

文本配图系统的设计与实现

张明西, 乐水波, 李学民, 董一鹏
(上海理工大学, 上海 200093)

摘要: **目的** 设计并开发文本配图系统, 实现面向文本数据的在线自动配图。 **方法** 基于图片和文本之间的描述关系构建“图片-标签”二分网络, 然后基于“图片-标签”的二分网络, 利用重启随机游走模型进行图片与标签之间的相关性计算。采用 TextRank 模型提取关键字, 并将关键字构成的集合作为查询, 将关键字视为标签。基于离线计算结果, 在线整合标签与图片之间的相关性, 得到文本与图片的相关性。依据相关性由大到小进行排序, 并返回前 k 个最相关的图片。 **结果** 实验结果表明, 前 5 个返回结果的 MAP 值能够达到 0.839, 能够准确地返回用户期望的图片。 **结论** 系统能够依据输入文本进行准确的图片匹配。

关键词: TF-IDF 模型; 文本配图; 重启随机游走; TextRank 模型

中图分类号: TP317.4 **文献标识码:** A **文章编号:** 1001-3563(2020)19-0252-07

DOI: 10.19554/j.cnki.1001-3563.2020.19.036

Design and Implementation of Picture Matching System for Text

ZHANG Ming-xi, YUE Shui-bo, LI Xue-min, DONG Yi-peng

(University of Shanghai for Science and Technology, Shanghai 200093, China)

ABSTRACT: The paper aims to design and develop a text matching picture system to realize on-line automatic picture matching facing text data. Based on the description relationship between pictures and text, a “picture-tag” bipartite network was built. And then based on the “picture-tag” bipartite network, the relevance scores between pictures and tags was computed by random walk with restart. The keywords of the input text were extracted by TextRank model and the set composed of key words was queried with the key words as tags. Based on the off-line calculation result, the relevance scores between text and pictures were obtained by integrating the relevance scores between pictures and tags. And the top k most relevant pictures were returned according to the relevance scores. The experimental results show that the MAP score of the top 5 five returned pictures can reach 0.839, which can accurately return the expected pictures. The system can accurately match pictures for input text.

KEY WORDS: TF-IDF model; text matching picture; random walk with restart; TextRank model

互联网和移动技术的出现促使包装印刷行业发生根本转变, 内容呈现方式由传统的纸质文本发展为图文并茂、多媒介融合的电子文本或网络文本。在互联网传媒时代, 读者习惯在移动手机端获取信息, 各

种 APP 软件将生活与工作连接在一起, 如“微信”、“微博”、“知乎”及“今日头条”等。用户在移动端阅读信息的习惯更加侧重“读图”而不是“读字”, 网络电子文本的快速且准确配图已是内容生产者的竞争点。图片

收稿日期: 2020-03-03

基金项目: 国家自然科学基金 (62002225); 上海市自然科学基金 (16ZR1422800); 上海理工大学国家级项目培育基金 (16HJPY-QN04); 国家新闻出版广电总局重点实验室招标课题 (ZBKT201809)

作者简介: 张明西 (1985—), 男, 博士, 上海理工大学副教授、硕导, 主要研究方向为图文信息处理、数据挖掘。

能够丰富文章的内容,增强版面的视觉效果,图片的作用在数字出版行业也愈发重要,读者在面对相同主题内容的文章时,符合文章要求的图片是必不可少的,符合主题内容且独特的配图会给读者留下深刻的印象。

传统文本配图主要是依靠人工进行图片的选择,编辑人员需要阅读和理解文本,再选择合适的图片,因此,编辑人员需要花费大量时间和精力判断图片和文本的匹配程度^[1]。尽管一些常用的搜索引擎能够有效减轻用户的工作量,但是大多数搜索引擎对于文本输入有一定限制。比如,百度对于输入文本有 38 个汉字的限制。当使用百度等搜索引擎对文本内容进行搜索时,字符长度受到限制,需要通过人工方式来从文本中提取关键字,才能实现图片与文本的匹配。

近年来,研究者在文本匹配方面做了大量工作。Yang 等^[2]涵盖了自然语言推理的数据集上提出一种快速的文本匹配模型方法。Fang 等^[3]提出从维吾尔族生物医学图像中提取文本,并在特定的词典中匹配文本进行语义分析。He 等^[4]利用蒙特卡洛树搜索增强马尔可夫决策来解决文本匹配问题。Lee 等^[5]研究图像文本匹配和图像字幕的问题,从图像和文本数据结合,研究图像与描述性句子相关联的问题。Gong 等^[6]将图像与描述性句子嵌入共同的潜在空间中,提高了检索图像描述的准确性。邓佩等^[7]通过形容词来建立语义的情感空间解释图像情感和图像认知的问题。Karpathy 等^[8]提出了一种全卷积本地化网络架构用来解决图片与文字之间描述的任务。Johnson 等^[9]提出一个图像与语言之间的描述模型,利用图像数据及其句子描述来了解语言和图像视觉之间的关系。

显然,现有研究工作很少涉及文本配图方面的研究,难以满足现实应用中对于文本配图方面的迫切需求。为此,文中结合重启随机游走模型在“昵图”网抓取的数据集上开展文本配图研究。对于用户输入文本进行分词、去除停用词等处理,采用 TextRank 模型提取关键字,并将关键字构成的集合作为查询。基于离线计算结果,整合关键字与图片之间的相关性,

得出文本与图片之间的相关性,进而设计并实现一种面向文本的配图系统。配图系统能够应用于日常办公,帮助用户快速查找合适的图片,节省用户在配图工作上的时间和精力,实现高效办公。

1 系统流程的框架

文本配图系统的流程框架见图 1,主要分为离线和在线 2 个阶段。其中,离线阶段主要负责对于原始数据的处理及图片与标签的相关性计算,具体步骤为:数据预处理,对图片与标签组成的数据集进行预处理,主要是使用 TF-IDF 模型构建图片与标签的二分网络,并设置阈值去除网络中关系较弱的联系;构建“图片-标签”网络,依据标签与图片之间的“描述”关系构建“图片-标签”网络;图片与标签的相关性计算,利用重启随机游走模型离线计算得到图片与标签之间的相关性。

在线阶段主要是为用户输入文本进行配图,具体步骤:提取关键字,对用户的输入文本提取关键字,采用 TextRank 模型进行关键字提取,并将关键字集合作为查询交给系统进行处理,并将关键字视为标签;文本与图片的相关性计算:整合标签与图片之间的相关性,得出文本与图片之间的相关性;图片排序并输出,将图片按照相关性由大到小排序,并将最相关的若干图片返回给用户。

图片匹配系统在实际场景中具有广泛的应用前景,例如网络日志配图、文本封面推荐、自媒体文章配图等。以网络日志配图为例,系统基于用户提供的日志内容,通过去噪、关键字选择、相关性计算等操作就能够从海量素材中找到合适当前日志的图片。在文本封面推荐方面,系统文本内容查找最相关的图片,并推荐给用户作为候选封面。在自媒体文章撰写时,通常会有着文本配图的需求,高效、准确的文字配图能够显著提升文章撰写者的工作效率和读者的阅读体验。显然,通过文本配图系统能够进行准确的图片匹配,节省用户在实际场景中寻找图片的时间和精力,提高用户的工作效率。

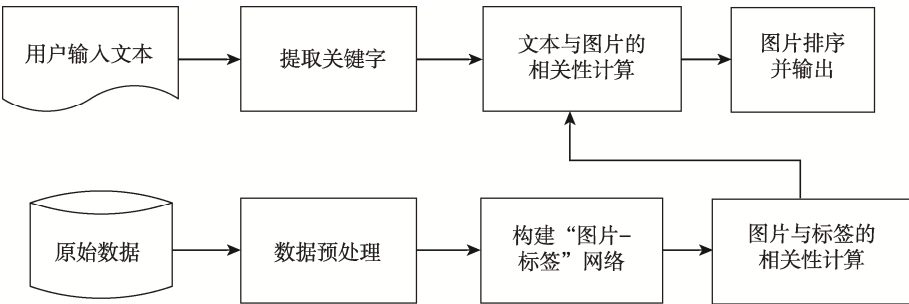


图 1 系统的流程框架
Fig.1 System process framework

2 数据预处理

数据集的组成是通过网络爬虫获取的图片与标签的联系,数据集中存在大量相关性较弱的噪音词会影响图片匹配的准确性,需对其进行预处理。数据预处理主要包括构建二分网络,优化二分网络。该过程通过图片与标签的联系构建二分网络,再使用 TF-IDF 模型去除相关性较弱的噪音词,减少计算成本开销。

2.1 构建二分网络

图片与标签的联系能够直接被二分网络描述,成为“图片-标签”网络,记为 $G=\langle V, E \rangle$, $V=V_p \cup V_t$ 表示图片与标签 2 种节点类型的集合, V_p 为图片的顶点集合, V_t 为标签的顶点集合, E 为用节点对表示的图片和标签的边集,边的权重为 2 种节点之间的权值。

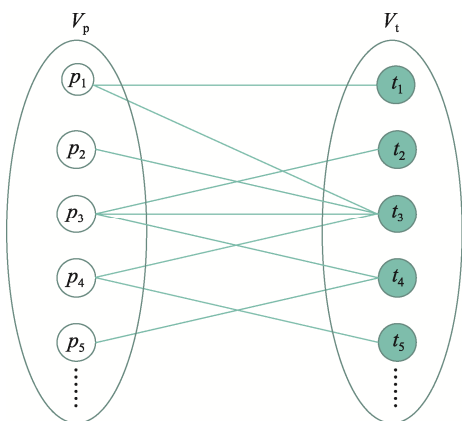


图2 “图片-标签”的二分网络
Fig.2 "Picture-Tag" bipartite network

“图片-标签”网络的示例见图 2。图片 p_1 与标签 t_1, t_3 相连,而与 t_2, t_4, t_5 没有边,表示该图片包含 t_1 和 t_3 的标签特征;标签 t_3 与图片 p_1, p_2, p_3, p_4 相连,表示这 4 个图片都含有标签 t_3 所代表的特征;图片 p_3 与标签 t_2, t_3, t_4 , 图片 p_5 与标签 t_4 相连。

2.2 利用 TF-IDF 获取权值及优化网络

TF-IDF 模型是一种统计方法,是一种用于信息检索与数据挖掘的加权技术^[10-11],用以评估某一术语对于语料库中的各文档的重要程度。在文中 TF-IDF 模型用来计算图片与标签之间的权值。

词频 (Term Frequency, TF) 指某一词在相应图片标签中出现的次数。 f_t 为标签 t 在图片 p 相应标签中出现的次数, m_t 为图片 p 所有标签出现的次数总和。标签 t 在图片 p 中的 TF 值计算为:

$$F_t = \frac{f_t}{m_t} \quad (1)$$

逆文档频率 (Inverse Document Frequency, IDF)

指如果某一词在不同图片标签中出现的频率越低, IDF 越大,说明该词区分类别的能力越好。用 $|N|$ 表示图片素材库中图片的总数, $|Q|$ 表示图片素材库中包含标签 t 的图片总数, 标签 t 的 IDF 值计算如下:

$$D_t = \lg \frac{|N|}{(|Q|+1)} \quad (2)$$

如果某一词在一张图片标签中出现的频率高,并且在其他图片标签中很少出现,则认为这个词具有很好的区分类别能力。TF-IDF 是 TF 与 IDF 的乘积,乘积越大表示这个词更能体现图片的个性。TF-IDF 的计算公式为:

$$T_t = F_t \times D_t \quad (3)$$

图片上传者为获取更多的检索次数,可能故意设置一些无用标签或者无意中增加一些停用词。图片的标签中包含越多类标签,其信息价值就会越低,同时也会增加相关性计算的开销。系统通过设置阈值的方式来去除这类噪音标签,达到优化网络的目的。用 Y_p 表示去除图片 p 中噪音标签而设置的阈值,计算的公式为:

$$Y_p = ((T_p)_{\max} - (T_p)_{\min}) \times \mu + (T_p)_{\min} \quad (4)$$

式中: $(T_p)_{\max}$ 为当前图片 p 所有标签中最大的 TFIDF 值; $\mu \in [0,1]$ 是用来控制去除数据的范围; $(T_p)_{\min}$ 为当前图片 p 所有标签中最小值的 TFIDF 值。如果当前标签 t 在图片 p 的 TFIDF 值小于 Y_p , 保留该标签 t ; 否则, 去除标签 t 。通过设置阈值,实现“图片-标签”网络的优化。参数 μ 取值越大,产生的阈值就越大,“图片-标签”网络也就会越稀疏,反之亦然。

3 图片与标签的离线计算

文本配图的关键是依据重启随机游走模型进行图片与标签之间的相关性计算。重启随机游走模型广泛应用于排序学习、实体推荐、图像标注和页面信息检索等方面^[12-13], 主要用来计算网络节点之间的相关性。相对于传统距离计算方法,重启随机游走模型能够利用网络全局信息计算节点相关性^[14-17]。

对于给定的“图片-标签”网络,在迭代次数为 l 时,节点之间的随机游走距离的矩阵形式定义为:

$$R^l = \sum_{r=1}^l c(1-c)^{r-1} P^r \quad (5)$$

式中: $c \in (0,1)$ 为重启概率; P 为“图片-标签”网络的转移概率矩阵。根据式 (5), 得到重启随机游走模型的递归形式为:

$$R^l = c(1-c)^l P^l + R^{l-1} \quad (6)$$

在网络中以概率 $(1-c)$ 随机游走到当前节点的邻居节点,以概率 c 随机跳跃到图中的初始节点重新开始游走,得到网络中每个节点被访问到的概率。用上一步获得矩阵作为下一次重启随机游走的初始输入,

反复迭代此过程,直至达到收敛状态^[18]。在收敛状态下,矩阵元素表示图片与标签的相关性。

4 在线查询处理

在线查询的过程中,系统首先对用户输入文本进行关键字的提取,将关键字构成的集合作为查询输入,并将关键字视为标签。基于离线计算的结果,整合标签与图片之间的相关性,得出文本与图片之间的相关性,再依据相关性排序返回相关的图片。

4.1 关键字提取

文本具有关键字多、内容较长等特点,仅使用分词、去停用词等过滤处理不能准确判断文本的主题,影响配图的准确性。采用 TextRank 模型从文本中提取多个词语作为网络节点,并构建仅包含词的网络,通过迭代计算将词语按重要性排序,保留得分比较高的关键字,去除得分较低的关键字。

关键字的提取是从文本中提取主题最相关的词语,关键字在数据库、信息检索和数据挖掘等领域的有着重要作用。TextRank 模型基本思想来源于 Google 的 PageRank^[19-20],它使用局部词之间的关系对关键字进行排序,并直接从文本中提取关键字。TextRank 模型不需要事先对多个文本进行学习训练,仅对单个文本进行分析就能够提取该文本的关键字,它具有实现简单,无监督,语言相关性弱等优点。

用户提供的输入文本记为集合 T ,集合元素为构成文本的关键字。依据关键字中的词序构建词图,该词图能够表示为有向图 $G'=\langle V', E' \rangle$,其中, V' 表示节点的集合, E' 表示边的集合。节点 v_i 和 v_j 之间的权值记为 w_{ij} 。 $I(v_i)$ 表示指向 v_i 的节点集合, $O(v_j)$ 表示被节点 v_j 指向的节点集合。TextRank 模型基于词图计算网络中每个节点的重要程度,得到每个文本信息单元的得分值。对于节点 v_i , TextRank 的计算函数如下:

$$S(v_i) = (1-d) + d \times \sum_{j \in I(v_i)} \frac{w_{ji}}{\sum_{v_k \in O(v_j)} w_{jk}} S(v_j) \quad (7)$$

式中: d 为阻尼系数,又称衰减系数,取值范围为 0 到 1 之间,一般取值为 0.85^[21]。

在关键字的提取过程中,关键字的边由共现关系决定。输入文本 T 中的句子过长,经过分词处理会产生数量较多的词集,影响算法的时间开销和计算的准确率,因此,关键字的共现关系又受到滑动窗口大小 W 的影响,窗口大小 W 表示每次最多出现的词语,通过每次向右滑动建立关键字之间的共现关系。 W 值过大或者过小,都会影响词图边的数量,最终影响配图的准确率。

对于输入文本 T ,基于 TextRank 计算结果,按

照分值对所有关键字进行排序,保留前 n 个得分比较高的关键字,去除得分较低的关键字,再将处理后的关键字构成的集合作为查询提交给系统,并将关键字视为标签。

4.2 文本与图片的相关性计算

基于离线计算结果,通过整合图片与标签之间的相关性来计算文本与图片之间的相关性。 R'_{tp} 表示在矩阵 R' 中标签 t 与图片 p 的相关性,文本 T 与图片 p 之间的相关性计算公式为:

$$H_{tp} = \sum_{t \in T} R'_{tp} S(t) \quad (8)$$

基于式(8),系统首先在线计算文本与图片之间的相关性,然后将图片按相关性由大到小排序,最后返回相关性最大的前 k 个图片。

5 实验结果与分析

5.1 实验设置

实验使用的 CPU 是 Intel(R) Core(TM) i7-9750H,内存是 16 GB。所有算法由 Java 编写,开发的环境为 Eclipse Java 2019。实验使用 python 爬虫的方法抓取“昵图”网(www.nipic.com)的网络数据集进行“图片-标签”网络构建,包含 37 221 张图片,58 725 个标签,图片和标签之间共有 676 965 条“描述”关系。

5.2 评估方法

实验过程中,由用户输入 20 个文本进行实验效果评估。采用平均精度均值(Mean Average Precision, MAP)来评估配图结果的准确性,定义为:

$$M = \frac{1}{Z} \sum_{i=1}^Z L_k(T_i) \quad (9)$$

式中: Z 为用于文本的数量; $L_k(T_i)$ 为文本 T_i 的前 k 个图片的平均精度函数。 $L_k(T_i)$ 定义为:

$$L_k(T_i) = \frac{1}{k} \sum_{j=1}^k L_k(T_i, p_j) \quad (10)$$

式中: p_j 为 T_i 结果序列中第 j 个位置对应的图片; $L_k(T_i, p_j)$ 为文本 T_i 和图片 p_j 的相关程度。 $L_k(T_i, p_j)$ 由用户根据返回结果的相关性进行打分,打分的取值范围为 0 到 1,打分越高表示两者之间的相关性越高。

5.3 不同参数设置对应的 MAP 值

不同参数 μ 对应的 MAP 值见图 3,其中 $k=5, 10, 30$ 。设置参数 μ 是去除网络中图片与标签较弱的关联,减少噪音词对计算的影响。在 μ 从 0 至 0.5 变化时,MAP 值基本维持不变;当 μ 从 0.5 至 0.8 变化时,MAP 值快速上升;当 μ 继续增长时,图片标签过于稀少,导致 MAP 值呈现快速下降的变化。基于图中的变化结果,后续实验中选择 $\mu=0.8$ 进行

效果测试。

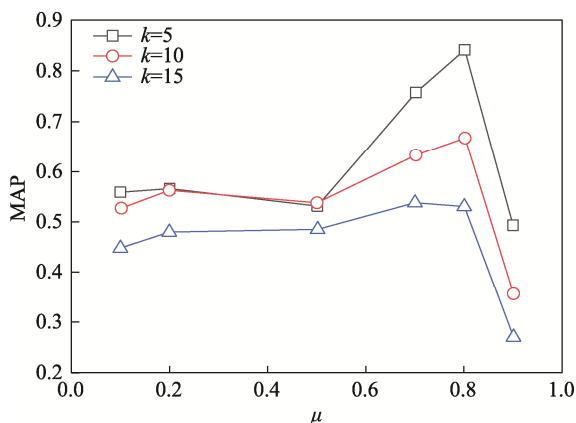


图3 不同参数 μ 对应的MAP值

Fig.3 MAP values corresponding to different parameters μ

不同重启概率 c 对应的MAP值见图4,其中 $k=5, 10, 30$ 。图片与标签的相关性矩阵受到重启概率 c 的影响。结果显示,随着 c 的增长MAP值呈现缓慢上升后急速下降的趋势,并在 $c=0.8$ 时取值最高,随后呈现快速下降趋势。

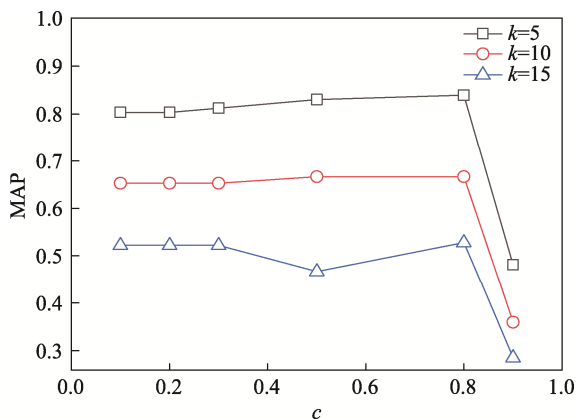


图4 不同重启概率 c 对应的MAP值

Fig.4 MAP value corresponding to different restart probability c

不同迭代次数 l 对应的MAP值见图5,其中 $k=5, 10, 30$ 。MAP值随着迭代次数的增加而不断增加,最终趋于稳定状态,说明采用重启随机游走模型能够保证效果的稳定性。当迭代次数为2时MAP基本趋于稳定,再增加迭代的步长不再有明显的变化,这表明所提方法能够快速达到稳定状态。

不同关键字数量 n 对应的MAP值见图6,其中 $k=5, 10, 30$ 。当 n 从1变化至3时,MAP值快速地增加。当 n 大于3时,MAP值开始轻微地下降,并基本趋于稳定。结果表明,关键字数量为3时,返回结果的效果最好。尽管关键字数量增加会带来一些非重要的文本数据参与查询,但系统采用了TextRank值作为系数来平衡不同重要程度的关键字影响,因此

MAP值并没有明显降低。

不同窗口大小 W 对应的MAP值见图7,其中 $k=5, 10, 30$ 。窗口大小是TextRank在构建“词图”时的关键指标,用来确定2个词对应节点间有无边相连。窗口为3时,MAP值最高,然后呈现下降的趋势。

不同 k 对应的MAP值见图8,其中 $W=1, 2, 3$ 。随着 k 的增加,MAP值呈现下降的趋势,这是因为

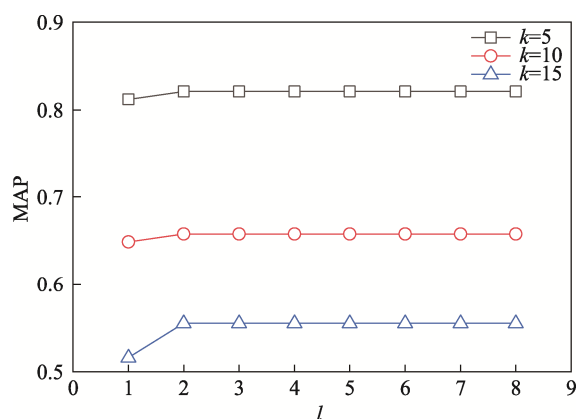


图5 不同迭代次数 l 对应的MAP值

Fig.5 MAP values for different iteration steps l

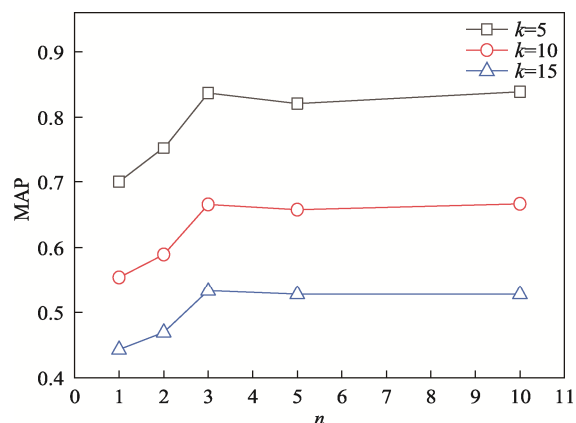


图6 不同关键字数量 n 对应的MAP值

Fig.6 MAP values for different keywords n

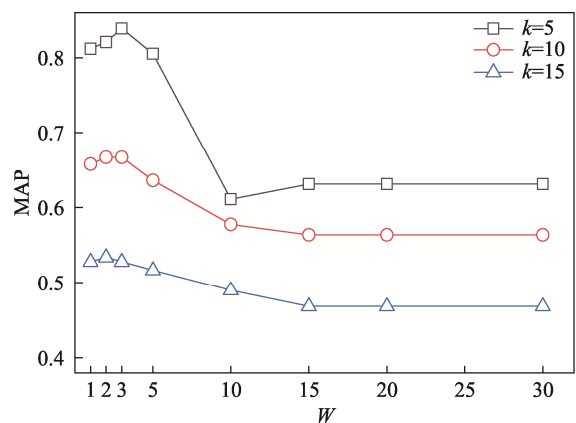


图7 不同窗口大小 W 对应的MAP值

Fig.7 MAP values corresponding to different window sizes W

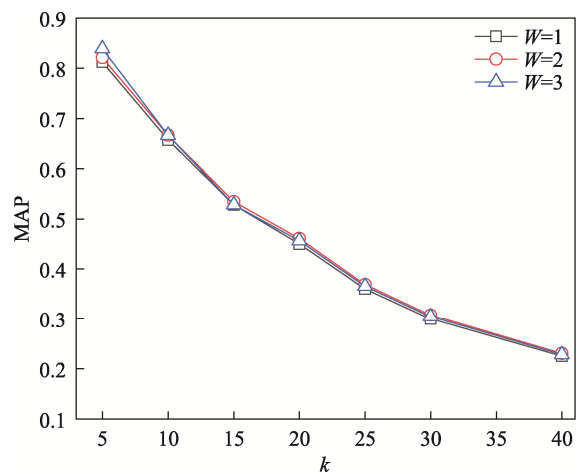


图 8 不同 k 对应的 MAP 值
Fig.8 MAP values corresponding to different k

排序越靠前的图片与输入文本的相关性越大，所对应的 MAP 值越高，反之亦然。实验结果表明所提方法的返回结果中的图片具有合理的排序。

5.4 实例研究

为了直观感受实验效果，随机抽选 3 个文本进行实例展示。文本实例如下所述。

文本 1:那时我们去的时候都已经下午五点多了，不一会儿东方明珠亮了，夜景来了，美丽的上海夜景也是那么的美?站在了东方明珠下面就是不一样的感觉，自己也有有一天在东方明珠下溜着玩了，太幸运了……

文本 2:摩天轮慢慢地升空，我往远处看：高大的楼房变得那么渺小，上面的窗户也只有指甲盖那么小，我又往下看：草坪就像一个“烧饼，”上面还有一些花纹呢！摩天轮已经升到了最高处，再往下看，咦？怎么有这么多“蚂蚁？”哦，原来是人呀……

文本 3:牙齿不仅能咀嚼食物、帮助发音，而且对面容的美有很大影响。由于牙齿和牙槽骨的支持，牙弓形态和咬合关系的正常，才会使人的面部和唇颊部显得丰满……

输入文本的返回结果见表 1。其中，文本 1 内容主要是关于“东方明珠”夜景的描述。在返回结果中，排序为 1 的是关于“东方明珠”夜景的图片，非常符合要求；排序为 2 的是关于“东方明珠”晨曦的图片，仅在天空方面有所不同；排序为 3 的是关于“东方明珠”和“外滩”夜景的图片，仅在图片细节方面有所不同；类似地，排序为 4—5 都是关于“东方明珠”的图片，但在拍摄时间方面有所不同。文本 2 内容主要是关于乘坐“摩天轮”游玩的图片，返回结果中所有图片的主题都是“摩天轮”的风景，很明显与文本内容一致。类似地，文本 3 的返回结果也与用户期望一致。通过分析实例结果发现，系统能够准确地返回用户期望的图片。

表 1 不同文本实例的返回图片序列
Tab.1 Returned sorted pictures for different texts

排序	文本 1	文本 2	文本 3
1			
2			
3			
4			
5			

6 结语

文中设计并开发了一种基于重启随机游走模型的文本配图系统，实现了面向文本数据的在线配图功能。通过网络爬虫获取图片与标签组成的数据集，使用重启随机游走模型得到图片与标签收敛的相关性。利用 TextRank 模型进行关键字提取，并将关键字集合作为查询。通过整合关键字与图片之间的相关性来查找相关图片。在真实数据集上开展大量实验，结果表明所提方法能够实现文本的自动配图。

与现有文字配图、图片搜索等方法相比，文中所提方法具有更加高效和灵活的优势。例如，在文档编辑方面，常见的文本配图为人工配图，这种方法依赖于编辑人员的经验判断进行配图，不仅涉及的素材有限，而且浪费大量的时间精力。所提方法素材来源为开发的互联网环境，通过高效的计算过程能够为用户快速、准确地选择图片。另外，现有搜索引擎也具有依据文本进行图片搜索的功能，但是搜索引擎为了提高查询速度，存在关键字长度的限制，忽视用户对长文本配图的需求，而且还需要用户人工选择合适的关键字进行查询。与搜索引擎不同，文本配图系统中的文本长度不受限制，并基于 TextRank 筛选文本的关键字，能够依据用户不同需求在线推送符合用户期望的图片。

为了进一步完善文本配图的问题，未来研究工作将着重从以下几个方面开展。首先探索如何降低图片与标签的离线计算离线时间开销，以适应大规模数据处理需求。其次，研究提高在线查找的效率，减少相关性在线计算的时间开销，为用户提供更好的查询体验。最后，还将进一步探索文本配图在具体场景中的

应用,例如网页配图、新闻配图、用户表情包推荐等实际场景。

参考文献:

- [1] 杨洋, 项辉宇, 薛真, 等. 模式图像的特征提取与配准方法[J]. 包装工程, 2017, 37(1): 185—190.
YANG Yang, XIANG Hui-yu, XUE Zhen, et al. Feature Extraction and Registration Method of Pattern Image[J]. Packaging Engineering, 2017, 37(1): 185—190.
- [2] YANG Run-qi, ZHANG Jian-hai, GAO Xing, et al. Simple and Effective Text Matching with Richer Alignment Features[J]. Association for Computational Linguistics, 2019, 1: 4699—4709.
- [3] FANG Shan-cheng, XIE Hong-tao, CHEN Zhi-neng, et al. Uyghur Text Matching in Graphic Images for Bio-medical Semantic Analysis[J]. Neuroinformatics, 2018, 16(3/4): 445—455.
- [4] HE Yi-xuan, TAO Shu-chang, XU Jun, et al. Text Matching with Monte Carlo Tree Search[C]// China Conference on Information Retrieval Springer, Cham, 2018: 41—52.
- [5] LEE K H, HAMID P, CHEN Xi, et al. Learning Visual Relation Priors for Image-text Matching and Image Captioning with Neural Scene Graph Generators[J]. Computing Research Repository, 2019(1): 13—29.
- [6] GONG Y, WANG L, HODOSH M, et al. Improving Image-Sentence Embeddings Using Large Weakly Annotated Photo Collections[C]// European Conference on Computer Vision Springer, Cham, 2014: 529—545.
- [7] 邓珮, 孙朔. 基于因子分析的图像语义研究[J]. 包装工程, 2017, 37(20): 124—127.
DENG Pei, SUN Shuo. Research on Image Semantics Based on Factor Analysis[J]. Packaging Engineering, 2017, 37(20): 124—127.
- [8] KARPATY A, LI Fei-fei. Deep Visual-Semantic Alignments for Generating Image Descriptions[J]. IEEE Transactions on Pattern Analysis & Machine Intelligence, 2014, 39(4): 664—676.
- [9] JOHNSON J, KARPATY A, LI Fei-fei. Dense Cap: Fully Convolutional Localization Networks for Dense Captioning[J]. Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2016(1): 4565—4574.
- [10] GUO Ai-zhang, YANG Tao. Research and Improvement of Feature Terms Weight Based on TF-IDF Algorithm[C]// IEEE, 2016: 415—419.
- [11] LI Ying-ying, SHEN Bo. Research on Sentiment Analysis of Microblogging Based on LSA and TF-IDF[C]// IEEE International Conference on Computer and Communications, 2017: 2584—2588.
- [12] 贺婉莹, 杨建林. 基于随机游走模型的排序学习方法[J]. 数据分析与知识发现, 2017, 1(12): 41—48.
HE Wan-ying, YANG Jian-lin. Ranking Learning Method Based on Random Walk Model[J]. Data Analysis and Knowledge Discovery, 2017, 1(12): 41—48.
- [13] JUNG J, PARK N, SAEL L, et al. BePI: Fast and Memory-efficient Method for Billion-scale Random Walk with Restart[C]// ACM International Conference, 2017: 789—804.
- [14] TONG Hang-hang, FALOUTSOS C, PAN Jia-yu. Fast Random Walk with Restart and Its Applications[J]. International Conference on Data Mining, 2006(1): 613—622.
- [15] YU A W, MAMOULIS N, SU Hao. Reverse Top-k Search Using Random Walk with Restart[J]. Proceedings of the VLDB Endowment, 2014, 7(5): 401—412.
- [16] JUNG J, JIN W, SAEL L, et al. Personalized Ranking in Signed Networks Using Signed Random Walk with Restart[C]// International Conference on Data Mining, 2016: 973—978.
- [17] JUNG J, SHIN K, SAEL L, et al. Random Walk with Restart on Large Graphs Using Block Elimination[J]. ACM Trans Database Syst, 2016, 41(2): 1—12.
- [18] YU Wei-ren, JULIE A, MCCANN. Random Walk with Restart over Dynamic Graphs[C]// International Conference on Data Mining, 2016: 589—598.
- [19] WANG Si-bo, TANG You-ze, XIAO Xiao-kui, et al. HubPPR: Effective Indexing for Approximate Personalized Pagerank[J]. PVLDB, 2016, 10(3): 205—216.
- [20] LI Guang-yi, WANG Hou-feng. Improved Automatic Keyword Extraction Based on TextRank Using Domain Knowledge[J]. Communications in Computer & Information Science, 2014, 496: 403—413.
- [21] 蒲梅, 周枫, 周晶晶, 等. 基于加权 TextRank 的新闻关键事件主题句提取[J]. 计算机工程, 2017, 43(8): 219—224.
PU Mei, ZHOU Feng, ZHOU Jing-jing, et al. Topic Sentence Extraction of Key News Events Based on Weighted TextRank[J]. Computer Engineering, 2017, 43(8): 219—224.