

一种权衡性能与隐私保护的推荐算法

马黛露丝^{1,2}, 朱海萍^{1,2}, 田锋^{1,2}, 冯沛^{1,2}, 陈妍^{1,2}, 计湘婷³, 李玉杰^{1,4}

(1. 西安交通大学电子与信息学部, 710049, 西安; 2. 西安交通大学陕西省天地网技术重点实验室, 710049, 西安;
3. 百度时代网络技术(北京)有限公司, 100085, 北京; 4. 南京理工大学计算机科学与工程学院, 210094, 南京)

摘要: 针对推荐系统中的隐私保护问题,提出一种隐私保护参数与推荐精度的均衡优化模型。以在线学习资源推荐系统为例建立矩阵分解模型,分别在数据输入和模型训练两个模块引入差分隐私噪声扰动,研究隐私保护参数 ϵ 与推荐精度的关系,并针对在线学习推荐系统数据的隐式反馈特点,提出基于资源热度负采样算法。基于西安交通大学网络学院数据集,利用百度飞桨深度学习平台,使用均方根误差作为推荐精度评价指标进行实验,结果表明:对数据基于资源热度负采样后,推荐精度更高;推荐精度与 ϵ 的倒数成正比关系,输入扰动和模型扰动分别在均方根误差取值不高于2.0和1.3且 ϵ 取值7和3时均衡效果最优;在相同 ϵ 下,当 $\epsilon \leq 5$ 时,模型扰动的差分隐私推荐算法推荐精度高于输入扰动的差分隐私推荐算法。

关键词: 推荐系统;差分隐私保护;矩阵分解;负采样

中图分类号: TP391 **文献标志码:** A

DOI: 10.7652/xjtuxb202107013 **文章编号:** 0253-987X(2021)07-0117-07



OSID 码

A Recommendation Algorithm Trading off Performance and Privacy Protection

MA Dailusi^{1,2}, ZHU Haiping^{1,2}, TIAN Feng^{1,2}, FENG Pei^{1,2},
CHEN Yan^{1,2}, JI Xiangting³, LI Yujie^{1,4}

(1. Faculty of Electronic and Information Engineering, Xi'an Jiaotong University, Xi'an 710049, China;
2. Shaanxi Province Key Laboratory of Satellite and Terrestrial Network Tech. R&D, Xi'an Jiaotong University,
Xi'an 710049, China; 3. Baidu.com Times Technology (Beijing) Co., Ltd., Beijing 100085, China;
4. School of Computer Science and Engineering, Nanjing University of Technology, Nanjing 210094, China)

Abstract: Aiming at the problem of privacy protection in recommender systems, a balanced optimization model between privacy protection parameters and recommendation accuracy is proposed. Taking online learning resource recommender system as an example, this paper builds a matrix factorization model and studies the relationship between privacy protection parameters and recommendation accuracy after introducing differential privacy noise into data input module and model training module. According to the implicit feedback characteristics of online learning resource recommender system data, a negative sampling algorithm by resource-popularity is proposed. The experiment is based on the original and balanced data from the Network College of Xi'an Jiaotong University using Baidu PaddlePaddle platform, and the root mean square error is used as an evaluation index to measure the recommendation accuracy. The results show that negative sampling makes higher prediction accuracy than the original. And the recommendation

收稿日期: 2020-11-24。 作者简介: 马黛露丝(1999—),女,硕士生;田锋(通信作者),男,教授。 基金项目: 国家自然科学基金资助项目(61877048);陕西省自然科学基金基础研究计划资助项目(2020JM-070);百度科研合作项目。

网络出版时间: 2021-03-10

网络出版地址: <http://kns.cnki.net/kcms/detail/61.1069.T.20210309.1603.008.html>

prediction accuracy is directly proportional to the reciprocal of differential privacy protection parameter. When RMSEs are less than 2.0 and 1.3 and the privacy protection parameter is 7 and 3, respectively, both the input-based algorithm and the model-based algorithm achieve best balance. Moreover, when the privacy protection parameter is no more than 5, the model-based algorithm has a higher recommendation accuracy than the input-based algorithm.

Keywords: recommendation system; differential privacy protection; matrix factorization; negative sampling

推荐算法在电商系统、学习平台等的应用为用户高效、便捷的使用提供了有力保障。但是,由于推荐的结果建立在用户各类数据之上,推荐系统也存在隐私泄露的风险。以在线学习为例,根据推荐系统提供的推荐列表,攻击者可以推测出不同学习者的偏好,再结合学习日志等辅助信息,甚至可以反向推测出特定偏好对应的学习者,从而获取其姓名、学校、专业、学号、客户端 IP 地址等真实身份信息^[1]。虽然用户的姓名、电话等信息往往不作为推荐系统使用的特征,不影响推荐精度,但是若攻击者反推得到用户偏好,进而反推得到对应该偏好的用户,便会造成用户真实身份信息及偏好数据的隐私泄露,相比于电商系统中的虚拟账户信息,此类包含真实身份的信息需要的隐私保护程度更高,因此在线学习平台的隐私保护越来越受到公众的关注。在线学习平台大多缺少显式评分,存在数据不平衡问题,并且大多使用基于隐式反馈的推荐系统,此类系统相比于显式系统,在数据处理、用户偏好建模等阶段更为复杂,实施代价更大^[2-4],因此如何在隐私保护程度与推荐精度之间均衡十分重要。

针对推荐系统的隐私保护安全问题,国内外学者的研究主要分为数据加密技术、基于数据模糊技术和差分隐私保护技术^[5-6]。差分隐私保护技术在保证算法对特定统计结果的输出概率不发生显著变化的前提下,利用算法对数据集的统计结果随机化处理,保护原始数据信息,其差分隐私保护参数能够描述使用的随机算法给出的最高隐私保护水平^[7]。差分隐私已被应用到矩阵分解推荐模型中,评估隐私保护和推荐效率的权衡效果^[8]。之后的研究将差分隐私技术与贝叶斯后验采样融合,能够得到推荐准确率更高的差分隐私推荐框架^[9]。Meng 等通过在敏感和非敏感数据的训练集中引入不同的噪声强度,保护了社交推荐系统的隐私^[10]。现有的针对在线学习平台的基于邻域的差分隐私保护推荐算法^[11],在计算用户或项目相似度矩阵时,会因稠密的矩阵临时存储表而占用相当的内存,而矩阵分解

的推荐算法在训练过程中不需要存储与维护相关表,节省内存,如在包括 4 多万用户的 Netflix Prize 数据集中,使用基于邻域的推荐方法,大约需要 30 GB 内存,而矩阵分解推荐方法只需要 4 GB 内存。

本文针对在线学习资源推荐平台隐式反馈的性质,提出基于资源热度负采样算法,解决隐式反馈系统中数据不平衡的问题,使用差分隐私保护参数 ϵ 描述隐私保护程度分级,提出性能与隐私保护均衡的推荐算法,研究矩阵分解的差分隐私保护算法中隐私保护参数与推荐精度的关系。

1 基于矩阵分解推荐模型的隐私泄露风险分析

在推理攻击与重构攻击^[12]两种攻击模型的攻击下,矩阵分解模型存在隐私泄露的风险。其中,推理攻击通常被用于推断某个模型的训练集中是否包含某个体的评级;重构攻击是根据目标受害者的一些背景信息,预测其敏感特征的准确值。当两种攻击结合起来^[13],在基于矩阵分解的推荐过程中,攻击者首先借推理攻击推断数据集中是否包括目标攻击用户的评级,其次通过重构攻击,利用已知的部分评级反向预测受害用户的潜在特征,便能发现用户的偏好信息以及其他潜在敏感信息^[14]。

矩阵分解推荐模型中攻击者的反向预测过程如文献[15]所述,受害者用户 u_1 与攻击者共享更新后的项目因子 V ,若攻击者掌握的背景信息为受害者非敏感的偏好评级 R_{i2} ,则可以通过 $u_1 = v_2 R_{i2}^{-1}$ 反推出受害者向量 u_1 ,从而进一步根据 $R_{i1} = u_1 V$ 获取受害者的其余敏感的偏好评级,造成用户向量、偏好向量等隐私泄露。

2 性能与隐私保护均衡的推荐算法

根据上述对推荐中矩阵分解模型隐私泄露风险的分析,本文在所提出的性能与隐私保护均衡的推荐算法中,使用基于矩阵分解的差分隐私保护推荐算法^[6]作为基础推荐模型,分析其中隐私保护参数

与推荐精度的关系,以在线学习资源推荐平台为例,研究框图如图 1 所示。因矩阵分解的目标优化函数为非凸函数,对求解结果的敏感度较高,在输出模块引入噪声会使预测结果的可用性大大降低,因此本文只分析矩阵分解模型中分别在数据输入和模型训练模块引入差分隐私保护方法的结果。在数据准备部分,本文针对在线学习用户数据的隐式反馈特点,提出基于资源热度负采样算法,对推荐模型的输入数据进行样本平衡处理。

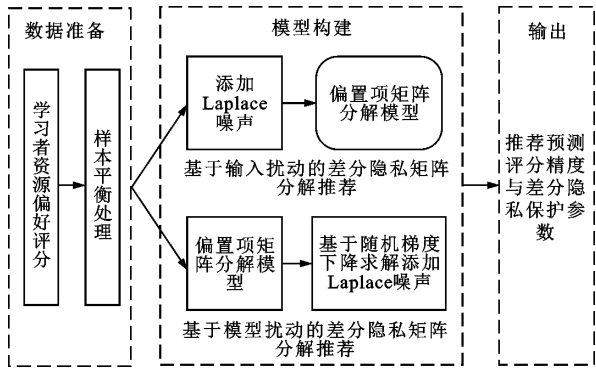


图 1 性能与隐私保护均衡的推荐算法研究框架

Fig. 1 Research framework of recommendation algorithm balancing performance and privacy protection

2.1 基于输入扰动的差分隐私矩阵分解推荐算法

基于输入扰动的差分隐私矩阵分解推荐算法是基于差分隐私保护的原理,在推荐模型训练之前对构建的隐式反馈评分引入噪声,之后再构建矩阵分解模型。对于样本平衡处理后得到的用户项目偏好评分矩阵 R 使用下式引入噪声,即

$$\left. \begin{aligned} R &= (r_{ij}) \\ f(R) &= R \\ \Delta f &= r_{\max} - r_{\min} \\ \Lambda(R) &= f(R) + \text{Laplace}\left(\frac{\Delta f}{\epsilon_1}\right) \end{aligned} \right\} \quad (1)$$

式中: Δf 表示敏感度。根据差分隐私保护的性质,最终得到的偏好评分矩阵 $\Lambda(R)$ 满足 ϵ 差分隐私保护。

2.2 基于模型扰动的差分隐私矩阵分解推荐算法

基于模型扰动的差分隐私矩阵分解保护推荐算法在梯度下降的每一步迭代中,对误差函数进行 Laplace 噪声的添加,此过程也称梯度加扰^[16]。同时,每一次迭代算法都满足 ϵ/e 差分隐私,根据差分隐私保护的组合性质,则在 e 次迭代计算后,最终整个偏置项矩阵分解满足 $\epsilon = e\epsilon/e$ 差分隐私保护。

2.3 基于资源热度负采样算法

隐式反馈推荐系统中只有用户的操作历史行为,缺少用户的负反馈样本,导致的样本不平衡问题会影响到推荐精度,为此参考已有的引入负样本的 4 种策略^[13],结合汤普森采样算法^[17]在推荐系统中的应用,针对在线学习资源推荐系统,提出基于资源热度负采样算法,其中资源热度与资源 item 出现的次数成正比,热门资源是指资源热度大于所设阈值的资源。使用基于资源热度负采样算法使得负采样后的平衡数据满足以下 3 点特征:

- (1)对于每个学习者,负采样后保证正负样本的数目均衡;
- (2)对于每个学习者,进行负样本的选择时,负样本取自同一课程下的资源库中;
- (3)在进行负采样的时候,在学习者没有操作的热门资源中,利用 β 分布随机选取资源。

基于资源热度负采样算法的具体步骤如算法 1 所述,在同类课程下, $item_list_all$ 是所有视频资源的列表, $item_list$ 则是代表学习者已经有过操作行为的资源视频的集合。

算法 1 基于资源热度负采样算法:

输入学习者编号 $user_id$,学习者有操作行为的某课程视频资源列表 $item_list$,某课程所有视频资源列表 $item_list_all$

输出正负样本均衡的样本数据 $sample$

```
1: sample = {}
2: for item in item_list do
3:     sample[item] = user_item_rating
    // 正样本不变
4: end for
5: candidate_list = Thompson(item_list_all)
    // 在热门资源中使用汤普森采样获得资源列表
6: for item in candidate_list do // 负采样
7:     if item in sample then
8:         continue
9:     end if
10:    sample[item] = 0 // 负样本评分置零
11:    n += 1
12:    if n = len(item_list) then
13:        break // 保证正负样本数均衡
14:    end if
15: end for
16: return sample
```

3 实验与结果分析

3.1 实验设计

3.1.1 实验数据集 本文主要依托西安交通大学网络学院学习平台(以下简称网络学院),选取网络学院计算机专业第4学期的《操作系统原理》《计算机网络原理》《数据结构》《计算机组成原理》《编译原理》和《Java语言》6门课程,分析在线学习者与视频学习资源的交互行为,将其总结为学习者的视频学习次数偏好、视频学习时长偏好、视频学习暂停拖动次数偏好3种操作行为,以三者的算数加权作为学习者对课程视频的评分,具体表示如下

$$\left. \begin{aligned} p_f(u_i, v_k) &= \frac{f(u_i, v_k)}{\max_{1 \leq j \leq N} f(u_j, v_k)} \\ p_d(u_i, v_k) &= \begin{cases} \frac{d(u_i, v_k)}{3d_{v_k}}, & d(u_i, v_k) \leq 3d_{v_k} \\ 1, & d(u_i, v_k) > 3d_{v_k} \end{cases} \\ p_{pd}(u_i, v_k) &= \frac{p(u_i, v_k) + g(u_i, v_k)}{\max_{1 \leq j \leq N} (p(u_j, v_k) + g(u_j, v_k))} \\ p(u_i, v_k) &= \alpha p_f(u_i, v_k) + \beta p_d(u_i, v_k) + \gamma p_{pd}(u_i, v_k) \end{aligned} \right\} (2)$$

式中: $\alpha=1; \beta=5; \gamma=4^{[11]}$; $p_f(u_i, v_k) \in [0, 1]$ 为视频学习次数偏好,由学习者 u_i 观看视频 v_k 的累计次数 $f(u_i, v_k)$ 计算得到; $p_d(u_i, v_k) \in [0, 1]$ 为视频学习时长偏好; $d(u_i, v_k)$ 为累计时长; d_{v_k} 为原始时长; $p_{pd}(u_i, v_k) \in [0, 1]$ 为视频学习暂停拖动次数偏好^[18]; $p(u_i, v_k)$ 为 u_i 暂停 v_k 的累计次数; $g(u_i, v_k)$ 为 u_i 拖动 v_k 的累计次数。对于给定的学习者与给定课程,可以得到基于学习行为偏好的评分,从而得到基于学习者学习行为偏好的学习者课程视频评分矩阵,以下统称为学习者学习资源评分矩阵,并以此作为后续实验的基础。

3.1.2 实验结果评价指标 依据所选取的网络学院的99 732条视频资源日志作为原始数据,计算得到学习者对资源的偏好评分矩阵,使用五折交叉验证方法选取测试集与训练集,实验应用百度飞桨深度学习框架平台完成。

本文实验首先选取不同的隐私保护参数,计算最终算法得到的推荐精度,探讨隐私保护参数与推荐精度的关系;其次,比较是否使用基于资源热度负采样算法进行数据处理的推荐精度,分析样本平衡性对模型的影响。实验采用的算法性能评测指标为均方根误差 r ,这类指标常被用于评估推荐算法的

预测评分精度(以下简称推荐精度)。 r 越小,预测的评分越准确,计算式如下

$$r = \frac{\sqrt{\sum_{u,i \in T} (r_{ui} - \hat{r}_{ui})^2}}{|T|} \quad (3)$$

式中: r_{ui} 代表用户 u 对项目 i 的真实评分; $\hat{r}_{ui}$ 表示预测评分; T 代表测试集。未选取准确度、召回率、命中率、平均绝对误差等其他指标^[18],是因为 r 中平方项的存在使得误差大的项和误差小的项之间差异变大,更具参考价值^[19]。

3.2 输入扰动差分隐私推荐算法实验与结果分析

3.2.1 实验设计 为验证在输入扰动模型下所提资源热度负采样算法的有效性,以及探讨隐私保护参数与推荐精度的关系,采用的对比实验有不添加噪声的基于随机梯度下降法求解的基本矩阵分解推荐算法(Clean MF)、不添加噪声的偏置项矩阵分解推荐算法(Clean Biased MF),以及添加噪声的基本矩阵分解推荐算法(INR MF)。其中,Clean MF用来对比偏置优化后的矩阵分解推荐算法与 Clean Biased MF 的推荐精度,INR MF用来与 INR Biased MF 对比在不同隐私保护参数下的推荐精度损失程度。

3.2.2 实验结果分析 选取不同的隐私保护参数,根据 Laplace 噪声的生成方法可视化其噪声分布发现,当 ϵ 的取值越小,所添加的噪声分布越分散且噪声值大,此时算法的隐私保护程度比较高;反之,当 ϵ 取值越大,所添加的噪声趋于 0,此时算法所损失的效用性最小,但相应的隐私保护程度也降低。

(1)算法推荐精度与隐私保护参数的关系分析。分析在不同的隐私保护参数下推荐算法的推荐精度,其中参数设置如下:隐式分解维度为 3,正则化参数为 0.1,迭代次数为 30,梯度下降学习速率为 0.02。

实验结果如图 2 所示,其中红色曲线明显低于绿色曲线,表示在无噪声添加的情况下,考虑用户以及项目的偏置因素会使得推荐精度更高。在进行噪声的添加后,如蓝色与橘色曲线所示,在基于输入扰动的情况下,偏置项矩阵分解对噪声的敏感度更低。橘色曲线表示随着 ϵ 的增加,带偏置项矩阵分解的推荐 r 逐渐降低,并在 $\epsilon=15$ 时趋近于无噪声的推荐精度,这与我们的理论分析也相符合,即 $r \propto 1/\epsilon$,因此针对不同的推荐系统隐私保护程度需求以及推荐精度要求,可以选取对应的 ϵ 。

(2)样本的平衡性对于算法推荐精度的影响。

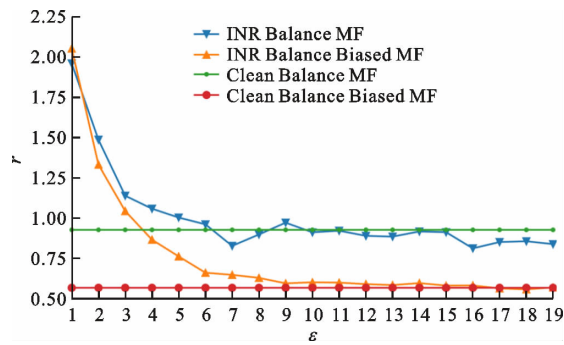


图 2 平衡样本中输入扰动时不同算法的推荐精度比较

Fig. 2 Comparisons of recommendation accuracy of different algorithms with input disturbances in balanced samples

在经过样本平衡处理的原始评分矩阵上进行相同的推荐实验,结果如图 3 所示,与图 2 所得均方根误差曲线趋势相同。

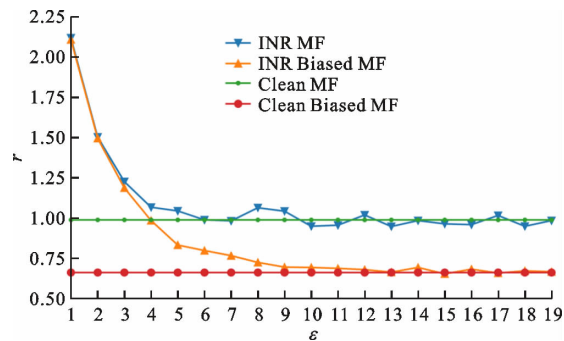


图 3 非平衡样本中输入扰动时不同算法的推荐精度比较

Fig. 3 Comparisons of recommendation accuracy of different algorithms with input disturbances in unbalanced samples

输入扰动下平衡样本与非平衡样本上的实验结果如图 4 所示。由图可知,无论是否添加噪声扰动,平衡样本下的推荐 r 都低于非平衡样本下的,并且 ϵ 取值越大影响越明显,说明在隐式反馈的矩阵分解推荐中进行样本平衡处理是必要的。

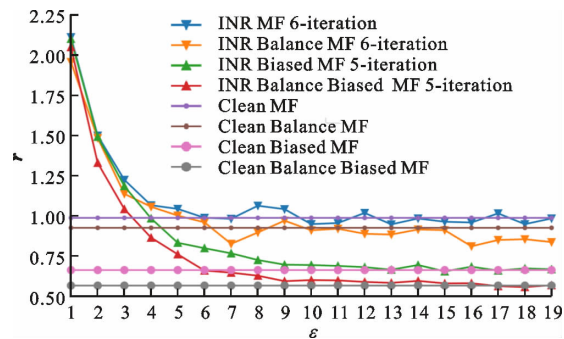


图 4 输入扰动下样本平衡性对推荐精度的影响

Fig. 4 Influence of sample balance on recommendation accuracy under input disturbance

3.3 模型扰动差分隐私推荐算法实验与结果分析

3.3.1 实验设计 为验证在模型扰动算法下的结果,采用的对比实验有 Clean MF、clean Biased MF、添加噪声的基本矩阵分解算法(MNR MF)、带偏置项矩阵分解算法(MNR Biased MF)以及文献[20]中的差分隐私保护全局平均值(Global average)和差分隐私保护项目平均值(Item average)两种方法。

3.3.2 实验结果分析

(1)算法推荐精度与隐私保护参数的关系分析。不同 ϵ 下推荐算法精度的比较如图 5 所示。实验中相关参数的设置如下:隐式分解维度为 3,正则化参数为 0.05,梯度下降学习速率为 0.01。由图 5 可知,当 $\epsilon \geq 5$ 时,本文所使用的 MNR Biased MF 推荐精度高于 Global average 和 Item average 的推荐精度。观察图 5 中 MNR MF 与 MNR Biased MF 的推荐精度与隐私保护参数的关系发现: $\epsilon = 1$ 时,推荐精度相近; $\epsilon < 1$ 时,即隐私保护程度比较高的时候,MNR MF 算法推荐精度更高; $\epsilon > 1$ 时,MNR Biased MF 算法推荐精度更高。这表示 MNR Biased MF 的矩阵分解对过高的噪声引入敏感性较高,但在隐私保护程度适中($\epsilon \geq 5$)的情况下,推荐精度在不损失过多时远高于 MNR MF。

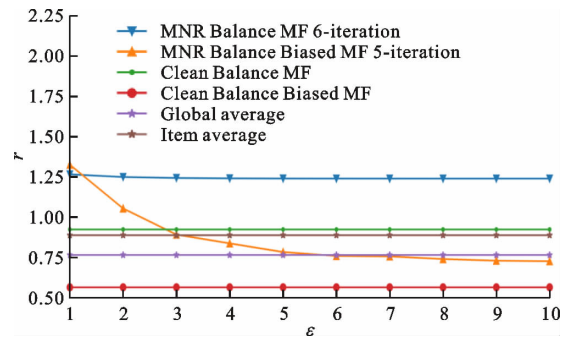


图 5 平衡样本中模型扰动时不同算法的推荐精度比较

Fig. 5 Comparisons of recommendation accuracy of different algorithms with model disturbances in balanced samples

(2)样本的平衡性对于算法推荐精度的影响。在经过样本平衡处理的原始评分矩阵上进行相同的推荐实验,结果如图 6 所示,与图 5 所得均方根误差曲线趋势相同。

样本平衡与非平衡下的实验结果如图 7 所示。无论是 MNR MF 还是 MNR Biased MF,样本平衡后的推荐精度都更高。对于 MNR Biased MF,当 $\epsilon \leq 4$ 时样本平衡处理的效果更加明显,而对于 MNR MF, $\epsilon \leq 2$ 时样本平衡处理的效果更加明显。这表

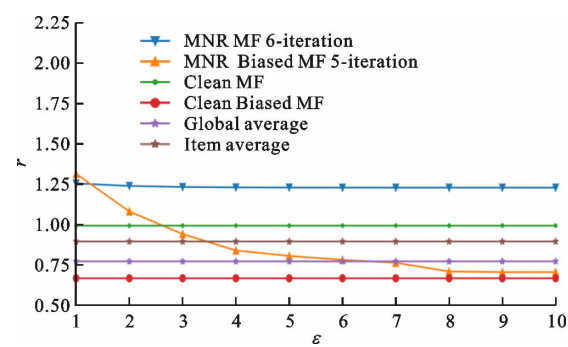


图 6 非平衡样本下模型扰动时不同算法的推荐精度比较
Fig. 6 Comparisons of recommendation accuracy of different algorithms for model disturbances in unbalanced samples

明样本平衡性对于 MNR MF 的影响更显著,因此实际中采用本节提出的 MNR Biased MF 能更好地均衡推荐精度与隐私保护程度。

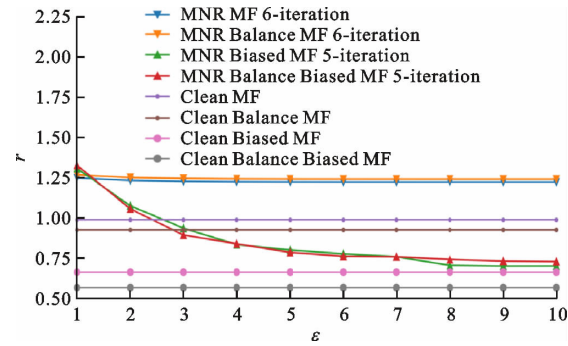


图 7 模型扰动下样本平衡性对推荐精度的影响
Fig. 7 Influence of sample balance on recommendation accuracy under model disturbance

3.4 均衡优化与对比分析

根据上述实验中发现的 ϵ 越小隐私保护程度越高, r 越小推荐精度越好, 以及 r 与 $1/\epsilon$ 正相关的关系, 可对 min-max 标准化后的两个变量构造优化目标如下

$$\min(w_1 r + w_2 \epsilon), \quad w_1 + w_2 = 1 \tag{4}$$

经实验发现, 当 $w_1 < 0.6$ 时, r 精度差, ϵ 倾向取最小值; 当 $w_1 = 0.6$, 输入扰动下在 $\epsilon = 7$ 时, 模型扰动下在 $\epsilon = 3$ 时, 平衡样本与样本非平衡实验场景下的推荐精度与隐私保护程度达到均衡最优。

同时对比分析发现: 在不同的 ϵ 下, 如表 1 所示, 相同隐私保护程度下无论样本平衡与否, 在 ϵ 取值较小时, 基于模型扰动的推荐精度优于基于输入扰动的, 但当 ϵ 取较大值 ($\epsilon > 5$) 时, 基于输入扰动的推荐精度优于基于模型扰动的。

表 1 模型扰动算法与输入扰动算法推荐预测评分精度比较

Table 1 Comparisons of recommendation prediction scoring accuracy between model-based and input-based algorithms

不同 ϵ 下的推荐算法	r	
	非平衡样本下	平衡样本下
INR ($\epsilon=1$) MF	2.116 7	1.961 6
MNR ($\epsilon=1$) MF	1.318 1	1.290 6
INR ($\epsilon=5$) MF	1.044 9	1.003 1
MNR ($\epsilon=5$) MF	1.286 8	1.261 1
INR ($\epsilon=10$) MF	0.949 8	0.911 2
MNR ($\epsilon=10$) MF	1.286 1	1.260 5
INR ($\epsilon=1$) Biased MF	2.106 0	2.054 2
MNR ($\epsilon=1$) Biased MF	1.344 7	1.382 6
INR ($\epsilon=5$) Biased MF	0.833 4	0.761 4
MNR ($\epsilon=5$) Biased MF	0.787 4	0.785 6
INR ($\epsilon=10$) Biased MF	0.694 5	0.600 6
MNR ($\epsilon=10$) Biased MF	0.708 3	0.699 8

4 结论及展望

根据本文实验结果可知, 矩阵分解的差分隐私保护推荐方案中, 无论是输入扰动还是模型扰动, 样本平衡处理后推荐精度会提高, 其对应的推荐精度与差分隐私保护参数的倒数均呈正比关系, 特别地, 输入扰动与模型扰动分别在 ϵ 取值 7 和 3 时, 推荐精度与隐私保护程度达到均衡。当隐私保护程度较高, 即 $\epsilon \leq 5$ 时, 基于模型扰动的差分隐私矩阵分解推荐在相同的隐私保护程度下, 优于基于输入扰动的; 当隐私保护程度较低, 即 $\epsilon > 5$ 时, 基于输入扰动的差分隐私矩阵分解推荐在相同的隐私保护程度下, 优于基于模型扰动的。根据模型扰动算法对系统模型维护有更高要求的特点可得, 对于隐私保护程度要求很高的推荐系统, 建议采用基于模型扰动的矩阵分解差分隐私推荐, 对于隐私保护程度要求适中的推荐系统, 建议采用维护成本较低的基于输入扰动的矩阵分解差分隐私推荐算法。

目前, 差分隐私保护与在线学习资源推荐相结合的研究还相对较少, 本文所提出的方案是一次很有意义的尝试, 有较强的实际应用价值, 但仍有优化和完善的空间, 下一步的工作将研究在增量隐式反馈数据下推荐算法的隐私保护问题。

参考文献:

- [1] CALANDRINO J, KILZER A, NARAYANAN A, et al. You might also like: privacy risks of collaborative filtering [C]//2011 IEEE Symposium on Security and Privacy. Piscataway, NJ, USA: IEEE, 2011: 231-246.
- [2] ABEL F, DELDJOO Y, KOHLSDORF D, et al. Recsyschallenge 2017: offline and online evaluation [C]//Eleventh ACM Conference on Recommender Systems. New York, USA: ACM, 2017: 372-373.
- [3] KREN M, KOS A, ZHANG Y, et al. Public interest analysis based on implicit feedback of IPTV users [J]. IEEE Transactions on Industrial Informatics, 2017, 13(4): 2077-2086.
- [4] ZHANG J, CHONG W, LIU O, et al. Improving video recommendation systems from implicit feedback in the e-marketing environment [C]//International Multi Conference of Engineers and Computer Scientists. Hong Kong, China: IAENG, 2017: 723-727.
- [5] 周俊,董晓蕾,曹珍富.推荐系统的隐私保护研究进展[J].计算机研究与发展,2019,56(10):2033-2048.
ZHOU Jun, DONG Xiaolei, CAO Zhenfu. Research progress of privacy protection in recommendation system [J]. Computer Research and Development, 2019, 56(10): 2033-2048.
- [6] 谭作文,张连福.机器学习隐私保护研究综述[J].软件学报,2020,31(7):2127-2156.
TAN Zuowen, ZHANG Lianfu. A survey of privacy protection in machine learning [J]. Journal of Software, 2020, 31(7): 2127-2156.
- [7] 熊平,朱天清,王晓峰.差分隐私保护及其应用[J].计算机学报,2014,37(1):101-122.
XIONG Ping, ZHU Tianqing, WANG Xiaofeng. Differential privacy protection and its application [J]. Journal of Computer Science, 2014, 37(1): 101-122.
- [8] RENDLE S, KRICHENE W, ZHANG L, et al. Neural collaborative filtering vs matrix factorization revisited [C]//Fourteenth ACM Conference on Recommender Systems. New York, USA: ACM, 2020: 240-248.
- [9] STEPHANE M, FALTINGS B, SCHICKEL V. Context-tree recommendation vs matrix-factorization: algorithm selection and live users evaluation [J]. Proceedings of the AAAI Conference on Artificial Intelligence, 2019(33): 9534-9540.
- [10] MENG Xuying, WANG Suhang, SHU Kai, et al. Personalized privacy preserving social recommendation [C]//Proceedings of the 32nd AAAI Conference on Artificial Intelligence. New York, USA: ACM, 2018: 1-8.
- [11] FENG Pei, ZHU Haiping, LIU Yu, et al. Differential privacy protection recommendation algorithm based on student learning behavior [C]//2018 IEEE 15th International Conference on e-Business Engineering (ICEBE). Piscataway, NJ, USA: IEEE, 2018: 285-288.
- [12] 陆艺,曹健.面向隐式反馈的推荐系统研究现状与趋势[J].计算机科学,2016,43(4):7-15.
LU Yi, CAO Jian. Research status and trend of implicit feedback-oriented recommendation system [J]. Computer Science, 2016, 43(4): 7-15.
- [13] RICCI F, ROKACH L, SHAPIRA B, et al. Recommendation system: technology, evaluation and efficient algorithm [M]. Beijing, China: Machinery Industry Press, 2015.
- [14] KOREN Y, BELL R, VOLINSKY C. Matrix factorization techniques for recommender systems [J]. Computer, 2009, 42(8): 30-37.
- [15] SHOKRI R, STRONATI M, SONG C, et al. Membership inference attacks against machine learning models [J]. Cryptography and Security, 2017, 41(5): 3-18.
- [16] FREDRIKSON M, JHA S, RISTENPART T. Model inversion attacks that exploit confidence information and basic countermeasures [C]//The 22nd ACM SIGSAC Conference. New York, USA: ACM, 2015: 1322-1333.
- [17] KAWALE J, BUI H, KVETON B, et al. Efficient Thompson sampling for online matrix-factorization recommendation [C]//Neural Information Processing Systems (NIPS 2015). Cambridge, MA, USA: MIT Press, 2015: 1297-1305.
- [18] ZHU Haiping, LIU Yu, TIAN Feng, et al. A cross-curriculum video recommendation algorithm based on a video-associated knowledge map [J]. IEEE Access, 2018, 6(1): 57562-57571.
- [19] 郭贵冰.推荐系统进展:方法与技术[M].北京:科学出版社,2019:37-38.
- [20] BERLIOZ A, FRIEDMAN A, KAAFAR M, et al. Applying differential privacy to matrix factorization [C]//Proceedings of the 9th ACM Conference on Recommender Systems. New York, USA: ACM, 2015: 107-114.

(编辑 荆树蓉 刘杨)