

Sentiment Classification

- Classify documents (e.g., reviews) based on the overall sentiments expressed by authors
- Similar but different from topic-based text classification
 - In topic-based text classification, topic words are important.
 - In sentiment classification, sentiment words are more important, e.g., great, excellent, horrible, bad, worst, etc.

Experimental Set-up

- Domain
 - movie review
- Data Source
 - Internet Movie Database (IMDb)
- PreProcess
 - Ratings were automatically extracted and converted into one of three categories: positive, negative or neutral
- Data set:
 - Randomly selected 700 positive-sentiment and 700 negative-sentiment documents
 - Then, they divided this data into three equal-sized folds, maintaining balanced class distributions in each fold

Support Vector Machines

- Highly effective at traditional text categorization
- Find a hyperplane $\vec{w}$ not only separates the document vectors in one class from those in the other, but for which the separation, or margin, is as large as possible
 - $\vec{w} := \sum_i \alpha_i \vec{e}_i, \alpha_i \geq 0$
 - $c_j \in \{-1, 1\}$
 - α_i is obtained by solving a dual optimization problem

Results

	Features	# of features	frequency of percentage	NI	ME	SVS
(1)	unigram	16145	freq	78.7	N/A	72.8
(2)	unigram	16145	freq	81.1	82.9	75.1
(3)	unigram	22339	freq	80.6	80.8	82.7
(4)	unigram+bigram	22339	freq	77.5	77.4	75.1
(5)	unigram+bigram	22339	freq	81.1	80.4	82.1
(6)	unigram+POS	16695	freq	77.0	77.7	75.1
(7)	unigram+POS	2483	freq	80.3	81.0	81.4
(8)	unigram+positional	22339	pos	81.3	80.1	81.6

Figure 3: Average three-fold cross-validation accuracies, in percent. Boldface: best performance for a given setting (row). Recall that our baseline results ranged from 50% to 60%.

Thumbs up? Sentiment Classification using Machine Learning Technique

Bo Pang, Lillian lee and Shivakumar Vaithyanathan
EMNLP 2002

Presented by
Lu Wang

Problem description

- Aim in this work:
 - Whether it suffices to treat sentiment classification simply as a special case of topic-based categorization
 - Whether special sentiment-categorization methods need to be developed

Maximum Entropy

- $p_{d,c} := \frac{1}{Z(d)} \exp(\sum_{i=1}^I F_{i,c}(d, c))$
- $F_{i,c}$ is a normalization function
- $F_{i,c}$ is a feature/class function for feature f_i and class c , defined as follows:

$$F_{i,c}(d, c') := \begin{cases} 1, & \text{if } (d, c') = c \\ 0, & \text{otherwise} \end{cases}$$
- $\lambda_{i,c}$ are feature-weight parameters
- MaxEnt makes no assumption about the relationships between features

Evaluation

- The tag “**NOT_**” was added to every word between a negation word (“not”, “isn’t”, “didn’t”) and the first punctuation mark following the negation word
 - Don’tlike →NOT_like
- **Feature**
 - Focused on features based on unigrams and bigrams
 - 16,165 unigrams appearing at least 4 times in the 1400-document corpus
 - 16,165 most often occurring bigrams in the same data(they didn’t add negation tags to the bigrams)

Baseline

- **Assumption**
 - There are certain words people tend to use to express strong sentiments, so that it might suffice to simply produce a list of such words by introspection and rely on them alone to classify the texts
 - They did classification by counting the number of positive and negative words in the documents according to the word lists

Machine Learning Methods

- Standard bag-of-features framework
- Let $\{f_1, \dots, f_m\}$ be a predefined set of m features that can appear in a document
- Let $n_i(d)$ be the number of times f_i occurs in document d .
- Then, each document d is represented by the document vector $\vec{d} = (n_1(d), n_2(d), \dots, n_m(d))$

Baseline

- Random-choice baseline: 50%

	Proposed word lists	Yes	Accuracy
Human 1	positive: dazzling, brilliant, phenomenal, excellent, fantastic negative: sick, terrible, awful, unsuitable, ludicrous	58%	75%
Human 2	positive: gripping, mesmerizing, riveting, spectacular, cool, awesome, thrilling, badass, excellent, moving, exciting negative: boring, lackluster, mediocre, tedious, stupid, slow	64%	80%

Figure 1: Baseline results for human word lists. Data: 700 positive and 700 negative reviews.

	Proposed word lists	Accuracy	Two
Human 3 + state	positive: <i>love, wonderful, best, great, superb, still, beautiful</i> negative: <i>bad, worst, stupid, waste, boring, y, f</i>	69%	16%

Figure 3: Results for baseline using introspection and simple statistics of the data (including test data).

Naive Bayes

- Assign to a given document d the class
 - $c^* = \arg \max_c P(c|d)$
 - $P(c|d) = \frac{P(c)P(d|c)}{P(d)}$
 - $P_{\theta}(c|d) = \frac{P_{\theta}(c) \prod_{i=1}^n P_{\theta}(x_i|c)^{x_i^{(d)}}}{P_{\theta}(d)}$
- Conditional independence assumption between the features

Conclusion

- Naive Bayes tends to do the worst and SVMs tend to do the best
- Unigram presence information is the most effective
- "unwarted expectation", many words indicative of the opposite sentiment to that of the entire review
 - "This film should be brilliant." It sounds like a great plot; the actors are first grade, and the supporting cast is good as well, and Stallone is attempting to deliver a good performance. However, it can't hold up"
- Some form of discourse analysis is necessary