

Learning Personalized Preference of Strong and Weak Ties for Social Recommendation

Xin Wang^{†‡}

Steven C.H. Hoi[§]

Martin Ester[‡]

Jiajun Bu[†]

Chun Chen[†]

[†]Zhejiang Provincial Key Laboratory of Service Robot, College of Computer Science, Zhejiang University, China

[‡]School of Computing Science, Simon Fraser University, Burnaby, B.C., Canada

[§]School of Information Systems, Singapore Management University, Singapore

[†]{xinwang,bjj,chenc}@zju.edu.cn [‡]{xwa49,ester}@cs.sfu.ca [§]chhoi@smu.edu.sg

ABSTRACT

Recent years have seen a surge of research on social recommendation techniques for improving recommender systems due to the growing influence of social networks to our daily life. The intuition of social recommendation is that users tend to show affinities with items favored by their social ties due to social influence. Despite the extensive studies, no existing work has attempted to distinguish and learn the personalized preferences between strong and weak ties, two important terms widely used in social sciences, for each individual in social recommendation. In this paper, we first highlight the importance of different types of ties in social relations originated from social sciences, and then propose a novel social recommendation method based on a new Probabilistic Matrix Factorization model that incorporates the distinction of strong and weak ties for improving recommendation performance. The proposed method is capable of simultaneously classifying different types of social ties in a social network w.r.t. optimal recommendation accuracy, and learning a personalized tie type preference for each user in addition to other parameters. We conduct extensive experiments on four real-world datasets by comparing our method with state-of-the-art approaches, and find encouraging results that validate the efficacy of the proposed method in exploiting the personalized preferences of strong and weak ties for social recommendation.

Keywords

Social Recommendation; Personalization; Strong and Weak Ties; User Behavior Modeling

1. INTRODUCTION

Recommender systems have saturated into our daily life — we experience recommendations when we see “More Items to Consider” or “Inspired by Your Shopping Trends” on Amazon and “People You May Know” on Facebook (i.e., friend recommendation [45]) — other popular online web services such as eBay, Netflix and LinkedIn etc. also provide users with the recommendation

features. Thus algorithmic recommendation [25, 37] has become a necessary mechanism for many online web services which recommend items such as music, movies or books to users. These online web services normally make recommendations based on collaborative filtering which suggests items favored by similar users. Representative collaborative filtering algorithms include low-rank matrix factorization. However, most recommender systems suffer from the *data sparsity* problem, where the number of items consumed by a user (e.g., giving a rating) is often very small compared to the total number of items (usually hundreds of thousands to millions or even billions in web-scale applications).

The *data sparsity* issue can significantly affect the performance of model-based collaborative filtering methods such as low-rank matrix factorization mainly because of two reasons: the “overfitting” problem where insufficient data is available for training models, and the “cold start” problem in which recommender systems fail to make recommendations for new users when there is no historical behavior data to be collected. To resolve the data sparsity challenge, one promising direction is resorting to *social recommendation* where the data sparsity is tackled by utilizing the rapidly growing social network information in recommender systems [44, 14, 15, 26, 29, 28, 43, 46, 39].

On the other hand, despite quite a lot of literature studies attempting to explore tie strength prediction in demographic data [34] and social media [33, 8, 40, 3, 7, 32, 2, 23, 16, 41], all but one of the existing social recommendation methods fail to distinguish different types of social ties for pairs of connected users. In social sciences, Granovetter [10] introduces different types of social ties (strong, weak, and absent), and concludes that weak ties are actually the most important reason for new information or innovations to spread over social networks. Based on Granovetter’s statement, the model proposed by Wang et al. [39] is the only one among those existing social recommendation approaches that pays attention to the important distinctions between strong and weak ties. Nevertheless, Wang et al. simply assume every individual has the same preference for strong and weak ties — either everyone prefers strong ties to weak ties or everyone prefers weak ties to strong ties. In practice, different users may have different preferences for strong and weak ties, e.g., one may trust strong ties more than weak ties and others may behave opposite. Thus Wang’s model suffers from the limitation that no personalized preferences of strong and weak ties can be learned. As such, although Wang’s model addresses the concern that lacking the distinctions for different social ties may significantly limit the potential of social recommendation, we argue that ignoring the personalized tie type preference for each individual tends to result in sub-optimal solutions as well.

©2017 International World Wide Web Conference Committee (IW3C2), published under Creative Commons CC BY 4.0 License.
WWW 2017, April 3–7, 2017, Perth, Australia.
ACM 978-1-4503-4913-0/17/04.
DOI: <http://dx.doi.org/10.1145/3038912.3052556>



Therefore, inspired by the claims in social sciences and the promising results in Wang’s work [39], we investigate whether distinguishing and learning the personalized tie type preference for each individual would improve the prediction accuracy of social recommendation. However, there exist several challenges for the combination of personalized tie type preferences and social recommendation. First, how to effectively identify each type of social tie (“strong” or “weak”) in a given social network? Sociologists [10, 9] typically assume the *dyadic hypothesis*: the strength of a tie is determined solely by the interpersonal relationship between two individuals, irrespective of the rest of the network. For example, Granovetter uses the frequency of interactions to classify strong and weak ties [9], that is, if two persons meet each other at least once a week, then their tie is deemed strong; if the frequency is more than once a year but less than once a week, then the tie is weak. This is simple and intuitive, but requires user activity data which is not publicly available in modern online social networks because of security and privacy concerns¹. Second, assuming there is a reliable method for differentiating between strong and weak ties, how can we efficaciously combine it with existing social recommendation approaches such as *Social Matrix Factorization* (SMF) [15] to improve the accuracy? Third, different people may have different preferences for strong and weak ties, and thus how do we learn a personalized tie type preference for each of them?

To handle these challenges, we first adopt Jaccard’s coefficient [13] to compute the social tie strength [24, 31]. Naturally, Jaccard’s coefficient captures the extent to which those users’ friendship circles overlap, making itself a feature *intrinsic* to the network topology, and requiring no additional data to compute. Our choice is supported by the studies on a large-scale mobile call graph by Onnela et al. [31], which show that (i) tie strength is partially determined by the network structure relatively local to the tie and (ii) the stronger the tie between two users, the more their friends overlap. We define ties as strong if their Jaccard’s coefficient is above some threshold, and weak otherwise. We would like to point out that the optimal threshold (w.r.t. recommendation accuracy) will be learnt from the data. Furthermore, we exclude absent ties in our model because they do not play an important role as indicated in Granovetter’s work. We distinguish strong and weak ties by thresholding Jaccard’s coefficient between two users, while Granovetter thresholds the number of interactions between two users.

We then propose the *Personalized Social Tie Preference Matrix Factorization* (PTPMF) method, a novel probabilistic matrix factorization based model that simultaneously (i) classifies strong and weak ties w.r.t. optimal recommendation accuracy and (ii) learns a personalized preference between strong and weak ties for each user in addition to other parameters. More precisely, we employ gradient descent to learn the best (w.r.t. recommendation accuracy) threshold of tie strength (above which a tie is strong; otherwise weak) and the personalized tie type preference for each user as well as other parameters such as the latent feature vectors for users and items.

This work makes the following three contributions:

- We recognize the importance of strong and weak ties in social relations as motivated by the sociology literature, and incorporate the notion of strong and weak ties into probabilistic matrix factorization for social recommendation.
- We present a novel algorithm to simultaneously learn user-specific preferences for strong and weak ties, the optimal (w.r.t. recommendation accuracy) threshold for classifying strong and weak ties, as well as other model parameters.

¹https://en.wikipedia.org/wiki/Privacy_concerns_with_social_networking_services

- We conduct extensive experiments on four real-world public datasets and show that our proposed method significantly outperforms the existing methods in various evaluation metrics such as RMSE, MAE etc.

The remainder of this paper is organized as follows: we review related work in Section 2. Section 3 discusses the effects of strong and weak social ties that are evident in the sociology literature, and proposes to incorporate these notions into social recommendation. Section 4 gives a detailed formation of our proposed *Personalized Social Tie Preference Matrix Factorization* (PTPMF) model, followed by a description of model inferences for PTPMF in Section 5. Section 6 presents our experiments, compares our approach with baseline recommendation methods and comments on their performances for both all users and cold-start users in terms of various evaluation metrics. Finally, we conclude our work and point out some potential future work for further investigation in Section 7.

2. RELATED WORK

In this section, we review three major categories of related work in recommender systems and social ties studies.

Collaborative Filtering. When it comes to recommender systems, collaborative filtering is one of the most popular algorithmic solutions so far, which makes recommendations based on users’ past behaviors such as ratings, clicks, purchases and favorites etc. Further, low rank matrix factorization is among the most effective methods for collaborative filtering, and there is a large body of work on using matrix factorization for collaborative filtering [30, 36, 21, 12, 20, 38]. As a general treatment, Koren [22] gives a systematic introduction to the application of matrix factorization to recommender systems. Among the literature of matrix factorization, Salakhutdinov and Mnih [30] propose a probabilistic version of matrix factorization (PMF) which assumes a Gaussian distribution on the initializations of latent feature vectors, making the model more robust towards the problem of overfitting and linearly scalable with the number of observations at the same time. However, these matrix factorization based models still suffer from the data sparsity and cold start problems, which gives rise to *social recommendation*.

Social Recommendation. The fact that cold start problem has always been an important factor to deteriorate the performance of collaborative filtering motivates the advent of work on *social recommendation*, which utilizes social information among users to improve the performances of recommender systems. Indeed, social influence tends to have strong effects in changing human behaviours [19, 4], such as adopting new opinions, technologies, and products. This has stimulated the study of *social recommendation*, which aims to leverage social network information to help mitigate the “cold-start” problem in collaborative filtering [43, 44, 46, 42, 15, 28, 29, 26, 27, 14, 39], in the hope that the resulting recommendations will have better quality and higher relevance to users who have given little feedback to the system. In particular, Ma et al. [28] propose a probabilistic matrix factorization model which factorizes user-item rating matrix and user-user linkage matrix simultaneously. They later present another probabilistic matrix factorization model which aggregates a user’s own rating and her friends’ ratings to predict the target user’s final rating on an item. In [15], Jamali and Ester introduce a novel probabilistic matrix factorization model based on the assumption that users’ latent feature vectors are dependent on their social ties’. Wang et al. [39] are the first to try integrating the concepts of strong and weak ties into social recommendation through presenting a more fine-grained categorization of user-item feedback for Bayesian Personalized Ranking (BPR) [35]

by leveraging the knowledge of tie strength and tie types. However, they assume a global rather than personalized preference between strong and weak ties. In other words, their proposed model assumes either all individuals prefer strong ties to weak ties or all individuals prefer weak ties to strong ties, which ignores the fact that different people may have different preferences for strong and weak ties (i.e., some prefer strong ties over weak ties while some others prefer weak ties over strong ties). Our proposed method addresses the limitation in Wang et al.'s work by learning a personalized tie type preference for each individual. In general, the model introduced in [39] conceptually becomes a special case of our proposed method when we assume everyone has the same preference for strong and weak ties.

Social Ties in Social Media. Different types of social ties have attracted lots of interests from researchers in social sciences [9, 10, 5, 18], followed by some recent work which pays attention to tie strength in demographic data [34] and social media [33, 8, 40, 3, 7, 32, 47, 2, 16, 41]. In particular, Gilbert et al. [8] bridge the gap between social theory and social practice through predicting interpersonal tie strength with social media and conducting user-study based experiments over 2000 social media ties. Wu et al. [40] propose a regression analysis to discover two different types of closeness (i.e., professional and personal) for employees in an IBM enterprise social network. Panovich et al. [32] later carry out an investigation related to different roles of tie strength in question and answer online networks by taking advantage of Wu's approach.

In summary, no work so far brings the learning of personalized tie type preference to social recommendation. This is no surprise, since the combination is very specific.

3. STRONG AND WEAK TIES

Speaking of interpersonal ties, Granovetter may probably be the first one who comes into our mind. Granovetter, in his book *Getting a job: A study of contacts and careers* [9], conducts a survey among 282 professional, technical, and managerial workers in Newton, Massachusetts and reports that personal contact is the predominant method of finding out about jobs. The result of his survey shows that nearly 56% of his respondents used personal contacts to find a job while 18.8% used formal means and 18.8% used direct applications instead. Besides, Granovetter's research also demonstrates that most respondents prefer the use of personal contacts to other means and that using personal contacts can lead to a higher level of job satisfaction and income. Thus it will be interesting to explore the important role social influence plays in people's decision making process which does not necessarily need to be limited to an employee's decision about changing a job.

Social influence takes effect through a social network which consists of people and interpersonal ties connecting these people in the network. Granovetter, in his other work [10], introduces different types of interpersonal ties (e.g. strong tie, weak tie and absent tie) and concludes that weak ties are the most important source for new information or innovations to reach distant parts of the network. Again, different ties between the job changer and the contact person who provided the necessary information are analyzed and the strength and importance of weak ties in occupational mobility are shown in [9]. In the late 1960's and early 1970's when the Internet had not come into existence, tie strength was measured in terms of how often they saw the contact person during the period of the job transition, using the following measurement:

- **Often:** at least once a week
- **Occasionally:** more than once a year but less than twice a week

- **Rarely:** once a year or less

In the age of information, social media and online social networks are playing crucial roles in the establishment of social networks. We are able to know new friends and form new relationships/ties through the Internet without necessarily meeting them face to face. Just as Kavanaugh et al. [18] state, the appearance of the Internet has helped to strengthen weak ties and increase their numbers across social groups. Though the importance of weak ties has been exposed to us by sociologists, it is not wise to ignore the roles strong ties play in our lives because strong ties should intuitively be more trustworthy than weak ties. On the other side, different individuals may have different relative degree of trust for their strong and weak ties — one may trust his/her strong ties (or weak ties) more than one another. Thus an interesting and challenging question is that how to learn these user-specific (and perhaps different) preferences for different types of ties. This being the case, considering both strong and weak ties in social recommendation, then optimally distinguishing them w.r.t recommendation accuracy and finally learning a user-specific personalized tie type preference become three key parts of an appropriate solution to improve social recommendation.

In this section we will present how the notion of strong/weak ties and the thresholding strategy are incorporated into social recommendation. We leave the remaining two parts to section 4 for more concrete descriptions. In order that the distinction between strong and weak ties can be incorporated into social recommendation, we will need to be able to define and compute tie strength, and then classify ties. Several potential options seem to serve as adequate candidates. First, as mentioned in Section 1, sociologists use dyadic measures such as frequency of interactions [9]. However, this method is not generally applicable due to lack of necessary data. An alternative approach relies on community detection. Specifically, it first runs a community detection algorithm to partition the network $\mathcal{G} = (U, E)$ into several subgraphs. Then, for each edge $(u, v) \in \mathcal{E}$, if u and v belong to the same subgraph, then it is classified as a strong tie; otherwise a weak tie. However, a key issue is that although numerous community detection algorithms exist [6], they tend to produce (very) different clusterings, and it is unclear how to decide which one to use. Furthermore, if a "bad" partitioning (w.r.t. prediction accuracy) is produced and given to the recommender system as input, it would be very difficult for the recommender system to recover. In other words, the quality of recommendation would depend on an *exogenous* community detection algorithm that the recommender system *has no control over*. Hence, this approach is undesirable.

In light of the above, we resort to node-similarity metrics that measure neighborhood overlap of two nodes in the network. The study of Onnela et al. [31] provides empirical confirmation of this intuition: they find that (i) tie strength is in part determined by the local network structure and (ii) the stronger the tie between two users, the more their friends overlap. In addition, unlike frequency of interactions, node-similarity metrics are intrinsic to the network, requiring no additional data to compute. Also, unlike the community detection based approach, we still get to choose a tie classification method that best serves the interest of the recommender system.

More specifically, we use Jaccard's coefficient [13], a simple measure that effectively captures neighborhood overlap. Let $\text{strength}(u, v)$ denote the tie strength for any $(u, v) \in \mathcal{E}$. We have:

$$\text{strength}(u, v) =_{\text{def}} \frac{|\mathcal{N}_u \cap \mathcal{N}_v|}{|\mathcal{N}_u \cup \mathcal{N}_v|} \quad (\text{Jaccard}), \quad (1)$$

where $\mathcal{N}_u \subseteq \mathcal{U}$ (resp. $\mathcal{N}_v \subseteq \mathcal{U}$) denotes the set of ties of u (resp. v). If $\mathcal{N}_u = \mathcal{N}_v = \emptyset$ (i.e., both u and v are singleton nodes), then simply define $\text{strength}(u, v) = 0$. By definition, all strengths as defined in Equation (1) fall into the interval $[0, 1]$. This definition has natural probabilistic interpretations: Given two arbitrary users u and v , their Jaccard's coefficient is equal to the probability that a randomly chosen tie of u (resp. v) is also a tie of v (resp. u) [24].

Thresholding. To distinguish between strong and weak ties, we adopt a simple *thresholding* method. For a given social network graph $\mathcal{G}$, let $\theta_{\mathcal{G}} \in [0, 1)$ denote the threshold of tie strength such that

$$(u, v) \text{ is } \begin{cases} \text{strong,} & \text{if } \text{strength}(u, v) > \theta_{\mathcal{G}}; \\ \text{weak,} & \text{if } \text{strength}(u, v) \leq \theta_{\mathcal{G}}. \end{cases} \quad (2)$$

Let $\mathcal{W}_u =_{\text{def}} \{v \in \mathcal{U} : (u, v) \in \mathcal{E} \wedge \text{strength}(u, v) \leq \theta_{\mathcal{G}}\}$ denote the set of all weak ties of u . Similarly, $\mathcal{S}_u =_{\text{def}} \{v \in \mathcal{U} : (u, v) \in \mathcal{E} \wedge \text{strength}(u, v) > \theta_{\mathcal{G}}\}$ denotes the set of all strong ties of u . Clearly, $\mathcal{W}_u \cap \mathcal{S}_u = \emptyset$ and $\mathcal{W}_u \cup \mathcal{S}_u = \mathcal{N}_u$.

The value of $\theta_{\mathcal{G}}$ in our proposed approach is *not* hardwired, but rather is left for our model to learn (Section 4), such that the resulting classification of strong and weak ties in $\mathcal{G}$, together with other learned parameters of the model, leads to the best accuracy of recommendations. We conclude this section by pointing out that Granovetter and we both threshold strong and weak ties, we utilize Jaccard's coefficient (degree of connectivity between users) to do the thresholding while Granovetter resorts to the number of interactions between users instead.

4. PERSONALIZED TIE PREFERENCE MATRIX FACTORIZATION FOR SOCIAL RECOMMENDATION

In this section, we present the proposed new model of *Personalized Tie Preference Matrix Factorization* (PTPMF) for social recommendation in detail. Before introducing PTPMF, we will first briefly explain some background knowledge of the classical *Probabilistic Matrix Factorization* (PMF) and of another popular social recommendation model known as *Social Matrix Factorization* (SMF).

4.1 Probabilistic Matrix Factorization

In recommender systems, we are given a set of users $\mathcal{U}$ and a set of items $\mathcal{I}$, as well as a $|\mathcal{U}| \times |\mathcal{I}|$ rating matrix R whose non-empty (observed) entries R_{ui} represent the feedbacks (e.g., ratings, clicks etc.) of user $u \in \mathcal{U}$ for item $i \in \mathcal{I}$. When it comes to social recommendation, another $|\mathcal{U}| \times |\mathcal{U}|$ social tie matrix T whose non-empty entries T_{uv} denote $u \in \mathcal{U}$ and $v \in \mathcal{U}$ are ties, may also be necessary. The task is to predict the missing values in R , i.e., given a user $v \in \mathcal{U}$ and an item $j \in \mathcal{I}$ for which R_{vj} is unknown, we predict the rating of v for j using observed values in R and T (if available).

A matrix factorization model assumes the rating matrix R can be approximated by a multiplication of d -rank factors,

$$R \approx U^T V, \quad (3)$$

where $U \in \mathbb{R}^{d \times |\mathcal{U}|}$ and $V \in \mathbb{R}^{d \times |\mathcal{I}|}$. Normally d is far less than both $|\mathcal{U}|$ and $|\mathcal{I}|$. Thus given a user u and an item i , the rating R_{ui} of u for i can be approximated by the dot product of user latent feature vector U_u and item latent feature V_i ,

$$R_{ui} \approx U_u^T V_i, \quad (4)$$

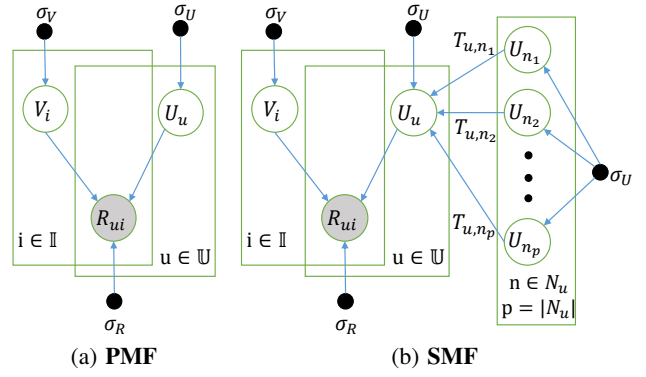


Figure 1: Graphical models of PMF and SMF

where $U_u \in \mathbb{R}^{d \times 1}$ is the u th column of U and $V_i \in \mathbb{R}^{d \times 1}$ is the i th column of V . For ease of notation, we let $|\mathcal{U}| = N$ and $|\mathcal{I}| = M$ in the remaining of the paper.

Later, the probabilistic version of matrix factorization, i.e., *Probabilistic Matrix Factorization* (PMF), is introduced in [30], based on the assumption that the rating R_{ui} follows a normal distribution whose mean is some function of $U_u^T V_i$. The conditional probability of the observed ratings is:

$$p(R|U, V, \sigma_R^2) = \prod_{u=1}^N \prod_{i=1}^M \left[\mathcal{N}(R_{ui} | g(U_u^T V_i), \sigma_R^2) \right]^{I_{ui}^R}, \quad (5)$$

where $\mathcal{N}(x|\mu, \sigma^2)$ is the normal distribution with mean μ and variance σ^2 . If u has rated i , then the indicator function I_{ui}^R equals to 1, otherwise equals to 0. $g(\cdot)$ is the sigmoid function, i.e., $g(x) = \frac{1}{1+e^{-x}}$, which bounds the range of $U_u^T V_i$ within $[0, 1]$. Moreover, U_u and V_i are both subject to a zero mean normal distribution. Thus the conditional probabilities of user and item latent feature vectors are:

$$\begin{aligned} p(U|\sigma_U^2) &= \prod_{u=1}^N \mathcal{N}(U_u | 0, \sigma_U^2 \mathbf{I}) \\ p(V|\sigma_V^2) &= \prod_{i=1}^M \mathcal{N}(V_i | 0, \sigma_V^2 \mathbf{I}), \end{aligned} \quad (6)$$

where $\mathbf{I}$ is the identity matrix. Therefore, the posterior probability of the latent variables U and V can be calculated through a Bayesian inference,

$$\begin{aligned} &p(U, V | R, \sigma_R^2, \sigma_U^2, \sigma_V^2) \\ &\propto p(R|U, V, \sigma_R^2) p(U|\sigma_U^2) p(V|\sigma_V^2) \\ &= \prod_{u=1}^N \prod_{i=1}^M \left[\mathcal{N}(R_{ui} | g(U_u^T V_i), \sigma_R^2) \right]^{I_{ui}^R} \\ &\quad \times \prod_{u=1}^N \mathcal{N}(U_u | 0, \sigma_U^2 \mathbf{I}) \times \prod_{i=1}^M \mathcal{N}(V_i | 0, \sigma_V^2 \mathbf{I}). \end{aligned} \quad (7)$$

The graphical model of PMF is demonstrated in Figure 1(a) and readers may refer to [30] for more details.

4.2 Social Matrix Factorization

There has been some work on social recommendation, among which Jamali and Ester [15] present a well-known social recommendation model called *Social Matrix Factorization* (SMF) that

incorporates trust propagation into probabilistic matrix factorization, assuming that the rating behaviour of a user u will be affected by his social ties N_u through social influence. In SMF, the latent feature vector of user u depends on the latent feature vectors of u 's social ties n , i.e., $n \in N_u$. As is shown by the graphical model of SMF in Figure 1(b),

$$U_u = \frac{\sum_{n \in N_u} T_{un} U_n}{|N_u|},$$

where U_u is u 's latent feature vector and N_u is the set of social ties of user u . T_{un} is either 1 or 0, indicating u and n are "ties" or "not ties".

The posterior probability of user and item latent feature vectors in SMF, given the observed ratings and social ties as well as the hyperparameters, is shown in (8).

$$\begin{aligned} & p(U, V | R, T, \sigma_R^2, \sigma_T^2, \sigma_U^2, \sigma_V^2) \\ & \propto p(R | U, V, \sigma_R^2) p(U | T, \sigma_T^2) p(V | \sigma_V^2) \\ & = \prod_{u=1}^N \prod_{i=1}^M \left[\mathcal{N}(R_{ui} | g(U_u^T V_i), \sigma_R^2) \right]^{I_{ui}^R} \\ & \times \prod_{u=1}^N \mathcal{N}(U_u | \sum_{k \in N_u} T_{uk} U_k, \sigma_U^2 \mathbf{I}) \\ & \times \prod_{u=1}^N \mathcal{N}(U_u | 0, \sigma_U^2 \mathbf{I}) \times \prod_{i=1}^M \mathcal{N}(V_i | 0, \sigma_V^2 \mathbf{I}). \end{aligned} \quad (8)$$

The main idea in (8) and Figure 1(b) is that the latent feature vectors of users should be similar to the latent feature vectors of their social ties. We refer readers to [15] for more details.

4.3 The PTPMF Model

We divide social ties into two groups: strong ties and weak ties. People usually tend to share more common intrinsic properties with their strong ties while they are more likely to be exposed to new information through their weak ties. Both strong ties and weak ties are important in terms of social influence while they play different roles in affecting people. For an individual user, strong ties tend to be more similar to her, on the other hand, weak ties may provide her with more valuable information which can not be obtained from strong ties. Based on this assumption, we propose our approach, PTPMF, to utilize the different roles of strong and weak ties when making recommendations. Besides, by introducing two additional parameters, θ_G and B_u , PTPMF is capable of learning the optimal (w.r.t. recommendation accuracy) threshold for classifying strong and weak ties, user-specific preferences between strong and weak ties as well as other parameters at the same time.

Figure 2 presents the graphical model of PTPMF. We introduce a random variable θ_G for the threshold classifying strong and weak ties. S_u and W_u are the sets of strong and weak ties of user u respectively, classified according to (2). Due to different roles of strong and weak ties in affecting users' rating behaviors, we introduce two new random variables, U_u^s and U_u^w , as strong-tie and weak-tie latent feature vectors for each user u . The strong-tie (resp. weak-tie) latent feature vector of u is dependent on the latent feature vectors of all u 's strong ties (resp. weak-ties). This influence is modeled as follows:

$$U_u^s = \frac{\sum_{s \in S_u} T_{us} U_s}{\sum_{s \in S_u} T_{us}} \quad \text{and} \quad U_u^w = \frac{\sum_{w \in W_u} T_{uw} U_w}{\sum_{w \in W_u} T_{uw}},$$

where $T_{uv} = \text{strength}(u, v)$ is the tie strength between u and v defined in (1), different from SMF in which T is a Boolean vari-

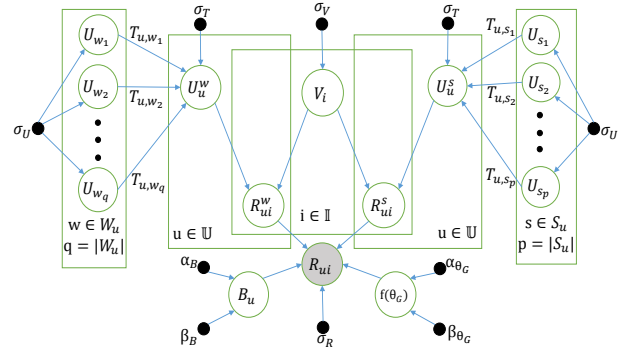


Figure 2: Graphical model of the proposed PTPMF

able. We normalize the tie strength of u and her social ties so that $\sum_{s \in S_u} T_{us} = 1$ and $\sum_{w \in W_u} T_{uw} = 1$. Now the conditional probability of weak-tie and strong-tie latent feature vectors, U_u^w and U_u^s , becomes:

$$\begin{aligned} & p(U^w, U^s | T, U, \sigma_T^2) \\ & = \prod_{u=1}^N \mathcal{N}(U_u^w | \sum_{k \in W_u} T_{uk} U_k, \sigma_T^2 \mathbf{I}) \\ & \times \prod_{u=1}^N \mathcal{N}(U_u^s | \sum_{k \in S_u} T_{uk} U_k, \sigma_T^2 \mathbf{I}). \end{aligned} \quad (9)$$

The dot product of U_u^w (resp. U_u^s) and item latent feature vector V_i then determines u 's weak-tie generated rating on item i (resp. u 's strong-tie generated rating on item i), denoted by R_{ui}^w (resp. R_{ui}^s). Different from SMF, PTPMF further enables the learning of a personalized preference between strong and weak ties for each user through introducing another new variable, B_u , as the probability that u prefers strong ties instead. To generate u 's final rating for item i , PTPMF puts more emphasis on her weak-tie generated rating R_{ui}^w with probability B_u , and on her strong-tie generated rating R_{ui}^s with probability $1 - B_u$ (more details to be discussed below). Thus the conditional probability of the observed ratings can be expressed as:

$$\begin{aligned} & p(R | U^w, U^s, V, B, \theta_G, T, \sigma_R^2) \\ & = \prod_{u=1}^N \prod_{i=1}^M \left[\mathcal{N} \left(R_{ui} | g \left(B_u \left[f(\theta_G) U_u^{wT} V_i + (1 - f(\theta_G)) U_u^{sT} V_i \right] \right. \right. \right. \\ & \quad \left. \left. \left. + (1 - B_u) \left[(1 - f(\theta_G)) U_u^{wT} V_i + f(\theta_G) U_u^{sT} V_i \right] \right), \sigma_R^2 \right) \right]^{I_{ui}^R}, \end{aligned} \quad (10)$$

where $g(\cdot)$ is the sigmoid function, i.e., $g(x) = \frac{1}{1+e^{-x}}$, and $f(\theta_G) = g((t_s - \theta_G)(\theta_G - t_w)) \geq 0.5$, given t_s, t_w as the average tie strength of strong ties and weak ties respectively. The underlying intuition is that when a threshold θ_G gives a small degree of separation, t_s and t_w will be close to θ_G , $f(\theta_G)$ will then be close to 0.5, indicating very few distinctions between strong and weak ties. Similarly, a larger degree of separation results in more distinctions between strong and weak ties in our model. When u prefers weak ties, more weight (i.e., $f(\theta_G) \geq 0.5$) will be given to her weak-tie generated rating (i.e., $U_u^{wT} V_i$), less weight (i.e., $1 - f(\theta_G) \leq 0.5$) will be given to her strong-tie generated rating (i.e., $U_u^{sT} V_i$) and vice versa. Moreover, how much weight to give is dependent upon how well the current threshold, θ_G , classifies strong and weak ties –

a larger degree of separation given by θ_G will result in more weight being given to the preferred tie type.

We assume θ_G and B follow a Beta distribution so that both of them lie in $[0, 1]$. Also, U and V follow the same zero mean normal distribution in (6). Through a Bayesian inference, the posterior probability of all model parameters, given the observed ratings and social ties as well as the hyperparameters, is shown in (11).

$$\begin{aligned}
& p(U^w, U^s, U, V, B, \theta_G | R, T, \sigma_R^2, \sigma_T^2, \sigma_U^2, \sigma_V^2) \\
& \propto p(R | U^w, U^s, V, B, \theta_G, T, \sigma_R^2) p(U^w, U^s | T, U, \sigma_T^2) \\
& \quad p(U | \sigma_U^2) p(V | \sigma_V^2) p(\theta_G | \alpha_{\theta_G}, \beta_{\theta_G}) p(B | \alpha_B, \beta_B) \\
& = \prod_{u=1}^N \prod_{i=1}^M \\
& \quad \left[\mathcal{N} \left(R_{ui} | g \left(B_u \left[f(\theta_G) U_u^{wT} V_i + (1 - f(\theta_G)) U_u^{sT} V_i \right] \right. \right. \right. \\
& \quad \left. \left. \left. + (1 - B_u) \left[(1 - f(\theta_G)) U_u^{wT} V_i + f(\theta_G) U_u^{sT} V_i \right] \right), \sigma_R^2 \right) \right]^{I_{ui}^R} \\
& \quad \times \prod_{u=1}^N \mathcal{N} \left(U_u^w | \sum_{k \in W_u} T_{uk} U_k, \sigma_T^2 \mathbf{I} \right) \\
& \quad \times \prod_{u=1}^N \mathcal{N} \left(U_u^s | \sum_{k \in S_u} T_{uk} U_k, \sigma_T^2 \mathbf{I} \right) \\
& \quad \times \prod_{u=1}^N \mathcal{N} (U_u | 0, \sigma_U^2 \mathbf{I}) \times \prod_{i=1}^M \mathcal{N} (V_i | 0, \sigma_V^2 \mathbf{I}) \\
& \quad \times \text{Beta}(\theta_G | \alpha_{\theta_G}, \beta_{\theta_G}) \times \prod_{u=1}^N \text{Beta}(B_u | \alpha_B, \beta_B). \tag{11}
\end{aligned}$$

Compared to SMF, our PTPMF model shown in (11) and Figure 2 treats strong and weak ties separately, learns the optimal (w.r.t. recommendation accuracy) threshold for distinguishing strong and weak ties. In addition, our PTPMF is able to learn a personalized tie preference (denoted as B_u) for each user u . Our goal is to learn $U, U^w, U^s, V, B, \theta_G$ which maximize the posterior probability shown in (11).

5. PARAMETER LEARNING

We learn the parameters of PTPMF using maximum a posteriori (MAP) inference. Taking the $\ln$ on both sides of (11), we are maximizing the following objective function:

$$\begin{aligned}
& \ln p(U^w, U^s, U, V, B, \theta_G | R, T, \sigma_R^2, \sigma_T^2, \sigma_U^2, \sigma_V^2) \\
& = -\frac{1}{2\delta_R^2} \sum_{u=1}^N \sum_{i=1}^M I_{ui}^R \left(R_{ui} - g(\mu_{R_{ui}}) \right)^2 \\
& \quad - \frac{1}{2\delta_U^2} \sum_{u=1}^N U_u^T U_u - \frac{1}{\delta_V^2} \sum_{i=1}^M V_i^T V_i \\
& \quad - \frac{1}{\delta_T^2} \sum_{u=1}^N \left((U_u^w - \sum_{k \in W_u} T_{uk} U_k)^T (U_u^w - \sum_{k \in W_u} T_{uk} U_k) \right) \\
& \quad - \frac{1}{\delta_T^2} \sum_{u=1}^N \left((U_u^s - \sum_{k \in S_u} T_{uk} U_k)^T (U_u^s - \sum_{k \in S_u} T_{uk} U_k) \right) \\
& \quad + \sum_{u=1}^N \left((\alpha_B - 1) \ln B_u + (\beta_B - 1) \ln(1 - B_u) \right) \\
& \quad + (\alpha_{\theta_G} - 1) \ln \theta_G + (\beta_{\theta_G} - 1) \ln(1 - \theta_G)
\end{aligned}$$

$$\begin{aligned}
& -\frac{1}{2} \left((N \cdot K) \ln \delta_U^2 + (M \cdot K) \ln \delta_V^2 + (2N \cdot K) \ln \delta_T^2 \right) \\
& -\frac{1}{2} \left(\sum_{u=1}^N \sum_{i=1}^M I_{ui}^R \ln \delta_R^2 - N \ln B(\alpha_B, \beta_B) - \ln B(\alpha_{\theta_G}, \beta_{\theta_G}) \right) \\
& + \text{Constant}, \tag{12}
\end{aligned}$$

where

$$\begin{aligned}
\mu_{R_{ui}} & = B_u \left(f(\theta_G) U_u^{wT} V_i + (1 - f(\theta_G)) U_u^{sT} V_i \right) \\
& + (1 - B_u) \left((1 - f(\theta_G)) U_u^{wT} V_i + f(\theta_G) U_u^{sT} V_i \right), \tag{13}
\end{aligned}$$

and $B(\cdot, \cdot)$ is the beta function:

$$B(x, y) = \int_0^1 t^{x-1} (1-t)^{y-1} dt. \tag{14}$$

Fixing the Gaussian noise variance and beta shape parameters, maximizing the log-posterior in (12) over $U^w, U^s, U, V, B, \theta_G$ is equivalent to minimizing the following objective function:

$$\begin{aligned}
& \mathcal{L}(R, T, U^w, U^s, U, V, B, \theta_G) \\
& = \frac{1}{2} \sum_{u=1}^N \sum_{i=1}^M I_{ui}^R \left(R_{ui} - g(\mu_{R_{ui}}) \right)^2 \\
& + \frac{\lambda_U}{2} \sum_{u=1}^N U_u^T U_u + \frac{\lambda_V}{2} \sum_{i=1}^M V_i^T V_i \\
& + \frac{\lambda_T}{2} \sum_{u=1}^N \left((U_u^w - \sum_{k \in W_u} T_{uk} U_k)^T (U_u^w - \sum_{k \in W_u} T_{uk} U_k) \right) \\
& + \frac{\lambda_T}{2} \sum_{u=1}^N \left((U_u^s - \sum_{k \in S_u} T_{uk} U_k)^T (U_u^s - \sum_{k \in S_u} T_{uk} U_k) \right) \\
& - \lambda_B \sum_{u=1}^N \left((\alpha_B - 1) \ln B_u + (\beta_B - 1) \ln(1 - B_u) \right) \\
& - \lambda_{\theta_G} \left((\alpha_{\theta_G} - 1) \ln \theta_G + (\beta_{\theta_G} - 1) \ln(1 - \theta_G) \right), \tag{15}
\end{aligned}$$

where $\lambda_U = \frac{\delta_R^2}{\delta_U^2}$, $\lambda_V = \frac{\delta_R^2}{\delta_V^2}$, $\lambda_T = \frac{\delta_R^2}{\delta_T^2}$ and $\lambda_B = \lambda_{\theta_G} = \delta_R^2$.

A local minimum of (15) can be found by taking the derivative and performing gradient descent on $U^w, U^s, U, V, B, \theta_G$ separately. The corresponding partial derivative with respect to each model parameter is shown as follows:

$$\begin{aligned}
\frac{\partial \mathcal{L}}{\partial U_u^s} & = \sum_{i=1}^M I_{ui}^R \left(g(\mu_{R_{ui}}) - R_{ui} \right) g'(\mu_{R_{ui}}) \\
& \quad \left(B_u + f(\theta_G) - 2B_u f(\theta_G) \right) V_i \\
& \quad + \lambda_T \left(U_u^s - \sum_{k \in S_u} T_{uk} U_k \right), \tag{16}
\end{aligned}$$

$$\begin{aligned}
\frac{\partial \mathcal{L}}{\partial U_u^w} & = \sum_{i=1}^M I_{ui}^R \left(g(\mu_{R_{ui}}) - R_{ui} \right) g'(\mu_{R_{ui}}) \\
& \quad \left(1 - f(\theta_G) - B_u + 2B_u f(\theta_G) \right) V_i \\
& \quad + \lambda_T \left(U_u^w - \sum_{k \in W_u} T_{uk} U_k \right), \tag{17}
\end{aligned}$$

$$\begin{aligned}
\frac{\partial \mathcal{L}}{\partial U_u} & = \lambda_U U_u - \lambda_T \sum_{v|u \in W_v} T_{vu} (U_v^w - \sum_{k \in W_v} T_{vk} U_k) \\
& - \lambda_T \sum_{v|u \in S_v} T_{vu} (U_v^s - \sum_{k \in S_v} T_{vk} U_k), \tag{18}
\end{aligned}$$

$$\frac{\partial \mathcal{L}}{\partial V_i} = \sum_{u=1}^N I_{ui}^R (g(\mu_{R_{ui}}) - R_{ui}) g'(\mu_{R_{ui}}) \left((B_u + f(\theta_G) - 2B_u f(\theta_G)) U_u^s + (1 - f(\theta_G) - B_u + 2B_u f(\theta_G)) U_u^w \right) + \lambda_V V_i, \quad (19)$$

$$\frac{\partial \mathcal{L}}{\partial B_u} = \sum_{i=1}^M I_{ui}^R (g(\mu_{R_{ui}}) - R_{ui}) g'(\mu_{R_{ui}}) \left((2f(\theta_G) - 1) U_u^w + (1 - 2f(\theta_G)) U_u^s \right) V_i - \lambda_B \left(\frac{\alpha_B - 1}{B_u} - \frac{\beta_B - 1}{1 - B_u} \right), \quad (20)$$

$$\frac{\partial \mathcal{L}}{\partial \theta_G} = (t_s + t_w - 2) g'((t_s - \theta_G)(\theta_G - t_w)) \sum_{u=1}^N \sum_{i=1}^M I_{ui}^R (g(\mu_{R_{ui}}) - R_{ui}) g'(\mu_{R_{ui}}) \left((2B_u - 1) U_u^w + (1 - 2B_u) U_u^s \right) V_i - \lambda_{\theta_G} \left(\frac{\alpha_{\theta_G} - 1}{\theta_G} - \frac{\beta_{\theta_G} - 1}{1 - \theta_G} \right). \quad (21)$$

The update is done using standard gradient descent:

$$x^{(t+1)} = x^{(t)} + \eta^{(t)} \cdot \frac{\partial \mathcal{L}}{\partial x}(x^{(t)}), \quad (22)$$

where η is the learning rate and $x \in \{U^w, U^s, U, V, B, \theta_G\}$ denotes any model parameter. Finally, the algorithm terminates when the absolute difference between the losses in two consecutive iterations is less than 10^{-5} .

We note that in order to avoid overfitting, our proposed model has the standard regularization terms (L2 norm) for user latent feature vectors ($\sum U_u^T U_u$) and item latent feature vectors ($\sum V_i^T V_i$) in the third line of (15). Since the weak tie and strong tie latent feature vectors depend on the user latent feature vectors, these additional parameters in our model are also indirectly regularized.

6. EMPIRICAL EVALUATION

In this section, we report the results of our experiments on four real-world public datasets and compare the performance of our PTPMF model with different baseline methods in terms of various evaluation metrics. Our experiments aim to examine if incorporating the new concepts of distinguishing strong and weak ties is able to improve the recommendation accuracy as measured by MAE / RMSE (how close the predicted ratings are to the real ones) and Precision@K / Recall@K (accuracy for top-K recommendations), and how significant are the improvements achieved if any.

6.1 Experimental Settings

Datasets. We use the following four real-world datasets.

- *Flixster*. The Flixster dataset² containing information of user-movie ratings and user-user friendships from Flixster, an American social movie site for discovering new movies (<http://www.flixster.com/>).
- *CiaoDVD*. This public dataset contains trust relationships among users as well as their ratings on DVDs and was crawled from the entire category of DVDs of a UK DVD community website (<http://dvd.ciao.co.uk>) in December, 2013 [11].

²<http://www.cs.ubc.ca/~jamalim/datasets/>

- *Douban*. This public dataset³ is extracted from the Chinese Douban movie forum (<http://movie.douban.com/>), which contains user-user friendships and user-movie ratings.
- *Epinions*. This is the Epinions dataset⁴ which consists of user-user trust relationships and user-item ratings from Epinions (<http://www.epinions.com/>).

The statistics of these data sets are summarized in Table 1.

	<i>Flixster</i>	<i>CiaoDVD</i>	<i>Douban</i>	<i>Epinions</i>
#users	76013	1881	64642	31117
#items	48516	12900	56005	139057
#non-zeros	7350235	33510	9133529	654103
#ties (edges)	1209962	15155	1390960	410570

Table 1: Overview of datasets (#non-zeros means the number of user-item pairs that have feedback)

For all the datasets, we randomly choose 80% of each user's ratings for training, leaving the remainder for testing. We split the portion of the 80% of the dataset (i.e., the training set) into five equal sub-datasets for 5-fold cross validation. During the training and validation phase, each time we use one of the five sub-datasets for validation and the remaining for training. We repeat this procedure five times so that all five sub-datasets can be used for validation. And we pick the parameter values having the best average performance. Then we evaluate different models on the 20% of the dataset left for testing (i.e., the test set).

Methods Compared. In order to show the performance improvement of our PTPMF method, we will compare our method with some state-of-art approaches which consist of non-personalized non-social methods, personalized non-social methods and personalized social methods. Thus, the following nine recommendation methods, including eight baselines, are tested.

- *PTPMF*. Our proposed PTPMF model, which is a personalized social recommendation approach by exploiting social ties.
- *TrustMF*. A personalized social method originally proposed by Yang et al. [42], which is capable of handling trust propagation among users.
- *SMF*. This is a personalized social approach [15] which assumes that users' latent feature vectors are dependent on those of their ties.
- *SoReg*. The individual-based regularization model with Pearson Correlation Coefficient (PCC) which outperforms its other variants, as indicated in [29]. This is a personalized social method.
- *STE*. Another personalized social method proposed by Ma et al. [26] which aggregates a user's own rating and her friends' ratings to predict the target user's final rating on an item.
- *SoRec*. The probabilistic matrix factorization model proposed by Ma et al. [28] which factorizes user-item rating matrix and user-user linkage matrix simultaneously. This is also a personalized social method.
- *PMF*. The classic personalized non-social probabilistic matrix factorization model first introduced in [30].

³<https://www.cse.cuhk.edu.hk/irwin.king.new/pub/data/douban>

⁴http://www.trustlet.org/wiki/Epinions_dataset

		UserMean	ItemMean	PMF	SoRec	STE	SMF	SoReg	TrustMF	PTPMF
Flixster	MAE	0.840127	0.853447	0.801346	0.795724	0.770012	0.749708	0.758309	0.792434	0.715910
	Impv	14.7%	16.1%	10.7%	10.0%	7.03%	4.51%	5.59%	9.66%	—
	RMSE	1.061324	1.074465	1.012973	1.008995	0.974290	0.952560	0.960418	1.001670	0.914541
	Impv	13.8%	14.9%	9.72%	9.36%	6.13%	3.99%	4.78%	8.70%	—
CiaoDVD	MAE	0.904175	0.894703	0.876668	0.830892	0.834754	0.829287	0.865684	0.828039	0.789901
	Impv	12.6%	11.7%	9.90%	4.93%	5.37%	4.75%	8.75%	4.61%	—
	RMSE	1.133421	1.195009	1.106291	1.088455	1.088869	1.109867	1.124882	1.087434	1.019105
	Impv	10.1%	14.7%	7.88%	6.37%	6.41%	8.18%	9.40%	6.28%	—
Douban	MAE	0.685375	0.627068	0.569055	0.568788	0.554951	0.554731	0.554378	0.569364	0.542439
	Impv	20.9%	13.5%	4.68%	4.63%	2.25%	2.22%	2.15%	4.73%	—
	RMSE	0.852284	0.783605	0.720964	0.719435	0.716873	0.717495	0.700033	0.720403	0.686182
	Impv	19.5%	12.4%	4.82%	4.62%	4.28%	4.36%	1.98%	4.75%	—
Epinions	MAE	0.969965	0.988781	0.916315	0.900854	0.882172	0.870062	0.897915	0.864209	0.822009
	Impv	15.2%	16.9%	10.3%	8.75%	6.82%	5.52%	8.45%	4.88%	—
	RMSE	1.170197	1.189446	1.137936	1.127590	1.120252	1.119862	1.121258	1.107024	1.060537
	Impv	9.37%	10.8%	6.80%	5.95%	5.33%	5.30%	5.42%	4.20%	—

Table 2: MAE and RMSE on all users (boldface font denotes the winner in that row)

- *UserMean*. A non-personalized non-social baseline, which makes use of the average ratings of users to predict missing values.
- *ItemMean*. Another non-personalized non-social baseline, utilizing the average ratings of each items to make predictions.

All experiments are conducted on a platform with 2.3 GHz Intel Core i7 CPU and 16 GB 1600 MHz DDR3 memory. We use grid search and 5-fold cross validation to find the best parameters. For example, we set $\lambda_U = \lambda_V = 0.001$ after exploring each value in (0.001, 0.0025, 0.005, 0.0075, 0.01, 0.025, 0.05, 0.075, 0.1) with cross validation and set $\lambda_B = \lambda_\theta = 0.00001$ in a similar way. The latent factor dimension is set to 10 for all models (if applicable). The learning rate of gradient descent (i.e., η) is set to 0.05 for θ_G and 0.001 for other parameters. For baselines, we adopt either the optimal parameters reported in the original paper or the best we can obtain in our experiments.

Evaluation Metrics. We use four metrics, i.e., Mean Absolute Error (MAE), Root Mean Square Error (RMSE), Recall and Precision, to measure the recommendation accuracy of our PTPMF model in comparison with other recommendation approaches.

- *Mean Absolute Error*.

$$\text{MAE} = \frac{\sum_{i,j} |R_{ij} - \hat{R}_{ij}|}{N}.$$

- *Root Mean Square Error*.

$$\text{RMSE} = \sqrt{\frac{\sum_{i,j} (R_{ij} - \hat{R}_{ij})^2}{N}}.$$

where R_{ij} is the rating that user i gives to item j (original rating) and $\hat{R}_{ij}$ is the predicted rating of user i for item j . N is the number of ratings in test set.

- *Recall@K*.

This metric quantifies the fraction of consumed items that are in the top- K ranking list sorted by their estimated rankings. For each user u we define $S(K; u)$ as the set of already-consumed items in the test set that appear in the top- K list and $S(u)$ as the set of all items consumed by this user in the test set. Then, we have

$$\text{Recall@K}(u) = \frac{|S(K; u)|}{|S(u)|}.$$

- *Precision@K*.

This measures the fraction of the top- K items that are indeed

consumed by the user in the test set:

$$\text{Precision@K}(u) = \frac{|S(K; u)|}{K}.$$

6.2 Experimental Results

Table 2 presents the performances of all nine recommendation methods on all four datasets, in terms of MAE and RMSE. We also present the percentage increase of PTPMF over each baseline right under its corresponding MAE and RMSE values and boldface font denotes the winner in each row. We would like to point out that, due to the randomness in data splitting and model initialization as well as differences in data preprocessing, our results for some baselines are slightly different from the results reported in the original papers. Among the eight baselines, UserMean and ItemMean are non-personalized methods which do not take social information into account; PMF is a personalized non-social model; the remainder are personalized approaches which also take social information into consideration. We observe from Table 2 that the personalized non-social method (PMF) outperforms the non-personalized non-social methods (UserMean and ItemMean), which shows the advantage of a personalized strategy. Moreover, through taking extra social network information into consideration, personalized social methods (SoRec, STE, SMF, SoReg and TrustMF) achieve a performance boost over the personalized non-social method (PMF), consistent with the assumption in the social recommendation literature that social information can help improve recommender systems. Finally, we observe that PTPMF consistently outperforms all eight baselines on all datasets for both metrics, demonstrating the benefit of the distinction and thresholding of different tie types, as well as learning a personalized tie preference for each user. Due to the randomness in data splitting, model initialization and even data preprocessing, our results for some baselines may not be exactly the same as reported in the original work, though given our best efforts to diminish the variances.

Recall and Precision. Figure 3 depicts Recall (X -axis) vs. Precision (Y -axis) of the seven recommendation methods. We exclude the two naive methods (UserMean and ItemMean) for the sake of clarity of the figures. Data points from left to right on each line were calculated at different values of K , ranging from 5 to 50. Clearly, the closer the line is to the top right corner, the better the algorithm is, indicating that both recall and precision are high. We observe that PTPMF again clearly outperforms all baselines. Besides, Figure 3 also demonstrates the trade-off between recall and precision, i.e., as K increases, recall will go up while precision will go down.

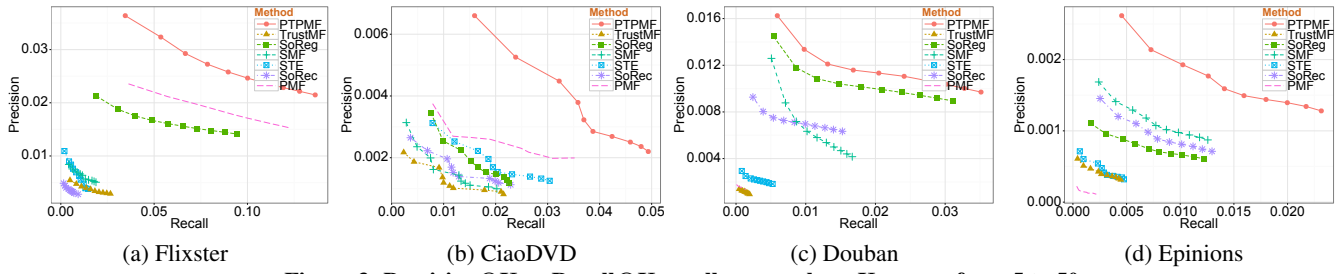


Figure 3: Precision@K vs Recall@K on all users, where K ranges from 5 to 50

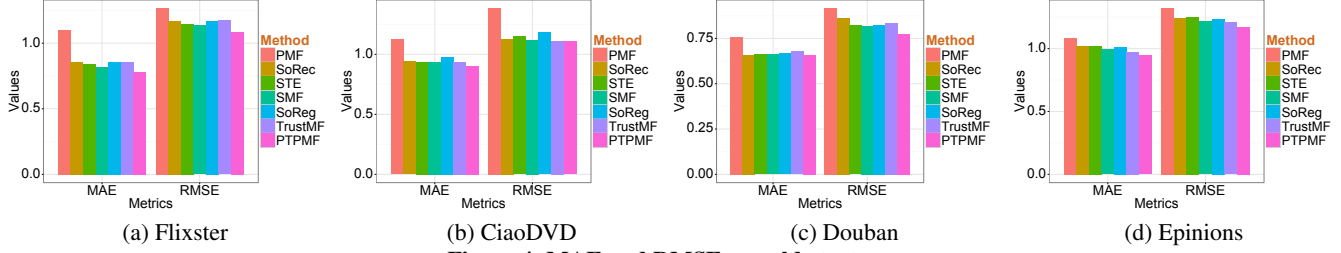


Figure 4: MAE and RMSE on cold-start users

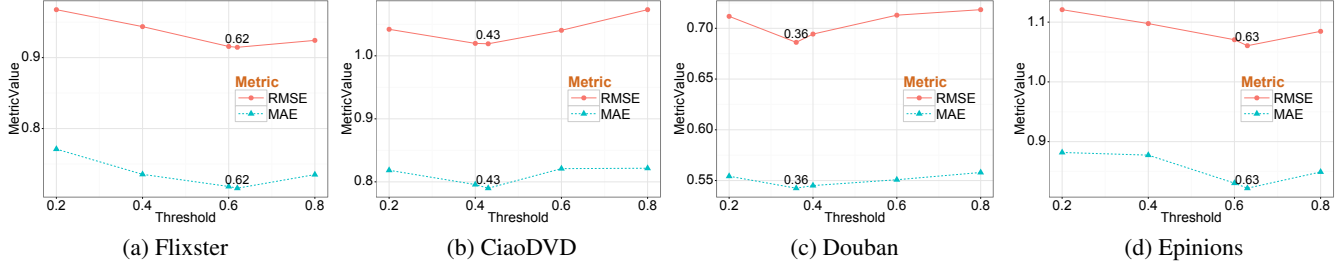


Figure 5: MAE and RMSE for several pre-fixed thresholds and our learned thresholds, with the numbers (and the corresponding points below them) denoting our learned threshold values

Comparisons on Cold-Start Users. We further drill down to the *cold-start* users. As is common practice, we define users that rated less than five items as cold-start. Figure 4 shows the performances of various methods on cold start users. It is well known that the social recommendation methods are superior to their non-social competitors particularly for cold-start users. The results in Figure 4 verify this – all social recommendation methods significantly outperform PMF in terms of both MAE and RMSE. Furthermore, our PTPMF model again beats other social recommendation baselines.

Learned threshold vs. Fixed threshold. Last but not least, we compare the results from our learned thresholds with those from several pre-fixed thresholds in Figure 5 in order to prove that the threshold learning does contribute to the accuracy of the recommendations. For each dataset, we set θ_G to be four fixed values, i.e., 0.2, 0.4, 0.6, 0.8. We then compare the results obtained through fixing θ_G with that obtained from dynamically learning the threshold. Figure 5 demonstrates that the best results are achieved by the dynamically learned thresholds in terms of both MAE and RMSE. We remark that the thresholds learned from different datasets vary greatly, which is another supporting argument for learning the thresholds from the data.

In summary, we compare PTPMF with various kinds of baselines including non-personalized non-social methods, personalized non-social methods and personalized social methods in terms of both rating prediction and top- K ranking evaluation metrics. We conclude from the above extensive experiments that our proposed model, PTPMF, is an effective social recommendation method given its better performance over other baselines on both all users and cold-start users.

7. CONCLUSIONS

In this paper, inspired by the seminal work in social science [10, 9], we start from recognizing the important roles of different tie types in social relations and present a novel social recommendation model, a non-trivial extension to probabilistic matrix factorization, to incorporate the personalized preference of strong and weak ties into social recommendation. Our proposed method, PTPMF, is capable of simultaneously classifying strong and weak ties w.r.t. recommendation accuracy in a social network, and learning a personalized tie type preference for each individual as well as other model parameters.

We carry out thorough experiments on four real-world datasets to demonstrate the gains of our proposed method. The experimental results show that PTPMF provides the best accuracy in various metrics, demonstrating that learning user-specific preferences for different types of ties in social recommendation does help to improve the performance.

One interesting direction for future work is to find a personalized threshold of classifying strong and weak ties for each user, though it can be challenging due to the sparsity of data. Further, we did not examine other node similarity metrics such as Adamic-Adar [1] or Katz [17] in this work and it is also interesting to explore different node similarity metrics.

ACKNOWLEDGMENTS

This research is supported by the National Research Foundation, Prime Minister’s Office, Singapore under its International Research Centres in Singapore Funding Initiative, and also supported by the National Science Foundation of China (Grant Number: 61173186).

8. REFERENCES

- [1] L. A. Adamic and E. Adar. Friends and neighbors on the web. *Social networks*, 25(3):211–230, 2003.
- [2] V. Arnaboldi, A. Guazzini, and A. Passarella. Egocentric online social networks: Analysis of key features and prediction of tie strength in facebook. *Computer Communications*, 36(10):1130–1144, 2013.
- [3] S. Berkovsky, J. Freyne, and G. Smith. Personalized network updates: increasing social interactions and contributions in social networks. In *User Modeling, Adaptation, and Personalization*, pages 1–13. Springer, 2012.
- [4] R. M. Bond, C. J. Fariss, J. J. Jones, A. D. I. Kramer, C. Marlow, J. E. Settle, and J. H. Fowler. A 61-million-person experiment in social influence and political mobilization. *Nature*, 489(7415):295–298, 09 2012.
- [5] N. A. Christakis and J. H. Fowler. Connected: how your friends' friends' friends affect everything you feel, think, and do. *New York, NY: Little, Brown, and Company*, 2009.
- [6] S. Fortunato. Community detection in graphs. *Physics Reports*, 486(3):75–174, 2010.
- [7] E. Gilbert. Predicting tie strength in a new medium. In *Proceedings of the ACM 2012 conference on Computer Supported Cooperative Work*, pages 1047–1056. ACM, 2012.
- [8] E. Gilbert and K. Karahalios. Predicting tie strength with social media. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, pages 211–220. ACM, 2009.
- [9] M. Granovetter. *Getting a job: A study of contacts and careers*. University of Chicago Press, 1995.
- [10] M. S. Granovetter. The strength of weak ties. *American journal of sociology*, pages 1360–1380, 1973.
- [11] G. Guo, J. Zhang, and N. Yorke-Smith. A novel bayesian similarity measure for recommender systems. In *Proceedings of the 23rd International Joint Conference on Artificial Intelligence (IJCAI)*, pages 2619–2625, 2013.
- [12] Y. Hu, Y. Koren, and C. Volinsky. Collaborative filtering for implicit feedback datasets. In *Data Mining, 2008. ICDM'08. Eighth IEEE International Conference on*, pages 263–272. Ieee, 2008.
- [13] P. Jaccard. *Distribution de la Flore Alpine: dans le Bassin des dranses et dans quelques régions voisines*. Rouge, 1901.
- [14] M. Jamali and M. Ester. Trustwalker: a random walk model for combining trust-based and item-based recommendation. In *Proceedings of the 15th ACM SIGKDD international conference on Knowledge discovery and data mining*, pages 397–406. ACM, 2009.
- [15] M. Jamali and M. Ester. A matrix factorization technique with trust propagation for recommendation in social networks. In *Proceedings of the fourth ACM conference on Recommender systems*, pages 135–142. ACM, 2010.
- [16] I. Kahanda and J. Neville. Using transactional information to predict link strength in online social networks. *ICWSM*, 9:74–81, 2009.
- [17] L. Katz. A new status index derived from sociometric analysis. *Psychometrika*, 18(1):39–43, 1953.
- [18] A. L. Kavanaugh, D. D. Reese, J. M. Carroll, and M. B. Rosson. Weak ties in networked communities. *The Information Society*, 21(2):119–131, 2005.
- [19] D. Kempe, J. M. Kleinberg, and É. Tardos. Maximizing the spread of influence through a social network. In *KDD*, pages 137–146, 2003.
- [20] Y. Koren. Factorization meets the neighborhood: a multifaceted collaborative filtering model. In *Proceedings of the 14th ACM SIGKDD international conference on Knowledge discovery and data mining*, pages 426–434. ACM, 2008.
- [21] Y. Koren. Collaborative filtering with temporal dynamics. *Communications of the ACM*, 53(4):89–97, 2010.
- [22] Y. Koren, R. Bell, and C. Volinsky. Matrix factorization techniques for recommender systems. *Computer*, (8):30–37, 2009.
- [23] J. Li, X. Hu, J. Tang, and H. Liu. Unsupervised streaming feature selection in social media. In *Proceedings of the 24th ACM International Conference on Conference on Information and Knowledge Management*, pages 1041–1050. ACM, 2015.
- [24] D. Liben-Nowell and J. Kleinberg. The link-prediction problem for social networks. *Journal of the American society for information science and technology*, 58(7):1019–1031, 2007.
- [25] C. Liu, T. Jin, S. C. Hoi, P. Zhao, and J. Sun. Collaborative topic regression for online recommender systems: an online and bayesian approach. *Machine Learning*, 2017.
- [26] H. Ma, I. King, and M. R. Lyu. Learning to recommend with social trust ensemble. In *Proceedings of the 32nd international ACM SIGIR conference on Research and development in information retrieval*, pages 203–210. ACM, 2009.
- [27] H. Ma, I. King, and M. R. Lyu. Learning to recommend with explicit and implicit social relations. *ACM TIST*, 2(3):29, 2011.
- [28] H. Ma, H. Yang, M. R. Lyu, and I. King. Sorec: social recommendation using probabilistic matrix factorization. In *Proceedings of the 17th ACM conference on Information and knowledge management*, pages 931–940. ACM, 2008.
- [29] H. Ma, D. Zhou, C. Liu, M. R. Lyu, and I. King. Recommender systems with social regularization. In *Proceedings of the fourth ACM international conference on Web search and data mining*, pages 287–296. ACM, 2011.
- [30] A. Mnih and R. Salakhutdinov. Probabilistic matrix factorization. In *Advances in neural information processing systems*, pages 1257–1264, 2007.
- [31] J.-P. Onnela, J. Saramäki, J. Hyvönen, G. Szabó, D. Lazer, K. Kaski, J. Kertész, and A.-L. Barabási. Structure and tie strengths in mobile communication networks. *Proceedings of the National Academy of Sciences*, 104(18):7332–7336, 2007.
- [32] K. Panovich, R. Miller, and D. Karger. Tie strength in question & answer on social network sites. In *Proceedings of the ACM 2012 conference on computer supported cooperative work*, pages 1057–1066. ACM, 2012.
- [33] A. Petróczy, T. Nepusz, and F. Bazsó. Measuring tie-strength in virtual social networks. *Connections*, 27(2):39–52, 2007.
- [34] R. Reagans. Preferences, identity, and competition: Predicting tie strength from demographic data. *Management Science*, 51(9):1374–1383, 2005.
- [35] S. Rendle, C. Freudenthaler, Z. Gantner, and L. Schmidt-Thieme. Bpr: Bayesian personalized ranking from implicit feedback. In *Proceedings of the twenty-fifth conference on uncertainty in artificial intelligence*, pages 452–461. AUAI Press, 2009.
- [36] R. Salakhutdinov and A. Mnih. Bayesian probabilistic matrix factorization using markov chain monte carlo. In *Proceedings of the 25th international conference on Machine learning*, pages 880–887. ACM, 2008.
- [37] J. Wang, S. C. Hoi, P. Zhao, and Z.-Y. Liu. Online multi-task collaborative filtering for on-the-fly recommender systems. In *Proceedings of the 7th ACM conference on Recommender systems*, pages 237–244. ACM, 2013.
- [38] X. Wang, R. Donaldson, C. Nell, P. Gorniak, M. Ester, and J. Bu. Recommending groups to users using user-group engagement and time-dependent matrix factorization. In *Thirtieth AAAI Conference on Artificial Intelligence*, 2016.
- [39] X. Wang, W. Lu, M. Ester, C. Wang, and C. Chen. Social recommendation with strong and weak ties. In *Proceedings of the 25th ACM International on Conference on Information and Knowledge Management*, pages 5–14. ACM, 2016.
- [40] A. Wu, J. M. DiMicco, and D. R. Millen. Detecting professional versus personal closeness using an enterprise social network site. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, pages 1955–1964. ACM, 2010.
- [41] R. Xiang, J. Neville, and M. Rogati. Modeling relationship strength in online social networks. In *Proceedings of the 19th international conference on World wide web*, pages 981–990. ACM, 2010.
- [42] B. Yang, Y. Lei, D. Liu, and J. Liu. Social collaborative filtering by trust. In *Proceedings of the Twenty-Third international joint conference on Artificial Intelligence*, pages 2747–2753. AAAI Press, 2013.
- [43] S.-H. Yang, B. Long, A. Smola, N. Sadagopan, Z. Zheng, and H. Zha. Like like alike: joint friendship and interest propagation in social networks. In *Proceedings of the 20th international conference on World wide web*, pages 537–546. ACM, 2011.
- [44] M. Ye, X. Liu, and W.-C. Lee. Exploring social influence for recommendation: a generative model approach. In *Proceedings of the 35th international ACM SIGIR conference on Research and development in information retrieval*, pages 671–680. ACM, 2012.
- [45] Z. Yu, C. Wang, J. Bu, X. Wang, Y. Wu, and C. Chen. Friend recommendation with content spread enhancement in social networks. *Information Sciences*, 309:102–118, 2015.
- [46] T. Zhao, J. McAuley, and I. King. Leveraging social connections to improve personalized ranking for collaborative filtering. In *Proceedings of the 23rd ACM International Conference on Conference on Information and Knowledge Management*, pages 261–270. ACM, 2014.
- [47] J. Zhuang, T. Mei, S. C. Hoi, X.-S. Hua, and S. Li. Modeling social strength in social media community via kernel-based learning. In *Proceedings of the 19th ACM international conference on Multimedia*, pages 113–122. ACM, 2011.