

研究前沿识别方法探析^{*}

宫小翠 赵迎光 安新颖

(中国医学科学院医学信息研究所 北京 100020)

〔摘要〕 总结研究前沿的识别方法,包括基于文献计量学的识别方法和基于计算机的自动、半自动化方法,指出各自的优缺点,提出研究前沿识别应利用语义网络等工具向更深的粒度、更高的准确度方向发展。

〔关键词〕 研究前沿;发展趋势;文献计量学;LDA;网络主题

〔中图分类号〕 R-058 〔文献标识码〕 A 〔DOI〕 10.3969/j.issn.1673-6036.2015.09.011

Exploration and Analysis of Identification Methods for Research Fronts GONG Xiao-cui, ZHAO Ying-guang, AN Xin-ying, Institute of Medical Information, Chinese Academy of Medical Sciences, Beijing 100020, China

〔Abstract〕 The paper summarizes identification methods for research fronts, including the method based on bibliometrics and the automatic and semi-automatic methods based on computers. Their respective advantages and disadvantages are noted and it is suggested that such tools as semantic network should be utilized to identify research fronts in a deeper and more accurate manner.

〔Keywords〕 Research fronts; Development trend; Literature metrology; LDA; Web topics

1 引言

研究前沿是具有较高发展潜力的知识,是指伴随着某一领域的重要事件的发生而大量出现的知识,其实质就是某一个研究领域内处于领先地位的成果和思想。这一概念最早是 Price^[1]为了描述某一领域的瞬时性特征而提出的,认为在一个给定的研究领域内,科学家积极引用的近期文献的集合所表

征的研究领域就是研究前沿。之后这一概念又产生了多种描述,但至今尚未形成明确、统一的定义。在科学领域中,技术、方法以及科学概念的知识数不胜数,研究前沿不断产生、成熟、平稳和衰退。Upham 等^[2]认为,成功的研究前沿会产生两种结果:成长为独立的研究领域或者被其他领域吸收。成为独立的研究领域表现在本身的发文量增加,被其他的领域吸收表现在引用量增加。多学科性显著的前沿内生性强,多学科性不显著的前沿引用量增加,内生性不强^[2]。

研究前沿知识这种循环演化也是科学技术研究中推动科技创新的力量之所在。从浩瀚的科技信息文献中,识别研究前沿,不仅可以帮助科研人员全面、及时、准确地发现研究领域潜在的知识,而且可以帮助决策人员快速制定相关政策,合理分配科研资源。高潜力的研究前沿挖掘还可能带来新的科技发现,给社会带来利益,吸引更多的科学家参与

〔修回日期〕 2015-05-06

〔作者简介〕 宫小翠,在读硕士研究生;通讯作者:安新颖,博士,副研究员,发表论文 20 余篇。

〔基金项目〕 国家自然科学基金项目“基于语义的医学领域前沿知识发现及演化机制研究”(项目编号:71303259);教育部人文社会科学研究青年基金项目“基于知识组织体系的科技文献新主题监测研究”(项目编号:11YJC870001)。

研究,产生更多的科学发现。

研究前沿挖掘技术在国内的研究较多,相对比较成熟,而在国内的研究目前还比较少。研究前沿的识别方法主要为基于文献计量学的方法和基于计算机技术的自动半自动化方法。

2 基于文献计量学的方法

2.1 基于引用关系的方法

基于引用关系的方法主要包括共被引、文献耦合和直接引用的方法。Small 等^[3]用共被引方法分析了有机膜传感器领域的新主题突显、发展、消亡的过程。李柏洲等^[4]采用国际领先的知识分子结构法进行文献收集与萃取,该方法集成作者同被引分析 (Author Co-citation Analysis, ACA) 与探路者网络 (Pathfinder Network, PFNET),通过引入信息学方法分析相关知识单元间联系的文献,引用和被引用的文献以地图的形式表示并形成树桩或链状等不同形态的网络结构,揭示文献间隐含的关系。Schiebel^[5]认为可以从文献耦合聚类的文献簇中识别研究前沿,共被引文献簇中识别知识基础,提出模仿地理地图的二维和三维图像探测研究前沿和研究基础的可视化方法。万小军等提出一种自动评估引用强度的方法,利用 LibSVM 等工具将引文长度、引用密度、是否自引、施引位置、提及次数、被引与施引的时间间隔考虑在内,对引用文献进行打分,此方法可用于引用关系的前期处理,将打分低微的文献不考虑在内。探测研究前沿以共被引和文献耦合方法居多,直接引用方法则较少较新。基于引用关系的方法识别研究前沿具有可行性,但存在时间的滞后性这一显著的缺点。

2.2 基于内容词的方法

基于内容词的方法包括基于词频和基于共词的方法。基于词频的方法通过识别文献中的爆发词,根据爆发词的时间分布和变化特征识别前沿知识发展趋势。基于共词的方法是对一组词的两两统计它们在同一篇文献中出现的次数,对这些词进行聚类分析,进而分析这些词所代表的学科和主题的结构

变化,是一种运用相对较多的研究前沿识别方法。Ohniwa 等^[6]选取增长率高的 MeSH 术语,用共词的方法将其分组,通过不同的时间窗比较探究生命科学领域的研究前沿^[6]。基于内容词的角度探究前沿知识只是基于单词的共现角度,很难从语义角度进行前沿知识的揭示。

2.3 混合方法

2.3.1 引文与共词的方法相结合 这样更具有优势,可以弥补二者的不足。Chen^[7]开发的 CiteSpace II 将引文与词检测方法相结合,用爆发词检测、引文聚类视图、关键点文献发掘和时区视图方法揭示研究前沿^[7]。混合方法的研究还包括用共被引分析得到高被引核心文献,再用共词分析界定前沿领域,或是基于文献耦合的引文-文本相结合的方法等。

2.3.2 作者耦合分析、 $h-b$ 指数和计算参考文献相对年等 也是掌握研究前沿较为常用的方法。马瑞敏等^[8]用作者耦合方法对图书情报学的知识结构进行可视化分析,证明作者耦合分析能较好地挖掘一个学科的前沿知识结构。杨露^[9]提出基于 $h-b$ 指数和 m 值识别新的热点话题的方法,某一领域有 $h-b$ 篇论文的每篇被引次数不少于 $h-b$ 次,再用 $h-b$ 除以自科学家发表第一篇论文起计算的年数得到 m 值,较大的 m 值表现该领域的论文迄今时间较短,但引用较多可能为新兴领域。计算参考文献中近两年发表的文献占的比例和引文的平均年份也能粗略地估计此文献是否是较新的科研成果。

2.3.3 基于文献计量学方法 比较单薄,不能从多个指标维度系统分析构建热点识别方法,这些方法虽然从发表量、引用频次等各种角度在数量上进行统计分析,但是反映的只是知识演化的整体过程和宏观走向,无法体现知识体系演化的历史阶段特征和微观特征。基于计算机的全自动化或半自动化系统,在研究前沿的探测上相对比较高效,且目前正处起步阶段,有比较大的发展空间和比较好的应用前景。

3 基于计算机的全自动或半自动化方法

3.1 文献中前沿话题识别

文本挖掘是利用自动或半自动方法发现文本中

潜在的、先前未知的知识的过程。文本挖掘是一个多学科交叉领域,融合了数据挖掘、自然语言处理、机器学习和信息检索方面的知识。

3.1.1 研究现状 文本挖掘中的自动探测方法大体上可以分为主题特征表示、主题识别、主题判定3个阶段。国内相关研究比较少,国外关于学科趋势的研究不少,热点及趋势分析只是其研究中的分支,并没有被系统地研究和分析。R. Swan 和 D. Jensen 开发出了 Time - Mines 系统,利用信息抽取技术、自然语言处理技术来抽取有时间标签的自由文本数据,并采用假设检验技术来判定给定时间框中最相关的主题,主题是否是新兴交由用户判断。Havre 等^[10]开发的可视化主题演化系统 Theme River,每条河流代表一个主题,河流用不同的颜色表示,河流的宽度代表主题强度,可以直观地观测主题强度随时间的变化。Patent Miner 系统旨在用动态生成的 SQL 查询发现专利数据中的发展趋势,系统与包含所有美国专利的 IBM DB2 数据库相连,主要包含两个组件:基于序列模式挖掘的短语识别和基于形状查询的趋势预测。Rapid Miner^[11]系统是一个机器学习和数据挖掘试验环境,它允许试验由大量任意嵌套的操作构成,提供一个图形用户接口(Graphical User Interface, GUI)设计分析管道,由 GUI 产生 XML 文件,XML 文件中定义了用户希望对数据实施的分析操作,这个文件然后由 Rapid Miner 自动读取。Weka 是一套用 JAVA 编写的流行的机器学习软件,Weka 工作台包括一批可视化工具和数据分析算法,并有一些图形用户接口,方便用户实现数据分析预测功能。各种文本挖掘过程在进行数据预处理后,然后进行聚类、分类等过程,在识别研究前沿方面,语义关联关系识别是非常重要的。语义关联关系识别需要考虑两个方面的问题:命名实体识别和实体关系识别,基于规则和统计是两种常用的方法,基于规则的方法需要许多的领域专家参与制定规则。目前广泛使用的统计方法有隐马尔科夫模型、最大熵模型以及条件随机场模型,但是随着大量新术语的不断出现和语义关系的不断丰富,给单纯的基于规则和统计的方法带来了极大的挑战。

3.1.2 主题模型 文本挖掘技术是帮助科研人员从海量文献中快速发现新兴主题的途径之一,而主题模型作为一套新的能够对文献资源语义抽取的算法^[12],提供了一种新的解决问题的方法,并成为国际研究的热点。主题模型是文本降维技术的一种。Hofmann^[13]等提出 PLSI 模型,认为一篇文档由条件依赖于文档的多个主题组成,表示主题的单词服从于主题的多项式分布;但模型的参数随着文集增长呈线性增长,即出现过拟合问题。之后, Blei 等^[14]提出隐含狄利克雷分布(Latent Dirichlet Allocation, LDA)模型,认为文档由服从多项分布的主题构成,每个主题由服从多项分布的单词构成。

基于 LDA 话题演化主要采用了3种方法:(1)时间演化话题(Topic Over Time, TOT)^[15],该模型把连续的时间信息引入到 LDA 模型中,对共现词和文档的时间共同建模,这样模型的时间是连续的,不会出现在离散时间的方法中时间粒度选取的问题,但此模型仅考虑话题强度的变化趋势,而忽略了话题内容的变化;另外,该模型必须一次对所有文档运用,无法对新文档扩展。(2) LDA 后离散方法,在整个文集上应用 LDA 模型,然后根据文档的时间戳划分子集,再用子集的话题分布表征该段时间的话题,该方法认为不同时间的话题数不变,难以探测新话题产生与旧话题消失。(3)先离散分析方法,在建模前,根据时间段划分文档,然后依次处理每个时间窗口上的文本集合,最终形成话题随时间的演化,在离散分析的方法中,下一时刻的模型参数往往依赖于当前时刻或前几个时刻的模型参数的后验或者模型输出结果,如动态话题模型(Dynamic Topic Model, DTM)^[16]和在线话题模型(Online - LDA, OLDA)^[17]。

在数据挖掘方面,对 LDA 模型已有大量研究,产生各种变体。结合 AT、TOT 模型史庆伟等^[18]提出作者主题演化模型(Author - Topic over Time, ATot),用主题 - 词项分布与作者 - 主题分布分别用来描述主题随时间的变化规律和作者研究兴趣的变化规律。将引文分析引入主题模型中,改进对主题演化的预测正得到越来越多人的关注,Nallapati 等^[19]提出将施引文献和被引文献构成一个文献对,

将文献对放在主题模型中联合建模,主题之间的影响是服从于一个超参数决定的概率分布。叶春蕾等^[20]在贝叶斯概率模型中引入引文因素,将引文作为模型的一项参数,与关键词共同建立引文-主题模型,并通过试验表明该模型能全面、深入地识别科技文献中的主题内容。曹建平等^[21]根据时效性较强的在线的文本数据流,提出一种基于 LDA 的双轨道在线主题演化模型(BPE-OLDA),在下一时间片生成文本时考虑文本的内容遗传和强度遗传,估算模型参数对 Gibbs 采样算法进行了简化,试验证明该模型在提取时效性较强的文本数据流的主题方面具有明显的效果。LDA 主要是基于 Kullback Leibler (KL)^[22]相对熵或 Jensen-Shannon (JS)^[23]来计算主题间的相似性,用改进的 Z-Score 方法计算主题随时间的偏移反应主题演化的情况,从而发现主题演化中的主题遗传和主题变异^[24]。虽然该方法可以探测到主题间一对多,多对多的演化关系,但反映的仍是宏观走向,不能体现知识演化的微观特征。

3.2 网络中新话题检测方法

3.2.1 基于传统的话题检测方法 基于关键词的方法对于热门话题挖掘是较好的方法,但是对于新兴话题的识别效果则不是很好。Cataldi 等^[26]把在一段时间内突发出现的频次较高的词称为新兴关键词,他们识别新兴关键词并利用共词的方法找到与它们高频共现的词从而发现新兴话题。

3.2.2 基于词典学习和非概率矩阵分解的方法 Saha 和 Sindhwani^[26]采用一个非负因式分解的方法学习社交媒体文本流中的新兴话题,当连续时间戳方面的主题矩阵连续性考虑在内时,会表现出一个更好的话题建模性能。

3.2.3 从多个维度来分析热点话题 姜晓伟等^[27]首先在网上搜集并格式化出现感兴趣的词的微博,对于这些微博中的所有词汇,综合考虑影响力、突发性和相关性 3 个要素对其重要性进行评估;其次用含有同一关键词的微博的集合为输入文档训练 LDA 模型;然后通过对主题关键词的概率分布的推导,实现对词的聚类 and 主题的挖掘。

3.2.4 综合考虑多个主客体 通过分析访问者留言及链接关系等来发现热点话题,这些方法都使得结果更加客观。如湛志群等^[28]在采用共词分析和 Bisecting K-means 聚类算法检测网络论坛热点话题基础上,提出了一个综合考虑话题帖子篇数与帖子热度的热点话题关注度计算方法。陈立伟等^[29]将各向异性扩散技术引入词网,在词网中体现相关词语的影响,同时保护主题间的边界,提出有限记忆和被动冷却机制,利用有限的存储空间对词网进行部分索引,不扫描和处理不活动词语,实现热门主题及其词语的快速访问,利用有限的存储资源记录互联网文字中包含的主题,对于未知主题可以自动识别和记录,对相关主题自动联想^[29]。

4 结论

研究前沿识别方面的研究在国内研究的比较少,在国外相对比较多,大多基于统计学和计量学进行识别,但是该方法角度比较单一,没能从多维度、深角度对研究前沿知识进行识别,前沿知识粒度识别不够精深。基于计算机的自动半自动化方法正在兴起,但仍处于起步阶段,LDA 模型是近年来研究的热点,由于 LDA 本身是概率式的产生模式,存在不确定性,主题演化的粒度描述同样存在不够专深的问题。网络中热点话题的识别方法研究比较多,对于文献前沿热点挖掘有许多值得借鉴的地方。

虽然研究前沿识别方面已有一定的成果,但是还有以下几个问题需要解决:(1)研究前沿的定义没有形成统一的共识,需要有一个系统的揭示和综合的解释。研究前沿的定义及判定标准随所采用方法的不同而不同,虽然已有研究者尝试设计一套指标来辅助判定研究前沿,但公认的客观可信赖的指标体系还有待进一步研究。(2)对研究前沿的特征没有形成客观统一的描述。国内关于研究前沿特征的研究还比较少,关于研究前沿基本特征的描述还没有形成客观、科学的模型,尽管有学者指出研究前沿的一些特点和规律,但是比较系统、客观的描述还没有形成,需要进一步研究。(3)研究前沿的

解读缺乏一定的语义环境。国内研究前沿的发现大部分是基于统计的方法, 这样对于研究前沿的解读带来一定困难。如果利用一些语义网络工具, 不仅可以提高前沿监测的准确度, 而且还可以从更深的粒度对研究前沿进行解读。

在未来的研究中, 应尝试利用生命周期、知识传播、系统动力学等理论, 开展研究前沿识别研究, 同时将语义网络应用到研究前沿发现的重要环节, 为知识单元内部关系识别、知识群间关系识别以及知识演化分析等方面提供技术支持, 从而提高研究前沿发现的效率和精度。

参考文献

- Price D D. Networks of science papers [J]. Science, 1965, (149): 510-515.
- Uphams, Small H. Emerging Research Fronts in Science and Technology: patterns of new knowledge development [J]. Scientometrics, 2010, 83 (1): 15-38.
- Small H, Upham P. Citation structure of an Emerging Research Area on the Verge of Application [J]. Scientometrics, 2009, 79 (2): 365-375.
- 李柏洲, 赵健宇, 袭希, 等. 基于知识分子结构法的知识管理研究主题演化趋势分析 [J]. 研究与发展管理, 2014, 26 (2): 59-75.
- Schiebel E. Visualization of Research Fronts and Knowledge Bases by Three-dimensional Areal Densities of Bibliographically Coupled Publications and Co-citations [J]. Scientometrics, 2012, (91): 557-566.
- Ohniwa R, Hibino A, Takeyasu K. Trends in Research Foci in Life Science Fields Over the Last 30 Years Monitored by Emerging Topics [J]. Scientometrics, 2010, (85): 111-127.
- Chen C. CiteSpace II: detecting and visualizing emerging trends and transient patterns in scientific Literature [J]. Journal of the American Society for Information Science and Technology, 2006, 57 (3): 359-377.
- 马瑞敏, 倪超群. 作者耦合分析: 一种新学科知识结构发现方法的探索性研究 [J]. 中国图书馆学报, 2012, 38 (198): 4-11.
- 杨露. h-b 指数: 一种确定学术研究热点的新方法 [J]. 四川教育学院报, 2012, 28 (5): 24-26.
- Havre S, Hetzier B, Nowell L. Theme River: visualizing thematic changes in large document collections [J]. IEEE Transactions on Visualization and Computer Graphics, 2002, 8 (1): 9-20.
- Ganesh M S, Reddy CHP, Manikandan, et al. TDPA: trend detection and predictive analytics [J]. International Journal on Computer Science and Engineering, 2011, 3 (3): 1033-1039.
- Blei D M. Probabilistic Topic Models [J]. Communications of the ACM, 2012, 55 (4): 77-84.
- Hofmann T. Probabilistic Latent Semantic Indexing [C]. New York: Proceedings of the 22nd Annual International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR99). 1999: 50-57.
- Blei D M, Ng A Y, Jordan M I. Latent Dirichlet Allocation [J]. Journal of Machine Learning Research, 2003, (3): 993-1022.
- Wang X, McCallum A. Topic Over Time: a non-markov continuous-time model of topical trends [C]. Philadelphia, USA: ACM SIGKDD-2006, 2006: 424-433.
- Blei D M, Lafferty J D. Dynamic topic models [C]. Pittsburgh, Pennsylvania: Proceedings of the 23rd International Conference on Machine Learning, 2006: 113-120.
- Alsumait L, Barbara D, Domeniconi C. On-line LDA: adaptive topic models for mining text streams with applications to topic detection and tracking [C]. Pisa, Italy: In IC-DM, 2008: 3-12.
- 史庆伟, 李艳妮, 郭朋亮. 科技文献中作者研究兴趣动态发现 [J]. 计算机应用, 2013, 33 (11): 3080-3083.
- Nallapati R M, Ahmed A, Xing E P, et al. Joint Latent Topic Models for Text and Citations [C]. New York: Proceedings of the 14th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD08), 2008: 542-550.
- 叶春蕾, 冷伏海. 基于引文-主题概率模型的科技文献主题识别方法研究 [J]. 信息系统, 2013, 36 (9): 100-103.
- 曹建平, 王晖, 夏友清, 等. 基于 LDA 双通道在线主题演化模型 [J]. 自动化学报, 2014, 40 (12): 2877-2886.
- 李保利, 杨星. 基于 LDA 模型和话题过滤的研究主题演化分析 [J]. 小型微型计算机系统, 2012, 33 (12): 2738-2743.

(下转第 64 页)

- 14 Basu A, Kumar BSV. International Collaboration in Indian Scientific Papers [J]. *Scientometrics*, 2000, 48 (3): 381 – 402.
- 15 Glänzel W, de Lange C. A Distributional Approach to Multinationality Measures of International Scientific Collaboration [J]. *Scientometrics*, 2002, 54 (1): 75 – 89.
- 16 钟镇. 学术论文合著规模与引文影响力的相关性分析——以图书情报学为例 [J]. *情报杂志*, 2013, 32 (12): 127 – 131, 164.
- 17 余新丽, 赵文华, 杨颀. 我国研究型大学国际合作论文的现状与趋势分析——以上海交通大学为例 [J]. *中国高教研究*, 2012, (8): 30 – 34, 39.
- 18 Teodorescu D, Andrei T. The growth of international collaboration in East European scholarly communities: a bibliometric analysis of journal articles published between 1989 and 2009 [J]. *Scientometrics*, 2011, 89 (2): 711 – 722.
- 19 Terekhov AI. Scientific Collaboration in the Sphere of Carbon Nanostructures in the Mirror of Bibliometric Analysis [J]. *Scientometrics*, 2014, 84 (4): 265 – 271.
- 20 温芳芳, 李佳靓. 中国情报学期刊论文合著现象分析——基于五种情报学核心期刊的统计分析 [J]. *情报杂志*, 2011, 30 (8): 55 – 60.
- 21 Miquel JF, The Changing Patterns of International Collaboration in Universities [J]. *Scientometrics*, 1991, (59): 141 – 151.
- 22 郑如青, 张琰. 北京大学科研国际合作的成效与发展对策 [J]. *北京大学学报: 自然科学版*, 2010, 46 (5): 851 – 854.
- 23 范爱红, 战玉华, 杨芳, 等. 国内外研究型大学国际合作论文的比较研究 [J]. *情报杂志*, 2013, 32 (11): 59 – 63.

(上接第 51 页)

- 23 李湘东, 张娇, 袁满. 基于 LDA 模型的科技期刊主题演化研究 [J]. *情报杂志*, 2014, 33 (7): 115 – 121.
- 24 崔凯, 周斌, 贾焰, 等. 一种基于 LDA 的在线主题演化挖掘模型 [J]. *计算机科学*, 2010, 37 (11): 156 – 159.
- 25 Cataldi M, Caro L D, Schifanella C. Emerging Topic Detection on Twitter Based on Temporal and Social Terms Evaluation [C]. New York: Proceedings of the 10th International Workshop on Multimedia Data Mining, 2010.
- 26 Saha A, Sindhwani. Learning Evolving and Emerging Topics in Social Media: a dynamic nmf approach with temporal regularization [C]. Seattle Washington: Proceedings of the 5th ACM International Conference on Web Search and Data Mining, 2012.
- 27 姜晓伟, 王建民, 丁贵广. 基于主题模型的微博重要话题发现与排序方法 [J]. *计算机研究与发展*, 2013, (50): 179 – 185.
- 28 湛志群, 徐宁, 王荣波. 基于主题演化图的网络论坛热点跟踪 [J]. *情报科学*, 2013, 31 (3): 147 – 150.
- 29 陈立伟, 谢朝阳, 唐权华. 基于各向异性热度扩散的主题检测方法 [J]. *计算机工程与设计*, 2014, 35 (8): 2886 – 2889.

(上接第 55 页)

- 5 Pak T R, Roth F R. ChromoZoom: aflexible, fluid, web – based genome browser [J]. *Bioinformatics*, 2013, 29 (3): 384 – 386.
- 6 Frant L, Goldstein B, Ma Y, et al. MedlinePlus Mobile: consumer health information on – the – go [J]. *IT Prof*, 2012, 14 (3): 44 – 49
- 7 Vercruyssen, Venkatesan A, Kuiper M. OLSVis: an animated, interactive visual browser for bio – ontologies [J]. *BMC Bioinformatics*, 2012, (13): 116.
- 8 HTML5 TEST [EB/OL]. [2015 – 01 – 10]. <http://html5test.com/compare/browser/chrome06.html>.
- 9 深度分析 HTML5 在移动开发方面的发展状况 [EB/OL]. [2015 – 01 – 10]. <http://www.evget.com/zh – CN/info/catalog/16374.html>.
- 10 主流浏览器 CSS3/HTML5 兼容性清单 [EB/OL]. [2015 – 01 – 10]. <http://tools.yesky.com/233/30105733.shtml>.
- 11 侯震, 侯丽, 李姣. 医学美术作品的创作与管理 [J]. *中国医学教育技术*, 2013, 27 (5): 614 – 617
- 12 Sinha A U, Armstrong S A. iCanPlot: visual exploration of high – throughput omics data using interactive canvas plotting [J]. *PLoS One*, 2012, 7 (2): e31690.
- 13 Leon P S, Knock S A, Woodman M M. The Virtual Brain: a simulator of primate brain network dynamics [J]. *Front Neuroinform*, 2013, (7): 10.