

基于对抗性学习的协同过滤推荐算法

詹皖江¹ 洪植林¹ 方路平¹ 吴哲夫¹ 吕跃华²

1 浙江工业大学信息学院 杭州 310023

2 浙江省科技信息研究院 杭州 310006

(wzf@zjut.edu.cn)

摘要 推荐系统能够根据用户的兴趣特点和购买行为,向用户推荐感兴趣的信息和商品。随着用户生成内容 UGC 逐渐成为当前 Web 应用的主流,基于 UGC 的推荐也得到了广泛关注。区别于传统推荐中用户与物品的二元交互,有的 UGC 推荐采用协同过滤方法,提出了消费者、物品和生产者的三元交互,进而提高了推荐准确度,但大多算法都集中在推荐的性能而忽略了对鲁棒性的研究。因此,通过结合对抗性学习和协同过滤的思想,提出了一种基于对抗性学习的协同过滤推荐算法。首先在三元关系模型参数上加入对抗性扰动,使模型的性能降至最差,与此同时使用对抗性学习的方法训练模型,以达到提高推荐模型鲁棒性的目的;其次设计了一种高效的算法用于求解模型所需的参数;最后在 Reddit 和 Pinterest 两个公共数据集上进行测试。实验结果表明:1)在相同参数设置下,与现有算法相比,所提方法的 AUC、Precision 和 Recall 指标均有明显的提高,验证了其可行性与有效性;2)该算法不仅增强了推荐性能,还提高了模型的鲁棒性。

关键词: 推荐系统;UGC;矩阵分解;对抗性学习;协同过滤

中图法分类号 TP301.6

Collaborative Filtering Recommendation Algorithm Based on Adversarial Learning

ZHAN Wan-jiang¹, HONG Zhi-lin¹, FANG Lu-ping¹, WU Zhe-fu¹ and LYU Yue-hua²

1 College of Information Engineering, Zhejiang University of Technology, Hangzhou 310023, China

2 Zhejiang Institute of Science and Technology Information, Hangzhou 310006, China

Abstract The recommendation system can recommend relevant information and commodities to the user according to the user's hobbies and purchase behavior. As user-generated content UGC gradually becomes the mainstream of current Web applications, recommendations based on UGC have also received widespread attention. Different from the binary interaction between user and item in traditional recommendation, the existing UGC recommendation adopts collaborative filtering method to propose a ternary interaction between consumer, item and producer, thereby improving the accuracy of recommendation, but most of the algorithms focus on the recommended performance and ignore the research on robustness. Therefore, by combining the ideas of adversarial learning and collaborative filtering, a collaborative filtering recommendation algorithm based on adversarial learning is proposed. First, the adversarial disturbance is added to the ternary relationship model parameters to make the performance of the model the worst. At the same time, the adversarial learning method is used to train the model to achieve the purpose of improving the robustness of the recommendation model. Secondly, an efficient algorithm is designed used to transform the parameters required by the model. Finally, it is tested on two public data sets generated by Reddit and Pinterest. The experimental results show that under the same parameter settings, compared with the existing algorithms, the AUC, Precision and Recall indicators of the proposed algorithm have been significantly improved, verifying its feasibility and effectiveness. The algorithm not only enhances the recommendation performance, but also improves the robustness of the model.

Keywords Recommendation system, User-generated content, Matrix factorization, Adversarial learning, Collaborative filtering

1 引言

随着互联网技术的发展,数据量呈爆炸式增长,信息过载成为一个不可避免的问题,推荐系统成为解决问题的关键技术之一。通过为用户提供个性化的内容建议,推荐系统不仅

可以改善用户体验,而且可以通过增加流量为企业带来可观的商业价值^[1]。一方面,协同过滤作为一种广泛应用的推荐方法正逐渐成为学术界和工业界研究的热点^[2-3];另一方面,随着 Web2.0 的发展,UGC(User Generated Content)平台受到越来越多的关注,它拥有额外的动态模型,为推荐系统的发

到稿日期:2020-06-12 返修日期:2020-11-26

基金项目:浙江省自然科学基金(LY18F010025)

This work was supported by the Natural Science Foundation of Zhejiang Province(LY18F010025).

通信作者:吕跃华(lyh@zjinfo.gov.cn)

展提供了新方向。CPR(Consumer and Producer Recommendation)是一种用于 UGC 平台的推荐算法,受到了广泛关注^[4]。

对抗性机器学习^[5]的最新进展表明,许多先进的分类器在受到外界对抗样本攻击时会错误地对图片进行分类,该结果指出了仅在静态标记数据上训练模型存在局限性。为了克服这种局限性并提高模型的鲁棒性,研究人员开发了对抗性训练方法以实现在添加了对抗样本的分类器上仍能正确地对图片进行分类的目的。

对于推荐系统而言,同样存在类似的现象,但是直接从图像域嫁接生成对抗样本的方法是不可行的,因为推荐模型的输入大多是离散的,而将噪声应用于离散特征没有意义,还可能改变样本语义。为了解决这个问题,考虑在更深的层次上探索推荐模型的鲁棒性。CPR 是一种具有竞争力的协同过滤推荐算法,但是易受参数的对抗性扰动影响,因此本文选取 CPR 作为基础模型,采用对抗性学习的方法来展开对其鲁棒性的研究。

本文的贡献有以下 3 个方面:1)分析了三元关系模型存在脆弱性问题;2)提出了一种基于对抗性学习的协同过滤推荐算法,该算法不仅提高了模型的推荐性能,还增强了模型的鲁棒性;3)在真实的数据集上进行实验,验证了本文算法的有效性。

2 相关工作

2.1 基于用户的协同过滤推荐算法

基于用户的协同过滤推荐算法(User-based Collaborative Filtering, UCF)已经在学术界和工业界进行了广泛的研究。带有显式反馈(如用户评分)的 UCF 任务直接反映了用户对物品的偏好,通常被表述为一个评分预测问题^[6-7],其目标是最小化实际评分和相应预测评分之间的总体误差。其中,矩阵分解(Matrix Factorization, MF)作为一种具有简单性和有效性的协同过滤推荐算法得到了广泛应用^[8]。为了进一步提高 MF 在评分预测中的性能,研究者提出了有偏 MF 的方法。

此外,隐式反馈(如查看、点击)也是推荐系统上用户-物品交互的主要方式,许多方法都是基于隐式反馈^[9-11]而提出的。带有隐式反馈的 UCF 通常被视为 top-N 推荐任务^[12],它向潜在用户提供了一个简短的排序列表。从技术上讲,评分预测和 top-N 推荐任务的主要区别在于模型的优化方式不同。前者通常只对已知的评分构建一个回归损失函数进行优化,而后者则需要考虑剩余的数据(实际负反馈和缺失数据的混合),然而这些数据在显式反馈模型中往往被忽略。

2.2 对抗性学习

在正常的监督训练过程中,分类器容易受到对抗性样本的攻击,这揭示了泛化过程中存在着模型不稳定的问题。为了解决这一问题,研究人员提出了对抗性训练方法,该方法通过动态生成的对抗样本来优化训练过程^[13]。学习这些对抗样本可以看作是使训练过程规范化的一种方法。

在快速增长的对抗性学习推动下,研究人员试图制造对抗性噪声来攻击机器学习模型并尝试使用对抗性学习的方法来改进模型。在计算机视觉和深度学习领域中,基于 GAN

的框架^[14]主要包含一个生成模型(generator)和一个判别模型(discriminator),其中生成模型的目的是从有标记和无标记的数据中训练出一个模型,而判别模型则试图通过对抗性的方式生成差异样本。

2.3 基于 GAN 的协同过滤推荐算法

由于 GAN 已经展示了它能从大量数据中学习的能力和潜力,因此它在最新研究的推荐系统中被广泛应用。IR-GAN^[15]和 GraphGAN^[16]实现了 GAN 在协同过滤推荐中的成功应用,其主要思路是给定一个用户,让生成器生成用户可能会购买的物品并期望它能够捕捉到用户对物品的真实偏好,并且让判别器能够准确区分出该物品。针对上述方法在离散项索引生成上的局限性,Chae 等^[17]进一步提出了 CF-GAN(Collaborative Filtering Based on Generative Adversarial Networks)技术,通过采用 vector-wise 的训练方式,对一个给定的用户,由生成器生成其购买向量,而由判别器来判别该向量的真实性。

3 基于对抗性学习的协同过滤推荐

3.1 问题描述

基于 UGC 平台推荐算法的关键属性是用户不仅提供关于物品的反馈,而且所有物品都是由用户自己创建的。我们使用 U 和 I 来分别表示用户和物品的集合,对于每个用户 u 使用 I_u^+ 来表示用户提供的所有正面反馈物品,物品 $i \in I$ 由用户 $p_i \in U$ 生产。

定义集合 $C \subset U$ 和 $P \subset U$ 分别表示消费者和生产者,此外 $C \cup P = U$, $PS = C \cap P$ 表示用户既是生产者又是消费者,其中, $|PS|/|U|$ 的比率表示消费者和生产者群体重叠的程度。

3.2 CPR 模型

有偏 MF 广泛地运用于推荐系统中^[18]。具体地说,它在用户与物品潜在向量上的内积添加了偏置项来对用户-物品交互进行建模。

$$\hat{x}_{ui} = \alpha + \beta_u + \beta_i + \langle \gamma_u, \gamma_i \rangle \quad (1)$$

尽管有偏 MF 在建模用户-物品交互方面显示了强大的性能,但并没有完全建模消费者 u 、物品 i 和生产者 p 之间的三元交互。因此,CPR 通过因子分解的方法来捕捉这种交互性:

$$\hat{x}_{ui} = \alpha + \beta_u + \beta_i + \langle \gamma_u^c, \gamma_i \rangle + \langle \gamma_u^c, \gamma_i^p \rangle \quad (2)$$

其中, γ_u^c 和 γ_u^p 分别表示用户的两个角色(消费者和生产者)。用户的消费者嵌入和生产者嵌入源自单个核心嵌入 γ_u 和两个转换矩阵:

$$\gamma_u^c = W^c \gamma_u \quad (3)$$

$$\gamma_u^p = W^p \gamma_u \quad (4)$$

其中, $W^c, W^p \in \mathbb{R}^{K \times K}$ 表示转换矩阵,使用它们将一个用户的核心嵌入投影到两个角色嵌入中。使用这种处理方式的主要原因是:1)用户之间的“跟随”关系是不对称的,无法用同质嵌入的内积建模;2)当用户扮演不同的角色时,他们可能表现出不同的行为。

在提出的偏好预测模型的基础上,文献[4]采用了贝叶斯个性化排序 BPR 准则^[19]对 CPR 模型进行了隐式反馈的优

化。其目标是为每个用户观察到的反馈排名进行近似优化。给定一个训练实例 $(u, i, j) \in D$:

$$D = \{(u, i, j) | u \in U \wedge i \in I_u^+ \wedge j \in I \setminus I_u^+\} \quad (5)$$

其中, $i \in I_u^+$ 表示用户对物品给出了反馈, $j \in I \setminus I_u^+$ 表示没有给出反馈。对于用户 u , 预测模型应该给物品 i 分配一个比物品 j 更大的偏好分数。因此, BPR 定义了偏好评分之间的差异:

$$\hat{x}_{uij} = \hat{x}_{ui} - \hat{x}_{uj} \quad (6)$$

CPR 算法最终通过最大化后验来优化排名:

$$L_{BPR}(\theta) = \sum_{(u, i, j) \in D} -\ln \sigma(\hat{x}_{uij}) - \lambda \|\theta\|^2 \quad (7)$$

其中, $\theta = \{\gamma_u, \gamma_i, \beta_j, W^c, W^p\}$ 包括所有的模型参数; $\sigma(\cdot)$ 是 sigmoid 函数; λ 表示正则化参数, 防止过拟合。

3.3 基于对抗性学习的 CPR 推荐模型

本文算法与基于 GAN 的方法都采用了对抗性学习的思想, 本文算法与 He 等^[20]提出的算法类似, 通过在模型参数上注入对抗性扰动来量化其参数扰动下 CPR 模型的损失, 并使用对抗性学习的方法来训练优化模型参数, 得到一个更具鲁棒性和更高准确度的预测模型。而后者旨在用 GAN 框架来解决协同过滤推荐的问题, 其生成器生成的是用户或物品的购买向量, 判别器将这些向量进行混合, 并尽可能正确地对其进行识别, 该方法表明了协同过滤推荐结合 GAN 框架的有效性。

3.3.1 模型架构

本文在 CPR 模型的基础上进行改进, 其预测模型可以表示为:

$$\hat{x}_{ui} = \alpha + \beta_u + \beta_i + \langle \gamma_u^c + \Delta_u^c, \gamma_i + \Delta_i \rangle + \langle \gamma_u^p + \Delta_u^p, \gamma_i^p + \Delta_i^p \rangle \quad (8)$$

其中, 扰动向量 Δ 与它对应的潜在向量耦合, 如 Δ_u^c 表示消费者的扰动向量。此外, 对抗性扰动的目标是对模型产生最大影响, 也被称为最坏情况扰动^[21]。因此, 我们通过最大化 BPR 损失来找到最佳对抗性扰动:

$$\Delta_{adv} = \arg \max_{\Delta} L_{BPR}(\theta + \Delta) \quad (9)$$

其中, $\|\Delta\| \leq \epsilon$, ϵ 控制对抗扰动的大小, $\|\cdot\|$ 表示 L_2 范数, θ 是模型的参数。

因此, 本文的目的是设计一个新目标函数, 既适合个性化排名, 又可以抵抗对抗性干扰。结合式(7)和式(9)最小化对抗 BPR 损失:

$$L_{ACPR}(\theta) = L_{BPR}(\theta) + \alpha L_{BPR}(\theta + \Delta_{adv}) \quad (10)$$

其中, α 控制对抗性扰动对模型优化的影响, 在极端情况下 ($\alpha=0$) 提出的模型变为原来式(7)中的 BPR 框架。因此, 本文算法可以看作是现有 CPR 模型的推广, 但同时考虑了模型的鲁棒性。

3.3.2 模型优化

由于中间变量 Δ 使目标函数最大化, 而目标函数被 θ 最小化, 因此可以将等式(10)中的优化公式转化为极小极大目标函数:

$$\theta^*, \Delta^* = \arg \min_{\theta} \max_{\Delta, \|\Delta\| \leq \epsilon} L_{BPR}(\theta) + \alpha L_{BPR}(\theta + \Delta) \quad (11)$$

其中, 模型参数 θ 的优化是最小化参与者, 而对抗性扰动 Δ 是最大化参与者。两名参与者交替玩极大极小游戏直到收敛, 我们通过以下方法来解决极小极大优化问题。

(1) 更新 Δ : 给定一个训练实例 (u, i, j) , 可以通过最大化训练数据的 BPR 损失来获得最优扰动 Δ :

$$\max_{\Delta, \|\Delta\| \leq \epsilon} l_{adv}(\Delta) = -\alpha \ln \sigma(\hat{A}_{uij}(\hat{\theta} + \Delta)) \quad (12)$$

其中, $\hat{A}_{uij}(\hat{\theta} + \Delta) = \bar{x}_{ui}(\hat{\theta} + \Delta) - \bar{x}_{uj}(\hat{\theta} + \Delta)$, $\hat{\theta}$ 表示当前模型的常量参数。由于 l_{adv} 的非线性和 ϵ 的约束优化, 我们采用快速梯度符号法将 Δ 附近的目标函数近似为线性函数。可以通过将变量朝其梯度方向移动来获得最优解 Δ :

$$\frac{\partial l_{adv}(\Delta)}{\partial \Delta} = -\alpha \cdot \sigma(-\hat{A}_{uij}(\hat{\theta} + \Delta)) \frac{\partial \hat{A}_{uij}(\hat{\theta} + \Delta)}{\partial \Delta} \quad (13)$$

式(13)具有最大范数约束 $\|\Delta\| \leq \epsilon$, 可以通过式(14)得出 Δ_{adv} 的最优解:

$$\Delta_{adv} = \epsilon \frac{\Gamma}{\|\Gamma\|}, \Gamma = \frac{\partial l_{adv}(\Delta)}{\partial \Delta} \quad (14)$$

根据式(14)可以将式(8)的扰动向量梯度表示为:

$$\frac{\partial \hat{A}_{uij}(\hat{\theta} + \Delta)}{\partial \Delta_u^c} = \gamma_i + \Delta_i + \gamma_u^p + \Delta_u^p - \gamma_j - \Delta_j \quad (15)$$

$$\frac{\partial \hat{A}_{uij}(\hat{\theta} + \Delta)}{\partial \Delta_u^p} = \gamma_u^c + \Delta_u^c \quad (16)$$

$$\frac{\partial \hat{A}_{uij}(\hat{\theta} + \Delta)}{\partial \Delta_i} = \gamma_u^c + \Delta_u^c \quad (17)$$

$$\frac{\partial \hat{A}_{uij}(\hat{\theta} + \Delta)}{\partial \Delta_j} = -\gamma_u^c - \Delta_u^c \quad (18)$$

(2) 更新 θ : 可以通过最小化训练数据的 BPR 损失来获得模型参数 θ :

$$\min_{\theta} l_{ACPR}(\theta) = -\ln \sigma(\hat{A}_{uij}(\theta)) - \alpha \ln \sigma(\hat{A}_{uij}(\theta + \Delta_{adv})) + \lambda \|\theta\|_F^2 \quad (19)$$

其中, Δ_{adv} 是式(14)计算出的常量。关于 θ 的导数为:

$$\frac{\partial l_{ACPR}(\theta)}{\partial \theta} = -\sigma(-\hat{A}_{uij}(\theta)) \frac{\partial \hat{A}_{uij}(\theta)}{\partial \theta} + 2\lambda\theta - \alpha \cdot \sigma(-\hat{A}_{uij}(\theta + \Delta_{adv})) \frac{\partial \hat{A}_{uij}(\theta + \Delta_{adv})}{\partial \theta} \quad (20)$$

在训练过程中, 使用随机梯度下降 (Stochastic Gradient Descend, SGD) 最小化目标函数, 然后通过反向传播更新参数:

$$\theta \leftarrow \theta - \eta \frac{\partial l_{ACPR}(\theta)}{\partial \theta} \quad (21)$$

其中, η 表示学习率。

(3) 初始化: 值得一提的是, 模型参数 θ 的初始化不是随机初始化, 而是通过优化 BPR 得到的。这是因为只有当模型参数开始与数据过度匹配时, 才有必要添加对抗性扰动。当模型欠拟合时, 正常的训练过程就足以获得较好的参数。除了使用 BPR 进行预训练外, 另一种可行的策略是动态调整 ϵ 来控制训练过程中的扰动级别。

根据上述分析, 给出了本文所用算法, 具体描述如算法 1 所示。

算法 1 SGD Learning algorithm

输入: 数据集 D , 潜在因子维度 K , L_2 正则化系数 λ , 对抗噪声 ϵ , 对抗

正则化超参数 α , 学习率 η

输出: 模型参数 θ

1. Initialize θ from original BPR by solving Eq. (7);
2. Initialize Δ randomly, such that $\|\Delta\| \leq \epsilon$;
3. while Stopping criteria is not met do:
4. Randomly draw (u, i, j) from D ;
- // 构造对抗性扰动
5. Update Δ_{adv} by Eq. (14);
- // 更新模型参数
6. Update θ by Eq. (21);
7. end
8. return θ

4 实验及结果分析

4.1 数据集

实验采用来自 UGC 平台上的两个公共数据集 Pinterest 和 Reddit 来评价本文算法的性能。

Pinterest 是一个主要基于图像内容发现的应用程序,包括了用户浏览、上传、喜欢和保存的图像。实验采用 2013 年 1 月抓取的数据集^[22],其中包括 89 万个用户,240 万幅图片和 5600 万个行为(喜欢和保存),将喜欢和保存行为都视为隐式反馈,每个物品都与一个上传者相关联。

Reddit 是一个讨论网站,涵盖了新闻、科学、电影等话题,用户可以提交内容和评论意见。实验所使用的数据集包括了 2017 年 3 月 Reddit 上所有的提交和评论,包含了 130 万个用户、960 万次提交和 4800 万条评论。将每个提交视为一个物品,并将评论的行为视为隐式反馈,每个提交都与一个用户相关联。

两个数据集中的用户、物品行为记录数的详细信息如表 1 所列。

表 1 数据集的详细信息(预处理后)

Table 1 Details in datasets(after preprocessed)

	Pinterest	Reddit
# users ($ U $)	134 747	52 654
# item ($ I $)	201 792	336 743
# actions ($\sum_{u \in U} I_u^+ $)	690 506	1 786 032
# consumer ratio($ C / U $)	93.65%	99.60%
# producer ratio($ P / U $)	80.76%	87.24%
# prosumer ratio($ PS / U $)	74.42%	86.85%

4.2 评价指标

为了评估所提算法的性能,本文采用了 3 个常用的评估指标,分别是 AUC, Precision 和 Recall。

$$AUC = \frac{1}{|U|} \sum_{u \in U} \frac{1}{|D_u|} \sum_{(i,j) \in D_u} \xi(\bar{x}_{ui} - \bar{x}_{uj})$$

(22)

$$D_u = \{(i, j) | (u, i) \in \Gamma_u \wedge (u, j) \notin (P_u \cup V_u \cup \Gamma_u)\}$$

$$Precision = \frac{\sum_u |R(u) \cap T(u)|}{\sum_u |R(u)|}$$

(23)

$$Recall = \frac{\sum_u |R(u) \cap T(u)|}{\sum_u |T(u)|}$$

(24)

其中, $\xi(\cdot)$ 表示指示函数, $R(u)$ 表示测试集上用户的物品集合, $T(u)$ 表示通过推荐算法最终给用户推荐的物品集合, $|R(u)|$ 表示用户推荐的物品数量, $|T(u)|$ 表示用户喜欢的物品数量。

4.3 实验结果

4.3.1 比较算法

为了验证本文算法的有效性,我们将其与以下算法进行了比较。

(1) POP 是一种非个性化的方法,用于对个性化推荐的性能进行基准测试,它根据物品的受欢迎程度(通过训练数据中的交互次数)对项目进行排序。

(2) MF-BPR 是一种协同过滤方法,使用 BPR 目标函数优化 MF,该方法将用户和物品分别以嵌入矩阵的形式表现出来,并用它们的内积来建模用户对一个物品的偏好程度。

(3) FM^[23] 提供了一种通用的因子分解方法,可用于对用户、物品及其特性之间的交互进行建模。

$$\hat{x}_{ui} = \alpha + \beta_u + \beta_i + \beta_p + \langle \gamma_u, \gamma_i \rangle + \langle \gamma_u^c, \gamma_i^c \rangle + \langle \gamma_i^p, \gamma_u^p \rangle$$

(25)

(4) CPR 是为 UGC 平台推荐内容而定制的算法,从相同核心用户嵌入派生出两个角色嵌入并通过消费者、物品和生产者之间的三元关系来预测每个交互行为。

(5) ACPR (Adversarial Consumer and Producer Recommendation) 是本文提出的一种增强模型鲁棒性的算法,融合了 CPR 推荐的三元交互模型和对抗性学习。

4.3.2 模型的鲁棒性对比

为了验证使用对抗性学习的效果,我们选用模型的鲁棒性来进行分析对比。首先,使用 BPR 优化训练 CPR 直至其基本达到收敛;然后,将收敛的 CPR 参数用来初始化 ACPR 并进行迭代训练至其收敛。本文采用了 Reddit 数据集来进行模型鲁棒性的对比,实验结果如表 2 所列。

表 2 对抗性干扰下 CPR 和 ACPR 的性能下降比率(AUC)

Table 2 Performance drop ratio(AUC) of CPR and ACPR in presence of adversarial perturbations

(单位: %)

Algorithm	$\epsilon = 0.6$	$\epsilon = 0.7$	$\epsilon = 0.8$
CPR	-20.7	-24.2	-24.7
ACPR	-0.4	-4.3	-11.0

从表 2 中可知,当 CPR 和 ACPR 加入相同大小的扰动时, CPR 性能下降的幅度大于 ACPR,这表明了 ACPR 模型相比较于基础模型 CPR 更能抵御来自外界的恶意攻击,对于对抗性扰动的敏感性降低,鲁棒性增强。

从另一个方面看,随着 ϵ 的增大,模型性能受影响程度也有所增大,但 ACPR 性能受影响程度总是小于 CPR,这也表明了 ACPR 的鲁棒性优于 CPR,对抗性学习确实能改善模型的鲁棒性。

4.3.3 参数设置

实验中主要参数的选定范围如下: 1) 潜在因子 K 的范围为 $[10, 20, 40, 60, 80]$; 2) 通过在验证集上进行网格搜索来调整 L_2 , 正则化参数为 $[0.0001, 0.001, 0.1]$; 3) 学习率范围为 $[0.001, 0.01, 0.1, 0.5]$; 4) 扰动大小 ϵ 的范围为 $[0.1, 0.3, 0.5, 0.7, 0.9, 1]$; 5) 对抗性正则化项 α 的范围为 $[0.001, 0.01, 0.1, 1, 10, 100, 1000]$ 。

本文实验的运行环境为: Intel(R) Core(TM) i7-8700 CPU, GPU NVIDIA GeForce RTX 2080, Windows10 64bit 操作系统, 实验使用了 Python, Tensorflow 实现相关推荐算法。

(1)潜在因子 K

由图 1 可以看出,当潜在因子 K 为 20 时,ACPR 算法的性能达到最大,Pinterest 和 Reddit 的 AUC 指标分别为 0.728 和 0.931。综合考虑效果后,选取 ACPR 算法的潜在因子 K 为 20。

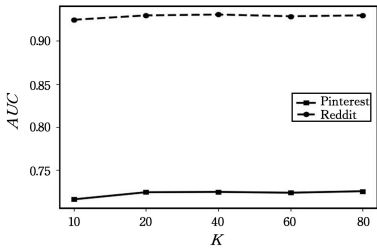


图 1 潜在因子 K 对 ACPR 算法的影响

Fig. 1 Influence of latent factor K on ACPR algorithm

(2)学习率 η

图 2 给出了在两个数据集上,ACPR 算法在不同学习率 η 下的 AUC 指标变化情况。

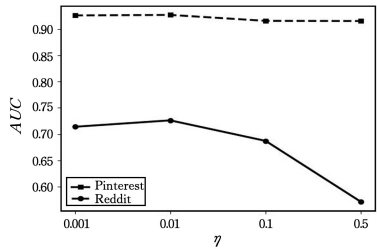


图 2 学习率 η 对算法的影响

Fig. 2 Influence of η on ACPR algorithm

由图 2 可以看出,两个数据集在学习率 η 为 0.01 时,ACPR 算法的 AUC 达到最大,分别为 0.726 和 0.927。因此,选取 ACPR 算法的学习率为 0.01。

(3)对抗扰动 ϵ

我们研究了超参数 ϵ 在 Pinterest 数据集上的影响。在实验中,我们固定 α 为 1,从图 3 可以看出当 ϵ 接近于 0 时,ACPR 的行为类似于 CPR;当 $\epsilon \leq 0.5$ 时,ACPR 的性能趋于平稳;当 ϵ 值超过 0.5 时 ACPR 的性能急剧下降,这说明了过大的扰动会破坏模型参数的学习过程。

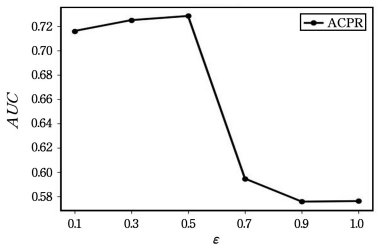


图 3 对抗噪声大小 ϵ 的影响

Fig. 3 Impact of ϵ on ACPR algorithm

(4)对抗正则项 α

从图 3 中可以看出, $\epsilon=0.5$ 时 ACPR 的性能最好。因此,我们固定 ϵ 为 0.5,改变对抗正则项 α 。由图 4 可以看出,当 $\alpha=1$ 时,ACPR 的 AUC 达到最优;在 $\alpha=1$ 的左边增加 α 将提高算法性能,在 $\alpha=1$ 的右边增加 α 将显著降低算法性能;当 $\alpha \leq 0.1$ 时,其 AUC 值变化较小,即推荐精度趋于平稳,

这说明算法在 $\alpha \leq 0.1$ 时对该超参数的选择不敏感。

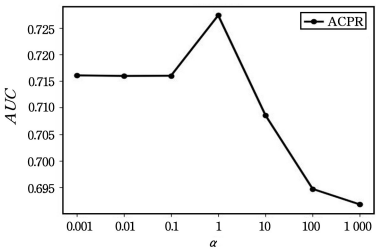


图 4 对抗正则项 α 的影响

Fig. 4 Impact of α on ACPR algorithm

4.3.4 不同算法的实验结果与分析

表 3—表 5 分别以不同的指标比较了 ACPR 算法和对比的基线算法分布在 Reddit 和 Pinterest 两个数据集上的性能,其中 $RI(R)$ 和 $RI(P)$ 分别表示相对于基线算法的改进程度。

表 3 算法对比的实验结果(AUC)

Table 3 Experimental results for algorithm comparison (AUC)

Algorithm	Pinterest	$RI(P)/\%$	Reddit	$RI(R)/\%$
POP	0.613	15.3	0.640	31.2
MF-BPR	0.653	9.8	0.827	11.1
FM	0.687	5.1	0.893	4.0
CPR	0.715	1.3	0.918	1.3
ACPR	0.724	—	0.930	—

表 4 算法对比的实验结果(Precision)

Table 4 Experimental results for algorithm comparison (Precision)

Algorithm	Pinterest(P@5/10)	Reddit(P@5/10)
POP	0.152	0.138
MF-BPR	0.297	0.236
FM	0.331	0.262
CPR	0.419	0.302
ACPR	0.455	0.373

表 5 算法对比的实验结果(Recall)

Table 5 Experimental results for algorithm comparison (Recall)

Algorithm	Pinterest(P@5/10)	Reddit(P@5/10)
POP	0.056	0.097
MF-BPR	0.073	0.155
FM	0.082	0.177
CPR	0.104	0.203
ACPR	0.135	0.286

从表 3—表 5 可以发现:1)Reddit 的整体性能优于 Pinterest,主要原因是 Reddit 用户更倾向于重复使用同一生产者的内容,这使得其行为更容易预测。2)个性化推荐方法的性能优于非个性化推荐的方法。MF-BPR 是一种应用隐式反馈的个性化排序,其结果优于非个性化的 POP 排序方法,这是因为个性化排序会对产品进行建模,而 POP 则是直接根据流行度来进行推荐排名。3)FM/CPR 比 MF-BPR 更准确,这说明了在 UGC 平台上考虑生产者信息的重要性。4)ACPR 的性能优于 CPR,说明加入对抗性学习不仅能改善模型的鲁棒性,还可以提高模型的推荐性能。5)本文提出的 ACPR 算法推荐效果优于上述的其他基线算法,从而验证了本文算法的有效性。

综上所述,本文提出的算法在这两个数据集上均表现出较高的推荐性能。

结束语 本文提出了结合协同过滤推荐和对抗性学习的改进 CPR 推荐算法。研究表明,采用成对学习方法的 BPR 优化的模型,其参数容易受到对抗性干扰,这暗示了使用 BPR

优化的模型存在脆弱性问题。为了学习更强大的个性化排名模型,对 CPR 进行了对抗性训练。本文详细阐述了构建对抗性扰动的方法,并使用基于 SGD 的通用学习算法来优化 CPR。本文所提算法在推荐性能和鲁棒性上都有着一定的优势,但仍存在以下不足:1)算法针对的是浅 MF 模型的嵌入层,并没有在深层隐藏层上使用对抗性训练;2)算法仅适合于 UGC 平台上获取的数据集;3)算法增加了模型在训练时的时间复杂度。

在未来工作中,我们计划进一步克服该算法的上述局限性,将对抗性学习算法扩展到其他模型上,如探索在神经网络 CF 模型上使用对抗性学习,以进一步提高物品推荐的性能。

参 考 文 献

[1] SMITH B,LINDEN G. Two Decades of Recommender Systems at Amazon. com[J]. IEEE Internet Computing, 2017,21(3): 12-18.

[2] SCHAFER J B,FRANKOWSKI D,HERLOCKER J,et al. Collaborative Filtering Recommender Systems[M]// The Adaptive Web. Berlin:Springer, 2007:291-324.

[3] MA W K,LI G,LI Z Y,et al. A Top-N Personalized Recommendation Algorithm Based on Tags[J]. Computer Science, 2019, 46(S2):224-229.

[4] KANG W C,MCAULEY J. Learning Consumer and Producer Embeddings for User-Generated Content Recommendation [C]//Proceedings of the 12th ACM Conference on Recommender Systems. 2018:407-411.

[5] SZEGEDY C,ZAREMBA W,SUTSKEVER I,et al. Intriguing Properties of Neural Networks [J/OL]. Computer Science, 2013. <http://arxiv.org/pdf/1312.6199.pdf>.

[6] ZHANG Z P,GUO X L. Optimized Collaborative Filtering Recommendation Algorithm Based on Item Rating Prediction[J]. Computer Application Research,2008,25(9):2658-2660.

[7] Sarwar B,Karypis G,Konstan J,et al. Item Based Collaborative Filtering Recommendation Algorithms[C]// Proceedings of the 10th International Conference on World Wide Web. 2001:285-295.

[8] HE X,ZHANG H,KAN M Y,et al. Fast Matrix Factorization for Online Recommendation with Implicit Feedback[C]// Proceedings of the 39th International ACM SIGIR Conference on Research and Development in Information Retrieval. 2016:549-558.

[9] HU G N,DAI X Y,QIU F Y,et al. Collaborative Filtering with Topic and Social Latent Factors Incorporating Implicit Feedback [J]. ACM Transactions on Knowledge Discovery from Data (TKDD),2018,12(2):1-30.

[10] WANG C,ZHU H,ZHU C,et al. SetRank: A Setwise Bayesian Approach for Collaborative Ranking from Implicit Feedback[J/OL]. 2020. <http://arxiv.org/pdf/2002.09841.pdf>.

[11] YU W H,ZHANG H D,HE X N,et al. Aesthetic-based Clothing Recommendation[C]// the 2018 World Wide Web Conference. 2018:649-658.

[12] LI X,JIANG M,HONG H,et al. A Time-aware Personalized Point-of-Interest Recommendation via High-order Tensor Factorization[J]. ACM Transactions on Information Systems(TOIS), 2017,35(4):1-23.

[13] WU Y,BAMMAN D,RUSSELL S. Adversarial Training for Relation Extraction[C]// Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing. 2017:1778-1783.

[14] LEE D,KIM J,MOON W J,et al. CollaGAN: Collaborative GAN for Missing Image Data Imputation[C]// Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition. 2019:2487-2496.

[15] WANG J,YU L,ZHANG W,et al. IRGAN: A Minimax Game for Unifying Generative and Discriminative Information Retrieval Models[C]// Proceedings of the 40th International ACM SIGIR Conference on Research and Development in Information Retrieval. 2017:515-524.

[16] WANG H,JIA W,WANG J,et al. GraphGAN: Graph Representation Learning with Generative Adversarial Nets[J]. IEEE Transactions on Knowledge and Data Engineering, 2017(99): 2508-2515.

[17] CHAE D K,KANG J S,KIM S W,et al. CfGAN: A Generic Collaborative Filtering Framework Based on Generative Adversarial Networks[C]// Proceedings of the 27th ACM International Conference on Information and Knowledge Management. 2018:137-146.

[18] KOREN Y. Factorization Meets the Neighborhood: A Multifaceted Collaborative Filtering Model [C] // Proceedings of the 14th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining. 2008:426-434.

[19] RENDLE S,FREUDENTHALER C,GANTNER Z,et al. BPR: Bayesian Personalized Ranking from Implicit Feedback [C] // Proceedings of the 25th Conference on Uncertainty in Artificial Intelligence. 2009:452-461.

[20] HE X N,HE Z,DU X,et al. Adversarial Personalized Ranking for Recommendation[C]// The 41st International ACM SIGIR Conference on Research &. Development in Information Retrieval. 2018:355-364.

[21] GOODFELLOW I J,SHLENS J,SZEGEDY C. Explaining and Harnessing Adversarial Examples[J]. arXiv:1412.6572,2014.

[22] ZHONG C T,SHAH S,SASTRY N,et al. Sharing the Loves: Understanding the How and Why of Online Content Curation [C]// Proceedings of Seventh International AAAI Conference on Weblogs and Social Media Sharing. 2013:659-667.

[23] RENDLE S. Factorization Machines[C]// 2010 IEEE International Conference on Data Mining. IEEE,2010:995-1000.



ZHAN Wan-jiang, born in 1994, post-graduate. His main research interests include recommendation system and data mining.



LYU Yue-hua, born in 1978, master. His main research interests include data mining and artificial intelligence.