

Analysis of an Advanced Sentiment Classification Model Using Sentiment140 Dataset with Semantic Text Normalization and Comparative Machine Learning

Deepti Gupta¹, Himashu Gautam²

¹M. Tech Scholar, Department of Computer Science & Engineering, Compucom Institute of Technology & Management, Jaipur

²Assistant Professor, Department of Computer Science & Engineering, Compucom Institute of Technology & Management, Jaipur

Abstract: Sentiment analysis of social media text is still an important but difficult problem in Natural Language Processing (NLP) because of the highly informal and noisy text nature. This paper proposes the design and evaluation of an advanced sentiment classification system designed for the Sentiment140 dataset which contains 1.6 million labeled tweets. The work introduces a complex pipeline for semantic text normalization including slang expansion, negation handling, URL/mention removal, and domain-specific tokenization before the feature extraction step. Three classical machine learning classifiers (Logistic Regression (LR), Support Vector Machine (SVM), Naive Bayes (NB)) and two ensemble methods (Random Forest (RF), Gradient Boosting (GB)) are considered in the experimental part. The BoW representation, TF-IDF weighting of unigrams and bigrams, and combination of word-level n-grams are used as the features in the experiments. The results show that the SVM classifier using TF-IDF bigram features produces the best classification accuracy of 87.63%, surpassing the baseline NB classifier by 9.4 percentage points. In general, the presented normalization pipeline consistently leads to better results with respect to all classifiers, with the average improvement being 4.2% in terms of F1-score. This research also includes an ablation study that focuses on the importance of each of the normalization steps and proves the advantage of using linear kernel SVM for solving problems related to high-dimensional sparse feature sets.

Keywords: Sentiment Analysis, Sentiment140, Text Normalization, Machine Learning, Support Vector Machine, TF-IDF, Natural Language Processing, Twitter Sentiment, Comparative Classification

1. Introduction

The emergence of social media websites, especially micro-blogging sites such as Twitter, now called X, has led to the creation of vast amounts of text data that contains opinions. With more than 500 million tweets being posted per day, this database constitutes one of the most informative sources of up-to-date sentiment data that can be tapped for the purposes of automation in order to benefit brand management, political sentiment research, public health, financial forecasting, and customer experience management.

The problem of automated opinion mining involves the identification of the orientation contained in a certain text sample (in terms of positive, negative, or neutral sentiment). While the problem has already been successfully addressed within the scope of formal text corpora such as product reviews and news articles, the peculiarities of social media languages are still posing significant difficulties. Acronyms, slang, phonetic spelling, emoticons, hashtags, and user mentions constitute the core elements of this language which makes NLP processing ineffective without adaptation to the domain.

Sentiment140 [Go et al., 2009], developed using distant supervision based on emoticons as proxy for polarity, has been one of the most popular benchmarks for this type of tasks, providing 1.6 million labeled tweets. Nevertheless, the noise of the original dataset requires the implementation of solid preprocessing techniques before feature

extraction. Existing studies tend to use general text cleaning algorithms or task-specific normalization procedures without quantitatively analyzing the effect of each individual step of this process.

2. Related Work

2.1 Early Approaches to Sentiment Classification

The original work by Pang et al. [2002] revealed that machine learning techniques, in particular, SVM and Naive Bayes classifiers, outperformed basic bag-of-words baselines in sentiment classification task for movie reviews. The work by Turney [2002] showed that unsupervised sentiment orientation scoring with the help of pointwise mutual information approach was able to perform equally well, making possible the use of dictionaries.

Go et al. [2009] were the first to introduce distant supervision in Twitter sentiment classification, creating the Sentiment140 dataset and showing that Naive Bayes, MaxEnt, and SVM classifiers reached accuracy from 81% to 83% on held-out data. The mentioned research became the main benchmark that needed to be overcome by other works via better preprocessing and feature extraction.

2.2 Text Normalization for Social Media NLP

Han and Baldwin [2011] showed that normalization of tweets via string similarity and context-based approaches drastically cuts down OOV rate and enhances tagging performance downstream. The model was capable of normalizing tweets with 75.5% accuracy. According to Eisenstein [2013], variations of social media are structured and systematic, and he suggested normalization that would be sociolinguistically-aware as opposed to correction by force.

Zhang et al. [2020] offered an end-to-end normalization pipeline based on character-level sequence-to-sequence models, reaching a new level of normalization benchmarking results. Nevertheless, normalization by deep learning approach is characterized by significant computational complexity.

2.3 Feature Representation in Sentiment Classification

The text representation is still an essential factor that influences the classification performance. Bag-of-Words and TF-IDF are simple and efficient text representations to use in combination with linear classifiers [Sebastiani, 2002]. N-gram features, especially bigrams, allow one to consider limited context and significantly increase the accuracy of sentiment classification in comparison with unigrams only [Wang and Manning, 2012].

Distributed word representations proposed by Mikolov et al. [2013] with Word2Vec made it possible to obtain more generalizable text representations. Pre-trained language models like GloVe [Pennington et al., 2014], FastText [Bojanowski et al., 2017], and BERT [Devlin et al., 2019] have gradually increased the performance level for sentiment classification task. Nevertheless, they require much more computing power and are not interpretable.

2.4 Machine Learning Classifiers for Sentiment Analysis

Linear classifiers, such as Logistic Regression and SVM with linear kernel, reliably achieve competitive results on high dimensional text classification problems due to their ability to work well on sparse TF-IDF features [Joachims, 1998]. Naïve Bayes, although based on conditional independence assumption, is still an effective and explainable model [McCallum and Nigam, 1998]. The use of ensemble techniques such as Random Forest and Gradient Boosting provided improvements in stability and performance but decreased the model explainability [Oza, 2005; Chen and Guestrin, 2016].

Recent benchmark studies conducted using Sentiment140 corpus [Neethu and Rajasree, 2013; Agarwal et al., 2011; Kouloumpis et al., 2011] show that preprocessing quality is a critical factor of classifier performance,

playing a bigger role than the choice of the classifier itself. This conclusion provides motivation for our normalization-oriented approach.

3. Dataset and Methodology

3.1 Sentiment140 Dataset Overview

The Sentiment140 dataset [Go et al., 2009] comprises 1,600,000 tweets labeled manually for sentiment polarity. Tweets that had positive emoticons (:), :-), :D) were labeled as positive (class 4), whereas tweets with negative emoticons (:(), :(-), :'()) were labeled as negative (class 0). The dataset is balanced with 800,000 instances per class, meaning that class weighting will not be required during training.

The dataset contains six attributes: polarity, tweet ID, date, query flag, username, and tweet text. For the current analysis, only two attributes were considered, namely polarity and tweet text. Polarity was binarized in the following way: 0 -> 0 (negative), 4 -> 1 (positive). Stratified 80/10/10 train/validation/test split was performed, resulting in 1,280,000/160,000/160,000 instances.

Table 1: Sentiment140 dataset partition statistics.

Split	Samples	Positive / Negative
Training Set	1,280,000	640,000 / 640,000
Validation Set	160,000	80,000 / 80,000
Test Set	160,000	80,000 / 80,000
Total	1,600,000	800,000 / 800,000

3.2 Semantic Text Normalization Pipeline

The suggested normalization pipeline works in the form of a chain of deterministic transformations, which are applied to each raw tweet before its processing. The pipeline is made up of 8 steps, where each step removes a specific kind of Twitter noise:

Step 1: URLs and Mentions Removal

URL links (http/https patterns) and Twitter mentions (@username) are removed by means of regular expressions as they do not have any sentiment value. The hashtag symbol (#) is removed while retaining the text of the tag itself as the content of the hashtags often includes words that have sentiment value (e.g., #awful, #greatmovie).

Step 2: Contractions Expansion

English contractions are expanded by applying a handpicked lookup dictionary consisting of 120 entries (e.g., "can't" → "cannot", "I'm" → "I am", "won't" → "will not"). This step is required for proper negation scope detection in Step 4.

Step 3: Slang and Abbreviations Normalization

An internet slang lexicon of 3,400 words has been compiled from various open sources and manually corrected in case of ambiguity. Slang tokens (e.g., "lol" → "laugh out loud", "omg" → "oh my god", "tbh" → "to be honest") are expanded to their standard counterparts. This step is necessary for proper negation scope detection in step 4.



Stage 4: Negation Handling

Negation scope identification is carried out by suffixing a "_NEG" to all words in the sequence starting from the negation trigger word (not, no, never, n't, cannot) until the next punctuation boundary. Thus, the sentence "I do not like this movie" gets transformed into "I do not like_NEG this_NEG movie_NEG". This method was proposed by Pang et al. [2002] to successfully change the polarity of the context to vocabulary terms.

Stage 5: Repeated Character Handling

Emphatic elongated words (e.g., "sooooooo", "amazingggggg") are processed using the maximum two-repeat rule, where any character repeated more than twice is represented as two occurrences only (e.g., "sooooo" → "soo"). This ensures that vocabulary does not get fragmented and allows preserving the emphatic information via the two-repeat feature.

Stage 6: Emoticon and Emoji Handling

A set of 180 textual emoticons and 420 Unicode emoji were represented via descriptive sentiment terms (e.g., ":)") → "POSITIVE_EMOTICON", ":(" → "NEGATIVE_EMOTICON", laughing emoji → "LAUGHING").

Stage 7: Noise Removal and Lowercasing

Residual special characters, numbers, and HTML entities are filtered out. The text is converted to lowercase. Overabundant white spaces are standardized to one white space. Words whose length is smaller than two are stripped because they usually do not carry any sentiment.

Stage 8: Tokenization and Stopword Elimination

Tokenization takes place via white space splitting after the normalization stage. Domain-specific stopwords elimination is used based on the NLTK English stopwords list enriched with Twitter-specific filler words. Note that negation words ("not," "never," "no") are not included in stopwords elimination.

3.3 Feature Extraction

Three representations of features were developed and evaluated separately:

- BoW (Bag of Words): Binary occurrence vectors based on the vocabulary of only the top 50,000 unigrams.
- TF-IDF Unigrams: TF-IDF vectors based on the top 75,000 unigrams vocabulary, employing sub-linear TF scaling.
- TF-IDF Bigrams: TF-IDF weightings based on a vocabulary of the top 100,000 features (including unigrams and bigrams).

Feature matrices were stored in sparse CSR (Compressed Sparse Row) representation due to high dimensionality and sparsity of text-based features. Dimensionality reduction techniques were not used since linear classifiers work best in the original feature space [Wang and Manning, 2012].

3.4 Classification Models

Five machine learning classifiers were implemented and evaluated:

- Logistic Regression (LR): L2-regularized logistic regression with lbfgs solver, C=1.0, maximum 1000 iterations.

- Support Vector Machine (SVM): Linear kernel SVM via LinearSVC implementation, $C=1.0$, using one-vs-rest strategy for binary classification. Linear kernel was selected based on prior literature demonstrating its superiority in high-dimensional sparse text spaces.
- Naive Bayes (NB): Multinomial Naive Bayes with Laplace smoothing ($\alpha=1.0$), appropriate for TF-IDF and count-based feature representations.
- Random Forest (RF): Ensemble of 200 decision trees with maximum depth 50, using Gini impurity criterion and bootstrapped samples.
- Gradient Boosting (GB): XGBoost implementation with 150 estimators, learning rate 0.1, maximum depth 6, and subsampling ratio 0.8.

All models were implemented using the scikit-learn library (version 1.3.0) in Python 3.10, except Gradient Boosting which used XGBoost 1.7.6. Hyperparameters were tuned via 5-fold cross-validation on the training set using grid search over predefined parameter grids.

4. Experimental Setup

4.1 Hardware and Software Environment

All experiments were conducted on a server equipped with an Intel Xeon Gold 6230R CPU (26 cores, 2.1 GHz), 256 GB DDR4 RAM, and an NVIDIA Tesla V100 32 GB GPU. The operating system was Ubuntu 20.04 LTS. Python 3.10.6 was the primary programming environment. Key libraries included scikit-learn 1.3.0, NLTK 3.8.1, XGBoost 1.7.6, NumPy 1.24.3, Pandas 2.0.2, and SciPy 1.10.1.

4.2 Evaluation Metrics

For the balanced class distribution present in the Sentiment140 dataset, accuracy is an appropriate evaluation metric. In order to obtain a more detailed understanding of the classifier performance, the following metrics have been calculated using the test dataset:

- Accuracy: Fraction of correct predictions.
- Precision (macro-average): Precision averaged over two sentiment classes.
- Recall (macro-average): Recall averaged over two sentiment classes.
- F1-score (macro-average): Balanced performance metric based on precision and recall.
- AUC-ROC: Discriminatory power of the classifier.
- Training Time: Time taken to train the model on the training dataset.

4.3 Baseline Comparison

Two baseline experiments have been conducted, including:

- Naive Bayes using Bag of Words on raw (un-normalized) text, which is the simplest data pre-processing method;
- Experiments conducted according to Sentiment140 [Go et al., 2009].

5. Results and Discussion

5.1 Overall Classification Performance

The performance of all five models on three feature types is summarized in Table 2 for the normalized test set of Sentiment140 dataset. Performance is calculated on the test set containing 160,000 data instances.

Table 2: Classification performance comparison on the Sentiment140 test set (normalized text). Best results in bold.

Model	Feature	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)
Naive Bayes	BoW	78.24	78.11	78.24	78.17
Naive Bayes	TF-IDF Uni	80.56	80.49	80.56	80.51
Naive Bayes	TF-IDF Bi	81.33	81.27	81.33	81.30
Logistic Regression	BoW	82.17	82.09	82.17	82.13
Logistic Regression	TF-IDF Uni	84.78	84.71	84.78	84.74
Logistic Regression	TF-IDF Bi	85.94	85.89	85.94	85.91
SVM (Linear)	BoW	83.41	83.36	83.41	83.38
SVM (Linear)	TF-IDF Uni	86.29	86.22	86.29	86.25
SVM (Linear)	TF-IDF Bi	87.63	87.58	87.63	87.60
Random Forest	TF-IDF Bi	83.07	82.99	83.07	83.03
Gradient Boosting	TF-IDF Bi	84.52	84.47	84.52	84.49

5.2 Key Findings

The SVM with linear kernel and TF-IDF bigram features performs best on all metrics, scoring 87.63% in accuracy and 87.60% in macro-F1-score. This is an improvement of 6.30 percentage points compared to the Naive Bayes BoW baseline (78.24%), and is better than the best performance reported in the original Sentiment140 paper (MaxEnt: 83.0%) by 4.63 points. Some key findings from these experiments include the following:

Effectiveness of Feature Representation: TF-IDF bigrams significantly outperform both TF-IDF unigrams and BoW on all classifiers, thus establishing the necessity of context in sentiment disambiguation. Bigrams provide valuable information about phrases like "not good", "very happy", and "doesn't work" which are completely lost when using unigrams. The improvement in accuracy from unigrams to bigrams averages 1.2%.

Comparison between Linear & Ensemble Classifiers: Particularly interesting is that while the linear classifier of SVM performs better than Random Forest and Gradient Boosting even though they have more complicated models. The intuition behind that surprising observation can be easily found in the "curse of dimensionality" – high-dimensional and sparse TF-IDF space is linearly separable. And tree-based ensembles cannot use the sparse representation efficiently due to their nature and cannot do axis-aligned splits in linearly separable high dimensional space.

Performance of Logistic Regression: The Logistic Regression shows the performance (85.94%) that is quite close to SVM's (87.63%), proving that both are good linear classifiers for the task. The margin between the two (1.69%) is explained by the margin-maximizing objective function of SVM, which gives better generalization due to structural risk minimization.

5.3 Impact of Normalization Pipeline

Table 3 shows the performance with and without the normalization pipeline, keeping constant the feature representation (TF-IDF bigrams) and the classifier (SVM).

Table 3: Incremental impact of each normalization stage (SVM + TF-IDF Bigrams).

Condition	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)
Raw Text (No Normalization)	81.44	81.38	81.44	81.41
+ URL/Mention Removal	82.91	82.85	82.91	82.88
+ Contraction Expansion	83.67	83.62	83.67	83.64
+ Slang Normalization	84.89	84.83	84.89	84.86
+ Negation Handling	85.92	85.87	85.92	85.89
+ Repeated Char Norm.	86.44	86.39	86.44	86.41
+ Emoticon Processing	87.11	87.06	87.11	87.08
Full Pipeline (All Steps)	87.63	87.58	87.63	87.60

The contribution from the normalization pipeline towards the overall accuracy improvement is 6.19%. The significant individual steps in this process include the use of slang normalization (1.22%) and negation treatment (1.03%). This is indicative of the importance of expansion of informal language and detection of polarity change

contexts in sentiment analysis on Twitter. Even the last step in normalization involving emoticon treatment has an effect of 0.52%.

5.4 Computational Efficiency Analysis

Table 4 shows training times and inference throughputs for all classifier configurations as part of practical deployment decisions.

Table 4: Computational efficiency of evaluated classifiers (training on 1.28M samples, TF-IDF Bigram features)

Model	Feature	Train Time (min)	Inference (tweets/sec)
Naive Bayes	TF-IDF Bi	3.2	142,000
Logistic Regression	TF-IDF Bi	28.4	98,000
SVM (Linear)	TF-IDF Bi	34.7	87,000
Random Forest	TF-IDF Bi	218.3	4,200
Gradient Boosting	TF-IDF Bi	387.1	2,800

Linear classifiers such as NB, LR and SVM perform training and inference extremely fast as compared to ensemble models. NB takes only 3.2 minutes for training and performs inference at a throughput of 142,000 tweets/sec and hence emerges as the best option for those applications which are highly sensitive to latency with moderate accuracy expectations. SVM has the highest accuracy with a minimal computational overhead as compared to ensemble algorithms, and a good inference throughput of 87,000 tweets/sec.

6. Feature Importance Analysis

6.1 Vocabulary Size Sensitivity

To analyze how the size of the vocabulary influences the SVM classification results, the tests have been performed using the vocabulary with maximum TF-IDF bigrams of the size ranging from 20,000 to 150,000 features. The improvement is fast up to 75,000 features (86.94%) and the gain becomes negligible after 100,000 features (87.63% for 100K features versus 87.71% for 150K features – almost imperceptible difference). Hence, 100,000 features are

6.2 N-gram Range Analysis

Table 5 studies the impact of choosing n-grams on the SVM classification accuracy using TF-IDF features.

Table 5: N-gram range sensitivity for SVM with TF-IDF features.

N-gram Range	Vocabulary Size	Accuracy (%)
Unigrams only (1,1)	75,000	86.29
Bigrams only (2,2)	75,000	83.41

N-gram Range	Vocabulary Size	Accuracy (%)
Unigrams + Bigrams (1,2)	100,000	87.63
Unigrams + Trigrams (1,3)	120,000	87.28
Unigrams + Bigrams + Trigrams (1,3)	150,000	87.71

Using bigrams alone (83.41%) is less effective than using unigrams (86.29%), which validates the notion that individual words are better at sentiment classification than phrases. The most optimal compromise between the two would be to use both unigrams and bigrams (87.63%), and using trigrams further improves performance marginally by just 0.08% but doubles the dictionary size.

6.3 Regularization Sensitivity (SVM)

The C constant used in SVM determines the bias-variance tradeoff. A grid search using values $C = \{0.01, 0.1, 0.5, 1.0, 2.0, 5.0, 10.0\}$ showed that the accuracy was highest when $C = 1.0$ (87.63%), with symmetrical decreases on both sides, indicating that the default value of C was well-chosen for this task. Using higher values of C led to overfitting ($C = 5.0$: 87.21%,

7. Discussion

7.1 Interpretability and Practical Deployment

A Linear SVM model offers a large scope for interpretability by virtue of feature weight analysis. An analysis of top 50 weights of positive and negative bigrams shows linguistic feasibility in terms of sentiment discrimination. Positive features include tokens like “love_you”, “POSITIVE_EMOTICON”, “so_happy”, “great_time” and “thank_you”. Negative features include “not_working”, “NEGATIVE_EMOTICON”, “so_sad”, “hate_this”, and “not_good”. The semantic relevance of these features is a validation of the model itself and helps in debugging as well as auditing of biases.

7.2 Limitations of the Present Study

There are several limitations to the generalization of the current results. The first limitation is associated with the distant supervision approach used in labeling the Sentiment140 dataset, which allows for scalability but produces label noise as not all tweets with positive emoticons convey the same degree of sentiment positivity. Label noise in the form of the ratio of 8-12% limits the maximum possible accuracy that can be achieved using only preprocessing and modeling improvements.

The second limitation refers to the two-class problem formulation in which tweets are classified into either positive or negative sentiments, while ignoring the neutral sentiment category used in social media posts. To incorporate the neutral class in the classification problem, a new dataset would have to be created.

Third, domain generalizability of the preprocessing pipeline to other social media sites (Reddit, Facebook, Instagram) has not been tested yet.

7.3 Comparison with Deep Learning Approaches

Modern deep learning techniques like transformers fine-tuned on Twitter data (BERTweet [Nguyen et al., 2020] and TweetBERT), have yielded accuracies above 91% on the Sentiment140 test dataset. However, their performances incur significant increases in cost: BERTweet needs to be run on specialized hardware for GPU-

accelerated inference and fine-tuning, and is 2–3 orders of magnitude slower than the linear SVM technique investigated here.

The SVM technique discussed in this paper provides an optimal choice for companies working with limited budgets and/or need for explainability. The difference in accuracy of 3.4–4.0% between linear SVM and fine-tuned transformers has to be put into context.

8. Conclusion

This research paper provided the design and evaluation of an advanced framework for sentiment classification based on the Sentiment140 dataset, implemented through a new 8-step semantic text normalization technique and a detailed machine learning comparison. The results obtained are as follows:

- The suggested normalization pipeline provides a total accuracy boost of 6.19% compared to the baseline, with the most significant contributions made by slang normalization and negation techniques.
- Linear SVM with bigram features extracted with TF-IDF provides the highest possible performance of 87.63% accuracy and 87.60% F1-score among all tested classifiers, including ensemble methods.
- Bigram-unigrams provide the optimal n-gram range, capturing sentiment information from both word and phrase levels, but not expanding vocabulary significantly like the use of trigrams.
- Linear classifiers show dramatically better efficiency, with SVM providing an inference speed of 87,000 tweets per second against 2,800 tweets per second for Gradient Boosting, allowing real-time stream processing.
- The learned SVM model proves highly interpretable, with weights assigned to the feature map semantically coherent sentiment concepts.

Three major areas will be considered in future research. First, the use of domain-specific word embeddings (such as GloVe trained on Twitter datasets) alongside the normalization pipeline in a hybrid representation approach. Second, adaptation for multi-class sentiment classification involving a neutral class. Third, the generalization of the normalization pipeline to different social media platforms and languages, specifically, the use of code-mixed Hindi-English language.

References:

- [1] Agarwal, A., Xie, B., Vovsha, I., Rambow, O., & Passonneau, R. (2011). Sentiment analysis of Twitter data. In Proceedings of the Workshop on Language in Social Media (LSM 2011), pp. 30–38. Association for Computational Linguistics.
- [2] Bojanowski, P., Grave, E., Joulin, A., & Mikolov, T. (2017). Enriching word vectors with subword information. Transactions of the Association for Computational Linguistics, 5, 135–146.
- [3] Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. In Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD '16), pp. 785–794.
- [4] Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. In Proceedings of NAACL-HLT 2019, pp. 4171–4186.
- [5] Eisenstein, J. (2013). What to do about bad language on the internet. In Proceedings of NAACL-HLT 2013, pp. 359–369.
- [6] Go, A., Bhayani, R., & Huang, L. (2009). Twitter Sentiment Classification using Distant Supervision. CS224N Project Report, Stanford University. (Sentiment140 Dataset).

-
- [7] [7] Han, B., & Baldwin, T. (2011). Lexical normalisation of short text messages: Makn sens a #twitter. In Proceedings of ACL-HLT 2011, pp. 368–378.
 - [8] [8] Joachims, T. (1998). Text categorization with support vector machines: Learning with many relevant features. In Proceedings of ECML-98, pp. 137–142. Springer.
 - [9] [9] Kouloumpis, E., Wilson, T., & Moore, J. (2011). Twitter sentiment analysis: The good the bad and the OMG! In Proceedings of the 5th International AAAI Conference on Weblogs and Social Media (ICWSM 2011), pp. 538–541.
 - [10] [10] McCallum, A., & Nigam, K. (1998). A comparison of event models for Naive Bayes text classification. In AAAI-98 Workshop on Learning for Text Categorization, pp. 41–48.
 - [11] [11] Mikolov, T., Chen, K., Corrado, G., & Dean, J. (2013). Efficient estimation of word representations in vector space. In Proceedings of ICLR 2013.
 - [12] [12] Neethu, M. S., & Rajasree, R. (2013). Sentiment analysis in Twitter using machine learning techniques. In Proceedings of the 4th International Conference on Computing, Communications and Networking Technologies (ICCCNT), pp. 1–5. IEEE.
 - [13] [13] Nguyen, D. Q., Vu, T., & Nguyen, A. T. (2020). BERTweet: A pre-trained language model for English tweets. In Proceedings of EMNLP 2020: System Demonstrations, pp. 9–14.
 - [14] [14] Oza, N. C. (2005). Online bagging and boosting. In Proceedings of the 2005 IEEE International Conference on Systems, Man and Cybernetics, Vol. 3, pp. 2340–2345.
 - [15] [15] Pang, B., Lee, L., & Vaithyanathan, S. (2002). Thumbs up? Sentiment classification using machine learning techniques. In Proceedings of EMNLP 2002, pp. 79–86.
 - [16] [16] Pennington, J., Socher, R., & Manning, C. D. (2014). GloVe: Global vectors for word representation. In Proceedings of EMNLP 2014, pp. 1532–1543.
 - [17] [17] Sebastiani, F. (2002). Machine learning in automated text categorization. ACM Computing Surveys, 34(1), 1–47.
 - [18] [18] Turney, P. D. (2002). Thumbs up or thumbs down? Semantic orientation applied to unsupervised classification of reviews. In Proceedings of ACL 2002, pp. 417–424.
 - [19] [19] Wang, S., & Manning, C. D. (2012). Baselines and bigrams: Simple, good sentiment and topic classification. In Proceedings of ACL 2012 (Short Papers), pp. 90–94.
 - [20] [20] Zhang, S., Takamura, H., & Okumura, M. (2020). Text normalization for social media: An empirical study with sequence-to-sequence models. In Proceedings of COLING 2020, pp. 5200–5211.