

Exploiting User Preference and Mobile Peer Influence for Human Mobility Annotation

RENJUN HU, Beihang University & Beijing Advanced Innovation Center for Big Data and Brain Computing

YANCHI LIU, NEC Labs

YANYAN LI and JINGBO ZHOU, National Engineering Laboratory of Deep Learning Technology and Application

SHUAI MA, Beihang University & Beijing Advanced Innovation Center for Big Data and Brain Computing

HUI XIONG, Rutgers University

Human mobility annotation aims to assign mobility records the corresponding visiting Point-of-Interests (POIs). It is one of the most fundamental problems for understanding human mobile behaviors. In literature, many efforts have been devoted to annotating mobility records in a pointwise or trajectory-wise manner. However, the user preference factor is not fully explored and, worse still, the mobile peer influence factor has never been integrated. To this end, in this article, we propose a novel framework, named JEPPI, to jointly exploit user preference and mobile peer influence to tackle the problem. In our JEPPI, we first unify the two distinct factors in a behavior-driven user-POI graph. This graph enables us to model user preference with user-POI visiting relationships, and model two types of mobile peer influence with co-location and co-visiting peer relationships, respectively. Moreover, we devise an equivalence-emphasizing metric to reduce redundancy in the second-order co-visiting peer influence. In addition, a mutual augmentation learning approach is proposed to preserve the latent structures of various factors exploited. Notably, our learning approach preserves all factors in a shared representation space such that user preference is learned with mobile peer influence being considered at the same time, and vice versa. In this way, the different factors are mutually augmented and semantically integrated to enhance human mobility annotation. Finally, using two large-scale real-world datasets, we conduct extensive experiments to demonstrate the superiority of our approach compared with the state-of-the-art annotation methods.

CCS Concepts: • **Information systems** → **Mobile information processing systems**;

Additional Key Words and Phrases: Human mobility annotation, Point-of-Interest, network embedding, mobile analysis

This work is supported in part by the National Key R&D Program of China (Grant: 2018AAA0102301), NSFC (Grant: 71531001, 61925203, U1636210, and 61421003), and Beijing Advanced Innovation Center for Big Data and Brain Computing. Authors' addresses: R. Hu and S. Ma, SKLSDE Lab, Beihang University & Beijing Advanced Innovation Center for Big Data and Brain Computing, No. 37 Xueyuan Road, Haidian District, Beijing, China, 100191; emails: {hurenjun, mashuai}@buaa.edu.cn; Y. Liu, NEC Labs, 4 Independence Way, Princeton, NJ, 08540; email: yanchi@nec-labs.com; Y. Li and J. Zhou, Business Intelligence Lab, Baidu Research, National Engineering Laboratory of Deep Learning Technology and Application, No. 10 Xibeiwang East Road, Haidian District, Beijing, China, 100193; emails: {liyanyanliyanyan, zhoujingbo}@baidu.com; H. Xiong, Rutgers University, 1 Washington Park, Newark, NJ, 07102; email: hxiong@rutgers.edu. Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

© 2020 Association for Computing Machinery.

1556-4681/2020/09-ART67 \$15.00

<https://doi.org/10.1145/3406600>

ACM Reference format:

Renjun Hu, Yanchi Liu, Yanyan Li, Jingbo Zhou, Shuai Ma, and Hui Xiong. 2020. Exploiting User Preference and Mobile Peer Influence for Human Mobility Annotation. *ACM Trans. Knowl. Discov. Data* 14, 6, Article 67 (September 2020), 18 pages.
<https://doi.org/10.1145/3406600>

1 INTRODUCTION

A fundamental problem in human mobile behavior understanding is mobility annotation, which assigns mobility records the corresponding visiting Point-of-Interests (POIs). Figure 1 illustrates such an example, where mobility records of a user are represented as a series of spatiotemporal points and we aim to assign POIs to records that are produced when the user is visiting these POIs. Since POIs carry rich semantics, identifying the visiting POIs is crucial to bridge the gap between the semantic-free mobility data and the deep understanding of human behaviors. Afterward, the semantic information of POIs can facilitate a lot of mobility- and POI-related applications, ranging from visit purpose detection [25] and demand modeling [16] to POI and trip recommendation [2].

Recently, human mobility annotation has attracted attention from both academia and industry. The industrial solutions require much user participation and are usually insufficient. Existing solutions in literature either ignore [1, 19, 21] or could not fully explore [28, 31] the personalized preference for visiting POIs. In addition, all these approaches do not consider the effect of user interactions in mobile behaviors, *i.e.*, mobile peer influence. Please note that the mobile peer influence refers to the true influence on real life POI-visiting behaviors between people, and is different from the peer influence in cyberspace like social media.

On the other hand, incorporating user preference and mobile peer influence has great potential for the task. Our first observation is that user preference has complex and significant impacts on choosing POIs. These preference patterns usually remain latent and hard to be captured by explicit statistics or rules developed in [19, 28]. More importantly, modeling mobile peer influence between users can very likely provide extra clues as POI-visiting behaviors have intrinsic social attributes, *e.g.*, people might visit POIs together with or recommended by others. To illustrate our insight, we present a real-world example in Figure 2, in which four mobility records of user u_0 are to be annotated with POIs p_0 – p_3 . The advantages of combining the two factors are three-fold:

- (i) The visiting relationship (u_0, p_1) is the clue for annotating with p_1 and the peer relationships between u_0 and u_1/u_2 consolidate the evidence given the preference of u_1 and u_2 for p_1 .
- (ii) We do not have direct clues for p_2 and p_3 . However, the peer relationships between u_0 and u_3 – u_5 can provide some evidence if user preference and mobile peer influence are jointly considered.
- (iii) It is difficult to annotate the mobility record with p_0 . In this case, a promising choice is to infer a POI according to the preference of other users “similar” to u_0 .

To this end, in this article, we propose a novel framework, named JEPPI, to Jointly Exploit user Preference and mobile Peer Influence for human mobility annotation. Simultaneously exploiting these two factors in an effective way is challenging both because user preference usually differs from person to person and because peer influence remains unexplored. To unify the two distinct factors, we construct a behavior-driven user-POI graph from mobility data (see Figure 2 for an example). More specifically, we model user preference with user-POI visiting relationships and model two types of mobile peer influence with co-location and co-visiting peer relationships, respectively. Note that co-visiting mobile peer influence is essentially a kind of second-order

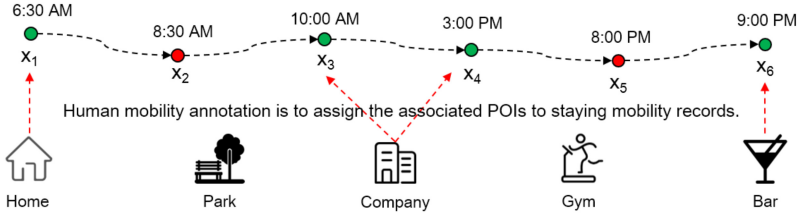


Fig. 1. An example of human mobility annotation, where green (x_1 , x_3 , x_4 , and x_6) and red (x_2 and x_5) points represent mobility records produced when the user is at staying and moving statuses, respectively.

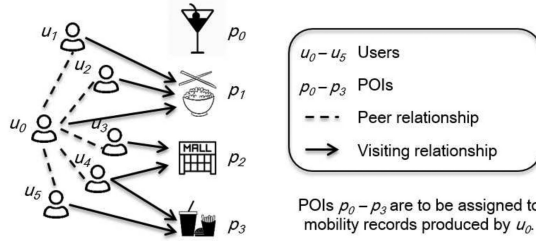


Fig. 2. Illustration of exploiting user preference and mobile peer influence for human mobility annotation.

proximity. To reduce redundancy, we devise an equivalence-emphasizing metric which emphasizes peer influence between users that have equivalent POI-visiting patterns. We then develop a mutual augmentation learning approach which assigns latent representations to users and POIs, and optimizes the shared representations iteratively and sequentially with the two factors. That is, we preserve user preference with mobile peer influence being considered, and vice versa. With this unique design, these factors are mutually augmented and semantically integrated. We finally combine the representations and a simple yet effective Bayesian candidate generation method for human mobility annotation. To summarize, we make the following contributions in this article:

- We identify an important mobile peer influence factor in human mobile behaviors and inventively propose to jointly exploit user preference and mobile peer influence for the fundamental and practical human mobility annotation task.
- We propose the JEPPI framework to tackle the problem. It utilizes mutual augmentation learning on a behavior-driven user-POI graph to organically integrate user preference and two types of mobile peer influence for boosting annotation effectiveness.
- We conduct extensive experiments on two large real-world datasets to demonstrate the superiority of JEPPI, compared with state-of-the-art competitors.

The remaining of the article is organized as follows. We discuss related work in Section 2 and formalize our problem in Section 3. The details of our JEPPI framework are described in Section 4, followed by experiments in Section 5 and conclusion in Section 6.

2 RELATED WORK

Mobile behavior understanding. Mobile analytics can greatly improve our understanding of human mobility patterns, such as sequential pattern [20] and periodic pattern [36]. Moreover, it is also helpful to obtain deep insights into urban dynamics. For instance, [11, 23, 26] study mobility along with social ties from a spatiotemporal perspective. People's daily need information could be identified from mobility data to enable POI demand modeling [16]. Hu et al. [9] develop a

user decision profiling model to identify the key factors behind people's decisions on choosing POIs, which directly facilitates to understand human mobile behaviors. Enriching mobility data with semantics is also among popular, e.g., [7, 25, 29] and ours. Wu et al. [29] study semantic annotation such that each mobility record is enhanced with semantics retrieved from social media. Wang et al. [25] detect trip purpose of taxi trajectories. User interest discovery is explored in [7], which also utilizes the rich semantics contained in user-generated contents.

Different from the above studies, our work aims to identify POIs associated with mobility records. The mobility data is thus enriched with semantics embedded in POIs, which can facilitate a wide range of mobility- and POI-related applications.

Human mobility annotation is an important step towards better understanding of mobility, e.g., a rich mobility dataset has been constructed recently [17], which contains ground-truth POIs as labels. Location-based services (such as Foursquare¹) and crowdsourcing platforms [34] provide a partial annotation solution, by allowing people to share their locations along with the visiting POIs. The problem has also been investigated in literature either in a pointwise manner such that each mobility record is processed independently [1, 19, 21], or in a trajectory-wise manner that considers a trajectory produced by a single user at one time [28, 31]. These approaches have already exploited user preference in a number of ways. For instance, Shaw et al. [19] design a set of features for candidate POIs, which include the number of historical visits by a user to a POI for modeling preference, and annotate mobility records through learning-to-rank. Markov random field could also capture preference by enforcing consistency in individuals' POI-visiting behaviors [28]. Besides preference, Hidden Markov model is utilized to capture transitional patterns [31]. However, mobile peer influence that we consider remains ignored in all existing methods. Moreover, we propose a mutual augmentation learning approach to jointly exploiting user preference and mobile peer influence, which is totally different from the previously-used models.

On the other hand, human mobility annotation bears similarity with POI or location recommendation that has been extensively studied [5, 13, 15, 30, 32, 37]. Note that only [13, 32] consider peer influence. However, the peer influence is drawn from explicit social relationships and is learned separately, whereas we collect peer relationships from user mobile behaviors and jointly learn user preference and mobile peer influence for mutual augmentation.

Network embedding has become an emerging technique for network-centric analysis [3]. Typical embedding models are mainly based on random walks, e.g., node2vec [6], Graphgan [24], and metatpath2vec [4], and based on pairwise proximities, e.g., LINE [22]. Notably, Qiu et al. [18] unifies a number of approaches in matrix factorization. Embedding models are also proposed to preserve advanced network information, e.g., asymmetric proximity [38], community structures [27], and dynamism [39]. This technique has been successfully applied in various tasks such as anomaly detection [8] and transportation recommendation [14]. Our embedding model is essentially based on LINE [22] for proximity preserving, while it adopts a mutual augmentation design.

3 PROBLEM DEFINITION

In this section, we introduce the related concepts and formalize our problem. Let a spatiotemporal point $x = (l, t)$ be a pair of location l (e.g., latitude and longitude) and timestamp t .

Definition 1 (Mobility Records). The mobility records $\mathcal{M}_u$ of user u are a series of ordered spatiotemporal points generated by user u , i.e., $\mathcal{M}_u = [x_1, x_2, \dots, x_L]$, where L is the length of $\mathcal{M}_u$ and the timestamps satisfy $t_1 < t_2 < \dots < t_L$.

¹<https://foursquare.com/>.

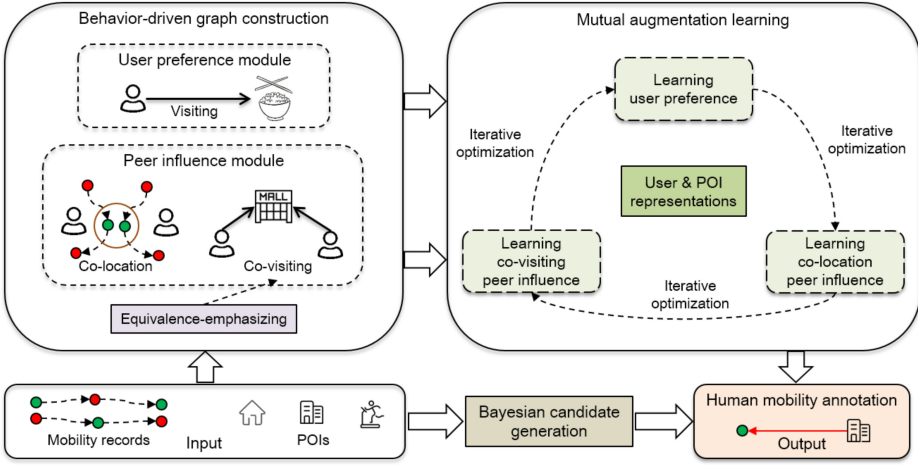


Fig. 3. Overview of our JEPPI framework.

Basically, human mobility annotation aims to assign POIs to mobility records. In some cases, *e.g.*, on Foursquare, each record is generated when people visit POIs and can be annotated directly. While in other scenarios, *e.g.*, with raw Global Positioning System data, mobility records include moving spatiotemporal points that are generated when people move from one place to another. Obviously, moving points do not correspond to any meaningful POIs and should be filtered out before annotation. This can be achieved by detecting stays [10, 12] in mobility records and only annotating records included in stays. Each stay represents a continuous subseries of spatiotemporal points that are both spatially- and temporally-close and span for a period of time. Following [10], stay detection for a single user can be done in linear, *i.e.*, $O(L)$, time by enumerating a beginning point (say x_i) and finding the longest subseries $[x_i, \dots, x_j]$ that satisfies the stay conditions.

Hereinafter, without loss of generality, we assume that each mobility record is associated with a POI. For instance, the mobility record x_1 in Figure 1 is generated when u is at home, x_3 and x_4 are produced when u works at the company, and, finally, x_6 is associated with a bar. These POIs, as well as the embedded semantics, are informative for understanding people's mobility patterns. In practice, the POIs of a fraction of mobility records are easy to obtain, *e.g.*, people may occasionally check-in at certain POIs on social media. Based on the known POIs, we are looking forward to annotating the rest records. Formally, our problem is as below.

PROBLEM 1 (HUMAN MOBILITY ANNOTATION). *Given a set of mobility records $\mathcal{M}_u$ of users $u \in \mathcal{U}$ and a set $\mathcal{P}$ of POIs, human mobility annotation aims to annotate each record, that has unknown visiting POIs, with the top- k most likely visiting POIs $\mathcal{P}_s \subset \mathcal{P}$.*

4 METHODOLOGY

To tackle the problem, we propose a JEPPI framework which integrates user preference and mobile peer influence with mutual augmentation learning.

4.1 Framework Overview

The overview of JEPPI is illustrated in Figure 3. It consists of the following five main components:

- *The input* of JEPPI is a set of mobility records and a set of POIs.
- *Behavior-driven graph construction* derives user preference and mobile peer influence from users' mobile visiting, co-location, and co-visiting behaviors, and encodes them in a

- user-POI graph. An equivalence-emphasizing metric is devised to quantify co-visiting peer influence. This component enables to exploit the distinct factors in a unified way.
- *Mutual augmentation learning* is the core of JEPPI which integrates the various factors into a shared representation space. It assigns a latent representation vector to each user and each POI, and updates the representations sequentially and iteratively with user preference and two types of mobile peer influence.
 - *Bayesian candidate generation* fuses the spatial influence and POI prior with a simple Bayesian method to identify candidate POIs for mobility records. In this way, these basic factors are further incorporated for human mobility annotation.
 - *Human mobility annotation* combines the learned representations and candidate POIs to obtain the output, *i.e.*, associated POIs of mobility records.

In the following, we introduce these components in details.

4.2 Behavior-driven Graph Construction

We construct a heterogeneous user-POI graph $G = (\mathcal{U} \cup \mathcal{P}, \mathcal{E}_v \cup \mathcal{E}_{col} \cup \mathcal{E}_{cov}, w)$. In this graph, each node denotes either a user or a POI. The visiting relationships $\mathcal{E}_v \subset \mathcal{U} \times \mathcal{P}$ encode users' preference for POIs, while $\mathcal{E}_{col}, \mathcal{E}_{cov} \subset \mathcal{U} \times \mathcal{U}$ are two sets of peer relationships reflecting mobile peer influence between people. Finally, w is the edge weight function.

More specifically, we construct G in a behavior-driven manner. Recall that a fraction of mobility records are annotated with POIs in advance. Relationships in $\mathcal{E}_v$ are then derived from annotated records. Suppose that a record of user u is associated with POI p . We then form an edge $(u, p) \in \mathcal{E}_v$ to record the fact that user u has visited POI p and, thus, has preference for p . Intuitively, the more frequently u visits p , the stronger the preference is. Hence edge weight $w(u, p)$ is introduced to distinguish the various levels of preference, and $w(u, p)$ is set to the number of visits from user u to POI p . We note that, although visiting behaviors are usually sparse, it is very likely that a large proportion of weights $w(u, p)$ will exceed 1 due to the periodicity of human mobility [36].

While numerous peer relationships are formed in online social media, they do not necessarily lead to peer influence in real life POI-visiting behaviors. On the other hand, mobility data itself is a good source for mining mobile peer influence. Thus, we also turn to directly collect peer relationships from user mobile behaviors. Here we consider two types of behaviors reflecting mobile peer influence, *i.e.*, co-location and co-visiting, and derive $\mathcal{E}_{col}$ and $\mathcal{E}_{cov}$, respectively. Basically, there is an edge $(u, u') \in \mathcal{E}_{col}/\mathcal{E}_{cov}$ if both u and u' have located at the same places or visited the same POIs. It is worth noting that co-visiting peer influence is essentially a kind of second-order proximity for visiting preference. To reduce redundancy, we devise an equivalence-emphasizing metric which assigns a high peer influence weight to two users when they have equivalent/similar POI-visiting patterns. We put off the details of peer relationship construction and the metric to Section 4.3.2 since they are deeply coupled with mobile peer influence learning.

4.3 Mutual Augmentation Learning

After graph construction, the various hidden factors are integrated via mutual augmentation learning. To be specific, all users and POIs are embedded in a shared d -dimensional space $\mathbb{R}^d$. Figure 4 presents an example of user and POI representations ($d = 2$). Each representation can be regarded as a point in the hidden space. Recall in Figure 1 that the user (say u_1) has mobility records associated with POIs *Home*, *Company*, and *Bar*. Correspondingly, u_1 and the three POIs are placed together in the representation space, as bounded by a dashed circle in Figure 4. In addition, assuming users u_1 and u_2 have mobile peer influence on each other, they are also close in the space. User and POI representations, *i.e.*, their positions in the space, are not explicitly interpretable. However,

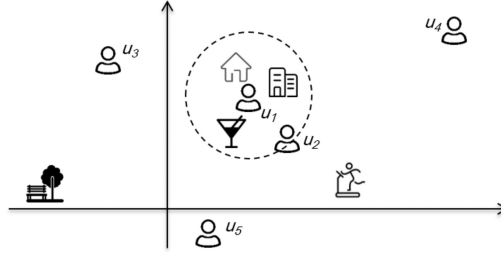


Fig. 4. Illustration of user and POI representations in a two-dimensional hidden space.

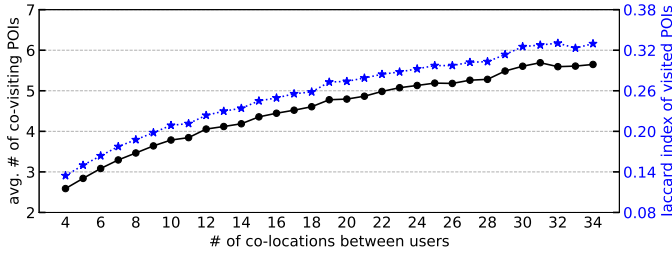


Fig. 5. The average number of co-visiting POIs and the corresponding Jaccard index increase with the number of co-locations between users on our BEIJING data.

their pairwise proximity, *e.g.*, inner product, captures certain information. That is, the proximity of user representations indicates the strength of peer influence, while the proximity of user and POI representations indicates the strength of preference.

4.3.1 Learning User Preference. We learn user preference by preserving the structure of $\mathcal{E}_v$. Formally, let $\mathbf{u}$ and $\mathbf{p}$ denote the representation vectors of user u and POI p , respectively. For each $(u, p) \in \mathcal{E}_v$, we minimize the following KL-divergence and negative sampling-based objective [22]:

$$O_1(u, p) = -\log \sigma(\mathbf{u} \cdot \mathbf{p}) - \sum_{i=1}^K \mathbb{E}_{p' \sim D_p} \log \sigma(-\mathbf{u} \cdot \mathbf{p}'), \quad (1)$$

where $\sigma(x) = 1/(1 + e^{-x})$ is the sigmoid function, K is the number of negative samples, and D_p is a discrete distribution for generating random POI samples p' . We empirically evaluate the two popular sampling distributions, *i.e.*, uniform [38] and node-degree based [22], and find that the uniform distribution works better for our task. Hence, we adopt $D_p(p) = 1/|\mathcal{P}|$.

In practice, the model is usually trained with stochastic gradient descent. We iteratively draw a subset of both positive samples $(u, p) \in \mathcal{E}_v$ and random negative samples (u, p') , and update the representations by minimizing Equation (1). By drawing positive (u, p) according to $w(u, p)$, the users and POIs having more visiting records are placed closer.

4.3.2 Learning Mobile Peer Influence. We incorporate two types of mobile peer influence in G . The co-location peer influence assumes that a pair of users u and u' have influence on each other if u and u' have co-locations. The rationale is that there usually exists a positive correlation between co-locations and commonness of POI-visiting behaviors. We have also verified this with statistics from our data reported in Figure 5. Specifically, let $\mathcal{P}_u$ and $\mathcal{P}_{u'}$ denote the sets of POIs visited by users u and u' , respectively. We use the number of co-visiting POIs (*i.e.*, $|\mathcal{P}_u \cap \mathcal{P}_{u'}|$), as well as the Jaccard index (*i.e.*, $|\mathcal{P}_u \cap \mathcal{P}_{u'}|/|\mathcal{P}_u \cup \mathcal{P}_{u'}|$), to evaluate the commonness of POI-visiting behaviors. Figure 5 shows that both the two metrics increase with the number of co-locations between users.

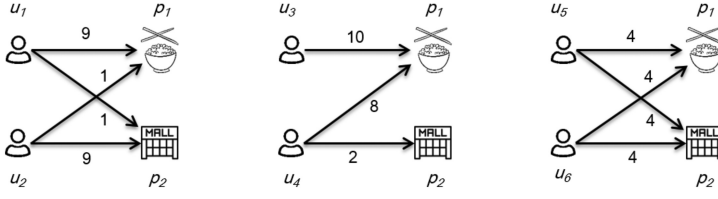


Fig. 6. Toy examples for co-visiting peer influence.

The increased commonness ensures that the knowledge of u can help to annotate the mobility records of u' if u and u' have a proper number of co-locations.

Formally, we say that users u and u' have a co-location if there exist $x_i = (l_i, t_i)$ of $\mathcal{M}_u$ and $x'_j = (l'_j, t'_j)$ of $\mathcal{M}_{u'}$ such that l_i and l'_j are spatially close (e.g., within 100 meters) and t_i and t'_j are temporally close (e.g., within 1 hour). We then form an edge $(u, u') \in \mathcal{E}_{col}$ to record the co-location peer relationship between u and u' . Similarly, edge weight $w(u, u')$, i.e., the number of co-locations of u and u' , is introduced to distinguish the different strength of influence.

Modeling co-location peer influence alone may suffer from data sparsity. On the other hand, a pair of users also bear similarity if they have visited the same POIs. The above two observations motivate us to further exploit co-visiting peer influence inspired by co-visiting behaviors. When defining the new type of mobile peer influence, we consider the following. (i) It might be misleading to evaluate the strength of co-visiting peer influence with the numbers of co-visiting POIs. For instance, in Figure 6, the influence between users u_3 and u_4 is obviously stronger than u_1 and u_2 , even though u_3 and u_4 have less number of co-visiting POIs. Note that u_3 and u_4 share a strong common interest to POI p_1 while u_1 and u_2 do not have such a common interest. (ii) The second-order co-visiting peer influence should avoid high redundancy with the first-order user preference. Again in Figure 6, the proximity between u_3 and u_4 has already been well preserved by learning user preference with $w(u_3, p_1)$ and $w(u_4, p_1)$. In contrast, the one between u_5 and u_6 should be further emphasized via co-visiting peer influence.

To cope with the above requirements, we devise a novel equivalence-emphasizing metric to evaluate the strength of co-visiting peer influence. Its main idea is to assign a high weight to a pair of users if they have co-visited POIs for a number of times and their roles (e.g., regular or rare visitor) w.r.t. these POIs remain largely equivalent. More specifically, we include an edge $(u, u') \in \mathcal{E}_{cov}$ if u and u' have co-visited at least one POI. The weight $w(u, u')^2$ is determined by:

$$w(u, u') = \sum_p \min\{w(u, p), w(u', p)\} \cdot \sum_{p, w(u, p) > 0, w(u', p) > 0} \frac{\min\left\{\frac{w(u, p)}{W(u)}, \frac{w(u', p)}{W(u')}\right\}}{\max\left\{\frac{w(u, p)}{W(u)}, \frac{w(u', p)}{W(u')}\right\}}. \quad (2)$$

Recall that $w(u, p)$ is the number of visits from user u to POI p , and $W(u) = \sum_p w(u, p)$. The first sum operator in Equation (2) counts the number of co-visiting times between two users and the second sum operator quantifies the equivalence of two users in terms of POI-visiting. Here $w(u, p)/W(u)$ can be interpreted as the role of u to p : ~ 1 for regular and ~ 0 for rare visitor. When the roles of u and u' to p are equivalent or similar, POI p has a large contribution to the equivalence, regardless of the values of $w(u, p)$ and $w(u', p)$. With Equation (2), we have $w(u_1, u_2) = 0.44$, $w(u_3, u_4) = 6.4$ and $w(u_5, u_6) = 8$ in Figure 6, which satisfies our requirements.

²To make our notation simple, we use $w(u, u')$ to denote the weights of both types of peer relationships. It should be easy to distinguish between them from the context.

ALGORITHM 1: Mutual augmentation learning

Input: Graph G , number d of dimensions, trade-off parameters λ_1, λ_2 , learning rate α , batch size B , number K of negative samples, and iteration threshold N_i ;

Output: User and POI representation vectors $\mathbf{u}$ and $\mathbf{p}$;

```

1 Initialize entries of  $\mathbf{u}$  and  $\mathbf{p}$  from  $\text{unif}(-0.5/d, 0.5/d)$ ;
2  $iter \leftarrow 1$ ;
3 repeat
4   Sample a batch of  $B$  visiting relationships  $(u, p)$  from  $\mathcal{E}_v$ ;
5   foreach sampled  $(u, p)$  do
6      $\mathbf{u} \leftarrow \mathbf{u} - (1 - \frac{iter}{N_i})\alpha(\sigma(\mathbf{u} \cdot \mathbf{p}) - 1)\mathbf{p}$ ;  $\mathbf{p} \leftarrow \mathbf{p} - (1 - \frac{iter}{N_i})\alpha(\sigma(\mathbf{u} \cdot \mathbf{p}) - 1)\mathbf{u}$ ;
7     Sample  $K$  negative POIs  $p' \sim D_P$ ;
8     foreach sampled  $p'$  do
9        $\mathbf{u} \leftarrow \mathbf{u} - (1 - \frac{iter}{N_i})\alpha\sigma(\mathbf{u} \cdot \mathbf{p}')\mathbf{p}'$ ;  $\mathbf{p}' \leftarrow \mathbf{p}' - (1 - \frac{iter}{N_i})\alpha\sigma(\mathbf{u} \cdot \mathbf{p}')\mathbf{u}$ ;
10    Sample a batch of  $B$  peer relationships  $(u, u')$  from  $\mathcal{E}_{col}$ ;
11    foreach sampled  $(u, u')$  do
12       $\mathbf{u} \leftarrow \mathbf{u} - (1 - \frac{iter}{N_i})\alpha\lambda_1(\sigma(\mathbf{u} \cdot \mathbf{u}') - 1)\mathbf{u}'$ ;  $\mathbf{u}' \leftarrow \mathbf{u}' - (1 - \frac{iter}{N_i})\alpha\lambda_1(\sigma(\mathbf{u} \cdot \mathbf{u}') - 1)\mathbf{u}$ ;
13      Sample  $K$  negative users  $v \sim D_U$ ;
14      foreach sampled  $v$  do
15         $\mathbf{u} \leftarrow \mathbf{u} - (1 - \frac{iter}{N_i})\alpha\lambda_1\sigma(\mathbf{u} \cdot \mathbf{v})\mathbf{v}$ ;  $\mathbf{v} \leftarrow \mathbf{v} - (1 - \frac{iter}{N_i})\alpha\lambda_1\sigma(\mathbf{u} \cdot \mathbf{v})\mathbf{u}$ ;
16    Sample a batch of  $B$  peer relationships  $(u, u')$  from  $\mathcal{E}_{cov}$ ;
17    foreach sampled  $(u, u')$  do
18       $\mathbf{u} \leftarrow \mathbf{u} - (1 - \frac{iter}{N_i})\alpha\lambda_2(\sigma(\mathbf{u} \cdot \mathbf{u}') - 1)\mathbf{u}'$ ;  $\mathbf{u}' \leftarrow \mathbf{u}' - (1 - \frac{iter}{N_i})\alpha\lambda_2(\sigma(\mathbf{u} \cdot \mathbf{u}') - 1)\mathbf{u}$ ;
19      Sample  $K$  negative users  $v \sim D_U$ ;
20      foreach sampled  $v$  do
21         $\mathbf{u} \leftarrow \mathbf{u} - (1 - \frac{iter}{N_i})\alpha\lambda_2\sigma(\mathbf{u} \cdot \mathbf{v})\mathbf{v}$ ;  $\mathbf{v} \leftarrow \mathbf{v} - (1 - \frac{iter}{N_i})\alpha\lambda_2\sigma(\mathbf{u} \cdot \mathbf{v})\mathbf{u}$ ;
22     $iter \leftarrow iter + 1$ ;
23 until  $iter > N_i$ ;
24 return  $\{\mathbf{u}\}, \{\mathbf{p}\}$ ;

```

Based on $\mathcal{E}_{col}$ and $\mathcal{E}_{cov}$, we preserve the two types of mobile peer influence. The objective is:

$$O_2(u, u') = -\log \sigma(\mathbf{u} \cdot \mathbf{u}') - \sum_{i=1}^K \mathbb{E}_{v \sim D_U} \log \sigma(-\mathbf{u} \cdot \mathbf{v}). \quad (3)$$

Here $D_U(u) = 1/|\mathcal{U}|$ is the discrete uniform distribution that generates negative user samples v . Again, each trained positive (u, u') is drawn according to $w(u, u')$.

4.3.3 Mutual Augmentation Learning. The overall objective sums over all relationships in G :

$$O = \sum_{(u,p) \in \mathcal{E}_v} O_1(u, p) + \lambda_1 \sum_{(u,u') \in \mathcal{E}_{col}} O_2(u, u') + \lambda_2 \sum_{(u,u') \in \mathcal{E}_{cov}} O_2(u, u'). \quad (4)$$

where trade-off parameters λ_1 and λ_2 are introduced to regularize the importance of user preference and two types of mobile peer influence.

We now explain the mutual augmentation learning process presented in Algorithm 1. It takes as input the constructed graph G and related training parameters, and outputs user and POI representations. It first initializes all entries of representations with a uniform distribution on the interval $(-0.5/d, 0.5/d)$ (line 1). Afterward, it iteratively and sequentially optimizes O with the user preference and mobile peer influence factors (lines 2–23). For each factor, it first samples a batch of B

positive relationships such that the probability of selecting each (u, p) or (u, u') is proportional to $w(u, p)$ or $w(u, u')$. For each positive sample, it then samples K negative users or POIs and updates representations to minimize O with both positive and negative samples.

From the above analysis, it is easy to see that the time and space complexities of the algorithm are $O(N_iBKd)$ and $O(d(|\mathcal{U}| + |\mathcal{P}|))$, respectively.

4.4 Annotation

Human mobility annotation should also consider basic factors like geographical distance and POI prior. We use a Bayesian method to integrate these factors. Specifically, for each mobility record $x_i = (l_i, t_i)$, we first retrieve POIs within 500 meters [28] from l_i . Among the retrieved POIs, we then identify the top- N_c candidate POIs with a simple Bayesian formula. The probability that the mobility record should be assigned POI p can be expressed according to Bayes' theorem as follows:

$$Pr(p|l_i) \propto Pr(l_i|p)Pr(p). \quad (5)$$

Following [28], we adopt an exponential function with a negative exponent to model $Pr(l_i|p)$:

$$Pr(l_i|p) \propto \exp\{-\phi \cdot \text{dist}(l_i, p.l)\}, \quad (6)$$

where $p.l$ is the location of p and $\phi \geq 0$ is a parameter controlling the probability decay. Such a parameter offers robustness to the Bayesian method, *i.e.*, a low ϕ can weaken spatial impact if positioning accuracy is low. In practice, it suffices to choose ϕ within $[0.001, 0.1]$: the probability halves every (693, 6.9) meters with (0.001, 0.1), respectively. We use the numbers Q_p of user interactions with POI p , *e.g.*, map search queries, for evaluating the prior probability $Pr(p)$:

$$Pr(p) \propto \log(Q_p + 1) + 1. \quad (7)$$

Note that POIs having more user interactions are usually more popular for visiting.

For each of the N_c candidate POIs p , we combine the learned representations to obtain the final annotation likelihood score:

$$\text{score}(p) = \exp\{\mathbf{u} \cdot \mathbf{p}\} \cdot Pr(p|\bar{l}_{ij}). \quad (8)$$

Finally, the k candidate POIs having the top- k highest $\text{score}(p)$ are assigned to the mobility record.

Remarks. Temporal influence also remains important for POI visiting behaviors. Accordingly, we can incorporate a temporal factor term in the Bayesian formula in Equation (5) by modeling the activeness of each POI at different time with user interaction data. However, we find that this term barely has improvement upon the considered factors. Thus, we remove it for simplicity.

5 EXPERIMENTS

Using two large-scale real-world datasets, we conduct extensive experiments to demonstrate the superiority of our JEPPI framework for human mobility annotation. We compare the effectiveness and efficiency of JEPPI with five baseline methods and three variants of JEPPI. We also empirically evaluate the convergence and parameter sensitivity of JEPPI and present a case study.

5.1 Experimental Setting

Datasets. We chose two datasets to test our approach.

(1) BEIJING is produced based on the logs from a commercial map service platform. The ground-truth POIs of mobility records were determined (i) with an empirical rule combining both mobility and map search query data and (ii) by human experts. We filtered out users who had visited less than 10 POIs and used map queries as the user interaction data for POI prior.

Table 1. Statistics of Data

Description	BEIJING	NYC
time spanning	6/1/18~6/30/18	4/12/12~2/16/13
# of users	23,227	1,083
# of POIs	873,814	318,162
# of mobility records	702,858	154,883
# of visiting relationships	88,351*	18,581*
% of visiting relationships with $w(u, p) > 1$	25.3%*	21.3%*
# of co-location peer relationships	87,156	6,782
# of co-visiting peer relationships	234,085*	43,309*

* the average number on ten user-POI graphs.

(2) NYC is produced based on a public Foursquare check-in dataset [33]. We treated each check-in as a mobility record and collected POIs as well as the numbers of likes to POIs with Foursquare developers APIs (Application Programming Interfaces).³ The numbers of likes were used to estimate POI prior. The original check-in locations were very close to ground-truth POIs and a simple distance-first strategy already achieved an accuracy of 0.364. Following [28], we injected noises drawn from a Gaussian distribution with ($\mu = 0, \sigma = 2 \times 10^{-4}$) to the check-in locations.

Mobility records are randomly split into 20%-10%-70% for training-validation-testing. Note that in practice a majority of records do not have associated POIs. Table 1 lists the statistics of our data.

Evaluation metric. We adopted the Recall@k metric to evaluate the performance, which is the fraction of test mobility records whose ground-truth visiting POIs are recalled by their respective top-k annotated POIs over all test mobility records:

$$\text{Recall@k} = \frac{\text{\# of test mobility records correctly annotated by their top-}k \text{ predictive POIs}}{\text{total \# of test mobility records}}. \quad (9)$$

Under a fixed k , an approach with a higher Recall@k is more effective in the sense that the associated POIs of more mobility records are revealed. Note that all test mobility records are annotated and, hence, Recall@k is sometimes referred to as Accuracy@k [30].

Baselines. We compared our JEPPI with five baselines and three simplified variants.

- Dist only utilizes spatial information and returns the top- k nearest POIs given a location.
- Bayes returns the top- k POIs ranked by the Bayesian Equation (5) given a mobility record.
- LTR trains a LambdaMART learning-to-rank model for ranking POIs near a location [19]. We used six of the nine developed features, excluding *Creator*, *Mayor*, and *Friends-Here-Now* that are not available in our data.
- MRF constructs a Markov random field [28], which also captures user preference by enforcing consistency in spatially- or temporally-close records. We used the supervised version and only reported Recall@k with $k = 1$ since it assigns exactly one POI to each record.
- GE is a classic non-sequential POI recommendation approach [30] which learns graph-based POI embeddings to preserve POI-POI, POI-word and POI-time relationships.
- JEP, JEPPI(col), JEPPI(cov) are variants of our complete JEPPI. The simplest JEP considers user preference only, i.e., $\lambda_1 = \lambda_2 = 0$. Based on JEP, JEPPI(col) and JEPPI(cov) further incorporate co-location and co-visiting peer influence, respectively.

³<https://developer.foursquare.com/>.

Table 2. Effectiveness Evaluation (Mean Recall@k $\pm$ Standard Deviation)

Dataset	Method	$k = 1$	$k = 2$	$k = 3$	$k = 4$
BEIJING	Dist	$0.055 \pm 2e-4$	$0.103 \pm 3e-4$	$0.145 \pm 4e-4$	$0.183 \pm 4e-4$
	Bayes	$0.159 \pm 4e-4$	$0.233 \pm 3e-4$	$0.282 \pm 3e-4$	$0.321 \pm 4e-4$
	LTR	$0.211 \pm 4e-2$	$0.276 \pm 5e-2$	$0.306 \pm 5e-2$	$0.326 \pm 5e-2$
	MRF	$0.195 \pm 5e-4$	/	/	/
	GE	$0.261 \pm 6e-4$	$0.349 \pm 7e-4$	$0.387 \pm 7e-4$	$0.407 \pm 7e-4$
	JEPPI	$0.316 \pm 6e-4*$	$0.420 \pm 5e-4*$	$0.469 \pm 7e-4*$	$0.495 \pm 6e-4*$
	JEP	$0.242 \pm 6e-4$	$0.325 \pm 5e-4$	$0.375 \pm 7e-4$	$0.409 \pm 7e-4$
	JEPPI(col)	$0.314 \pm 6e-4$	$0.419 \pm 5e-4$	$0.468 \pm 6e-4$	$0.494 \pm 6e-4$
NYC	JEPPI(cov)	$0.295 \pm 7e-4$	$0.398 \pm 8e-4$	$0.450 \pm 9e-4$	$0.480 \pm 8e-4$
	Dist	$0.253 \pm 8e-4$	$0.395 \pm 7e-4$	$0.486 \pm 8e-4$	$0.555 \pm 9e-4$
	Bayes	$0.437 \pm 5e-4$	$0.581 \pm 6e-4$	$0.666 \pm 8e-4$	$0.714 \pm 8e-4$
	LTR	$0.370 \pm 4e-2$	$0.502 \pm 4e-2$	$0.581 \pm 4e-2$	$0.638 \pm 3e-2$
	MRF	$0.409 \pm 3e-3$	/	/	/
	GE	$0.435 \pm 6e-3$	$0.574 \pm 2e-3$	$0.644 \pm 2e-3$	$0.683 \pm 2e-3$
	JEPPI	$0.547 \pm 1e-3*$	$0.656 \pm 1e-3*$	$0.705 \pm 1e-3*$	$0.733 \pm 9e-4*$
	JEP	$0.520 \pm 2e-3$	$0.639 \pm 2e-3$	$0.694 \pm 1e-3$	$0.725 \pm 8e-4$
	JEPPI(col)	$0.546 \pm 2e-3$	$0.655 \pm 1e-3$	$0.704 \pm 9e-4$	$0.732 \pm 5e-4$
	JEPPI(cov)	$0.539 \pm 1e-3$	$0.649 \pm 1e-3$	$0.699 \pm 8e-4$	$0.727 \pm 1e-3$

* statistically significant (paired t-test p-level less than 0.01) compared with the best baseline.

Parameters and Implementation. All algorithms were implemented with C++, except for LTR using the Java RankLib.⁴ We tuned Bayesian factor ϕ in $\{0.001, 0.003, 0.01, 0.03, 0.1\}$, trade-off parameter λ_1, λ_2 in $\{0.1, 0.3, 1, 3\}$, iteration threshold N_i in $\{10^4, 2 \times 10^4, 10^5, 2 \times 10^5\}$, learning rate α in $\{0.01, 0.025, 0.05\}$, and the parameters of baselines on the validation set.

According to our tests, we set $\phi = (0.003, 0.03)$ on (BEIJING, NYC) for Bayes. Note that the location information on NYC was more accurate than BEIJING. For LTR, we applied a Z-score normalization on features and the shrinkage parameter was optimized to 0.6. For both MRF and GE, it sufficed to utilize their recommended parameters. For our JEPPI, we chose $\lambda_1 = 3, \lambda_2 = 0.3, N_i = 2 \times 10^5$ on BEIJING and $\lambda_1 = \lambda_2 = 1, N_i = 2 \times 10^4$ on NYC by default, respectively. Besides, we chose $\phi = 0.01, \alpha = 0.025, N_c = 20, d = 96, B = 64$, and $K = 5$ by default. Note that, if not specified explicitly, all parameters of JEPPI were fixed to their default values in our tests. When quantity measures were evaluated, the test was repeated over 10 times using different train-validation-test splits and the average is reported.

5.2 Experimental Results

Exp-1: Effectiveness. In the first set of tests, we evaluate the overall effectiveness of JEPPI. We aim to answer the following three questions: (Q1) how JEPPI performs for human mobility annotation; (Q2) how the user preference and mobile peer influence factors impact on the effectiveness; and (Q3) whether the equivalence-emphasizing metric improves the effectiveness.

Exp-1.1: Overall performance. We first compared the Recall@k of our JEPPI with baselines Dist, Bayes, LTR, MRF, and GE. The results are reported in Table 2, and we find the following.

First, spatial information alone is insufficient for this task and Dist performs the worst. By further combining POI prior, Bayes makes a noticeable improvement over Dist. Recall that we use

⁴<https://sourceforge.net/p/lemur/wiki/RankLib/>.

Table 3. Effectiveness Study of the Equivalence-emphasizing Metric
(Mean Recall@k $\pm$ Standard Deviation)

Dataset	$w(u, u')$ in Equation (2)	$k = 1$	$k = 2$	$k = 3$	$k = 4$
BEIJING	# of co-visiting POIs	$0.271 \pm 1e-3$	$0.361 \pm 2e-3$	$0.410 \pm 1e-3$	$0.443 \pm 1e-3$
	# of co-visiting times	$0.285 \pm 8e-4$	$0.380 \pm 8e-4$	$0.429 \pm 9e-4$	$0.460 \pm 9e-4$
	ours	$0.295 \pm 7e-4$	$0.398 \pm 8e-4$	$0.450 \pm 9e-4$	$0.480 \pm 8e-4$
NYC	# of co-visiting POIs	$0.540 \pm 1e-3$	$0.650 \pm 2e-3$	$0.700 \pm 7e-4$	$0.728 \pm 1e-3$
	# of co-visiting times	$0.540 \pm 2e-3$	$0.650 \pm 1e-3$	$0.700 \pm 8e-4$	$0.728 \pm 8e-4$
	ours	$0.539 \pm 1e-3$	$0.649 \pm 1e-3$	$0.699 \pm 8e-4$	$0.727 \pm 1e-3$

the numbers of map queries and likes to model POI prior, and it turns out that both of them are good choices. However, Bayes still remains too simple for the task.

Second, the more complex LTR, MRF, and GE approaches obtain fairly better results, especially on BEIJING. Despite their advanced strategies, they all suffer from some issues. For instance, LTR simply uses the number of historical visits by a user to a POI to model preference, which suffers greatly from data sparsity. Besides, the assumption of MRF, *i.e.*, the associated POIs of spatially- or temporally-close mobility records of the same users are consistent, is too strong. Finally, GE does not explicitly model users. Instead, it combines all the visited POIs by a user to represent the user, which might lose many unique characteristics.

Third, jointly exploiting user preference and mobile peer influence via mutual augmentation learning is effective. Our full JEPPI consistently outperforms all tested baselines on both datasets and the improvement is statistically significant. Indeed, on BEIJING and NYC, the Recall@k of JEPPI is on average (293%, 75%, 52%, 62%, 21%) and (65%, 12%, 29%, 34%, 14%) higher than (Dist, Bayes, LTR, MRF, and GE) with all tested k , respectively. The improvement is more remarkable on BEIJING than NYC because the spatial information on NYC ensures all methods a good performance. Another difference between the two datasets is that POI prior is very effective on NYC, as a result of which Bayes alone is comparable to the more complex LTR, MRF, and GE.

Exp-1.2: Ablation study. To evaluate the impacts of user preference and mobile peer influence, we further computed the Recall@k of JEP, JEPPI(col) and JEPPI(cov), reported in Table 2.

Comparing JEP with Bayes, we can see that the user preference factor can consistently improve the effectiveness. Indeed, the Recall@k of JEP is on average (38%, 8.6%) higher than Bayes on (BEIJING, NYC). In addition, each of the two types of mobile peer influence can further enhance annotation effectiveness. The Recall@k of (JEPPI(col), JEPPI(cov)) is on average (26%, 20%) and (2.5%, 1.6%) higher than JEP on BEIJING and NYC, respectively. Note that, given the adequate user preference information on NYC, it is not easy for the mobile peer influence to obtain further improvement. Finally, combining the two types of mobile peer influence together, our full JEPPI is slightly better than both JEPPI(col) and JEPPI(cov) on BEIJING and NYC.

Exp-1.3: Study on the equivalence-emphasizing metric. To evaluate the impacts of our metric, we computed the Recall@k of JEPPI(cov) with three choices of co-visiting weights $w(u, u')$ in Equation (2): (i) the number of co-visiting POIs, *i.e.*, $w(u, u') = |\{p \mid w(u, p) \cdot w(u', p) > 0\}|$, (ii) the number of co-visiting times *i.e.*, $w(u, u') = \sum_p \min\{w(u, p), w(u', p)\}$, and (iii) our Equation (2) which is the equivalence-emphasizing version of co-visiting times. The results are reported in Table 3.

As can be seen, the different choices of $w(u, u')$ give almost the same results on NYC. This should not be surprising since co-visiting mobile peer influence itself only has minor improvement on NYC compared with JEP (Table 2). The Recall@k on BEIJING is effectively influenced. As expected, the number of co-visiting POIs is not a good choice, compared with the number of co-visiting

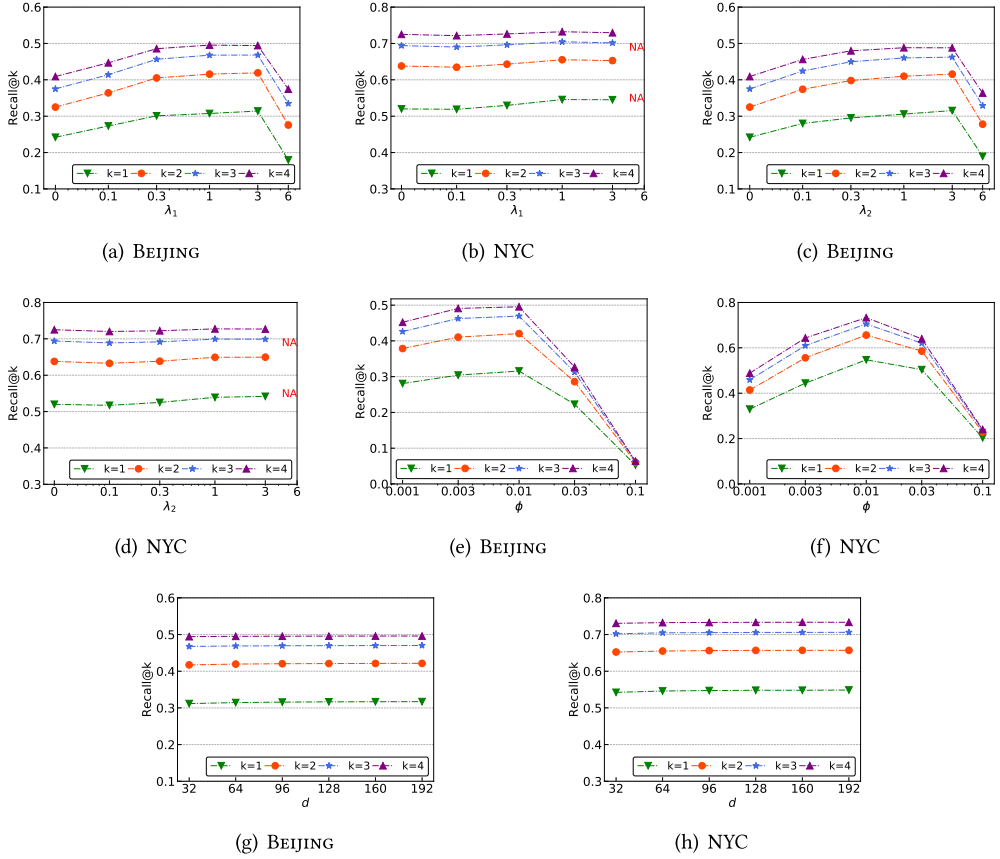


Fig. 7. Parameter sensitivity.

times. Moreover, the equivalence-emphasizing strategy further boosts the effectiveness. Indeed, (with, without) equivalence-emphasizing, the Recall@k of JEPPI(cov) is on average (10%, 20%) higher than JEP on BEIJING, *i.e.*, the improvement is doubled by our metric.

Exp-2: Parameter sensitivity. In the second set of tests, we evaluate the parameter sensitivity of JEPPI *w.r.t.* four parameters: trade-off parameters λ_1 and λ_2 , Bayesian factor ϕ , and number d of dimensions. Note that all parameters were fixed to their default values if not specified.

Exp-2.1: Impacts of λ_1 , *i.e.*, co-location mobile peer influence. To evaluate the impacts of λ_1 , we varied λ_1 from 0 to 6, fixed $\lambda_2 = 0$, and tested the Recall@k. The results are reported in Figure 7(a) and (b). When varying λ_1 , the Recall@k of JEPPI first increases and then decreases with the increment of λ_1 on both datasets. Note that JEPPI does not converge on NYC with $\lambda_1 = 6$. With less adequate user preference information on BEIJING, λ_1 has larger impacts on BEIJING than NYC. The best improvement is achieved when $1 \leq \lambda_1 \leq 3$ on both datasets.

Exp-2.2: Impacts of λ_2 , *i.e.*, co-visiting mobile peer influence. To evaluate the impacts of λ_2 , we varied λ_2 from 0 to 6, fixed $\lambda_1 = 0$, and tested the Recall@k. The results are reported in Figure 7(c) and (d). When varying λ_2 , the Recall@k of JEPPI also first increases and then decreases on both datasets. Again, JEPPI has convergence issue on NYC with $\lambda_2 = 6$ and is more sensitive to λ_2 on BEIJING. The best improvement is also achieved when $1 \leq \lambda_2 \leq 3$. From the above, it should be easy to choose a reasonable value for both λ_1 and λ_2 .

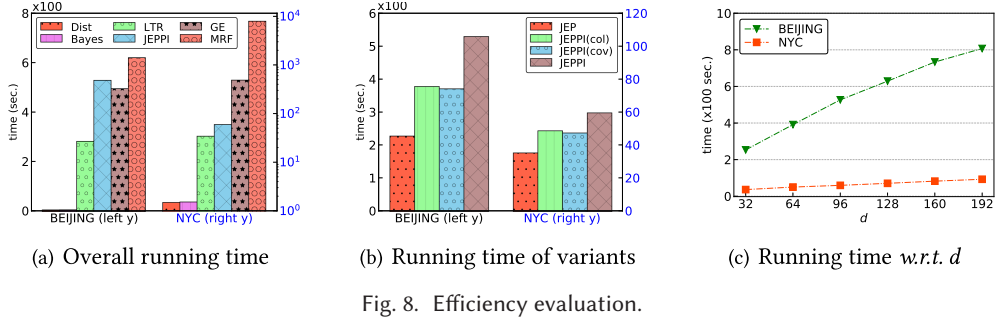


Fig. 8. Efficiency evaluation.

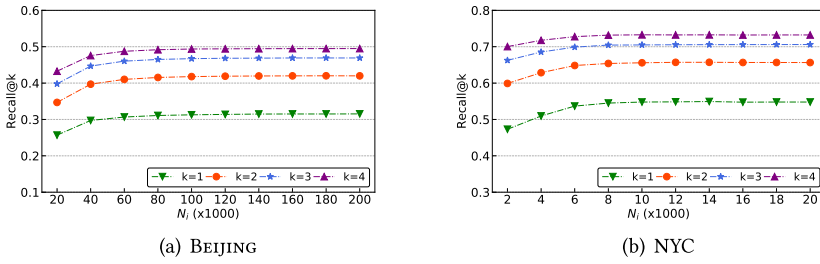


Fig. 9. Convergence study.

Exp-2.3: Impacts of Bayesian factor ϕ . To evaluate the impacts of ϕ , we varied ϕ from 0.001 to 0.1 and tested the Recall@k. The results are reported in Figure 7(e) and (f). When varying ϕ , the Recall@k of JEPI first increases and then decreases on both datasets for all tested k . In other words, it suffices to choose a good ϕ from [0.001, 0.1]. Besides, a moderate $\phi = 0.01$ can work well for both inaccurate and accurate positioning, i.e., on BEIJING and NYC.

Exp-2.4 Impacts of the number of dimensions. To evaluate the impacts of d , we varied d from 32 to 192 and tested the Recall@k. The results are reported in Figures 7(g) & 7(h). When varying d , the Recall@k keeps stable. Our approach is very robust to d .

Exp-3: Efficiency. In the third set of tests, we evaluate the efficiency of JEPI shown in Figure 8. Overall, Dist and Bayes are the most efficient due to their simple strategies (Figure 8(a)). Algorithm LTR also runs faster than JEPI. Note that LTR is a learning-to-rank approach with six features. On the other hand, JEPI is comparable with GE in efficiency and consistently faster than MRF. For the variants of JEPI, the running time increases if more factors are exploited (Figure 8(b)). The complete JEPI can finish the tests in (527, 60) seconds on (BEIJING, NYC). Finally, the running time of JEPI increases linearly with the number d of dimensions (Figure 8(c)). To conclude, the extra running time of JEPI is affordable for achieving better effectiveness.

Exp-4: Convergence. Our JEPI iteratively optimizes the objective with various factors. In the fourth set of tests, we empirically evaluate the convergence of JEPI by testing the Recall@k with different number N_i of iterations. The results are reported in Figure 9, where the Recall@k first increases at the beginning and then becomes stable on both datasets. We find that, with proper trade-off parameters, JEPI can converge after (100K, 10K) iterations on (BEIJING, NYC).

Exp-5. Case study. We finally present a case study in Figure 10 to illustrate how the idea of jointly exploiting user preference and mobile peer influence works in a real-life example. Recall that user u_0 generates four mobility records x_1-x_4 and we are to annotate these records with the POIs p_0-p_3 . We compare the top-five POIs identified by JEP and JEPI.

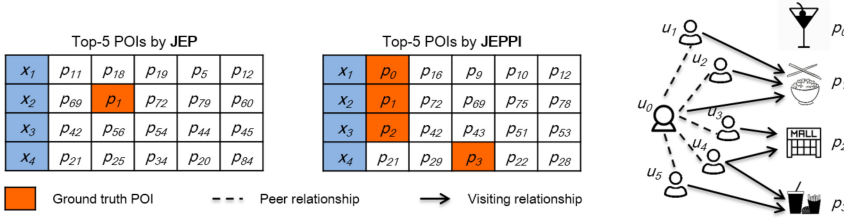


Fig. 10. Case study.

Learning user preference, JEP ranks POI p_1 of x_2 in the second place based on visiting relationship (u_0, p_1) . By further incorporating the mobile peer influence, JEPPI significantly refines the results. (i) For x_2 , it further ranks POI p_1 at the top-1 place by learning the mobile peer influence between u_0 and u_1/u_2 as well as u_1/u_2 's preference for p_1 . (ii) For x_3 and x_4 , it successfully ranks their ground-truth visiting POIs at top places, also by learning the preference of u_3-u_5 for p_2 and p_3 . (iii) Finally, for x_1 , it still finds out the true visiting POI p_0 even though there are no direct clues. This is because JEPPI can indirectly learn the preference of “similar” people of u_0 through mobile peer relationships.

6 CONCLUDING REMARKS

In this article, we proposed a new framework, named JEPPI, to jointly exploit user preference and mobile peer influence for human mobility annotation. The two distinct factors were unified in a behavior-driven user-POI graph via users' visiting, co-location, and co-visiting behaviors, respectively. An equivalence-emphasizing metric was proposed for evaluating the second-order co-visiting peer influence. We learned latent user and POI representations with mutual augmentation learning to preserve user preference and mobile peer influence. Moreover, a Bayesian method was used to further fuse spatial influence and POI prior. Finally, extensive experiments were conducted on two real-world datasets to demonstrate the superiority of our approach as well as the advantages of jointly exploiting user preference and mobile peer influence.

A couple of issues are worth further investigation. First, we are to explore peer influence from social media, like [35], for human mobility annotation and use it to enrich our behavior-driven mobile peer influence if possible. Second, we are to evaluate the utility of the annotated human mobility in real-world applications.

ACKNOWLEDGMENTS

For any correspondence, please refer to Shuai Ma and Hui Xiong.

REFERENCES

- [1] Luis Otavio Alvares, Vania Bogorny, Bart Kuijpers, Jose Antonio Fernandes de Macedo, Bart Moelans, and Alejandro Vaisman. 2007. A model for enriching trajectories with semantic geographical information. In *Proceedings of the 15th ACM International Symposium on Geographic Information Systems (ACM-GIS'07)*. ACM, 22.
- [2] Ramesh Baral, S. S. Iyengar, Xiaolong Zhu, Tao Li, and Pawel Sniatala. 2019. HiRecS: A hierarchical contextual location recommendation system. *IEEE Transactions on Computational Social Systems* 6, 5 (2019), 1020–1037.
- [3] Peng Cui, Xiao Wang, Jian Pei, and Wenwu Zhu. 2019. A survey on network embedding. *IEEE Transactions on Knowledge and Data Engineering* 31, 5 (2019), 833–852.
- [4] Yuxiao Dong, Nitesh V. Chawla, and Ananthram Swami. 2017. metapath2vec: Scalable representation learning for heterogeneous networks. In *Proceedings of the 23rd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*.

- [5] Jean-Benoit Griesner, Tael Abdessalem, and Hubert Naacke. 2015. POI recommendation: Towards fused matrix factorization with geographical and temporal influences. In *Proceedings of the 9th ACM Conference on Recommender Systems*.
- [6] Aditya Grover and Jure Leskovec. 2016. node2vec: Scalable feature learning for networks. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*.
- [7] Long Guo, Dongxiang Zhang, Yuan Wang, Huayu Wu, Bin Cui, and Kian-Lee Tan. 2018. CO2: Inferring personal interests from raw footprints by connecting the offline world with the online world. *ACM Transactions on Information Systems* 36, 3 (2018), 31:1–31:29.
- [8] Renjun Hu, Charu C. Aggarwal, Shuai Ma, and Jinpeng Huai. 2016. An embedding approach to anomaly detection. In *Proceedings of the 32nd IEEE International Conference on Data Engineering*.
- [9] Renjun Hu, Xinjiang Lu, Chuanren Liu, Yanyan Li, Hao Liu, Jingjing Gu, Shuai Ma, and Hui Xiong. 2019. Why we go where we go: Profiling user decisions on choosing POIs. In *Proceedings of the 29th International Joint Conference on Artificial Intelligence (IJCAI'19)*.
- [10] Renjun Hu, Jingbo Zhou, Xinjiang Lu, Hengshu Zhu, Shuai Ma, and Hui Xiong. 2020. NCF: A neural context fusion approach to raw mobility annotation. *IEEE Transactions on Mobile Computing* (2020), 1–1.
- [11] Cong Li, Shumin Zhang, and Xiang Li. 2019. Can multiple social ties help improve human location prediction? *Physica A: Statistical Mechanics and its Applications* 525 (2019), 1276–1288.
- [12] Quannan Li, Yu Zheng, Xing Xie, Yukun Chen, Wenyu Liu, and Wei-Ying Ma. 2008. Mining user similarity based on location history. In *Proceedings of the 16th ACM SIGSPATIAL International Conference on Advances in Geographic Information Systems*.
- [13] Defu Lian, Xing Xie, Vincent W. Zheng, Nicholas Jing Yuan, Fuzheng Zhang, and Enhong Chen. 2015. CEPR: A collaborative exploration and periodically returning model for location prediction. *ACM Transactions on Intelligent Systems and Technology* 6, 1 (03 2015), 1–27.
- [14] Hao Liu, Ting Li, Renjun Hu, Yanjie Fu, and Jingjing Guand Hui Xiong. 2019. Joint representation learning for multi-modal transportation recommendation. In *Proceedings of the 33rd AAAI Conference on Artificial Intelligence, AAAI*.
- [15] Qiang Liu, Shu Wu, Liang Wang, and Tieniu Tan. 2016. Predicting the next location: A recurrent model with spatial and temporal contexts. In *Proceedings of the 13th AAAI Conference on Artificial Intelligence*.
- [16] Yanchi Liu, Chuanren Liu, Xinjiang Lu, Mingfei Teng, Hengshu Zhu, and Hui Xiong. 2017. Point-of-interest demand modeling with human mobility patterns. In *Proceedings of the 23rd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*.
- [17] Arielle Moro, Vaibhav Kulkarni, Pierre-Adrien Ghiringhelli, Bertil Chapuis, and Benoît Garbinato. 2019. Breadcrumbs: A feature rich mobility dataset with point of interest annotation. In *Proceedings of the 27th ACM International Conference on Advances in Geographic Information Systems (SIGSPATIAL)*. 508–511.
- [18] Jiezhong Qiu, Yuxiao Dong, Hao Ma, Jian Li, Kuansan Wang, and Jie Tang. 2018. Network embedding as matrix factorization: Unifying DeepWalk, LINE, PTE, and node2vec. In *Proceedings of the 11th ACM International Conference on Web Search and Data Mining*.
- [19] Blake Shaw, Jon Shea, Siddhartha Sinha, and Andrew Hogue. 2013. Learning to rank for spatiotemporal search. In *Proceedings of the 6th ACM International Conference on Web Search and Data Mining*.
- [20] Hongzhi Shi, Chao Zhang, Quanming Yao, Yong Li, Funing Sun, and Depeng Jin. 2019. State-sharing sparse hidden markov models for personalized sequences. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*.
- [21] Stefano Spaccapietra, Christine Parent, Maria Luisa Damiani, Jose Antonio de Macedo, Fabio Porto, and Christelle Vangenot. 2008. A conceptual view on trajectories. *Data & Knowledge Engineering* 65, 1 (2008), 126–146.
- [22] Jian Tang, Meng Qu, Mingzhe Wang, Ming Zhang, Jun Yan, and Qiaozhu Mei. 2015. LINE: Large-scale information network embedding. In *Proceedings of the 24th International Conference on World Wide Web*.
- [23] Jameson L. Toole, Carlos Herrera-Yaque, Christian M. Schneider, and Marta C. González. 2015. Coupling human mobility and social ties. *Journal of The Royal Society Interface* 12, 105 (2015), 20141128.
- [24] Hongwei Wang, Jia Wang, Jialin Wang, Miao Zhao, Weinan Zhang, Fuzheng Zhang, Xing Xie, and Minyi Guo. 2018. Graphgan: Graph representation learning with generative adversarial nets. In *Proceedings of the 32nd AAAI Conference on Artificial Intelligence*.
- [25] Pengfei Wang, Yanjie Fu, Quannan Liu, Wenqing Hu, and Charu Aggarwal. 2017. Human mobility synchronization and trip purpose detection with mixture of hawkes processes. In *Proceedings of the 23rd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*.
- [26] S. Wang, S. Nepal, R. Sinnott, and C. Rudolph. 2019. P-STM: Privacy-protected social tie mining of individual trajectories. In *Proceedings of the IEEE International Conference on Web Services (ICWS'19)*.
- [27] Xiao Wang, Peng Cui, Jing Wang, Jian Pei, Wenwu Zhu, and Shiqiang Yang. 2017. Community preserving network embedding. In *Proceedings of the 31st AAAI Conference on Artificial Intelligence*.

- [28] Fei Wu and Zhenhui Li. 2016. Where did you go: Personalized annotation of mobility records. In *Proceedings of the 25th ACM International Conference on Information and Knowledge Management (CIKM'16)*.
- [29] Fei Wu, Zhenhui Li, Wang-Chien Lee, Hongjian Wang, and Zhuojie Huang. 2015. Semantic annotation of mobility data using social media. In *Proceedings of the 24th International Conference on World Wide Web (WWW'15)*.
- [30] Min Xie, Hongzhi Yin, Hao Wang, Fanjiang Xu, Weitong Chen, and Sen Wang. 2016. Learning graph-based POI embedding for location-based recommendation. In *Proceedings of the 25th ACM International Conference on Information and Knowledge Management*.
- [31] Zhixian Yan, Dipanjan Chakraborty, Christine Parent, Stefano Spaccapietra, and Karl Aberer. 2013. Semantic trajectories: Mobility data computation and annotation. *ACM Transactions on Intelligent Systems and Technology* 4, 3 (2013), 49:1–49:38.
- [32] Cheng Yang, Maosong Sun, Wayne Xin Zhao, Zhiyuan Liu, and Edward Y. Chang. 2017. A neural network approach to jointly modeling social networks and mobile trajectories. *ACM Transactions on Information Systems* 35, 4 (2017), 36:1–36:28.
- [33] Dingqi Yang, Daqing Zhang, Vincent W. Zheng, and Zhiyong Yu. 2015. Modeling user activity preference by leveraging user spatial temporal characteristics in LBSNs. *IEEE Transactions on Systems, Man, and Cybernetics: Systems* 45, 1 (2015), 129–142.
- [34] Jingru Yang, Ju Fan, Zhewei Wei, Guoliang Li, Tongyu Liu, and Xiaoyong Du. 2018. Cost-effective data annotation using game-based crowdsourcing. *Proceedings of the VLDB Endowment* 12, 1 (2018), 57–70.
- [35] Hongzhi Yin, Qinyong Wang, Kai Zheng, Zhixu Li, Jiali Yang, and Xiaofang Zhou. 2019. Social influence-based group representation learning for group recommendation. In *Proceedings of the 35th IEEE International Conference on Data Engineering, ICDE*.
- [36] Quan Yuan, Wei Zhang, Chao Zhang, Xinhe Geng, Gao Cong, and Jiawei Han. 2017. PRED: Periodic region detection for mobility modeling of social media users. In *Proceedings of the 10th ACM International Conference on Web Search and Data Mining (WSDM'17)*.
- [37] Pengpeng Zhao, Haifeng Zhu, Yanchi Liu, Jiajie Xu, Zhixu Li, Fuzhen Zhuang, Victor S. Sheng, and Xiaofang Zhou. 2019. Where to go next: A spatio-temporal gated network for next poi recommendation. In *Proceedings of the 33rd AAAI Conference on Artificial Intelligence*.
- [38] Chang Zhou, Yuqiong Liu, Xiaofei Liu, Zhongyi Liu, and Jun Gao. 2017. Scalable graph embedding for asymmetric proximity. In *Proceedings of the 31st AAAI Conference on Artificial Intelligence*.
- [39] Dingyuan Zhu, Peng Cui, Ziwei Zhang, Jian Pei, and Wenwu Zhu. 2018. High-order proximity preserved embedding for dynamic networks. *IEEE Transactions on Knowledge and Data Engineering* 30, 11 (2018), 2134–2144.

Received October 2019; revised March 2020; accepted June 2020