

# FGRec: A Fine-Grained Point-of-Interest Recommendation Framework by Capturing Intrinsic Influences

Yijun Su<sup>1,3</sup>, Jia-Dong Zhang<sup>4</sup>, Xiang Li<sup>1,2,3</sup>, Daren Zha<sup>3\*</sup>, Ji Xiang<sup>3</sup>, Wei Tang<sup>1,2,3</sup>, Neng Gao<sup>2,3</sup>

<sup>1</sup>*School of Cyber Security, University of Chinese Academy of Sciences, Beijing, China*

<sup>2</sup>*State Key Laboratory of Information Security, Chinese Academy of Sciences, Beijing, China*

<sup>3</sup>*Institute of Information Engineering, Chinese Academy of Sciences, Beijing, China*

<sup>4</sup>*Department of Computer Science, City University of Hong Kong, Hong Kong, China*

{suyijun, lixiang9015, zhadaren, xiangji, tangwei, gaoneng} @iie.ac.cn

jzhang26@cityu.edu.hk

**Abstract**—Point-of-interest (POI) recommendation has become an important service to help users discover attractive locations. A variety of available check-in data make it possible to build a personalized POI recommender system, but the extreme sparsity of check-in data poses a severe challenge for POI recommendation. Recent studies mainly utilize *social information, categorical information and/or geographical information* to supplement the highly sparse check-in data. However, these studies often apply shallow methods for the extra information and provide considerably limited improvements on POI recommendation. In this paper, we propose a fine-grained POI recommendation framework, called FGRec to capture the intrinsic influences of social, categorical and geographical information on the check-in behaviors of users. First, we study the social influence in depth by exploiting the multi-hop social friends and top- $n$  nearest neighbor friends, not only the direct friends (i.e., 1-hop friends). Second, we investigate the categorical influence by factorizing both user-POI and user-category matrices simultaneously over the same user embedding space, rather than simply using the popularity of POI categories. Third, we explore the geographical influence by integrating two types of distance (i.e., *the distance between user homes and POIs and the distance among POIs*) into a unified probability distribution over check-in POIs, instead of modeling them separately. Finally, experimental results on two large-scale real-world datasets demonstrate the effectiveness and superiority of the proposed method.

**Index Terms**—POI recommendation, Intrinsic influence, Location-based social network

## I. INTRODUCTION

Recent years have witnessed the rapid prevalence of location-based social networks (LBSNs), such as Foursquare, Yelp and Facebook Places. These LBSNs have attracted millions of users to check in point-of-interests (POIs), e.g., restaurants, cinemas and tourists spots, and share their experiences of visiting these POIs with friends. LBSNs have accumulated various data including *historical check-ins of users on POIs, social links between users, categories of POIs, and geographical information of user homes and POIs*. These data bring unparalleled opportunities for developing a personalized POI

recommender system, which is a crucial demand in location-based services [1], [2].

It is still a challenging task to build an effective POI recommender system in LBSNs, because the recommendation performance is severely affected by the extreme sparsity of check-in data. To address this challenge, current studies mainly utilize **three types of information to supplement the highly sparse check-in data**: (1) **Social information**. Most works incorporate social links between users into POI recommendation based on collaborative filtering (CF) methods, e.g., friend-based CF [3], [4], [5], matrix factorization with social regularization [6], and friend-based matrix factorization [7]. However, these CF methods purely focus on using the direct friends of users (i.e., 1-hop friends) and overlook the effect of multi-hop friends. For example, the friends of friends (i.e., 2-hop friends) of a user may also affect the check-in behaviors of the user. (2) **Category information**. Existing methods [8], [9], [10], [7] often aggregate the popularity of categories for users or POIs and integrate the popularity into CF for deriving the similarity of users or the relevance score of users on POIs. Nonetheless, these existing methods are relatively simple in utilizing the category information and may not capture the preference of a user on a given POI category. (3) **Geographical information**. Most methods [3], [6], [5], [11], [12], [13], [14] employ the geographical information by estimating the check-in probability distribution *over the distance between user homes and POIs and/or over the distance among POIs*. Nevertheless, these methods model the two distance distributions separately and may not catch the interaction of the two types of distance.

To address the limitations of existing methods using social, categorical and geographical information, in this paper we concentrate on deeply modeling the intrinsic influences of the three types of information. (1) **Social influence**. Beside using the direct friends of a user, this study also leverages multi-hop social friends and top- $n$  nearest neighbor friends, i.e., a set of users who are geographically close to the user's home and have not socially connected with this user. To

\* Corresponding author.

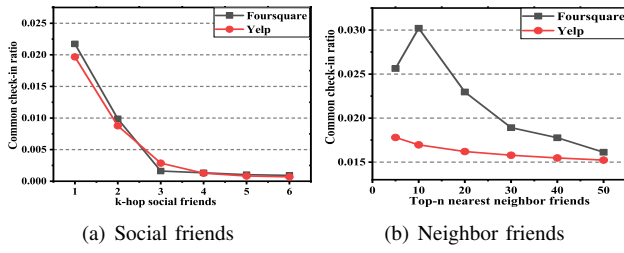


Fig. 1. Common check-in ratio between users

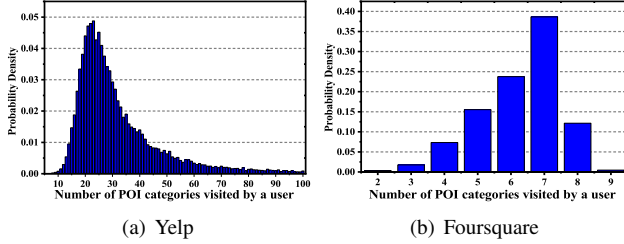


Fig. 2. Distribution of POI categories

clarify this motivation, we analyze the check-in data collected from Foursquare and Yelp (see TABLE III in Section IV for details) and plot the common check-in ratio between users and their friends in Fig. 1; the common check-in ratio is measured by  $\frac{|\mathcal{L}_i \cap \mathcal{L}_j|}{|\mathcal{L}_i \cup \mathcal{L}_j|}$ , where  $\mathcal{L}_i$  and  $\mathcal{L}_j$  are the set of check-in locations of user  $i$  and user  $j$ , respectively. From Fig. 1(a), the shorter the hop between users, the higher the common check-in ratio. Obviously, the  $k$ -hop friendship ( $k > 1$ ) has an impact on users' check-in behaviors. From Fig. 1(b), the closer to the user, the higher the common check-in ratio, which indicates that the neighbor friends have a greater impact than the distant friends. (2) **Categorical influence.** To investigate the impact of POI categories on users' check-in behaviors, we plot the distribution on the number of POI categories visited by a user on the two check-in datasets in Fig. 2. Based on Fig. 2, the number of POI categories visited by most users is relatively small and concentrated, which indicates that users have unique preferences on POI categories. In other words, users usually like several fixed categories of POIs. (3) **Geographical influence.** To study the influence of the two types of geographical distance (i.e., *the distance between user homes and POIs* and *the distance among POIs*), we estimate the distribution of activity ranges of users in the two datasets. Specifically, Fig. 3(a) depicts the distribution of the activity range measured by the max distance of the user's home to her visited POIs, while Fig. 3(b) shows the distribution of the activity range measured by the max distance of all pairs of POIs visited by the same user. According to Fig. 3, most users' activity ranges are relatively centralized, which is in line with the human mobility pattern because people tend to check in nearby locations [3], [6], [15], [16].

Thus, we are motivated to propose a **Fine-Grained POI Recommendation** framework by capturing the three types of intrinsic influences (i.e., social, categorical and geographical influences), called **FGRec** consisting of three modules corre-

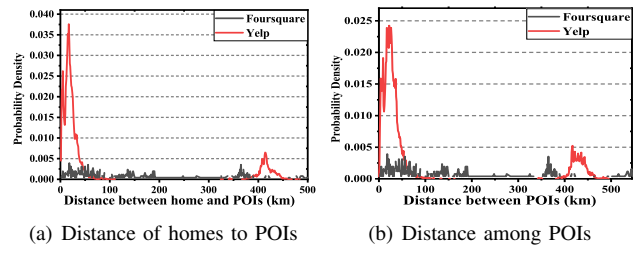


Fig. 3. Distribution of users' activity ranges

spondingly: (1) **Social module.** FGRec adaptively captures the impact of multi-hop social friends and top- $n$  nearest neighbor friends by designing a novel collective friends model (**CFM**); CFM combines friend-based CF with a network representation learning technique which captures the friendship between any two users by embedding complex social relations into a low-dimensional vector space. (2) **Category module.** FGRec leverages the categorical influence by devising a joint Poisson factor (**JPF**) model for simultaneously factorizing *user-POI matrix* and *user-category matrix* over the same user embedding space. The user-category matrix has much higher density than the user-POI matrix and thus greatly enhances the learning on the user embeddings. (3) **Geography module.** FGRec takes full advantage of the geographical influence by developing a personalized Gaussian kernel model (**GKM**); GKM estimates a unified probability distribution on the two types of distance, i.e., *the distance between user homes and POIs* and *the distance among POIs*.

The main contributions of this study can be summarized:

- We deeply study the impact of multi-hop social friends and top- $n$  nearest neighbor friends on the user check-in behaviors. Modeling such fine-grained social correlations effectively alleviates the sparsity problem and assists to make accurate recommendations.
- We concurrently factorize user-POI matrix and user-category matrix by sharing the user embedding space. In comparison with the extremely sparse user-POI matrix, the higher density of the user-category matrix significantly enhances the learning on the user embeddings.
- We unify a probability model for *the distance between user homes and POIs* and *the distance among POIs*. The unified model can capture the interaction of the two types of distance and improve the geographical influence modeling for POI recommendation.
- We conduct extensive experiments on two large-scale real-world datasets to evaluate the performance of our framework. The experimental results show that our framework outperforms other state-of-the-art methods.

The rest of this paper is organized as follows. Section II reviews the related work. Section III presents the proposed model in detail. Section IV reports the experimental results. Finally, we draw some conclusions of this study in Section V.

## II. RELATED WORK

Many traditional collaborative filtering (CF) technologies have been proposed for POI recommendation. Memory-based CF is very popular and well-known. It can be classified into user-based CF and item-based CF. User-based [3] first finds similar users to the target user based on their historical check-ins using a similarity measure, such as Cosine similarity or Pearson correlation. Then the preference from the user on an unvisited POI can be derived by computing a weighted combination of historical check-ins on the same item from similar users. In contrast, item-based works according to the user's preferences on other similar items. Matrix factorization (MF) [6], [11], [7], [17], [18], [19], [1], [20], [21] has also become a popular model in POI recommendation as it can learn the latent factors that represent users' inherent preferences over an item's multiple dimensions.

Efforts have also been made to utilize social, categorical and geographical influences for improving the performance of recommendation. (1) **Social influence.** The information of friends is widely used in POI recommender systems [3], [22], [4], [23], [6], [24]. Friend-based CF [3] and MF with social regularization [6] are two effective CF algorithms in LBSNs. Besides, Zhang et al. [8] designed a model to estimate the social check-in frequency by using a power-law distribution. Li et al. [7] developed a POI recommendation framework that mainly exploits potential locations learned from users' friends to improve recommendation accuracy. Manotumruksa et al. [12] proposed a personalized ranking framework with social correlations and geographical influences for POI recommendation. However, existing methods of modeling social influence [3], [8], [12], [6] overlook the impact of multi-hop social friends and top- $n$  nearest neighbor friends. (2) **Categorical influence.** Categorical influences are often used to express users' preferences for POI categories in existing works [8], [10], [7], [25]. Zhang et al. [8] employed a power-law distribution to model the popularity of POI categories for capturing the categorical influence. Li et al. [25] transformed one-hot representation of POI categories to a latent vector as its embedding for next POI recommendation. These existing methods of utilizing category information [8], [10], [7], [25] are relatively simple, and few of these methods directly capture the preference of the user on a given POI category. (3) **Geographical influence.** Some studies [3], [6], [5], [26], [27], [28], [29], [30], [31], [32] have pointed out that geographical influence can be used to improve the performance of POI recommendation. In particular, several representative models, such as power-law distribution [3], multi-center Gaussian distribution [6], and kernel density estimation [5], are proposed to capture the geographical influence in POI recommendation. Most of the above algorithms do not simultaneously consider the two types of geographical distance (i.e., *the distance between user homes and POIs* and *the distance among POIs*), which are important for determining whether users travel far or not.

In this paper, we address the aforementioned limitations by deeply modeling the three kinds of information. Our

TABLE I  
MATHEMATICAL NOTATIONS

| Symbol     | Size         | Meaning                                        |
|------------|--------------|------------------------------------------------|
| $F^x$      | $m \times n$ | user-POI check-in matrix                       |
| $F^y$      | $m \times q$ | user-category check-in matrix                  |
| $C^x$      | $m \times n$ | user-POI expected count matrix                 |
| $C^y$      | $m \times q$ | user-category expected count matrix            |
| $f_{ij}^x$ | $\mathbb{R}$ | check-in frequency of user $i$ on POI $j$      |
| $f_{ic}^y$ | $\mathbb{R}$ | check-in frequency of user $i$ on category $c$ |
| $c_{ij}^x$ | $\mathbb{R}$ | expected count of user $i$ on POI $j$          |
| $c_{ic}^y$ | $\mathbb{R}$ | expected count of user $i$ on category $c$     |

framework differs from the above approaches in three aspects. First, we propose a novel collective friends model in capturing the social influence, which takes full advantage of the impact of multi-hop social friends and top- $n$  nearest neighbor friends. Second, we design a joint Poisson factor model to learn categorical influence by simultaneously factorizing user-POI matrix and user-category matrix by sharing the same user embedding space. Third, we present a personalized Gaussian kernel model to integrate *the distance between user homes and POIs* and *the distance among POIs* into a unified model.

## III. PROPOSED MODEL

In this section, we first define the problem of POI recommendation, then present the unified POI recommendation framework, and finally propose fine-grained models for capturing intrinsic influences.

### A. Problem Definition

We define  $\mathcal{U} = \{u_1, u_2, \dots, u_m\}$  as a set of users, where each user  $u_i$  checked in a set of POIs  $\mathcal{L}_i$ , denoted as  $\mathcal{L}_i = \{l_1, l_2, \dots, l_n\}$ , where each POI has a geographical coordinate  $l_j = \{\text{lon}_j, \text{lat}_j\}$  on the longitude and latitude and a set of categories  $\mathcal{Z}_j = \{z_1, z_2, \dots, z_q\}$ . For convenience, we term  $i$  as user  $u_i$ ,  $j$  as POI  $l_j$  and  $c$  as category  $z_c$ , unless stated otherwise. Each check-in action  $(i, j, t)$  indicates user  $i$  checks in location  $j$  at time  $t$  and all check-in actions are expressed as  $C$ . In addition, we use  $SF_i$  and  $NF_i$  to represent the user's social and neighbor friends, respectively. Some key notations used in this paper are listed in TABLE I.

Given the geographical coordinates of POIs, user-POI check-in matrix  $F^x$ , user-category check-in matrix  $F^y$ , social friends  $SF_i$ , and neighbor friends  $NF_i$ , the task of POI recommendation is to predict the preference score  $r_{i,j}$  for user  $i$  to an unvisited POI  $j$ , and then return the top- $N$  POIs with the highest recommendation score  $r_{i,j}$  to user  $i$ .

### B. Architecture

To better address the challenge of recommendation arising from the check-in data sparsity, we propose the fine-grained POI recommendation framework (named **FGRec**) by integrating the three types of intrinsic influences (i.e., social, categorical and geographical influences). Fig. 4 depicts the overall architecture of FGRec consisting of social, category and geography modules based on the collective friends model

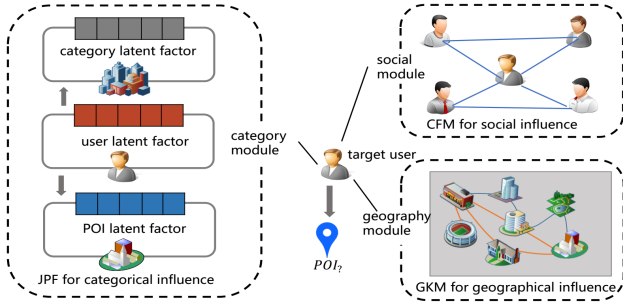


Fig. 4. The architecture framework of FGRec

(CFM), joint Poisson factor (JPF) model and Gaussian kernel model (GKM), respectively.

As shown in the previous studies [4], [8], [5], the product rule is simple and effective in integrating different factors, in which it is not required to normalize each factor because the normalization cannot affect the ranking of results. In this paper, we apply the product rule to integrate the above intrinsic influences. In the future work, we will explore more comprehensive integration methods, e.g., factorization machine and deep learning techniques, which is not the focus of this paper.

The recommendation score  $r_{i,j}$  for user  $i$  to the unvisited POI  $j$  can ultimately be expressed as follows:

$$r_{i,j} = p_{JPF}(i,j) \cdot p_{CFM}(i,j) \cdot p_{GKM}(i,j), \quad (1)$$

where  $p_{JPF}(i,j)$ ,  $p_{CFM}(i,j)$ ,  $p_{GKM}(i,j)$  are the relevant recommendation scores based on social, categorical and geographical influences, respectively.

### C. Modeling Social Influence

In LBSNs, the social link represents the social friendship of users. In this paper, we design a novel collective friends model (CFM), which combines friend-based collaborative filtering with a network representation learning technique, to capture the social influence by deeply leveraging the information of multi-hop social friends and top- $n$  nearest neighbor friends.

The major novelty of the proposed model lies in adaptively capturing the impact of multi-hop social friends and top- $n$  nearest neighbor friends for POI recommendation. Two types of friend correlation weights can be automatically learned. Based on the weighted method, the relevant recommendation score based on social influence can be written as:

$$p_{CFM}(i,j) = \psi_1 p(j|SF_i) + \psi_2 p(j|NF_i), \quad (2)$$

where  $\psi_1$  and  $\psi_2$  are correlation weights, and  $p(j|SF_i)$  and  $p(j|NF_i)$  are friend influence strengths. Since the two kinds of friends may overlap, the sum of weights is not equal 1.

Accordingly, the process of CFM consists of three steps: *correlation weight estimation*, *influence strength computation*, and *parameter inference*.

**Step 1: Correlation weight estimation.** We first investigate the check-in patterns between users and their two types of friends. We further plot the distribution of co-occurrences as the number of observed check-ins of all users increases in

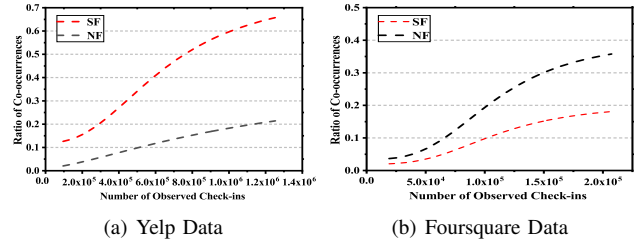


Fig. 5. Distribution of Co-occurrences.

TABLE II  
CHECK-IN FEATURES

| Feature    | Description                         |
|------------|-------------------------------------|
| $N_X$      | Number of $X$                       |
| $N_X^c$    | Number of check-ins from $X$        |
| $N_X^{cc}$ | Number of common check-ins from $X$ |
| $N_X^{nc}$ | Number of unique check-ins from $X$ |

Fig. 5. The  $x$ -axis represents the number of observed check-ins of all users, and the  $y$ -axis denotes the ratio of co-occurrences. It can be seen from the figure that the curve of co-occurrences increases with the increase of observed check-ins at the beginning, and eventually tends to stabilize. The reason for this trend may come from two parts: (1) in the early time, users have few friends and historical records when they are LBSN new users; (2) as time goes on, users' friends and co-occurrences will increase, and finally become stable. Hence, it motivates us to set the correlation weight as an activation function, which considers a set of features capturing check-in behaviors between users and their friends:

$$\psi_o = f(\mathbf{w}_o^T \mathbf{f}_o), 0 \leq \psi_o \leq 1, \quad (3)$$

where  $o \in \{1, 2\}$  ( $o = 1$  for social friends (SF) and  $o = 2$  for neighbor friends (NF)),  $\mathbf{w}_o$  is a weight vector, and  $\mathbf{f}_o$  is a feature vector. In this study, we define four features in TABLE II, where  $X$  is SF or NF,  $N_X$  is the number of  $X$ ,  $N_X^c$  represents the number of check-ins from  $X$ ,  $N_X^{cc}$  denotes the total number of common check-ins between the user and  $X$  and  $N_X^{nc}$  indicates the number of unique check-ins from  $X$ .

$f(\bullet)$  is a real-valued and differentiable function that guarantees the range of  $\psi_o$  limited in  $[0, 1]$ . In this case, a sigmoid function is often used, which can approximately capture the characteristic:

$$f(\mathbf{w}_o^T \mathbf{f}_o) = \frac{1}{1 + \exp(-\mathbf{w}_o^T \mathbf{f}_o)}, \quad (4)$$

**Step 2: Influence strength computation.** The influence strength can be computed by the combination of the friend-based collaborative filtering [3] and the network representation learning technique [33]. For social friends, we have

$$p(j|SF_i) = \frac{\sum_{s \in SF_i} SI_{i,s} \cdot f_{s,j}^x}{\sum_{s \in SF_i} SI_{i,s}}, \quad (5)$$

where  $SI_{i,s}$  is the friend similarity between the user  $i$  and the friend  $s$ , and  $f_{s,j}^x$  is the check-in frequency of friends  $s$

on POI  $j$ . We define  $SI_{i,s}$  using a Gaussian kernel function, which fully captures the impact of multi-hop social friends:

$$SI_{i,s} = \exp(-\frac{\|\mathbf{g}_i - \mathbf{g}_s\|^2}{\sigma^2}), \quad (6)$$

where  $\sigma$  is a scale parameter that can be tuned by a local scaling technique,  $\mathbf{g}_i \in \mathbb{R}^k$  and  $\mathbf{g}_s \in \mathbb{R}^k$  are two low-dimensional vectors learned by the network representation learning technique, which can effectively learn user node representations by capturing the impact of multi-hop social friends. In this paper, we employ the Struct2vec representation learning algorithm [33]. This way is friendly to the user with fewer friends.

For neighbor friends, the definition of  $p(j|NF_i)$  is similar to Equation (5). We calculate the friend similarity with  $SN_{i,s} = \frac{|\mathcal{L}_i \cap \mathcal{L}_s|}{|\mathcal{L}_i \cup \mathcal{L}_s|}$ , where  $\mathcal{L}_i$  and  $\mathcal{L}_s$  are the set of check-ins of  $i$  and  $s$ , respectively. Here, we choose the top- $n$  nearest neighbor friends as  $NF_i$ .

**Step 3: Parameter inference.** Let  $\Theta = \{\mathbf{w}_1, \mathbf{w}_2\}$  denotes all parameters to be learned, where  $\mathbf{w}_1$  is the weight parameter vector of social friends, and  $\mathbf{w}_2$  represents the weight parameter vector of neighbor friends. In CFM, the product of check-in probabilities over on the whole set of user-POI interaction can be defined as follows:

$$P(C|\Theta) = \prod_{(i,j,t) \in C} p_{CFM}(i,j), \quad (7)$$

All parameters are learned by using maximum likelihood estimation (MLE) method, which can be converted to the following minimization problem.

$$\min \sum_{i,j,t \in C} -\ln P(C|\Theta) + \lambda (\|\mathbf{w}_1\|_2^2 + \|\mathbf{w}_2\|_2^2), \quad (8)$$

where  $\lambda$  is regularization parameter that avoids overfitting. In this paper, a Stochastic Gradient Descent (SGD) method is employed to solve Equation (8).

#### D. Modeling Categorical Influence

The user's category preference for POIs typically reveals some intrinsic characteristics of the user. For instance, a foodie is very interested in POIs related to food. To directly model the categorical influence, we propose a joint Poisson factor (JPF) model that simultaneously factorizes user-POI matrix  $F^x$  and user-category matrix  $F^y$  by sharing the user embedding space. In this model, the accuracy of learning user embedding can be improved by introducing the user-category matrix with much higher density. These studies [34], [6] have verified the fact that Gaussian distribution outputs poor performance when applied to the check-in frequency data. Thus, we turn to Poisson distribution. To demonstrate its effectiveness, we plot the check-in frequency distribution of a randomly selected user in Yelp and Foursquare in Fig. 6. From the figure, we observe that Poisson distribution is more suitable for fitting check-in frequency data. This observation is consistent with [6].

More specifically, for each observed check-in frequency  $f_{ij}^x$  in  $F^x$ , we assume that it follows the Poisson distribution

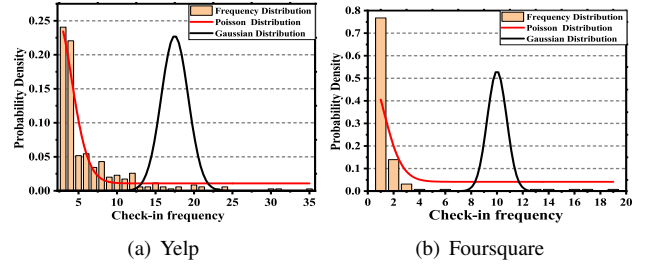


Fig. 6. The check-in frequency distribution in Yelp and Foursquare

with the mean  $c_{ij}^x$  in  $C^x$ :  $f_{ij}^x \sim \text{Poisson}(c_{ij}^x)$ . In the same way, we can get  $f_{ic}^y \sim \text{Poisson}(c_{ic}^y)$ . The matrix  $C^x$  is factorized into two matrices  $U \in \mathbb{R}^{m \times d}$  and  $L \in \mathbb{R}^{n \times d}$ , and  $C^y$  is decomposed into  $U \in \mathbb{R}^{m \times d}$  and  $Z \in \mathbb{R}^{q \times d}$ . Column vectors  $\mathbf{u}_i$ ,  $\mathbf{l}_j$  and  $\mathbf{z}_c$  represent user-specific, POI-specific and category-specific latent feature vectors, respectively. Each element  $u_{ik}$  in  $\mathbf{u}_i$ ,  $l_{jk}$  in  $\mathbf{l}_j$ , and  $z_{ck}$  in  $\mathbf{z}_c$  are assumed to follow the Gamma distribution as the empirical priors [6]:  $u_{ik} \sim \text{Gamma}(\alpha_U, \beta_U)$ ,  $l_{jk} \sim \text{Gamma}(\alpha_L, \beta_L)$ ,  $z_{ck} \sim \text{Gamma}(\alpha_Z, \beta_Z)$ , where  $\alpha_U, \alpha_L, \alpha_Z, \beta_U, \beta_L, \beta_Z > 0$ . Further, the expected counts  $c_{ic}^y = \sum_{k=1}^d u_{ik} z_{ck} = \mathbf{u}_i^T \mathbf{z}_c$  and  $c_{ij}^x = \sum_{k=1}^d u_{ik} l_{jk} = \mathbf{u}_i^T \mathbf{l}_j$ .

The log of posterior distribution over  $U, L$  and  $Z$  can be derived by utilizing the method of maximum a posterior (MAP) [34]:

$$p(U, L, Z | C^x, C^y, \alpha_U, \alpha_L, \alpha_Z, \beta_U, \beta_L, \beta_Z) \propto p(F^x | C^x) p(F^y | C^y) p(U | \alpha_U, \beta_U) p(L | \alpha_L, \beta_L) p(Z | \alpha_Z, \beta_Z), \quad (9)$$

Finally, we get the objective function of JPF model:

$$\begin{aligned} L(U, L, Z; F^x, F^y) = & \sum_{i=1}^m \sum_{j=1}^n (f_{ij}^x \ln c_{ij}^x - c_{ij}^x) \\ & + \sum_{i=1}^m \sum_{c=1}^q (f_{ic}^y \ln c_{ic}^y - c_{ic}^y) \\ & + \sum_{i=1}^m \sum_{k=1}^d ((\alpha_U - 1) \ln(u_{ik}/\beta_U) - u_{ik}/\beta_U) \\ & + \sum_{j=1}^n \sum_{k=1}^d ((\alpha_L - 1) \ln(l_{jk}/\beta_L) - l_{jk}/\beta_L) \\ & + \sum_{c=1}^q \sum_{k=1}^d ((\alpha_Z - 1) \ln(z_{ck}/\beta_Z) - z_{ck}/\beta_Z), \end{aligned} \quad (10)$$

where parameters  $U, L$  and  $Z$ , can be learned by minimizing  $L(U, L, Z; F^x, F^y)$ . We use the stochastic gradient ascent to update  $u_{ik}$ ,  $l_{jk}$ , and  $z_{ck}$ .

The predicted score based on categorical influence can be computed as follows:

$$p_{JPF}(i, j) = g \left( \left( \delta + \frac{1}{|\mathcal{Z}_j|} \sum_{c \in \mathcal{Z}_j} \mathbf{u}_i^T \mathbf{z}_c \right) \mathbf{u}_i^T \mathbf{l}_j \right), \quad (11)$$

where  $\mathbf{u}_i^T \mathbf{l}_j$  is the preference of user  $i$  for POI  $j$ ,  $\mathbf{u}_i^T \mathbf{z}_c$  denotes the preference of user  $i$  for category  $c$ ,  $\delta$  is a weight tuning



TABLE III  
STATISTICAL INFORMATION OF THE TWO DATASETS

| Statistical item        | Yelp    | Foursquare |
|-------------------------|---------|------------|
| Number of users         | 30,887  | 2,551      |
| Number of POIs          | 18,995  | 13,474     |
| Number of categories    | 624     | 10         |
| Number of check-ins     | 860,888 | 124,933    |
| Number of social links  | 265,533 | 32,512     |
| User-POI matrix density | 0.14%   | 0.291%     |

parameter and  $g(x) = 1/(1+\exp(-x))$  is the logistic function, which bounds the range of predictions to  $[0,1]$ .

#### E. Modeling Geographical Influence

The user's activity range is a key factor in determining whether the user travels far or not. In this paper, there are two types of geographical distance to be taken into consideration for capturing the geographical influence. (1) **Home-POI**: the geographical distance between users' home and POIs. The reason we consider is that the distance restricts users' check-in activity ranges. (2) **POI-POI**: the distance among POIs. The behind reason is that users are very interested in nearby POIs of a POI they liked (i.e., geographical clustering phenomenon or geographical proximity among POIs), even if it is far away from home. In view of the above geographical nature, we design a personalized Gaussian kernel model (GKM) to estimate a unified probability distribution by capturing the interaction of the two types of distance.

Therefore, the geographical relevance score based on geographical influence can be obtained:

$$p_{GKM}(i, j) = \frac{1}{\text{dist}(j, h_i)} \frac{\sum_{l_k \in \mathcal{L}_i} \exp(-\frac{\Upsilon_i}{2} \|l_j - l_k\|^2)}{\sum_{l_r \in \mathcal{L}} \sum_{l_k \in \mathcal{L}_i} \exp(-\frac{\Upsilon_i}{2} \|l_r - l_k\|^2)}, \quad (12)$$

where  $\|\cdot\|$  denotes the Euclidean norm in the geographical space,  $\sum_{l_r \in \mathcal{L}} \sum_{l_k \in \mathcal{L}_i} \exp(-\frac{\Upsilon_i}{2} \|l_r - l_k\|^2)$  is the normalization constant,  $\text{dist}(j, h_i)$  is the distance between user's home  $h_i$  and POI  $j$ ,  $\mathcal{L}$  represents a set of all POIs in LBSNs, and  $\Upsilon_i$  is an adaptive bandwidth that depicts user  $i$  activity area:  $\Upsilon_i = \max \{\|l_k - h_i\|^2\}, l_k \in \mathcal{L}_i$ . The first term of Equation (12) is the check-in distance cost between the user's home and the POI. The second indicates the contribution of user's mobility patterns. That is, the user is likely to choose nearby POIs of the POI checked in.

### IV. EXPERIMENTS

#### A. Datasets

In this paper, we use two large-scale real-world datasets: Yelp [22] and Foursquare [7], which are publicly available on the web. In the two datasets, each check-in record includes a user identity, location identity and check-in timestamp. In addition, each location is associated with its latitude, longitude and category information. The datasets also provide social links between users. We adopt the recursive grid method [35] to estimate users' home location because top- $n$  nearest neighbor friends are used in our model. We empirically filter out

those users who have fewer than 10 check-in POIs and those POIs which are visited by less than 10 users. The statistics of the datasets are shown in TABLE III. In our experiments, we divide each dataset into training set, validation set, and testing set in terms of the check-in time instead of choosing a random partition method. For each user, the earliest 70% check-ins are selected for training, the most recent 20% check-ins for testing, and the next 10% for validation.

#### B. Evaluation Metrics

We employ four standard metrics to quantify the recommendation performance: precision (Pre@ $N$ ), recall (Rec@ $N$ ), mean average precision (MAP@ $N$ ) and normalized discounted cumulative gain (NDCG@ $N$ ) [22]. Pre@ $N$  refers to the ratio of recovered POIs to the top- $N$  recommended POIs, Rec@ $N$  measures the ratio of recovered POIs to the set of visited POIs in the testing data and MAP@ $N$  defines the arithmetic mean of top- $N$  average precision over all users. NDCG@ $N$  measures the quality of ranking recommended. Its value is 0 to 1 and the higher value means better recommendation results.

#### C. Compared Methods

To illustrate the effectiveness of our recommendation framework, we compare it with the following state-of-the-art methods.

- **PMF**: PMF [36] is a matrix factorization model that employs Gaussian distribution on check-in data.
- **PFM**: PFM [34] is a probabilistic factor method that applies Poisson distribution to check-in data.
- **Geosoca**: Geosoca [8] is a personalized POI recommendation method consisting of three modules **Geo**, **So** and **Ca**, which are used to capture geographical, social and categorical correlations, respectively.
- **iGSLR**: iGSLR [5] is a well-known POI recommendation algorithm for capturing geographical and social influences.
- **GS2D**: GS2D [4] is a typical recommendation model that utilizes social and geographical influences.
- **SG**: SG [3] is a fusion approach for exploiting social and geographical influences.
- **FMFMGM**: FMFMGM [6] is a fused matrix factorization framework with the multi-center Gaussian model.
- **GE**: GE [31] is a graph-based POI recommendation method that can capture multi influences simultaneously.
- **EDHG**: EDHG [37] is a heterogeneous graph-based recommendation method that learns correlations between users and POIs.

#### D. Parameter settings

For all baselines, we select the optimal parameter configuration reported in their studies. In our experiments, parameters  $\mathbf{w}_1$  and  $\mathbf{w}_2$  are automatically learned from check-in data according to Equation (8), and the remaining parameters are tuned through cross-validation. The regularization parameter  $\lambda$  in Equation (8) is set to 0.05 and the scale parameter  $\sigma$

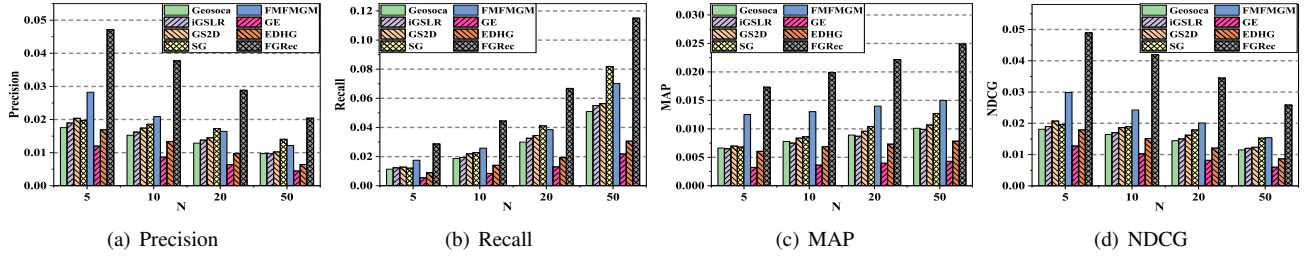


Fig. 7. Varying top- $N$  on Foursquare

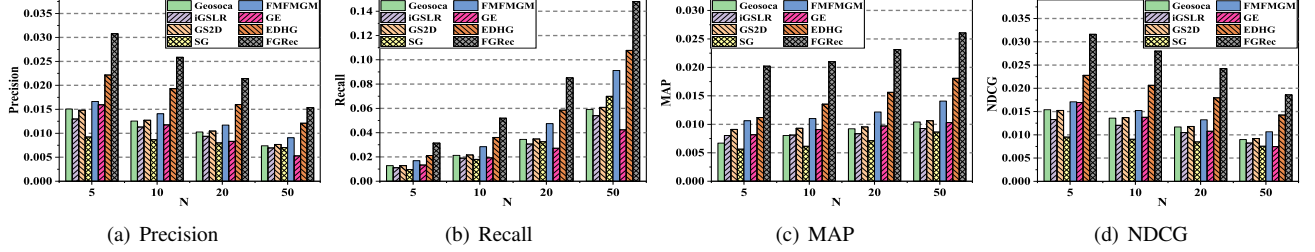


Fig. 8. Varying top- $N$  on Yelp

in Equation (6) is set to 0.1. In Yelp dataset, parameters  $\delta$  in Equation (11),  $d$  in Equation (10) and top- $n$  of neighbor friends are set to 4, 50 and 5,  $\alpha_U = \alpha_L = \alpha_Z = 40$  and  $\beta_U = \beta_L = \beta_Z = 0.2$ , in Equation (10). In Foursquare dataset, parameters  $\delta$ ,  $d$  and top- $n$  are set to 2, 40 and 10,  $\alpha_U = \alpha_L = \alpha_Z = 40$  and  $\beta_U = \beta_L = \beta_Z = 0.35$ . The effect of latent factor dimension  $d$  will be detailed later.

## E. Experimental Results

1) *Performance Comparison:* Fig. 7 and 8 depict the overall performance comparison on Foursquare and Yelp, respectively. From the results on both datasets, our framework FGRRec always achieves the best performance. Fig. 7 shows Pre@ $N$ , Rec@ $N$ , MAP@ $N$  and NDCG@ $N$  of all methods on Foursquare dataset. By observing the results, we find that FMFMGM is significantly better than other baselines. This may be attributed to the multi-center check-in distribution utilized in FMFMGM. iGSLR and GS2D have almost equivalent performance. Because the basic ideas behind them are the same, and they all use the kernel density estimation technique. GE has the worst performance. The reason is that it suffers from the check-ins data sparsity. Our framework FGRRec significantly outperforms the second best approach FMFMGM. The reasons are two-fold: one is that the accuracy of leaning user embeddings is improved in the category module; another is that the impact of users' activity ranges is captured in the geography module. Compared to Geosoca using the same information, our framework also presents absolute superiority. For instance, the Pre@5, the Rec@5, the MAP@5 and the NDCG@5 are improved by around 168%, 152%, 146% and 105%, respectively. The reasons are threefold: (1) FGRRec takes full advantage of the impact of multi-hop social friends and top- $n$  nearest neighbor friends, while Geosoca only employs 1-hop social friends' check-ins. (2) Simultaneously factorizing user-POI matrix and user-category matrix can enhance the

accuracy of learning user embeddings. When the category information is sparse, Geosoca that uses a power-law distribution is greatly affected. However, our category module can maintain good robustness because it exploits two different types of check-in sources to mutual compensate. (3) FGRRec fully takes two important geographical distances (i.e., *the distance between users' home and POIs* and *the distance among POIs*) and users' activity ranges into account in the geography module, but Geosoca only considers the distance among POIs. One obvious difference between Fig. 7 and Fig. 8 is that, FMFMGM has the best performance in Fig. 7, while EDHG is the best in Fig. 8. The reason is that less social information on Foursquare data largely affects the performance of EDHG. The rest of the results in Fig. 7 is roughly similar to that in Fig. 8.

2) *Performance of Modules in FGRRec:* Here we study the recommendation quality of the three modules. Due to limited space and similar results, we only present the results on Pre@ $N$  and Rec@ $N$  in Fig. 9. In order to demonstrate the effectiveness of three modules (JPF, CFM and GKM) of FGRRec, PMF, PFM, Ca, So and Geo are added here for comparison. Based on the results, we can see that JPF performs better than PMF and PFM. The possible reason is that the category information assists to improve the performance. PMF reports the lowest performance among all methods, because it is developed for explicit feedback data, and not suitable for the implicit feedback data. By comparing JPF to Ca, we can see the performance of JPF is better than that of Ca with an absolute advantage. One possible explanation is that JPF enhances the learning on the user embeddings by jointly factorizing user-POI and user-category matrices, while Ca only relies on the categorical popularity. Once category data are sparse, its performance will be unstable.

We further observe that CFM gives the best performance among the three modules in FGRRec. This tells us that infor-

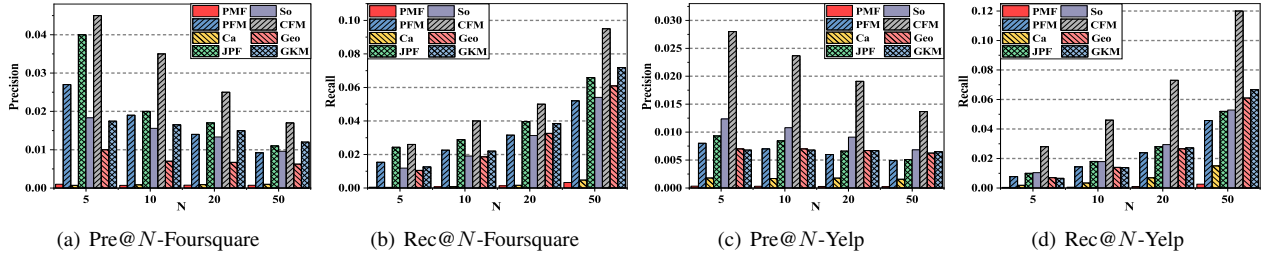


Fig. 9. Recommendation performance of three modules (JPF, CFM and GKM) in FGRec

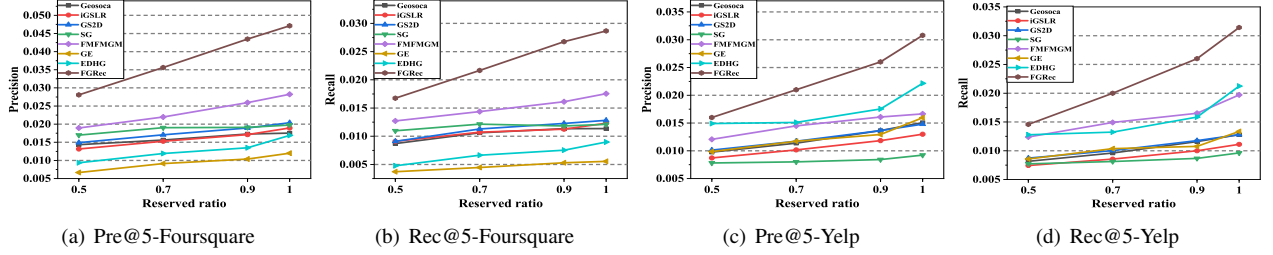


Fig. 10. Impact of data sparsity

mation of social friends and neighbor friends affects users' check-in choices and plays a key role in enhancing the quality of POI recommendation. This finding is consistent with [7]. The performance of So is still not as good as CFM. One possible reason is that CFM fully exploits the information of multi-hop social friends and top- $n$  nearest neighbor friends, while So model only considers 1-hop friends. We also find that the performances of GKM are slightly superior to Geo. One possible explanation is that GKM integrates the distance between user homes and POIs and the distance among POIs into a unified model. The interaction of two types of distance in GKM assists to make more accurate recommendation. Furthermore, GKM exhibits the worst recommendation quality across three modules of our framework. We think that its instability results in the poor performance. This situation is reasonable in reality, because the distance among POIs could exceed the ranges of users' activity.

3) *Impact of Data Sparsity*: In order to verify the effectiveness of our model on the sparse data, we randomly reserve  $x\%$  ( $x = 50, 70, 90, 100$ ) of check-ins from each user's visited records. This way generates the check-in data with different sparsity. The smaller the reserved ratio  $x$  is, the sparser the check-in data are. Fig. 10 shows the Pre@5 and the Rec@5 of all the methods on Yelp and Foursquare data under different data sparsities. From Fig. 10, we can observe that the performance of all methods is increasing with the increase of the density of check-in data (from left to right). This is because all algorithms rely on the number of users' check-ins. When the data become sparser, FGRec always performs better than the second best result given by FMFMGM or EDHG in terms of Pre@5 and the Rec@5. This shows that our framework is very effective in compensating the check-in data sparsity by taking full advantage of the information of social and geographical. Thus, FGRec shows greater strength than all

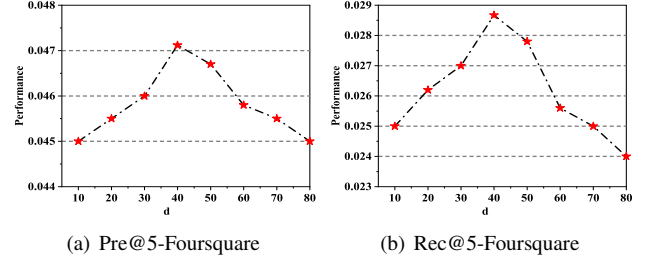


Fig. 11. The influence of latent factor dimension  $d$  on Foursquare data

state-of-the-art methods under various data sparsity scenarios.

4) *Study of Influence of Latent Factor Dimension  $d$* : Here, we study the influence of  $d$ . In our experiment, we set  $d$  to 10, 20, 30, 40, 50, 60, 70, and 80, respectively. Due to limited space, we only show the recommendation performance of FGRec on Foursquare dataset. Fig.11 shows that the recommended quality for different values of  $d$ . Based on the results, we can find that the performances in the Pre@5 and Rec@5 have similar behaviour with the varying value of  $d$ . It is observed that the performance increases with the increase of the  $d$  at the beginning, then hits the highest recommended quality when  $d = 40$ , and eventually tends to decline. So we finally choose the optimal parameter  $d=40$ .

## V. CONCLUSIONS

In this paper, we propose a fine-grained framework (FGRec) for POI recommendation, which addresses the challenge arising from the user-POI matrix sparsity. In our framework, we model the user's preference for an unvisited POI by simultaneously considering three types of intrinsic influences (i.e., social, categorical and geographical influences): (1) We propose the novel collective friends model that adaptively



captures the impact of multi-hop social friends and top- $n$  nearest neighbor friends. (2) We design the joint Poisson factor model that simultaneously factorizes user-POI matrix and user-category matrix by sharing the user embedding space to capture the categorical influence. (3) To acquire the geographical influence, we present the personalized Gaussian kernel model by fully leveraging users' activity ranges from the perspective of home-POI and POI-POI. Extensive experimental results on Foursquare and Yelp clearly show that FGRec significantly outperforms the state-of-the-art methods in terms of various metrics.

## ACKNOWLEDGMENT

This work is supported by the National Key Research and Development Program of China, and National Natural Science Foundation of China (No. U163620068).

## REFERENCES

- [1] D. Lian, Y. Ge, F. Zhang, N. J. Yuan, X. Xie, T. Zhou, and Y. Rui, "Scalable content-aware collaborative filtering for location recommendation," *IEEE TKDE*, vol. 30, no. 6, pp. 1122–1135, 2018.
- [2] S. Zhao, T. Zhao, I. King, and M. R. Lyu, "Geo-teaser: Geo-temporal sequential embedding rank for point-of-interest recommendation," in *WWW*, 2017, pp. 153–162.
- [3] M. Ye, P. Yin, W.-C. Lee, and D.-L. Lee, "Exploiting geographical influence for collaborative point-of-interest recommendation," in *ACM SIGIR*, 2011, pp. 325–334.
- [4] J.-D. Zhang, C.-Y. Chow, and Y. Li, "Lore: Exploiting sequential influence for location recommendations," in *ACM SIGSPATIAL*, 2014, pp. 103–112.
- [5] J.-D. Zhang and C.-Y. Chow, "igslr: Personalized geo-social location recommendation: A kernel density estimation approach," in *ACM SIGSPATIAL*, 2013, pp. 334–343.
- [6] C. Cheng, H. Yang, I. King, and M. R. Lyu, "A unified point-of-interest recommendation framework in location-based social networks," *ACM TIST*, vol. 8, no. 1, pp. 10:1–10:21, 2016.
- [7] H. Li, Y. Ge, R. Hong, and H. Zhu, "Point-of-interest recommendations: Learning potential check-ins from friends," in *ACM KDD*, 2016, pp. 975–984.
- [8] J.-D. Zhang and C.-Y. Chow, "Geosoca: Exploiting geographical, social and categorical correlations for point-of-interest recommendations," in *ACM SIGIR*, 2015, pp. 443–452.
- [9] X. Liu, Y. Liu, K. Aberer, and C. Miao, "Personalized point-of-interest recommendation by mining users' preference transition," in *ACM CIKM*, 2013, pp. 733–738.
- [10] J. He, X. Li, and L. Liao, "Category-aware next point-of-interest recommendation via listwise bayesian personalized ranking," in *IJCAI*, 2017, pp. 1837–1843.
- [11] Y. Liu, W. Wei, A. Sun, and C. Miao, "Exploiting geographical neighborhood characteristics for location recommendation," in *ACM CIKM*, 2014, pp. 739–748.
- [12] J. Manotumruksa, C. Macdonald, and I. Ounis, "A personalised ranking framework with multiple sampling criteria for venue recommendation," in *ACM CIKM*, 2017, pp. 1469–1478.
- [13] X. Li, G. Cong, X.-L. Li, T.-A. N. Pham, and S. Krishnaswamy, "Rank-geofm: A ranking based geographical factorization method for point of interest recommendation," in *ACM SIGIR*, 2015, pp. 433–442.
- [14] F. Yuan, G. Guo, J. Jose, L. Chen, and H. Yu, "Joint geo-spatial preference and pairwise ranking for point-of-interest recommendation," in *IEEE ICTAI*, 2016, pp. 46–53.
- [15] H. Yin, B. Cui, X. Zhou, W. Wang, Z. Huang, and S. Sadiq, "Joint modeling of user check-in behaviors for real-time point-of-interest recommendation," *ACM TOIS*, vol. 35, no. 2, pp. 11:1–11:44, 2016.
- [16] B. Chang, Y. Park, D. Park, S. Kim, and J. Kang, "Content-aware hierarchical point-of-interest embedding model for successive POI recommendation," in *IJCAI*, 2018, pp. 3301–3307.
- [17] H. Li, R. Hong, Z. Wu, and Y. Ge, "A spatial-temporal probabilistic matrix factorization model for point-of-interest recommendation," in *SDM*, 2016, pp. 117–125.
- [18] S. Wang, Y. Wang, J. Tang, K. Shu, S. Ranganath, and H. Liu, "What your images reveal: Exploiting visual contents for point-of-interest recommendation," in *WWW*, 2017, pp. 391–400.
- [19] J. He, X. Li, L. Liao, D. Song, and W. K. Cheung, "Inferring a personalized next point-of-interest recommendation model with latent behavior patterns," in *AAAI*, 2016, pp. 137–143.
- [20] D. Lian, C. Zhao, X. Xie, G. Sun, E. Chen, and Y. Rui, "Geomf: joint geographical modeling and matrix factorization for point-of-interest recommendation," in *ACM KDD*, 2014, pp. 831–840.
- [21] D. Lian, K. Zheng, Y. Ge, L. Cao, E. Chen, and X. Xie, "Geomf++: Scalable location recommendation via joint geographical modeling and matrix factorization," *ACM TOIS*, vol. 36, no. 3, pp. 33:1–33:29, 2018.
- [22] Y. Liu, T.-A. N. Pham, G. Cong, and Q. Yuan, "An experimental evaluation of point-of-interest recommendation in location-based social networks," *VLDB*, pp. 1010–1021, 2017.
- [23] C. Yang, L. Bai, C. Zhang, Q. Yuan, and J. Han, "Bridging collaborative filtering and semi-supervised learning: A neural approach for poi recommendation," in *ACM KDD*, 2017, pp. 1245–1254.
- [24] W. Zhang and J. Wang, "Location and time aware social collaborative retrieval for new successive point-of-interest recommendation," in *ACM CIKM*, 2015, pp. 1221–1230.
- [25] R. Li, Y. Shen, and Y. Zhu, "Next point-of-interest recommendation with temporal and multi-level context attention," in *IEEE ICDM*, 2018, pp. 1110–1115.
- [26] W. Liu, Z.-J. Wang, B. Yao, and J. Yin, "Geo-alm: Poi recommendation by fusing geographical information and adversarial learning mechanism," in *IJCAI*, 7 2019, pp. 1807–1813.
- [27] C. Ma, Y. Zhang, Q. Wang, and X. Liu, "Point-of-interest recommendation: Exploiting self-attentive autoencoders with neighbor-aware influence," in *ACM CIKM*, 2018, pp. 697–706.
- [28] H. Wang, H. Shen, W. Ouyang, and X. Cheng, "Exploiting poi-specific geographical influence for point-of-interest recommendation," in *IJCAI*, 2018, pp. 3877–3883.
- [29] H. Yin, W. Wang, H. Wang, L. Chen, and X. Zhou, "Spatial-aware hierarchical collaborative deep learning for poi recommendation," *IEEE TKDE*, vol. 29, no. 11, pp. 2537–2551, 2017.
- [30] W. Wang, H. Yin, X. Du, Q. V. H. Nguyen, and X. Zhou, "TPM: A temporal personalized model for spatial item recommendation," *ACM TIST*, vol. 9, no. 6, pp. 61:1–61:25, 2018.
- [31] M. Xie, H. Yin, H. Wang, F. Xu, W. Chen, and S. Wang, "Learning graph-based poi embedding for location-based recommendation," in *ACM CIKM*, 2016, pp. 15–24.
- [32] T. Qian, B. Liu, Q. V. H. Nguyen, and H. Yin, "Spatiotemporal representation learning for translation-based poi recommendation," *ACM TOIS*, vol. 37, no. 2, pp. 18:1–18:24, 2019.
- [33] L. F. R. Ribeiro, P. H. P. Saverese, and D. R. Figueiredo, "Struc2vec: Learning node representations from structural identity," in *ACM KDD*, 2017, pp. 385–394.
- [34] H. Ma, C. Liu, I. King, and M. R. Lyu, "Probabilistic factor models for web site recommendation," in *ACM SIGIR*, 2011, pp. 265–274.
- [35] Z. Cheng, J. Caverlee, K. Lee, and D. Z. Sui, "Exploring millions of footprints in location sharing services," in *ICWSM*, 2011.
- [36] R. Salakhutdinov and A. Mnih, "Probabilistic matrix factorization," in *NIPS*, 2007, pp. 1257–1264.
- [37] M. Hang, I. Pytlarz, and J. Neville, "Exploring student check-in behavior for improved point-of-interest prediction," in *KDD*, 2018, pp. 321–330.