

基于行为模式的用户声誉度量方法研究

王洁¹, 刘建国^{2,3}

(1. 上海理工大学 管理学院, 上海 200093; 2. 上海财经大学 会计与财务研究院, 上海 200433;

3. 复旦大学 全球传播全媒体研究院, 上海 200433)

摘要: 在线评级系统由于水军和恶意打分者的存在而无法对商品给出客观评价, 因此, 建立一个基于打分行为的声誉度量模型对于在线评级系统的健康发展至关重要。现有的用户声誉度量方法仅依靠用户评分和商品质量之间的差异进行计算, 忽略了用户的行为模式。将用户的评分偏差和行为模式相结合, 提出了一种新的声誉度量算法, 该算法不仅考虑了用户打分频率的极值, 还考虑了用户打分总次数。在两个实证数据集上的实验结果表明, 新算法对随机打分的识别准确率相较于经典算法最高可以提高 17%, 对于解决冷启动和鲁棒性问题具有更好的表现。

关键词: 用户声誉; 在线评级系统; 行为模式

中图分类号: G 35 文献标志码: A

Measurement of user reputation via users' behavior patterns

WANG Jie¹, LIU Jianguo^{2,3}

(1. Business School, University of Shanghai for Science and Technology, Shanghai 200093, China; 2. Institute of Accounting and Finance, Shanghai University of Finance and Economics, Shanghai 200433, China; 3. Institute for Global Communications and Integrated Media, Fudan University, Shanghai 200433, China)

Abstract: Online rating systems are unable to provide objective evaluations of products due to the presence of water armies and malicious raters. Therefore, it is crucial to establish a reputation measurement model based on rating behavior for the healthy development of online rating systems. Existing user reputation measurement methods only take into account the difference between the user's rating information and product quality, regardless of user rating behavior patterns. Combining user rating bias and behavior patterns, a new reputation measurement algorithm was proposed, and the algorithm considered not only the extremes of user rating frequency, but also the total number of user ratings. The extensive experimental results for two empirical datasets show that the accuracy of the new algorithm for identifying random ratings can be improved by up to 17% compared to the classical algorithm, and has better performance for solving cold start and robustness problems.

Keywords: user reputation; online rating systems; behavior pattern

收稿日期: 2023-01-06

基金项目: 国家自然科学基金资助项目 (72171150, 71771152, 61773248, 71901144)

第一作者: 王洁 (1997-), 女, 硕士研究生。研究方向: 社会网络分析、大数据分析。E-mail: 913658433@qq.com

通信作者: 刘建国 (1979-), 男, 教授。研究方向: 社交大数据、商务智能、金融科技。E-mail: liujg004@ustc.edu.cn

互联网的快速发展导致人们对网络技术的依赖程度越来越高,同时也带动了在线评级系统的发展^[1]。用户可以很容易地在互联网上获取商品和服务的相关信息。与此同时,信息超载的问题也逐渐暴露出来^[2-3]。为此,用户面对大量商品而无从选择时,往往会根据评级系统提供的商品评价结果来作出消费决策^[4-6]。然而,由于在线评级系统的虚拟性,商家和消费者之间存在着严重的信息不对称,导致电子商务平台出现严重的信用问题^[7-8]。主要原因在于不是所有用户都会对商品给予合理的评价,他们可能由于判断力不佳,或者出于某些利益需求而给出极高或极低的分值^[9]。而这些不合理的评分结果会误导其他用户选择低质量商品、错过高质量商品^[10-12]。评分结果与实际不符,恶意评级用户难以识别,这些因素导致相应的电子商务平台逐渐失去消费者的信任,丧失竞争力^[13]。因此,建立一个高效可靠的声誉体系,根据用户的打分行为度量用户声誉,对在线评级系统有着非常重要的意义^[14-16]。

1 相关工作

关于在线评级系统中的用户声誉问题,学者们从不同角度提出了多种度量方法。Laureti 等^[17]认为商品质量与用户评分差值越大,用户声誉越低,提出了 IR(iterative refinement)算法。IR 算法将初始为 1 的用户声誉作为权重计算商品质量,商品质量与评分差值的倒数作为用户声誉,而后两者不断迭代至算法收敛。Zhou 等^[18]引入量化两变量相关程度的皮尔森相关系数,提出了 CR(correlation-based ranking)算法,其基本思想是打分和商品质量相似度越高的用户得到的声誉值也越高。在此基础上,考虑到系统中恶意用户的存在,Liao 等^[19]提出了声誉值再分配迭代算法,简称 IARR(iterative ranking algorithm with reputation redistribution)算法,引入惩罚因子后又提出了 IARR2 算法。两种算法增强了系统中高声誉或打分次数多的用户即高活跃度用户的影响力,这类用户被确定为系统中的高声誉用户。后来,Gao 等^[20]提出了一种基于群组的 GR(group-based ranking)算法,算法通过用户的群组行为来评估用户的声誉,即如果用户总是属于大评级组,他们就会被赋予较高的声誉分数。在引入迭代机制后,Gao 等^[21]提出了 IGR(iterative group-based ranking)算法。

Fu 等^[22]认为用户有自己的打分偏好,引入了基于用户评分偏差的 IGDR(iterative group-based and difference ranking)算法,认为用户评分的偏差越小声誉越高。Dai 等^[23]认为每个人的打分偏好不一样,有的习惯打高分,有的习惯打低分,所以需要将用户给出的评级范围映射到相同的评级标准上,提出了 PGR(improved group-based rating method based on the user preference)算法。Liu 等^[24]基于贝叶斯分析,提出了一种无参数的在线用户声誉排序 BR(parameter-free algorithm based on the Bayesian analysis)算法,该算法基于用户评价与所有用户评价主要部分一致的概率来计算用户声誉。Lee 等^[25]提出了一种基于偏差的排序方法,根据质量分类的准确性来分配每个用户的声誉指标,从用户评价标准函数的几个似是而非的假设出发,成功地推导出了一个新的声誉指数,用来衡量用户评价的统计意义,简称 DR(deviation-based spam-filtering)算法。Sun 等^[26]研究发现:正常用户的评分偏差较小,且评分呈峰值分布;相反,恶意用户通常给出有偏见的评级,他们的评级分数几乎不遵循一个已知的模式。基于此,他们提出了通过评级统计模式评估在线评级系统用户声誉的 IOR(iterative optimization ranking)算法。之前的算法都局限于商品质量分与用户打分的相似度对比,而 IOR 算法则开创了新的思路,考虑用户所有评分的分布模式,使得该算法在大量恶意用户的进攻下仍旧能保持很好的鲁棒性。

然而,IOR 算法仅仅考虑了用户最大打分频率和最小打分频率的统计分布,忽略了用户打分的整体分布特征。算法的核心要素是用户的打分模式,打分总次数和打分频率的极值对于度量用户声誉也都具有重要作用。基于此,本文提出了考虑用户打分模式的 BPR(ranking method based on user rating patterns)算法,并通过不同实证数据集上的实验结果验证了该算法的有效性。

2 基于打分模式的算法

本文提出的基于用户行为模式的 BPR 算法,其中心思想是正常用户打分时经常喜欢给出自己习惯打的分值,而其他分值则较少给出,即用户的打分行为符合峰值分布。与正常用户不同,恶意用户给商品打分通常带有主观偏见,并且打分的行为模式不同于正常用户。为此,本文度量用

户评分和商品质量之间的差异,进而考察用户的打分行为模式是否符合峰值分布,以区分正常用户与恶意用户。在此基础上,本文提出了BPR算法来度量商品的质量和用户的可信度,以下为具体的算法步骤。

Step 1 统计每个用户对于不同分值的打分次数。

$$A'_{i\beta} = \{O_i | r_{i\alpha} = \beta\}, \beta = 1, 2, 3, 4, 5 \quad (1)$$

式中: $r_{i\alpha}$ 为用户 i 对商品 α 给出的评分; β 为评分等级,此处以五级评分制为例; O_i 为用户 i 评分过的商品集合; $A'_{i\beta}$ 为用户 i 给出 β 评分的次数。

Step 2 计算用户的打分次数极差占总评分次数的比例 P_i , 因为恶意用户只给出最大和最小评分,所以他们通常比正常用户拥有更大的最小评分次数 $r_{i,\min}$, 为此,在分母中加入 $r_{i,\min}$, 让恶意用户拥有更小的 P_i 值。

$$P_i = \frac{r_{i,\max} - r_{i,\min}}{r_{i,\sum} + r_{i,\min}} \quad (2)$$

式中: $r_{i,\max}$ 和 $r_{i,\min}$ 分别代表用户 i 给出的打分种类中最多和最少的次数; $r_{i,\sum}$ 代表用户 i 的总评分次数。

通常认为,评分的统计模式符合峰值分布的正常用户会有较大的 P_i 值,而不符合峰值分布的用户会有较小的 P_i 值,本文将这类用户称为失真评级者。

Step 3 将 P_i 作为用户的初始声誉值,代入式(3)计算出商品的质量分数。

$$Q_\alpha = \frac{\sum_{i \in U_\alpha} R_i r_{i\alpha}}{\sum_{i \in U_\alpha} R_i} \quad (3)$$

式中: Q_α 表示商品 α 的质量分数; U_α 表示给商品 α 打分的用户集合; R_i 表示用户 i 的声誉值。

Step 4 通过用户打分和商品质量分数之差的绝对值计算用户打分的偏差向量 D_i 为

$$D_i = |A_i - Q_i| \quad (4)$$

式中: A_i 为用户 i 的评分矩阵; Q_i 为用户 i 每条评分对应的商品的质量分数矩阵。

Step 5 计算用户偏差向量的均值。

$$E_i = \frac{\sum_{\alpha \in O_i} d_{i\alpha}}{|D_i|} \quad (5)$$

式中: $d_{i\alpha}$ 为用户 i 对商品 α 的评分与其质量分数的差; $|D_i|$ 为偏差向量 D_i 内的元素个数,即用户 i 的打分次数。

Step 6 根据式(6)计算用户声誉,再将得到的声誉值代入式(3),得到商品最终的质量分数。正常用户通常具有较小的评分偏差 E_i 和较大的 P_i 值,而失真评级者与正常用户相比,表现出相反的情况,用户 i 的声誉 R_i 可以由 E_i 和 P_i 决定。

$$R_i = \frac{1}{E_i} + P_i \quad (6)$$

通过人工加入失真评级者,基于用户的打分模式,运用算法对所有用户的声誉进行评价,失真评级者应该具有更低的声誉值。图1给出了BPR算法的小样本示例。

图1中: $U = \{U_1, U_2, U_3, U_4, U_5\}$ 代表用户集合, $O = \{O_1, O_2, O_3, O_4\}$ 代表商品集合,二者的连边表示用户对商品的评分,这三者组成了一个加权的用户-商品二分网络,可以用于表示在线评级系统; A 为用户的评分矩阵; A' 为用户的评分次数矩阵,矩阵内元素记录了用户对不同种类分数的打分次数;矩阵 P 记录了每个用户的 P 值; Q 为商品的

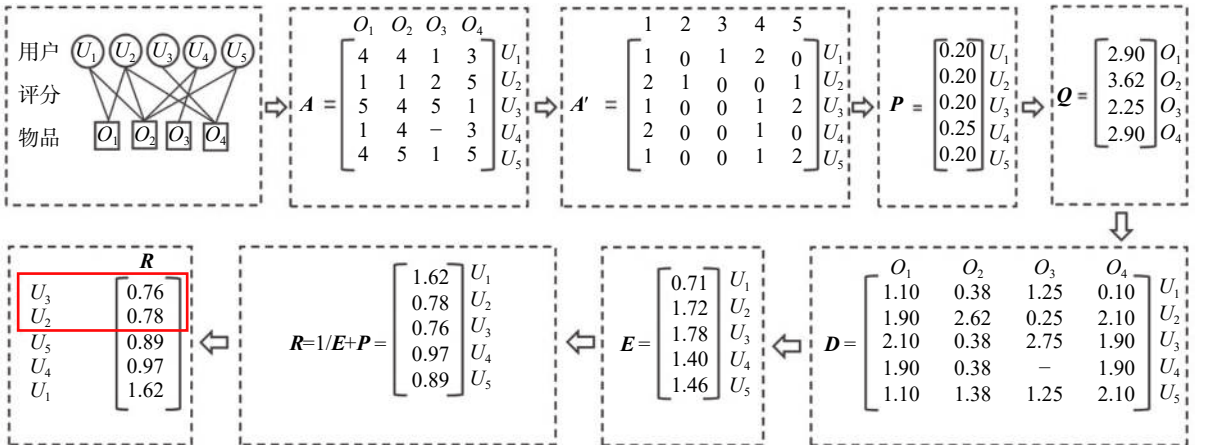


图1 BPR 算法小样本示例

Fig.1 Small sample of the BPR algorithm

质量分数矩阵; $\mathbf{D}$ 为用户评分与商品质量的差异矩阵; $\mathbf{E}$ 为用户评分偏差的均值矩阵, 为矩阵 $\mathbf{D}$ 内每行元素的均值; $\mathbf{R}$ 代表用户的声誉分数矩阵; “-”表示数据不存在。示例最后对用户声誉从小到大进行排序, 设置阈值 L 为 2, 则声誉最低的两名用户被视为算法识别出的失真评级者。

3 实验结果分析

3.1 数据集来源

MovieLens 25M 是发布于 2019 年 12 月的数据集, 由 GroupLens (<https://grouplens.org/datasets/movielens/>) 提供, 拥有 2500 万条数据可供用户随机挑选 (所有用户的评分次数均大于 20), 评分制度为 0.5~5.0 分的十级评分制。为了评估 BPR 方法的有效性, 本文使用两个真实数据集 MovieLens 10W 和 MovieLens 100W 来验证所提新算法的性能。MovieLens 10W 和 MovieLens 100W 分别为 MovieLens 25M 数据集中提取的前 10 万条数据和前 100 万条数据。实验所用的两个数据集的基本统计数据汇总在表 1 中, 所有数据均为网站提供的原始数据, 未经过处理。

3.2 用户打分模式

MovieLens 25M 数据集中大部分的用户是正常用户, 因此, 应用该数据集对用户行为模式进行探索得出的结论可以看作是正常用户的行为模式。本文将 MovieLens 25M 数据集中的 2500 万条数据分为 25 份子数据集, 每份包含 100 万条数据, 并以数字 1~25 命名。同一个用户可能被分割在相邻数据集的首尾部分, 为此, 首先删掉每个数据集的首尾两个用户。然后, 删掉评分种类小于 2 的用户, 计算用户的最大和最小打分频次占自身总评分次数的比例 (Max1 和 Min1), 在每个数据集内对所有用户求均值。接着, 删掉评分种类小于 4 的用户, 计算用户的第二大和第二小打分频次占自身总评分次数的比例 (Max2 和 Min2), 在每个数据集内对所有用户求均值。最后, 删掉评分种类小于 6 的用户, 计算用户的第三大和第三小打分频次占自身总评分次数的比例 (Max3 和 Min3), 在每个数据集内对所有用户求均值。

实验结果如图 2 所示, MovieLens 25M 中正常用户的评分行为符合峰值分布, 即用户会对习惯的少数分值打出较多的次数, 而较少打出其他分

表 1 数据集的统计特征

数据集	用户数	商品数	评分数	稀疏度
MovieLens 10W	757	9786	100000	0.0135
MovieLens 100W	6747	21952	1000000	0.0068

值。值得注意的是, 此处 Min3 比 Min2 的数值小, 这是因为计算 Min3 的数据集中删掉了评分种类小于 6 的用户, 而评分种类越多, 每种评分所占的比例自然越低。

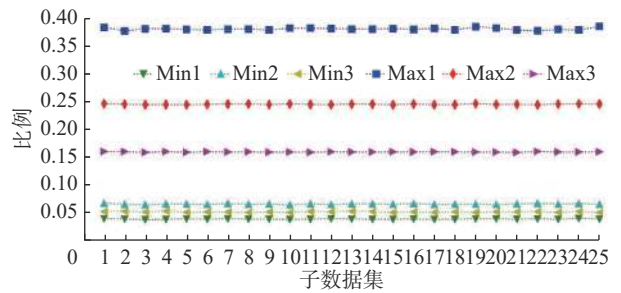


图 2 在 MovieLens 25M 的 25 个子数据集中用户的打分频率分布图

Fig.2 Frequency distributions of the user ratings in 25 different sub-datasets of the MovieLens 25M

此外, 本文还统计了 MovieLens 100W 数据集内所有用户不同种类分数打分频率的散点图, 如图 3 所示。

3.3 失真评级者

真实的评级系统中广泛存在着两种失真的评级, 即恶意评级和随机评级。恶意评级者定义为那些只对商品给予最高或最低评分的用户; 随机评级者定义为对商品没有偏好并随机打分的用户。这些异常用户的打分行为会导致商品质量分数的提高或降低, 偏离商品真实的质量。

在实验中, 随机选择数据集中的 d 个用户, 并修改他们的评分为失真的评级。计算不同算法对这些用户的召回率, 然后恢复数据集数据, 再次随机选择 d 个用户, 重复同样的操作, 每个数据集的实验重复 100 次。在 MovieLens 10W 和 MovieLens 100W 数据集中生成人造失真评级者时, 恶意评级者的打分为 0.5 或 5.0, 随机评级者的打分为集合 {0.5, 1.0, 1.5, 2.0, 2.5, 3.0, 3.5, 4.0, 4.5, 5.0} 中的任意一个数。

3.4 评价指标

实验使用召回率 $R(L)$ 和 ROC (receiver operating characteristic) 曲线下的面积 AUC (area under curve) 值来测量本文算法的性能。召回率衡量在从小到大排序的用户声誉排名列表中, 前

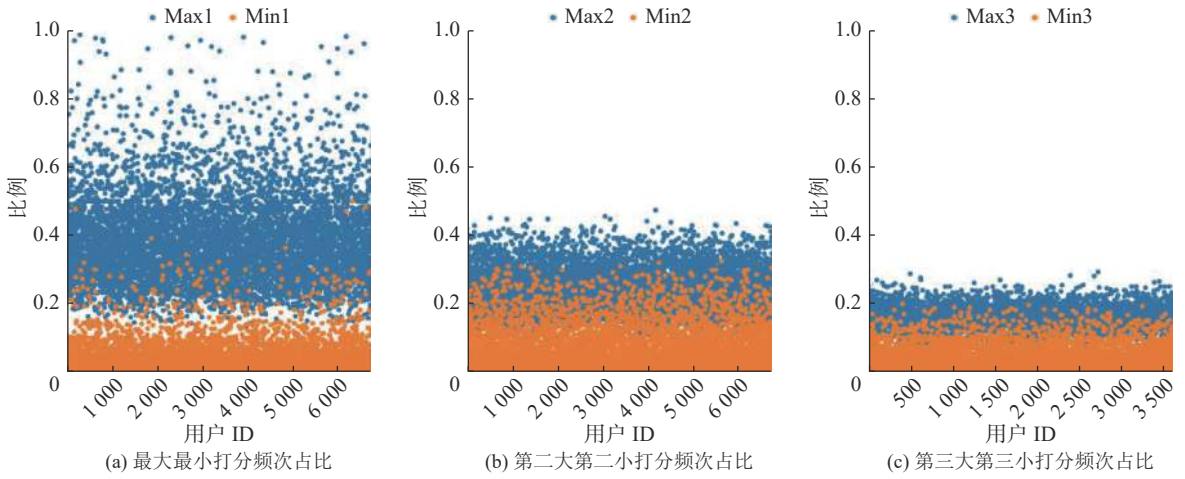


图3 MovieLens 100W数据集内所有用户不同打分频次占比

Fig.3 Proportions of the user rating frequencies for the MovieLens 100W dataset

L 位用户里可以检测到的失真评级者数量占数据集中实际添加的失真评级者数量的比例。召回率 $R(L)$ 为

$$R(L) = \frac{d'(L)}{d} \quad (7)$$

式中： $d'(L)$ 表示该方法检测到的失真评级者数量； d 表示数据集中实际添加的失真评级者数量。召回率 $R(L)$ 的范围为 $[0,1]$ ，较高的召回率表示声誉排名拥有较高的准确度。

ROC 曲线下的面积 AUC 值是衡量整个排行榜排序情况的指标，它可以解释为随机选择的失真评级者声誉值低于随机选择的正常用户声誉值的概率。本文在失真评级者和正常用户里各随机选取一人并进行比较，实验重复 N 次(本文 $N=10000$)，然后统计失真评级者的声誉值低于或等于正常用户的次数。AUC 的值为

$$S_{AUC} = \frac{N' + 0.5N''}{N} \quad (8)$$

式中： S_{AUC} 表示 AUC 的值； N 表示比较的总次数； N' 表示失真评级者的声誉值低于正常用户声誉值的次数； N'' 表示失真评级者的声誉值等于正常用户声誉值的次数。

S_{AUC} 的值越高，表明该方法的效果越好，如果 S_{AUC} 的值为 0.5，则表明该方法对所有用户的声誉进行了随机排序。

3.5 实验结果分析

基于 BPR 算法公式，计算 Movielens 10W 和 Movielens 100W 数据集的用户声誉分数，并按从小到大的顺序排列，将声誉最低的前 L 位用户视为算法检测出的失真评级者。接下来，选择 BR(基于贝叶斯的排名)、PGR(考虑用户偏好的基

于组的排名)和 IOR(迭代优化排名)这 3 种算法与本文的 BPR 算法进行比较。

一个好的声誉度量方法不受数据集大小的影响，当数据量非常大且失真评级者数量非常少时仍可以准确识别出失真评级者。因此，首先在这两个数据集上添加 50 个相同类型的失真评级者，然后通过不同算法给出的声誉分数列表计算出失真评级者的召回率 $R(L)$ 。实验结果如图 4 所示，图中每条线均为 100 次独立实验的均值。

从图 4 中可以发现，无论对于恶意评级用户还是随机评级用户，BPR 和 IOR 算法的识别召回率都高于其他算法。在图 4(a) 和 (b) 中，BPR 和 IOR 算法的召回率基本持平。图 4(b) 中， $L=50$ 时，BPR 算法相比于 IOR 算法，计算所得的恶意评级用户的召回率要高出 4%。在图 4(c) 和 (d) 中可以发现，本文提出的算法在 MovieLens 25M 数据集上识别随机评级者时具有显著优势。图 4(c) 中， $L=50$ 时，BPR 算法相比于 IOR 算法，计算所得的随机评级用户的召回率要高出 15%。图 4(d) 中： $L=100$ 时，BPR 算法的随机评级用户的召回率比 IOR 算法的要高出 21%； $L=250$ 时，BPR 算法的召回率达到了 0.96。

随着用户量的日益上涨和新产品的不断上市，用户的评分行为急剧增加，在线评级系统的数据体量从开放之日起随时间不断变大。因此，面对庞大体量的数据集，算法的运行时间，即它的效率，也至关重要。本文统计了图 4 中每个算法分别在人工添加恶意评级和随机评级的数据集上，产出用户声誉排名列表的运行时间，并取两者平均值，具体结果见图 5。该时间不包含导入数据和人工生成失真评级者的时间，只统计算法运

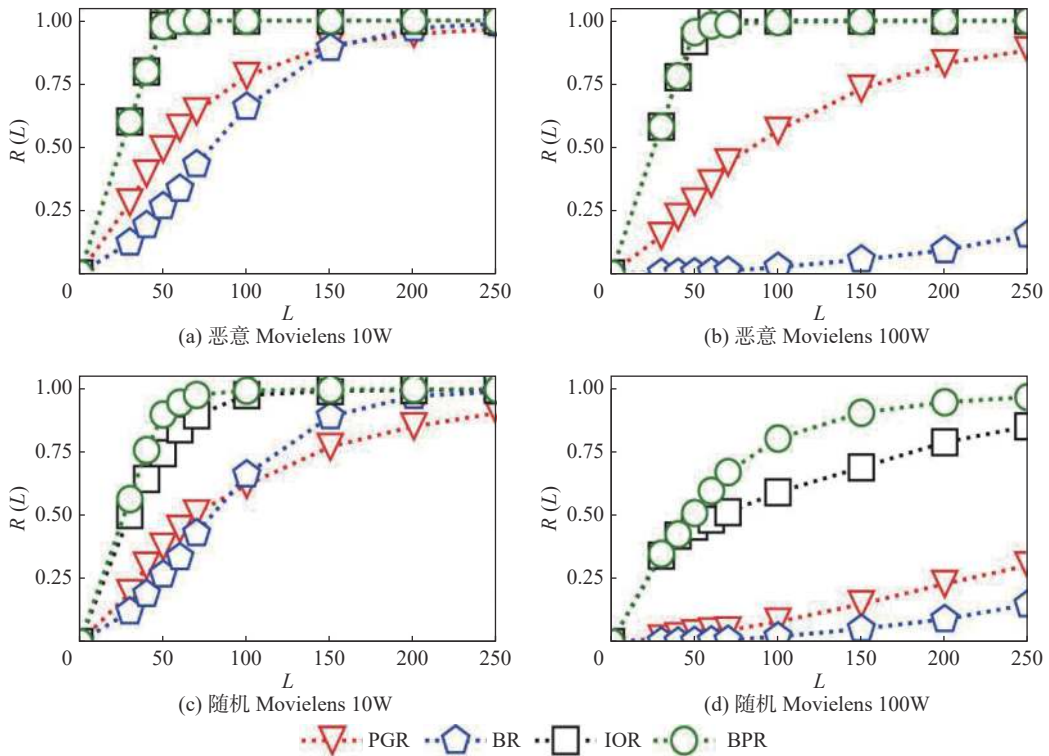


图 4 BPR 和其他算法在 Movielens 10W 和 Movielens 100W 数据集上的召回率

Fig.4 Recall rate of the BPR and other algorithms for the Movielens 10W and Movielens 100W datasets

行并导出用户声誉排名列表的时间。柱子上的数字表示运行时间的具体数值,该数值受不同机器性能的影响会有所不同,但不影响各个算法之间的比较。

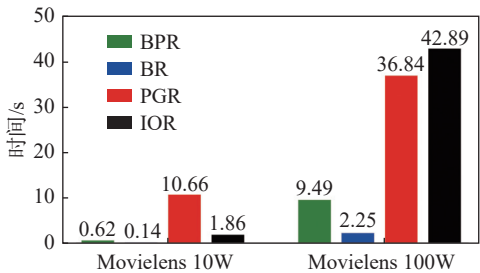


图 5 不同算法在 Movielens 10W 和 Movielens 100W 数据集上的运行时间

Fig.5 Computation time of different algorithms for the Movielens 10W and Movielens 100W datasets

因为 IOR 算法需要迭代至前后两次的商品质量分数差异小于 10^{-4} ,更重要的是理论上无法保证所有的数据都能够收敛。而 BPR 算法不需要迭代,所以算法的运行时间明显低于召回率同样较高的 IOR 算法。从图 5 可以看出,相较于基于迭代的 IOR, PGR 等算法,本文提出的 BPR 算法的计算复杂性大大低于已有的迭代算法。结合图 4 可以发现,本文提出的算法既可以提高计算所得

的召回率,也可以大幅度节省计算时间。

接下来,实验验证了算法对大量失真评级者攻击时的鲁棒性。实验设置添加人造失真评级者的数量占数据集中用户总数的比例为 $\{0.05, 0.10, 0.15, 0.20, 0.25, 0.30, 0.35, 0.40, 0.45, 0.50\}$,设置每个节点的 L 值与添加的人造失真评级者的数量相等。图 6 给出了在不同比例的失真评级者攻击评级系统时,不同算法所得的召回率变化趋势,图中每个点均为 100 次独立实验的均值。

从图 6(a) 和 (b) 可以发现, PGR 算法的召回率随系统中添加的失真评级者数量的增加而一直降低,而 BPR 和 IOR 算法的召回率始终保持在 0.97 以上。从图 6(c) 和 (d) 可以看出,随着失真评级者的增加,所有方法的召回率都在小幅增加,而 BPR 算法的召回率一直高于 IOR, BR 和 PGR 算法。从图 6(c) 可以发现:失真评级者比例为 0.05 时, BPR 算法的召回率比 IOR 的高 17%;当失真评级者比例为 0.2~0.5 时, BPR 算法的召回率比 IOR 的平均高 7%。从图 6(d) 可以发现:当失真评级者比例为 0.05 时, BPR 算法的召回率比 IOR 的高 14%;失真评级者比例为 0.2~0.5 时, BPR 算法的召回率比 IOR 的平均高 7%。上述实验结果表明,当系统中存在大量失真评级者时,本文所提出的 BPR 算法的召回率比 IOR, BR 和 PGR 算

法的都要好。

图 7 给出了不同算法的 AUC 值，图中每个点

均为 100 次独立实验的均值。在图 7(a) 和 (b) 中，当恶意评级用户的数量增加时，BPR 算法和 IOR

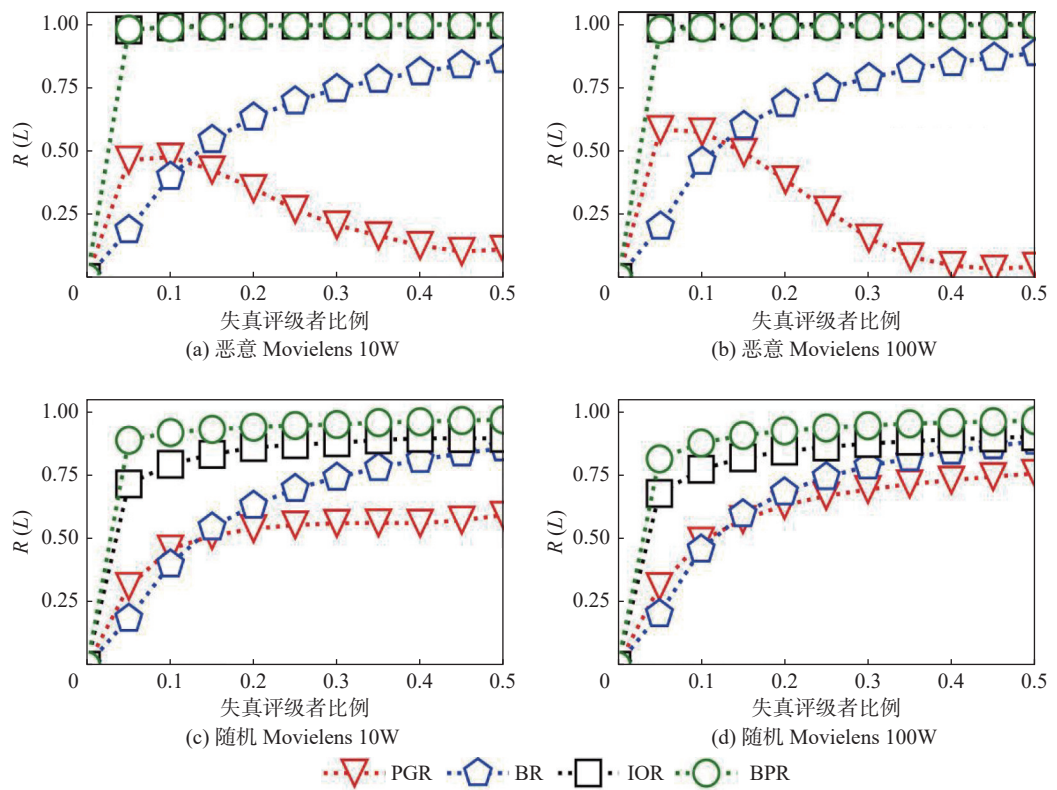


图 6 不同算法对于不同比例的恶意评级和随机评级用户的召回率

Fig.6 Recall rate of different algorithms for different ratios of the malicious and random rating users

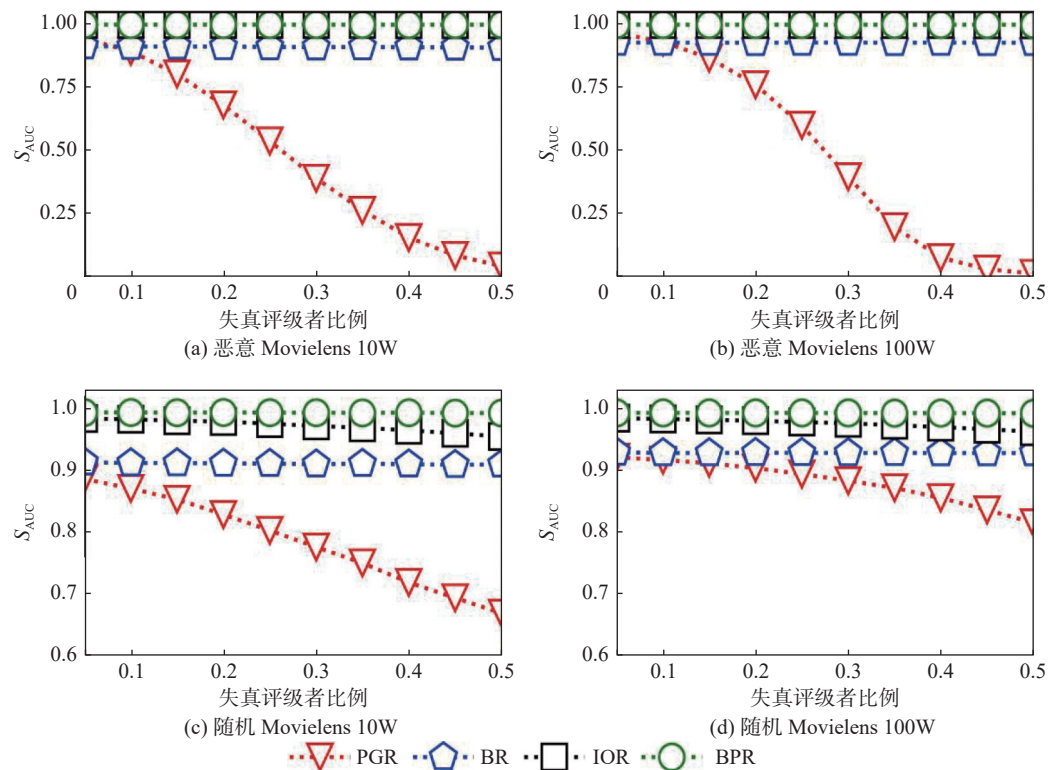


图 7 不同算法对于不同比例的恶意评级和随机评级用户的 AUC 值

Fig.7 AUC values of different algorithms for different ratios of the malicious and random rating users

算法的 AUC 值都稳定地接近 1，显著高于 BR 算法。而 PGR 算法的 AUC 值则随系统中恶意评级用户比例参数的增加而不断降低。从图 7(c) 和 (d) 可以发现，BPR 算法的 AUC 值稳定为 0.99，而 IOR 算法的 AUC 值则下降。

当用户新加入在线评级系统时，因为系统缺乏他们的评分数据，所以可能难以确认他们的身份，这即为冷启动问题。后续实验验证不同算法在冷启动问题上的表现。因为 MovieLens 25M 数据集上没有评分次数小于 20 的用户，所以将数据

集中评分次数小于 25 的用户设置为失真评级者，并计算召回率，如图 8 所示。两个数据集上评分次数少于 25 次的用户数量分别为 69 和 714 人，图中每条线均为 100 次独立实验的均值。每个子图中横坐标的两个节点分别为该数据集上添加的失真评级者数量的 1 倍和 2 倍。在图 8(a) 和 (b) 中，BPR 和 IOR 算法的召回率基本持平，稳定在 0.96 以上。在图 8(c) 和 (d) 中，阈值 L 等于失真评级者数量时，BPR 算法的召回率比 IOR 算法的高约 6%。

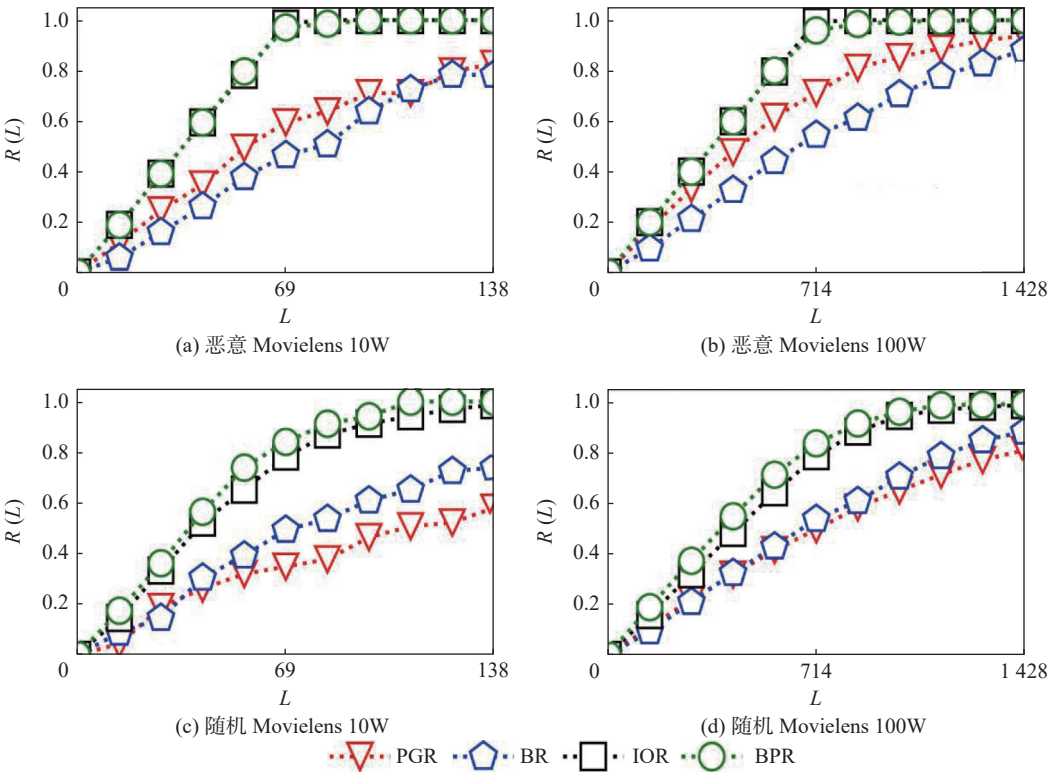


图 8 不同算法在不同参数 L 下评分次数少于 25 次的用户召回率

Fig.8 User recall rate with different algorithms rating less than 25 times under different parameters L

4 结 论

准确识别在线评级系统中用户的声誉，进而识别正常用户和给出恶意评级与随机评级的失真评级者，对于在线评级平台的健康发展具有重要意义。在线评级系统中存在失真评级者给出的随机或恶意评分，扭曲了正常的商品分数，影响了消费者的决策。因此，为在线评级系统设计一个可以准确识别失真评级者声誉的排名方法，可以使其更公正地给出商品评分，增加正常用户对平台的信任。实证分析 MovieLens 25M 上的 2500 万条真实数据发现，除了最高频和最低频打分，用

户的整体打分为具有非常稳定的行为模式，即不同频次的打分具有非常稳定的分布概率。考虑用户的打分模式，本文提出了基于用户行为模式的在线用户声誉度量方法，不仅考虑了用户的最高最低打分频次，还考虑了用户打分的总次数。不同真实数据集上的实验结果表明，本文算法计算所得的召回率、AUC 值等指标均优于其他算法，更重要的是本文提出的算法具有更低的计算复杂度和更高的准确率。进一步，对算法的鲁棒性和冷启动等问题进行实证分析后发现，本文提出的 BPR 算法在解决鲁棒性和冷启动问题中同样优于其他算法。

在未来的工作中,该领域还可从以下几方面开展进一步的研究:a.现实世界中,一些人为了自身的利益,通常会在一个连续时间段内雇佣大量水军进行虚假评分,造成某些或某个商品的实际质量与评分不一致的情况出现,因此,如何有效识别失真评级者团体将会是一个符合现实需求的研究方向。b.由于用户的审美会随流行趋势的变化而变化,为此,在统计评分结果时,可以考虑对不同时间段的评分给予不同的权重。c.大量恶意攻击发生时,失真评级者往往因为占据了“话语权”而难以被识别,所以还应通过寻找恶意评级者的固有属性而不是仅仅依靠其评分,来更有效地将其识别出来。

参考文献:

- [1] LI M, JIANG Y X, DI Z R. Characterizing the reputation of evaluators using vectors in the object feature space[J]. *Expert Systems with Applications*, 2022, 201: 117136.
- [2] ZENG A, VIDMER A, MEDO M, et al. Information filtering by similarity-preferential diffusion processes[J]. *Europhysics Letters*, 2014, 105(5): 58002.
- [3] ZHANG F G, ZENG A. Improving information filtering via network manipulation[J]. *Europhysics Letters*, 2012, 100(5): 58005.
- [4] ZHAO Y, WANG L, TANG H J, et al. Electronic word-of-mouth and consumer purchase intentions in social e-commerce[J]. *Electronic Commerce Research and Applications*, 2020, 41: 100980.
- [5] WU X, LIAO H, TANG M. Decision making towards large-scale alternatives from multiple online platforms by a multivariate time-series-based method[J]. *Expert Systems with Applications*, 2023, 212: 118838.
- [6] ESPOSITO C, GALLI A, MOSCATO V, et al. Multi-criteria assessment of user trust in social reviewing systems with subjective logic fusion[J]. *Information Fusion*, 2022, 77: 1–18.
- [7] WANG L, WAN J, ZHANG Y Q, et al. Trustworthiness two-way games via margin policy in e-commerce platforms[J]. *Applied Intelligence*, 2022, 52(3): 2671–2689.
- [8] UREÑA R, KOU G, DONG Y C, et al. A review on trust propagation and opinion dynamics in social networks and group decision making frameworks[J]. *Information Sciences*, 2019, 478: 461–475.
- [9] CHUNG C Y, HSU P Y, HUANG S H. βP : a novel approach to filter out malicious rating profiles from recommender systems[J]. *Decision Support Systems*, 2013, 55(1): 314–325.
- [10] ZHANG Y L, GUO Q, NI J, et al. Memory effect of the online rating for movies[J]. *Physica A*, 2015, 417:

- 261–266.
- [11] YANG Z M, ZHANG Z K, ZHOU T. Anchoring bias in online voting[J]. *Europhysics Letters*, 2012, 100(6): 68002.
- [12] SIDDIQUI S, FAISAL M S, KHURRAM S, et al. Quality prediction of wearable apps in the Google play store[J]. *Intelligent Automation & Soft Computing*, 2022, 32(2): 877–892.
- [13] AL-ADHAILEH M H, ALSAADE F W. Detecting and analysing fake opinions using artificial intelligence algorithms[J]. *Intelligent Automation & Soft Computing*, 2022, 32(1): 643–655.
- [14] ZENG A, CIMINI G. Removing spurious interactions in complex networks[J]. *Physical Review E*, 2012, 85(3): 036101.
- [15] ALLAHBAKHS M, IGNJATOVIC A, MOTAHARINEZHAD H R, et al. Robust evaluation of products and reviewers in social rating systems[J]. *World Wide Web*, 2015, 18(1): 73–109.
- [16] 刘晓露, 贾书伟, 刘建国, 等. 基于 Skyline Query 的高声誉用户识别方法研究 [J]. *复杂系统与复杂性科学*, 2018, 15(2): 62–70.
- [17] LAURETI P, MORET L, ZHANG Y C, et al. Information filtering via iterative refinement[J]. *Europhysics Letters*, 2006, 75(6): 1006–1012.
- [18] ZHOU Y B, LEI T, ZHOU T. A robust ranking algorithm to spamming[J]. *Europhysics Letters*, 2011, 94(4): 48002.
- [19] LIAO H, ZENG A, XIAO R, et al. Ranking reputation and quality in online rating systems[J]. *PLoS One*, 2014, 9(5): e97146.
- [20] GAO J, DONG Y W, SHANG M S, et al. Group-based ranking method for online rating systems with spamming attacks[J]. *Europhysics Letters*, 2015, 110(2): 28003.
- [21] GAO J, ZHOU T. Evaluating user reputation in online rating systems via an iterative group-based ranking method[J]. *Physica A*, 2017, 473: 546–560.
- [22] FU Q Y, REN J F, SUN H L. Iterative group-based and difference ranking method for online rating systems with spamming attacks[J]. *International Journal of Modern Physics C*, 2021, 32(5): 2150059.
- [23] DAI L, GUO Q, LIU X L, et al. Identifying online user reputation in terms of user preference[J]. *Physica A*, 2018, 494: 403–409.
- [24] LIU X L, LIU J G, YANG K, et al. Identifying online user reputation of user–object bipartite networks[J]. *Physica A*, 2017, 467: 508–516.
- [25] LEE D, LEE M J, KIM B J. Deviation-based spam-filtering method via stochastic approach[J]. *Europhysics Letters*, 2018, 121(6): 68004.
- [26] SUN H L, LIANG K P, LIAO H, et al. Evaluating user reputation of online rating systems by rating statistical patterns[J]. *Knowledge-Based Systems*, 2021, 219: 106895.