

Research Article

Abstractive Summarization of Long Educational Transcripts in the Era of GenAI

Deepa Fernandes and Rupali Sunil Wagh

Department of Computer Science, Christ University, Bangalore, India

Article history

Received: 17-12-2025

Revised: 15-04-2026

Accepted: 04-08-2026

Corresponding Author:

Deepa Fernandes
Department of Computer
Science, Christ University,
Bangalore, India
Email:
n.deepa@res.christuniversity.in

Abstract: The surge in online video data presents the imperative requirement of summarizing it for information compression, efficient understanding, and to support decision-making. Recent research studies have predominantly utilized structured input data for summarization. This paper explores the current state-of-the-art models in summarizing unstructured, conversational, and long transcript data. We further apply a transfer learning approach using a fine-tuned BART model, pre-trained on the SAMSum dataset. Hyperparameter tuning and selective layer freezing are applied to optimise model performance. This study focuses on abstractive summarization of long video transcript data. Our methodology integrates ChatGPT to generate reference abstractive summaries from long transcript data. A semi-automated pipeline using TextRank is proposed for reference summary generation. The proposed fine-tuned model shows a significant increase in ROUGE scores over the baseline model. The findings suggest that the transfer learning approach is effective for abstractive summarization of long transcript data in real-world conversational domains.

Keywords: NLP, Abstractive Text Summarization, Transfer Learning, BART, Long Transcripts

Introduction

Studies show that online video use has increased over the past few years, especially during and after the COVID-19 pandemic. There is an increase in the number of participants in virtual meetings via the videoconferencing app daily, from 10 million in December 2019 to 300 million four months later (Aiken, 2020). Ever since the COVID-19 pandemic, the educational relevance of video-based learning materials has gained additional traction, with schooling, courses in universities, and courses in adult education centers during national lockdowns (Chick et al., 2020; Merkt et al., 2022). A surge in online learning facilities has increased the generation of educational videos. Such educational videos vary in duration, content, and presentation styles. For instance, lecture videos contain PowerPoint slides and have long durations, in some cases more than 30 minutes (Hamid and Samad, 2015; Seidel and De, 2024). Given the substantial volume of educational video data accumulated, the requirement for content comprehension and concise understanding emerges.

Manual summarization in such cases is difficult and time-consuming. Hence, the utility of automatic text summarization tools has gained prominence.

In the field of NLP, automatic summarization emerges as pivotal, characterised by a multitude of complexities involving both understanding language, such as identifying essential content elements, and generating content, which entails consolidating and reformulating recognised content to create a summary. Artificial Intelligence (AI) has captured significant attention across diverse sectors for its capacity to produce creative and realistic summary outputs. Several AI tools, such as ChatGPT, BART, and T5, have broadened access to Large Language Models, thereby enabling the production of content that resembles human language (Lin et al., 2024). This study investigates the efficacy of state-of-the-art models like BART and T5 in summarizing long video lecture transcripts.

Video transcripts contribute as valuable resources for academic research as text-based representations of spoken content. They are rich sources of information, aiding researchers in analyzing verbal interactions,

discourse patterns, and language use with greater precision and efficiency. In domains such as education, media studies, linguistics, and human-computer interaction, transcripts enable qualitative and quantitative analyses that would be difficult to perform with raw video alone. Additionally, reproducibility and transparency are enhanced through transcripts, which allow other researchers to verify findings and conduct secondary analyses. As the use of multimedia content continues to grow, the integration of video transcripts into research workflows has become critical for data-driven insights and information retrieval (Elahi et al., 2024; Navarrete et al., 2025).

Summarization approaches are broadly classified as extractive and abstractive (El-Kassas et al., 2020). In extractive summarization, all sentences are scored, and high-ranking sentences are extracted as the summary. Whereas abstractive summarization is more creative, with the inclusion of words or sentences that are not part of the source text. Thus, abstractive summarization is computationally more complex than the extractive approach. We explore the abstractive summarization approach for a comprehensive understanding of video transcript data in this study.

The unique features of video transcript data utilized in this study include the length of transcripts, interactiveness, and unstructuredness. Application of summarization models on raw video transcript data resulted in moderate summaries, as shown in Section 4. A possible solution is to create a custom-designed model for the proposed dataset. However, this approach is impractical due to its significant computational complexity and time-consuming nature. Alternatively, recent research highlights the potential of transfer learning techniques as a viable solution. The Transfer Learning approach leverages knowledge from pre-trained models and results in enhanced summary output. This approach has proved to be beneficial in research scenarios with experimental new datasets that can be small and unlabeled (He et al., 2024; Mishra et al., 2023). A pre-trained BART model on the SAMSum dataset (Gliwa et al., 2019) is applied to our transcript dataset.

This paper proposes an implementable solution for summarizing long academic interactive transcripts. The main contributions of this paper include:

- a) Application of state-of-the-art summarization models on a long interactive transcript dataset
- b) Generation of reference summaries using ChatGPT
- c) Evaluation and results analysis using ROUGE scores
- d) Application of the transfer learning approach using hyperparameter tuning and fine-tuning of the BART model, and

- e) Evaluation of summary results via ROUGE scores

Related Work

Automatic text summarization approaches have been applied extensively on news datasets (Jiang and Dreyer, 2024), legal documents (Tyss et al., 2025), research papers, and technical scientific papers (Achkar et al., 2024; Takeshita et al., 2024). The data utilized in the referenced research is structured or sectioned for a better understanding of content and interrelations between sections. For example, scientific papers are structured into sections such as the introduction, experimental details, and conclusion. Target datasets organised in a structure enhance key information identification and leverage their summarization. Recent text summarization approaches have switched from pure text to diverse data like videos and multimodal inputs (Jangra and Hasanuzzaman, 2023). Video data summarization has been made achievable by employing expert annotations, attention mechanisms, and a pointer network. The complexity of the data collected is increased by multimodal data, where interconnections between text, image, audio, and video are being studied for summarization. Video data acts as a rich source of information in this context.

Alternatively, video data can be summarized utilising its transcriptions. For understanding, we broadly classify the transcriptions into two types: Pure text and interactive transcript data. Pure text transcripts are human-generated, noise-free, with clear descriptions and structured content. Pure text transcripts focus more on conveying the videos' subjective content than on interactions. Conversely, video interactive transcript data contains conversations between two or more users in real time (Liu et al., 2024; Lundberg et al., 2022). As a result, interactive transcript data contains interactions, explanations of certain topics, and also noise. In this study, we focus on interactive transcripts and the challenges that emerge in the summarization of interactive transcript data.

Interesting research has been done on the summarization of interactive transcript data. Call transcripts (Ahmed et al., 2024), dialogue summarization (Lundberg et al., 2022), and meeting summarization (Jin et al., 2025) take into account interactive transcript data. Automatic summarization of call transcripts utilizes textual descriptions of audio recordings. DialogueSum consists of real-life scenarios with a dataset containing 13,460 dialogues (Chen et al., 2021). These, however, are not a single long transcript; they consist of snippets of short dialogues between multiple people. Meeting summarization is done on datasets like AMI (Carletta et al., 2005; Lundberg et al., 2022), ICSI, and Committee meetings (Zhang et al., 2022), which employ annotators and come up with query-summary pairs (Zhong et al., 2021). Though annotations promote effective and faster

summarization, it is an expensive affair.

In this paper, we focus on the video transcript data related to education. Several video conferencing apps, such as Zoom, Google Meet, and Webex, offer video recording. We utilize the Webex platform (Cisco Systems, 2025), which contains saved online classroom recordings, with the option to obtain transcriptions in real time. These lecture recordings are rich in information about the subject matter. However, the transcripts also pose technical challenges due to noise, variability, and the need for robust natural language processing techniques. The transcripts are in the form of conversations or dialogues and lack a structured form. Additional challenges related to the dataset used in the study are: 1. The interactive transcript has multiple speaker interactions, 2. The educational video is 45 minutes long, and the transcript datasets average 15126 tokens.

To the best of our knowledge, studies on long-form academic video transcripts are still in their early stages, with limited research in this field. We apply the state-of-the-art models BART and T5 for abstractive transcript summarization (Golia and Kalita, 2023; Lundberg et al., 2022; Saxena and El-Haj, 2023) for abstractive summarization of long transcript data.

Since our target data for the study is unlabeled and small, we come up with two possibilities for interactive transcript data summarization: 1. Creating a new summarization model, which can be time-consuming and computationally complex. 2. Utilizing a pre-trained model on a similar dataset, also known as a transfer learning approach. Transfer learning is an emerging paradigm in natural language processing for its ability to significantly reduce training time and annotation requirements (Goyal et al., 2022; Shen and Wang, 2024; Bhukya and Bhukya, 2024; Song et al., 2024). We utilize the knowledge of a pre-trained model on a similar dataset to improve summarization on our dataset. We employ the BART model trained on the SAMSum dataset on our interactive transcript data. The summarization results are examined in Section 4.

Evaluation of transcript summaries presents novel challenges due to the lack of human-generated or gold summaries for comparison. Acquiring gold standard reference human summaries for extended transcript data is a labour-intensive, cumbersome, and costly process (Abishek et al., 2025; Suarez et al., 2020; Pilault et al., 2020). Alternatively, we utilize ChatGPT to aid our evaluation of the generation of reference summaries, as detailed in Section 3.2. ChatGPT, much like BART and T5, exhibits versatility across a spectrum of text-based tasks. With its robust architecture, ChatGPT adeptly handles tasks ranging from summarization to translation, demonstrating proficiency in generating concise and accurate summaries of lengthy text (Grassini, 2023; Soni and Wade, 2023).

ROUGE metric (Barbella and Tortora, 2022; Kumar Bhukya and Bhukya, 2024; Lin, 2004; Lundberg et al., 2022) is applied to quantitatively evaluate model-generated summaries of transcript data against summaries generated from ChatGPT. We further display the results in Section 4. To further assess the summary qualitatively, the involvement of human experts is a promising direction for future work.

We extend our experiments by applying the transfer learning approach to adapt pre-trained summarization models to our interactive transcript data. The absence of gold summaries in our dataset posed a significant challenge for both model training and evaluation. To address this, we collaborated with domain experts who assisted in creating reference summaries aided by ChatGPT. The process of generating reference summaries is provided in Section 3. We further aim to automate manual transcript reduction by integrating TextRank into our methodology. These reference summaries act as an initial basis to facilitate supervised fine-tuning of our summarization models.

While prior work has achieved significant success in short-text summarization, there remains a gap in effectively adapting these methods to large-interactive-transcript data without an extensive annotated resource, a challenge our research seeks to address.

Materials and Methods

In this section, we describe the dataset, state-of-the-art abstractive summarization models, application of transfer learning, and implementation details of the research study.

Interactive Transcript Dataset

Transcription is the first critical and essential step in our methodology. Many online transcription tools were explored, such as Google Docs (Google, 2025b), ASR via Android Phone (Google, 2025a), and the HappyScribe website (HappyScribe, 2025). Although these online tools are freely available, they are often time-consuming and prone to noise. Additionally, the Webex online platform is utilized by many educational institutions for online class recordings (Cisco Systems, 2025). Webex offers built-in transcription features that facilitate real-time speech-to-text conversion and the retrieval of accessible transcriptions. For this study, we obtained the class recording transcripts via the Webex transcription service (Fernandes and Wagh, 2022).

The length of the academic video transcripts is due to classroom interactions, including lecture explanations, student questions, debates, and discussions. The average length of online class video recordings utilized in this study varies between 45 and 90 minutes. This duration is comparatively longer than other research studies (Shen et al., 2022; Kawamura and Rekimoto, 2024). Our

dataset consists of two main domains, namely Law and Physics. The sub-details of the transcript dataset are provided in Table 1. The average number of tokens in Law transcripts is 10583, and in Physics it is 15126.

Integration of ChatGPT

The transcript data in this study are unlabeled and lack prior reference summaries. We integrate OpenAI's Chat Generative Pre-trained Transformer (ChatGPT) into our workflow to generate reference summaries of our transcripts. While several studies have assessed ChatGPT's performance on established Natural Language Processing (NLP) tasks (Grassini, 2023; Soni and Wade, 2023), the associated challenges in its usage cannot be disregarded. A key limitation in providing long video transcript data as input to the free version of ChatGPT is a 4,096-token limit (Sun et al., 2023). It limits the amount of labelled data usable for in-context learning, unlike supervised methods, which can leverage full datasets. To mitigate this limitation, we involved human experts to filter and limit the volume of transcript input. We provide uniform instructions to ChatGPT for generating a summary. The process for creating a reference summary is illustrated in Figure 1. These initial summaries generated via ChatGPT were adopted as reference summaries for preliminary experimentation with summarization models. For clarity in terminology, we refer to these reference summaries as "Human-assisted

reference summaries".

As an alternative approach, we explore automation of the manual process involved in transcript reduction. The TextRank algorithm was observed to have certain advantages, such as being unsupervised, graph-based centrality, low redundancy, and simplicity (Arora et al., 2023). We integrate the graph-based TextRank approach into our workflow. TextRank treats sentences as nodes in a graph, and edges represent similarity. The algorithm applies a variant of PageRank; hence, the more central sentences get higher scores (Zieve et al., 2023). We utilize the graph-based centrality of TextRank to obtain ranked sentences within ChatGPT's token limit, based on similarity. The TextRank process flow shown in Figure 2 is included. We refer to these reference summaries as "TextRank-integrated-reference summaries".

Preliminary Experimentation for Abstractive Transcript Summary Generation

Abstractive summarization helps in humanly describing the text, capturing the core meaning of the text in a compact, readable, and fluent approach. It is computationally more complex than the extractive technique. Abstractive summaries are hence generated by the model, where the tokenizer encodes and decodes the input and output sequences, and parameters control the generation process to balance fluency and relevance.

Table 1: Description of Transcript Dataset with subcategories in the law and physics domains

Domain	Sub-Category	Video Duration (Max)	Transcript length (Average length)
Law	Jurisprudence	01:41:32.279	19737
	Research Methodology	01:24:21.869	20918
	Renewable Energy	01:25:48.510	14255
	Statistical Physics	01:36:06.538	16734
Physics	Quantum Mechanics	01:02:02.652	11598
	Solid State Physics	01:31:02.699	17241
	Electromagnetic Studies	00:52:11.130	13910

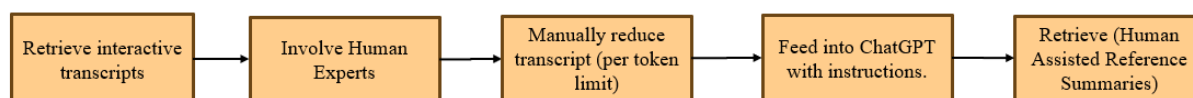


Fig. 1: Approach to retrieve reference summaries for the transcript data utilized in the research study with integration of ChatGPT

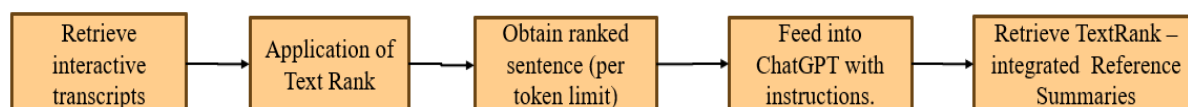


Fig. 2: Approach to retrieve reference summaries for transcripts using TextRank with integration of ChatGPT

We utilize state-of-the-art models like the T5 transformer and the BART model for abstractive summarization. We used Hugging Face Transformers (version 4.31.0) to access pre-trained models such as BART and T5. The methodology, as illustrated in Figure 3, outlines the step-by-step process of our approach.

BART (Bidirectional and Auto-Regressive Transformers) model algorithmically leverages a multi-layer bidirectional Transformer encoder-decoder architecture for sequence-to-sequence tasks (Lundberg et al., 2022; Saxena and El-Haj, 2023). During training, BART employs a denoising autoencoder objective, where corrupted input sequences are reconstructed. This training approach includes tasks such as masked language modelling and sequence shuffling. BART utilizes self-attention mechanisms to capture long-range dependencies and contextual information efficiently. Additionally, it incorporates positional encodings to maintain sequential information within the input embeddings.

T5 (Text-To-Text Transfer Transformer) model primarily employs a variant of the Transformer architecture for text-to-text tasks. It utilizes the standard Transformer algorithms, including self-attention mechanisms, multi-head attention, position-wise feedforward networks, and layer normalization (Kumar Bhukya and Bhukya, 2024; Saxena and El-Haj, 2023). T5 introduces a text-to-text framework in which all tasks, including text generation, classification, and translation, are cast into a unified format.

The above state-of-the-art models are applied to the interactive transcript dataset, and model-generated summaries are obtained. Given ChatGPT's growing popularity and ease of use, we integrate it into our methodology to generate reference summaries. ROUGE scores were applied to evaluate the model-generated summaries. The results are shown in Table 2 in Section 4. When applying ROUGE metrics, the BART model achieves higher scores. Following this preliminary

experimentation, the BART model was subsequently utilized to apply a transfer learning approach.

Transfer Learning

Transfer learning is a machine learning technique where a model trained on one task is repurposed or adapted for a related task. In transfer learning, the pre-trained model's knowledge, often gained from large datasets, is transferred to a new model to help it learn faster and with less data (Shen and Wang, 2024; Song et al., 2024). This approach is particularly beneficial when the new task involves a limited dataset or when computational resources are constrained in our research study.

Analogous to our interactive transcript dataset, we chose the SAMSum dataset, which consists of conversations and summaries created by linguists similar to those typically written daily. The SAMSum dataset consists of 16,369 conversations and was made publicly available by the Samsung R&D Institute, Poland (Gliwa et al., 2019). The training and validation datasets contain a dialogue, a gold summary, and an identifier. The training dataset included 14732 conversations, including both formal and semi-formal dialogues. The summaries are brief texts written in the third-person perspective. The key fields in the SAMSum dataset are: ID: A unique identifier for each conversation; dialogue: The actual multi-turn conversation between two or more participants, formatted like a chat (with speaker names); and summary: A human-written abstractive summary of the dialogue, capturing the main points.

The SAMSum dataset has been considered for this research study due to its features like speaker interactions, dialogue dynamics, and human-written summaries of realistic conversational dialogues. These features closely resemble lecture transcript data in real-world applications. Additionally, with its high-quality annotations and manageable size, it is effective for both benchmarking and fine-tuning models.

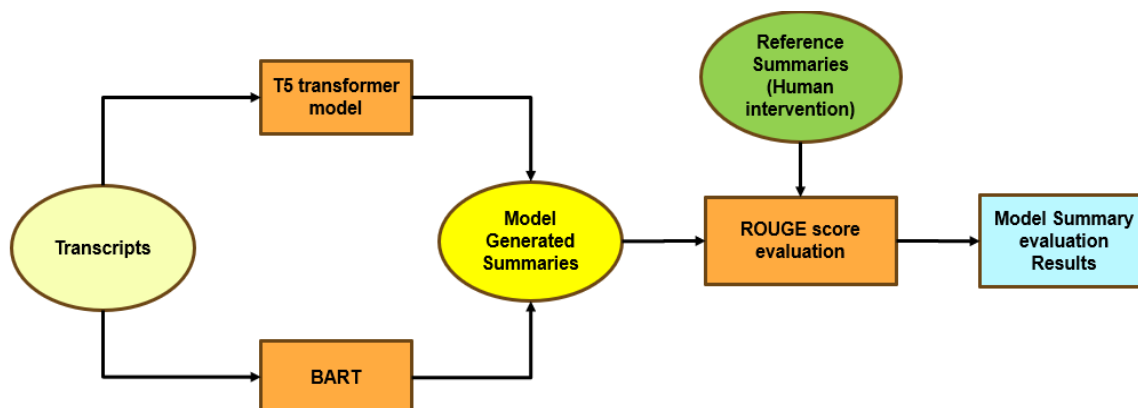


Fig. 3: Overview of process flow for generating and evaluating transcript summaries

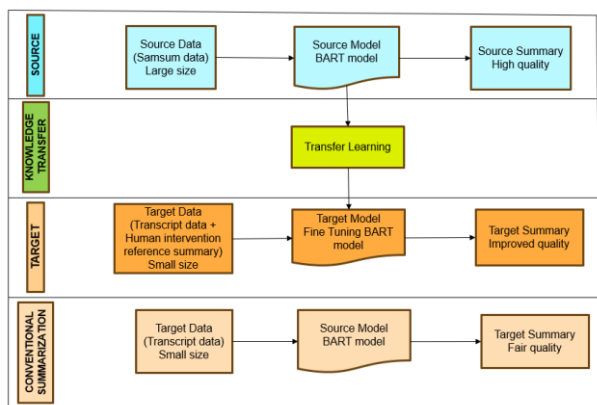


Fig. 4: The overview of the Transfer learning approach applying the BART model

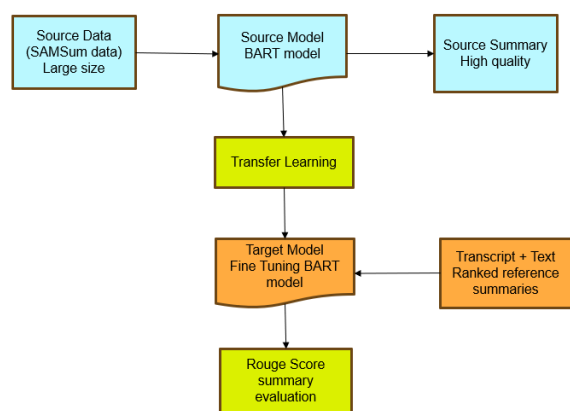


Fig. 5: Transfer Learning Approach utilizing the TextRank-integrated-reference summaries

We train the BART model on the SAMSum dataset mentioned. The pre-trained BART model's knowledge is applied to our interactive transcript dataset. The key fields in the interactive transcript dataset are: ID: A unique identifier for each transcript; text: the lecture transcript as a collection of sentences spoken during class; and summary: A human-written abstractive summary of the transcript, capturing the main points.

The process flow for our transfer learning approach with human-assisted ChatGPT summaries is shown in Figure 4. The sample-generated summaries are presented in Tables 5 and 6 in Section 4. The process flow for our transfer learning approach on TextRank-integrated-reference summaries is displayed in Figure 5.

Hyperparameter Tuning

Hyperparameter Tuning enhances the summarization process by optimizing the model for better summary coherence and content selection (Liu et al., 2024; Velayutham and Swarnakar, 2025). The initial phase involves analyzing the outcomes of hyperparameter fine-tuning of the pre-trained BART model on a new interactive transcript dataset, as described in Section 3.1. The hyperparameter tuning aims to observe and/or improve the effectiveness of BART on interactive transcript summarization. During hyperparameter tuning, we experimented with various learning rates. The most significant parameter values with the highest model performance across training iterations for the BART model trained on human-assisted summaries are a learning rate = $1e-5$ and a batch size = 2 for 3 epochs. Whereas for Text-Rank integrated summaries, the model was optimised at a learning rate = $5e-5$, a batch size = 1 for 3 epochs. Results, including ROUGE scores, are provided in Section 4.

Fine-Tuning BART Model

Fine-tuning involves selectively freezing and training specific layers of a pre-trained model on task-specific summarization data to retain general language understanding while improving summarization performance (Shen and Wang, 2024; Song et al., 2024). During the first stage of fine-tuning, the encoder parameters are frozen, restricting gradient updates to the decoder and the task-specific head. This allows the model to retain the rich pre-trained encoder representations while gradually adapting the decoder to the characteristics of the target transcript dataset. In the second stage, the encoder is unfrozen, enabling full model fine-tuning.

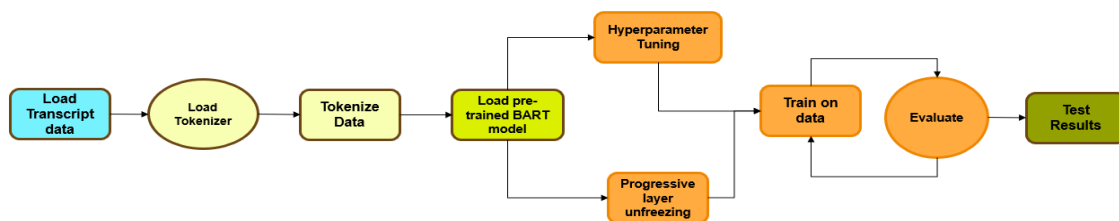


Fig. 6: Transfer Learning Approach with fine-tuning the BART model on the video transcript dataset

At this point, both encoder and decoder parameters are updated jointly, allowing the entire model to adapt more deeply to task-specific features and thereby improving overall performance. The overall fine-tuning process is illustrated in Figure 6, and its results with the updated parameters are given in Table 4 in the Results Section.

Evaluation Procedure

ROUGE (Recall-Oriented Understudy for Gisting Evaluation) is one of the widely used benchmarking metrics for evaluating text summarization tasks. ROUGE is a set of metrics used to evaluate automatic summarization by comparing system-generated summaries with reference summaries (Barbella and Tortora, 2022; Lin, 2004; Lundberg et al., 2022). It primarily measures n-gram overlap, with common variants including ROUGE-1, ROUGE-2, and ROUGE-L for unigram, bigram, and longest common subsequence matching. Its widespread adoption is due to its simplicity, interpretability, and effectiveness in capturing content overlap between generated and reference summaries. The interactive transcript summaries were evaluated on ROUGE-1, ROUGE-2, and ROUGE-L.

Results

Preliminary Experimentation

We evaluate the summaries generated by the state-of-the-art T5 transformer and BART models during preliminary experimentation. As shown in Table 2, the abstractive summaries generated by the BART model achieve higher ROUGE scores than those generated by the T5 transformer model.

Baseline Model

We chose the fine-tuned BART model on call transcripts from a recent research study as a comparative baseline (Ahmed et al., 2024). The study employs the pre-trained BART-large-xsum model, originally trained on the XSum dataset, as a baseline for abstractive summarization. The model is fine-tuned on domain-

specific call center transcripts, using GPT-3-generated summaries as reference targets. The fine-tuned BART model outperforms other variants in ROUGE metrics, demonstrating improved adaptation to conversational and ASR-transcribed data. For nomenclature, we refer to this model as the Baseline model. Table 2 contains the ROUGE scores of the baseline model. Call transcripts and lecture transcripts share similarities in scattered topics across utterances and noise; however, lecture transcripts prove to be heavy with technical terms, domain specifications, explanations, debates, discussions, and tutor instructions.

Transfer Learning

The BART model was also used for transfer learning. BART was initially trained on the SAMSum dataset. The Transfer Learning approach was applied by initializing our model with the pre-trained BART weights. The model was fine-tuned on the interactive transcript dataset for the abstractive summarization task. For nomenclature, we use model names such as "BART model trained on human-assisted summary" and "BART model trained on Text-Rank integrated summary". Table 3 shows the ROUGE scores obtained by our trained models.

Hyperparameter tuning is the process of optimising the settings that control how a machine learning model learns, such as the learning rate, batch size, number of layers, and dropout rate. The purpose is to identify a combination that maximizes model performance on validation data. The optimal configuration (learning rate = $1e-5$, batch size = 2, epochs = 3) yielded the best performance with human-assisted summaries, achieving improvements in ROUGE scores (ROUGE-1: 44.19, ROUGE-2: 15.09, ROUGE-L: 25.4) compared to the baseline. While using Text-Rank integrated summaries, the configuration (learning rate = $5e-5$, batch size = 1, epochs = 3) achieved ROUGE scores (ROUGE-1: 33.32, ROUGE-2: 7.94, ROUGE-L: 17.86). These results indicate that careful tuning significantly enhances model performance.

Table 2: Rouge scores of the preliminary experimentation and baseline BART model fine-tuned on call transcripts

Dataset	Models	Rouge1	Rouge2	RougeL
Proposed Human-assisted summaries	T5 model	18.12	4.23	15.21
	BART	22.31	11.41	15.3
Call Transcripts (Baseline)	BART pretrained (Baseline)	24.5	5.7	16.6
	BART fine-tuned (Baseline)	37.8	16.8	31

Table 3: ROUGE scores and performance of the BART model on human-assisted and Text-Rank integrated summaries

Parameter	Model	Rouge1	Rouge2	RougeL	RougeLSum
Number of beams = 4, Length Penalty = 2.0, Max Length = 1024, Min Length = 50, Early Stopping = True	BART trained on Human-assisted summary	20.56	4.53	13.52	13.54
	BART trained on Text-Rank integrated summary	19.5	3.35	12.16	12.23

Table 4: Results of ROUGE scores on fine-tuning the BART model on the two types of reference summaries

Reference summary	Process	Rouge1	Rouge2	RougeL	RougeLsum	Loss
Human-assisted summaries	Transfer Learning (Stage 1)	41.55	13.27	24.03	24.1	Validation loss = 2.86
	Transfer Learning (Stage 2)	43.35	13.68	24.71	24.72	Validation loss = 3.13
TextRank integrated summaries	Transfer Learning (Stage 1)	33.96	8.2	18.13	18.11	Validation loss = 2.75
	Transfer Learning (Stage 2)	34.58	6.91	17.74	17.66	Train loss = 1.225 Validation loss = 2.86

Table 5: Generated transcript summary outputs by the baseline BART model per training with a human-assisted reference summary

Summarization approaches	Generated transcript summary
Reference Summary (Human-assisted)	In the transcript, the professor elaborates on the concept of energy bands in the context of periodic potential. He explains how these bands arise from the electron wave functions and discusses the implications of standing and traveling waves in relation to electron behavior in solids. The professor emphasizes that at certain potentials, the wave functions may not exist, leading to regions where reflections occur. He introduces the idea of constructing standing waves from traveling waves, noting the relationship between potential energy changes and the electron density distribution. The discussion transitions to the mathematical representation of periodic potential, with the professor explaining how it can be expressed using Fourier series. He describes the significance of coefficients in this expansion and their role in modifying energy calculations. Overall, the transcript covers the foundational concepts necessary to understand the connection between periodic potentials, energy bands, and electron behavior in solid-state physics, while also indicating the upcoming mathematical approaches to be explored.
Baseline Model(human-assisted summary)	There is a band gap in the electrons. The reason could be due to the effect of the "patriotic potential". The band is influenced by the fact that electrons are in a weak periodic potential and can't be bound to the potential.
1st Hyperparameter Tuning	In the transcript, the professor elaborates on the concept of electron wave functions and their properties, focusing on the effects of the periodic potential on electron wave functions. The discussion begins with a review of nearly free electron models, which assume that electrons are not bound to potentials but are influenced by the effect of the wave function at certain points. The professor explains that when an electron is in a weak periodic potential, it cannot escape from the wave field, leading to the formation of a band gap. The importance of this band gap in determining electron beam propagation is emphasized, with the professor noting that it could indicate the presence of an electron wave field that propagates with respect to a single wave vector.
Transfer Learning (Stage 1)	In the transcript, the professor discusses various types of electron models, focusing on nearly free and nearly free electron models. The discussion begins with a review of the nearly free model, which assumes that electrons are completely free from the periodic potential and cannot be bound to potentials. The professor explains that the electron is influenced by the effect of periodic potentials, which restrict access to certain regions of the electron's wave vector, leading to the formation of band gaps. The conversation then shifts to the nearly electron-induced band gap, which arises when electrons are bound to a certain region of the wave vector and cannot freely move within it. The importance of this restriction is highlighted, as it limits the ability of electrons to propagate freely.
Transfer Learning (Stage2 - full)	In the transcript, the professor discusses the significance of periodic potentials in quantum mechanics, specifically focusing on electron wave functions and the relationships between wave functions. He introduces the nearly free electron model and the concept of a band, explaining that the presence of gaps in the wave function indicates that electrons cannot be bound to potentials and must be derived from partially filled bands. The professor emphasizes the importance of considering the behavior of electrons within these gaps, particularly when considering cases where the electron wave function cannot be expressed as a function of position, which could lead to false positives. The discussion then shifts to the effects of translational symmetry and the relationship between position and wave function probabilities, which are related to the number of electrons per unit cell.

In the next step, a two-stage transfer learning strategy is applied. Stage 1 of transfer learning involves freezing the majority of the pre-trained model's parameters and fine-tuning only the final layers on the transcript-specific dataset. This approach improves the linguistic knowledge

acquired during large-scale pretraining while mitigating the risk of overfitting on limited data. Subsequently, in Stage 2 of transfer learning, staged unfreezing is performed. In this phase, all layers are gradually trained, and the entire model is fine-tuned. This enabled the model

to adapt to domain-specific features, leading to improved performance and consistency in summary generation. The results of the fine-tuned model, Stage 1 and Stage 2, are displayed in Table 4. The fine-tuning parameters are a

learning rate of $5e-5$, a batch size of 2, and 3 epochs of training.

The generated summaries from different approaches are displayed in Tables 5 and 6.

Table 6: Generated transcript summary outputs by baseline BART model per training with TextRank-integrated reference summary.

Summarization approaches	Generated transcript summary
Reference Summary (TextRank integrated)	The lecture stresses that jurisprudence must be studied as a philosophical and conceptual discipline, not merely by memorizing legal rules. It contrasts natural law, which asks how law ought to be, with positive/analytical law, which studies law as it is. Thinkers like Bentham, Austin, Fuller, Plato, and Aristotle illustrate different views on law, morality, the state, and human rights. Central ideas of ontology (the nature of reality) and epistemology (how reality is known) are introduced, showing how perspectives shape legal understanding. Plato treating reality as beyond appearances versus Aristotle grounding it in sensory experience. Jurisprudence is best learned through active discussion and reasoning, since it connects law with philosophy, sociology, psychology, and politics. The lecture also explains how law regulates behavior through rights and duties, the limits of enforcement, and debates about the authority of the state. Overall, it highlights jurisprudence as essential for sharpening legal thinking, understanding reforms, and linking law to morality and society.
Baseline model (TextRank integrated) Hyperparameter Tuning 1	Good morning, everyone. Alina will tell you what happened in the last class and where they discussed it. The class started with a discussion about the social contract and ended with a small discussion about marriage and homosexuality. They also talked about the role of the state in modern society. The lecture provided detailed instructions for the upcoming legal research methodology exam and related assignments, highlighting the need for clarity, precision, and analytical depth within strict word limits. The exam will consist of three compulsory questions: The first will focus on analyzing a news item by identifying rights, relevant laws, loopholes, data sources, and counterarguments; the second will involve a literature review of a given article by summarizing the author's arguments; and the third will be on constitutional law, where students will examine the text of the law of unintended consequences, distinguishing between natural law, which focuses on how laws ought to be, and legal fiction, which examines how laws actually are. The instructor emphasized the importance of clarity and precision.
Transfer Learning (Stage1)	The discussion focuses on the concept of the social contract and the social contract in law, focusing on how it affects individuals, institutions, and society as a whole. The social contract is a contract between a person and the state, and it relates to rights and responsibilities in society. It includes the right to life in the body, the right not to be killed, and the duty not to take life, among other rights and duties. The discussion also addresses the role of the state in social contract enforcement, highlighting how it can be used to coerce people into compliance with laws and social norms. Students are encouraged to present their research papers, objectives, methods, and counterarguments.
Transfer Learning (Stage2)	The discussion explores the concept of legal rights, their enforcement, and their relationship with the state. It highlights how the right to life of a human being is limited by law and how this affects individuals' rights and liberties, including the ability to exercise them in private, as well as in public spheres like employment, housing, education, and health, where the state has the power to enforce rights and duties under the law but cannot legally control individuals or groups. It also highlights the importance of the social contract itself and how it affects individuals and groups, particularly in cases of discrimination against LGBT people. The discussion highlights how social contract enforcement is influenced by historical, cultural, and legal contexts.

Discussion

The ROUGE scores showed a reasonable improvement relative to baseline models, highlighting the effectiveness of hyperparameter tuning and the transfer learning approach. The baseline scores have been provided in Table 2. During hyperparameter tuning, the BART model trained on human-assisted summaries yields a 20% increase in ROUGE-1 and a 9% gain in ROUGE-2 and ROUGE-L, compared to the baseline pretrained BART. Compared to the baseline fine-tuned BART model, the hyperparameter-tuned model achieves a 6% increase in ROUGE-1. The proposed transfer learning approach improved model

performance by leveraging knowledge from a pretrained model. The proposed BART fine-tuned model on human-assisted summaries increased ROUGE-1, ROUGE-2, and ROUGE-L scores by 19%, 8%, and 8%, respectively, compared to the baseline pretrained model. ROUGE-1 shows a 6% increase over the baseline fine-tuned model.

The hyperparameter-tuned model on TextRank-integrated summaries achieves 9% gains in ROUGE-1 and 2% in ROUGE-2 over the baseline pretrained BART. The proposed fine-tuned model with TextRank-integrated summaries increased ROUGE-1, ROUGE-2, and ROUGE-L scores by 9, 2 and 2%, respectively, compared to the baseline pretrained BART.

Future work may explore additional fine-tuning strategies to further enhance the quality of the generated summaries. While this study's results demonstrate improvements in the quality of transcript summaries using the transfer learning approach, several avenues for future research remain. One promising direction is the exploration of alternative pre-trained models, such as PEGASUS or GPT-based architectures, which may offer improved performance in diverse summarization tasks. Additionally, to address the challenge of manual annotation, future work could investigate the development of semi-automated annotation tools, enabling faster and more efficient creation of high-quality datasets. Another potential area for improvement is detecting repetitive and non-beneficial utterances in transcripts to condense the length of domain-specific transcripts.

The transfer learning approach significantly reduces computational cost compared to training models from scratch. Extensive training on massive hardware setups is replaced by a pretrained BART model. The key requirement evolves into fine-tuning the model on task-specific datasets. This significantly lowers training time, memory requirements, and energy consumption. Though the inference cost may increase given the model's autoregressive nature of generation, overall, transfer learning offers a practical trade-off by minimizing development cost while maintaining strong summarization quality (Cleeman et al., 2023; Iman et al., 2023).

The lecture transcript data used in this study are complex due to their domain richness, length, and unstructuredness. Attaining gold summaries for such long transcripts is both time-consuming and challenging. The use of tools like ChatGPT provides the necessary momentum to the summarization pipeline. However, the ChatGPT-generated summaries may exhibit a certain degree of generalisation bias and implicit user bias in language generation. While such bias is often considered acceptable in assistive or exploratory settings, it is important to acknowledge and mitigate these biases through evaluation metrics, human oversight, and transparency in system design (Hake et al., 2024; Smith, 2024; Chen et al., 2025).

Conclusion

This study set out to summarize a long transcript dataset. In this study, we demonstrated an abstractive approach to summarize a lengthy transcript dataset. ChatGPT is integrated into our methodology for generating reference summaries. We automate the manual transcript reduction process using TextRank. Both reference summaries are utilized for preliminary experimentation and a transfer learning approach. Among state-of-the-art models, BART produces higher-quality video transcript summaries. To

improve summary results with reduced computational time, a transfer learning approach is applied to a small, lengthy transcript dataset. BART is initially pre-trained on the SAMSum dataset and subsequently fine-tuned on a long transcript dataset. The results show that staged unfreezing significantly improves summary quality, with decent improvement in ROUGE scores compared to baseline models. A modest improvement in ROUGE scores is observed compared to the baseline BART model; the results indicate potential for further enhancement through additional fine-tuning or model optimization. To further improve summary quality, human evaluations can be incorporated in future work. Although the proposed approach performed well, the experiments were restricted to a small dataset. Thus, further validation on larger and multilingual datasets is required. Future research could investigate real-world applications in low-resource settings. Finally, investigating the feasibility of real-time summarization for streaming data could further enhance the model's utility in dynamic environments.

Acknowledgment

The authors gratefully acknowledge Christ University, Bangalore, Karnataka, India, for providing the necessary resources, access to relevant materials and data sources, and the infrastructure that supported the successful completion of this work.

Funding Information

The authors have not received any financial support or funding.

Author's Contributions

The authors jointly conceptualized the research problem, designed the study, and analyzed the results. Deepa Fernandes implemented the methodology, conducted the experiments, and wrote and edited the original draft. Dr. Rupali Wagh supervised the research, reviewed, and edited the draft. Both authors read and approved the final manuscript.

Ethics

This article is original and contains unpublished material. Both authors have read and approved the manuscript, and there are no ethical issues involved.

References

- Abishek, R., Shivani, A., & Sanjay, S. (2025). Single-Document Abstractive Text Summarization: A Systematic Literature Review. *ACM Computing Surveys*, 57(3), 1–37.
<https://doi.org/10.1145/3700639>

- Achkar, P., Gollub, T., & Potthast, M. (2024). Ask, Retrieve, Summarize: A Modular Pipeline for Scientific Literature Summarization. *CEUR Workshop Proceedings*, 05.
- Ahmed, I., Zhou, Y., Sharma, N., & Hosier, J. (2024). Text Summarization for Call Center Transcripts. *Intelligent Systems and Applications*, 822, 542–551. https://doi.org/10.1007/978-3-031-47721-8_36
- Aiken, A. (2020). Zooming in on privacy concerns: Video app Zoom is surging in popularity. In our rush to stay connected, we need to make security checks and not reveal more than we think. *Index on Censorship*, 49(2), 24–27. <https://doi.org/10.1177/0306422020935792>
- Arora, M., Mudgil, P., Sharma, U., Chopra, C., & Singh, N. H. (2023). Evaluation of text summarization techniques in healthcare domain: Pharmaceutical drug feedback. *Intelligent Decision Technologies*, 17(4), 1309–1322. <https://doi.org/10.3233/idt-230129>
- Barbella, M., & Tortora, G. (2022). Rouge Metric Evaluation for Text Summarization Techniques. *SSRN Electronic Journal*, 31. <https://doi.org/10.2139/ssrn.4120317>
- Bhukya, V. K., & Bhukya, U. (2024). Abstractive Text Summarisation using T5 Transformer Architecture with analysis. *Research Square*. <https://doi.org/10.21203/rs.3.rs-4986903/v1>
- Chen, X., Wang, T., Zhou, J., Song, Z., Gao, X., & Zhang, X. (2025). Evaluating and mitigating bias in AI-based medical text generation. *Nature Computational Science*, 5(5), 388–396. <https://doi.org/10.1038/s43588-025-00789-7>
- Chen, Y., Liu, Y., Chen, L., & Zhang, Y. (2021). DialogSum: A Real-Life Scenario Dialogue Summarization Dataset. *Findings of the Association for Computational Linguistics: ACL-IJCNLP 2021*, 5062–5074. <https://doi.org/10.18653/v1/2021.findings-acl.449>
- Chick, R. C., Clifton, G. T., Peace, K. M., Propper, B. W., Hale, D. F., Alseidi, A. A., & Vreeland, T. J. (2020). Using Technology to Maintain the Education of Residents During the COVID-19 Pandemic. *Journal of Surgical Education*, 77(4), 729–732. <https://doi.org/10.1016/j.jsurg.2020.03.018>
- Cisco Systems, Inc. (2025). Cisco Webex. *Cisco Webex Website*. <https://www.webex.com>
- Cleeman, J., Agrawala, K., & Malhotra, R. (2023). Accelerated and inexpensive Machine Learning for manufacturing processes with incomplete mechanistic knowledge. *Manufacturing Letters*, 37, 53–56. <https://doi.org/10.1016/j.mfglet.2023.07.017>
- Elahi, E., Morato, J., & Iglesias, A. (2024). Improving Web Readability Using Video Content: A Relevance-Based Approach. *Applied Sciences*, 14(23), 11055. <https://doi.org/10.3390/app142311055>
- El-Kassas, W. S., Salama, C. R., Rafea, A. A., & Mohamed, H. K. (2021). Automatic text summarization: A comprehensive survey. *Expert Systems with Applications*, 165, 113679. <https://doi.org/10.1016/j.eswa.2020.113679>
- Fernandes, D., & Wagh, R. S. (2022). Transforming online class recording into useful information repositories using NLP methods: An Empirical Study. *2022 4th International Conference on Advances in Computing, Communication Control and Networking (ICAC3N)*, 589–594. <https://doi.org/10.1109/icac3n56670.2022.10074025>
- Gliwa, B., Mochol, I., Biesek, M., & Wawer, A. (2019). SAMSum Corpus: A Human-annotated Dialogue Dataset for Abstractive Summarization. *Proceedings of the 2nd Workshop on New Frontiers in Summarization*, 70–79. <https://doi.org/10.18653/v1/d19-5409>
- Golia, L., & Kalita, J. (2023). Action-Item-Driven Summarization of Long Meeting Transcripts. *Proceedings of the 2023 7th International Conference on Natural Language Processing and Information Retrieval*, 91–98. <https://doi.org/10.1145/3639233.3639253>
- Google. (2025a). Android Automatic Speech Recognition (ASR). *Android Developers Documentation*. <https://developer.android.com/reference/android/speech/SpeechRecognizer>
- Google. (2025b). Google Docs – Online Document & PDF Editor. *Google Workspace Website*. https://workspace.google.com/intl/en_in/products/docs/
- Goyal, T., Xu, J., Li, J. J., & Durrett, G. (2022). Training Dynamics for Text Summarization Models. *Findings of the Association for Computational Linguistics: ACL 2022*, 2061–2073. <https://doi.org/10.18653/v1/2022.findings-acl.163>
- Grassini, S. (2023). Shaping the Future of Education: Exploring the Potential and Consequences of AI and ChatGPT in Educational Settings. *Education Sciences*, 13(7), 692. <https://doi.org/10.3390/educsci13070692>
- Hake, J., Crowley, M., Coy, A., Shanks, D., Eoff, A., Kirmer-Voss, K., Dhanda, G., & Parente, D. J. (2024). Quality, Accuracy, and Bias in ChatGPT-Based Summarization of Medical Abstracts. *The Annals of Family Medicine*, 22(2), 113–120. <https://doi.org/10.1370/afm.3075>
- Hamid, O. H., & Samad, A. E. (2015). The Blurred Line between “Long” and “Short”: How the Length of Video Lectures Affects the Viewing Behavior of E-Learners. *Online*, 6(3), 32–37.
- HappyScribe. (2025). HappyScribe – AI Transcription, Subtitles & Translation Platform. *HappyScribe Website*.

- He, J., Su, H., Cai, J., Shalymov, I., Song, H., & Mansour, S. (2024). Semi-Supervised Dialogue Abstractive Summarization via High-Quality Pseudolabel Selection. *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, 5976–5996.
<https://doi.org/10.18653/v1/2024.naacl-long.333>
- Iman, M., Arabnia, H. R., & Rasheed, K. (2023). A Review of Deep Transfer Learning and Recent Advancements. *Technologies*, 11(2), 40.
<https://doi.org/10.3390/technologies11020040>
- Jangra, A., & Hasanuzzaman, M. (2023). 296 A Survey on Multi-modal Summarization. *Journal of the ACM*, 37, 296.
- Jiang, X., & Dreyer, M. (2024). CCSum: A Large-Scale and High-Quality Dataset for Abstractive News Summarization. *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, 7306–7336.
<https://doi.org/10.18653/v1/2024.naacl-long.406>
- Jin, Y., Shi, Q., & Liu, Q. (2025). CFAS: consensus-focused abstractive meeting summarization through multi-party discourse modeling. *Journal of King Saud University Computer and Information Sciences*, 37(7), 187.
<https://doi.org/10.1007/s44443-025-00210-3>
- Kawamura, K., & Rekimoto, J. (2024). FastPerson: Enhancing Video-Based Learning through Video Summarization that Preserves Linguistic and Visual Contexts. *Proceedings of the Augmented Humans International Conference 2024*, 205–216.
<https://doi.org/10.1145/3652920.3652922>
- Lin, C.-Y. (2004). ROUGE: A Package for Automatic Evaluation of Summaries. *Text Summarization Branches Out*, 74–81.
- Lin, Z., Chen, S., Wang, N., & Li, H. (2024). Evaluating Text Summarization Generated by Popular AI Tools. *Proceedings of the 2nd International Seminar on Artificial Intelligence, Networking and Information Technology*, 350–354.
<https://doi.org/10.5220/0012283800003807>
- Liu, M., Ma, Z., Li, J., Cheng Wu, Y., & Wang, X. (2024). Deep-Learning-Based Pre-Training and Refined Tuning for Web Summarization Software. *IEEE Access*, 12, 92120–92129.
<https://doi.org/10.1109/access.2024.3423662>
- Lundberg, C., Viñuela, L. S., & Biales, S. (2022). Dialogue Summarization using BART. *Proceedings of the 15th International Conference on Natural Language Generation: Generation Challenges*, 121–125.
- Merkt, M., Hoppe, A., Bruns, G., Ewerth, R., & Huff, M. (2022). Pushing the button: Why do learners pause online videos? *Computers & Education*, 176, 104355.
<https://doi.org/10.1016/j.compedu.2021.104355>
- Mishra, N., Sahu, G., Calixto, I., Abu-Hanna, A., & Laradji, I. (2023). LLM aided semi-supervision for efficient Extractive Dialog Summarization. *Findings of the Association for Computational Linguistics: EMNLP 2023*, 10002–10009.
<https://doi.org/10.18653/v1/2023.findings-emnlp.670>
- Navarrete, E., Nehring, A., Schanze, S., Ewerth, R., & Hoppe, A. (2025). A Closer Look into Recent Video-based Learning Research: A Comprehensive Review of Video Characteristics, Tools, Technologies, and Learning Effectiveness. *International Journal of Artificial Intelligence in Education*, 35(4), 1631–1694.
<https://doi.org/10.1007/s40593-025-00481-x>
- Pilault, J., Li, R., Subramanian, S., & Pal, C. (2020). On Extractive and Abstractive Neural Document Summarization with Transformer Language Models. *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 9308–9319.
<https://doi.org/10.18653/v1/2020.emnlp-main.748>
- Saxena, P., & El-Haj, M. (2023). Exploring Abstractive Text Summarisation for Podcasts: A Comparative Study of BART and T5 Models. *Proceedings of the Conference Recent Advances in Natural Language Processing - Large Language Models for Natural Language Processings*, 1023–1033.
https://doi.org/10.26615/978-954-452-092-2_110
- Seidel, N. (2024). Short, Long, and Segmented Learning Videos: From YouTube Practice to Enhanced Video Players. *Technology, Knowledge and Learning*, 29(4), 1965–1991.
<https://doi.org/10.1007/s10758-024-09745-2>
- Shen, H., Zhang, H., Yuan, Y., & Gao, X. (2022). Lecture Video Automatic Summarization System Based on DBNet and Kalman Filtering. *Computational Intelligence and Neuroscience*, 1–10.
- Shen, J., & Wang, Y. (2024). Leveraging Parameter-Efficient Fine-Tuning for Multilingual Abstractive Summarization. *Natural Language Processing and Chinese Computing*, 15361, 293–303.
https://doi.org/10.1007/978-981-97-9437-9_23
- Song, J., Park, N., Hwang, B., Yun, J., Joe, S., Gwon, Y., & Yoon, S. (2024). Entity-level Factual Adaptiveness of Fine-tuning based Abstractive Summarization Models. *Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics (Volume 1: Long Papers)*, 915–929.
<https://doi.org/10.18653/v1/2024.eacl-long.55>

- Smith, J. M. (2024). "I'm Sorry, but I Can't Assist": Bias in Generative AI. *Proceedings of the 2024 on RESPECT Annual Conference*, 75–80.
<https://doi.org/10.1145/3653666.3656065>
- Soni, M., & Wade, V. (2023). Comparing Abstractive Summaries Generated by ChatGPT to Real Summaries Through Blinded Reviewers and Text Classification Algorithms. *ArXiv*.
<https://doi.org/10.48550/arXiv.2303.17650>
- Suarez, J. B., Martínez, P., & Martínez, J. L. (2020). Combining Financial Word Embeddings and Knowledge-Based Features for Financial Text Summarization: UC3M-MC System at FNS-2020. *Proceedings of the 1st Joint Workshop on Financial Narrative Processing and MultiLing Financial Summarisation*, 112–117.
- Sun, X., Dong, L., Li, X., Wan, Z., Wang, S., Zhang, T., Li, J., Cheng, F., Lyu, L., Wu, F., & Wang, G. (2023). Pushing the Limits of ChatGPT on NLP Tasks. *ArXiv*.
<https://doi.org/10.48550/arXiv.2306.09719>
- Takeshita, S., Green, T., Reinig, I., Eckert, K., & Ponzetto, S. (2024). ACLSum: A New Dataset for Aspect-based Summarization of Scientific Publications. *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, 6660–6675.
<https://doi.org/10.18653/v1/2024.naacl-long.371>
- Tyss, S., Jia, C., Goroncy, P., & Grabmair, M. (2025). RELexED: Retrieval-Enhanced Legal Summarization with Exemplar Diversity. *Findings of the Association for Computational Linguistics: NAACL 2025*, 427–434.
<https://doi.org/10.18653/v1/2025.findings-naacl.26>
- Velayutham, V. M., & Swarnakar, S. (2025). Pegasus-CNN/dailymail fine-tuning for domain-specific summarization: performance evaluation and hyperparameter optimization. *Iran Journal of Computer Science*, 8(4), 2007–2022.
<https://doi.org/10.1007/s42044-025-00299-9>
- Zhang, Y., Ni, A., Mao, Z., Wu, C. H., Zhu, C., Deb, B., Awadallah, A., Radev, D., & Zhang, R. (2022). SummN: A Multi-Stage Summarization Framework for Long Input Dialogues and Documents: A Multi-Stage Summarization Framework for Long Input Dialogues and Documents. *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 1592–1604. <https://doi.org/10.18653/v1/2022.acl-long.112>
- Zhong, M., Yin, D., Yu, T., Zaidi, A., Mutuma, M., Jha, R., Awadallah, A. H., Celikyilmaz, A., Liu, Y., Qiu, X., & Radev, D. (2021). QMSum: A New Benchmark for Query-based Multi-domain Meeting Summarization. *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, 5905–5921.
<https://doi.org/10.18653/v1/2021.naacl-main.472>
- Zieve, M., Gregor, A., Stokbaek, F. J., Lewis, H., Mendoza, E. M., & Ahmadnia, B. (2023). Systematic TextRank Optimization in Extractive Summarization. *Proceedings of the 14th International Conference on Recent Advances in Natural Language Processing*, 1274–1281.
https://doi.org/10.26615/978-954-452-092-2_135