



研究与开发

基于学生知识追踪的多指标习题推荐算法

诸葛斌, 尹正虎, 斯文学, 颜蕾, 董黎刚, 蒋献

(浙江工商大学信息与电子工程学院, 浙江 杭州 310020)

摘要: 个性化习题推荐是教育信息化时代的重要课题, 传统的习题推荐算法忽略了学生在学习过程中的遗忘规律, 未能充分挖掘学生的知识掌握水平和相似学生的共性特征, 推荐习题的评价指标单一, 推荐习题的新颖度和多样性不足, 不能合理地促进学生对新知识的学习或帮助学生查缺补漏。针对上述缺陷, 提出一种基于学生知识追踪的多指标习题推荐方法, 该方法分为习题初筛和再过滤两个模块, 围绕习题推荐的新颖度、难度以及多样性 3 个指标展开研究, 首先构造了一个结合学生遗忘规律的知识概率预测 (student forgetting behavior based knowledge concept coverage prediction, SF-KCCP) 模型, 保证推荐习题的新颖性; 接着基于动态键值的知识追踪 (dynamic key-value memory networks for knowledge tracing, DKVMN) 模型精准挖掘学生的知识概念掌握水平, 以保证推荐合适难度的习题; 最后将基于用户的协同过滤 (user-based collaborative filtering, UserCF) 算法融入再过滤模块, 利用学生群体之间的相似性实现推荐结果的多样性。通过大量的实验证明, 所提方法比一些现有的基线模型取得了更好的性能。

关键词: 深度学习; 习题推荐; 知识追踪; 协同过滤

中图分类号: TP301

文献标志码: A

doi: 10.11959/j.issn.1000-0801.2022234

Student knowledge tracking based multi-indicator exercise recommendation algorithm

ZHUGE Bin, YIN Zhenghu, SI Wenxue, YAN Lei, DONG Ligang, JIANG Xian

College of Information and Electronic Engineering, Zhejiang Gongshang University, Hangzhou 310020, China

Abstract: Personalized exercise recommendation was an important topic in the era of education informatization, the forgetting laws of students in the learning process were ignored by the traditional problem recommendation algorithm, which failed to fully tap the students' knowledge mastery level and the common characteristics of similar students, insufficient, could not reasonably promote students' learning of new knowledge or help students find and fill omissions. In view of the above defects, a multi-index exercise recommendation method based on student knowledge

收稿日期: 2022-03-21; 修回日期: 2022-06-27

通信作者: 蒋献, jiangxianxd123@126.com

基金项目: 浙江省重点研发计划项目 (No.2021C01036); 浙江省自然科学基金资助项目 (No.LY18F010006); 浙江省新型网络标准与应用技术重点实验室项目 (No.2013E10012); 浙江工商大学高等教育研究课题 (No.Xgy21012)

Foundation Items: The Key Research and Development Program of Zhejiang Province (No.2021C01036), Natural Science Foundation of Zhejiang Province (No.LY18F010006), Zhejiang Key Laboratory of New Network Standards and Application Technology (No.2013E10012), Higher Education Research Project of Zhejiang Gongshang University (No.Xgy21012)



tracking was proposed, which was divided into two modules: preliminary screening and re-filtering of exercises, focusing on the novelty, difficulty and diversity of exercise recommendation. Firstly, a knowledge probability prediction (SF-KCCP) model combined with students' forgetting law was constructed to ensure the novelty of the recommended exercises. Then, students' knowledge and concept mastery level was accurately excavated based on the dynamic key-value knowledge tracking (DKVMN) model to ensure that exercises of appropriate difficulty were recommended. Finally, the user-based collaborative filtering (UserCF) algorithm was integrated into the re-filtering module, and the similarity between student groups was used to achieve the diversity of recommendation results. The proposed method is demonstrated by extensive experiments to achieve better performance than some existing baseline models.

Key words: deep learning, exercise recommendation, knowledge tracking, collaborative filtering

0 引言

随着教育信息化和移动端应用的加速普及,在线教育打破了传统教学的时空壁垒,学习者可以随时随地进行线上学习,海量的习题和教学资源能够充分满足不同层次学习者的学习需求,吸引大批用户来到线上平台学习。在这个信息爆炸的时代,线上教育资源也在呈指数级增长,然而这些资源的质量参差不齐,极大地增加了学习者选择学习资源的难度,而个性化推荐技术恰能解决这一难题。因此,如何利用推荐技术从海量的线上资源中为学习者提供精准、个性化的习题推荐,让学习者能更高效地学习,是在线教育急需解决的问题。

传统的习题推荐技术主要分为两类,一类是基于内容和基于邻域的协同过滤推荐^[1-2],这是推荐系统中应用最为广泛的推荐算法,主要根据学生的历史学习记录,找到与目标学生兴趣相似的近邻学生,无须考虑习题的具体属性。其中 Li 等^[3]搭建了一个在线教育平台,通过协同过滤算法为用户推荐个性化课程;吴云峰等^[4]针对协同过滤推荐法中用户的历史信息不足等问题,提出基于多分类器的迁移 Bagging 习题推荐算法;Segal^[5]等设计了一种针对学生的个性化教育算法,首先通过 UserCF 对相似学生进行排序,再为目标学生的习题构建一个“难度”排名,可辅助教师为学生布置个性化作业和考试。虽然基于 UserCF 的推荐方

法结构简单、可解释性强,但这类算法忽略了不同学生知识掌握水平的差异,只能粗粒度地对学生进行分类,无法满足智慧教育的个性化习题推荐要求。

另一类是基于认知诊断的推荐技术,如 Ozaki^[6]提出了传统 DINA 模型,该方法将习题的知识点考查矩阵作为输入,得到学生的认知掌握向量,然后根据学生认知情况向目标学生个性化推荐习题。Wang 等^[7]提出了一个通用神经认知诊断(neural cognitive diagnosis, NeuralCD)框架,结合神经网络学习复杂的运动交互,将学生和练习投射到因子向量上,并利用多个神经层对其交互进行建模,以获得准确且可解释的诊断结果。Liu 等^[8]提出了一种模糊认知诊断方法框架(fuzzy cognitive diagnosis framework, FuzzyCDF),用于考生对客观和主观问题的认知建模,结合模糊集理论和教育假设来模拟考生的行为,通过考虑失误和猜测因素对每个问题进行评分。项目反应理论(item response theory, IRT)是一种经典的认知诊断方法,它可以为分析学生表现提供可解释的参数(即学生潜在特质、问题辨别和难度)。Cheng 等^[9]提出了一个通用的深度项目反应理论(deep item response theory, DIRT)框架,利用深度学习提取问题文本的语义表征,从而增强传统的认知诊断 IRT。基于认知诊断的方法能够考虑学生的知识状态,再基于学生个人的知识水平做出推荐,但是没有将相似学生的信息加入参考,

忽略了学习的群体性, 推荐的多样性受限。

为了弥补上述两种算法的不足, 本文提出了一种基于学生知识追踪的多指标习题推荐 (student knowledge tracking based multi-indicator exercise recommendation, SKT-MIER) 方法, 首先对学生知识概念学习情况进行深入研究, 利用 LSTM 神经网络对学生知识点学习范围进行预测, 以便对未答知识点或答错知识点的相关习题进行优先推荐; 利用更先进的深度知识追踪模型对学生的知识点掌握情况进行建模, 从而得到学生当前知识状态, 预测其在未来学习中的表现; 其次融入 UserCF 算法, 结合学生相似群体之间的学习情况, 从而提高推荐的多样性。综上所述, 本文的主要创新和贡献如下:

- 提出了一个结合学生遗忘规律的知识概率预测模型 SF-KCCP, 得到学生的知识点学习范围向量, 保证推荐习题的新颖性, 有利于学生学习新知识;
- 基于 DKVMN 模型, 得到学生的知识点掌握水平向量, 确保推荐习题难度适当, 有利于激发学生的学习热情;
- 将上述两组向量值为计算依据, 利用协同过滤算法得到相似学生, 结合学生自身和其相似学生的知识水平进行推荐, 提高习题推荐的多样性, 有利于帮助学生未掌握的知识点进行查漏补缺;
- 在 4 个真实开源的数据集上对本文提出的推荐模型进行评估, 并与其他主流模型进行对比, SKT-MIER 在相应的评价指标上均有所提高。

1 深度知识追踪概述

知识追踪是根据学生过去的答题情况对学生的知识掌握情况进行建模, 从而得到学生当前知识状态表示的一种技术。早期的知识追踪模型是基于一阶马尔可夫模型的, 但其有着明显的缺陷,

随着技术的进步, 有研究将深度学习 (deep learning, DL) 技术引入知识追踪领域, 最早基于 DL 的知识追踪模型来自美国斯坦福大学的深度知识追踪 (deep knowledge tracing, DKT)^[10], 文献[10]提出利用递归神经网络 (recursive neural network, RNN) 对学生的学习情况进行建模, 将学生过去的学习活动压缩在隐藏层中, 映射学生对不同知识概念的掌握水平, 以此预测学生未来的答题正确率。DKT 在不需要过多专家经验和大量特征工程的情况下, 精度优于传统方法。针对 DKT 算法预测结果的一致性与波动性问题, Yeung 等^[11]提出了在 DKT 算法的损失函数中加入 3 个正则化项的方法, 增强了算法预测的精度, 使知识状态转移更为平滑。Zhang 等^[12]在练习层面加入了更多的特征, 并使用自动编码器将高维信息转换为低维特征, 提升了模型性能。为了更好地把握学生做题情况的历史关联性, 神经教学代理^[13] (neural pedagogical agent, NPA) 模型利用双向长短期记忆网络对学生知识状态进行建模, 该网络配备了一个注意层用来对学生的学习历史进行重点分析, 从而得出更准确的预测。

除了基于 RNN 架构的知识追踪模型, Zhang 等^[14]提出了基于动态键值对网络的知识追踪模型 DKVMN。它借鉴了记忆增强神经网络的思想, 依赖于外部存储器中的读写操作, 比 RNN 和 LSTM 更适合对大量的学生答题记录进行建模, 解决了 RNN 类模型难以评估学生对每个知识点的掌握程度的问题。DKVMN 模型网络结构如图 1 所示, 其包含一个读过程和一个写过程, 使用键值对存储器结构而不是单个矩阵。其中包含了一个静态的键 (key) 矩阵 M^k , 用来存储所有习题包含的所有知识点, 是不可变的; 一个动态的值 (value) 矩阵 M^v , 根据擦除向量和增加向量存储并更新学生对于各个知识点的掌握情况, 是可变的。

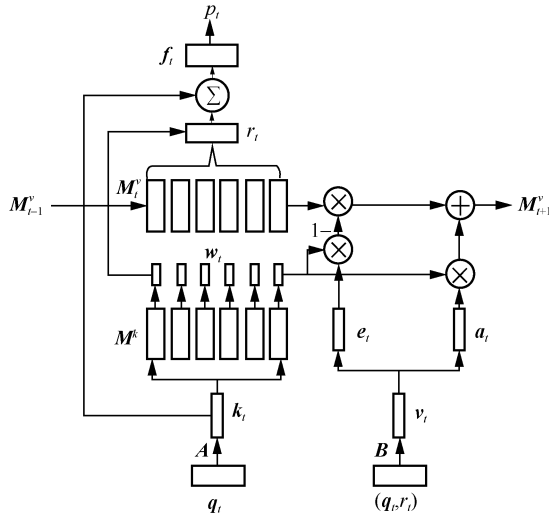


图1 DKVMN 模型网络结构

在每个时间戳，DKVMN 模型有两组输入，第一组输入为离散的练习题标签 q_t ，对应输出正确答案 q_t 的概率 $p(r_t=1|q_t)$ ，第二组输入为题目和回答元组 (q_t, r_t) ，用来更新下个时间步的值矩阵。其中， q_t 属于具有若干个不同练习的集合 Q ， r_t 为该题是否正确回答的二值表示。假设 Q 包含 m 个知识点概念 $\{k_1, k_2, k_3, \dots, k_m\}$ ，那么这些概念存储在键矩阵 M^k （大小为 $m \times d_k$ ）中。学生对每个知识点的掌握程度，即知识点掌握水平状态 $\{S_t^1, S_t^2, S_t^3, \dots, S_t^m\}$ 则会存储在值矩阵 M_t^v 中（大小为 $m \times d_v$ ），显然该矩阵随时间变化，输入学生的每一条做题记录作为更新依据。 k_t 为习题嵌入向量，由输入 q_t 和嵌入矩阵 A 相乘得到。 k_t 与 M^k 矩阵中每个知识点向量进行内积，再通过激活函数 Softmax 计算权重向量 w_t ，它代表了习题和每个概念之间的关系：

$$w_t(i) = \text{softmax}(k_t^T M^k(i)) \quad (1)$$

其中， $\text{softmax}(z_i) = \frac{e^{z_i}}{\sum_j e^{z_j}}$ ，读过程和写过程都

使用此权重向量，表示当前习题与键矩阵中每个知识点之间的相关性。 v_t 为学生知识增长向量，由输入元组 (q_t, r_t) 和嵌入矩阵 B 相乘得到，经过擦除和增加操作，对值矩阵进行更新。

(1) 读过程

读过程可以实现对学生答题情况的预测，当练习题 q_t 来临时，首先通过计算权重向量 w_t 得到该习题包含哪些知识点，继而将 w_t 与值矩阵求内积，得到该学生对于答对这道题所需的知识点的掌握情况，如果掌握了大部分知识点，则估计该学生答对这道题的可能性较大，具体计算式如下：

$$r_t = \sum_{i=1}^m w_t(i) M_t^v(i) \quad (2)$$

尽管有些题目包含的知识点内容相同，但每个习题其自身的难度不一，因此将学生对练习的掌握程度 r_t 和输入练习嵌入向量 k_t 连接，再通过具有 tanh 激活的完全连接层获得总结向量 f_t ，就包含了习题本身的难度：

$$f_t = \tanh(W_1^T [r_t, k_t] + b_1) \quad (3)$$

其中， W_1^T 表示待训练的权重矩阵并转置， b_1 表示待训练的偏置参数向量， $\tanh(z_i) = \frac{e^{z_i} - e^{-z_i}}{e^{z_i} + e^{-z_i}}$ ，再将 f_t 经过一个具有 Sigmoid 激活函数的全连接层预测学生答对概率 p_t ，它是一个标量，表示正确答案 q_t 的概率：

$$p_t = \sigma[W_2^T f_t + b_2] \quad (4)$$

(2) 写过程

写过程根据学生的实际答题情况实现对学生知识点掌握情况的更新，体现在值矩阵的动态改变上。输入答题记录元组 (q_t, r_t) 得到学生的知识增长向量 v_t ，在对值矩阵进行更新之前，先模仿 LSTM 的输入门和遗忘门进行记忆擦除和新增操作，从而更符合现实中学生的遗忘规律和记忆强化规律。

计算擦除向量 e_t ：

$$e_t = \sigma(E^T v_t + b_e) \quad (5)$$

其中， E 为变换矩阵，使得 e_t 成为一组列向量且每个元素取值均在 (0,1) 内。上一时间步的值矩阵经过擦除之后如下：

$$\tilde{\mathbf{M}}_t^v(i) = \mathbf{M}_{t-1}^v(i)[1 - \mathbf{w}_t(i)\mathbf{e}_t] \quad (6)$$

计算更新向量 $\mathbf{a}_t$ ，其中， $\mathbf{D}^T$ 表示需训练的权重矩阵， $\mathbf{b}_a$ 表示需训练的偏置参数向量：

$$\mathbf{a}_t = \tanh(\mathbf{D}^T \mathbf{v}_t + \mathbf{b}_a)^T \quad (7)$$

其中， $\mathbf{a}_t$ 为行向量，经过更新之后得到下一个时间步的值矩阵：

$$\mathbf{M}_{t+1}^v(i) = \tilde{\mathbf{M}}_t^v(i) + \mathbf{w}_t(i)\mathbf{a}_t \quad (8)$$

如式 (9) 所示，在训练过程中，嵌入矩阵 $\mathbf{A}$ 和 $\mathbf{B}$ 以及其他参数和 $\mathbf{M}^k$ 、 $\mathbf{M}_t^v$ 的初始值通过最小化 p_t 和真实标签 r_t 之间的标准交叉熵损失来共同学习：

$$\mathcal{L} = -\sum_i (r_i \log p_i + (1 - r_i) \log(1 - p_i)) \quad (9)$$

2 基于学生知识追踪的多指标习题推荐方法

本文提出的多指标的习题推荐系统框架如图 2 所示，采用一种先筛选后过滤的模式，首先根据 SF-KCCP 模型从原始题库中筛选符合当前学生知识点学习范围的习题，得到预备题库。在此基础上，采用基于用户的协同过滤算法过滤，将学生知识点学习范围向量和学生知识点掌握程度向量作为相似学生的评判指标，然后计算学生间的余弦相似度，选择相似度排名大于一定阈值（通常为 0.95）的学生作为相似用户，计算相似用户的平均知识点掌握程度，最后根据设置的期望难度，过滤得到推荐习题集。

2.1 习题初筛模块

对于目标学生在某个做题时间 t ，首先通过初筛模块生成符合学生知识点学习范围的习题库 $\mathbf{EB}_p$ 。习题初筛结构如图 3 所示。首先将学生做题记录 $\mathbf{X}_t$ 输入 SF-KCCP 模型对学生已学知识点情况进行建模，预测下一时刻学生需要学习的知识点向量 $\mathbf{P}^F(\mathbf{k}_t^c)$ ，并将其与原始习题库中每道习题真实包含的知识点向量 $\mathbf{K}_i$ 进行计算，得到它们的相似性分数 ω_i ，相似性分数越高，越符合学生需要学习的知识范围，最终选择符合要求的前 N 个习题 e 组成 $\mathbf{EB}_p$ 。算法 1 描述了 $\mathbf{EB}_p$ 的计算过程。相似性分数计算式如下：

$$\omega_i = \cos \theta[\mathbf{P}^F(\mathbf{k}_t^c), \mathbf{K}_i] \quad (10)$$

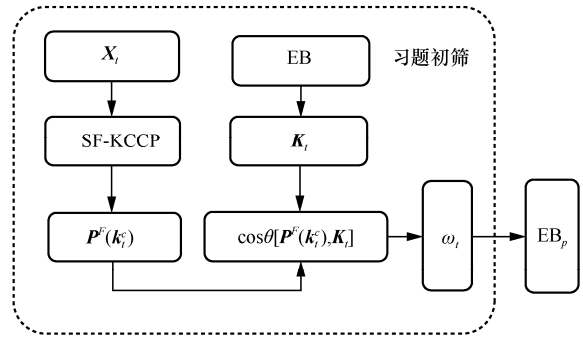


图 3 习题初筛结构

2.1.1 融入遗忘规律的知识点学习范围预测模型

著名的艾宾浩斯遗忘曲线理论展示了学生对所学知识的遗忘规律^[15]，遗忘的速度呈现先快后慢的趋势。遗忘带来的影响是学生知识点学习范

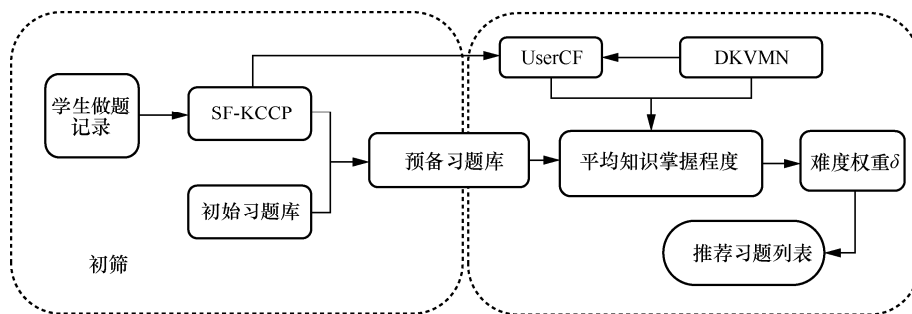


图 2 多指标的习题推荐系统框架

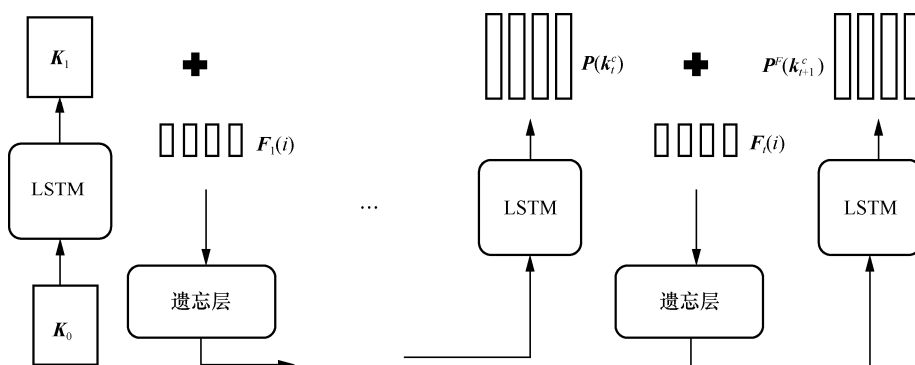


图4 SF-KCCP 模型

围的缩小, 之前学过的知识点也会随着时间的流逝被逐渐遗忘, 从而转变为未学知识点。影响遗忘的具体因素有学生重复学习知识点的次数与距离上次学习该知识点的时间间隔。SF-KCCP 模型如图 4 所示, 其由本文将遗忘规律定量地融入预测模型中而得到。

2.1.2 知识点学习范围预测

利用知识概念从 0 到 t 的出现顺序来预测每个知识概念在 $t+1$ 时刻的出现概率, 因此知识点学习范围预测是一个基于顺序的预测问题, 可以采用 LSTM 获得序列预测模型。基于 LSTM 的知识点学习范围预测模型如图 5 所示, 其中标记为 A 的圆角矩形表示模型在时刻 t 的状态。输入矩形 K_i 表示在时刻 t 出现的知识概念, 从做题记录 X_i 中得到; 输出矩形 K_{t+1} 表示通过模型预测得到的在时间 $t+1$ 知识概念出现的概率。

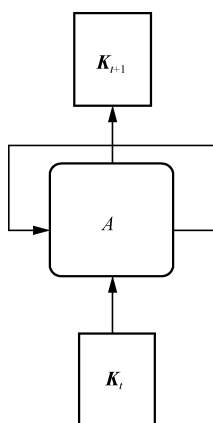


图5 基于 LSTM 的知识点学习范围预测模型

在被输入网络之前, 先对 K_i 进行独热编码^[16], 结果表示为 $\varphi(K_i)$ 。模型的输出 K_{t+1} 是一个维度等于所有知识概念数量的向量, 表示对下一时刻知识概念出现的预测, 而 $t+1$ 时刻模型的输入记为 $\varphi(K_{t+1})$ 即真实的知识点学习向量。依据 LSTM 模型, 在一个训练周期 T 结束后对知识点学习范围预测问题进行建模, 使用二元交叉熵损失构造一个损失函数对模型进行训练。

$$\mathcal{L} = \sum_{t=0}^T \ell(K_{t+1} \cdot \varphi(K_{t+1}), 1) \quad (11)$$

最终模型的输出是一个向量, 其长度等于课程中所有知识概念的数量, 表示为:

$$\mathcal{F}(K) = [\mathcal{F}(k_1), \mathcal{F}(k_2), \dots, \mathcal{F}(k_m)] \quad (12)$$

其中, $\mathcal{F}(k_i)$ 表示第 i 个知识概念在 $t+1$ 时刻出现的概率。为了满足推荐的新颖性, 让已经回答正确的知识概念在以后的练习中尽量少出现, 模型赋予学生答题过程中未答知识点和回答错误的知识点更大的权重, 即在 LSTM 网络的输出层添加一个权重向量, 其长度等于知识概念的数量:

$$\omega(K) = [\omega(k_1), \omega(k_2), \dots, \omega(k_m)] \quad (13)$$

$\omega(k_i)$ 可以表示为:

$$\omega(K_i) = \begin{cases} 1 - \frac{r_i}{c_i}, & c_i \neq 0 \\ 1, & c_i = 0 \end{cases} \quad (14)$$

其中, r_i 表示某学生在答题时间段 T 内正确回答知识点 K_i 的次数, c_i 表示 T 时间段内知识点 K_i 出现的次数。通过 LSTM 模型的输出与权重向量的

结合可以得到下次练习中知识概念的出现概率向量, 即 $P(k_{t+1}^c)$ 。

$$P(k_{t+1}^c) = \mathcal{F}(K^c) \cdot \omega(K^c) \quad (15)$$

2.1.3 融入学生遗忘规律

将距离上次学习相同知识点的时间间隔因素记为 SK (same knowledge), 将重复学习相同知识点的次数因素记为 RT (repeat times), 这两个标量组合得到向量 $F_t(i) = [SK_t(i), RT_t(i)]$, 表示学生对知识点 i 的遗忘程度。受 LSTM 网络的启发, 采用全连接网络计算记忆擦除向量和记忆更新向量来拟合遗忘行为, 与第 2.1.2 节中得到的 $P(k_{t+1}^c)$ 相结合得到新的知识点学习范围 $P^F(k_{t+1}^c)$ 。

遗忘处理过程如图 6 所示, 将影响学生遗忘的因素 $F_t(i)$ 与 t 时刻学生的知识点学习范围嵌入矩阵 $P(k_t^c)$ 进行遗忘处理, 得到学生本次学习开始时的知识点学习范围嵌入矩阵 $P^F(k_t^c)$ 。受 LSTM 中遗忘门与输入门的启发, 根据影响知识遗忘的因素 $F_t(i)$ 更新 $P^F(k_t^c)$ 时, 首先要擦除 $P(k_t^c)$ 中原有的信息, 再写入信息。其中, 擦除过程控制学生知识点学习范围的衰减, 写入过程控制学生知识点学习范围的更新。

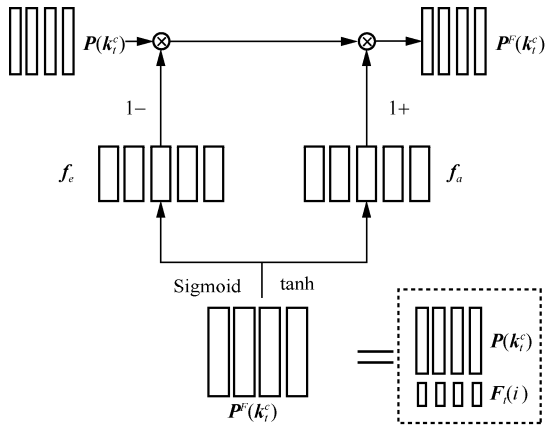


图 6 遗忘处理过程

利用一个带有 Sigmoid 激活函数的全连接层实现擦除过程, 具体计算式为:

$$f_e(i) = \sigma[W_e F_t(i) + b_e] \quad (16)$$

其中, σ 表示 Sigmoid 激活函数, W_e 为权重

矩阵, 形状是 $(d_v + d_c) \times d_v$, 全连接层的偏置向量 b_e 为 d_v 维。

利用一个带有 tanh 激活函数的全连接层实现更新过程, 具体计算式为:

$$f_a(i) = \sigma[W_a F_t(i) + b_a] \quad (17)$$

根据擦除过程与更新过程对 $P(k_t^c)$ 更新得到 $P^F(k_t^c)$:

$$P^F(k_t^c) = P(k_t^c)(1 - f_e(i))(1 + f_a(i)) \quad (18)$$

习题初筛模块的算法描述如下。

算法 1 习题的初筛算法

输入: X_t, EB, N {做题记录, 习题库, 需要的题数}

输出: EB_p {预备习题库}

$P^F(k_t^c) \leftarrow \text{SF-KCCP}(X_t)$ {计算下次需要学习的只是概念向量}

for $i \leq \text{size}(EB)$ do

$K_t^i \leftarrow \text{get}(EB, i)$ {从题库中获取习题 i 的只是概念向量}

$\omega_t^i \leftarrow \cos \theta[P^F(k_t^c), K_t^i]$ {计算相似分数}

end for

sort(EB, ω_t) {根据相似分数 ω_t 对习题排序}

$EB_p \leftarrow \text{Select}(EB, N)$ {筛选前 N 道满足条件的习题}

2.2 习题再过滤模块

2.2.1 计算相似学生用户

本节介绍了如何基于 SF-KCCP 模型与 DKVMN 模型得到目标学生的相似用户 S' , 再通过相似用户计算每个目标学生的平均知识掌握程度 g_i 。算法 2 描述了 g_i 的计算过程。其中, $P^F(k^c)(m \times n)$ 表示确定时刻学生的知识点学习范围矩阵, m 表示学生个数, n 表示知识点数量; $p(m \times z)$ 表示学生知识掌握水平矩阵, z 表示题目数量。假设学生 i 与学生 j 的 $P_i^F(k^c) = \{a_1, a_2, \dots, a_n\}$, $p_i = \{b_1, b_2, \dots, b_z\}$; $P_j^F(k^c) = \{c_1, c_2, \dots, c_n\}$,



$p_j = \{d_1, d_2, \dots, d_z\}$, 则学生 i 与学生 j 的余弦相似度计算式如下:

$$\cos \theta = \frac{\sum_{i=1}^n a_i \cdot c_i}{2\sqrt{\sum_{i=1}^n (a_i)^2} \cdot \sqrt{\sum_{i=1}^n (c_i)^2}} + \frac{\sum_{i=1}^z b_i \cdot d_i}{2\sqrt{\sum_{i=1}^z (b_i)^2} \cdot \sqrt{\sum_{i=1}^z (d_i)^2}} \quad (19)$$

对于目标学生 i , 其最终的平均知识掌握程度计算式如下:

$$g_i = \rho(p_i) + (1 - \rho) \cdot \text{average}(p_{s'}) \quad (20)$$

其中, $\text{average}(p_{s'})$ 为与学生 i 相似的所有学生的知识掌握水平向量平均值, g_i 由学生个人的知识水平 p_i 和相似学生的共性知识水平 $\text{average}(p_{s'})$ 按一定的比例相加得到, 比例由参数 ρ 调节, ρ 的取值范围为 $[0, 1]$ 。 ρ 越大, 学生个人知识水平产生的影响越大, 这里 $\rho = 0.7$ 。

计算学生的平均知识掌握程度算法描述如下。

算法 2 学生的平均知识掌握程度算法

输入: $P^F(k^c)(m \times n)$, $p(m \times z)$, S {学生的知识学习范围矩阵, 知识掌握水平矩阵, 学生}

输出: g {所有学生的平均知识掌握水平}

for $i \leq \text{size}(S)$ **do**

for $j \leq \text{size}(S)$ **do**

$$\cos_j \theta \leftarrow \frac{\sum_{i=1}^n a_i \cdot c_i}{2\sqrt{\sum_{i=1}^n (a_i)^2} \cdot \sqrt{\sum_{i=1}^n (c_i)^2}} + \frac{\sum_{i=1}^z b_i \cdot d_i}{2\sqrt{\sum_{i=1}^z (b_i)^2} \cdot \sqrt{\sum_{i=1}^z (d_i)^2}} \quad \{\text{计算学生的余弦相似度}\}$$

end for

$\text{sort}(S, \cos \theta)$ {根据 $\cos \theta$ 对学生排序}

$S' \leftarrow \text{Select}(S, \cos \theta \geq 0.95)$ {选择相似度大于 0.95 的学生}

$g_i \leftarrow \rho(p_i) + (1 - \rho) \cdot \text{average}(p_{s'})$ {计算学生 i 的平均知识掌握水平}

end for

2.2.2 再过滤算法

在进行习题推荐时, 考虑学生个人的知识掌握水平既可以使推荐结果更具针对性, 又能对推荐题目的难度把控更到位。而加入 UserCF 算法可以结合学生相似群体之间的学习情况, 提高推荐结果的多样性。

为了达到上述目的, 习题再过滤结构如图 7 所示, 设置习题再过滤模块用于对初筛后的预备题库进行二次过滤。 g_i 表示学生 i 的平均知识掌握程度向量, 每一位的值越大代表学生 i 答对该题的概率越大, 即该题对该学生来说难度越小, 因此可以使用 $1 - g_i$ 表示学生 i 的习题预测难度向量。其中 δ_i 表示推荐题目的期望难度, 它可以是一个适当的难度范围, 表示向学生推荐正确概率为 δ_i 的题目。标有 Distance 的圆角矩形是一个计算习题预测难度和期望难度 δ_i 之间距离的过程, 预测难度即 $1 - g_i$ 。在获得每个练习的预测难度与期望难度之间的距离 $|\delta_i - (1 - g_i)|$ 后, 通过设定最大权重值 Ω_i 对练习进行排序, 并选择所有小于 Ω_i 的 n 个练习组成最终的输出练习集 $\{re_1, re_2, \dots, re_n\}$ 。

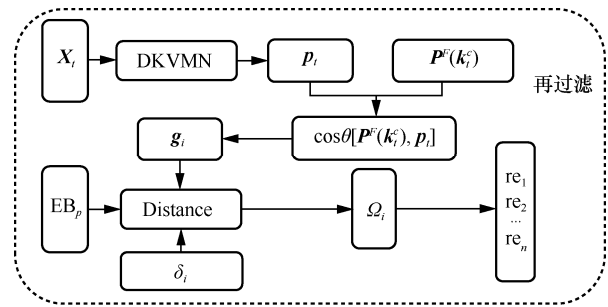


图 7 习题再过滤结构

3 实验分析与讨论

3.1 数据集

本文的实验部分使用了 3 个数据集作对比分析, 分别是 ASSISTment2009、Statics 2011 以及本

文作者团队自主开发的智慧在线教育系统 DouDouYun 的数据集, 3 个数据集的缩写分别是 ASST09、STATICS、DouDY, 数据集简介见表 1。

表 1 数据集简介

数据集	学生数/个	习题数/道	知识概念数/个	答题记录数/个
ASSISTment2009	4 151	110	378	325 637
Statics 2011	333	1 223	1 362	189 297
DouDouYun	8 333	1 287	300	5 166 342

3.2 对比模型

实验将本文提出的推荐算法与以下 3 个基线模型进行对比, 分别是基于学生的协同过滤^[17] (student-based collaborative filter, SB-CF)、基于习题的协同过滤^[18] (exercise-based collaborative filter, EB-CF) 以及马骁睿等^[19]提出的基于协同过滤的深度知识追踪 (deep knowledge tracking based on collaborative filtering, DKT-CF) 模型方法。

(1) SB-CF

该模型参考协同过滤的思想, 为特定用户寻找感兴趣的内容, 在习题推荐中基于学生相似度进行推荐, 根据学生的做题记录构建学生之间的相似度矩阵, 然后识别得到与目标学生答题情况最为相似的前 10 名学生, 再从每个相似学生的答题记录中抽取适合难度的题目进行推荐。

(2) EB-CF

该模型参考物品之间相似度进行推荐的思想, 利用项目的内在品质或者固有属性进行推荐, 在习题推荐中根据学生的习题作答情况为每个习题设置难度权重, 然后计算习题相似矩阵并从中提取和已做练习相似的习题, 再根据想要的习题难度权重进行推荐。

(3) DKT-CF

该模型通过深度知识追踪模型 DKT 对学生知识状态建模, 再结合相似学生的信息进行协同过滤推荐, 使推荐结果既考虑学生个人的知识状

态又考虑群体学生学习的共性, 提高了推荐结果的可解释性和准确性。

3.3 评价指标

本文结合多种算法实现了一个多指标的习题推荐系统, 多指标指的是练习的难度、新颖度和多样性。其中, 习题难度会影响学生答题的正确率, 从而影响学生的使用体验, 所以难度对于推荐习题的质量把控起到最为关键的作用^[20]; 推荐习题的新颖程度能够给学生提供一种惊喜感, 有助于学生对新知识点的学习; 多样性则能够方便学生查漏补缺, 对于所有知识点进行一个更为全面的学习^[21]。

(1) 准确性

在实际应用中, 可以通过答题准确性 (accuracy) 评估习题难度。如式 (21) 所示, r_i^{re} 表示学生对所推荐习题的作答结果 (正确为 1, 不正确为 0), $\sum_{i=1}^n N(r_i^{\text{re}} = 1)$ 表示在推荐的 n 道习题中学生作答正确的题数, 做对题数除以总题数表示答对概率, 通过判断答对概率与设定值 δ (在实验中, 该值为 0.7) 的接近程度来衡量推荐效果, 越接近则推荐效果越好。

$$\text{Accuracy} = 1 - \left| \delta - \frac{\sum_{i=1}^n N(r_i^{\text{re}} = 1)}{n} \right| \quad (21)$$

(2) 新颖性

第 2 个指标为习题新颖性 (Novelty)^[22], 若推荐习题中包含越多学生未答或答错知识点, 则推荐的习题新颖性越高。在介绍新颖性计算式之前先了解一下 Jaccard 系数的概念, Jaccard 相似系数^[23]用于比较有限样本集之间的相似性与差异性, 系数值越大, 样本相似度越高。下列计算式中 $e(k_i) = 1$ 表示习题 $e(k)$ 包含第 i 个知识点, $E(k)$ 表示习题 $e(k)$ 包含的所有知识概念; 相对的 $E(k)^r$ 表示正确回答的习题 $e(k)^r$ 包含的知识点; $1 - \text{Jaccsim}[E_i(k), E_i(k)^r]$ 为 Jaccard



距离的计算式,表示未答或答错知识点占有知识点的比重,通过对推荐习题列表中的每一道题进行 Jaccard 距离加权平均值计算就可以得到推荐习题的新颖性值。

$$E(k) = \{k_i | e(k_i) = 1\} \quad (22)$$

$$E(k)^r = \{k_i | e(k_i)^r = 1\} \quad (23)$$

$$\text{Novelty} = \frac{\sum_{i=1}^n 1 - \text{Jaccsim}[E_i(k), E_i(k)^r]}{n} \quad (24)$$

(3) 多样性

第 3 个指标为习题多样性 (Diversity), 多样性表示推荐练习所包含知识点间的差异, 用两个习题的知识点向量之间的余弦相似度进行定量表示^[24]。首先计算推荐列表中每道题之间的知识点差异度, 再对推荐的 n 道题进行差异度均值计算得到多样性值, 计算式如下:

$$\text{different}[e(k)_i, e(k)_j] = 1 - \frac{e(k)_i \cdot e(k)_j}{e(k)_i e(k)_j} \quad (25)$$

$$\text{Diversity} = \frac{\sum_{i=1, j=1, i \neq j}^n \text{different}[e(k)_i, e(k)_j]}{n} \quad (26)$$

3.4 实验及结果分析

3.4.1 实验环境

实验环境操作系统为 Windows10 64 bit, 语言为 Python3.6, 采用 PyTorch1.8.1 版本, GPU 型号为 RTX 2080 Ti。

3.4.2 模型参数设置

在 DKVMN 模型训练阶段, 学习速率和迭代次数设置为 0.005 和 50, 使用 Adam 优化算法,

使用 Glorot 初始化所有可训练的参数^[25]。对于模型批尺寸 (batchsize) 的设置, 通过在测试数据集上的平均 AUC 得分来确定, AUC 的值越大表明模型的性能越强, 不同批尺寸下 AUC 值对比结果见表 2。ASSIST2009 和 DouDouYun 数据集的批尺寸为 32 时, AUC 值最高; STATICS 数据集的批尺寸为 8 时, AUC 值最高, 这是由于该数据集缺乏随机噪声引起的。通过模型在测试数据集上的平均 AUC 得分来确定知识点嵌入向量的维度 d_k 和学生知识掌握程度嵌入向量的维度 d_v 的超参数设置。为了减少参数数量, 一般设 $d_k = d_v = d$ ^[26]。

3.4.3 模型预测性能对比分析

实验采用五重交叉验证的形式, 每轮将数据集随机分成两份, 其中 70% 作为模型的训练数据, 10% 作为验证数据, 剩余 20% 的数据集用作整个推荐模型的测试数据。为了度量 DKVMN 算法的性能并与现有知识追踪模型——贝叶斯知识追踪 (Bayesian knowledge tracking, BKT)、DKT 进行比较, 本文采用知识追踪领域的标准度量 AUC (曲线下面积)、ACC (准确率)、RMSE (均方根误差) 作为性能衡量指标, AUC 和 ACC 的值越大表明模型的预测性能越强, RMSE 的值越小表明模型的预测性能越好。模型预测性能对比见表 3, 给出了 DKVMN 与另外两个模型在 3 个数据集的平均 AUC 值、平均 ACC 值和平均 RMSE 值的对比结果。表 3 表明 DKVMN 在 3 个数据集上的平均 AUC、ACC、RMSE 都优于另外两种知识追踪模型, 这也证明了 DKVMN 在预测学生答题表现上性能优于现有模型。BKT 在 3 个模型中性能

表 2 不同批尺寸下 AUC 值对比结果

模型	ASST2009		STATICS2011		DouDouYun	
	批尺寸	AUC	批尺寸	AUC	批尺寸	AUC
DKVMN	8	0.791	8	0.859	8	0.782
	16	0.824	16	0.820	16	0.837
	32	0.845	32	0.789	32	0.852
	64	0.813	64	0.764	64	0.819

表 3 模型预测性能对比

模型	AUC			ACC			RMSE		
	ASST09	STATICS	DouDY	ASST09	STATICS	DouDY	ASST09	STATICS	DouDY
BKT	0.713	0.742	0.725	0.718	0.735	0.729	0.457	0.426	0.446
DKT	0.769	0.814	0.786	0.762	0.801	0.783	0.446	0.415	0.430
DKVMN	0.845	0.859	0.852	0.819	0.837	0.831	0.427	0.398	0.421

最低, 这表明对学生的知识掌握状态直接二值化是有缺陷的。DKT 的预测性能也低于 DKVMN, 这是因为其使用 RNN 虽然能够获得学生整体的知识水平状态, 却无法准确预测对某个知识点的掌握水平。可以看到在不同的数据集上模型的精确度不同且差异较大, 这与数据集中包含的题库类别、单位学生做题数及单位习题包含知识概念数息息相关。随着预测精度的提高, 推荐模型可以更好地确保推荐的效果, 因此采用上述训练得到的最优模型进行组合得到本文提出的 SKT-MIER 习题推荐模型。

3.4.4 实验结果

为了比较各习题推荐模型的效果, 实验使用上文提到的 4 个数据集, 其中 20% 的学生数据用于测试推荐系统的效果。首先将这 20% 的学生做题记录输入 SF-KCCP 模型, 得到学生知识点学习

范围向量, 根据范围向量得到预备习题库。在习题再过滤阶段, 根据设置的期望难度 δ_i 和预测难度 g_i 之间的距离对预备习题库进行重排序, 选择前 n 道习题进行推荐。

习题推荐的准确性见表 4, 习题推荐多样性见表 5, 习题推荐的新颖性见表 6, 分别展示了上述 4 种模型在推题准确性、新颖性和多样性上的结果。Mean 值表示推荐列表中相应指标的平均值, 该值越接近 1, 性能越好。Std.D 表示各指标的标准差, 该值越接近 0, 推荐效果越稳定。

表 4 反映了 4 种模型在 4 个数据集上推荐习题的准确性, SKT-MIER 模型相较其他 3 种模型的推题准确度是最高的。在 DouDouYun 数据集中, 模型 DKT-CF 的稳定性略优于本文所提模型, 这是

表 4 习题推荐的准确性

准确性	ASST2009		STATICS2011		DouDouYun	
	Mean	Std.D	Mean	Std.D	Mean	Std.D
SB-CF	0.493	0.213	0.688	0.210	0.708	0.202
EB-CF	0.655	0.101	0.390	0.213	0.533	0.255
DKT-CF	0.796	0.121	0.748	0.135	0.772	0.113
SKT-MIER	0.894	0.076	0.864	0.074	0.847	0.122

表 5 习题推荐的多样性

多样性	ASST2009		STATICS2011		DouDouYun	
	Mean	Std.D	Mean	Std.D	Mean	Std.D
SB-CF	0.678	0.162	0.649	0.138	0.661	0.202
EB-CF	0.745	0.147	0.776	0.102	0.428	0.174
DKT-CF	0.857	0.084	0.847	0.079	0.670	0.219
SKT-MIER	0.889	0.089	0.899	0.063	0.847	0.121



表6 习题推荐的新颖性

新颖性	ASST2009		STATICS2011		DouDouYun	
	Mean	Std.D	Mean	Std.D	Mean	Std.D
SB-CF	0.386	0.211	0.720	0.158	0.733	0.115
EB-CF	0.599	0.157	0.561	0.148	0.720	0.120
DKT-CF	0.685	0.212	0.827	0.113	0.826	0.104
SKT-MIER	0.863	0.077	0.899	0.077	0.928	0.073

由于该数据集学生数较多,且 SKT-MIER 在计算学生 i 的平均知识掌握水平 $\bar{g}_i$ 时,需要考虑相似学生的知识掌握水平。这里设置了较大的相似学生的权重系数 ρ , 所以推荐的稳定性受到影响。

表 5 反映了 4 种模型在不同数据集推荐习题多样性上的差异,可以发现,当数据集的知识概念较少时,生成的推荐习题集的多样性也会降低。在 ASST2009 数据集中,模型 DKT-CF 的稳定性略优于本文所提模型,此处与表 4 的原因相同,因为 SKT-MIER 设置了较大的相似学生的权重系数 ρ , 所以推荐的稳定性受到影响。

表 6 反映了 4 种模型在 3 个数据集上推荐习题新颖性上的差异,习题数量越多则习题的新颖性越高。

习题推荐准确性如图 8 所示,习题推荐多样性如图 9 所示,习题推荐新颖性如图 10 所示,图 8~图 10 中不同的折线表示不同模型在推题准确性、多样性和新颖性上的结果,横坐标为推题组数,每组推荐 10 道题。

从图 8 的实验结果可知,传统的习题推荐算法 SB-CF 和 EB-CF 在推荐的准确性上比较差,而 DKT-CF 和 SKT-MIER 模型的推题准确性有显著的提高,这证明了 DKT 在提高推荐准确性上的有效性。SKT-MIER 采用了基于增强神经网络的 DKVMN 知识追踪网络,可以更精准地挖掘学生的知识掌握水平,所以习题推荐的准确性明显优于其他 3 个模型。

习题推荐多样性如图 9 所示。值得注意的是本文提出的 SKT-MIER 模型在 3 个数据集的多样性结果比 DKT-CF 模型平均高出 8.7%,这是因为本文采

用了学生知识点学习范围向量和学生知识掌握水平向量作为双重指标判定相似学生,可选学生范围更广,可供推荐的知识点更加多元,所以推荐的习题更加多样化。

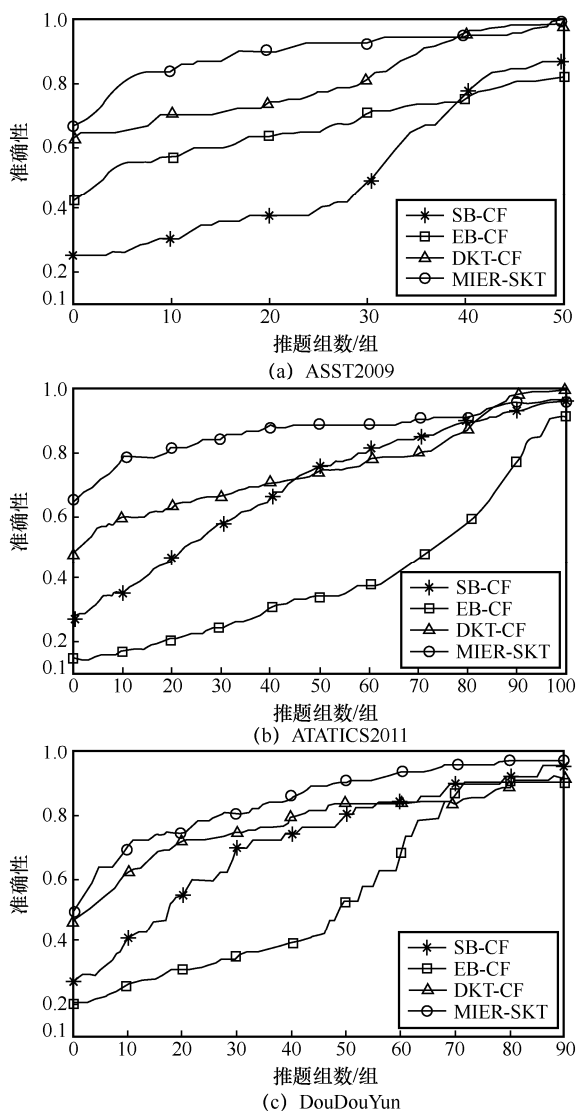


图8 习题推荐准确性

习题推荐新颖性如图 10 所示, 显示了 SKT-MIER 模型在推题新颖性方面有较好的表现。传统推荐算法忽略了学习中的遗忘特性, 在不断推题过程中无法根据学生知识掌握情况的改变而变化, 导致推题的新颖性不足。学生的遗忘行为会导致先前学过的知识点在一定程度上转变为未学知识点, 因此本文在 SKT-MIER 中融入了学生的遗忘规律, 使推荐习题中包含更多未学或遗忘的知识点, 所以在推题新颖性上优于传统模型。综上所述, 本文所提方法比其他 3 种基线模型有明显的优势。

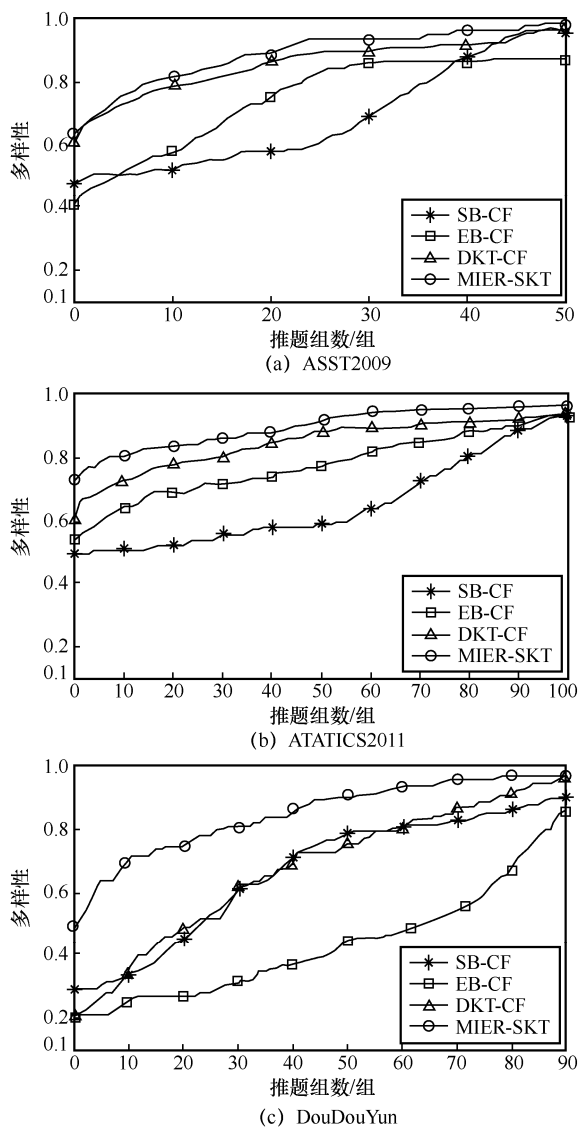


图9 习题推荐多样性

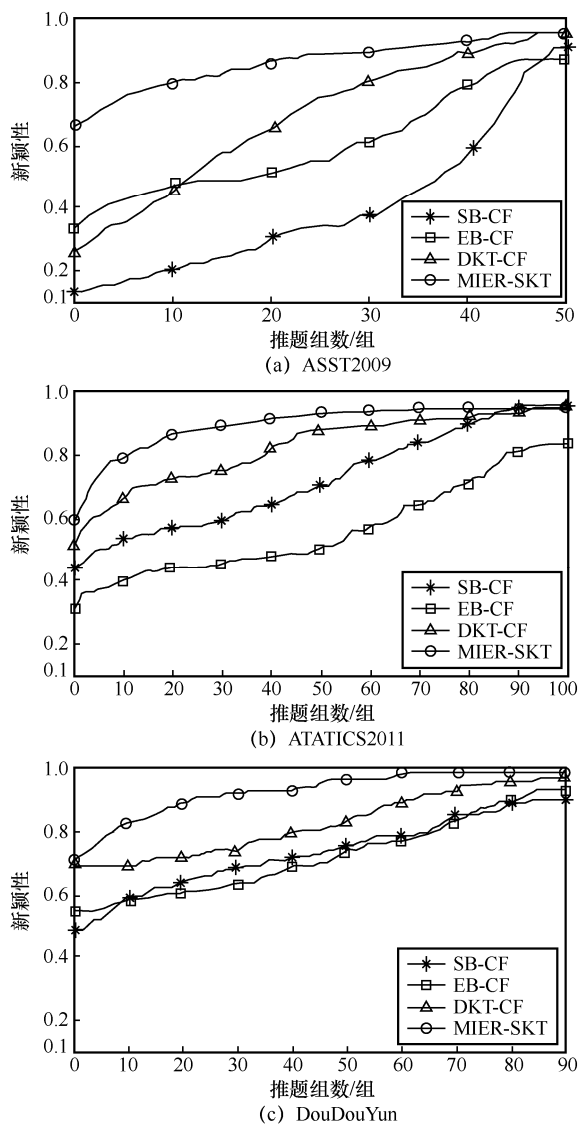


图10 习题推荐新颖性

4 结束语

针对传统的习题推荐算法忽略了学生的遗忘规律, 不能合理地将学生知识掌握水平与相似学生的共性特征进行结合, 推荐习题的评价指标单一, 新颖性和多样性不足, 不能合理地促进学生对新知识的学习或帮助学生查缺补漏等问题, 本文结合深度知识追踪模型和基于用户的协同过滤算法, 提出一种基于学生知识追踪的多指标习题推荐方法。该方法首先从学生已经学习的知识点范围出发, 提出了基于学生遗忘规律的知识点学习范围预测模型,



从而满足习题推荐新颖性,合理地促进学生学习新知识;然后从学生对知识点的掌握程度出发,利用 DKVMN 模型追踪学生的知识掌握水平,从而实现习题的个性化精准推荐,合理地激发学生的学习热情;最后从学生群体之间的相似性出发,将 UserCF 算法融入再过滤模块,实现推荐习题的多样性,合理地帮助学生查缺补漏。最后本文作者团队在几个经典教育数据集上验证了 SKT-MIER 方法的优势,证实了该方法的有效性及其合理性。

虽然本文提出的习题推荐模型优于多个现有方法,但仍有许多可以改进的地方,在未来可以朝以下几个方向深入研究。

(1) 对题目文本进行语义分析

建立学生行为和习题内容之间的语义联系,进一步解决推荐时的冷启动^[27]问题。

(2) 进行多类型习题推荐研究

前大多数模型的研究以选择题居多,答题结果也只有对或错两类,在实际情况中习题类型丰富多样,有选择填空、问答以及计算题等,答题结果并非只有 0 和 1 的情况,对于多样化的习题类型,知识图谱^[28]技术会是一个不错的选择。

(3) 提高推荐算法的效率

为了充分挖掘学生行为数据,实现精细化推荐,本文研究极大地提高了算法复杂度,在未来可采用更先进的模型或优化流程简化模型复杂度,实现快速推荐。

(4) 探索知识追踪多场景应用

在实际学习中除了习题、知识点、答题结果等数据,还有如做题尝试次数、题目解析查看情况以及题目收藏情况等标签,当前对于这些标签的使用场景研究还不够深入,未来可持续探索。

参考文献:

- [1] BOBADILLA J, ALONSO S, HERNANDO A. Deep learning architecture for collaborative filtering recommender systems[J]. Applied Sciences, 2020, 10(7): 2441.
- [2] NGUYEN L V, HONG M S, JUNG J J, et al. Cognitive similarity-based collaborative filtering recommendation system[J]. Applied Sciences, 2020, 10(12): 4183.
- [3] LI J, YE Z. Course recommendations in online education based on collaborative filtering recommendation algorithm[J]. Complexity, 2020: 6619249.
- [4] 吴云峰, 冯筠, 孙霞, 等. 基于多分类器的迁移 Bagging 习题推荐[J]. 计算机应用, 2013, 33(7): 1950-1954.
WU Y F, FENG Y, SUN X, et al. Online transfer-Bagging question recommendation based on hybrid classifiers[J]. Journal of Computer Applications, 2013, 33(7): 1950-1954.
- [5] SEGAL A, KATZIR Z, GAL K, et al. Edurank: a collaborative filtering approach to personalization in e-learning[C]//Proceedings of Conference on Educational Data Mining. [S.l.:s.n.], 2014.
- [6] OZAKI K. DINA models for multiple-choice items with few parameters: considering incorrect answers[J]. Applied Psychological Measurement, 2015, 39(6): 431-447.
- [7] WANG F, LIU Q, CHEN E H, et al. Neural cognitive diagnosis for intelligent education systems[J]. Proceedings of the AAAI Conference on Artificial Intelligence, 2020, 34(4): 6153-6161.
- [8] LIU Q, WU R Z, CHEN E H, et al. Fuzzy cognitive diagnosis for modelling examinee performance[J]. ACM Transactions on Intelligent Systems and Technology, 2018, 9(4): 1-26.
- [9] CHENG S, LIU Q, CHEN E H, et al. DIRT: deep learning enhanced item response theory for cognitive diagnosis[C]//Proceedings of the 28th ACM International Conference on Information and Knowledge Management. [S.l.:s.n.], 2019: 2397-2400.
- [10] PIECH C, BASSEN J, HUANG J, et al. Deep knowledge tracing[J]. Advances in Neural Information Processing Systems, 2015(28).
- [11] YEUNG C K, YEUNG D Y. Addressing two problems in deep knowledge tracing via prediction-consistent regularization[C]//Proceedings of the 5th Annual ACM Conference on Learning at Scale. New York: ACM Press, 2018: 1-10.
- [12] ZHANG L, XIONG X L, ZHAO S Y, et al. Incorporating rich features into deep knowledge tracing[C]//Proceedings of the 4th (2017) ACM Conference on Learning @ Scale. New York: ACM Press, 2017: 169-172.
- [13] LEE Y, CHOI Y, CHO J, et al. Creating a neural pedagogical agent by jointly learning to review and assess[J]. arXiv preprint arXiv:1906.10910, 2019.
- [14] ZHANG J, SHI X, KING I, et al. Dynamic key-value memory networks for knowledge tracing[C]//Proceedings of the 26th International Conference on World Wide Web. [S.l.:s.n.], 2017: 765-774.
- [15] ZHOU G R, ZHU X Q, SONG C R, et al. Deep interest network for click-through rate prediction[C]//Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining. New York: ACM Press, 2018: 1059-1068.
- [16] 申情, 郭文宾, 楼俊钢, 等. 考虑多层次潜在特征的个性化推荐模型[J]. 电信科学, 2022, 38(2): 71-83.
SHEN Q, GUO W B, LOU J G, et al. Personalized recommendation model with multi-level latent features[J]. Telecommunications Science, 2022, 38(2): 71-83.
- [17] LIU G, HAO T. User-based question recommendation for ques-

- tion answering system[J]. International Journal of Information and Education Technology, 2012, 2(3): 243.
- [18] WEI S, YE N, ZHANG S, et al. Item-based collaborative filtering recommendation algorithm combining item category with interestingness measure[C]//Proceedings of 2012 International Conference on Computer Science and Service System. Piscataway: IEEE Press, 2012: 2038-2041.
- [19] 马骁睿, 徐圆, 朱群雄. 一种结合深度知识追踪的个性化习题推荐方法[J]. 小型微型计算机系统, 2020, 41(5): 990-995.
- MA X R, XU Y, ZHU Q X. Personalized exercises recommendation method based on deep knowledge tracing[J]. Journal of Chinese Computer Systems, 2020, 41(5): 990-995.
- [20] LV P, WANG X, XU J, et al. Utilizing knowledge graph and student testing behavior data for personalized exercise recommendation[C]//Proceedings of ACM Turing Celebration Conference-China. New York: ACM Press, 2018: 53-59.
- [21] PARDOS Z A, HEFFERNAN T. KT-IDEM: Introducing item difficulty to the knowledge tracing model[C]//Proceedings of International Conference on User Modeling, Adaptation, and Personalization. Heidelberg: Springer, 2011: 243-254.
- [22] KAMINSKAS M, BRIDGE D. Diversity, serendipity, novelty, and coverage: a survey and empirical analysis of beyond-accuracy objectives in recommender systems[J]. ACM Transactions on Interactive Intelligent Systems, 2017, 7(1): 2.
- [23] NEDUNGADI P, REMYA M S. Incorporating forgetting in the personalized, clustered, bayesian knowledge tracing (PC-BKT) model[C]//Proceedings of 2015 International Conference on Cognitive Computing and Information Processing(CCIP). Piscataway: IEEE Press, 2015: 1-5.
- [24] WANG Z, ZHU J L, LI X, et al. Structured knowledge tracing models for student assessment on coursera[C]//Proceedings of the 3rd (2016) ACM Conference on Learning @ Scale. New York: ACM Press, 2016: 209-212.
- [25] NAGATANI K, ZHANG Q, SATO M, et al. Augmenting knowledge tracing by considering forgetting behavior[C]//Proceedings of WWW '19: The World Wide Web Conference. [S.l.:s.n.], 2019: 3101-3107.
- [26] HUO Y J, WONG D F, NI L M, et al. Knowledge modeling via contextualized representations for LSTM-based personalized exercise recommendation[J]. Information Sciences, 2020 (523): 266-278.
- [27] ZHAO J J, BHATT S, THILLE C, et al. Cold start knowledge tracing with attentive neural turing machine[C]//Proceedings of the Seventh ACM Conference on Learning @ Scale. New York: ACM Press, 2020: 333-336.
- [28] 周晶, 孙喜民, 于晓昆, 等. 知识图谱与数据应用: 智能推荐[J]. 电信科学, 2019, 35(8): 165-172.
- ZHOU J, SUN X M, YU X K, et al. Knowledge graph and data application: intelligent recommendation[J]. Telecommunications Science, 2019, 35(8): 165-172.

[作者简介]



诸葛斌(1976-), 男, 博士, 浙江工商大学信息与电子工程学院教授, 主要研究方向为网络和通信技术、互联网技术和网络安全。



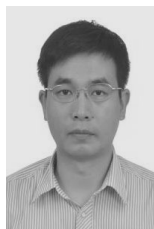
尹正虎(1997-), 男, 浙江工商大学信息与电子工程学院硕士生, 主要研究方向为在线教育、自然语言处理和个性化推荐。



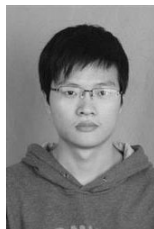
斯文学(1996-), 女, 浙江工商大学信息与电子工程学院硕士生, 主要研究方向为在线教育、数据分析和个性化推荐。



颜蕾(1996-), 女, 浙江工商大学信息与电子工程学院硕士生, 主要研究方向为教育数据挖掘、深度学习和机器学习。



董黎刚(1972-), 男, 博士, 浙江工商大学信息与电子工程学院教授, IEEE 和 IEEE-CS 成员, 中国电子学会高级会员, 浙江计算机学会理事, 主要研究方向为智能网络、在线教育。



蒋献(1988-), 男, 浙江工商大学信息与电子工程学院讲师、实验员, 主要研究方向为在线教育。