

How to Cite:

Ingle, A. (2022). Exploring a collaborative filtering algorithm-based optimization approach for internet marketing strategy. *International Journal of Health Sciences*, 6(S3), 4257–4269. <https://doi.org/10.53730/ijhs.v6nS3.6790>

Exploring a collaborative filtering algorithm-based optimization approach for internet marketing strategy

Dr. Arun Ingle

Director, DKKKP's Maharashtra Institute of Management, Kalamb-Walchandnager, Pune, India

Abstract---In the beyond couple of many presentations and articles has been given throughout the years concerning goods that sell well on the internet. On current hit records and the racks of the largest web-based businesses, books, CDs, recordings, PCs, and home electrical gadgets are among the items found. There are technologies available, as outlined in this paper that can assist in the effective internet marketing of these things and sales can be supported via recommender systems. Market-leading US dotcom companies have used the cooperative filtering method, which is also depicted in this research report, with 17) success. Incredible outcome lately. On account of said showcasing achievement, the presence of European proposal frameworks can be expected very soon. To take care of these issues, we propose an incline one calculation in view of the combination of believed information and client comparability, which can be ought to choose confided in information. Second, we ought to work out the likeness between clients. Third, we really want to add this likeness to the weight component of the better slant one calculation, and afterward, we get the last suggestion condition. We have completed various examinations with the Amazon dataset, and the outcomes demonstrate that our recommender calculation performs more precisely than the customary incline one algorithm.deployed in different recommender frameworks. This calculation includes three systems.

Keywords---online marketing, collaborative filtering algorithm, recommender system.

Introduction

In like manner business action, the hypothetical premise of recommender frameworks has been known for quite a while. "The Chef's Favorites" on a café

menu or a specialist's paper on vehicle wash administrations are two instances of normal recommender frameworks. While meandering the racks at a book shop, experiencing book suggestions on the covers is normal. Celebrities (commentators, experts, columnists, editors, etc) are often approached to audit books in the abstract world, subsequently helping or affecting purchasers in their choice. For quite a while, this recommender procedure has been a generally involved promoting device in the corporate world. Local area suggestion apparatuses, notwithstanding individual and master proposal methodologies, are notable in showcasing practice (Nikolaeva-Sriram 2006). A lot of customer decision information is evaluated, arranged, and afterward put into a simple to-impart structure to disconnect and apply these advances. As a general rule, this is the methodology of top records or notoriety appraisals that might be found in all types of media (films, books, deals and hit records in music). These individual or local area based recommender approaches are normal mass advertising instruments. The theme is the way to assemble customized recommendations in view of individual preferences, and whether recommender approaches can be formed into recommender frameworks.

Movie Lens (www.movielens.org), an individual film recommender, is maybe the most notable of the local area destinations that utilization an online cooperative separating process. It is a genuine illustration of how cooperative separating functions. To utilize this framework, clients should initially finish an internet based enrollment structure in which they should rate films that they have recently found (for this situation a 5-grade scale is utilized). The framework then differentiates the appraisals (for example a client's profile) with those of different clients. Basically, the calculation utilizes film appraisals to track down the nearest neighbor with the most practically identical profile to that of the client. Obviously, a reliable scale is expected for profile examination (Cosley et al 2003), as well as a program that does multivariable factual estimations (relationship math, bunch investigation with gatherings) (Goldberg et al 2001).

Recommender system

Web based advertising frameworks have become more reasonable and inescapable in our regular routines as the web and brilliant gadgets have filled in prevalence. There are a wide range of kinds of items in a comprehensive web based shopping website like Alibaba, Amazon, and others, in this way it very well may be challenging for a client to choose a reasonable thing among all the others, which can influence the client's craving to buy, bringing down exchange deals. Therefore, for the clients and undertakings of a comprehensive web based advertising site, a satisfactory recommender framework will be basic. To work on the presentation of an online business framework, a proposal framework is used, which depends on buying data in most existing frameworks. A recommender framework assembles data from a client and proposes things from among the current items that it figures the client will esteem the most.

Suggestion Systems are programming apparatuses and methodology that are utilized to consequently offer items to clients in light of their inclinations. The ideas introduced endeavored to furnish clients with an assortment of dynamic choices. Information recuperation (IR), AI, choice emotionally supportive networks

(DSS), and text classification are for the most part teaches where proposal frameworks are utilized. These frameworks are utilized to adapt to the issue of data over-burden (IO) by prescribing possibly critical or important things to clients. They've demonstrated to be viable IO handling answers for online clients, and they've filled in prominence as quite possibly the most famous and strong web based advertising stage. Many existing suggestion frameworks utilize the Collaborative Filtering Algorithm (CF), which is generally utilized in web based business. They are profoundly solid and effective methods, as per a few eminent internet showcasing organizations.

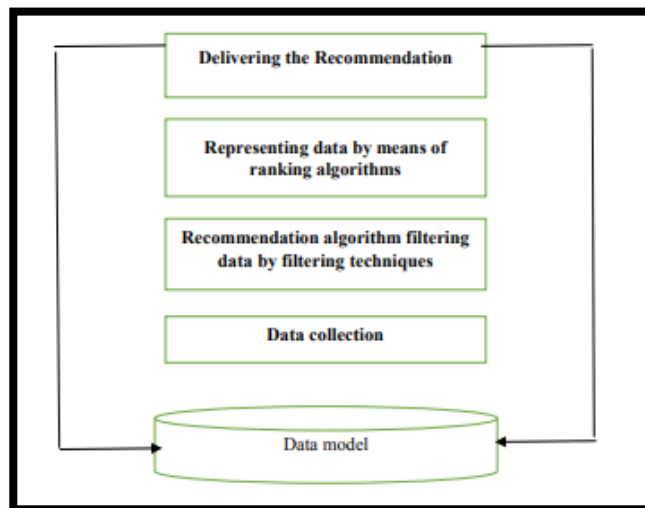


Fig 1. Recommendation framework work

There are six unique sorts of recommender frameworks utilized in the media and diversion business.

Collaborative recommender system

The idea depends on the possibility that things that different clients with comparable interests cherished already are prescribed to the objective clients. The comparability in discernment between at least two clients is determined involving the similitude in the client's earlier scores. All CF calculations have the ability to use earlier client scores to suggest or gauge new things that countless shoppers would appreciate. The veritable hypothesis is profoundly dependent on the idea of likeness between clients or things, with the comparability between earlier inclinations or appraisals communicated as an element of custom. Client-based and thing based calculations are two essential options in contrast to CF calculations.

Sifting by agreement Many Online showcasing proposal frameworks use suggestion calculations, which are a kind of individualized recommender technique. It's a framework in light of three standards: individuals have comparable preferences and considerations, their top picks and considerations are predictable, and their decisions might be reasoned from their authentic top

choices. Subsequently, the cooperative calculation is worked around clients' activities to find direct neighbors for one another and expect their inclinations in light of their neighbors' advantages or inclinations.

Amazon utilized cooperative sifting to make item suggestions to customers [10]. A significant number of these worked on cooperative separating proposal techniques are devoted to building suggestion frameworks, and they can be assembled into two methodologies: client based and thing based. To uncover connections between divergent things, thing based CF investigations the client thing grid first, then, at that point, computes suggestions for clients randomly utilizing these relationships [11]. Cold beginning, information sparsely, and lacking adaptability are on the whole issues with the underlying adaptations of CF. The essential guideline of User-based shared sifting of ideas depends on client similitude, which is estimated and analyzed between target clients and different clients in light of client conduct inclinations. [12]. The framework will actually want to suggest the client item valued by their neighbor once the objective client's neighbor has been recognized. These neighbors are considered regular while looking to suggest objects; this type of tantamount client is ordinarily alluded to as the closest neighbor. [14]. In any case, the standard CF strategy can choose inadequately agent clients as the ongoing client's neighbors, bringing about deficiently right future rules. [13]. In any case, as the requirement for exactness develops, suggestion calculations become more perplexing and challenging to fathom. Accordingly, an effective however easy to-execute technique is required. At the point when a lot of information are involved, even CF calculations experience the ill effects of high computational utilization and grid intricacy, as well as the need to recalculate each time the items are refreshed. Therefore, there is a ton of exploration in view of CF calculations to create and acquire spread suggestion frameworks.

Model-based and memory-based collaborative filtering

The CF Algorithm memory-based and the CF Algorithm model-based cooperative it were examined to channel techniques. Memory-based calculations decide the comparability between two clients by looking at their positioning on an assortment of things. Adaptability and sparsity are the two most normal kinds of issues. Model-based approaches have been introduced as a choice to address these difficulties, however these methodologies seem to restrict purchaser decision.

Model-based CF

This kind of CF can in like manner be utilized to demonstrate huge worth over a memory-based technique to ability, however as of not long ago, a similar degree of accuracy was not accessible. [16,17]. It utilizes an excited learning way to deal with acquire a probabilistic strategy for two assignments: foreseeing or suggesting content, and it pre-computes an information model, for example (client information or thing information), with scores for those things in the proposed structure. AI procedures, like Bayesian organizations, grouping, and rule-based approaches, have generally been used before model-based sifting models. The examination of two every now and again utilized productive calculations, one-

sided network factorization and normal lattice factorization, both utilizing stochastic inclination plummet, was given by khishigsuren davagdorj, kwang ho park, and keun ho ryes (sgd). We ran tests on two true open datasets: book crossing and film focal point 100 k, and looked at the outcomes utilizing two measurements: root mean square blunder (rimes) and mean outright mistake (MAE) (made). In both trial and genuine world datasets, one-sided lattice factorization utilizing the sag method brought about a huge lift in proposal exactness for rating expectation. For the book crossing dataset, one-sided framework factorization decreased the rimes by 25.78 percent and the made by 19.69 percent, while for the film focal point 100 k dataset, the rimes was diminished by 19.69 percent and the mea was decreased by 14.08 percent. While looking at the results of various datasets, one-sided grid factorization with sag produces lower forecast blunder, true to form. [Figure 18a]. Numerous computations are performed, which raises method intricacy as a drawback, however works on gauge precision as an advantage.

Memory-based algorithms

These Approaches are more normal in fictitious works than model-based approaches, and this methodology is expected to authorize a broad memory. It has developed into a notable CF plan. Numerous web based business frameworks, especially Amazon, have coordinated it in an intriguing way. All computations are basically required to be postponed until a gauge or suggestion is required. Moreover, for memory-based calculations, pre-estimation isn't needed, and no disconnected plan is utilized. Therefore, in memory-based CF procedures, all subtleties connecting with the latest exchange records are right away open to deliver an estimate or idea, which can be a significant benefit [20]. Memory-based CF approaches generally utilize factual techniques in light of a background marked by shared appraisals to find K-closest neighbors to online clients or objective items. For this situation, each neighbor is allotted a weight in light of their separation from the internet based buyer, and the calculation then joins the top choices of the closest neighbors to construct a proposal for the objective client. As a matter of fact, as a neighbor-based CF, memory-based CF has been all the more ordinarily connected with showing its whole dependence on the k-closest neighbor approach. A memory-based CF technique is ordinarily used to develop two key NNH approaches: client based closest area and thing based closest area.

Ahmed Zahir, Yuyu Yuan, and Krishna Moniz introduced a paper in which they offered another trust-based system called AgreeRel Trust to conquer the troubles related with forecast precision based certainty extraction techniques. Dissimilar to accuracy based strategies, this strategy doesn't need the assessment of the principal gauge, and the trust relationship is more significant. Aggregate game plans between any two clients and their connected exercises are converged to get the relationship of certainty. They tried our procedure on three public informational indexes and contrasted the forecast exactness with a notable cooperative sifting strategy to perceive how successful it was. The aftereffects of the trials show that our methodology beats different procedures fundamentally. This procedure enjoys the benefit of lessening high processing expenses and

information sparsity, however it has the impediment of just giving restricted express criticism.

User-based neighborhood

Client based area moves toward first figure out who in the objective client evaluations had similar pattern, then, at that point, utilize similar client appraisals to create gauges and afterward make ideas. This strategy for producing the rating for a functioning client's unrated thing midpoints the evaluations of the closest neighbors for that particular thing. To make more valid forecasts, loads are given to neighbor positioning scores in light of their similarity to the objective client. This method is relegated loads to give a more definite forecast of encompassing qualities in light of their similarity to the dynamic client.

Ahmed Shabbier, Md. Mahamudul Hasan, Md. Mahamudul Hasan, Md. Mahamudul Hasan, Md. Armful Islam Mali and Shabbir Ahmed proposed a technique to join specific customary comparability estimations to uncover three classes of associated clients that are basically the same, very unique, and normal comparative. They're additionally making another proportion of likeness that will prove to be useful when a proficient normal of indistinguishable client matches is required. At last, the proposed Recommendation Process is scrutinized through experimentation. Genuine information from Movielens and Epinions are incorporated. Therefore, individuals might reach the resolution that the proposed closeness measure gives the way to a more extensive way to deal with the proposed client based cooperative separating framework and outflanks other standard comparability measurements. [24]. The utilization of likeness calculations in this approach increments framework intricacy (tedious), which is respected a weakness; by the by, these computations have the advantage of upgrading conjecture exactness.

Item-based neighborhood

Client based approaches are changed into thing based closest neighbor calculations, which create expectations in light of thing similitude. A thing based approach is utilized to exploit object closeness. This technique thinks about the Goal Object to an assortment of things surveyed by a client and decides that they are (To choose whether to recommend it to the shopper,). We introduced another information model in light of client assumptions to Amman Jabakji, Hasan Da g, and others to expand the precision of thing related ideas utilizing the Apache Mahout module. They additionally give subtleties of how this model chips away at a dataset acquired from Amazon. Our information show that the proposed system will bring about impressive upgrades in idea adequacy. [26]. with no criticism from clients, the technique might experience a virus start trouble, yet then again, suggestion precision improves, which is a benefit.

Collaborative filtering solutions

Cooperative sifting calculations can be joined with an assortment of IT frameworks There are manual cooperative sifting frameworks in which individuals make or solicitation such suggestions, yet most of business applications are

computerized frameworks that gather, store, and examine client inclination information; they then, at that point, search for clients with comparative preferences and suggest specific items in light of the data acquired. Client inclinations are the main thing that drives cooperative separating frameworks these client inclinations address individual preferences, however they likewise - incidentally - give the informational index expected to track down the nearest neighbors. These are exchange information gathered during buying; notwithstanding standard deals information (like what is bought, when, and for what cost), online deals additionally consider how long clients spend on a page, what they see, print out, save, and even the way that they assess specific things. Cooperative separating approaches can be utilized once appraisals and additionally inclinations for a specific buyer bunch have been gained. Getting back to the past model, the framework can likewise give proposals about regardless of whether a specific film merits seeing, yet there is additionally the choice of getting suggestions from the framework with practically no intercession. Adopting these different strategies into account, John Riedl and Eric Vrooman distinguished three fundamental groupings of recommender frameworks and three separate cooperative sifting philosophies in light of their capacities: pullactive CF, push-dynamic CF, and computerized CF. Since the three frameworks produce various sources of info, they might be used in a similar association or on similar site on occasion.

Pull-active CF

The client is a functioning member during the time spent the framework making suggestions (in light of questions) because of their solicitation utilizing pull-dynamic cooperative sifting applications. This recommender works by finding out about other clients' inclinations or individual interests inside a local area, and afterward looking for others' ideas and comments while searching for an answer for an issue or assignment. Embroidered artwork was the first broadly utilized electronic cooperative separating framework. Embroidered artwork was made as an exploration project by Xerox PARC with the basic role of framing workgroups to assist with picking which articles (for the most part on electronic announcement sheets) merit perusing. Woven artwork clients could submit remarks to articles, and different clients could have the framework search for texts that met specific models, like an article's key terms (utilizing data recovery and sifting), others' comments, for sure others did with a specific article.

Push-active CF

In regular work, messages are often sent with just a contraction of FYI; to other people who may be intrigued. Web clients are more likely than excluded from networking messages, in which jokes are sent to companions and associates who, preferably, share their awareness of what's actually funny. PUSH-ACTIVE CF deals with the idea that clients can undoubtedly advance (push) content to others utilizing a product assuming they think it fascinating or valuable. Lotus Research's David Malts and Kate Ehrlich made the principal CF model of this kind.

Automated CF

The primary contrast between computerized CF and PUSH and PULL dynamic CF is that mechanized CF gets data on client inclinations, looks at and investigations it, and afterward communicates it to clients, though the last option require human interest (Ahn 2006). GroupLens was a precursor in mechanized CF innovation.

Collaborative filtering on the Internet

For researchers, financial analysts, government officials, and any other person who can get to the web, the Internet has opened up new open doors. An ever increasing number of information is amassing in this non-progressive organization. Today, it's best not to express "promoting gathering" into an internet searcher's inquiry field on the grounds that going through each of the outcomes could require over seven days (at the time this article was being composed, Google displayed around 5,040,000 connections). Consequently, while utilizing the Internet nowadays, there are two factors to battle with: the volume and trustworthiness of data. A couple of years prior, an organization's top administration may as yet peruse an everyday leaflet made by a press-checking administration. Attempting to appreciate the information on the Internet these days, whether by catchphrase search or just item data, is a troublesome endeavor. Regardless of whether time and assets aren't an issue, there's generally the topic of the quality or trustworthiness of material got on the web, which is troublesome, in the event that certainly feasible, to make due. For instance, there is an enormous modern organization with a few hundred workers that has had some work posting for a project supervisor on its site for a long time. It's difficult to tell whether they haven't had the option to fill this post yet or whether they just neglected to refresh their site experiencing the same thing.

Cooperative separating calculations could be compelling for deciding a site's subject utilizing watchwords, yet additionally topics, quality, taste, or fields of interest. Perhaps the main benefit of carrying on with work on the Internet is that each visit includes a two-way association between the guest and the site proprietor. Guests constantly give data about themselves while survey an organization's data (Vandermerwe 2000), however not for the more extensive public, but rather for the site's proprietor. At the point when a guest asks requests, places orders, finishes up an enrollment structure or a survey, or keeps in touch with the firm or a conversation bunch about their considerations on an item or a particular issue, data is gathered. Extra information is assembled by doing factual investigations of visits (e.g., breaking down the log record), which essentially supports the development of a coherent construction for the site and the estimation of its ubiquity. To summaries, automated collaborative filtering is based on information gathered from past human-machine interactions. Automated collaborative filtering systems, in their most basic version, maintain note of every item a user scores, including how much he or she likes it. The system then identifies which customers can "predict" other people's preferences based on their own. It will eventually try to suggest new things based on these forecasts.

Research Methodology

Statistical accuracy metrics compare the numerical deviation of the projected ratings from the actual user ratings to determine the accuracy of a prediction system. The prediction rating has a number of accuracy indicators. The concept of accuracy metrics is straightforward: calculate the difference between the predicted and actual rating. The mean absolute error (MAE) is the most common metric (Gong and Ye 2009). In the test dataset, the MAE primarily determines the average absolute error between the prediction rating and the actual rating. The less the MAE, the more accurate the predictions will be, allowing for more accurate suggestions. Assume that the actual rating set is $r_1, r_2, \dots, r_N$, and that the forecast rating set is $p_1, p_2, \dots, p_N$, and that the MAE is:

$$MAE = \frac{\sum_{i=1}^N |p_i - r_i|}{N}$$

The root mean square error (RMSE) is more strict than the MAE since it increases the punishment (square punishment) for inaccurate prediction ratings, making system evaluation more difficult. The lower the RMSE, the more accurate the predictions will be, allowing for more accurate recommendations. The RMSE is calculated as follows:

$$RMSE = \frac{\sqrt{\sum_{i=1}^N (p_i - r_i)^2}}{N}$$

Data analysis

The comparison between the slope one algorithm based on trusted data and the traditional algorithm

The trustworthy ratio of user ratings is represented by r . The following table takes into account trusted data that is bigger than the table's trusted ratio (the last column $r = \text{Null}$, which is the usual algorithm that ignores the trusted ratio). The following are the results of the rating prediction accuracy: First, we'll go through how to choose a reliable ratio. When the trusted ratio is near to 0, the forecast is inaccurate, indicating that our data contains a lot of low-trusted ratio data. As a result, we only consider data with a trusted ratio greater than 0.5, and the outcome is still not very good. As a result, we chose data with a high trusted ratio, such as those with a trusted ratio larger than 0.8. Without taking into consideration the trusted ratio, the result shows that the prediction accuracy is greater than the data. Of course, the best forecast accuracy is when the trusted ratio is 1.

Table 2 The result of MAE comparison between the algorithm based on trusted data and the traditional algorithm

r	>0	>0.5	>0.8	=1	Null
MAE	1.389	1.214	0.776	0.659	0.967

Table 2 demonstrates that the higher the trusted ratio, the lower the MAE, which is supported by the assumption that trusted data are correct. At the same time, the traditional algorithm's MAE is 0.967 without taking into account the trusted data, and it can be shown that the prediction accuracy is higher than the traditional algorithm only when the trusted ratio is more than 0.8. The main reason, according to the survey research, is that people don't care for this kind of helpful clicking behavior, and there is an increase in internet fraudulent users, so entirely reliable data is unusual. Most crucially, when the trusted ratio is 1, the rating forecasting accuracy might improve precision by 31.9 percent, which is a significant boost over the previous approach.

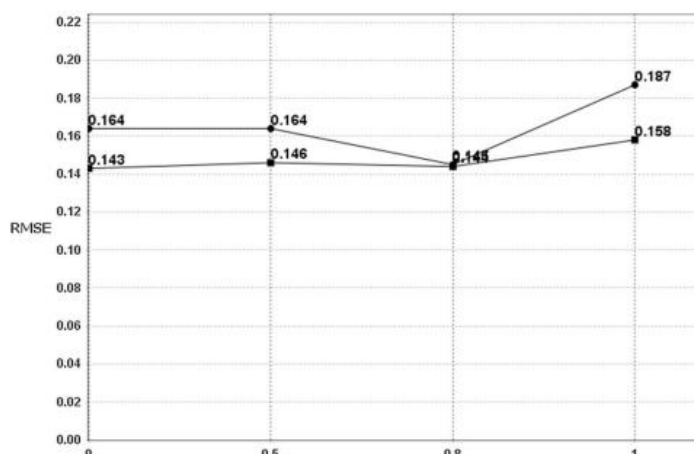
The Comparison between the slope one algorithm based on the fusion of trusted data and similarity and the algorithm based on trusted data using RMSE

The RMSE in Table 4 is based on trusted data, whereas the surmise is based on a combination of trusted data and similarity. When utilizing RMSE as an indicator, Table 4 demonstrates the forecast accuracy of the three types of algorithms. The dataset is different when the trusted ratio (r) is different. The size of the dataset shrinks as r grows, especially when r is close to 1. As a result, the RMSE does not properly reflect the declining trend as r grows. The slope one algorithm based on the fusion of trusted data and similarity, however, is clearly better than the slope one method based on the trusted data, as shown in Fig. 3. That is, the RMSE of the slope one algorithm based on the fusion of trusted data and similarity is smaller than the RMSE of the slope one method based on the trusted data when the trusted ratio is the same. In a nutshell, these tests reveal that the slope one algorithm, which is based on the merger of trusted data and similarity, is the most effective (Fig. 4).

Table 4 The result of RMSE comparison between the slope one algorithm based on the fusion of trusted data and similarity and the algorithm based on the trusted data

r	>0	>0.5	>0.8	=1	Null
srrRMSE	0.143	0.146	0.144	0.158	0.049
RMSE	0.164	0.164	0.145	0.187	0.057

Fig. 4 The result of the RMSE comparison between the slope one algorithm based on the fusion of trusted data and similarity and the algorithm based on the trusted data



Result and Discussion

Furthermore, when discussing the recommendation system, we will be confronted with a massive amount of data (Wu et al. 2016). It is impossible to provide correct recommendations, at least not precise enough, if we do not have adequate data as input. To put it another way, more data leads to a stronger recommendation effect. When we have access to a significant amount of user data, we must apply the same recommendation algorithm to a larger dataset. With such a large dataset, computing recommendation results will take a long time. Users will be too anxious to wait for our recommendations if we spend too much time on them, which is a nightmare for recommendation applications. When dealing with large amounts of data, cloud computing (Voorsluys et al. 2011; Li et al. 2018; Gao et al. 2018; Tian et al. 2018) is a frequent approach. It uses a large number of computers to do the actual computing function in parallel. The entire computing job is broken into several jobs that can be completed on thousands of machines at the same time using this strategy. As you can expect, this type of computing will drastically reduce the amount of time it takes to get good recommendation results. As a result, if we can combine cloud computing with the recommendation algorithm, we may be able to get a leg up on the competition in terms of computing speed. Scalability is a major issue with collaborative filtering. The cost of computation for CF will be quite high when the dataset volume is very large.

Conclusion

To alleviate the problem of information overload, which has become a potential concern for many Internet users, it is required to rank, filter, and effectively distribute relevant information on the Internet, where the amount of possibilities is overwhelming. RS overcomes this obstacle by analyzing a large number of dynamically generated data and providing clients with personalized information and services. Using web marketing recommender tactics will only help you sell more of your favorite books, CDs, movies, and home goods. This marketing plan might work for any web-based business that offers wine, chocolate, or clothing. Click-and-mortar businesses can utilize recommender techniques to develop

community ideas (top lists like Wine of the Week) based on purchase data, and after analyzing user profiles, they can create personalized offers for their customers (using the method of closest neighbors). We run our experiment on a subset of Amazon's item rating dataset and assess it in four ways. To begin, we compared the slope one algorithm, which is based on reliable data, to the classic approach. Second, we looked into the differences between the slope one algorithm, which is based on the fusion of trusted data and similarity, and the MAE-based algorithm, which is based on trusted data. Finally, the slope one algorithm based on the fusion of trusted data and similarity and the algorithm based on trusted data using RMSE are compared. Finally, we compared our slope one technique based on user similarity to other algorithms using different dataset sizes. The experimental results reveal that the slope one algorithm based on the fusion of trusted data and user similarity outperforms the classic slope one algorithm in terms of prediction accuracy. The prediction accuracy will improve greatly if we can improve people's subjective voting behaviour and identify fraudulent internet users, and we will be able to give more accurate recommendation services for users. Furthermore, we may offer more comprehensive recommendation services based on other data kinds (McCauley et al. 2015). On the other hand, alternative methods used in recommendation systems, such as semi supervised feature analysis, may be considered (Chang and Yang 2017).

References

1. Park, D.H.; Kim, H.K.; Choi, I.Y.; Kim, J.K. A literature review and classification of recommender systems research. *Expert Syst. Appl.* 2012, 39, 10059–10072. [CrossRef]
2. Yang, W.S.; Cheng, H.C.; Dia, J.B. A location-aware recommender system for mobile shopping environments. *Expert Syst. Appl.* 2008, 34, 437–445. [CrossRef]
3. Aimeur, E.; Brassard, G.; Fernandez, J.; Onana, F. ALAMBIC: A privacy-preserving recommender system for electronic commerce. *Int. J. Inf. Sec.* 2008, 7, 307–334. [CrossRef]
4. Baltrunas, L.; Ludwig, B.; Peer, S.; Ricci, F. Context relevance assessment and exploitation in mobile recommender systems. *Pers. Ubiquitous Comput. PUC 2012*, 16, 1–20. [CrossRef]
5. Kompan, M.; Bielíková, M. Content-Based News Recommendation. In *E-Commerce and Web Technologies*; Buccafurri, F., Semeraro, G., Eds.; Lecture Notes in Business Information Processing; Springer: Berlin/Heidelberg, Germany, 2010; pp. 61–72. [CrossRef]
6. Yang, Z.; Wu, B.; Zheng, K.; Wang, X.; Lei, L. A Survey of Collaborative Filtering-Based Recommender Systems for Mobile Internet Applications. *IEEE Access* 2016, 4, 3273–3287. [CrossRef]
7. Woerndl, W.; Schueller, C.; Wojtech, R. A Hybrid Recommender System for Context-aware Recommendations of Mobile Applications. In *Proceedings of the 2007 IEEE 23rd International Conference on Data Engineering Workshop*, Istanbul, Turkey, 17–20 April 2007; pp. 871–878. [CrossRef]
8. Chen, C.C.; Wan, Y.H.; Chung, M.C.; Sun, Y.C. An effective recommendation method for cold start new users using trust and distrust networks. *Inf. Sci.* 2013, 224, 19–36. [CrossRef]

9. Nikolakopoulos, A.N.; Kouneli, M.A.; Garofalakis, J.D. Hierarchical Itemspace Rank: Exploiting hierarchy to alleviate sparsity in ranking-based recommendation. *Neurocomputing* 2015, 163, 126–136. [CrossRef]
10. Gunawardana, A.; Shane, G. A Survey of Accuracy Evaluation Metrics of Recommendation Tasks. *J. Mach. Learn. Res.* 2009, 10, 2935–2962.
11. Beal, J.; Langer, S.; Genzmehr, M.; Gap, B.; Breiting, C.; Nürnberger, A. Research paper recommender system evaluation: A quantitative literature survey. In *Proceedings of the International Workshop on Reproducibility and Replication in Recommender Systems Evaluation, RepSys 2013*, Hong Kong, China, 12 October 2013; pp. 15–22. [CrossRef]
12. Ström, R.; Vendel, M.; Bredican, J. Mobile marketing: A literature review on its value for consumers and retailers. *J. Retail. Consum. Serv.* 2014, 21, 1001–1012. [CrossRef]
13. Mulcahy, R.F.; Riedel, A.S. ‘Touch it, swipe it, shake it’: Does the emergence of haptic touch in mobile retailing advertising improve its effectiveness? *J. Retail. Consum. Serv.* 2018, in press. [CrossRef]
14. Li, Y.M.; Yeh, Y.S. Increasing trust in mobile commerce through design aesthetics. *Comput. Hum. Behav.* 2010, 26, 673–684. [CrossRef]
15. Cyr, D.; Head, M.; Ivanov, A. Design aesthetics leading to m-loyalty in mobile commerce. *Inf. Manag.* 2006, 43, 950–963. [CrossRef]