn from this work and outlines potential future work.

2 PRELIMINARIES

2.1 Datasets

Twitter Emotion Corpus (TEC) dataset consist of 21,051 emotional self-labeled tweets [6]. Only tweets that include one of the basic six emotions hash-tags were compiled. A detailed description of the dataset is out of the scope of this paper, while interested readers can refer to [6] for all relevant information. We note that the distribution of the emotions in the TEC dataset is imbalanced. For example, joy has the highest percentage by 39.1% while the percentage of disgust is less than 4%. Moreover, for the ground truth, the hash-tags that have been used as labels were removed from the tweets in the dataset thus it can be used for training and testing.

2.2 Benchmark Approach

For each emotion of the basic six emotions Mohammad [6] developed a binary classifier to identify emotions using WEKA [2]. For example, Fear-NotFear classify whether a tweet express fear or not. Sequential Minimal Optimization (SMO) was the algorithm chosen in [6] to be used with SVM; binary features are used to classify whether unigrams are present or absent. Unigrams that occurred less than two times were discarded. The author reports the mean results for 10-fold cross-validation, and we follow this approach.

3 CMA-ES

The covariance matrix adaptation evolution strategy (CMA-ES) is a stochastic optimization algorithm proposed by Hansen et al, [3]. The CMA-ES is an iterative evolutionary algorithm that is based on a set of solutions for continuous optimization. It generates λ

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 ACM 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.

GECCO '19 Companion, July 13–17, 2019, Prague, Czech Republic

© 2019 Association for Computing Machinery.

ACM ISBN 978-1-4503-6748-6/19/07...\$15.00

<https://doi.org/10.1145/3319619.3322057>

new vector solutions from a multivariate Gaussian distribution according to: $x_i \sim N(m_k, \sigma_k^2 C_k)$ for $i = 1, \dots, \lambda$ with mean m_k , covariance matrix C_k and step size σ_k . In this work, we use CMA-ES as a black-box optimisation method in order to search through a high-dimensional search space without using gradient information.

4 EVOLUTIONARY SEARCH EXPERIMENT

Using the TEC dataset, in this experiment we will try to find the “idealised” vectors that will perform optimally in emotion detection using the CMA-ES as an Evolutionary algorithm. Euclidean Distance is used to calculate the distance between the tweets and the “idealised” vectors. Ten-fold cross validation is followed as in the benchmark approach [6]. The general procedure is as follows: The vectors are retrieved for each word of a tweet from Google’s pre-trained Word2Vec model [5]. Then the mean of the vectors is calculated for each tweet. The dataset (the mean vector for each tweet) is split into ten groups. For each unique group: 1- Take one group as the test dataset. 2- Take the remaining groups as a training dataset. 3- Fit a CMA-ES model of the training set and evaluate it on the test set. The CMA-ES generations is set to 200 and the population to 100 and the dimensionality to 1800 because a vector of 300 dimensions is needed for each emotion category: anger, disgust, fear, joy, sadness and surprise. Initialize the σ step size with 0.1. For each generation the exact algorithm is presented in Algorithm 1. while, computational time is 10s per generation.

Algorithm 1: Evolutionary Search

```

for each generation do
  select mini batches (8000) randomly from training data.
  for each element in population do
    calculate euclidean distance between the “idealised”
    vector and the mean vector of each tweet of the mini
    batch.
    set the emotion of the tweet based on the shortest
    distance.
    calculate the F1 score.
  end
  update the parameters for next generation based on the
  results and the current “idealised” vector.
  if generation is even then
    use the “idealised” vector to evaluate the model on the
    test data set.
    retain evaluation scores.
  end
end

```

Table 1 shows the average F1 and the error bounds for the 95%th confidence interval for the ten-fold cross validation of the last generation test of the evolutionary search experiment. It can be seen that the evolutionary search results exceeded the benchmark

Table 1: Experiments results

F1	anger	disgust	fear	joy	sadness	surprise
benchmark	27.9	18.7	50.6	62.4	38.7	45
evolutionary search	33.82 ± 1.87	29.07 ± 3.16	52.30 ± 1.24	66.88 ± 0.56	43.31 ± 0.79	47.24 ± 0.87

approach [6]. The evolutionary search results achieved the best results over all emotions, Table 1.

5 DISCUSSION

Table 1 demonstrates that the results of the evolutionary search experiment have surpassed all previous published benchmarks. As mentioned in the datasets section, disgust has very low percentage in the dataset and, as expected, it has the larger confidence intervals, see table 1. In contrast, joy has the largest percentage in the dataset therefore all runs end up in more or less similar results, see Table 1.

6 CONCLUSIONS

Our approach using the CMA-ES to continuously update the “idealised” vector for each emotion in order to be used for emotion detection using the Word2Vec embedding boosts the results in all fronts. To the best of our knowledge, this work is the first to detect emotions from tweets by adopting an evolutionary strategy with word embeddings. For future work, we are aiming to explore using different evolutionary strategies and techniques in a similar manner - this work forms a baseline. Furthermore, the same evolutionary search method could be used to learn transformations in other data beyond tweets e.g. images in order to detect emotions for instance. Finally, we plan on undertaking further analysis and try to identify which words correspond to or are closer to what the evolutionary algorithm discovered in the embeddings search space.

ACKNOWLEDGMENTS

We acknowledge support from grant number ES/L011859/1, from The Business and Local Government Data Research Centre, funded by the Economic and Social Research Council to provide researchers and analysts with secure data services.

REFERENCES

- [1] Rakesh C Balabantaray, Mudasir Mohammad, and Nibha Sharma. 2012. Multi-class twitter emotion classification: A new approach. *International Journal of Applied Information Systems* 4, 1 (2012), 48–53.
- [2] Mark Hall, Eibe Frank, Geoffrey Holmes, Bernhard Pfahringer, Peter Reutemann, and Ian H Witten. 2009. The WEKA data mining software: an update. *ACM SIGKDD explorations newsletter* 11, 1 (2009), 10–18.
- [3] Nikolaus Hansen, Sibylle D Müller, and Petros Koumoutsakos. 2003. Reducing the time complexity of the derandomized evolution strategy with covariance matrix adaptation (CMA-ES). *Evolutionary computation* 11, 1 (2003), 1–18.
- [4] Maryam Hasan, Emmanuel Agu, and Elke Rundensteiner. 2014. Using hashtags as labels for supervised learning of emotions in twitter messages. In *ACM SIGKDD Workshop on Health Informatics, New York, USA*.
- [5] Tomas Mikolov, Ilya Sutskever, Kai Chen, Greg S Corrado, and Jeff Dean. 2013. Distributed representations of words and phrases and their compositionality. In *Advances in neural information processing systems*. 3111–3119.
- [6] Saif M Mohammad. 2012. # Emotional tweets. In *Proceedings of the First Joint Conference on Lexical and Computational Semantics-Volume 1: Proceedings of the main conference and the shared task, and Volume 2: Proceedings of the Sixth International Workshop on Semantic Evaluation*. Association for Computational Linguistics, 246–255.
- [7] Wenbo Wang, Lu Chen, Krishnaprasad Thirunarayan, and Amit P Sheth. 2012. Harnessing twitter “big data” for automatic emotion identification. In *Privacy, Security, Risk and Trust (PASSAT), 2012 International Conference on and 2012 International Conference on Social Computing (SocialCom)*. IEEE, 587–592.