

Cornell Belief and Sentiment System at TAC 2017

Kai Sun

Department of Computer Science
Cornell University
Ithaca, NY 14853
ks985@cornell.edu

Claire Cardie

Department of Computer Science
Cornell University
Ithaca, NY 14853
cardie@cs.cornell.edu

Abstract

In this paper we describe the 2017 system of the CornMich team for the TAC Belief and Sentiment (BeSt) task for Chinese.

1 Introduction

The Cornell-Michigan team (aka CornMich) submitted 3 Chinese runs for NIST TAC Belief and Sentiment (BeSt) Track 2017. We employed a rule-based system for belief, and a hybrid system for sentiment. In the following sections, we present the design for belief and sentiment respectively.

2 Belief

Our system is based on the approach of the majority baseline of the BeSt Evaluation 2016 (Rambow et al., 2016).

2.1 Target Extraction

Our system regards every relation/event in the ERE input as a target with type="cb" and polarity="pos".

2.2 Source Extraction

Given a target, our system looks for the post/article where the mention text/trigger of the target first appears, and uses the post author as the source.

2.3 Submissions

Confidence is set to 1 for every belief. 3 runs are identical.

3 Sentiment

Our system is built upon Cornell's Chinese sentiment system at TAC 2016 (Niculae et al., 2016).

3.1 Target Extraction

Figure 1 illustrates the flow of sentiment target extraction. The input is any target candidate (i.e. any entity, relation, or event in the ERE input), and the output is the polarity of the input (i.e. positive, negative, or none). Below is a brief description of each module.

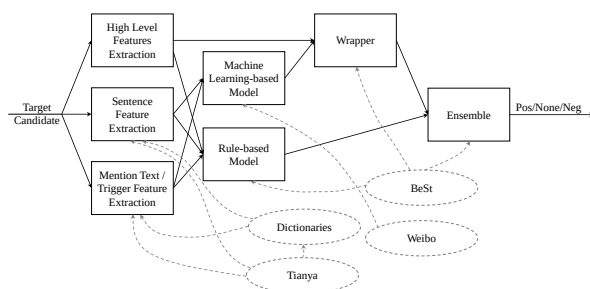


Figure 1: Sentiment Target Extraction. Rectangles and dashed circles represent modules and data respectively.

The three *feature extraction* modules extract word vectors (trained with posts crawled from Tianya), part-of-speech tags, character/word/phrase level sentiment-related information from dictionaries¹, indicators of ERE, the target candidate type, the document genre, the number of entities in the sentence that the target candidate belongs to etc..

The *machine learning-based model* is a neural network for sentence-level sentiment analysis. The neural network is composed of a single LSTM layer, an average pooling layer followed by a softmax layer. It is trained with about 4k sentences from Weibo with polarity annotated. The sentiments of the mention text/trigger and the sentence

¹While most dictionaries are collected from the Internet, we have one induced by semi-supervised learning from Tianya posts (i.e. the dashed arrow from *Tianya* to *Dictionaries* in Figure 1).

output by this model are combined by the *wrapper*.

The *rule-based model* is model with a bunch of genre-specific and domain-specific rules. Specifically, for newswire, it has rules for dealing with diplomatic language. for discussion forum, it has rules for dealing with informal languages, especially morphs in Chinese Internet Slang.

In the last step, the *rule-based model* and the *wrapper* are tuned on BeSt 16 gold data separately, and then ensembled together.

3.2 Source Extraction

For discussion forum, sentiment source extraction is the same as belief source extraction. For newswire, our system outputs up to 2 sources for each target. One is the post author, which is the same as belief source extraction. The other is the first entity (if exists) on the left side of the reporting speech (such as “say”, “mention”) around which the mention text/trigger of the target appears.

3.3 Submissions

We have 7 versions of the system optimized by different F_β measure, and confidence of a sentiment is set based on how many versions by which the sentiment is reported. Run 1, 2, 3 are focused on good F_1 , precision, recall respectively.

References

- Vlad Niculae, Kai Sun, Xilun Chen, Yao Cheng, Xinya Du, Esin Durmus, Arzoo Katiyar, and Claire Cardie. Cornell belief and sentiment system at TAC 2016. In *TAC*, 2016.
- Owen Rambow, Meenakshi Alagesan, Michael Arrigo, Daniel Bauer, Claire Cardie, Adam Dalton, Hoa Dang, Mona Diab, Greg Dubbin, Jason Duncan, Gregorios Katsios, Axinia Radeva, Tomek Strzalkowski, and Jennifer Tracey. The 2016 TAC KBP BeSt evaluation. In *TAC*, 2016.