



Reviews Are Not Equally Important: Predicting the Helpfulness of Online
Reviews

Yang Liu
Xiangji Huang
Aijun An
Xiaohui Yu

Technical Report CSE-2008-05

July 12, 2008

Department of Computer Science and Engineering
4700 Keele Street Toronto, Ontario M3J 1P3 Canada

Reviews Are Not Equally Important: Predicting the Helpfulness of Online Reviews

Yang Liu, Xiangji Huang, Aijun An, Xiaohui Yu
York University
{yliu, jhuang, aan, xhyu}@cse.yorku.ca

Abstract

Online reviews provide a valuable resource for potential customers to make purchase decisions. However, the sheer volume of available reviews as well as the large variations in the review quality present a big impediment to the effective use of the reviews, as the most helpful reviews may be buried in the large amount of reviews of low qualities. The goal of this paper is to develop models and algorithms for predicting the helpfulness of reviews, which provides the basis for discovering the most helpful reviews for given products. We first show that the helpfulness of a review depends on three important factors: the reviewer's expertise, the writing style of the review, and the timeliness of the review. Based on the analysis of those factors, we present HelpMeter, a nonlinear regression model for helpfulness prediction. Our empirical study on the IMDB movie reviews dataset demonstrates that the proposed approach is highly effective.

1 Introduction

The increasing impact of the Internet has dramatically changed the way that people shop for goods. More and more people are now gravitating to reading products reviews prior to making purchasing decisions. Such reviews have become an indispensable component of e-commerce Websites such as Amazon¹, and they are also available through dedicated Websites such as CNET² and IMDB³. While reading reviews can help the potential customers make informed decisions, in many cases the large quantity of reviews available for a product can be overwhelming and actually impede the customers' ability to evaluate the product. This is further aggravated by the fact that the quality of the online reviews tends to be very uneven, ranging from excellent detailed opinions to simple repetition of product specifications to (in the worst case) pure spams. As a consequence, potential consumers have to sift through a large number of reviews in order to form an unbiased judgment regarding the product.

Reluctant to do so, some customers may choose to simply use the average numerical customer ratings provided by the Websites as a measure of the product quality. Although the average rating represents the collective sentiment towards a product to a certain extent, it does not fully reflect the richness of the opinions expressed in the reviews and may not convey sufficient information. Moreover, recent work has shown that the majority of online reviewers tend to give an extremely high or low rating on a product to attract public attention [4], which makes the average rating an unreliable measure.

To alleviate this problem, many Websites are now allowing readers of a review to indicate whether they think that review is helpful by voting for or against it, and a tally (or score) is provided in the form of “100 out of 150 people found the following review helpful”. The reviews can be sorted according to their helpfulness using those scores. Although this is certainly an improvement in the right direction, there are still important issues to be addressed. For example,

¹<http://www.amazon.com>

²<http://www.cnet.com>

³<http://www.imdb.com>

- For newly posted reviews, most likely no vote or only a few votes have been cast, and therefore, identifying their helpfulness is difficult.
- Presenting the reviews ranked by their user-voted helpfulness scores may create situations of “monopoly” in that only the highest ranked reviews get viewed, leaving no opportunities for the newly published yet un-voted reviews to show up on users’ radar.
- In some cases, reviews can be incorrectly labeled as helpful or not helpful due to spam voting [11].

In these scenarios, it will be highly desirable to have a way to predict the helpfulness of the given reviews. The predicted helpfulness scores can then be used to address the above problems either directly or indirectly, by combining with existing user votes (if there is any).

This paper is concerned with the problem of automatically evaluating the *helpfulness* of reviews and consequently identifying the most helpful reviews for a particular product. Previous research on review mining has focused on answering questions like “*What do people think of the product?*” [4, 24, 28], “*How would users’ evaluation affect the sales of a certain product?*” [1, 6, 19], and “*How to understand and summarize the reviews with minimum human efforts?*” [9, 31], but only a few studies explicitly consider the problem of evaluating the quality of product reviews [6, 14, 30]. In those studies, a variety of the factors that may affect the helpfulness of reviews are explored, but most of them are related to the contents of the reviews only. Some important factors which are essential to the helpfulness prediction, such as the expertise of the reviewers and the timeliness of the reviews, are still missing. Furthermore, most of those studies rely on off-the-shelf solutions, such as SVM and logistic regression, to model the factors, which may not account for the unique characteristics of each individual factor.

In this paper, we address those problems, and develop a novel model for predicting the helpfulness of reviews. Our model is based on a thorough analysis of some major factors that may affect the helpfulness of a review, such as areas of expertise of the reviewer, the writing styles, the timeliness of the reviews, etc. We provide a detailed analysis of those factors and explain their effects on the helpfulness of reviews. We then develop a non-linear regression model based on radial basis functions that takes the most important factors into consideration, serving as a basis for helpfulness prediction. Extensive experiments were conducted on the IMDB dataset, demonstrating the effectiveness of the proposed approach.

To make our discussions and results more concrete, in this paper we use movie reviews in the past two years (2006-2007) collected from the IMDB Website as a case study. However, our approach is general enough to be easily adapted to handling other types of online reviews. Extensive experiments were conducted on the IMDB dataset, demonstrating the effectiveness of the proposed approach.

Equipped with the proposed model, the e-commerce Websites or review Websites can greatly improve the way reviews are presented. For example, the reviews can be presented in the descending order of their helpfulness without suffering from the aforementioned “monopoly” problem, because even if there is no vote for the newly posted review, we can still estimate its helpfulness with our method, and use it for helpfulness ranking. The proposed model can also help detect spam voting by providing an independent alternative helpfulness rating on the reviews.

2 Problem Definition and Observations

In this section, we first formally define the problem of helpfulness prediction, and then analyze the factors that may affect the helpfulness of a review, which will provide the basis for the proposal of the model in the next section.

2.1 Problem definition

The goal of this research is to develop a model that can accurately predict the helpfulness of a review. For a given review, its “helpfulness” H is defined as the expected fraction of people who will find the review helpful. That is, H is a number falling in the range $[0, 1]$, and greater values of H imply higher helpfulness.

As in any prediction tasks, the prediction model will be obtained based on available training data, which consist of reviews and related product information. Let the set of reviewers (authors of the reviews) concerned be $\mathcal{U}$, the set of movies be $\mathcal{M}$, the set of reviews be $\mathcal{D}$, then each review can be represented as a quadruple $R = (u, d, m, t)$, where $u \in \mathcal{U}$ denotes the reviewer, $d \in \mathcal{D}$ represents the review, $m \in \mathcal{M}$ represents the movie for which the review

Factor	Correlation Coefficient
Comment length	0.1396
Polarity	0.0947
Movie average rating	0.0070
Number of responses	0.1829
Review subjectivity	0.0307

Table 1. Correlation Coefficient to Helpfulness Rating

is written, and t indicates the number of days elapsed from the movie release to the time the review is published. For each movie in $\mathcal{M}$, assume that the genres it falls in are also available.

The helpfulness of a review in the training data can be approximated by the tally attached to that review, which takes the form of “ x out of y people found the following review helpful”. That is, $H = \frac{x}{y}$. As an effective indicator of the public opinions, this evaluation metric has also been widely adopted in previous product review helpfulness studies [14, 30]. To maintain the robustness of the prediction model, in this study, we only consider reviews with at least 10 votes, i.e., $y \geq 10$.

2.2 Observations

In order to develop an effective model for helpfulness prediction, we must carefully analyze the important factors that may affect a review’s helpfulness rating. To this end, we have examined the reviews on several popular Websites, including CNET, Amazon, and IMDB, and conducted preliminary experiments to evaluate the various factors involved. We first considered a collection of possible factors that may affect the helpfulness values including length of the review, polarity of the review, the number of responses the review received, the subjectivity of the review, and the average rating of all reviews on the movie. Most of these factors have been studied in previous literatures to measure the helpfulness of product reviews, e.g., reviews on digital cameras and MP3 players, from commercial website, e.g., Amazon and Ebay [14, 6]. In this study, we want to examine their effectiveness in the context of movie reviews.

To achieve this, we first examine the distribution of various factors, and show the results in Figure 1. It is clear to see that

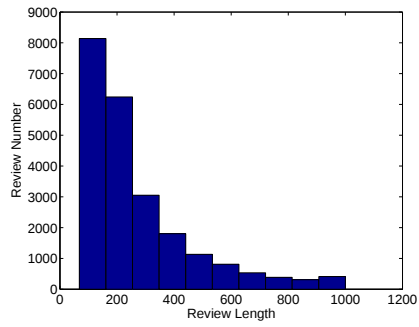
1. Most of the movie reviews in our dataset contain 100-200 sentences, and the number of reviews exponentially decreases as the increase of review length.
2. The polarity, movie rating, and subjectivity in our movie review data generally follow a normal distribution.
3. Only a very few reviews can get a large number of responses.

We then use the correlation coefficients defined as,

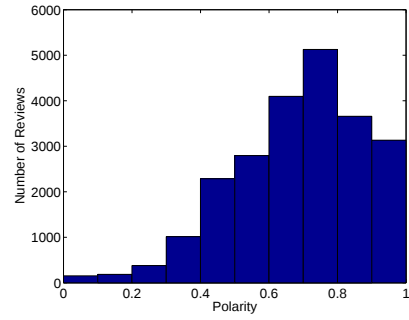
$$\rho_{X,Y} = \frac{E(XY) - E(X)E(Y)}{\sqrt{E(X^2) - E^2(X)} \sqrt{E(Y^2) - E^2(Y)}}, \quad (1)$$

to approximate the strength of a linear relationship between the real helpfulness ratings and each individual factor, where X represents the vector of helpfulness ratings in movie review dataset, and Y corresponds to the vector of each single factor. Note that we use the number of sentences in a review to estimate the length of a comment in this study. The polarity of a review is considered to be the ratio of the number of positive words to the number of all appraisal words in a post. In addition, the number of responses of a review is estimated as the number of people who rated the review, i.e., the value of y in “ x out of y people found the following review helpful”. The review subjectivity is calculated as the ratio of the number of appraisal words to the number of sentences in a given post. Finally, the results of their correlations are shown in Table 1.

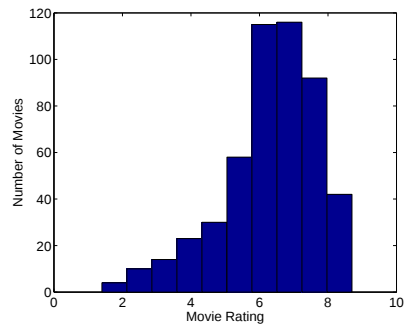
Surprisingly, we notice that none of the above factors demonstrate strong correlation to the helpfulness ratings. This might because the task of understanding movie reviews is inherently difficult [17], and it present challenges that can not be easily addressed with simple content-based factors. To solve this problem, in this study, we investigate other factors that might be of interest, and our efforts reveal that the following are among the most important factors.



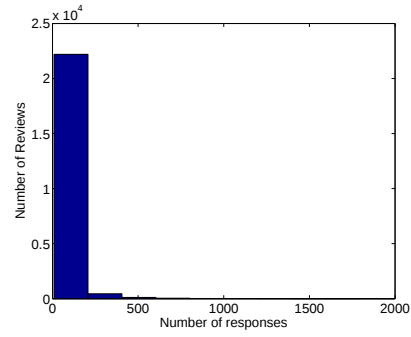
(a) Distribution of review length



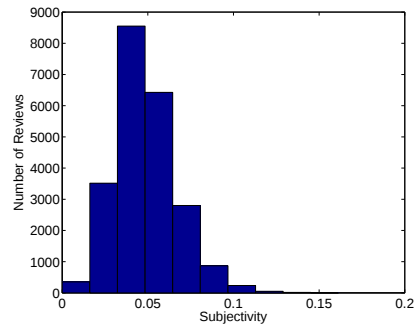
(b) Distribution of polarity



(c) Distribution of movie rating



(d) Distribution of response



(e) Distribution of subjectivity

Figure 1. The distribution various factors

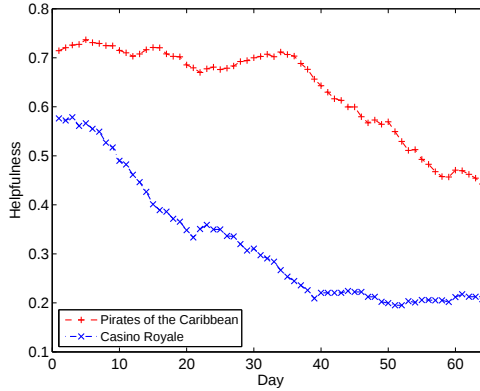


Figure 2. An example of review helpfulness vs. time of review (number of days since the release of the movie). The average helpfulness numbers presented here are 14-day moving averages in order to smooth-out the short-term irregularities and show the overall trend.

1. **Reviewer Expertise:** Product reviews often involve personal experience, thoughts, and concerns. Also, it is common that different reviewers demonstrate expertise on different types of products. For example, one reviewer may be very familiar with the pros and cons of the digital cameras on the market, but may not have much experience with LCD TVs, whereas another reviewer may be an expert on LCD TVs but on not digital cameras. Naturally, we expect that the digital camera reviews written by the first reviewer to be more helpful than those by the second reviewer, and vice versa.

The same observation can be made on other types of reviews, e.g., movie reviews. For example, reviewers fond of science fictions are likely to be familiar with and produce good reviews on sci-fi movies like *Star Wars* and *The Matrix*, but may be less proficient in writing reviews for *American Zeitgeist* and *An Inconvenient Truth*, which fall into the category of documentaries. Those preferences and expertise might be well reflected through reviews they compose, which we must take into consideration when building the prediction model.

2. **Writing Style:** Due to the large variation of the reviewers' background and language skills, the online reviews are of dramatically different qualities. Some reviews are highly readable and therefore tend to be more helpful, whereas some reviews are either lengthy but with few sentences containing author's opinions, or snappy but filled with insulting remarks. Simply ignoring such differences in readability and style may produce misleading estimates of the review quality. A proper representation of such difference must be identified and factored into the prediction model.
3. **Timeliness:** In addition to the available review content, most online reviews are also associated with a particular time stamp, which indicates when the review is posted. In general, the helpfulness of a review may significantly depend on when it is published. For instance, research shows that a quarter of a motion picture's total revenue comes from the first two weeks [5], which means a timely review might be especially valuable for users seeking opinions about the movie. As a concrete example, Figure 2 shows the average helpfulness of reviews versus the time the reviews are published (number of days since the release) for two movies, *Pirates of the Caribbean (Dead Man's Chest)* and *Casino Royal*. It is evident that the general trend is that the average helpfulness of movie reviews declines as time passes by. In addition, some previous studies made similar observations that product reviews written early tend to get more user attention on e-commerce websites, such as Amazon [11]. This further confirms our hypothesis that timely reviews are usually more helpful.

Other factors, such as server-side weblogs indicating how many users read but did not respond to the helpfulness question, might also facilitate the prediction. However they are not considered in our study due to the data availability issue.

3 HelpMeter: A Model for Helpfulness Prediction

Based on the observations from the previous section, we propose a model that accounts for these three important factors. Once trained, this model can be used for predicting the helpfulness of a given review. In the following discussions, we will use the IMDB movie data as a case study, although the model can be easily applied to other types of review data.

Since radial basis functions (RBF) are used in the modeling of both the expertise factor and the writing style factor (in the next subsection), a brief introduction of RBF is in order. After that, we will analyze how to model each factor mentioned in the previous section, and then present the non-linear regression model with all those factors incorporated, followed by a description of the training algorithm.

3.1 Radial basis functions

Function approximation is an important component to solving prediction problems defined over both continuous and discrete spaces. A powerful function approximator will not only accurately represent a value for a state it has experienced, but also generalize values to nearby states it has not experienced before. The most common type of approximator is the linear approximator. It has the benefit of being straightforward and involving lower computational cost, but it is obviously unreliable if the true relation between the inputs and the output is nonlinear. One then has to rely on non-linear approximators, such as RBF.

Radial basis functions have the advantage of being much simpler than other popular function approximators, such as multilayer perceptron neural networks, but still serving as a universal function approximator. They are generally used when local properties of the functional relationship needs to be captured, as is the case in the modeling of reviewer expertise and writing style. Due to its high flexibility, radial basis functions have been widely used in many areas, including finance[16], health care [29], and image processing [3]. Nonetheless, to the best of our knowledge, we are the first to model the multiple factors that may affect review helpfulness using a RBF non-linear approximator for review mining.

A radial basis function is a real-valued function whose value depends only on the distance of the input vector $\mathbf{x}$ from some *center* point $\boldsymbol{\mu}$. In the most general form, the RBF $\phi(\mathbf{x}|\boldsymbol{\mu}, \boldsymbol{\Sigma}) = f((\mathbf{x} - \boldsymbol{\mu})^T \boldsymbol{\Sigma}^{-1}(\mathbf{x} - \boldsymbol{\mu}))$, where f is the function used (Gaussian, Cauchy, etc.) and $\boldsymbol{\Sigma}$ is the metric. The term $(\mathbf{x} - \boldsymbol{\mu})^T \boldsymbol{\Sigma}^{-1}(\mathbf{x} - \boldsymbol{\mu})$ represents the distance between the input $\mathbf{x}$ and the center $\boldsymbol{\mu}$ in the metric defined by $\boldsymbol{\Sigma}$. Here, we choose the distance metric to be Euclidean. In this case, $\boldsymbol{\Sigma} = \sigma^2 \mathbf{I}$ for some scalar radius σ . Hence,

$$\phi(\mathbf{x}|\boldsymbol{\mu}, \boldsymbol{\Sigma}) = f\left(\frac{(\mathbf{x} - \boldsymbol{\mu})^T(\mathbf{x} - \boldsymbol{\mu})}{\sigma^2}\right). \quad (2)$$

The function f can take various forms. In this study, we choose the commonly used Gaussian RBF: $f\left(\frac{(\mathbf{x} - \boldsymbol{\mu})^T(\mathbf{x} - \boldsymbol{\mu})}{\sigma^2}\right) = e^{-\frac{(\mathbf{x} - \boldsymbol{\mu})^T(\mathbf{x} - \boldsymbol{\mu})}{\sigma^2}}$, where σ is also called the spread of the RBF. Intuitively, the further away $\mathbf{x}$ is from the center $\boldsymbol{\mu}$, the smaller the function value is, and the function peaks at the center when $\mathbf{x} = \boldsymbol{\mu}$. In addition, the value of the spread σ determines the “tightness” of the RBF, i.e., how fast the function value falls off when the input $\mathbf{x}$ gets further away from the center.

Multiple RBFs can be combined to build up function approximations of the form

$$g(\mathbf{x}) = \sum_{i=1}^k a_i \phi(\mathbf{x}|\boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i), \quad (3)$$

where the approximation function $g(\mathbf{x})$ is represented as a weighted sum of k radial basis functions, each with a different center $\boldsymbol{\mu}_i$, a metric $\boldsymbol{\Sigma}_i$, and a weight a_i . Such function approximation models are sometimes referred to in the literature as radial basis function networks. Figure 3 shows an example of using 3 radial basis functions to approximate a function. In this example, one would like to fit a function to the scattered data points. Although in our model for helpfulness prediction, we will deal with multi-dimensional input, for illustration purpose, this example deals with one-dimensional input. The fitted function represented by the solid line can be obtained by taking the weighted sum of the three individual RBFs. Fitting the data with the function involves determining the centers and spreads of the RBFs as well as the weight of each RBF.

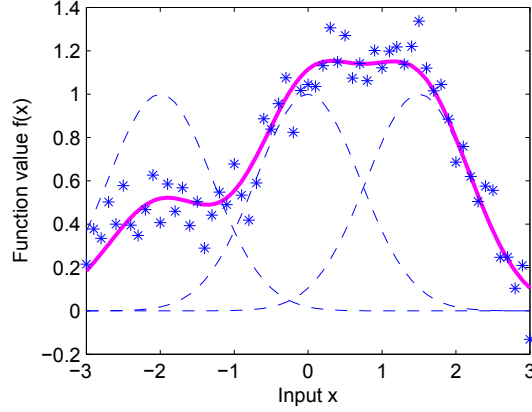


Figure 3. Using radial basis functions for function approximation. The solid line represents the fitted function, and the three RBFs are plotted with dotted lines.

3.2 Modeling expertise

As explained in Section 2, the helpfulness of a review depends in part on the level of expertise of the reviewer on the product (movie) concerned. For example, for a given reviewer, if his past reviews on a certain set of movies (denoted by $\mathcal{A}$) are rated very high while his reviews on some other movies (denoted by $\mathcal{B}$) are very low, then we have reasons to expect that a new review by this reviewer will be considered more helpful if the movie concerned is more similar to the movies in $\mathcal{A}$ than to those in $\mathcal{B}$.

In order to quantify the “similarity”, we first need to choose the right features to represent each movie. To this end, we use the genres provided by IMDB to represent each movie. As an example, the movie *Casino Royal* is labeled by IMDB as “Action”, “Adventure”, and “Thriller”, which can be used to represent the movie for our purpose. Formally, each movie is represented by a m -dimensional vector $\mathbf{x} = (x_1, x_2, \dots, x_m)$, where m is the number of different genres available for all movies. Each dimension corresponds to one genre, and x_i ($1 \leq i \leq m$) takes the value of $\frac{1}{l}$, if the movie belongs to the corresponding genre (where l is the number of genres the movie falls into), and 0 otherwise. Note that due to the normalization factor l , $\sum_{i=1}^m x_i = 1$.

The next step is to measure the similarity of a given movie to movies that have been reviewed by the same reviewer, and relate this measure to the helpfulness score. We choose to approximate the relationship using RBFs. If we were to predict the helpfulness of a review based solely on the reviewer expertise factor, then we would fit the following regression model on the training data.

$$\hat{H}_1 = \sum_{i=1}^{k_1} u_i \phi(\mathbf{x} | \boldsymbol{\mu}_i, \sigma_i), \quad (4)$$

where $\hat{H}_1$ is the estimated helpfulness score, $\mathbf{x}$ is the feature vector representing the movie, k_1 is the number of centers in the RBF network, $\boldsymbol{\mu}_i$ and σ_i are the center and spread of the i -th RBF respectively, and u_i is the weight of the i -th RBF.

RBFs are a particular good choice for modeling expertise in that when we represent each movie using a feature vector based on its genres, each center can be considered as corresponding to one “cluster” of movies that are similar to each other in terms of their genres. The helpfulness of a given movie is thus the weighted sum of the distance between the movie to those centers. In this way, the reviewer’s expertise on different clusters of movies can be naturally captured in that similar movies will have similar distances to the centers and therefore have similar helpfulness scores.

3.3 Modeling writing style

In Section 2, we explained that the writing style of reviews is an important factor in determining whether a review is helpful and should be taken into consideration in our model, and this observation is also supported by the previous study [31]. In fact, shallow syntactical features like part-of-speech provide more predicting powers than deeper features at the lexical level. Thus, we choose to label the part-of-speech of the words contained in the reviews with a fixed set of tags using LingPipe⁴, a suite of Java libraries for the linguistic analysis of natural language. For our purpose, we consider the tags that can potentially contribute to the differentiation of writing style due to their implication of the subjectivity/objectivity of a review. The tags chosen include:

1. qualifiers (e.g., quite, rather, enough),
2. modal auxiliaries (e.g., can, should, will),
3. nominal pronouns (e.g., everybody, nothing),
4. comparative and superlative adverbs (e.g., harder, faster, most prominent),
5. comparative and superlative adjectives (e.g., bigger, chief, top, largest),
6. proper nouns (e.g., Caribbean, Snoopy),
7. interjections/exclamations (e.g., ouch, well),
8. *wh*-determiners (e.g., what, which), *wh*-pronouns (who, whose, which), and *wh*-adverbs (e.g., how, where, when).

For each review, we parse it using the LingPipe tagger, and count the number of words with each of the above 8 tags. Those counts are further normalized by dividing them with the word count of the review. The resulting 8 numbers form a vector, denoted by $\mathbf{y}$, with each number corresponding to one dimension. This vector $\mathbf{y}$ is used as a representation of the review for the purpose of modeling writing styles.

We again use a radial basis function network to model the relationship between the feature vector $\mathbf{y}$ and the helpfulness of the review, with each RBF explaining part of the functional relationship, and the weights indicating the contribution of each RBF. Formally, if we were to predict the helpfulness solely based on the writing style, the regression model we would like to use is

$$\hat{H}_2 = \sum_{i=1}^{k_2} v_i \psi(\mathbf{y} | \boldsymbol{\nu}_i, \xi_i), \quad (5)$$

where $\hat{H}_2$ is the estimated helpfulness, v_i , $\boldsymbol{\nu}_i$, and ξ_i are the weight, center, and the spread of the i -th RBF respectively, and k_2 is the number of RBFs.

3.4 Modeling timeliness

Our analysis in Section 2 has shown that there is a strong correlation between the helpfulness of a review and when it is published. Having observed the trend for a large number of movies, we hypothesize that the helpfulness of a movie review is subject to exponential decay with respect to time. Therefore, we propose the following model for movie reviews if the prediction of helpfulness were to be done only based on the timeliness:

$$\hat{H}_3 = e^{-\beta(t-t_0)+d}, \quad (6)$$

where $\hat{H}_3$ is the estimated helpfulness, t_0 is the release time of the movie, t is the time when the review is published, and β and d are parameters in the model which are to be estimated. Intuitively, β controls the rate of decay in the helpfulness as we move further away in time from the movie release.

Note that the relationship between the helpfulness of a review and the time when the review is published depends on the type of the review subjects. For movies, books and many other products that do not evolve or

⁴<http://alias-i.com/lingpipe/>

change themselves after they become existent, timely reviews are usually more helpful. However, for some other types of subjects (such as hotels) which can improve or degrade over time, a more recent review (i.e., reviews closer to the current time) may be more helpful than earlier reviews since certain aspects of the subject may change over time, which makes the earlier reviews outdated. The timeliness models for the two types of subjects should be different. In this paper, we focus on modeling the review helpfulness for the first type of subjects. Nonetheless, we can easily adapt to modeling the timeliness of the second type by assigning higher weights to the current time, e.g., using logarithmic functions.

3.5 The complete model

Now that we have built the regression model for each individual factor, we are ready to propose the complete model, which we call the HelpMeter model, to incorporate all of the above factors. The idea is to consider the helpfulness score a weighted sum of the three individual models, as shown below:

$$\hat{H} = p \sum_{i=1}^{k_1} u_i \phi(\mathbf{x}|\boldsymbol{\mu}_i, \sigma_i) + q \sum_{i=1}^{k_2} v_i \psi(\mathbf{y}|\boldsymbol{\nu}_i, \xi_i) + r \cdot e^{-\beta(t-t_0)+d}, \quad (7)$$

where p , q , and r are the weights of the three components. Note that the above equation can be further simplified, as the weights p , q , and r can be “absorbed” by the individual components. For example, $p \sum_{i=1}^{k_1} u_i \phi(\mathbf{x}|\boldsymbol{\mu}_i, \sigma_i)$ can be rewritten as $\sum_{i=1}^{k_1} u'_i \phi(\mathbf{x}|\boldsymbol{\mu}_i, \sigma_i)$, where $u'_i = p \cdot u_i$, and $r \cdot e^{-\beta(t-t_0)+d}$ can be rewritten as $w \cdot e^{-\beta(t-t_0)}$, where $w = r \cdot e^d$. Therefore, the model can be written in a more concise form:

$$\hat{H} = \sum_{i=1}^{k_1} u_i \phi(\mathbf{x}|\boldsymbol{\mu}_i, \sigma_i) + \sum_{i=1}^{k_2} v_i \psi(\mathbf{y}|\boldsymbol{\nu}_i, \xi_i) + w \cdot e^{-\beta(t-t_0)}, \quad (8)$$

where the notations u_i and v_i are overloaded for the sake of brevity, with them actually referring to u'_i and v'_i as defined above.

The model given in Equation 8 makes it possible to capture all of the factors discussed in this section, with the weights $\{u_i\}_{i=1}^{k_1}$, $\{v_i\}_{i=1}^{k_2}$ and w controlling the “contribution” of each factor to the helpfulness score.

3.6 Parameter estimation

We now develop the algorithm that can be used to estimate the model parameters based on the training data (movie reviews). Assume that the training data consists of N reviews, and for each review j ($1 \leq j \leq N$), $\mathbf{x}_j$, $\mathbf{y}_j$, and t_j can be obtained, as well as the true helpfulness score H_j . The set of parameters in the model include

1. the weights $\{u_i\}_{i=1}^{k_1}$, $\{v_i\}_{i=1}^{k_2}$, and w ;
2. the centers $\{\boldsymbol{\mu}_i\}_{i=1}^{k_1}$ and $\{\boldsymbol{\nu}_i\}_{i=1}^{k_2}$,
3. the spreads $\{\sigma_i\}_{i=1}^{k_1}$ and $\{\xi_i\}_{i=1}^{k_2}$, and
4. the decay rate β .

The values of k_1 and k_2 are supplied by the user.

The goal of training is to estimate the parameters such that the sum of squared error (SSE) between the true values and the model output values is minimized, i.e., we would like to minimize

$$\varepsilon = \frac{1}{2} \sum_{j=1}^N \delta^2, \quad (9)$$

where $\delta = H_j - \hat{H}_j$. The optimization can be done through the method of steepest descent. By computing the partial derivatives of Equation 9, we can apply the following rules to iteratively update the values of the parameters as follows.

Let $\{\eta_u, \eta_v, \eta_w \dots\}$ be the user-defined learning rate for parameters $\{u_i, v_i, w \dots\}$ in the model.

1. For the weights, we have

$$\begin{aligned}
u_i^{new} &= u_i^{old} - \eta_u \frac{\partial \varepsilon}{\partial u_i} = u_i^{old} - \eta_u \sum_{j=1}^N \delta \phi(\mathbf{x}_j | \boldsymbol{\mu}_i^{old}, \sigma_i^{old}), \\
v_i^{new} &= v_i^{old} - \eta_v \frac{\partial \varepsilon}{\partial v_i} = v_i^{old} - \eta_v \sum_{j=1}^N \delta \psi(\mathbf{y}_j | \boldsymbol{\nu}_i^{old}, \xi_i^{old}), \\
w^{new} &= w^{old} - \eta_w \frac{\partial \varepsilon}{\partial w} = w^{old} - \eta_w \sum_{j=1}^N \delta e^{-\beta^{old}(t_j - t_{0,j})};
\end{aligned}$$

2. For the centers, we have

$$\begin{aligned}
\boldsymbol{\mu}^{new} &= \boldsymbol{\mu}^{old} - \eta_{\boldsymbol{\mu}} \frac{\partial \varepsilon}{\partial \boldsymbol{\mu}} \\
&= \boldsymbol{\mu}^{old} - 2\eta_{\boldsymbol{\mu}} \boldsymbol{\mu}^{old} \sum_{j=1}^N \delta \phi(\mathbf{x}_j | \boldsymbol{\mu}_i^{old}, \sigma_i^{old}) \frac{\mathbf{x}_j - \boldsymbol{\mu}_i^{old}}{(\sigma_i^{old})^2}, \\
\boldsymbol{\nu}^{new} &= \boldsymbol{\nu}^{old} - \eta_{\boldsymbol{\nu}} \frac{\partial \varepsilon}{\partial \boldsymbol{\nu}} \\
&= \boldsymbol{\nu}^{old} - 2\eta_{\boldsymbol{\nu}} \boldsymbol{\nu}^{old} \sum_{j=1}^N \delta \psi(\mathbf{y}_j | \boldsymbol{\nu}_i^{old}, \xi_i^{old}) \frac{\mathbf{y}_j - \boldsymbol{\nu}_i^{old}}{(\xi_i^{old})^2};
\end{aligned}$$

3. For the spreads, let $\omega = \frac{1}{\sigma^2}$, and $\zeta = \frac{1}{\xi^2}$, and we have

$$\begin{aligned}
\omega^{new} &= \omega^{old} - \eta_{\omega} \frac{\partial \varepsilon}{\partial \omega} \\
&= \omega^{old} + \eta_{\omega} u_i^{old} \sum_{j=1}^N \delta \phi(\mathbf{x}_j | \boldsymbol{\mu}_i^{old}, \omega_i^{old}) \\
&\quad \cdot (\mathbf{x}_j - \boldsymbol{\mu}_i^{old})^T (\mathbf{x}_j - \boldsymbol{\mu}_i^{old}) \\
\zeta^{new} &= \zeta^{old} - \eta_{\zeta} \frac{\partial \varepsilon}{\partial \zeta} \\
&= \zeta^{old} + \eta_{\zeta} v_i^{old} \sum_{j=1}^N \delta \psi(\mathbf{y}_j | \boldsymbol{\nu}_i^{old}, \zeta_i^{old}) \\
&\quad \cdot (\mathbf{y}_j - \boldsymbol{\nu}_i^{old})^T (\mathbf{y}_j - \boldsymbol{\nu}_i^{old})
\end{aligned}$$

4. Finally, for the decay rate β , we have

$$\beta^{new} = \beta^{old} - \eta_{\beta} \frac{\partial \varepsilon}{\partial \beta} = \beta^{old} + \eta_{\beta} w^{old} \sum_{j=1}^N (t_j - t_{0,j}) e^{-\beta(t_j - t_{0,j})}$$

4 Empirical Study

We conducted extensive experiments on the IMDB data set to evaluate the effectiveness of the proposed Help-Meter model and study the behavior of the model as we change the user-tunable parameters.

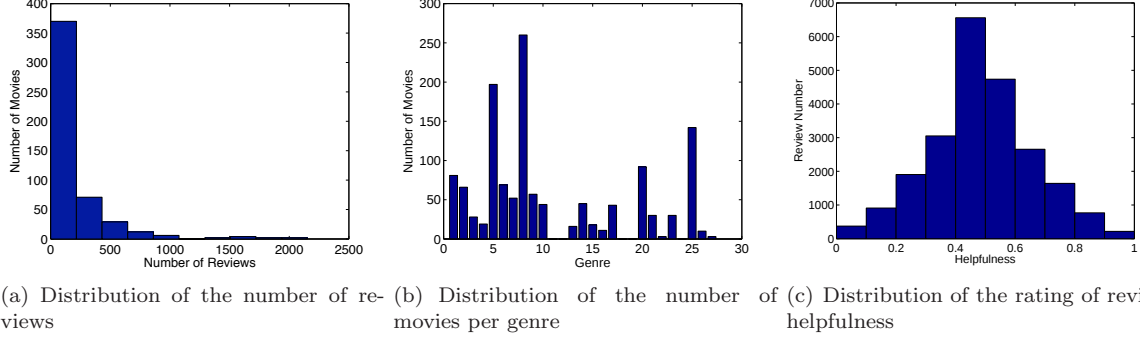


Figure 4. The distribution properties of the data

4.1 Experiment settings

The movie review data set was obtained from the publicly accessible IMDB Website. Specifically, we collected the reviews for 504 movies released in the United States during the period from January 6, 2006 to November 21, 2007. We intentionally selected the time that is not very close to the present time in the hope that the voting of helpfulness has stabilized, as less and less reviews are expected to appear as time increases across the whole time span. To model reviewer expertise, we also collected the genre labels for each movie, and the total genre number is 27. In total, 94,919 reviews were collected, and the number of review entries collected for each movie ranges from 2,152 (for *Superman Returns* [2006]) to 2 (for *Absolute Wilson* [2006]). Those reviews were posted by 56,588 different reviewers. Note that we only collected reviews posted by reviewers from the US as it helps to ensure the consistency in the release time (it is common for a movie to be released on different dates in different countries).

Figure 4 (a), (b), and (c) show the distributions of the number of reviews available for movies, the number of movies per genre, the ratings of helpfulness respectively. To ensure the robustness of the prediction model, we only use the reviews with at least 10 votes, and reviewers with at least 25 posts. Also, for the purpose of training and testing, only the reviews with a helpfulness score available are used. The number of such movie reviews is 22,819. The movie information (genres for each movie) and the review data are indexed using Apache Lucene⁵. For each review, its feature vectors are obtained as described in Section 3, and we use 10-fold cross validation to evaluate our approach.

4.2 Evaluation Metrics

We evaluate the effectiveness of the proposed model using two metrics as we anticipate that the model will be used in different ways. First, the model can be used to predict the helpfulness of reviews directly, so we would like to measure the deviation of the predicted value from the true value. We call this a prediction problem. Second, the model can be also used to help retrieve only those reviews that are considered helpful, i.e., the reviews having a predicted helpfulness higher than a certain threshold. We call this a classification (or retrieval) problem.

Two metrics, which were used in previous literatures [30, 6], are adopted to evaluate the predication accuracy in those two scenarios respectively, namely, the *Mean Squared Error* (MSE) (for the prediction problem) and the *F-measure* (for the classification problem). Specifically, for each review in the test set, we make a prediction for its helpfulness and compute the squared deviation between the predicated value and the true helpfulness. MSE is defined as the sum of all the deviations divided by the total number of predictions. That is,

$$MSE = \frac{1}{n} \sum_{i=1}^n (H_i - \hat{H}_i)^2. \quad (10)$$

where n is the number of reviews in the test set. Note that lower MSE values indicate higher prediction accuracy. To measure the performance using F-measure, we consider a review as helpful if its helpfulness score is greater

⁵<http://lucene.apache.org>

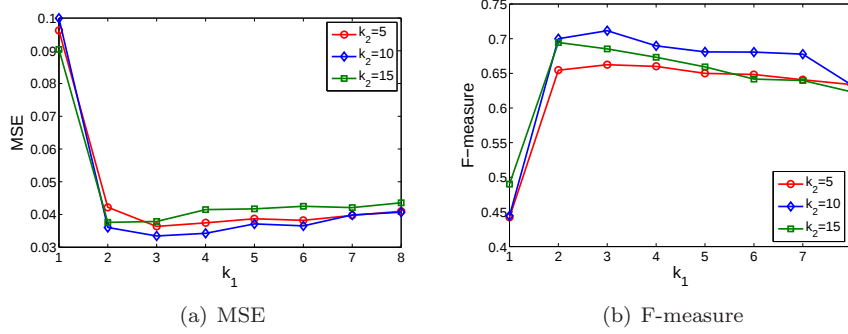


Figure 5. Effect of k_1 on prolific reviewers

than a given threshold θ . In our experiments, we first set $\theta = 0.6$, which has been adopted in previous studies [6], and then investigate how the F-measure will change with various θ in Section 4.7. Let $\hat{\mathcal{X}}$ be the set of reviews that are predicted to be helpful, and $\mathcal{X}$ be the set of reviews that are truly helpful. Then,

$$\text{Precision} = \frac{|\mathcal{X} \cap \hat{\mathcal{X}}|}{|\hat{\mathcal{X}}|}, \text{Recall} = \frac{|\mathcal{X} \cap \hat{\mathcal{X}}|}{|\mathcal{X}|}, \text{and,}$$

$$F = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}.$$

Note that larger values of F-measure indicate better accuracy.

4.3 Parameter selections

In the HelpMeter model, there are two user-chosen parameters that provide the flexibility to fine tune the model for optimal performance, i.e., the number of RBFs in the RBF network k_1 and k_2 . We now study how the choice of these parameter values affects the prediction accuracy.

We first vary the values of k_1 , and observe from Figure 5 that there is a large improvement in accuracy when k_1 increases from 1 to 2, and the model achieves its best performance of $\text{MSE} = 0.0332$ and $\text{F-measure} = 0.7382$ at $k_1 = 3$. This implies that introducing multiple components to analyze the reviewer expertise can greatly improve the prediction accuracy. However, after k_1 past a threshold, the accuracy tends to decrease. This might be due to over-fitting the training data with more RBFs. Nonetheless, the accuracy remains stable for a wide range of k_1 values, indicating the insensitivity of the model with respect to the choice of k_1 values. It is also worth noting that the trend in accuracy remains the same regardless of the choice of k_2 .

Similarly, we fix the values of k_1 , and vary k_2 from 1 to 12. As shown in Figure 6 (b), there is also an optimal choice of k_2 , which is 10. Similar to the case of k_1 , the accuracy remains quite stable over a wide range of k_2 , which again demonstrates that the model is not sensitive to the choice of parameter values. Note that we let $k_1 \in [2, 4]$ in this experiment, as we believe that a normal author is likely to be specialized in writing reviews of only a few types of movies.

4.4 Effects of individual factors

In our study, three factors that may affect the review helpfulness are considered, and we propose a non-linear regression model to incorporate them into one model. Here, we study how the three factors affect the prediction of helpfulness individually. That is, how would the model perform if we choose to use only one of the factors for prediction? In Section 3, we discussed three models (Equations 4, 5, and 6) corresponding to the three factors. For the experiments, we train the three individual models as presented in Section 3 with the corresponding feature vectors and measure the accuracy of each one. In particular, we let $k_1 = 3$ and $k_2 = 10$ in this experiment, and the results are shown in Table 2.

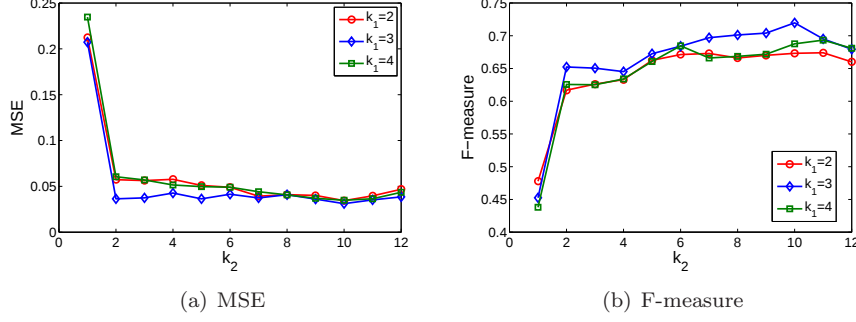


Figure 6. Effect of k_2 on prolific reviewers

Component	MSE	F-measure
reviewer expertise	0.1912	0.5179
writing style	0.1937	0.2433
timeliness	0.1004	0.6482
all of the three	0.0332	0.7382

Table 2. Individual effects on helpfulness prediction

Apparently, considering the timeliness factor only yields the best results in both MSE and F-measure among the three factors. This coincides with our intuition that a timely review can be very helpful for customers to evaluate the product of interest. Only considering the writing style gives the worst performance of the three, implying that it has less predictive power compared with reviewer expertise and timeliness. Note that the results obtained by considering only one factor are not as good as considering all the factors together.

4.5 Comparison with alternative methods

To evaluate the effectiveness of the proposed algorithm, in this section, we compare the performance of our prediction model with those of three popular regression methods, i.e., linear regression (LR), support vector machine regression, and multilayer perceptron model. Specifically, as the output of each regression model is a real number, we adopt *MSE* as the evaluation metric in these experiments.

4.5.1 Linear Regression

To demonstrate the effectiveness of our proposed model, we first compare it against a baseline model that use linear regression (LR). For each review, we obtain the feature vectors $(\mathbf{x}, \mathbf{y}, t)$ corresponding to each factor in the same way as described in Section 3 and concatenate them together to form one vector $\mathbf{r}$. Then the linear regression model can be written as

$$\hat{H}_t = \beta^T \mathbf{r} + b, \quad (11)$$

where β is the coefficient vector and b is the intercept. This model can be fit to the training data using standard linear least squares method. We let $k_1 = 3$, and $k_2 = 10$ in this experiment, and the result is shown in Table 3.

4.5.2 Support Vector Machine Regression

We then compare the performance of our model with that of the state-of-the-art Support Vector Machine regression method, which has been widely applied in previous studies [14, 18, 30]. For each review, we build a feature vector in the same way as described above, and adopt the same settings as presented in [30].

Method	MSE
Linear Regression	0.2861
SVM Regression	0.1295
MLP	0.0847
HelpMeter	0.0332

Table 3. Comparison with other methods

4.5.3 Multilayer Perceptrons

We next compare our method with a popular type of neural networks, namely, the multilayer perceptron (MLP) model. Different from standard linear perceptrons, an MLP uses multiple layers of neurons with nonlinear activation functions, and thus is expected to be more powerful for distinguishing non-linearly separable data.

In this experiment, we first formulate the feature vector for each review as described in Section 4.5.1, and then construct two double-layer networks for prolific users and non-prolific users respectively. We finally use the most commonly used Batch Gradient Descent backpropagation algorithm [7] as our training function, and obtain the network parameters by feeding the training data. In the experiment, we varied the number of neurons in the first and second layers, and observed that typically the best results are obtained when the numbers of neurons in the two layers are 3 and 1 respectively, and therefore the results under those settings are presented here.

As shown in Table 3, our proposed HelpMeter model shows clear advantage over the popular SVM regression model, the baseline linear regression method, as well as the MLP. We believe that the difference in accuracy is due to the fact that the various factors that affect the helpfulness rating have their own unique characteristics; therefore should be modeled separately. HelpMeter is more capable of such a task.

4.6 Prolific vs non-prolific reviewers

Recall that in modeling the reviewer expertise as described in Section 3.2, we rely on the genres of the movies the reviewer has commented and the corresponding helpfulness scores. This requires sufficient past reviews of the reviewer in order to achieve meaningful results. In reality, some reviewers may have written only a few or none reviews, or the reviews a reviewer has written may not be present due to data availability issues. We therefore make the distinction between prolific reviewers and non-prolific reviewers and revise the model correspondingly. We call a reviewer a prolific reviewer if the number of reviews authored by him/her in the data set exceeds a certain threshold T , and non-prolific otherwise. For prolific users, we simply use the model described in Section 3.5, whereas for non-prolific ones, we need to drop the first term regarding reviewer expertise in the model, as we do not have sufficient grounds to make meaningful inference in that regard. In that case, the model becomes

$$\hat{H} = \sum_{i=1}^{k_2} v_i \psi(\mathbf{y} | \boldsymbol{\nu}_i, \xi_i) + w \cdot e^{-\beta(t-t_0)}. \quad (12)$$

Note that since the above model does not involve any information regarding individual reviewers, a common model can be trained for all of the non-prolific reviewers. The parameter estimation can be done using the update formulae presented in Section 3.6. It is worth pointing out that the distinction between prolific and non-prolific reviewers is due to data availability; we do not assume that the reviews written by prolific reviewers are more helpful than those written by non-prolific reviewers.

To further investigate how the values of T affect the performance of helpfulness prediction, in this experiment, we use a threshold T to distinguish a prolific reviewer from a non-prolific reviewer, based on how many reviews in the data are authored by that reviewer. We train different models for the two types of reviewers. With fixed values of k_1 and k_2 ($k_1 = 3$, and $k_2 = 10$), we vary T , and observe the changes in accuracy. Similar trends can be observed for other values of k_1 and k_2 . As shown in Table 4, as T increases from 10 to 30, the prediction performance improves in both F-measure (for the classification problem) and MSE (for the prediction problem), and at $T = 30$, it achieves the best accuracy with F-measure=0.7382 and MSE=0.0332. This implies that accumulating more reviews for a given author allows our model to better capture the effects that influence the helpfulness, which leads to more accurate predictions. In addition, the accuracy for prolific reviewers is much

T		MSE	F-measure	Reviews#	Reviewers#
10	P	0.0486	0.6717	2378	109
	N	0.0768	0.4307	20441	17266
15	P	0.0392	0.6748	1814	65
	N	0.0632	0.4307	21005	17310
20	P	0.0386	0.6886	1258	33
	N	0.0668	0.4364	21561	17342
25	P	0.0354	0.6989	1079	25
	N	0.0661	0.4363	21740	17350
30	P	0.0332	0.7382	912	19
	N	0.0658	0.4365	21907	17356

Table 4. Effect of T . N and P refer to non-prolific and prolific reviewers respectively.

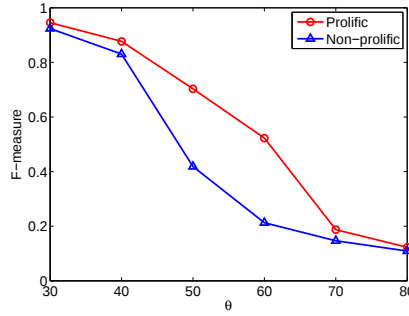


Figure 7. Effect of θ

superior to that for non-prolific reviewers across different values of T , indicating the effectiveness of the “reviewer expertise” factor in the model for prolific reviewers. Note that less training data become available as the increase of T , because relatively fewer reviewers are likely to post reviews of a large number. In this experiment, we stop at $T = 30$ to make our evaluation robust.

4.7 Effect of θ

In classifying a review as helpful or not helpful, we use a fixed threshold $\theta = 0.6$ in previous experiments. In this experiment, we examine the effect of the value of θ on the accuracy, and demonstrate the result in Figure 7. Clearly, smaller θ values tend to lead to better accuracy. This is as expected, because a larger θ means less reviews can be classified as helpful, and is therefore more restrictive, making accurate classifications more difficult.

5 Related Work

5.1 Review mining

With the rapid growth of online reviews, automatic review mining has attracted a lot of research attention. Early work in this area was primarily focused on determining the semantic orientation of reviews. Among them, some of the studies attempt to learn a positive/negative classifier at the document level [24, 23], while others work at a finer level and use words as the classification subject [13, 28]. Pang et al. [24] employ three machine learning approaches (Naive Bayes, Maximum Entropy, and Support Vector Machines) to label the polarity of IMDB movie reviews. In a follow up work, they propose to firstly extract the subjective portion of text with a graph min-cut algorithm, and then feed them into the sentiment classifier [23]. There are also studies that work at a finer level and use words as the classification subject. They classify words into two groups, “good” and “bad”, and then use certain functions to estimate the overall score for the documents. Kamps et al. [13] propose to evaluate the

semantic distance from a word to good/bad with WordNet. Turney [28] measures the strength of sentiment by the difference of the mutual information (PMI) between the given phrase and “excellent” and the PMI between the given phrase and “poor”. Pushing further from the explicit two-class classification problem, Liu et al. [9] build a framework to compare consumer opinions of competing products using multiple feature dimensions. After deducting supervised rules from product reviews, the strength and weakness of the product are visualized with an “Opinion Observer”. Liu et al. [19] assume that sentiment consists of multiple hidden aspects, and use a probability model to quantitatively measure the relationship between sentiment aspects and reviews, as well as sentiment aspects and words.

Compared to sentiment mining, identifying the quality of online reviews has received relatively less attention. A few recent studies along this direction attempt to detect the spams or low-quality posts that exist in online reviews. Jindal et al. [11] present a categorization of review spams, and propose some novel strategies to detect different types of spams. Liu et al. [18] propose a classification-based approach to discriminate the low-quality reviews from others, in the hope that such a filtering strategy can be incorporated to enhance the task of opinion summarization. Our work can be considered complimentary to those studies in that the spam filtering model can be used as a preprocessing step in our approach. Besides, there are also studies focusing on investigating how the different content features may affect the quality of reviews [6, 14, 30].

In our work, we not only investigate the impact of the basic semantic content, such as syntactical features, but introduce other major factors that may effect the review helpfulness rating, e.g., the reviewers expertise and the timeliness. In addition, we explore the possibility of developing a prediction model that can capture the distinctive characteristics of various factors instead of relying on the simple off-the-shelf solutions.

5.2 Ranking and Recommender Systems

Recommender systems have emerged as an important solution to the information overload problem where people find it more and more difficult to identify the useful information effectively. Studies in this area can generally be divided into three directions: content-based filtering, collaborative filtering, and hybrid systems. Content-based recommenders rely on rich content descriptions of behavioral user data to infer their interest, which places significant engineering challenge for researchers as the required domain may not be readily available or easy to maintain. As an alternative, collaborative filterings (CF) take the rating data as input, and applying data mining or machine learning approaches to discover usage patterns that represent the user models. When a new user comes to the site, his/her activity will be matched against these patterns to find like-minded users and select possible interesting items as recommendations.

Various CF algorithms ranging from typical nearest-neighbour methods [26] to more complex probabilistic based methods [8, 25] have been designed to identify users of similar interests. A few variations and hybrid methods that combines both content information and collaborative filtering have also been proposed [2, 10].

Similar to other recommender systems, our work also accounts for textual contents and peer votes in those reviews to effectively construct and evaluate the prediction model; however, the objective of this paper is to investigate the helpfulness of movie reviews, which is different from the above work.

5.3 Authority and importance mining

Identifying the quality of Web documents has received a lot of attention, particularly because of its application to search engines. PageRank and HITS are two popular link-based ranking algorithms to determine the importance of web pages [15, 22]. The HITS algorithm is based on the observation that a good hub usually points to good authorities and a good authority usually points to good hubs. The Pagerank algorithm doesn’t distinguish hub and authority pages. Instead, it estimates the importance of the web page’s neighbours, and the authority of the page is considered proportional to this value.

Motivated by this idea, various algorithms have been proposed to discover the authorities or leaders in the Web domain. Jurczyk et al. [12] study the link structure in the question answering (QA) community to discover authoritative users in topical categories. Song et al. [27] develop the InfluenceRank algorithm to identify influential opinion leaders from blogosphere. Nakajima et al. [21] characterize bloggers based on their roles according to previous blog threads. Based on permulink analysis, people in blogosphere are divided into agitators who stimulate discussions and summarizers who summarize discussions.

Note that our approach is different from above methods in that we use semantic information of web document rather than link structures for evaluating the helpfulness of online reviews.

6 Conclusions and Future Work

In this paper, we have considered the important problem of predicting the helpfulness of reviews. We provided a detailed analysis of the major factors affecting the helpfulness of a review, and proposed a nonlinear model based on radial basis functions for helpfulness prediction. Extensive experiments on the IMDB data set have confirmed the effectiveness of the proposed model.

Our study in this paper has focused on the movie reviews, but our approach is general enough to be easily adapted to other domains as well. For example, if we would like to handle product reviews on Amazon or CNET, we can simply replace the genres of movies with the categories of products, and the writing style and timeliness can still be modeled in the similar way as described in Section 3.

This study presents the first step in modeling the helpfulness of reviews. For future work, we plan to study the related ranking problem, i.e., how do we rank the reviews based on the helpfulness? One way to do this is to rank the reviews based on their predicted helpfulness, but we can also develop a model to directly predict the set of most helpful reviews. Another possible direction for future work is to incorporate existing votes as an indicator of the future helpfulness, and build an adaptive model which can automatically update the predication value of helpfulness as new reviews come in. In the adaptive model, other possible timeliness models can also be investigated to accommodate the situation that some products reviews, such as those on hotels, may evolve with the changes of product quality over time. In that scenario, a more recent review (i.e., reviews closer to the current time) may be more helpful than earlier reviews since certain aspects of the subject may change over time, which makes the earlier reviews outdated; we thus need to model the factor of timeliness differently. Besides, we also plan to incorporate collaborative filtering methods, such as [8, 10], to help build a personalized helpfulness prediction model.

References

- [1] Nikolay Archak, Anindya Ghose, and Panagiotis G. Ipeirotis. Show me the money!: deriving the pricing power of product features by mining consumer reviews. In *KDD*, pages 56–65, 2007.
- [2] Justin Basilico and Thomas Hofmann. Unifying collaborative and content-based filtering. In *ICML*, page 9, 2004.
- [3] J. C. Carr, R. K. Beatson, J. B. Cherrie, T. J. Mitchell, W. R. Fright, B. C. McCallum, and T. R. Evans. Reconstruction and representation of 3D objects with radial basis functions. In *SIGGRAPH '01*, pages 67–76, 2001.
- [4] Kushal Dave, Steve Lawrence, and David M. Pennock. Mining the peanut gallery: opinion extraction and semantic classification of product reviews. In *WWW*, pages 519–528, 2003.
- [5] Chrysanthos Dellarocas, Neveen Farag Awad, and Xiaoquan (Michael) Zhang. Exploring the value of online reviews to organizations: Implications for revenue forecasting and planning. In *ICIS*, pages 379–386, 2004.
- [6] Anindya Ghose and Panagiotis G. Ipeirotis. Designing novel review ranking systems: predicting the usefulness and impact of reviews. In *ICEC*, pages 303–310, 2007.
- [7] M.T. Hagan, H.B. Demuth, and M.H. Beale. *Neural Network Design*. MA: PWS Publishing, 1996.
- [8] Thomas Hofmann and Jan Puzicha. Latent class models for collaborative filtering. In *IJCAI*, pages 688–693, 1999.
- [9] Mingqing Hu and Bing Liu. Mining and summarizing customer reviews. In *KDD*, pages 168–177, 2004.
- [10] X. Jin, Y. Zhou, and B. Mobasher. A maximum entropy web recommendation system: combining collaborative and content features. In *KDD*, pages 612–617, 2005.

- [11] Nitin Jindal and Bing Liu. Opinion spam and analysis. In *WSDM*, pages 219–230, 2008.
- [12] Pawel Jurczyk and Eugene Agichtein. Discovering authorities in question answer communities by using link analysis. In *CIKM*, pages 919–922, 2007.
- [13] Jaap Kamps and Maarten Marx. Words with attitude. In *Proc. of the First International Conference on Global WordNet*, pages 332–341, 2002.
- [14] S.-M. Kim, P. Pantel, T. Chklovski, and M. Pennacchiotti. Automatically assessing review helpfulness. In *EMNLP*, pages 423–430, 2006.
- [15] Jon M. Kleinberg. Authoritative sources in a hyperlinked environment. *J. ACM*, 46(5):604–632, 1999.
- [16] A. Lendasse, J. Lee, E. de Bodt, V. Wertz, and M. Verleysen. *Connectionist Approaches in Economics and Management Sciences*, chapter Approximation by Radial-Basis Function networks - Application to option pricing, pages 203–214. Kluwer academic publishers, 2003.
- [17] Bing Liu, Mingqing Hu, and Junsheng Cheng. Opinion observer: analyzing and comparing opinions on the web. In *WWW*, pages 342–351, 2005.
- [18] Jingjing Liu, Yunbo Cao, Chin-Yew Lin, Yalou Huang, and Ming Zhou. Low-quality product review detection in opinion summarization. In *EMNLP*, pages 334–342, 2007.
- [19] Yang Liu, Xiangji Huang, Aijun An, and Xiaohui Yu. ARSA: a sentiment-aware model for predicting sales performance using blogs. In *SIGIR*, pages 607–614, 2007.
- [20] Yang Liu, Xiangji Huang, Aijun An, and Xiaohui Yu. Mining helpfulness of online reviews. Technical Report CSE-2008-05, Department of Computer Science and Engineering, York University, 2008.
- [21] Shinsuke Nakajima, Junichi Tatemura, Yoichiro Hino, Yoshinori Hara, and Katsumi Tanaka. Discovering important bloggers based on analyzing blog threads. In *WWW*, 2005.
- [22] Lawrence Page, Sergey Brin, Rajeev Motwani, and Terry Winograd. The pagerank citation ranking: Bringing order to the web. Technical report, Stanford Digital Library Technologies Project, 1998.
- [23] Bo Pang and Lillian Lee. A sentimental education: sentiment analysis using subjectivity summarization based on minimum cuts. In *ACL*, page 271, 2004.
- [24] Bo Pang, Lillian Lee, and Shivakumar Vaithyanathan. Thumbs up? sentiment classification using machine learning techniques. In *EMNLP*, 2002.
- [25] Alexandrin Popescul, Lyle H. Ungar, David M. Pennock, and Steve Lawrence. Probabilistic models for unified collaborative and content-based recommendation in sparse-data environments. In *UAI*, pages 437–444, 2001.
- [26] Badrul Sarwar, George Karypis, Joseph Konstan, and John Reidl. Item-based collaborative filtering recommendation algorithms. In *WWW*, pages 285–295, 2001.
- [27] Xiaodan Song, Yun Chi, Koji Hino, and Belle Tseng. Identifying opinion leaders in the blogosphere. In *CIKM*, pages 971–974, 2007.
- [28] Peter D. Turney. Thumbs up or thumbs down?: semantic orientation applied to unsupervised classification of reviews. In *ACL*, pages 417–424, 2001.
- [29] Hong-Qiang Wang and D.S.Huang. Non-linear cancer classification using a modified radial basis function classification algorithm. *Journal of Biomedical Science*, 12(5):819–826, 2005.
- [30] Zhu Zhang and Balaji Varadarajan. Utility scoring of product reviews. In *CIKM*, pages 51–57, 2006.
- [31] Li Zhuang, Feng Jing, and Xiao-Yan Zhu. Movie review mining and summarization. In *CIKM*, pages 43–50, 2006.