

利用双主题表征的涉案微博评价对象识别方法^{*}

相艳^{1,2}, 余正涛^{1,2}, 郭军军^{1,2}, 黄于欣^{1,2}, 线岩团^{1,2}

¹(昆明理工大学 信息工程与自动化学院, 云南 昆明 650504)

²(云南省人工智能重点实验室(昆明理工大学), 云南 昆明 650504)

通信作者: 余正涛, E-mail: ztyu@hotmail.com



摘要: 微博评价对象识别是涉案网络舆情分析的基础. 目前基于主题表征的评价对象识别方法需要预设固定的主题数目, 且最终评价对象识别依赖人工推断. 针对此问题, 提出一种弱监督涉案微博评价对象识别方法, 仅采用少量标签评论即可实现对评价对象的自动识别. 具体实现思路为: 首先基于变分双主题表征网络对评论进行两次编码和重构, 获得丰富的主题特征; 然后, 利用少量标签评论, 引导主题表征网络自动判别评价对象类别; 最后采用联合训练策略, 对双主题表征的重构损失与评价对象分类损失进行联合调优, 最终实现对评价对象的自动分类和评价对象词项的挖掘. 在涉案舆情的两个数据集上进行了实验, 结果表明, 所提出的模型在评价对象分类、评价对象词项的主题连贯性和多样性等方面均优于几个基线模型.

关键词: 评价对象识别; 变分编码; 主题模型; 弱监督学习; 涉案舆情

中图法分类号: TP18

中文引用格式: 相艳, 余正涛, 郭军军, 黄于欣, 线岩团. 利用双主题表征的涉案微博评价对象识别方法. 软件学报. <http://www.jos.org.cn/1000-9825/6408.htm>

英文引用格式: Xiang Y, Yu ZT, Guo JJ, Huang YX, Xian YT. Identification Method of Microblog Opinion Targets Involved in Cases Based on Dual Topic Representation. Ruan Jian Xue Bao/Journal of Software (in Chinese). <http://www.jos.org.cn/1000-9825/6408.htm>

Identification Method of Microblog Opinion Targets Involved in Cases Based on Dual Topic Representation

XIANG Yan^{1,2}, YU Zheng-Tao^{1,2}, GUO Jun-Jun^{1,2}, HUANG Yu-Xin^{1,2}, XIAN Yan-Tuan^{1,2}

¹(Faculty of Information Engineering and Automation, Kunming University of Science and Technology, Kunming 650504, China)

²(Yunnan Key Laboratory of Artificial Intelligence (Kunming University of Science and Technology), Kunming 650504, China)

Abstract: The identification of opinion targets in microblog is the basis of analyzing network public opinion involved in cases. At present, the identification method of opinion targets based on topic representation needs to preset a fixed number of topics, and the final results rely on artificial inference. In order to solve these problems, this study proposes a weak supervision method, which only uses a small number of labelled comments to automatically identify the opinion targets in microblog. The specific implementation is as follows. Firstly, the comments are encoded and reconstructed twice based on the variational dual topic representation network to obtain rich topic features. Secondly, a small number of labelled comments are used to guide the topic representation network to automatically identify the opinion targets. Finally, the reconstruction loss of double topic representation and the classification loss of opinion targets identification are optimized together by the joint training strategy, to classify comments of opinion targets automatically and mine target terms. Experiments are carried out on two data sets of microblogs involved in cases. The results show that the proposed model outperforms several baseline models in the classification of opinion targets, topic coherence, and diversity of target terms.

Key words: opinion targets identification; variational encoder; topic model; weak supervision learning; public opinion involved in cases

* 基金项目: 国家重点研发计划 (2018YFC0830105, 2018YFC0830101, 2018YFC0830100); 云南省重大科技专项计划 (202002AD080001); 云南省基础研究专项面上项目 (202001AT070047, 202001AT070046)

收稿时间: 2020-08-16; 修改时间: 2021-02-04; 采用时间: 2021-06-28

案件相关的负面突发事件通常会引发网友在互联网微博热议,并在短时间内形成传播快、范围广的热点话题,进而产生涉案网络舆情。从大量评论语料中识别出涉案舆情所关心的评价对象,如法律机构、当事人、媒体等,是舆情分析和态势评估等任务的基础。涉案微博评价对象识别的具体任务为:(1)从评论语料中挖掘评价对象词项,并将含义相近的评价对象词项聚集到相应的类别中;(2)将评论句判别为某个评价对象类别。如图1为“奔驰女车主维权案”的评价对象识别框架。对于该案件的微博评论,首先从语料中挖掘“监管部门”“弱势群体”“奔驰”等评价对象词项,并将“监管部门”“工商”和“税务局”等词项聚集为一个评价对象类别“法律机构”。基于此,可以将评论“市场的监管部门就是一个摆设”一句分类至“法律机构”类别。

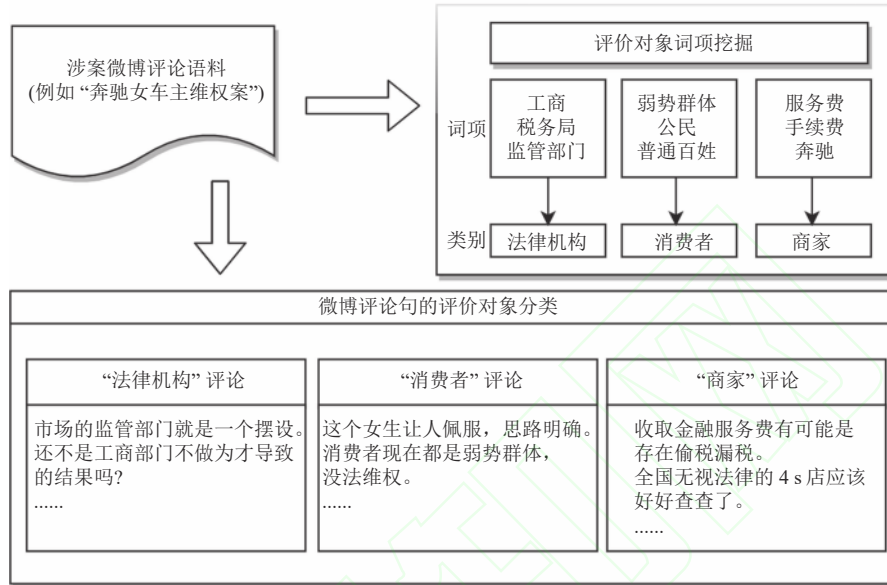


图1 涉案微博评价对象识别示例

在涉案微博评论中,同一案件的评价对象词项较为多样和复杂,例如图1中,“消费者”这一评价对象类别有“弱势群体”“普通百姓”“公民”等差异较大的词项;同时,不同案件的评价对象类别也有区别。针对微博评价对象的识别,目前大都基于主题表征的方法。传统的主题模型是将每类评价对象建模为一个主题,而评论被建模为这些主题的混合,如潜在Dirichlet分布模型(latent Dirichlet allocation, LDA)^[1]。文献[2-4]则改进和扩展了LDA,并将其用于评价对象识别。但这类主题模型仅限于应用到正式和句法良好的长文档中,比如新闻报道和科技文本。当处理涉案微博评论时,由于文本较短和表述复杂导致的数据稀疏问题,会影响这类模型的评价对象识别性能。

近年来,基于深度学习框架的神经主题模型^[5-7]得到了较好的发展,其中He等人^[7]提出了基于注意力的自编码模型(attention-based aspect extraction, ABAE),该模型利用数据集上预训练的词向量来获取词共现的分布,并基于自编码的框架来预测句子的评价对象概率分布,从而识别评价对象。与传统的基于多项式词分布的主题模型相比,基于连续空间构建的神经主题模型可以更好的处理低频词,从而在短文本评价对象识别任务中,取得比LDA等传统主题模型更好的识别效果。但是,这类神经主题模型用于涉案微博评价对象识别仍然存在以下不足:(1)模型只对文本进行一次重构,并固定了主题数目,这限制了模型对主题表征的学习。(2)模型可以获取若干组词项来表示不同评价对象类别,但某组词项究竟表示哪类评价对象则需要人工推断。如果某组词项难以推断,则会直接影响句子分类结果。

针对以上两个问题,本文提出一种基于变分双主题表征的弱监督评价对象识别模型,以更好的实现评价对象词项挖掘和句子自动分类。本文的主要贡献有:

(1)提出一种变分双主题表征网络,通过对微博评论语料进行辅助主题和核心主题的变分编码与重构,使模型能学习到更全面的主题表征,从而更好的进行评价对象词项挖掘和句子分类。

(2) 提出一种弱监督评价对象识别框架, 联合优化少量标签评论的分类损失和大量无标签样本的重构损失, 从而使主题表征网络学习自动判别评价对象类别. 该弱监督学习框架避免了用评价对象词项来推断类别的人工过程, 因此可以实现评论句的自动分类.

(3) 在案件舆情领域的两个不同的数据集上验证了本文模型. 与多个无监督或弱监督的评价对象识别模型相比, 本文模型在评价对象分类方面具有更优异的性能, 同时能发现更多样和更有价值的评价对象词项.

1 相关工作

目前主流的评价对象识别方法可以分为有监督方法、无监督方法和弱监督方法.

有监督方法借鉴了评论文本实体属性抽取的思路, 目标是找出评价对象在文本中的位置, 主要有基于规则或特征的方法^[8-10]和基于序列标注的方法^[11-14]. 序列标注的方法中, 基于条件随机场的模型^[12,13]因有效利用了相邻词标签的依赖关系而取得了不错的结果. Xu 等人^[14]则以卷积神经网络 (convolutional neural networks, CNN) 为基本框架, 融合了通用领域和特定领域的嵌入信息, 取得了比之前模型更好的抽取结果. 在前人工作基础上, 文献^[15,16]进一步利用了观点词的信息, 对评价对象和观点词进行联合建模, 从而改善了评价对象抽取性能. 上述有监督模型依赖于大量手动标注的数据, 但对于微博评论而言, 不同案件的评价对象表述通常较为复杂和多变, 数据标注成本较高, 同时监督模型对新案件数据的泛化能力也较差. 以上问题限制了有监督模型在涉案微博评价对象抽取中的使用.

事实上, 微博评价对象识别更适合采用无监督的方法. 无监督方法通过对评论文本进行特征聚类, 获取表征评价对象的词项, 进而实现评价对象的识别. 其中一类方法主要基于文本统计特征的方法. 例如文献^[17]采用 TF-IDF 特征获取候选特征词集, 并基于期望最大化算法挖掘特征词项, 进而识别出评论文本中的评论对象. 另一类则是以主题表征为基础的方法^[17,18]. 传统的主题模型大都基于吉布斯采样或变分期望值最大化算法^[1,17], 较适合处理长文本, 但在微博评论这样的短文本上难以取得较好的效果. 为了解决这一问题, 一些研究工作将单词嵌入^[19-22]和预训练知识^[23]等外部表征融入主题建模过程中, 实现评价对象抽取. 另一些工作则考虑了丰富短文本的上下文, 如 Biterm 主题模型 (biterm topic model, BTM)^[24,25], 该模型将一个文本扩展到一个 biterm 集合, 集合中包含文本中出现的任意两个不同单词的所有组合. BTM 的建模方式可以缓解短文本主题建模的问题, 但在稀疏数据设置中仍然存在不足.

随着深度学习技术的不断发展, 基于深度神经网络框架的无监督主题模型在评价对象抽取任务上取得较好效果. 其中变分推断自编码模型 (auto-encoding variational Bayes, VAEs)^[5]使用一个推断网络直接将文档映射到它的后验分布, 该方法不仅可以提升模型的鲁棒性, 还有效的降低了计算成本. 更进一步地, 文献^[26]将 VAEs 和 BTM 进行结合, 提出了一种 GraphBTM 方法, 通过图卷积网络嵌入聚合的 biterm 图, 进一步缓解数据稀疏的问题. 文献^[27]同时对文本进行序列编码和 VAEs 主题建模, 综合利用了文本的全局表示和局部表示, 因此能更好的抽取到低频的评价对象词项. 文献^[28]也利用类似思想, 用神经网络来编码文本语义特征, 提升主题模型的语义连贯性. Dieng 等人^[29]则提出一种嵌入主题模型 (embedded topic model, ETM), 将传统主题模型与词嵌入结合在一起, 使用分类分布对每个词进行建模, 该方法可以学习到更连贯的语言模式和准确的单词分布. 此外, He 等人^[7]提出的 ABAE 模型通过最小化句子重构错误来训练参数, 在两个产品评论数据集上取得了较好的评价对象分类结果, 并挖掘出具有较好主题连贯性的词项.

不同于无监督的方法, 弱监督方法可以融合领域知识来进行建模, 因此可以在一定程度上缓解无监督方法对于特定评价对象抽取的不足. 通常而言, 领域知识可以是一组指定的评价对象种子词, 例如 Lu 等人^[30]在概率主题模型中使用种子词作为领域先验知识, Angelidis 等人^[31]则提出一种多种子评价对象抽取模型 (multiseed aspect extractor, MATE), 为每个特定的评价对象选取若干种子词. 这类弱监督模型虽然可以获取特定词项, 但在将评论句区分为不同评价对象类别方面同样存在人工判别困难的问题.

综上所述, 有监督模型需要对每个案件数据进行大量标注, 人工成本较高; 现有的无监督和弱监督模型虽然可以快速获取评价对象词项, 但在评价对象分类中仍然需要人工判断每个主题所对应的类别. 为此, 本文提出一种弱

监督的评价对象识别方法, 通过对评论句进行两次不同的编码和重构, 同时利用少量分类标签数据引导, 使模型能自动分类评价对象, 挖掘评价对象词项.

2 基于变分双主题表征的评价对象识别模型

本文模型的整体思路包括如下 3 个部分: 首先基于无监督变分双主题表征网络对微博评论进行编码和重构, 在相同的向量空间中建立词向量、句向量和主题表征向量之间的关系; 然后基于少量标签样本的类别信息对评价对象识别过程进行微调; 最后联合训练无监督变分双主题重构损失和有监督标签分类损失, 实现对评论句的自动分类. 模型整体结构如图 2 所示.

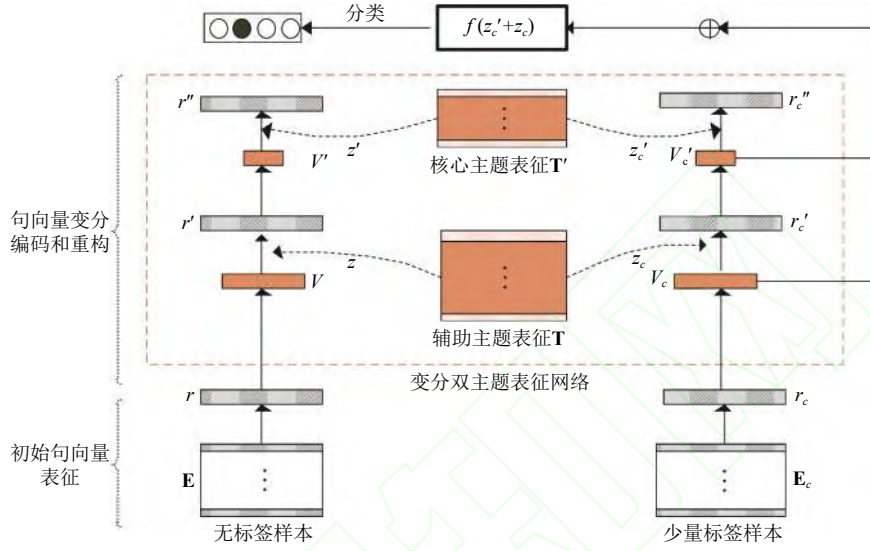


图 2 本文模型结构图

2.1 基于变分双主题表征网络的句向量编码和重构

变分双主题表征网络可以实现对微博评论的无监督编码和重构, 包括: 句向量始表征和句向量变分编码与重构两个主要模块.

2.1.1 初始句向量表征

首先预训练数据集的词向量, 得到每个词的词向量. 用 $e_{x_i} \in \mathbb{R}^D$ 表示输入句子 x 中第 i 个单词的词向量, 则句子嵌入为 $\mathbf{E} = e_{x_1} \oplus e_{x_2} \oplus \dots \oplus e_{x_n}$, $\oplus$ 是拼接操作, $\mathbf{E} \in \mathbb{R}^{n \times D}$, n 是句子长度, D 是词向量维度. 利用注意力来计算句子的原始评价对象向量, 具体操作如下所示.

$$A = (\mathbf{E}\mathbf{M} + b\mathbf{u}^T)\mathbf{E}^T \quad (1)$$

$$\theta_i = \frac{1}{n} \sum_{j=1}^n A_{ij} \quad (2)$$

$$\alpha_i = \frac{\exp(\theta_i)}{\sum_{i=1}^n \exp(\theta_i)} \quad (3)$$

$$r = \sum_{i=1}^n \alpha_i e_{x_i} \quad (4)$$

其中, $\mathbf{M} \in \mathbb{R}^{D \times D}$, $b \in \mathbb{R}^n$ 为待优化的参数, $\mathbf{u} \in \mathbb{R}^D$ 是值全为 1 的向量, $A \in \mathbb{R}^{n \times n}$ 是自注意力矩阵, θ_i 是句子中第 i 个词的权重, 而 α_i 是第 i 个词的归一化权重. 通过注意力操作, 输入句子被表示为初始的句向量 $r \in \mathbb{R}^D$, 它更多地关

注与评价对象相关的单词.

2.1.2 句向量变分编码和重构

将初始句向量进行两次变分编码和重构, 包括基于辅助主题表征的编码和重构, 以及基于核心主题表征的编码和重构, 进而得到句子的辅助主题分布和核心主题分布. 辅助主题表征的主题数目设置为较核心主题表征更大的值, 因此辅助主题向量代表向量空间中较小的主题聚类簇; 核心主题表征则对应于较大的主题聚类簇. 基于不同大小聚类簇的编码与重构能使句子学到更多的主题特征.

(1) 基于辅助主题表征的编码与重构

将初始句向量 r 用变分网络编码为隐向量 $z \in \mathbb{R}^K$, z 为 K 维的辅助主题分布, 其中的某个值 z_l 表示输入句子 x 属于第 l 个评价对象的概率. 变分编码结构如图 3 所示. 假设 z 服从正态分布 $N(\mu, \sigma)$, 则:

$$z = \mu + \sigma \odot \varepsilon \quad (5)$$

其中, $\mu = d_1(r)$, $\log \sigma = d_2(r)$, d_1 和 d_2 为两个线性变换层, ε 为服从正态分布的随机值.

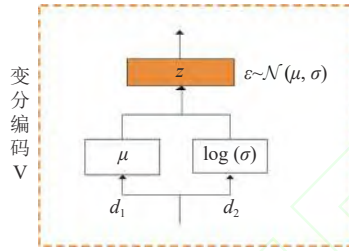


图 3 变分编码结构图

用 $t_l \in \mathbb{R}^D$ 表示数据集中第 l 个评价对象的 D 维向量, 则辅助主题表征为 $\mathbf{T} = t_1 \oplus t_2 \oplus \dots \oplus t_K$, $\mathbf{T} \in \mathbb{R}^{K \times D}$, 辅助句向量可以由 p_z 和 $\mathbf{T}$ 进行重构, 过程如下所示.

$$p_z = \text{softmax}(z) \quad (6)$$

$$r' = \mathbf{T}^T p_z \quad (7)$$

其中, r' 为重构的辅助句向量.

(2) 基于核心主题表征的编码与重构

将辅助句向量 r' 用如图 3 所示的变分编码为隐向量 $z' \in \mathbb{R}^{K'}$, z' 代表 K' 维的核心概率分布, 则分布中的某个值 z'_l 表示输入句子 x 属于第 l 个核心评价对象的概率. 同样假设 z' 服从正态分布 $N(\mu', \sigma')$, 则

$$z' = \mu' + \sigma' \odot \varepsilon \quad (8)$$

其中, $\mu' = d_1'(r')$, $\log \sigma' = d_2'(r')$, d_1' 和 d_2' 为两个线性变换层, ε 为服从正态分布的随机值.

用 $t'_l \in \mathbb{R}^D$ 表示数据集中第 l 个核心主题的 D 维向量, 则核心主题表征为 $\mathbf{T}' = t'_1 \oplus t'_2 \oplus \dots \oplus t'_{K'}$, $\mathbf{T}' \in \mathbb{R}^{K' \times D}$, 句子的核心句向量可以由 $p_{z'}$ 和 $\mathbf{T}'$ 进行重构.

$$p_{z'} = \text{softmax}(z') \quad (9)$$

$$r'' = \mathbf{T}'^T p_{z'} \quad (10)$$

其中, r'' 为第 2 次重构的核心句向量.

2.1.3 无监督编码重构的损失

根据以上步骤, 得到输入句子的 3 个表征, 即初始句向量 r , 重构的句向量 r' 和 r'' . 对于输入句子, 从数据集中随机采样 num 个句子作为负样本, 将每个负样本向量用其平均词向量 n_i 来表征. 如果模型两次重构的句向量较为合理, 则它们与初始句向量应该较为相似, 而和负样本向量尽量不同. 因此, 第 j 个句子的损失可使用铰链损失 $J_j(\theta)$, 即最大化 r' 、 r'' 和 r 之间的内积, 同时最小化 r' 、 r'' 和负样本之间的内积, 计算公式如下.

$$J_j(\theta) = \sum_{i=1}^{num} \max(0, 1 - \lambda r r' - r r'' + \lambda r' n_i + r'' n_i) \quad (11)$$

其中, λ 是一个超参数, 用于控制辅助主题重构的权重.

将数据集中所有句子的重构损失 $J_j(\theta)$ 加和, 得到模型的无监督重构损失 $J(\theta)$.

2.2 基于少量标签数据的评价对象类别预测

对于少量的标签数据, 同样用第 2.1 节的方法进行句子编码和重构. 其中的注意力层、句子两次重构所用到的辅助主题表征 $\mathbf{T}$ 和核心主题表征 $\mathbf{T}'$ 是与无标签数据共享参数的, 而两次变分编码所用到的线性变换层 d_{1c} 、 d'_{1c} 、 d_{2c} 、 d'_{2c} 则与无标签数据不同. 将标签数据的辅助主题分布 z_c 和核心主题分布 z'_c 进行拼接, $z_{c_all} = z_c \oplus z'_c$, $z_{c_all} \in \mathbb{R}^{K+K'}$, 之后将拼接的特征用于分类, 计算出标签数据属于评价对象类别的概率 r . 计算公式为:

$$r = z_{c_all} W_c + b_c \quad (12)$$

用 Softmax 对 r 进行归一化, 得到模型所预测的评价对象类别 y :

$$y = \text{softmax}(r) \quad (13)$$

最后分类损失采用交叉熵代价函数计算:

$$J_c(\theta) = \sum_{i=1}^C g_i \log(y_i) \quad (14)$$

其中, g_i 表示真实的评价对象类别标签.

2.3 损失函数的联合优化

通过最小化无监督重构损失 $J(\theta)$, 可以优化主题表征网络参数; 通过最小化分类损失 $J_c(\theta)$, 则可以优化模型的分类网络参数. 考虑到两个优化目标互有影响, 因此, 本文采用联合训练策略, 同时优化重构损失 $J(\theta)$ 和分类损失 $J_c(\theta)$. 此外, 评价对象类型可能遭遇冗余问题, 因此参考文献 [7] 的做法, 在损失函数中加入两个正则项, 确保评价对象的多样性.

$$V'(\theta) = \|\mathbf{T}'_n \cdot \mathbf{T}'_n^T - \mathbf{I}\| \quad (15)$$

$$V(\theta) = \|\mathbf{T}_n \cdot \mathbf{T}_n^T - \mathbf{I}\| \quad (16)$$

其中, $\mathbf{I}$ 是单位矩阵, $\mathbf{T}'_n$ 是 $\mathbf{T}'$ 的行归一化, $\mathbf{T}''_n$ 和 $\mathbf{T}''$ 也是如此. 当任意两个不同行向量的内积为零时, V' 和 V'' 达到它们的最小值. 因此, 正则化项鼓励主题表征的各行向量之间的正交性, 并惩罚不同行向量之间的冗余. 最终的目标函数 $L(\theta)$ 为:

$$L(\theta) = J(\theta) + \alpha J_c(\theta) + \beta V'(\theta) + \beta V''(\theta) \quad (17)$$

其中, α 是控制分类损失权重的超参数, β 是控制评价对象多样性权重的超参数.

模型学习目标是通过优化参数来最小化目标函数 $L(\theta)$. 模型训练完成后, 可以通过公式 (13) 中的分类概率将测试句子分类到对应的评价对象类别, 并选择词向量最接近于核心主题表征中某个行向量的前 n 个词作为对应评价对象类别的词汇.

3 实验与分析

3.1 数据集

本文采集了 2 个案件的新浪微博评论数据集, 进行模型训练和评估. 3 名经过培训的专业人员同时给评论打标签, 将评论标记为涉案舆情较重要的评价对象类别, 最终选择标签一致的评论. 数据集基本信息如表 1 所示. 案件 1 为奔驰女车主维权案, 数据集包含 44907 条无标签样本, 有 4 种标注的评价对象类别, 分别为法律机构、商家(当事人)、消费者(当事人)、其他, 标签样本共 1925 条. 案件 2 为重庆公交车坠江案, 数据集包含 23705 条无标签样本, 有 4 种手动标注的评价对象类别, 分别为政府机构、公交司机(当事人)、媒体、其他, 标签样本共 1660 条. 两个数据集均划分 80% 的标签样本作为最终分类性能评估的测试集.

3.2 实验参数设置

采用 Skipgram 模型^[32]初始化数据集的词向量, 将向量维度设置为 200, 窗口大小设置为 5, 负样本大小设置

为 5. 使用 k 均值聚类将词向量聚类为不同的簇, 并取簇的质心向量初始化主题表征中的主题向量, 其他模型参数随机初始化. 模型训练过程中词向量固定不动, 使用 Adam 优化器, 学习率为 0.001, 为避免过拟合, 采用了 dropout 策略. 通过实验比较, 将核心主题数目设置为 10 个, 辅助主题数目设置为 20 个, 公式 (11) 中的负样本数目为 20. 超参数 α 设置为 1, β 设置为 0.1, λ 设置为 0.5.

表 1 实验数据集

数据集	无标签样本	标签样本	词表大小
案件1: 奔驰女车主维权案	44 907	法律机构: 865 商家(当事人): 640 消费者(当事人): 290 其他: 130	45 023
案件2: 重庆公交车坠江案	23 705	政府机构: 286 公交司机(当事人): 564 媒体: 662 其他: 146	17 017

3.3 基线模型

为了验证本文模型的性能, 将其与一些基线进行了比较. 第 1 类是无监督模型, 包括 LDA、BTM、ETM、ABAE 和 Ours_unlabelled, 需要通过人工推断评价对象的类别标签来对句子进行分类. 第 2 类是弱监督主题模型, 包括 MATE 和 ABAE_lablled. 为了公平比较不同模型在评价对象词项挖掘方面的性能, 所有基线模型的主题数目设为与本文模型的核心主题数目相同, 即设为 10.

- LDA^[1]: 采用 LDA 的标准实现, 参数推理采用的是 Gibbs 采样, 每句评论都作为一个单独的文档处理. 设置 $\alpha=0.05$ 和 $\beta=0.1$, 并运行 100 次 Gibbs 采样迭代.

- BTM^[24]: 专门为短文本设计的 Biterm 主题模型, 通过生成无序的词对共现来缓解短文本中的数据稀疏问题. 对于 BTM, 设置 $\alpha = 50/K$ 和 $\beta = 0.005$.

- ETM^[29]: 是一种将传统主题模型与词嵌入结合在一起的模型, 其主题词分布是词嵌入及其嵌入的指定主题之间的内积. ETM 的词嵌入维度设为 200, 使用 Adam 优化器, 学习率为 0.005.

- ABAE^[7]: 一种无监督的神经主题模型. ABAE 的词向量维度为 200, 使用 Adam 优化器, 学习率为 0.001.

- Ours_unlabelled: 将本文模型中标签数据的分类损失去除, 只用重构损失进行参数优化. 最终评论的评价对象类别用其核心主题分布中的最大概率值进行判断. 其余模型参数设置与本文模型一致.

- MATE^[32]: 是在 ABAE 模型基础上改进的一种弱监督学习模型, 为每个感兴趣的评价对象选取若干种子词, 使用相应种子词嵌入的平均值初始化评价对象嵌入, 并在整个训练过程中固定. 该模型也需要通过人工推断评价对象类别标签来对句子进行分类. 本文实验中, 为每类评价对象选取 10 个种子词. 其余模型参数设置与 ABAE 一致.

- ABAE_lablled: 以 ABAE 为基本框架, 对有标签和无标签样本进行编码和重构, 并以标签样本的主题分布为分类特征. 采用与本文相同的模式, 联合训练重构损失和分类损失. 模型训练完成后可直接进行评价对象分类. 其模型参数设置与 ABAE 一致.

3.4 实验结果

3.4.1 评价对象类别推断

首先比较了传统主题模型 BTM、基于词嵌入的主题模型 ETM 和本文模型挖掘到的案件 1 中表征主题的前 10 个 (top10) 评价对象词项, 如表 2 所示. 每个模型有 10 个主题, 表 2 中列举了 6 个主题.

评价对象类别推断针对的是建模的 10 个主题, 根据每个主题的若干个词项人工判别其所对应的评价对象类别, 不同主题可能被判别为同一个评价对象类别. 从表 2 可以看出, 几种主题模型都存在难以根据评价对象词项推断出评价对象类别的情况, 如 BTM 对应于“法律机构”类别的主题 1, ETM 对应于“消费者”类别的主题 6, 以及本文模型对应于“商家”类别的主题 4, 代表词项的指示性都不强, 极有可能推断错误. 这 3 种模型还有其他主题的代

表词项未在表 2 中列出,也难以推断评价对象类别.这是所有的无监督主题模型用于评价对象判别时的主要问题.此外,相比其他模型,本文模型所挖掘到的同类评价对象代表词项更为相似,更容易推断出评价对象类别.这是因为本文模型利用了主题向量和词向量在向量空间中的关系,相近的词更容易聚集为一类主题.在案件 2 数据集上有类似的实验结果.这些获得的评价对象词项的主题连贯性和多样性在第 3.4.3 节中说明.

表 2 案件 1 的评价对象词项

推断的评价对象类别	模型	top10的评价对象词项
法律机构	本文模型	主题1: 纪委 市长 工商部门 热线 没人接 组建 公安局 消协 物价局 陕西省 主题2: 处罚 法规 违反 违法行为 经营 给予 三倍 甩锅 重罚涉嫌
	BTM	主题1: 消费者 维权 部门 回复 中国 国家 法律 奔驰 政府 社会 主题2: 西安 回复 陕西 违法 成立 年 政府 维权 中国 地方
	ETM	主题1: 维权 消费者 部门 国家 法律 政府 老百姓 相关 监管部门 商家 主题2: 西安 央视 盗抢 高太多 探访 连夜 全国 陕西 工商局 成立
商家	本文模型	主题3: 手续费 服务费 按揭 贷款 交 费用 利息 一万 买房子 全款 主题4: 转向 方向盘 功能 车体 失控 偷换 灯亮 开着 公里 附送
	BTM	主题3: 服务费 金融 回复 贷款 说 奔驰 买收 买车 销售 主题4: 奔驰 万 车 漏油 买 回复 版 换 发动机 说
	ETM	主题3: 服务费 金融 4s店 销售 买车 贷款 收取 汽车 费用 手续费 主题4: 奔驰 漏油 事件 奔驰车 品牌 发动机 宝马 66 探访 德国
消费者	本文模型	主题5: 弱势群体 公民 太难 艰难 底层 普通百姓 践踏 薄弱 精神 内需 主题6: 句句 有理有据 条理清晰 在理 思路清晰 赞 姐姐 清晰 哑口无言 音频
	BTM	主题5: 奔驰 车主 回复 女 说 维权 真的 消费者 小姐姐 中国 主题6: 回复 奔驰 说哭 坐在 女子 维权 问 禅师 110
	ETM	主题5: 真的 感觉 说话 哈哈 逻辑 女高管 清晰 有理有据 高管 本来 主题6: 车主 希望 事情 小姐姐 解决 这件 支持 关注 道歉 声援

3.4.2 评价对象分类

对于句子的评价对象分类,基线模型先根据句子主题分布的最高概率值得到其对应的主题,再根据主题与评价对象类别之间的映射关系将评价对象类别分配给句子.例如,对于“工商执法等部门介入只是退款退车”一句,如果主题分布最大值是表 2 中主题 2,而主题 2 被人工推断为“法律机构”类别,那么该句即被标记为“法律机构”类别.对于本文模型和 ABAE_labelled,则通过句子的最大分类概率得到其评价对象类别.本文评估了句子的评价对象分类在两个案件数据集上的性能.评估指标是精度 (precision, P)、召回率 (recall, R) 和 $F1$ 值,结果见表 3 和表 4.对于本文模型和 ABAE_labelled,表中列出的结果是使用数据集有标签样本的约 12% 用于分类训练得到.

表 3 不同模型对于案件 1 的评价对象分类结果

评价对象类别	法律机构			商家			消费者			weighted-average		
模型	<i>P</i>	<i>R</i>	<i>F1</i>	<i>P</i>	<i>R</i>	<i>F1</i>	<i>P</i>	<i>R</i>	<i>F1</i>	<i>P</i>	<i>R</i>	<i>F1</i>
LDA	0.737	0.576	0.647	0.513	0.606	0.556	0.512	0.441	0.474	0.621	0.565	0.586
BTM	0.798	0.85	0.823	0.546	0.331	0.412	0.464	0.855	0.602	0.654	0.666	0.641
ETM	0.762	0.653	0.703	0.664	0.666	0.665	0.627	0.476	0.541	0.705	0.629	0.663
ABAE	0.731	0.665	0.661	0.732	0.638	0.682	0.498	0.731	0.592	0.694	0.666	0.657
Ours_unlabelled	0.781	0.734	0.757	0.859	0.475	0.612	0.498	0.731	0.592	0.763	0.641	0.678
MATE	0.781	0.734	0.757	0.668	0.591	0.627	0.392	0.738	0.512	0.678	0.684	0.671
ABAE-labeled	0.853	0.702	0.77	0.716	0.703	0.709	0.441	0.763	0.559	0.751	0.71	0.721
本文模型	0.854	0.81	0.831	0.782	0.753	0.768	0.645	0.814	0.72	0.795	0.79	0.791

观察表 3 和表 4 的结果,可以有以下分析: (1) 从 3 个评价对象类别的 weighted-average 评估指标结果来看,无监督的 LDA、BTM、ETM 和 ABAE、Ours_unlabeled 等未使用标签数据的主题模型需要通过代表词项来推断评价对象类别,影响了分类性能,表现都不太理想.需要说明的是,表中列出的结果是从人工推断的结果中选取了每种模型多组测试的最优结果. LDA 的分类结果在两个数据集上都最差, ETM 稍好. Ours_unlabeled 相比 ABAE 有

0.02 和 0.03 的 $F1$ 值提升, 说明本文模型通过两次变分编码和解码能得到更好的核心主题分布, 有利于评价对象分类. (2) 对于案件 1 数据集, MATE 的性能有所提升, 但在案件 2 数据集上则反而性能下降, 原因可能与种子词的选择质量有关. 本文提出的弱监督方式也是一种先验领域知识的指导, 但只需要人工标注少量句子标签. 根据评论判别评价对象类别相比根据整个语料挑选评价对象种子词来说更加直观和便利, 且分类结果提升较明显. 相比 MATE, 本文模型的加权平均 $F1$ 值在两个数据集上分别提升了 0.13 和 0.176. ABAE_labelled 在 ABAE 基础上加入标签样本训练分类器, 相比原来的 ABAE 模型也有较大的提升, 两个数据集的加权宏平均 $F1$ 值相比 ABAE 分别提升了 0.064 和 0.088, 再次证明了本文提出的利用少量有标签样本进行评价对象类别指导的有效性. (3) 对于 3 种不同的评价对象类别, 本文模型和 ABAE-labelled 通常会在标签样本多的类别上具有更高的分类性能, 而 LDA、BTM、ETM、ABAE 和 Ours_unlabelled 等未使用标签数据的主题模型则没有这个规律, 分类性能主要取决于人工推断的标签和评论主题分布的合理性.

表 4 不同模型对于案件 2 的评价对象分类结果

评价对象类别	政府机构			公交司机			媒体			weighted-average		
模型	P	R	$F1$	P	R	$F1$	P	R	$F1$	P	R	$F1$
LDA	0.771	0.525	0.625	0.508	0.264	0.347	0.825	0.685	0.748	0.745	0.545	0.626
BTM	0.513	0.549	0.531	0.869	0.507	0.64	0.657	0.849	0.74	0.643	0.672	0.643
ETM	0.763	0.306	0.437	0.683	0.868	0.765	0.747	0.822	0.783	0.741	0.638	0.65
ABAE	0.513	0.556	0.534	0.49	0.528	0.508	0.759	0.846	0.8	0.616	0.677	0.645
Ours_unlabelled	0.513	0.549	0.531	0.869	0.507	0.64	0.757	0.837	0.795	0.687	0.667	0.667
MATE	0.513	0.549	0.531	0.624	0.668	0.645	0.654	0.84	0.735	0.596	0.698	0.642
ABAE-labelled	0.628	0.848	0.721	0.776	0.759	0.768	0.714	0.753	0.728	0.694	0.79	0.733
本文模型	0.754	0.799	0.776	0.853	0.806	0.829	0.832	0.867	0.849	0.807	0.83	0.818

3.4.3 主题连贯性和多样性

采用主题连贯性来评估模型挖掘评价对象词项的性能, 这是一种基于单词共现的度量评价对象质量的方法^[33], 定义为:

$$coh(t, V^{(t)}) = \frac{2}{M(M+1)} \sum_{m=2}^M \sum_{l=1}^{m-1} \log \frac{D(v_m^{(t)}, v_l^{(t)}) + 1}{D(v_l^{(t)})} \quad (18)$$

其中, $V^{(t)}$ 表示在主题 t 中包含有 M 个概率最大的代表性词项. $v_m^{(t)}$ 和 $v_l^{(t)}$ 是 $V^{(t)}$ 中的第 m 个和第 l 个词, $D(v_l^{(t)})$ 是数据集中包含 $v_l^{(t)}$ 的句子数, $D(v_m^{(t)}, v_l^{(t)})$ 是同时包含 $v_m^{(t)}$ 和 $v_l^{(t)}$ 的句子数. 通常, 主题连贯性得分越高, 主题的语义相关性越强. 用以下公式计算平均主题连贯性得分:

$$COH = \sum_{t=1}^K coh(t, V^{(t)}) \quad (19)$$

所有模型的主题数 K 设为 10, 分别计算前 10 个 (top 10) 至前 50 个 (top 50) 词项的平均主题连贯性得分, 如图 4 所示.

可以看出案件 1 的主题连贯性不如案件 2, 但两个数据集的总体差距不大. 此外, 所有模型在前 10 个至 40 个词项上的主题连贯性表现没有较大差异, 而在前 50 个词项的主题连贯性上差距明显, 其中 BTM 最好, ETM 次之, 包括本文模型在内的其余模型差距不大.

此外, 结合主题多样性来评估模型. 主题多样性定义为在评价对象词项中不重复的词项占有所有词项的百分比^[33]. 多样性接近 0 表示词项冗余, 多样性接近 1 表示词项更加多样. 两个数据集的多样性结果如表 5 和表 6 所示.

可以看出 BTM 虽然具有最好的主题连贯性, 但多样性是最差的, 即 BTM 抽取到的重复词项较多. 这个特性也在表 2 中有所体现, 例如 BTM 在几个主题中重复抽取到“中国”“奔驰”等词项. LDA、ABAE、MATE 和本文模型的主题多样性较好, 说明这几种模型挖掘到的评价对象词项最为丰富. 尤其是 ABAE、MATE 和本文模型的主题多样性基本在 0.99 以上, 说明这类以词嵌入为基础模型能挖掘出更具多样性和独特性的词项.

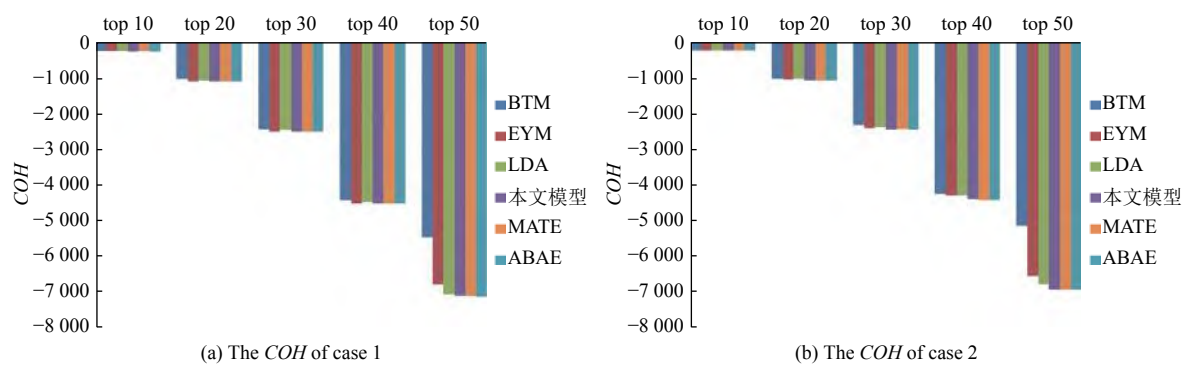


图 4 不同评价对象词项的平均主题连贯性得分

表 5 不同模型的案件 1 主题多样性

词项	LDA	BTM	ETM	ABAE	MATE	本文模型
top 10	0.95	0.64	0.95	0.99	0.99	0.99
top 20	0.945	0.685	0.93	0.99	0.98	0.995
top 30	0.937	0.673	0.91	0.993	0.99	0.997
top 40	0.932	0.6825	0.8825	0.997	0.997	0.997
top 50	0.938	0.606	0.846	0.998	0.99	0.998

表 6 不同模型的案件 2 主题多样性

词项	LDA	BTM	ETM	ABAE	MATE	本文模型
top 10	0.94	0.71	0.97	1	1	1
top 20	0.935	0.655	0.95	1	1	1
top 30	0.913	0.656	0.903	1	0.99	0.997
top 40	0.916	0.64	0.885	0.99	0.99	0.997
top 50	0.918	0.498	0.858	0.994	0.994	0.994

3.4.4 不同训练样本数量的结果比较

对于两个数据集的标签样本, 分别随机选取占标签样本总数的 6%、9%、12%、15% 进行训练, 用划分好的测试集测试. 所得分类结果如图 5 所示.

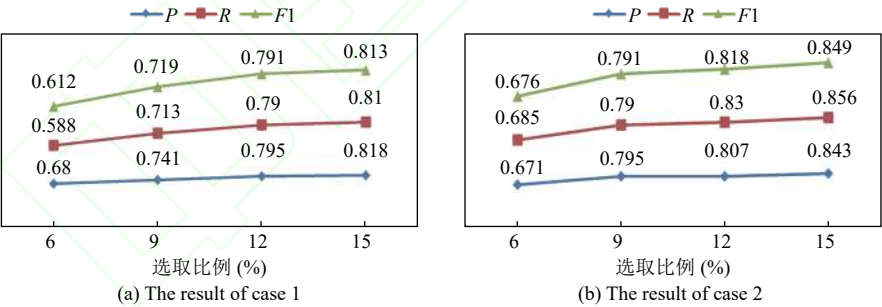


图 5 不同训练样本数量的分类结果

可以看出, 随着标签样本的增加, 精确率 P 、召回率 R 和 $F1$ 值都明显增加, 其中 R 的提升最为显著, 说明随着标签样本分类信息的加入, 模型能从数据集中找出更多特定评价对象的评论. 当使用了 6% 左右的有标签样本 (约 100 个样本) 时, 模型在两个数据集上即可达到 0.612 和 0.676 的 $F1$ 值, 这已经在案件 2 数据集上稍优于表 2 中大多数无监督模型. 随着标签样本的增加, 模型分类性能不断提升. 当使用了 12% 左右的有标签样本 (约 200 个样本) 时, 模型的 $F1$ 值已经达到 0.79 和 0.818, 相比表 2 中的所有其他模型具有明显的优势. 继续增加至 15% 的

有标签样本, 模型 $F1$ 值提升至 0.813 和 0.849. 总体来说, 本文模型只需要少量的标签样本即可达到较高的分类性能.

3.4.5 双主题表征有效性实验

对本文所提出的双主题表征网络进行了消融实验. 将本文模型的辅助主题重构去除, 即模型只对句向量进行一次重构, 学习一个主题表征, 标签样本也只使用一种主题分布作为分类特征. 对于一次重构学习的主题表征, 分别设置主题数 K 为 10, 20 和 30, 而本文完整模型的核心主题数为 10, 辅助主题数为 20. 比较结果如图 6 所示, 其中案件 1 使用了 12% 的标签样本, 案件 2 使用了 15% 的标签样本.

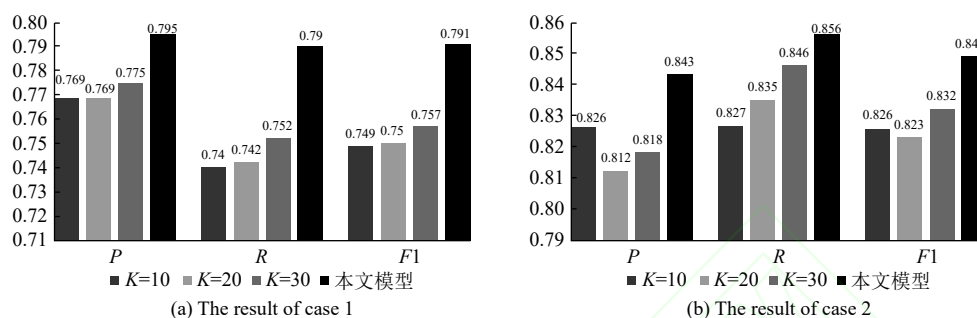


图 6 一次重构模型与完整模型的分类结果比较

从图 6 可以看出, 对于只进行一次重构的模型, 在案件 1 数据集上, 随着主题数增加, 分类结果的 $F1$ 值略有提升; 而在案件 2 数据集上, 主题数为 20 时精确率 P 值较低, 使得 $F1$ 值也低于主题数 10 的结果. 本文完整模型相比只进行一次重构的模型, 在 3 个评价指标上都有明显提升, 其中在案件 1 数据集上 $F1$ 值提升了 0.04 左右, 在案件 2 数据集上 $F1$ 值提升了 0.02 左右. 以上结果说明通过辅助主题重构学习到的主题分布对于评价对象分类有较好的作用. 本文模型中, 辅助主题数目设置为核心主题数目的倍数, 因此, 核心主题向量对应于向量空间中相对较大的聚类簇 (主题数目少), 而辅助主题向量对应为小的聚类簇 (主题数目大). 这样, 同时利用辅助主题和核心主题可以缓解 ABAE 等主题模型中主题数目固定带来的不同大小聚类簇特性学习不充分的问题.

3.4.6 辅助主题挖掘的评价对象词项

上述实验表明, 通过辅助主题表征矩阵, 提高了模型的分类性能. 实际上, 通过辅助主题矩阵, 还可以获得额外的评价对象词项. 表 7 列出了案件 2 的辅助主题向量推断的 2 组评价对象词项, 分别为公交车坠江案中的案件刑罚和当事人. 这两组评价对象词项是其他主题模型都没挖掘到的词项, 对于全面掌握案件的评价对象而言很有价值.

表 7 使用辅助主题向量抽取的评价对象词项

top 10 的评价对象词项	评价对象
刑法判 重罚 有期徒刑 触犯 十年 刑事责任 行政拘留 无期徒刑 严重后果	案件刑罚
乘客 几秒钟 肢体 还击 争吵 劝阻 动手 争执 肢体冲突 上前	乘客 (案件当事人)

4 结 论

针对目前主流的基于深度主题表征的评价对象识别方法需要预设固定的主题数目, 且最终评价对象识别依赖人工推断问题, 本文提出了一种基于变分双主题表征的弱监督评价对象识别模型, 结合了两个不同的主题表征来重构句子表示, 同时基于少量标签样本的类别信息, 能较好的将评论句自动分类为评价对象类别, 挖掘评价对象词项. 相比其他无监督主题模型, 本文模型通过有效利用少量有标签样本的类别信息, 能使模型准确预测评价对象类别; 相比需要挑选种子词的弱监督主题模型, 本文模型标注句子评价对象类别的方式更容易实现, 分类性能更好. 同时, 所提出的两次变分编码和重构, 能使模型学习到更合理的主题表征, 从而提高分类性能; 通过核心主题表征

得到兼具较好主题连贯性和多样性的评价对象代表词项, 并通过辅助主题表征得到额外的评价对象词项. 未来, 计划在更多的数据集中实践所提出的模型.

References:

- [1] Blei DM, Ng AY, Jordan MI. Latent dirichlet allocation. *Journal of Machine Learning Research*, 2003, 3: 993–1022. [doi: [10.5555/944919.944937](https://doi.org/10.5555/944919.944937)]
- [2] Fei G, Chen ZY, Liu B. Review topic discovery with phrases using the pólya urn model. In: *Proc. of the 25th Int'l Conf. on Computational Linguistics: Technical Papers*. Dublin: ACL, 2014. 667–676.
- [3] Mukherjee A, Liu B. Aspect extraction through semi-supervised modeling. In: *Proc. of the 50th Annual Meeting of the Association for Computational Linguistics*. Jeju Island: ACL, 2012. 339–348.
- [4] Zhao X, Jiang J, Yan HF, Li XM. Jointly modeling aspects and opinions with a MaxEnt-LDA hybrid. In: *Proc. of the 2010 Conf. on Empirical Methods in Natural Language Processing*. Cambridge: ACL, 2010. 56–65.
- [5] Miao YS, Grefenstette E, Blunsom P. Discovering discrete latent topics with neural variational inference. In: *Proc. of the 34th Int'l Conf. on Machine Learning*. Sydney: PMLR, 2017. 2410–2419.
- [6] Srivastava A, Sutton C. Autoencoding variational inference for topic models. In: *Proc. of the 5th Int'l Conf. on Learning Representations*. Toulon: OpenReview.net, 2017.
- [7] He RD, Lee WS, Ng HT, Dahlmeier D. An unsupervised neural attention model for aspect extraction. In: *Proc. of the 55th Annual Meeting of the Association for Computational Linguistics*. Vancouver: ACL, 2017. 388–397. [doi: [10.18653/v1/P17-1036](https://doi.org/10.18653/v1/P17-1036)]
- [8] Hu MQ, Liu B. Mining and summarizing customer reviews. In: *Proc. of the 10th ACM SIGKDD Int'l Conf. on Knowledge Discovery and Data Mining*. Seattle: ACM, 2004. 168–177. [doi: [10.1145/1014052.1014073](https://doi.org/10.1145/1014052.1014073)]
- [9] Zhuang L, Jing F, Zhu XY. Movie review mining and summarization. In: *Proc. of the 15th ACM Int'l Conf. on Information and Knowledge Management*. Arlington: ACM, 2006. 43–50. [doi: [10.1145/1183614.1183625](https://doi.org/10.1145/1183614.1183625)]
- [10] Liu Q, Gao ZQ, Liu B, Zhang YL. Automated rule selection for aspect extraction in opinion mining. In: *Proc. of the 24th Int'l Conf. on Artificial Intelligence*. Buenos Aires: AAAI Press, 2015. 1291–1297. [doi: [10.5555/2832415.2832429](https://doi.org/10.5555/2832415.2832429)]
- [11] Wang WY, Pan SJ. Recursive neural structural correspondence network for cross-domain aspect and opinion co-extraction. In: *Proc. of the 56th Annual Meeting of the Association for Computational Linguistics*. Melbourne: ACL, 2018. 2171–2181. [doi: [10.18653/v1/P18-1202](https://doi.org/10.18653/v1/P18-1202)]
- [12] Wang WY, Pan SJ, Dahlmeier D, Xiao XK. Recursive neural conditional random fields for aspect-based sentiment analysis. In: *Proc. of the 2016 Conf. on Empirical Methods in Natural Language Processing*. Austin: ACL, 2016. 616–626. [doi: [10.18653/v1/d16-1059](https://doi.org/10.18653/v1/d16-1059)]
- [13] Chernyshevich M. IHS R&D Belarus: Cross-domain extraction of product features using conditional random fields. In: *Proc. of the 8th Int'l Workshop on Semantic Evaluation*. Dublin: The Association for Computer Linguistic, 2014. 309–313.
- [14] Xu H, Liu B, Shu L, Philip SY. Double embeddings and CNN-based sequence labeling for aspect extraction. In: *Proc. of the 56th Annual Meeting of the Association for Computational Linguistics*. Melbourne: ACL, 2018: 592–598. [doi: [10.18653/v1/P18-2094](https://doi.org/10.18653/v1/P18-2094)]
- [15] Li X, Lam W. Deep multi-task learning for aspect term extraction with memory interaction. In: *Proc. of the 2017 Conf. on Empirical Methods in Natural Language Processing*. Copenhagen: ACL, 2017. 2886–2892. [doi: [10.18653/v1/d17-1310](https://doi.org/10.18653/v1/d17-1310)]
- [16] Qiu G, Liu B, Bu JJ, Chen C. Opinion word expansion and target extraction through double propagation. *Computational Linguistics*, 2011, 37(1): 9–27. [doi: [10.1162/coli_a_00034](https://doi.org/10.1162/coli_a_00034)]
- [17] Das A, Kannan A. Discovering topical aspects in microblogs. In: *Proc. of the 25th Int'l Conf. on Computational Linguistics: Technical Papers*. Dublin: ACL, 2014. 860–871.
- [18] Salakhutdinov R, Hinton G. Replicated softmax: An undirected topic model. In: *Proc of the 22nd Int'l Conf. on Neural Information Processing Systems*. Vancouver: Curran Associates, Inc. , 2009. 1607–1614. [doi: [10.5555/2984093.2984273](https://doi.org/10.5555/2984093.2984273)]
- [19] Nguyen DQ, Billingsley R, Du L, Johnson M. Improving topic models with latent feature word representations. *Trans. of the Association for Computational Linguistics*, 2015, 3: 299–313. [doi: [10.1162/tac1_a_00140](https://doi.org/10.1162/tac1_a_00140)]
- [20] Li CL, Duan Y, Wang HR, Zhang ZQ, Sun AX, Ma ZY. Enhancing topic modeling for short texts with auxiliary word embeddings. *ACM Trans. on Information Systems*, 2018, 36(2): 11. [doi: [10.1145/3091108](https://doi.org/10.1145/3091108)]
- [21] Shi B, Lam W, Jameel S, Schockaert S, Lai KP. Jointly learning word embeddings and latent topics. In: *Proc. of the 40th Int'l ACM SIGIR Conf. on Research and Development in Information Retrieval*. Shinjuku: ACM, 2017. 375–384. [doi: [10.1145/3077136.3080806](https://doi.org/10.1145/3077136.3080806)]
- [22] Das R, Zaheer M, Dyer C. Gaussian LDA for topic models with word embeddings. In: *Proc. of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th Int'l Joint Conf. on Natural Language Processing*. Beijing: ACL, 2015. 795–804.

- [doi: [10.3115/v1/p15-1077](https://doi.org/10.3115/v1/p15-1077)]
- [23] Gupta P, Chaudhary Y, Buettner F, Schütze H. Document informed neural autoregressive topic models with distributional prior. In: Proc. of the 33rd AAAI Conf. on Artificial Intelligence. Honolulu: AAAI Press, 2019. 6505–6512. [doi: [10.1609/aaai.v33i01.33016505](https://doi.org/10.1609/aaai.v33i01.33016505)]
- [24] Yan XH, Guo JF, Lan YY, Cheng XQ. A biterm topic model for short texts. In: Proc. of the 22nd Int'l Conf. on World Wide Web. Rio de Janeiro: ACM, 2013. 1445–1456. [doi: [10.1145/2488388.2488514](https://doi.org/10.1145/2488388.2488514)]
- [25] Cheng XQ, Yan XH, Lan YY, Guo JF. BTM: Topic modeling over short texts. IEEE Trans. on Knowledge and Data Engineering, 2014, 26(12): 2928–2941. [doi: [10.1109/TKDE.2014.2313872](https://doi.org/10.1109/TKDE.2014.2313872)]
- [26] Zhu Q, Feng Z, Li XL. GraphBTM: Graph enhanced autoencoded variational inference for biterm topic model. In: Proc. of the 2018 Conf. on Empirical Methods in Natural Language Processing. Brussels: ACL, 2018. 4663–4672. [doi: [10.18653/v1/D18-1495](https://doi.org/10.18653/v1/D18-1495)]
- [27] Liao M, Li J, Zhang HS, Wang LZ, Wu XX, Wong KF. Coupling global and local context for unsupervised aspect extraction. In: Proc. of the 2019 Conf. on Empirical Methods in Natural Language Processing and the 9th Int'l Joint Conf. on Natural Language Processing. Hong Kong: ACL, 2019. 4579–4589. [doi: [10.18653/v1/D19-1465](https://doi.org/10.18653/v1/D19-1465)]
- [28] Peng M, Yang SX, Zhu JH. Semantic enhanced topic modeling by bi-directional LSTM. Journal of Chinese Information Processing, 2018, 32(4): 40–49 (in Chinese with English abstract). [doi: [10.3969/j.issn.1003-0077.2018.04.005](https://doi.org/10.3969/j.issn.1003-0077.2018.04.005)]
- [29] Dieng AB, Ruiz FJR, Blei DM. Topic modeling in embedding spaces. Trans. of the Association for Computational Linguistics, 2020, 8: 439–453. [doi: [10.1162/tac1_a_00325](https://doi.org/10.1162/tac1_a_00325)]
- [30] Lu B, Ott M, Cardie C, Tsou BK. Multi-aspect sentiment analysis with topic models. In: Proc. of the 11th IEEE Int'l Conf. on Data Mining Workshops. Vancouver: IEEE, 2011. 81–88. [doi: [10.1109/ICDMW.2011.125](https://doi.org/10.1109/ICDMW.2011.125)]
- [31] Lund J, Cook C, Seppi K, Boyd-Graber J. Tandem anchoring: A multiword anchor approach for interactive topic modeling. In: Proc. of the 55th Annual Meeting of the Association for Computational Linguistics. Vancouver: ACL, 2017: 896–905. [doi: [10.18653/v1/P17-1083](https://doi.org/10.18653/v1/P17-1083)]
- [32] Mikolov T, Yih WT, Zweig G. Linguistic regularities in continuous space word representations. In: Proc. of the 2013 Conf. of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies. Atlanta: ACL, 2013. 746–751.
- [33] Mimno D, Wallach HM, Talley E, Leenders M, McCallum A. Optimizing semantic coherence in topic models. In: Proc. of the Conf. on Empirical Methods in Natural Language Processing. Edinburgh: ACL, 2011. 262–272. [doi: [10.5555/2145432.2145462](https://doi.org/10.5555/2145432.2145462)]

附中文参考文献:

- [28] 彭敏, 杨绍雄, 朱佳晖. 基于双向LSTM语义强化的主题建模. 中文信息学报, 2018, 32(4): 40–49. [doi: [10.3969/j.issn.1003-0077.2018.04.005](https://doi.org/10.3969/j.issn.1003-0077.2018.04.005)]



相艳(1979—), 女, 博士, 副教授, 主要研究领域为自然语言处理, 情感计算.



黄于欣(1983—), 男, 博士, 副教授, 主要研究领域为自然语言处理, 文本摘要.



余正涛(1970—), 男, 博士, 教授, CCF 高级会员, 主要研究领域为自然语言处理, 神经机器翻译, 信息检索.



线岩团(1981—), 男, 博士, 副教授, 主要研究领域为自然语言处理, 信息检索.



郭军军(1987—), 男, 博士, 副教授, CCF 专业会员, 主要研究领域为自然语言处理, 神经机器翻译, 多模态情感分析.