

Sentiment Prediction for Market Volatility

Niraj Patel¹

¹ST. Clair College

Publication Date: 2025/04/16

Abstract: This project presents an automated framework for generating sentiment metrics from SEC 10-K filings, aiming to predict stock market returns and volatility at the sector, portfolio, and firm levels. The system comprises two core models: an SEC Filing Extraction Model, which preprocesses filings, and a Supervised Lexicon Learning Model, which analyzes sentiment using a four-step process. This includes identifying sentiment-related words, assigning predictive weights, aggregating sentiment scores, and applying the Kalman Filter for trend analysis. Empirical results demonstrate the effectiveness of sentiment metrics from 10-K filings, particularly the Item 1A risk factor section, in forecasting market movements.

How to Cite: Niraj Patel. (2025). Sentiment Prediction for Market Volatility. *International Journal of Innovative Science and Research Technology*, 10(2), 2531-2555. <https://doi.org/10.38124/ijisrt/25feb804>

I. INTRODUCTION

➤ Motivation

Now is the age of Artificial Intelligence and Big data. With the advance of computational powers, a large amount of various data, such as text, video, and audio have been used for scientific analysis. Among a myriad of data forms, textual data has gotten the fastest attention in the social science academic field. Textual data's numerical representation for statistical analysis in nature is extremely high in dimensions that empirical study seeking its textual richness should face its dimensionality challenges. Machine learning will be employed to extract richer meaning from textual data for predictive analysis in a high-dimensional data environment [1].

In finance, textual data is commonly employed for predicting market movements[1], [2], [3], [4], [5], [6]. In stock prediction, textual analysis of market sentiment has shown notable success. News data was employed to analyse sentiments in the prediction of short-term stock price movements [3]. Similarly, social media textual data was utilised by integrating social media sentiment and AI [5]. Also, annual report data was used for stock market forecasting [6]. Likewise, we could find myriad types of textual data were used in predicting market movement. Then, what type of textual data can be informative for such a purpose? If you want to invest in a public company in the United States, where can you begin your investment journey?

There are myriad ways to begin your investment in a public company. However, what if you do not know about a firm at all in which you would invest or what if you do just partially know the firm? Then you should first know the firm correctly. If so, where we can find reliable and trustworthy information about a firm? You can find a rich deposit of reliable knowledge on Form 10-K filing. The Form 10-K filing has been mandated by the Securities and Exchange

Commission (SEC) since 1934. It has its origins in the Section 13 or 15(d) of the Securities Exchange Act of 1934. The 10-K is a comprehensive official document offering a through overview of a company's business, its' potential challenges, and its financial performance through the fiscal year. A company's leadership, in the 10-K, provides their perspective on the business outcomes and the factors influencing them [7]. Furthermore, several studies found that 10-K filings offered predictive power to stock price prediction[8], [9], [10], [11], [12]. [8] found a shred of evidence that the market reacted to 10-K filings in a statistically significant way. [10], [9] showed there was a correlation between the complexity of 10-K filing and stock price volatility. [11], [12] found a positive correlation between 10-K filing and stock price through cutting-edge AI methodologies. Hence, investors should pay attention to 10-K filings.

Since the beginning of 2005, a new section has been required to be included in all firms 'annual filings by the SEC. The section called "Section 1A of the Annual Report on Form 10-K" discusses "the most significant factors that make the company speculative or risky" [13]. Prior to this alteration, companies were only obligated to provide this information in their registration when issuing their equity or debt securities. Also, some companies voluntarily offer risk disclosures in the section called "Management's Discussion and Analysis of Financial Condition and Results of Operations(MD&A)"

Opponents of the new disclosure requirements argue that risk factor disclosures are unlikely to offer valuable information. First, risk disclosure can be biased. Managers might resist disclosing negative information about their business or career incentives [14], [15], [16], [17], [18]. Second, managers' overconfidence could make them perceive less risk or overconfident managers could have the illusion that they can effectively manage the risks confronting their firms [19]. Third, managers tended to disclose all possible

risks and uncertainties without making precise predictions or providing detailed financial assessments. This practice stems from the fact that companies were not required to predict the possibility that a disclosed risk would actually materialise. Furthermore, there was no obligation for firms to specify the financial influence that a disclosed risk might have on their present or future financial statements [20]. Since 2010, the SEC has warned companies to “avoid generic risk factor disclosure that could apply to any company” [21] and has continuously pushed the precise risk factor disclosures through the comment letter process [22]. Recently, the SEC has been demanding the explicit and specific inclusion of both cyber security risk factor disclosure and climate change risk factor disclosure [23], [24].

Notably, numerous studies reach an opposite conclusion by providing evidence that the risk factor disclosures, in fact, provide valuable information [25], [26], [27], [28], [29], [20], [19]. The disclosures reflect the genuine risks confronting their firms. Also, it might help investors assess the volatility of a firm’s cash flows, and tax-related risk factor disclosures offer details about the level of a firm’s future cash flows, helping investors incorporate this information into current stock prices [25]. The newly created risk factor disclosures also show a correlation with conventional asset pricing risk factors, indicating that the disclosed factors are valuable for assessing overall risk [26]. cyber security risk factor disclosures increase the risk of a company’s stock price declining in the future [27]. Furthermore, it has been revealed that the length of the disclosure is associated with market reaction. Lengthier risk factor disclosures have a negative correlation with market reactions [25]. More detailed disclosures tend to generate more profound market reactions [28]. Alterations in the length of risk disclosures also can influence an investor’s risk perceptions [29]. From these observations, although disclosures are occasionally seen as generic [20] or susceptible to bias [19], they still provide risk-related information that can assist investors and affect stock values.

➤ Project Objective

Given the informational value of 10-K filings, as introduced in the motivation section, our project aims to achieve the following main goal:

- Developing an automated pipeline to generate sentiment scores from 10-K reports of firms in the technology industry for market movement prediction of a firm, i.e., return and volatility.

The Following are the Sub-Steps for the Main Goal:

- *Extraction of 10-K fillings, followed by extraction of risk factors from the extracted 10-K fillings. The list of 10-K fillings was based on the Invesco QQQ Trust Series 1(10-K).*

Our pipeline will collect an annual disclosure from the SEC’s Electronic Data Gathering, Analysis, and Retrieval(EDGAR) system and extract the risk factor section of each report. It also will be able to automatically collect the

latest 10-K filings for prompt generation of sentiment information.

- *Generate Various Sentiment Scores from the Extracted 10-K Reports.*

In this research, our study aims to generate sentiment metrics for market movement prediction of a firm, i.e., return and volatility, from both the entire 10-K report and risk factor disclosures. This is the main goal of my project. Recently, a new model has been proposed to generate strong indicators for predicting price reactions to new information. This model [1] generated a powerful return-predictive signal from their supervised sentiment text model with news data. Our research methodology for generating sentiment scores referred to the methods suggested in [1]. However, our study will show an innovative improvement compared to [1]. [1] employed news articles only to generate return-predictive signals, while our study will use 10-K reports, especially focusing on the risk factor section, to generate volatility-predictive signals as well as return-predictive signals. Furthermore, the predictive sentiment signals will be generated in three different stakeholder levels such as a sector, a portfolio, and an individual firm. As far as we are aware, no research applies supervised sentiment learning with 10-K reports for predicting volatility, nor is there any that offers a comparative study of pre-established and acquired sentiments by using 10-K reports for returns or volatility predictions.

- *Build an Automated Pipeline on the Airflow Framework.*

10-K filings should be released annually, but the publication dates of it are different to each firm. In the case of 100 firms listed in the QQQ for our study, for instance, a firm’s 10-K is released almost every single day within that year. Due to that, in order to offer prompt sentiment information, our system should update the latest 10-K filings daily or at least monthly.

- *Evaluate Sentiment Scores in the Context of Prediction for Returns and Volatility.*

After we achieve our main goal, we will evaluate our generated sentiment scores quantitatively and qualitatively. For quantitative analysis, we will use Pearson correlation (check Appendix E for the formula) to find a correlation between the metrics we generated. For qualitative analysis, we will employ the top 15 most influential words we extracted from the process of sentiment score generation. The most impactful words will be used to generalise a topic which may affect a sentiment.

➤ Contribution

We have completed the main goal, including all sub-goals, we set in the previous project objective section.

Our contributions are shown as follows:

- *Developed an Automated System for Generating Sentiment Analysis Metrics, Comprising Two Main Models:*
- ✓ The 10-K Filing Extraction Model, which automatically retrieves 10-K filings from the Electronic Data Gathering,

Analysis, and Retrieval (EDGAR) system, including the extraction of the Item 1A risk factor disclosure section. This data feeds into the Sentiment Score Prediction Model. Referenced in [30].

- ✓ The Sentiment Score Prediction Model, which evaluates sentiment across three levels of stakeholders: sector (e.g., Technology), portfolio (e.g., top 10 firms), and individual company (e.g., Nvidia). Referenced in [1]. Unlike the model in [1], our model demonstrates results across three levels of stakeholders to maximize the model's applicability.
- *Formulated 12 Sentiment Metrics across these Three Levels:*
 - ✓ Sector Level: Developed sector sentiment metrics from 10-K filings, labelling based on QQQ's return and volatility.
 - ✓ Portfolio Level: Formulated portfolio sentiment metrics using 10-K filings, labelling based on portfolio returns and volatility.
 - ✓ Company Level: Established company-specific sentiment metrics through 10-K filings, labelling based on company returns and volatility.
 - ✓ Generated sentiment scores for 10-K/Item 1A using the methodology in [1]. This approach is novel, as no prior study has applied this methodology to 10-K filings/Item 1A risk factors.
 - ✓ Utilized return/volatility to train sentiment scores, which is a unique approach. No other study, including [1], has used volatility to train sentiment signals. Our model produced both volatility-predictive and return-predictive sentiment signals.
 - ✓ Assessed the performance of the Sentiment Score Prediction Model.
 - ✓ Conducted both quantitative (correlation analysis) and qualitative (most influential words to return/volatility) analysis to critically evaluate the derived sentiment metrics. These evaluations help investors improve their understanding of market or firm fundamentals.
 - ✓ Connected to the Airflow server for automation updates. Our system automatically updates the latest 10-K filings, facilitating seamless data preparation and providing the latest sentiment scores for three stakeholder levels. Investors can receive prompt sentiment information for a sector, portfolio, and a firm. Detailed explanation available in Appendix H.

➤ *Outline of the Project*

The thesis structure is outlined below:

- Chapter 1 - Introduction provided motivation for this study, including study objectives and contributions
- Chapter 2 - Related Works explained related works about various textual sentiment analysis approaches as well as its relevant background information.
- Chapter 3 - 10-K filings Extraction Model explained the 10-K filing acquisition process as well as the detailed extraction process of Item 1A risk factor section.
- Chapter 4 - Methodology detailed a comprehensive explanation of the sentiment score prediction model.

- Chapter 5 - Experiment Results critically evaluated all generated sentiment metrics through quantitative (i.e. correlation analysis) and qualitative analysis (i.e. most influential word set).
- Chapter 6 - Conclusion summarised the results of our suggested model and suggested to investors about model usage. Limitations and future works were mentioned.

II. RELATED WORKS

➤ *Textual Sentiment Analysis in Finance*

In the field of finance, sentiment analysis has grown in significance due to its wide range of applications. Through sentiment analysis, we can analyse, interpret, and derive insights from large volumes of financial data. One popular use case of sentiment analysis is stock market movement prediction [2], [3], [4], [6], [5], [31], [32]. In the context of sentiment analysis in stock market prediction, two types of sentiments have been researched: investor sentiment and textual sentiment. Investor sentiment refers to an investor's subjective sentiment on firms or markets. The second type of sentiment is textual sentiment or text-based sentiment. It indicates the level of positive or negative sentiment in texts. In some studies [1], [33], for instance, the term 'tone' (i.e. positive or negative) in a corporate disclosure means sentiment. Investor sentiment and textual sentiment are fundamentally different. Investor sentiment encompasses the subjective assessments and behavioural traits of investors. In contrast, textual sentiment may incorporate investment sentiment but also covers a more objective representation of the state within companies, institutions, and markets [34]. In the context of textual sentiment analysis, various sources have been used, such as news articles, social media data, or corporate disclosures. News articles and social media data are used to analyse the short-term effects of the sentiment on market variables like return, volatility, and stock volume [3], [5]. Corporate disclosure usually focuses on finding the relationship between sentiment and firm performance [35], [36]. In this paper, we focused on text-based sentiment using a corporate disclosure, analysing its impact not only on firm performance but also on portfolios and the market.

Furthermore, many previous studies to extract textual sentiments depended on pre-defined sentiment dictionaries instead of using statistical text analysis. Pre-defined sentiment dictionary is created through the dictionary-based approach. This approach utilises a mapping algorithm through which a computer programme reads text and categorises words, phrases, or sentences into predefined categories [37]. There are major pre-defined sentiment dictionaries for textual sentiment analysis in finance. The first one is the General Inquirer (GI) or Harvard IV-4 dictionary (HIV4). The second one is the DICTION dictionary. The two dictionaries have been widely used for financial analysis [38], [39], [40], [41], [42], [43]. However, both the GI/Harvard and DICTON, being general English linguistic dictionaries, offer inaccuracies in a financial context. To overcome this issue, Loughran and McDonald recreated the LM dictionary specific to the finance domain. Since the LM dictionary was developed for 10-K sentiment assessment, we used this dictionary as our benchmark to

evaluate our estimated sentiments. The LM was used by many researchers [44], [45], [46], [47], [48].

➤ *Textual Sentiment Analysis with AI in Finance*

With the increase in computing power and the development of cutting-edge AI methodologies, AI technology has been applied to textual sentiment analysis in finance. In classical machine learning, some previous papers utilised off-the-shelf machine learning techniques like Support Vector Machine(SVM), Naïve Bayes, Decision Trees, or artificial neural networks(ANN) to control the curse of high dimensionality in textual sentiment in finance context [49], [37], [50], [51], [52]. For instance, [52] extracted sentiments from media textual data(e.g. StockTwits) through a variety of machine-learning binary classifiers. They found that the SVM classifier achieved higher accuracy than Decision Trees and Naïve Bayes. However, classical machine-learning techniques could not capture a sentence's complex features and contextual information. These tasks require deep-learning techniques, which facilitate understanding sequential information, complex feature extraction, and location identification [53].

Deep learning, a branch of machine learning, employs multi-layered neural networks, called deep neural networks, for extracting complex features [54]. In textual sentiment analysis, deep learning can be useful for generating learning patterns and learning contextual information in a sentence [55]. Many studies showed the effectiveness of deep learning models, such as recurrent neural networks(RNN) [56], [57], convolutional neural networks(CNN) [58], [59], [60], and attention mechanisms [53], [61] for sentiment analysis in finance. Recently, transformer architectures such as BERT or RoBERTa have shown superior performance in sentiment analysis in the finance domain [62]. However, transformers' superior performance comes at a cost. They demand extensive data and computational power for training and testing. Moreover, they require considerable prediction times, rendering them less viable for real-time applications or environments with constrained processing capabilities [63]. Also, the transformer is perceived as a 'black box' due to its uninterpretable complex internal mechanisms [64].

Recently, [1] suggested an interesting and innovative sentiment model with a good performance. The suggested model is a lexicon learning model to find the correlation between the sentiments from firm-specific news and returns. This model has five virtues. First, this model is a transparent and simple supervised lexicon learning model. It needs only basic econometric methods, such as correlation analysis and maximum likelihood estimation. Hence, this approach is entirely 'white box'. Secondly, this model requires minimal computing resources. It only takes a matter of minutes to handle millions of documents on a laptop computer. Thirdly, this model has a broad range of scalability. Unlike existing lexicon-based models that depend on a pre-existing sentiment dictionary, this model is able to use various types of textual data in the finance domain without a pre-defined dictionary. Fourthly, this supervised model does not require manual labelling labour to train the sentiment model. The model is free from the expensive expense of a significant amount of

manual labelling labour. With minimal, reliable, and clever assumptions, the labelling mechanism works in an automated process. Finally, by training on returns from sampled companies to analyse sentiments, this method is likely more effective at predicting stock returns compared to others.

In this paper, we gained access to the model's other four benefits by leveraging its scalability. In other words, unlike the model using news articles to generate return-predictive sentiment signals, we used informative 10-K filings(plus, Item 1A risk factor section) to produce volatility-predictive sentiment signals as well as return one instead. No prior study applied the model's methodology to the 10-K filings and risk factor section. Also, even in this methodology, they did not use volatility to train sentiment signals. Moreover, we broadened the range of model application levels from a portfolio level to a sector level and a firm level to maximise the model's applicability.

III. 10-K FILINGS EXTRACTION MODEL

In this research, we utilised textual data for financial analysis. Among the myriad kinds of textual data available, we selected Form 10-K filings, which contain some of the richest information about firms. This paper collected the entire 10-K filings for predicting sentiment scores while collecting the Item 1A risk factor section of the filing to extract more informative features. Form 10-K filings are the official documents that all publicly traded firms in the United States are required to submit to the SEC. The SEC enforces very strict rules regarding the content and structure of the information required in Form 10-K filings. These filings contain no pictures or charts. A well-structured 10-K is divided into five separate parts. The first three parts offer a concise summary of the firm's primary business activities, including its services and products; enumerate every risk encountered by the firm; and provide detailed financial information about the firm over the last five years. The fourth part delivers a senior management analysis of its financial results. The final part includes the actual financial figures; the firm's audited financial statements, which consist of the income statement, balance sheets, and statement of cash flows. A detailed description of the 10-K structure can be found in Appendix A. Hence, the Form 10-K exhibits uniform structures. Thanks to these organised structures, we can algorithmically extract information on the filings through. [65], [30]

A. 10-K Filings Collection

Form 10-K filings can officially be found on the SEC website, where they are available to the public. The SEC provides 10-K filings in HTML or TXT formats through its database system, known as the Electronic Data Gathering, Analysis, and Retrieval(EDGAR). This system offers various official filings, including 10-K, 10-Q, 8-K, and 6-K, among others. For our research, we were able to collect most of the 10-K filings from firms in the technology sector on the EDGAR. EDGAR offers files in both TXT and HTML formats, but since most of the 10-K reports were easily accessible in HTML format, our algorithms focused exclusively on extracting HTML files. For instance, EDGAR

has 8140 filings in HTML format of all firms in the S&P 500, whereas, there are only 48 filings in TXT format [30]. We considered the portfolio of Invesco QQQ Trust Series 1, an exchange-traded fund (QQQ or QQQ ETF), to compile a list of tech firms representing the technology sector in the United States. The QQQ is designed to passively track the Nasdaq 100 Index, which conventionally represents the technology sector. More details of the QQQ portfolio can be found in Portfolio Appendix C.

Form 10-K filings can be easily collected by using a Central Index Key (CIK) in EDGAR. A CIK is a number given to an individual or company by the SEC, used to identify the ownership of a filing. Initially, we collected the CIKs list of firms listed in QQQ to facilitate data retrieval through the RSS Feeds provided by EDGAR. It is important to consider the crawler's limitation, which is a maximum request rate of ten requests per second, to ensure equitable access. In this report, we specifically focused on Item 1A, the risk factor, as well as the entirety of the 10-K filing for our purpose. As the mandatory disclosure of Item 1A as a separate section of the filing has been required since 2005, our crawler collected 10-K forms spanning nearly 17 years, from January 2006 to December 2023, for every firm listed in QQQ. In total, we collected 1383 filings, primarily stored in HTML. You can find an example case in Appendix B.

B. Risk Factor Extraction Process

In this study, our focus was on Item 1A, the Risk Factor section, to attain richer information. Initially, we extracted all contents from the Item 1A Risk Factor section. We separately extracted each risk heading along with its corresponding content and then aggregated these elements together. To extract only the risk factor section, we used a well-organised structure of 10-K. You can find an example case in Appendix B.

The Form 10-Ks feature well-structured formats in HTML. We observed the following regularities in the structure of the 10-Ks:

- The heading of a risk is followed by the explanations of the risk (Check the example in Appendix B).
- There is a summary of the Risk Factor section at the beginning of the section.
- Non-informative elements are located in the footers, including page numbers, copyright information, or disclaimers
- Reiterated expressions appear in a certain section; specifically, "Item 1A Risk Factor (Continued)."
- Diverse layouts exist within a report, with variations in fonts, headings, and overall formatting.

With the aid of the regularity in the filing's structure in HTML, our algorithms can accurately locate and extract the risk factor section. The process of extracting risk factors from HTML involves four steps:

- Page Footers Removal
- Headings Extraction
- Risk Factor Section Detection
- Titles Extraction

C. Page Footers Removal

There is repetitive information in page footers of HTML files. Initially, we removed "Item 1A. Risk Factors(Continued)" by using RegExr 4 at Table. We observed that page footers are typically found near horizontal lines marked by '<hr>' or '<div>' tags with 'page-break-after: always' RegExr 2 at Table. They were then replaced with a 'split of pages' marker by using the BeautifulSoup library. Our algorithm subsequently identified and removed non-informative page footers- both textual and numerical- located near these markers by scanning both forward and backward.

D. Headings Extraction

Headings were extracted by detecting tags typically associated with headings, identified by their styles or attributes. The algorithm uses these attributes at Table to identify headings and encapsulate their contents within '<heading>' and '</heading>' markers. These identified headings are then readily identifiable for further analysis steps. Additionally, we removed the '<heading>' and '</heading>' tags when the heading contained five words or fewer.

E. Risk Factor Section Detection

The risk factors section is extracted by identifying the heading "Item 1A - Risk Factor" RegExr 6,7 at Table, and the subsequent heading RegExr 8-16 at Table. The contents of the section are located between the heading and the subsequent headings. Typically, if the pattern "Item 1A -Risk Factor" appears in the heading, it may contain extra spaces, line breaks or variations. Conversely, if it is not in the heading, the pattern is consistently enclosed in special quotation marks, either "“", "”", or """. The algorithm, improved by [65], detects the positions of the several types of headings by iterating through the regex pattern at Table I. Additionally, after extracting the contents of the section, the algorithm checks if the extracted content typically exceeds 1000 characters in length.

F. Titles Extraction

The extracted headings include both titles (such as "Risks Related to Legal, Regulatory and Compliance Matters; FORWARD-LOOKING STATEMENT") and headings (like "We may be unable to adequately protect our proprietary intellectual property rights, which may limit our ability to compete effectively."). We noted that titles have all words starting with capital letters, after removing stopwords. Thus, for headings fitting this pattern, we replace their markers with '<title>' and '</title>'.

Table 1 Heading Pattern

Sample	Description
font-weight: bold	Bold heading with style attribute
font-weight: 700	Bold heading with style attribute
	Bold tag
strong;strong;	Strong tag
text-decoration: underline	Underlined heading with style attribute
font-style: italic	Italic heading with style attribute
<i></i>	Italic tag
	Emphasized tag

G. Preliminary Analysis

The data extraction algorithm successfully collected the 10-K filings from 94 firms out of 100 firms listed in the QQQ ETF directly from the SEC EDGAR database. The six firms that are not included are foreign-based and instead issued 20-F filings. In total, the algorithm gathered 1397 filings, demonstrating its effectiveness for 10-K filings. However, some documents do not follow the standard regularity, including having an unstructured format or placing the risk factor section in other sections. The number of non-standardised documents is 63 out of the 1397 documents. Thus, we collected 1334 filings in total, excluding 63 fillings. In these cases, our algorithm was not able to accurately identify and collect the relevant information. Despite this fact, its design allows for adaptability to collect various types of SEC filings, making it a scalable tool for financial research that requires textual report analysis.

IV. METHODOLOGY

In this section, we introduced the supervised learning model for 10-K sentiment analysis and explained how the model functioned. Section 4.1, which was referred to as [1], demonstrated how the model generated a sentiment score of a 10-K filing by using the return at the time of the filing's publication as a label. In Section 4.2, we described in detail the model that used volatility as a label to estimate sentiment and its adaptation process. In Section 4.3, to represent the macroscopic sentiment trend of the technology sector, we introduced the Kalman filter and its process for removing noise in the context of 10-K sentiment analysis.

➤ A Sentiment Score Prediction Model

➤ Notation

To establish notation for the probabilistic sentiment model, we considered n as a set of 10-K filings and V as a vocabulary consisting of m words. The $i = 1, \dots, n$ represented the index of 10-K filings. It represented both the filing's publication date and the company to which it related. The word counts were recorded in a vector $d_i \in \mathbb{R}_+^m$, where $d_{i,j}$ represented the number of times word j occurred in 10-K filing i . We defined $D \in \mathbb{R}_+^{n \times m}$ as a document-term matrix, with $D = [d_1, \dots, d_n]$, representing word counts in each document. d_i was the i -th row of D , and the indices of columns were listed in the set S , which was a subset of vocabulary, i.e., $S \subseteq V$. $D_{[S],i}$ was the submatrix of the i -th filing. $d_{[S],i}$ was the word count vector in subset S for the i -th filing.

We labelled filing i with the associated time series variable y_i , either return or volatility, on the publication date of the filing. In this project, we assumed that each filing had a sentiment score, denoted by $p_i \in [0, 1]$. In the case of return fitting, a high score suggested that the filing had a predominantly positive tone in the report. However, it implies neither positive nor negative in the context of volatility fitting, as volatility signifies market uncertainty. Therefore, in the volatility setting, a high score indicates high market uncertainty and a low score suggests low market uncertainty.

➤ Model Assumptions

• Assumption 1:

We assumed that vocabulary V consisted of a set S of sentiment-charged words and a set N of neutral words, i.e., $V = S \cup N$. These sets were mutually exclusive, i.e. $S \cap N = \emptyset$. Furthermore, we posited the set of sentimentcharged words affected the tone of a filing, whereas the set of neutral words did not influence its tone, i.e. $d_{[S],i} \sqsupseteq p_i$ and $d_{[N],i} \perp p_i$. The sentiment word count was independent of the neural word count, i.e., $d_{[S],i} \perp d_{[N],i}$, implying that the model did not include the neutral words.

• Assumption 2:

We assumed that the sentiment-charged word counts $d_{[S],i}$ were produced by a mixture multinomial distribution:

$$d_{[S],i} \sim \text{Multinomial}(s_i, p_i O^+ + (1 - p_i) O^-) \quad [1]$$

Where s_i was the total count of sentiment-charged words in the i -th filing, i.e., $s_i = \sum_{w \in S} d_{w,i}$. This determined the scale of the multinomial. We then modelled $O \in \mathbb{R}_{+}^{|S| \times 2}$ matrix, representing the probabilities of individual word counts using a mixture model that incorporated positive and negative topics. The model included O_+ , a vector of $|S|$ non-negative elements with a unit ℓ^1 -norm, representing a probability distribution across words, such that $\sum_{w \in S} o_{+,w} = 1$, where $o_{+,w} \in O_+$ and $w \in |S|$. O_+ symbolised a 'positive sentiment topic' and represented the expected distribution of word frequencies in a filing with the highest possible positive sentiment, where the probability p_i equals 1. Similarly, O_- symbolised a 'negative sentiment topic' that represented the distribution of word frequencies in a filing with the most pronounced negative sentiment, where the probability p_i equals 0. The sentiment score p_i , where $0 < p_i < 1$, was a mixture coefficient of two sentiment topics.

• *Assumption 3:*

Finally, we assumed that the sentiment score fully encapsulated the information within a filing that affected the dependent variables, i.e., $y_i | p_i \perp\!\!\!\perp d_i$. This implied that sentiment score could primarily be used as a feature for return or volatility prediction models.

➤ *Model*

The model, incorporating these three assumptions, consisted of three steps for predicting the sentiment score of a 10-K filing. First, we extracted the set of sentiment-charged words, denoted as $\hat{S}^n$. Second, we estimated the probabilistic distribution matrix of positive-negative topic parameters over words, which is $\hat{O} = [\hat{O}_+, \hat{O}_-]$. Third, we predicted the sentiment scores $\hat{p}_i$ of a 10-K filing using penalised maximum likelihood estimation. Each step was described in detail in Section 4.1.3 to Section 4.1.4. The model overview can be found in.

• *Extracting Sentiment-Charged Words:*

In a 10-K filing, sentiment-neutral words were likely to predominate in terms of the number of words and total counts. This dominance tended to introduce noise and could make the extraction of sentiment-charged words from the entire vocabulary computationally burdensome, especially if sentiment-neutral words were not selectively excluded. Hence, in our model, we filtered out the set of sentiment-neutral words and focused solely on the subset of sentiment-charged words, estimating topic parameters for this subset alone. To achieve this, we utilised realised stock returns as a label (Check Return y_i , or R_t calculation) (Volatility was also used as a label, with the fitting process described in Section 4.2 in detail). The indicator function, represented by $\mathbb{I}\{a\}$, was defined as 1 when condition a was met, and 0 otherwise. Consequently, $\mathbb{I}\{d_{i,w}>0\}$ represented the presence of word w in the i -th filing, and $\mathbb{I}\{y_i>0\}$ denoted that the return variable associated with the i -th filing was positive. For each word $w \in V$, we could compute the frequency of word w in 10-K filings with positive returns relative to its overall frequency across all filings using the following equation:

$$f_w = \frac{\sum_{i=1}^n \mathbb{I}\{d_{w,i}>0\} * \mathbb{I}\{y_i>0\}}{\sum_{i=1}^n \mathbb{I}\{d_{w,i}>0\}} \quad [2]$$

This measure signified the word's sentiment tone to return values. For any given word w , if filings that contained it generally coincide with positive rather than negative returns, w was tagged with a positive sentiment. Conversely, if occurrences of w tended more often to align with negative returns, it was considered to have a negative return.

Subsequently, we assessed f_w against appropriate thresholds. If f_w was approximately 0.5, this suggested that the word was sentiment-neutral and belonged to the set N . To differentiate sentiment-charged words, we defined α_+ and α_- within the interval $(0, 0.5]$ to filter out sentiment-neutral words. A word was classified as a positive sentiment word if $f_w > 0.5 + \alpha_+$ and as a negative sentiment word if $f_w < 0.5 + \alpha_-$. We also defined a third threshold, $k \in \mathbb{N}$, which pertains to the count of word w across all filings. The threshold k acted

as a minimum frequency requirement to mitigate the impact of rare words, which could be sources of noise. For instance, a term appearing only once in all filings would result in an f_w of either 0 or 1, thus being noisy and unreliable itself. By applying the threshold k , we filtered out such noisy instances, requiring that a word appeared more than k times to be considered: $\sum_{i=1}^n \mathbb{I}\{d_{w,i}>0\} > k$. With these conditions, our extracted set $\hat{S}$ was defined by.

$$\hat{S} = (\{w : f_w > 0.5 + \alpha^+\} \cup \{w : f_w < 0.5 - \alpha^-\}) \cap \left\{ w : \sum_{i=1}^n \mathbb{I}\{d_{w,i}>0\} > k \right\} \quad [3]$$

• **Return or Calculation*

When using a return as our label, we used the publication time t over the days $t-1$ to $t+1$, as suggested in [1]. Using only the return at the publication time t can produce a noisy signal. A return may slowly respond to the content of 10-K filings, so the market may need the time to reflect the released 10-K filing of a firm to its stock price. Additionally, some contents of 10-K filings can already be reflected in its return as the 10-K filings, annually released, can contain some repetitive contents. To mitigate the noisy signal, we used open-close log returns R_t :

$$R_t = \log \left(\frac{P_{(t-1)c}}{P_{(t+1)o}} \right) \quad [4]$$

Where P refers to the equity price at a given time. The market open time at day t was denoted t_o , and the market close time at day t was denoted t_c . We used open-close return to make a symmetric alignment with the 10-K filing's published date.

• *Estimating Probabilistic Distribution of Topic Parameters:*

In this process, our goal was to estimate the topic parameters for a report. $O_i = [O_i^+, O_i^-]$ referred to a 1×2 vector, with each element corresponding to the estimated probabilistic distribution of positive topic or negative topic for a filing, respectively. To obtain $O = [O^+, O^-]$, we associated the sentiment expressed in a 10-K filing with stock returns or volatility (Fitting on volatility will be explained on Section 4.2). This approach comes from the assumption that stock returns or volatility at the publication date of a 10-K filing represented the sentiment for it. Note that we did not have direct sentiment scores for the filings, thus we estimated the sentiment proxy by using the standardised ranks of stock returns as a label in the equation:

$$p_i^{\frac{rank(y_i)}{n}} \quad [5]$$

Where the expression defined $\tilde{p}_i$, the estimated sentiment proxy for the i -th filing, as the rank of y_i , i.e., a stock return or volatility (Fitting on volatility will be explained on

Section 4.2), divided by n , the total number of filings. The rank of y_i was sorted in ascending order.

With these estimated sentiment proxies, we had a matrix $\hat{W}$ containing two rows for each filing: one for the estimated positive sentiment proxy $\tilde{p}_i$ and one for the estimated negative sentiment proxy $1 - \tilde{p}_i$ in the equation:

$$\hat{W} = \begin{bmatrix} \tilde{p}_1 & \tilde{p}_2 & \dots & \tilde{p}_n \\ 1 - \tilde{p}_1 & 1 - \tilde{p}_2 & \dots & 1 - \tilde{p}_n \end{bmatrix} \quad [6]$$

Subsequently, we adjusted the counts of sentiment-charged words by dividing each count by the total sentiment-charged word count for the filing like in the equation:

$$\hat{h}_i = \frac{d_{[\hat{S}],i}}{\hat{s}_i}, \text{ where } \hat{s}_i = \sum_{w \in \hat{S}} d_{w,i}, \quad [7]$$

Where $\hat{h}_i$ was the counts of sentiment-charged words. This was defined as the ratio of $d_{[\hat{S}],i}$, each word count within the subset $[\hat{S}]$ of the i -th filing, to $\hat{s}_i$, which was the total sentiment-charged word count for the filing. These relative term counts were collected in a matrix $\hat{H} = [\hat{h}_1, \dots, \hat{h}_n]$.

In this process, we aimed to estimate a matrix $\hat{O}$, which contained parameters that referred to the probability distribution of positive and negative sentiments. Then, we could estimate $\hat{O}$ with a regression, using matrix $\hat{H}$ as the predicted outcome and matrix $\hat{W}$ as the predictor like in the equation:

$$\hat{O} = \hat{H} \hat{W}^T (\hat{W} \hat{W}^T)^{-1} \quad [8]$$

In the final step, to ensure the estimates correspond to a probability distribution, we corrected any negative entries by resetting them to zero and then re-normalised each column so that their totals were equal to one. This process produced a revised matrix, but to simplify notation, we reused $\hat{O}$ for the resulting matrix, and we labelled its first and second columns as $\hat{O}_+$ and $\hat{O}_-$, respectively.

➤ Scoring New Filing

Now that we constructed estimators $\hat{S}$ and $\hat{O}$. Also, from the mixed multinomial distribution in Assumption 4.1.2.2, we could estimate the filing's count vector, which is $d_{[S]}$. Given all estimates $\hat{S}$, $\hat{O}$, and $d_{[S]}$, we could estimate the best probable sentiment score p by Maximum Likelihood Estimation (MLE):

$$\hat{p} = \arg \max_{p \in [0,1]} \left\{ \hat{s}^{-1} \sum_{w \in \hat{S}} d_{i,w} \log \left(p \hat{O}_{+,w} + (1-p) \hat{O}_{-,w} \right) + \lambda \log(p(1-p)) \right\} \quad [9]$$

Where $\hat{s}$ represented the total count of words from the set $\hat{S}$ in the new filing, while $d_{i,w}$, $\hat{O}_{+,w}$, and $\hat{O}_{-,w}$ referred to

the w -th elements of their respective vectors, and λ was a positive constant used to adjust the model. In the MLE, $\lambda \log(p(1-p))$ was a penalty term to avoid overfitting when there were reports with few sentiment-charged words. For instance, if the filing had a limited number of positive sentiment-charged words without negative words, the model believed the filing had a positive tone even if the filing just only contained a few positive words. The term served as a regularising factor that pushed the estimated sentiments toward 0.5, indicative of a neutral sentiment. In other words, the penalty terms nudged the model to be more conservative when scoring new filings.

➤ Volatility Label

Section 4.1 introduced the sentiment prediction model with the firm return as y_i , as in [1]. Notably, [1] suggested the model is universally adaptable. However, its core assumptions and methodologies might not suit every forecasting objective such as using volatility as a label. While the current model with a return label fits into a binary or discrete framework, volatility does not naturally fit into a binary or discrete framework. Due to that, we should do adaptation to use volatility as a label for the model.

In the current narrative of returns prediction, we could employ returns as a label corresponding to binary sentiment topics, i.e., positive and negative, because one either made a loss or a profit. However, this classification was not clear in the context of volatility. The adaptation for volatility was to set a threshold θ and we labelled all values above this θ point as high volatility and those below it as low volatility. We then replaced 0 by the θ in Equation 4.2 and the value 0.5 by quantile q in Equation 4.3. Note that there were no right threshold or quantile, but we should keep in mind that our choice of hyper-parameter would impact the outcome of the model.

Volatility is, in nature, an asymmetric variable; thus, getting a ranking of the volatility like in Equation 4.4 will lose substantial informational value to estimate the sentiment score. Subsequently, normalising the volatility can be an appropriate alternative as it preserves asymmetry:

$$p_i^* = \frac{y_i - \min(y)}{\max(y) - \min(y)} \quad [10]$$

Where $\hat{p}_i^*$ is greater than or equal to 0 and less than or equal to 1. However, the normalised volatility itself was not able to catch the outlier of the market movement despite making it symmetric around zero, leading to poor predictive performance. Thus, we used volatility values over multiple days for the robust labels. In practice, we averaged the volatility over three days. We used the intra-daily range among other standard volatility proxies such as squared returns, or the realised volatility as it is a less noisy volatility proxy, leading to less distortion [66], [67]. The intra-daily log range is defined as:

$$RG_t = \max_{\tau} \log P_{\tau} - \min_{\tau} \log P_{\tau}, \tau[to, tc] \quad [11]$$

- The Volatility Proxy V_t was then Computed Like this:

$$\tilde{V}_t = \frac{RG_t^2}{4\log(2)}, \quad [12]$$

Where the intra-daily log range was squared and divided by the adjustment factor, $\frac{1}{4\log(2)}$. The adjustment factor is used to correct potential bias that may occur in the data generation process when we assume that the process follows Brownian motion with drift Parkinson, 1980. Then, we attained the averaged volatility over three days for a more reliable label:

$$V_t = \frac{1}{3} (\tilde{V}_{t-1} + \tilde{V}_t + \tilde{V}_{t+1}) \quad [13]$$

Which is the common approach to calculating multi-day aggregation of daily volatility proxies [68].

➤ Kalman Filter

In this paper, we generated sentiment metrics of both the technology sector and a single firm (e.g. Nvidia) with 10-K filings over a certain time. But, the estimation of the industry's time-varying sentiment measures should be very noisy. To obtain robust sentiment features on time series, we adapted the Kalman filter to smooth the time-varying noisy signals. In the context of smoothing sentiments of 10-K filings across the industry, we referred to [69] to adapt the Kalman filter to our context. Following the adapted Kalman filter in our context of work.

$$\mu_{t+1} = \mu_t + \eta_t, \eta_t \sim N(0, \sigma^2 \eta) \quad [14]$$

$$v_t = \mu_t + \epsilon_t, \epsilon_t \sim N(0, \sigma^2 \epsilon) \quad [15]$$

Firstly, we assumed that there were the observed sentiment score v_t and an unobserved sentiment score - the state variable μ_t in the Kalman filter framework. We extracted the unobserved sentiment score from the publication-date-based sentiment scores via Equation 4.14 and Equation 4.15. Equation 4.14 referred to the prediction model and Equation 4.15 referred to the correction model. The prediction model updated the unobserved state, and the correction model represented the observed sentiment cores that were affected by the state variable. The Kalman filter worked recursively through prediction and correction. It predicted the unobserved state estimates from the previous weighted variable. You can find a more detailed explanation for this Kalman filter model in [70]. In our study, we employed the Kalman smoother, instead of the Kalman filter, as the former one was likely to be more proper to show a retrospective trend to analyse the industry-level sentiments. Both the Kalman filter and the Kalman smoother are very similar in terms of estimating the unobserved sentiment score. However, the Kalman filter estimates the unobserved sentiment score at time t given observations up to and including time t , whereas the Kalman smoother estimates the hidden sentiment score based on all observations for $t = 1, \dots, T$. In this paper, we called the Kalman smoother as the Kalman filter for convenience.

To adapt the Kalman filter in our context, our sentiment scores should be converted to time series data, because there were multiple sentiment scores on the same publication data of 10-K filings. In order to convert it to time series data, we aggregated sentiment scores on the same date over all firms like $\tilde{p}_t = \frac{1}{|A_t|} \sum_{i \in A_t} \tilde{p}_i$, where A_t referred to the set of 10-K filings published on day t . Subsequently, v_t was substituted by $\tilde{p}_t$ in Equation 4.9 and we obtained the filtered industry-level sentiment scores $\tilde{p}_t$ after applying the Kalman filter. In other words, this approach implied that all 10-K filings published on the same day were treated equally. This approach may enhance the model's ability to estimate sentiment scores more accurately. However, it comes with the trade-off of transforming the score into an aggregate metric of all 10-K publications within a given day.

V. EXPERIMENT RESULTS

In this paper, we aimed to generate various sentiment scores from our suggested sentiment prediction model. To achieve that, we generated sentiment scores at three stakeholder levels: Sector level, Portfolio level, and Company level. At the sector level, we generated scores from the model where it aggregated all filings of companies from the technology sector. At the portfolio level, we generated scores for a specific set of companies. In the practice of our study, we aggregated filings of the top 10 firms listed in QQQ. Finally, at the company level, we aggregated all filings of a firm to generate sentiment scores. We generated Nvidia's sentiment metrics in our study. All models for each level commonly used either the entire 10-K or only the risk factor section, by labelling with return or volatility. That is each level generated four types of sentiment scores. Three levels were multiplied by four types of sentiment scores. In total, we generated 12 types of sentiment scores (check Contribution). Furthermore, for analysis of both a sector level and portfolio level, we applied the Kalman filter to mitigate noise.

In this chapter, we showed the statistics overview of experimental results to introduce what we found. Then, we evaluated sentiment scores we generated at three stakeholder levels through correlation analysis and qualitative analysis with the most influential words.

➤ Descriptive Statistics

This section shows descriptive statistics of sentiment scores we generated at the three different stakeholder levels. We show the number of observations for each model and the number of sentiment words used for each model. We also calculate the mean, Standard Deviation (SD), and maximum and minimum sentiment scores for each model.

➤ Evaluation Methods

➤ Correlation Analysis

We used the Pearson correlation to evaluate our estimated sentiment scores. The Pearson correlation coefficient could have two assumptions [71]. The first assumption is two variables are continuous. Secondly, two variables are assumed to have a linear relation. Based on these

assumptions, we could analyse our sentiment scores with the Pearson correlation. We used four variables (i.e. PRET, PVOL, PLM, and Stock) for correlation analysis. The correlation of variable pairs was calculated through the Pearson correlation formula in Appendix E, and their linear correlation could be checked through a significance test [72]. In a significance test, 0.05 is conventionally used as a significance level. However, a long list of scientists recently proposed a new p-value threshold (i.e. 0.005) to improve the reproducibility as the conventionally used threshold leads to a high rate of considerable false positives, even without taking into account any issues related to the experiment, procedure, or documentation [73]. In the practice of our study, we significantly increased our experiment's reliability by setting both 0.05 and 0.005.

➤ Qualitative Analysis with Most Influential Words

In the process of predicting sentiment scores for all models we have introduced so far, we could extract the top 15 influential words from $\tilde{p}^{RET}$, and $\tilde{p}^{VOL}$ for each model. The detailed extraction process can be found in Equation 4.3. With the extracted set, we could infer a topic for the analysis of three stakeholder levels.

When interpreting the set of the extracted words, the interpretations might be interpreted in various ways. The objective of the word extractions was to offer the most influential keywords used for financial analysis. We assumed the value of the words would be synergised with domain knowledge. Also, using generative AI could offer a good reference to our word set for in-depth and insightful analysis.

➤ Evaluation Results

➤ Sector Level (i.e. Technology Sector)

In a technology sector analysis, we computed the sector sentiment scores with all firms' filing, additionally considering the entire 10-K filing or only the Item 1A section, respectively. To improve readability, we introduced a table at the top right for all the score graph figures to identify easily which models you are referring to. Moreover, we employed a Kalman filter to represent a reliable sentiment trend for a sector analysis [70]. An industry-level analysis generated a considerable amount of signal noise. We could control these noisy signals by using the Kalman filter.

• Sector Correlation Analysis:

Table 2 Sentiment token Statistics with 10-K and Risk Factor

	# of Total Tokens (# of Sentiment Tokens)	Mean	SD	Max/Min
$\tilde{p}^{RET}$ with 10-K (Figure X)	15,863 (8,725)	0.48	0.28	0.92/0.08
$\tilde{p}^{VOL}$ with 10-K (Figure X)	15,863 (9,518)	0.43	0.24	0.92/0.08
$\tilde{p}^{RET}$ with Risk factor (Figure X)	9,030 (5,147)	0.46	0.27	0.92/0.08
$\tilde{p}^{VOL}$ with Risk factor (Figure X)	9,030 (3,792)	0.42	0.25	0.92/0.08

Table 3 Sentiment token Statistics with 10-K and Risk Factor

	# of Total Tokens (# of Sentiment Tokens)	Mean	SD	Max/Min
$\tilde{p}^{RET}$ with 10-K	7,040 (6,434)	0.42	0.30	0.92/0.08
$\tilde{p}^{VOL}$ with 10-K	7,040 (6,985)	0.43	0.33	0.89/0.08
$\tilde{p}^{RET}$ with Risk factor	4,265 (4,052)	0.42	0.29	0.92/0.08
$\tilde{p}^{VOL}$ with Risk factor	4,265 (4,005)	0.44	0.31	0.88/0.08

Table 4 Sentiment token Statistics with 10-K and Risk Factor Signal Noise.

	# of Total Tokens (# of Sentiment Tokens)	Mean	SD	Max/Min
$\tilde{p}^{RET}$ with 10-K	3,861 (3,861)	0.49	0.32	0.89/0.13
$\tilde{p}^{VOL}$ with 10-K	3,861 (3,861)	0.39	0.21	0.86/0.13
$\tilde{p}^{RET}$ with Risk factor	2,201 (2,201)	0.49	0.32	0.90/0.11
$\tilde{p}^{VOL}$ with Risk factor	2,201 (2,201)	0.41	0.18	0.80/0.17

➤ The Sector Sentiment Model with 10-K Filing (Figure)

Figure showed the predicted sentiment scores for the technology sector. These scores were estimated with almost all firms' 10-K filings listed in the QQQ (i.e. 94 firms out of 100), and we used the entire 10-K filings to calculate the scores. Note that the sentiment scores labelled with return and volatility in Figure are filtered. Upon reviewing graphs in and , it was evident that both unfiltered return sentiment and volatility sentiment contained significant noise. Therefore, filtered versions were chosen to more reliably depict the sector's trend. Other models, which will be introduced in the following parts, also used the filtered one.

We calculated two types of average loss for the same windows: one between the estimated sentiment score $\tilde{p}^{RET}$ and the normalized rank of return, and the other between the estimated sentiment score $\tilde{p}^{VOL}$ and the normalized rank of volatility. Furthermore, we calculated how well our sentiment analysis models can predict the sentiment of the financial markets based on return or volatility. For this model (Figure), the loss of the model $\tilde{p}^{RET}$ was 0.25, and the accuracy rate was 78%. It showed $\tilde{p}^{RET}$ was strongly well predicted compared to the QQQ return given a window. It meant the $\tilde{p}^{RET}$ represented the sentiment of the technology sector.

Additionally, the $\tilde{p}^{VOL}$ sentiment of this model showed stronger prediction accuracy. The $\tilde{p}^{VOL}$ model performed a 92% accuracy rate with 0.22 loss.

To evaluate our model critically, we selected the LM as our benchmark as it was created for 10-K sentiment assessment. In Figure , the $\tilde{p}^{LM}$ score showed a steady decrease until 2018, followed by a significant and gradual decline, with intermittent fluctuations due to noise. This trend occurred despite the technology sector's rapid growth from 2018 until just before the COVID-19 pandemic, indicating a generally negative direction in the model's movements. This was because the predominance of negative over positive vocabulary in the LM dictionary (2,355 negatives vs 354 positives out of 86,531 ~ words) suggested a bias towards negative sentiment in the $\tilde{p}^{LM}$ score, as approximately 85% of the relevant vocabulary is negative. The rest, 83,822 words, were neutral and excluded from the sentiment analysis. This imbalance likely caused the $\tilde{p}^{LM}$ score to trend negatively. Additionally, the sector sentiment prediction model took 37 seconds to execute.

These tables G.1, and Table 5.1) represented Pearson correlation coefficients between the sentiment estimates including QQQ's stock price, both filtered and unfiltered. Upon analyzing G.1, we found that, with the exception of the VOL and LM pair, there was no linear correlation between any pairs of unfiltered sentiment scores. However, analysis of filtered sentiment scores (Table 5.1) revealed more

significant correlations compared to unfiltered scores. All other pairs, with the exception of the LM and Stock, exhibited weak correlations. Specifically, an r value of -.245 between $\tilde{p}^{RET}$ and $\tilde{p}^{VOL}$ indicated a weak negative correlation, suggesting that positive sentiment in RET generally corresponded to negative sentiment in VOL, and vice versa. Furthermore, an r value of +.296 between $\tilde{p}^{RET}$ and Stock implies a weak positive correlation, indicating that positive sector sentiment in RET was associated with rising QQQ stock prices, and vice versa. In Figure , despite oscillated fluctuations due to noise from the sector, $\tilde{p}^{RET}$ showed a slow increase during the observed window. Notably, the sentiment trend in 2023 for $\tilde{p}^{RET}$ closely mirrored the QQQ stock performance in the same year. The case of $\tilde{p}^{VOL}$ and Stock ($r=+.121$) showed also a positive correlation, but weaker. Moreover, both $r_{\tilde{p}^{RET}, \tilde{p}^{LM}} (= -.359)$ and $r_{\tilde{p}^{VOL}, \tilde{p}^{LM}} (= -.173)$ showed a weak negative correlation. , all pairs exhibited low p -values, with those in Table being significantly lower. This suggests that the sample results—specifically, the correlation coefficients—provide sufficient evidence to reject the null hypothesis from the entire population [74]. In our case, the null hypothesis was that there was no correlation between the pair, whereas the alternative hypothesis was that there was a correlation between them. However, the lower p -value of all pairs does not measure the probability that the alternative hypothesis is true. The p value is not an absolute index for arguing that a hypothesis is true. Instead, the lower p -value shows our experiments have statistical significance [75], [76], [77].

Table 5 Filtered, A Sector Sentiment Correlation with 10-K Filing

$\tilde{p}^i, \tilde{p}^j$	RET	VOL	LM	Stock
RET	1	1	1	1
VOL LM	**-.245	**-.173	**-.0910	
Stock	**-.0359	**0.121		
	**0.296			

- Note: * p -value < 0.05, ** p -value < 0.005

➤ The Sector Sentiment Model with Only-Risk-Factor

The sentiment scores in were predicted based on the risk factor section. Note again that only showed the filtered one (the unfiltered can be found at.

The PRET model accuracy was 73% with the 0.25 loss. The risk factor $\tilde{p}^{RET}$ model accuracy is 5% lower than the entire 10-K $\tilde{p}^{RET}$ model, which is 78%. Moreover, the risk factor $\tilde{p}^{VOL}$ model showed a stronger model performance. It performed 90% accuracy with a 0.22 loss. Comparing the 10-K sector ~ PLM and the risk factor $\tilde{p}^{LM}$, the latter is generally lower than the former. This pattern suggested that the $\tilde{p}^{LM}$ score, influenced by an LM dictionary dominated by negative words and a risk factor section pessimistically toned, inherently leaned towards a negative sentiment [25], [78].

When comparing the training data for both models (Figure , and), it's noted that the aggregated word count from all QQQ firms' filings was 15,863 post-preprocessing, and the count from just the risk factor sections of these filings was 9,030 after undergoing the same preprocessing. This observation highlighted that the risk

factor sections, despite typically constituting only 10% to 15% of the entire filing, contained a significant portion of the relevant words to the model's sentiments. It showed that a risk factor section could offer informational value for textual analysis in the finance domain. Also, the minor differences in accuracy between the models could suggest that the risk factor section provided a valuable textual context in predicting sentiment scores for both $\tilde{p}^{RET}$ and $\tilde{p}^{VOL}$. However, upon comparing both models, it was observed that all scores from the model analyzing risk factor sections were lower than those from the comprehensive model. It indicated that the tone of a risk factor section is pessimistic/negative as other studies also empirically found [25], [78]. The training time for this model, additionally, was also reduced by 40% (i.e. 23s).

Table showed the Pearson correlation coefficients of the filtered sentiment sector model with only the risk factor. Compared to the sector model with the entire 10-K filing, shown in Table 5.1, the sector risk factor model showed, in general, lower r values for all pairs, and some pairs showed altered sign such as the pair of $\tilde{p}^{RET}$ and $\tilde{p}^{LM}$, the pair of $\tilde{p}^{RET}$ and Stock, and the pair of $\tilde{p}^{VOL}$ and $\tilde{p}^{LM}$. Except for the pair of $\tilde{p}^{RET}$ and $\tilde{p}^{LM}$ which did not show statistical significance,

the $r = -0.376$ from P^{RET} and Stock showed a weak negative correlation in the risk factor model. The more positive tone the sector has, the lower the price the sector has. This outcome was in contrast to that of the model (Table 5.1) and also our intuition where a good sentiment is related to a higher return. Other pairs did not show a meaningful correlation as they were close to 0.

- Sector Most Influential Words:

➤ *The Sector Sentiment Model with 10-K Filing (Figure)*

The following words were the most influential words in P^{RET} sentiment score prediction, positively or negatively, respectively.

From the positive influential words in S_+^{RET} , we could categorise a few themes that could impact the technology sector's positive sentiment tone. The first theme can be 'Innovation and Development'. The words such as musk, maximiser, searcher, multifaceted, and incisive could be interpreted as innovation-relevant vocabulary. Maximiser and searcher could refer to a figure to lead technology innovation or development. Interestingly, the word, musk, seemed to indicate "Elon Musk", who is one of the figures to represent innovation under the fourth industrial revolution era. Moreover, the words (i.e. multifaceted, incisive) could indicate a capability or capacity (or both) required for innovation. In Figure , the technology sector has been gradually and significantly developed during the given period. This remarkable development could be associated with innovations. In this context, words such as musk, torque, and bolivia were additionally associated with innovation, specifically indicating an electric car or robotics. Elon musk, CEO of Tesla, leads the American firm known for its electric vehicles and robotics innovations. The torque word could be related to an electric vehicle or robot. Bolivia is home to the world's largest lithium deposits [79]. Lithium is a key component of electric car batteries. Furthermore, the words, including gilt, div (we interpreted as dividend), memoranda, and stance, could be interpreted as finance-relevant terms. Good financial health could be an important factor in economic growth [80]. Hence, these finance terms represent 'Financial Health'. The emphasis on improving the financial health of each tech firm might impact their stock positively, leading to a rise in the sector return.

We could infer a theme from the negative impactful words in S_-^{RET} . 'Technological Difficulty' could be a theme that makes the sector tone negatively. The words, such as tango, transformer, scalar, infer, and tera could represent a firm's technological difficulties. For instance, Tango was Google's Augmented Reality deprecated project due to its technological issue [81], [82]. Actually, Google's Gemini glitches cost Google 90 billion dollars in stock loss in a single day. Furthermore, transformer is a significant natural language process (NLP) model architecture for generative AI. scalar, an element used to define a vector space in machine learning, can refer to the number of parameters for an AI model, which is associated with model performance. Tera referred to a terabyte, representing a large dataset.

Two themes could be inferred from the extracted words to increase market uncertainty. The first theme could be 'COVID 19' from words, including lancet and forearm. lancet is a major medical journal. forearm is a body part relevant to a vaccine. Both could be related to the COVID-19. We secondly could infer 'Innovation and Growth'. The words contained fang, bully, amenable, killing, granularity, koa, and renter. fang referred to FANG, which is an acronym for major technology firms in the US. They were known for their volatile stock prices. Words like granularity, koa, and renter could be associated with a tech firm's development. The word renter indicated the concept of the sharing economy. A representative example of this sharing economy could be cloud service. koa represented a web framework for Node.js, which highlighted the back-end technology.

In fact, we could interpret the extracted words in various ways. One theme we could infer was 'Environmental Issues' from words, including swine, laryngoscope, farming, moss, and subway. swine could refer to swine flu. laryngoscope could be related to COVID-19. This flu and the pandemic could increase or decrease the volatility of the technology sector. Moreover, words such as farming, moss, and subway could be associated with modernisation or urbanisation for sustainable development. Or, they could be related to climate change. A few studies supported that environmental issues were correlated with volatility [83], [84].

➤ *The Sector Sentiment Model with Only-Risk-Factor*

These words affected the sector sentiment positively. The word set showed a similar theme compared to the previous one (check $\tilde{S}_+^{RET}$ of the 10-K sector model in here). The first theme, 'Innovation and Development,' could be inferred from the words, including artificially, mechanic, excellence, and explorer. artificially represents Artificial Intelligence (AI), which is a cornerstone of the fourth industrial revolution. Explorer could be interpreted in the same context of maximizer, searcher in the previous one (check $\tilde{S}_+^{RET}$ of the 10-K sector model in here). The second theme was 'Financial Health'. Words, such as roe, stance, escheat, wrongfully, and ruble could indicate the theme. roe referred to ROE (i.e. Return on Equity). roe, stance, and wrongfully could symbolise firms' financial condition. Escheat could convey asset retention. Also, ruble could indicate global market interactions. The ruble is the currency of the Russian Federation.

There are a few words that could be interpreted as similar to the 'Technological difficulty' theme mentioned previously (check $\tilde{S}_-^{RET}$ of the 10-K sector model in here). The words included biogen, surviving, kinase, ramification, and chemotherapy. These words could specifically indicate biotechnology. biogen, kinase, and chemotherapy are terminology used in biotechnology. Ramification and surviving could be interpreted as a difficulty to develop biotechnology. Interestingly, this $\tilde{S}_-^{RET}$ model trained with the risk factor section extracted the word covid explicitly, which rapidly dropped the sector sentiment with the reduction of stock price.

A few words, such as pressing, prominently, and endocrinologist could show the COVID-19 relevance. Also, Spoof appeared again. It could symbolise cybersecurity, including impropriety, programmatic, and erroneously. Interestingly, clothing and polymeric were extracted. They could indicate wearable technology in health care.

➤ *K. Portfolio Level (i.e. Top10 Firms Portfolio)*

We generated sentiment scores at a portfolio level. The portfolio consists of the top 10 firms equally listed in the QQQ fund. Top 10 firms denoted the top 10 most invested companies from 1 to 10 in QQQ, as the 50 percentile of the QQQ fund. Note again that the Kalman filter also was applied here.

• Portfolio Correlation Analysis:

➤ *The Top10 Sentiment Model with 10-K Filing*

Indicated the predicted sentiment score for the top10 portfolio. These scores were estimated based on the top 10 firms' entire 10-K filing for a given window from 2006 to 2023. The $\hat{p}^{RET}$ model performed a 76% accuracy rate with 0.22 loss. The previously mentioned sector models had 78% and 72%, respectively. The $\hat{p}^{VOL}$ model decreased its accuracy rate to 78% with the 0.20 loss (: 92%, : 90%), but it still predicted the sector's sentiment labelled with volatility well. The training time was slightly faster, reduced to 13 seconds, as we only used the 10-K filings of the top 10 firms.

Compared to the sector's sentiment score (Figure), the top10 10-K model () showed a significantly far clearer trend while noisy signals were removed again through the Kalman filter (the unfiltered model can be found at). We found extremely interesting a few phenomena in this model. Firstly, we observed the $\hat{p}^{RET}$ strongly mirrored the top 10 stock price movement. We could see that the $\hat{p}^{RET}$ had a steady and moderate increase from 2006 to around 2016. Then, beginning around 2018, $\hat{p}^{RET}$ showed a rapidly sharp rise, but even after the outbreak of COVID-19. Very similarly, the Top 10 stock rose sharply over a few years in early 2018 and showed a steep drop from 2021 to the beginning of 2023. Then, it increased sharply again. The steep drop from 2021 to 2023 could be stemmed from the COVID-19 pandemic. Several studies argued that quantitative easing(QE) seemed to be significantly and positively related to a higher recovery of the USA's stock market. This could explain the QQQ's sharp increase after the outbreak of COVID 19 [85], [86], [87], [88], [89], [90], [91], [92]. However, the post-COVID-19 surge in sentiment suggested the $\hat{p}^{RET}$ model did not accurately gauge event significance, a factor likely considered in 10-K filings. This limitation could arise from managers' ignorance about an event's impact at filing time or the model's inability to quantify how much sentiment is driven by the event. Essentially, the model treated significant terms like "COVID-19" or "restrictions" merely as text, without evaluating their actual importance or sentiment contribution.

Secondly, the $\hat{p}^{VOL}$ model showed interesting points. The $\hat{p}^{VOL}$ model revealed trends in market volatility, with a sharp increase until around 2010, followed by a decline until

around 2016, likely influenced by the 2007-2008 financial crisis and subsequently moderated by quantitative easing (QE) policies. Economists credited the Federal Reserve's fiscal stimulus for mitigating the crisis's effects [85], [86], [93], [88], [87], [89], [94]. A similar volatility movement happened again during the outbreak of COVID-19. Furthermore, this $\hat{p}^{VOL}$ model indicated a decalcomania-like movement with $\hat{p}^{LM}$, where increases in volatility seemed to be associated with decreases in $\hat{p}^{LM}$ and vice versa. Despite this, $\hat{p}^{LM}$ failed to accurately reflect significant market events, such as the 2008 financial crisis or the COVID-19 pandemic, showing no substantial change in its trend from 2018 to 2024. This lack of responsiveness suggested the $\hat{p}^{LM}$ model's inability to dynamically mirror the market, likely due to its training dataset being heavily skewed towards negative terms.

Table 5 showed the Pearson correlation coefficients of the filtered top10 sentiment model with the entire 10-K filings. All r values indicated a strong correlation with rigorous statistical significance. Compared to Figure , the top 10 firms revealed clearer correlations by removing the noise associated with the other 90 firms. A strong positive correlation of $r = +.905$ between $\hat{p}^{RET}$ and $\hat{p}^{VOL}$ indicated that higher sector positivity was linked with greater uncertainty. Additionally, $r_{\hat{p}^{RET}, Stock} = +.956$ showed that increases in sector stock prices correlated with more positive sentiment, and similarly, a strong positive correlation existed between $\hat{p}^{VOL}$ and stock prices, suggesting that stock prices rose with increased market uncertainty. Strong negative correlations of $r_{\hat{p}^{RET}, \hat{p}^{LM}} = -.922$ and $r_{\hat{p}^{VOL}, \hat{p}^{LM}} = -.903$ indicated that decreases in $\hat{p}^{LM}$ were associated with the increased market sentiment and uncertainty, and vice versa. Also, $\hat{p}^{LM}$ had a strong negative correlation with Stock.

➤ *The Top10 Sentiment Model with Only-Risk-Factor*

Showed the top 10 firms' estimated sentiment scores by training only the risk factor section. Compared to the 10 K top10 model (), this model() performance was decreased slightly to 71% with the 0.23 loss ($\hat{p}^{RET}$) and 74% with the 0.21 loss ($\hat{p}^{VOL}$). This could be because of the smaller size of the vocabulary of the risk factor than the entire 10-K filings. Then, the training time also was reduced to 8 seconds. Due to the same reason, we observed that the $\hat{p}^{RET}$ model steadily increased. However, this model did not reflect the market situation well. This model gradually increased during the technology industry boom time and very slightly decreased after the outbreak of COVID-19. Also, the $\hat{p}^{VOL}$ model showed the market uncertainty steadily fell during the same period. The $\hat{p}^{RET}$ model dynamically, albeit minimally, reflected market changes, while the $\hat{p}^{VOL}$ model struggled due to its focus on the risk factor section dominated by negative tones. When using the model trained on risk factors, we must remember it might offer a pessimistic market view. Furthermore, the $\hat{p}^{LM}$ model indicated a slow, moderate shift towards negativity, showing it lacks dynamic market responsiveness, similar to patterns seen in the earlier $\hat{p}^{LM}$ model(). This negative bias distorted the analysis of the portfolio level.

When examining the Pearson correlation Table, the correlations including $\hat{p}^{VOL}$ showed differently compared to the 10-K top 10 sector model (Table F.2). We observed that all correlation, such as $r_{pVOL, pRET}$ ($=+.145$), $r_{pVOL, pLM}$ ($=.349$), and $r_{pVOL, Stock}$ ($=+.185$), became to have a weak linear correlation. The $r_{pVOL, pRET}$ and the r_{even} were close to zero, implying that they lost their correlation. We could assume that the $\hat{p}^{VOL}$ score was not calculated correctly in the model trained with the risk factor section.

- Portfolio Most Influential Words:

➤ *The Top10 Sentiment Model with 10-K Filing ()*

The word sets extracted from the top10 10-K model seemed to indicate the 'Electric cars'. The 10-K model includes words like roadster, musk, supercharger, cluster, roof, sedan, and dangerous. These words explicitly indicated 'Electric cars'. Notably, the same theme commonly appeared in the following top10 risk-factor model. From the observation, we could infer that 'Electric cars' have a strong correlation with the technology sector return.

The 10-K word set to affect return negatively seemed to reveal generic words, hindering catching a theme. However, the following S_{-RET} in the top10 risk factor model seemed to highlight a distinct theme.

The positive word set had a common theme, 'Electric cars', with S_{+RET} in the top10 10-K model. The words indicating the topic include autonomous, neural, pilot ford, berlin, and race. autonomous, neural, and pilot could refer to self-driving, while ford, berline, and race could represent an automobile industry. From this observation, 'Electric cars' also was the topic to increase the technology market's uncertainty as well as the market's positivity.

The words from S_{-VOL} such as restrain, resilience, tester, and insecure could indicate a topic, 'Stability and Resilience'. The continuous efforts (tester) of a tech firm to overcome technological difficulties (insecure, restrain) could improve firm stock's stability (resilience), contributing to lower volatility.

➤ *The Top10 Sentiment Model with Only-Risk-Factor*

From the words such as battery, solar, musk, screen, roadster, and motor. This set could indicate 'Electric Cars'. Interestingly, this topic appeared in the top10 10-K model to increase return and volatility, suggesting investors should carefully pay attention to the electric car industry. It could be the next alpha(α) signal to beat the market.

The theme "Supply chain risk", for instance, could be revealed from words like subcontractor, entrust, dependence, and nationwide. Several studies showed the risk of supply chains for high-technology industries severely increased due to geopolitical disruptions [95], [96], [97], [98]. subcontractor, entrust, dependence could indicate firms' dependency (dependence) on their productions and services. nationwide could represent supply chain risk occurring globally.

The following word set extracted from the risk factor increased or decreased the portfolio uncertainty.

The S_{+VOL} could indicate a theme, 'Innovation and development' to increase market uncertainty. circuit, generating could symbolise generative AI, and its semiconductor chip, GPU. virtual could refer to the metaverse. transact could indicate transaction technology like blockchain or digital payment systems. Several studies found that firms' R&D investment in these fields, being represented by innovation, has a strong positive relationship with volatility [99], [100], [101]. The topic, 'Innovation', from the model supported the argument of these studies.

The S_{-VOL} could refer to a theme, 'Infrastructure and Services', to decrease market uncertainty. depot, host could indicate a data centre and its hosting service. Also, sea could be related to undersea cables for data transmission. club, membership can represent a social infrastructure (or network) for tech operations. Several studies found digital and traditional infrastructure investments are key components to stabilising and driving forward economic growth as well as reducing market uncertainty [102], [103], [104]. These studies support the model's argument that 'Infrastructure and Services' could decrease market uncertainty.

➤ *Company Level (i.e. Nvidia)*

- Nvidia Correlation Analysis:

➤ *Nvidia Sentiment Model with 10-K filings*

Represented Nvidia's predicted sentiment scores trained with the entire 10-K filing. The $\hat{p}^{RET}$ model performed a 75% with 0.19 loss and the $\hat{p}^{VOL}$ showed a 69% accuracy with 0.25 loss. Our sentiment prediction model predicted a single firm's sentiment fairly well. It took 7 seconds.

The $\hat{p}^{RET}$ model seemed to reveal a periodic pattern of the sentiment tone. Its sentiment tone dramatically became positive from 2006 to 2011 and vastly became negative from 2014 to 2020. In the same period, Nvidia's stock price remained the same. Then, it sharply and rapidly became positive. The stock was hugely impacted by the outbreak of COVID-19 and rebounded due to the US government's aggressive QE policy and Nvidia's dominant position in the GPU market in the fourth industrial revolution. We assumed that Nvidia's periodic rise and fall in their sentiment might have stemmed from the firm's internal optimism and the market not responding to that optimism. Furthermore, the $\hat{p}^{VOL}$ model showed a similar movement to the $\hat{p}^{RET}$ model. This model could capture the aftermath of COVID-19, but it is not that great. As our model was not trained with 2024 data, the $\hat{p}^{VOL}$ could not capture the current soaring as well as $\hat{p}^{RET}$, and $\hat{p}^{LM}$. The $\hat{p}^{LM}$, interestingly, showed a decalcomania-like pattern to the $\hat{p}^{RET}$. Also, they could capture COVID-19's aftermath.

Revealed all r values (i.e. all correlations) showed the same movement compared to the 10-K top10 model's correlation (Table F.2), but represented a lower correlation power. For instance, Nvidia showed a moderate positive

correlation at $r_{pRET, pVOL}(=+.612)$, whereas the 10-K top10 model (Table F.2) showed a strong positive correlation at $r_{pRET, pVOL}(=+.905)$.

➤ *Nvidia Sentiment Model with Only-Risk-Factor*

Indicated Nvidia sentiment prediction score trained with the risk factor. Compared to the 10-K Nvidia sentiment model (), the $\hat{p}^{RET}$ showed almost the same performance, and the $\hat{p}^{VOL}$ increased 6% accuracy. The prediction took less time. It took 4 seconds as fewer tokens were used (i.e. 2,201). In the same comparison between and , the $\hat{p}^{RET}$ and the $\hat{p}^{VOL}$ showed a similar movement, whereas the risk factor $\hat{p}^{LM}$ model showed a gradually falling movement. We assumed this was also because the pLM model dataset was predominantly filled with negative words, as previously mentioned in .

Revealed a meaningful correlation compared to the 10-K Nvidia model showed a slightly stronger correlation but with more rigorous statistical significance in both $r_{pVOL, pRET}$ and r^{VOL} , Stock. The correlation results of the 10-K Nvidia model were not actually reliable as almost all correlations did not show statistical significance, except for the RET and VOL pair. For your consideration, both $\hat{p}^{LM}$'s correlations were less informative for our project. Even pLM's movements inaccurately mirrored Nvidia stock's behavior.

- Nvidia Most Influential Words:

➤ *Nvidia Sentiment Model with 10-K filing*

The following words were the most impactful words in Nvidia's $\hat{p}^{RET}$ sentiment score prediction, positively or negatively, respectively.

The word set to impact positively on Nvidia sentiment could be interpreted into two categories. The first category could be 'Innovation and Development'. This theme appeared again in the sector analysis. Words such as talent, thermal, failing, strict, and connect could be related to a firm's innovation. Talent could indicate that talented employees were drivers of innovation. Thermal could convey innovative GPU development, which is Nvidia's main product, as temperature control is required for chip's higher performance [105]. Also, words, including chain, favour, exclusion, and tender could symbolise the second theme, 'GPU market domination'. tender could represent the tendering process, which is a key component of the supply chain management system. We could interpret its GPU market domination as being caused by its innovation and strategic supply chain management, including tender processing, leading to the exclusive market position [106], [107].

The word set to influence the return negatively could be interpreted in a various way. ChatGPT4([108]) suggested that 'Market Challenges' could have a negative impact on Nvidia's sentiment. The relevant words for the theme included redemption, grid, initiative, petition, eastern, big, retrospective, branded, revolving, enactment, reducing, drone, joint, and pursuit. ChatGPT4([108]) argued that they reflected various operational and market challenges. For Nvidia, this could entail navigating energy regulations (grid),

responding to legal and regulatory petitions, while managing its brand in a competitive market big, eastern by adapting to laws affecting its business model initiative or product offerings (enactment).

We could interpret "Nvidia, Venture Capitalist" from words such as starting, attempt, and derive. Nvidia currently is ramping up its venture investment [109], [110], [111]. Venture investment in a startup (starting) that attempts to derive innovation could be risky and subsequently could increase volatility.

The following words mostly influenced the reduction of Nvidia's market uncertainty.

The words, such as virus, grown, returned, settle, and near, occurrence, could refer to "Stability from Alleviating the after math of COVID-19". They may hint at recovery(settle, returned, near, occurrence) from setbacks(virus, grown), signalling the reduction of volatility.

➤ *Nvidia Sentiment Model with Only-Risk-Factor ()*

The following positive word set seemed to show more accurate information to identify a positive topic on return.

For example, the word set, including card, mineral, original, thermal, chain, and intangible, could symbolise 'Graphics Processing Unit(GPU)', which has been the cash cow for Nvidia. Those words seemed to directly be indicative of GPU. card could represent a graphic card. mineral could indicate core raw materials for GPU manufacturing [112]. chain could refer to the supply chain of Nvidia that had a dominant position in the GPU market. thermal could indicate the innovative development of GPU. original and intangible could indicate Nvidia GPU's originality and its intellectual property.

Note that the words were extracted from the risk factor section where negative words were prevalent. This section contained valuable information to analyse firms' risks. So, we assume that the risk section could offer better informative insight for analysing their negative sentiment. From words like console and video, we could deduce the 'Video Game Industry' topic and its negative impact on return. Actual historical revenue data of Nvidia supported our interpretation. The gaming segment's revenue has been reduced [113].

The theme "Regulation and social responsibility on AI" emerged from words, including social, responsible, severe, restricted, and procedure. Given that the Environmental, Social, and Governance (ESG) ratings effectively influence return and volatility, ESG scores become important financial indicators [114], [115], [116]. Consequently, we could interpret from the word set that the importance(severe) of self regulation(restricted) for its product manufacturing procedure as part of their social responsibility (social, and responsible) on AI.

The words relevant to the 'External risks' theme included worm, hacker, fault, occurrence, confidential, and geographic. Cyber security, as an external risk, could be

deduced from words like worm, hacker, fault, occurrence, and confidential. Also, geographic could indicate a geopolitical issue as the US government is restricting cutting-edge AI chip export to China [117]. While cyber security or geopolitical issues could heighten volatility, we could interpret that Nvidia's effective risk management strategies, in the context of low volatility, mitigated its volatility. However, note that the model could be biased as managers' biases were reflected in the training data, a risk factor. Just because managers' arguments were positive to those issues, that does not mean risks did not disappear[118].

VI. CONCLUSIONS

A. Conclusions

Our suggested models estimated sentiment scores at three stakeholder levels: sector level, portfolio level, and company level. Also, each model was trained with either the entire 10-K filings or the risk factor section.

At the sector level, the sector model trained with the entire 10-K provided more informative sentiments than the risk-factor centred sector model. The 10-K model outperformed the risk factor model for both the return-labelled sentiment(i.e. p^{RET}) and volatility-labelled sentiment(i.e. p^{VOL}). The accuracies for p^{RET} or p^{VOL} were higher by 5% and 2% respectively. The 10-K model took more time for training though. Notably, p^{VOL} , in the 10-K sector model, showed periodic fluctuation trends in volatility, whereas p^{RET} was less fluctuated. In terms of correlation analysis, the 10-K sector showed a stronger correlation than the risk-factor sector model. For the qualitative analysis with the most impactful words, both models seemed to exhibit similar topics. The baseline model (i.e. p^{LM}) did show a distorted performance compared to both the 10-K sector model and the risk factor sector model. This was because the baseline model was biased to be negative as the training set was predominately filled with negative tokens. It is noteworthy that the Kalman filter should be applied for sector analysis to control noise.

At the portfolio level, the model based on the entire 10-K generally revealed more informative sentiments than the model with the risk factor. The 10-K portfolio model outperformed the risk factor portfolio model for both $\tilde{p}^{RET}$ and $\tilde{p}^{VOL}$. The accuracies for $\tilde{p}^{RET}$ or $\tilde{p}^{VOL}$ were higher by 5% and 4% respectively. The training time took longer than the risk-factor model. When it comes to correlation analysis, the 10-K sector depicted a way stronger correlation for all pairs. In the 10-K portfolio model, $\tilde{p}^{RET}$ strongly mirrored the top 10 portfolio stock price movement. $\tilde{p}^{RET}$ seemed to show no significant response to macroeconomic external factors such as the 2008 financial crisis and COVID-19. On the other hand, $\tilde{p}^{VOL}$ revealed trends in volatility response to the external factors. For the influential word analysis, the risk factor portfolio model seemed to highlight a distinct topic. Thus, we recommended employing two models(i.e. both the 10-K one and the risk factor one) for a portfolio-level analysis. Again, it is remarkable that the Kalman filter was also required for a portfolio analysis.

At the company level, the company-specific model using risk factors yielded sentiments with greater informational value than the model focusing on the complete 10-K documents. The risk factor model outperformed the 10-K-centred model For both p^{RET} and p^{VOL} . The accuracy levels for p^{RET} and p^{VOL} saw increases of 1% and 6%, respectively. The training time was shorter. For both the risk-factor company model and the 10-K company model, p^{RET} and p^{VOL} exhibited a similar movement. For correlation analysis, the model centred on risk factors demonstrated a marginally higher correlation, accompanied by more robust statistical significance. In the most influential word analysis, both models provided a distinct and detailed word, leading to capturing the theme. Hence, at the company level, we advised using the risk-factor-centred model for trend and correlation analysis while suggesting the use of both models for word analysis.

VII. LIMITATIONS AND FUTURE WORKS

Our suggested models can not currently capture meaningful phrase words as the model essentially calculates sentiment-Charged words based on a bag of words. For instance, the model will extract the phrase word 'chef executive officers(CEO)' in a word base separately and then evaluate whether each word can be sentiment-charged words concerning the dependent variables. In this word-based separation process, it loses the original meaning, which is CEO. Hence, in the future, we suggest to add a function that contains contextual information on the model. Bi-gram or n-gram can be examples. Industrial sentiment score prediction does not consider the allocation proportion of the QQQ portfolio. In the QQQ ETF fund, the top 10 firms take around 45% allocation proportion of the total in 2023, and the rest of 90 firms take the rest of 55% proportion. The portfolio is reconstructed annually. So, to predict a robust industrial-level sentiment score, the scores should consider the portfolio allocation proportion. In our industrial sentiment trending analysis, however, all sentiment scores are equally considered as the portfolio rebalancing data is restricted to attain. For instance, the return-labelled sentiment score of Apple Inc. in 2023 is 0.006, whereas Amazon.com Inc.'s sentiment in 2023 is -0.009. We used these scores without weighting their portfolio proportion. In 2023, the QQQ allocated Apple at 9.22% and Amazon at 4.83%. Thus, the weighted sentiment scores(i.e. $0.05532 = 0.006 * 9.22$ for apple and $-0.04347 = -0.009 * 4.83$ for Amazon) are required for a robust sentiment score calculation. In the future, we suggest that the sentiment score at the date should consider the allocation weight of the portfolio of the same date. Our model can not adapt to the latest firms' return or volatility because our model calculation only works with the filings released at the publication date. In other words, the predicted sentiment score is an annual data point at which the filing is released. If we have more latest textual information to represent a firm, our model would generate more recent sentiment scores. We suggest using 10-Q filling, which is a comprehensive report of a firm like 10-K but must be submitted quarterly.

REFERENCES

- [1]. Z. Ke, B. T. Kelly, and D. Xiu, "Predicting returns with text data," University of Chicago, Becker Friedman Institute for Economics Working Paper, no. 2019-69, September 2020, yale ICF Working Paper No. 2019-10, Chicago Booth Research Paper No. 20-37. [Online]. Available: <https://ssrn.com/abstract=3389884>
- [2]. R. P. Schumaker and H. Chen, "Textual analysis of stock market prediction using breaking financial news: The azfin text system," *ACM Transactions on Information Systems*, vol. 27, no. 2, p. 12, Mar 2009. [Online]. Available: <https://doi.org/10.1145/1462198.1462204>
- [3]. D. Shah, H. Isah, and F. Zulkernine, "Predicting the effects of news sentiments on the stock market," in *2018 IEEE International Conference on Big Data (Big Data)*. IEEE, 2018, pp. 4705–4708.
- [4]. J. Maqbool, P. Aggarwal, R. Kaur, A. Mittal, and I. A. Ganaie, "Stock prediction by integrating sentiment scores of financial news and mlp-regressor: A machine learning approach," *Procedia computer Science*, vol. 218, pp. 1067–1078, 2023. [Online]. Available: <https://doi.org/10.1016/j.procs.2023.01.086>
- [5]. Z. Wang, Z. Hu, F. Li et al., "Learning-based stock trending prediction by incorporating technical indicators and social media sentiment," *Cognitive Computation*, vol. 15, pp. 1092–1102, 2023, Received September 15, 2022; Accepted February 9, 2023; Published March 9, 2023; Issue Date May 1, 2023. [Online]. Available: <https://doi.org/10.1007/s12559-023-10125-8>
- [6]. A. Glodd and D. Hristova, "Extraction of forward-looking financial information for stock price prediction from annual reports using Nlp techniques," in *Proceedings of the 56th Hawaii International conference on System Sciences*, 2023, p. 10, rights: Attribution NonCommercial-NoDerivatives 4.0 International. [Online]. Available: <https://hdl.handle.net/10125/103313>
- [7]. United States Securities and Exchange Commission, "Form 10-k," 2024. [Online]. Available: <https://www.sec.gov/files/form10-k.pdf>
- [8]. S. Asthana and S. Balsam, "The effect of edgar on the market reaction to 10-k filings," *Journal of Accounting and Public policy*, vol. 20, no. 4-5, pp. 349–372, 2001. [Online]. Available: [https://doi.org/10.1016/S0278-4254\(01\)00035-7](https://doi.org/10.1016/S0278-4254(01)00035-7)
- [9]. H. You and X. Zhang, "Financial reporting complexity and investor underreaction to 10-k information," *Review of Accounting studies*, vol. 14, pp. 559–586, 2009. [Online]. Available: <https://doi.org/10.1007/s11142-008-9083-2>
- [10]. C. Kim, K. Wang, and L. Zhang, "Readability of 10-k reports and stock price crash risk," *Contemporary Accounting Research*, vol. 36, pp. 1184–1216, 2019. [Online]. Available: <https://doi.org/10.1111/1911-3846.12452>
- [11]. S. Blomme and J. Dedeyne, "Predicting the effect of 10-k, 10-q and 8-k company reports on abnormal stock returns using finbert nlp methods," University of Ghent, Tech. Rep., 2020. [Online]. Available: https://libstore.ugent.be/fulltxt/RUG01/002/837/812/RUG01-002837812_2020_0001_AC.pdf
- [12]. K. R. Jönsson and J. B. Jako, "Predicting stock performance using 10-k Filings: A natural language processing approach employing convolutional neural networks," *Copenhagen Business School, Tech. Rep.*, 2020.
- [13]. "Securities and exchange commission final rule, release no. 33-8591(fr-75)," <http://sec.gov/rules/final/33-8591.pdf>, 2005.
- [14]. R. Watts and J. Zimmerman, *Positive accounting theory*. Englewood Cliffs, NJ: Prentice Hall, 1986.
- [15]. T. Scott, "Incentives and disincentives for financial disclosure: Voluntary disclosure of defined benefit pension plan information by canadian firms," *The Accounting Review*, vol. 69, pp. 26–43, 1994.
- [16]. T. Fields, T. Lys, and L. Vincent, "Empirical research on accounting choice," *Journal of Accounting and Economics*, vol. 31, pp. 255–307, 2001.
- [17]. S. P. Kothari, X. Li, and J. Short, "The effect of disclosures by management, analysts, and business press on cost of capital, return volatility, and analyst forecasts: a study using content analysis," *The Accounting Review*, vol. 84, pp. 1639–1670, 2009.
- [18]. S. P. Kothari, S. Shu, and P. Wysocki, "Do managers withhold bad news?" *Journal of Accounting Research*, vol. 47, pp. 241–276, 2009.
- [19]. S. Chang, "Risk factor disclosures and ceo overconfidence," *College of Business Administration Seoul National University*, 2019.
- [20]. Reuters, "'refco risks boiler-plate disclosure.' By scott malone," <http://w4.stern.nyu.edu/news/news.cfm?docid=5094>, 2005.
- [21]. SEC, "Form 10-k instructions," <http://www.sec.gov/about/forms/Form10-k.pdf>, 2010.
- [22]. C. Magazine, "Sec pushes companies for more risk information," August 2, 2010, 2010.
- [23]. D. Kingsley, M. Solomon, and K. Jaconi, "Sec risk factor disclosure rules," <https://corpgov.law.harvard.edu/2021/12/22/Sec-risk-factor-disclosure-rules/>, 2021.
- [24]. H. M. Peirce, "Sec harming investors and helping hackers: Statement on cybersecurity risk management, strategy, governance, and incident disclosure," 2023.
- [25]. J. L. Campbell, H. Chen, D. Dhaliwal, H. Lu, and L. Steele, "The information content of mandatory risk factor disclosures in corporate filings," *Review of Accounting Studies*, vol. 19, no. 1, pp. 396–455, 2014.
- [26]. R. Israelsen, "Tell it like it is: disclosed risks and factor portfolios," 2014, working paper.
- [27]. V. Song, H. CAVUSOGLU, G. M. Lee, and M. L. Z. Ma, "It risk factor disclosure and stock price crashes," in *Proceedings of the 53rd Hawaii international Conference on System Sciences*, 2020.
- [28]. O. Hope, D. Hu, and H. Lu, "The benefits of specific risk-factor Disclosures," *Review of Accounting Studies*, 2016, forthcoming.
- [29]. T. Kravet and V. Muslu, "Textual risk disclosures and investors' risk Perceptions," *Review of Accounting Studies*, vol. 18, no. 4, pp. 1088–1122, 2013.

- [30]. J. Sha, "A data pipeline framework for automated extraction of risk factors from sec-10k filings," School of Informatics, 2023.
- [31]. J. Smailovic, M. Gracar, N. Lavrac, and M. Znidaršič, "Predictive sentiment analysis of tweets: A stock market application," in Proc. Int. Workshop Hum.-Comput. Interact. Knowl. Discovery Complex Unstructured Big Data, 2013, pp. 77–88.
- [32]. R. Ren, D. D. Wu, and T. Liu, "Forecasting stock market movement direction using sentiment analysis and support vector machine," IEEE systems Journal, vol. 13, no. 1, pp. 760–770, Mar 2019.
- [33]. N. Jegadeesh and D. Wu, "Word power: A new approach for content Analysis," Journal of Financial Economics, vol. 110, pp. 712–729, 2013.
- [34]. C. Kearney and S. Liu, "Textual sentiment in finance: A survey of methods and models," International Review of Financial Analysis, Vol. 33, pp. 171–185, 2014.
- [35]. Y. Guo and L. Zhou, "Textual tone in corporate financial disclosures: A survey of the literature," International Journal of Disclosure and Governance, vol. 17, pp. 101–110, 2020, received October 25, 2019; Published June 27, 2020; Issue Date September 1, 2020.
- [36]. S. Che et al., "Anticipating corporate financial performance from Ceo letters utilizing sentiment analysis," Mathematical Problems in Engineering, vol. 2020, 2020.
- [37]. F. LI, "The information content of forward-looking statements in corporate filings—a naïve bayesian machine learning approach," Journal of Accounting Research, vol. 48, pp. 1049–1102, 2010.
- [38]. P. Tetlock, "Giving content to investor sentiment: The role of media in the stock market," Journal of Finance, vol. 62, pp. 1139–1168, 2007.
- [39]. J. Engelberg, "Costly information processing: Evidence from information announcements," in AFA 2009 San Francisco Meetings Paper, 2008, Available at SSRN: <http://ssrn.com/abstract=1107998>.
- [40]. J. Engelberg, A. Reed, and M. Ringgenber, "How are shorts informed? Short sellers, news, and information processing," Journal of Financial Economics, vol. 105, no. 2, pp. 260–278, 2012.
- [41]. S. Ferris, G. Hao, and M. Liao, "The effect of issuer conservatism on Ipo pricing and performance," Review of Finance, vol. 7, no. 3, pp. 993–1027, 2013.
- [42]. E. Henry and A. J. Leone, "Measuring qualitative information in capital markets research," April 2009, available at SSRN: <https://ssrn.com/Abstract=1470807> or <http://dx.doi.org/10.2139/ssrn.1470807>.
- [43]. A. K. Davis, W. Ge, D. Matsumoto, and J. L. Zhang, "The effect of manager-specific optimism on the tone of earnings conference calls," January 13 2014, cAAA Annual Conference 2012, Available at SSRN: <https://ssrn.com/abstract=1982259> or <http://dx.doi.org/10.2139/Ssrn.1982259>.
- [44]. J. Doran, D. Peterson, and S. Price, "Earnings conference call content and stock price: The case of reits," Journal of Real Estate Finance and Economics, vol. 45, no. 2, pp. 402–434, 2010.
- [45]. X. Huang, S. H. Teoh, and Y. Zhang, "Tone management," August 21 2013, the Accounting Review, Forthcoming, Available at SSRN: <https://ssrn.com/abstract=1960376> or <http://dx.doi.org/10.2139/ssrn.1960376>
- [46]. N. J. Ferguson, D. Philip, H. Lam, and J. M. Guo, "Media content and stock returns: The predictive power of press," pp. 1–31, May 27 2015, Available at SSRN: <https://ssrn.com/abstract=2611046>.
- [47]. H. Chen, P. De, Y. J. Hu, and B.-H. Hwang, "Wisdom of crowds: The value of stock opinions transmitted through social media," Review of Financial Studies (RFS), Forthcoming, available at SSRN: <https://ssrn.com/abstract=1807265> or <http://dx.doi.org/10.2139/ssrn.1807265>.
- [48]. T. Loughran and B. McDonald, "Ipo first-day returns, offer price revisions, volatility, and form s-1 language," Journal of Financial Economics, vol. 109, pp. 307–326, 2013. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0304405X13000603>
- [49]. A. H. Huang, A. Zang, and R. Zheng, "Evidence on the information content of text in analyst reports," Accounting Review, April 2014, Forthcoming. [Online]. Available: <https://ssrn.com/abstract=1888724>
- [50]. N. Li, X. Liang, X. Li, C. Wang, and D. D. Wu, "Network environment and financial risk using machine learning and sentiment analysis," Human and Ecological Risk Assessment: An International Journal, Vol. 15, no. 2, pp. 227–252, 2009.
- [51]. S. Shuhidan, S. Hamidi, S. Kazemian, S. Shuhidan, and M. Ismail, "Sentiment analysis for financial news headlines using machine learning algorithm," in Proceedings of the 7th International Conference on Kansei Engineering and Emotion Research 2018. KEER 2018. Advances in Intelligent Systems and Computing, vol. 739. Singapore: Springer, 2018.
- [52]. G. Wang, T. Wang, B. Wang, D. Sambasivan, Z. Zhang, H. Zheng, B. Zhao, and S. Barbara, "Crowds on wall street: Extracting value from collaborative investing platforms," 03 2015.
- [53]. S. Sohngir, D. Wang, A. Pomeranets, and T. Khoshgoftaar, "Big data: deep learning for financial sentiment analysis," Journal of Big Data, Vol. 5, no. 1, p. 3, Dec 2018.
- [54]. IBM, "Deep learning," <https://www.ibm.com/topics/deep-learning>, accessed: 2024-04-01.
- [55]. L. Zhang, S. Wang, and B. Liu, "Deep learning for sentiment analysis: A survey," WIREs Data Mining and Knowledge Discovery, vol. 8, no.E1253, 2018.
- [56]. D. Tang, B. Qin, and T. Liu, "Document modeling with gated recurrent neural network for sentiment classification," in Proceedings of the conference on Empirical Methods in Natural Language Processing, 2015, pp. 1422–1432.
- [57]. K. S. Tai, R. Socher, and C. Manning, "Improved semantic representations from tree-structured long

- short-term memory networks,” <http://arxiv.org/abs/1503.00075>, 2015.
- [58]. Y. Kim, “Convolutional neural networks for sentence classification,” <http://arxiv.org/abs/1408.5882>, 2014.
- [59]. X. Zhang, J. Zhao, and Y. LeCun, “Character-level convolutional networks for text classification,” in *Proceedings of the Advances in neural Information Processing Systems*, 2015, pp. 649–657.
- [60]. R. Johnson and T. Zhang, “Deep pyramid convolutional neural networks for text categorization,” in *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Long Papers)*, vol. 1, 2017, pp. 562–570.
- [61]. Z. Yang, D. Yang, C. Dyer, X. He, A. Smola, and E. Hovy, “Hierarchical attention networks for document classification,” in *Proceedings of the conference of the North American Chapter of the Association for computational Linguistics: Human Language Technologies*, 2016, pp. 1480–1489.
- [62]. K. Mishev, A. Gjorgjevikj, I. Vodenska, L. Chitkushev, and D. Trajanov, “Evaluation of sentiment analysis in finance: From lexicons to transformers,” *IEEE Access*, vol. 8, pp. 131 662–131 682, 2020.
- [63]. M. Rizinski, H. Peshov, K. Mishev, M. Jovanovik, and D. Trajanov, “Sentiment analysis in finance: From transformers back to explainable lexicons (xlex),” *IEEE Access*, vol. 12, pp. 7170–7198.
- [64]. J. Dobson, “On reading and interpreting black box deep neural networks,” *International Journal of Digital Humanities*, vol. 5, pp. 431–449, 2023.
- [65]. J. Hering, “The annual report algorithm: Retrieval of financial statements and extraction of textual information,” Friedrich-Alexander-Universität Erlangen-Nürnberg, Lange Gasse 20, 90403 Nuremberg, Germany, 11 “2016, first Version: October 4, 2016. Current Version: November 28, 2016.
- [66]. S. Alizadeh, M. W. Brandt, and F. X. Diebold, “Range-based estimation of stochastic volatility models,” *The Journal of Finance*, vol. 57, no. 3, pp. 1047–1091, 2002.
- [67]. A. J. Patton, “Volatility forecast comparison using imperfect volatility proxies,” *Journal of Econometrics*, vol. 160, no. 1, pp. 246–256, 2011.
- [68]. F. Corsi, “A simple approximate long-memory model of realized volatility,” *Journal of Financial Econometrics*, vol. 7, no. 2, pp. 174–196, 2009.
- [69]. S. Borovkova and P. Lammers, “Sector news sentiment indices,” Available at SSRN 3080318, 2017.
- [70]. J. Durbin and S. J. Koopman, *Time series analysis by state space methods*. OUP Oxford, 2012, vol. 38.
- [71]. Statistics Solutions, “Pearson correlation assumptions,” 2023, accessed: 24-March-2024. [Online]. Available: <https://www.statisticssolutions.com/pearson-correlation-assumptions/>
- [72]. Newcastle University, “Critical region and confidence interval,” <https://www.ncl.ac.uk/webtemplate/ask-assets/external/maths-resources/statistics/hypothesis-testing/critical-region-and-confidence-interval.html>, 2024, accessed: 24-03-2024.
- [73]. D. J. Benjamin, J. Berger, M. Johannesson, B. A. Nosek, E. Wagen-Makers, R. Berk et al., “Redefine statistical significance,” Jul 2017, <https://doi.org/10.31234/osf.io/mky9j>.
- [74]. M. B. Editor. (2014) How to correctly interpret p values. Topics: Hypothesis Testing. [Online]. Available: <https://blog.minitab.com/en/Adventures-in-statistics-2/how-to-correctly-interpret-p-values>
- [75]. J. Frost. (2014) Why are p values misinterpreted so frequently? [Online]. Available: <https://statisticsbyjim.com/hypothesis-testing/p-values-misinterpreted/>
- [76]. A. S. Association, “American statistical association releases statement on statistical significance and p-values,” 2016, for more information: Ron Wasserstein, ron@amstat.org. [Online]. Available: <https://www.Amstat.org/asa/files/pdfs/P-ValueStatement.pdf>
- [77]. J. Park. (2016) Don’t be overwhelmed by p-value. [Online]. Available: <https://boxnwhis.kr/2016/04/15/dont-be-overwhelmed-by-pvalue.html>
- [78]. J. J. Filzen, “The information content of risk factor disclosures in quarterly reports,” *Accounting Horizons*, vol. 29, no. 4, pp. 887–916, Dec 2015. [Online]. Available: <https://doi.org/10.2308/acch-51175>
- [79]. L. Chan and N. Devia-Valbuena. (2023, Dec.) In the global rush for lithium, bolivia is at a crossroads. Accessed: 2024-03- 26. [Online]. Available: <https://www.usip.org/publications/2023/12/Global-rush-lithium-bolivia-crossroads>
- [80]. R. S. Koijen, T. J. Philipson, and H. Uhlig, “Financial health economics,” *Econometrica*, vol. 84, pp. 195–242, 2016.
- [81]. J. Kastrenakes. (2017, Dec.) Google’s project tango is shutting Down because arcore is already here. Accessed: 2024-03-26. [Online]. Available: <https://www.theverge.com/2017/12/15/16782556/Project-tango-google-shutting-down-drcore-augmented-reality>
- [82]. J. D’Onfro. (2017, Dec.) Google is going to shut down Its once-hyped augmented reality project, tango. Accessed: 2024-03-26. [Online]. Available: <https://www.cnn.com/2017/12/15/Google-kills-project-tango-ar-project.html>
- [83]. M. Majeed and M. Mazhar, “Environmental degradation and output Volatility: A global perspective,” *Pakistan Journal of Commerce and Social Science*, vol. 13, pp. 180–208, 03 2019.
- [84]. M. Majeed, M. Mazhar, and S. Sabir, “Environmental quality and Output volatility: the case of south asian economies,” *Environmental Science and Pollution Research*, vol. 28, pp. 31 276–31 288, 2021. [Online]. Available: <https://doi.org/10.1007/s11356-021-12659-6>
- [85]. J. Gagnon, M. Raskin, J. Remache, and B. Sack, “Large-scale asset purchases by the federal reserve:

- Did they work?" Federal Reserve Bank of New York, Staff Report 441, March 2010.
- [86]. H. Chen, V. Curdia, and A. Ferrero, "The macroeconomic effects of 'Large-scale asset purchase programs,'" Federal Reserve Bank of New York, Staff Report 527, December 2011.
- [87]. V. Curdia, "How stimulatory are large-scale asset purchases?" FRBSF Economic Letter, 08 2013.
- [88]. S. Gilchrist and E. Zakrajsek, "The impact of the federal reserve's 'Large-scale asset purchase programs on corporate credit risk," *Journal Of Money, Credit and Banking*, vol. 45, pp. 29–57, 2013.
- [89]. G. Wang, "The effects of quantitative easing announcements on the Mortgage market: An event study approach," *International Journal of Financial Studies*, vol. 7, no. 1, p. 9, 2019. [Online]. Available: <https://doi.org/10.3390/ijfs7010009>
- [90]. S. Sunder, "How did the u.s. stock market recover from the covid-19 Contagion?" *Mind Soc*, vol. 20, pp. 261–263, 2021. [Online]. Available: <https://doi.org/10.1007/s11299-020-00238-0>
- [91]. U. Seven and F. Yilmaz, "World equity markets and covid-19: 'Immediate response and recovery prospects,'" *Research in International Business and Finance*, vol. 56, p. 101349, 2021. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0275531920309570>
- [92]. P.-O. Gourinchas, Kalemli-Özcan, V. Penciakova, and N. Sander, "Fiscal policy in the age of covid: Does it 'get in all of the cracks?," *National Bureau of Economic Research, Working Paper 29293*, September 2021. [Online]. Available: <http://www.nber.org/papers/w29293>
- [93]. J. C. Stein, "Evaluating large-scale asset purchases," October 2012. [Online]. Available: <https://www.federalreserve.gov/newsevents/speech/stein20121011a.htm>
- [94]. S. Luck and T. Zimmermann, "Ten years later-did qe work?" *May Org/2019/05/ten-years-laterdid-qe-work/*
- [95]. The Economist, "America's war on huawei nears its endgame," *The Economist*, July 2020. [Online]. Available: <https://www.economist.com/Briefing/2020/07/16/americas-war-on-huawei-nears-its-endgame>
- [96]. C. Ting-Fang and L. Li, "Us-china tech war: Beijing's secret chipmaking champions," *Nikkei Asia*, May 2021. [Online]. Available: <https://asia.nikkei.com/Spotlight/The-Big-Story/US-China-tech-war-Beijing-s-secret-chipmaking-champions>
- [97]. J. Lewis, *Learning the superior techniques of the barbarians China's pursuit of semiconductor independence*. Washington: CSIS, 2019.
- [98]. D. Wu, Y. Lee, and Y. Ngui, "Chip shortage set to worsen as covid rampages through malaysia," *Bloomberg*, August 2021. [Online]. Available: <https://www.bloomberg.com/news/articles/2021-08-23/chip-shortage-set-to-worsen-as-covid-rampages-through-malaysia?sref=TtblOutP>
- [99]. M. Mazzucato and M. Tancioni, "R&d, patents and stock return volatility," *J Evol Econ*, vol. 22, pp. 811–832, 2012. [Online]. Available: <https://doi.org/10.1007/s00191-012-0289-x>
- [100]. S. Gharbi, J.-M. Sahut, and F. Teulon, "R&d investments and high-tech firms' stock return volatility," *Technological Forecasting and Social Change*, vol. 88, pp. 306–312, 2014. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0040162513002643>
- [101]. B. Hai, Q. Gao, X. Yin, and J. Chen, "R&d volatility and market value: the role of executive overconfidence," *Chinese Management Studies*, vol. 14, no. 2, pp. 411–431, 2020. [Online]. Available: <https://doi.org/10.1108/CMS-05-2019-0170>
- [102]. M. Brinkman and V. Sarma, "Infrastructure investing will never be the same," *McKinsey & Company*, 8 2022, article.
- [103]. F. Blanc-Brude, W. Schmundt, T. Bumberger, R. Friedrich, B. Georgii, A. Gupta, L. Lum, and M. Wilms, "Infrastructure strategy 2022: A pivot to the digital frontier," *Boston Consulting Group*, 3 2022, article.
- [104]. L. Gunnion, "Infrastructure investment: An economist's view from the ground up," 7 2021, article.
- [105]. A. Prakash, H. Amrouch, M. Shafique, T. Mitra, and J. Henkel, "Improving mobile gaming performance through cooperative cpu-gpu thermal management," in *Proceedings of the 53rd Annual Design Automation Conference*, ser. DAC '16. New York, NY, USA: Association for Computing Machinery, 2016. [Online]. Available: <https://doi.org/10.1145/2897937.2898031>
- [106]. K. Leswing, "Meet the \$10,000 nvidia chip powering the race for a.i." <https://www.cnn.com/2023/02/23/nvidias-a100-is-the-10000-chip-powering-the-race-for-ai.html>, 2023, accessed: 2024-03-27. [nvidia.com/en-gb/industries/retail/supply-chain-management/](https://www.nvidia.com/en-gb/industries/retail/supply-chain-management/), 2024, accessed: 2024-03-27.
- [107]. "Nvidia industries: Retail - supply chain management," <https://www.cnn.com/2023/02/23/nvidias-a100-is-the-10000-chip-powering-the-race-for-ai.html>, 2023, accessed: 2024-03-27.
- [108]. OpenAI, "Chatgpt (4) [large language model]," <https://chat.openai.com>, 2024.
- [109]. "Nvidia for startups: Venture capital," <https://www.nvidia.com/en-gb/startups/venture-capital/>, 2024, accessed: 2024-03-27.
- [110]. "Nvidia investments," <https://blogs.nvidia.com/blog/nvidia-investments/>, 2024, accessed: 2024-03-27.
- [111]. PYMNTS, "Nvidia invests in 35 Ai in companies 2023," <https://www.pymnts.com/artificial-intelligence-2/2023/nvidia-invests-in-35-ai-companies-in-2023/>, December 2023, accessed: 2024-03-27. [Online]. Available: https://euromines.org/files/key_value_chain_euromines_2020_electronics_euromines_final.pdf
- [112]. Euromines, "Key value chain electronics euromines," 2020. [euromines_2020_electronics_euromines_final.pdf](https://euromines.org/files/key_value_chain_euromines_2020_electronics_euromines_final.pdf)
- [113]. Nvidia, "2023 annual report," 2023. [Online]. Available: https://s201.q4cdn.com/141608511/files/doc_financials/2023-Annual-Report-1.pdf
- [114]. M. La Torre, F. Mango, A. Cafaro, and S. Leo, "Does the esg index affect stock return? evidence from the

- eurostox50," Sustainability, vol. 12, no. 16, p. 6387, 2020. [Online]. Available: <https://doi.org/10.3390/su12166387>
- [115]. G. Capelle-Blancard and A. Petit, "Every little helps? esg news and stock market reaction," Journal of Business Ethics, vol. 157, pp. 543–565, 2019. [Online]. Available: <https://doi.org/10.1007/s10551-017-3667-3>
- [116]. G. Giese, L.-E. Lee, D. Melas, Z. Nagy, and L. Nishikawa, "Foundations of esg investing: How esg affects equity valuation, risk, and performance," The Journal of Portfolio Management, vol. 45, pp. 69-83, 2019.
- [117]. D. Sevastopulo and Q. Liu, "Us-china doomsday scenario not likely to happen, says nvidia's jensen huang," Financial Times, 10 2023. [Online]. Available: <https://www.ft.com/content/>
- [118]. Y. Yu, "Us-china doomsday scenario not likely to happen, says nvidia's jensen huang," Nikkei Asian Review, 03 2024. [Online]. Available: <https://asia.nikkei.com/Business/Tech/Semiconductors/U.S.-China-doomsday-scenario-not-likely-to-happen-Nvidia-s-Jensen-Huang>
- [119]. Securities and Exchange Commission, "Form 10-k," <https://www.sec.gov/files/form10-k.pdf>, 2024, annual Report Pursuant to Section 13 or 15(d) of The Securities Exchange Act of 1934, MB Number: 3235-0063, Expires: December 31, 2026, Estimated average burden hours per response: 2,249.36.
- [120]. Wiki, "Form 10-k," https://en.wikipedia.org/wiki/Form_10-K, 2024,
- [121]. Analytics Vidhya, "Stemming vs lemmatization in nlp: Must know differences," 2022, accessed: 22-03-2024. textual analysis, dictionaries, and 10-ks," The Journal of Finance, vol. 66, no. 1, pp. 35-65, 2011.
- [122]. D. Garcia, "Sentiment during recessions," The Journal of Finance, vol. 68, no. 3, pp. 1267-1300, 2013.
- [123]. T. Loughran and B. McDonald, "When is a liability not a liability?"
- [124]. Wikipedia, "Pearson correlation coefficient - Wikipedia, the free encyclopedia," 2024, [Online; accessed 24-March-2024]. [Online]. Available: https://en.wikipedia.org/wiki/Pearson_correlation_coefficient

APPENDIX

➤ *Part 1*

- **Item 1: Business** - This section describes the business of the company: its main products or services, subsidiaries, and markets in which it operates. It may also contain recent events, competition, regulation, and labour problems.
- **Item 1A: Risk Factors** - This section provides risks and uncertainties, likely external effects, or possible failures that could affect their financial performance. Risk factors are generally enumerated based on their importance.
- **Item 1B: Unresolved Staff Comments** - This section offers an explanation of any issues raised by the SEC staff on the previous reports if these issues have not been resolved afterwards.
- **Item 2: Properties** - This section only lays out the company's significant physical properties, not intellectual or intangible property.
- **Item 3: Legal Proceedings** - This section discloses any significant ongoing lawsuit or other legal proceeding.
- **Item 4 - [RESERVED]**

➤ *Part 2*

- **Item 5: Market for Registrant's Common Equity, Related Stockholder Matters and Issuer Purchases of Equity Securities** - This section discloses their performance in the stock market, dividends, and repurchases of their own stocks.
- **Item 6 - [RESERVED]**
- **Item 7: Management's Discussion and Analysis of Financial Condition and Results of Operations (MD&A)** - This section discusses the firm's management for its financial performance, challenges, chances, and future outlook.
- **Item 7A: Quantitative and Qualitative Disclosures about Market Risks** - This section shows the company's exposure to market risks.
- **Item 8: Financial Statement and Supplementary Data** - This section offers the audited financial statements. It contains the balance sheet, income statement, cash flow statement, and footnotes.
- **Item 9: Changes in and Disagreements with Accountants Accounting and Financial Disclosure** - Companies address any alterations in or disputes with their accountants over financial reporting.
- **Item 9A: Controls and Procedures** - This section details the company's procedures for disclosure controls and its internal control mechanisms for financial reporting. **Item 9B: Other Information** - This section offer additional information that does not align with the conte other sections.

➤ *Part 3*

- **Item 10: Directors, Executive Officers and Corporate Governance** – This section delves into the specifics of the company's leadership, their respective roles, and the practices in place for corporate governance.
- **Item 11: Executive Compensation** – This section addresses compensation policies and programmes, and the compensation Oo top executives.
- **Item 12: Security Ownership of Certain Beneficial Owners and Management and Related Stockholder Matters** – This section offers major shareholders' ownership of the company's stock as well as insiders.
- **Item 13: Certain Relationships and Related Transactions, and Director Independence**- This section encompasses transactions involving directors, executives, and their affiliates, along with details regarding the independence of directors.
- **Item 14: Principal Accountant Fees and Services** – The section details the fees charged by the company's auditors for their services.

➤ *Part 4*

- **Item 15: Exhibits, Financial Statement Schedules** – This section contains a list of the financial statements and exhibits.[119], [120]

➤ *10-K Filing (e.g Nvidia 2010-03-18)*

- [heading]Our Company[/heading]

NVIDIA Corporation helped awaken the world to the power of computer graphics when it invented the graphics processor unit, or GPU, in 1999. Expertise in programmable GPUs has led to breakthroughs in parallel processing which make supercomputing inexpensive and widely accessible.

Heading]ITEM 7. MANAGEMENT'S DISCUSSION AND ANALYSIS OF FINANCIAL CONDITION AND RESULTS OF OPERATIONS[/heading] The following discussion and analysis of our financial condition and results of operations should be read

in conjunction with "Item 1A. Risk Factors", "Item 6. Selected Financial Data", our Consolidated Financial Statements and related Notes thereto

[Heading]Overview[/heading] [heading]Our Company[/heading] NVIDIA Corporation helped awaken the world to the power of computer graphics when it invented the graphics processor unit, or GPU, in 1999. Expertise in programmable GPUs has led to breakthroughs in parallel processing which make super- computing inexpensive and widely accessible. We serve the entertainment and consumer market with our GeForce graphics products, the professional design and visualization market with our Quadro graphics products, the high-performance computing market with our Tesla computing solutions products, and the mobile computing market with our Tegra system-on-a-chip products.

➤ Risk Factor Section (e.g Nvidia 2023-02-24)

- [Title]Risk Factors Summary[/title]
- [Heading]Risks Related to Our Industry and Mar- kets[/heading]

Failure to meet the evolving needs of our industry and markets may adversely impact our financial results. Competition in our current and target markets could cause us to lose market share and revenue.

- [Heading]Risks Related to Demand, Supply and Manufactur- ing[/heading]

Failure to estimate customer demand properly has led and could lead to mismatches between supply and demand. Dependency on third-party suppliers and their technology reduces our control over product quantity and quality, manufacturing yields, development, enhancement, and product delivery schedules and could harm our business. Defects in our products have caused and could cause us to incur significant expenses to remediate and can damage our business.

➤ Preprocessing

Through the 10-K filings extraction model at Chapter 3, we finally collected almost all 10-K filings listed in the QQQ from 2006 to 2023, which is 1383 filings as well as the text files with only the Item 1A risk factor section extracted. With these files, we preprocessed them before feeding them into our prediction model. Firstly, we split all sentences of a filing into words. Secondly, we removed non-alphabetic tokens from the texts such as numbers, proper nouns, and special characters(i.e. punctuations). For instance, Item 8 of a 10-K filing contains many numbers as a financial statement of a company is included in that section. As numbers were not necessary for textual analysis, we removed them. Thirdly, we removed stop words such as "is", "the", or "and", as they do not contain informative information. In the final step of the preprocessing, we selected lemmatisation, instead of stemming. Although lemmatisation is computationally more expensive than stemming, it tends to capture the more accurate base form of a word through linguistic analysis (i.e. considering Parts-of-speech) [121].

➤ Hyper-Parameter Setting

We set the hyper-parameters equal for the models to ensure a fair comparison between the sentiments with return labels and those with volatility labels for a sector level, a portfolio level, and a company level.

In Section 4.2, we replaced 0 by the θ in Equation 4.2, and the value 0.5 by quantile q in Equation 4.3. θ is a threshold and we used it to define high and low volatility. To figure out the balanced and q hyper-parameter for both the technology sector and a firm(in this paper, Nvidia), we selected the value of θ and q from, on the training windows. In the technology sector from, we practically set θ as the 65th percentile of the distribution and q as 0.65. It implies we defined the volatility above the 65 quantile as high volatility, whereas the volatility below the 65 quantile is low volatility. Note that we set the 65th percentile and 0.65 for θ and q for both QQQ volatility itself and the volatilities of the top 10 firms in QQQ, respectively. Empirically, the 3-day volatility for both QQQ and the top 10 showed a similar trend. We can see the peaks from the graph, referring to the financial crisis of 2008 and the COVID- 19 pandemic around 2020. In the case of a firm(i.e. Nvidia) from, θ was set as the 65th percentile of the distribution, and q was 0.65. Similar to the QQQ volatility graph, Nvidia experienced high levels of volatility during the 2008 financial crisis and COVID-19 showed higher volatility movement in general during the training windows.

Furthermore, returns for both the technology sector and a firm showed a balanced proportion at the median of three- day returns. Hence, we set 0.5 at Equation 4.3 for both the technology sector in and a firm return in.

In Equation 4.3, we specifically defined α and k . Note that α was a threshold to filter out sentiment-neutral words. We set α^+ and α^- within the interval $(0, 0.5]$ such that each of the positive word sets and the negative word sets includes 100 words. Moreover, note that $k \in \mathbb{N}$ was another threshold to relate to the count of word w across all filings such that we used k as a minimum frequency requirement to reduce the influence of rare words. We set k as the 90 percentile quantile of the term frequency distribution. It means we ignored words that appear less than the 90 percentile quantile in a filing.

In Equation 4.8, we also defined λ . It was a positive constant And used in a penalty term to adjust our model. The penalty term was used to avoid the model overfitting when few sentiment- charged words appear in the filing. Without the penalty term, the models can consider the filing that contains few negative word but does not include positive as a negative filing. However, just because the model contains few negative sentiment-charged words without positive words, that does not mean the filing has a negative tone. To control this phenomenon, we set the penalty coefficient λ to 0.1 in the penalty term.

➤ Baseline

The purpose of the paper was to predict the sentiment score for both a technology sector and a firm from the contents of 10-K filings. To achieve that, we could infer the sentiment scores for $\hat{p}^{\text{RET}}$ and $\hat{p}^{\text{VOL}}$ by labelling with return and volatility, respectively. Then, we introduced a baseline sentiment score to compare our sentiment scores with it.

Our baseline model to calculate baseline sentiment scores was suggested by [122], and used the dictionary created by [123]. Our baseline Loughran and McDonald(LM) sentiment score, $\hat{p}^{\text{LM}}$, is computed as

$$\hat{p}_i^{\text{LM}} = \left(\sum_{w \in \text{LM}^+} d_{i,w} - \sum_{w \in \text{LM}^-} d_{i,w} \right) \left(\sum_{w \in V} d_{i,w} \right)^{-1} \quad [16]$$

Where LM^+ refers to the positive word lists and LM^- refers to the negative words list of Loughran and McDonald's dictionary. To attain the base sentiment scores, we calculated the difference between the number of positive words and the number of negative words, and then we divided this result by the total number of words in each 10-K filing.

To explain more of the LM's dictionary for our robust evaluation, the LM dictionary is traditionally used for financial analysis. LM offered an improved textual dictionary for a better accurate financial analysis in financial documents. Existing the financial dictionary of 10-K filings based on the Harvard dictionary misclassified the negative words in a financial context. Three-fourths of negative words in the 10-K filings do not carry a negative connotation in financial reports. To this issue, LM suggested three approaches. Firstly, LM created a refined list of words that more accurately reflects negative sentiment in a financial context, by analysing every word that appears in at least 5% of the SEC's 10-K filings. Secondly, LM introduced a term weighting scheme that controls the influence of frequently mentioned words and amplifies the significance of rarer terms, thus mitigating the misclassification of words. Finally, they added five other word classifications(e.g., positive, uncertainty, litigious, strong modal, and weak modal words). They found that these new classifications can be linked to market reactions, volatility in stock returns, unexpected earnings, and trading volumes [123]. In our study, we used only negative and positive categories to adjust the financial dictionary for our model.

➤ Portfolio

In this study, we formed a portfolio to evaluate the technology sector sentiment scores. A portfolio can be changed to any portfolio for financial analysis. In practice of our study, we selected Invesco QQQ Trust Series 1 an exchange-traded fund (QQQ or QQQ ETF). This passive fund(i.e., our portfolio) tracks the Nasdaq 100 Index, which consists of shares from 100 of the largest and most innovative non-financial firms listed on the Nasdaq stock exchange. Holdings in QQQ are predominantly in large-cap technology firms, accounting for 60% of the portfolio. As such, the QQQ is conventionally considered as a technology sector fund. The top 10 holdings represent a 50% allocation of the portfolio, with 9 out of 10 firms being in the tech sector. To represent the technology sector in the US, we formed two portfolios from the QQQ fund. Firstly, we formed the portfolio, which has the exact same allocation proportion as the QQQ fund itself as of 2023. This allocation proportion is annually corrected so that our current portfolio reflects 2023 allocation data. The actual 2023 portfolio allocation can be found in Appendix D. This portfolio is used to compare the sector sentiment scores.

The second portfolio was constructed with the top 10 firms, considering the asset allocation ratio of the first portfolio. In other words, the second portfolio consisted of the top 10 firms, accounting for 50% of the QQQ portfolio, according to the proportion invested in the first portfolio. This portfolio was computed as:

$$A_j = \frac{w_j}{\sum_{j=1}^{10} w_j} \quad [17]$$

Where j denoted a firm invested at the j -th ranking in the 2023 portfolio, and w_j represented the portfolio weight of the j -th firm. A_j referred to the allocated proportion of the j -th firm in the 2023 portfolio.

In the case of a single firm evaluation, we did not form a portfolio for it. Instead, we followed the firm's stock market price.

➤ Definition

$$r = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum (x_i - \bar{x})^2 \sum (y_i - \bar{y})^2}} \quad [18]$$

- r refers to the correlation coefficient.
- $\bar{x}_i$ and $\bar{y}_i$ refer to the individual sample points for variables x and y , respectively.
- $\bar{x}$ and $\bar{y}$ represent the mean values of the samples for x and y . [124]

➤ *The Sector Sentiment Model with Only-Risk-Factor*

➤ *The Top10 Sentiment Model with 10-K*

Our system can update the latest SEC filings, facilitating data preparation(i.e. preprocessing, extracting the risk factor section, and creating a document-term dictionary). Basically, our system can collect every type of SEC filing. In practice, we collected 10-K filings, followed by the extraction of the risk factor section and the creation of a document-term dictionary. As mentioned in Experiment Result, we generated sentiment metrics for three different stakeholders with either the entire 10-K filing or the risk factor section, respectively. To do that, we need 6 types of a document-term dictionary. One of the dictionaries, for instance, an individual firm(e.g. Nvidia)'s document-term dictionary from the complete 10-K filings.

Our system can automatically create 6 types of a documentterm dictionary with the latest 10-K filings. This automation process works on Apache Airflow, an open-source platform designed to manage workflow processes in data engineering pipelines. Our airflow automation system consists of three dags corresponding to each stakeholder level. Note that a Directed Acyclic Graph(DAG) refers to a collection of all the tasks you want to execute, arranged to reflect their relationships and dependencies. Each dag creates the corresponding documentterm dictionary for our Sentiment Score Prediction Model. As mentioned in Project Objective, the publication dates of 10-K filing are various to each firm although it should be released nearly every day throughout that year. Thus, we scheduled our dags differently. Sector-level dag(i.e. Technology sector from the QQQ) updates daily, Portfolio-level dags(i.e. Top10) updates monthly, and Firm-level dag update yearly. You can check data workflow in.