

引入激活加权策略的分组排序学习方法

李玉轩^{1,2+}, 洪学海^{1,3}, 汪 洋¹, 唐正正^{1,2}, 班 艳¹

1. 中国科学院 计算机网络信息中心, 北京 100190

2. 中国科学院大学, 北京 100049

3. 中国科学院 计算技术研究所, 北京 100190

+ 通信作者 E-mail: liyuxuan@cnic.cn

摘 要: 排序学习(LtR)将有监督机器学习技术(SML)用于解决排序问题,旨在给出输入文档列表的相关度更优化的排序结果。此前关于深度排序模型的研究,对于列表内文档的相关度计算彼此独立,缺乏考虑文档之间的相互作用。近年来一些新方法致力于挖掘文档之间的相互影响,如分组评分法(GSF),通过学习多元变量评分函数来联合判断文档相关性,但大多忽略了文档间相互影响的差异性,同时增加了很大的计算代价。针对此问题,提出了一种带权重的分组深度排序模型(W-GSF)。该方法借鉴推荐领域的深度兴趣网络,引入其根据候选商品调整历史行为序列权重的思想,在排序学习中多元评分法基础上,以多层前馈神经网络为主体结构,并在输入端加入激活单元,利用神经网络自适应学习调整输入的多元变量的权重,来挖掘交叉文档关系的差异性。在公共基准数据集MSLR上的实验验证了该方法的有效性,相比基线排序模型,激活策略的引入带来了排序指标上的明显提升,同时相对于同等效果的排序方法计算量大幅降低。

关键词: 排序学习(LtR); 分组评分法(GSF); 深度神经网络; 深度兴趣网络

文献标志码: A **中图分类号:** TP391

Groupwise Learning to Rank Algorithm with Introduction of Activated Weighting

LI Yuxuan^{1,2+}, HONG Xuehai^{1,3}, WANG Yang¹, TANG Zhengzheng^{1,2}, BAN Yan¹

1. Computer Network Information Center, China Academy of Sciences, Beijing 100190, China

2. University of China Academy of Sciences, Beijing 100049, China

3. Institute of Computing Technology, Chinese Academy of Sciences, Beijing 100190, China

Abstract: Learning to rank (LtR) applies supervised machine learning (SML) technologies to the ranking problems, aiming at optimizing the relevance of input document list. As regard to previous studies on the deep ranking model, the calculation of the relevance of the documents in the list is independent of each other, which lacks consideration of document interactions. In recent years, some new methods are devoted to mining the interaction between documents, such as groupwise scoring function (GSF), which learns multivariate scoring function to jointly judge the correlation, but most of these methods ignore the differences of the interaction between documents, and bring high calculation cost at the same time. In order to solve this problem, this paper proposes a weighted groupwise deep ranking model (W-GSF). In view of the deep interest network in the field of recommendation, this paper introduces the idea of adjusting the weight of historical behavior sequence according to the candidate products. On the basis of multivariate scoring method in learning to rank field, this method uses multi-layer feed forward neural networks as main structure, and adds an activation unit into it before the input module, taking advantage of neural networks to adjust the weight of input multiple variables adaptively, so as to mine the differences of cross document

relationship. Experiments on the public benchmark dataset MSLR verify the effectiveness of the method. Compared with baseline ranking models, the introduction of activation strategy brings a significant improvement of ranking metrics, and the computational complexity is greatly reduced compared with the same effect learning to rank methods.

Key words: learning to rank (LtR); groupwise scoring function (GSF); deep neural network; deep interest network

随着信息量级的增长,如何获取更符合检索目的的高质量信息,已成为信息检索(information retrieval, IR)领域中的重要课题。排序学习(learning to rank, LtR)^[1]则用于在搜索过程中,对召回的候选文档列表进行精确排序。由于其效果显著,排序学习方法已逐渐应用到实际的搜索场景当中。其任务是对于给定的查询,对一系列文档进行相关度打分,进而按分数高低给出排序后的列表,这意味着相关度越高,文档的排名应越靠前。不同于其他的分类和回归问题,排序学习并不关注文档的具体评分,而更加关注个体之间的相对顺序。

排序学习算法通常以查询信息和文档的向量化特征作为输入,通过监督学习的方法,最优化目标函数,得到将特征向量映射到实值的评分函数,并以此作为排序依据。按输入或按损失函数的不同形式可以将排序学习分为单文档法(pointwise)^[2]、文档对法(pairwise)^[3]和文档列表法(listwise)^[4-5]。新提出的序列文档法^[6]、分组评分法(groupwise scoring function, GSF)^[7]等都属于这三种方法的变种。

现存大部分的排序学习方法虽然有效,但往往局限于单变量的评分函数,即列表中文档的相关性计算相互独立,缺乏考虑文档之间的相互影响。而用户在与检索结果的交互过程中,会表现出强的对比选择倾向。相关研究工作已证明通过对比文档得出的偏好判断比直接判断更有效^[8],当用户行为被一种相对的方式建模时,模型有更好的预测能力^[9]。因此,列表中文档的相关性排序位置除取决于本身因素之外,还受其上下文特征的影响。Ai等人^[7]基于上述论点,认为文档的相关性得分通过与列表中其他项在特征级别上联合计算,提出了多元变量的相关度评分模型GSF,即将文档进行分组,组内文档的特征向量共同作用,并利用深度神经网络(deep neural network, DNN)自动挖掘跨文档信息。但是不足之处在于当分组较大时,模型复杂度和计算量大大增加,导致该方法在对于时效性要求较高的搜索推荐场景中缺乏吸引力。其次,GSF的分组操作对于组内的每个文档同等对待,忽略了文档间相互影响的差异

性。举例来讲,用户在对比返回结果时,更多的是对同类型文档的择优,而差异较大但又同时与查询信息相关的文档间交叉影响较小。换句话说,不同文档对单个样本的相关度评价影响是不同的。基于这样的推断假设,同时也启发于推荐领域的深度兴趣网络^[10],本文在GSF中引入激活单元对多元变量进行加权,并用神经网络自适应学习权重,文中表示为W-GSF(weighted-GSF)。通过在公共基准的排序学习数据集上的对比实验,证明该方法相对于基线模型有更好的效果。同时,对于达到同样水准的排序学习方法,本文模型在计算代价上优势明显。

1 相关工作

1.1 排序学习

排序学习方法用于推断给定列表的排名顺序,包含评分和排序两个过程,其目标是构造一个评分函数,通过该函数可以归纳出文档列表的排序。该领域内的大部分研究可以从两方面进行分类:评分函数和损失函数的构造。不同的评分函数即不同的模型结构,包括线性模型^[11]、梯度提升树^[3,12]、支持向量机^[13]和神经网络^[6-7]等;损失函数的构造针对定义问题的形式和解决问题的思路,例如单文档法、文档对法和文档列表法。

单文档方法将单篇文档的特征向量映射为得分;文档对法则关注在给定两个文档时判断相对顺序^[14];文档列表法则将单次查询对应的整个文档列表作为整体输入来训练,可以直接优化核心排序指标,如NDCG(normalized discounted cumulative gain)^[15]等。近几年,得益于深度学习的长足发展,深度排序模型层出不穷。许多研究已经表明,单文档的方法缺乏考虑文档之间的相互关系,对于相关性评价的建模不够完善。文档对方法以LambdaMART^[3]为代表,巧妙地转化文档位置顺序关系为概率模型,使之可以直接利用机器学习模型进行学习。文档列表法对于相关性建模更加全面,但是往往计算复杂度高,且具有完整排序信息的数据较难获取,实际应用中立落地场景较少。

1.2 GSF 简介

排序学习领域最近提出的多元变量评分模型 GSF^[7]是一种通过将文档进行分组,学习交叉文档关系的深度排序模型,不同于文档对法和文档列表法专注于损失函数的构造,GSF 是一种对模型结构提出改进的新方法。

图 1 给出了 GSF 的模型示意图。示例样本为含有三个文档的列表,分组大小为 2,分组策略给出了文档列表所有可能的二元组。作为输入,分组内的文档特征向量会做拼接操作成为单个向量,每个分组产生的向量分别输入多层前馈神经网络,前向计算后输出对应的文档评分。其过程相对于传统的单元变量评分函数,简而言之就是将 $g(\cdot)$ 从单文档映射到单实数,改进为多文档映射到多实数。由于一个文档会出现在多个分组之中,故在经过 $g(\cdot)$ 输出各分组评分之后,需要将多次评分结果累和,作为最终输出的相关度评分。分组策略保证公平性,各个文档出现次数相同,故不会影响最终评价的相对顺序。

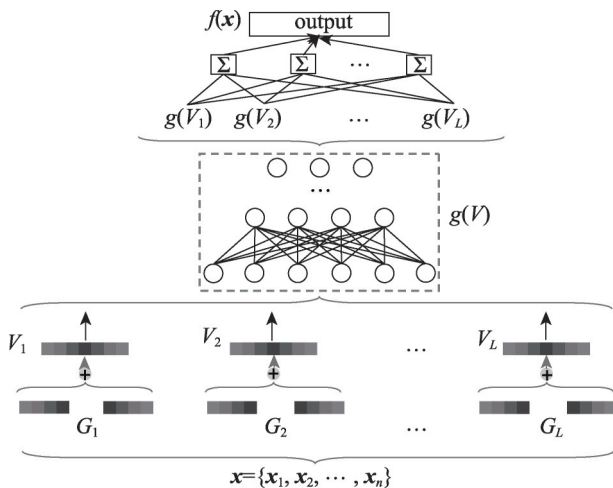


图1 基本的GSF模型结构(以只含有三个文档的列表为例)

Fig.1 Structure of basic GSF model, with a sample which only has three documents

1.3 深度兴趣网络

深度兴趣网络(deep interest network, DIN)^[10]用于解决推荐系统领域的点击率预估问题。传统点击率预估模型以 Embedding+DNN 结构为基础,对于用户的历史行为数据不加区分地向量化,这导致在判断不同的候选商品时,用户向量的表示相同,限制了模型的表达能力。

为了增加模型的表达能力,DIN 认为历史的行为

序列中不同商品对于判断候选商品的点击率影响存在差异,故在传统点击率预估模型的基础上,加入了激活单元,可以根据候选商品对历史兴趣中的商品赋予不同权重,使得模型在对用户向量化表达的过程中,能根据候选商品的变化自适应地调整。这实际上是一种与注意力机制相似的思想。排序学习的很多方法已经应用到了各种推荐场景中^[16],本文尝试将此推荐模型的思想用于改进排序学习模型。

2 排序学习公式化定义

本章将对排序学习问题进行公式化定义,与域内常用的表示形式保持一致性。排序学习中,训练集可以表示为 $\Psi = \{x, y\} \in X^n \times R^n$,其中 x 是一个由 n 个元素 $x_i, 1 \leq i \leq n$ 组成的向量, y 为含有 n 个实数的标签向量,其元素 y_i 代表对应位置的 x_i 的真实相关度, X 为特征取值空间。在主流的数据集上以及实际的场景下, x_i 通常表示一条查询信息和对应单个文档的组合特征向量,即数据集的单条样本包含了一条查询信息及一个带相关度标签的文档列表。

排序学习的主要目标是利用这样的数据集,学习一个评分函数,该评分函数可以将任意的 x_i 映射到实值,并且可以在训练集上最小化经验性的损失函数。即找到 $f: X^n \rightarrow R^n$,最小化:

$$\mathcal{L}(f) = \frac{1}{|\Psi|} \sum_{(x, y) \in \Psi} \ell(y, f(x)) \quad (1)$$

其中, ℓ 是待定义的损失函数。

上述框架为排序学习方法的通用形式,各种排序学习算法差别主要在于评分函数的构造和损失函数的定义。之前的很多研究中使用了不同形式的损失函数^[1],但现存大部分的算法在评价相关度时局限于单元变量的评分函数 $f: X \rightarrow R$,即 $f(x)$ 的各个维度是对文档 x_i 独立计算得到的。最近提出的分组文档法将文档列表进行分组,并对分组内的文档联合评分,即 $f: X^m \rightarrow R^m$,其中 m 为分组大小,该方法隶属多元评分函数的范畴。本文的评分函数以 GSF 为基础,同时在训练时会引入激活单元以调整各个不同文档的权重,用于体现交叉影响的差异性,则最小化的目标函数表示为:

$$\mathcal{L}(f) = \frac{1}{|\Psi|} \sum_{(x, y) \in \Psi} \ell(y, f(\alpha_1 x_1, \alpha_2 x_2, \dots, \alpha_i x_i, \dots, \alpha_n x_n)) \quad (2)$$

实际的计算过程中,单个训练样本的损失函数是先分组后求和计算。下一章中将详细介绍如何利用模型来最小化该目标函数。同时为了表述简洁,

后续章节将使用 GSF- m 表示分组大小为 m 的分组文档法, W-GSF 则表示本文提出的引入了激活加权机制的分组文档法。

3 W-GSF 排序学习模型

模型以 GSF^[7] 结构为基础, 其改进自文献[17]的上下文列表文档法, 用分组的形式挖掘文档之间的相互影响并用于协同评分。图 1 给出了 GSF 的简单示例, 神经网络 $g(\cdot)$ 将所有分组分别作为输入, 输出对应的初步评分, 再按照相应文档进行求和得到各自的最终得分。但如果在训练过程中, 输入样本分组集合中的全部元素, 即所有可能的分组情况, 其时间复杂度会非常高。故考虑对分组集合作下采样, 沿用 GSF 中的方法将 $\mathbf{x}$ 随机打散后, 用固定大小的滑动窗口取其所有连续子序列。该采样方法保证了每个文档出现在 m 个分组中, 降低了时间复杂度, 且实践证明其效果近似输入整个分组集合。

3.1 引入激活单元

文档的相关性应当取决于其本身的特征和列表内的样本分布情况。举个简单的例子, 当用户在搜索关键字“卷积神经网络”时, 如果返回的结果都是较近的, 则用户关注点会集中于最新的研究成果; 而如果返回结果都是经典的文献, 则被引量或作者知名度等可能会成为影响用户选择的侧重点。采用分组文档的方法时, 应当考虑到其他文档对于某个文档评分的影响, 且不同文档带来的影响也是不同的。考虑在上述的搜索示例中, 如果返回结果是两种情况的混合, 则用户可能会对两种情况同时判断: 在较新的文档选择创新性高的, 在经典文档中选择更权威的。这就涉及到在评分函数建模时, 对文档重要性的挖掘。

受启发于深度兴趣网络^[10], 本文引入激活单元来对 GSF 作出上述改进。如图 2(a), 激活单元以一个文档对的特征向量为原始输入, 分别为主要文档和次要文档, 其输出为单个实值, 用于在评价文档的相关性时对次要文档进行加权。除了两个文档的特征向量, 输入还包含了两个特征向量的差值, 三个部分拼接为整体输入后续的全连接神经网络。由于实际情况中特征数据值的分布并不固定, 如果采用如 PReLU^[18] 的激活函数时, 由于分裂点固定在了零点, 会增加拟合样本的难度。为了解决该问题, 这里的激活函数沿用了 Dice^[10]:

$$f(s) = p(s) \times s + (1 - p(s)) \times \beta s \quad (3)$$

$$p(s) = \frac{1}{1 + e^{\frac{s - E[s]}{\sqrt{\text{Var}[s] + \varepsilon}}}} \quad (4)$$

这里 s 是待激活的值, $p(\cdot)$ 可以看作 $f(\cdot)$ 的控制函数, 即用来决定选择 $f(s) = s$ 和 $f(s) = \alpha s$ 两个不同频道的函数。 $E[s]$ 和 $\text{Var}[s]$ 是每个小批量输入数据的均值和方差。 ε 是一个小的定值, 在此设置为 10^{-8} 。此时, 激活单元的激活过程可以表示为:

$$W = AU(x_{\text{main}}, x_{\text{minor}}) \quad (5)$$

3.2 W-GSF 的模型结构

实际场景中文档特征通常高维且稀疏, 对比基于树的排序学习方法^[12], 深度神经网络在稀疏高维的数据集上可以表现出更好的拟合效果, 故本文采用神经网络作为主体结构。排序学习的数据形式为一条查询对应一个文档列表。作为基本模型的输入, 首先需要对文档列表进行分组, 单条样本的第 i 个分组为:

$$G_i = \text{concat}(\mathbf{x}_i^1, \mathbf{x}_i^2, \dots, \mathbf{x}_i^m) \quad (6)$$

其中, m 为分组的大小, 即将组内文档的特征向量作拼接操作。如图 2 所示, 网络的主体结构为包含三个隐藏层的多层前馈神经网络, 原始输入在进入该神经网络之前, 会通过激活单元对输入向量进行部分激活加权, 将其转变为:

$$\mathbf{h}_0 = \text{concat}(\mathbf{x}_{\text{main}}, \alpha_2 \mathbf{x}_2, \dots, \alpha_m \mathbf{x}_m) \quad (7)$$

其中, α 为次要文档的激活值。每个隐藏层含有两个待学习参数 $\mathbf{w}_k$ 和 $\mathbf{b}_k$, 分别表示第 k 层的权值矩阵和偏置向量, 则第 k 层网络的输出为:

$$\mathbf{h}_k = \sigma(\mathbf{w}_k^T \mathbf{h}_{k-1} + \mathbf{b}_k), k \in \{1, 2, 3\} \quad (8)$$

其中, σ 是非线性的激活函数, 这里使用的是整流线性单元 ReLU: $\sigma(t) = \max(t, 0)$ 。基于此, 主体结构网络的输出可以表示为如下形式:

$$g(\mathbf{x}) = \sigma(\mathbf{w}_o^T \mathbf{h}_3 + \mathbf{b}_o) \quad (9)$$

其中, $\mathbf{w}_o$ 和 $\mathbf{b}_o$ 为输出层的权值矩阵和偏置。输出结果为 m 维的向量, 代表的是在当前分组下文档的相关度得分。

由于每个文档会出现在多个分组之中, 文档最终的相关度得分来源于所有包含该文档的分组对应的输出。不妨令 $\Pi_m(\mathbf{x})$ 为样本 $\mathbf{x}$ 的所有大小为 m 的分组组成的集合, π_k 为其中一个元素, 当文档 x_i 属于 π_k 时, 该文档得分有一部分来自于 $g(\pi_k)$ 。便于表示, 首先定义一个符号函数:

$$\text{sign}(\pi, \mathbf{x}) = \begin{cases} 1, & \mathbf{x} \in \pi \\ 0, & \mathbf{x} \notin \pi \end{cases} \quad (10)$$

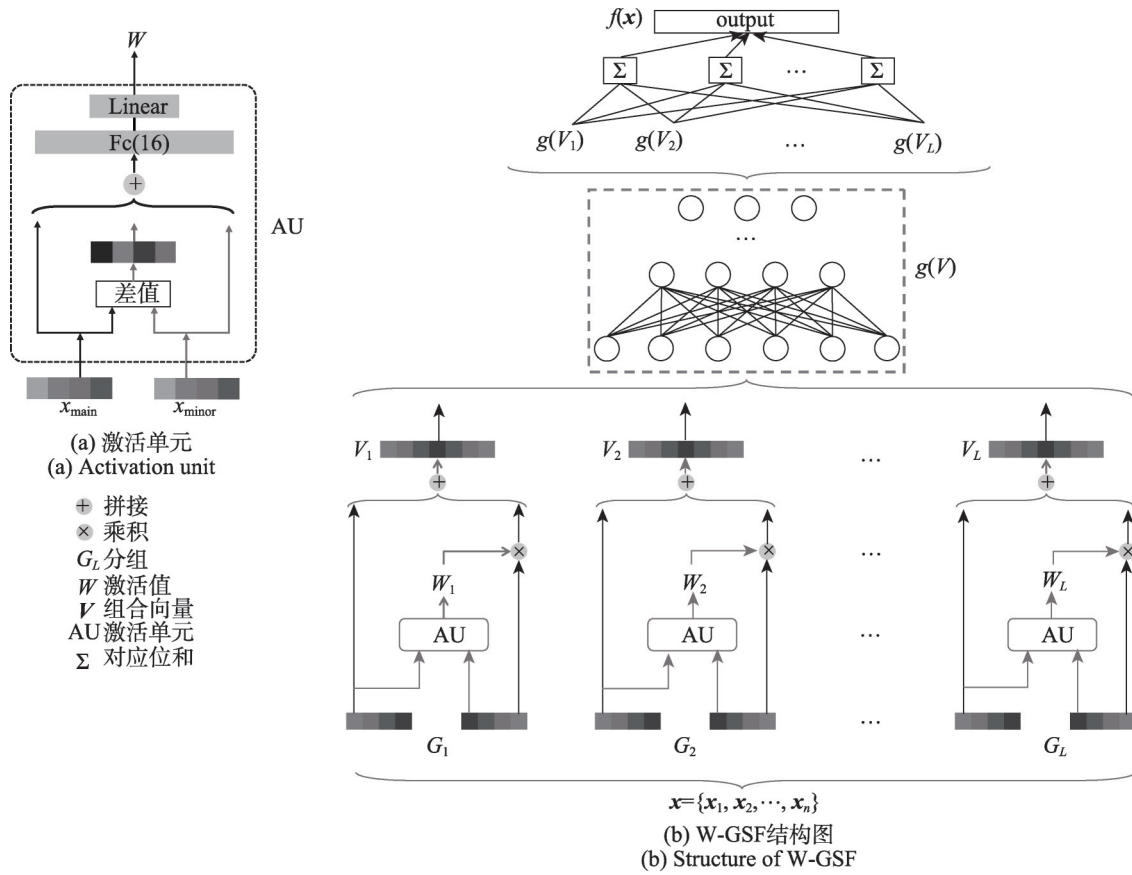


图2 W-GSF与激活单元的结构

Fig.2 Structure of W-GSF and activation unit

则W-GSF对文档 x_i 的输出为:

$$f(x_i) = \sum_{\pi_k \in \Pi_n(x)} \text{sign}(\pi_k, x_i) * g(\pi_k) |_{\pi_k^{-1}(x_i)} \quad (11)$$

这里用 $\pi_k^{-1}(x_i)$ 代表文档 x_i 出现在群组 π_k 中的位置。上述即为W-GSF的结构和整体工作流程,实验中网络权重初始化为均值为0,标准差为1的截断正态分布。

3.3 损失函数

在使用反向传播方法训练前述网络结构前,需要定义损失函数。即定义式(1)中的 $l(\cdot)$ 。在W-GSF的网络框架下同样可以使用多种损失函数。在实验中发现Softmax交叉熵损失函数在该模型下会比较有效,其形式为:

$$\ell(y; \hat{y}) = - \sum_{i=1}^n \frac{y_i}{\sum_i y_i} \lg p_i \quad (12)$$

其中, y 为样本的真实标签, $\hat{y}$ 为模型的输出值。上式相当于原始的标签值进行了归一化,可以将归一化值看作文档排名更高的概率,其中 p_i 为模型输出的归一化:

$$p_i = \frac{\exp(\hat{y}_i)}{\sum_i \exp(\hat{y}_i)}, 1 \leq i \leq n \quad (13)$$

该操作的目的是去除标签的实际值大小带来的偏差影响,在将文档得分转化为一个较小的值的同时,保持相对排名顺序不变。实验中使用小批量梯度下降法AdaGrad来优化该损失函数。

4 实验分析与结果

4.1 数据集

实验中采用的数据集为微软MSLR-Web30K以及它的随机采样版本MSLR-Web10K^[19],分别含有30 000条和10 000条查询。数据集由查询和一系列的文档特征向量组成,每一行代表一个查询-文档对,并给出了相关性评价的标签。相关性从0(不相关)到5(完全相关)代表从低到高的相关程度,由微软搜索引擎Bing的日志数据统计解析而出。在训练中丢弃了没有相关文档的查询项。每一条查询样本为136维的特征向量,包含了查询信息、文档特征以及交互特征,形式上既有离散的分类特征,也有连续的例如

BM25等稠密特征,代表了研究领域内常用的基本特征。该数据集被划分成5个子集 $\{S_1, S_2, S_3, S_4, S_5\}$,按照不同的训练集、验证集和测试集划分情况形成了5个不同的文件夹,其划分占比为3:1:1。因此不同的文件夹数据相同只是分配的子集不同,不失一般性,在实验中使用的划分为 $\{(S_1, S_2, S_3), (S_4), (S_5)\}$ 。

4.2 排序学习方法与评价指标

在上述数据集上,实验对比了本文提出的模型和各种现有的排序学习方法,包括基于支持向量机、梯度提升树和基于深度神经网络的多种方法。RankingSVM^[20]作为经典方法参考, LambdaMART^[3]作为树模型方法的对比, RankNet^[3]和GSF^[7]作为神经网络方法的参考基线,详细信息如表1所列。

表1 基线排序学习方法

Table 1 Baseline learning-to-rank models

方法	具体描述
RankingSVM	经典的基于支持向量机的文档对排序方法
LambdaMART	基于树的方法,善于处理稠密特征向量,多个公开数据集上的SOTA技术
RankNet	经典的文档对方法,基于神经网络模型,可以在损失函数上乘上 λ -权重变成文档列表法
GSF- m	群组大小为 m 的分组文档法

实验中RankingSVM目标函数的正则化系数设置为20.0,并使用径向基核函数。LambdaMART基于LightGBM^[21]实现,训练了500棵回归树,每棵树的叶子节点为255,其他的参数设置与文献[7]保持一致。该模型在很多排序学习公开数据集上都是表现最优秀的算法之一,此处作为参考以比较不同类型的排序学习方法之间的差异。RankNet和GSF均基于神经网络,实现时利用了广泛使用深度学习库TensorFlow。RankNet是一个多层的前向全连接神经网络,是文档对方法的基本模型,在本文中该方法和GSF使用了相同的隐藏层参数设置,均为三层(64,32,16)全连接深度神经网络,并在每层使用了批量归一化,激活函数为整流线性单元。除了二者的输入形式差异,使用的损失函数也有所不同:RankNet使用了在排序学习领域的经典逻辑斯蒂损失函数^[3],而GSF使用的是Softmax交叉熵^[22]。实验中这三种模型的超参在沿用原文设置的基础上,进行了优化微调。

对于本文提出的W-GSF模型,其主体结构与GSF一致,额外增加了激活单元对文档调整加权。激活单元含有一层尺寸为(16)的隐藏层,输出单个实值权重。值得一提的是,激活过程为两两文档的交

互,因此其天然地适合分组大小为2的情况,而对于大分组的GSF,简单的依次激活加权会使得文档之间的交叉影响变得混乱。另外由于本身GSF方法中大尺寸的分组会使得训练复杂度较高,而加入激活单元后大分组带来的计算代价增长倍率更高。故实验中W-GSF仅关注两两文档之间的交叉关系,使其加权操作简洁合理且增加的计算量在可接受范围内,并期望带来排序指标上的提升。W-GSF在实验中小批量训练的批量大小设置为128,模型学习率为0.01。

实验在Web30K和Web10K两个数据集上分别训练了上述五种模型。评价指标使用了排序推荐场景中常用的指标NDCG^[15],即归一化折损累计增益。该指标的优点在于可以对变长的排序列表给出有意义的评价,同时可以保证当高相关度的文档排在整个列表中靠前的位置时,排序列表的得分较高。

4.3 结果对比分析

模型训练时使用小批量的梯度下降法来优化损失函数,由于每个查询包含的文档数量及相关度的分布并不一致,故在GSF和W-GSF中计算整体损失之前会对小批量样本去偏,即对分组损失值按照其关联的所有文档相关度总和进行加权。在MSLR的两个数据集上,不同的模型效果如表2所示。给出的指标均为在对应数据上的测试集结果,除LambdaMART外,指标由训练过程中在验证集上NDCG@5指标达到最高时的模型计算得出,数据展示的是10次实验的均值,置信区间为95%。

表2给出了GSF分别在分组大小为2和64时的模型效果,且实验中经过调参优化,相对于原文中报告的结果^[7],一定程度提升了在MSLR数据集上的表现。从表中可以看出,本文模型指标远超经典算法RankingSVM以及神经网络排序模型RankNet。W-GSF在Web30K数据集上的实验表明,在NDCG@5处相对于该两种模型分别提升了9.7个百分点和2.65个百分点,在Web10K上该指标的提升分别为7.79个百分点和1.76个百分点。

W-GSF本质是经GSF-2的模型结构改进而来,故在对比指标时更关注其与GSF-2的关系。从表中注意到相对于GSF-2模型,增加了激活单元的W-GSF在NDCG指标上同样全面提升,在Web10K数据集上,W-GSF分别在NDCG@1,5处提升了0.73个百分点和0.39个百分点,对于Web30K数据集则分别在NDCG@1,5处提升了1.04个百分点和1.08个百分

表2 MSLR数据集上模型的NDCG指标对比

Table 2 Comparison of models' test NDCG on MSLR dataset

%

Dataset	Model	NDCG@1	NDCG@3	NDCG@5	NDCG@10
Web10K	RankingSVM	36.22(± 0.12)	—	34.91(± 0.09)	31.07(± 0.09)
	RankNet	39.74(± 0.05)	—	40.94(± 0.07)	43.37(± 0.05)
	LambdaMART	46.20 (± 0.13)	—	46.23 (± 0.09)	48.33 (± 0.11)
	GSF-2	41.79(± 0.14)	42.13(± 0.09)	42.31(± 0.11)	44.79(± 0.09)
	GSF-64	42.75(± 0.14)	42.61(± 0.07)	42.65(± 0.12)	44.87(± 0.13)
	W-GSF	42.52(± 0.09)	42.57(± 0.09)	42.70 (± 0.09)	44.91 (± 0.09)
Web30K	RankingSVM	37.33(± 0.09)	—	35.07(± 0.13)	32.61(± 0.17)
	RankNet	40.74(± 0.11)	—	42.12(± 0.05)	44.73(± 0.06)
	LambdaMART	50.33 (± 0.07)	—	49.20 (± 0.09)	51.05 (± 0.11)
	GSF-2	43.37(± 0.11)	43.41(± 0.09)	43.69(± 0.11)	46.22(± 0.09)
	GSF-64	44.43(± 0.09)	44.52(± 0.13)	44.89(± 0.12)	46.84(± 0.10)
	W-GSF	44.41(± 0.07)	44.54 (± 0.09)	44.77(± 0.09)	46.88 (± 0.09)

点。另外对比GSF-64可以看出,虽然W-GSF在最前端的排序结果稍差,但在NDCG@10处表现更好,即在评价的头部序列拉长时整体的排序效果占优势。而由于GSF-64的高指标表现计算代价很大,相对而言本文的模型在计算效率上更有优势。

除此之外可以发现,LambdaMART作为基于树的方法在实验中表现较好,其主要原因是MSLR数据集大部分的特征为稠密连续特征,离散的稀疏类别特征较少,这也是公开排序学习数据集的普遍现象。而提升树恰好擅长处理该类特征,相对的,神经网络模型可以更好地处理稀疏特征向量。这一结论由此前大量的相关工作^[6-7]给出,即使其基于神经网络的排序方法在MSLR数据集上的表现不如LambdaMART,但是在业务场景下,数据集往往远比MSLR大很多且特征维度更高,当这种高维稀疏特征用于树模型时会产生过拟合的情况,降低了模型的泛化能力,而神经网络模型由于通常会带有正则化防止过拟合,故在最终指标上的表现往往呈现相反的趋势。

4.4 模型效率对比

通过挖掘文档之间的相互影响来提升排序指标本身是一件是十分耗费计算量的事情,当在GSF中试图同时挖掘较多文档的交叉关系时,扩大分组尺寸带来的收益与增加的计算量不成正比,综合性价比不高,尤其是在搜索这种对响应时间要求较高的场景下。

为了对比模型之间的性价比,除了使用NDCG@10作为效果提升的指标外,将FLOPs作为衡量W-GSF、GSF-2和GSF-64的计算代价指标。FLOPs代表模型的理论浮点运算量,其大小不仅和模型的输

入形式有关,还和模型结构有关。对比运算量时,不关心训练过程,而只关注对于同样的一条样本,不同模型经过一次前向计算给出相关度评分所需要的浮点计算次数,加法和乘法各算一次。对于全连接层,其FLOPs大小为:

$$FLOPs = (2 \times I - 1) \times O \quad (14)$$

其中, I 和 O 分别代表输入输出的维度,若考虑偏置则不需要减1。同时计算时忽略了激活函数的影响。

根据式(14),计算了如图1不同分组和图2所示的共三种模型的计算量FLOPs,并在表3给出结果。计算条件为在同样的一条由查询和文档列表组成的样本下,不同模型运行一次所消耗的浮点计算量,并假设该样本含有100个文档。以GSF-2作为基准线, $\Delta NDCG@10$ 为相对基准线的提升百分比, $\Delta FLOPs$ 为相对基准线的倍率。

表3 W-GSF与GSF在MSLR-Web30K上的计算量对比

Table 3 Comparison of calculational cost of W-GSF and GSF on MSLR-Web30K

Model	$\Delta FLOPs$	$\Delta NDCG@10\%$
GSF-2	—	—
GSF-64	$\times 28.10$	+1.34
W-GSF	$\times 1.33$	+1.43

从表3可以看出,GSF-64对于同样的样本,消耗的浮点计算量是GSF-2的28倍以上,而W-GSF在排序指标达到了相同水准的情况下,相对基准线只增加了约30%的计算量。总体而言,本文模型在改进GSF-2时,以增加少量的计算量带来了可观的效果提

升,相对于GSF-64性价比更高。该计算量上的优势使得其在众多业务场景中有比大分组GSF更多的应用空间。

5 结束语

本文将推荐领域里深度兴趣网络的思路迁移到排序学习中,结合GSF分组文档法的思想提出了带有激活加权策略的W-GSF,通过学习不同文档之间的差异性影响对交叉文档关系有了更深层次的挖掘。在公开数据集的一系列实验证明了其在指标上带来的提升,同时相对原模型只增加了较少的计算量,相比大分组尺寸的GSF方法,本文模型更有性价比。但该方法仍然有不足之处,对于大分组的方法,激活加权的过无法自然地推广拓展,文档之间的交叉影响差异性会使得这一过程变得混乱,单个文档的最终激活值难以界定。一种可能的解决方法是将单个文档与分组内其他所有文档交叉,最终权重为各激活权重之和,但这样会带来计算量上的大幅增加。另一种方法是将评分函数再改回单元变量的形式,只是在评价一个文档的相关度评分时,将同列表内的其他文档作加权求和后作为辅助特征帮助评分。除此之外,对于主体神经网络结构的纵向加深和横向尺寸扩展也是进一步改进该方法的可行路径。将上述几种思路作为未来的进一步工作方向。

参考文献:

- [1] LIU T Y. Learning to rank for information retrieval[C]//Proceeding of the 33rd International ACM SIGIR Conference on Research and Development in Information Retrieval, Geneva, Jul 19-23, 2010. New York: ACM, 2010: 904.
- [2] CHEN W, LIU T Y, LAN Y Y, et al. Ranking measures and loss functions in learning to rank[C]//Advances in Neural Information Processing Systems 22, Vancouver, Dec 7-10, 2009. Red Hook: Curran Associates, 2009: 315-323.
- [3] BURGESS J C. From RankNet to LambdaRank to LambdaMART: an overview: MSR-TR-2010-82[R]. 2010.
- [4] CAO Z, QIN T, LIU T Y, et al. Learning to rank: from pairwise approach to listwise approach[C]//Proceedings of the 24th International Conference on Machine Learning, Corvallis, Jun 20-24, 2007. New York: ACM, 2007: 129-136.
- [5] 曹军梅, 马乐荣. 卷积重提取特征的文档列表排序学习方法[J]. 中文信息学报, 2020, 34(8): 86-93.
CAO J M, MA L R. Listwise reranking via convolutional re-extracted features[J]. Journal of Chinese Information Processing, 2020, 34(8): 86-93.
- [6] WANG R X, FANG K, ZHOU R K, et al. SERank: optimize sequencewise learning to rank using squeeze-and-excitation network[J]. arXiv:2006.04084, 2020.
- [7] AI Q Y, WANG X H, BRUCH S, et al. Learning groupwise multivariate scoring functions using deep neural networks [C]//Proceedings of the 2019 ACM SIGIR International Conference on Theory of Information Retrieval, Santa Clara, Oct 2-5, 2019. New York: ACM, 2019: 85-92.
- [8] YE P, DOERMANN D. Combining preference and absolute judgements in a crowd-sourced setting[C]//Proceedings of the 2013 ICML Workshop on Machine Learning Meets Crowd-sourcing, Atlanta, Jun 16-21, 2013: 1-7.
- [9] JOACHIMS T, GRANKA L A, PAN B, et al. Accurately interpreting clickthrough data as implicit feedback[C]//Proceedings of the 28th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval, Salvador, Aug 15-19, 2005. New York: ACM, 2005: 154-161.
- [10] 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 and Data Mining, London, Aug 19-23, 2018. New York: ACM, 2018: 1059-1068.
- [11] JOACHIMS T. Training linear SVMs in linear time[C]//Proceedings of the 12th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, Philadelphia, Aug 20-23, 2006. New York: ACM, 2006: 217-226.
- [12] FRIEDMAN J H. Greedy function approximation: a gradient boosting machine[J]. Annals of Statistics, 2001, 29(5): 1189-1232.
- [13] JOACHIMS T. Optimizing search engines using clickthrough data[C]//Proceedings of the 12th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, Philadelphia, Aug 20-23, 2006. New York: ACM, 2006: 133-142.
- [14] 熊李艳, 陈晓霞, 钟茂生, 等. 基于PairWise排序学习算法研究综述[J]. 科学技术与工程, 2017, 17(21): 184-190.
XIONG L Y, CHEN X X, ZHONG M S, et al. Survey on PairWise of the learning to rank[J]. Science Technology and Engineering, 2017, 17(21): 184-190.
- [15] TAYLOR M J, GUIVER J, ROBERTSON S, et al. SoftRank: optimizing non-smooth rank metrics[C]//Proceedings of the 2008 International Conference on Web Search and Web Data Mining, Palo Alto, Feb 11-12, 2008. New York: ACM, 2008: 77-86.
- [16] 黄震华, 张佳雯, 田春岐, 等. 基于排序学习的推荐算法研究综述[J]. 软件学报, 2016, 27(3): 691-713.
HUANG Z H, ZHANG J W, TIAN C Q, et al. Survey on learning-to-rank based recommendation algorithms[J]. Journal

of Software, 2016, 27(3): 691-713.

- [17] AI Q Y, BI K P, GUO J F, et al. Learning a deep listwise context model for ranking refinement[C]//Proceedings of the 41st International ACM SIGIR Conference on Research and Development in Information Retrieval, Ann Arbor, Jul 8-12, 2018. New York: ACM, 2018: 135-144.
- [18] HE K M, ZHANG X Y, REN S Q, et al. Delving deep into rectifiers: surpassing human-level performance on imagenet classification[C]//Proceedings of the 2015 IEEE International Conference on Computer Vision, Santiago, Dec 7-13, 2015. Washington: IEEE Computer Society, 2015: 1026-1034.
- [19] QIN T, LIU T Y. Introducing LETOR 4.0 datasets[J]. arXiv: 1306.2597, 2013.
- [20] JOACHIMS T. A support vector method for multivariate performance measures[C]//Proceedings of the 22nd International Conference on Machine Learning, Bonn, Aug 7-11, 2005. New York: ACM, 2005: 377-384.
- [21] KE G L, MENG Q, FINLEY T, et al. LightGBM: a highly efficient gradient boosting decision tree[C]//Advances in Neural Information Processing Systems 30, Long Beach, Dec 4-9, 2017. Red Hook: Curran Associates, 2017: 3146-3154.
- [22] BRUCH S, WANG X H, BENDERSKY M, et al. An analysis of the softmax cross entropy loss for learning-to-rank with binary relevance[C]//Proceedings of the 2019 ACM SIGIR International Conference on Theory of Information Retrieval, Santa Clara, Oct 2-5, 2019. New York: ACM, 2019: 75-78.



李玉轩(1996—),男,江苏宿迁人,硕士研究生,主要研究方向为机器学习、信息检索等。
LI Yuxuan, born in 1996, M.S. candidate. His research interests include machine learning, information retrieval, etc.



洪学海(1967—),男,安徽巢湖人,博士,研究员,博士生导师,CCF杰出会员,主要研究方向为高性能计算、人工智能、大数据等。

HONG Xuehai, born in 1967, Ph.D., professor, Ph.D. supervisor, outstanding member of CCF. His research interests include high performance computing, artificial intelligence, big data, etc.



汪洋(1984—),男,湖北武汉人,博士,高级工程师,硕士生导师,中国科学院计算机网络信息中心信息化发展战略与评估中心主任,主要研究方向为大数据分析、态势感知系统、信息化发展战略研究等。

WANG Yang, born in 1984, Ph.D., senior engineer, M.S. supervisor. His research interests include big data analysis, situational awareness system, strategy research on informatization development, etc.



唐正正(1992—),男,安徽怀远人,博士研究生,主要研究方向为机器学习、数据挖掘、图表示学习等。

TANG Zhengzheng, born in 1992, Ph.D. candidate. His research interests include machine learning, data mining, graph representation learning, etc.



班艳(1981—),女,北京人,高级工程师,主要研究方向为信息化发展战略研究。

BAN Yan, born in 1981, senior engineer. Her research interest is strategy research on informatization development.