

Postprint: Citation Sentiment Recognition Based on Citation Content Analysis

Authors: Liao Junhua, Liu Ziqiang, Bai Rujiang, Chen Junying

Date: 2023-08-27T00:00:00+00:00

Abstract

[Purpose/Significance] To address the challenge of automatically identifying citation sentiment in academic papers, this study proposes a novel identification methodology based on citation content analysis accompanied by visual presentation, thereby overcoming the limitation that simple citation frequency metrics fail to distinguish among different citation sentiments. [Method/Process] Initially, regular expressions are employed to extract citation content information from full-text papers; subsequently, the TF-IDF algorithm is utilized to identify citation sentiment feature words, which are combined with a sentiment lexicon to perform citation sentiment identification on citation content using sentiment analysis techniques; finally, visualization tools are applied to demonstrate the overall distribution of citation sentiment. [Results/Conclusion] The proposed method effectively identifies citation sentiment within the anti-aging research paper dataset. Experimental results reveal that in this field, positive citations constitute 21% of total citations, neutral citations account for 78%, and negative citations represent merely 1%. Compared with traditional citation networks, visualization graphs based on citation sentiment can effectively discern the distribution patterns of different citation sentiments across the entire dataset.

Full Text

Preamble

Citation Sentiment Recognition Based on Citation Content Analysis

Liao Junhua¹, Liu Ziqiang^{2,3}, Bai Rujiang¹, Chen Junying¹

¹Institute of Scientific and Technical Information, Shandong University of Technology, Zibo 255049

²Chengdu Library and Information Center, Chinese Academy of Sciences, Chengdu 610041

³University of Chinese Academy of Sciences, Beijing 100190

Abstract

[Purpose/Significance] This paper proposes a citation sentiment identification method based on citation content analysis and visualizes the results to overcome the limitation of simple citation frequency metrics that cannot distinguish different citation sentiments. **[Method/Process]** First, regular expressions are used to extract citation content information from full-text papers. Then, the TF-IDF algorithm is employed to select sentiment feature words, which are combined with a sentiment dictionary to identify citation sentiments using sentiment analysis techniques. Finally, visualization tools are used to display the overall distribution of citation sentiments. **[Result/Conclusion]** The method effectively identifies citation sentiments in the anti-aging research domain. Experimental results show that positive citations account for 21% of total citations, neutral citations account for 78%, and negative citations account for only 1%. Compared with traditional citation networks, the sentiment-based visualization map can effectively identify the distribution of different citation sentiments across the entire dataset.

Keywords: citation content analysis; citation sentiment; sentiment analysis; visualization

2 Related Research

Citation content analysis research emerged as early as the 1970s. M. J. Moravcsik et al. conducted detailed analyses of citation sentiment tendencies, functions, and importance by interpreting citation content and context, demonstrating the necessity of citation content analysis [4]. Subsequently, H. Small used manual reading and summarization methods to analyze citation content of highly cited papers in chemistry, considering citation content as conceptual symbols representing viewpoints in citing literature. As citation content analysis research progressed, scholars attempted to improve citation analysis methods based on citation counts by incorporating citation content analysis [5]. In 1980, H. Small et al. proposed a co-citation clustering analysis method combined with citation content analysis, which first analyzed the evolution of a discipline based on co-citation clustering, then used subject terms and phrases extracted from citation content to characterize specific citation content, thereby analyzing the topics of co-citation clusters and revealing deep associations among co-cited literature. This approach enhanced the cognitive value of co-citation links. Through empirical research in the recombinant DNA field, they demonstrated that this method was significant for exploring the development and evolution of disciplinary fields [6]. However, due to limitations in full-text database quality and computer technology at the time, researchers primarily relied on manual reading and summarization for citation content analysis, which was difficult to apply to large samples and suffered from strong subjectivity, limiting the further development of citation content analysis.

With the development of natural language processing technology and full-text

databases, citation content analysis has gained new momentum and injected vitality into traditional citation analysis. In 2012, Y. Ding proposed a citation content analysis (CCA) research framework, arguing that citation content analysis represents the next generation of citation analysis and can expand and deepen citation analysis research and applications [7]. Zhu Qingsong and Leng Fuhai conducted citation content mining and analysis on highly cited papers in the carbon nanotube field, using the C-value algorithm to identify research topics in citation content. Their study showed that compared with topic identification based on titles and abstracts, citation content analysis produced results more aligned with paper content and better revealed semantic associations between cited and citing literature, suggesting that citation content analysis is an important supplement to traditional citation count-based analysis [8]. Lu Wei et al. argued that to better support semantic relationship mining in literature, natural language processing and machine learning techniques should be introduced into citation content analysis, and proposed a citation content annotation framework [9]. Zhao Rongying et al. considered citation content analysis a new development in citation analysis that can more accurately measure and evaluate the influence of cited authors and journals, and provide insights into authors' citation motivations, greatly benefiting scientometrics and the science of science [10]. Building on this, Zhao Rongying et al. in 2016 proposed a location-based co-citation analysis framework combined with citation content analysis, demonstrating that this approach was significantly superior to traditional co-citation analysis methods [11].

In terms of citation sentiment type research, as early as 1962, E. Garfield identified the limitations of citation frequency analysis, noted the diversity of citation sentiments, and summarized 15 types of citation sentiments including paying tribute to pioneers and peers [12]. Subsequently, V. Cano and other experts also pointed out that due to the complexity of citation behavior, citation sentiments are not always positive (negative citations, false citations, etc. exist), and simple citation counts are insufficient to measure academic influence [13-17]. Therefore, accurately and efficiently identifying citation sentiment tendencies and distinguishing positive from negative citations can effectively improve the quality of evaluation based on simple citation frequency.

Regarding citation sentiment identification methods, M. J. Moravcsik et al. studied citation sentiment through manual reading of full citation texts, dividing citation sentiment into five dimensions including positive and negative citations [4]. In the same year, based on Moravcsik's research, D. E. Chubin et al. studied the auxiliary and alternative roles of citation content analysis for bibliographic analysis, using classification trees for citation sentiment analysis with the first level divided into positive and negative citation subtrees [18]. Overall, due to limitations in information technology, citation sentiment identification primarily relied on questionnaire surveys and manual reading, suffering from low efficiency and strong subjectivity.

In recent years, citation sentiment identification research based on natural lan-

guage processing technology has made progress. S. Teufel et al. proposed a supervised machine learning framework for automatic citation sentiment classification, using sentiment analysis techniques to classify citation functions (sentiments) into four categories: weakness, strength, contrast, and neutral, demonstrating that sentiment analysis technology can accurately and effectively identify citation sentiment [19]. Liu Shengbin et al. proposed a citation sentiment identification method based on data mining technology, using the PubMed full-text database as a data source, employing citation content semantic structure and feature words to identify citation sentiment (positive, negative, and neutral), and built a citation evaluation platform based on citation content [20]. Compared with survey and manual reading methods, natural language processing-based citation sentiment identification improves analysis efficiency and objectivity, though the analysis of results remains relatively superficial, requiring further research on how to effectively utilize citation sentiment identification results.

In summary, traditional citation sentiment identification methods primarily rely on questionnaire surveys and manual reading, suffering from low efficiency and strong subjectivity. Statistical natural language processing-based methods struggle to effectively analyze citation sentiment identification results, reducing accuracy and effectiveness to some extent. Particularly in feature lexicon construction, previous research mainly used existing sentiment dictionaries for polarity discrimination, but scientific paper citation sentiment differs significantly from general sentiment dictionaries. Therefore, to improve current citation sentiment identification research, this paper proposes a citation sentiment identification method based on citation content analysis. By using tf-idf combined with part-of-speech tagging to construct a scientific literature citation sentiment lexicon, we propose a citation sentiment discrimination model to achieve citation sentiment polarity identification. Accurate identification of citation sentiment in papers can provide data support for distinguishing different citation behaviors in citation frequency-based bibliometrics.

3 Methodology Framework

To accurately and effectively identify citation sentiment, drawing on existing citation motivation identification theories and methods and combining data mining, sentiment analysis, and visualization techniques, we propose a citation sentiment identification method based on citation content analysis (see Figure 1 [Figure 1: see original paper]). This method uses full-text citation data as the research object. First, regular expression technology extracts citation information from full-text papers. Then, feature dictionaries and sentiment analysis techniques identify citation sentiment. Finally, visualization analysis methods analyze and display citation sentiment identification results.

Step 1: Dataset Construction. Retrieve relevant domain research papers from full-text databases, use Python crawler programs to obtain XML-format full-text data, and save it to a local computer.

Step 2: Citation Content Extraction. Write full-text information extraction programs on the Python platform to extract metadata and citation content information for both citing and cited literature, and write it to CSV files for subsequent analysis.

Step 3: Citation Sentiment Identification. Use the tf-idf method to select feature words from citation content, combine them with a sentiment dictionary, and use sentiment analysis technology to analyze and identify citation sentiment in citation content.

Step 4: Citation Sentiment Visualization Analysis. Based on citation sentiment identification results, construct a complex network dataset containing citation sentiment to reveal the distribution of citation sentiment.

3.1 Citation Content Extraction

Paper data format significantly affects citation content extraction effectiveness. PDF-format full-text data is difficult to parse and has poor readability, resulting in low extraction accuracy and difficulty processing large samples. Compared with unstructured full-text data such as PDF, XML full-text data provides detailed annotations for papers (preprocessing full-text data, labeling figures, tables, citation content, citation positions, etc.), facilitating computer-based extraction and analysis of large-sample citation content information.

Using the XML full-text of “Neuroprotective and Anti-Aging Potentials of Essential Oils from Aromatic and Medicinal Plants” from the PubMed database as an example, its structure was analyzed (see Figure 2 [Figure 2: see original paper]). The XML full-text data includes title, journal, author, abstract, figures, tables, citation content, citation positions, and other information. Python programs can extract citing literature metadata (... , 28611658), cited literature tags and citation content (ref-content, Author, pub-date), and cited literature tags and bibliographic information (... , ...).

Extracting ref-content (Author, pub-date) tags from full-text is the key and challenging aspect. Considering authors’ non-standard and inconsistent use of reference numbers during paper writing, the following two situations must be considered when constructing citation content extraction rules:

- (1) **Citing only one reference:** “The EOs are abundant in flowers, leaves, barks and are usually isolated via hydro-distillation, cold pressing methods (Edris, 2007).” When a sentence contains only one ref-content tag (one reference), the entire sentence containing the ref-content tag is extracted as citation content.
- (2) **Citing two or more references:** “The use of EOs as therapeutic remedy is very ancient and in the bible (Dumas and Newhouse, 2011), EOs were considered as spiritual, mental and physical healing agents (Guenther, 1950). Thus, a boost in the cholinergic tone may potentially regress the cognitive function.” When a sentence contains two or more ref-content tags

(two or more references), the citation sentence is divided using adjacent pub-date tags as markers, with the smallest unit being a clause.

3.2 Citation Sentiment Identification

3.2.1 Citation Sentiment Category Division For over half a century, numerous experts including E. Garfield, H. Small, and W. Shadish [21] have conducted in-depth research on citation sentiment category division. From a methodological perspective, approaches have gradually shifted from questionnaire surveys and manual reading to natural language processing, sentiment analysis, and visualization methods. From a category division perspective, citation sentiment categories have become increasingly concise and generalized to facilitate large-sample data analysis. This simplification has two main reasons: first, the transformation of citation sentiment identification methods makes it difficult to accurately identify overly complex citation sentiments; second, the explosive growth of citation analysis data volume means overly detailed category division reduces analysis efficiency and accuracy.

Therefore, drawing on the citation sentiment categories of S. Teufel, Liu Shengbin, and other scholars [19-20], this paper divides citation sentiment into three categories: positive citation, negative citation, and neutral citation.

3.2.2 Citation Sentiment Identification Based on Feature Words and Sentiment Dictionary Due to the normative and scientific nature of journal papers, strongly emotional statements rarely appear in full texts. Compared with microblog and forum data, journal papers lack sentiment vocabulary, making it difficult to judge citation sentiment through simple sentiment words. Therefore, using sentiment analysis technology alone for citation sentiment identification has limitations [22-24]. This paper uses tf-idf with part-of-speech tagging to filter polar feature words and combines them with a general sentiment dictionary for citation sentiment identification.

Citation content primarily consists of content words and feature words. Content words are the main carriers of information in citations, mainly manifested as nouns. Feature words express sentiment and state in citations, mainly manifested as adjectives, verbs, and conjunctions. By analyzing feature words, we can effectively understand sentence sentiment. Therefore, we first use Stanford POSTagger [25] for part-of-speech tagging, annotating feature words from four categories: adjectives (JJ), verbs (VB), adverbs (RB), and conjunctions (CC). For example, the sentence “The objective of anti-aging medicine is to live as long as possible in good health.” is tagged as: The_{DT} objective_{NN} of_{IN} anti-aging_{NN} medicine_{NN} is_{VBZ} to_{TO} live_{VB} as_{RB} long_{RB} as_{IN} possible_{JJ} in_{IN} good_{JJ} health_{NN}.

Then, based on the TF-IDF algorithm, adjectives (JJ), verbs (VB), adverbs (RB), and conjunctions (CC) are selected as feature words. The TF-IDF algorithm formula is:

$$TF-IDF_{i,j} = TF_{i,j} \times IDF_{i,j} \quad (TF_{i,j} = \frac{n_{i,j}}{\sum_k n_{k,j}}, \quad IDF_{i,j} = \log \frac{|D|}{1 + |D_{t_i}|})$$

where $n_{i,j}$ refers to the frequency of term t_i in document d_j , and $\sum_k n_{k,j}$ refers to the total number of terms in document d_j ; $|D|$ refers to the total number of documents, and $|D_{t_i}|$ refers to the number of documents containing term t_i (the 1+ in the denominator prevents division by zero).

After selecting feature words using the TF-IDF algorithm, we combine them with HowNet to construct a sentiment dictionary and assign sentiment values (see Table 1). If a term is not in the HowNet sentiment dictionary, we manually annotate and organize it to form the final citation sentiment dictionary. Using formula (2), we judge citation sentiment in citation content to improve identification accuracy.

Table 1: Citation Sentiment and Sentiment Dictionary

Sentiment Type	Example Words	Sentiment Intensity
Positive	accurate better complete convinced agreement proud...	0.5
Negative	burden contrast degrade lack poor worse however regret...	-0.5
Neutral	apply describe discuss publish use...	0

The sentiment intensity E is calculated using formula (2):

$$E(S_i) = \frac{\sum_{i=1}^n E(S_i)}{n}$$

where E represents the sentiment intensity value of the citation sentence (sum of sentiment values of all sentiment words in the sentence), and $E(S_i)$ represents the sentiment intensity value of word S_i in the citation sentence.

For example, to calculate the sentiment intensity of the citation sentence “However, this algorithm is very slow and has been outperformed by more recent methods,” based on the constructed sentiment dictionary, we can identify the sentiment words “However,” “slow,” and “outperformed.” Using formula (2), the sentiment intensity value is $E = [(-1) + (-1) + (-1)]/3 = -1$, indicating a negative citation.

3.2.3 Citation Sentiment Visualization Since identified citation sentiment results cannot holistically indicate the distribution of different sentiment literature in the dataset, we need to add sentiment markers to traditional citation network visualization analysis methods and use social network visualization software such as Gephi [26] for citation sentiment visualization analysis.

The specific steps for constructing citation sentiment visualization maps are:

- (1) **Data Conversion.** Convert the citing literature ID, cited literature ID, and citation sentiment from the citation sentiment identification result table into edge table data format recognizable by Gephi, using Source, Target, and Weight labels respectively. Positive, negative, and neutral citations are marked as 1, -1, and 0, then saved as .csv files for subsequent processing.
- (2) **Initial Citation Network Visualization Map Construction.** Import the edge table data generated in the previous step into Gephi to calculate and parse edge and node data, generating an initial citation network map. Then customize the map layout (node and edge sizes and colors, etc.).
- (3) **Citation Sentiment-Labeled Network Visualization Map Construction.** Based on the initial citation network visualization map and citation sentiment data, further adjust the layout to construct a citation network visualization map with sentiment labels. Specifically, edges with Weight=1, -1, and 0 are represented by red, green, and yellow directed lines indicating positive, negative, and neutral citations respectively.

4 Experiments and Analysis

4.1 Experimental Environment and Dataset Construction

- (1) **Hardware:** Windows 10 system, i5-2450 CPU, 8GB RAM, 500GB Hard Drive.
- (2) **Software:** Python, KNIME, Stanford POSTagger, Gephi.
- (3) **Dataset Construction.** Using XML-format full-text papers in the anti-aging domain from the PubMed biomedical database as the research object. The search query (TITLE:“anti-aging” OR ABSTRACT:“anti-aging” OR KW:“anti-aging”) was used to retrieve 1,135 relevant papers from PubMed up to December 31, 2016. A Python crawler program was written using PubMed’s OpenURL interface to crawl and save XML full-text data to a local computer.

Based on the analysis in Section 3.1, we extracted the relationships between citing and cited literature and the corresponding citation content in full texts, obtaining 45,527 citation sentences. We manually annotated the sentiment polarity of 2,000 sentences (see Figure 3 [Figure 3: see original paper]), where article_{id1} represents the citing literature, article_{id2} represents the cited literature, ref_{content} represents citation content, section represents citation location, and pub_{date} represents the publication date of cited literature.

4.2 Citation Sentiment Identification

4.2.1 Feature Word Selection Based on TF-IDF After completing citation content extraction, Stanford POSTagger was used to annotate part-of-speech in citation content. Then TF-IDF was used to select high-weight feature

words (adjectives JJ, verbs VB, adverbs RB, and conjunctions CC) (partial results in Figure 4 [Figure 4: see original paper]). For 45,629 citation sentences, 1,197,191 words were tagged for part-of-speech. Data mining software KNIME was used to calculate TF, IDF, and TF-IDF values, generating feature word lists based on TF-IDF values for adjectives (JJ), verbs (VB), adverbs (RB), and conjunctions (CC) (see Table 2).

4.2.2 Citation Sentiment Identification Based on Sentiment Dictionary Combining sentiment vocabulary from the basic HowNet dictionary with selected feature words, we constructed an anti-aging domain sentiment dictionary. The specific steps were: first, manually annotate the top 300 feature words by TF-IDF value for adjectives (JJ), verbs (VB), adverbs (RB), and conjunctions (CC), dividing them into positive, negative, and neutral categories; second, combine them with sentiment vocabulary from HowNet, deduplicate and merge the two word lists; finally, obtain the anti-aging domain sentiment dictionary needed for this experiment, containing 4,570 positive sentiment words, 4,363 negative sentiment words, and 235 neutral sentiment words (see Figure 5 [Figure 5: see original paper]).

Based on the proposed citation sentiment calculation model, sentiment polarity discrimination experiments were conducted on the KNIME data mining platform using manually annotated citation polarity datasets. Two experimental groups were designed: Group 1 used sentiment vocabulary from the basic HowNet dictionary as feature words, while Group 2 used the sentiment dictionary constructed in this paper. Both groups applied the sentiment discrimination method proposed in Section 3.2.2 to 45,527 citation sentences from PubMed's anti-aging domain (see Figure 6 [Figure 6: see original paper] for the experimental workflow).

For result evaluation, precision, recall, and F1-score metrics were used. Experimental results are shown in Table 3 :

Table 3: Experimental Results Analysis

Sentiment Type	Precision	Recall	F1-Score
	Exp 1	Exp 2	Exp 1
Positive	77.31%	78.62%	79.32%
Negative	81.20%	83.00%	83.20%
Neutral	76.16%	76.31%	75.36%

Table 3 shows that Group 2 achieved performance improvements in precision, recall, and F1-score across all categories, with more significant improvements in negative citation discrimination. This demonstrates that the citation sentiment dictionary constructed in this paper effectively improves identification performance. The Group 2 results show positive citations account for 20.74% of total

citations, neutral citations account for 77.82%, and negative citations account for only 1.43% (see Table 4):

Table 4: Citation Sentiment Percentage Statistics

Sentiment Type	Percentage
Positive	20.74%
Negative	1.43%
Neutral	77.82%

Figure 7 [Figure 7: see original paper] shows detailed experimental results in the KNIME platform, where positive citations (POS), negative citations (NEG), and neutral citations (NEU) are displayed in red, green, and yellow respectively.

4.3 Citation Sentiment Visualization The proposed citation sentiment visualization analysis adds sentiment markers to traditional citation network visualization methods to construct citation network maps that identify citation sentiment, thereby effectively discovering the distribution of different citation sentiments across the entire dataset.

The specific steps for constructing citation sentiment visualization maps are:

- (1) Remove nodes with zero citations (remaining: 1,127 nodes, 1,191 edges) to obtain the initial anti-aging citation network (see Figure 8 [Figure 8: see original paper]). In Figure 8, blue nodes represent papers, node size is proportional to citation count, and node labels are paper PMIDs (PubMed's unique ID for each paper). For visualization optimization, only the top 20 most-cited paper IDs are displayed. Yellow directed lines indicate citation direction ($A \rightarrow B$ means B cites A).
- (2) Based on citation sentiment identification results, add sentiment markers to the initial citation network to obtain the anti-aging citation network with sentiment labels (see Figure 9 [Figure 9: see original paper]). Red directed lines indicate positive citations, green indicate negative citations, and yellow indicate neutral citations.

Compared with traditional citation networks, Figure 9 can effectively identify the distribution of different citation sentiments across the entire dataset. It enables further analysis of positive and negative citation chains, providing support for citation content analysis-based bibliometrics. For example, node 10380075 has a negative citation (10380075 $\rightarrow$ 22500228, highlighted by purple box in Figure 8). To verify the accuracy of this sentiment identification result and the effectiveness of the visualization map, we retrieved papers with PMID=22500228 and PMID=10380075 from PubMed: “Insulin in central nervous system: more than just a peripheral hormone” and “Phosphatidylinositol 3-kinase-mediated regulation of neuronal apoptosis and necrosis by insulin and IGF-I.” The citation content was located: “In contrast, Ryu et al. [160] failed to show protection by

IGF-1 against excitotoxic or oxidative stress-induced necrosis, despite a decrement in neuronal apoptosis.”

Manual reading of this citation content (“In contrast, despite reduced neuronal apoptosis, Ryu et al. failed to demonstrate that IGF-1 can protect against excitotoxicity or oxidative stress-induced necrosis”) reveals it is a negative citation, using Ryu et al.’s research to contrast and highlight the value and innovation of the authors’ discovery about insulin and IGF-I’s role in weakening retinal and brain neuronal apoptosis under conditions such as oxidative stress. Therefore, this citation sentiment identification result is accurate, and the citation sentiment visualization analysis is feasible and effective.

References

- [1] Zhao Rongying, Wang Jianpin. A comparative study of citation content analysis and citation bibliographic analysis [J]. Library and Information Service, 2017, 61(10): 110-115.
- [2] Office of National Science and Technology Awards. Introduction to National Natural Science Award [EB/OL]. [2017-07-27]. <http://www.nosta.gov.cn/web/detail.aspx?menuID=158&contentID=158>
- [3] Hu Zhigang, Chen Chaomei, Liu Zeyuan, et al. Design and implementation of a citation analysis system based on XML full-text data [J]. New Technology of Library and Information Service, 2012(11): 72-77.
- [4] MORAVCSIK M J, MURUGESAN P. Some results on the function and quality of citations [J]. Social studies of science, 1975, 5(1): 86-92.
- [5] SMALL H G. Cited documents as concept symbols [J]. Social studies of science, 1978, 8(3): 327-340.
- [6] SMALL H G, GREENLEE E. Citation context analysis of a co-citation cluster: recombinant-DNA [J]. Scientometrics, 1980, 2(4): 277-301.
- [7] DING Y. Content-based citation analysis: the next generation in citation analysis [EB/OL]. [2012-09-26]. <http://www.lis.illinois.edu/events/2012/09/26/content-based-citation-analysis-next-generation-citation-analysis>.
- [8] Zhu Qingsong, Leng Fuhai. Research on topic identification of highly cited papers based on citation content analysis [J]. Journal of Library Science in China, 2014, 40(1): 30-49.
- [9] Lu Wei, Meng Rui, Liu Xingbang. Research on citation content annotation framework for citation relationships [J]. Journal of Library Science in China, 2014, 40(6): 93-104.
- [10] Zhao Rongying, Zeng Xianqin, Chen Bikun. Full-text citation analysis: a new development in citation analysis [J]. Library and Information Service, 2014, 58(9): 129-135.

- [11] Zhao Rongying, Guo Fengjiao, Zeng Xianqin. Empirical research on location-based co-citation analysis [J]. Journal of the China Society for Scientific and Technical Information, 2016, 35(5): 492-500.
- [12] GARFIELD E. Can citation indexing be automated? [J]. Essays of an information scientist, 1962, 1: 84-90.
- [13] CANO V. Citation behavior: classification, utility, and location [J]. Journal of the American Society for Information Science, 1989, 40(4): 284-290.
- [14] LIU M X. Progress in documentation: the complexities of citation practice: a review of citation studies [J]. Journal of documentation, 1993, 49(4): 370-408.
- [15] CASE D O, HIGGINS G M. How can we investigate citation behavior? A study of reasons for citing literature in communication [J]. Journal of the American Society for Information Science, 2000, 51(7): 635-645.
- [16] KESSLER M M. Bibliographic coupling between scientific papers [J]. American documentation wiley online library, 1963, 14(1): 10-25.
- [17] SMALL H G. Co-citation in the scientific literature: a new measure of the relationship between two documents [J]. Journal of the American Society for Information Science, 1973, 24(4): 265-269.
- [18] CHUBIN D E, MOITRA S D. Content analysis of references: adjunct or alternative to citation counting? [J]. Social studies of science, 1975, 5(4): 423-441.
- [19] TEUFEL S, SIDDHARTHAN A, DAN T. Automatic classification of citation function [C]//Conference on empirical methods in natural language processing, 2006, 14(1): 103-110.
- [20] Liu Shengbin, Ding Kun, Zhang Chunbo. Research on citation evaluation based on citation content nature [J]. Information Studies: Theory & Application, 2015, 38(3): 77-81.
- [21] SHADISH W R, TOLLIVER D, GRAY M, et al. Author judgments about works they cite: three studies from psychology journals [J]. Social studies of science, 1995, 25(3): 477-498.
- [22] Verlic M, Stiglic G, Kocbek S, et al. Sentiment in science: a case study of CBM contributions in years 2003 to 2007 [C]//Computer-based medical systems, 2008 CBMS'08 21st IEEE international symposium on. Albuquerque, Jyväskylä, Finland: IEEE, 2008: 138-143.
- [23] SMALL H. Interpreting maps of science using citation context sentiments: a preliminary investigation [J]. Scientometrics, 2011, 87(2): 373-388.
- [24] BONACICH P B. Factoring and weighting approaches to status scores and clique identification [J]. Journal of mathematical sociology, 1972, 2(1): 113-120.
- [25] TOUTANOVA K, MANNING C D. Enriching the knowledge sources used in a maximum entropy part-of-speech tagger [C]//Joint sigdat conference on

empirical methods in natural language processing. Hong Kong: ACM, 2000, 25(6): 63-70.

[26] JACOMY M, BASTIAN M, HEYMANN S. Gephi: an open source software for exploring and manipulating networks [C]//International conference on weblogs & social media. San Jose: The AAAI Press, 2009: 361-362.

Author Contributions:

Liao Junhua: Research proposition and conceptual design

Liu Ziqiang: Data analysis and paper writing

Bai Rujiang: Framework formulation and data collection

Chen Junying: Data collection and preprocessing

Note: Figure translations are in progress. See original paper for figures.

Source: ChinaXiv — Machine translation. Verify with original.