

TCFD 報告與文字分析

摘要

本文目的在於：運用文字探勘工具，分析氣候相關財務揭露報告書（簡稱 TCFD 報告書），作為評鑑本國銀行業與保險業 TCFD 報告書的參考。我們以人工標記篩選出各公司報告書中，有關治理、策略、風險管理、指標與目標等四個面向的敘述，藉由文字分析挑選適合的變數，搭配各公司的相關變數，透過迴歸分析建構評分模型。然後依據本國銀行、壽險、產險公司 2022 年 TCFD 報告書之人工評審分數，示範如何利用探索性資料分析，挑選重要變數，並透過交叉驗證測試迴歸模型的穩定性。研究結果發現本文所提出的迴歸分析模型與人工計分的結果相當契合，其中外部專家評審分數的迴歸模型的估計效果較佳。

關鍵詞：氣候變遷、氣候相關財務揭露報告書、文字分析、探索性資料分析、迴歸分析

TCFD Reporting and Text Mining

Abstract

The aim of this research is to employ text mining tools to analyze TCFD reports issued by domestic banks and insurers for evaluation purposes. We label descriptions from these financial firms' TCFD reports relating to governance, strategy, risk management, and metrics and targets. Using text mining and regression analyses, we first select variables and construct scoring models. We then demonstrate how to use exploratory data analysis to select variables and test stability of regression models through cross-validation. We find that our regression results are consistent with manual evaluation results, in particular those from experts and scholars.

Keywords : Climate change, TCFD reporting, text mining, exploratory data analysis, regression analysis

一、緒論

依據世界經濟論壇（World Economic Forum，2024）所發布 2024 年全球風險報告（The Global Risks Report），在未來 10 年內，全世界前 4 大風險分別是：極端氣候、自然系統的重大改變、生物多樣性損失與生態系統崩潰、自然資源短缺。這些環境變遷風險可能造成的危害，已經讓環境變遷議題成為社會各界最關心的議題之一。一般認為氣候變遷，特別是溫室氣體（Greenhouse Gas）的大量排放，係造成環境變遷的重要原因之一。例如：與 18 世紀中葉（第一次工業革命）相比，二氧化碳濃度至今已經增加了 50%（Eggleton，2013）。這些溫室氣體的增加，導致全球暖化（Friedlingstein 等人，2022）。依據聯合國跨政府氣候變化委員會（Intergovernmental Panel on Climate Change）於 2023 年發布研究報告，在過去 200 年間，全球氣溫已較工業化時代初期升溫攝氏 1.1 度。

氣候報導（Climate Reporting）是處理氣候變遷危機的第一步。為了讓企業在進行氣候揭露時有所依循，金融穩定委員會（Financial Stability Board）的氣候相關財務揭露工作小組（Task Force on Climate-Related Financial Disclosures，簡稱 TCFD）於 2017 年 6 月發布揭露框架建議，公布氣候相關財務揭露建議（Recommendations on Climate-related Financial Disclosures），希冀企業能依據相關建議進行揭露，並作為公司財務性報告（Financial reports）的一部份。TCFD 旋於 2021 年 10 月公布氣候相關財務揭露建議附錄（2021 TCFD Annex）的最新版本，取代 2017 年關於執行氣候相關財務揭露工作小組建議之附錄（2017 TCFD Annex）。TCFD 揭露建議被英國商業、能源與產業策略部門描述為「企業分析、理解及最終揭露氣候相關財務資訊的最有效框架之一」。由於氣候風險的資訊對於投資決策和風險定價至關重要，愈來愈多投資者對公司施加壓力，要求發布全面氣候相關財務揭露的報告（Krueger 等人，2020）。

為呼應 TCFD 公布氣候相關財務揭露建議，我國金融監督管理委員會（以下簡稱金管會）於 2021 年發布「本國銀行氣候風險財務揭露指引」，¹及「保險業氣候相關風險財務揭露指引」，²要求本國銀行與保險業者，自 2023 年起，每年 6

¹ 金管會發布「本國銀行氣候風險財務揭露指引」，銀行業者應自 112 年起揭露氣候相關風險財務資訊。參照：

https://www.fsc.gov.tw/ch/home.jsp?id=96&parentpath=0,2&mcustomize=news_view.jsp&dataserno=202111300008&dtable=News

² 金管會發布「保險業氣候相關風險財務揭露指引」，保險業者應自 112 年起揭露氣候相關風險財務資訊。參照：

月底前，應揭露氣候相關風險財務資訊，將相關資訊納入永續報告書或公布於公司網站。

我們利用政治大學企業永續管理研究中心（以下簡稱政大永續中心）第一屆氣候相關財務揭露（TCFD）報告書評鑑資料，並使用文字探勘工具，分析氣候相關財務揭露報告書（簡稱 TCFD 報告書），作為評鑑銀行保險業 TCFD 報告書的參考。我們示範如何利用探索性資料分析，挑選重要變數，並透過交叉驗證測試迴歸模型的穩定性。研究發現：本文所提出的迴歸分析模型與人工計分的結果相當契合，特別是外部專家評審分數的迴歸模型。

本文的貢獻有二，可以分成學術與實務方面。在學術方面，目前氣候揭露的研究仍然不算多，尤其是以文字分析（Text Mining）的角度切入。即便有也多半係藉由關鍵詞及其使用比例，來捕捉企業界認為的氣候變遷風險（Sautner 等人，2022）。過去文獻很少以文本語境（Context）的觀點考慮詞彙本身可能的模糊性，或是探究詞彙在文句中的角色及詞彙間的關聯（Varini 等人，2020）。推究過去較少文獻以文字分析來研究 TCFD 報告書的主要原因有三：TCFD 框架缺乏一致評分方法、TCFD 建議的揭露事項不夠具體、氣候變遷相關開放資料庫有限。本文的研究結果，可以彌補相關研究的文獻空缺；在實務方面，我們的研究結果顯示迴歸模型相當程度可以反應評審計分的想法，可以作為檢核人工計分的輔助工具，包括檢查是否有誤植分數等問題，因此本研究有其實務應用的重要性。

本文以 2023 年本國銀行 38 家、21 家壽險公司、23 家產險公司的報告書為例，使用迴歸模型檢視內部人工、外部專家評審的評分結果，並提供研究心得及相關建議。由於報告文字屬於非結構資料（Unstructured Data），需要先經過資料結構化（Structurization）定義變數，本文將大略說明如何結構化文字資料，以及將其轉化為反映寫作風格的可能作法。本文套用探索性資料分析（Exploratory Data Analysis, EDA）驗證結構化是否有效，作為挑選迴歸模型解釋變數的依據；另外，除了常見的迴歸分析，本文也考慮幾種常見的機器學習（Machine Learning）模型，評估本文作法的優劣。

二、文獻回顧與研究方法

最近幾年 BERT (Bidirectional Encoder Representations from Transformers, 譯為「基於變換器的雙向編碼器表示技術」) 等自然語言處理 (Natural Language Processing) 工具的引進, 學者嘗試捕捉文本的詞彙語境語義, 使得氣候相關文本的研究大幅超越傳統方法 (Luccioni 等人, 2020)。然而, 氣候變遷是一個複雜、快速變化且模糊的議題, 這些語言模型需要大量的訓練資料, 但氣候相關開放資料非常有限, 使得 TCFD 的文字分析進展緩慢。本研究以人工標記我國金融保險業 38 家本國銀行、21 家壽險及 23 家產險公司的永續報告書為依據, 套用統計方法論建立 TCFD 的分析模型, 希冀可作為研究臺灣發展氣候變遷之參考。因為 TCFD 報告評鑑在我國屬於新型態研究, 對於建立分析模型和挑選關鍵資訊的建立不多見, 尤其是哪些資訊是判斷 TCFD 報告書內容的關鍵, 以下將介紹如何挑選重要文字變數, 包括如何格式化文字資料。

資料分析通常可分為兩個步驟: EDA 和 CDA (Confirmatory Data Analysis, 驗證性資料分析, 簡稱 CDA)。EDA 目的在於從各面向探索資料的關鍵特性, 並以圖表等視覺化 (Data Visualization) 方式呈現, 協助研究者在 CDA 流程中挑選較為合適變數及其函數型態; CDA 則是根據 EDA 選擇的變數, 建構能夠詮釋研究目標的模型, 使其兼顧實質的詮釋、模型的準確性。EDA 通常會檢視資料資本特性, 包括各變數平均數及變異數、兩個變數間的關係等, 從中觀察、推敲哪些變數較為適合置入 CDA 模型, 但這些工作項目不容易掌握重點, 非常需要依賴使用者的經驗及判斷, 因此在 EDA 經常被輕忽、甚至完全省略, 導致不少研究者認為模型可以決定一切, 忘記了「垃圾進、垃圾出」(Garbage in, garbage out) 這個常識, 忽視輸入錯誤、無意義資料的可能性。

EDA 對於文字、影音等非結構資料之分析, 格外重要。以文字的分析為例, 如何詮釋其意義取決於使用者目的、習慣、文字特性等因素, 或許可以找出常見作法及慣例, 但未必能夠訂出明文規則, 必須視研究目標調整分析步驟及作法。現代中文的書寫屬於白話文, 多半透過雙字或多字組成的詞彙表達意義 (何立行、余清祥、鄭文惠, 2014), 因此中文文本分析通常會經過斷詞, 但字彙有時也能反映寫作風格, 不能完全摒棄。常見的中文斷詞工具包括 Jieba 和 CKIP (余清祥、葉昱廷, 2020), Jieba 軟體根據《人民日報》等簡體中文文本為資料集, 透過字典樹 (Trie) 等電腦演算法進行訓練, 普遍認為 Jieba 斷詞執行速度快, 但用於正體中文的精確性較低, 我們測試了自行匯入詞典檔, 可改善臺灣正體中文文稿的斷詞成效。CKIP 則是由中央研究院中文詞知識庫小組 (CKIP Lab) 所開發

的斷詞套件，近年推出以 BiLSTM 架構訓練的 CKIP Tagger (Li 等人, 2020) 及 CKIP Transformers 等版本，雖然執行速度比先前版本改善不少但仍然不如 Jieba，不過新版本具有專有名詞辨識（或實體辨識，Named Entity Recognition, NER），而且在正體中文的斷詞準確率略優於 Jieba。考量執行速度、軟體使用的普遍性，本文使用 Jieba 斷詞分析雙字詞及多字詞的使用特性。

關於文字變數的結構化，大致可分為兩個方向：一則是整體的用字及風格，計算描述文字資料的重要特徵，像是資料分配的期望值、變異數之類的統計量；另一則是逐句解析字、詞間的關聯，可借用相關係數之類的測量值，評估哪些字詞會共伴出現（或可稱為關鍵詞叢），進一步探索語言背後的語境意義。由於字詞關聯牽涉層面較為廣泛，限於篇幅，本文只以整體用字風格為目標，介紹常見的幾種文字特色的測量值，說明這些數值與人工計分間的關係。由於文字屬於類別資料 (Categorical data)，以不同類別標示較為合適，也就是將「字」、「詞」（或是雙字詞、多字詞）視為基本單位，計算 TCFD 報告之字/詞總數及其不同類別數（種類），並考量字詞的幾種常見不均度指標，包括 TTR (Type-Token-Ratio，相異字數; Templin, 1957)、熵 (Entropy; Shannon, 1948)、Simpson 指標 (Simpson, 1949)，以及詞向量 (Word to Vector; Mikolov 等人, 2013)。

如果將文章比擬為生態圈，可套用生物多樣性 (Species Diversity) 描述寫作風格，其中 TTR、熵、Simpson 指標可用於測量文字豐富度及不均度。TTR 又稱為生字率，測量在文章中出現了多種不同字彙，但中文常見字彙在四千個左右，文章篇幅愈長、TTR 數值也愈小，因此會設定固定字數下的字彙總數，這種測量方式也稱為標準化相異字數比例 (Standardized TTR)，本文設定隨機抽樣 500 字得出之平均字彙數。熵和 Simpson 指標可反映字詞使用比例，若 p_i 為第 i 個字彙或詞彙的使用比例，定義熵 $= -\sum_i p_i \times \log(p_i)$ 及 Simpson 指標 $\theta_s = \sum_i p_i^2$ 測量字彙使用的不均度，其中熵愈小愈不平均、Simpson 指標則是愈大愈不平均，這兩個指標的數值也會受到字數影響，同樣隨機抽樣 500 字估計這兩個指標。

詞向量也稱為詞嵌入 (Word Embeddings)，是將詞彙轉換成數字向量，藉此將語意關聯性、相似性比較高的字詞歸納成同一類，接著就能進行語意搜尋、文本分類等之分析，進一步捕捉文字背後的語意訊息。詞向量透過類神經網路 (Neural Network, NNs)，主要有 CBOW 與 Skip-gram 兩種模型 (Tomas 等人, 2013)，直觀而言，Skip-gram 目的在於輸入字詞、預測上下文，而 CBOW 則是給定上下文、預測輸入字詞。本文使用的詞向量軟體為 Word2Vec，算是 Skip-gram

模型的副產品，其中 Skip-gram 將輸入的字詞轉換成詞向量的階段，即是 Word2Vec。

本文主要以迴歸模型評估外部評審/內部人工計分，透過探索性資料分析檢視與人工計分比較有關聯的變數，接著將這些變數置入評分模型，藉此評估能否以量化分析的角度評估人工計分。評分模型以常見的迴歸模型為主，可兼顧解釋變數的詮釋、整體模型的效果，而且操作相對簡單、比較容易被大家接受，本文也會將幾種機器學習模型的結果作為比較對象。迴歸模型廣為大眾使用，通常以估計值與觀察值的平方差異（或最小平方法，Least Squares）作為估計標的，但銀行保險業 TCFD 報告的樣本數不多，為了避免出現過度配適（Overfitting，或譯為過度擬合），通常會在最小平方法框架下加上某個懲罰項。懲罰項之目的希望模型盡量精簡，避免模型包含不必要的解釋變數或函數型態，近年懲罰最大概似估計量（Penalized Maximum Likelihood Estimator）就是典型範例，通常可改善最大概似估計量的偏誤。

本文選用 LASSO 迴歸模型，這個方法是 Tibshirani 在 1996 年提出，其中 LASSO 譯為「最小絕對壓縮挑選機制」（Least Absolute Shrinkage and Selection Operator），透過最小化以下計算式估計迴歸參數：

$$\text{Objective Function} = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 + \lambda \sum_{j=1}^p |\beta_j| \quad (1)$$

上式 λ 為調整係數及迴歸模型定義如下：

$$Y = \beta_0 + \beta_1 X_1 + \dots + \beta_p X_p + \varepsilon \quad (2)$$

β_0 、 β_1 、...、 β_p 為迴歸參數， X_1 、...、 X_p 為解釋變數， Y 是目標變數， ε 為誤差項。

LASSO Regression之 λ 值可透過交叉驗證（Cross-validation）決定，本文資料分析以學術界常用的R及Python為主，這兩個軟體的 λ 預設值為0.1。

本文以 LASSO Regression 為主要分析工具，同時也考量三種常見的機器學習模型：隨機森林（Random Forest）、極限梯度提升（eXtreme Gradient Boosting, XGBoost）、支持向量機（Support Vector Machine, SVM）。前兩種方法屬於集成學習（Ensemble Learning），藉由結合多個機器學習運算法來提高模型的準確度和穩定性，想法類似「三個臭皮匠勝過一個諸葛亮」，結合不同模型可突破單一模型的限制，截長補短、而且通常模型較為簡單（Breiman, 2001；Chen and Guestrin, 2016）。SVM 的想法恰恰相反，其原則在於尋找一個邊界線或超平面（hyperplane），

有效地分隔不同類別的觀察值 (Cortes and Vapnik, 1995)。SVM 在處理複雜、高維度資料多半會有不錯的結果，但也因為運算量大需要更長的計算時間，而且經常會出現過度配適的問題。

三、探索性資料分析³

本研究主要以 38 家本國銀行、21 家壽險公司及 23 家產險公司的 2022 年 TCFD 報告書或永續報告書為研究分析對象，其中 76% 的本國銀行業單獨公布 TCFD 報告，而壽險業與產險業則分別有 67% 與 49% (表 1)。倘若公司沒有出版獨立的 TCFD 報告書，而將其氣候相關財務資訊揭露於永續報告書，就以永續報告書中的氣候相關財務揭露資訊，內外部評審會依照 TCFD 或永續報告書作為評選依據，本文的分析採用與政大永續中心同樣標準。為了簡化文字說明，本文後續會以 TCFD 報告書代表銀行、壽險、產險業者的 TCFD 或公布於永續報告書的氣候相關財務揭露資訊。

表 1、氣候相關財務揭露資訊揭露來源

來源	銀行	壽險	產險
TCFD 報告書	29	14	11
永續報告書	9	7	12
總計	38	21	23

政大永續中心根據 2021 TCFD 氣候相關財務揭露建議與金管會年發布的「本國銀行業氣候風險財務揭露指引」及「保險業氣候相關風險財務揭露指引」，針對 38 家本國銀行、21 家壽險及 23 家產險公司，分組進行氣候相關風險財務揭露評鑑。此評鑑的性質係屬遵循績效 (Compliance Performance) 的評鑑，與目前國內評鑑企業社會責任、ESG 及永續績效的主要獎項，包括：遠見雜誌所舉辦之「遠見 ESG 企業永續獎」、天下雜誌的「天下永續公民獎」、臺灣永續能源研究基金會的「TCSA 台灣企業永續獎」、研訓院、保發中心及證基會「永續金融評鑑」等，有所不同；該評鑑採用內部人工編碼 (Coding)、文字分析、外部專家學者等三面向進行評鑑。內部人工編碼評鑑由該中心 3 位同仁負責評鑑所有組別；

³ 本文第三與五單元的部分圖表內容，取自作者參與之教育部高教深耕計畫執行成果「2023 TCFD 氣候相關財務揭露年鑑」。

外部專家學者評鑑評審則由各大學專業教師遴聘，分成銀行業、壽險業、產險業三組，每組成員 6 至 7 人。⁴其中內部、外部評審分數先計算平均，之後作為本文分析的目標變數，以迴歸分析及機器學習模型比較估計結果。

圖 1 為內部人工編碼與外部專家學者評鑑分數的箱型圖 (Boxplot)。明顯可見外部專家的分數平均數比較高，而且專家人數比較多、分數卻相對集中。這個現象相當有趣，若以變異係數 (Coefficient of Variation, CV)⁵ 更容易看出變異程度的差異，外部分數的 CV 平均不到內部的 1/3 (表 2)。觀察值變異程度大者，由迴歸模型估計出的震盪幅度也會較大，亦即估計準確率通常也會比較大，下一節會再檢視這個現象是否成立。

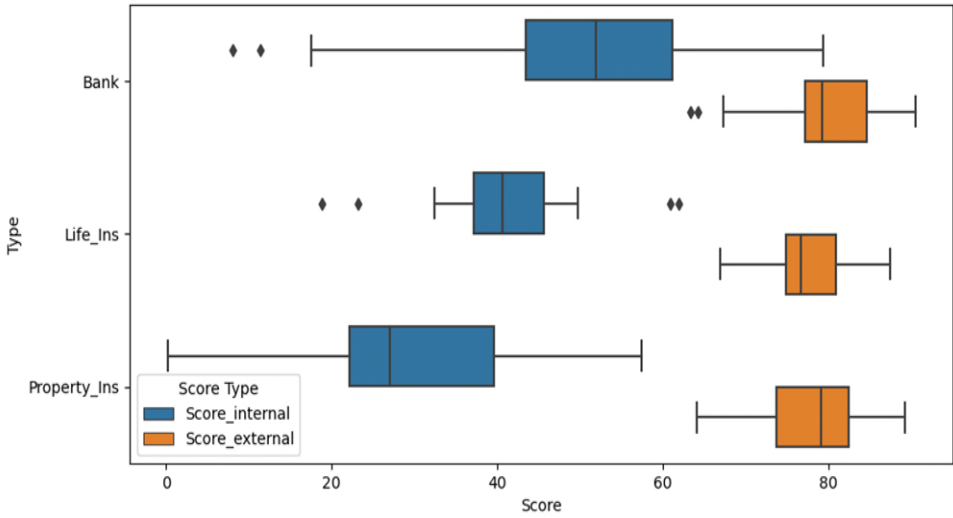


圖 1、內部人工編碼與外部專家學者評鑑分數分佈

表 2、內部人工編碼與外部專家學者評鑑分數的變異係數

專家分數	內部人工編碼評鑑	外部專家學者評鑑
------	----------	----------

⁴ 依據政大企業永續管理中心所出版的年鑑說明，人工編碼小組主要係依據 2021 TCFD Implementing Guidance 建構題組與題項，對受評鑑機構就各大核心要素下各題項作出 0 至 3 分的評分。在專家學者評審方面，則同樣主要依據 2021 TCFD Implementing Guidance，就受評鑑機構提交之報告書，在四大核心要素的揭露品質，填入 60 至 100 分的分數，並取簡單平均作為總分。該中心另外邀請專家學者，針對情境分析內容揭露之品質，在 60 至 100 分的範圍內評分。由於不同 Coding 小組成員和學者專家評審的評分尺度有異，為避免評鑑結果受極端值影響，最後會將 Coding 小組成員及學者專家評審評分高低轉化為排名，計算排名平均值，並將此平均值予以排序。數字越小者，評鑑總排名越佳。關於政大永續管理中心的評鑑方法論，參閱許永明、余清祥、蕭景元、陳仁帥、李榕時、萬宸宇 (2024) 或該中心的網站 (https://cbs.nccu.edu.tw/zh_tw/TaskFocusonclimaterelatedFinancialDisclosures)。

⁵ CV = 標準差/期望值。

本國銀行業	0.25	0.08
壽險業	0.24	0.07
產險業	0.49	0.09

本文以文字分析的角度評估 TCFD 報告書評分結果，由於評分是以治理、策略、風險管理、指標與目標等四個核心要素內容為主，我們先從 TCFD 及永續報告書中篩選四個要素的相關描述，再經過斷詞等處理程序。在此先描述報告書的文字特性，如表 3 所示。本國銀行、壽險、產險業的 TCFD 用字總數的差異非常大，亦即報告書的篇幅也不同。以銀行業的字數最多，產險業則最少。再者，TCFD 之治理、策略、風險管理、指標與目標等四個核心要素，因為在報告中扮演不同角色，文字風格自然也有差異。無論用字（詞）總數或是相異字（詞）彙都明顯不同，其中「策略」的文字總數及多樣性最高、「風險」次之、「治理」則最低，無論本國銀行業、壽險業、產險業都有相同現象。四個核心要素的標點符號使用比例很接近，而且標點類型都是十餘種。不過本國銀行業 TCFD 報告書每個句子的平均字數稍微少一些；產險業報告書的句子則比較長一些。每句字數的計算依據五種斷句標點符號「，。；！？」：逗號、句號、分號、驚嘆號、問號，將介於兩個斷句標點符號之間的文字視為一句。⁶

表 3、TCFD 報告書的文字基本資訊

指標	本國銀行業				壽險業				產險業			
	治理	策略	風險	指標	治理	策略	風險	指標	治理	策略	風險	指標
字總數	29,007	105,204	82,546	46,833	15,323	50,036	36,471	16,059	6,506	22,260	19,675	8,615
每家平均	763	2,769	2,172	1,232	730	2,383	1,737	765	283	968	855	375
字種類	838	1,396	1,229	1,203	752	1,215	1,072	859	583	967	925	731
詞總數	11,995	42,114	33,464	18,365	6,245	20,158	14,777	6,246	2,637	8,983	7,911	3,310
詞種類	2,021	6,560	5,125	3,978	1,376	3,951	3,019	1,724	800	2,142	1,970	1,156
句數	1,419	5,656	4,257	2,590	710	2,421	1,814	759	284	977	892	395
句長	22.52	21.58	21.98	21.45	24.19	23.91	22.86	25.27	26.28	25.91	24.76	27.53
標點總數	2,479	9,789	7,377	4,345	1,241	4,491	3,238	1,510	552	1,906	1,547	854
標點種類	12	15	12	11	13	15	14	11	13	15	12	11
標點比例	0.085	0.093	0.089	0.093	0.081	0.090	0.089	0.094	0.085	0.086	0.079	0.099

⁶ 句子長度是一種判斷寫作風格的方法，像是現代中文大多使用白話文，句長通常比 1918 年五四運動之前的文言文長一些（何立行等人，2014）。

本文以文字多樣性指標作為迴歸模型的解釋變數，在此說明這些指標的判斷方式。表 4 列出隨機抽樣 10,000 個字、100 次電腦模擬的四個核心要素之平均 TTR、Entropy、Simpson 三個多樣性數值，包括以字、詞為計算單位的結果。以隨機抽樣計算多樣性的原因，在於文本篇幅差異不小（如表 3 所示）。然而，由於有常用字彙之故，多樣性未必會隨著字數而增加，因此相異字數 TTR 會隨著字數增加而降低。為了降低字數的影響，通常會隨機抽出某個固定字數，再比較這些字中有多少不同相異字（或字彙）。本文從 TCFD 報告書隨機抽出一萬字，抽出可放回，藉此比較不同公司報告書的文字多樣性。相對而言，以詞為單位的計算結果差異較大，「治理」的多樣性最低、不均度最高，其中屬於多樣性指標的相異字（詞）數 TTR 最小，不均度指標的 Simpson 指標數值最大（或是 Entropy 最小）。在三個產業中，以產險業報告書的多樣性最低、不均度最高。

表 4、本國銀行業、壽險業、產險業 TCFD 報告書的文字多樣性

企業別	要素	字			詞		
		TTR	Entropy	Simpson	TTR	Entropy	Simpson
本國銀行業	治理	648	5.489	0.008	1428	5.802	0.013
	策略	865	5.947	0.004	2490	6.685	0.005
	風險	808	5.830	0.006	2223	6.438	0.008
	指標	834	5.945	0.004	2298	6.764	0.003
壽險業	治理	625	5.456	0.008	1194	5.690	0.013
	策略	827	5.871	0.005	2226	6.539	0.006
	風險	775	5.783	0.006	1973	6.351	0.009
	指標	725	5.832	0.005	1507	6.436	0.004
產險業	治理	546	5.403	0.009	788	5.539	0.014
	策略	769	5.773	0.006	1692	6.266	0.008
	風險	744	5.695	0.008	1617	6.157	0.012
	指標	665	5.725	0.005	1120	6.263	0.004

接下來透過 EDA 挑選與目標變數（評審分數）有關的變數，包括各家銀行及保險業者相關的公司變數，以及 TCFD 報告的文字風格變數，前者為具有固定格式的結構資料，可透過公開資訊搜尋；後者需將報告書的文字結構化，再轉換為具有固定格式的變數。本文考量的迴歸分析變數（附錄 A），本國銀行/壽險/產

險各有 56/54/54 個變數，如果變數與目標變數（人工計分）有關，則可將其置入迴歸模型，在此以幾個變數為例，說明迴歸模型可能會包括這些變數。圖 2 為公司資本額與內部人工計分，顯示兩者有明顯的正向關係：資本額愈高、評審分數也愈高；圖 3 為最常用 500 字 Entropy 與內部人工計分，兩者也有明顯的正向關係：Entropy 愈大、代表用字愈均勻；圖 4 則為相異字比例 TTR 與外部評審分數，同樣也有明顯正向關係，TTR 反映字彙的豐富度。上述這兩個圖形顯示專家學者評審似乎對於寫作風格有一致想法。

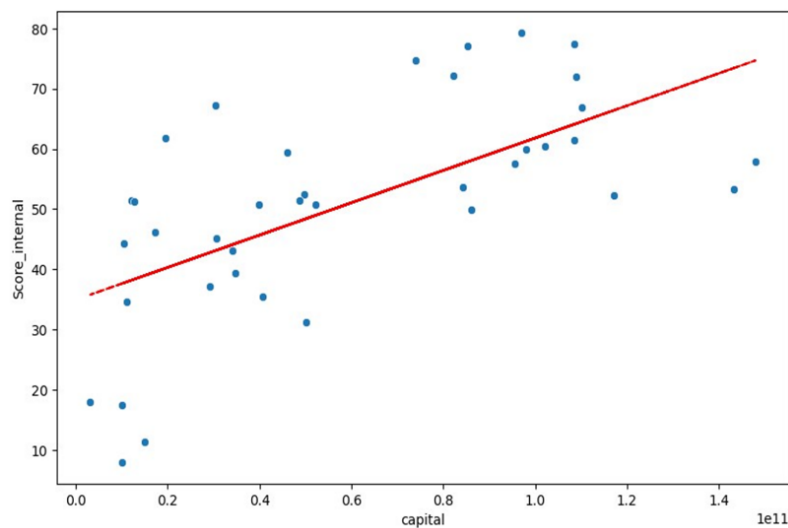


圖 2、內部人工計分與公司資本額的關係圖

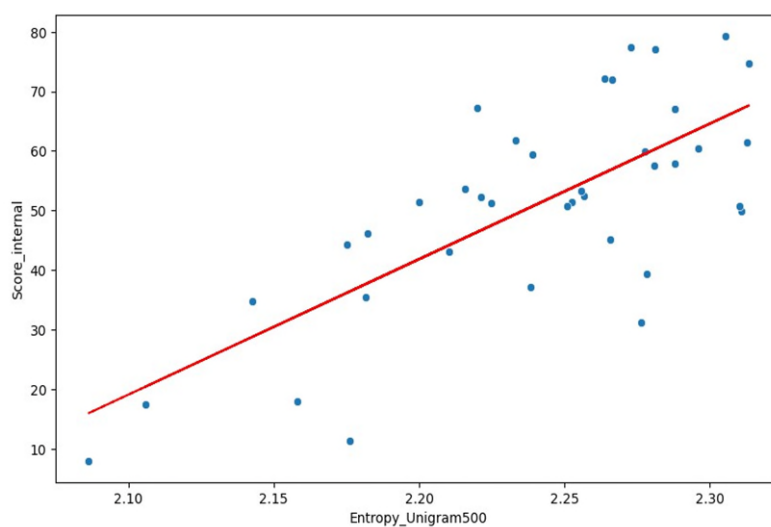


圖 3、內部人工計分與最常用 500 字 Entropy 關係圖

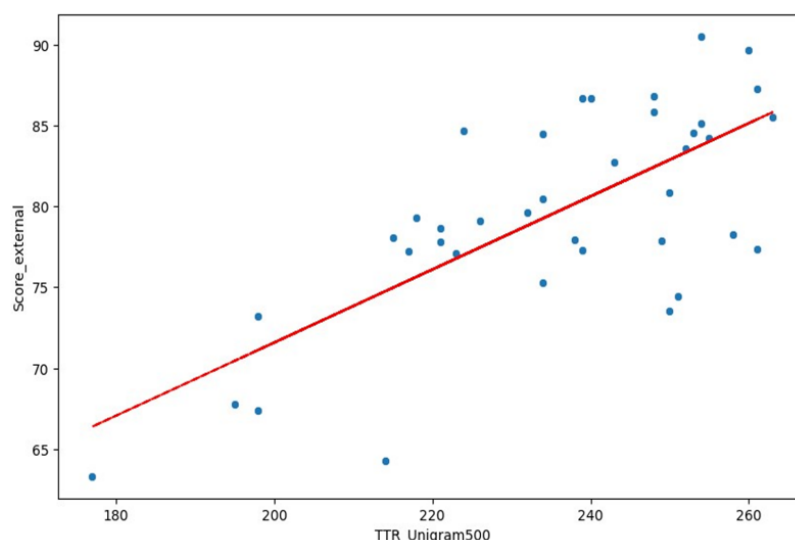


圖 4、外部評審分數與 TTR 關係圖

四、迴歸分析

為了探索公司變數、文字變數與內外部評審分數的關聯，將分成三種變數組合測試迴歸分析：「公司變數」、「文字變數」、「公司＋文字變數」，並呈現迴歸分析的判定係數 R^2 、顯著變數個數。再者，我們也考量不同統計分析結果，包括迴歸分析的殘差檢查，藉此檢視結果是否符合模型的假設，確保分析結果的可靠性。

無論本國銀行、壽險、產險業報告書的內外部評審分數，文字、公司變數各有所長，結合兩者更能捕捉內外部評審的想法，亦即「公司＋文字變數」組合最佳，有最高的 $\text{Adj-}R^2$ 值，如表 5 所示。 $\text{Adj-}R^2$ 為調整後判定係數 (Adjusted R^2)，定義為：

$$\text{Adjusted } R^2 = 1 - \frac{n-1}{n-k-1} (1 - R^2) \quad (3)$$

其中 n 為觀察值個數， k 為解釋變數個數，原意在於懲罰過多的解釋變數，當觀察值個數 (n 值) 不多、卻有較多解釋變數 (k 值) 時與原先 R^2 的差異較大。外部評審分數的迴歸分析結果較佳， $\text{Adj-}R^2$ 至少都有 0.65，其數值都優於內部評審分數，內部、外部評審分數模型的顯著變數估計值及其 p 值等資訊參閱附錄 B。其中，外部評審的迴歸分析結果較佳與變異係數的數值類似，通常變異數值愈大、調整後判定係數 $\text{Adj-}R^2$ 愈小，透過變異係數更可看出差異，意義類似機器學習的信號雜訊比 (Signal-to-noise Ratio，或簡稱為 S/N Ratio)，表 3 的 CV 結果與迴歸模型的 $\text{Adj-}R^2$ 趨勢一致。此外，銀行評審分數的迴歸分析結果較佳，而且三個

產業的迴歸模型中顯著變數又以文字變數較多。

表 5、TCFD 評審分數的迴歸分析結果

	內部			外部		
	顯著變數 個數	Adj-R ²	變異係數	顯著變數 個數	Adj-R ²	變異係數
銀行	9(6)	0.83	0.25	9(5)	0.87	0.08
壽險	3(2)	0.59	0.24	7(4)	0.65	0.07
產險	4(3)	0.58	0.49	4(3)	0.67	0.09

註：括號內為文字變數個數。

迴歸模型是否可用，除了考量上述判定係數、顯著變數之類的資訊，也需注意迴歸分析背後的假設條件，因為挑選解釋變數的方法需要這些條件的配合，或稱為殘差檢查，通常聚焦於誤差是否符合殘差常態性、同質性、模型適合度等統計檢定。本文殘差檢定的選擇包括：常態性檢定 Kolmogorov-Smirnov test、殘差變異數同質性檢定 Goldfeld-Quandt test、模型適合度檢定 Rainbow test。表 5 的六個迴歸模型全部滿足這些檢定，意謂迴歸分析是可行的量化分析模型。除了殘差檢定，我們也檢視迴歸分析結果有沒有產生異常的離群值（Outliers），離群值可作為檢查迴歸分析是否合理，以及觀察值是否有值得注意的觀察值，表 5 中的模型也沒有產生異常的離群值。

本文選擇迴歸模型的主因除了操作簡單、廣為大眾接受外，還考量整個模型的可解釋性，亦即可找出哪些變數與評審分數關係密切。表 5 中迴歸模型的顯著變數包括：銀行上市櫃狀態、（字彙或詞彙）Entropy、（字彙或詞彙）TTR 等，這些變數與評審分數高度相關，可進一步探索背後的相關因素，例如：我們猜測上市櫃公司的比較有規範，易於遵循 TCFD 報告書的建議內容。另外，除了迴歸模型外，我們也嘗試不同的分析模型，包括隨機森林、極限梯度提升、支持向量機，這些模型也能提供不錯的分析結果，本文將比較迴歸分析與這些機器學習模型的分析結果。限於篇幅，本文僅以外部專家評審的分數為例，比較迴歸模型及機器學習模型的結果差異。

本文考慮兩種解釋變數的組合：全部解釋變數（銀行/壽險/產險各有 56/54/54 個變數）、表 5 迴歸模型中的顯著解釋變數，透過交叉驗證比較迴歸模型、機器

學習模型。交叉驗證將資料分成訓練集（Training set）、測試集（Testing set），將訓練集資料得出之模型代入測試集，並以測試集的估計誤差作為評估標準。本文以均方誤差（Root Mean Squared Error, RMSE）為比較依據，也就是估計值的偏誤(Bias)平方、變異數兩者之加總開根號，亦即 $RMSE = \sqrt{(Bias)^2 + Variance}$ ，或是

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (\hat{y}_i - y_i)^2} \quad (4)$$

其中 y_i 為第 i 個TCFD評審分數、 $\hat{y}_i$ 為第 i 個模型估計值、 n 為觀察值個數。另外，以大約4：1比例隨機將資料分成訓練集及測試集，重複這個步驟500次，分別計算500次模擬的RMSE的平均數及標準差，評估LASSO迴歸模型、三種機器學習模型（Random Forest, XGBoost, SVM）的估計成效。交叉驗證通常以測試集的誤差為模型優劣的判斷依據，有時也會考慮訓練集、測試集的誤差值，兩者相差太多代表模型有過度配適的疑慮。

銀行、壽險、產險資料的500次模擬均方誤差可參閱表6、表7、表8，表內數值為訓練集及測試集的平均誤差（AVE）、誤差標準誤（S.D.），這些表格的數值可作為選擇最後模型的依據，用意在於檢視模型是否存在過擬合（Overfitting）的問題。一如預期，機器學習模型在這三個表格的訓練集RMSE誤差相對較小，XGBoost有最小的訓練誤差，數值非常接近0（除了表8產險業）。測試集的RMSE誤差出現非常有趣的現象，變數愈多未必會有誤差比較小，無論是銀行業、壽險業或是產險業，使用較少變數的迴歸模型有最小RMSE，其數值低於代入所有變數的三種機器學習模型。換言之，常見的迴歸模型不但使用較少解釋變數，而且保留變數的可詮釋性，同時模型誤差也足以與機器學習模型抗衡，亦即前一節提及的傳統EDA探索分析仍舊有其必要。另外，如果比較訓練集、測試集兩者的RMSE，使用顯著變數的迴歸模型（表格右側）在銀行、壽險、產險都有最小差異，其平均數值都在2之下，也就是說迴歸模型比較不會有過度配適之虞。反觀Random Forest、XGBoost、SVM這三種機器學習模型的訓練集、測試集，差異多半會大於2，有時甚至大於5或6，差異數值偏大的模型有過擬合的疑慮。

表 6、銀行業外部專家分數模型

Features	All (56)		Proposed (9)	
RMSE	Training	Testing	Training	Testing

	AVE	S.D.	AVE	S.D.	AVE	S.D.	AVE	S.D.
Regression	2.39	0.17	4.08	0.83	2.74	0.18	<u>3.51</u>	0.79
Random Forest	1.78	0.19	3.66	1.01	1.99	0.19	4.02	1.08
XGBoost	0.03	0.00	3.56	1.11	0.05	0.01	4.49	1.10
SVM	1.60	0.32	<u>3.53</u>	1.15	2.12	0.40	3.74	1.11

註：粗體下線者為測試結果最佳之模型，括號內為解釋變數個數。

表 7、壽險業外部專家分數模型

Features	All (54)				Proposed (7)			
RMSE	Training		Testing		Training		Testing	
	AVE	S.D.	AVE	S.D.	AVE	S.D.	AVE	S.D.
Regression	1.67	0.32	6.63	1.64	3.45	0.43	<u>4.55</u>	1.35
Random Forest	3.85	0.37	<u>4.72</u>	1.44	3.99	0.39	<u>4.55</u>	1.42
XGBoost	0.01	0.00	5.69	1.46	0.05	0.02	6.35	1.46
SVM	3.23	0.50	5.49	2.70	3.83	0.51	8.05	5.64

註：粗體下線者為測試結果最佳之模型，括號內為解釋變數個數。

表 8、產險業外部專家分數模型

Features	All (54)				Proposed (4)			
RMSE	Training		Testing		Training		Testing	
	AVE	S.D.	AVE	S.D.	AVE	S.D.	AVE	S.D.
Regression	2.14	0.37	7.41	2.74	3.39	0.24	<u>5.24</u>	1.97
Random Forest	2.93	0.25	<u>5.65</u>	1.44	3.48	0.30	5.85	1.63
XGBoost	1.25	0.26	6.50	1.54	2.13	0.31	5.94	1.58
SVM	4.64	0.50	6.29	1.86	6.17	0.46	6.69	1.82

註：粗體下線者為測試結果最佳之模型，括號內為解釋變數個數。

從表 6、7 與 8 可以看出，銀行業模型的 RMSE 幾乎都比壽險與產險業小，代表銀行業的模型 R^2 最大，亦即估計誤差最小，預測最準確。其經濟意義上的解釋與討論如下：回顧本文研究目的之一在於建立數位模型，探討內外部人工及專家如何評分 TCFD 報告書，嘗試詮釋評審閱讀報告書的觀點。不過，雖然迴歸模型的 R^2 愈大，代表數位模型愈能揣摩評審的觀點，但不等同於評審就是以這些變數評分。因此，銀行業迴歸模型的 R^2 最大，可能原因之一是該模型選擇的解釋變數比較能反映評分標準；相對而言，壽險業、產險業迴歸模型選擇的解釋變數相關性較低，或許加入新的解釋變數會提高這兩個產業的 R^2 。另一個可能性是樣本數的大小，根據表 3 與 4 呈現的 TCFD 報告書之文字特性，每家銀行報告書的字詞總數、種類都較多且豐富，亦即文字資料的樣本數較大，因此無論評審計分、模型分析都比較能夠清楚區隔報告書的不同，因此銀行業模型有較高 R^2 。另

一種可能是資料品質，出版獨立 TCFD 業者以銀行業比例最高（表 1），因此報告書比較符合金融穩定委員會氣候相關財務揭露工作小組所公布之氣候相關財務揭露建議及附錄。

五、結論與建議

本文藉由文字分析的技術，解構政大永續管理中心所舉辦 TCFD 報告書評鑑分數，分析各公司氣候風險報告的寫作風格，區隔四個核心要素之角色與差異，研究結果可以作為協助評審及考核 TCFD 報告書的參考。我們研究結果顯示本國銀行、壽險、產險公司的迴歸模型，皆為可接受的模型，其中文字相關變數佔挑選總變數的一半以上，也顯示文字特徵的價值。所有迴歸模型的殘差分析都符合要求，並以外部評審計分的銀行模型有最小估計誤差，內部人工計分的壽險模型有最大估計誤差。另外，本文以交叉驗證驗證迴歸模型的穩健性，500 次電腦模擬的結果中，迴歸模型不但誤差小於三種機器學習模型（Random Forest、XGBoost、SVM），而且還能保有各個解釋變數的實質詮釋。

本研究原始動機在於協助專家判讀 TCFD 報告，提供模型及方法以建立量化評鑑機制，在分析文字的過程中，我們發現文字使用的多樣性、不均度與報告書的分數有關，像是字彙（及詞彙）愈豐富、分數也愈高，顯示寫作風格與分數高低關係密切。文字分析多半會經過前處理（Pre-process），根據每種文字特性篩選出重要特徵，我國 TCFD 報告以白話中文撰寫，分析前需經過斷詞處理，這與英文報告分析通常會經過字根化（stemming）類似。我們覺得本文的分析理念及步驟可應用於其他國家財務機構的報告書，尤其是選擇文字重要特性的作法可套至其他語言，也就是如何定義文字相關變數（或是資料結構化，Structurization），包括字詞 TTR 及熵之類的多樣性、不均度指標，以及常見於文字分析的詞向量。

受限於本文篇幅及資料數量，本文沒有說明 TCFD 報告的寫作風格分析。常見的風格分析包括關鍵詞及關鍵詞叢，由於現代中文多為白話文，而白話文通常以雙字詞（或多字詞）做為敘事單位，透過常見的關鍵詞及關鍵詞組合往往可一窺文章的大意。例如：「氣候」、「風險」、「管理」是本國銀行、壽險、產險報告書使用比例最高的三個詞彙，表 9 列出四個核心要素中，這三個詞彙之後最常出現的詞彙，似乎都在「氣候」、「風險」、「管理」等詞彙打轉，如為其他詞彙則加上陰影，明顯可見本國銀行、壽險、產險 TCFD 報告書的詞彙略嫌單薄。或許因為氣候相關財務揭露相關指引對業者而言，仍屬於新的領域，各公司尚未熟悉

政府規範的內容及其形式，僅能就熟悉已知的角度切入。未來政府機關在訂定相關揭露指引時，或許可參考其他國家，提供業界比較具體的報告規範。

表 9、四個核心要素的關聯詞組

先行詞	治理	策略	風險	指標
氣候	風險	風險	風險	風險
	管理	情境	管理	目標
	相關	相關	相關	相關
風險	管理	氣候	管理	氣候
	氣候	影響	氣候	管理
	相關	情境	評估	產業
管理	風險	風險	風險	氣候
	氣候	氣候	氣候	風險
	相關	相關	相關	目標

關鍵詞組的分析也可藉由實體辨識(Named Entity Recognition, NER)模型，針對 TCFD 報告書語境開發專屬實體標註系統，如治理要素提取「永續」、「董事會」、「委員」，指標要素：「目標」、「溫室氣體」、「排放」，甚至可將分類與標註結果與資料視覺化工具相結合，協助專家快速檢索並分析報告的重點內容。亦可導入大型語言模型，進行關鍵詞的研究，例如：利用 BERT 語境分析(Contextual Analysis)，評估文本內部的語意連貫性，判斷是否存在矛盾或無關內容。

我們的研究結果顯示迴歸模型可大略推敲出評審計分的想法，可作為檢核人工計分的輔助工具，包括檢查是否有誤植分數等問題，提醒評審及主辦單位，避免出現錯誤。雖然本文的迴歸分析結果還算可以接受，但只使用一年的 TCFD 報告書，分析時並未個別考量四個核心要素的內容，而且也沒有將關鍵詞叢等寫作風格列為解釋變數。未來研究者可以蒐集更多年度的資料，同時分別處理四個核心要素的人工計分結果，逐一檢視每個要素及其迴歸模型，尤其著重於報告文字的格式化，相信更能確定量化分析模型可作為人工評鑑的參考。

參考文獻

一、中文文獻

1. 何立行、余清祥、鄭文惠 (2014)「從文言到白話：《新青年》雜誌語言變化統計研究」，《東亞觀念史集刊》，第 7 期，頁 427-454。
2. 交通部中央氣象署 <http://www.cwa.gov.tw>
3. 余清祥、葉昱廷 (2020)「以文字探勘技術分析臺灣四大報文字風格」，《數位典藏與數位人文》，第 6 期，67-94。
4. 許永明、余清祥、蕭景元、陳仁帥、李榕時、萬宸宇 (2024)「2023 TCFD 氣候相關財務揭露年鑑」，國立政治大學企業永續管理研究中心出版。

二、英文文獻

1. Breiman, L. (2001) “Random Forests”, *Machine Learning*, 45(1): 5–32.
2. Chen, T. and Guestrin, C. (2016) “XGBoost: A Scalable Tree Boosting System”, *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. <https://arxiv.org/abs/1603.02754>
3. Cortes, C. and Vapnik, V. (1995) “Support-vector Networks”, *Machine Learning*, 20(3): 273–297.
4. Eggleton, T. (2013) *A Short Introduction to Climate Change*, Cambridge University Press. ISBN 9781107618763.
5. Friedlingstein, P., O’Sullivan, M., Jones, M.W., Andrew, R.M., Gregor, L., Hauck, J., Le Quéré, C., Luijkx, I.T., Olsen, A., Peters, G.P., Peters, W., Pongratz, J., Schwingshackl, C., Sitch, S., and Canadell, J.G. (2022) “Global Carbon Budget 2022”, *Earth System Science Data*, 14 (11): 4811–4900.
6. Intergovernmental Panel on Climate Change (2023) *Climate Change 2021: The Physical Science Basis*. <https://www.ipcc.ch/report/ar6/wg1/>
7. Krueger, P., Sautner, Z., and Starks, L. T. (2020) “The Importance of Climate Risks for Institutional Investors”, *The Review of Financial Studies*, 33(3): 1067–1111.
8. Li, P.H., Fu, T.J., and Ma, W.Y. (2020) “Why Attention? Analyze BiLSTM Deficiency and Its Remedies in the Case of NER”, in *Proceedings of the AAAI*

- Conference on Artificial Intelligence*, 34(5): 8236–8244.
9. Luccioni, A., Baylor, E., and Duchene, N. (2020) “Analyzing Sustainability Reports Using Natural Language Processing”, *Tackling Climate Change workshop, NeurIPS 2020*. <https://arxiv.org/abs/2011.08073>
 10. Mikolov, T., Chen, K., Corrado, G., and Dean, J. (2013) “Efficient Estimation of Word Representations in Vector Space”, *Proceedings of Workshop at ICLR*. <https://doi.org/10.48550/arXiv.1301.3781>
 11. Sautner, Z., van Lent, L., Vilkov, G., and Zhang, R. (2022) “Firm-level Climate Change Exposure”, *Journal of Finance*, 78(3): 1449–1498.
 12. Shannon, C.E. (1948) “A Mathematical Theory of Communication”, *Bell System Technical Journal*, 27(3): 379–423.
 13. Simpson, E.H. (1949) “Measurement of Diversity”, *Nature*, 163(4148), 688.
 14. Templin, M.C. (1957) *Certain Language Skills in Children; Their Development and Interrelationships*. *Institute of Child Welfare, University of Minnesota Press*.
 15. Tibshirani, R. (1996) “Regression Shrinkage and Selection via the Lasso”, *Journal of the Royal Statistical Society: Series B (Methodological)*, 58(1): 267–288.
 16. Varini, F.S., Boyd-Graber, J., Ciaramita, M., and Leippold, M. (2020) “ClimaText: A Dataset for Climate Change Topic Detection”, in *Tackling Climate Change with Machine Learning (Climate Change AI) Workshop* at NeurIPS.
 17. World Economic Forum (2024) “The Global Risks Report”, <https://www.weforum.org/publications/global-risks-report-2024/>

附錄 A、迴歸分析變數表

附表 A-1、文字及公司變數

序號	本國銀行變數名稱	序號	壽險、產險變數名稱
1	<u>銀行狀態</u>	1	有金控母公司
2	有金控母公司	2	<u>金控母公司是公股</u>
3	<u>銀行上市櫃狀態</u>	3	資本額
4	<u>公股</u>	4	TCFD_會計師確信
5	資本額	5	TCFD_BSI 查核
6	TCFD_會計師確信	6	溫室氣體驗證
7	TCFD_BSI 查核	7	有委任第三方驗證
8	溫室氣體驗證	8	總字數
9	有委任第三方驗證	9	不同字數
10	總字數	10	隨機抽樣 500 字之平均 TTR
11	不同字數	11	隨機抽樣 1000 字之平均 TTR
12	隨機抽樣 500 字之平均 TTR	12	隨機抽樣 2000 字之平均 TTR
13	隨機抽樣 1000 字之平均 TTR	13	隨機抽樣 3000 字之平均 TTR
14	隨機抽樣 2000 字之平均 TTR	14	隨機抽樣 4000 字之平均 TTR
15	隨機抽樣 3000 字之平均 TTR	15	隨機抽樣 5000 字之平均 TTR
16	隨機抽樣 4000 字之平均 TTR	16	隨機抽樣 500 字之平均 Entropy
17	隨機抽樣 5000 字之平均 TTR	17	隨機抽樣 1000 字之平均 Entropy
18	隨機抽樣 500 字之平均 Entropy	18	隨機抽樣 2000 字之平均 Entropy
19	隨機抽樣 1000 字之平均 Entropy	19	隨機抽樣 3000 字之平均 Entropy
20	隨機抽樣 2000 字之平均 Entropy	20	隨機抽樣 4000 字之平均 Entropy
21	隨機抽樣 3000 字之平均 Entropy	21	隨機抽樣 5000 字之平均 Entropy
22	隨機抽樣 4000 字之平均 Entropy	22	隨機抽樣 500 字之平均 Simpson
23	隨機抽樣 5000 字之平均 Entropy	23	隨機抽樣 1000 字之平均 Simpson
24	隨機抽樣 500 字之平均 Simpson	24	隨機抽樣 2000 字之平均 Simpson
25	隨機抽樣 1000 字之平均 Simpson	25	隨機抽樣 3000 字之平均 Simpson
26	隨機抽樣 2000 字之平均 Simpson	26	隨機抽樣 4000 字之平均 Simpson
27	隨機抽樣 3000 字之平均 Simpson	27	隨機抽樣 5000 字之平均 Simpson
28	隨機抽樣 4000 字之平均 Simpson	28	總詞數
29	隨機抽樣 5000 字之平均 Simpson	29	不同詞數
30	總詞數	30	隨機抽樣 500 詞之平均 TTR
31	不同詞數	31	隨機抽樣 1000 詞之平均 TTR
32	隨機抽樣 500 詞之平均 TTR	32	隨機抽樣 2000 詞之平均 TTR
33	隨機抽樣 1000 詞之平均 TTR	33	隨機抽樣 3000 詞之平均 TTR
34	隨機抽樣 2000 詞之平均 TTR	34	隨機抽樣 4000 詞之平均 TTR

35	隨機抽樣 3000 詞之平均 TTR	35	隨機抽樣 5000 詞之平均 TTR
36	隨機抽樣 4000 詞之平均 TTR	36	隨機抽樣 500 詞之平均 Entropy
37	隨機抽樣 5000 詞之平均 TTR	37	隨機抽樣 1000 詞之平均 Entropy
38	隨機抽樣 500 詞之平均 Entropy	38	隨機抽樣 2000 詞之平均 Entropy
39	隨機抽樣 1000 詞之平均 Entropy	39	隨機抽樣 3000 詞之平均 Entropy
40	隨機抽樣 2000 詞之平均 Entropy	40	隨機抽樣 4000 詞之平均 Entropy
41	隨機抽樣 3000 詞之平均 Entropy	41	隨機抽樣 5000 詞之平均 Entropy
42	隨機抽樣 4000 詞之平均 Entropy	42	隨機抽樣 500 詞之平均 Simpson
43	隨機抽樣 5000 詞之平均 Entropy	43	隨機抽樣 1000 詞之平均 Simpson
44	隨機抽樣 500 詞之平均 Simpson	44	隨機抽樣 2000 詞之平均 Simpson
45	隨機抽樣 1000 詞之平均 Simpson	45	隨機抽樣 3000 詞之平均 Simpson
46	隨機抽樣 2000 詞之平均 Simpson	46	隨機抽樣 4000 詞之平均 Simpson
47	隨機抽樣 3000 詞之平均 Simpson	47	隨機抽樣 5000 詞之平均 Simpson
48	隨機抽樣 4000 詞之平均 Simpson	48	全部詞彙詞向量
49	隨機抽樣 5000 詞之平均 Simpson	49	全部詞彙詞向量_主成分分析
50	全部詞彙詞向量	50	前 500 大詞彙詞向量
51	全部詞彙詞向量_主成分分析	51	前 100 大詞彙詞向量
52	前 500 大詞彙詞向量	52	前 10 大詞彙詞向量
53	前 100 大詞彙詞向量	53	文字源於 TCFD 報告書
54	前 10 大詞彙詞向量	54	文字源於 ESG 報告書
55	文字源於 TCFD 報告書		
56	文字源於 ESG 報告書		

註：粗體下線者為其他產業沒有之解釋變數。

附錄 B、迴歸分析估計結果

表 B-1、銀行文字迴歸模型

內部分數 R squared : 0.87				外部分數 R squared : 0.90			
變數名稱	係數	t 值	p 值	變數名稱	係數	t 值	p 值
常數	2465.02	5.47	0.00	常數	865.41	4.68	0.00
詞向量	4.15	2.19	0.04	詞向量	1.46	2.20	0.04
字平均 Entropy	-1113.52	-5.49	0.00	字平均 Entropy	-237.75	-3.86	0.00
詞平均 TTR	0.39	2.57	0.02	詞平均 TTR	0.43	3.45	0.00
詞平均 Simpson	-11790.00	-5.66	0.00	詞平均 Simpson	-4028.96	-4.36	0.00
銀行狀態	13.07	4.37	0.00	銀行狀態	3.72	3.53	0.00
總字數	0.00	-2.61	0.01	詞平均 Entropy	-149.86	-1.99	0.06
不同字數	0.13	2.83	0.01	資本額	0.00	2.27	0.03
公股	12.73	3.60	0.00	TCFD_BSI 查核	-5.36	-2.76	0.01
溫室氣體驗證	17.95	3.37	0.00	有任第三方驗證	6.75	4.34	0.00

表 B-2、壽險文字迴歸模型

內部分數 R squared : 0.65				外部分數 R squared: 0.77			
變數名稱	係數	t 值	p 值	變數名稱	係數	t 值	p 值
常數	4115.68	3.18	0.01	常數	2779.37	3.84	0.00
字平均 Entropy	-2198.66	-3.35	0.01	字平均 Entropy	-1450.85	-3.83	0.00
字平均 Simpson	-26320.00	-2.66	0.02	字平均 Simpson	-21030.00	-3.80	0.00
字平均 TTR	4.68	3.98	0.00	字平均 TTR	8.60	4.35	0.00
有金控母公司	11.30	2.76	0.02	有金控母公司	-5.06	-2.13	0.05
溫室氣體驗證	12.80	2.33	0.04	金控母公司是公股	3.43	1.87	0.08
不同字數	-0.07	-1.85	0.09	有任第三方驗證	3.82	4.24	0.00
				詞平均 TTR	-0.69	-3.74	0.00

表 B-3、產險文字迴歸模型

內部分數 R squared : 0.66				外部分數 R squared : 0.73			
變數名稱	係數	t 值	p 值	變數名稱	係數	t 值	p 值
常數	-75.44	-2.48	0.02	常數	29.13	2.05	0.06
詞向量	12.46	3.62	0.00	詞向量	5.64	4.04	0.00
詞平均 TTR	0.17	2.96	0.01	詞平均 TTR	0.10	4.28	0.00
有金控母公司	11.77	2.36	0.03	有金控母公司	5.14	2.45	0.03
總詞數	0.00	-1.93	0.07	不同詞數	0.00	-2.86	0.01

註：平均 Entropy、Simpson、TTR 皆隨機抽樣 500 個字或詞。