

基于话题模型的科技文献话题发现和趋势分析*

贺亮, 李芳

(上海交通大学 计算机科学与工程系, 上海 200240)

摘要: 自动挖掘科技文献话题, 总结发展趋势及最新研究动态, 有助于科技工作者的研究。本文提出一种话题发现和趋势分析的方法, 该方法首先利用 LDA 话题模型抽取科技文献的话题, 然后计算话题的强度和影响力, 最后针对热门和冷门话题以及影响力高和影响力低的话题, 进行了趋势分析。本文提出的话题强度和影响力计算方法, 可以针对任何文集。对 ACL 论文集的实验, 显示了计算语言学领域过去的发展状况。和其他方法的对比实验, 也验证了本文提出的话题强度和影响力的计算方法是正确和可行的。

关键词: 话题模型; 趋势分析

中图分类号: TP391

文献标识码: A

Topic Discovery and Trend Analysis in Scientific Literature Based on Topic Model

HE Liang, LI Fang

(Dept. of Computer Science & Engineering, Shanghai Jiao Tong University, Shanghai 200240, China)

Abstract: Automatically extracting topics from scientific literature and finding the research trends will help researchers a lot. In this paper, we use LDA model to generate topics from the scientific literature; then calculate the strength and impact of the topic; finally, find the trends of the hot topics vs. cold topics, high vs. low impact topics. The method of calculating topic strength and impact is suitable for any document. The experiments on ACL anthology have shown the research trend in computational linguistics. And the contrast experiment also proved the proposed calculating method.

Keywords: Topic Model, Trend Analysis

1 引言

在这个信息爆炸的时代, 科学技术的发展也日新月异, 对于科技工作者来说, 需要快速的获取相关领域的最新研究动态。为了了解最新的研究工作, 科技工作者会关注该领域的关键问题, 这些问题都用到了什么样的技术, 在众多的技术中, 哪些是目前的研究热点, 哪些逐渐被人们淡忘。因此, 对于科学技术趋势的自动分析研究, 旨在帮助科学工作者从大量的学术会议和科技文献中提取出有用的信息, 具有重要的现实意义。

要进行趋势分析, 首先需要从大量的语料集中提取出潜在的语义信息, 亦可称之为话题。传统的 VSM 模型使用关键字来表示话题。但这种表达方式比较局限于对文档贡献较大的词, 很多用于表示文档的词语, 由于存在二义性, 对于文档的语义上的描述, 效果往往差强人意。为了克服 VSM 模型的这些缺点, 有学者提出了语义模型[1,2]。首先是 LSI 模型[1], 可以利用 SVD 技术对文本降维; 进一步, 在 LSI 模型中引入概率模型, 得到 pLSI 模型[2], 该模型是生成模型, 它假设每篇文档是由多项式随机变量(话题)混合而成, 而文档中每个词, 由一个话题产生, 文档中不同的词可有不同的话题生成。但是 pLSI 模型参数数量随着文集增长而线性增长, 并且对于没有观测到的文本没有很好的预测。Blei[3]提出的 LDA 模型可以挖掘大规模语料的语义信息, 是

* 收稿日期: 定稿日期:

基金项目: 国家自然科学基金项目 (60873134)

作者简介: 贺亮(1987—), 男, 硕士, 主要研究领域为自然语言处理, 信息检索与信息抽取; 李芳(1963—), 女, 博士, 副教授, 主要研究领域为自然语言处理, 信息检索与信息抽取。

机器学习、信息检索等领域很流行的一个模型。LDA 模型继承了 pLSI 模型的所有优点，可以很好的产生话题的分布；同时，LDA 模型的参数数量也不会随着文集增长而线性增长，有很好的泛化能力。因此，本文将采用基于 LDA 模型的方法对科技文献进行分析研究。

本文的主要贡献包括两个方面：

- 利用话题模型即 LDA 模型对语料建模，挖掘出该领域中的研究热点及相关技术，提出一个针对话题的热门程度和影响力衡量标准；
- 基于话题的强度，研究这些子领域以及技术在整个时间段上的趋势变化。

本文的组织结构如下：第二章介绍相关的工作，在第三章给出了我们的研究方法，第四章是实验结果和分析，第五章为结论及展望。

2 相关工作

目前对于科技文献的研究，主要利用了科技文献的作者、文本信息、引用信息和时间信息，去进行话题的发现和趋势的分析工作。

首先是发现话题 (*topic*)，即是挖掘文献中的隐含的语义信息。目前主要有两类方法可以发掘话题，第一类利用话题模型进行话题发现，这里话题的定义是一组词的概率分布。根据文集的文本信息可以利用 LDA 以及其拓展模型 (CTM、DTM 等) 进行建模[4,5]，发现话题；如果结合作者信息，有作者话题模型 (ATM) 以及其拓展模型 (ACT、TATM 等) [6-8]，通过对该模型的推导可以得到每个作者在话题空间上的分布，通过分析该分布就可以了解在某一特定领域 (话题) 都有哪些专家，以及这些专家关注的研究领域 (话题) 是什么；结合文献引用信息，既考虑到了文献间引用关系对生成过程中的影响，有继承话题模型 (ITM) [9]。第二类方法则通过构造网络图，利用文献的文本信息以及文献间的引用信息进行话题发现。有学者使用词组 (*term*) 来表示话题，然后利用词组 (*term*) 在文集分布关系并结合文集之间的引用关系发现话题[10]。

从文集中发掘出话题信息后，就可以在这话题空间上进一步分析这些话题的特点。有学者利用 LDA 对文集建模得到的话题空间，再加入文献之间引用的信息，去研究话题的特性。这些特性有话题的影响因子，用于衡量话题对文档的影响；有话题的影响多样性，衡量话题的影响范围；有话题的年龄，衡量话题的新旧程度；还有话题的转移度，衡量话题之间相互的影响[11]。

更进一步，加入时间的信息，进行话题的趋势分析。有学者利用话题的后验概率去定义话题的强度，通过计算每个时间点上的强度得到其强度的趋势变化[12,13]，对这些话题的趋势变化进行分析，以获得科技发展的一些特点，例如一些技术的应用走向，是偏向理论性的研究还是偏向于实际应用等等[13]。斯坦福大学的一个开源话题建模工具¹ (*tmt*) 也是基于这种方法进行分析，通过简单地统计不同时间段词频能得到话题内容随时间的变化。有学者使用分时间段进行话题建模，考虑各个时间段话题之间关联的方法，可以从内容上去分析话题的变化趋势[5,9,14]。有学者在作者话题模型的基础上，加入时间信息，利用话题与作者间对应关系，从而可以分析这些作者的研究兴趣如何随时间推移而变化[8]。

为了提出一种方法能够针对任何文集，例如新闻报道[14]，数字文献等，我们只考虑文献的时间和文本信息，忽略作者和引用信息。采用 LDA 话题模型，找到潜在话题，借鉴文献[11-15]对话题的强度和影响力这两个特性进行研究，提出了不同的计算公式，通过这两个特性的分析可以找到热点话题和有影响力的话题，然后根据话题的强度再对它们进行趋势分析。

3 研究方法

首先对文本集合应用 LDA 建模，抽取文章的话题，然后，定量分析话题的强度和影响力，提供一套可靠有效的评价标准，最后对热点话题和有影响力话题进行趋势分析。话题强度主要描述了话题的受关注度，比如说，讨论某话题的文章数越多，就说明该话题的强度越高，可以认为

¹ <http://www-nlp.stanford.edu/software/tmt/tmt-0.2/>

是热门话题。话题的影响力则是指当前话题对其他话题的影响力，如果一个话题对多个话题都有一定程度的影响，该话题可以认为是具有影响力的话题。

3.1 话题建模

首先，表 1 列出了本文使用的符号。

表 1 文中使用到的符号

符号	符号的描述
α	LDA 模型的 Dirichlet 先验参数，表示文档-话题分布的先验
β	LDA 模型的 Dirichlet 先验参数，表示话题-词分布的先验
θ_d	文档 d 的话题多项式分布
φ_z	话题 z 的词表多项式分布
K	话题个数
D^t	t 时间内所有的文档集合
D_z^t	t 时间内话题 z 的支持文档集合
$S(z, t)$	话题 z 在 t 时间段的文档支持率
$I(z)$	话题 z 的影响力

LDA 模型是一个生成概率模型，是三层的变参数层次贝叶斯模型，首先假设词由话题的概率分布混合产生，而每个话题是在词汇表上的一个多项式分布；其次假设文档是潜在话题的概率分布的混合；最后针对每个文档从 Dirichlet 分布中抽样产生该文档包含的话题比例，结合话题和词的概率分布生成该文档中的每一个词汇。该模型描述文档的生成过程，有以下步骤：

- 1) 对于每个文档 d，根据 $\theta_d \sim \text{Dir}(\alpha)$ ，得到多项式分布参数 θ_d
- 2) 对于每个话题 topic z，根据 $\varphi_z \sim \text{Dir}(\beta)$ ，得到多项式分布参数 φ_z
- 3) 对文档 d 中的第 i 个词 w_i :
 - a) 根据多项式分布 $z \sim \text{Mult}(\theta_d)$ ，得到话题 z
 - b) 根据多项式分布 $w_i \sim \text{Mult}(\varphi_z)$ ，得到词 w_i

3.2 话题强度计算

在一段时期内，如果文集中大多数文档都是关于某一个话题的，那么该话题是热门的。谈及该话题的文档数越多，就说明话题越热门。一般地，话题的热门程度通常使用话题强度进行量化。话题强度描述了一个话题的受关注程度，本文使用文档支持率作为话题强度的表示，具体定义如下：

根据 LDA 话题抽取的结果，我们知道一个文档上话题的分布并不均匀，也就是说文档对于每个话题的贡献度不同。也就是说，针对一个话题，有的文档属于重要文档，有的文档对于该话题并不是很重要。综上，我们定义话题的支持文档如下：假设某一文档 d 中有至少 10% 的词是由话题 z 生成的，那么该文档是话题 z 的支持文档。根据该定义，一篇文档可以支持多个话题。

话题 z 在时间间隔 t 的文档支持率计算公式如下：

$$S(z, t) = \frac{|D_z^t|}{|D^t|} \quad (1)$$

其中， $|D_z^t|$ 表示 t 时间段话题 z 的支持文档个数， $|D^t|$ 表示该时间段文档的总数。

3.3 话题影响力计算

话题的影响力使用其影响的多样性 (Impact Diversity) 来衡量。我们基于这样的假设，一个话题在产生之后，可能会对之后的时间段的话题有影响，这种影响将通过文档之间的关联来体现，如果前一时间段 t 话题 z 的支持文档 d 与后一时间段 t' 话题 z' 的支持文档 d' 是关联的，那么可以

认为话题 z 对话题 z' 有一定的影响作用。

计算影响力时需要统计属于不同话题的文章之间的关联数量。每篇文章可表示为在话题空间上的分布 $\theta(n_1, n_2, \dots, n_k)$, n_k 表示话题 k 出现在该文档中的概率, 通过计算话题空间上分布的 JS 距离 (Jensen-Shannon divergence) 来判断文章之间是否关联。假设时间段 t 话题 z 的支持文档 d 与后一时间段 t' 话题 z' 的支持文档 d' , 在话题空间中的分布分别为 θ_d 和 $\theta_{d'}$, 则它们的 JS 距离计算公式如下:

$$D_{JS}(\theta_d || \theta_{d'}) = \frac{1}{2} (D_{KL}(\theta_d || \theta_m) + D_{KL}(\theta_{d'} || \theta_m)) \quad (2)$$

$$\text{其中 } \theta_m = \frac{1}{2} (\theta_d + \theta_{d'})$$

话题之间的影响作用可以使用这些话题的支持文档关联数量来计量, 我们定义话题 z 对话题 z' 的影响程度为话题 z 对 z' 的影响作用占所有话题对 z' 影响作用的比例, 提出一个计算话题影响程度的公式如下:

$$P_z(z') = \frac{\sum_t \sum_{t' > t} \# \text{relations between } D_{z'}^{t'} \text{ and } D_z^t}{\sum_t \sum_{t' > t} \# \text{relations between } D^{t'} \text{ and } D_z^t} \quad (3)$$

其中, 分子表示话题 z 的支持文档与后续所有时间段的话题 z' 的支持文档关联数量, 分母表示话题 z 的支持文档与后续所有时间段的文档关联数量。

为了计量一个话题对其他所有话题的影响程度, 我们定义话题 z 的影响力为话题 z 对所有话题的影响程度的熵, 计算公式如下:

$$I(z) = H(P_z) = - \sum_{z'} P_z(z') \log P_z(z') \quad (4)$$

通过该公式计算出一个话题的影响力越大, 说明它的影响范围越广; 反之则说明它的影响范围较狭隘。

4 实验结果与分析

实验主要包括三个方面, 一是热门话题的实验, 采用文献[14]提出的系统作为对比; 二是研究话题的影响力, 采用文献[7]提出的方法作为对比; 三是研究它们随时间变化的趋势, 采用斯坦福大学的 TMT 分析工具作为对比。

4.1 实验数据

ACL 论文集 (ACL Anthology)² 作为实验的数据集, 它包括 1985 年至 2009 年的 ACL、COLING、EACL、EMNLP 等众多会议, 总共 11072 篇文章。以上语料只取标题和摘要, 并过滤停用词、低频词等。本实验利用 Gibbs Sampling 方法进行参数的推理。实验使用了开源的 Gibbs Sampling 工具³, 模型参数 α , β 分别设置为 50/K 和 0.01, 话题个数 K 设为 100。

4.2 热门话题分析

通过公式 (1) 计算话题每年的强度, 比较话题的强度, 可以发现每年的热门话题。表 2 展示了 2006 年至 2009 年每年最热门的五个话题, 话题名称均为人工标签。

表 2. 2006 年至 2009 年热门话题

2006		2007		2008		2009	
话题	强度	话题	强度	话题	强度	话题	强度
Dependency Parsing	0.063	Stat. MT	0.110	Stat. MT	0.111	Stat. MT	0.121
Stat. MT	0.062	Dependency Parsing	0.075	Sentiment	0.072	Dependency Parsing	0.062

² <http://www.aclweb.org/anthology/>

³ <http://gibbslda.sourceforge.net/>

Q&A System	0.055	WSD	0.061	Dependency Parsing	0.055	Sentiment	0.056
Named entity	0.054	Sentiment	0.048	Named entity	0.051	CRF	0.051
Info. Retrieval	0.050	Named entity	0.046	Summarization	0.046	Semantic Role	0.046

从表 2 可以看到，基于统计的机器翻译（Stat. MT）是近几年来最热门的话题。众所周知，自从统计技术在机器翻译领域取得成效后，人们对其的研究热情一直未减。统计技术也同样应用于计算语言学的其他方面，如依存关系句法分析（Dependency Parsing），热门程度仅次于基于统计的机器翻译。值得一提的还有情感分析（Sentiment）在近年的研究热度迅速提升。

为了对以上结果进行验证，我们选择使用文献[16]提出的一种基于句法分析（Parsing）和语义元组提取（Semantic Tuple Extraction）方法专门针对 ACL 论文集进行分析的 Searchbench 系统。将我们得到的话题在该系统中查询，得到每年的文章数量，除以当年文章总量得到话题权重，通过比较话题权重可以得到每年的热门话题，与我们的结果进行对比。表 3 展示了 ACL-Searchbench 系统得到的结果。

表 3. 2006 年至 2009 年热门话题（ACL-Searchbench）

2006		2007		2008		2009	
话题	权重	话题	权重	话题	权重	话题	权重
Stat. MT	0.039	Stat. MT	0.061	Stat. MT	0.046	Stat. MT	0.042
Named entity	0.030	Info. Retrieval	0.044	Semantic Role	0.030	Semantic Role	0.037
Semantic Role	0.028	WSD	0.035	Summarization	0.019	Info. Retrieval	0.034
Sentiment	0.016	Dependency	0.028	Sentiment	0.018	Dependency	0.026
Dependency	0.015	Clustering	0.022	Dependency Parsing	0.017	Sentiment	0.025

通过与 Searchbench 系统得到的结果对比，可以看到找到的热门话题大体上是一致的，只存在少量的话题或者是位置排名的差异，这说明了我们的方法是有效的。

4.3 话题影响力分析

文献[7]提出一种话题影响力的计算方法，它利用文档之间引用关系计算话题间影响概率，再计算这些影响概率的熵值，作为话题影响力。该方法作为 Baseline 与我们的方法进行对比。

首先使用公式（2）计算文档之间的关联度，阈值定为 0.07，然后，利用公式（3）（4）计算话题的影响力。表 4 分别列出了影响力前五和后五的话题。

表 4 话题影响力得分情况

我们的方法		Baseline-Impact	
话题	影响力	话题	影响力
Kernel Method	5.17	SVM	7.00
Probabilistic Model	5.14	Tagging	6.99
SVM	5.09	Probabilistic Model	6.74
Sentiment	5.08	Sentiment	6.73
Ngram Model	4.99	Kernel Method	6.70
Stat. MT	4.26	WSD	5.94
Spell Correction	4.24	Stat. MT	5.86
Speech Recognition	4.12	Summarization	5.68
WSD	3.40	Word Segmentation	5.56
Word Segmentation	3.05	Morphology	5.48

结果显示了影响力高的话题都是一些使用比较广泛的技术，例如核方法（Kernel Method）、支持向量机（SVM）等在数据挖掘、机器学习领域很流行的分类技术，它们在计算语言学领域也

发挥着很大的作用。而影响力较小的话题都是一些偏应用方面的领域，比如说机器翻译、词义消歧（WSD）以及分词（Word Segmentation）等，这些领域的特点是比较专一，影响面比较窄。

实验结果与 Baseline-Impact 方法的结果大体一致，虽然我们的方法计算量比 Baseline-Impact 大，但是不需要额外的文档之间相互引用的信息，可以应用于任何文档集合。

4.4 话题趋势分析

本小节的实验是利用话题逐年的强度变化来分析话题的变化趋势，这些话题包括热门话题，冷门话题，影响力大的以及影响力小的话题，以此了解计算语言学领域近二十多年发展情况。我们使用斯坦福大学提供的一个开源话题建模工具（TMT）作为 baseline 方法对 ACL 文集进行建模分析，与我们的方法得到的实验结果进行对比以及验证。

首先来看最近几年的热门话题的强度变化趋势。从图 1 可以看出基于统计方法的机器翻译技术作为最热门的话题从 1999 年开始，进入了一个飞跃上升的阶段。出现这个变化的原因，就是在 1999 年出现了一个机器翻译的热潮，其最主要的特征是基于统计的方法在这一领域开始占据主导地位，机器翻译的质量出现了一个跨越式提高。这股热潮持续至今，仍未现衰减之势。同时，基于统计的句法分析的强度也随着这股热潮不断提升。而情感分析在 2000 年前一直都是比较冷门的话题，但现今研究者对它的青睐不断增加。

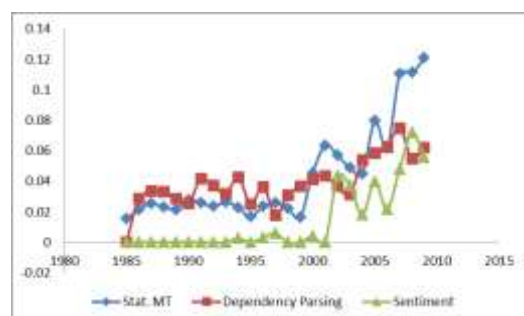


图 1 热门话题强度变化趋势

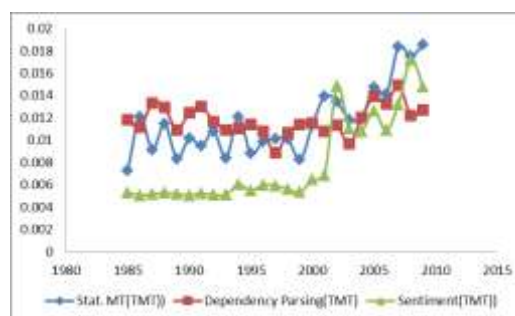


图 2 热门话题强度变化趋势(baseline-TMT)

根据实验结果，图 3 列出了一些冷门技术的变化趋势，包括语言识别（Speech Recognition）和联并方法（Unification）。

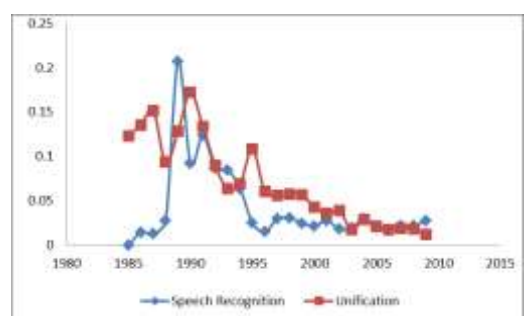


图 3 冷门话题强度变化趋势

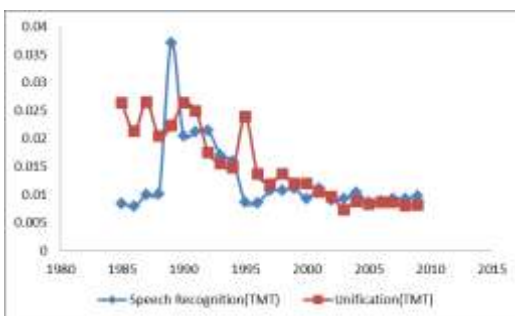


图 4 冷门话题强度变化趋势(baseline-TMT)

联并方法是八十年代末九十年代初的研究热点，其后渐渐地淡出了研究者的视线。而语音识别技术的变化趋势比较奇特，它在 1989 年至 1994 年有一个爆发式的高峰。究其缘由，是因为这几年举办的 DARPA 语音及自然语言研讨会（DARPA Speech and Natural Language Workshop），这些研讨会产生了大量这方面技术的研究论文，而之后该技术的研究就进入低谷。

通过对热门话题和冷门话题的趋势分析，可以看到统计技术的兴起对这些热门话题的强度上升起了很大的推动作用；另一方面，冷门话题的下降趋势也有不同的表现形式，有的是缓慢下降，有的是急速下降。

接下来看影响力比较高的话题变化趋势情况，见图 5。

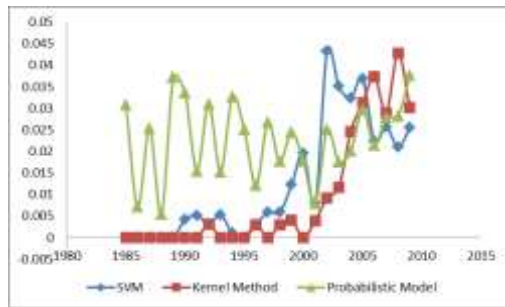


图5 影响力高的话题强度变化趋势

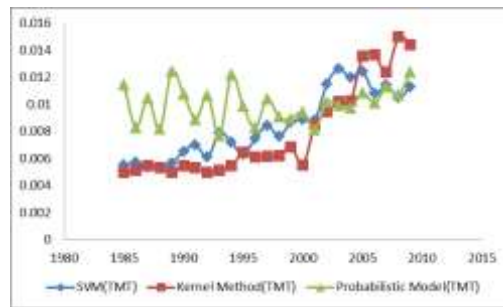


图6 影响力高的话题强度变化趋势 (baseline-TMT)

这几个话题都是一些流行的技术，首先是概率模型 (Probabilistic Model)，它在计算语言学的领域一直都是比较主流的技术，它的强度变动在 2000 年前呈波动形式，之后呈上升趋势。而支持向量机和核方法在九十年代末开始兴起，此后也越来越受到研究者重视，保持着上升的形式，成为了计算语言学领域中比较重要的分析方法。

而影响力较低的话题比较偏应用，趋势变化没有固定的特点，从图 7 可以看到，有的呈现上升趋势，例如基于统计的机器翻译；有的呈现下降趋势，例如语音识别。

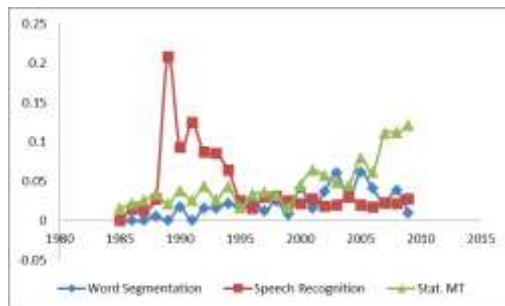


图7 影响力低的话题强度变化趋势

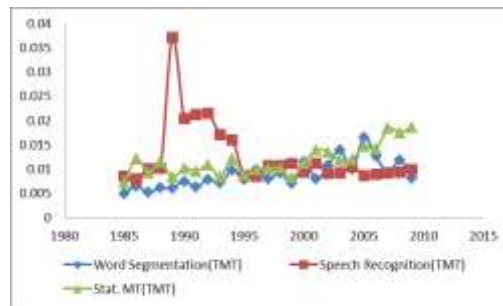


图8 影响力低的话题强度变化趋势 (baseline-TMT)

通过对影响力大的和影响力小的话题进行趋势分析，可以发现它们的强度变化趋势与影响力大小是无关的，这也说明了话题强度和话题影响力这两个指标是相互独立的两个标准，可以从不同方面去描述话题的特性。

通过与 baseline-TMT 方法得到结果进行对比，我们发现这些话题的趋势跟我们的方法得到的趋势大体一致，且在变化方向上是一样的。这也验证了我们方法的正确性和有效性。而在某些话题例如情感分析话题 (sentiment)，我们的方法得到它在 2000 年前的大部分年份强度都为零，说明当时还没产生这个话题，与实际情况相符。这也说明了我们的方法更为精确。

5 结论与展望

本文利用话题模型对科技文献进行建模分析。首先使用 LDA 话题建模，发现文集中隐含的话题。接着，使用两个指标——话题强度和话题影响力去研究话题的特性。同时，对这些研究领域或技术受关注程度随时间变化的趋势进行分析，发现它们的变化特点。

通过分析实验结果，可以发现利用话题模型能够从大量文献中发掘出有意义的信息。实验结果与实际情况相符合，说明我们的方法对科技文献的分析是行之有效的。以下是对 ACL 论文集分析研究得到的一些结论：

- 1) 最近比较热门的研究领域包括机器翻译、句法分析以及情感分析等；
- 2) 理论型的技术（例如核方法、概率模型）往往有较大的影响范围，可能会应用到多个子领域，而应用型的研究领域（例如机器翻译）的影响范围比较窄；
- 3) 通过趋势分析，可以了解计算语言学近二十多年来的发展情况，包括统计技术的流行大大促进了机器翻译和句法分析的研究，语音识别技术的研究热潮兴起与回落，联并语法

研究的逐步衰落等等。

今后的工作将考虑如何进一步挖掘话题的特点，更好地探索话题之间的关联。另外，从更多的角度去分析话题的变化趋势，比如从内容上分析话题在各个时间段的特点。

参考文献:

- [1] S.Deerwester, S.Dumais, T.Landauer, G.Fumas, and R.Harshman. Indexing by Latent Semantic Analysis[C], Journal of the American Society of Information Science, 41(6):391-407, 1990
- [2] T.Hofmann. Probabilistic Latent Semantic Indexing[C]. Proceedings of the Twenty-Second Annual International SIGIR Conference, 1999.
- [3] D.M.Blei, A.Y.Ng, and M.I.Jordan. Latent Dirichlet Allocation[C]. The Journal of Machine Learning Research, 2003, vol.3, pp.993-1022.
- [4] D.M.Blei, J.D.Lafferty. A Correlated Topic Model of Science[C]. The Annals of Applied Statistics 2007, Vol.1, No.1, 17-35.
- [5] D.M.Blei and J.D.Lafferty. Dynamic Topic Model[C]. In International conference on Machine Learning, 2006, pp.113-120.
- [6] M. Rosen-Zvi, T. Griffiths, M. Steyvers, and P. Smyth. The Author-Topic Model for Authors and Documents[C]. In Proceedings of the 20th Conference on Uncertainty in Artificial Intelligence, 2004.
- [7] Jie Tang, Jing Zhang, Limin Yao, Juanzi Li, Li Zhang, and Zhong Su. AmetMiner: Extraction and Mining of Academic Social Networks[C]. Proceedings of the Fourteenth ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (SIGKDD).2008, pp.990-998.
- [8] A.Daud, Juanzi Li, Lizhu Zhou, and F.Muhammad, Exploiting Temporal Authors Interests via Temporal-Author-Topic Modeling[C]. Proceedings of ADMA, 2009, LNAI 5687, pp.435-443.
- [9] Q.He, B.Chen, J.Pei, B.Qiu, P.Mitra and C.L.Giles. Detecting Topic Evolution in Scientific Literature: How Can Citations Help[C]. Proceeding of CIKM, 2009, pp.957-966.
- [10] Y.Jo, C.Lagoze, C. L.Giles. Detecting Research Topics via the Correlation between Graphs and Texts[C]. Proceedings of KDD, 2007, pp.370-379.
- [11] G.S.Mann, D.Mimno, and A.McCallum. Bibliometric Impact Measures Leveraging Topic Analysis[C]. Proceedings of the 6th ACM/IEEE-CS joint conference on Digital libraries , 2006.
- [12] T.L.Griffiths and M.Steyvers. Finding Scientific Topics[C]. Proceeding of the National Academy of Science, 2004, pp.5228-5235.
- [13] D.Hall, D.Jurafsky, and C.D.Manning. Studying the History of Ideas Using Topic Models [C]. Proceedings of the 2008 Conference on Empirical Methods in Natural Language Processing.2008, pp.363-371.
- [14] 楚克明, 李芳. 基于 LDA 话题关联的话题演化[J]. 上海交通大学学报, 2010, 44(11): 1501-1506.
- [15] 单斌, 李芳. 基于 LDA 话题演化研究方法综述[J]. 中文信息学报, 2010, 24(6):43-49.
- [16] Ulrich Schaefer Bernd Kiefer Christian Spurk Jörg Steffen Rui Wang. The ACL Anthology Searchbench[C]. Proceedings of the ACL-HLT 2011 System Demonstrations, pages 7-13.