

THE PENNSYLVANIA STATE UNIVERSITY
SCHREYER HONORS COLLEGE

Predicting the U.S. Unemployment Rate Using Natural Language Processing Models

SANCHITA BHUSARI
SPRING 2025

A thesis
submitted in partial fulfillment
of the requirements
for baccalaureate degrees
in Finance & Statistical Modeling Data Sciences
with honors in Finance

Reviewed and approved* by the following:

Christoph Hinkelmann
Professor of Finance
Thesis Supervisor

Brian Davis
Professor of Finance
Thesis Honors Adviser

* Electronic approvals are on file.

ABSTRACT

This research project investigates the predictability of the unemployment rate during economically uncertain periods through the analysis of Federal Open Market Committee (FOMC) meeting minutes, using Natural Language Processing (NLP) techniques such as FinBERT and Principal Component Regression (PCR) modeling. Traditionally, the relationship between inflation and unemployment follows an inverse pattern, where shifts in employment levels influence consumer strength and pricing power. However, recent economic conditions—shaped by post-pandemic inflationary pressures, geopolitical events, and monetary policy interventions—have complicated this relationship, challenging existing predictive models. This study aims to bridge the predictive gap by developing a model that translates the sentiment and tone of FOMC communications into unemployment rate forecasts. The methodology proposed involves text scraping of Federal Reserve meeting minutes from key periods of economic uncertainty (2003-2022), applying sentiment analysis to assign sentiment scores, and conducting regression analysis to assess the sentiment-unemployment relationship.

By training a PCR model on historical unemployment data, this study seeks to create a nowcasting tool to forecast short-term unemployment trends. This project ultimately aims to identify how Federal Reserve communications reflect and potentially forecast labor market responses, offering a vital resource for anticipating economic shifts in real time.

TABLE OF CONTENTS

LIST OF FIGURES	iv
LIST OF TABLES	v
Chapter 1 Introduction	1
Chapter 2 Related Work.....	5
Previous Unemployment Models.....	5
Nowcasting via Social Media.....	7
FOMC Statement Usage in Text Analysis Tools.....	9
Neural Network Models.....	10
LLaMA-3 Model.....	12
FinBERT Model.....	13
Chapter 3 Data	16
Chapter 4 Methodology	19
Model Comparison.....	19
Sentiment Analysis.....	20
Data Integration.....	22
Regression Analysis.....	22
Evaluation Metrics.....	25
Chapter 5 Results & Evaluation.....	26
Model Performance Comparison	26
Analysis of Results.....	31
Economic Predication Application.....	32
Chapter 6 Discussion & Future Work.....	36
Code Appendix	40
Dataset Appendix.....	42
BIBLIOGRAPHY.....	43

LIST OF FIGURES

Figure 1. U.S. Unemployment Rate (2003-2022).....	17
Figure 2. Histogram of Positiveness Sentiment Score.....	21
Figure 3. Correlation Graph of Linear Regression Model.....	27
Figure 4. Correlation Graph of Random Forest Regression Model.....	28
Figure 5. Correlation Graph of KNN Regression Model.....	29
Figure 6. Correlation Graph of Gradient Boosting Regression Model.....	29
Figure 7. Correlation Graph of Light Gradient Boosting Regression Model.....	30
Figure 8. Correlation Graph of Decision Tree Regression Model.....	30
Figure 9. Correlation Graph of XGBoost Regression Model.....	31
Figure 10. Actual vs. Predicted Unemployment Rate through XGBoost Regression Model.....	33

LIST OF TABLES

Table 1. Success Metrics Table.....	27
Table 2. Unemployment Forecasting Probability Table	34

ACKNOWLEDGEMENTS

I would like to thank my Thesis Advisor and Thesis Supervisor, Dr. Davis and Dr. Hinkelmann, for their continued support and guidance throughout my research process. From the earliest stages of project ideation to the final stages of writing, their guidance has been integral in shaping this work. I am incredibly grateful for their willingness to consistently make time to meet with me, whether it was to brainstorm ideas, provide insightful feedback, or simply help me get a better grasp on the implementation aspect of my project. Additionally, I would like to thank all of the Schreyer team for their continuous support and willingness to answer any questions throughout this thesis process as well.

Chapter 1

Introduction

Predicting the unemployment rate during periods of economic uncertainty is a challenging problem with significant implications for policymakers and businesses. Understanding how the unemployment rate responds to shifts in economic conditions, particularly in response to the Federal Reserve's actions, is crucial for decision-making. This is especially important in periods of inflationary pressure or when the economy is undergoing complex disruptions. Accurate predictions of unemployment trends possibly guide more effective monetary policies and better economic planning. This project looks to explore more about the relationship between the unemployment rate and the Federal Reserve's actions, such as rate hikes/cuts through quantitative tightening/easing. It has been established that usually inflation and unemployment have an inverse relationship. To simplify this relationship, this means that as the unemployment rate goes down, the consumer will become stronger (given more consumers have jobs) and this will result in inflation rising as sellers receive the pricing power to charge more. On the flip side, if unemployment is rising this means that the consumers are weak and are unable to afford to buy more expensive items, so a seller must reduce prices to entice consumers to buy from them (giving consumers pricing power).

As touched upon above, current research on unemployment and inflation typically explores the inverse relationship between these two economic indicators. It is well-established that as unemployment falls, consumer strength rises, often leading to higher inflation as demand increases and sellers gain pricing power. Conversely, rising unemployment tends to reduce consumer purchasing power, which puts downward pressure on inflation. However, existing

works generally rely on historical economic relationships and do not account for more complex factors influencing the economy.

Post-pandemic this entire relationship became far more complex as inflation reached record highs given the housing shortage, high energy costs due to the Russia/Ukraine crisis, and increased federal spending. The Federal Reserve raised interest rates throughout 2022 and into the first half of 2023, yet there was not a linear relationship between the unemployment rate and inflation, as there was not an immediate change in unemployment after the Federal Reserve's actions. The debate that continues to occur regarding this topic is how can the Federal Reserve continue to combat inflation while not completely tanking the jobs market in the process, and is a soft landing actually achievable? Throughout the past almost two years, there have been times where the unemployment rate's reactions to the Federal Reserve's actions were not quite explainable. For example, despite the interest rates maintaining in the 5.00-5.25% range from the beginning of 2024 to September 2024, the unemployment rate was still ticking up and the jobs report results were not easily predictable due to this. Hence, another question that remains unsettled is how exactly can one predict how the unemployment rate will react to the Federal Reserve's actions?

This leads to the contribution of this project; building a model to predict the unemployment rate based on FOMC meeting minutes. There is currently no way to predict the unemployment rate from Federal Reserve communications (besides a reliance on the pre-established relationship between unemployment and inflation), and given the uncertain economic environment, many feel uninformed about the job market. Additionally, a tool like this would have been useful throughout the past two years when there was not a direct relationship between

the Federal Reserve's actions and the job market's response, as there were so many different factors at play.

The topic of this thesis surrounds predicting the unemployment rate during uncertain economic periods using Natural Language Processors (NLPs) such as a Large Language Model for extracting information and conducting Sentiment Analysis (FinBERT) and various regression models. The FinBERT model is well-suited for financial texts, and by extracting sentiment scores from FOMC meeting minutes, it will provide valuable insights into the Fed's stance and its likely effects on the economy. The sentiment score will be compared to historical unemployment data, and then various regression models will be implemented to assess the relationship between FOMC sentiment and the unemployment rate two months later (to allow for time for the meeting to influence the market but before the next FOMC meeting). This approach is expected to help provide more timely and informed predictions by directly linking the sentiment of Federal Reserve communications to economic outcomes.

To evaluate this approach, there will be web scraping of Federal Reserve meeting minutes from 2003 to 2022, a period marked by notable economic uncertainty. By applying FinBERT to analyze sentiment, there will be a comparison of sentiment scores with historical unemployment data. Success metrics such as accuracy, R^2 score, and p-values will be calculated to assess the model's power in terms of establishing a relationship between the FOMC sentiment scores and the unemployment rates. Additionally, seven different regression models will be employed to assess which method is the best to understand this relationship and compare success metrics accordingly.

Chapter 2

Related Work

The prediction of the U.S. unemployment rate continues to grow in importance as the economy continues to navigate uncertain periods. Accurate unemployment predictions are crucial for economic planning, policy development, and workforce management. Researchers have explored various approaches to predict the unemployment rate, including text analysis tools that leverage data from social media platforms and publicly available websites. More recently, studies have investigated the use of neural networks and Large Language Models as predictive tools for this purpose. This chapter will discuss the previous studies that contributed to the research within this paper and provide further context regarding the problem at hand. These studies cover previous unemployment models explored by economists, nowcasting techniques through social media, sentiment analysis methods, applications of Llama-3 models, neural network models, and FinBERT (the model employed in this research).

Sub-Chapter 1: Previous Unemployment Models

Understanding previous models that economists have used to predict and analyze unemployment is crucial for developing improved prediction methods. Bande et al. (2010) critique the traditional Natural Rate of Unemployment (NRU) hypothesis and introduce Chain Reaction Theory (CRT) as an alternative framework for explaining unemployment dynamics [1]. The NRU model assumes that unemployment will stabilize at a fixed equilibrium rate in the long run, regardless of monetary policy or inflation trends. The authors argue that this view oversimplifies the relationship between inflation and unemployment, overlooking the influence

of frictional growth — the culmination of persistent effects of nominal frictions such as wage rigidity and delayed adjustments.

To support this argument, the authors propose Chain Reaction Theory (CRT), which suggests that unemployment and inflation remain interconnected due to lagged adjustments and economic shocks. CRT introduces the notion that economic shocks, such as monetary policy changes, can initiate chain reactions that sustain deviations in unemployment from its perceived "natural rate" for prolonged periods.

The authors provide well-rounded evidence from multiple economies, including Spain, the EU, and the United States and the findings reveal that unemployment rates do not consistently stabilize at fixed NRU levels, and instead, exhibit prolonged deviations caused by inflation changes and wage dynamics. The authors also emphasize the policy implications of the findings. They caution policymakers against relying solely on NRU-based models when designing monetary and fiscal policies. Instead, they advocate for policy frameworks that account for delayed economic effects and recognize the sustained relationships between inflation, wage dynamics, and unemployment.

This study is especially relevant to forecasting models for unemployment prediction as it highlights the need for approaches to capture long-term dependencies and dynamic interactions. For the research at hand, it provides a basis of understanding of prediction models that economists have previously used, which assists with developing a strong context surrounding what factors influence the unemployment rate and what textual content to utilize for sentiment analysis.

Sub-Chapter 2: Nowcasting via Social Media

Bukenya et al. (2022) explore a unique method for predicting the unemployment rate during the COVID-19 pandemic by leveraging social media data, specifically Twitter, as a primary information source [3]. The purpose of the study was to address the gap in unemployment data availability in South Africa when traditional census data collection was disrupted during the pandemic. To achieve this, the researchers collected relevant data using a Twitter API, conducted sentiment analysis on employment-related tweets, and employed Principal Component Regression (PCR) to estimate the unemployment rate. The researchers found that public sentiment derived from tweets had a strong negative correlation with the unemployment rate — as unemployment increased, tweet sentiment became more negative. The PCR model achieved a R^2 -score of 0.93, demonstrating strong predictive power. This study highlights the potential of social media data as a timely and effective tool for tracking labor market trends, particularly during uncertain economic periods. While this research effectively utilized sentiment analysis and PCR modeling, its methodology was applied specifically to South Africa's unemployment trends using Twitter data. This study similarly utilizes text analysis and sentiment scoring to predict economic indicators. However, this research differentiates itself by applying these methods to Federal Open Market Committee (FOMC) communications rather than social media data, focusing on how the sentiment within these official communications can nowcast the U.S. unemployment rate [12].

Continuing with research exploring the use of social media data for nowcasting unemployment rates, Blanchflower and Bryson (2021) propose an alternative method for predicting unemployment trends through qualitative social survey data in the "Economics of

Walking About" (EWA) framework [2]. This approach leverages consumer and employer sentiment data to forecast economic downturns and changes in unemployment rates. Unlike traditional economic models that emphasize structural shocks as unpredictable events, the EWA model emphasizes the predictive power of real-time social perceptions as an early warning system for labor market shifts.

The authors utilize a comprehensive dataset covering 29 European countries over a 439-month period (1985–2021), integrating over 10,000 observations from the Joint EU Harmonized Programme of Business and Consumer Surveys [2]. The dataset captures consumer perceptions about unemployment risk, economic stability, and household finances. The study's findings demonstrate that individuals' fear of unemployment consistently predicts changes in unemployment rates up to 12 months in advance. This predictive power remained strong even after controlling for country-specific effects and lagged unemployment data. Similarly, business sentiment, particularly employers' outlook on hiring conditions, was also found to be a reliable predictor of labor market trends. In addition, the researchers found that individuals' respective assessments of a country's economic stability and personal financial situations served as strong indicators of future unemployment trends.

The study highlights that qualitative social data successfully signaled economic downturns during major crises, including the Great Recession and the economic slowdown leading up to the COVID-19 pandemic. In addition to its predictive success, the EWA framework underscores the importance of capturing collective economic sentiment. The authors argue that these "on-the-ground" insights provide timely signals for policymakers, who may otherwise rely too heavily on delayed economic indicators. The paper advocates for integrating

qualitative sentiment data into mainstream forecasting models to improve early detection of economic downturns. This finding highlights the importance of this research paper, by touching on the necessity of further exploration surrounding sentiment analysis for economic forecasting.

Sub-Chapter 3: FOMC Statement Usage in Text Analysis Tools

Shapiro and Wilson (2022) investigate central bank objectives by applying sentiment analysis to Federal Open Market Committee (FOMC) meeting transcripts [18]. This research seeks to estimate the Federal Reserve's loss function to determine what economic variables the FOMC prioritizes, challenging the traditional assumption that central banks maintain quadratic preferences over inflation and economic slack.

To achieve this, the researchers employed a sentiment analysis approach using the Loughran-McDonald (LM) financial-specific dictionary to assess negativity in FOMC transcripts from 1976 to 2015. The negativity score was calculated by subtracting the fraction of positive words from the fraction of negative words, ensuring the measure was independent of overall emotivity. This automated method was validated through human evaluations, demonstrating strong reliability.

In addition to its findings, this study also provides valuable insights into effectively working with FOMC transcripts in a text analysis tool, demonstrating best practices for data collection, filtering relevant remarks, and selecting appropriate language models for financial text.

The study's findings revealed that the FOMC's implicit inflation target was approximately 1.25%—below the officially announced 2% target established in 2012. Furthermore, the analysis showed that output growth and financial conditions played a greater

role in influencing FOMC sentiment than economic slack. The authors concluded that the FOMC's focus on output growth and financial stability, alongside its implicit inflation target, may require policymakers and economists to reconsider traditional assumptions about central bank objectives.

This study's application of text analysis methods demonstrates how sentiment analysis can be a powerful tool in examining economic decision-making. This study successfully implemented FOMC communications as an input within text analysis tools to generate viable conclusions, providing a template for the research at hand [16]. While this paper discovered overarching central bank objectives, it did not offer a quantitative predictive tool for an economic metric, which is this paper's purpose.

Sub-Chapter 4: Neural Network Models

Kaouk and Alvares (2022) examine the implementation of advanced Natural Language Processing (NLP) techniques to enhance comment analysis in the federal rulemaking process [4][11]. This study aimed to improve the efficiency of public comment analysis, a process that is typically resource-intensive and time-consuming. The researchers collaborated with organizations such as the EPA and FCC to develop standardized NLP tools designed to streamline comment analysis across various government agencies. The tools developed in this project were capable of categorizing comments by topic, identifying duplicates, and grouping semantically similar comments for more efficient review. To support this, neural network models were trained on extensive text data, providing a versatile solution that could be customized for different agencies.

The pilot project successfully demonstrated the time-saving potential of these tools, reducing comment processing time by approximately 4.5 hours per 100 comments through deduplication and similarity tools and saving an additional 8 hours per 100 comments with the hierarchical latent Dirichlet allocation (hLDA) topic modeling tool. The study concluded that these NLP-based solutions provided agencies with improved insights and faster response times to public comments, while also reducing development costs by minimizing duplicated tool creation. While this research demonstrated considerable success in enhancing comment analysis for federal agencies, its focus remained on public commentary rather than financial or economic forecasting. Nevertheless, this work is significant as it highlights the versatility of NLP tools in text analysis applications and assisted when weighing choices for the methodology of this paper.

Building on the growing exploration of neural network models for unemployment prediction, Mero et al. (2024) propose a novel method for predicting the unemployment rate by integrating Long Short-Term Memory (LSTM) networks with Genetic Algorithms (GA) [14]. The primary objective of the study was to enhance the accuracy of unemployment rate predictions, particularly in Ecuador, by combining the sequential data-handling capabilities of LSTM networks with the optimization efficiency of GAs. The hybrid model, referred to as GA-LSTM, aimed to address challenges in parameter determination, which is crucial for improving recurrent neural network performance.

The methodology involved utilizing the GA heuristic search method to optimize key LSTM parameters, such as the time window size and the number of LSTM units. Data preprocessing steps included data normalization and outlier treatment to ensure data quality. The study leveraged economic indicators such as inflation, GDP, gross fixed capital formation

(GFCF), and the national minimum wage as predictive features, with unemployment rates as the target variable.

The results demonstrated that the GA-LSTM model outperformed traditional models such as BiLSTM and GRU in various performance metrics, including Mean Squared Error (MSE), Mean Absolute Error (MAE), and Mean Absolute Percentage Error (MAPE). The hybrid model's ability to optimize hyperparameters effectively contributed to its predictive accuracy, making GA-LSTM an effective tool for addressing the complexities of unemployment forecasting in dynamic economic environments. This research highlights the potential of combining deep learning techniques with evolutionary algorithms to improve economic prediction models and assisted with the choice of evaluation metrics for this paper's study.

Sub-Chapter 5: LLaMA-3 Model

Mai et al. (2024) explore the application of LLaMA-3, an up-and-coming large language model, for financial sentiment analysis [13]. The study investigates LLaMA-3's ability to extract sentiment from financial text data and compares its performance with baseline models. The researchers sought to evaluate whether LLaMA-3's extensive pre-training enables it to effectively interpret finance-specific language and outperform traditional sentiment analysis methods.

The methodology employed by the researchers involved fine-tuning LLaMA-3 on financial text data and evaluating its performance using precision, recall, and F1 scores. The results demonstrated that LLaMA-3 significantly outperformed the majority and random baselines, achieving an F1 score of 0.67 overall. The model demonstrated notable strengths in identifying neutral sentiment, showing consistent performance across various input lengths and

text structures. However, the study also identified that LLaMA-3's accuracy declined when analyzing longer financial texts, suggesting limitations in its ability to capture nuanced sentiment shifts in complex content.

This research contributes valuable insights into the practical applications of large language models for financial text analysis. The study's findings align with broader research highlighting the strengths of transformer-based models in extracting sentiment from financial narratives. Furthermore, this study offers useful strategies for leveraging LLaMA-3 in text analysis tools, particularly in finance-based contexts where sentiment detection is essential. While FinBERT has been established as a strong performer in financial sentiment analysis, this paper highlights LLaMA-3's potential to complement existing tools by enhancing performance on shorter financial texts and offering improved versatility for diverse NLP applications. This paper's findings were very valuable in approaching the methodology design and model choice for the project at hand and served as one of the model options.

Sub-Chapter 6: FinBERT Model

Huang et al. (2022) introduce FinBERT, a language model specifically designed for analyzing financial text data [9]. FinBERT is adapted from Google's BERT model and trained on a large corpus of financial texts, including corporate filings, analyst reports, and earnings call transcripts, to improve its understanding of financial language. The primary objective of the study was to enhance sentiment classification accuracy and improve information extraction from financial documents, addressing limitations seen in traditional models like the Loughran-McDonald dictionary and other machine learning algorithms.

To develop FinBERT, the authors followed the BERT model's pretraining and fine-tuning procedure, applying financial-specific texts to enhance model performance. The researchers compared FinBERT's performance against various NLP techniques, including support vector machines (SVM), random forest (RF), convolutional neural networks (CNN), and long short-term memory (LSTM). FinBERT achieved an overall sentiment classification accuracy of 88.2%, outperforming other models such as LSTM (76.3%), CNN (75.1%), and even the standard BERT model (85.0%). Notably, FinBERT exhibited superior performance in identifying negative sentiment, achieving 89.7% accuracy in this category, which was considerably higher than its competitors.

The research further demonstrated that FinBERT maintains high accuracy even when trained on smaller datasets, showcasing its efficiency in data-scarce scenarios. This characteristic highlights FinBERT's adaptability and effectiveness in real-world applications where labeled data may be limited. Additionally, the study explored FinBERT's role in environmental, social, and governance (ESG) discussions, achieving an accuracy rate of 89.5% in classifying ESG-related text.

This study is significant as it emphasizes the value of domain-specific training for improving NLP performance in specialized fields such as finance. FinBERT's design and methodology resulted in high interpretability and accuracy due to its expertise regarding financial literacy. The outperformance of FinBERT in comparison to other LSTM and Neural Network models assisted when crafting the methodology for the research at hand and choosing an NLP model to utilize.

By building on these related works and addressing the gaps in the existing literature, this project aims to provide a more accurate and timelier tool for predicting unemployment trends in response to Federal Reserve sentiment.

Chapter 3

Data

To effectively analyze and predict the unemployment rate, data spanning from 2003 to 2022 was utilized. This range was selected to encompass both stable economic periods and major economic disruptions, such as the 2008 financial crisis and the COVID-19 pandemic. By including a diverse range of economic conditions, the dataset is well-equipped to capture various patterns in economic behavior and unemployment trends.

To gather data from Federal Open Market Committee (FOMC) meeting minutes, a web scraping method was implemented using the `FederalReserveMins` class from the `FedTools` library [8]. This tool was chosen for its efficiency in automating the extraction of structured data from the Federal Reserve's website. The scraping process involved constructing links between 2003 and 2022, retrieving articles, and compiling the extracted data into a structured `DataFrame`. This automated process ensured comprehensive coverage of FOMC minutes, minimizing manual effort and maximizing efficiency. The `thread_num = 10` parameter was included to enable multithreading, enhancing the speed of the scraping process. The code utilized is included in the Code Appendix (pg. 41) for further reference.

To extend off the FOMC data, monthly unemployment rate data was collected from the Federal Reserve Economic Data (FRED) platform, specifically from the `UNRATE` series available at FRED [19]. This dataset was chosen due to its established credibility and frequent use in economic research. The data represents seasonally adjusted monthly unemployment rates in the United States, ensuring timely and accurate coverage of the FOMC data.

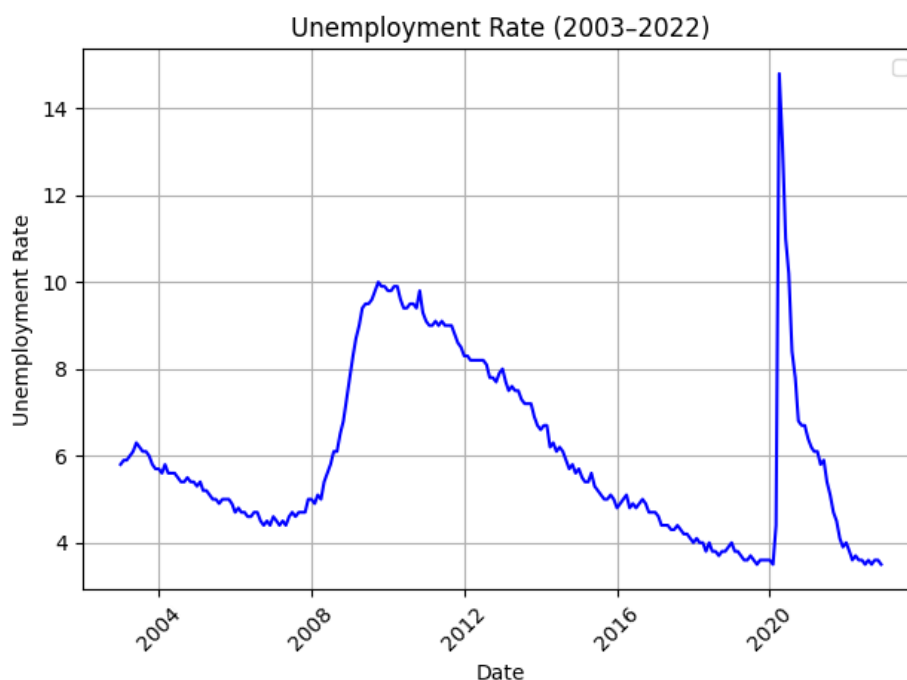


Figure 1. U.S. Unemployment Rate (2003-2022)

Figure 1 shows the U.S. unemployment rate from 2003 to 2022. There are three noticeable areas within the chart: the peak during the Great Recession around 2009, a gradual decline until 2020, and a sharp spike in 2020 due to the COVID-19 pandemic. The rate dropped quickly after the 2020 peak, followed by a relatively stable period with slight fluctuations. Overall, the chart highlights how major economic events impact labor market trends.

This project operates under several key baseline assumptions. First, it is assumed that the Federal Open Market Committee (FOMC)'s language in its meeting minutes reflects real-time economic assessments that are predictive of short-term labor market outcomes. This presumes a direct relationship between the sentiment of the meeting minutes and subsequent economic conditions. Second, the consistency of language and sentiment indicators within FOMC meeting minutes over time is assumed, enabling the effective training of the natural language processing (NLP) model. Finally, it is assumed that sufficient historical data is used for both FOMC meeting

minutes and unemployment rates, providing an efficient foundation for training and testing the NLP model to ensure reliable and meaningful results.

Chapter 4

Methodology

This section outlines the steps taken to develop the prediction model for the U.S. unemployment rate using sentiment analysis of FOMC meeting minutes. Sentiment analysis was conducted on the text, assigning a “positiveness” sentiment score to each communication. These sentiment scores were combined with unemployment data to develop multiple regression models aimed at evaluating the relationship between sentiment and unemployment rates. The most accurate regression model was subsequently applied to examine the relationship between predicted unemployment rates and actual historical unemployment data.

Sub-Chapter 1: Model Comparison

When selecting a model for text-based economic forecasting, various options were considered, including FinBERT, LSTM networks, and LLMs such as LLaMA-3. Each model offers distinct strengths and limitations, making the choice of model highly dependent on the specific objectives of this research project. LSTM networks are well-suited for sequential data processing and have demonstrated strong performance in time-series forecasting tasks, particularly for predicting economic trends based on historical data. However, LSTM models often require extensive fine-tuning to achieve optimal performance and may struggle to capture the broader contextual understanding needed when analyzing complex financial language. On the other hand, LLaMA-3, as a large language model, excels in text generation, document summarization, and conversational tasks. While LLaMA-3 offers versatility and deep contextual analysis, it is resource-intensive and often requires significant fine-tuning for specialized domains such as finance.

Given the focus of this research on sentiment analysis of FOMC communications, FinBERT was ultimately selected for its domain-specific training and superior performance in financial text analysis. As a model fine-tuned on financial data, FinBERT is optimized to understand economic terminology, policy language, and financial sentiment, making it particularly well-suited for this project. Its strong performance in identifying negative sentiment ensures that fluctuations in FOMC language can be effectively captured and interpreted as indicators of economic priorities and outlook. Additionally, FinBERT's efficiency in sentiment classification reduces the need for extensive retraining or data preprocessing when applied to financial texts. By leveraging FinBERT's specialized capabilities, this research aims to extract meaningful insights from FOMC communications to improve unemployment rate forecasting.

Sub-Chapter 2: Sentiment Analysis

The first step involved performing sentiment analysis on the collected FOMC meeting minutes. To achieve this, FinBERT's pre-trained classifier and tokenizer are imported into the coding environment from Yi Yang's *Fin-BERT tone* online source, which are then used to construct an NLP pipeline for sentiment analysis [20]. The dataset containing FOMC meeting minutes, indexed by dates, is processed in three main iterations corresponding to different timeframes for computational ease.

The sentiment analysis process involved iterating through each sentence within the meeting minutes. Each FOMC statement was first tokenized into individual sentences. Sentences exceeding 500 characters were excluded to ensure that only concise and interpretable content was analyzed. For each valid sentence, the pre-trained FinBERT model classified the text's sentiment as either Positive, Negative, or Neutral. This formula tracks the cumulative count of

positive and negative statements, normalizing the score by the total number of sentences. This score served as the primary independent variable for predicting unemployment rates. The positiveness score for each statement was calculated by taking the number of positive sentences minus the number of negative sentences and dividing that value by the total number of valid sentences in the statement:

$$\text{Positiveness Score} = \frac{\# \text{ of Positive Sentences} - \# \text{ of Negative Sentences}}{\text{Total \# of Sentences}}$$

The overall "positiveness" score for each statement is computed as the average sentiment score across all sentences. The histogram below shows a distribution of the positiveness scores across the FOMC communications [20].

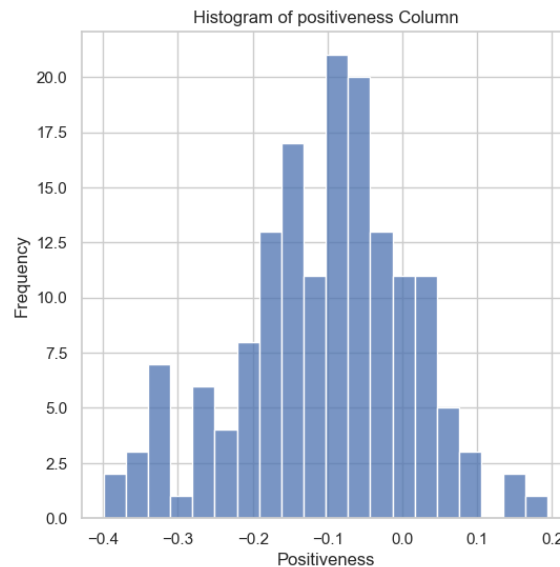


Figure 2. Histogram of Positiveness Sentiment Score

As displayed in Figure 2, the distribution appears to be approximately normal, with a central peak around -0.1, suggesting that the majority of the communications lean slightly negative. The data is moderately symmetrical, though there is a slight skew to the left, with a

heavier tail extending toward more negative sentiment values (around -0.3). This skew indicates that negative sentiment was somewhat more frequent or pronounced in the dataset. Additionally, fewer observations show strongly positive sentiment, which aligns with expectations given the typically cautious and risk-aware language used in FOMC communications. The clustering around slightly negative scores may reflect the Federal Reserve's tendency to emphasize economic concerns, inflation risks, or market uncertainty rather than expressing overtly positive sentiment.

Sub-Chapter 3: Data Integration

To combine the sentiment scores with corresponding unemployment data, monthly unemployment rates were retrieved from the FRED platform via the UNRATE series. Since FOMC meetings are not held consistently every month, sentiment scores were offset by two months to align with the expected lag in economic impact. This offset was implemented to account for the time it typically takes for policy sentiment to influence unemployment trends. Aligning these datasets ensured that the sentiment data appropriately reflected conditions to then cause changes in the unemployment rate. The integrated dataset paired each monthly unemployment value with the corresponding sentiment score (offset by two months), creating a comprehensive time series suitable for regression analysis.

Sub-Chapter 4: Regression Analysis

The integrated dataset was then employed into a regression analysis aimed to explore the relationship between the calculated positiveness scores from FOMC Meeting Minutes and the unemployment rate. By examining this relationship through multiple regression models, this analysis sought to determine which method most effectively captures the underlying patterns and

dependencies. Implementing a variety of regression models was crucial to identifying the most accurate and reliable predictor. Given that economic trends are often influenced by non-linear relationships and external factors, utilizing multiple models allows for a more comprehensive evaluation of potential forecasting approaches. The following section outlines the seven regression models employed, describing the respective methodologies and theoretical advantages.

The XGBoost Regression Model was employed as one of the regression techniques for predicting the unemployment rate. This algorithm is designed for structured data and efficiently captures non-linear relationships by iteratively improving weak learners. The model's strength lies in its ability to minimize errors through gradient boosting, making it an ideal choice for modeling complex economic relationships.

The Gradient Boosting Regression Model was included to compare performance with XGBoost. Although similar in its boosting methodology, this model allows for variations in parameter tuning that may improve prediction results. Its emphasis on gradually correcting model weaknesses improves its performance.

The Light Gradient Boosting Model (LightGBM) was utilized as one of the regression techniques for predicting the unemployment rate. By employing a leaf-wise growth strategy, LightGBM efficiently identifies complex patterns while minimizing loss. Its inclusion in the analysis provided a valuable comparison with other boosting models like XGBoost and Gradient Boosting Regression.

The Random Forest Regression Model was implemented to evaluate the performance of ensemble learning methods that combine multiple decision trees. This method aggregates

predictions across multiple trees to reduce overfitting and improve model effectiveness. By averaging predictions, Random Forest provides stability when modeling noisy data, making it an effective alternative to boosting models.

The K-Nearest Neighbors (KNN) Regression Model was introduced to determine if simpler, non-parametric approaches could effectively capture the relationship between sentiment scores and unemployment rates. By predicting values based on the average of the nearest data points, KNN is highly dependent on data distribution and was tested to assess its suitability for this task.

The Decision Tree Regression Model was employed to explore how well simple decision rules could predict unemployment rates. Decision trees recursively split data into branches based on feature values, making them interpretable yet prone to overfitting. This model was included to compare its performance with ensemble methods such as Random Forest and XGBoost.

The Linear Regression Model was included to serve as a baseline for assessing model complexity and effectiveness. As a fundamental regression method, it assumes a linear relationship between independent and dependent variables. While simpler than boosting models, it provides a useful benchmark for evaluating model performance.

Sub-Chapter 5: Evaluation Metrics

To assess the performance of the regression models, three evaluation metrics were employed. Mean Squared Error (MSE) was selected as a primary measure to evaluate the average squared difference between predicted and actual values. MSE places greater emphasis on larger errors, making it ideal for assessing precision in unemployment forecasting. The MSE formula is given by:

$$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

where y_i are the observed values, $\hat{y}_i$ are the predicted values, and n is the number of observations.

R-squared (R^2) was used to evaluate the proportion of the variance in the dependent variable that can be explained by the independent variable. Higher R-squared values indicate stronger model performance, providing insight into how well the chosen models fit the data. The R-squared formula is given by:

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2}$$

where $\bar{y}$ is the mean of the observed values.

Mean Absolute Percentage Error (MAPE) was employed to measure the average percentage error between predicted and actual values. By calculating percentage deviations, MAPE offers a practical measure of prediction accuracy and is particularly useful in economic forecasting contexts. The MAPE formula is given by:

$$MAPE = \frac{100}{n} \sum_{i=1}^n \left| \frac{y_i - \hat{y}_i}{y_i} \right|$$

By considering all metrics, the analysis ensures a comprehensive understanding of each regression model's strengths and weaknesses, ultimately selecting the approach that best balances predictive accuracy and interpretability.

Chapter 5

Results & Evaluation

This section presents the results obtained from the regression models tested to predict the U.S. unemployment rate using sentiment analysis of FOMC Meeting Minutes. The findings highlight the performance of each model, emphasizing the success of the XGBoost model over alternative approaches. The section also discusses the key insights gained from examining the relationship between sentiment scores and the unemployment rate, as well as the value of integrating financial sentiment data for economic prediction.

Sub-Chapter 1: Model Performance Comparison

The results of this study demonstrate that the XGBoost regression model outperformed all other tested models, including Gradient Boosting, Random Forest, Decision Tree, K-Nearest Neighbors (KNN), and Linear Regression. Each model was tuned for optimal performance using hyperparameter tuning methods to ensure that the results were as accurate as possible.

The success of XGBoost can be attributed to its implementation of gradient-boosted decision trees, which are optimized for both speed and performance. The model builds a sequence of decision trees, where each subsequent tree corrects the errors made by the previous trees. This iterative learning process allows XGBoost to excel in handling complex data patterns and reducing prediction errors. By assigning greater weight to data points that previous models predicted poorly, XGBoost effectively improved its predictions with each iteration.

Table 1 summarizes the results of each model based on three evaluation metrics: Mean Squared Error (MSE), R-squared (R^2), and Mean Absolute Percentage Error (MAPE). These

metrics provide insight into each model's predictive accuracy, ability to explain variance, and percentage error rate in predictions.

Model Name	MAPE	Accuracy	MSE	R ²
Linear Regression	28.24%	71.74%	3.634	0.1022
Random Forest	18.64%	81.36%	1.927	0.5238
KNN	25.29%	74.71%	3.091	0.2357
Gradient Boosting	20.25%	79.75%	1.887	0.5388
Light Gradient Boosting	25.68%	74.32%	3.096	0.2349
Decision Tree	26.65%	73.35%	3.229	0.2021
XGBoost Regressor	14.72%	85.28%	1.363	0.6632

Table 1. Success Metrics Table

Correlation plots were generated for each of the seven regression models to best compare the predicted unemployment rate results of the regression model with the actual unemployment rate. Please see below for discussion on each of the seven models:

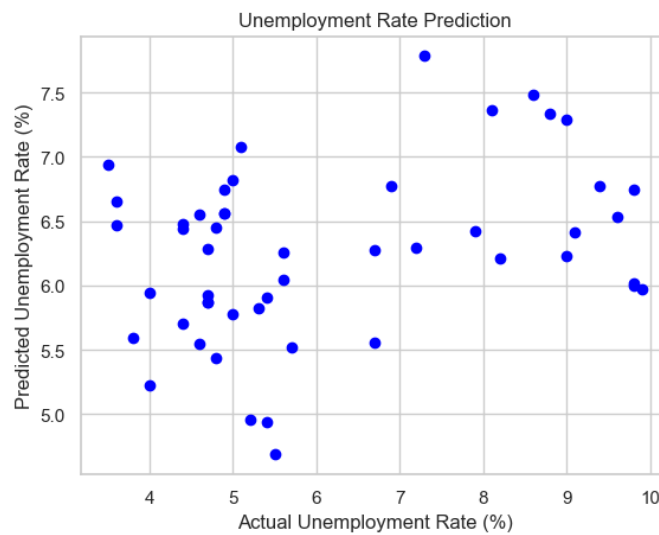


Figure 3. Correlation Graph of Linear Regression Model

Linear Regression: The data points are widely scattered with no clear linear pattern, suggesting that the model struggled to establish a meaningful correlation between the predicted and actual unemployment rates. This aligns with the model's relatively poor performance metrics noted in Table 1.

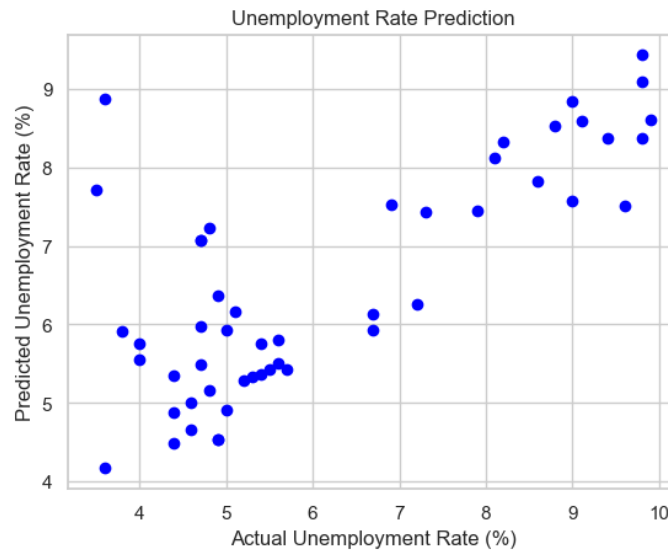


Figure 4. Correlation Graph of Random Forest Regression Model

Random Forest: While some clustering appears around the diagonal (indicating a improved correlation), there are a few points that are still dispersed, reflecting moderate predictive power but limited precision. Random Forest was the second best performer, as displayed by a linear diagonal pattern being formed.

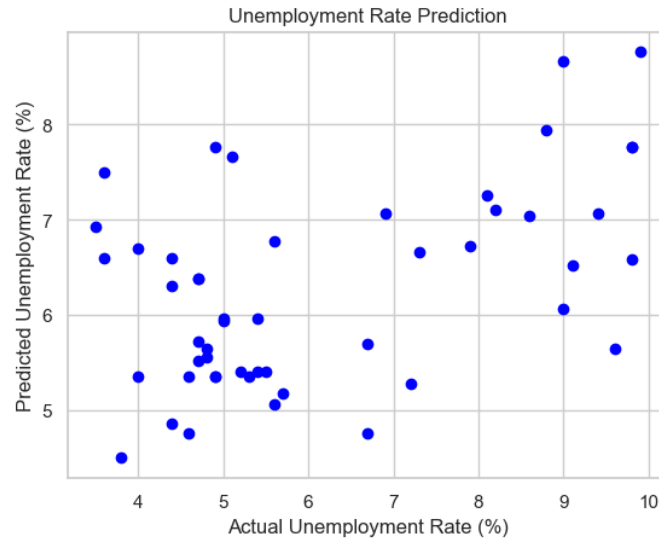


Figure 5. Correlation Graph of KNN Regression Model

K-Nearest Neighbors (KNN): The scatter pattern indicates substantial noise, with no clear trend between the predicted and actual values. This suggests that the KNN model underperformed, likely due to its sensitivity to data variability and its dependence on local relationships, which may not effectively capture complex economic patterns.

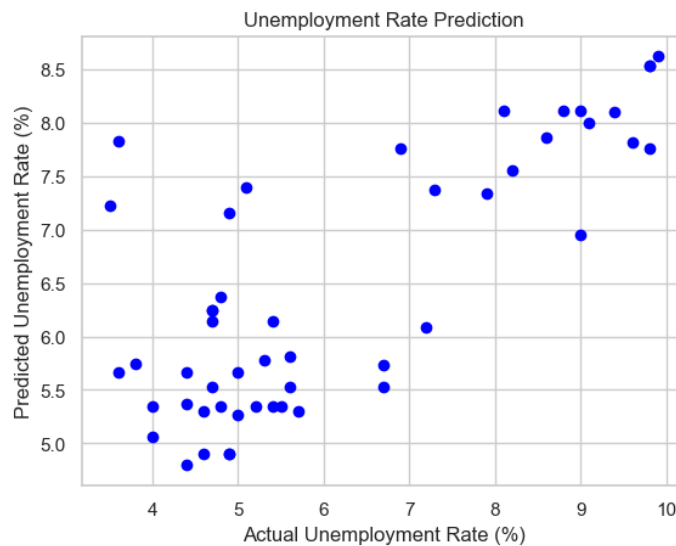


Figure 6. Correlation Graph of Gradient Boosting Regression Model

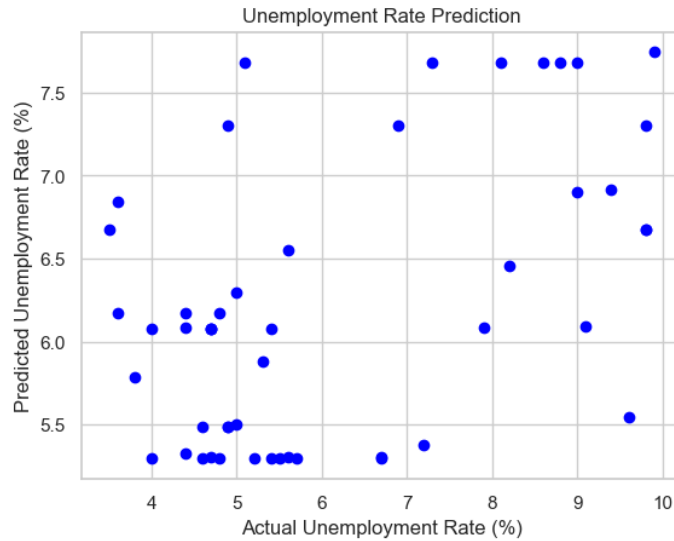


Figure 7. Correlation Graph of Light Gradient Boosting Regression Model

Gradient Boosting and Light Gradient Boosting: These models show modest improvements in alignment with the diagonal line, indicating a stronger correlation with actual unemployment rates. LGBM performed significantly worse than Gradient Boosting, displaying horizontal patterns and limited predictive abilities.



Figure 8. Correlation Graph of Decision Tree Regression Model

Decision Tree: This model exhibits the most disorganized pattern, with predicted values clustering horizontally and failing to track the actual unemployment rates effectively. This result aligns with the Decision Tree model's tendency to overfit the training data, making it less adaptable to unseen data.

Sub-Chapter 2: Analysis of Results

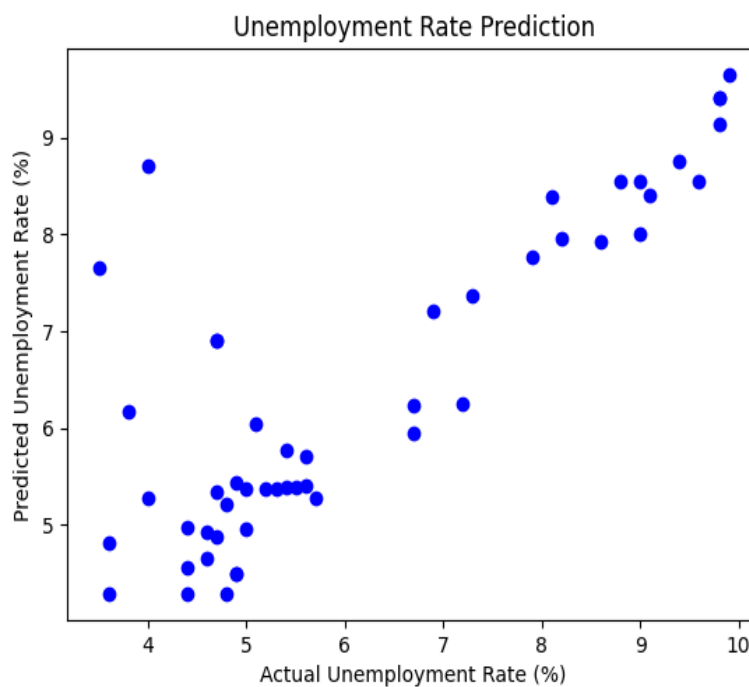


Figure 9. Correlation Graph of XGBoost Regression Model

Looking at Figure 9, the clustering of points along a diagonal trajectory indicates that the model effectively captures the relationship between input features and unemployment rates. While some dispersion remains, particularly in the lower and mid-range unemployment values, the XGBoost model clearly outperforms alternative models in aligning predicted values with actual data.

The superior performance of the XGBoost model can be attributed to its advanced implementation of the gradient boosting framework, which iteratively corrects errors to enhance

accuracy. By effectively managing overfitting through its regularization parameters and being specifically tailored for larger datasets, XGBoost achieved the lowest MSE and highest R^2 among the models.

The Random Forest model, while performing moderately well, lagged behind XGBoost, emphasizing the advantage of boosting methods in refining model predictions, with Gradient Boosting also performing moderately well. Decision Tree, KNN, and Linear Regression models performed notably worse, highlighting limitations in capturing the complexities of economic data with non-linear dependencies.

Sub-Chapter 3: Economic Prediction Application

This study demonstrates that incorporating sentiment analysis of FOMC meeting minutes offers valuable insights for economic forecasting. The calculated positiveness scores provided a meaningful measurement for assessing market conditions and economic outlook. Aligning these scores with unemployment rate data allowed the models to identify patterns that traditional economic indicators might overlook. The results suggest that positive or negative sentiment expressed in FOMC communications has a measurable impact on future labor market trends.

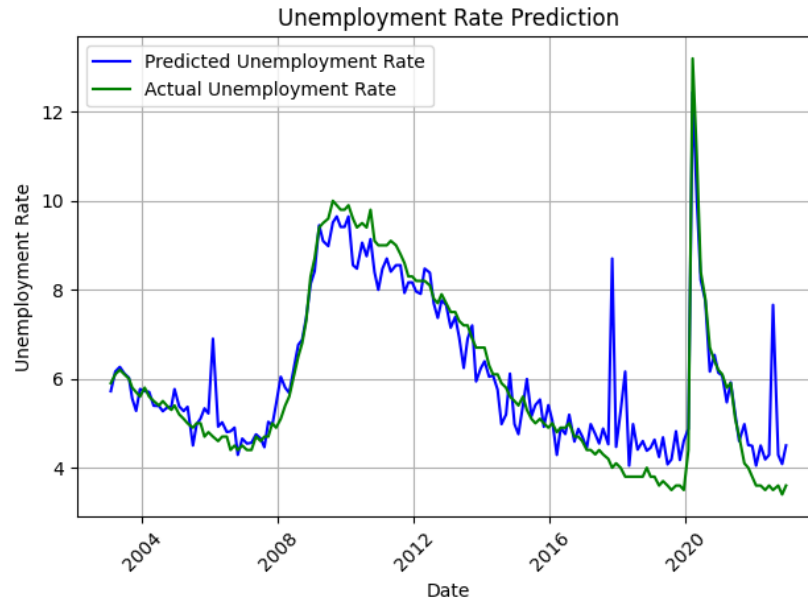


Figure 10. Actual vs. Predicted Unemployment Rate through XGBoost Regression Model

Figure 10 presents the actual and predicted unemployment rates using the XGBoost Regression Model, illustrating the model's ability to track key economic trends displaying future applicability. The graph demonstrates that XGBoost effectively captures the general trajectory of the unemployment rate, particularly during periods of economic stability. The model successfully follows the downward trend from the early 2000s to 2008 and mirrors the sharp increase in unemployment during the 2008 financial crisis and the economic instability following the COVID-19 pandemic, reflecting the importance of tracking sentiment via FOMC communications.

The key findings from this study include the strong performance of the XGBoost model, which produced the most accurate unemployment predictions across all tested metrics. The Random Forest model also demonstrated solid results, while the Gradient Boosting model showed moderate effectiveness. Traditional models such as Decision Tree, KNN, and Linear

Regression proved less effective for this specific economic prediction task. The integration of sentiment analysis significantly improved model performance, suggesting that textual data from FOMC communications provides valuable information for understanding labor market trends and improving unemployment forecasts.

The integration of sentiment analysis into the prediction models was particularly effective because FOMC communications are uniquely suited to reflect economic assessments and policy directions. By tracking these public communications, the sentiment analysis process was able to capture the Federal Reserve’s perspective on market stability, labor trends, and inflation risks. This added dimension of qualitative data improved model performance, providing evidence that textual analysis can enhance economic forecasting efforts.

	Metric Value When FOMC is Negative (Positiveness < -0.2)	Average Unemployment Rate Under All Sentiment
Probability of UNRATE Decrease	77.36%	6.26
Probability of UNRATE Increase	22.64%	

Table 2. Unemployment Forecasting Probability Table

Looking at Table 2, the observed relationship between negative FOMC sentiment and subsequent decreases in the unemployment rate (with a 77.36% probability) appears counterintuitive at first glance, as one might expect a pessimistic tone from policymakers to signal deteriorating labor market conditions. However, after delving deeper into this relationship, there are several plausible explanations and findings that can be derived from this data. First, the FOMC's use of negative language may be proactive rather than reactive, reflecting anticipatory concerns about economic risks rather than current economic weakness. In such cases, a cautious tone could prompt preemptive monetary policy actions—such as interest rate reductions or asset

purchases—that stimulate economic activity and improve employment outcomes in the following months. Second, the observed time lag between sentiment and unemployment rate measurements may capture the effects of these policy responses, rather than immediate macroeconomic conditions. This suggests that negative sentiment may function as an early signal for accommodative interventions that ultimately support job growth. Third, the sentiment scores derived from textual analysis may not align perfectly with real-time macroeconomic data. The FOMC may adopt a cautious communicative tone in response to uncertainty, geopolitical tensions, or global economic developments, even when domestic labor market indicators remain stable or improving. Thus, Table 2 emphasizes that the presence of negative sentiment in FOMC communications does not necessarily imply worsening unemployment and may in fact precede or coincide with labor market improvement.

Chapter 6

Discussion & Future Work

This study examined the application of machine learning models in combination with sentiment analysis from FOMC meeting minutes to predict the U.S. unemployment rate. By leveraging FinBERT for financial sentiment analysis and integrating this data into various regression models, the study identified the XGBoost Regression Model as the most effective method for predicting unemployment trends.

The integration of sentiment analysis from FOMC meeting minutes offered a unique advantage by extracting economic insights directly from policy discussions. The findings demonstrated that positive and negative sentiment scores, calculated using FinBERT, provide valuable insights into future unemployment rates. This approach bridges the gap between traditional economic models and emerging text-based machine learning techniques. While this study successfully demonstrated the predictive power of FinBERT-based sentiment analysis, opportunities remain for enhancing these models and expanding the scope of economic prediction frameworks.

While the models presented in this study performed effectively, several limitations highlight areas for improvement. One key limitation was the reliance on a two-month lag to align sentiment scores with unemployment data. While this approach demonstrates strong predictive power in the short term, economic trends often develop over longer time periods. Extending the study to explore three-month, six-month, or one-year unemployment forecasts may provide additional insights into how FOMC sentiment shapes long-term economic outcomes. Furthermore, FinBERT's reliance on pre-trained financial text embeddings may introduce

limitations when analyzing policy language that evolves over time. As FOMC language shifts with emerging economic conditions, FinBERT's understanding of context may degrade. Exploring alternative models that leverage evolving language trends may help address this challenge.

Building on the limitations identified in this study, future research on hybrid models specifically could expand the application of text-based economic forecasting. Incorporating Neural Network Models such as Recurrent Neural Networks (RNNs), LSTMs, and Gated Recurrent Units (GRUs) offers a powerful opportunity to improve prediction accuracy. These models specialize in handling sequential data, which is highly relevant for analyzing economic trends that rely on temporal patterns. By retaining information across multiple time intervals, LSTM models, in particular, excel at identifying long-term dependencies. Implementing LSTM-based models could improve the prediction of unemployment rates, particularly when forecasting trends over several months or even years. Additionally, Transformer-based models, such as LLaMA-3, present an opportunity for improving sentiment extraction and prediction capabilities. LLaMA-3's architecture allows for deeper contextual understanding of text, outperforming earlier NLP models in understanding complex language patterns. By building a hybrid model potentially incorporating LLaMA-3/ LSTM, this study's methodology could be enhanced by capturing nuanced sentiment shifts in FOMC language that may better predict economic outcomes.

Another promising avenue for future research involves incorporating Genetic Algorithms (GAs) to enhance model performance. GAs mimic the process of natural selection by iteratively refining model parameters to identify the most optimal solution. By applying GAs to optimize

the parameters of machine learning models such as XGBoost, Gradient Boosting, and Random Forest, researchers could further improve model accuracy. Additionally, GAs may assist in selecting key sentiment features, ensuring the most relevant data points are emphasized in economic forecasting. The integration of GAs with NLP models could also improve the sentiment score calculation process. Rather than relying on static weights for positive and negative sentiment scores, a GA could dynamically adjust weight parameters based on changing economic conditions, improving prediction outcomes.

Incorporating additional data sources into the prediction framework presents another opportunity for improving model performance. While this study focused on integrating unemployment rates and FOMC sentiment scores, additional variables such as inflation trends, consumer spending patterns, and stock market volatility could provide deeper insights into economic conditions. Including these metrics could improve model performance, especially when paired with advanced machine learning frameworks that manage high-dimensional data. Future research could also investigate expanding textual data sources by including speeches from Federal Reserve officials, economic reports, and media coverage of monetary policy. These diverse text sources could complement FOMC data, providing a broader understanding of economic sentiment.

The field of NLPs continues to evolve rapidly, presenting exciting opportunities for future research. Recently introduced architectures such as MPT-30B, Falcon-40B, and GPT-4 offer powerful language processing capabilities that may surpass FinBERT's financial language model. These models specialize in analyzing nuanced language patterns, capturing shifting economic landscapes, and improving text classification. Future work could evaluate these

emerging NLP models as potential replacements for FinBERT in economic forecasting applications.

To improve the practical application of this research, future work could also explore the development of a real-time prediction system. By leveraging web-scraping tools and automated text analysis pipelines, future research could develop a system that continuously monitors FOMC statements, economic reports, and market conditions to deliver live unemployment forecasts. Such systems would enhance the accessibility and responsiveness of economic prediction models, ensuring that stakeholders and policymakers receive timely insights and make this project even more useful.

This study successfully demonstrated the predictive power of sentiment analysis combined with machine learning techniques for forecasting unemployment rates. The results highlight the value of integrating FOMC communication data with machine learning models such as XGBoost to improve economic prediction accuracy. While this research achieved promising results, there remains considerable potential for enhancement through alternative machine learning models, broader data integration, and improved interpretability techniques. By exploring the application of neural networks, genetic algorithms, and advanced NLP models such as LLaMA-3, future research can further improve the predictive performance and robustness of economic forecasting models. These advancements have the potential to revolutionize the integration of sentiment analysis into economic prediction frameworks, enabling more accurate and actionable insights for policymakers, investors, and researchers.

Code Appendix

FinBERT Code:

```
# Import the pretrained classifier and tokenizer of Bert from FinBERT online source
finbert = BertForSequenceClassification.from_pretrained('yiyanghkust/finbert-tone', num_labels=3)
tokenizer = BertTokenizer.from_pretrained('yiyanghkust/finbert-tone')

# Build up the NLP pipeline with the pretrained classifier model and tokenizer from above
nlp = pipeline("sentiment-analysis", model=finbert, tokenizer=tokenizer)
```

Positiveness Score Calculation Code (for one subset of Data):

```
date_list = dataset0313.index.tolist()
data_positiveness0313 = dataset0313
for date in date_list:
    sentence_list = tokenize.sent_tokenize(dataset.loc[date, 'Federal_Reserve_Mins'])
    positive_sent=0
    sentence_list=[sent for sent in sentence_list if len(sent) <= 500]
    for sentence in sentence_list:
        nlp_result = pd.DataFrame(nlp(sentence))
        positive_sent += 1 if nlp_result.loc[0, "label"] == "Positive" else -1 if nlp_result.loc[0, "label"] == "Negative" else 0
    positiveness = positive_sent/len(sentence_list)
    data_positiveness0313.loc[date, "positiveness"] = positiveness
data_positiveness0313.to_csv("positiveness0313.csv")
data_positiveness0313
```

Data Integration Code:

```
from dateutil.relativedelta import relativedelta

# Convert dates to datetime format
data_positiveness['date'] = pd.to_datetime(data_positiveness['date'])
unrate_data['observation_date'] = pd.to_datetime(unrate_data['observation_date'])

# Add two months to FOMC meeting dates
data_positiveness['corresponding_date'] = data_positiveness['date'].apply(lambda x: x + relativedelta(months=2))
# data_positiveness['corresponding_date'] = data_positiveness['date'] + relativedelta(months=2)

# Extract only year and month for merging
data_positiveness['merge_key'] = data_positiveness['corresponding_date'].dt.to_period('M')
unrate_data['merge_key'] = unrate_data['observation_date'].dt.to_period('M')

# Merge based on the "merge_key"
merged_data = pd.merge(data_positiveness, unrate_data, on='merge_key', how='left')

# Drop unnecessary columns if needed
merged_data = merged_data.drop(columns=['corresponding_date', 'merge_key', 'observation_date'])
merged_data
```

XGBoost Regression Code:

```
from xgboost import XGBRegressor
X = merged_data.drop(columns=['date', 'Federal_Reserve_Mins', 'UNRATE'])
y = merged_data['UNRATE']

X_train, X_test, y_train, y_test = train_test_split(X, y, test_size=0.2, random_state=42)

# Model initialization and training
model = XGBRegressor(n_estimators=100, learning_rate=0.1)

model.fit(X_train, y_train)

y_pred = model.predict(X_test)

mse = mean_squared_error(y_test, y_pred)
r2 = r2_score(y_test, y_pred)
mape = mean_absolute_percentage_error(y_test, y_pred)

# Success metrics
print(f'Mean Squared Error: {mse}')
print(f'R-squared: {r2}')
print(f'Mean Absolute Percentage Error (MAPE): {mape * 100:.2f}%')
```

Prediction Comparison of XGBoost Predictions and Actual Unemployment Data Code:

```
from xgboost import XGBRegressor
from sklearn.model_selection import train_test_split
from sklearn.metrics import mean_squared_error, r2_score, mean_absolute_percentage_error

X = merged_data.drop(columns=['date', 'Federal_Reserve_Mins', 'UNRATE'])
y = merged_data['UNRATE']

X_train, X_test, y_train, y_test = train_test_split(X, y, test_size=0.2, random_state=42)

# Model initialization and training
model = XGBRegressor(n_estimators=100, learning_rate=0.1)
model.fit(X_train, y_train)

# Predict for the entire dataset
y_pred_full = model.predict(X)

y_pred_test = model.predict(X_test)

# Success metrics
mse = mean_squared_error(y_test, y_pred_test)
r2 = r2_score(y_test, y_pred_test)
mape = mean_absolute_percentage_error(y_test, y_pred_test)

print(f'Mean Squared Error (Test): {mse}')
print(f'R-squared (Test): {r2}')
print(f'Mean Absolute Percentage Error (Test): {mape * 100:.2f}%')

# y_pred_full now matches the length of merged_data["date"]
```

Dataset Appendix

Example of Merged Dataset:

Date	Federal_Reserve_Mins Script	Positiveness	Act. UNRATE
2003-01-29	The Federal Reserve, the central bank of...	0.045643	5.9
2003-03-18	A meeting of the Federal Open Market...	-0.174863	6.1
2003-05-06	A meeting of the Federal Open Market...	-0.250891	6.2
2003-06-25	A meeting of the Federal Open Market...	-0.092429	6.1
2003-08-12	A meeting of the Federal Open Market...	-0.042229	6.0
...	...	...	...
2022-06-15	The Federal Reserve, the central bank of...	-0.147727	3.6
2022-07-27	The Federal Reserve, the central bank of...	-0.236934	3.5
2022-09-21	The Federal Reserve, the central bank of...	-0.147410	3.6
2022-11-02	The Federal Reserve, the central bank of...	-0.133588	3.4
2022-12-14	The Federal Reserve, the central bank of...	-0.167969	3.6

BIBLIOGRAPHY

- [1] Bande, Roberto, and Marika Karanassou. "The Natural Rate of Unemployment Hypothesis and the Evolution of Regional Disparities in Spanish Unemployment." *Urban Studies*, vol. 50, no. 10, 2013, pp. 2044–62. *JSTOR*, <http://www.jstor.org/stable/26144590>. Accessed 1 Apr. 2025.
- [2] Blanchflower, David G. & Bryson, Alex, 2021. "The Economics of Walking About and Predicting Unemployment," GLO Discussion Paper Series 922, Global Labor Organization (GLO).
- [3] Bukenya, Jetro, et al. "Nowcasting Unemployment Rate During the COVID-19 Pandemic Using Twitter Data: The Case of South Africa." *Frontiers in Public Health*, vol. 10, 2022, <https://doi.org/10.3389/fpubh.2022.952363>.
- [4] "CDOC Recommendations Report Comment Analysis." Federal Data Strategy, Dec. 2022, https://resources.data.gov/assets/documents/CDOC_Recommendations_Report_Comment_Analysis_FINAL.pdf.
- [5] Chakraborty, Tanujit, et al. "Unemployment Rate Forecasting: A Hybrid Approach - Computational Economics." SpringerLink, Springer US, 19 Aug. 2020, link.springer.com/article/10.1007/s10614-020-10040-2.
- [6] Dahir, Ahmad, et al. "Similarity Analysis of Federal Reserve Statements Using Document Embeddings: The Great Recession vs. COVID-19." *SN Operations Research Forum*, vol. 3, no. 3, 2022, <https://doi.org/10.1007/s43546-022-00248-9>.
- [7] Doh, Taeyoung, et al. Deciphering Federal Reserve Communication via Text ..., 20 Oct. 2020, www.bancaditalia.it/pubblicazioni/altri-atti-convegni/2020-bi-frb-nontraditional-data/doh_paper.pdf.
- [8] Federal Open Market Committee. (n.d.). Meeting calendars and information. The Fed - Meeting calendars and information. <https://www.federalreserve.gov/monetarypolicy/fomccalendars.htm>

- [9] Huang, A., Wang, H., & Yang, Y. (2023). FinBERT: A Large Language Model for Extracting Information from Financial Text. Retrieved July 5, 2023, from https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3910214FedTools. (20223). Retrieved July 5, 2023, from <https://pypi.org/project/Fedtools/>.
- [10] Khademi, Mohsen, et al. "ARIMAX and ARX Models with Social Media Information to Predict Unemployment Rate." *Journal of Advanced Management Science*, vol. 3, no. 4, 2015, pp. 362–369, <https://www.joams.com/uploadfile/2015/0914/20150914022716863.pdf>.
- [11] Kaouk, T., & Alvares, C. (2021, July). CDO. Federal CDO Council - CDO Council pilot uses NLP for Comment Analysis. <https://www.cdo.gov/news/comment-analysis/>.
- [12] Lee, Kyunghoon, and Jaewon Lee. "Predicting the Unemployment Rate Using Social Media Analysis." *Journal of Information Processing Systems*, vol. 12, no. 3, 2016, pp. 582–589, https://s3.ap-northeast-2.amazonaws.com/journal-home/journal/jips/fullText/26/jips_582.pdf.
- [13] Mai, Zhelu & Zhang, Jinran & Xu, Zhuoer & Xiao, Zhaomin. (2024). Financial Sentiment Analysis Meets LLaMA 3: A Comprehensive Analysis. 171-175. 10.1145/3696271.3696299.
- [14] Mero, Kevin, et al. ‘Unemployment Rate Prediction Using a Hybrid Model of Recurrent Neural Networks and Genetic Algorithms’. *Applied Sciences (Basel, Switzerland)*, vol. 14, no. 8, MDPI AG, Apr. 2024, p. 3174, <https://doi.org/10.3390/app14083174>.
- [15] Mohamed, Khalid, et al. "NLP-Based Bi-Directional Recommendation System: Towards Recommending Jobs to Job Seekers and Resumes to Recruiters." *Machines*, vol. 6, no. 4, 2022, <https://doi.org/10.3390/machines6040147>.
- [16] Patterson, Adam. “Examining Text Consistency amongst FOMC Statements and Minutes Subtext Using Deep Learning.” *Issues in Information Systems*, 2024, iacis.org/iis/2024/4_iis_2024_252-260.pdf.

- [17] Smith, Larry. "The Great Recession: A Statistical Analysis of Its Effects on Unemployment." *Journal of Business and Economics Research*, vol. 14, no. 5, 2015, pp. 301–312, <https://go.gale.com/ps/i.do?id=GALE%7CA430547672&sid=googleScholar&v=2.1&it=r&linkaccess=abs&issn=1931907X&p=AONE&sw=w>.
- [18] Shapiro, Adam, and Daniel Wilson. "Taking the Fed at Its Word: A New Approach to Estimating ..." Federal Reserve Bank of San Francisco, www.frbsf.org/wp-content/uploads/wp2019-02.pdf.
- [19] U.S. Bureau of Labor Statistics, Unemployment Rate [UNRATE], retrieved from FRED, Federal Reserve Bank of St. Louis; <https://fred.stlouisfed.org/series/UNRATE>, March 31, 2025.
- [20] Yi, Y. (2020, June 10). FinBERT-tone. Retrieved July 5, 2023, from <https://huggingface.co/yiyanghkust/finbert-tone>.