

Finding Domain-based Expert for Improving Collaborative Filtering Algorithm

Guo Wu¹, Yujiu Yang² and Wenhuan Liu³

¹ Key Laboratory of Broadband Network and Multimedia, Graduate School at Shenzhen, Tsinghua University
Shenzhen, Guangdong Province, China
wuguoissac@gmail.com

² Key Laboratory of Broadband Network and Multimedia, Graduate School at Shenzhen, Tsinghua University
Shenzhen, Guangdong Province, China
yang.yujiu@sz.tsinghua.edu.cn

³ Key Laboratory of Broadband Network and Multimedia, Graduate School at Shenzhen, Tsinghua University
Shenzhen, Guangdong Province, China
liuwh@sz.tsinghua.edu.cn

Abstract

Traditional neighborhood-based collaborative filtering algorithms are widely used in recommender system field for its accuracy, interpretability and operability. In this paper, we introduce expert user model into collaborative filtering and determine authoritative expert users via expert finding methods in large corpus. We propose a method to produce predications for target user. Instead of the similarity between normal users and target user, we determine target user's neighborhood based on the similarity between expert users and target user. Experiments on Amazon datasets show that our method has better performance than neighborhood-based collaborative filtering on recommendation accuracy, novelty and calculation efficiency.

Keywords: Domain, Expert, Recommender System, Collaborative Filtering.

1. Introduction

As an active information filtering method, recommender system is widely applied in different fields. A recommender system is a system that models users' interests and preferences based on their historical records. Hence, the system could recommend users goods or social elements which are previously unknown. The core of a recommender system is recommender algorithm. The leading algorithms nowadays are content-based recommender algorithm, collaborative filtering algorithm, mixed recommender algorithm and social recommender algorithm. Among the algorithms, neighbor-based collaborative filtering (NCF) algorithm is the most commonly used one. NCF algorithm chooses k nearest neighbors of the target user and makes recommendations for the target user according to the weighted average of the k neighbors. NCF algorithm achieves high

recommendation accuracy and interpretability. However, the algorithm faces problems like result convergence, data sparsity and low efficiency as well.

The fundamental idea of NCF is that similar users have same preferences. However, in some cases, even similar users have different tastes and recommend the same product is not good for exploring new interests of the target user. In practice, people like to listen to the opinions from experts (e.g., when choosing which university to enter, we listen to the opinions from experts in education rather than from parents or relatives). We assume there exist a user group called domain expert users in the recommender system. The domain expert users have professional knowledge in specific fields, know most products in the recommender system, sensitive to unpopular products and are more trusted comparing to ordinary users. Then based on common sense, a user would prefer to choose products which are chosen by the expert who is most similar to the user.

Based on the idea mentioned above, we propose a domain-based expert selection algorithm as well as a method for retrieving expert users. We design and conduct a series of experiments, showing that comparing to traditional algorithms our algorithm has a better performance in accuracy, diversity and novelty.

The left sections are organized as follows. 2) We review related works in recommender system field. 3) We give the definition of domain-based expert users and propose a method to retrieve expert users. We improve the NCF algorithm based on the expert users. 4) We conduct experiments using Amazon's data and show the superiority of our algorithm. 5) We conclude our result and discuss future research interests.

2. Related Works

User-based collaborative filtering (UCF) algorithm searches users who are most similar to the target user (nearest neighbors) and make prediction on the target user based on the weighted average score of the nearest neighbors [1]. Linden et al. [2] propose an item-based collaborative filtering (ICF) algorithm. The ICF algorithm searches similar items instead of similar users for recommendation. In practice, the data scale of items is much larger than that of the users. Moreover, item data has a relatively low frequency of data change is low and higher interpretability. Therefore, the ICF algorithm is widely used in the industry. However, the efficiency and operability of the two similarity-based algorithms mentioned above is far from satisfaction, due to the large scale of data searching. Meanwhile, the efficiency of the two algorithms is sensitive to noise data and sparsity.

For dealing with such problems, John O'Donovan et al. [3] propose an algorithm that makes recommendations based on the credibility rather than similarity between users. J. Cho et al. [4] propose expert-user circle with taking both credibility-based and similarity-based user circles into consideration. They make recommendations for users on online education and digital comic sites through similarity-based user circles and expert-user circles. Here, expert is defined as a user who has a large number of activity records in a certain field so that he/she could make appropriate recommendations for other users. Furthermore, X. Amatriain et al. [5] propose an expert-opinion-based collaborative filtering algorithm, where an expert is defined as an individual who can make appropriate recommendation. They select and gather film critics as expert set from a film review website in America called Rotten Tomatoes. Then, they make recommendations for local users according to the experts' opinions. Afterwards, many researchers use expert-based or mixed similarity-expert-based collaborative filtering to deal with the problems faced by traditional collaborative filtering algorithms. Maria Dima et al. [6] give a new definition for expert, which is based on their professionalism and contribution. Professionalism is measured by user's knowledge towards product characteristics. The expert-based recommendation solves the cold boot problem for new users and new products. Weiliang Kong et al. [7] define expert according to user's activeness and influence in a certain field. They design a clustering-based expert selecting method. Qiang Liu et al. [8] obtain abstract star users via large scale of user data training. The method improves the accuracy and efficiency of the recommender system. Julian McAuley et al. [9] point out that hot products may not be the best recommendations for users. Each user starts as amateur

and finally becomes an expert. It is a dynamic process. At last, most users would choose the products that are best for themselves. They model the dynamic process to gain better results by introducing the potential factor recommender system and investigating the professionalism of users.

The works listed above all try to improve the performance of traditional algorithms through changing the definition of expert users. Expert Finding (EF) problem is a hot topic in information retrieval field as well. Given a specific task and a set of experts, one needs to find an appropriate expert to finish the task. Early expert finding problems are solved by building databases of employees' knowledge and skills [10]. However, building a database requires a lot of effort and the costs money and time. Considering the disadvantages, Balog et al. [11] propose a statistic language model based expert finding algorithm which statistically models the candidate text set and the candidate expert set. They calculate the correlation between the candidate expert and the task and define the correlation as the professionalism of the candidate. Petkova et al. [12] compare several commonly used language models in expert finding and conclude that the difference between these models is mainly on the independence assumptions of the models. Statistic-language-model-based expert finding has a high requirement on data set. In practice, it is not easy to get access to such high quality data. Li et al. [13] introduce explicit semantic analysis (ESA) method into traditional language models. They overcome the disadvantage that searching keywords must be included in the candidate text by calculating the semantic distance between the keyword and the text. So far, expert finding method is very popular in enterprise searching field, but has not yet been applied in recommender system field. We improve the recommendation results by defining expert user in recommender system with the method of expert finding.

3. Domain-based Expert Model

3.1 Definition of Domain-based Expert User

In practice, professional knowledge is domain-based as well as the experts. A good example is that we listen to university professors' opinions rather infant teachers' when choosing a university to study in. We divide experts into different levels based on the breadth of their knowledge: profession experts, who know almost everything in a certain profession, lack of depth of knowledge, e.g., experts on classical literature; domain experts, who has deep knowledge in a certain field of a

profession, balanced in breadth and depth of knowledge, e.g., experts on Chinese fictions in Ming and Qing Dynasties; specialist, who has a very deep understanding of a certain problem, but lack of breadth of knowledge, e.g., experts on A Dream of Red Mansions. Profession experts lack detailed information of a certain product. Hence, they are not suitable for providing personalized recommendations. In contrast, specialists do not know much about the whole system, and therefore their opinions are limited. The knowledge of domain experts is wide and deep enough to make personalized recommendations for users. Thus, domain experts are the experts who we are looking for in this paper.

We define domain-based expert users as follow. A domain-based expert user e is an advanced user who has certain professional knowledge in a certain field d in recommender system R . Domain-based expert users are selected from the recommender system users. Thus, they are a part of the recommender system. We find out the domain-based expert user e from the entire user set U via the level of expertise function $\text{domain_expertise}(e, d)$. The general mathematical formula is

$$\forall u \in U, e = \arg \max_d \text{domain_expertise}(u, d) \quad (1)$$

3.2 Finding Domain-based Expert

Professional knowledge is the understanding level of products in a certain field, which has various measurements. In the field of information retrieval, level of expertise is measured by the relativity between a candidate expert and keywords in a certain searching field through searching candidate texts. Such searching method is called the Expert Finding (EF) method. In this paper, we apply the EF method in the dataset of a recommender system for searching domain-based expert users.

Definition of expert finding: expert finding is to find expert users for every domain in a given set of candidate expert users $X = [x_1, \dots, x_2]$, with corresponding text sets DOC and domain sets of the expert users. Given the searching keyword Q , calculating the relativity between a candidate user and a specific domain is mathematically equivalent to calculating the probability that the candidate belongs to the specific domain, denoted by $P(x_i | Q)$.

According to the Bayesian Criterion, we have

$$P(x_i | Q) = \frac{P(Q | x_i) P(x_i)}{P(Q)} \propto P(Q | x_i) P(x_i) \quad (2)$$

where $P(x_i)$ is the general expertise of user x_i . It is independent from the searching keywords and therefore can be offset by setting an initial value. The problem of solving $P(x_i | Q)$ reduces to solving $P(Q | x_i)$. We solve the

problem using language models. There are two popular methods of expert finding: user introduction centered method, direct professional knowledge modeling based on connecting texts and candidate experts; text centered method, first figure out the texts that related to searching Q , then find out the related expert through the texts. The mathematical formulas of the two methods are as follows.

$$P(Q | x_i) = \prod_{t_i \in Q} P(t_i | x_i) = \prod_{t_i \in Q} \left[\sum_{doc_j \in \text{DOC}} P(t_i | doc_j) P(doc_j | x_i) \right] \quad (3)$$

$$P(Q | x_i) = \sum_{doc_j \in \text{DOC}} P(Q | doc_j) P(doc_j | x_i) \quad (4)$$

$$= \sum_{doc_j \in \text{DOC}} \left[\prod_{t_i \in Q} P(t_i | doc_j) \right] P(doc_j | x_i)$$

where t_i is the i th word of searching keyword Q , $P(t_i | doc_j)$ is the probability of t_i given the word distribution of text and $P(doc_j | x_i)$ denotes the relativity between text doc_j and candidate expert x_i . The value of $P(doc_j | x_i)$ can be 0 or 1, depending on whether x_i is the author of doc_j .

To the expert finding technique in information retrieval field, $P(x_i | Q)$ represents the relativity between the target user and searching keyword Q . We need to know the expert-text set for using this method. Whereas in a recommender system with comment information, we can choose $P(x_i | Q)$ as the indicator of the level of expertise of target user e in a certain domain d (that is, $\text{domain_expertise}(e, d)$) by setting the elements in the domain and the user comments as the searching keyword and expert-text set, respectively, and assuming that the relativity between the target user and searching reflects the user's understanding level toward the searching.

3.3 Collaborative Filtering based on Domain-based Expert

The traditional neighbor-based collaborative filtering algorithm uses the weighted average score of the neighbor users as the score prediction for the target user (Fig.1 a). The formula is as follow.

$$r_{ui} = \frac{\sum_{v \in U} r_{vi} \text{sim}(v, u)}{\sum_{v \in U} \text{sim}(v, u)} \quad (5)$$

where r_{ui} denotes the product i 's score given by the target user u and $\text{sim}(v, u)$ represents the similarity between user u and v . There are many measurements of the

similarity. In this paper, we choose the cosine similarity as the measurement.

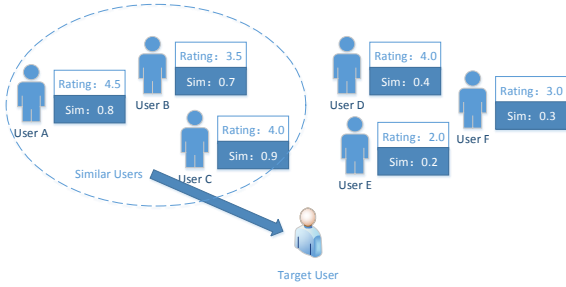


Fig.1 a Neighbor-CF

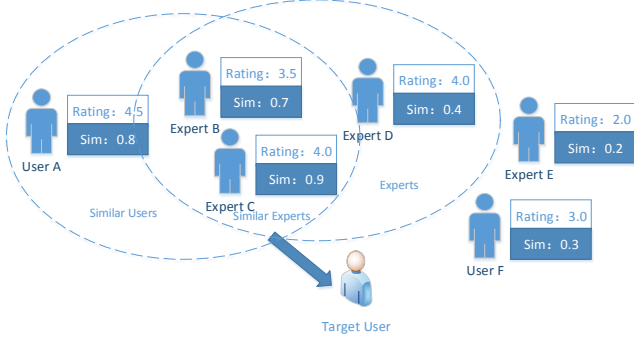


Fig.1 b %Expert-CF

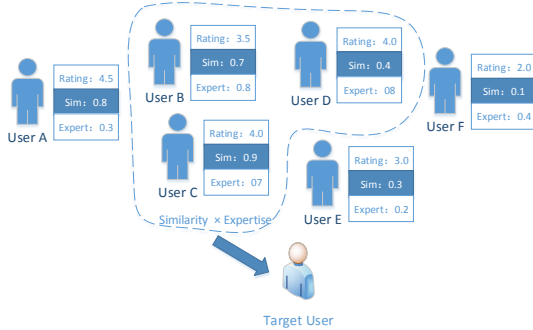


Fig.1 c Expertise-CF

After obtaining the user's level of expertise in the recommender system through expert finding, we can introduce the level of expertise into the traditional collaborative filtering model with the following two methods.

- 1) Domain-expert-user-based collaborative filtering (%expert-CF): set a threshold level value m . Select the top $m\%$ users in the level of expertise as expert users to generate the expert-user set E . When predicting the score of the target user, select similar users from the expert-user set instead of the entire system to calculate the weighted average score (Fig.1 b) as follow.

$$r_{ui} = \sigma_u + \frac{\sum_{j=1}^n \sum_{e \in E_d^j} (r_{ei} - \sigma_e) \text{sim}(e, u)}{\sum_{e \in E} \text{sim}(e, u)} \quad (6)$$

where σ_u is the average score of all users in the recommender system and $e \in E_d^j$ denotes user e is an expert user in domain j .

- 2) Level-of-expertise-based collaborative filtering (expertise-CF): calculate the score using users' level of expertise directly without generating expert-user set (Fig.1 c) as follow.

$$r_{ui} = \frac{\sum_{j=1}^n \sum_{v \in V_d^j} r_{vi} \text{sim}(v, u) \delta_{vd}}{\sum_{v \in U} \text{sim}(v, u) \delta_{vd}} \quad (7)$$

where v_d^j denotes user j 's score record in domain d and denotes the level of expertise of user v in domain d .

With the two methods, we can predict the score of the target user using user-rating matrices and user profiles.

4. Experiment

4.1 Dataset Description

For most of the current online system applications such as e-commerce sites and movie rating sites, users of the system can comment on the merchandise. So the recommendation systems not only has the user scoring matrix but also contains a wealth of textual information. The Amazon website comment dataset¹ we adopt here is like that. The entire dataset contains a total of 35 million product reviews in 18 years from 1995 to 2013. These reviews contain user information, product information, ratings and review text with titles. It is a five-star rating system, in which five star stands for 'very good' and a star stands for 'very bad' and the smallest unit is a half star. According to the types of products this dataset is divided into many categories. In this article we have chosen movie products, electronics products and automotive supplies for experiments. The details are shown in Table 1:

Table 1: Dataset Statistics

	Movies	Electronic	Automotive	TOTAL
#users	1318175	884175	133526	2335876
#items	235042	96643	47577	379262
#reviews	856772	1371574	188728	10128029
average words	146.98	108.30	72.45	

¹ <http://snap.stanford.edu/data>

From the above table it can be seen that the data size of three categories of is very large. The data size of movie products is the largest followed by electronics products and the data size of automotive supplies is the smallest. The average words of each type is quite different. The average words of movie products is the biggest followed by electronics products and automotive supplies. The reviews with more words are more valuable and can be helpful for us to get the user's preferences and professional level.

4.2 Domain-based Expert Group Analysis

According to section 3.2, we first use users and review information from the three datasets to build a user-document set and then generate the corresponding query keywords Q according to the different domains of the dataset. By retrieval in the user-document set using the language model, we can get the correlation between the users and the domain topics and then obtain the users' professional level.

Table 2: Query Keywords for Finding Experts

datasets	domains
Movies	Action Adventure Animation Comedy Crime Documentary Drama Family Fantasy Horror Music Musical Mystery Romance Sci-Fi Short Thriller War Western
Electronic	Phone Camera DV USB TV DVD Walkman iPhone iPad Audivox Cyper Airconditoner Microwave Radio Razor Xbox Van
Automotive	Pinstriping Tape Automotive Enthusiast Merchandise Interior Trim Products Paint, Body & Trim Products Trim Automotive Enthusiast Vehicle Accessories Automotive Enthusiast Garage & Shop Automotive Enthusiast Apparel Automotive Enthusiast Bags & Accessories Automotive Enthusiast Collectibles

To implement the expert-CF algorithm, we set the threshold m as 5 and the top 5% user as an expert user collection. The words of the common users in the three datasets are 72.45 146.98 and 108.30, respectively. In contrast, the words of the according experts are 185.31, 119.23 and 91.57 respectively. The reviews of experts are obviously longer than the common users which shows expert users in these areas have a more in-depth understanding. In addition, the expert users and ordinary users also have different distributions of ratings. The

distribution for Amazon movies data set is shown in Figure 2. We can see that the common users' average ratings distribution is more wide, while the expert users' average ratings distribution is more concentrated, indicating that expert users have the consistency.

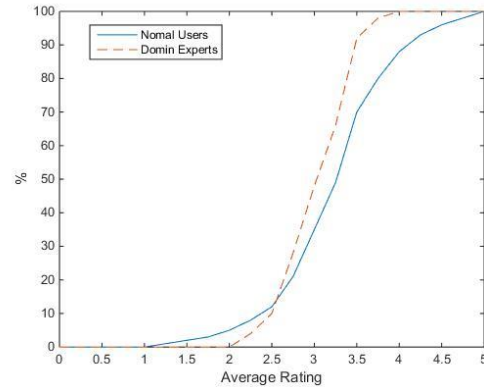


Fig. 2 Distribution of the Average Rating

4.3 Results

We chose the following three methods to predict the recommended score and used Neighbor-CF as a baseline to make comparison:

- Neighbor-CF: traditional collaborative filtering algorithm which use the neighboring users' weighted score as a predictor score.
- Expert-CF: described in Formula 2 in section 3.2, calculating the expert user dataset and calculating the similarity between the target user and expert users.
- Expertise-CF: described in Formula 3 in section 3.2, directly using the professional level bas weights to predict the target score

In the experiment, each data set is divided into two parts: 80% of the part is the training set, and the remaining 20% is the test set. In evaluation, we consider results from two aspects: the conventional accuracy indicator; the user experience indicators such as diversity, coverage, novelty and other indicators.

Accuracy index is the most common indicator for recommendation evaluation and can be divided into two types according to different application scenarios. In the scenario of predicting scores, we choose the mean absolute error MAE value as an index, the smaller the value, the better; In the scenario of Top-N recommendation, we select top-10 accuracy (10@Precision) and top-10 recall rate (10@Recall) as indicators, the higher the value, the better. The results of accuracy is shown in Table 3.

Table 3: Accuracy Performance of Our Approach

	Movies			Electronic			Automotive		
	MAE	10@P	10@R	MAE	10@P	10@R	MAE	10@P	10@R
Neighbor-CF	0.716	0.075	0.156	0.675	0.062	0.137	0.585	0.083	0.161
Expert-CF	0.682	0.113	0.163	0.647	0.104	0.162	0.533	0.125	0.188
Expertise-CF	0.701	0.101	0.184	0.655	0.126	0.174	0.541	0.121	0.152

By the table, we can see that: 1) the MAE indicator of our two methods are better than traditional collaborative filtering algorithm, which demonstrates that our approach has a more accurate prediction score to predict. Among all the methods, the Expert-CF has the most outstanding performance. 2) Our two approaches is also superior to the traditional collaborative filtering method on the precision and recall indicators. 3) Recommended result in Automotive data sets and Electronic data sets are better than Movies datasets, not only because the user's subjectivity in these two data sets is more weaker than the movie merchandise trade, but also for the user's score more discrete. 4) The improvement brought by our method is significantly greater in Automotive data set and Electronic data set. Because the two types of commodities require a deeper knowledge, the recommended method based on domain experts get a better performance. 5) The introduction of the collection of patent expert and professional level to these two methods will lead to different results, but the difference is not noticeable.

The indicator of precision is usually used to demonstrate the improvement of the whole system, however, it does not ensure the better experience for the user. Because the only pursuit of precision will lead to "right but meaningless recommendation" such as commodities that are very popular but have been acquainted well by the user. So we consider other indicators which involved diversity, coverage[14], novelty[15] to verify our method. Diversity means we should spare no effort to make the list of recommendation diversified and make recommendation about related domains commodities. Coverage means the ability for the system to discover low-frequency items, and the measurement for the ratio of the quantity of recommend items of the system to the overall quantity of the items. Novelty means the ability for the system to recommend unheard commodities. The result of the evaluation of these indicators are shown in Fig.3. As we can see from the figure, our method exceeds the traditional collaborative filtering on coverage and diversity, especially a huge improvement on novelty. It demonstrate our method tend to provide user with more various, more novelty items, not just recommend the popular commodities. So we can greatly improve the experience for the user.

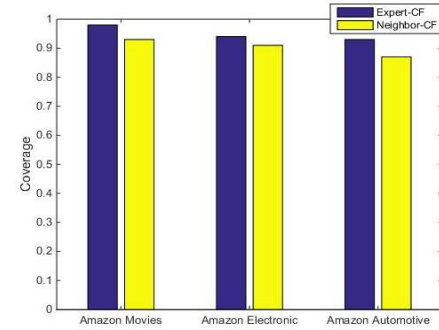


Fig. 3 a Coverage on Movies, Electronic and Automotive

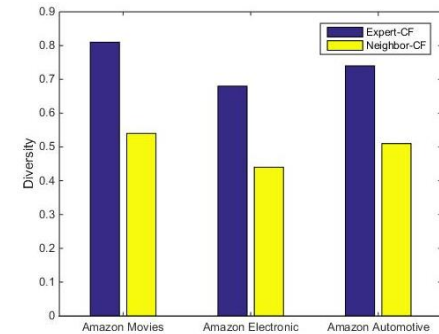


Fig. 3 b Diversity on Movies, Electronic and Automotive

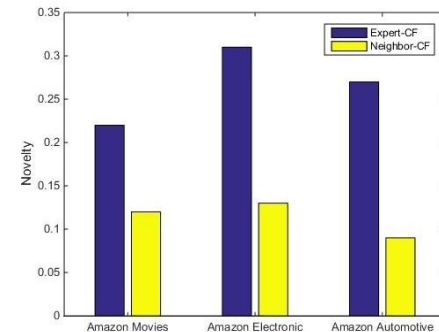


Fig. 3 c Novelty on Movies, Electronic and Automotive

When compared with traditional collaborative filtering, our method not only perform well in the indicators mentioned above, but also greatly improved the efficiency and the scale of computation because we have adopted much smaller collection than the ordinary collection. We will discuss the influence of the scale of expert collection for the result of recommendation. The Fig.4 shows the value of MAE in the algorithm Expert-CF varies greatly, when different choices are made on the ratio of the user to be treated as expert. We can reach a conclusion that as the

collection of user expert become larger, the value of MAE will increase and the recommendation will be more accurate. The value will reach the minimum when the ratio is 11%, and no conspicuous improvement even though the ratio increase. When the ratio is about 2.5%, the value of MAE of our method has surpassed the optimal result 0.716 of the traditional collaborative filtering. From all of the mentioned experiment results above, we can infer two conclusions. First, even though in the pursuit of the optimal result, the scale of the collection patent user is only 11% of the scale of the original user---the scales has been largely decreased. Second, if we are not so sensitive of the optimal result, we can also adopt smaller scale of collection of user patent which will be much more efficient in computation.

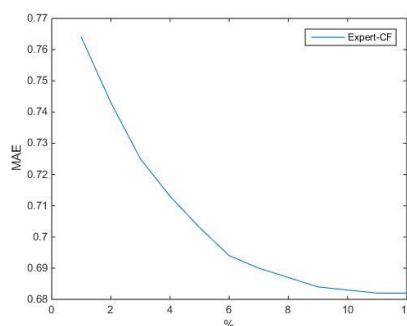


Fig. 4 the Effect of Expert Group Size on the Results

5. Conclusions

To solve the problems of data sparsity, low efficiency and result convergence faced by traditional collaborative filtering algorithms, we innovatively propose a domain-expert-based model to improve the performance of collaborative filtering algorithms. A domain expert user is a core user who has professional knowledge in a certain domain in the recommender system and therefore can make recommendations for ordinary users. We design a method to find expert users in a certain domain by making use of the expert-finding techniques in the information retrieval field as well as the rich comment-text data in the recommender system. Moreover, we propose two methods to improve the performance of the collaborative filtering model with introducing the measurements of user expertise levels. We conduct experiments on three big data sets, showing that our methods achieve a significant improvement not only in the algorithm accuracy, but also the diversity, coverage and novelty. The result indicates that the recommendations from our new methods are diversified and innovative. In addition, we improve the computing efficiency by reducing data sparsity with using small scale expert-user data.

For future research, we may consider dynamic domain expert models. We investigate how the level of expertise dynamically changes over time to further improve our recommendation results.

References

- [1] Breese, John S., David Heckerman, and Carl Kadie. "Empirical analysis of predictive algorithms for collaborative filtering." Proceedings of the Fourteenth conference on Uncertainty in artificial intelligence. Morgan Kaufmann Publishers Inc., 1998.
- [2] G. Linden, B. Smith, and J. York, "Amazon.com Recommendations: Item-to-item Collaborative Filtering", IEEE Internet computing, vol. 7, no. 1, pp. 76–80, 2003.
- [3] O'Donovan, John, and Barry Smyth. "Trust in recommender systems." Proceedings of the 10th international conference on Intelligent user interfaces. ACM, 2005.
- [4] Cho, Jinhyung, Kwiseok Kwon, and Yongtae Park. "Collaborative filtering using dual information sources." Intelligent Systems, IEEE 22.3 (2007): 30-38.
- [5] X. Amatriain, N. Lathia, J. Pujol, H. Kwak, and N. Oliver, "The Wisdom of the Few," in Proceedings of the 2009 ACM SIGIR Conference on Research and Development in Information Retrieval. ACM, 2009, pp. 532–539.
- [6] Dima, Maria, et al. "Expert based prediction of user preferences." Semantic Media Adaptation and Personalization (SMAP), 2010 5th International Workshop on. IEEE, 2010.
- [7] Kong, Weiliang, et al. "Collaborative Filtering Algorithm Incorporated with Clusterbased Expert Selection." National Engineering Research Center for E-learning, Central China Normal University (2012).
- [8] Liu, Qiang, Bingfei Cheng, and Congfu Xu. "Collaborative Filtering Based on Star Users." Tools with Artificial Intelligence (ICTAI), 2011 23rd IEEE International Conference on. IEEE, 2011.
- [9] McAuley, Julian John, and Jure Leskovec. "From amateurs to connoisseurs: modeling the evolution of user expertise through online reviews." Proceedings of the 22nd international conference on World Wide Web. International World Wide Web Conferences Steering Committee, 2013.
- [10] Maron, M. E., Sean Curry, and Paul Thompson. "An inductive search system: Theory, design, and implementation." Systems, Man and Cybernetics, IEEE Transactions on 16.1 (1986): 21-28.
- [11] Balog, Krisztian, Leif Azzopardi, and Maarten De Rijke. "Formal models for expert finding in enterprise corpora." Proceedings of the 29th annual international ACM SIGIR conference on Research and development in information retrieval. ACM, 2006.
- [12] Petkova, D., and W. B. Croft. "Generalizing the language modeling framework for named entity retrieval." Proc. of SIGIR. Vol. 7. 2007.
- [13] Li, Xiu, et al. "A Service Mode of Expert Finding in Social Network." Service Sciences (ICSS), 2013 International Conference on. IEEE, 2013.

- [14]Ge, Mouzhi, Carla Delgado-Battenfeld, and Dietmar Jannach. "Beyond accuracy: evaluating recommender systems by coverage and serendipity." Proceedings of the fourth ACM conference on Recommender systems. ACM, 2010.
- [15]Celma, Òscar, and Perfecto Herrera. "A new approach to evaluating novel recommendations." Proceedings of the 2008 ACM conference on Recommender systems. ACM, 2008.