

臺灣漁業犯罪 AI 量刑分類與 Q&A 問答效能之研究

The Study on AI Sentencing Classification and Q&A Performance for Fisheries Crime in Taiwan

吳國清 *

林永清 **

Wu, Kuo-Ching Lin, Yeong-Ching

摘要

本研究以司法院官網有關河川海洋漁業裁判書公開資料，透過內容分析法並撰寫程式進行處理與分析，探索資料表達與深度學習演算法對量刑分類與 Q&A 問答效能之作用為何。研究結果顯示：(1) 從深度學習觀點，機器不會去剖析文字語義的；(2) 以數字取代文字資料表達對於機器學習是可行且有效的；(3) 不同演算法對於量刑分類效能有顯著差異，但不同資料表達對效能改善則不顯著。其中以卷積神經網路效能最佳，其次是以變換器為基礎的雙向編碼器（BERT）；(4) 不同資料表達對於 Q&A 問答效能有極為顯著的差異。在沒有圖形處理器（GPU）情況下，數值向量資料表達的執行效能約為現今市場處理模式的 2 倍，執行時間約縮減一半。

關鍵字：人工智慧、海洋資源、聊天機器人、生態保育、深度學習。

Abstract

This study explores the role of data representation and deep learning algorithms in sentencing classification and Q&A effectiveness using open data from judicial judgments related to river and marine fisheries on the website of the Judicial Yuan. Content analysis and programming were employed for data processing and analysis. The research findings indicate that: (1) From the perspective of deep learning, machines do not analyze textual semantics; (2) Replacing textual data representation with numerical data is feasible and effective for machine learning; (3) Different algorithms exhibit significant differences in sentencing classification performance, but the improvement in performance due to

* 中央警察大學資訊管理學系兼任教授（已退休）。研究專長：程式設計、統計分析、人工智慧。

** 國立澎湖科技大學資訊管理系專任副教授。研究專長：電子商務、資訊管理、統計分析。

different data representations is not significant. Among them, Convolutional Neural Networks (CNN) performed the best, followed by BERT; (4) Different data representations show significant differences in Q&A performances. In the absence of GPUs, vector data representation performs approximately twice as well as current market processing models, with execution time reduced by half.

Key words: Artificial Intelligence, Marine Resources, Chatbot, Ecological Conservation, Deep Learning

壹、前言

人類所賴以生存的地球，海洋面積將近占有 71%。人類生活離不開海洋，且無時無刻都享受著海洋所帶來的便利，國際商品因海洋托運而流通即為顯例。海洋的獨有特徵，對人類產生了巨大公共利益，但社會個體難以從私益角度去維護浩瀚無邊的海洋。自古以來人類依賴海洋、讚嘆海洋具有容納百川與萬物的屬性，但過去人類很少關注到排放或傾倒入河川海洋的污染物或廢棄物，會破壞海洋生態與魚類繁殖生存。近年來，隨著社會發展、科技進步，人類保護海洋河川生態環境的意識逐漸強烈，氣候的惡化與海洋自然資源耗損與生態環境的加速破壞，才讓人類體驗到問題的嚴重性，進而開始考慮如何合理使用與治理海洋資源，並保護海洋生態（湯二子，2022）。在我國《國家海洋政策綱領》中早已揭櫫臺灣是一個海洋國家，確立了以海洋立國的整體發展方針。因此，如何透過海洋教育，培育國民關注海洋、保護海洋的態度是一項嚴肅的重要議題（王嘉陵，2018）。近幾年來，不論是政府、民間企業或學術界，通力合作積極推動相關政策與研究計畫，並發展人工智慧在各種領域的實務應用。其中，屬於人工智慧的聊天機器人（Chatbot），更讓全世界人士所喜愛使用與讚譽，如 OpenAI 的 ChatGPT、Google 的 Gemini 或 Microsoft 的 Copilot 等等。基於此，本文擬以司法院法學資料（裁判書）系統中的海洋與河川生態保育相關漁業犯罪裁判書，作為本研究之研究用材料，針對臺灣漁業犯罪 AI 量刑分類與 Q&A 問答之效能進行研究。透過內容分析法與實驗法，進行統計檢定，希望達到下列目的：

一、從深度學習觀點，探究機器是否依據詞性標註（Part-Of-Speech Tagging），將文章句子中，透過語法結構與詞類、語料庫給予標註出來，再去剖析與了解

人類文字語義呢？即探討自然語言處理是採用詞性學習（POS Learning），或是為特徵學習（Feature Learning）呢？

二、試圖驗證選擇不同資料表達是否會使機器（AI 演算法）失去學習作用？

三、從資料科學與 AI 機器學習觀點，分析不同資料表達對於漁業犯罪 AI 量刑分類與 Q&A 問答之效能是否發生作用，即是否有顯著的差異呢？

貳、文獻探討

一、臺灣海洋資源與保護

地球上百分之七十的面積屬於海洋（王嘉陵，2018），它關係者地球多數生物的生存，也是蘊含龐大的生物棲息哺育場所，同時也會影響全球氣候變遷（葉俊榮，2005）。海洋議題與科技研發息息相關，不論是產業發展或是環境保護都是國家未來的核心政策與議題，而海洋治理除了著重使用開發外，現今對於海洋關懷則更強調生態保育及資源保護（葉俊榮，2005）。自 2001 年以來，我國政府相繼發布《海洋白皮書》、《國家海洋政策綱領》及《海洋教育政策白皮書》，宣示我國開始從大陸國家轉型成海洋國家，並明訂具體步驟與方針，其中，「海洋教育」可說是我國轉向「海洋立國」的重要基礎（周祝瑛，2011）。例如，政府於 2007 年制定的《海洋教育政策白皮書》中主張培育具有海洋公民素養的國民是未來海洋教育的重要任務。近幾年來，伴隨著國家發展，政府開始重視海洋教育的推展（王嘉陵，2018）。

根據黃明和（2012）研究指出，漁業資源係具有再生、洄游、多樣，以及資源共有等特質之自然資源，同時也是供給相對稀有、報酬遞減、漁業水域利用競爭與外部性的經濟資源。魚是確保人類均衡營養與良好健康狀況所需蛋白與必需微量元素的極珍貴資源，全世界人口動物蛋白攝取量中，有 16.6%是來自水產品。然而，海洋看似浩瀚，但其漁業生物資源卻是相當有限，而且生態系統更是十分脆弱。人類的漁業活動業已對海洋生態系造成嚴重影響，已是不爭的事實。目前，全球海洋漁業漁獲量正持續下降，而海洋漁業資源在管理上正面臨一些挑戰與問題，例如，經濟魚類資源，持續面臨捕撈壓力過大、生存環境惡化，以及海洋環境污染日益嚴重等問題皆值得關注。

海洋生態（Marine Ecology）關注在海洋環境中的各種生物與其所處的相互關係所構成的生態系統。它包括海洋中的各種生物種類、它們之間的食物鏈、生態

平衡以及它們對環境的相互作用；而海洋生態保育（Marine Conservation）則關注在保護與維護海洋環境與其中的生態系統，以確保海洋生物多樣性、生態平衡與環境的可持續性，其實質的保育工作內容包括減少污染、限制過度捕撈、保護瀕危物種、維護海洋生態圈等措施而言。換言之，海洋生態保育是一項致力於保護與維護海洋環境、生物多樣性與生態系統的努力，它的目標是確保海洋生態系統的健康與穩定，以促進可持續的漁業、維護生態平衡，並保護威脅的物種。

海洋生態保育之首要工作是優先改善海洋廢棄物的污染問題，而海洋廢棄物可概分為陸源及海源，其中 80%以上為陸地人為或其它活動所產生的。我國於 2018 年 4 月成立了「中華民國海洋委員會」作為海洋事務的最高主管機關，翌年（2019）11 月公布《海洋基本法》，奠定了我國海洋事務之政策方向，藉以更積極推動海洋保護（孫榮聰，2022）。另外，立法院於 2023 年 5 月 31 日通過《海洋污染防治法》部分條文修正，加強地方政府面對海洋污染的應變能力，明確化相關權責；跟進國際趨勢的快速變化與提升國人生態環境保護的意識，充分展現出我國政府為強化防治海洋污染，以利有效管理海洋防污事務，進而維護生態環境永續發展的決心（黃雅琪，2023）。

二、海洋生態系統破壞

在人類科技進步的同時，也造成海洋生物過度利用，致使世界各國紛紛進行海洋資源保育的工作（陳光雄，2005）。人類對海洋造成的破壞是一個嚴重的環境問題，對海洋生態系統與全球環境產生了負面影響。根據黃明和（2012）研究指出，臺灣沿、近海漁業正面臨著下列問題，包括（1）新海洋法公約架構下，造成公海漁場大幅削減及專屬經濟海域內新漁業秩序尚未完整建立；（2）沿、近海漁捕能力過剩，削減不易；（3）漁業環境及資源易受自然及人為因素（包括過度捕撈、棲地破壞、環境污染及氣候變遷等）之影響；（4）支撐沿、近海漁業之保障體系薄弱；（5）漁業使用漁法結構之不合理；（6）海域多元利用競合衝擊；以及（7）採取引發爭議之漁場改造及漁業資源強化措施。

此外，氣候的變遷對海洋生態的影響甚鉅，氣候的變遷對海洋生態、漁業資源與漁民生計都會有舉足輕重的影響，因此對於氣候研究相對應的預測與評估也是海洋生態保育重要的討論議題（劉家良，2010）。根據李國添（2010）研究指出，由於全球暖化、南北極冰層之融化、二氧化碳氣體增加，導致：（1）海平面上升、海水酸化及海洋物種生態改變，且生產力下降；（2）生物之成長週期改變，海洋

及淡水生態系之食物網結構紊亂，進而造成漁業捕獲量與養殖業生產量之減少與不安定。另外，氣候暖化所造成的極端海象現象，在臺灣發生的頻率變高，也給臺灣海洋生態系統與各類養殖業帶來毀滅性之災害。整體而言，氣候變遷所導致的漁業生態衝擊現象，不僅影響臺灣國民食用魚品質，也引發食品安全、健康與居家風險；氣候變遷導致海洋生態系的改變，過去 20 多年來，臺灣漁獲魚種改變及漁獲品質有下降之趨勢；另外，漁民生計的不穩定性，國人食用魚與取得管道也間接受到衝擊。

我國是漁業大國，有關氣候變遷對漁業部門之衝擊及調適或減緩策略，應及早因應。根據 2023 年行政院農業委員會漁業署中華民國臺閩地區漁業統計年報，從表 1 與表 2 顯示，臺灣近五年（2018-2022）各類漁業漁產量（公噸）與漁產值（千元），已明顯逐年呈現直線式遞減現象。

表 1 臺灣近五年（2018-2022）漁產量

場域類別	漁產量（公噸）				
	2022 年	2021 年	2020 年	2019 年	2018 年
遠洋漁業	475,111	531,972	431,902	560,744	621,317
近海漁業	115,325	148,374	150,013	153,102	153,529
沿岸漁業	19,863	20,701	24,549	31,004	26,393
內陸漁撈業	126	143	56	156	3,738
海面養殖業	19,168	20,833	22,413	22,699	24,973
內陸養殖業	245,103	253,977	256,109	269,808	259,433
總計	874,696	976,000	885,041	1,037,512	1,089,382

表 2 臺灣近五年（2018-2022）漁產值

場域類別	漁產值（千元）				
	2022 年	2021 年	2020 年	2019 年	2018 年
遠洋漁業	35,498,911	32,658,442	25,560,838	33,840,306	35,740,304
近海漁業	9,606,447	9,796,239	10,317,881	12,300,101	12,690,104
沿岸漁業	3,352,285	3,286,909	3,406,916	4,438,938	3,605,905
內陸漁撈業	16,254	18,925	3,380	17,299	227,501
海面養殖業	4,639,083	4,932,702	4,767,242	4,672,802	5,254,914
內陸養殖業	28,904,941	27,186,990	27,159,795	31,353,171	31,819,553
總計	82,017,923	77,880,206	71,216,052	86,622,617	89,338,279

三、海洋資源與生態保育相關研究

邵廣昭、林幸助（2010）研究指出氣候變遷影響到海洋生物多樣性的主要因素，包括有水溫上升、海水酸化、海平面上升、洋流型態改變、氣候異常等。文中根據國外文獻，就各種氣候變遷型態對全球海洋生態之影響、各種氣候變遷型態對臺灣海洋生態的可能衝擊及其衝擊程度、各國因應氣候變遷採用之主要調適策略主題提出研究結果，最後，提出臺灣已規劃或已實施之調適策略，以及因應氣候變遷待解決或加強之課題：（1）維護健全海洋生態體系，提昇因應氣候變遷能力；（2）加強海洋生物多樣性因應氣候變遷之研究；（3）加速海洋生物多樣性監測、評估與資訊流通。

呂學榮（2010）研究指出氣候變遷對海洋生態、漁業及漁民的影響充滿不確定性，並歸納氣候變遷對全球海洋漁業生態資源主要衝擊包括八大項：（1）漁業生產力下降及捕獲量降低；（2）捕獲量變動加劇；（3）漁場分佈的改變；（4）漁業成本上升與獲利下降；（5）洪氾、暴潮與海面上升對沿岸、河岸及三角洲漁村與設施的直接衝擊；（6）漁業作業風險的增加；（7）貿易與市場的衝擊，以及（8）生物資源遷移導致的新漁民加入等。文中並就氣候變遷型態對臺灣海洋漁業的可能衝擊及衝擊程度，綜合歸納提出四點可能的衝擊，（1）漁場分佈的改變及捕獲量降低；（2）捕獲量變動加劇；（3）漁撈作業的影響；（4）漁業設施直接衝擊。其中有關捕獲量變動加劇部分，歸因於受到過度捕撈、棲地破壞、污染及溫排水、外來種入侵等人為因素影響，使臺灣海洋生物多樣性遭受破壞。

劉家良（2010）透過行政院農委會漁業署各年份的漁業年報，統整各年份漁業的產量與價值，分出遠洋漁業、近海漁業、沿岸漁業、海面養殖、內陸漁撈、內陸養殖等六類漁業魚種，進而計算出各年份中各漁業及各個魚種的產量、價值、價格、價值百分比、價格百分比。最後經由需求分析估算出臺灣地區的漁業需求函數，繼而進一步分析與評估氣候變遷對於臺灣地區的漁業經濟影響。

由「財團法人綠色和平基金會」與「荒野保護協會」共同發起「臺灣海岸廢棄物快篩計畫」，此為臺灣第一次以系統性、科學化方式調查海岸廢棄物，於 2018 年 7 月開始，將環繞臺灣本島約 1,210 公里的海岸進行快篩工作。它是以視覺評估方式，將海岸分為五大區域，記錄臺灣遭海洋廢棄物污染之分布狀況。該計畫調查發現臺灣海岸廢棄物以各種塑膠垃圾為最大宗，然漁業廢棄物如漁網、繩索及各種浮具的數量、體積可觀，亦不能忽視（孫榮聰，2022）。

黃明和（2012）研究指出，具有高度經濟利用價值的任何一種漁業資源或魚

群，都會受到其自身再生能力的最大極限，以及其所處環境包容力的制約；由於海洋漁業資源在利用上所面臨著：（1）共同財產的悲劇；（2）囚犯的兩難困境；及（3）個人的理性造成整體的非理性結局等現象，無法得到解決，致使得更多的經濟魚類資源，面臨捕撈壓力過大、生存環境惡化及海洋環境污染日益嚴重而導致枯竭的困境。在研究方法上，黃明和（2012）首先根據漁業資源的特性與利用規律，比較漁業管理上傳統漁業管理方法與漁業生態系方法的差異，並根據聯合國糧農組織（Food and Agriculture Organization of the United Nations, FAO）所倡導實施「漁業生態系方法」（Ecosystem Approach to Fisheries, EAF）之原則及操作流程，檢討臺灣周邊海域實施 EAF 的作法並提出應該擴大漁業政策架構及其管理分類，進一步地全面推動實施以生態系為基礎的漁業管理的建議。

王嘉陵（2018）從深層生態學的方向論述人與海洋之間的關係，認為海洋是整體生態環境的一環；人與海洋之間是和諧共存的關係；人類應養成不過度消耗海洋資源的生活型態。楊憶婷、葉庭光（2022）透過問卷調查方式，研究發現父母在海洋科學態度構面對於孩子學習海洋科學有很大的影響，特別是父母對海洋永續發展的責任心與孩子學習海洋科學有顯著的高相關性。

陳光雄（2005）研究指出，增加大眾認知（Awareness）是促進民眾參與行動的重要手段，有了民眾的積極參與，海洋資源方得以改善，改善的成效良否，將影響民眾長期支持保育的態度。增加大眾認知必須由加強保育教育、環境教育和環境體驗三方面著手，培養民眾負責任的環境行為。為了達到「健康海洋」的願景，政府與民眾必須訂定共同的政策目標，政府、非政府組織（Non-Government Organization, NGO）、企業、社區和個人必須分工合作，營造海洋資源保育的學習環境，增進民眾認知，才能促成民眾參與海洋資源保育行動。

綜合上述文獻回顧，可歸納出人類對海洋生態破壞方式大致分為：（1）非法捕撈。例如，未經合法許可或超過限額捕撈海洋生物；（2）海洋棄置垃圾。例如，將垃圾、塑料等廢棄物，或捕撈用具，隨意丟棄在海洋中；（3）破壞海洋生態系統。例如，填海造地、砍伐紅樹林等；（4）不當漁業實踐。例如，使用非法或無法永續的漁獲方法，如以爆炸或漁刺網、投毒等方式漁撈；（5）海邊濫建。例如，在海岸邊未經主管機關環評通過而逕行濫蓋房子、飯店營業，供遊客休憩、度假，快速破壞海洋生態；（6）違法採礦。例如，在海洋中進行非法採礦或鑽油活動，可能導致生態破壞和污染；（7）非法破壞生態保護區。例如，進入、開發或破壞受保護的海洋生態區域（如珊瑚礁區），破壞重要的自然資源和生物棲息地等。

四、臺灣海洋生態犯罪行為與政府因應策略

（一）漁獲量銳減與海洋生態破壞

近年來臺灣漁獲量驟減，是海洋環境被破壞所發出的警訊，而毒、電、炸等非法漁撈行為，更是造成海洋資源加速枯竭的重大原因（蔣謙正，2023）。根據統計，臺灣海洋生態保育主要的犯罪行為包括以下幾種：（1）非法捕撈海洋生物，尤其是瀕危或受保護物種；（2）排放有害廢物和污染物質到海洋中，造成海洋生態污染，例如，排放工業廢水、船隻排放、油污染等；（3）破壞珊瑚礁是對海洋生態系統的重大破壞，例如，遊客潛水、船隻錨泊或採石等行為皆屬之；（4）非法捕獵和貿易稀有或受保護的海洋生物，如瀕危的鮎魚、海馬等行為；（5）盜竊海洋資源，例如，珍珠、貝類、珍貴礦物等，也被視為犯罪行為屬之。由此可見，臺灣的海洋生態保育面臨著一些主要的犯罪行為，包括非法捕撈、違法投棄廢棄物、破壞珊瑚礁和海岸線、以及非法捕捉瀕危物種等。這些行為對海洋生態環境造成嚴重影響，需要透過法規、監管和教育來加以防範和打擊。

（二）提升組織層級與修改法令

2023 年對我國環境保護與海洋生態保育來說，是個關鍵年與重要里程碑。2023 年立法院通過《環境部組織法》，確立「環保署」升格為「環境部」，環境部組織法於 2023 年 5 月 24 日總統華總一義字第 11200043181 號令公布之（環境部，2023）。立法院於同年 5 月 31 日通過《海洋污染防治法》修正案。環境部轄下設立 4 個三級行政機關，分別是氣候變遷署、資源循環署、化學物質管理署、環境管理署。氣候變遷署負責規劃與執行溫室氣體減量及氣候變遷調適事項；資源循環署負責規劃與執行廢棄物源頭減量、資源回收與循環利用及清除處理事項；化學物質管理署負責規劃與執行化學物質管理、災害預防及應變事項；至於環境管理署則負責規劃與執行環境管理與執法、土壤及地下水污染整治事項（謝文哲，2023）。

關於立法院通過《海洋污染防治法》修正案明文規定，為防治海洋污染，針對造成海洋污染行為的罰鍰，從原先最高 150 萬元大幅上修至 1 億元；同時明定中央主管機關應發布《海洋污染防治白皮書》，可向潛在污染行為者徵收「海洋污染防治費」，納入海洋污染防治基金。《海洋污染防治法》經修正後，明定中央主管機關可向潛在的污染行為者，徵收海洋污染

防治費，納入海洋污染防治基金，而海洋污染防治基金則應專供全國海洋污染防治與應變措施、清除、處理、補償及其他有關海洋污染防治工作；強化中央主管機關應處海洋污染事件的應處，主管機關必要時，得逕行採取應變措施、清除及處理污染，相關費用由污染行為人負擔。

（三）政府推動漁業政策

考量整體漁業大環境的改變，政府在沿、近海漁業管理上曾揭示要發展「資源合理利用」與「生態永續」兩者並重的優質漁業，並提出生態海洋漁業願景，向全民承諾兼顧生產、生活、生態之「三生」概念，對珊瑚、魷魷、飛魚（卵）及鯨鯊等漁業，採取逐年執行「禁止或嚴格管制具危害資源或破壞生態」之漁業政策，以期落實生態漁業，進而有助於海洋生態系之平衡；另外，政府也先後採取了如降低漁撈努力量（包括漁船限建、漁船收購、獎勵休魚等）、推動資源培育工作（包括種苗放流、投放人工魚礁等）、推動棲地環境維護工作（主要為漁業資源保育區之設置）、強化保育與管理作為（包括拖網漁業管制、刺網漁業管制、魷魷漁業管制、飛魚卵漁業管制、珊瑚漁業管制、燈火漁業管制、漁業科學證據和技術之相關研究與資料蒐集等）、輔導傳統漁業轉型、擴大教育宣導工作等林林種種的傳統漁業管理作為（黃明和，2012）。

（四）落實犯罪執法

在漁業犯罪執法方面，相關主管和執法機關（包括海洋委員會、內政部警政署），應落實《漁業法》第五章保育與管理、第六章漁業發展，以及第 48 條第一項明定「採捕水產動植物，不得以左列方法為之：一、使用毒物。二、使用炸藥或其他爆裂物。三、使用電氣或其他麻醉物。」等嚴格執法，以確保「資源合理利用」與「生態永續」並重的優質漁業發展。

五、人工智慧與機器學習探討

（一）機器學習

人類學習主要透過「樣式」（Patterns）之識別、記憶、歸納、推理，然後建構一套的思維法則，且在人類所有學習中有 98%是透過圖形識別（Pattern Recognition）途徑（Roiger et al., 2003）。事物的（規則性）重複出現與不斷地學習，將可增強（Reinforcement）人類對它的學習效果。屬於

資訊科學範疇的人工智慧（Artificial Intelligence, AI）試圖將機器（電腦）訓練為能像人類一樣具有思考與學習。經由人類經驗、知識與資訊科學的演算法（程式碼），以讓機器模仿人類的學習行為，累積更多的智慧。因此，學習是人類或機器欲達到具有智慧或知識發現過程之必經途徑。簡言之，人工智慧係指人類智慧賦予機器智慧。

人工智慧的機器學習主要分成傳統學習與深度學習。凡具規則性、特徵性與反覆出現性等事物者，皆適合機器學習對象。然而，現今機器（電腦）大都採用數位來達到存取與複雜運算之目的。問題是，人類對於所發生的現象或所接觸到的事物，包括語言文字，不可能以機器儲存方式來表示與呈現。因此，如何將任何資料（含文字、圖形、音頻、視頻或影像等）轉換成數字（編碼），然後再提供給中央處理器（Central Processing Unit, CPU）、圖形處理器（Graphics Processing Unit, GPU）或張量處理器（Tensor Processing Unit, TPU）處理，是機器學習所要面臨到的挑戰；另一學習挑戰就是學習時間，不過這個問題，隨著資訊科學的進步，及 GPUs、Tensors（Cloud TPU）、專屬電腦語言與學習演算法等技術突破，而獲得顯著的改善，並已廣泛應用在各領域上。然而，雖然解決了機器學習的算力，卻又伴隨著耗費大量電力和機器散熱等問題。隨著雲端運算環境所提供的儲存、算力等大幅提升，機器散熱的改善，加上 5G 的高速網路，使得分散於各地（各國）的巨量資料（Big Data）得以在人類可容忍時間內去大量匯集與進行機器學習。如今，線上聊天機器人的誕生與普及，以及大型語言模型（Large Language Model, LLM）結合了檢索增強產生（或稱檢索增強生成）（Retrieval-Augmented Generation, RAG）的廣泛應用，便是明證。

（二）深度學習

深度學習（Deep Learning），也稱為深度結構化學習，是基於具有表示學習的人工神經網路的廣泛機器學習（Machine Learning, ML）方法系列的一部分。學習策略可以是監督、半監督、非監督或增強等策略。它允許由多個處理層組成的計算模型，來學習具有多個抽象級別的資料表示，且使用「反向傳播演算法」（Backpropagation Algorithm）來指示著機器應如何更改其內部參數。現今，大家所熟知與使用的 LLM 即為一種深度學習模型，透過匯集了巨量的文字資料與學習知識。聊天機器人可以了解使用者所輸

入的提問語義，並從大量的文章、影音、書籍中學習單詞和句子之間的關係，然後回答問題、翻譯、產生文字或聲音、解決數學（如代數和微積分）或基礎統計分析等。然而，這麼強大的功能，背後需要許多人力、人才、時間和巨額成本的。

（三）卷積神經網路

卷積神經網路（Convolutional Neural Network, CNN），或稱深度卷積網路（Deep Convolutional Nets），是一種前饋神經網路，其主要構想是來自動物視覺皮層[大腦視覺]組織的神經連接（傳導）方式。類神經元可以回應一部分覆蓋範圍內的周圍單元，單個神經元只對有限區域內的刺激作出反應，不同神經元的感知區域相互重疊從而覆蓋整個視野。它的運作原理主要包括卷積層（Convolution Layer）、池化層（Pooling Layer）與全連接層（Fully Connected Layer）。卷積神經網路模型為一種受到生物啟發的模型，二維 CNN 已被用於圖形識別上，例如人臉識別和手寫數字識別（Matsugu et al., 2003）；一維 CNN 是 CNN 的一種變體，用於處理一維序列資料，包括語音辨識、自然語言（文字）處理和時間序列預測等；三維 CNN 則已成功應用於 3D 影像識別與醫學影像診斷，如磁共振造影（Magnetic Resonance Imaging, MRI）（Bernal et al., 2019）。

（四）雙向長短記憶網路

長短記憶網路（Long Short-Term Memory, LSTM）屬於循環（或遞歸）神經網路（Recurrent Neural Network, RNN）的一種。LSTM 適合用於時序資料的建模，如文字資料。儘管 LSTM 解決了 RNN 的長期依賴問題，但在實際應用中，發現其訓練時間過長、參數數量太多、容易造成模型過度配適（Overfitting）現象，以及計算資源消耗太大等方面的缺點。

另一種選擇，雙向長短記憶網路（Bi-directional Long Short-Term Memory, BiLSTM）是一種序列處理模型（Sequence Processing Model），由兩個長短期記憶網路（LSTM）組成。此演算法是透過連接兩個 RNN 的輸出來完成的，其工作原理如下：（Geeksforgeeks, 2023）

1. 前向 LSTM：它以正向連續處理輸入序列，從左到右抓取句子與分析文字脈絡。
2. 反向 LSTM：它以反向連續處理輸入序列，從右到左抓取句子與分析文字

脈絡。

3. 它使用了有限序列來預測或標記序列的每一個元素，並易於捕捉文脈資訊（Contextual Information Capture），能夠同時從正向和反向兩個方向學習脈絡資訊。這使得 BiLSTM 能夠更好地理解句子的語義。
4. 特徵學習（Feature Learning）：BiLSTM 可以自動從輸入序列中學習相關特徵，無需手動進行特徵工程。透過雙向處理輸入序列，模型可以從資料中擷取複雜的樣式（Pattern）與表達（Representation），從而在各種任務中獲得更好的效能（Performance）。

由此可見，BiLSTM 是自然語言處理領域的重要模型之一，其主因是有效增加網路可利用的（長短期記憶）資訊量、抓取長距離依賴關係、擷取有意義特徵、防止模型過度配適，可進行平行計算（Parallel Computation）（如採用 Tensorflow），同時比起整合移動平均自動迴歸模型（Autoregressive Integrated Moving Average, ARIMA）或 LSTM 模型有著更好的預測能力（Siami-Namini et al., 2019）等優點。透過同時考慮句子中詞語或字（Words）的前後關係，可以更好地理解文字脈絡。BiLSTM 的主要貢獻在於可對句子進行建模與獲得所有單字（Word）的完整序列資訊（Zhang et al., 2015），並且常用於自然語言處理（Natural Language Processing, NLP）任務中，例如文字分類、機器翻譯、語音辨識等。

（五）BERT

CNN 在處理影像、視頻與音頻等方面帶來了突破，而遞歸網路（Recurrent Nets）則在文字與語音等序列（句子）資料處理上表現突出（Lecun et al., 2015）。一直到了 2017 年，深度學習演算法出現了突破性的進展，就是由 Google 研究團隊，Vaswani 等人在 2017 年第三十一屆神經資訊處理系統研討會中所發表的"Attention Is All You Need"文章（Vaswani et al., 2017）。其做法是，透過一種「注意機制」（Attention Mechanism）將編碼器（Encoder）與解碼器（Decoder）相連接起來，以構成一個嶄新的簡單網路架構（Network Architecture），稱為轉換器（Transformer）。它只關注在這個注意機制上，完全不需要遞歸或卷積，並在各國語言機器翻譯品質上表現更好，包括中英、英德與英法等語言文字翻譯任務。

轉換器是一個避免使用到遞歸模型結構，而完全依靠注意機制來得出

輸入與輸出之間的全域性依賴關係（Global Dependencies）。它允許更多的平行化（Parallelization）處理，可在多顆 GPUs 與 CPU 進行平行工作，且經過較短時間的訓練與學習後，就能達到語文翻譯的良好品質效果。因而，轉換器是一種完全依靠自我注意（Self-Attention）來計算其輸入與輸出的特徵（Feature）或表達式（Representation），而不使用到依序列對齊做法的 RNN 或 CNN 的轉化模型。

在轉換器模型應用上，以轉換器基礎的雙向編碼器表達（Bidirectional Encoder Representation from Transformers, BERT）最具代表性。這個 BERT 轉換器是用於自然語言處理預先訓練的一種技術。BERT 是在 2018 年，由 Google 研究團隊 Devlin 等人所提出（Devlin, et al., 2019），是一個預先訓練的語言表達模型（Language Representation Model）。它不採用傳統單向語言模型或把兩個單向語言模型以淺層連接方法進行預先訓練，而是採用新的標記（New Token）或遮蔽式語言模型（Masked Language Model, MLM）方法，以產生深度的雙向語言表達式。BERT 也是一種自我監督學習（Self-supervised Learning）。BERT 模型構想源自 Vaswani 等人（2017）"Attention Is All You Need"文章提到的轉換器（Transformer）。雙向意指它在處理一個單字（Word）時，會考慮到該字的前面與後面單字的資訊，從而獲取文字脈絡（Text Context）或字脈特徵。

由於 BERT 採用 MLM 對雙向的轉換器進行預先訓練，以產出深層的雙向語言表達式，並經由預先訓練後，只需要新增一個額外的輸出層來進行微調（Fine-tuning）任務。其中[CLS]與[SEP]為演算法主動加入的特殊字符（或稱標記）（Tokens），用於識別句子（Sentence）的起始與終止。因此，它在自然語言處理（Natural Language Processing, NLP）上表現極為出色。然而，它也相對付出一些代價，包括事先下載與儲存空間較大的套件與檔案，如"bert-base-uncased"或"bert-base-chinese"。例如，以遮蔽語言建模進行預先訓練模型之 BERT-base 110 million 與 BERT-large 340 million 等大量參數，容易導致主記憶體不足（Out of Memory），學習時間相當長，而常須借助 GPU 或 TPU，及很耗費電力。

參、研究方法

一、研究方法

內容分析法是一種對文件（如歷史資料、犯罪案件等）內容進行質性與量性分析的研究方法，其目的是經由文件內容本質的事實（Fact）、樣式（Pattern）或趨勢（Trend），試圖發現文件內容的隱性或潛藏價值，並透過內容分析來詮釋它們的價值意涵。基於此，本研究採用內容分析法，結合機器學習與專業統計等軟體，撰寫程式來進行處理與分析。

二、資料來源與研究設備

本研究的研究資料來源是透過司法院網站的【裁判書系統】（網址 <https://judgment.judicial.gov.tw/FJUD/default.aspx>）以人工方式逐一下載歷年與海洋生態保護、違反漁業法相關之司法裁判書，共計 735 件 pdf 檔案。其研究過程中所使用的軟、硬體設備及其用途說明如下：

- （一）MAXQDA：以內容分析法，探索質性資料本質與樣式萃取（Pattern Extraction）。
- （二）電腦語言：程式碼開發與機器學習。Python、Tensorflow、Torch、STATA（do and ado 語言）（StataSE17 with python）、FAISS、Mongo（NoSQL）。
- （三）MongoDB：裁判書檔案儲存與作為 Q&A 問答之回應句子資料。
- （四）FAISS：將嵌入值經轉換向量值，儲存在 FAISS 的檔案中，供 Q&A 問答處理。
- （五）STATA：資料之清洗、萃取、量性資料轉換成質性資料與統計分析。
- （六）硬體：Intel Core i7 中央處理器 1 顆 + NVIDIA Geforce RTX 3050 圖形處理器 1 顆+ 16 GBRAM。
- （七）作業環境：Windows 11，64 位元。

肆、研究結果

一、AI 犯罪量刑分類效能分析

本研究從司法院法學資料檢索系統網站，點選裁判書查詢（<https://judgment.judicial.gov.tw/FJUD/default.aspx>），並由研究人員逐一下載相關

違反我國《漁業法》之論罪科刑法條案件，包括：第 48 條與第 60 條等在內的案件。有關下載裁判書 pdf 檔案內容格式大致可分成兩種，如圖 1（A）與圖 1（B）所示。其中圖 1（A）在匯入與資料清洗上，比圖 1（B）來得複雜許多，因為它含有列編號（01, 02, ...）。這些都是雜訊資料，必須加以去除之，否則無法進行深度學習。此外，在文件資料萃取上，各地方法院的資料表達不盡相同，增加了程式對資料處理的難度。例如，「事實及理由」、「事實」、「犯罪事實」……等。

01	
	臺灣屏東地方法院刑事簡易判決
02	108年度原簡字第143號
03	聲 請 人 臺灣屏東地方檢察署檢察官
04	被 告 鳳羽軒
05	富亞玲
06	上列被告等因違反漁業法案件，經檢察官聲請以簡易判決處刑（
07	108 年度偵字第8326號），本院判決如下：
08	主 文
09	鳳羽軒共同犯漁業法第六十條第一項之非法採捕水產動物罪，累
10	犯，處拘役參拾日，如易科罰金，以新臺幣壹仟元折算壹日。
11	富亞玲共同犯漁業法第六十條第一項之非法採捕水產動物罪，處
12	拘役貳拾日，如易科罰金，以新臺幣壹仟元折算壹日。
13	事實及理由
14	一、鳳羽軒、富亞玲均明知非基於試驗研究目的，復未經主管機
15	關許可，採捕水產動物，不得使用電氣方法為之，竟共同基
16	於非法採捕水產動物之犯意聯絡，於民國108年9月20日15
17	時許，在屏東縣車城鄉台26線8.5公里處（竹坑溪），由鳳
18	羽軒負責以電瓶連結電擊金屬棒電蝦後，再以網子撈起放在
19	岸邊，富亞玲則負責將蝦類撿起裝袋之方式，共同非法採捕
20	水產動物過山蝦30隻（已原地放回）。嗣於同日17時5分許

圖 1（A） 原始裁判書 pdf 檔案內容格式（一）

司法院法學資料檢索系統		匯出時間：112/05/20 01:51
裁判字號：臺灣澎湖地方法院 105 年訴字第 26 號刑事判決		
裁判日期：民國 106 年 01 月 18 日		
裁判案由：違反漁業法		
臺灣澎湖地方法院刑事判決 105年度訴字第26號		
公 訴 人 臺灣澎湖地方法院檢察署檢察官		
被 告 陳怡銘		
陳志達		
王榮進		
上列被告因違反漁業法案件，經檢察官提起公訴（105年度偵字第411號），被告於準備程序中均就被訴事實為有罪之陳述，經本院合議庭裁定由受命法官獨任進行簡式審判程序，並判決如下：		
主 文		
陳怡銘共同犯漁業法第六十條第一項之非法採捕水產動物罪，處有期徒刑壹年參月，緩刑伍年，並向公庫支付新臺幣貳拾萬元。		
陳志達、王榮進共同犯漁業法第六十條第一項之非法採捕水產動物罪，均處有期徒刑壹年，均緩刑肆年，並各向公庫支付新臺幣拾萬元。		
扣案之電槍壹支、安培錶壹個、電池壹顆、變電器壹個、電池外殼壹個、漁獲肆拾壹公斤均沒收。		
事 實		
一、陳怡銘為澎湖籍CT2-6289「瑞安六號」漁船之實際所有人兼		

圖 1（B） 原始裁判書 pdf 檔案內容格式（二）

（一）處理方法

有關本研究探索 AI 犯罪相關量刑效能上，採取方法如下：

1. 量刑分類：在 AI 犯罪量刑分類方面，本研究分為{罰金(=0)，拘役(=1)，有期徒刑(=2)}等三種量刑。
2. 分類學習：用以處理犯罪量刑分類深度學習演算法，包括前饋神經網路的一維卷積神經網路（One-Dimensional Convolutional Neural Network, 1D CNN），雙向長短記憶網路（BiLSTM），以及從轉換器雙向編碼器表達（Bidirectional Encoder Representations from Transformers, BERT）等三種。現今，這些演算法常被用於自然語言處理（Natural Language Processing, NLP）。
3. 資料表達：本研究分為{文字性，編碼（數字）性}等兩種資料表達。

（二）處理程序

有關本研究對於 AI 犯罪相關量刑分類效能之處理程序如下：

1. 匯入與儲存。撰寫 python 程式，將全部法院裁判書 pdf 檔案匯入。去除

- 雜訊字後，以原始文字性資料分別儲存到 MongoDB 資料庫與 CSV 檔案格式中。其中這個資料庫的資料集檔案將永久被保留且不會變更它。
2. 並行處理。(1) 透過 STATA 專業統計軟體 (STAT/SE 17 版)，將 CSV 檔案匯入到 STATA 環境中。(2) 另開啟質性資料分析軟體——MAXQDA (2022 版)，將全部法院裁判書 pdf 檔案匯入到 MAXQDA 「Windows」視窗中。(3) 透過 MAXQDA 與人工來檢視這些裁判書內容的錯別字 (如補撈、電亟等)、樣式 (Patterns) 等，然後到 STATA 的「Do-file Editor」程式編輯區中撰寫 do 程式碼，以更有效率地來進行文件的清洗工作。意指 MAXQDA 提供瀏覽內容，STATA 用於同步撰寫與執行程式碼、修刪內容等。其中 STATA 的 local 指令功能相當強大，可簡化程式撰寫的複雜度，這是 python 所無法做到的。當完成文件內容清洗工作後，將正確結果儲存到「法院裁判書.csv」檔案格式 (UTF-8 碼) 中。
 3. 建立樣式取代表。開啟 Microsoft Excel，建立一個含有 {樣式, 取代} 兩個欄位名稱。其中 [樣式] 欄位可為規則表達式 (Regular Expression) 或停用詞，如「捕[漁撈獲]」；[取代] 欄位內容將是用來準備取代 [樣式] 欄位內容的。若屬於停用詞者，則 [取代] 儲存格內容填入 "NA"，以便後續資料轉換時，再以空字串取代之。然後儲存到「Stop_Pattern_Replaced_Words.csv」檔案。
 4. 取得資料一致性。透過 Python，同時匯入「法院裁判書.csv」與「Stop_Pattern_Replaced_Words.csv」檔案，以便進行全部裁判書內容的取代。為了避免「Stop_Pattern_Replaced_Words.csv」內容的重複性或編排的不一致發生，我們採用 Python 的 re.split ()、strip ()、lower () 等函式，以及 set 型態來維持它們的唯一性與資料表達的一致性。然後再度存回「Stop_Pattern_Replaced_Words.csv」檔案中。
 5. 進行資料轉換。這是一項相當耗費人力與時間的工作，大約占全部工作量的 70%。首先，開啟 MAXQDA 與匯入全部法院裁判書 (pdf)，一方面透過【Analysis > Text Search & Autocode】功能，另一方面則同步撰寫 Python 程式碼，這兩者是以「並行作業」方式，將質性的關鍵字、樣式或代碼 (Code) 等資料，由 Python 程式轉換成量性資料，其中這些關鍵字、樣式或代碼等名稱即作為欄位名稱。例如，罰金新臺幣參萬伍仟元 (國字)，將轉換成：35000 (數字) 儲存到 [罰金] 欄位中。裁判書中有載

- 明"吳郭魚"者，我們就賦予[吳郭魚]欄位給予 1，否則給予 0。處理與儲存單位是個很重要的概念。例如，有期徒刑的處理單位可能為年月、年、月等三種情況，我們必須換算成以月為單位來儲存之，比如國字「參年肆月」將被換算成數字「40」。然後，將它儲存到「裁判書_量性資料.csv」（UTF-8 碼）。此外，產出[法院_犯罪行為]與[目標變數]兩個欄位與其內容。前者儲存全部已被清洗過的裁判書內容，後者儲存{0（罰金），1（拘役），2（有期徒刑）}。這個[目標變數]即作為監督式學習策略的 Target，{0, 1, 2}即為它的標籤值（Labels）。檔案名稱為「深度學習_資料集.csv」。
6. 資料表達。在 Python 編程環境中，匯入「深度學習_資料集.csv」，撰寫一個 class，其類別名稱：資料表達策略。它含有三個核心函式：第一者為 `def Stopwords_Pattern_Processing (self)`，它是用來根據已建立好的「Stop_Pattern_Replaced_Words.csv」檔案停用詞與樣式對全部裁判書進行取代工作。其次為 `def Character_Oriented (self)`（參見公式 (1)）。第三者為 `def Encoding_Oriented (self)`（參見公式 (2)）。
7. 執行效能比較。其主要目的有二：一者是分類效能（Classification Performance）的比較；另一者是 CPU+GPU 的執行效能（Execution Performance）的比較。當然這些比較性，只適用於檔案或文件內容絕大部分或全部是以中文或全形字的資料表達為主。有關處理資料表達用公式如下：

$$\Phi_{list} = \sum_{i=0}^{N-1} \prod_{j=0}^{m-1} D_i (C_j || space) \quad (1)$$

$$\Psi_{list} = \sum_{i=0}^{N-1} \prod_{j=0}^{m-1} D_i (E_j || space) \quad (2)$$

其中，(1) || 表示字串的串接（Concatenation）符號。每一個字元 C 或編碼 E 後面串接著一個半形空白字元之目的，在於受到這些深度學習演算法要求每一個標記（Token）長度 $|T| \geq 2$ 的限制條件；(2) Σ 為 sublist 的添加（Append）符號， Π 為 string 的合併（Merge）符號；(3) N 為全部裁判書的總件數，它是固定值；(4) m 為每一裁判書所含有的總字數，它是可變值；(5) D（Document）代表正被處理的裁判書；(6) C 為任何字元，可能為中文、英文、數字或任何符號。E 為 C 的相對應編碼值（Encoding Values）

或經數學函數計算所得數值，它由一組數字所組成的，須具唯一存在性，相同的字元必須對應到同一個編碼值或數值。如 UTF-8, Unicode Little-Endian or Unicode Big-Endian, Big-5, Index, 或者經由線性數學函數 $AX+B$ 轉換後所得值等。這些函式執行所需時間必須納入到整體效能考量中。舉例來說，句子："被告因違反漁業法等案件，"，其處理過程與結果如圖 2 所示：

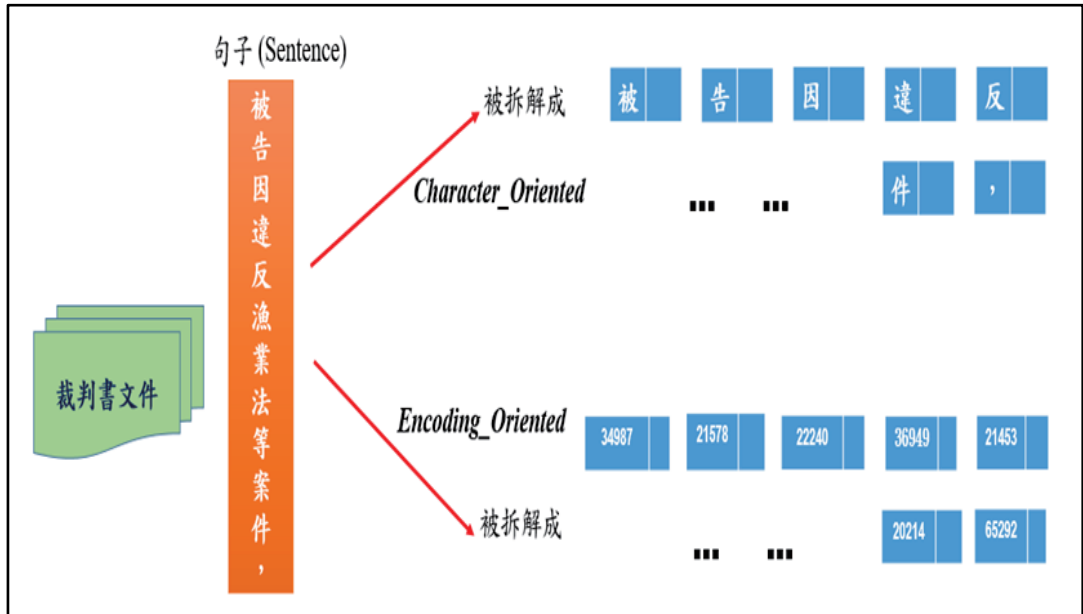


圖 2 文件拆解句子後的資料表達處理方式

註：Unicode(被) = 34987，Unicode(告) = 21578，以此類推。當然也可採用 Big5。

一般情況下， $Time(\Psi_{list}) > Time(\Phi_{list})$ ，因為編碼值必須透過計算而得的。故而我們合理推測，這兩者在不同或同一深度學習演算法過程中所需執行時間有可能是不同的，而且它們的整體執行時間可能也會有所差異的。因而在同一個深度學習演算法，且參數、超參數與隨機取樣等相同條件下，本研究提出研究假設檢定一：

$$H_0: \text{Total}_{\text{ExecutionTime}}(\Psi_{list}) = \text{Total}_{\text{ExecutionTime}}(\Phi_{list})$$

$$H_1: \text{Total}_{\text{ExecutionTime}}(\Psi_{list}) \neq \text{Total}_{\text{ExecutionTime}}(\Phi_{list})$$

（三）處理結果

本研究採用了三種深度學習演算法，每一種演算法須分別處理字元導向與編碼導向的兩種不同資料表達，每一種資料表達則被反覆執行 5 次，因此，共計執行 30 次。這些實驗皆在電腦使用 AC 市電，而非在 DC 蓄電池情況下來完成的，因為前者比後者在執行時間上會縮短一些。有關執行結果如表 3 所示。從表 3 統計數據得知：

1. 在效能表現上，雖然 CNN 的準確率（82.52%）比 BERT（81.138%）略高些，但 BERT 的整體執行時間（444.274 秒）約為 CNN（8.766 秒）的 51 倍；而 BiLSTM 的整體表現平平，約為 75%。
2. 在資料表達對效能作用上，字元導向與編碼導向差異很小。本研究透過 ANOVA 變異數分析結果顯示，即使編碼導向整體執行時間略高於字元導向，但並未達統計 0.05 顯著水準，資料表達對 Total Execution Time 的 ANOVA 變異數 F 機率值為 0.9552。因此，本研究的假設檢定一結果是接受 H_0 。這意味著，標記器（Tokenizer）對於這兩種資料表達在處理速度上差距甚小。
3. 模型預測三種量刑能力上，CNN、BERT 與 BiLSTM 表現並不一致。CNN 排名：precision_{罰金}>precision_{有期徒刑}>precision_{拘役}；BERT 排名：precision_{罰金}>precision_{拘役}>precision_{有期徒刑}。BiLSTM 則以拘役表現最好。這種差異性可能來自樣本數太少，也可能來自各級法院的法官對這三種量刑裁判不一致所產生的。例如，應判決有期徒刑案件，結果法官處以拘役論斷。
4. 表中未呈現 BERT Encoding-Oriented 效能表現數據。這是因為它在編碼資料表達的 Training loss 曲線變化上，呈現鋸齒狀起伏變化，不因回合數（Epochs）的學習次數增加而趨於平穩（收斂）狀態，且整體準確率低於 63%（未達 75%以上的預測價值），意指 BERT 可能不適合用於全部以數字資料表達的字串資料類型（String Data Type），但也可能只是不適合本研究所收集到的裁判書。本研究認為，主要受到標記器（Tokenizer）只能接受輸入文字性資料類型的限制，使得我們不能使用數值（如整數或浮點數）資料類型輸入到標記器與處理。
5. 機器學習非依據詞性標註去剖析與了解人類所用語文（文字）語義的，意指它不採用「詞性學習法」，取而代之的是採用「特徵學習法」。關於這一點，我們已選擇兩種不同資料表達方式（含文字與數字）獲得證明，

因機器學習結果仍不會失去學習作用。

6. BERT 深度學習所需訓練時間相當長，很消耗電力（算力）和記憶體資源（龐大參數的數量），相較於 CNN Conv1D 效能表現就相當出色，其整體準確率將近 83%，整體執行時間不到 9 秒。建議同時使用 GPU 處理器。

表 3 CNN、BiLSTM 與 BERT 在裁判書量刑分類效能的表現（N=75）

深度學習 演算法	效能指標	量刑 等級	Precision (%)	Recall (%)	F1-Score (%)
CNN Conv1D Character-Oriented	Accuracy=82.52% Total Execution Time= 8.766 秒	罰金	100	37.8	54.4
		拘役	81.4	98.8	89.4
		有期徒刑	83.2	60	69.8
CNN Conv1D Encoding-Oriented	Accuracy=82.70% Total Execution Time=8.778 秒	罰金	100	47.4	63.6
		拘役	81.2	99	89
		有期徒刑	85.2	56.8	68.2
BiLSTM Character-Oriented	Accuracy=74.94% Total Execution Time=63.972 秒	罰金	71	53	59.6
		拘役	78.6	86	82
		有期徒刑	70	60.8	64
BiLSTM Encoding-Oriented	Accuracy=73.734% Total Execution Time=64.146 秒	罰金	63.8	73.4	68.4
		拘役	76.2	84.2	79.6
		有期徒刑	70	52.8	59.8
BERT Character-Oriented	Accuracy=81.138% Total Execution Time=444.274 秒	罰金	100	56	72
		拘役	86	88	86.6
		有期徒刑	70.6	74.6	72.2

註：（1）BERT *Encoding-Oriented* 因受到標記器必須輸入文字的限制而無法使用，故不列入比較。（2）表中各值為每一個演算法重複執行 5 次後取其平均值。

前揭表 3 是從機器學習觀點，即透過深度學習來獲得分類效能表現，但是對於它們彼此之間的統計意義仍然未知。基於此，本研究擬從統計學觀點，來更深入了解它們的實質影響與其彼此之間的交互效果（Interaction Effect）。做法上，本研究透過 STATA 統計軟體所提供的 MANOVA 指令進行多因子多變量變異數分析（Multivariate ANalysis Of VAriance, MANOVA）。其中多因子包括{演算法，資料表達，量刑}，相依變數為{Precision, Recall, F1_Score, Accuracy, Total_Execution_Time}，觀察樣本數為 75 筆。

由表 4 的 MANOVA 統計結果得知：

1. 在個別效果檢定方面，不同的資料表達並不會影響到各種效能的表現。然而，不同的演算法與不同量刑等級，對效能表現有著極為顯著的差異（參見前揭表 3）。
2. 在交互效果檢定方面，（1）量刑同時對演算法、資料表達具有相當顯著的交互作用。不同演算法影響到效能表現十分顯著，然而，資料表達對量刑的交互作用（效果）雖然達到統計顯著與其解釋意義（機率值小於 0.05），但仍然遠不如演算法的作用。（2）演算法、資料表達、量刑三者同時交互作用，雖然 Roy's largest root 已達到小於 0.05 的顯著水準，但其 Wilks' lambda 機率值卻大於 0.05。在實務應用上，仍以 Wilks' lambda 作為判斷依據。因此，我們難以論斷這三個因子同時交互作用將對效能表現有著顯著的作用。

表 4 CNN、BiLSTM 與 BERT 對量刑分類效能多變量變異數分析（N=75）

Source	Statistic		df	F (df1, df2)		F	Prob>F
Model	W	0.00000	14	70	270.7	67.06	0.0000 a
	R	1000.000		14	60	4328.59	0.0000 u
演算法	W	0.0001	2	10	112	1016.97	0.0000 e
	R	841.9657		5	57	9598.41	0.0000 u
資料表達	W	0.9162	1	5	56	1.02	0.4122 e
	R	0.0915		5	56	1.02	0.4122 e
量刑	W	0.036	2	10	112	47.8	0.0000 e
	R	8.1466		5	57	92.87	0.0000 u
演算法#資料表達	W	0.9109	1	5	56	1.1	0.3732 e
	R	0.0978		5	56	1.1	0.3732 e
演算法#量刑	W	0.1412	4	20	186.7	7.51	0.0000 a
	R	2.6624		5	59	31.42	0.0000 u
資料表達#量刑	W	0.7002	2	10	112	2.18	0.0236 e
	R	0.4116		5	57	4.69	0.0012 u
演算法#資料表達 #量刑	W	0.8072	2	10	112	1.27	0.2582 e
	R	0.2383		5	57	2.72	0.0286 u
Residual			60				
Total			74				

註：e=exact, a=approximate, u=upper bound on F; W=Wilks' lambda, R=Roy's largest root.

二、AI 的 Q&A 問答效能分析

ChatGPT (Chat Generative Pre-trained Transformer)，稱作「聊天產生預先訓練轉換器」，是由 OpenAI 公司所開發的人工智慧聊天機器人程式名稱，於 2022 年 11 月推出。該程式使用 GPT-3.5 (<https://chat.openai.com/>) (免費)、GPT-4 (付費) 架構的大型語言模型並以強化學習訓練。ChatGPT 提供人類所熟知的自然語言文字或語音與機器人進行 Q&A 問答或互動，自動產生文字、自動問答、自動摘要等多種任務。此外，Google 於 2023 年 3 月 21 日推出實驗版的 Bard 與 ChatGPT 相抗衡的 AI 聊天機器人，後來更名成 Gemini (<https://gemini.google.com/>)。

(一) 處理方法

本研究為能實現臺灣「企業 AI 落地化」起見，撰寫了為數不少的 Python 電腦程式碼，及其存取 MongoDB 與 FAISS 等資料庫用的專屬語法程式碼，來實作一個簡易型 AI 機器人 Q&A 問答雛型。

由前一節「AI 犯罪量刑分類效能分析」中得知，BERT 對於效能預測上是採用 Tokenizer，它不能以向量或陣列之數字資料類型作為輸入值。然而，當我們想要讓 BERT 的編碼器產出嵌入值時，就可改採用 SentenceTransformers，它的輸入資料就可不用經過預先資料標記化 (Tokenization) 的處理程序，故不用受到必須事先為文字資料類型的限制。基於此，本研究提出下列的四種不同處理方式 (Processing Approach)：

1. 字元導向 (Character-Oriented) 且字串資料表達+CPU。
2. 編碼導向 (Encoding-Oriented) 且字串資料表達+CPU。
3. 編碼導向 (Encoding-Oriented) 且數字向量資料表達+CPU。
4. 編碼導向 (Encoding-Oriented) 且數字向量資料表達+CPU+GPU。

其中第一種處理方式即為目前資料科學家在處理 AI 所最廣泛採用的模式，剩餘三種是由本研究所提出的，提出此構想源自中文的特性與嵌入維度 (Embedding Dimensions) 所持有的嵌入[特徵]值的唯一性，其目的在於探索這三種是否可縮短執行時間。此外，理論上，圖形處理單元 (GPU) 在提升 AI 整體執行效能有著極大的作用，至於其效果為何，則有待統計檢定方能得知。

CUDA 是由 Nvidia 所發展出來的一個平行計算平臺 (Parallel Computing Platform)。軟體開發人員可在這平臺上使用 GPU+CPU 與 python 所支持的

Tensorflow 與 Pytorch 等電腦語言，進行相關人工智慧的程式設計，其目的在於縮短整個程式執行所需時間。值得一提的是，GPU 是不能在沒有 CPU 環境下去獨自執行程式，單獨作業的。

為了確保試驗結果可比較性起見，本研究對每一種處理方式各執行 10 次，故可收集到 40 筆 Q&A 問題的所需執行時間數據，並且提出下列的研究假設檢定：

假設檢定二： H_0 : 資料表達與Q&A執行時間無關

H_1 : 資料表達與Q&A執行時間可能有關

假設檢定三： H_0 : 機器處理器與Q&A執行時間無關

H_1 : 機器處理器與Q&A執行時間可能有關

(二) 處理程序

1. 資料下載。將從司法院網站的【裁判書系統】下載有關與海洋生態保護相關的違反漁業法裁判書共計 735 件 pdf 檔案，以作為本研究實作的資料集。然後將它們儲存到 MongoDB 資料庫中。
2. 拆解與儲存。載入資料庫全部文件集 (Collection)，進行文件斷句 (非為斷詞) 處理任務。其處理單位: Sentence，拆解用分隔符號為[,;。! ?]，刪除 Sentence 長度小於 2 後，得到句子總數量為 70,273 句，然後以 S (ObjectId (序號), 句子) 配對方式儲存到 MongoDB 資料庫中。其中 [ObjectId] 欄位值，可由資料庫系統自動產生，亦可由設計人員自訂。但是，它的儲存格值必須滿足唯一性與不可存在缺失或空白值 (NIL/NULL)。
3. 句子資料表達。本研究將以句子作為拆解單位，將產出的全部句子 (=70,273) 分別進行三種不同資料表達 (參見圖 3) 與儲存。
4. 標記編碼。本研究採用 BERT 的 SentenceTransformer ('all-MiniLM-L6-v2') 來建立 Transformer 模型。它的 embedding dimensions = 384，即以 384 個特徵 (嵌入值) 來記錄著這 735 篇裁判書計 70,273 句子的重要文脈 (Text Context) 與樣式 (Patterns)。其中 SentenceTransformers 是來自 Sentence-BERT 的一個 Python 框架，用於句子、文字或影像等嵌入值的產出。有關建構嵌入與嵌入維度的主記憶體運作情形，如圖 4 所示。在圖中右側嵌入值個數總與即為嵌入維度的數量，計 384 個 32 位元儲存

空間，裡頭儲存著浮點數值。相同的標記（Token）持有相同的維度（特徵或嵌入）值。此外，圖中的 index 為 FAISS 向量資料的索引空間對應值。這樣就能把 BERT SentenceTransformers 所建立的嵌入值與向量資料檔案所建立與儲存到硬碟中的 index 相連結或映射（Mapping）起來。這是相當重要的概念。

5. 提問內容處理。使用者輸入一串文字（例如，"被告駕駛漁船出海，使用炸藥及引爆之雷管，以炸魚方式捕魚。"），同樣去進行斷句處理與使用 SentenceTransformers 以產出相對應嵌入值——Promt_Encode。然後將這些值放入向量檔案已建立好 index 的搜尋函式中，再去呼叫向量資料庫 FAISS 的 index.search (Promt_Encode, k) 函式進行內積搜尋。其中 k 為想要回應給使用者的句子數量。
6. 問與答相似度比較。依 FAISS (Facebook AI Similarity Search)¹ 所產出相似度最高的前 k 個 index 值，以作為回答使用者提問 (Prompt) 的結果。這些索引值 ([ObjectId]) 不得空缺值，並具唯一性。它們必須事先被轉換成 MongoDB 的 ObjectId 值，這些值可由 MongoDB 自動提供或使用者自訂內容與其長度。一般預設長度與內容為 24 個十六進制字元的字串所組成，例如，ObjectId ('4aba160ee23f6b543e000002')，然後去查詢 MongoDB 資料庫中符合 ObjectId 值條件且回應其相對應的句子，乃因一般向量資料庫或 FAISS 所儲存內容皆為一長串的向量數值（嵌入值）而非文字，當然人們是看不懂它們所表達的語義或意義的。

¹ 向量資料庫 (Vector Database) 為一提供完整的儲存、管理與搜尋等向量儲存體。FAISS 未提供管理功能，故被定位為一種程式庫 (Library)，並作為建立向量資料庫的有用工具。亦有人視它為一簡易型向量資料庫。請參見 https://bard.google.com/chat/226cb7872568e12e?hl=zh_tw。

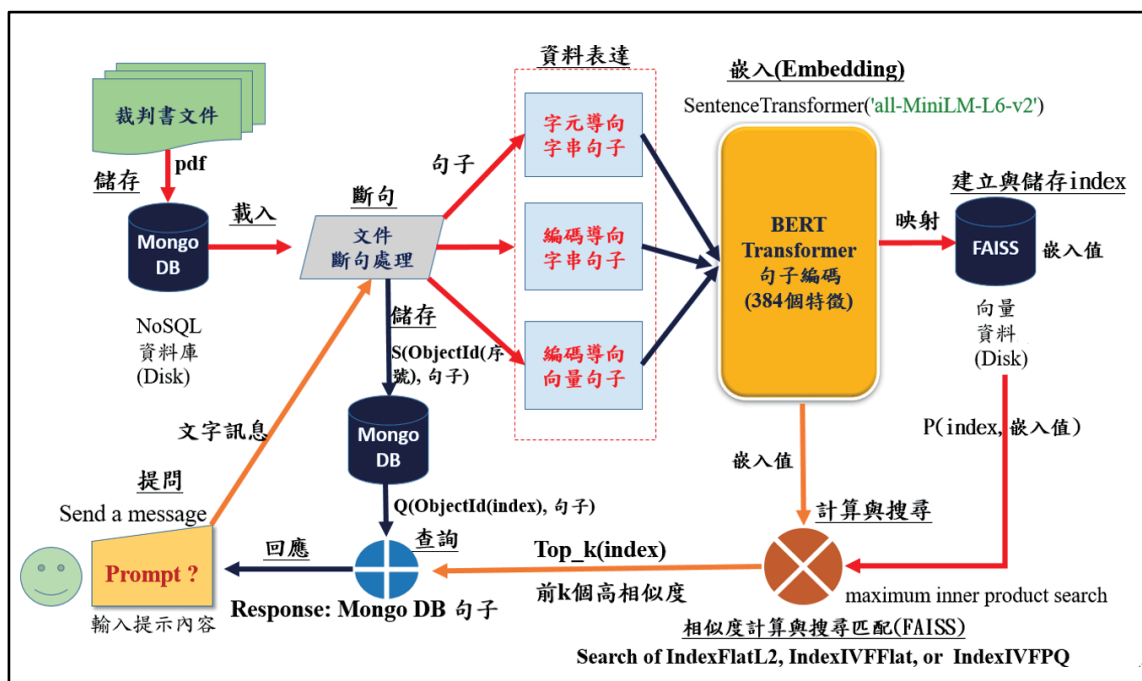


圖 3 簡易型 AI 機器人 Q&A 問答的處理過程

註：在 S (ObjectId (序號), 句子) 中，序號採用本研究所自訂的 12 個位元組與其 ASCII 值。

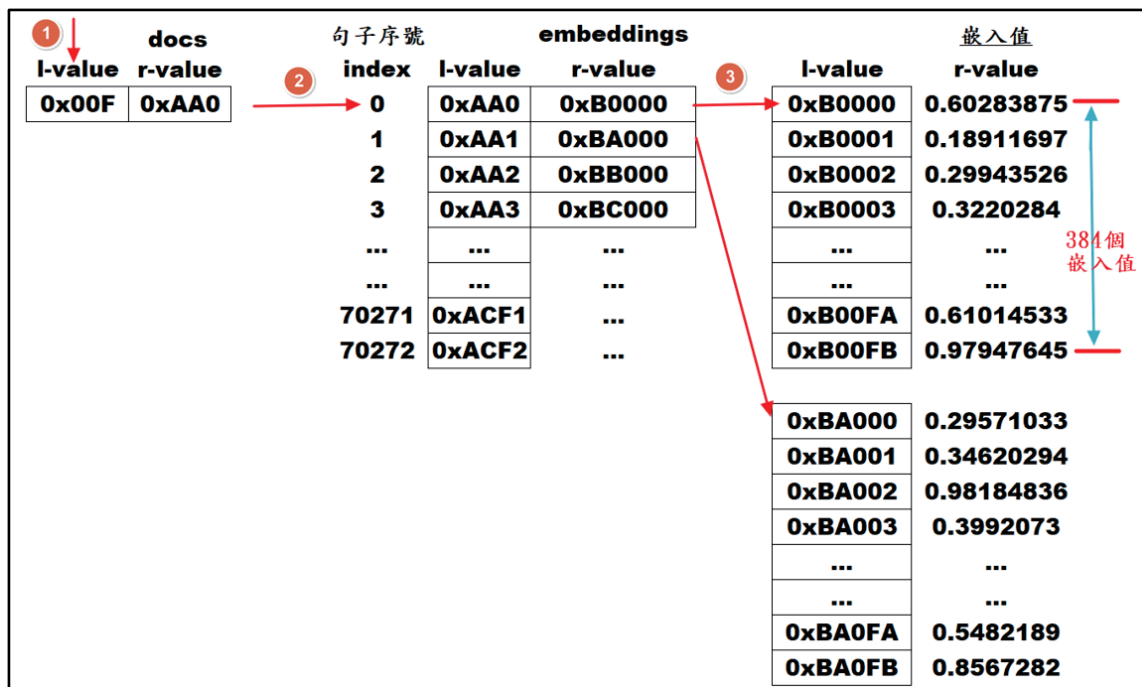


圖 4 總計 70,273 個句子在主記憶體中嵌入過程的示意圖

(三) 處理結果

經由 STATA 統計軟體的統計結果得知：

1. 在敘述統計方面，分別完成第一到第四種處理方式平均所需執行時間為 143.75 秒、446.46 秒、71.27 秒、15.04 秒。可見以向量資料表達+CPU+GPU 表現最佳，其次是向量資料表達+CPU，再來是字元導向字串資料表達+CPU，表現最差的是編碼導向字串資料表達+CPU。若以最佳時間作為參考標準，則第三種為第四種的 4.7 倍，第一種為第四種的 9.6 倍，第二種為第四種的 29.7 倍。如以（近似）概估它們比值的話，第四種：第三種：第一種：第二種 = 1:5:10:30。當開發人員電腦在沒有使用 GPU 設備情況下，仍比現行企業或學術界處理模式可減少一半的執行時間；如果具有 GPU 情況下，可縮短所需時間約 10 倍；有 GPU 比沒有 GPU 者縮短 5 倍；當採用編碼且字串時，最為不利，這可能與 BERT 演算法中的標記（Token）順序安排與存取邏輯等因素有關。
2. 在 ANOVA 變異數分析方面，從表 5 統計結果得知，不同資料表達，以及不同的機器設備，皆會對執行時間產生極為顯著的差異（機率值遠小於 0.05），且這樣的統計解釋效果具消滅誤差比例（Proportionate Reduction in Error, PRE）高達 0.996，近乎完美。因而，研究假設檢定二與研究假設檢定三，皆拒絕 H_0 。這些證據顯示，數字向量資料表達對於縮短類似 OpenAI 的 ChatGPT 或是 Google 的 Bard 等中文為主的聊天機器人所需執行時間是具有正面效果的。即使已經訓練好的大型語言模型（Large Language Model, LLM）結果已儲存在向量資料庫或其他類型資料庫成立下，使用者使用 ChatGPT 或 Bard 所輸入的提示（Prompt）內容仍需經過句子拆解與編碼程序所需執行時間。
3. 研究發現，執行 BERT SentenceTransformers 句子編碼所需時間約占掉全部執行時間的 99%以上，意味著 AI 機器學習是很花人力、時間、資源（含設備之算力、電力與記憶空間等）與成本。

表 5 四種處理方式對程式執行時間的變異數分析 (N=40)

Source	Partial SS	df	MS	F	Prob>F
Model	1108880.1	3	369626.7	3475.21	0.0000
資料表達	792050.47	2	396025.2	3723.4	0.0000
機器	15904.8	1	15904.8	149.54	0.0000
Residual	3829	36	106.3611		
Total	1112709.1	39	28531		
Number of obs=40, R-squared=0.9966, Root MSE=10.3132, Adj R-squared= 0.9963.					

伍、結論與建議

一、結論

本研究採用違反漁業法之司法裁判書作為實驗用材料，透過多種電腦語言與軟體，撰寫了相當多的程式碼，並得到下列之結論：

- (一) 從深度學習（演算法）觀之，機器非以詞性（POS）標註去剖析與了解人類文字語義的；意味著 AI 是採用「特徵學習」，而非以「詞性學習」。然而，從人類語文認知論之，機器學習產出結果會讓人類認為 AI 真的了解句子意涵或語義的。
- (二) 編碼導向（如文字編碼系統、Lookup table 中的 index、雜湊值或其他數學公式等）的資料表達，只要能產出文字的唯一性，自可適用於深度學習演算法。
- (三) 採用字元導向或文字編碼（數值）導向的資料表達，並不會影響到監督式的學習效果。
- (四) 在整體效能上，不論採用字元導向或文字編碼（或字元經使用者自訂數學函數轉換成數值）導向的資料表達，CNN Conv1D 表現十分出色。整體準確率將近 83%，整體執行時間不到 9.0 秒。BERT 雖然在準確率也表現很好（超過 81%），但執行時間約為 CNN 的 51 倍。
- (五) 在量刑分類效能（準確度）方面，CNN Conv1D，以預測罰金能力最佳，其次是有期徒刑，最後為拘役。此外，不同演算法也會影響到量刑分類效能。
- (六) 在 AI 的 Q&A 問答效能上，本研究採用不同於現今市場流行的字元導向資

料表達，事實證明文字編碼（或經數學函數所得數值）導向資料表達很適合 Q&A 問答應用上。

- （七）在向量資料搜尋處理上，FAISS 比 annlite 更易於處理搜尋與結果的回應。
- （八）在文件內容大部分或全部為中文情況下，應用在 AI 的 Q&A 問答效能表現上，本研究所提出國內首創的編碼（或數學函數）導向資料表達，經由統計檢定結果發現，在沒有使用 GPU 設備情況下，比現行企業或學術界處理模式，可減少一半執行時間；而有 GPU 比沒有 GPU 之執行時間可縮短將近 5 倍。這也是本研究的最大貢獻。這對於國內企業界推動 AI 實務應用中文落地化將有所助益的。

二、建議

從此次研究過程中發現，有相當多時間花在清洗資料工作上。基於此，為使後續相關主題研究更為順利起見，其建議如下：

- （一）為使我國在 AI 實務應用與中文化更能加速發展起見，由政府資料開放平臺（<https://data.gov.tw/>）所公開的政府開放資料，行政院應要求各級政府對提供的資料更為豐富、詳實與正確。比如多個欄位出現不少空白內容，或者一個檔案只存在一個連結網址。多年來政府一直喊出「數位轉型」口號，建議將開放資料改由資料庫處理、儲存與 CSV 格式匯出等。唯有如此，才能提升我國政府機關的資訊素養與資料品質。
- （二）目前商用聊天機器人，如 OpenAI 的 ChatGPT、Google 的 Gemini 或 Microsoft Copilot 等，仍然存在中國用語與文字。倘若中央政府不去重視此問題，那麼可預測在國人大量使用聊天機器人情況下，將會看到中國文字普遍在台灣的傳播媒體、政府刊物、教科書和學術文章中。例如，英文字的 Text，中國譯成「文本」，台灣應譯成「文字」，但「文本」一詞在國內學術界卻普遍被使用。
- （三）我國地方法院對於刑事案件裁判書仍存在待改善之處，包括：（1）裁判書寫作格式之嚴謹性和統一性仍不如公文寫作格式；（2）仍有少數的法院名稱誤植；（3）少數裁判書「主文」內容過於簡要；（4）「拘役」與「有期徒刑」難以透過深度學習加以明顯區隔，致使它們的預測效能不理想；（5）出現異常字或全形空白字，如「參」字；（6）證據能力和證據之證明力的事實陳述過於簡要。

(四) 期待借助本研究的編碼(或數學函數計算)的數值資料表達,以提升我國在 LLM+RAG 處理過程中的執行效能,縮短機器學習時間,節省電費。

致謝

本研究的完成要感謝國科會 112 年度深耕、創新、永續：建構離島社會發展新典範 III(社會共創篇 2/3)(計畫編號 NSTC112-2420-H-346-001)、國立澎湖科技大學高教深耕計畫(計畫編號 112G004-1)經費補助;另外,感謝葉品恩、蔡嘉欽、徐勝揚、林郁凱、溫鎮洲等位同學,協助問卷調查與資料收集。

參考文獻

一、中文

- 王嘉陵(2018),〈海洋教育：從深層生態學省思人與海洋之間的關係〉,教育脈動, 13, 1-8。
- 行政院農業委員會漁業署(2023),〈2018-2022 年中華民國 109 年臺閩地區漁業統計年報〉,行政院農業委員會。
- 呂學榮(2010),〈第三章：海洋漁業因應氣候變遷衝擊之調適策略〉,行政院農業委員會漁業署委辦計畫研究報告,計畫案號：L9910309。
- 李國添(2010),〈因應氣候變遷臺灣漁業產業之策略調適探討〉,行政院農業委員會漁業署委辦計畫研究報告。
- 李國添(2010),〈第一章：因應氣候變遷臺灣漁業產業之策略調適探討〉,行政院農業委員會漁業署委辦計畫研究報告,計畫案號：L9910309。
- 周祝瑛(2011),〈臺灣海洋教育之回顧與展望〉,海洋事務與政策評論,創刊號, 43-64。
- 孫榮聰(2022),〈海平面下的視角——從水肺潛水員角度談「限塑」對海洋生態保育之重要性〉,臺灣博物季刊, 41(3), 20-29。
- 陳光雄(2005),〈民眾參與海洋資源保育行動之策略〉,研考雙月刊, 29(4), 43-55。
- 湯二子(2022),〈海洋自然資源使用與生態環境保護之公益訴訟檢察〉,湖北經濟學院學報(人文社會科學), 19(2), 97-100。
- 黃明和(2012),〈臺灣周邊海域實施漁業生態系方法的管理思維：與談〉,海洋事

務與政策評論，1（1），37-60。

黃雅琪（2023），〈立法院三讀通過《海污法》修法：海洋污染罰鍰最高 1 億元〉，
鏡週刊。網址：<https://reurl.cc/L67q2K>。

楊憶婷、葉庭光（2022），〈博物館推動海洋教育新方向——家長海洋科學態度與
子女學習海洋科學的相關性〉，科技博物，26（2），5-29。

葉俊榮（2005），〈建構海洋臺灣發展藍圖〉，研考雙月刊，29（4），8-21。

劉家良（2010），〈氣候變遷對臺灣沿岸漁業之經濟影響〉，臺北大學自然資源與環
境管理研究所碩士論文，1-70。

二、西文

Bernal, J., Kushibar, K., Asfaw, D. S., Valverde S., Oliver A., Martí R., and Lladó, X.
(2019). “Deep Convolutional Neural Networks for Brain Image Analysis on
Magnetic Resonance Imaging: A Review,” *Artificial Intelligence in Medicine*,
Vol 95, 64-81.

Costanza, R., Andrade, F., Antunes, P., Belt, Marjan van den, and et al.(1998). Principles
for Sustainable Governance of the Oceans. *Science*, 281, 198-199.

Devlin, J., Chang, Ming-Wei, Lee, K., and Toutanova, K. (2019). “BERT: Pre-training
of Deep Bidirectional Transformers for Language Understanding,” *Proceedings
of the 2019 Conference of the North American Chapter of the Association for
Computational Linguistics: Human Language Technologies*. Association for
Computational Linguistics, 4171-4186.

Geeksforgeeks (2023). “Bidirectional LSTM in NLP,” [https://www.geeksforgeeks.org/
bidirectional-lstm-in-nlp/](https://www.geeksforgeeks.org/bidirectional-lstm-in-nlp/).

Lecun, Y., Bengio, Y., Hinton, G. (2015). “Deep learning,” *Nature*, Volume 521, Issue
7553, 436-444.

Matsugu. M., Mori, K., Mitari, Y., and Kaneda, Y. (2003). “Subject Independent Facial
Expression Recognition with Robust Face Detection Using a Convolutional
Neural Network,” *Neural Networks*, 16, 555-559.

Roiger, Richard J., Geatz, Michael W. (2003). *Data Mining: A Tutorial-Based Primer*.
Pearson Education, Inc., 4-5.

Siarni-Namini, S., Tavakoli, N., and Namin, A. S. (2019). “The Performance of LSTM

and BiLSTM in Forecasting Time Series,” 2019 IEEE International Conference on Big Data (Big Data), Los Angeles, CA, USA, 3285-3292.

Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., and Kaiser, Ł. (2017). “Attention Is All You Need,” 31st Conference on Neural Information Processing Systems (NIPS 2017), Long Beach, CA, USA., 5998-6008. <https://arxiv.org/pdf/1706.03762.pdf>.

Zhang, S., Zheng, D., Hu, X., and Yang M. (2015). “Bidirectional Long Short-Term Memory Networks for Relation Classification,” Proceedings of the 29th Pacific Asia Conference on Language, Information and Computation, 73-78.