

DOI:10.3876/j.issn.1000-1980.2019.06.013

融合位置信息和物品流行度的协同过滤算法

谢修娟¹,莫凌飞²,李香菊¹,陈永³

- (1.东南大学成贤学院计算机工程系,江苏南京210088;
- 2.东南大学仪器科学与工程学院,江苏南京211189;
- 3.金斗云信息科技有限公司,江苏南京211189)

摘要:针对绝大多数用户消费习惯对地理位置的敏感性,以及推荐过程中的“长尾效应”,提出融合位置信息和物品流行度的协同过滤算法。对传统的协同过滤算法作出2点改进:第一,将用户兴趣偏好与位置偏好相结合,提出一种新的基于地理位置的用户相似度计算方法;第二,在预测评分时,引入物品流行度权重,合理地调整流行物品和长尾物品的推荐期望值。使用Foursquare数据集作为实验数据集,与相关算法进行对比实验。结果表明,改进算法能有效提高推荐的精度和推荐结果的多样性。

关键词:协同过滤;地理位置;推荐多样性;兴趣偏好;位置偏好;物品流行度

中图分类号:TP301 文献标志码:A 文章编号:1000-1980(2019)06-0568-06

Collaborative filtering algorithm combining location information and item popularity

XIE Xiujuan¹, MO Lingfei², LI Xiangju¹, CHEN Yong³

- (1. Department of Computer Engineering, Southeast University Chengxian College, Nanjing 210088, China;
2. College of Instrument Science and Engineering, Southeast University, Nanjing 211189, China;
3. Action Cloud Information Technology Co., Ltd., Nanjing 211189, China)

Abstract: Aiming at the sensitivity of most users' consumption habits to the geographical location and the "long tail effect" in the recommendation process, a collaborative filtering algorithm is proposed combining the location information and the item popularity. In this study, two improvements are made to the conventional collaborative filtering algorithm. At first, a new calculating method of user similarity based on the geographic location is proposed by combining the preference of user interest and the position preference. Then, in predicting, the recommended expectations of the popular items and the long-tailed items are reasonably adjusted by introducing the popularity of goods weight. The Foursquare data set is used as the experimental data set in this work, and experiments are compared with related algorithms. The experimental results show that the improved algorithm can effectively improve the precision and the diversity of the recommended results.

Key words: collaborative filtering; geographic location; recommended diversity; interest preference; position preference; item popularity

随着移动互联网的快速发展和智能手机的日益普及,O2O(on line to offline)的电子商务模式越来越受到大众的欢迎。用户通过电子商务网站线上购买服务,然后在线下实体店完成消费,十分快捷便利。用户线

基金项目:国家自然科学基金(61603091);江苏省教育信息化研究立项课题(20180054);全国高等院校计算机基础教育研究会教学研究项目(2018-AFCEC-210)

作者简介:谢修娟(1981—),女,讲师,硕士,主要从事信息提取与分析研究。E-mail:xiujuan_xie@126.com

引用本文:谢修娟,莫凌飞,李香菊,等.融合位置信息和物品流行度的协同过滤算法[J].河海大学学报(自然科学版),2019,47(6):568-573.
XIE Xiujuan, MO Lingfei, LI Xiangju, et al. Collaborative filtering algorithm combining location information and item popularity[J]. Journal of Hohai University(Natural Sciences), 2019, 47(6): 568-573.

下消费时,有一个很重要的因素——地理位置,现实生活里,绝大多数用户的消费习惯均具有位置敏感性,比如:用户倾向于在居住地附近消费,或者习惯于在自己或朋友认可度比较高的几个固定地点进行消费。所以要提高用户个性化推荐的准确度,必须考虑地理位置信息。

在重视推荐准确度的同时,还要注重多样性、新颖性、覆盖率等因素对推荐质量的影响^[1]。近些年,推荐过程中的“长尾效应”^[2]被广泛关注,推荐结果往往偏向于热门物品,如何增加用户感兴趣的长尾物品的推荐几率,提高推荐结果的多样性,提升用户的消费体验,已然成为推荐系统中一个新的研究热点。

为此,本文提出融合地理位置信息和物品流行度的协同过滤算法,充分考虑用户消费敏感的地理位置因素对推荐结果准确度的影响,同时根据物品的流行度对 TOP-N 推荐列表进行重排,增加非热门物品的推荐几率。

1 现有研究

协同过滤的概念是在 1992 年由 Coldberg、Nieols、Oki 和 Terry 首次提出,应用于 Tapestry 系统。传统的协同过滤算法分为基于记忆的方法和基于模型的方法^[3]。基于记忆的方法是通过用户对物品的评分数据,计算用户或物品之间的相似度,而后根据相似的近邻来进行推荐,该方法主要包括 user-based KNN^[4]和 item-based KNN^[5]两类。基于模型的方法是通过机器学习训练出有效的推荐模型,并应用此模型进行预测,该方法主要包括贝叶斯层次模型^[6]、聚类模型^[7]和主题模型^[8]等。

传统的协同过滤算法由于存在数据稀疏性、冷启动等问题,推荐精度受到制约。于是,研究者们着手利用用户社交关系、用户上下文信息等元素,多角度、多维度地改进原有算法。

在用户社交关系方面, Ma 等^[9]提出使用用户之间的社会关系矩阵对目标函数进行约束,并给出了融合社会关系信息的概率矩阵分解模型,以提高推荐精度;刘贵松等^[10]基于最短路径的信任关系计算用户相似度,并在社交网络的应用中对关系链的深度进行限制,以提高推荐效率。

在用户上下文信息方面,蒋翠清等^[11]提出基于用户时空相似性的地点推荐算法,将用户签到的时间信息和空间信息融入用户相似度计算中,改进邻近用户搜索质量,以提高推荐效果;孙光福等^[12]通过时序信息构建消费网络图,而后计算相似度并集成到概率矩阵分解中,以提高推荐精度;申艳光等^[13]将包含时间、地点和活动状态的情景信息引入到协同过滤推荐中,通过情景预过滤,对“用户-项目-情景”三维模型降维,缓解了数据稀疏性问题,以提高推荐的精确度。

推荐系统如果仅把推荐精确度作为唯一目标,一些长尾物品可能永远无法得到推荐,长此以往,推荐性能将会受到影响,因此,多样化推荐开始被强调。王付强等^[14]提出一种基于物品推荐期望的 TOP-N 推荐方法,通过控制全体物品的推荐期望,达到提高推荐总体多样性的目的;颢悦等^[15]使用主题模型对用户关注对象建模,将相似对象聚为一类,选取代表性标签进行融合排序,从而得到能够反映出用户兴趣特点且多样化的人物标签推荐结果。

基于以上研究,本文对传统的协同过滤算法分别在推荐精度和推荐多样性方面作出两点改进:第一,将用户兴趣偏好与位置偏好相结合,提出一种新的基于地理位置的用户相似度计算方法,提高近邻集的质量,从而提升推荐精度;第二,在预测评分时,引入物品流行度权重,合理地调整流行物品和长尾物品的推荐期望值,达到平衡多样性的推荐目的。

2 改进的协同过滤算法

2.1 基于地理位置的用户相似性计算

传统的协同过滤算法仅依靠用户-物品评分数据,计算用户的相似性,从而确定近邻,近邻的质量直接影响推荐精确度,因此,用户相似性的计算是协同推荐的关键步骤。大量的调查研究发现^[11,14],用户消费行为一般呈现区域集中的趋势,超过一半用户的消费区域范围在 3 km 以内,为此本文将地理位置信息融进用户的相似度计算中,先确定目标用户的频繁消费区域,而后计算其他用户的历史消费位置与目标用户的频繁消费区域之间的关系,两者之间重合度越高,欧式距离越近,则用户越相似,推荐成功率就越高。基于位置的用户相似性计算分为 3 步进行。

第一步:对目标用户 u 的消费位置进行 K-means 聚类,对聚类结果再进行筛选,得到用户 u 的若干频繁

消费区域。

这一步的大致过程是:(a)确定聚类数 K ,选取目标用户 u 的消费次数最高的消费地点作为第一个初始聚类中心,并从剩余的点中选择离第一个点最远的位置作为第二个初始聚类中心,再从余下的点中选择离第一、第二 2 个点距离之和最大的点作为第三个初始聚类中心,以此类推,直至得到 K 个初始聚类中心;(b)计算目标用户 u 的每个消费地点与每个初始聚类中心的距离,将其归类到最近的聚类中心所在的簇中;(c)重新计算每个簇中心,进行下一轮聚类,反复迭代,直至准则函数收敛,聚类结束,得到 K 个聚类簇;(d)计算每个聚类簇中目标用户 u 的消费次数,保留消费次数高于用户 u 平均消费次数的簇,作为目标用户 u 的频繁消费区域,记为 $F=\{f_1,f_2,\cdots,f_m\}$,每个频繁消费区域的中心集合记为 $C=\{c_1,c_2,\cdots,c_m\}$,每个频繁消费区域中最远的点与簇中心的距离,视作频繁消费区域半径,记为 $R=\{r_1,r_2,\cdots,r_m\}$ 。用户 u 平均消费次数的计算公式如下:

$$t = \frac{\sum_{i=1}^n t_i}{n} \tag{1}$$

式中: t_i ——目标用户 u 在第 i 个地点消费的次数; n ——目标用户 u 的消费地点的总数量。

第二步:计算其他用户 v 的历史消费地点与目标用户 u 频繁消费区域的重合次数(n_{uv}),以及用户 v 的所有历史消费位置与目标用户 u 的频繁消费区域中心的最小距离的平均值(d_{uv}),从而计算用户 u 和用户 v 的地理位置上的相似性($s(u,v)_{loc}$)。

n_{uv} 的统计方法是:计算用户 v 的第 i 个消费地点与目标用户 u 的第 $j(0 \leq j \leq m)$ 频繁消费区域中心 c_j 的欧式距离,若距离小于该频繁消费区域半径 r_j ,则认为用户 v 的第 i 个消费地点属于目标用户 u 的第 j 个频繁消费区域,则 n_{uv} 加 1,继续判断第 $i+1$ 个消费地点。

d_{uv} 的计算公式如下:

$$d_{uv} = \frac{\sum_{v_j \in v} \min_{c_i \in C} d(v_j, c_i)}{n_v} \tag{2}$$

式中: v_j ——用户 v 的第 j 个消费地点; $d(v_j, c_i)$ ——用户 v 的第 j 个消费地点和目标用户 u 的第 i 个频繁区域中心的距离; n_v ——用户 v 的消费地点的总数量。

$s(u,v)_{loc}$ 的计算公式如下:

$$s(u,v)_{loc} = e^{-d_{uv}}(1 + e^{-n_{uv}})^{-1} \tag{3}$$

第三步:将用户间基于地理位置的相似性,与传统的用户余弦相似性加权,得到最终的用户相似性。

$$s(u,v) = \lambda s(u,v)_{loc} + (1 - \lambda) s(u,v)_{cos} \tag{4}$$

$$s(u,v)_{cos} = \frac{\sum_{a \in I} (r_{ua} - \bar{r}_u)(r_{va} - \bar{r}_v)}{\sqrt{\sum_{a \in I} (r_{ua} - \bar{r}_u)^2 (r_{va} - \bar{r}_v)^2}} \tag{5}$$

式中: λ ——地理位置的相似度所占的权重,其取值范围为 $[0,1]$; I ——所有的物品集合; r_{ua} 、 r_{va} ——用户 u 和用户 v 对物品 a 的评分; $\bar{r}_u$ 、 $\bar{r}_v$ ——用户 u 和 v 对各自已评分物品的平均评分。

本文提出的基于地理位置的用户相似性,不仅考虑了用户消费行为的兴趣偏好,还考虑了用户对地理位置的偏好,根据不同的应用场景,可以动态调整两种偏好权重,得到一个最优结果。

2.2 融入物品流行度的 Top-N 推荐

传统的协同过滤算法,根据近邻的偏好,预测目标用户对未评分物品的评分值,生成推荐排序列表,前 N 项即为目标用户的 Top- N 推荐集。这一过程中没有考虑物品流行度影响因素,多数情况下,热门物品由于在近邻中的消费频次高,故比较容易获得较高的预测评分而被推荐,但这或许仅仅是受流行度影响所致,而并非目标用户的个人兴趣。相反,一些冷门但却是目标用户感兴趣的物品往往被淹没在长尾中。例如,一个流行度很高的物品 A,以及一个流行度较低的物品 B,物品 A 有 10 个近邻对其作出评价,物品 B 只有 2 个近邻对其作出评分,根据近邻得到的预测评分结果是:A 出现在 Top- N 推荐集中, B 没有被推荐。事实上,对目标用户而言,物品 B 更具有推荐价值,因为它的 2 个近邻评分更能反映近邻的兴趣所向,而不是受流行趋势的影响。

为了挖掘用户感兴趣的长尾物品,在用户 u 对物品 A 的预测评分($p(u,A)$)的基础上,引入物品流行度权重(w_A),对初始推荐列表按照 $p(u,A)w_A$ 进行重排,惩罚流行度对于物品推荐期望的影响,最后从重排后的推荐列表中选择前 N 项推荐给目标用户。 w_A 以及 $p(u,A)$ 的计算公式如下:

$$w_A = \frac{1}{\lg(1 + N_A)}$$

(6)

式中: N_A ——物品 A 被用户评价的总次数。

长尾物品评价次数较少, N_A 的最小值为 1,此时该物品的权重 w_A 最大,值为 1;而流行物品,评价次数较多, N_A 较大,权重反而小。

$$p(u,A) = \bar{r}_u + \frac{\sum_{v \in G} s(u,v)(r_{v,A} - \bar{r}_v)}{\sum_{v \in N} s(u,v)}$$

(7)

式中: G ——目标用户 u 的近邻集; $r_{v,A}$ ——用户 v 对物品 A 的评分。

2.3 算法过程描述

本文提出的融合位置信息和物品流行度的协同过滤算法,输入信息包括用户物品、评分、商品、位置信息,目标用户 u ,消费区域聚类数 K ;输出为目标用户 u 的 N 个物品推荐集。该算法的基本步骤为:(a)使用 2.1 中的知识,计算目标用户 u 与其他用户之间的基于地理位置的相似度;(b)对相似度进行排序,为 u 选择相似度最高的近邻集;(c)查找 u 的近邻集中未评分的物品,使用式(7)计算 u 对每个未评分物品的预测评分;(d)使用式(6)计算每个未评分物品的流行度权重,并与步骤(c)中的预测评分进行乘积,得到一个新的推荐排序列表;(e)筛选前 N 项作为用户 u 的推荐集。

3 实验与分析

3.1 数据集

实验所使用的是 Foursquare 数据集中纽约 2012—2013 年间的用户消费数据。该数据集包含 2 156 个用户,31 786 个地点,227 428 个用户签到记录,由于数据稀疏度非常高,过滤掉签到次数低于 5 的用户和地点,最后得到 1 061 个用户,4 981 个地点,79 532 个用户签到记录,经过处理后的数据依然比较稀疏,稀疏度约为 97.5%。由于在计算用户相似度时需要用到用户对物品的评分数据,本文运用以下映射方法将签到次数转换为评分值:签到 0 次,对应评分为 0;签到 1 次,对应评分为 1;签到 2 次,对应评分为 3;签到 3 次或以上,对应评分为 5,最高评分为 5。实验时,在数据集中随机选取 80% 作为训练集,剩余 20% 作为测试集。

3.2 评价标准

本文采用 F 度量值作为推荐准确性评测的指标,采用覆盖率 C 作为推荐多样性评测的指标。 F -Measure 方法兼顾了召回率 R 和准确率 P 这 2 个指标,召回率是指测试集中用户喜欢的物品出现在 Top- N 推荐列表中的比例,准确率是指推荐列表中用户喜欢的物品所占的比例,覆盖率是指推荐给所有用户的物品数占总物品数的比例。计算公式分别如下:

$$R = \frac{\sum_{u \in U} |L(u) \cap B(u)|}{\sum_{u \in U} |B(u)|} \quad P = \frac{\sum_{u \in U} |L(u) \cap B(u)|}{\sum_{u \in U} |L(u)|}$$

(8)

$$F = \frac{2RP}{R + P} \quad C = \frac{|U_{u \in U} L(u)|}{|I|}$$

(9)

式中: U ——用户集; $L(u)$ ——用户 u 的 Top- N 推荐列表;
 $B(u)$ ——测试集中用户 u 喜欢的物品集。

3.3 实验结果及分析

3.3.1 位置信息对用户偏好的影响

提取实验数据集中的每个用户的消费地理位置信息,统计并分析用户的频繁消费区域半径,如图 1 所示统计数据显示用户的消费地点都比较集中,绝大多数用户都会有自己的频繁消费区域,从图 1 中可以看出,55% 以上的用户频繁区域半

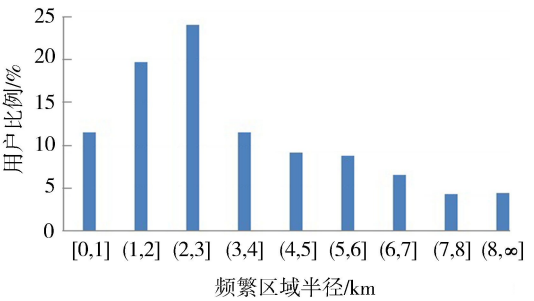


图 1 用户消费位置分析

Fig. 1 Analysis of users' consumption locations

径在 3 km 以内,用户频繁区域半径超过 8 km 的用户比例不足 5%,从而呈现出用户对地理位置的一种偏好,这也从一个侧面说明了使用地理位置信息进行用户个性化推荐切合现实情况,具有实际意义。

3.3.2 参数训练

改进算法中影响实验结果的参数有 3 个:初始聚类数 K 、地理位置的相似度权重 λ 和近邻数 z 。实验中,通过多次测试来寻找最佳的参数设置,从而得到最优 F 值。

将近邻数定为典型值 25,消费区域聚类数 K 分别取 2、3、4、5,地理位置的相似度权重 λ 从 0 开始,以 0.2 的步长增加到 1, λ 值为 0 意味着算法仅考虑用户间的余弦相似度,即传统的协同过滤算法中相似度的计算方法, λ 值为 1 意味着算法只考虑用户间地理位置的相似度。表 1 是本文算法在上述取值情况下的 F 值。

通过观察表 1 的数据,不难发现,当 $\lambda=0.6$ 时,各种 K 值下的推荐结果的 F 值基本都取了最大值,当 $K=3$ 时,随着 λ 值的变化,整个算法的 F 值比较稳定,幅度变化不大。在 $K=3, \lambda=0.6$ 时,算法的 F 值达到最佳,表明推荐精度综合效果最好。

3.3.3 推荐效果对比

为验证改进算法的有效性,将本文算法(LP-CF)的 K 和 λ 分别取固定值 3 和 0.6,邻近用户数在 5~40 之间取值,与传统的基于用户的协同过滤算法(CF),以及文献[11]中的仅空间感知下用户相似性的基于用户的协同过滤算法(UG-CF),进行对比实验,图 2 是各方法的 F 值以及 C 值的对比结果。

表 1 K 和 λ 的不同取值所对应的 F 值
Table 1 F-Measure value corresponding to different values of K and λ

λ	F			
	$K=2$	$K=3$	$K=4$	$K=5$
0	0.025 432	0.037 821	0.031 842	0.030 145
0.2	0.028 121	0.038 123	0.035 456	0.034 123
0.4	0.029 324	0.038 478	0.039 131	0.035 103
0.6	0.031 465	0.040 381	0.039 457	0.035 872
0.8	0.028 653	0.039 294	0.038 792	0.032 781
1.0	0.029 178	0.037 579	0.035 433	0.033 148

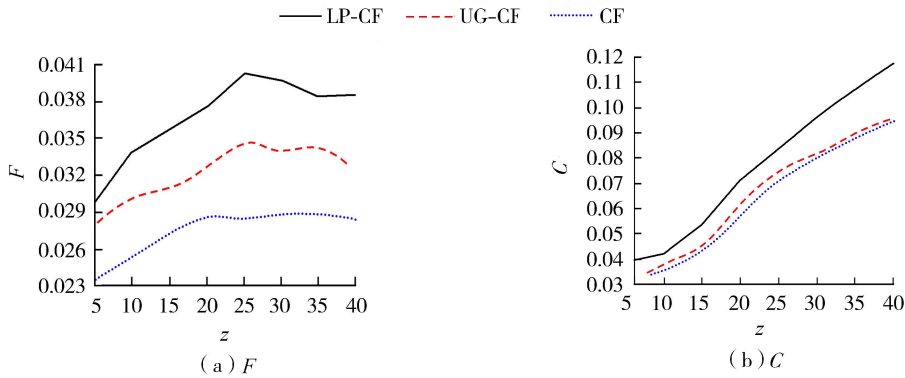


图 2 3 种方法的 C 值及 F 值对比结果

Fig. 2 Comparisons of F-Measure value and Coverage value between three methods

从图 2(a)中不难看出,LP-CF 算法和 UG-CF 算法相较于传统的 CF 算法都有明显的效果提升,LP-CF 的平均 F 值为 0.765 09,比 CF 提高了约 21.73%,UG-CF 的平均 F 值为 0.064 528,比 CF 提高了约 2.66%,充分说明融合地理位置信息后能够有效提高算法的推荐准确度。另外,随着近邻数 z 取值的不断增加,3 种算法的 F 值都呈现递增趋势,当 z 取值超过 25 之后, F 值不再上升,反而有所下降,这主要是由于数据稀疏导致,也再次验证了参数训练中近邻数定为典型值 25 的合理性。

观察图 2(b),与 UG-CF 和 CF 这 2 种算法相比,本文算法 LP-CF 的覆盖率 C 值有显著提高,证明融入物品流行度的 Top- N 推荐方法能有效调节流行物品和长尾物品的推荐期望值,提高推荐的多样性。

综上所述,LP-CF 的 F 值和 C 值均高于算法 CF 和 UG-CF,这说明本文基于位置信息和物品流行度的协同过滤算法具有一定可行性,能够提高推荐精确度,提高推荐的多样性。

4 结 语

推荐精度和推荐多样性均是评价推荐质量的重要方面,本文提出的融合位置信息和物品流行度的协同过滤算法,对传统协同过滤算法分别在推荐精度和推荐多样性两个方面作出改进,将地理位置融进用户的相

似度计算中,同时使用物品流行度对 Top-*N* 推荐列表重排。最后将改进的算法在真实数据集 Foursquare 上进行实验,验证了本文算法的有效性。

参考文献:

[1] 肖文强,姚世军,吴善明.一种改进的 top-*N* 协同过滤推荐算法[J]. 计算机应用研究,2018,35(1):105-108,112. (XIAO Wenqiang,YAO Shijun,WU Shanming. Improved top-*N* collaborative filtering recommendation algorithm[J]. Application Research of Computers,2018,35(1):105-108,112. (in Chinese))

[2] 黄璐,林川杰,何军,等. 融合主题模型和协同过滤的多样化移动应用推荐[J]. 软件学报,2017,28(3):708-720. (HUANG Lu,LIN Chuanjie,HE Jun,et al. Diversified mobile app recommendation combining topic model and collaborative filtering[J]. Journal of Software,2017,28(3):708-720. (in Chinese))

[3] CHEN Tianqi,TANG Linpeng,LIU Qin et al. Combining factorization model and additive forest for collaborative followee recommendation[C]//KDD-cup Work Shop. [s. l.]:[s. n.],2012:108-115.

[4] RESNICK P,IACOVOUN,SUCHAK M, et al. GroupLens: an open architecture for collaborative filtering of net-news[C]//Proceedings of the ACM Conference on Computer Supported Cooperative Work. Chapel Hill:ACM,1994:175-186.

[5] LINDEN G,SMITH B,YORK J. Amazon. com recommendations: item-to-item collaborative filtering[J]. Internet Computing,IEEE,2003,7(1):76-80.

[6] KOREN Y. Factor in the neighbors:scalable and accurate collaborative filtering[J]. Acm Transactions on Knowledge Discovery from Data,2010,4(1):1-24.

[7] 王明文,陶红亮,熊小勇.双向聚类迭代的协同过滤推荐算法[J]. 中文信息学报,2008,22(4):61-65. (WANG Mingwen,TAO Hongliang,XIONG Xiaoyong. A collaborative filtering recommendation algorithm based on iterative bidirectional clustering [J]. Journal of Chinese Information Processing,2008,22(4):61-65. (in Chinese))

[8] BLEI D M,NG A Y,JORDAN M I. Latent dirichlet allocation[J]. Journal of Machine Learning Research,2012,3:993-1022.

[9] MA H,ZHOU D,LIU C, et al. Recommender systems with social regularization[C]//Forth International Conference on Web Search & Web Data Mining. Hong Kong: ACM,2011:287-296.

[10] 刘贵松,解修蕊,黄海波. 基于最短路径信任关系的推荐项目计算方法[J]. 电子科技大学学报,2014,43(2):162-166. (LIU Guisong,XIE Xiurui,HUANG Haibo. Fast computing method for items recommendation based on shortest-path trust relationship[J]. Journal of University of Electronic Science and Technology of China,2014,43(2):162-166. (in Chinese))

[11] 蒋翠清,疏得友,段锐. 基于用户时空相似性的地点推荐算法[J]. 计算机工程,2018,44(7):177-182. (JIANG Cuiqing,SHU Deyou,DUAN Rui. Location recommendation algorithm based on spatial-temporal similarity of user [J]. Computer Engineering,2018,44(7):177-182. (in Chinese))

[12] 孙光福,吴乐,刘淇,等. 基于时序行为的协同过滤推荐算法[J]. 软件学报,2013,24(11):2721-2733. (SUN Guangfu,WU Le,LIU Qi, et al. Recommendations based on collaborative filtering by exploiting sequential behaviors[J]. Journal of Software,2013,24(11):2721-2733. (in Chinese))

[13] 申艳光,郭高尚,吴晶晶. 结合情景和协同过滤的移动推荐算法[J]. 科学技术与工程,2014,14(8):49-52. (SHEN Yanguang,GUO Gaoshang,WU Jingjing. A context-aware collaborative filtering algorithm on mobile recommendation[J]. Science Technology and Engineering,2014,14(8):49-52. (in Chinese))

[14] 王付强,彭甫镛,丁小焕,等. 基于位置的非对称相似性度量的协同过滤推荐算法[J]. 计算机应用,2016,36(1):171-174. (WANG Fuqiang,PENG Furong,DING Xiaohuan,et al. Location-based asymmetric similarity for collaborative filtering recommendation algorithm[J]. Journal of Computer Applications,2016,36(1):171-174. (in Chinese))

[15] 颀悦,熊锦华,程学旗. 一种融合个性化与多样性的人物标签推荐方法[J]. 中文信息学报,2017,31(2):154-162. (ZHUAN Yue XIONG Jinhua CHENG Xueqi. User tag recommendation with personalization and diversity [J]. Journal of Chinese Information Processing,2017,31(2):154-162. (in Chinese))

(收稿日期:2018-05-08 编辑:张志琴)