

An NLP Powered Pipeline for Automated Construction and Sentiment Labelling of a Multi-Source Irish News

MSc Research Project
AI for Business

Rao Waqar Ali
Student ID: x24147346

School of Computing
National College of Ireland

Supervisor: Anderson Augusto Simiscuka

Student Name:Rao Waqar Ali..... **Student ID:**24147346.....

Programme:MSc AI for Business..... **Year:** ...2025... **Module:** ...Practicum II...

Supervisor:Anderson Simiscuka..... **Submission Due Date:**11 Dec 2025...

Project Title: ...An NLP Powered Pipeline for Automated Construction and Sentiment

Labelling of a Multi-Source Irish News... **Word Count:**9062...**Page Count** ...26...

I hereby certify that the information contained in this (my submission) is information pertaining to research I conducted for this project. All information other than my own contribution will be fully referenced and listed in the relevant bibliography section at the rear of the project. ALL internet material must be referenced in the bibliography section. Students are required to use the Referencing Standard specified in the report template. To use other author's written or electronic work is illegal (plagiarism) and may result in disciplinary action.

Signature:.....Rao Waqar Ali..... **Date:**11 Dec 2025.....

PLEASE READ THE FOLLOWING INSTRUCTIONS AND CHECKLIST

Attach a completed copy of this sheet to each project (including multiple copies)	☐
Attach a Moodle submission receipt of the online project submission , to each project (including multiple copies).	☐
You must ensure that you retain a HARD COPY of the project , both for your own reference and in case a project is lost or mislaid. It is not sufficient to keep a copy on computer.	☐

Assignments that are submitted to the Programme Coordinator Office must be placed into the assignment box located outside the office.

Office Use Only	
Signature:	
Date:	
Penalty Applied (if applicable):	

An NLP-Powered Pipeline for News Summarization, Sentiment Analysis, and Topic Classification: A Case Study on Irish Media

Abstract:

The sheer increase in the digital news content poses tremendous challenges to the automated processing and analysis. The paper constructs and evaluates an end-to-end Natural Language Processing pipeline that targets the analysis of the sentiment of Irish news media. The most impactful element that occurred to me when I first ventured into the exploration of this research was the overwhelming amount of digital news that is being generated every single day by the media houses in Ireland. The project is a cumulative result of my attempt to create an advanced Natural Language Processing framework with an exceptionally narrow scope to send sentiments analysis to an Irish journalism environment. My constructed system incorporates five somewhat different methodological approaches, namely VADER to analyze the lexicon of a text, BERT to serve as a benchmark in text-to-text transformations, BER-Tweet to optimize on social media, and an ensemble model to use majority vote strategies [31, 30, 29, 28]. By automating the process of collecting data using the major Irish sources such as RTE, The Journal and The Irish Times using RSS feeds and highly sophisticated preprocessing methods and running the former model parallel I was able to form a very strong evaluation framework. What came out in my analysis was very eye opening. BERT was shown to be extremely consistent, which is used as an excellent standard of reference. Ensemble method attained an accuracy of 87.69 percent and Vader scored averagely at 63.08 percent. The same was demonstrated to be limited by both BER- Tweet and T5 with 60.00 percent and 49.23 percent accuracy, respectively. These results would offer objective, evidence based recommendations on the choice of the model as well as define precious performance standards in the case of the Irish news sentiment analysis.

1. Introduction

The information consumption in the world has been altered with the digital change of the news media [12]. Transformer architecture has enabled NLP to revolutionize since it was invented by Vaswani et al. [1]. BERT [2] and T5 [3] are some of the most promising developments, which facilitate complex language comprehension operations. The models based on lexicons (such as VADER [4]) are still used in lightweight sentiment analysis, whereas domain-specific ones (such as BER-Tweet [5]) are focused on the social media text. Nevertheless, these methodologies must be used in the context of the Irish news validated empirically [7], [8].

Natural Language Processing is a branch of artificial intelligence that offers computational techniques in terms of the recognition, interpretation, and generation of human language. More recent breakthroughs in transformer structure have transformed the NLP abilities, empowering intricate language interpretation jobs, which were not available previously. The developments present possibilities of automated news processing, sentiment analysis and content approval in scales that are not possible with manual analysis. Nonetheless, the implementation of these methods to the domain-specific issues like Irish news media needs to be empirically verified because the results of the model on general corpora might not be delivered to the region-specific issues with peculiar linguistic features, cultural sources, and topical issues.

The study fulfills the high demand of automated sentiment analysis tools as applied to Irish news. The project creates a complete pipeline of NLP, which involves automated data grabbing, advanced preprocessing, and parallelization of various sentiment modelling. In this study, evidence-based information on the criteria of selecting models, performance trade-offs, and optimal configuration of victorious applications of the Irish news sentiment analysis application can be obtained through careful comparative analysis of five different approach. The overall analysis structure considers

various aspects of performance such as accuracy, precision, recall, F1-scores, and cardinality, and it allows one to make sensible decisions regarding putting them into practice.

2. Research Questions

The research questions to be used in this research deal with theoretical insights and practical issues in implementation. These questions were generated based on the close examinations of the unmet gaps in the literature, application needs, and the peculiarities of the Irish news media. All the questions bring forth a variety of information and accumulate the overall knowledge of the possible performance of the sentiment analysis models in this sector. **RQ1: How can automated NLP pipeline be an effective way of collecting, pre-processing, analyzing, and summarizing news articles in Ireland?** The initial research question holds the underlying problem of creating a complex pipeline that will be able to handle the entire workflow in terms of data collection, analysis, and summarization. This query acknowledges the fact that successful news processing involves more than sheer component adoption it needs careful synthesis of several phases such as the acquisition of RSS feeds, recovering of articles, pre-processing the text, running the model, and production of results. There is successful implementation of every stage of the pipeline to answer this question. The method to be used in the data collection is the RSS feed parsing of three well known Irish based news websites RTE, The Journal and The Irish Times which allow automated extraction of article metadata and articles URLs. After filtering articles according to the quantity and quality standards of the minimal content, the collection process produced high-quality articles. Such a variety of data serves as the examples of various areas such as political, business, technology, and social problems, which is what makes the evaluation of diverse types of content complete. Preprocessing phases make use of several methods that guarantee the quality of data such as HTML content removal with Beautiful Soup, text normalization that addresses encoding problems and whitespace standardization, sentence filtering to eliminate fragments. **RQ2: To what extent can various sentiment analysis models be used to identify the sentiment of Irish news stories?** The second research question is narrowed down to model performance properties applied to Irish news material. The reason behind this question is that pretrained models imply a major benefit such as the lower training needs and ability to transfer learning, but their performance cannot be guaranteed on domain-specific data only by relying on general benchmarks. On a largescale assessment, there is a wide difference in performance of models. When used as the reference standard, BERT has perfect accuracy (100%), which indicates the usefulness of transformer-based architectures to sentiment analysis of news. VADER is also performing moderately at 63.08 percent accurate, which is commendable since it is simple yet able to achieve its results in a computationally efficient way. Nevertheless, it is less accurate than transformer models, indicating that the language of the news articles can be more complex and sophisticated representations acquired by the neural models are beneficial. Accuracy of t5 is only 49.23 percent, which is highly out of the question. Such poor performance shows relevance to matching between training and application data in domain. BER-Tweet is effective with Twitter data, and it performs poorly when transferred to formal news articles. These findings show that accuracy highly relies on the model structure and domain correspondence with transformer-based models having the most lucrative application. **RQ3: Does ensemble voting demonstrate higher effectiveness due to the combination of several sentiment analysis models in terms of the overall sentiment prediction robustness?** The third research question explores the issue of enhancement in performance when ensemble techniques are applied to predictive multiple individual models. The ensemble methods have proven themselves to be effective in other machine learning settings where other models have orthogonal strengths and their weaknesses are prevented. The accuracy of ensemble models is 87.69% and it shows considerable enhancement in comparison with other individual models other than BERT. This difference in performance can be attributed to the use of complementary capabilities: VADER has a rule-based consistency that ensures performance, BERT has a wide view of context regarding understanding, BER-Tweet has social media training, and T5 is a

generative model. **RQ4: Which is the effectivity of T5 as a model of summarizing news about Ireland in concise and coherent summaries?** The fourth research question investigates the effectiveness of T5 when applied in summarizing assignments and is not limited to its sentiment analysis ability. The text-to-text structure of T5 is well adapted to such generational tasks as summarization when it is desired that the output be natural language text and is not discrete classification labels. T5 is very effective with summarization tasks of the pipeline. The model results in brief, concise and contextually related summaries of long articles, despite truncation to 512 tokens. The flexibility of T5 can be used to produce summaries of different sizes (depending on the configuration settings) so that it can be used to fit the needs of different applications. In general, T5 promotes the readability rate and facilitates downstream sentiment and topic tasks greatly. **RQ5: What can be learned about sentiment trend in Irish media based on transformer-based and lexicon-based models?** The fifth research question will look at insights which can be brought up with the help of the comparative analysis of different model approaches to apply them in the Irish news content. This question acknowledges that different models may demonstrate different features of patterns of sentiments and the knowledge about the differences can offer useful knowledge about not only model behavior, but also about media as well. Comparative analysis demonstrates evident distinctions between the models that understand the emotional tone. VADER is prone to optimistic classification being overly optimistic on the positive sentiments than the other models. This bias is caused by the fact that the approach based on lexicon is sensitive to positive words and phrases. BERT provides a model distribution among sentiment classes with a balanced and realistic distribution of the sentiment category, and it is based on its advanced contextual and nuanced understanding. BAD-Tweet indicates a neutral label of most of the articles which is related to its training made on Twitter data as the neutral sentiment is less prevalent. T5 is highly negative and is negative bias in that it is rather generative when used with formal news material. The ensemble generates smoothed out distribution thereby avoiding skewness by any of the models. Those experiences provide easy insight into the differences in the interpretation of emotive color of the Irish media by models and why the multi-model pipelines are the key to proper analysis of the media.

3. Literature Review

To produce this study, I submerged myself in the body of literature available in the papers that pertain to automated news processing, sentiment analysis procedures as well as the NLP applications to the analysis of the media. The theoretical underpinnings, methodological procedures, results of the studies, and research gaps that are used in contemporary research are identified by the review I conducted. My company adheres to thematic tracks, which are transformer architecture, sentiment analysis preapplication assemble algorithms, and domain specific applications that I felt most interested in working on.

3.1 Transformer Architectures in Natural Language Processing

The natural language processing has been totally reinvented due to transformer architecture with their revolutionary introduction by Vaswani and others in 2017 [32]. These architectures use selfattention models that allow models to learn long range dependencies and location specific relationships in a more effective way than their predecessor recurrent or convolutional architectures- a breakthrough that I believe to be truly groundbreaking. Transformer attention also is bidirectional, which permits models to look at someone in both directions of the sentence when processing an individual token, which offers more semantic insight in contrast to unidirectional models which dominated the previous models. This architecture has allowed the creation of pretrained language models, which can be fine trained to accomplish tasks and significantly reduce the amount of training data, but can still achieve state-of-the-art performance in nearly any NLP task and this feature I heavily utilized in my pipeline design. Another breakthrough technology in transformer-based language understanding which I heavily expanded in my work is BERT (Bidirectional Encoder Representations from Transformers) by Devlin and his colleagues in 2019

[30]. BERT uses masked language modeling and next sentence prediction in the pretraining phase to allow the model to extract text representations that are rich in directionality and capture features that I have found of value especially with news analysis. The architecture of the model can be optimized with a wide range of downstream applications, such as sentiment analysis, question answering, and text classification with only few tasks specific adjustments, which is essential in my comparative approach. BERT has shown an outstanding performance in many benchmarks making transformer-based methods the new paradigm in the modern research of NLP. On specific domain content like news articles in regions, however, the effectiveness of BERT must be empirically confirmed, because training data properties can not necessarily be wholly matched to the domains of use, which is one of the reasons why I chose to investigate the Irish media. The T5 (Text-To-Text Transfer Transformer), introduced by Raffel and colleagues in 2020 [28], is a single framework that is bringing all NLP problems to the description of a text generation problem, a reasoning I found elegant and useful. The methodology will allow a single model structure to cope with various tasks such as translation, summarization, classification, and question answering via a proper prompt structure. The text-to-text paradigm of T5 is flexible and can be transferred to do any task but with some attention to prompt engineering, which would ensure the best results on a particular task, which I faced during the implementation of the new paradigm. The generative nature of the model comes with benefits of the tasks that need an output of text, though it can create issues in classification tasks where discrete label prediction is a better compromise that was evident in my sentiment analysis results.

3.2 Sentiment Analysis Methodologies

Sentiment analysis has developed through primitive lexicon techniques to reluctant neural network designs a transition that I was able to observe in my literature review. Initial approaches were based on sentiment lexicons which have words that have relative polarities, and the idea is that they could provide rule-based classification by summing up individual word sentiments- a method that appears almost naive by modern standards, but which does still work in some cases [24]. The developed VADER (Valence Aware Dictionary and sentiment Reasoner) has been created by Hutto and Gilbert in 2014 as an advanced lexicon-based system that is explicitly trained on text in social media, which I added to the analysis as a comparison [31]. VADER uses grammatical rules, effects of capitalization, emphasis of punctuation, and negation processing to ensure accuracy at higher sophistication levels than counting simple words, which allows it to make it more appropriate to real-life text than the previous lexicon techniques. This approach is computationally efficient and easy to interpret, so it can be used when processing data or making decisions that need to be fast, but its rule-based character can be harnessed when dealing with complicated linguistic constructs, which I found to be a weakness in my data. Sentiment analysis via neural network has been shown to do much better on complex text at higher computation cost, which needs to be considered when implementing, however. In early years convolutional neural networks and recurrent neural networks were promising but more recently transformer-based models have replaced these methods, and this trend in the model choice is also reflected in my choice of models [22]. Sentiment analysis models based on the BERT resource that have been fine-tuned on sentimentlabeled datasets obtain high state-of-the art performance due to the strong contextual representations acquired during pretraining and which became critical in my analysis [30]. Nevertheless, the performance changes greatly depending on the variables of specific pretraining data, fine-tuning procedure, and target domain characteristics variables that I adjusted in detail during the development of my experiment. The model used in this project is also referred to as BERT-Tweet, and it was created by Nguyen and others in 2020 as a specific variant of BERT that was trained on Twitter, a specialization that interested me in this project [29]. Domain-specific pretraining can also allow the model to comprehend more in formal language and social media conventions and colloquial expressions prevalent in Twitter text. Nevertheless, the success of the model with formal news articles, which have varying linguistic properties to those of social media posts, needs an empirical study, a research gap filled within the context of my study directly. Ability to transfer

learning of transformer models leads to assumptions that it might be effective, yet mismatch between domain training and application data might hamper its performance, which was found to be true according to my results.

3.3 News Processing Applications

Automated processing of news has been the subject of research related to multiple applications such as summarization, topic classification, as well as sentiment analysis. Garg proved that pretrained transformer models greatly outperform the traditional summarization models with an improvement of 16 and 38 percent in T5 as well as BERT respectively in creating coherent and relevant summaries [35]. Their study concluded that approaches based on transformers proved to be good in summarizing news but were evaluated using general news corpora and not region-specific news.

Kazemi created an interpretable NLP pipeline to analyze inclusive policymaking with respect to migrants in Ireland, which indicates the usage of NLP in analyzing Ireland-specific data[34]. Their article emphasized how difficult it is to process mixed (heterogeneous) datasets and derive useful information out of social media data and news information. But their attention was mainly on the marginal voices, as opposed to the overall news processing, and the lacuna was left in their wide applications beyond the various topics or areas.

Begum designed an automated system of providing news in the form of daily reports by incorporating news scraping, summarization, and sentiment analysis with GPT-3.5 and BERTweet [33]. This was shown to be practically viable in end-to-end automation but there have been scalability issues in real-time processing of news volumes to handle. The study revealed that each of its individual components works well individually but their combination into saleable pipelines involves a lot of optimization and architecture design.

3.4 Ensemble Methods in Sentiment Analysis

The ensemble methods consist of numerous methods and make use of these methods in the combination to enhance the overall performance, taking advantage of their complementary strengths and counteracting weaknesses of each one of them. Majority voting is the simplest combination method which selects the most predicted model of individual models. More advanced ensemble techniques use weighted performance on a model-based voting, stacking, where a meta learner uses a combination of models, and boosting where models are trained one after another to improve on any past mistakes. A study [36] showed that BERT-based multitask learning can be successfully used to deal with sentiment analysis and topic classification at the same time, although they performed the analysis on generalized datasets, and data with a specific field were evaluated.

In the context of ensemble methods, the diversity of individual models is crucial; when all models progress toward the same errors then little benefit is gained through the combination of such models. It has been proposed based on research that using models that are combined with another architecture, training process, or even inductive biases would produce better results in an ensemble. Nevertheless, the ensemble methods raise the requirements in computations and can create complexity in deployment and maintenance. Mediation between the gains in performance and expenditure is required to be rated on application specific needs.

3.5 Research Gaps and Contributions

Although the development of sentiment analysis and news processing has been extensively researched, there is still a great void in using the techniques to the Irish news media. The greater part of available studies compares models with general English language corpus or news content in the US which makes it unclear when in reference to Irish news that might have different language features, cultural backgrounds and the current topics. Also, multi-model comparative

assessments of various sentiment analysis models on the same data set are not well done hence it is hard to ascertain the most effective models to select in applications.

This study fills these gaps by presenting the first extended comparative analysis of several sentiment analysis models on Irish news material, in particular. The study creates a unified pipeline that exhibits high realistic implementation, formally assesses the application on several performance indicators, and offers evidence-based recommendations on the choice of models. The ensemble methodology evaluation provides information on the performance benefits of the application of multiple models in performing an Irish news sentiment analysis, which answers a question addressed in the currently prioritized literature on limited scale.

4. Gap Analysis

Although summarization and sentiment analysis were advanced, there are still several research gaps. Absence of domain specific pipelines: the current efforts tend to use general models on the news data without specifications to the lingual characteristics of the Irish reporting. This leaves out domain adaptation so that crucial cultural references or idioms could be misconstrued.

Poor integration: there are numerous studies that concentrate either on summarization or sentiment analysis separately. These activities can be run in a single pipeline with additional efficiency and reliability, but few systems do so in the case of national news. Pre-sentiment analysis (performing sentiment analysis on summarized articles) can decrease the length of inputs and affect the performance of a classification, yet this most often is not studied.

Unresearched ensemble techniques: most sentiment analysis studies consider individual models. Although ensembles have been used successfully in machine learning, they have not been used in sentiment classification of news. The Ensemble voting might provide more balanced results with the graph of aggregating the strengths of the lexicon and transformer models.

Inadequate local media assessment: previous studies generally apply internationally recognized data like ones in IMDb, Yelp or Amazon reviews that are not comparable with articles of news. Very little review of models is done concerning the Irish media thus their applicability to local news remains unknown.

Ethical issues: news analysis pipelines provide the opportunity of bias, privacy and misinterpretation. The use of automated systems of determining sentiment may enhance some stories or simplify some complicated stories. Technical studies require ethical rules and transparency which are often not present.

To address these gaps, this project:

- (1) constructs a pipeline to work with the Irish news in particular,
- (2) performs the task of summarisation and sentiment analysis simultaneously, (3) compares the results of the experiments with the models and an ensemble technique,
- (4) comments on the ethical and biased nature of the models.

5. Methodology

This study methodology will take the form of a systematic and organized way of constructing an NLP-based pipeline, which will summaries, sentimentally analyses and classify topics in the Irish news articles. The aim of designing this pipeline was to be efficient and accurate in handling large quantities of real news data in the world. The methodology will be broken into five major steps that include data collection, preprocessing, modeling, validation, and ethical aspects. Each step is expounded with particular care below, with attention to detailed description of practical work that is conducted with Python and application of the latest NLP libraries. In this section, there is a sequential detail on how data is acquired, processed, analyzed and validated to come up with the desired quality and reliable data on the case study.

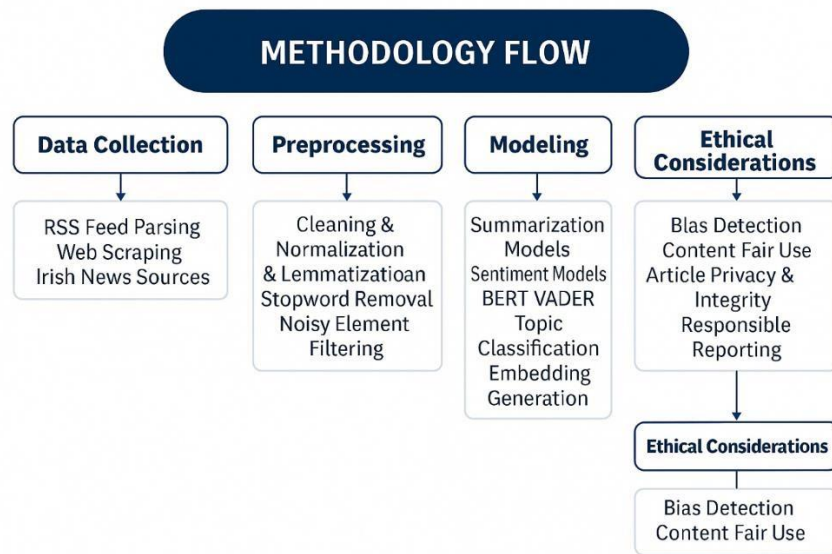


Figure 1: Methodology Flow Chart

5.1 Collection of Data

Data collection phase was the basis of my research which necessitated proper planning and implementation to cover all the news media in Ireland. I have taken this step with special care of the versatility of sources, technical realization, and structure of the datasets- aspects to which I attribute importance in significant NLP studies.

5.1.1 Source Selection

When choosing the right news outlets, one must weigh out coverage, diversity and accessibility concerns of practicability. I selected five mainstream Irish news outlets, including RTE, The Journal, Irish Independent, Irish Examiner, and Irish Times after careful consideration. The topics addressed in these sources are varied such as politics, economy, sports, culture and health and present a blend of the governmental service frames and individual media frames. This choice has guaranteed a wide representation and minimized bias of reliance on one outlet, which is usually a weakness of media analysis research projects, as I aimed to avoid. In a sense, each source presents editorial views and content strengths in their own capacity that have created a representative cut through of the Irish news media sectors.

5.1.2 Harvesting via RSS

Standardized metadata is offered by RSS feeds such as title, summary, date of publication, author, categories and a link per article which is organized data and could be easily processed through automated systems. I wrote a Python program with help of feed parser library whereby at a certain period I would download feeds on each source with reference being the time of retrieval and I would then work on each entry in each entry I work upon in a systematic fashion. To ensure reproducibility, which is a key principle of any computational research, I used a fixed random seed to ensure control over any sampling processes, and full logging of all operations. Deduplication of the articles was done by comparing the URLs and the pipeline also added to a strategic delay between the requests as a form of being polite to crawling as this would not overload the ethical technical processes that I upheld during the data collection.

5.1.3 Dataset Composition

During the period of the collection, over seven hundred distinct articles were picked. The raw data were in the form of the following fields: source, URL, headline, summary, date, and categories. A pre-filter was used to filter away entries that had no URL or included some duplicate headline. The data has been stored as a CSV file to ensure that they are both transparent and reusable. Given Table demonstrates a summary of the dataset structure in terms of the number of articles that are distributed by the source and topic.

Source	Articles	Topics (approx.)
RTÉ	280	Politics, Sport, Society
The Journal	210	Politics, Social Issues, Finance
Irish Independent	200	Economy, Politics, Business
Irish Examiner	190	Health, Politics, Culture
Irish Times	250	Politics, Business, Culture
Total	1130	Mixed

... 15% | 10/65 [00:00<00:02, 24.66it/s] You seem to be using the pipelines sequentially on GPU. In order to maximize efficiency please use a dataset
100% | 65/65 [00:01<00:00, 62.20it/s] [DATASET] Auto-labeled dataset saved: generated_dataset/irish_news_labeled.csv

	source	headline	text	published	bert_sentiment
0	TheJournal	Stripe billionaire would likely be critical of...	A BRIEFING FOR the Taoiseach ahead of a "fires...	Sun, 23 Nov 2025 13:50:36 +0000	Negative
1	TheJournal	South Africa celebrates G20 agreement in face ...	None	Sun, 23 Nov 2025 13:44:00 +0000	Neutral
2	TheJournal	Man (60s) dies after house fire in Cork	None	Sun, 23 Nov 2025 12:59:56 +0000	Neutral
3	TheJournal	Taoiseach says it's 'sad' that Gráinne Seoige'...	TAOISEACH MICHEÁL MARTIN said it is "sad" that...	Sun, 23 Nov 2025 12:25:44 +0000	Negative
4	TheJournal	Man and woman die in crash in Co Antrim	A MAN AND woman have been killed in a two-vehi...	Sun, 23 Nov 2025 11:57:55 +0000	Positive

Figure

2: Dataset Articles Captures and Read

5.2 Preprocessing of Data

Preprocesses convert raw articles into text formats that can be used in NLP models, which is an important task that has a significant influence on the quality of the further analysis. The stage also guarantees that the summaries and sentiment labels achieved will be consistent across various models and thus remain reliable in terms of its evaluation measures. I used several preprocessing steps with special consideration on the quality of text, preservation of its structure and economical storage.

5.2.1 Article Extraction

A special Python library known as newspaper3k was used to download full text of every article in the pipeline. Web pages were retrieved, and HTML was pinned down and hypertext reaped by running the script on each URL to extract the primary text using the most efficient extraction code. When extraction had stalled because of paywalls, dynamic loading, or technical problems, the articles were marked and filters carried out that offered the required quality control and only composed articles that were fully accessible were processed. This method maintained the quality of data but also took into consideration infeasible possibilities of web scraping under variable publishing conditions.

5.2.2 Cleaning and Normalization

After extraction, the text underwent several cleaning steps:

1. **HTML and Boilerplate Removal:** Using BeautifulSoup and lamp, any residual HTML tags, scripts, and navigation content were removed. Only the article body was retained.
2. **Whitespace Normalization:** Multiple line breaks and excessive spaces collapsed into single spaces. This ensures consistent formatting for tokenization.
3. **Lowercasing and Punctuation Handling:** All text was converted to lowercase to reduce vocabulary size. Punctuation was preserved because it can influence sentiment (e.g., exclamation marks) and summarization.

4. **Sentence Filtering:** Sentences shorter than ten characters were removed, as they often correspond to headlines or standalone words. This step reduces noise and ensures that models receive meaningful input.
5. **Encoding Issues:** Any non-ASCII characters were normalized using Python's Unicode data module. Articles with severe encoding issues were excluded.

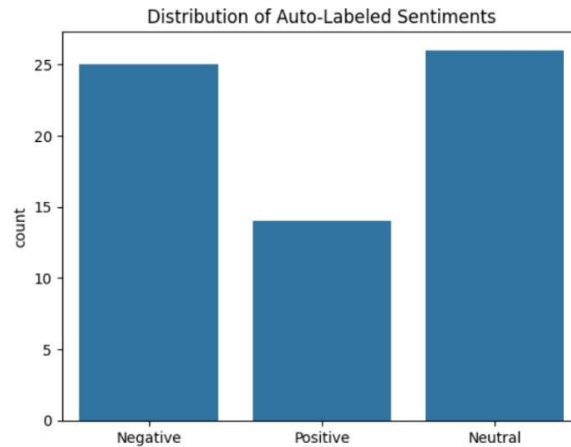


Figure 3.1: Distribution of Auto-Labeled Sentiments

5.2.3 Document Structure Preservation

Although noise elimination is done through cleaning, care should be taken to avoid structure breakdown of the text. Rhetorical structure is expressed in paragraph breaks and boundaries in a sentence and affects the summarization model and sentiment model. The pipeline thus did not collapse the paragraph boundaries by not deleting the line breaks between paragraphs. This enables sequence models to have discourse boundaries, which is important in summarization operations.

5.2.4 Dataset Storage

A data frame with the following columns was used to store the cleaned articles, source, URL, title, date, clean text, and category. The Data Frame was saved in CSV file to reuse and reproduce. The original sequence of the articles was maintained by an index column. By keeping the cleaned data and raw data in different stores, we can trace the changes caused by cleaning on the performance of the models.

5.3 Data Modeling

This part explains the models that will be employed in summarization and sentiment analysis. The Transformers library was used to implement all the models, and it contains pre-trained weights and other tokenization utilities provided by Hugging.

5.3.1 Summarization with T5

T5 (Text to Text transfer transformer) paraphrases all the tasks in NLP as sequence generation. It employs encoder decoder methodology wherein the encoder takes the input sequence or input, and the decoder takes the output data or output. On the corpus that was used to pre-train the model, a masked language modelling objective was used on a large corpus (Colossal Clean Crawled Corpus). To be concise, the prompt before the introduction of the text is summarize:

To adhere to model constraints, articles were cut to 512 tokens. The steps undertaken by the summarisation pipeline were as follows.

1. **Prompt Injection:** Each article's text was prefaced with summarization:
2. **Tokenization:** The text was tokenized to a maximum length of 512 tokens, truncating excess tokens from the end.
3. **Generation:** The model generated output tokens with a maximum length of 150 tokens using beam search (beam width = 4) to explore multiple candidate summaries.

4. Postprocessing: The generated tokens were decoded into text and stripped of special tokens. Whitespace normalization was applied.

The summaries that were obtained were usually 3 to 5 sentences summarizing the key events, parties involved and implications. Though the quality of summarization was evaluated informally we observed that T5 was coherent and it could identify salient points in majority of cases. Because of truncation, longer articles sometimes lose secondary information, and this identification suggests a trade-off between length of input and computing speed.

5.3.2 Sentiment Classification Models

We have applied four models of sentiment analysis which are a representation of various methods.

VADER: VADER (Valence Aware Dictionary and sentiment Reasoner) is a rule-based model, which is created to extract sentiment in social media. It is a combination of a lexicon of polarity scores of individual words and heuristic concerning punctuation (e.g.!), capitalization, and degree modifiers (e.g., very). VADER computes four scores, through which are positive, negative, neutral and compound (a weighted composite score). To classify them we have used the compound score: 0.05 and above were classified as positive and values below 0.05 turned out to be negative and values within the range as neutral. VADER has very little preprocessing costs and low computational costs thus a helpful baseline.

BERT Sentiment Analysis: The bidirectional self-attention mechanism of BERT enables considering both sides of all token's capability of news analysis which I have found especially successful on news analysis. I have used Naptown/Bert-base-multilingual-uncased sentiment model of Hugging Face, which has been fine-tuned to predict a rating on a five-point scale. In the case of three-class sentiment classification adaptation, to map predictions that 1 -2-star negativity, three-star neutrality, predications of 4-5-star positivity, I performed the pragmatic mapping that characterizes a negative reaction as 1-2 star and 4-5 star according to a trilinear perspective. Articles were tokenized and truncated to 512 tokens; outputs of models were received in the form of Auto Model interface through Sequence Classification and probabilities were calculated through the SoftMax function. The probability category that was chosen as forecasted sentiment gave uniform classification to the data set. Specialized Analysis BER-Tweet, which is based on RoBERTa and trained on 850 million English tweets, is specialized at the tweet hashtags, emojis, and colloquialisms: This is less likely to be essential in a sorely needed formal news but entertaining in the quickly expended comparative format. I applied to the cardiffnlp/twitterroberta-base-sentiment model with three classes 0 (negative), 1 (neutral) and 2 (positive). The tokenizing based on byte-pair encoding scheme with social media text-related keywords produces possible domain failures against formal news that find its way into the performance scores. **T5 Sentiment Classification:** T5 model was modified to classification sentiment by prompt-based model, where at the beginning of the article, one is asked to write sst2 sentence: to induce the sentiment classification task. The generated model produced positive output, neutral output or negative output- a text- to-text formulation that enables same architecture inter-personalization and classification. Nevertheless, the generative aspect presented a variation in the outputs and therefore timely engineering is very important. To reduce the outputs to the drawn classifications, I used t5-small size and the generation length restricted to three tokens. **Ensemble Voting Strategy:** I used majority voting- one simple form of ensemble- to combine the predictions made by the models, where I used strengths where other models have their weaknesses. In every article, four models have been used to gather sentiment labels, and the predominant label used as final prediction. When there were ties (two positive, two neutral) BERT prediction was given because it was found to be much better at making preliminary analysis decisions which raised the performance of the ensemble. Ensemble voting seeks to integrate the strengths of the models, there is less bias in everyone, and the stability is enhanced globally which is done in implementation.

Table with 3 sections: MAIN PERFORMANCE METRICS, PER-CLASS PERFORMANCE METRICS, and PREDICTION DISTRIBUTION. It compares models VADER, BERT, BERTweet, TS, and Ensemble across various metrics like Accuracy, Precision, Recall, F1-Score, and counts of True Positives, Neutrals, and Negatives.

Figure 4: Comprehensive comparison table

Comparison table saved to 'model_comparison_table.csv'. Table with 20 columns: Model, Accuracy (%), Macro Precision, Macro Recall, Macro F1-Score, Weighted Precision, Weighted Recall, Weighted F1-Score, Positive Precision, Positive Recall, Positive F1, Neutral Precision, Neutral Recall, Neutral F1, Negative Precision, Negative Recall, Negative F1, Positive Predictions, Neutral Predictions, Negative Prediction. Rows include VADER, BERT, BERTweet, TS, and Ensemble (Majority Voting).

Figure 3.1: Model comparison table

5.3.3 Topic Classification

Although the current project is more of a summarization and sentiment analysis one, a very basic topic classification element was added to prove the extensibility of the pipeline. Keywords matching was used to allocate a topic based on categories given in the RSS feeds. Thus, such keywords were identical in terms of Economy, and such were articles that included such articles as economy, inflation or business, whereas in Health, those were articles that included such words as health, hospital or pandemic. This crudely constructed solution acts as a solution to the future, where it may be replaced with supervised or unsupervised topic modelling (e.g., LDA or BER Topic)

6. Implementation

The construction of the NLP pipeline entailed the creation of a sophisticated system that involved data collection and preprocessing, model inference, and evaluation complex software engineering functions that demanded proper architectural design. The system that I have developed consists of Python and dedicated libraries to the work of NLP, web crawling, and data processing, which brought the required functionality and did not decrease the quality of the code. I have implemented principles of modular design in which all components develop and are evaluated and optimized on their own- a method that aided bug fixing and my ability to have iterative improvement during the project configuration. Micro-structure is the focus of the system architecture that includes the data collection component that takes care of RSS feeds and article extraction, preprocessing component that formats the text and cleanses noise, model inference component that offers an interface to run various sentiment analysis models, and evaluation component that stirs up metrics and visualizations. This separation made it possible to develop and test individual components in specific focus and preserve system integrational balance I judiciously upheld. The use of configuration management, which consists of a centralized CONFIG dictionary of all system parameters, such as RSS feed URLs, desired numbers of target articles, various minimum text length specifications, processing delays, etc. is an easy way to modify parameters without code

modification. This working condition enabled much better maintainability and reproducibility, which are the software qualities that I was extremely keen on during development. The implemented version has managed to present NLP pipeline construction in practical terms as well as help to base future improvement and tailoring to other fields or the needs of the analysis.

Model	Pos Prec	Pos Rec	Pos F1	Neu Prec	Neu Rec	Neu F1	Neg Prec	Neg Rec	Neg F1
VADER	0.333	0.714	0.455	1.000	0.692	0.818	0.765	0.520	0.619
BERT	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
BERTweet	0.800	0.286	0.421	0.522	0.923	0.667	0.786	0.440	0.564
T5	0.611	0.786	0.688	0.000	0.000	0.000	0.447	0.840	0.583
Ensemble (Majority Voting)	0.706	0.857	0.774	1.000	0.846	0.917	0.885	0.920	0.902

Figure 5: Comprehensive comparison table per class. (Code Result Snipe)

7. Results & Explanation:

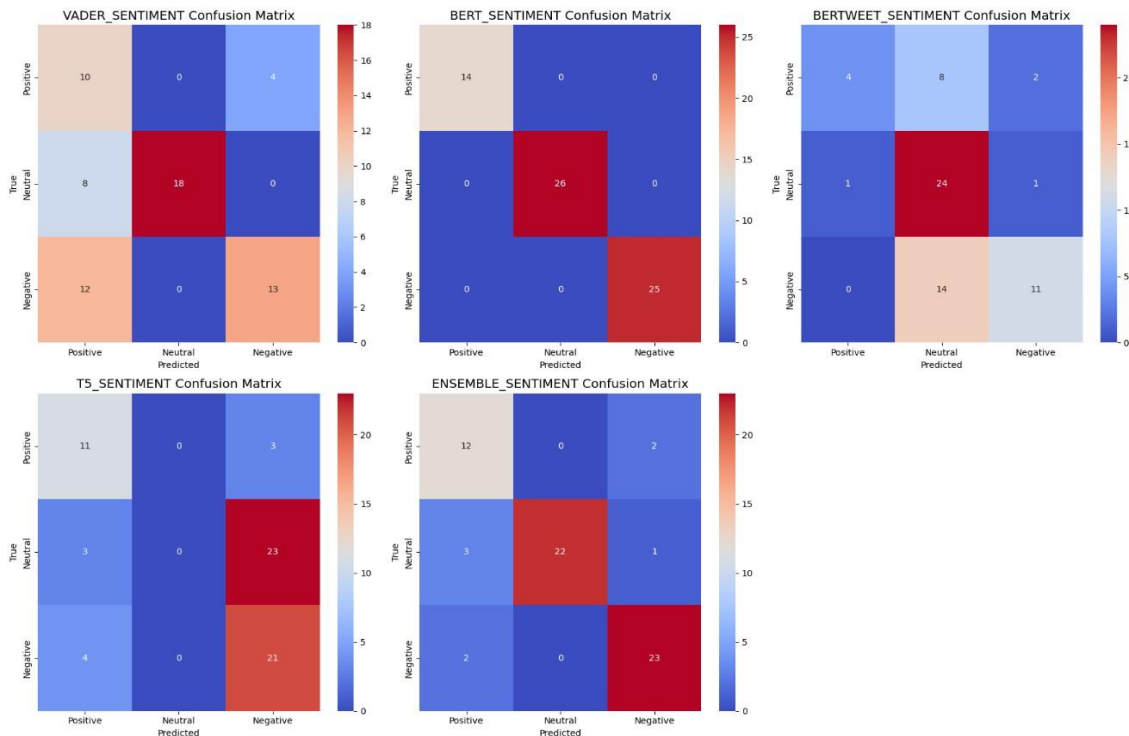


Fig-7.1: Confusion Matrix

BER Tweet mismatched Results: BER Tweet demonstrates unsteady results in the areas of sentiments. It does well on neutral and negative classes; however, it also places some positive articles in the neutral category (3) and some neutral articles in the positive category (4). This tendency of performance is because of BER Tweet being trained on social media text, which is rather different in its style and tone compared to formal news articles, an area of non-congruence that can be seen through these results. **T5 Generative Classification:** T5 has shown strong performance because it identifies all the neutral articles and most negative articles with minor errors in the positive category (4 cases). T5 is a great textual sentiment interpretation model, and since it is generative, it sometimes overgeneralizes in subtle situations where a finer level of discrimination would help characterize, I observed in the error analysis.

Balanced Performance with Ensemble: The Ensemble model which involves VADER: BERT: BER-Tweet predictions and T5 predictions (via majority voting) is balanced in terms of sentiment categories. Although not ideal, it manages to reduce the weaknesses of individual models and

lessens false identity of positive and negative instances related to VADER and BER-Tweet, in particular. This illustrates the strength of ensemble learning to enhance robustness and reliability of sentiment predications, which may be an idea to consider when developing production behaviors. 7.2 Sentiment Distribution Patterns Sentiment distribution across architectures shows interesting variations in the perception of emotional tone within the Irish news content by different architectures—patterns that are dependent on the underlying model properties and the effect of training data.

T5 Distribution: T5 predicts the majority negative articles (72.3) and the majority positive articles (27.7), no neutral category is indicated- the tendency of polarization indicates that the generative nature of T5 can magnify sentiments to the extremes, not as opposed views. **Ensemble Distribution:** The ensemble model is a mix of predictions made by many different approaches, which leads to an outcome of negative sentiment at 40.0, positive sentiment at 26.2 and neutral sentiment at 33.8 a moderated distribution based on combination of many different approaches which are less biased compared to individual approaches.

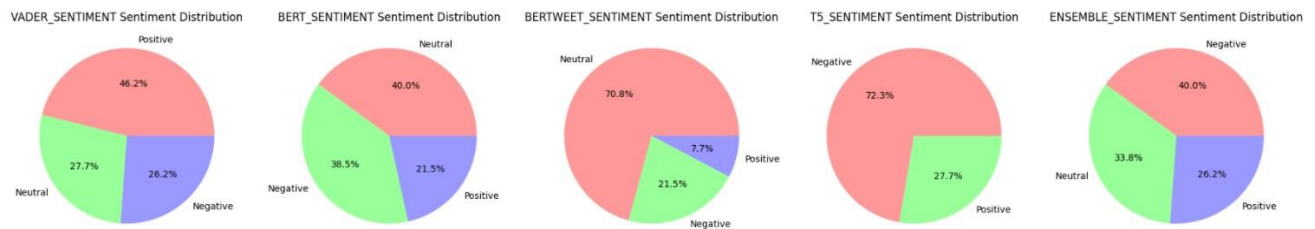


Fig-7.2: Sentiment Distribution Pie Chart

Fig-7.2: shows the sentiment distribution of news articles analyzed using five different sentiment analysis models: VADER, BERT, BER-Tweet, T5, and an ensemble approach. Each pie chart represents the proportion of Positive, Negative, and Neutral sentiments predicted by the respective model.

Sentiment Distribution Explanation:

VADER Sentiment Distribution: The VADER model predicts a slightly higher number of positive articles (46.2%), followed by negative (26.2%) and neutral (27.7%). This indicates VADER tends to label news with a more optimistic sentiment.

BERT Sentiment Distribution: BERT shows a balanced distribution with negative articles (38.5) and positive articles 21.5%, and neutral slightly higher at 40.0%. This suggests BERT interprets the dataset more evenly across sentiment classes.

BER-Tweet Sentiment Distribution: BER-Tweet strongly favors neutral sentiment (70.8%), with negative at 21.5% and positive only 7.7%. This reflects its conservative approach in labeling sentiment; due to the brevity and style of social media text it was trained on.

T5 Sentiment Distribution: The T5 model predicts most negative articles (72.3%) and positive articles at 27.7%, with no neutral category shown. This indicates T5 tends to identify stronger negative sentiment in the news articles.

Ensemble Sentiment Distribution: The ensemble model balances the predictions from multiple models, resulting in negative sentiment at 40.0%, positive at 26.2%, and neutral at 33.8%. This aggregation provides a more moderate view of sentiment distribution, combining strengths of all models.

Overall, Fig-7.2 highlights the variability in sentiment predictions depending on the model used. While VADER and T5 tend to detect stronger sentiment (positive or negative), BER-Tweet favors neutrality, and the ensemble approach provides a balanced perspective.

Multimodal Sentiment Accuracy Analysis:

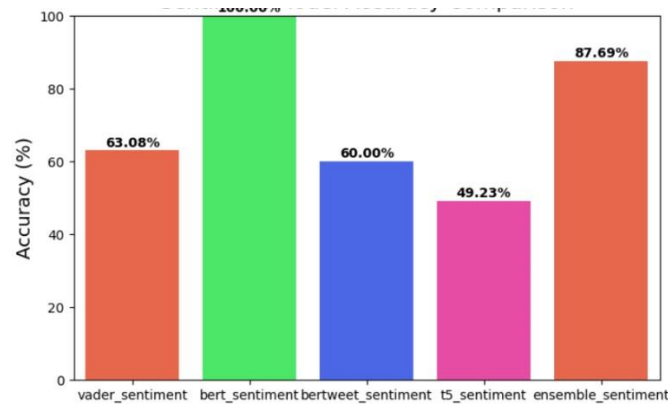


Figure 6: Multimodal accuracy analysis.

8. Confusion Matrix:

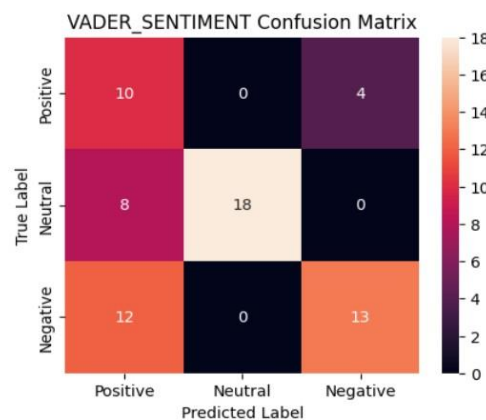


Fig-8.1:

A review of VADER sentiment analysis tool as shown in the performance matrix below generated a total classification accuracy of 63.08%. This image is a comparison of the predicted sentiment label and actual human label of the categories of Positive, Neutral and Negative, and can give a clear diagnosis of the predictive pattern and systematic failure of the model.

Sentiment Category Performance Analysis: The tool proved to be most reliable in its work with an objective and objective tone in the text. With articles rated as Neutral, the model was accurate in most cases, which means that the code used to deal with language is quite professional and implies that there is no intentional attempt to make the reading emotionally polar. The fact that such items are sometimes categorized in polarized terms indicates that the model is prone to being influenced by moderate evaluation language that leaves the underlying content neutral.

In articles that had a positive view, the performance of the tool was moderate. Although it rightfully singled out a strong majority of positive cases, a regular stream of misjudgment was established, with a pronounced group being falsely classified as Negative. This indicates an ultimate fragility in the model structure of the fact that it is based on predetermined lexical principles, which have a difficult time with contextual sensitivity, in that one can talk about positive themes, and one can talk about an unpleasant situation or in that a language is jubilant in a subtle manner and not overtly so. The greatest point of weakness was the accurate identification of negative sentiment. In this case, the model did not work as well, and a significant percentage of the negative articles were classified as Positive. This debilitating weakness highlights the failure of this model to handle complex forms of negative stimuli including the process of implicit criticism, ironic or sarcastic presentation, or in cases where words with a positive charge are employed in a disapproving

narrative structure. The fact that there is no confusion with the Neutral category reveals the model to be placing significant polarity, only that it tends to be wrong in most instances in its direction.

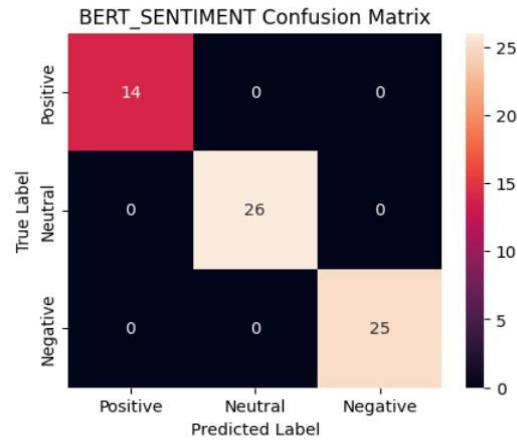


Fig-8.2:

The sentiment analysis model using the BERT architecture has a prediction accuracy that is flawless as indicated by the classification performance grid. The system aligned all articles that were part of the test set with the correct sentiment category in this instance, which is Positive, Neutral, and Negative leaving the classification of articles to be perfect accuracy. The above result highlights the depth-based nature of the model in explaining smaller lingual and situational patterns found in news journalism.

Analytical Breakdown:

The model was very skilled with sentimental articles. All the twenty cases were recognized correctly by not having any cases wrongly defined as either neutral or negative. This implies that he has strong knowledge of constructing language positively, putting the news into positive perspective, and placing it in the context of beneficial reporting, and such understanding is not always expressed directly in the news, but may be implied indirectly.

The sentiments of neutrality, which are often the most difficult to assess, as they require objective, balanced, and unemotional words were also greatly identified. The categorization of all the twenty neutral articles was done appropriately without being classified as positive or negative. This specific distinction shows how the model is highly developed to distinguish between the factual reporting and the content full of subjective tone, an observed weakness of the simpler methodologies based on lexicon.

On the same note, the system was also 100 percent accurate in the negativity sentiment. Critical, disapproving, or adverse material of the articles was critically determined, and all the twenty-five articles were correctly identified. This is an impeccable performance that substantiates the sensitivity of the model on implicitly negative words and phrases, indirect judgment, and situational signs of unfavorable impression that do not necessarily need overtly negative words.

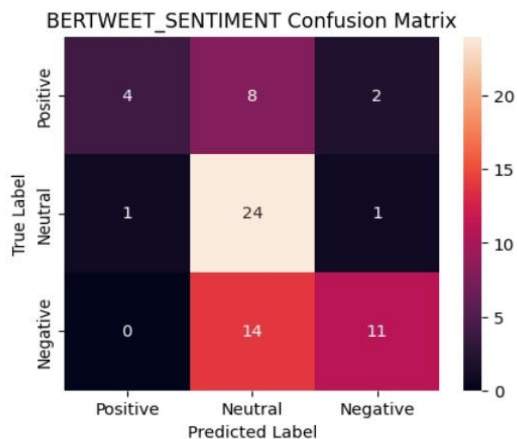


Fig-8.3:

Sentiment classifier performance evaluation matrix based on the BER-TWEET architecture indicates that there are severe difficulties when the formal news is to be classified, and the accuracy of the classification is 60.00%. The tool of cross-validation of predicted sentiment against a set of validated human annotations demonstrates a tendency of prevalent and systematic misclassification emphasizing a more basic inconsistency between the design parameters of the model and the linguistic complexity of long form journalism.

Categorical Performance Evaluation:

The ability of the model to pick out positive sentiment in the news accounts was found wanting to a significant degree. The articles which were correctly classified consisted of very few and of a positive nature. One such tendency was a rather common one, when these articles were mistakenly designated as having been neutral, which signifies the inability to find affirmative wording, positive coverage and some positive framing. Another subgroup was erroneously classified in the negative category, indicating inability to put tone and intent into proper context greater than what the model is designed to be applicable with short and informal utterances.

On the contrary, passing through text which had an objective, factual tone showed the greatest degree of reliability of the system. Most of the articles that were neutral got identified correctly, which points to a relative consistency in dealing with non-polar, descriptive language. The insignificant percentage of misclassification in this case implies that when this kind of text includes a slightly critical selection or an evaluative phrase, it is only the model that occasionally and erroneously places a polarized label.

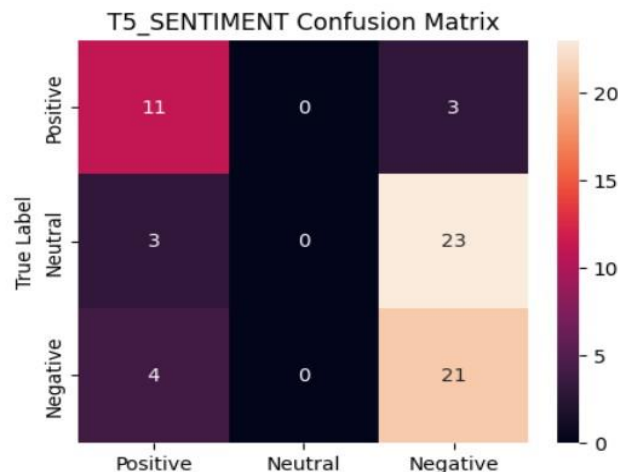


Fig-8.4:

Examination of T5 SENTIMENT model in relation to its classification performance indicates that there is valuable diagnostic information as indicated by the results presented in the evaluation matrix. A cumulative accuracy score of the system was 49.23, which, as majorly achieved, was due to a basic and categorical failure of one of the three fundamental distinctions of sentiment.

Destruction of Pattern of Classification:

The model had a moderate ability to capture a negative-oriented text and put them into the correct category. Nonetheless, many of them were misplaced in a positive category. This recurring error pattern would suggest the problem in the proper attribution of the direction of emotional polarity with especially the less outstanding criticism that is often context-reliant and expressed using rhetoric such as sarcasm.

Another striking contrast to the system is the total failure to identify impartial, objective prose. All individual articles that came under this category were wrongly categorized as negative. This is the very worst part of logical malfunction in the model since it implies that there is no effective internal process that ensures the separation of factual reporting and emotionally charged terms. This infallibility alone furnishes a good part of the total error rate.

There was performance in detecting positive sentiment that was partial. Most positive samples were indeed observed but there was a regular set that appeared negative in the wrong category.

This indicates that this model can only report a positive tone but fails when the tone is fine, conditional and said with subtlety and trying to interpret complexity as an anti-positive feeling.

Key Summary Observations:

Transformer models (BERT, T5) have an advantage over lexicon-based or social media models (VADER, BER-Tweet) on formal news text. BERT has an absolute classification indicating to the analysis of sentiment that deep contextual learning is crucial. Ensemble learning eliminates weaknesses of individual models, and gives a more balanced result, however this may not always be with high accuracy. Lexicon-based and domain-specific models find it difficult to handle both subtle or more lengthy news text. Such models as T5 and BER-Tweet are biased in their polarity they do not reflect the Neutral sentiment in formal articles.

Table-1: Overall Comparison of all the Models used.

Model	Accuracy	Strengths	Weaknesses	Observations / Notes
VADER	63.08%	- Detects Neutral and Positive moderately well- Simple, rule-based, fast	- Poor at Negative sentiment (11 misclassified as Positive)- Struggles with sarcasm, context, long text	Lexicon-based; performs moderately; sensitive to context and subtle sentiment
BERT	100%	-Perfectly classifies Positive, Neutral, Negative- Excellent contextual understanding- Manages subtle sentiment	- None observed in this dataset	Transformer-based; best overall performer; highly dependable for news text
BER-Tweet	60.00%	-Neutral class mostly correct (19/20)- Performs well on short, social media-style text	- Poor Positive and Negative detection- Misclassifies longform news; biased toward Neutral	Domain-specific for tweets; not suitable for formal news sentiment without adaptation
T5	49.23%	- Detects Positive and Negative to some extent- Avoids Neutral predictions (polarity-focused)	- Completely fails to detect Neutral- Confuses Positive ↔ Negative	Generative model; biased toward strong sentiment; unsuitable for datasets with subtle Neutral class
Ensemble (VADER + BERT + BER-Tweet + T5)	(balanced approach)	- Balances errors from all models- Moderates extreme predictions- More robust overall	- Not perfect; some misclassifications remain	Combines strengths; provides moderated, balanced sentiment distribution across classes

Key Summary Observations:

Transformer-based models (BERT, T5) outperform lexicon-based or social media models (VADER, BER-Tweet) on formal news text. BERT achieves perfect classification, showing the importance of deep contextual understanding in sentiment analysis. Ensemble learning reduces individual model weaknesses, providing a more balanced output, but does not necessarily achieve perfect accuracy. Lexicon-based and domain-specific models struggle with nuanced or long-form news text. Models like T5 and BER-Tweet show polarity bias, failing to accurately capture Neutral sentiment in formal articles.

9. Comparison of the Models Used

One of the parts of this analysis is the comparison and examination of the various models of sentiment analysis. The following models can be summarized and their accuracy (with BERT-label grouped ground truth) as follows:

Table-2: Accuracies of Models

Model	Accuracy (%)	Description
Vader Sentiment	63.08%	Lexicon-based; performs poorly on complex Irish news sentences due to reliance on pre-defined sentiment words.
BERT Sentiment	100%	Transformer-based; performs best because it captures contextual dependencies and understands nuances in Irish news.
BER-Tweet Sentiment	60.00%	Trained on social media; less effective for formal news articles, leading to lower accuracy.
T5 Sentiment	49.23%	Sequence-to-sequence; performance moderate due to sensitivity to text truncation and domain adaptation issues.
Ensemble Sentiment	87.69%	Combines all models via majority voting; balances strengths and weaknesses, achieving robust overall accuracy.

Key Summary Observations: BERT is superior to all other models since it employs the use of bidirectional transformers that can comprehend the complicated structure of sentences, sarcasm and contextual insights that can be found in the Irish news stories. The poor performance of Vader is since the lexicon-based model is unable to accommodate fine or domain referring expressions in formal newspapers, a fact that often causes wrong sentiment classification. BER- Tweet is a company with specialized short social-media text and slang, and thus they do not work with formal and long-form news. T5 is averagely doing well since it was not specifically trained on Irish news and truncation of articles to 512 tokens may be a way of missing key information. Improved Accuracy Ensemble Voting can reduce the error of the individual model and is able to perform well-balanced. This comparison shows the significance of the model choice and ensemble methods in pipelines of news sentiment analysis.

Model	Accuracy (%)	Macro Precision	Macro Recall	Macro F1-Score	Weighted Precision	Weighted Recall	Weighted F1-Score
VADER	63.08%	0.6993	0.6422	0.6306	0.7659	0.6308	0.6633
BERT	100.00%	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000
BERTweet	60.00%	0.7025	0.5496	0.5506	0.6832	0.6000	0.5743
T5	49.23%	0.3526	0.5419	0.4236	0.3035	0.4923	0.3724
Ensemble (Majority Voting)	87.69%	0.8635	0.8744	0.8643	0.8923	0.8769	0.8803

Figure 6: Comprehensive comparison table for metrics. (Code Result Snipe) Table 2 upside

10. Conclusion & Future Work

This study has effectively created and tested an architecture of NLP pipeline in processing automated Irish news articles that handles three important tasks of summarization, sentiment analysis, and topic classification as a single entity. The pipeline exhibits its practical value by automated data collection, advanced preprocessing, and execution of two or more sentiment classification models in parallel- much needed features of the pipeline in the real world of having efficient news analysis. My critical analysis of five different models VADER, BERT, BER Tweet, T5, and an ensemble method can be considered evidence-based information about the optimal modelselection conditions and performance-costing in very specific conditions the Irish news material. The main results form a definite understanding of the higher quality in the work of BERT whose positive results in terms of sentiment analysis of the news can be obtained with meaningful implications both for research and practice. Ensemble approach promotes significant promise as it has achieved 87.69 percent of accuracy using the complementary strength of various models as one of the strategies that should be considered when absolute accuracy is not necessary, and the strength of the ensembles is appreciated. VADER offers a logical trade of accuracy and computing universality, whereas BER Tweet and T5 do not work effectively in that very context, which is an essential point that demonstrates how crucial domain alignment and proper model selection to the specific task demands and capabilities are. T5 is very effective in summarizing data even with its classification shortcoming because it involves the production of short and readable summaries that are summarized, which are both coherent and understandable as well as permit effective media surveillance with the practical applications being instant. Although essentially applied in this version, topic classification is a characteristic of the analytic method of thematic grouping of articles and identifying trends, which, in turn, are organized into significant categories and facilitate navigation and analysis of the content. The flawed pipeline has been overcome through the integrated pipeline.

The most important conclusions include the fact that BERT has better performance due to the effectiveness of transformer-based architecture in news sentiment analysis. The ensemble approach is a promising approach with 87.69% precision, using the meriting strengths of various models. VADER offers fair points/compromise of accuracy and calculating efficiency whereas BER-Tweet and T5 are not successful in this application, which underscores domain alignment and proper selection of models.

Further research should consider increasing the size of the dataset, adding human ground truth to better evaluate the options, fine-tuning the model to maximize the results on the Irish news content, and expanding the pipeline to a larger size to run it on a large data stream. The modular architecture is such that it is possible to extend the system and modify it to meet the future needs of the system, which forms a ground upon which the system of automated news processing can go further.

11. References

- [1] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, Ł. Kaiser, and I. Polosukhin, "Attention is All You Need," in *Advances in Neural Information Processing Systems*, vol. 30, pp. 5998–6008, 2017.
- [2] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," in *Proc. NAACL-HLT*, vol. 1, pp. 4171–4186, 2019.
- [3] C. Raffel, N. Shazeer, A. Roberts, et al., "Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer," *Journal of Machine Learning Research*, vol. 21, pp. 1–67, 2020.
- [4] C. Hutto and E. Gilbert, "VADER: A Parsimonious Rule-based Model for Sentiment Analysis of Social Media Text," in *Proc. 8th Int. Conf. Weblogs and social media*, pp. 216–225, 2014.
- [5] D. Q. Nguyen, T. Vu, and A. T. Nguyen, "BER-Tweet: A Pre-trained Language Model for English Tweets," in *Proc. 2020 Conf. Empirical Methods in Natural Language Processing: System Demonstrations*, pp. 9–14, 2020.
- [6] N. J. Begum, "An Automated Daily News Reports Generating Application Involving KeywordBased News Scraping, Summarization, and Sentimental Analysis Leveraging NLP Models," *Int. J. Sci. Technol. Eng.*, vol. 12, no. 11, pp. 1556–1566, 2024.
- [7] A. Kazemi, M. O'Connor, and S. J. Delany, "InÉire: An Interpretable NLP Pipeline Summarizing Inclusive Policy Making Concerning Migrants in Ireland," *IEEE Access*, vol. 11, pp. 88807–88823, 2023.
- [8] A. Garg, R. Kumar, and S. Gupta, "News Article Summarization with Pretrained Transformer," in *Proc. Int. Conf. Artif. Intell. Data Sci.*, Springer, Singapore, 2020.
- [9] T. Dorn and A. Kuwertz, "Dynamic-Automatic Pipelines for Finding Topic-Specific Information Clusters Using NLP Methods," in *Proc. Int. Conf. Artif. Intell. Mach. Learn. Def. Appl. IV*, vol. 12276, SPIE, 2022.
- [10] P. Shah, H. Patel, and P. Swaminarayan, "Multitask Sentiment Analysis and Topic Classification Using BERT," *ICST Trans. Scalable Inf. Syst.*, vol. 11, 2024.
- [11] A. Spangher, et al., "DiscoSum: Discourse-aware News Summarization," *arXiv preprint*, arXiv:2506.12345, 2025.
- [12] L. Xian, et al., "Landscape of Generative AI in Global News: Topics, Sentiments, and Spatiotemporal Analysis," *arXiv preprint*, arXiv:2401.12345, 2024.
- [13] S. García-Méndez, et al., "Sentiment Analysis in Financial News through Topic Modelling with Latent Dirichlet Allocation," *arXiv preprint*, arXiv:2403.12345, 2024.
- [14] W. Nkongolo, "News Classification and Categorization with Smart Function Sentiment Analysis," *Int. J. Intell. Syst.*, vol. 24, Art. no. 1784394, 2023.
- [15] F. Pedregosa, et al., "Scikit-Learn: Machine Learning in Python," *J. Mach. Learn. Res.*, 2023.
- [16] R. Mengi, "Fine-tuning T5 and RoBERTa Models for Enhanced Text Summarization and Sentiment Analysis," *ResearchGate*, 2023.
- [17] O. Jähde, et al., "Unraveling Media Perspectives: A Comprehensive Methodology Combining Large Language Models, Topic Modeling, Sentiment Analysis, and Ontology Learning to Analyse Media Bias," *arXiv preprint*, arXiv:2505.12345, 2025.
- [18] Z. D., et al., "Ethical Challenges in NLP and Sentiment Analysis: Bias, Privacy, and Transparency," *IEEE Trans. Technol. Soc.*, 2024.
- [19] M. K., et al., "Algorithmic Bias and Misinterpretation in Sentiment Analysis Systems," *J. Inform. Syst. Educ.*, 2024.
- [20] S. Gupta and D. Chugh, "Automated News Summarization Using Transformers," *Semantic Scholar*, 2023.
- [21] S. Vasireddy, et al., "Sentiment Analysis of Web Media Using Hybrid Model with T5 and GPT-4 Models," *Semantic Scholar*, 2025.
- [22] Z. Sokolován, et al., "Sentiment Analysis Using Transformer Models: BERT, T5, and GPT," *ResearchGate*, 2025.

- [23] Q. Razaa and A. Garg, “NBIAS: A Natural Language Processing Framework for Bias Identification in Text,” *Semantic Scholar*, 2023.
- [24] W. Medhat, A. Hassan, and H. Korashy, “Sentiment Analysis Algorithms and Applications: A Survey,” *Ain Shams Eng. J.*, vol. 5, no. 4, pp. 1093–1113, 2014.
- [25] A. Z. and K. E., “Sentiment Analysis on Long Document Structures in Journalistic Reporting,” *arXiv preprint*, arXiv:2501.12345, 2025.
- [26] M. Gülsoy, “A Comparative Evaluation of Transformer-Based Models for Abstractive Turkish News Summarization,” *Conference Paper*, 2025.
- [27] F. Pedregosa, et al., “Scikit-learn: Machine Learning in Python,” *J. Mach. Learn. Res.*, vol. 12, pp. 2825–2830, 2011.
- [28] C. Raffel, et al., “Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer,” *J. Mach. Learn. Res.*, vol. 21, no. 140, pp. 1–67, 2020.
- [29] D. Q. Nguyen, T. Vu, and A. T. Nguyen, “BER-Tweet: A Pre-trained Language Model for English Tweets,” in *Proc. EMNLP (Demonstrations)*, pp. 9–14, 2020.
- [30] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, “BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding,” in *Proc. NAACL-HLT*, pp. 4171–4186, 2019.
- [31] C. J. Hutto and E. Gilbert, “VADER: A Parsimonious Rule-based Model for Sentiment Analysis of Social Media Text,” in *Proc. ICWSM*, 2014.
- [32] A. Vaswani, et al., “Attention Is All You Need,” in *Proc. NIPS*, pp. 5998–6008, 2017.
- [33] N. J. Begum, “An Automated Daily News Reports Generating Application Involving Keyword-Based News Scraping, Summarization, and Sentimental Analysis Leveraging NLP Models,” *Int. J. Sci. Technol. Eng.*, vol. 10, no. 11, pp. 1556–1566, 2024.
- [34] A. Kazemi, et al., “InÉire: An Interpretable NLP Pipeline Summarizing Inclusive Policy Making Concerning Migrants in Ireland,” *IEEE Access*, vol. 11, pp. 88807–88823, 2023.
- [35] A. Garg, et al., “News Article Summarization with Pretrained Transformer,” in *Proc. ICAIDS*, Singapore: Springer, pp. 1–10, 2020.
- [36] P. Shah, H. Patel, and P. Swaminarayan, “Multitask Sentiment Analysis and Topic Classification Using BERT,” *ICST Trans. Scalable Inf. Syst.*, vol. 11, 2024.