

RATING SYSTEM BY ANALYZING THE REVIEWS

P.RAJASEKHARAN¹ G THANUJA² JAINA SAICHAND³ G THANUJA⁴
^{1,2,3,4} ASST. PROFESSOR, DEPARTMENT OF ARTIFICIAL INTELLIGENCE,
^{1,2,3,4} SRI MITTAPALLI COLLEGE OF ENGINEERING

Abstract:

Consumers and businesses alike may benefit from this study's examination of the potential impact of review votes and user reviews on feature-level assessments of mobile devices. This grading system gives a fuller picture of the product than a product-level grading system. We know exactly what makes a product good or bad, therefore feature-level assessments are more detailed than product-level evaluations. Data on what consumers value most and least has always been important. It tells the manufacturer and the buyer about upcoming product upgrades and sales. Various categories of customers are drawn to distinct qualities. Consequently, feature-level assessments may help make customers' purchase selections more personalised. Using data collected from an online retailer (Amazon), we analyse user evaluations and review votes pertaining to various mobile devices. We do this by using a sentiment analysis that is feature-centric. When everything is said and done, we have ratings for 108 distinct features of more than 4000 distinct online-sold mobile phones. From one side, it helps manufacturers make better decisions about product enhancement, while from the other, it allows customers to make more personalized purchases. Recommendation systems, market research, and other areas might benefit from our work.

1. Introduction

The proliferation of the Internet and the rising busyness of people's lives have contributed to the widespread use of online shopping. Reviews written by other customers and posted online typically influence consumers' final purchase decisions. But most user reviews you'll find online are shallow, generic rants about products. Consequently, although product-level evaluations are useful for comparing various commodities, some people will still opt to buy certain things due to their distinctive qualities. If you want to know what other people think about a product's features, you usually have to read the comments section in its entirety. Choosing the best choice for a customer could be challenging when there are several possibilities for a single product (like a mobile phone). Such product-level evaluations also provide minimal space for interpretation, which is problematic from the manufacturer's perspective [3]. Manufacturers have a better understanding of how to produce their goods when feature-level evaluations are available. We could make use of all these benefits with a feature-level grading system.



Fig. 1. Product-level rating system versus feature-level rating system: while the first one is very generic, the latter is quite specific.

Customers may desire feature ratings in the same way they want product evaluations, but it's not a good idea since there might be too many features. It is far more practical to use current customer evaluations and review votes to provide feature-level ratings [4, 5]. The evaluations are in the form of sentences, and each word expresses an emotion, neutral or positive, from [6] to [9]. Their separability allows us to extract specific words connected to products and use them to derive sentiment assessments. We can build a feature-level rating system that can provide feature-level ratings by using the sentiment scores as a base and the review votes as a backbone. Figure 1 shows the result. However, there are just a handful of obstacles to overcome when developing a feature-level grading system. Identifying the features we want in a product should be our first priority. One further issue is that there are several ways to define a feature, and they all need to be merged into one. Step two involves cleaning up the data; for example, some review comments may not be in English or could have typos. The third step in assigning a fair rating to a product feature is figuring out how to include the review votes with the extracted sentiment ratings. Looking at the word frequency table for the whole set of customer evaluations for a product allows us to identify its features, such as mobile, up to a specific frequency.

2. Literature Survey

An extensive period has passed since sentiment analysis [11]-[18] was a popular area of research. In addition to e-commerce, it has uses in fields such as health care[19–21], politics[12–13], and sports [14–15]. It has been shown in [18] and [19] that a lot of information on the products may be gleaned from sentiment analysis of online reviews. Sentiment analysis of mobile brand mentions on Twitter was the main focus of Wiliam et al. [11]. Having said that, the studies don't focus in on specific traits, but rather examine mobility in general. This was tried by Nandal et al. [10], albeit with a limited set of commodities and features. In contrast to Sadhasivam and Kalivaradhan [11], who used ensembling, Nandal et al. [10] employed the supervised learning approach SVM. This post will use user reviews to rank 108 attributes of more than 4000 mobile phones sold online, an unprecedented effort. Additionally, we create unsupervised sentiment evaluations for sentences using the lexical technique, which is

neither supervised nor weakly supervised [13]-[15]. Our goals in studying digital cameras and TVs align with those of Zhang et al. [16]. Although it is a brief examination, ten aspects are covered. The only devices that Bafna and Toshniwal tested were an MP3 player, a Canon camera, and an iPhone 4s [18]. In contrast, before offering recommendations, we research and analyse over four thousand goods. But this is the only article that we are aware of that uses review votes as part of a feature-level evaluation.

3. Methodology

According to this method, there are three steps: The method's first four phases are feature selection, preprocessing, feature-based rating construction, and algorithm execution. This section provides a comprehensive overview of each level.

Feature Selection

Assume that we manually scan through the word frequency table of the whole set of customer review data for a product (a mobile device in this example) and compile a set of N feature-related words, denoted as $W = w_1, \dots, w_N$. Therefore, the traits that people care about the most could be located. If a phrase appears less than 0.02% of the time in the review data, it is likely that it is not a feature-related word and is therefore disregarded. The frequency of words associated to features is denoted by $Z = z_1$, whereas the remainder of the frequency distribution is denoted by N . Every word in W that is connected to a specific feature, including the feature itself, should be returned by a relationizer function. This is because all words linked to a feature should be integrated into a single feature. The relationizer function in question is known as $R(W, w_i)$. This relationizer function we're discussing here is completely manual, so bear that in mind. Our feature data set, denoted by the letter F , will be defined as a collection of connected feature words in the following way: In the first expression, $F = R(W, w_i)_{i=1, \dots, N}$ (1), the relationizer function is used to iteratively generate sets of related words from all the words in W . Since many related words will produce identical sets, we shall eliminate duplicates in order to divide the sets on the left. For convenience's sake, let's refer to the F feature data set as F and consider F_k to be the k th set of feature words. The whole set of feature words in any F_k may be located by selecting a single sample feature word. These individuals will be called feature keywords from now on. The most frequent feature word in the set is chosen as the representative (based on [18]) or feature keyword for the name of the set: Name the feature set after the term that occurs in it the most frequently in this instance. In the future, F_k may stand for both the keyword and the k th set of feature words, which is a group of related feature words.



Fig. 2. Preprocessing: useful characters retention, where unnecessary characters are removed in this illustration.

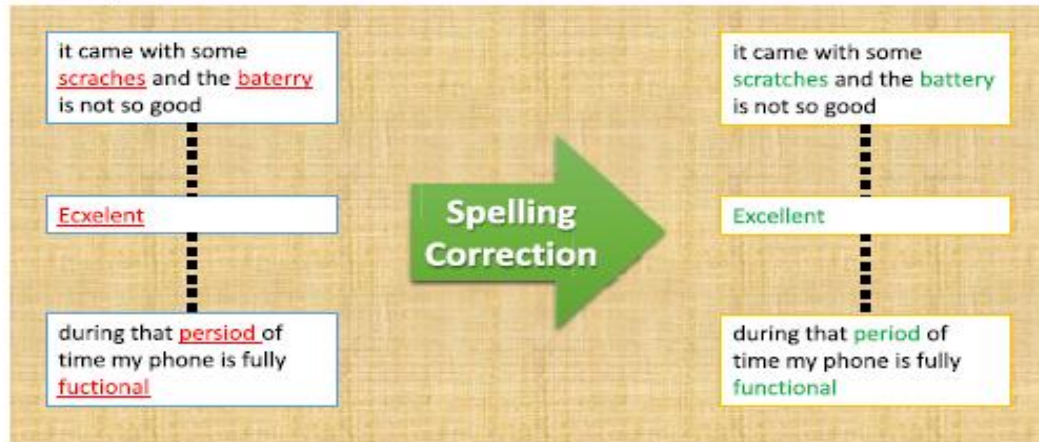


Fig. 3. Preprocessing: spelling correction, where misspelled words, such as “scraches,” “baterry,” “Excelent,” “period,” and “fuctional,” are corrected in this illustration.



Fig. 4. Preprocessing: keyword correction, where feature words, such as “speakers” and “photo,” are corrected to their respective feature keywords, “sound” and “pictures,” in this illustration.

Preprocessing

Unstructured data is common in review comments as customers post them publicly. To facilitate retrieval, disintegration, and correction of pertinent information, our preprocessing procedures aim to transform this unstructured input into structured data. Remember that after eliminating unnecessary characters, the remaining information is the important data. We made a mental note of how positive people were when reading the reviews to decide which features to include. People often utilise punctuation, emojis, and adjectives. Figure 2 shows that although punctuation, emoticons, and word structures are preserved, all other letters and numbers are deleted. We list all characters to be thorough. Following that, we will remove any blank items as they are not used by our feature-level grading system. Think about it this way: given a product (not an item), we may say that the review data is $D = C_1, C_2, \dots, C_m$, and that it contains m useful

remarks.

Data from consumer evaluations of the product in issue, including reviews of all mobile phones, was the last component of our feature selection process. However, data from a single mobile product review is represented by D.

Generation of Ratings Based on Specifications

We may now go through each of the remaining sentences, filtering out irrelevant ones, to determine whether any of them make reference to a certain trait, such as Fk. If that's the case, we may calculate the score based on the tone of the statement. In each of these passages, we reach this conclusion using sentiment analysis scores [19]. We use the code in [19] because it is versatile. Feel free to include emoticons into your research as well. Their services are used by us. For sentiment analysis, a compound score is required. It lies between one and one. The next step is to average the quality score across all product attributes based on user feedback. During this period, we amassed quite a bit of possessions. When considering who should hear the comment, we also look at the review's vote count. Based on these findings, we may infer that the reviews' claims are well-supported. I propose this next. (Cj) Display the total number of votes cast for Cj. We've arrived at the notion that, from this person's point of view, each sentence contributes equally to the total strength. Because of this, we adjust the tally of the original reviewers' self-votes by adding 1.

4. Results

Even though there aren't any objective, true feature ratings, we may nevertheless evaluate our method on a subset of the phone's testable capabilities. We have what it takes to pull this off. View our results on the go with the help of the built-in app named The bundle already includes general phone ratings. When evaluating the 'feature' of our phone, we take into account the weighted and averaged evaluations of individual customers. In order to compare the error metric numbers in Table with the truth-to-ground scores, we need to provide different information. Considering the average proposes an error measure based on the average error rate (MAER), it is astounding to see that the rating is just half a star, or 0.55 points, off from the real ratings. We have also made a confusion matrix for phone features. To get such a high total number of stars, we rounded our ratings to the closest whole number. This led to the present system of categorising phone ratings into five separate academic categories, ranging from one to five stars. Such a discordant results matrix has led us to cause further perplexity. Put simply, our system reliably predicts consumer satisfaction ratings for 265 mobile phones, with ratings that are all within one star of each other for 386 out of 141 mobile phones. Consequently, we employ this method 72.3% of the time to find out how accurate the proposed integer star rating is. However, if we can tolerate a one-star improvement in accuracy in the integer star rating prior to 94.83 percent of the time, Thus, it is reasonable to presume that the intended This method works well with the phone's advertised capability. The accurate number of mobile phones is 418, however there may be certain devices that don't allow our highlight word selection due to terminology related to the phone's operation, so it's important to keep that in mind.

5. Conclusion

We've come up with a system to rate cell phones according to several criteria. A total of 108 features have been considered based on reviews and comments from users. For the purpose of developing personalised services, we have the ability to rate up to four thousand distinct phones. Both product enhancement and purchasing decisions are critical. We were successful in achieving our goals because we first organised the unstructured data. We used that dataset to extract the sentences that comprise our feature. After that, we made the feature-level information public by evaluating the emotional nature of the words in question and using ratings. Our position is based on the number of features that each phone has, so we can recommend the best feature phones. With this information in hand, we honed our strategy for the so-called "phone" function, which views ratings as ground-level assessments that account for the whole customer. It has been shown that our team's method is commendable. The only result is an MAE of 0.555. That is, the quality is just enough for a half star. Applying our 52.3% accuracy yields the following outcome: The ability to accurately predict integer ratings is essential. On the other hand, we'd be willing to put up with an inaccurate integer rating of one star if it meant we could get an accuracy boost of 93.8 percent. There is no supervision whatsoever in the proposed technique. As it follows that we will endeavour to increase output via the application of a loosely monitored or loosely managed approach, we will need to make use of the information at our disposal in order to resolve the matter pertaining to all 108 of our distinct characteristics together.

References

- [1] V. Raghavan, G. Ver Steeg, A. Galstyan, and A. G. Tartakovsky, "Modeling temporal activity patterns in dynamic social networks," *IEEE Trans. Comput. Social Syst.*, vol. 1, no. 1, pp. 89–107, Mar. 2014.
- [2] A. Farasat, G. Gross, R. Nagi, and A. G. Nikolaev, "Social network analysis with data fusion," *IEEE Trans. Comput. Social Syst.*, vol. 3, no. 2, pp. 88–99, Jun. 2016.
- [3] Y. Zhu, D. Li, R. Yan, W. Wu, and Y. Bi, "Maximizing the influence and profit in social networks," *IEEE Trans. Comput. Social Syst.*, vol. 4, no. 3, pp. 54–64, Sep. 2017.
- [4] X. Yang *et al.*, "Collaborative filtering-based recommendation of online social voting," *IEEE Trans. Comput. Social Syst.*, vol. 4, no. 1, pp. 1–13, Mar. 2017.
- [5] F. S. N. Karan and S. Chakraborty, "Dynamics of a repulsive voter model," *IEEE Trans. Comput. Social Syst.*, vol. 3, no. 1, pp. 13–22, Mar. 2016.
- [6] F. Smarandache, M. Colhon, S. Vlăduțescu, and X. Negrea, "Wordlevel neutrosophic sentiment similarity," *Appl. Soft Comput.*, vol. 80, pp. 167–176, Jul. 2019.
- [7] K. Ravi, V. Ravi, and P. S. R. K. Prasad, "Fuzzy formal concept analysis based opinion mining for CRM in financial services," *Appl. Soft Comput.*, vol. 60, pp. 786–807, Nov. 2017.
- [8] J. Xu *et al.*, "Sentiment analysis of social images via hierarchical deep fusion of content and links," *Appl. Soft Comput.*, vol. 80, pp. 387–399, Jul. 2019.
- [9] F. H. Khan, U. Qamar, and S. Bashir, "SentiMI: Introducing point-wise mutual information with SentiWordNet to improve sentiment polarity detection," *Appl. Soft Comput.*, vol. 39, pp. 140–153, Feb. 2016.

- [10] K. R. Jerripothula, J. Cai, and J. Yuan, "Quality-guided fusion-based co-saliency estimation for image co-segmentation and colocalization," *IEEE Trans. Multimedia*, vol. 20, no. 9, pp. 2466–2477, Sep. 2018.
- [11] A. Wiliam, W. K. Sasmoko, and Y. Indrianti, "Sentiment analysis of social media engagement to purchasing intention," in *Proc. Understand. Digit. Ind. Conf. Manag. Digit. Ind., Technol. Entrepreneurship (CoMDITE)*. Bandung, Indonesia: Routledge, Jul. 2020, p. 362.
- [12] L. G. Singh, A. Anil, and S. R. Singh, "SHE: Sentiment hashtag embedding through multitask learning," *IEEE Trans. Comput. Social Syst.*, vol. 7, no. 2, pp. 417–424, Apr. 2020.
- [13] K.-P. Lin, Y.-W. Chang, C.-Y. Shen, and M.-C. Lin, "Leveraging online word of mouth for personalized app recommendation," *IEEE Trans. Comput. Social Syst.*, vol. 5, no. 4, pp. 1061–1070, Dec. 2018.
- [14] N. Bui, J. Yen, and V. Honavar, "Temporal causality analysis of sentiment change in a cancer survivor network," *IEEE Trans. Comput. Social Syst.*, vol. 3, no. 2, pp. 75–87, Jun. 2016.
- [15] R.-C. Chen and Hendry, "User rating classification via deep belief network learning and sentiment analysis," *IEEE Trans. Comput. Social Syst.*, vol. 6, no. 3, pp. 535–546, Jun. 2019.
- [Online]. Available: <https://ieeexplore.ieee.org/author/37086402578>
- [16] S. Kumar, K. De, and P. P. Roy, "Movie recommendation system using sentiment analysis from microblogging data," *IEEE Trans. Comput. Social Syst.*, early access, May 28, 2020, doi: 10.1109/TCSS.2020.2993585.
- [17] M. Ling, Q. Chen, Q. Sun, and Y. Jia, "Hybrid neural network for Sina Weibo sentiment analysis," *IEEE Trans. Comput. Social Syst.*, early access, Jun. 4, 2020, doi: 10.1109/TCSS.2020.2998092.
- [18] K. Chakraborty, S. Bhattacharyya, and R. Bag, "A survey of sentiment analysis from social media data," *IEEE Trans. Comput. Social Syst.*, vol. 7, no. 2, pp. 450–464, Apr. 2020.
- [19] F.-C. Yang, A. J. Lee, and S.-C. Kuo, "Mining health social media with sentiment analysis," *J. Med. Syst.*, vol. 40, no. 11, p. 236, 2016.
- [20] DVK, MASTHAN RAO KALE, "Utilising CNN and LSTM for Nutritional Deficiencies Identification in Grape Plant Leaves" *African Journal of Biological Sciences* 6 (3), 931-938
- [21] DVK, MASTHAN RAO KALE, "Examining Grape Plant Leaves using Deep Belief Networks and Soft Max Classifier to Identify Nutritional Deficiencies", *Neuroquantology* 20 (19), 5742-5754