

Automatic Generation of Information Briefs Integrating Selective Attention Decay Models: A Case Study of UNESCO Science and Technology Reports

Authors: Tian Wenbo, Song Peiyan, Wu Keying, Chaohui Feng, Song Peiyan

Date: 2022-09-01T00:00:00+00:00

Abstract

Purpose/Significance: Briefings constitute important intelligence products. UNESCO publishes a large volume of high-value professional literature. To meet users' demands for international professional knowledge, there is a need to rapidly generate information briefings and enhance intelligence service capabilities. **[Method/Process:** Based on the "selective attention attenuation" theory in cognitive science, this paper treats the generation of information briefings as a simulation of human cognitive information processing, and explores methods for implementing cross-lingual automatic summarization. First, grounded on the "attenuator" model in selective attention, an integrated design is developed across three levels—topic, topic sentence, and briefing—according to cognitive load capacity. Subsequently, KeyBERT and Transformer algorithms are employed to extract topic keywords and generate summaries from scientific and technical reports published by UNESCO, thereby achieving rapid generation of information briefings, with evaluation conducted using information entropy and ROUGE metrics. **[Results/Conclusion:** Experiments demonstrate certain advantages in information entropy and ROUGE-2, ROUGE-L values, indicating that the selective attention attenuation model can improve summarization effectiveness and cover core information in the text. The study further reveals that the close integration of cognitive science and computational models plays a significant role in enhancing the interpretability and scientific rigor of information briefings, and helps establish a computable and interpretable information briefing generation and knowledge service model.

Full Text

A Study on Automatic Information Briefing Generation Method Integrating Selective Attention Decay Model: The Case of UNESCO Science and Technology Reports

Tian Wenbo¹, Song Peiyan^{1*}, Wu Keying¹², Feng Chaohui^{1}

¹School of Business and Management, Jilin University, Changchun 130015, China

²School of Information Resource Management, Renmin University of China, Beijing 100872, China

Abstract

[Purpose/Significance] Information briefings are critical intelligence products, with abstracts and thematic keyword collections serving as essential components. UNESCO publishes a substantial volume of high-value scientific and technical reports. To meet users' demands for international expertise, there is a pressing need to rapidly generate information briefings and enhance intelligence service capabilities. **[Method/Process]** This paper investigates automatic abstract and thematic keyword generation methods by treating information briefing generation as a simulation of human cognitive information processing, grounded in the "Selective Attention Decay" theoretical model from cognitive science. First, supported by the "attenuator" model from selective attention theory, we designed an integrated framework across three hierarchical levels: abstracts, keywords, and briefings. Subsequently, we employed KeyBERT and Transformer algorithms to extract thematic keywords and generate abstracts from UNESCO science and technology reports, creating referenceable briefing intelligence products. The generated results were then evaluated using information entropy and ROUGE metrics. **[Results/Conclusions]** Experimental results evaluated via ROUGE-2 and ROUGE-L values demonstrate that the selective attention decay model can improve summarization effectiveness and cover core textual information. Further analysis from an information entropy perspective confirms that the automatically generated abstracts align with fundamental human cognitive levels. The study also reveals that close integration of cognitive science and computational models significantly enhances the interpretability and scientific rigor of information briefings, facilitating the development of a computable and interpretable model for information briefing generation and knowledge services.

Keywords: Knowledge Discovery; Selective Attention; Text Summarization; Topic Extraction

Classification Number: G353

1 Introduction

Information explosion has triggered information overload, challenging users' efficiency in acquiring information. Automatically summarizing massive information and delivering it to relevant institutions or users in the form of thematic briefings helps users grasp cutting-edge topics, stay informed of developments, and enable scientific decision-making. Reports issued by international organizations possess high authority and reference value, necessitating rapid monitoring and dynamic tracking to generate high-quality information briefings for relevant institutions—an endeavor with significant practical application value.

From a cognitive science perspective, information briefing generation is essentially a process of human information reprocessing. When reading texts, the human brain typically acquires limited key information within specific temporal and environmental constraints, allocating finite processing capacity to focal information—this is the selective attention mechanism. Humans can rationally allocate and fully utilize limited attentional resources to quickly and accurately filter high-value information from large volumes while discarding low-value information, achieving efficient information processing through a “principle of least effort.” This represents a survival mechanism formed through long-term evolution. The cognitive science community has proposed theoretical models such as the “filter model” and validated them through eye-tracking experiments, further confirming the important theoretical value of cognitive models for implementing summarization technology. Therefore, this paper takes information briefing generation as an application scenario, uses UNESCO data as an example, closely integrates the selective attention model with computational algorithms to simulate the intrinsic laws of human information processing, proposes and validates its role in information briefing generation, and provides strong support for improving the scientific rigor and interpretability of automatic summarization.

2 Related Research

2.1 Attention Models

Attention mechanisms have deeply integrated with deep learning technologies in recent years, advancing rapidly. Zeng Ziming et al. constructed a U-BiLSTM sentiment analysis model based on user attention mechanisms to analyze emotional evolution processes, demonstrating strong interpretability and accuracy while improving F1 scores and accuracy rates [1]. Zhou Ying et al. combined Long Short-Term Memory (LSTM) with attention mechanism models, using Huawei P10 flash memory gate incident as a case study to prove that sentiment analysis based on selective attention mechanisms can improve sentiment classification success rates and accurately extract emotional features, showing good application potential in short text analysis [2]. Hu Jiming et al. utilized the TextRank algorithm and CNN-BiLSTM-Attention ensemble model for policy text classification, enhancing classification efficiency and accuracy [3]. Research in this field primarily focuses on combining attention mechanisms with sentiment

analysis, text classification, and public opinion monitoring, which also provides insights for text summarization. Notably, Transformer offers certain advantages in computational efficiency and parallel processing, providing favorable technical conditions and methods for automatic summarization. A noteworthy new trend is that the combination of cognitive models and computational methods often yields more substantial scientific progress than pure algorithm optimization and parameter design. Future research should strengthen substantive integration with cognitive science to form more original research achievements.

2.2 Text Summarization

Existing research primarily employs machine learning algorithms, yielding numerous achievements in both computer science and information technology fields. Li Wei et al. proposed a Tibetan extractive summarization method combining TextRank algorithm with word embeddings, mapping each word in a sentence to high-dimensional word libraries to form sentence vectors for iteration, evaluating sentences and selecting high-scoring ones as summaries, thereby effectively improving summary quality [4]. Zhang Chengzhi et al. designed a book review summarization method based on fine-grained comment mining, providing multi-dimensional and fine-grained evaluations for book information [5]. Wang Xiaoyu improved traditional graph-based methods by enhancing text graph construction and word weighting approaches, enabling algorithms to generate semantic graphs with multiple attributes based on dependency relationships between sentence words, and proposed word weight calculation methods integrating keyword position information, conceptual hierarchy, and connection strength for importance ranking, selecting high-scoring nodes as keyword collections with certain improvements [6]. From a user perspective, existing summaries still have limitations in semantic coherence and readability. Generating topic-oriented new sentences by vectorizing words and sentences with different cognitive loads from a cognitive computing perspective holds promise for forming new technologies with greater cognitive explanatory power and accuracy.

2.3 Topic Extraction

Topics are essential components of information briefings. Currently, domestic topic extraction primarily employs topic models represented by Latent Dirichlet Allocation (LDA), using unsupervised learning to conduct semantic structure and clustering analysis on full texts to extract valuable topics and topic keyword distributions. For instance, Shi Jing et al. conducted text analysis based on LDA modeling for corpora and texts, with results significantly outperforming other methods [7]. Qu Jingye et al. applied LDA topic models to 梳理 (map out) the evolution of domestic information service research themes over the past 22 years, providing references and guidance for sustainable development in this field [8]. Common deep learning methods based on language models include NNLM, word2vec, Elmo, GPT, BERT, etc., among which BERT is a deep bidirectional Transformer pre-trained language model that represents

the culmination of embedding technologies such as NNLM, Word2vec, ELMO, and GPT. Li Songfan et al. proposed a BERT-based method for identifying cutting-edge research topics, achieving recognition of frontier research topics in the agricultural field [9]. How to integrate topic extraction with summarization generation through unified design requires further investigation.

Overall, academic research on briefing generation has yielded rich achievements, yet two prominent issues remain: first, there is abundant technical algorithm research but insufficient cognitive research, resulting in slightly inadequate theoretical explanatory power for summarization at a deep level; second, summarization and keyword extraction are often treated separately, when their intrinsic relationship should be processed integrally to improve briefing generation efficiency and consistency. Therefore, this paper introduces the selective attention “decay” model, combines it with KeyBERT and Transformer algorithms, proposes a feasible cognitive computing basis and implementation scheme for information briefings, and finally conducts quantitative evaluation and validation from two perspectives: information entropy and ROUGE values.

3 Theoretical Basis and Overall Framework Design

3.1 Selective Attention Decay Model and Cognitive Load Factors

The mechanism of information briefing generation can draw valuable insights from cognitive science research on “attention.” Attention is the ability to focus on specific stimuli. In most cases, the focusing characteristic of attention is associated with “selective attention” —the human capacity to focus or allocate attention to specific locations, subjects, or information. Cognitive psychologist Broadbent proposed the “Filter Theory,” suggesting that humans can filter out some information to allow other information to receive higher-level processing, thereby avoiding “information overload.” Treisman’s “Attenuator Theory” improves upon the filter model, positing that unattended information is not truly filtered but merely processed differently. Cognitive experiments further demonstrate that physical characteristics, language, and semantics can all capture attention and be used to absorb information. The model is shown in Figure 1 [Figure 1: see original paper]. Through the attenuator, various types of information have opportunities to be processed or filtered, ultimately forming memory. Cognitive load refers to the amount of cognitive resources required when performing a cognitive task. Familiar or simple tasks entail low cognitive load, while unfamiliar or difficult tasks require allocating more cognitive resources and thus have relatively higher cognitive complexity.

Figure 1. Theoretical Diagram of Selective Attention Decay Model

The selective attention decay model provides important theoretical support for automatic information briefing generation, manifested in three aspects: (1) The summarization process is essentially a process of humans allocating cognitive resources to filter information. Linguistic features and semantics are core bases for summarization. Beyond traditional summarization extraction based on exter-

nal features such as word frequency and position, semantics-centered extraction should be employed. (2) As semantic carriers, words and sentences impose different cognitive loads on humans, and their roles in information briefing generation should be considered comprehensively. Keywords mostly belong to low-load processing, while abstracts involve high cognitive load. Keywords serve as compatible-side interference items for abstracts, facilitating high-level processing and feature integration of information with lower cognitive complexity. (3) Information summarization often employs multiple parallel task execution to generate multiple parallel summary texts, yet obtaining information from unattended tasks remains possible. The ability to allocate attention can be trained through information volume or semantic consistency, gradually achieving automated processing. The framework designed in this paper was also developed under the guidance of this theory and implemented and validated using algorithms.

3.2 Framework Flowchart for Automatic Text Summarization

Based on the selective attention decay model theory described in Section 3.1, this paper designed the main framework for information briefing generation, as shown in Figure 2 [Figure 2: see original paper].

Figure 2. Automatic Summarization Framework Flowchart

The framework comprises three modules: data acquisition, text summarization and topic extraction, and briefing generation.

- (1) **Text Extraction:** Text extraction is the preparatory condition for briefing generation. First, information sources are selected according to requirements, target texts are identified, and information elements such as charts and formulas are removed. Second, the jieba Chinese word segmentation library is imported to preprocess texts by removing modal particles, punctuation marks, etc.
- (2) **Text Summarization and Topic Crawling:** Before computation, the input corpus is vectorized by importing the word2vec library. After generating new corpus through Transformer computation, evaluation metrics including ROUGE values and information entropy are applied. The results are further processed to form text summaries. After corpus vectorization, the KeyBERT algorithm library is called for keyword extraction to form keyword sets, which are then filtered and sorted to generate thematic keywords. Text summaries and thematic keywords constitute the main content of the briefing.

3.3 Computational Methods Based on Selective Attention Decay Model

This theory emphasizes that during text computation, computers can actively filter the importance of input text content, allocating more attention to more

important texts and screening out interfering low-load information. Therefore, based on this theory, this paper designed key methods for briefing generation using open-source algorithms.

3.3.1 Keyword Extraction Based on KeyBERT Thematic keywords are fundamental components of information briefings. They are words extracted from keywords that can reveal text structure and semantics, helping users grasp important information points, achieve overall understanding of main text content, and serve as clues for readers to comprehend summaries. KeyBERT is a compact and easy-to-use keyword extraction technique whose core remains based on selective attenuation attention. After vectorizing texts, it conducts weight calculation guided by semantic computation, using BERT embeddings and simple cosine similarity to create keywords or phrases most relevant to the document.

In this paper, we first utilize BERT to calculate document embedding values to obtain document-level vector representations. Then, word vectors are extracted for n-grams. Finally, cosine similarity is employed to identify keywords or key phrases most similar to the document, yielding keywords that best describe the entire document. Based on these keywords, further screening is conducted in conjunction with original text themes to remove weakly thematic expressions. The filtered results effectively represent text themes and match generative summarization results, providing readers with multi-dimensional key information.

3.3.2 Summary Sentence Generation Based on Transformer Model

The fundamental principle of Transformer involves adding weighted calculations to the encoder-decoder model to represent influence weights of each word and sentence. It constructs representations of relationships between sentences through attention training on the text itself [13]. Each Transformer layer contains a multi-head attention mechanism and a feed-forward neural network, as shown in Figure 3 [Figure 3: see original paper].

Figure 3. Transformer Model Diagram

In this paper, input texts are first converted into word vectors, then passed through encoder-decoder layers for attention scores calculation to construct sentence weight relationships, and finally output as summaries after softmax computation. Compared with previous research on extractive text summarization or deep learning summarization based on LSTM, Transformer not only increases parallel computational efficiency through repeated self-attention calculations but also improves computational performance and quality.

3.4 Briefing Quality Evaluation Metrics

3.4.1 Information Entropy for Briefing Measurement Information entropy measures the amount of information in a system. From a textual information perspective, higher entropy values indicate higher information content

in a text segment. Its applications are broad, including verifying text capacity required for expressing the same meaning across different languages and measuring information content in specific texts. The mathematical formula for information entropy can be expressed as [10]:

$$H(X) = - \sum_{i=1}^n p(x_i) \log_2 p(x_i)$$

where x represents a random variable with values $(x_1, x_2, \dots, x_n)$, and $p(x_i)$ denotes the probability of event x_i occurring. Generally, higher probability events contain less information, with $\sum p(x_i) = 1$, meaning the sum of probabilities for all random events equals 1. When applied to textual information calculation, this represents the probability of random information contained in a text segment. Considering that semantic computation involves numerous vocabulary items, we selected the bi-gram model for text information entropy calculation [11].

Using information entropy to judge text summary information content enables quantitative description of generative text summarization quality from an information volume dimension. This experiment constructed an information entropy calculation process using open-source code provided by the Python language. Generated summaries underwent word segmentation before text vectorization representation using the word2vec method, with final results calculated and output through the information entropy formula.

Literature review reveals that information entropy values vary across different languages (e.g., Chinese, English, French, Russian) and expression methods. Even within the same language, information entropy differs among different cognitive groups (e.g., scholars, children) describing the same content. Therefore, by calculating the information entropy value of summary results and comparing it with standard Chinese text information entropy metrics, we can verify whether the cognitive level of generative summarization results meets established standards.

3.4.2 ROUGE Values for Briefing Quality Evaluation ROUGE (Recall-Oriented Understudy for Gisting Evaluation) is a commonly used metric for evaluating text summarization quality. It judges summary quality by counting co-occurring words or sentences between reference texts and experimental texts. Calculating this metric requires manually generated expert summaries as standard reference sets for comparison with machine-generated summaries. The formula is expressed as:

$$ROUGE-N = \frac{\sum_{S \in \{\text{Reference Summaries}\}} \sum_{\text{gram}_n \in S} \text{Count}_{\text{match}}(\text{gram}_n)}{\sum_{S \in \{\text{Reference Summaries}\}} \sum_{\text{gram}_n \in S} \text{Count}(\text{gram}_n)}$$

ROUGE-N, where N represents N-gram, uses the number of N-grams in reference texts as the denominator and the number of N-grams in summary results as the numerator. ROUGE-1 and ROUGE-2 are typically used as evaluation metrics.

In addition to these two indicators, to enhance evaluation persuasiveness, this paper also references the longest common subsequence calculation method, namely ROUGE-L, with the calculation formula:

$$R_{\text{lcs}} = \frac{LCS(C, S)}{\text{len}(C)}, \quad P_{\text{lcs}} = \frac{LCS(C, S)}{\text{len}(S)}$$

$$F_{\text{lcs}} = \frac{(1 + \beta^2)R_{\text{lcs}}P_{\text{lcs}}}{R_{\text{lcs}} + \beta^2 P_{\text{lcs}}}$$

where C and S represent reference texts and experimental texts respectively, $LCS(C, S)$ denotes the length of the longest common subsequence between texts C and S , $\text{len}(S)$ and $\text{len}(C)$ represent the lengths of the two texts, R_{lcs} denotes recall rate, P_{lcs} denotes precision rate, and F_{lcs} is ROUGE-L. β represents a very large number; thus, after formula derivation, F_{lcs} is found to be approximately equal to R_{lcs} .

Since the summaries generated in this paper are primarily short, single-document summaries, we adopted ROUGE-2 and ROUGE-L as evaluation metrics.

4 Empirical Research

This study selected 50 education-related news articles or reports from the UNESCO website as experimental texts. We used open-source tools to crawl target texts, stored them in a database, and performed simple manual cleaning and pre-processing by removing modal particles, conjunctions, charts, complex formulas, and numbers, followed by text vectorization to facilitate subsequent computational processing. Finally, the text data underwent iterative computation using the constructed model to generate abstractive text summarization results.

4.1 Data Collection

Research data were sourced from the United Nations Educational, Scientific and Cultural Organization (UNESCO), with education reports selected as experimental texts for their readability and reference value. Following the flowchart shown in Section 3.1, we used Python's selenium library to select 50 representative UNESCO reports for preliminary manual preprocessing and database storage. Subsequently, the textual data in the database underwent word segmentation and stop-word removal to complete preprocessing.

4.2 Automatic Text Summarization Generation

This section validates the proposed framework through empirical testing. First, preprocessed texts were imported into Python for word segmentation using the jieba library, then converted into vector representations via word2vec for computational convenience. Guided by selective attention decay model theory, we utilized open-source tools to build Transformer [13] and KeyBERT algorithm frameworks. After preprocessing collected data into machine-recognizable word vectors, computations were performed through existing frameworks to generate text summaries and keyword sets separately. Building upon this, we established quantifiable information entropy and ROUGE values for evaluating text summarization results. Based on the keyword sets, keywords surrounding the text theme were extracted as thematic keywords and sorted according to text sequence, ultimately forming standardized information briefings. Partial experimental results are shown in Table 1.

Table 1. Partial Abstractive Text Summarization Results

Text 1: Abstract

UNESCO has launched a global “Happy Schools Initiative.” Schools should become venues that support social cohesion and create communities across differences. Schools should also cultivate students’ lifelong love of learning through happiness and engagement, rather than hindering learning by placing academic performance above all else, thereby damaging individual well-being. The Happy Schools Initiative is going global, with UNESCO launching it worldwide. It aims to help students and educators globally, initiated by UNESCO to assist children and educators. It will become a global model for schools and educators worldwide—an international model for international school education to help young people and educators across continents and countries, launched by the Global Education Research Institute for Asian Education, the Asian International Education Council, and the European Council of the World Education Council.

Text 1: Topics

UNESCO, worldwide, young people, happiness crisis, productivity, pressure, refugees, schools, improvement, students, cohesion, equality.

Text 2: Abstract

In 2017, globally, women accounted for less than one-quarter of those studying engineering, manufacturing and construction, or information and communication technologies (ICT) in over two-thirds of countries. In almost all education systems (87%), boys more frequently than girls answered that they wanted jobs involving mathematics. Boys’ and girls’ aspirations for mathematics-related work are closely related to their confidence in their abilities in the subject. This suggests that addressing girls’ confidence in science and mathematics should remain a concern for policymakers. The International Association for the Evaluation of Educational Achievement Compass points out that this may lead to fewer high-achieving girls entering STEM higher education fields. This education series briefing, published in IEA’s TIMSS 2019 data special issue with a

sample of 250,000 students, shows that more boys than girls in grade 8 want to pursue mathematics or science-related careers.

Text 2: Topics

UNESCO, countries, scientific research, statistics, one-quarter, girls, achievement, success, analysis, confidence, differences.

Text 3: Abstract

In Southeast Asia, if a child's temperature is 2 standard deviations above average, they are expected to lose 1.5 years of schooling. A 1.8°C increase in average annual temperature reduces academic performance by 1%, as do 6 days exceeding 32.2°C. Assuming optimal temperatures below 22°C, reducing classroom temperature from 30°C to 20°C would improve test scores by an average of 20%. Polluted air significantly reduces cognitive ability, possibly irreversibly due to suspected neurotoxicity. In Israel, even temporary exposure to dust reduces students' test scores and post-secondary education attainment. A cohort study of nearly 3,000 children in Barcelona, Spain found that after adjusting for socioeconomic status, children exposed to high pollution levels in Spain showed lower cognitive development growth than peers in less polluted schools.

Text 3: Topics

United Arab Emirates, high temperature, high pollution, average temperature, causes, test scores, secondary school, years, reduction, cognition, development.

The computational framework based on selective attention decay theory advances text summarization and topic extraction to semantic understanding levels, reorganizing texts to generate new corpus [14], resulting in natural summaries with strong explanatory power and more fluent textual coherence compared to extractive summarization.

By further processing results and calculating information entropy of text summaries, we evaluated the information content of summaries from 10 documents. Results are shown in Figure 4 [Figure 4: see original paper], where the horizontal axis represents ten calculation results, the vertical axis represents information entropy, and the top of each bar displays the information entropy value for each summary result.

Figure 4. Summary Quality Evaluation: Information Entropy

Through random entropy value calculation and statistics for ten summary results, the average information entropy of sampled summaries reached 6.852. Literature review and data investigation on Chinese text information entropy experiments confirm that this value meets Chinese text information entropy standards and aligns with human cognitive levels.

Additionally, this experiment further validated summary quality through ROUGE values. Verifying recall rates requires appropriate reference summaries. To enhance research objectivity, we manually generated summaries for ten extracted texts as reference answers to facilitate recall rate calculation.

Results are shown in Figure 5 [Figure 5: see original paper]. Referring to the information entropy formula in Section 3.4.2, this line chart displays word co-occurrence rates between manually generated reference summaries and experimentally obtained summaries. The trend shows that ROUGE-2 values are generally higher than ROUGE-L values because ROUGE-2 considers consecutive two-word vectors while ROUGE-L depends on consecutive multiple-word vectors.

Figure 5. Summary Quality Evaluation: ROUGE Values

Using ROUGE-2 and ROUGE-L as evaluation metrics, the standard assesses text recurrence rates between generative text summaries and reference human summaries. The average values for our sampled summaries are 0.432 and 0.367 for ROUGE-2 and ROUGE-L respectively, which are basically consistent with multiple studies such as “Research on Extractive Text Summarization Model Based on Maximum Marginal Relevance” [12], but with higher readability in generative summaries. The underlying reason is that generative text summarization guided by selective attention decay theory essentially achieves high cognitive load information processing effects through deep semantic computation and reorganizes these texts to generate concise, natural sentences.

4.3 Briefing Generation

Briefings consist of two components: thematic keywords and abstracts. Thematic keywords help readers quickly understand key information points of articles. The information provided by thematic keywords is relatively divergent, belonging to low cognitive load information processing, allowing readers to grasp main article threads through keywords. Abstracts are concentrated versions of article information containing specific details, enabling readers to understand article gist through reading [15]. Information processing through these two dimensions—summary generation and topic extraction—can adequately describe textual information.

Thematic keywords and text abstracts, as important components of briefing products, complement each other with strong connections in content and thematic grasp: abstracts are expanded and deepened versions of thematic keywords, while thematic keywords are core information of abstracts. Text abstracts and thematic keywords are computed in parallel and cross-referenced, with several thematic keywords capable of mining text key points to jointly constitute information briefings.

This paper integrates selective attention decay theory from cognitive science with information briefing generation methods, designing an information briefing generation framework that fuses keyword extraction and abstractive summarization. Using UNESCO data as an example and employing Transformer and KeyBERT methods for empirical research, experimental results demonstrate that the selective attention decay model can simultaneously consider keywords, sentences, and discourse, aligning with cognitive load levels of information pro-

cessing and exhibiting strong explanatory power and scientific rigor. Technically, evaluation using information entropy and ROUGE values generates information briefings with different granularities according to different cognitive load levels, offering certain application value for intelligence monitoring and decision support. Future research should further subdivide and explore influencing factors and mechanisms of selective attention to facilitate organic integration of cognitive science and intelligence technology, improving interpretability and accuracy of intelligence products—a direction worthy of exploration.

References

- [1] Zeng Ziming, Sun Jingjing. Research on Public Opinion Emotional Evolution in Public Health Emergencies Based on User Attention: The Case of COVID-19[J]. *Information Science*, 2021, 39(09): 11-17.
- [2] Zhou Ying, Liu Yue, Cai Jun. Microblog Sentiment Analysis Based on Attention Mechanism[J]. *Information Studies: Theory & Application*, 2018, 41(03): 89-94.
- [3] Hu Jiming, Fu Wenlin, Qian Wei, Tian Peilin. A Policy Text Classification Model Integrating Topic Model and Attention Mechanism[J]. *Information Studies: Theory & Application*, 2021, 44(07): 159-165.
- [4] Zhang Chengzhi, Tong Tiantian, Zhou Qingqing. Research on Automatic Book Review Summarization Based on Fine-Grained Review Mining[J]. *Journal of the China Society for Scientific and Technical Information*, 2021, 40(02): 163-172.
- [5] Wang Xiaoyu, Wang Fang. Keyword Extraction Algorithm for Paper Abstracts Based on Semantic Text Graphs[J]. *Journal of the China Society for Scientific and Technical Information*, 2021, 40(08): 854-868.
- [6] Shi Jing, Fan Meng, Li Wanlong. Topic Analysis Based on LDA Model[J]. *Acta Automatica Sinica*, 2009, 35(12): 1586-1592.
- [7] Qu Jingye, Chen Zhen, Hu Yinan. Comparative Study of Co-Word Analysis and LDA Topic Model in Text Theme Mining[J]. *Information Science*, 2018, 36(02): 18-23.
- [8] Li Songfan, Huang Yong, Yang Jinqing. Research on BERT-Based Frontier Research Topic Identification Method in Agriculture[J]. *Technology Intelligence Engineering*, 2021, 7(05): 100-114.
- [9] Liu Xiaojian, Xie Fei, Wu Xindong. Keyword Extraction Algorithm Based on Graph and LDA Topic Model[J]. *Journal of the China Society for Scientific and Technical Information*, 2016, 35(06): 664-672.
- [10] Ma Zheng. Measuring Academic Journals' Contribution to Information Entropy Changes in Knowledge Systems Based on "Counterfactual" Thinking[J].

Journal of the China Society for Scientific and Technical Information, 2022, 41(07): 745-761.

[11] Cheng Huiping, Cheng Yuqing. Personal Cloud Storage Security Risk Assessment Based on AHP and Information Entropy[J]. Information Science, 2018, 36(07): 145-151.

[12] Yu Chuanming, Guo Yajing, Zhu Xingyu, An Lu. Research on Extractive Text Summarization Model Based on Maximum Marginal Relevance[J]. Information Science, 2021, 39(02): 34-43.

[13] Vaswani A, Shazeer N, Parmar N, et al. Attention is all you need[EOL]. [2022-08-01] <https://dl.acm.org/doi/10.5555/3295222.3295349>.

[14] Niu Z, Zhong G, Yu H. A review on the attention mechanism of deep learning[J]. Neurocomputing, 2021, 452: 48-62.

[15] Liu Y, Lapata M. Text summarization with pretrained encoders[J]. arXiv:1908.08345, 2019.

Author Contributions

Author 1: Tian Wenbo—proposed research ideas, designed research methodology, wrote and revised the paper.

Author 2: Wu Keying—designed and implemented experiments, revised and improved the paper.

Author 3: Feng Chaohui—conducted experimental data analysis, revised and improved the paper.

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

Source: ChinaXiv—Machine translation. Verify with original.