

IFC-Bank of Italy Workshop on "Data science in central banking"

18-20 February 2025

Extracting economic issues from news data: with the help of generative AI¹

Younghwan Lee,
Bank of Korea

¹ This contribution was prepared for the workshop. The views expressed in this publication are those of the authors and do not necessarily represent the official views of the Committee, its members, or the BIS.

Extracting economic issues from news data: with the help of generative AI

Younghwan Lee¹

Abstract

This study proposes a method for extracting relevant economic issues in a timely manner and presenting them succinctly. Specifically, I introduce a method for directly assessing economic issues using news data. By closely tracking trending keywords in economic news articles, I identify key terms warranting attention during different periods. Based on these keywords and their contexts, several important issues driving the economic dynamics were inferred using the capabilities of generative AI. The primary advantage of this method is twofold. First, it provides policymakers with a powerful tool to directly focus on issues driving economic dynamics, unlike traditional approaches that indirectly identify the shocks behind economic variables. This is particularly important because policymakers deal with real issues and not mere economic variables. Second, the method employs algorithms that process vast volumes of data that are updated daily, making them timely, readily available, and potentially less prone to individual biases. The identified issues resulted primarily from the algorithm, with only minor refinements, allowing for replication with minimal human intervention. The usefulness of the proposed method is demonstrated through an example implementation using Korean news data from 2005.

Keywords: Economic Issue Extraction, Economic Trends Detection, Generative AI, Natural language processing

Table of Contents

Extracting economic issues from news data: with the help of generative AI	1
1. Introduction	2
2. Implementation	4
2.1 Sentiment Classification	5
2.2 Trending Keyword Extraction	7
2.3 Summarization	9
3. Concluding Remarks	10
References	11

¹ Gangneung Branch, Bank of Korea

E-mail: yhlee@bok.or.kr

Disclaimer: The views expressed in this article do not necessarily reflect those of Bank of Korea.

1. Introduction

One of the most important mandates of a national statistical institute is to keep policymakers informed of current economic issues, so that they can take timely and effective policy actions. Traditionally, economic indicators such as GDP, inflation, and sentiment indices have been regularly published and closely monitored to fulfill this mandate. This raises the fundamental question: Are traditional indicators sufficient to capture the evolving economic landscape? Should policymakers rely on well-defined statistics to understand the economy?

This is analogous to how doctors diagnose patients. Symptoms, such as body temperature and blood pressure, provide valuable information about a patient's condition. However, the ultimate goal of doctors is not just to relieve symptoms but to cure the underlying disease. Similarly, economic indicators reflect outcomes and not necessarily the root causes of economic issues that policymakers should truly be aware of. If we focus only on economic statistics, we risk confusing the indicator with the true underlying issue. To develop effective policies, we must move beyond superficial symptoms and directly analyze the economic forces shaping these indicators.

Despite the apparent need to assess trending economic issues directly, their linguistic nature poses a major challenge to econometricians and statisticians. Traditional statistics are inherently numeric, allowing for quantitative analysis such as direct comparison and aggregation. On the contrary, economic issues are often expressed in text rather than numbers. Consequently, conventional statistical tools cannot be applied directly.

For a long time, this was our good excuse—economic text was too complex to process effectively. However, recent breakthroughs in Natural Language Processing (NLP) and generative AI have eliminated this issue. These advancements have enabled us to analyze economic issues directly from text data in a systematic and scalable manner.

To assess key economic issues directly from a vast volume of text data, I present a method augmented by a Large Language Model (LLM), which is also known as generative AI. This method consists of three key steps: sentiment classification, trending keyword extraction, and summarization. First, each sentence in the collected articles is classified according to its revealed sentiment as positive, negative, or neutral. Second, trending keywords are extracted from sets of positive and negative sentences. That is, keywords are identified from negative sentences to monitor emerging negative issues and vice versa for positive sentences. Third, sentences containing each trending keyword are embedded into a semantic vector space and clustered based on their similarity. Each cluster is then summarized using a pretrained Generative AI model, producing concise descriptions of the key issues associated with each keyword.

One may wonder why summaries are not directly generated from large text corpora using large language models. However, this approach has significant practical limitations. General-purpose models are computationally intensive and impractical for daily, large-scale processing. More importantly, their output often lacks specificity and verifiability. For instance, when asked to explain a surge in negative sentiment in news articles, the model may return a generic statement such as "Macroeconomic conditions have worsened due to various factors," without identifying concrete

drivers or providing links to underlying evidence. These models are also susceptible to hallucinations and generate confident but unfounded explanations. Such shortcomings are critical for policy applications that demand transparency and accountability.

The proposed method, which is simple and effective, has three key advantages: (i) verifiability, (ii) flexibility, and (iii) ease of implementation. First, verifiability arises from the modular structure of the method. Each step—sentiment classification, trending keyword extraction, and clustering with summarization—produces intermediate outputs that can be inspected and traced back to the original text, ensuring transparency and accountability. Thus, policymakers and researchers can verify the final results down to the level of the original data — the raw sentences actually published in newspapers — ensuring full transparency and reliability.

Second, the method can be customized for specific tasks flexibly without significant modifications to the existing framework. A simple example is the use of a predefined set of keywords to monitor specific domains of economic activity, such as “exchange rate.” This small adjustment allows policymakers to remain informed of up-to-date issues related to exchange rate dynamics. A more intriguing application would be to use this method as a tool to monitor the transmission and interpretation of central bank communications in the media². By restricting the target sentences to those that mention or quote central bank statements, the framework can be adapted to track how policy messages are conveyed, distorted, or received by the public and market participants. This opens up the possibility of assessing not only the content but also the effectiveness of monetary policy communication.

Most importantly, the proposed method is easy to implement. Instead of reinventing the wheel, I assembled existing building blocks, such as the K-means clustering algorithm from the open-source FAISS library and the Word2Vec algorithm from Mikolov et al. (2013), with minimal modifications. In addition, the summarization step was powered by an open-source pretrained LLM, Llama 3.2. This approach allows practitioners to adopt and scale a system easily without the need for complex in-house model development, making it practical and accessible for real-time policy applications. Moreover, because the method relies on well-established algorithms with known computational costs, it can be efficiently executed on high-performance laptops, eliminating the need for specialized hardware.

The remainder of this paper is organized as follows. Section 2 details the implementation of the proposed method, outlining the conceptual framework and practical considerations. Each of the three main components—sentiment classification, keyword extraction, and clustering with summarization—is described, with attention paid to design choices and the rationale behind them. Section 3 concludes the paper with a discussion of remaining challenges and directions for future research.

² I am grateful to Bruno Tissot for bringing attention to the importance of central bank communication during the workshop discussion, which inspired the idea of applying this method to monitor how central banks’ messages are conveyed in the media.

2. Implementation

The core idea is to go beyond conventional numerical indicators and directly assess the underlying economic forces that drive fluctuations in a systematic and scalable manner. The proposed method builds a system that processes a large volume of economic news articles and returns concise summaries. These summaries capture the most relevant economic issues within a predefined target period, categorized by sentiment and thematic content. A three-step approach was adopted to ensure that the process remains verifiable, flexible, and easy to implement. Figure 1 briefly illustrates the overall procedure.

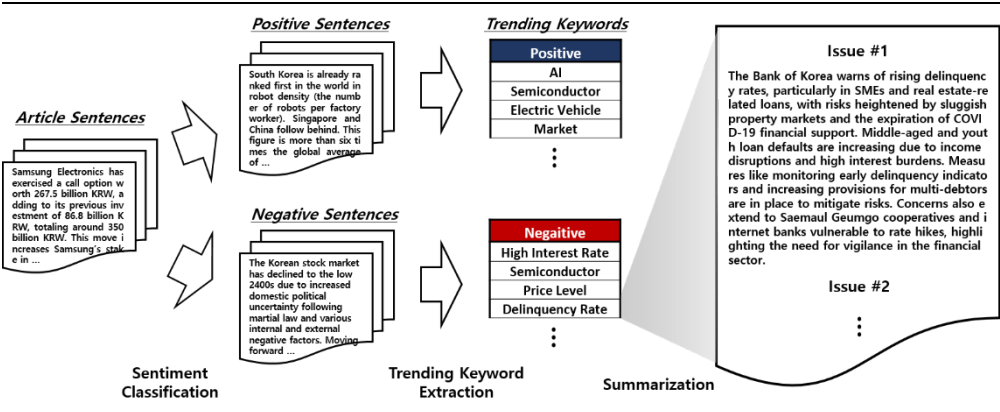
First, each sentence in the article corpus is classified by sentiment (positive, negative, or neutral) using a sentiment classifier trained on a large volume of human-labelled data. To properly assess the significance and economic impact of a factual event, the sentiment in which it is reported must be considered. For example, the statement “exports increased by 5%” may appear benign or even mildly favorable at first glance. However, if it frequently appears in a negative context, such as when driven by a sharp depreciation of the domestic currency, it should be interpreted with caution, as it may reflect broader economic concerns.

Second, trending keywords for the target period are extracted separately from the pools of positive and negative sentences. These keywords represent recurring themes or economic concepts gaining salience in the media, such as “inflation,” “COVID-19,” or “AI.” A trending score is defined and calculated for each keyword to measure its informativeness with respect to the target period. Keywords with high trending scores are selected as the trending keywords.

The final step groups sentences containing each keyword into semantically coherent clusters. Each cluster is then summarized using a pretrained generative model, yielding concise and readable descriptions of the underlying economic issues. This three-stage pipeline transforms raw text into structured outputs that are verifiable, policy-relevant, and easy to interpret. The following sections describe each step in greater detail.

Snapshot

Figure 1.



Source: Author’s illustration based on implementation results.

2.1 Sentiment Classification

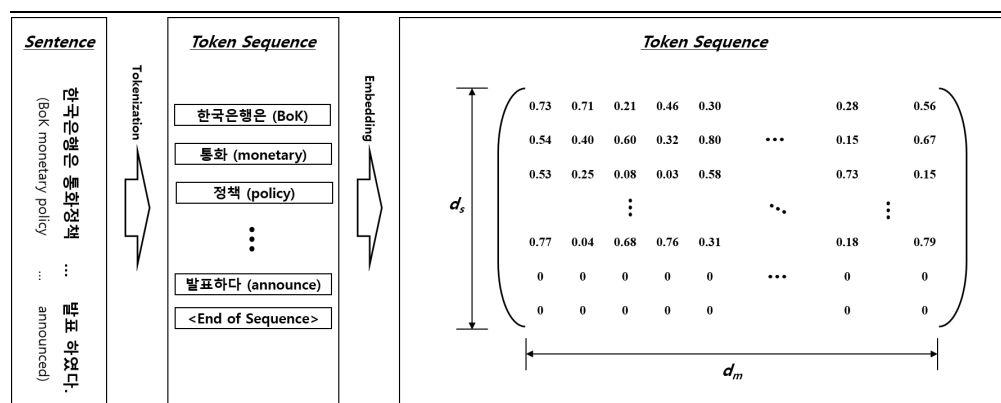
The sentiment classifier can be viewed as a function $C: D \rightarrow L$ that maps a sentence $d \in D$ to a sentiment label $l \in L$, where D denotes the set of article sentences and $L = \{\text{"positive", "negative", "neutral"}\}$. Rather than estimating C directly, I decompose it into two components: a preprocessing function $h: D \rightarrow E$ that transforms text into a numerical representation called embedding where E represent the embedding space, and a classification function $g: E \rightarrow L$ that assigns sentiment labels based on these representations. This modular design enhances flexibility in implementation and improvement. Using the preprocessing procedure, classification function g was calibrated using a labelled dataset consisting of approximately 400,000 sentences annotated by human operators.

The goal of preprocessing is to convert unstructured sentences into structured numerical representations that are both computationally tractable and semantically meaningful. This involves two key steps: tokenization, which breaks a sentence into smaller units (tokens), and embedding, which maps each token into a continuous vector space. The result is a matrix in $\mathbb{R}^{d_s \times d_m}$, where d_s is the maximum sequence length and d_m is the embedding dimension. Figure 2 illustrates this process.

Specifically, tokenization divides a sentence into smaller linguistic units, called tokens, using a pretrained tokenizer from the spaCy library³. These tokens typically represent words, punctuation marks, or other meaningful elements. For example, the sentence "The Bank of Korea announced changes to its monetary policy." is tokenized as ["The", "Bank", "of", "Korea", "announced", "changes", "to", "its", "monetary", "policy", "."]. This step structures the input for further processing while preserving syntactic and semantic content. Consequently, each sentence is transformed into a sequence of tokens that can be systematically analyzed and embedded into a numerical representation.

Preprocessing

Figure 2.



Source: Author's illustration

³ spaCy is a general-purpose natural language processing library with multilingual support. See Honnibal and Montani (2017) and <https://spacy.io> and

Following tokenization, each token is mapped to a vector in a continuous, real-valued space using the Word2Vec algorithm⁴. This neural embedding method captures semantic and syntactic relationships by training the model to predict the surrounding words, given a target word (skip-gram) or vice versa (CBOW). In the resulting semantic space, similar words are located close to one another—for example, "Bank of Korea" and "Fed" appear nearby due to their shared institutional role. Semantic proximity is typically measured using the cosine similarity, which quantifies the contextual resemblance between word vectors.

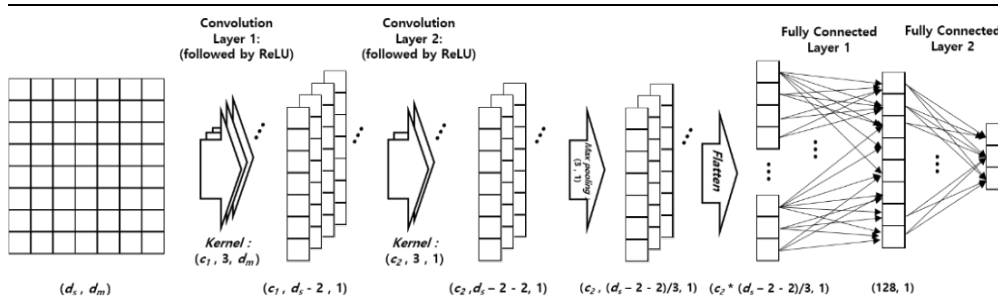
Given a maximum token sequence length d_s and embedding dimension d_m , this step yields a $d_s \times d_m$ real-valued matrix, where each row corresponds to the embedding vector of a token. If the sequence contains fewer than d_s tokens, zero-padding is applied to ensure consistent dimensionality across the inputs. For instance, the sequence ["Bank," "of," "Korea"] with $d_s = 6$ is padded to ["Bank," "of," "Korea," 0, 0, 0], with padding tokens mapped to zero vectors that contribute no semantic information.

Once the matrix representation of a sentence has been prepared, the next step is to construct a classification model $g: E \rightarrow L$ that maps it to a sentiment label. A Convolutional Neural Network (CNN) is employed for this purpose. A CNN interprets the matrix as a grayscale image, where d_s is the sequence length and d_m is the embedding dimension⁵. Similar to how CNNs identify features in images, the network applies filters to the sentence matrix to detect semantic structures that are useful for classifying sentiments.

The CNN used in this study comprises two convolutional layers that extract hierarchical features from the input. The first layer applies c_1 filters to capture localized semantic cues across token sequences, whereas the second layer builds higher-level abstractions from these features with c_2 number of channels. The max-pooling layer follows, selecting the most salient features and reducing the sequence length. The resulting feature maps are flattened and passed through two fully connected layers. The final output is a three-dimensional probability vector over the sentiment classes computed using the softmax function. Figure 3 shows the model architecture.

CNN Classifier

Figure 3.



Source: Author's illustration

Note: d_s , d_m , c_1 , and c_2 are set to 100, 128, 16, and 8, respectively.

⁴ For the details of the algorithm, see Mikolov et al. (2013).

⁵ For a comprehensive review of the CNN, see Li et al. (2021).

2.2 Trending Keyword Extraction

Trending keywords are defined as sets of terms that capture emerging or salient economic issues during the target period. In other words, they were the keywords that were most likely to come to mind when a specific period was recalled. To identify such keywords, I approach the problem from a dual perspective rather than solving the primal problem directly. Specifically, I ask which period is most likely to come to mind when a given keyword is recalled. Based on this perspective, a trending score is computed for each keyword, rooted in the Neyman–Pearson lemma, and used for keyword extraction.

For example, consider the following hypothetical scenario. Suppose you are presented with a set of keywords randomly selected from news articles originating at an unknown point in time. Based on these observed keywords, you are asked to determine whether the articles belong to a specific target period or base (normal) period. Under the formal hypothesis-testing framework, this task is equivalent to testing the null hypothesis that the keywords originate from the base period against the alternative that they originate from the target period.

The Neyman–Pearson lemma states that under certain conditions, the likelihood ratio test is uniformly most powerful for distinguishing between two simple hypotheses. Suppose the true probabilities under the null and alternative hypotheses are given as $\Pr[i | \text{base period}]$ and $\Pr[i | \text{target period}]$, respectively. In the aforementioned keyword-based hypothesis test context, the informativeness of observing keyword i to discern the null and alternative hypotheses is measured by the following score S_i :

$$S_i = \Pr[i | \text{target period}] \times \ln \left(\frac{\Pr[i | \text{target period}]}{\Pr[i | \text{base period}]} \right).$$

To assess the validity of the trending keyword extraction method, I compare the yearly trending keywords with the evolution of economic sentiment implied in news articles based on the sentiment classification results from the previous section. Specifically, a sentiment classifier is applied using 10,000 randomly sampled article sentences per day. Based on the number of positive and negative sentences obtained from the sentiment classification described above, the following index is constructed:

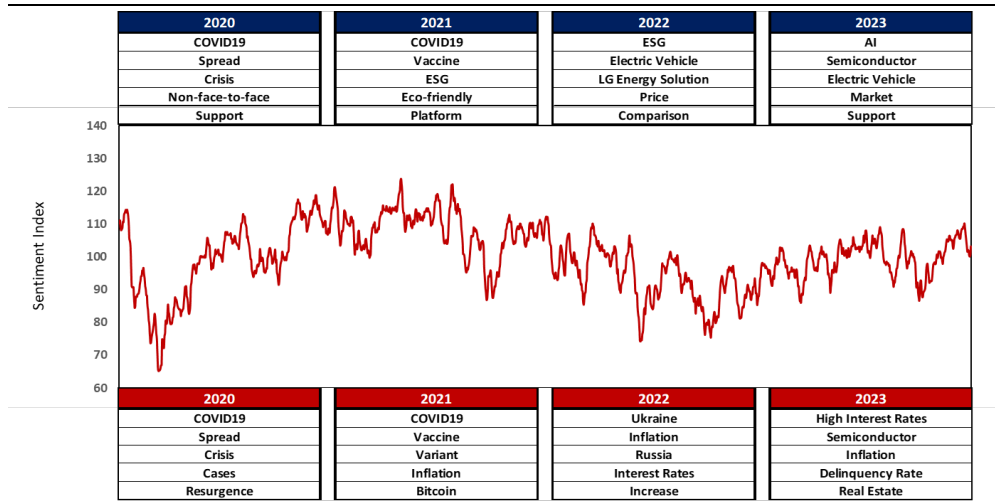
$$SI_t = \frac{\# \text{ of positive sentences} - \# \text{ of negative sentences}}{\# \text{ of positive sentences} + \# \text{ of negative sentences}}.$$

The sentiment index SI_t is normalized and standardized such that its mean and standard deviation are equal to 100 and 10, respectively.

Figure 4 illustrates the evolution of the sentiment index between 2020 and 2023, including the period of the COVID-19 outbreak and its aftermath (solid red line), alongside the corresponding yearly trending keywords extracted from positive (blue headings) and negative sentences (red headings). The figure conveys, in a single comprehensive view, both the trajectory of economic sentiment over time and the key issues underlying these shifts, offering a simultaneous understanding of how and why sentiment has evolved.

Trending Keywords

Figure 4.



Source: Author's illustration based on implementation results.

In 2020, the coronavirus disease (COVID-19) outbreak emerged as the most dominant economic and social issue in Korea. After the first confirmed case on January 20, sentiment began to decline sharply. This downward trend accelerated following the World Health Organization's declaration of the pandemic, reflecting heightened public anxiety and growing economic uncertainty.

The following year, sentiments recovered to some extent, supported by strong government intervention and growing optimism surrounding vaccine development. Although concerns such as the emergence of new variants and public debates over vaccination persisted, the tone of public discourse began to shift. New opportunities, particularly those related to ESG (Environmental, Social, and Governance) themes, have gained traction and were increasingly discussed in a positive light.

At the beginning of 2022, Russia's invasion of Ukraine had a profound and far-reaching impact on the global economy, including that of Korea. The resulting geopolitical instability, coupled with excess liquidity in financial markets, has led to a surge in inflation. In response, central banks worldwide, including the Bank of Korea, tightened their monetary policies through a series of interest rate hikes. These developments remained central throughout 2022 and continued to influence economic sentiment well into 2023.

Interest in ESG themes began to materialize more concretely, particularly through rising expectations for the electric vehicle and battery industries. Additionally, the launch of ChatGPT in early 2023 brought significant public attention to AI and semiconductor technologies, contributing to a rebound in positive market sentiment. Figure 4 aligns closely with the narrative described above, illustrating that the proposed method can serve as an effective tool for monitoring economic conditions in a structured and data-driven manner.

2.3 Summarization

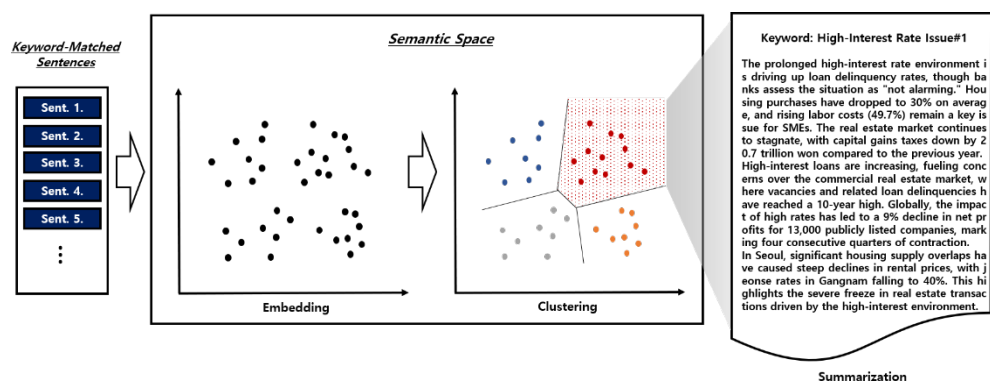
Although trending keywords help reveal the economic issues driving macroeconomic conditions, they often fail to provide a comprehensive perspective. Some keywords can oversimplify or even mislead users because they fail to capture the full range of implications. For instance, Figure 4 shows that the keyword “vaccine” was prominent in 2021 in both positive and negative contexts. However, the keyword alone does not clarify whether the discussion focuses on breakthroughs in vaccine development or concerns over side effects and public resistance.

The goal of the summarization step is to produce outputs that are succinct yet intellectually accessible while preserving the multidimensional perspectives and sentiments surrounding the trending keywords. It consists of three key components: embedding, clustering, and summarizing. First, each sentence containing a given keyword is embedded into a semantic space. In the system implemented in this study, the Word2Vec algorithm is employed for this task. Second, the embedded sentences are clustered using the K-means algorithm⁶. Because the semantic space preserves semantic similarity, clustering based on vector distance allows each cluster to be interpreted as distinct. Third, each cluster representing a specific issue is summarized using a Large Language Model (LLM). In this implementation, a version of LLaMA 3.2⁷ is used to generate these summaries.

The output of this process is presented in Figure 5, which illustrates how the system operates and the results produces. A set of positive sentences containing the keyword “high-interest rate” was processed using the aforementioned procedure. The system organizes these sentences by subject matter into issue clusters, and each cluster is then summarized into a concise description of the core issue. In experiments with various keywords across different periods, the method consistently succeeded in summarizing the underlying issues within several minutes on a personal laptop.

Trending Keywords

Figure 5.



Source: Author's illustration based on implementation results.

⁶ For the implementation, FAISS is used. For the details, see Johnson and Jégou (2017).

⁷ Specifically, llama-3.2-Korean-Blossom-3B is used. See Choi et. al (2024).

3. Concluding Remarks

Motivated by the desire to monitor economically significant issues in a timely manner, this study presents a method for extracting and summarizing such information from a large volume of news articles. Owing to recent advancements and the rapid commoditization of natural language processing (NLP) technology, the proposed method is easy to implement on a personal laptop without requiring significant investment in specialized hardware or expensive software. The ease of implementation and cost-efficiency of the proposed method lowers barriers for many central bankers and economists, helping alleviate their reluctance to expand their monitoring scope beyond traditional structured data. Despite its simplicity, this approach is versatile and can be adapted to various policy needs. For example, in environments where specific economic issues, such as bank stability, become prominent, this method can be used to monitor related issues in greater depth and in a timely manner. Additionally, with slight modifications, it can be employed for comprehensive monitoring of central bank communications, allowing policymakers to adjust the content or tone of their messaging by closely tracking how the media perceives and frames their communications.

Despite its contributions, this study had several limitations. Fine-tuning and data scalability are notable examples; however, other important challenges remain. Fine-tuning is required to tailor the system to specific use cases. The experimental results were generally robust, but the quality of the summarization was influenced by both the number of clusters and sentences included. Therefore, careful calibration of cluster counts relative to sentence volume is necessary. A more fundamental issue is whether the data scope can be expanded beyond news articles to include social media data. Social media can serve as a valuable complementary source for identifying public attention to economic issues owing to its immediacy and scale. However, unlike news articles written in formal or edited language, social media content tends to be noisy and linguistically inconsistent. To effectively leverage such diverse data sources, future research must develop methods capable of handling this linguistic variability and noise.

References

- Choi, C., Jeong, Y., Park, S., Won, I., Lim, H., Kim, S., Kang, Y., et al. (2024). Optimizing language augmentation for multilingual large language models: A case study on Korean. *arXiv Preprint, arXiv:2403.10882*. <https://doi.org/10.48550/arXiv.2403.10882>
- Honnibal, M., & Montani, I. (2017). spaCy 2: Natural language understanding with Bloom embeddings, convolutional neural networks and incremental parsing. <https://spacy.io>
- Johnson, J., Douze, M., & Jégou, H. (2019). Billion-scale similarity search with GPUs. *IEEE Transactions on Big Data*, 7(3), 535–547.
- Li, Z., Liu, F., Yang, W., Peng, S., & Zhou, J. (2021). A survey of convolutional neural networks: analysis, applications, and prospects. *IEEE transactions on neural networks and learning systems*, 33(12), 6999–7019.
- Mikolov, T., Chen, K., Corrado, G., & Dean, J. (2013). Efficient estimation of word representations in vector space. *arXiv preprint arXiv:1301.3781*.

Extracting Economic Issues from News Data: With the Help of Generative AI

Younghwan Lee

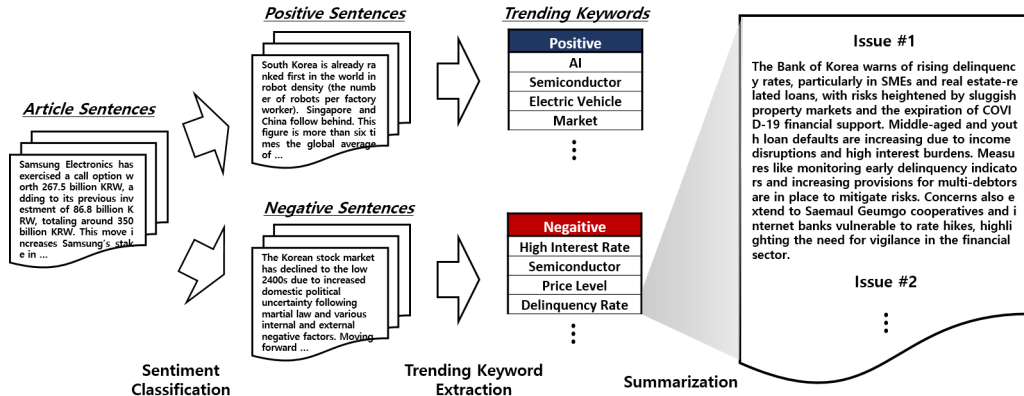
Bank of Korea

Disclaimer: The views expressed in this talk do not reflect necessarily those of Bank of Korea

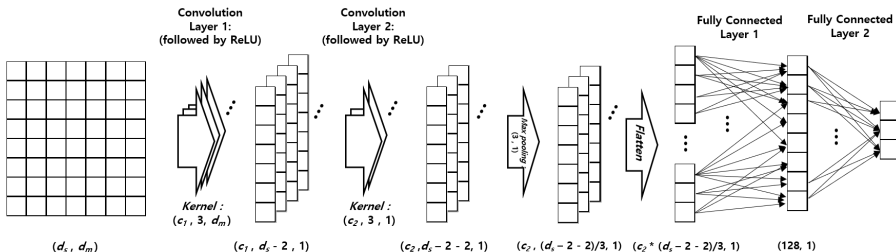
Motivation

- Traditional economic monitoring frameworks that rely on economic variables such as GDP and CSI might only be a second-best method.
 - Doctors treat diseases, not symptoms.
 - Similarly, policymakers must address economic issues, not just economic variables.
- However, directly monitoring economic issues is challenging in practice, as they are often linguistic rather than numeric.
- In this study, I present a methodology to incorporate economic issues into the monitoring framework that is:
 - i) Verifiable,
 - ii) Flexible,
 - iii) **Easy to implement.**

Snapshot

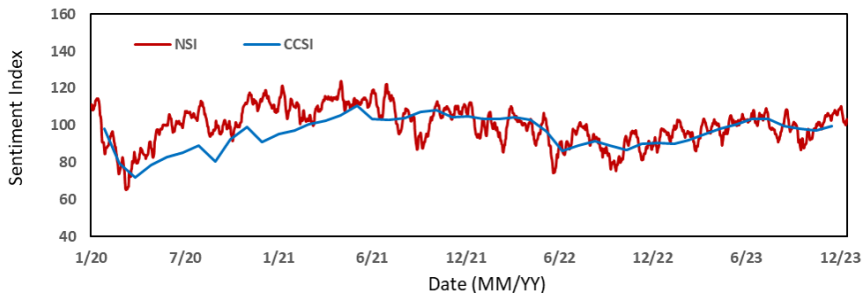


Classification: Methodology



- CNN (Convolutional Neural Network) is applied to classify sentiment into positive, neutral, and negative categories.
 - i) Tokens are generated using **spaCy** and embedded into 128-dimensional vectors via **Word2Vec**.
 - ii) Token vectors are concatenated into a 100×128 matrix, with zero-padding added if needed for consistent dimensions.

Classification: Result



- Sample 10,000 sentences daily to derive an index. The results show it is highly correlated with existing sentiment indices.

$$X_t = \frac{\# \text{ Positive Sentences} - \# \text{ Negative Sentences}}{\# \text{ Positive Sentences} + \# \text{ Negative Sentences}}$$

- Scale and translate X_t to resemble a standardized index (Mean = 100, Std. = 10).

Trending Keyword: Idea

- It is inefficient to directly work with a large set of sentences. We need a method to focus on a few important issues that characterize current economic conditions.
- **Trending Keywords** are defined as keywords extracted from news headlines during a specific time period that best describe the key issues of that period.
- Then, how can we capture the **Trending Keywords**?
- Suppose that you are provided with a collection of news headlines, and you are asked to guess when those headlines were published based on the keywords in the headlines.
- I define **Trending Keywords** as a set of keywords that are most useful for us to address the question above.

Trending Keywords: Methodology

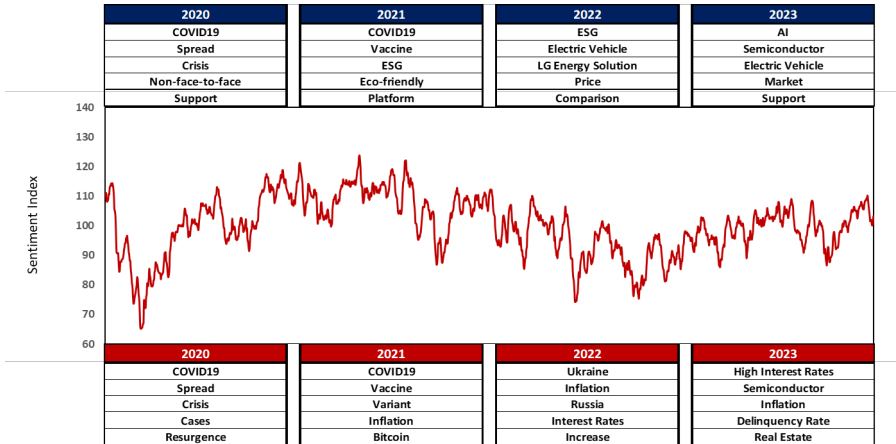
- The Neyman-Pearson lemma states that the log-likelihood test is the most powerful test. By utilizing the empirical likelihood ratio of keyword frequency, we can address the question at hand.
- For each keyword $i \in I$ define trending score S_i as follow:

$$S_i = \Pr[i|\text{target period}] \ln \frac{\Pr[i|\text{target period}]}{\Pr[i|\text{base period}]}.$$

- S_i increase when the keyword i
 - i) is encountered more frequently in the headlines during the target period,
 - ii) is relatively uncommon in the base period.
- We can extract **Trending Keywords** for the target period by ranking keywords according to their trending score.

Trending Keywords: Result

- **Trending Keywords** from 2020 summarize the captures important issues after COVID-19 outbreak.

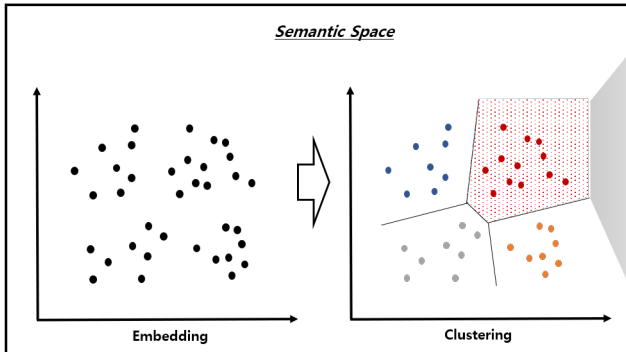


Summarization: Methodology

- The goal of summarization is to extract information comprehensively and deliver it succinctly.
 - Working with keywords is efficient but often lacks clarity and sufficient context.
 - Reviewing all keyword-matched sentences is inefficient as many convey similar meanings.
- Strategy:
 - Embed**: Transform keyword-matched sentences into semantic vectors (**Word2Vec**).
 - Cluster**: Group sentences by semantic similarity using K-means (via **FAISS**).
 - Summarize**: Generate concise summaries for each cluster with a pretrained LLM (**Llama 3.2**).

Summarization: Result

Keyword-Matched
Sentences



Keyword: High-Interest Rate Issue#1

The prolonged high-interest rate environment is driving up loan delinquency rates, though banks assess the situation as "not alarming." Housing purchases have dropped to 30% on average, and rising labor costs (49.7%) remain a key issue for SMEs. The real estate market continues to stagnate, with capital gains taxes down by 20.7 trillion won compared to the previous year. High-interest loans are increasing, fueling concerns over the commercial real estate market, where vacancies and related loan delinquencies have reached a 10-year high. Globally, the impact of high rates has led to a 9% decline in net profits for 13,000 publicly listed companies, marking four consecutive quarters of contraction. In Seoul, significant housing supply overlaps have caused steep declines in rental prices, with mortgage rates in Gangnam falling to 40%. This highlights the severe freeze in real estate transactions driven by the high-interest environment.

Summarization

Concluding Remarks

- This study proposes a verifiable, flexible, and easy-to-implement methodology to extract economic issues from a large volume of news articles.
 - With the help of generative AI and pretrained LLMs, the unstructured nature of economic issues is effectively addressed.
- Future Research:
 - Exploring how to predict next quarter's trending keywords and summarized issues.