



Advancements in Automatic Text Summarization using Natural Language Processing

¹Roha Irfan, ¹Rabia Tehseen, ²Anam Mustaqeem, ¹Ramsha Saeed, ¹Usman Aamer, ¹Jawad Hassan

¹Department of Computer Science, University of Central Punjab, Lahore, Pakistan

²Department of Software Engineering, University of Central Punjab, Lahore, Pakistan

*Correspondence: : rabia.tehseen@ucp.edu.pk

Citation | Irfan. R, Tehseen. R, Anam. M, Saeed. R, Aamer. U, Hassan. J, “Advancements in Automatic Text Summarization using Natural Language Processing”, IJIST, Vol. 07 Issue. 02 pp 1145-1460, June 2025

DOI | <https://doi.org/10.33411/ijist/20257211451160>

Received | April 10, 2025 **Revised** | May 31, 2025 **Accepted** | June 04, 2025 **Published** | June 06, 2025.

With the rapid expansion of data across various domains, the need for automated text summarization has become increasingly crucial. Given the overwhelming volume of textual and numerical data, effective summarization techniques are required to extract key information while preserving content integrity. Text summarization has been a subject of research for decades, with various approaches developed using natural language processing (NLP) and a combination of different algorithms. This paper is an SLR-type essay presenting the existing text summarization techniques and their evaluation. It covers the basic concepts behind extractive and abstractive summarization and how deep learning models could serve as a boost in the performance of summarization. The study goes on to investigate the present use of text summarization in different areas and investigates the various methodologies applied in this area. A total of twenty-four carefully selected research articles were being analyzed to identify key trends, challenges and limitations regarding text summarization techniques. It proposes a number of open research challenges with insight concerning possible future directions in text summarization.

Keywords: Natural language processing (NLP), Text summarization, Automatic text Summarization, Extractive Method (EXT), Abstractive Method (ABS), Deep learning).





Introduction:

Today, there is so much digital data that it is impossible to quantify it all, and therefore, there is a huge need for effective techniques for summarizing this data. Text summarization pursues a significant and fruitful application of Natural Language Processing (NLP) and aims at producing more concise yet informative summaries from large textual content without altering or distorting its meaning [1]. Manual summarization is usually very time-consuming and impractical in the present scenario with so much information available.

Automatic text summarization is particularly concerned with that which applies computational techniques to extraction or generation of a summary facilitating much faster information retrieval [2]. Over the course of time in the field, such a lot has changed from earlier statistical models to modern deep learning ones, having outlined applications such as news summarization, document condensation, and legal document processing to biomedical text summarization [3].

Exploration in this area is progressing towards building and developing various tools, techniques, and models over the years. There are mainly two types of approaches: Extractive summarization, in which important sentences are selected directly from the source text itself, and Abstractive summarization, which generates new sentences using its content reinterpretation [4]. While the methods of extraction have become widely studied and put into application, abstractive summarization offers a wider experimental challenge as it revolves around the task of natural language generation found in deep learning model [5]. Traditional approaches to natural language processing, including TF-IDF, Latent Semantic Analysis (LSA), and some graph-based ranking techniques (PageRank and its kind), have been used for extractive summarization; in contrast, some of the deep learning-based models that have been used more recently include RNNs, LSTMs, Transformers, and Reinforcement Learning models [6].

The paper reviewed deep-learning-based models for text summarization, emphasizing Transformer-based architectures, BERT, T5, and BART, which have shown substantial progress in generating coherent and context-rich summaries [7], [8]. These models rely on pretrained embedding schemes, attention mechanisms, and finetuning methods to provide better outcomes than traditional extractive methods [9]. Also explored are hybrid approaches that seek to integrate extractive and abstractive methods for improvement in summary quality and relevance [10].

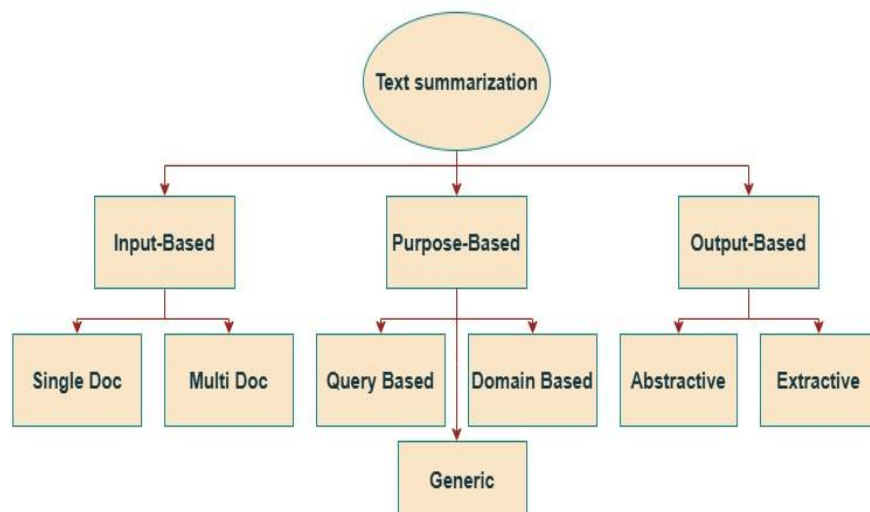


Figure 1. Classification of Text Summarization Techniques.

Figure 1 illustrates a structured classification of text summarization techniques based on three key dimensions: input, purpose and output. Input-based summarization is divided

into single-document and multi-document approaches, depending on the number of source texts. Purpose-based summarization includes query-based, domain-based, and generic types, reflecting the intent behind the summary. Lastly, output-based summarization is categorized into abstractive and extractive methods, where abstractive techniques generate new sentences while extractive methods select existing ones from the source text.

This classification offers a comprehensive view of how summarization strategies vary based on their context and implementation. A comparative analysis of various text summarization models highlights key strengths and limitations. Extractive methods are computationally efficient and maintain sentence integrity but often lack coherence in generated summaries. Abstractive models, powered by deep learning, offer humanlike summarization capabilities but require extensive training and large datasets [11]. Recent advancements in pretraining strategies, transfer learning, and knowledge distillation have improved the generalization capabilities of abstractive models, making them more effective in diverse application domains [12]. However, challenges such as redundancy, grammatical inconsistencies, and domain-specific adaptation remain areas of ongoing research [13]. As the image illustrates a hierarchical classification of text Summarization based on three key criteria: Input-Based, Purpose-Based, and Output-Based [14].

Literature Review:

In recent years, the area of text summarization has made significant strides, especially in extractive and abstractive summarization techniques. The extractive forms, which are based on the information retrieval paradigm of selecting key sentences from a document, are popular because of their simplicity and effectiveness [4]. These methods have been restricted by coherence and readability and have, therefore, motivated investigators to look toward abstractive summarization, which builds new sentences keeping the original meaning intact [7]. Backed by deep learning models, such as the sequence-to-sequence architecture, transformers, and reinforcement learning methods, abstractive summarization gained fresh momentum [8].

Recent works have been concentrated on enhancing summarization accuracy with deep learning-based techniques, such as Long Short-Term Memory (LSTM) networks, RNNs, and Transformer model [12]. UniLM and BERTSUM performed better in similar types of tasks for abstractive summarization [15]. And for extractive summarization, techniques, such as TextRank and Weighted Term Frequency, have been extensively implemented for various languages including Urdu [6]. Abstractive summarization, mimicking human cognition, generates new sentences encapsulating the essence of the original text. However, there has been an escalation in digitized information credibly warranting the automation of summarization methods for minimum retrieval [16].

Text summarization plays a crucial role in Natural Language Processing (NLP), addressing the growing demand for efficient information processing. As digital text continues to expand, automated summarization techniques extract meaningful insights from large textual data [14]. This sort of summarization is important in Natural Language Processing (NLP) because it has the potential to increase the power of information handling as these sources increase accessibility of digital text. Automated summarization methods will extract important bits of information from massive text. As far as developing extractive summarization, people are focusing on new statistical techniques and graph-based ranking methods, along with new deep learning methods that weight different text portions [17]. Despite its effectiveness, extractive summarization often lacks coherence, as it selects sentences without restructuring them for readability and logical flow [18].

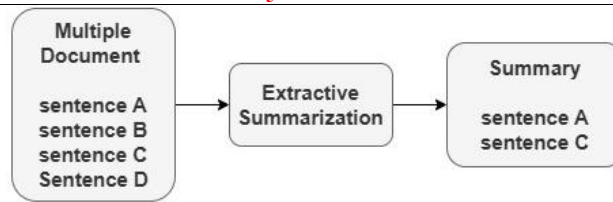


Figure 2. Extractive summarizationⁱ.

Figure 2 demonstrates the process of extractive summarization applied to multiple documents. In this approach, the system selects key sentences (such as Sentence A and Sentence C) directly from the original content without altering them. These selected sentences are then compiled to form a concise summary that retains the most important information from the source material. This method ensures the summary remains factually accurate and contextually grounded.

Deep learning models, RNNs, LSTMs, and Transformer-based models like BERTSUM and UniLM, apply attention mechanisms to produce fluent and contextually meaningful summaries [19]. Abstractive summarization is still computationally expensive and prone to a summary being inaccurate or misleading because it parses out every process of semantic understanding [20].

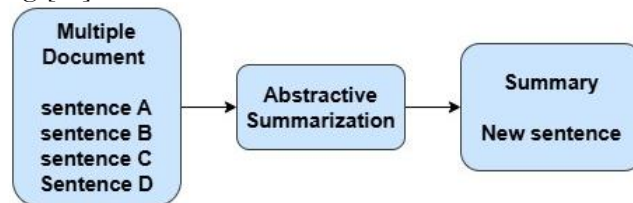


Figure 3. Abstractive Summarizationⁱⁱ.

Figure 3 illustrates the process of abstractive summarization from multiple documents. Unlike extractive methods, this approach does not directly select sentences from the source. Instead, it interprets the meaning of the content and generates a completely new sentence that conveys the core idea. This results in a summary that is more coherent, concise, and closer to how a human would write.

Recognizing the strengths and weaknesses of both approaches, researchers are increasingly focusing on hybrid models that combine extractive and abstractive techniques. These models use extraction-based pre-processing to identify key sentences, followed by abstractive techniques to refine and rephrase them for improved coherence [21]. Reinforcement learning and self-supervised learning methods have been explored to enhance summarization, allowing models to fine-tune their performance based on human-like evaluation metrics [22].

Future research should prioritize enhancing the linguistic quality of generated summaries while reducing computational complexity. Integrating domain-specific knowledge bases, multilingual support, and real-time summarization capabilities are promising areas of exploration [23]. Additionally, ethical considerations, such as bias mitigation and factual accuracy, will be critical for ensuring reliable automated summarization systems [24].

The reviewed literature spans multiple domains, including gamification for recruitment and training, IoT in smart farming, search result diversification, melanoma classification using deep learning, and advancements in automatic text summarization techniques. Each study provides valuable insights into its respective field, highlighting key challenges, technological advancements, and future research directions.

Gamification for Recruitment and Job Training Gamification:

Gamification has been increasingly adopted to enhance recruitment and job training by integrating gaming elements into non-gaming contexts. The reviewed study systematically analyzes gamification's role, emphasizing its ability to improve employee engagement and skill

acquisition [25]. However, challenges such as bias in recruitment processes and the need for more automated, adaptive gamification techniques remain key areas for future research [26].

IoT in Smart Farming:

The adoption of IoT in agriculture has revolutionized smart farming by enabling real-time monitoring, automation, and data-driven decision-making [27]. The reviewed study categorizes IoT technologies used in agriculture, emphasizing their role in resource optimization, precision farming, and security concerns. Future research should focus on integrating AI-driven predictive analytics with IoT to enhance sustainability and efficiency [28].

Search Result Diversification Techniques:

The increasing proliferation of web data necessitates effective search result diversification strategies to enhance user satisfaction. The reviewed study presents a taxonomy of diversification methods based on query types and algorithms, highlighting the balance between relevance and diversity [29]. Further refinement in hybrid models integrating machine learning with user behavior analysis is needed.

Melanoma Classification Using Deep Learning:

Deep learning has shown significant potential in melanoma classification, improving diagnostic accuracy and early detection [30], [31]. The reviewed study systematically examines CNN-based melanoma detection methods, comparing their effectiveness across various datasets. Challenges such as dataset biases, model interpretability, and the need for robust validation mechanisms remain key areas for improvement [32].

Automatic Text Summarization:

Advancements in automatic text summarization, particularly with large language models (LLMs), have transformed information extraction efficiency [33]. The reviewed studies provide a process-oriented framework covering extractive and abstractive techniques, optimization-based approaches, and LLM-driven summarization. Key challenges include computational efficiency, coherence in generated summaries, and evaluation metrics aligning with human judgment [34]. Future research should enhance fine-tuned LLMs for domain-specific summarization and real-time applications [35].

Recent studies have advanced automatic text summarization by leveraging large language models (LLMs) that integrate extractive and abstractive techniques, improving efficiency and quality [36]. Optimization methods now balance coherence, relevance, and conciseness while addressing computational and accuracy challenges [1]. Fine-tuning LLMs for domain-specific tasks, such as legal and medical texts, and real-time summarization of dynamic sources like social media, have become key research areas [37]. Hybrid models combining extractive filtering with abstractive rewriting show improved readability and coherence, enhanced by reinforcement learning and self-supervised training aligned with human evaluation metrics [38], [18]. Ethical concerns like bias mitigation and factual consistency are increasingly prioritized, especially in sensitive fields [39], [40].

Summarization technologies are expanding into related domains: gamification enhances training and recruitment engagement [41]; IoT in smart farming uses summarization for actionable insights [42]; search result diversification benefits from summary clustering [43]; and medical imaging combines summarization with deep learning for diagnostic support [44]. These developments highlight the rapid evolution of summarization as a multidisciplinary tool, with future work focusing on multilingual support, real-time processing, and robust evaluation frameworks [45].

Objectives:

The main goal of this study is as follows:

- To explore and identify various text summarization techniques within the field of Natural Language Processing.

- To provide a detailed overview of how different datasets have been utilized over time with various summarization methods.
- To identify the gaps and challenges encountered by different text summarization methods when summarizing text documents.

Novelty:

The novelty of this work lies in its comparative evaluation of extractive, abstractive, and hybrid approaches using deep learning models including transformers, reinforcement learning, and large language models (LLMs), across diverse datasets and application domains. Unlike prior reviews, this study highlights evolving trends, identifies key limitations, and proposes future research directions with a particular focus on multilingual support, ethical considerations, and domain-specific summarization, offering a holistic foundation for next-generation summarization systems.

Methodology:

A Systematic Literature Review (SLR) has been selected as the research methodology for this study as the workflow is shown in Figure 4. The primary objective of this analysis is to consolidate existing research contributions and develop a model capable of summarizing text across various application domains.

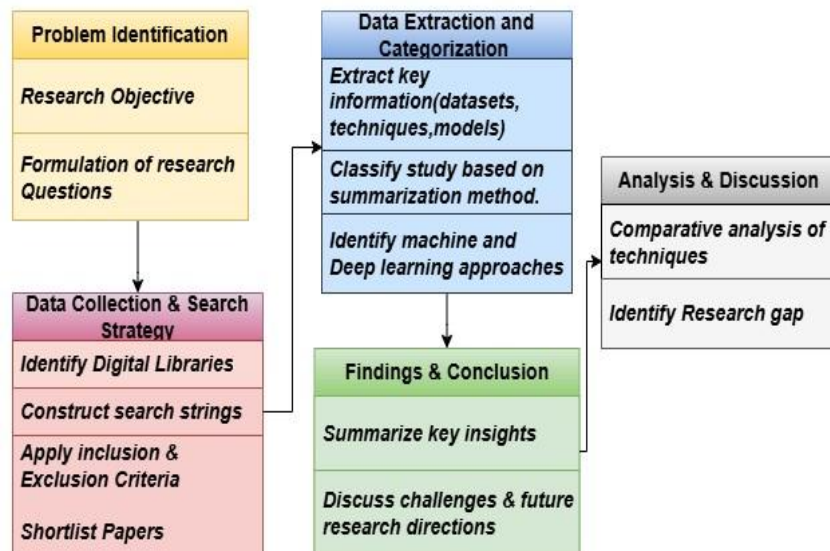


Figure 4. SLR Process Model.

To enhance the rigor and impact of our study, we have adopted the methodology outlined by [16] for structured study selection and result evaluation. Following the formulation of research questions, a well-defined search protocol was established. The detailed methodology for this systematic review is depicted in Figure 4.

Research Questions:

To effectively conduct this Systematic Literature Review (SLR), the key research questions have been carefully identified. Various questions have been formulated, specifically aligning with the focus of our study on Text Summarization. This SLR explores several critical research questions.

- 1: What are the various types of text summarization techniques applied across different domains?
- 2: Which datasets are most commonly used in text summarization?
- 3: What are the key issues and challenges associated with text summarization?
- 4: How can deep learning models be applied to text summarization?

This Systematic Literature Review explores several key research questions along with their underlying motivations:

To examine different text summarization techniques within the scope of the SLR.

To analyze various datasets utilized in the SLR and provide an in-depth discussion on them.

To identify and elaborate on research challenges and implementation issues in detail.

To investigate the applicability of deep learning models for text summarization.

Search Strategy:

The second phase of the Systematic Literature Review (SLR) involves formulating a search strategy and collecting relevant research articles within the field of text summarization.

Table 1. Terms and keywords used in search.

Repository	Search Keywords	Search Strings	No of papers
ACM Digital Library	("Text Summarization" OR "Automatic Summarization") AND ("Machine Learning" OR "Deep Learning" OR "NLP") AND ("Feature Extraction" OR "Text Processing" OR "Semantic Analysis")	("Text Summarization" OR "Automatic Summarization") AND ("Machine Learning" OR "Deep Learning" OR "NLP") AND ("Feature Extraction" OR "Text Processing" OR "Semantic Analysis")	4,500
Elsevier	TITLE-ABS-KEY ("Text Summarization" AND "Deep Learning" OR "Machine Learning") AND ("NLP" OR "Text Processing")	TITLE-ABS-KEY ("Text Summarization" AND "Deep Learning" OR "Machine Learning") AND ("NLP" OR "Text Processing")	6,000
Springer	("Automatic Summarization" OR "Text Summarization") AND ("NLP Applications" OR "Information Retrieval") AND ("Feature Extraction" OR "Semantic Analysis")	("Automatic Summarization" OR "Text Summarization") AND ("NLP Applications" OR "Information Retrieval") AND ("Feature Extraction" OR "Semantic Analysis")	3,800
IEEE Xplore	("Text Summarization" OR "Summarization Techniques") AND ("Natural Language Processing" OR "Text Mining") AND ("Deep Learning" OR "Machine Learning")	("Text Summarization" OR "Summarization Techniques") AND ("Natural Language Processing" OR "Text Mining") AND ("Deep Learning" OR "Machine Learning")	5,200
Google scholar	allintitle: ("Text Summarization" AND "Deep Learning" OR "Natural Language Processing")	allintitle: ("Text Summarization" AND "Deep Learning" OR "Natural Language Processing")	10,000

A keyword-based search was conducted using terms such as "text summarization techniques using NLP" to identify relevant studies. Articles were sourced from multiple digital libraries, including IEEE, Springer, ScienceDirect, and Google Scholar. A comprehensive internet search was also performed to ensure the inclusion of diverse and high-quality research papers. The selected articles, which form the foundation of this study, are summarized in Table 1.

Inclusion and Exclusion Criteria:

In this Systematic Literature Review (SLR), emphasis is placed on selecting high-quality research papers that explore various text summarization approaches, incorporating different techniques and methodologies. The following inclusion criteria have been established to ensure the relevance and quality of the selected studies.

- Research articles focused on text summarization across different languages.
- Studies presenting various methodologies applied to different types of datasets from 2015 to 2024.
- Research papers discussing text summarization techniques, including Machine Learning (ML), Natural Language Processing (NLP), and Deep Learning (DL).

Duplicate papers appearing across multiple sources were eliminated, ensuring that only those relevant to our study and search criteria were retained. Additionally, papers that did not align with our research focus were excluded. The following exclusion criteria have been established to refine the selection process.

- Research articles that do not focus on text summarization.
- Studies primarily centered on text classification methods and approaches.
- Articles that do not be from previous years.
- Papers not published in the English language were also excluded.

Quality criteria:

Assessing the quality of the selected studies is a critical step in the Systematic Literature Review (SLR) as given below in table 2. Each included study underwent a quality evaluation based on predefined criteria to ensure its relevance and reliability.

Table 2. Quality assessment score.

Criteria	Description	Rank	Score
Internal scoring			
a)	Did the abstract clearly define the method of proposed solution?	Yes Partially No	1.5 1 0
b)	Did the study show comparison of the particular method with previously defined methods?	Yes Partially No	1.5 1 0
c)	Was methodology clearly defined?	Yes Partially No	1.5 1 0
d)	Was the conclusion based on results?	Yes, Partially No	1.5 1 0
External scoring			
What is the ranking of the publication source?	Q1Q2Q3Q4		2
	Core ACore BCore C		1.5
			1
			1
			1.5
			0.5

Results and discussion:

In the process of conducting this systematic literature review, we analyzed various published studies. Based on this analysis, we formulated several research questions, which are explained in detail below.

1: What are the various types of text summarization techniques applied across different domains?

Depending on the application, various summarization techniques can be utilized and categorized accordingly, as illustrated in Figure 5. Apart from the commonly known abstractive and extractive methods, multiple other types of summaries exist that cater to specific use cases. A few of these are described below:

- **Detail-Oriented Summarization:** Summaries can be indicative or informative, depending on the level of detail. An indicative summary offers a brief glimpse into the content of a document, highlighting only its core message, much like the synopsis found on the back cover of a novel. In contrast, an informative summary provides a more condensed and detailed version of the full content, without simply skimming over the original material.

- **Content-Focused Summarization:** Based on content relevance, summaries may be generic or query-based. A generic summary delivers information solely from the author's perspective, treating all content equally and suitable for general use. On the other hand, a query-based summary is designed to answer specific questions posed by the user. It extracts information relevant to a user-defined query, tailoring the output to the reader's interests.
- **Input-Driven Summarization:** Depending on the source, summaries can be created from a single document or multiple documents. A single-document summary processes and summarizes text from one source, making it relatively straightforward. Conversely, multi-document summarization consolidates information from several related texts, which increases the complexity but allows for a broader synthesis of information.
- **Language-Based Summarization:** Summaries can also be monolingual or multilingual. A monolingual summary is restricted to processing and generating content in one language only, limiting its applicability across diverse linguistic datasets. In contrast, a multilingual summary supports input in multiple languages, although this significantly increases the complexity of implementation

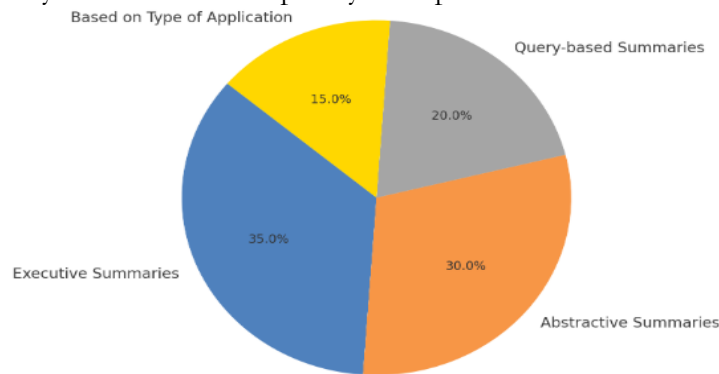


Figure 5. Studied type of summarization.

2: Which datasets are most commonly used in text summarization?

In this systematic literature review (SLR), research articles published between 2016 and 2024 were examined to understand the evolution of datasets and models used for text summarization. Over the years, various datasets have been introduced to support the training and evaluation of summarization systems. The availability of diverse and high-quality datasets has played a critical role in benchmarking models and fine-tuning their performance for different applications. Without structured and labeled data, it becomes challenging to validate model efficiency or optimize algorithmic parameters for varying use cases.

Recent advancements, particularly the emergence of transformer-based models and large language models (LLMs) like GPT, BART, and T5, demand large-scale datasets to achieve optimal performance. The following are some notable datasets widely utilized in the field:

- **DUC (Document Understanding Conference):** This early dataset includes approximately 500 news articles with human-authored summaries. While valuable for benchmarking, it is limited in size and not suitable for training modern, parameter-rich models. Thus, it is commonly used in combination with larger datasets to support evaluation tasks.
- **Gigaword:** A large-scale dataset comprising over 10 million news articles. It has been widely adopted for training abstractive summarization models, especially in the context of headline generation. Its scale makes it suitable for deep learning-based models, including LSTM and Transformer variants.

- **NEWSROOM:** Featuring more than 1.3 million articles, this corpus integrates both extractive and abstractive summarization techniques. Developed to reflect real-world editorial strategies, it provides high-quality summaries drawn from a variety of publishers and editorial styles, making it suitable for fine-tuning neural models.
 - **CAST Dataset:** Designed to assist with sentence-level decisions, the CAST corpus labels individual sentences as essential or non-essential. This binary classification structure aids in developing sophisticated sentence selection and compression algorithms, which are particularly useful in hybrid summarization models.
 - **CNN/Daily Mail:** One of the most widely used datasets in recent years, this collection of news stories and bullet-point summaries has become a standard benchmark for abstractive summarization tasks. It supports both supervised training and evaluation and has been instrumental in training models like PEGASUS and BART.
 - **XSum (Extreme Summarization):** A recent addition to the field, XSum provides one-sentence summaries for BBC articles, challenging models to generate highly concise yet informative outputs. It has been used extensively for evaluating LLMs and their performance in generating abstract, focused summaries.
 - **MultiLing and WikiLingua:** These multilingual datasets support cross-lingual summarization tasks, reflecting the growing interest in globalized content processing. They are used to evaluate the summarization capabilities of models in languages beyond English.
 - **Social media & Domain-Specific Datasets:** Text data from platforms such as Twitter (tweets), Reddit (threads) and specialized domains like biomedical (PubMed, arXiv) and legal texts are increasingly being incorporated into summarization research. These datasets enable fine-tuning of models for domain-specific summarization, offering more precise and relevant outputs.
 - As illustrated in Figure 6, datasets vary significantly depending on the source, structure, and intended summarization strategy—whether generic, query-based, extractive, abstractive, or hybrid. The current trend also emphasizes the use of synthetic dataset generation using LLMs and data augmentation strategies to overcome limitations in labeled data availability.
 - **3:** What are the key issues and challenges associated with text summarization?
- Text summarization faces numerous research challenges and practical implementation hurdles. One of the major difficulties arises in multi-document summarization, where issues such as data repetition, improper sentence sequencing, and grammatical inconsistencies can hinder the creation of a coherent and high-quality summary.

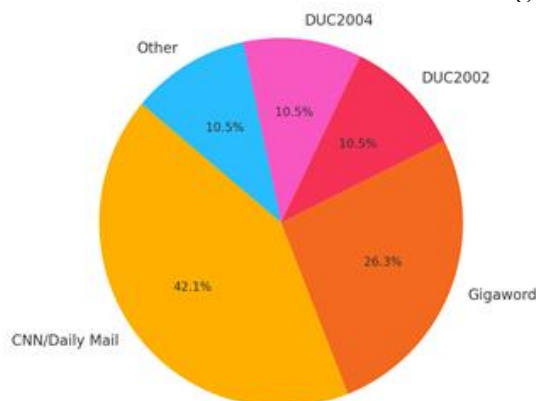


Figure 6. Datasets used over the years.

Additionally, the perceived quality of a summary often varies depending on the system or the individual evaluating it. Different users may prioritize different sets of sentences, leading

to subjective variations in what is considered essential information. Another frequent concern is the inclusion of irrelevant content, which undermines the core objective of summarization, to extract relevant and meaningful information.

Moreover, achieving full topic coverage is another challenge; summarizing all aspects of a document can introduce unnecessary repetition, while selective summarization may result in the omission of key points. The generation of effective summaries also heavily depends on identifying and utilizing high-quality keywords that accurately represent the main ideas. In recent years, the rise of transformer-based models like BERT, T5, and GPT has significantly improved the performance of text summarization systems. These models can understand contextual relationships better and generate more coherent summaries. However, they bring new challenges, such as the need for large computational resources, risks of generating hallucinated content (i.e., text that appears plausible but is factually incorrect), and difficulties in maintaining factual consistency.

Furthermore, despite these advancements, achieving domain-specific summarization such as in legal, medical, or financial texts, remains complex due to the specialized terminology and context required. Another emerging issue is the ethical aspect of summarization, especially in news or opinion summarization, where bias and misinformation can easily be amplified. Ensuring fairness, neutrality, and explaining ability in AI-generated summaries is now a growing area of concern. As real-time summarization becomes more common in applications like live feeds, social media, and customer service chats, latency and responsiveness also become critical factors, pushing the need for lightweight, efficient summarization models suitable for deployment on edge devices. Overall, while the field continues to evolve rapidly, researchers and developers must address both longstanding challenges and emerging issues to make text summarization more reliable, ethical, and practical across diverse real-world scenarios.

4: How can deep learning models be applied to text summarization?

Text summarization is typically accomplished through natural language processing (NLP) techniques, often relying on algorithms such as the PageRank algorithm. While these methods achieve the basic goals of summarization, they are limited in their ability to create new sentences or restructured paragraphs that do not already exist in the source text, unlike human-written summaries. Additionally, such approaches may result in grammatical inaccuracies. The use of deep learning has significantly enhanced the effectiveness and speed of summarization models. Deep learning models can generate comprehensive summaries that are both coherent and grammatically correct. When combined with fuzzy logic, the performance and precision of these models are further improved. During the training phase, the model learns to generate summaries based on input documents using deep learning techniques integrated with a fuzzy logic classifier.

The training phase in modern text summarization systems leverages advanced deep learning architectures, particularly transformer-based models like BERT, T5, and GPT to produce high-quality summaries. These systems are trained on large-scale datasets using supervised and reinforcement learning approaches. Recent developments also incorporate contextual embeddings, fine-tuning with transfer learning, and hybrid techniques involving fuzzy logic and attention mechanisms to enhance summary relevance and grammatical accuracy. The features used during training are extracted from various document characteristics and are crucial for improving model performance. As of 2025, these features include traditional linguistic indicators as well as semantic and contextual signals derived from pretrained language models.

Comparison of existing solutions:

This section highlights various techniques and algorithms employed by researchers to train models for generating summaries of lengthy articles. A comparative overview of these

methods, based on specific datasets and scholarly works reviewed over the years, is provided in Table 4.

Table 3. Training features

Token Importance Scoring	Assigning significance weights to words and sentences using attention maps
Title-Content semantic Alignment	Evaluating how well the content aligns with the title using embeddings
Sentence Embedding Similarity	Measuring similarity to the centroid using contextual sentence embeddings
Entity and Numeral Density	Counting named entities and numerical values to gauge importance
POS Tag Distribution	Analyzing part-of-speech patterns for syntactic relevance
Redundancy and similarity Clustering	Identifying repetitive or semantically similar points for filtering

Table 4. Comparison of techniques over years.

Years	Methods	Dataset	Evaluations
2016	Human Learning Optimization Algorithm	Text documents	Aimed to reduce sentence redundancy and highlight core ideas to produce concise summaries.
2017	K-nearest algorithm	Paragraph as input	Found useful in fields like medical and legal summarization due to strong similarity detection.
2017	Word vector embedding with neural networks	CNN News Corpus with abstractive summaries	Outperformed other models in ROUGE scores; captured semantic meaning better.
2018	Query-based extractive summarization using TF-IDF & fuzzy logic	DUC2007 corpus	Produced precise summaries (around 250 words) with improved matching to ROUGE benchmarks.
2020	Extractive summarization in Urdu using deep learning	Abstractive Urdu novel corpus	Addressed growing demand for regional-language summarization; popular in Urdu-speaking audience.
2021	Transformer-based Abstractive Summarization (BERTSUM, PEGASUS)	CNN/Daily Mail, XSUM	Models generated coherent, human-like summaries with state-of-the-art ROUGE performance.
2022	Reinforcement Learning-based Summarization with Reward Optimization	Multi-News, Reddit TIFU	Enhanced long-document summarization and factual consistency using policy gradient training.
2023	Multilingual Summarization with mT5	WikiLingua, MLSum	Enabled accurate summarization across 10+ languages; reduced dependency on English-only data.
2024	Zero-shot Summarization using GPT-4 and	Diverse web articles (no fine-tuning)	Generated high-quality summaries with zero-shot capabilities and instruction-based prompting.

	Instruction Tuning		
2025	Hybrid Extractive-Abstractive Pipeline with Knowledge Graph Integration	Scientific and medical research datasets (PubMed, arXiv)	Combines factual accuracy (from extractive) and fluency (from abstractive) using external graphs.

Challenges and Conclusion:

Text summarization continues to face various challenges and limitations, particularly when dealing with lengthy documents or multiple sources. One major issue is redundancy, where summaries often contain repetitive information that reduces clarity and conciseness. Minimizing redundancy by analyzing the similarity between elements in the original text can significantly improve summary quality. Another common problem is the inclusion of irrelevant content, which undermines the goal of summarization, to deliver meaningful and concise overviews. Including too many features from the original text may lead to off-topic or unnecessary information, so selecting the most relevant aspects is crucial. Additionally, achieving comprehensive topic coverage is difficult, especially in multi-document summarization. While trying to cover all topics, summaries may become cluttered with redundant or irrelevant details. General-purpose summaries aim for broader coverage, but this does not always guarantee better quality, whereas query-based summarization is more selective and focused. Recent summarization systems often struggle with balancing coverage and conciseness, particularly when dealing with diverse topics. A good summary should maintain coherence, logical flow, and readability to ensure the content is easy to understand. This systematic literature review (SLR) presented a range of summarization methods, including extractive, abstractive, and hybrid approaches, tailored to different application needs. The study analyzed 24 peer-reviewed articles using structured criteria and proposed a simplified framework for generating summaries without requiring domain-specific knowledge. With the explosive growth of digital content, summarization plays a vital role in saving users' time and improving access to key information. Although many algorithms and techniques have been developed, the generated summaries are not always perfectly accurate or contextually aligned with the source material. Existing models still face issues with consistency, coverage and relevance. As summarization remains an evolving field, continued research is essential, and future advancements are expected to lead to more refined, intelligent and domain-aware summarization systems.

References:

- [1] M. Gambhir and V. Gupta, "Recent automatic text summarization techniques: a survey," *Artif. Intell. Rev.*, vol. 47, no. 1, pp. 1–66, Jan. 2017, doi: 10.1007/S10462-016-9475-9/METRICS.
- [2] N. Moratanch and S. Chitrakala, "A survey on extractive text summarization," *Int. Conf. Comput. Commun. Signal Process. Spec. Focus IoT, ICCCSA 2017*, Jun. 2017, doi: 10.1109/ICCCSP.2017.7944061.
- [3] R. Mishra *et al.*, "Text summarization in the biomedical domain: A systematic review of recent research," *J. Biomed. Inform.*, vol. 52, pp. 457–467, Dec. 2014, doi: 10.1016/J.JBI.2014.06.009.
- [4] R. Nallapati, B. Zhou, C. dos Santos, Ç. Gulçehre, and B. Xiang, "Abstractive Text Summarization Using Sequence-to-Sequence RNNs and Beyond," *CoNLL 2016 - 20th SIGNLL Conf. Comput. Nat. Lang. Learn. Proc.*, pp. 280–290, Feb. 2016, doi: 10.18653/v1/k16-1028.
- [5] A. W. Palliyali, M. A. Al-Khalifa, S. Farooq, J. Abinahed, A. Al-Ansari, and A. Jaoua, "Comparative Study of Extractive Text Summarization Techniques," *Proc. IEEE/ACS Int. Conf. Comput. Syst. Appl. AICCSA*, vol. 2021-December, 2021, doi: 10.1109/AICCSA53542.2021.9686867.

- [6] R. Mihalcea and P. Tarau, "TextRank: Bringing Order into Text," 2004. Accessed: Jun. 05, 2025. [Online]. Available: <https://aclanthology.org/W04-3252/>
- [7] A. See, P. J. Liu, and C. D. Manning, "Get To The Point: Summarization with Pointer-Generator Networks," *ACL 2017 - 55th Annu. Meet. Assoc. Comput. Linguist. Proc. Conf. (Long Pap.*, vol. 1, pp. 1073–1083, Apr. 2017, doi: 10.18653/v1/P17-1099.
- [8] Y. Liu and M. Lapata, "Text Summarization with Pretrained Encoders," *EMNLP-IJCNLP 2019 - 2019 Conf. Empir. Methods Nat. Lang. Process. 9th Int. Jt. Conf. Nat. Lang. Process. Proc. Conf.*, pp. 3730–3740, Aug. 2019, doi: 10.18653/v1/d19-1387.
- [9] P. J. L. Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, "Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer," *arXiv:1910.10683*, 2023, doi: <https://doi.org/10.48550/arXiv.1910.10683>.
- [10] A. G. Aakash Sinha, Abhishek Yadav, "Extractive text summarization using neural networks," *arXiv:1802.10137*, 2018, doi: <https://doi.org/10.48550/arXiv.1802.10137>.
- [11] D. Suleiman and A. Awajan, "Deep Learning Based Abstractive Text Summarization: Approaches, Datasets, Evaluation Measures, and Challenges," *Math. Probl. Eng.*, vol. 2020, no. 1, p. 9365340, Jan. 2020, doi: 10.1155/2020/9365340.
- [12] L. Dong *et al.*, "Unified Language Model Pre-training for Natural Language Understanding and Generation," *Adv. Neural Inf. Process. Syst.*, vol. 32, May 2019, Accessed: Jun. 05, 2025. [Online]. Available: <https://arxiv.org/pdf/1905.03197>
- [13] M. L. Y. Zhang, "A neural approach to multidocument summarization," *Proc. ACL*, vol. 6231– 6240, 2019.
- [14] V. Gupta and G. S. Lehal, "A Survey of Text Summarization Extractive techniques," *J. Emerg. Technol. Web Intell.*, vol. 2, no. 3, pp. 258–268, Aug. 2010, doi: 10.4304/JETWI.2.3.258-268.
- [15] J. Devlin, M. W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of deep bidirectional transformers for language understanding," *NAACL HLT 2019 - 2019 Conf. North Am. Chapter Assoc. Comput. Linguist. Hum. Lang. Technol. - Proc. Conf.*, vol. 1, pp. 4171–4186, 2019.
- [16] N. Jayatilke and R. Weerasinghe, "A Hybrid Architecture with Efficient Fine Tuning for Abstractive Patent Document Summarization".
- [17] S. Narayan, S. B. Cohen, and M. Lapata, "Ranking Sentences for Extractive Summarization with Reinforcement Learning," *NAACL HLT 2018 - 2018 Conf. North Am. Chapter Assoc. Comput. Linguist. Hum. Lang. Technol. - Proc. Conf.*, vol. 1, pp. 1747–1759, 2018, doi: 10.18653/V1/N18-1158.
- [18] Q. Wang, P. Liu, Z. Zhu, H. Yin, Q. Zhang, and L. Zhang, "A Text Abstraction Summary Model Based on BERT Word Embedding and Reinforcement Learning," *Appl. Sci. 2019, Vol. 9, Page 4701*, vol. 9, no. 21, p. 4701, Nov. 2019, doi: 10.3390/APP9214701.
- [19] R. Kim, A. S. Angelidis, and M. Lapata, "Aspect-Controllable Opinion Summarization," pp. 6578–6593, Accessed: Jun. 05, 2025. [Online]. Available: <https://github.com/rktamplayo/AceSum>
- [20] R. Paulus, C. Xiong, and R. Socher, "A Deep Reinforced Model for Abstractive Summarization," *6th Int. Conf. Learn. Represent. ICLR 2018 - Conf. Track Proc.*, May 2017, Accessed: Jun. 05, 2025. [Online]. Available: <https://arxiv.org/pdf/1705.04304>
- [21] N. Shirish Keskar, B. Mccann, L. R. Varshney, C. Xiong, R. Socher, and S. Research, "CTRL: A Conditional Transformer Language Model for Controllable Generation," Sep. 2019, Accessed: Jun. 05, 2025. [Online]. Available: <https://arxiv.org/pdf/1909.05858>
- [22] S. Verma and V. Nidhi, "Extractive Summarization using Deep Learning," Aug. 2017,

- Accessed: Jun. 05, 2025. [Online]. Available: <https://arxiv.org/pdf/1708.04439>
- [23] A. R. Fabbri, W. Kryściński, B. McCann, C. Xiong, R. Socher, and D. Radev, "SummEval: Re-evaluating Summarization Evaluation," *Trans. Assoc. Comput. Linguist.*, vol. 9, pp. 391–409, Feb. 2021, doi: 10.1162/TACL_A_00373.
 - [24] D. Ognibene *et al.*, "Challenging social media threats using collective well-being-aware recommendation algorithms and an educational virtual companion," *Front. Artif. Intell.*, vol. 5, Jan. 2023, doi: 10.3389/FRAI.2022.654930.
 - [25] J. Hamari, J. Koivisto, and H. Sarsa, "Does Gamification Work? -- A Literature Review of Empirical Studies on Gamification," *2014 47th Hawaii Int. Conf. Syst. Sci.*, pp. 3025–3034, Jan. 2014, doi: 10.1109/HICSS.2014.377.
 - [26] I. Obaid, M. S. Farooq, and A. Abid, "Gamification for Recruitment and Job Training: Model, Taxonomy, and Challenges," *IEEE Access*, vol. 8, pp. 65164–65178, 2020, doi: 10.1109/ACCESS.2020.2984178.
 - [27] S. Terence and G. Purushothaman, "Systematic review of Internet of Things in smart farming," *Trans. Emerg. Telecommun. Technol.*, vol. 31, no. 6, Jun. 2020, doi: 10.1002/ETT.3958.
 - [28] V. Partel, S. Charan Kakarla, and Y. Ampatzidis, "Development and evaluation of a low-cost and smart technology for precision weed management utilizing artificial intelligence," *Comput. Electron. Agric.*, vol. 157, pp. 339–350, Feb. 2019, doi: 10.1016/j.compag.2018.12.048.
 - [29] H. Wu *et al.*, "Result Diversification in Search and Recommendation: A Survey," *IEEE Trans. Knowl. Data Eng.*, Dec. 2022, doi: 10.1109/TKDE.2024.3382262.
 - [30] M. S. Deiner, S. D. Mcleod, J. Wong, J. Chodosh, T. M. Lietman, and T. C. Porco, "Google Searches and Detection of Conjunctivitis Epidemics Worldwide," *Ophthalmology*, vol. 126, pp. 1219–1229, 2019, doi: 10.1016/j.ophtha.2019.04.008.
 - [31] H. A. Haenssle *et al.*, "Man against Machine: Diagnostic performance of a deep learning convolutional neural network for dermoscopic melanoma recognition in comparison to 58 dermatologists," *Ann. Oncol.*, vol. 29, no. 8, pp. 1836–1842, Aug. 2018, doi: 10.1093/annonc/mdy166.
 - [32] A. M. Ibrahim *et al.*, "Skin Cancer Classification Using Transfer Learning by VGG16 Architecture (Case Study on Kaggle Dataset)," *J. Intell. Learn. Syst. Appl.*, vol. 15, no. 3, pp. 67–75, Aug. 2023, doi: 10.4236/JILSA.2023.153005.
 - [33] M. Gambhir and V. Gupta, "Recent automatic text summarization techniques: a survey," *Artif. Intell. Rev.*, vol. 47, no. 1, pp. 1–66, Jan. 2017, doi: 10.1007/S10462-016-9475-9/METRICS.
 - [34] Z. Niu, G. Zhong, and H. Yu, "A review on the attention mechanism of deep learning," *Neurocomputing*, vol. 452, pp. 48–62, Sep. 2021, doi: 10.1016/J.NEUCOM.2021.03.091.
 - [35] S. . Satheeskumaran, Y. Zhang, V. E. Balas, T. Hong, and D. Pelusi, "Intelligent computing for sustainable development: first International Conference, ICICSD 2023, Hyderabad, India, August 25-26, 2023, revised selected papers. Part I," 2024.
 - [36] Y. Shen *et al.*, "ChatGPT and Other Large Language Models Are Double-edged Swords," *Radiology*, vol. 307, no. 2, p. 2023, Apr. 2023, doi: 10.1148/RADIOL.230163/ASSET/IMAGES/LARGE/RADIOL.230163.FIG1.JPEG.
 - [37] R. Girshick, J. Donahue, T. Darrell, J. Malik, U. C. Berkeley, and J. Malik, "1043.0690," *Proc. IEEE Comput. Soc. Conf. Comput. Vis. Pattern Recognit.*, vol. 1, p. 5000, Sep. 2014, doi: 10.1109/CVPR.2014.81.
 - [38] "(PDF) Real-Time Social Media Analytics with Deep Transformer Language Models: A Big Data Approach." Accessed: Jun. 17, 2025. [Online]. Available:

[https://www.researchgate.net/publication/350325104_Real-](https://www.researchgate.net/publication/350325104_Real-Time_Social_Media_Analytics_with_Deep_Transformer_Language_Models_A_Big_Data_Approach)

[Time_Social_Media_Analytics_with_Deep_Transformer_Language_Models_A_Big_Data_Approach](https://www.researchgate.net/publication/350325104_Real-Time_Social_Media_Analytics_with_Deep_Transformer_Language_Models_A_Big_Data_Approach)

- [39] N. Sourlos, J. Wang, Y. Nagaraj, P. van Ooijen, and R. Vliegenthart, "Possible Bias in Supervised Deep Learning Algorithms for CT Lung Nodule Detection and Classification," *Cancers (Basel)*, vol. 14, no. 16, pp. 1–15, 2022, doi: 10.3390/cancers14163867.
- [40] F. Nan *et al.*, "Improving Factual Consistency of Abstractive Summarization via Question Answering," *ACL-IJCNLP 2021 - 59th Annu. Meet. Assoc. Comput. Linguist. 11th Int. Jt. Conf. Nat. Lang. Process. Proc. Conf.*, pp. 6881–6894, May 2021, doi: 10.18653/v1/2021.acl-long.536.
- [41] T. H. Laine and R. S. N. Lindberg, "Designing Engaging Games for Education: A Systematic Literature Review on Game Motivators and Design Principles," *IEEE Trans. Learn. Technol.*, vol. 13, no. 4, pp. 804–821, Oct. 2020, doi: 10.1109/TLT.2020.3018503.
- [42] C. Anjanappa, S. Parameshwara, M. K. Vishwanath, M. Shrimali, and C. Ashwini, "AI and IoT based Garbage classification for the smart city using ESP32 cam," *Int. J. Health Sci. (Qassim)*, vol. 6, no. S3, pp. 4575–4585, May 2022, doi: 10.53730/IJHS.V6NS3.6905.
- [43] M. Drosou and E. Pitoura, "Search result diversification," *SIGMOD Rec.*, vol. 39, no. 1, pp. 41–47, 2010, doi: 10.1145/1860702.1860709.
- [44] R. D. Seeja and A. Suresh, "Deep Learning Based Skin Lesion Segmentation and Classification of Melanoma Using Support Vector Machine (SVM)," *Asian Pac. J. Cancer Prev.*, vol. 20, no. 5, p. 1555, 2019, doi: 10.31557/APJCP.2019.20.5.1555.
- [45] R. Aharoni *et al.*, "mFACE: Multilingual Summarization with Factual Consistency Evaluation," Dec. 2022, Accessed: Jun. 17, 2025. [Online]. Available: <https://arxiv.org/pdf/2212.10622v2>



Copyright © by authors and 50Sea. This work is licensed under Creative Commons Attribution 4.0 International License.