

Bilingual Summarization and Readability Optimization for Textile-Industry News Using the BART Model

Yuntao Li¹, Xueqin Li²

¹School of Foreign Languages, Mianyang Teachers' College, Mianyang 621006, Sichuan, China

²Mianyang Nanshan Middle School, Mianyang 621000, Sichuan, China

Corresponding author: Yuntao Li (e-mail: yt002025@163.com)

ABSTRACT To address translation inconsistencies, information omission, and poor readability in bilingual summarization of textile-industry news, this study proposes a Domain-Enhanced Bidirectional and Autoregressive Transformers (DEBART) framework based on the BART architecture. A high-quality Chinese-English parallel corpus is constructed, and domain-specific terminology knowledge is incorporated into the embedding layer through a gated-fusion mechanism to improve the accuracy and consistency of professional term translation. In addition, a coverage mechanism is introduced to suppress redundant content and enhance semantic completeness, while an FKGL-driven readability optimization module dynamically adjusts sentence structures and lexical complexity to improve fluency and reader accessibility. Experimental results demonstrate that the proposed model achieves ROUGE-1, ROUGE-2, and ROUGE-L scores of 46.25, 25.76, and 42.84, respectively, outperforming conventional sequence generation methods in both summarization quality and readability. The proposed framework provides an effective solution for bilingual information processing in the textile industry and offers valuable reference for intelligent multilingual information transmission, semantic communication, and knowledge dissemination in electromagnetic-enabled communication environments and future antenna-assisted information service systems, where accurate and efficient cross-language content delivery is increasingly important.

INDEX TERMS bilingual summarization, domain-enhanced bart, readability optimization, semantic communication, electromagnetic information transmission

I. INTRODUCTION

IN the context of globalization, the textile industry—an important sector of production and trade—has a growing demand for effective information dissemination. However, when disseminating bilingual news summaries in the textile industry, the differences in expression and professional terminology between Chinese and English make direct translations prone to missing information or misunderstandings, which affect the quality and accuracy of the summaries [1-3]. At the same time, existing methods have limitations in readability measurement, which reduce their practical effectiveness. These problems not only hinder the accurate transmission of information but also reduce the translation quality and readers' reading experience. In order to meet these challenges [4-6], we investigated bilingual summarization for textile-industry news based on BART and employ a readability optimization module to improve the efficiency of information communication. The DEBART model is based on the BART architecture and is domain-enhanced to better adapt to the characteristics of the textile industry. We introduce bilingual terminology dictionaries, integrate

textile-industry knowledge into the embedding layer, and apply a gated-fusion mechanism to enhance term recognition and translation, thereby mitigating term-meaning distortion. We design a readability enhancement module with FKGL as the core metric and integrate a multi-dimensional language optimization strategy.

The innovation value of this research is significant: (i) to build a high-quality bilingual corpus to support a domain-adaptive summarization model; (ii) to build a DEBART model to improve the accuracy and consistency of terminology translation and transmission; (iii) to introduce a coverage mechanism to improve information density and logical coherence; (iv) to build a readability optimization module to deeply integrate language quality control into the generation process ensuring accurate content and fluent language, enabling readers from different backgrounds to understand it easily.

II. RELATED WORK

In the field of bilingual news summarization research, with the continuous progress of natural language processing

technology, the topic has attracted extensive attention. In the research of text abstract technology, different scholars have proposed a variety of models and methods for different scenarios and problems. Pei et al. proposed a judicial text summary generation model based on the knowledge enhancement pre-training model, integrating judicial knowledge and combining comparative learning to improve the quality of the summary [7]; Wijayanti et al. investigated the application of Vecmap and BiVec bilingual word embedding methods in Indonesian–English [8]; Ning et al. proposed a keyword-based Chinese news summarization method using the seq2seq (Sequence-to-Sequence Learning) model that is prone to generating semantically irrelevant words in the abstract [9]; Zhang et al. proposed a method based on a maximum boundary-correlation algorithm and a lexical-semantic web-based summarization approach [10]; Singh et al. aimed to develop an end-to-end method that can generate short summaries and clear titles [11]. However, most of the current abstract generation methods based on the seq2seq model are still plagued by semantically irrelevant vocabulary, and there is less research on bilingual abstract generation of textile industry news.

In the field of abstract generation, the BART model shows strong potential with its two-way encoding and an autoregressive decoding mechanism [12,13]. Hartawan et al. used the BART model to generate abstract summaries from Indonesian texts [14]; Du proposed a dialogue summary model based on BART to promote the development of dialogue summary tasks [15]; Holle et al.'s research found that BART has a strong ability to denoise and generate coherent summaries in news document summary tasks [16]. However, due to the scarcity of corpus resources and poor model adaptability, the BART model generation results often have problems such as incomplete information and inaccurate terminology [17–19], and most research focuses on general fields. There is a lack of targeted optimization in vertical fields such as textile industry news [20]. In view of this, this article proposes the DEBART model, which is an improved version of the BART model. For domain enhancement, we integrate a professional corpus and terminology knowledge related to textile industry news to improve the model's understanding and processing power of texts in this field; at the same time, we introduce a readability optimization module to solve the shortcomings of existing methods in specific scenarios, and provide a more efficient solution for the generation of bilingual summaries in vertical fields.

Readability optimization is the key to improving the efficiency of bilingual news summary information communication and reader experience. Pantula et al. proposed a machine learning-based model to calculate the readability of web pages and use specific feature sets to classify the readability of web pages [21]; Eleyan et al. believed that the readability level can be measured by three formulas, of which the Flesch–Kincaid level can be improved by finding keywords and annotating statements [22]. These highlight the importance of FKGL in text readability optimization.

III. RESEARCH ON BILINGUAL ABSTRACT GENERATION AND READABILITY OPTIMIZATION OF TEXTILE INDUSTRY NEWS

CONSTRUCTION OF BILINGUAL CORPUS OF TEXTILE NEWS

In the task of bilingual abstract of textile industry news, the lack of high-quality field corpus has become a limiting factor [23]. Therefore, we have constructed a Chinese–English parallel abstract corpus in the textile vertical field. To obtain a professional corpus, we used customized web crawlers to collect content of authoritative sources in the global textile industry from 2018–2023. These sources include the Technical Bulletin of the International Federation of Textile Manufacturers, the Official Bulletin of the China Textile Industry Federation and professional media, etc., involving the research and development of new fiber materials and other sectors. We strictly complied with robots.txt rules and scraped only publicly accessible content. In order to avoid data leakage, a dual deduplication strategy based on URL hash and text semantic similarity was adopted to ensure that there was no overlap between the training set and the test set. At the same time, the data set was divided by hierarchical sampling to prevent information leakage and ensure that the model evaluation was fair and reliable. The collection of data was highly professional, and the coverage rate of professional terms was calculated through the constructed bilingual dictionary of terms in the textile field (15,800 in Chinese and 21,000 in English). The custom dictionary was loaded with Jieba for Chinese segmentation and NLTK for English, combined with the TextileGlossary terminology database. The statistical characteristics of the textile bilingual news corpus are shown in Table 1.

TABLE 1. Statistics of the bilingual textile-news corpuscategory

Category	Chinese corpus	English corpus
Number of news documents (articles)	482	434
Average abstract length	98.6 characters	72.3 words
Professional terminology coverage rate (%)	31.7	28.9
Theme distribution	New fiber materials: 32%; Intelligent manufacturing: 25%; Green printing and dyeing technology: 22%; Textile policies and regulations: 21%	New fiber materials: 28%; Intelligent manufacturing: 31%; Green printing and dyeing technology: 19%; Textile policies and regulations: 22%

As shown in Table 1, the corpus contains 482 Chinese news documents and 434 English documents. The average length of Chinese abstracts is 98.6 characters, and the average length of English abstracts is 72.3 words, reflecting the length difference when expressing the same content in Chinese and English. In terms of the coverage of professional terms, the Chinese corpus reaches 31.7%,

and the English corpus is 28.9%, indicating that there are many professional terms in textile industry news, which is a challenge for the summary generation model. We categorized news into core themes and industry standards (e.g., new fiber materials, intelligent manufacturing) via expert annotation. At the sentence level, we first screened candidates using a length-ratio threshold, then retained pairs with multilingual BERT similarity > 0.7 . This approach prevents misalignment caused by relying solely on average length ratios, ensuring reliable alignment. Three domain experts independently performed the annotation process following a standardized terminology guideline, achieving 95% consistency. Furthermore, the final corpus verification confirmed that this level of consistency was maintained, thus avoiding any issues related to inconsistent labeling.

The expert annotation process followed rigorous standards: A unified terminology translation standard was established before annotation, with a terminology database provided. Sentences were categorized into high, medium, and low difficulty levels according to the FKGL criteria. Experts underwent standardized training covering alignment rules and example exercises. During annotation, each expert underwent two rounds of review - first independent annotation followed by cross-verification of others' results. To ensure consistency, 10% of data were randomly selected for alignment verification. When disagreements arose between experts regarding sentence or paragraph alignment, a third senior expert arbitrated and documented the reasons for disputes. The arbitration success rate reached 95%, ensuring high accuracy and quality in annotation.

CORPUS DATA PREPROCESSING

The collected textile industry news data needs to be pre-processed and cleaned when constructing the data set. The original news text is often mixed with noise and non-text content. Affected by the characteristics of HTML (Hyper-Text Markup Language) pages, irrelevant information such as HTML tags, special characters, and line breaks also appear in the text, which needs to be processed uniformly. In processing textile-industry news, domain-specific segmentation is crucial. The Chinese text segmentation of textile industry news requires customized tools and domain dictionaries to accurately identify and segment professional terms [24-26]. We adopt the Jieba segmenter, load a special textile term dictionary, and use the Maximum Reverse-Matching algorithm [27-29] to ensure the complete segmentation of compound terms to avoid misclassifying "Sirofil core-spun yarn" as "Sirofil core/spun yarn." English text processing uses NLTK's standard tokenizer (Natural Language Toolkit), combined with the TextileGlossary term database, to annotate proper noun phrases. The Maximum Reverse-Matching algorithm is used to segment the input text into words [30,31], and the formula is expressed as:

$$\text{MaxMatch}(S) = \arg \max_{t \in \text{Dictionary}} (\text{Match}(S, t)) \quad (1)$$

where S is the text to be segmented; t is the term in the domain term dictionary; and $\text{Match}(S, t)$ represents the matching degree between text S and term t . The similarity matrix of Chinese and English textile terms and the matching pair analysis are shown in Figure 1.

As shown in Figure 1, the left y-axis represents English and the x-axis represents Chinese. Figure 1 intuitively shows the cosine similarity between the six Chinese and the corresponding English terms, and the color depth reflects the degree of similarity.

After word segmentation, text alignment and annotation are performed to ensure that the Chinese and English abstracts are semantically accurate [32-34]. We adopt an expert collaborative annotation mechanism, in which researchers with a textile engineering background extract key factual elements from the original news and construct a structured Chinese abstract template; then, a professional localization team performs terminology standardization and completes the English translation; finally, the bilingual editor-in-chief reviews the terminology consistency. Terms such as "warp-knitted elastic fabric" must be strictly corresponding, and readability classification labels must be annotated. Assuming that the Chinese and English term pair is (CC, EE), where CC is the Chinese term, and EE is the English term. The formula is as follows:

$$\text{Consistency}(C, E) = \text{True} \quad \text{if} \quad f(C) = f(E) \quad (2)$$

where $f(C)$ and $f(E)$ represent the standardized processing results of Chinese and English terms, respectively.

Classification is performed based on the readability score when generating readability grading labels. Setting the readability score to $\text{Label}(X)$ and defining T_1 , T_2 , and T_3 grading thresholds. The formula is:

$$\text{Label}(X) = \begin{cases} \text{High}, & \text{if } X \leq T_1, \\ \text{Medium}, & \text{if } T_1 < X \leq T_2, \\ \text{Low}, & \text{if } X \geq T_3. \end{cases} \quad (3)$$

where T_1 , T_2 , and T_3 are the thresholds for readability scores, and the corresponding labels (high, medium, and low) are obtained through readability analysis of the text.

DEBART MODEL CONSTRUCTION AND TERMINOLOGY ENHANCEMENT

Given the sector's globalization, automatic summarization of textile-industry news has become essential for efficient information transfer. However, when textile industry news is transmitted across languages, it often faces three major problems: terminology distortion, rigid sentence structure, and poor readability. These problems not only hinder the accurate transmission of information but also reduce the quality of translation and affect readers' understanding. The BART model structure combines BERT [35-37] and GPT, but unlike BERT, which only uses a single noise type,

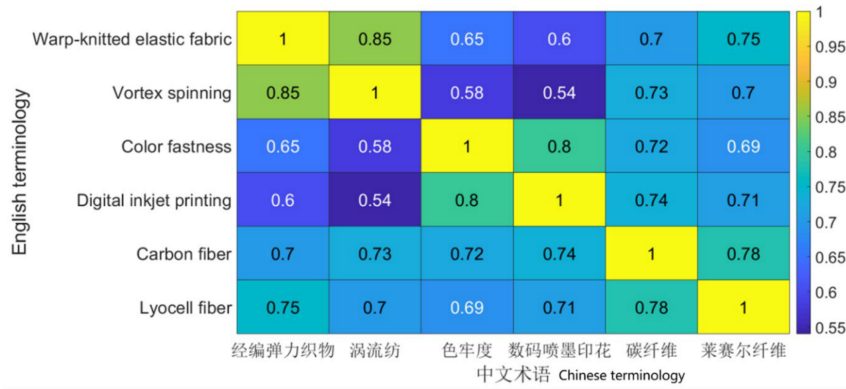


FIGURE 1. Similarity matrix and matching pair analysis of Chinese and English textile terms

BART employs multiple noise types on the encoder side. This study uses a large-scale pre-trained model based on BART, combining domain knowledge and structural innovation to construct the DEBART model to achieve semantic fidelity and expression optimization of bilingual summaries. DEBART is based on the BART pre-trained architecture, consisting of a 12-layer encoder and a 12-layer decoder, with a hidden layer dimension of 1024. In DEBART, the input sequence is first processed by a preprocessor to remove noise and irrelevant information before being input into the BART pre-trained model. The encoder uses a bidirectional Transformer structure to obtain rich contextual information about the input text.

The decoder uses a one-way autoregressive model to generate an output sequence that is asymmetric with respect to the input sequence, ensuring that the generated summary conforms to the grammar and expression habits of the target language [38,39]. During training, the model input contains the original text and the reference summary. The reference summary plays a guiding role and participates in summary generation through the autoregressive decoder. The objective is to minimize the negative log-likelihood function $L(\beta_c)$ between the generated result and the standard summary, that is:

$$L(\beta_c) = - \sum_{t=1}^T \log Q(y_t | x, y < t; \beta_c) \quad (4)$$

where x is the input sequence; $y < t$ is the output generated by the decoder before; y_t is the t -th word of the target summary; β_c is the model's parameter; and T is the length of the summary.

(1) Terminology consistency constraint Significant differences exist between Chinese and English terminology in textile industry news. Improper handling during bilingual abstract generation may lead to inaccurate or distorted translations, compromising the quality of abstracts. DEBART enhances the BART model by incorporating textile-specific term knowledge into its word-embedding layer, ensuring precise translation of technical terms and improving accuracy and consistency. To further optimize translation effectiveness, DEBART employs a gated-fusion mechanism

that combines term vectors with original word embeddings through weighted integration, enabling accurate processing of terms while avoiding distortion. The system first extracts term vectors from the dictionary, combines them with original word embeddings, then calculates weight coefficients based on contextual matching accuracy to dynamically adjust the fusion level. The BART vocabulary is used to convert the original text into a word-embedding vector, and then the word word-embedding and the positional encoding are added to obtain the original text feature Z_0^{enc} input to the encoder. The formula is as follows:

$$Z_0^{enc} = A_t + E_{pe}, \quad Z_0^{enc} \in \mathbb{R}^{N \times d_t} \quad (5)$$

where $A_t \in \mathbb{R}^{N \times d_t}$ is the word embeddings vector of the original text; N represents the sequence length of the original text; d_t represents the dimension of the word embeddings vector; and E_{pe} is the positional encoding of the original text.

(2) Coverage mechanism The original text of textile industry news has complicated sentence patterns and lengthy descriptions. Streamlining and transforming it into concise and smooth language is essential to improve the quality of the abstract. However, in the decoder-side interactive attention sublayer, the calculation of the attention mechanism at different times is independent of each other. The model may repeatedly focus attention on a certain location in several consecutive time slices, resulting in duplicate information in the generated summary text. To solve this problem, we introduce a coverage mechanism. First define the coverage vector with an initial zero vector, which represents the sum of the attention distribution at time t and reflects the degree of coverage of the input words. Then incorporate the coverage vector into the next attention calculation and modify the attention weight. When such a model generates new words, it will reduce attention to the area of interest, increase attention to previously unattended tokens/words, and effectively reduce duplicate words in the abstract.

Traditional PGN relies on coverage vectors to track source text positions, with the core mechanism being attention weight accumulation based on historical attention patterns to prevent redundant pointing to the same source term.

However, this approach is only applicable in single-language scenarios and depends on explicit position tracking of source text. DEBART's coverage mechanism introduces innovations by integrating bilingual term alignment features: On one hand, its coverage vector not only tracks generated English summary positions but also correlates Chinese term semantic coverage status through a gating fusion mechanism. For example, when generating "carbon fiber", it simultaneously marks the Chinese term "碳纤维" and its mapping relationship to avoid term ambiguity-induced duplication. On the other hand, dynamic weight adjustment implements priority rules for textile industry terminology, ensuring gradual decay of high-frequency professional term coverage while maintaining singular expressions of core terms during long-text generation.

To improve the efficiency of model optimization, this paper applies coverage loss in the loss function. In generating summaries of textile news text, if the model focuses too much on certain positions, it is easy to cause repetitive content in the generated summary, affecting diversity and information content. Coverage loss can effectively solve this problem. Assuming that at a certain time t , the model's attention weight for word w is a_t , and the sum of all attention to the word in the previous time (from 1 to $t-1$) is c_t . If at time t , the attention weight a_t is less than the sum of all attention to the word in the previous time c_t ; it means that at time t , word w is over-focused. At this time, its attention weight should be reduced to prevent repeated attention to the same position. Otherwise, if $a_t > c_t$, it means that at the previous time, the attention of word w is low, and the model should pay more attention to it. The change process of attention distribution after applying the coverage mechanism is shown in Figure 2.

(3) Beam search strategy In the process of bilingual summary generation, the decoding strategy significantly impacts the quality of the output results. To improve the language fluency and naturalness of the English summary while being faithful to the original information, this paper applies the beam search strategy in the model decoding stage to control the length and diversity of the generated text effectively. In the beam search algorithm, given the current candidate sequence set H_{t-1} , each time step t selects a new candidate word based on the candidate sequence at the previous moment. Assuming that the candidate word sequence at the current time step is H_t , its calculation formula is:

$$H_t = \arg \max_{h_t \in \{h_1, h_2, \dots, h_k\}} Q(m_t | h_{t-1}) * Q(h_{t-1}) \quad (6)$$

where h_t is the state of the current candidate sequence; $Q(m_t | h_{t-1})$ is the probability of generating the current word m_t given the previous sequence h_{t-1} ; $Q(h_{t-1})$ is the probability of the sequence at the previous moment; and k represents the number of current candidate sequences.

When generating English summaries, the model maintains a fixed beam width; that is, it simultaneously tracks multiple

most likely candidate sequences. By adjusting the beam width, a balance can be achieved between generation quality and computational complexity. A larger beam width is conducive to exploring more expressions and improving language diversity and logical coherence; a smaller beam width makes converging to a locally optimal solution easier. The changes in the length and diversity of the generated text under different beam widths are shown in Figure 3.

In Figure 3, the x-axis represents "beam width" (a dimensionless integer parameter), while the y-axis shows the "average length of generated text" measured in word count. The Distinct-2 metric quantifies textual diversity using the formula: Diversity = 1 - (Number of repeated binary grammatical structures / Total number of binary grammatical structures), with values ranging from [0,1]. Higher values indicate more diverse vocabulary and richer expression. The results demonstrate that as beam width increases, the average length of generated text gradually rises and stabilizes, suggesting that incorporating more candidate paths into the selection process enhances content information density or length. The diversity metric shows a rapid initial increase followed by stabilization, indicating that introducing more candidate paths improves the uniqueness and novelty of output text.

This paper sets the beam width to 9 during the DEBART model's decoding phase, as illustrated in Figure 3. The figure demonstrates that as beam width increases, the average text length and diversity initially rise rapidly before stabilizing. When beam width is smaller, the model can explore more paths to improve summary quality, though the improvement remains limited. At beam width 9, the process approaches a plateau phase. This configuration ensures both model-generated quality and diversity while controlling computational complexity and decoding time, achieving a balance between output effectiveness and efficiency.

READABILITY OPTIMIZATION MODULE CONSTRUCTION

Fluency and readability are essential when generating bilingual summaries of textile-industry news, especially for English outputs. Improving legibility substantially enhances communication effectiveness. To this end, this paper constructs a readability enhancement module that integrates multi-dimensional language optimization strategies, introduces FKGL as the core readability evaluation metric, and measures the reading difficulty by calculating the average sentence length and the number of syllables of the word, and incorporates FKGL into the generation control during training and inference, dynamically adjusting the sentence length and vocabulary complexity to ensure that the English abstract not only accurately conveys the original information, but also has good readability. For long and difficult sentences, the model prioritizes simplification or splitting to reduce the reading burden. In addition, this article also designs a sentence pattern conversion template and conjunctive guidance strategy to optimize the structure

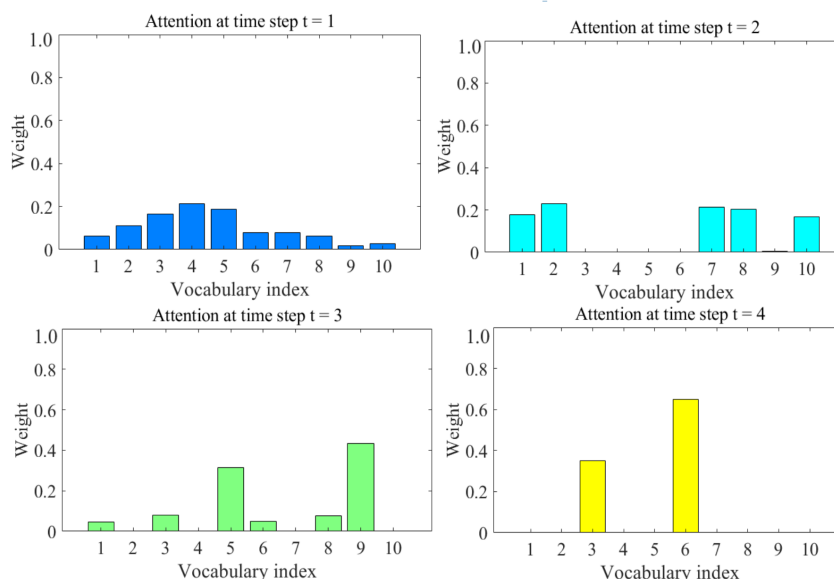


FIGURE 2. Changes in attention distribution after the application of the coverage mechanism

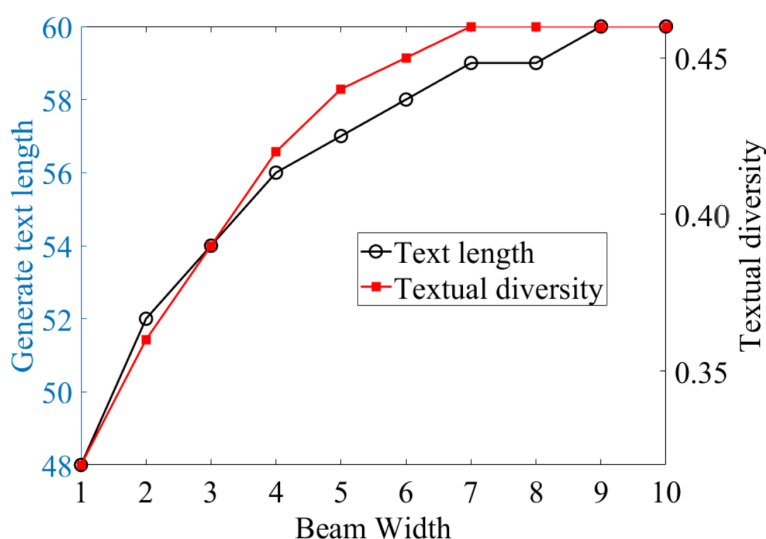


FIGURE 3. Line graph of generated text length and diversity under different beam widths

and fluency of the abstract. By adjusting the original sentence pattern and adding conjunctions, the description of technical parameters and professional terms is more concise and easier to understand, ensuring that the logical relationship between sentences is clear, significantly improving the quality of the abstract, and more in line with readers' reading habits. In the DEBART model, the FKGL algorithm not only serves as a readability evaluation metric but also actively participates in the generation process, enabling dynamic sentence simplification and conjunction addition. During summary generation, the model first calculates the FKGL score of the current text to assess complexity. If the score exceeds a preset threshold (e.g., 8 corresponding to medium difficulty), the sentence simplification mechanism

is activated. This mechanism applies gating and adjusts term vectors and word-embedding weights to prioritize simple sentence structures, such as splitting complex sentences. DEBART's conjunction addition strategy utilizes sentence transformation templates and attention weight adjustments. The model contains a built-in template library with more than 200 transformation rules, automatically matching templates when the FKGL score exceeds the threshold.

This study addresses Chinese text readability evaluation by developing a multi-dimensional assessment system that integrates syntactic complexity and term density, instead of directly adopting the FKGL metric. The system measures syntactic complexity through average sentence length and clause quantity, calculates term proportion using specialized

terminology dictionaries, and incorporates linguistic features like conjunction frequency and punctuation complexity. A Chinese readability scoring model was trained on expert-labeled corpora, with manual ratings from textile industry experts using a 5-point Likert scale. In DEBART's readability optimization, Chinese abstracts enhance information clarity through term consistency control and coverage mechanisms, while English abstracts explicitly employ FKGL metrics to dynamically simplify syntax and adjust vocabulary. This divide-and-conquer strategy for Chinese-English processing accommodates distinct linguistic characteristics, ensuring generated summaries achieve both semantic accuracy and reader-friendly readability.

IV. BILINGUAL ABSTRACT QUALITY ASSESSMENT SYSTEM

EXPERIMENTAL ENVIRONMENT

The specific experimental environment is shown in Table 2

TABLE 2. Experimental environment setup

Serial number	Parameter	Set up
1	Operating system	Ubuntu 18.04
2	CPU (central processing unit)	Intel i5-6200U
3	Memory	32 GB
4	Attention heads	16
5	Learning rate	2e-5
6	Dropout	0.1
7	Epochs	16
8	Batch size	32

As shown in Table 2, to ensure efficient and stable model training and evaluation, this experiment adopted Ubuntu 18.04 as the operating system platform. The hardware configuration utilized an Intel i5-6200U processor to provide sufficient computational power for data processing. The model architecture featured 16 attention heads to enhance feature capture capabilities, with a learning rate set at 2e-5 and a dropout rate of 0.1 implemented to prevent overfitting, thereby ensuring the model's generalization ability.

EXPERIMENTAL EVALUATION INDICATORS

ROUGE evaluates summarization by comparing overlaps between generated and reference summaries. In this study, we implemented standardized experimental protocols: Chinese text was segmented using Jieba, while English text utilized NLTK's standard segmentation tool. Evaluation retained stop words, with all ROUGE scores calculated using the F1 metric that balances recall and precision to maintain both information completeness and accuracy. The calculation formula of ROUGE-N is as follows:

$$\text{ROUGE-N} = \frac{\sum_{S \in \{Ref\}} \sum_{gram_N \in S} \text{Count}_{match}(gram_N)}{\sum_{S \in \{Ref\}} \sum_{gram_N \in S} \text{Count}(gram_N)} \quad (7)$$

where $gram_N$ denotes an n-gram; for example, $gram_1$ denotes a unigram. $\text{Count}_{match}(gram_N)$ represents an n-

gram subsequence in the model-generated summary that is the same as the reference summary. $\text{Count}(gram_N)$ represents an n-gram subsequence in the reference summary. Ref represents the reference summary content.

Because higher coverage of key information generally indicates better quality, we use key-information coverage as an evaluation metric to more accurately evaluate the summary generation effect. The calculation formula is as follows:

$$\text{Coverage}_{key}(D, S^*) = \frac{\sum_{g \in S^*} t(g \in key_D)}{|key_D|} \quad (8)$$

where D and S^* represent the summary content of the original document and the model generated, respectively; key_D represents the keyword information in document D ; t represents the indicator function.

This study integrates text salience analysis with domain-specific dictionary extraction for key information extraction. The TextRank algorithm is employed to calculate candidate word importance scores, while a textile terminology dictionary is applied to assign higher weights to specialized terms. The top 15% ranked vocabulary and terminology are selected as the keyword set, with exact matching implemented during coverage calculation. Experimental data are divided into training, validation, and test sets in a 7:2:1 ratio, incorporating both Chinese and English corpora. Chinese ROUGE evaluation combines Jieba segmentation with dictionary-based term segmentation, retaining stop words in both languages, with F1 scores serving as performance metrics. Human evaluations are conducted by seven experts using a five-point scale across four dimensions: terminological accuracy, fluency, coherence, and information completeness. The original data are divided as shown in Table 3.

TABLE 3. Original data partitioning

Language	Total documents	Training set	Verification set	Test set
Chinese	482	337	96	49
English	434	304	87	43
Total	916	641	183	92

EVALUATION METRIC ANALYSIS

To comprehensively evaluate the Chinese summary generation capabilities of the DEBART model in textile industry news, this study selected multiple representative traditional text generation models as comparative baselines, including RNN-context (Recurrent Neural Network with Context), BART, seq2seq, Bi-LSTM, BERT, and ASTCL (Abstract Summarization Model Based on Topic-awareness and Contrastive Learning). Among these, BERT was excluded due to its lack of autoregressive generation capability. Instead, we compared its extractive summarization approach—specifically, fine-tuning BERT to score sentence importance and select Top-N sentences for combined summaries—as a

robust extractive baseline. All baseline models were implemented using standard open-source code, with datasets partitioned proportionally from the same textile news corpus. Hyperparameters were rigorously tuned to ensure optimal performance and experimental fairness. Multi-dimensional analysis was conducted using core metrics such as ROUGE and Coverage, with results presented in Table 4.

TABLE 4. Performance comparison of different summary generation models

Model	ROUGE-1	ROUGE-2	ROUGE-L	Coverage
DEBART	46.25	25.76	42.84	26.82
RNN-context	39.82	19.44	36.75	18.41
BART	36.94	17.52	33.86	18.06
seq2seq	31.95	14.72	29.36	14.18
Bi-LSTM	38.11	18.27	35.03	17.69
BERT	35.28	16.65	32.41	16.07
ASTCL	41.67	20.53	38.12	21.44

Table 4 shows that in the task of generating news summaries in the textile industry, the DEBART model performed well and was significantly better than other models. It received the highest scores on the four indicators of ROUGE-1, ROUGE-2, ROUGE-L and Coverage, 46.25, 25.76, 42.84 and 26.82, respectively, highlighting the advantages of generation quality and information integrity. The scores of traditional models are generally low, reflecting the lack of processing complex semantics and long text. Although BART has strong semantic modeling ability, its score is still lower than DEBART, indicating that the improvements to the BART architecture proposed in this study have enhanced its ability to capture and express key information. DEBART scored high on ROUGE-1 and ROUGE-2, indicating that the vocabulary level of the generated results is highly consistent with the references, with accurate and natural expression; ROUGE-L indicates sentence structure and semantic coherence close to human-written summaries. The high coverage score indicates that while retaining key information, it can effectively reduce redundancy and omissions, which is of great significance in the generation of news summaries in the textile industry.

In order to comprehensively evaluate the performance of the DEBART model in the task of generating bilingual summaries of textile industry news and compare it with other mainstream sequence modeling methods, this study conducted 200 independent runs with different random seeds. Under a fixed dataset split, we initialized model parameters with different seeds and repeated training/evaluation to obtain a stable performance distribution, plotted the loss curve, and computed the average loss. In order to control the risk of overfitting and selective reporting, the training was strictly based on the performance of the verification set, and the results are based on the independent test set; the average performance of 200 experiments was taken, and the ranking of the results was consistent with a single fixed seed experiment, indicating that the performance advantage of the model is stable. The final conclusion is based on

the average performance of the model on a fixed test set to ensure that the results are repeatable and reliable. The experiment involved DEBART, RNN-context, BART and other models, and the results are shown in Figure 4.

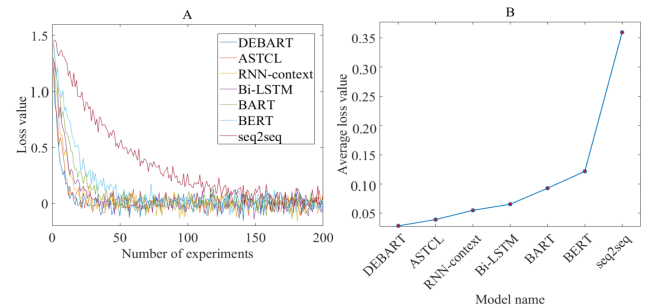


FIGURE 4. Comparison of loss curves and loss value rankings of different summary generation models. Figure 4A: Loss curve. Figure 4B: Loss value ranking

As shown in Figure 4A, the DEBART model rapidly decreases in the early iteration stage due to the application of the domain enhancement mechanism and coverage loss control strategy and remains stable in subsequent training, reflecting good learning efficiency and stability. Although models such as ASTCL and RNN-context also have a downward trend, the speed is slow, and the fluctuation is large, reflecting the limitations of processing professional terms and complex sentences. To highlight the differences in model performance, Figure 4B shows the average loss value of each model at the end of the training, which is calculated based on the results of 200 experiments and is sorted from smallest to largest by final loss, namely DEBART, ASTCL, RNN-context, Bi-LSTM, BART, BERT, and seq2seq, which is consistent with the expected performance ranking. To quantify the contributions of the core modules (terminology enhancement, coverage mechanism, readability optimization) in this study, we designed systematic ablation experiments. Conducted under standardized software-hardware environments with fixed random seeds to ensure reproducibility, the test sets were independently constructed and kept hidden during training parameter tuning. Representative baseline models were selected and hyperparameters were thoroughly optimized to ensure fair comparisons. The evaluation system integrated automated metrics with human assessments to comprehensively and objectively measure generated summary quality from multiple dimensions. Experimental results are presented in Table 5.

TABLE 5. Analysis of ablation experimental results of DEBART model

Model abbreviation	Model configuration description	ROUGE-L (F1)	Terminology accuracy(%)	FKGL (↓)	Manual readability(5-point system)
BART	Original BART model	33.86	76.2	8.3	3.5
+ Term	BART + Terminology enhancement (gated-fusion)	38.41	86.5	8.1	4.0
+Term+Cov	+ Term + coverage mechanism	40.62	90.8	7.8	4.2
+Term+Cov+FKGL	+Term +Cov + readability optimization (FKGL)	42.05	91.5	6.2	4.5
DEBART (Full)	+Term +Cov +FKGL +beam width = 9	42.84	95.3	6.2	4.8

Table 5 demonstrates that each core module of the DEBART model makes significant incremental contributions. Compared to the original BART baseline, progressive integration of modules yields systematic improvements: The term enhancement module significantly boosts ROUGE-L and term accuracy rates, highlighting the critical role of domain knowledge infusion in preserving professional terminology fidelity; The coverage mechanism further improves ROUGE-L and term accuracy to 90.8%, effectively suppressing redundancy while enhancing content density; Although the readability optimization module shows minimal improvement in ROUGE metrics, it substantially reduces FKGL values, elevates human readability scores, and enhances linguistic fluency; The fully optimized DEBART model incorporating beam width adjustment achieves optimal performance across all metrics, validating both the effectiveness of its design and the synergistic effects among components.

V. CONCLUSION

This research focuses on issues such as inconsistent translation of terms, missing information, and poor readability of English summaries in the generation of bilingual summaries of textile industry news, and proposes an optimization method based on DEBART. We built a high-quality Chinese and English textile news corpus, including professional terminology and industry knowledge, to provide data support for the model, and make up for the shortcomings of traditional methods in terminology processing. At the model level, DEBART injects domain term information on the basis of BART, and uses a gating mechanism to integrate term vectors and original word embeddings to improve the accuracy and consistency of term expression. We introduced a coverage mechanism, dynamically adjusted the distribution of attention, and suppressed the generation of repetitive words, and improved the completeness and coherence of the abstract. In the readability optimization stage, a readability

threshold based on FKGL was set to dynamically optimize the sentence structure and vocabulary complexity to ensure accurate semantics and smooth language. Experiments have shown that DEBART performs best on indicators such as ROUGE-1 and ROUGE-2, with the leading readability score and the lowest FKGL score, which effectively improves the consistency and naturalness of bilingual terminology, enhances the readability and user experience of English summaries, and provides a feasible path for automatic news summarization in the textile field.

AUTHOR CONTRIBUTIONS

Conceptualization – Yuntao Li; methodology – Yuntao Li and Xueqin Li; investigation – Xueqin Li; writing-original draft preparation – Yuntao Li, Xueqin Li. All authors have read and agreed to the published version of the manuscript.

CONFLICTS OF INTEREST

The authors declare no conflict of interest.

FUNDING

This research received no external funding.

ACKNOWLEDGEMENTS

Not applicable.

REFERENCES

- [1] H. Huang, Z. Chen, C. Xu, and X. Zhang, "Automatic summarization model of aerospace news based on domain concept graph," *Journal of Beijing University of Aeronautics and Astronautics*, vol. 50, no. 1, pp. 317-327, 2024, doi: 10.13700/j.bh.1001-5965.2022.0233.
- [2] M. A. Niculescu, S. Ruşeti, and M. Dascalu, "RoSummary: Control Tokens for Romanian News Summarization," *Algorithms*, vol. 15, no. 12, pp. 472, 2022, doi: 10.3390/a15120472.
- [3] Y. Kumar, K. Kaur, and S. Kaur, "Study of automatic text summarization approaches in different languages," *Artificial Intelligence Review*, vol. 54, no. 8, pp. 5897-5929, 2021, doi: 10.1007/s10462-021-09964-4.
- [4] Y. C. Chang, Y. W. Chiu, and T. W. Chuang, "Linguistic Pattern-Infused Dual-Channel Bidirectional Long Short-term Memory with Attention for Dengue Case Summary Generation from the Program for Monitoring Emerging Diseases-Mail Database: Algorithm Development Study," *JMIR Public Health and Surveillance*, vol. 8, no. 7, Art. no. e34583, 2022, doi: 10.2196/34583.
- [5] Y. Cao and Y. Xu, "Dual-Channel Text Summarization Generation Method Based on Graph Attention," *Computer Applications and Software*, vol. 41, no. 4, pp. 159-164, 241, 2024, doi: 10.3969/j.issn.1000-386x.2024.04.024.
- [6] R. Lin, C. Cheng, F. Lin, X. Peng, M. Lin, and H. Lin, "Method for Generating Complex Text Summary Based on Attention-copy Mechanism," *Computer & Digital Engineering*, vol. 49, no. 11, pp. 2292-2298, 2021, doi: 10.3969/j.issn.1672-9722.2021.11.023.
- [7] B. Pei, X. Li, K. Hu, and Z. Sun, "Judicial Text Summarization Based on Knowledge-enhanced Pretrained Language Models," *Science Technology and Engineering*, vol. 24, no. 20, pp. 8587-8597, 2024, doi: 10.12404/j.issn.1671-1815.2304954.
- [8] R. Wijayanti, M. L. Khodra, K. Surendro, and D. H. Widyantoro, "Learning bilingual word embedding for automatic text summarization in low resource language," *Journal of King Saud University-Computer and Information Sciences*, vol. 35, no. 4, pp. 224-235, 2023, doi: 10.1016/j.jksuci.2023.03.015.
- [9] S. Ning, X. Yan, G. Xu, F. Zhou, and L. Zhang, "Chinese news text abstractive summarization with keywords fusion," *Computer Engineering & Science*, vol. 42, no. 12, pp. 2265-2272, 2020, doi: 10.3969/j.issn.1007-130X.2020.12.021.
- [10] Q. Zhang, Y. Fan, and D. Jin, "Research on News Text Summarizations Generation Based on MMR and WordNet," *Journal of Southwest China Normal University (Natural Science Edition)*, vol. 48, no. 5, pp. 77-86, 2023, doi: 10.13718/j.cnki.xsxb.2023.05.011.
- [11] R. K. Singh, S. Khetarpaul, R. Gorantla, and S. G. Allada, "SHEG: Summarization and headline generation of news articles using deep learning," *Neural Computing and Applications*, vol. 33, no. 8, pp. 3251-3265, 2021, doi: 10.1007/s00521-020-05188-9.
- [12] M. La Quatra and L. Cagliero, "BART-IT: An Efficient Sequence-to-Sequence Model for Italian Text Summarization," *Future Internet*, vol. 15, no. 1, pp. 15, 2023, doi: 10.3390/fi15010015.
- [13] E. S. Kim, H. Yoo, and K. Chung, "Performance Improvement of Topic Modeling using BART based Document Summarization," *Journal of Internet Computing and Services*, vol. 25, no. 3, pp. 27-33, 2024, doi: 10.7472/jksii.2024.25.3.27.
- [14] G. Hartawan, D. S. Maylawati, and W. Uriawan, "Bidirectional and Auto-Regressive Transformer (BART) for Indonesian Abstractive Text Summarization," *Jurnal Informatika Polinema*, vol. 10, no. 4, pp. 535-542, 2024, doi: 10.33795/jip.v10i4.5242.
- [15] L. Du and B. A. R. T. Query-Based Dialogue Summarization Using, "Applied and Computational Engineering," 2024; 29:160-166, doi: 10.54254/2755-2721/29/20231149.
- [16] K. F. H. Holle, D. N. Munna, and E. W. Ekaputri, "Performance Evaluation of Transformer Models: Scratch, Bart, and Bert for News Document Summarization," *Jurnal Teknik Informatika (Jutif)*, vol. 6, no. 2, pp. 787-802, 2025, doi: 10.52436/1.jutif.2025.6.2.2534.
- [17] P. Wilman, T. Atara, and D. Suhartono, "Abstractive English Document Summarization Using BART Model with Chunk Method," *Procedia Computer Science*, vol. 245, pp. 1010-1019, 2024, doi: 10.1016/j.procs.2024.10.329.
- [18] E. Daraghmi, L. Atwe, and A. Jaber, "A Comparative Study of PEGASUS, BART, and T5 for Text Summarization Across Diverse Datasets," *Future Internet*, vol. 17, no. 9, pp. 389, 2025, doi: 10.3390/fi17090389.
- [19] M. N. Alzamzami and M. L. Khodra, "Abstract Meaning Representation Parser Development for Cross-lingual Indonesian-English with BART, Input Concatenation, and Dataset Augmentation," *International Journal on Electrical Engineering and Informatics*, vol. 16, no. 4, pp. 584-603, 2024, doi: 10.15676/ijeei.2024.16.4.5.
- [20] J. Li and K. K. Leonas, "Sustainability topic trends in the textile and apparel industry: A text mining-based magazine article analysis," *Journal of Fashion Marketing and Management: An International Journal*, vol. 26, no. 1, pp. 67-87, 2022, doi: 10.1108/JFMM-07-2020-0139.
- [21] M. Pantula and K. S. Kuppusamy, "A Machine Learning-Based Model to Evaluate Readability and Assess Grade Level for the Web Pages," *The Computer Journal*, vol. 65, no. 4, pp. 831-842, 2022, doi: 10.1093/comjnl/bxaa113.
- [22] D. Eleyan, A. Othman, and A. Eleyan, "Enhancing Software Comments Readability Using Flesch Reading Ease Score," *Information*, vol. 11, no. 9, pp. 430, 2020, doi: 10.3390/info11090430.
- [23] P. Donner, "Identifying constitutive articles of cumulative dissertation theses by bilingual text similarity," *Evaluation of similarity methods on a new short text task. Quantitative Science Studies*, vol. 2, no. 3, pp. 1071-1091, 2021, doi: 10.1162/qss_a_00152.
- [24] X. Zhao, J. Huang, J. Zhang, and Y. Song, "The Comprehensive Analysis of the Effect of Chinese Word Segmentation on Fuzzy-Based Classification Algorithms for Agricultural Questions," *International Journal of Fuzzy Systems*, vol. 26, no. 8, pp. 2726-2749, 2024, doi: 10.1007/s40815-024-01724-0.
- [25] Y.-K. Tsang, M. Yan, J. Pan, and M. Y. K. Chan, "A corpus of Chinese word segmentation agreement," *Behavior Research Methods*, vol. 57, no. 1, pp. 25, 2025, doi: 10.3758/s13428-024-02528-8.
- [26] S. Yang, "Investigating word segmentation of Chinese second language learners," *Reading and Writing*, vol. 34, no. 5, pp. 1273-1293, 2021, doi: 10.1007/s11145-020-10113-6.
- [27] M. Feldman and A. Szarf, "Maximum Matching Sans Maximal Matching: A New Approach for Finding Maximum Matchings in the Data Stream Model," *Algorithmica*, vol. 86, no. 4, pp. 1173-1209, 2024, doi: 10.1007/s00453-023-01190-4.
- [28] S. A. Almaaytah, "Arabic word tokenization system using the maximum matching model," *Edelweiss Applied Science and Technology*, vol. 8, no. 6, pp. 3210-3217, 2024, doi: 10.5521/4/25768484.v8i6.2682.
- [29] J. Pei, "A Dictionary-based Maximum Match Algorithm Via Statistical Information for Chinese Word Segmentation," *International Journal of Electronics and Information Engineering*, vol. 12, no. 1, pp. 24-33, 2020.
- [30] A. Tong, "Research on the Application of Natural Language Processing in Information Systems," *Creativity and Innovation (Creative Version)*, vol. 6, no. 6, pp. 37-39, <https://dx.doi.org/10.12184/wspcyycx2WSP2516-415507.20220606>, 2022.
- [31] Z. Zhang, W. Yu, G. Yan, and C. Yuan, "Chinese Word Segmentation Based on ACNNC Model," *Journal of Chinese Information Processing*, vol. 36, no. 8, pp. 12-19, 28, 2022, doi: 10.3969/j.issn.1003-0077.2022.08.002.
- [32] C. Meinecke, D. J. Wrisley, and S. Jänicke, "Explaining Semi-Supervised Text Alignment through Visualization," *IEEE Transactions on Visualization and Computer Graphics*, vol. 28, no. 12, pp. 4797-4809, 2022, doi: 10.1109/TVCG.2021.3105899.
- [33] T. Yousef and S. Janicke, "A Survey of Text Alignment Visualization," *IEEE Transactions on Visualization and Computer Graphics*, vol. 27, no. 2, pp. 1149-1159, 2021, doi: 10.1109/TVCG.2020.3028975.
- [34] S. Ge and R. Song, "English-Chinese Clause Alignment Corpus Tagging System Based on Component Sharing," *Journal of Chinese Information Processing*, vol. (6), pp. 27-35, 2020, doi: 10.3969/j.issn.1003-0077.2020.06.005.
- [35] P. Howlader, P. Paul, M. Madavi, L. Bewoor, and V. S. Deshpande, "Fine Tuning Transformer Based BERT Model for Generating the Automatic Book Summary," *International Journal on Recent and Innovation Trends in Computing and Communication*, vol. 10, no. 1s, pp. 347-352, <http://dx.doi.org/10.17762/ijritcc.v10i1s.5902>, 2022.
- [36] N. Sanchan, "Comparative Study on Automated Reference Summary Generation using BERT Models and ROUGE Score Assessment," *Journal of Current Science and Technology*, vol. 14, no. 2, pp. Article 26, 2024, doi: 10.59796/jest.V14N2.2024.26.
- [37] W. Nam, J. Lee, and B. Jang, "Text summarization of dialogue based on BERT," *Korean Society of Computer Information*, vol. 27, no. 8, pp. 41-47, 2022, doi: 10.9708/jksci.2022.27.08.041.
- [38] Z. Liu, K. Zhang, and H. Zhu, "Automatic Text Summary Generation Based on Transformer Model," *Computer & Digital Engineering*, vol. 52, no. 2, pp. 482-486, 527, 2024, doi: 10.3969/j.issn.1672-9722.2024.02.034.

- [39] J. Zhang and F. Zhao, "Summarization Generation Method for Long Text Based on Key Features," *Computer & Digital Engineering*, vol. 52, no. 5, pp. 1412-1417, 2024, doi: 10.3969/j.issn.1672-9722.2024.05.026.