

Weakly Supervised Deep Embedding Product Review for Sentiment Analysis

¹B. Sushma, ²B. Navyasree, ³B. Vyshnavi, ⁴C. Lakshmi prasanna

^{1,2,3,4} UG Student, Department of CSE,

Dr K V Subba Reddy College Of Engineering For Women, Kurnool, Andhra Pradesh, India

Article Info

Page Number: 454 - 459

Publication Issue:

Vol 69 No. 1 (2020)

Abstract

Product reviews are valuable for upcoming buyers in helping them make decisions. To this end, different opinion techniques have been proposed, where judging a review sentence's orientation (e.g. positive or negative) is one of their key challenges. Recently, deep learning has emerged as an effective means for solving sentiment classification problems. A neural network intrinsically learns a useful representation automatically without human efforts.

However, the success of deep learning highly relies on the availability of large-scale training data. We propose a novel deep learning framework for product review sentiment classification which employs prevalently available ratings as weak supervision signals.

The framework consists of two steps:

(1) learning a high-level representation (an embedding space) which captures the general sentiment distribution of sentences through rating information;

(2) adding a classification layer on top of the embedding layer and use labeled sentences for supervised fine-tuning. We explore two kinds of low-level network structure for modeling review sentences, namely, convolution feature extractors and long short-term memory

Article History

Article Received: 15 December 2019

Revised: 27 January 2020

Accepted: 22 February 2020

Publication: 28 March 2020

1. INTRODUCTION

Sentiment Analysis can be considered a classification process as illustrated in fig1. There are three main classification levels in SA: document-level, sentence-level, and aspect-level SA. Document-level SA aims to classify an opinion document as expressing a positive or negative opinion or sentiment. It considers the whole document a basic information unit (talking about one topic). Sentence-level SA aims to classify sentiment expressed in each sentence. The first step is to identify whether the sentence is subjective or objective. If the sentence is subjective, Sentence-level SA will determine whether the sentence expresses positive or negative opinions. Wilson et al. have pointed out that sentiment expressions are not necessarily subjective in nature. However, there is no fundamental difference between document and sentence level classifications because sentences are just short documents. Classifying text at the document level or at the sentence level does not provide the necessary detail needed opinions on all aspects. Popular sentiment classification methods generally fall into two categories: lexicon-based methods and machine learning methods. Lexicon-based methods typically take the tack of first constructing a sentiment lexicon of opinion words (e.g. “wonderful”, “disgusting”), and then design classification rules based on appeared opinion words and prior syntactic knowledge. Despite effectiveness, this kind of methods require substantial efforts in lexicon construction and

rule design. Furthermore, lexicon-based methods cannot well handle implicit opinions, i.e. objective statements such as “I bought the mattress a week ago, and a valley appeared today” . As pointed out in ,this is also an important form of opinions. Factual information is usually morehelpful than subjective feelings. Lexicon-based methods can only deal with implicit opinions in an ad-hoc way.

The first machine learning based sentiment classification work applied popular machine learning algorithms such as Naive Bayes to the problem. After that, most research in this direction revolved around feature engineering for better classification performance. Different kinds of features have been explored ,e.g. n-grams, Part-of-speech (POS) information and syntactic relations, etc. Feature engineering also costs a lot of human efforts,and a feature set suitable for one domain may not generate good performance for otherdomains. In recent years, deep learning has emerged as an effective means for solving sentiment classification problems. A deep neural network intrinsically learns a high level representation of the data, thus avoiding laborious work such as feature engineering. A second advantage is that deep models have exponentially stronger expressive power than shallow models. However, the success of deep learning heavily relies on the availability of large-scale training data. Labeling a large number of sentences is verylaborious.

For example: We are taking a review as an exam

2. LITERATUREREVIEW

We are employing Amazon Web Services, which belongs to the SAS (Software As Service) subcategory of cloud computing. There are some fundamental terms in it: Edge location, availability zone, and region. There are two or more availability zones in each region. The Content Delivery Network (CDN) endpoints for CloudFront are located at Edge Location, which is merely a data center in the Availability Zone. We are putting in place the Amazon Rekognition Service on AWS due to its built-in machine learning capabilities, which allow you to analyze millions of images and curate and organize enormous amounts of visual data.

A technology platform that is long-lasting and secure is Amazon Web Services. Amazon's data centers and services have multiple layers of physical and operational security to protect your data. In addition, AWS regularly audits its infrastructure to ensure its security. It provides complete privacy and security, as well as guarantees the availability, integrity, and confidentiality of your data. The cloud security that Amazon AWS offers is taken very seriously. The Amazon Detective, their most recent addition to security services, expedites and improves data investigations.

Utilizing More Saves Money: The more you use certain AWS services like S3 or data transfer OUT from EC2, the less you pay per gigabyte (GB). These are discounts based on volume that are beneficial over time. Free AWS Tier: Free access to more than 60 AWS services is provided when a new account is created. However, depending on the kind of product a company chooses to use, these free offers are further subdivided into three offers. Amazon Rekognition makes it simple to add images to applications. The Amazon Rekognition API is all that is required to enable the service to recognize images, people, text, and scenes. It can also identify content that is inappropriate. Additionally, Amazon Rekognition offers highly accurate capabilities for face comparison, face search, and facial analysis. For a wide range of applications, such as user

verification, cataloging, people counting, and public safety, faces can be detected, analyzed, and compared. Amazon Rekognition is built on the same tried-and-true deep learning technology that Amazon's computer vision scientists developed to analyze daily billions of images and videos. This technology is highly scalable. To use it, you don't need any machine learning expertise. Amazon Rekognition has a straightforward, user-friendly API that can quickly analyze any image file in Amazon S3. Amazon Rekognition is always learning from new data, and the service is constantly getting new labels and features for comparing faces

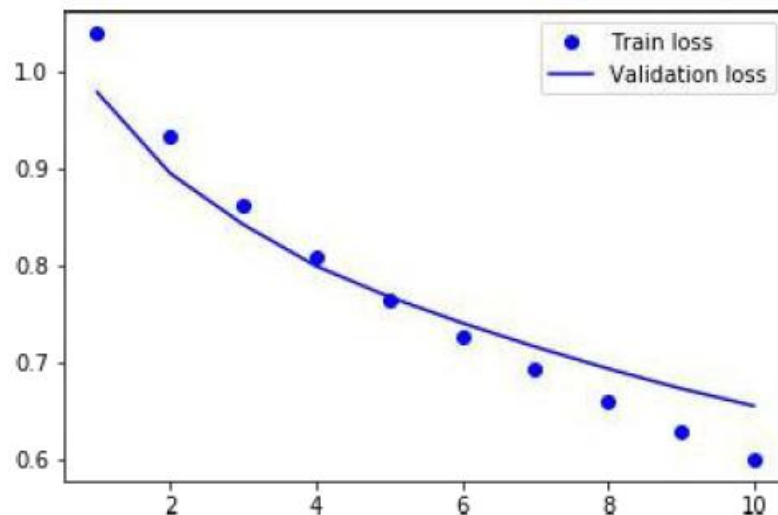


Fig.2 Existing System

3. PROPOSED SYSTEM

However, ratings are noisy labels for review sentences and would mislead classifier training if directly used in supervised training. In this paper, we adopt a simple rule to assign weak labels to sentences with 5-stars rating scale: where $l(s)$ denotes the weak sentiment label of sentence s . Note we follow previous works on aspect level sentiment analysis to only consider positive and negatives entiment labels. The reason is that when commenting on various aspects of a product, people hardly express neutral opinions. shows the percentages of wrong-labeled sentences by $l(s)$, estimated in our labeled review dataset. We can see the noise level is moderate but not ignorable. The general idea behind WDE is that we use large quantities of weakly labeled sentences to train a good embedding space so that a linear classifier would suffice to accurately make sentiment predictions. Here good embedding means in the space sentences with the same sentiment labels are close to one another, while those with different abels are kept away from each other. In the following, we first present the network architecture and explain the specific design choices for WDE-CNN and WDELSTM. Then we discuss how to train it with large-scale. The general architecture of the neural network designed for WDE is shown in Figure 3. At the first layer, the network takes a review sentence as input and extracts a fixed-length low-level feature vector from the sentence. Unlike many traditional methods for sentiment analysis, no feature engineering is required and the extractor is learned automatically. Specific implementation of the extractor will be discussed in the following for WDE-CNN and WDE-

LSTM respectively. The low-level feature vector is then passed through a hidden layer, adding sufficient nonlinearity, and the output is used to compute the embedding representation of the sentence. The embedding representation also takes the sentence's aspect contextual information into consideration. An aspect is a topic on which customers can comment with respect to a sort of entities. For instance, battery life is an aspect for cell phones. We use a learnable context vector to represent an aspect. The motivation for incorporating aspect information as the context of a sentence is that similar comments in different contexts could be of opposite orientations, e.g. "the screen is big" vs. "the size is big". In the weakly-supervised training phase, the goal is to learn embedding space which can properly reflect data's semantic distribution. Hence, the network used in this phase contains only the layers up to the embedding layer. Figure 6 illustrates the advantages of triplet-based training over pair-based training via an example. We use circles and triangles to represent sentences in and respectively. Black nodes denote wrong-labeled sentences. Since the majority of sentences are with correct labels, they would gather together in the training process. Wrong-labeled sentences would go towards the wrong clusters, but with slower speeds. In both training methods, undesirable moves could happen when wrong-labeled sentences are sampled. For clarity, we just show three such cases that are representative for respective methods.

The three cases in Figure 6(a) all result in undesirable moves: sentences with different orientations become closer (cases 1 and 2), while same-orientation sentences become more separated (case 3). In Figure 6(b), case 1 generates only undesirable moves: since s_3 (black triangle) is closer to s_1 (white circle) than s_2 (black circle), the algorithm will drag s_3 away from s_1 and drag s_2 toward s_1 . Cases 2 and 3 lead to a mixed behavior: one move is desirable while the other one is not. Therefore, cases 2 and 3 in Figure 6(b) are not as harmful as the case in Figure 6(a). Furthermore, in triplet-based training there will not be a move if the difference in distances exceeds the margin λ , since the derivative of weak becomes 0. This is useful in that we will not make things too bad. For example, in case 2 of Figure 6(b) s_2 is actually a negative sentence and should not be too close to s_1 . Notice s_3 is far away from s_1 . Hence, the distance difference may already exceed λ and there will be no move for this triplet. As a comparison, cases 1 and 2 in Figure 6(a) will continually move s_1 and s_2 toward each other until their distance becomes 0, which is the worst result.

Detects text in the input image and converts it into machine-readable text. Pass the input image as base64-encoded image bytes or as a reference to an image in an Amazon S3 bucket. If you use the AWS CLI to call Amazon Rekognition operations, you must pass it as a reference to an image in an Amazon S3 bucket. For the AWS CLI, passing image bytes is not supported. The image must be either a .png or .jpeg formatted file.

an image as input, the service detects the objects and scenes in the image and returns them along with a percent confidence score for each object and scene.

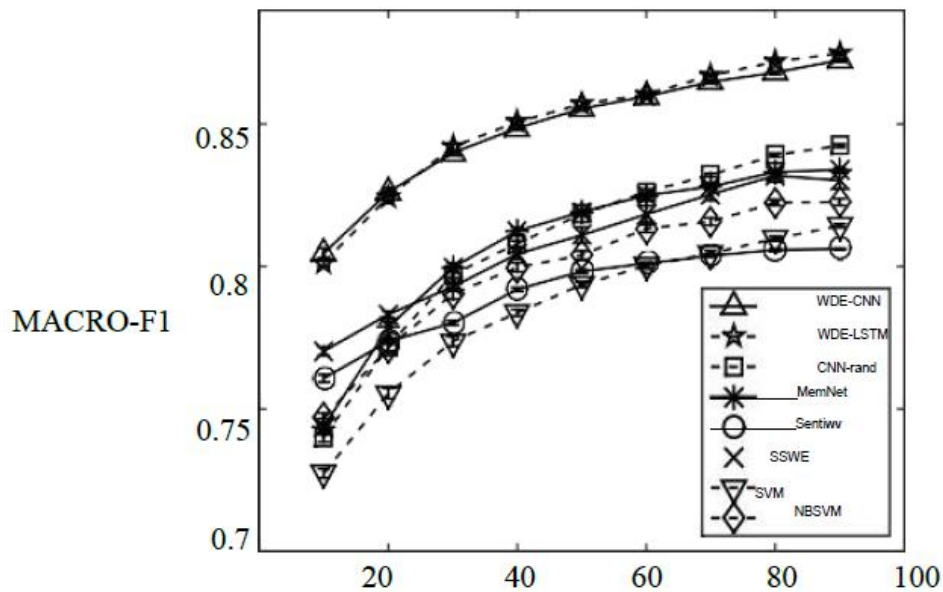


Fig.3 Proposed System

4.CONCLUSION

In this work we proposed a novel deep learning framework named Weakly supervised Deep Embedding for review sentence sentiment classification. WDE trains deep neural networks by exploiting rating information of reviews which is prevalently available on many merchant/review Websites. The training is a 2-step procedure: first we learn an embedding space which tries to capture the sentiment distribution of sentences by penalizing relative distances among sentences according to weak labels inferred formatting; then a soft max classifier is added on top of the embedding layer and we fine-tune the network by labeled data. Experiments on reviews collected from Amazon.com show that WDE is effective and outperforms baseline methods. Two specific instantiations of the framework, WDE-CNN and WDE-LSTM, are proposed. Compared to WDE-LSTM, WDE-CNN has fewer model parameters, and its computation is more easily parallelized on GPUs. Nevertheless, WDE-CNN cannot well handle long-term dependencies in sentences. WDE-LSTM is more capable of modeling the long-term dependencies in sentences, but it is less efficient than WDE-CNN and needs more training data. For future work, we plan to investigate how to combine different methods to generate better prediction performance. We will also try to apply WDE on other problems involving weak labels.

REFERENCES

1. Y. Bengio. Learning deep architectures for ai. Foundations and trends in machine Learning, 2(1):1 - 127, 2009.
2. Y. Bengio, A. Courville, and P. Vincent. Representation learning: A review and new perspectives. IEEE TPAMI, 35(8):1798 - 1828, 2013.

3. C. M. Bishop. Pattern recognition and machine learning. springer,2006.
4. L. Chen, J. Martineau, D. Cheng, and A. Sheth. Clustering for simultaneous extraction of aspects and features from reviews. In NAACL-HLT, pages 789 – 799,2016.
5. R. Collobert, J. Weston, L. Bottou, M. Karlen, K. Kavukcuoglu, and
6. P. Kuksa. Natural language processing (almost) from scratch. JMLR,12:2493 – 2537, 2011
7. <https://docs.aws.amazon.com/rekognition/index.html>.