



使用ChatGPT提升研究效率

謝秉均

國立陽明交通大學資訊工程學系副教授

壹、前言

談到大家熟悉的GPT-3和GPT-3.5，它們擁有約1,750億個參數。為了讓大家更直觀地理解這個數字的龐大，可以做一個不完全恰當的類比：將一個參數比作大腦中的神經元。一般來說，人類大腦約有860億個神經元，而大象則擁有約2,510億個神經元。因此，這些大型語言模型的參數量級與哺乳類動物大腦的神經元數量相當。

簡單來說，語言模型（Language Model）是一種可以預測下一個單詞的統計方法。它通過分析大量的文字數據，學習單詞之間的概率分布，從而能夠在給定一個句子的情況下，預測下一個單詞出現的可能性，專有名詞稱為「next-word prediction」。舉例來說，如果給語言模型輸入「Hello」，它會根據對語料庫的分析，預測下一個單詞出現的概率。例如，「Hello world」的概率可能為0.2，「Hello Kitty」的概率可能為0.1，「Hello, how are you?」的概率可能為0.7。最終輸出的單詞將從這些可能的單詞中根據概率進行抽樣。

其實，語言模型（Language Model）的概念並不新穎，早在ChatGPT出現之前，就已經廣泛應用於各種語言處理任務中。以下舉兩個例子：

一、輸入法推薦：在使用輸入法打字時，經常會看到下一個字的推薦選擇。例如，輸入「I'll meet you at the...」時，輸入法可能會建議「coffee」、「airport」或「office」等後續詞語。這是因為語言模型能夠根據輸入的字詞，預測下一個字出現的概率，並



以機率的方式提供推薦。

二、搜尋引擎提示：在使用Google搜尋時，當輸入「what is the」時，搜尋引擎會自動提示一些可能的後續詞語，例如「weather」、「meaning」或「capital」。這是因為語言模型能夠根據輸入的關鍵字，預測使用者可能正在搜尋的內容，並提供相關提示。

總而言之，ChatGPT所運用的語言模型技術並非全新，而是經過多年發展和驗證的成熟技術。其核心原理就是利用機率模型來預測下一個字出現的可能性，並以此為基礎提供各種語言處理功能。

貳、ChatGPT與一般的輸入法或搜尋引擎的差異

ChatGPT與傳統的輸入法或搜尋引擎相比，ChatGPT的序列到序列和自回歸特性使其能夠生成更具創造性和上下文相關的文字，以下說明：

- 一、序列到序列（Sequence-to-Sequence）：意味著它可以將輸入序列（例如句子）轉換為輸出序列（例如另一個句子）。例如，輸入「背誦第一定律」，ChatGPT嘗試生成機器人三大定律的第一定律——「A robot may not injure Humans」。
- 二、自回歸（Autoregressive）：意味著它逐字生成輸出，並根據先前生成的字詞預測下一個字詞。例如，輸入句子是「I am a」，ChatGPT可能會首先預測下一個字詞是「robot」，然後生成「I am a robot」。

為了訓練語言模型，我們會使用大量的人類語言數據，並將其視為一種「標準答案」。例如，可以從書籍、文章或網路中提取句子，然後將其中部分單詞刪除，讓語言模型嘗試猜測缺失的單詞。這種訓練方法稱為「self-supervised learning（自監督學習）」，因為它不需要人工標註的標籤。

然而，一些研究人員提出了一個問題：語言模型是否只是在背誦它們所看到的文字，而不是真正理解語言的含義？這種現象在2021年〈On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?¹〉文中稱為「Stochastic Parrots（隨機鸚鵡）」。意即語言模型可能會像鸚鵡一樣模仿人類的語言，但它可能無法理解句子或對話的真正含義。

1 Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). *On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?*. In FAccT '21: Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency, 610-23. Nueva York: Association for Computing Machinery. doi:10.1145/3442188.344592



雖然目前還沒有定論，但「隨機鸚鵡」現象是一個值得關注的問題。如果語言模型只是在背誦文字，那麼它的實際應用可能受到限制。因此，研究人員正在探索新的方法來訓練語言模型，使其能夠真正理解語言的含義。

即使大型語言模型像ChatGPT已經通過自監督學習掌握了大量的文本和語言規則，但這是否足以使其表現得像人類？答案是否定的。

構建大型語言模型的最終目標是使其能像人類一樣理解和使用語言。而要做到這一點，語言模型必須能夠理解人類普遍的概念，特別是抽象概念。例如，語言模型應該能夠理解「有趣」、「安全」、「安全感」甚至「愛」等概念。這些概念都是抽象的，無法用簡單的句子直接描述。

僅僅通過自監督學習，語言模型很難理解這些抽象概念。因為自監督學習的訓練數據中並沒有明確地定義這些概念。

參、ChatGPT如何學習抽象概念

對於抽象概念的學習，單一目標或目標函數的訓練方法往往難以奏效。因此，需要探索更有效的訓練策略。以下提出一個基於情境對比的訓練方法，以「愛」為例進行說明。

假設將語言模型視為一個孩子，並向其提出問題「愛是什麼？」。孩子可能無法給出明確的答案，但可以通過提供具體的情境來引導其理解。例如展示兩個情境，詢問哪個情境更符合「愛」的描述？

情境一：父母陪伴孩子進行親子共讀。

情境二：孩子獨自使用手機，用3C來育兒。

相信大家的直覺反應可能是情境一，因為親子共讀體現了父母對孩子的關愛和陪伴。

不斷提供類似的對比情境，也是現行訓練大型語言模型的做法，不斷地給語言模型終極二選一，然後告訴它，哪一個比較貼近我們要的概念，幫助語言模型理解抽象概念。這種方法可以有效地引導語言模型理解抽象概念的含義。如果對語言模型的訓練數據集感興趣，可參考Hugging Face數據集²。

2 Hugging Face資料集網站：https://huggingface.co/datasets/openai/webgpt_comparisons/viewer/default/train?p=4



肆、基於人類回饋的強化學習

介紹了基於人類回饋的強化學習（Reinforcement Learning From Human Feedback, 簡稱RLHF）所使用的終極二選一數據集。接下來，將探討如何使用RLHF方法來訓練語言模型。

強化學習（Reinforcement Learning, 簡稱RL）是一種機器學習方法，讓其透過嘗試錯誤來學習最佳策略。它會在環境中接收狀態和獎勵訊號，並根據策略採取動作。目標是找到使獎勵最大化的策略。以下以打磚塊遊戲為例說明RL如何訓練AI：

- 一、環境：打磚塊遊戲環境，包括球、球拍、磚塊等。
- 二、狀態：球、球拍的位置、磚塊的狀態等。
- 三、動作：球拍不動、向左、向右移動等。
- 四、獎勵：打掉磚塊獲得1分，未打掉磚塊獲得0分。
- 五、策略：根據球和球拍的位置，決定球拍的動作。

透過RL，AI可以學習最佳的策略，打掉更多的磚塊，獲得更高的分數。

伍、用RLHF訓練ChatGPT

一、如何基於RL來訓練ChatGPT

用RL訓練ChatGPT是一項複雜的挑戰，需要精心設計獎勵函數和訓練環境，透過遊戲分數作為獎勵來學習策略。

- (一) 定義狀態：使用者輸入的對話或問題（prompt）。
- (二) 定義動作：AI生成的回覆。
- (三) 定義策略：在不同狀態下，AI應如何生成回覆。
- (四) 定義獎勵：評估AI回覆的質量。
- (五) 定義訓練環境：提供大量對話或問題和回覆數據。

RL其實只是在建構獎勵與訓練環境這兩個要素，獎勵函數的設計是RL中的一大難題。如在CoastRunners遊戲中，如果只以遊戲分數作為獎勵，AI會發現不斷吃綠色物品可以獲得高分，而忽略了完成賽道的目標。由此可發現設計不當，AI可能會學習到投機取巧的策略，而非我們期望的目標。用RL來協助，或協助讓ChatGPT學習抽象概念過程中，設計獎勵函數並不容易。

二、如何透過終極二選一來訓練ChatGPT

RLHF是一種強大的訓練大型語言模型的方法，它已被證明可以有效地提高ChatGPT的能力。

- (一) 學習抽象概念：在文本摘要任務中，可以提供兩個摘要，並要求選擇其中一個。人類評估者會對摘要進行評分，並根據評分結果對模型進行獎勵或懲罰。透過這種方式，模型可以學習到哪些摘要是更好的，從而提高其文本摘要的能力。
- (二) 符合人類價值觀：在回答問題任務中，可以提供兩個答案，並要求選擇其中一個。人類評估者會對答案進行評估，並根據評估結果對模型進行獎勵或懲罰。透過這種方式，模型可以學習到哪些答案是更好的，從而提高其回答問題的能力，並使其答案更加符合人類的價值觀。

儘管RLHF在訓練ChatGPT方面取得了很大成功，但它也存在一些以下的侷限性。

- (一) 需要大量的人類評估者：為了訓練一個成功的RLHF模型，需要大量的人工標註數據。這意味著需要僱用大量的人類評估者，並對他們進行評估標準的培訓，要訓練一個很成功的模型，需要非常大量的人類評估者，這是在實際應用中需要考慮的因素。
- (二) 存在偏見風險：人工標註數據可能會存在偏見，這可能會導致模型也存在偏見。
- (三) 倫理問題：使用人工評估者可能會涉及一些倫理問題，例如低薪和工作條件惡劣。

陸、如何使ChatGPT來加速做研究

在對ChatGPT有了基本了解之後，探討如何使用它來加速研究之前。首先，要承認使用ChatGPT做研究一直存在一些質疑：ChatGPT容易產生幻覺（Hallucination），生成與事實不符的內容；ChatGPT推理能力可能有限，無法根據證據推論出合理的結論；還有一個是做研究的質疑，如「做研究是人類的強項，是人類心智的高度創意與冒險，是沒辦法被自動化的」。

一、ChatGPT幻覺問題

ChatGPT的幻覺問題一直是使用者關注的焦點。所謂幻覺，是指ChatGPT生成的內容與事實不符。ChatGPT產生幻覺的原因，可以從其訓練方式來理解。ChatGPT的訓練目標是預測下一個字是什麼。因此，在生成內容時，ChatGPT會根據前面已經生成的文字，預測接下來最有可能出現的字。



這種基於next-word prediction的訓練方式，容易導致ChatGPT產生幻覺。例如，如果ChatGPT要介紹一位臺灣資訊工程學者，它可能會預測接下來的文字應該是「任職於臺灣大學電機資訊學院資工系」，然而，這可能與事實不符。出現幻覺是大家預期的，為了解決幻覺問題，ChatGPT的開發人員採取了一系列措施，ChatGPT產生的幻覺已經有所減少。

如何解決大型語言模型產生的幻覺問題，一直是自然語言處理領域的熱門研究課題。目前最主流的做法之一是擷取增強生成RAG（Reterieval-Augmented Generation）。擷取增強生成是一種很有潛力的大型語言模型（Large Language Model，簡稱LLM）技術，可以有效解決幻覺問題，並提高大型語言模型的整體性能。然而，擷取增強生成也存在一些缺點，需要在實際應用中加以考慮。

擷取增強生成的基本原理是，讓大型語言模型在生成回應之前，先從外部知識庫中檢索相關資訊。這就好比讓學生在考試時，可以查閱資料輔助作答，從而提高答題的準確性。

擷取增強生成的檢索來源可以非常廣泛，包括網路資料、影片、電子郵件等。以下是一些知名的擷取增強生成應用案例：

- (一) Wolfram Alpha³整合：讓大型語言模型可以透過Wolfram Alpha平臺進行數學運算，解決其在數學方面能力不足的問題。
- (二) Perplexity⁴：一個結合搜尋引擎與大型語言模型的新創公司，其產品強調所有回應都必須有引用文獻佐證，以確保資訊的可靠性。

二、ChatGPT的推理能力

ChatGPT作為大型語言模型，在生成文字、翻譯語言等方面表現出色，但其推理能力一直備受關注。

(一) ChatGPT推理能力現況

研究人員針對ChatGPT的推理能力進行了評估，發現其在回答需要複雜推理的問題時表現不佳。例如，在被問及如何將沙發搬到屋頂上時，ChatGPT的回答往往不切實際，例如建議將沙發切割成小塊。

3 Wolfram Alpha網站：<https://www.wolframalpha.com/>

4 Perplexity網站：<https://www.perplexity.ai/>

原因在於ChatGPT的推理能力主要基於對常識的理解，而其資料集可能存在不足或偏差。此外，ChatGPT在處理複雜問題時容易陷入局部理解，無法進行全局思考。

(二) 提升ChatGPT推理能力的方法

針對ChatGPT推理能力不足的問題，研究人員提出了以下兩種提升方法：

1. 類比學習（In-Context Learning）：透過提供大量示範案例，讓ChatGPT學習如何從類似的問題中推導出答案。例如，在情感分析任務中，可以提供大量正向和負向文本範例，幫助ChatGPT理解情感的表達方式。
2. 步驟拆解提示（Chain-of-Thought Prompting）：將複雜的推理過程拆解為多個可操作的步驟，並明確提示ChatGPT每個步驟的含義。例如，在數學運算任務中，可以提供詳細的解題步驟，幫助ChatGPT理解解題思路。

這兩種方法的優點是無需重新訓練模型，即可在一定程度上提升其推理能力。

三、ChatGPT如何排除質疑並協助研究

許多人認為研究是人類心智的獨特能力，無法被自動化。然而，ChatGPT等大型語言模型的出現，為研究人員提供了一種新的工具，有可能部分自動化研究流程。研究的流程通常包含以下幾個關鍵步驟：

- (一) 發想研究題目：研究人員需要根據自身的興趣和專長，以及現有的研究成果，提出新的研究問題。
- (二) 做文獻回顧：研究人員需要查閱大量文獻，瞭解相關領域的研究進展，為自己的研究奠定基礎。
- (三) 想出解決方案：研究人員需要針對研究問題，提出創新的解決方案。
- (四) 做實驗：研究人員需要透過實驗或其他方法，收集數據驗證自己的解決方案。
- (五) 做數據分析：研究人員需要對收集到的數據進行分析，得出研究結論。
- (六) 寫論文投稿：研究人員需要將研究成果撰寫成論文，並投稿至學術期刊發表。
- (七) 回覆審稿人意見：研究人員需要在收到審稿人意見後，對論文進行修改完善。

隨著人工智能技術的發展，大型語言模型（LLM）逐漸融入各行各業，研究領域也不例外。LLM具有強大的信息處理和理解能力，可以為研究人員提供多方面的輔助。

在使用ChatGPT或其他LLM輔助研究時，應遵循以下原則：

- (一) 類比學習（In-Context Learning）：為LLM提供大量的示例，使其學習如何從類似的問



題中推導出答案。

- (二) 步驟拆解提示 (Chain-of-Thought Prompting)：將複雜的任務或問題拆解為多個可操作的步驟，並明確提示LLM每個步驟的含義。
- (三) 具體化問題：在問問題時，應盡可能具體、明確，以便LLM提供更準確的回答。
- (四) 嘗試不同的說法：由於LLM的訓練過程可能使其存在偏差，因此應嘗試不同的說法、不同的觀點來問問題。
- (五) 逐步解釋並提供引用：為LLM提供逐步解釋和引用文獻，可以減少其產生幻覺的可能性。

柒、ChatGPT協助研究的實際應用

一、ChatGPT協助發想研究題目

ChatGPT可以通過以下方式協助研究人員發想研究題目：

- (一) 提供豐富的知識庫：ChatGPT的訓練資料包含大量的網際網路信息，可以為研究人員提供廣泛的知識視野。
- (二) 激發新的思維模式：ChatGPT可以從不同的角度思考問題，幫助研究人員跳脫框架，激發新的創意。
- (三) 快速篩選潛在的研究方向：ChatGPT可以快速處理大量資訊，幫助研究人員快速篩選出潛在的研究方向。

以下以圖書館服務為例，使用ChatGPT發想研究題目：

- (一) 背景介紹：首先，向ChatGPT介紹自身的研究背景和興趣領域。例如，可以說：「我是一名研究型大學圖書館員，感興趣的議題是如何使用近五年的最新資通訊技術和AI技術來協助大學的研究生做研究。」
- (二) 提出具體問題：然後，提出具體的問題。例如，可以問：「請提供五個可能結合最新資通訊技術和AI技術的圖書館服務，並針對每一個服務都提供詳細的說明。」
- (三) 分析ChatGPT的回覆：ChatGPT可能會提供以下服務的建議：AI驅動的研究輔助平臺、智能推薦系統、自動化文獻分析與摘要生成、虛擬研究協作平臺、數字人文數據視覺化支持。
- (四) 追問並細化：針對感興趣的服務，可以進一步追問並細化。例如，可以問：「對於智

能推薦系統，我們想要做的一些事能不能具體提供可能的幾個步驟，來幫助圖書館可以在一年內就推行這個服務？」

- (五) 嘗試不同的說法：如果希望獲得更具創新性的研究題目，可以嘗試使用不同的說法。例如，可以問：「請提供五個最新的資通訊技術和AI技術的嶄新圖書館服務，而且這些服務從來沒有在任何現行圖書館推行過。」

ChatGPT等LLM可以為研究人員提供多方面的輔助，尤其是發想研究題目方面。通過遵循一定的原則和方法，可以有效利用LLM的優勢，提高研究效率和研究質量。

二、ChatGPT協助文獻回顧

在使用ChatGPT協助文獻回顧時，應提供明確的任務和研究問題：明確告知需要完成的任務是什麼，以及研究問題是什麼。

- (一) 制定具體的篩選標準：如果需要對文獻進行篩選，應制定具體的篩選標準。
- (二) 逐步細化要求：如果需要其回覆進行進一步的修改，應逐步細化要求，使其能夠理解目的。

以下以使用Web of Science資料庫進行文獻回顧為例，使用ChatGPT協助：

首先，向ChatGPT介紹研究背景和研究問題。例如，可以說：「我是一名研究型大學圖書館員，感興趣的議題是如何使用強化學習來進行風險控管的投資組合管理。請您根據此研究問題，開發一個演算法，在Web of Science上搜尋相關文獻。」

然後，根據研究問題的關鍵概念，制定步驟拆解提示，列出幾個可操作的步驟，讓ChatGPT容易達成目標。例如，將研究問題切成幾個主要的關鍵概念，再將每一個關鍵概念讓它去確認出幾個主要的關鍵字，可以透過組合這些關鍵字加上「AND」或「OR」運算子，組合這些關鍵概念並形成各種可能的條件。

透過上述步驟，ChatGPT可以知道用什麼關鍵字，及如何找到關鍵字去搜尋。接下來再寫出具體的研究問題，比如要請它找用強化學習來做風險控管的投資組合，這是個很具體的研究問題，善用步驟拆解提示（Chain-of-Thought Prompting），切成很多個可操作的步驟，之後可以發現ChatGPT回傳的內容較具體。

第二步篩選文獻，假設前一個例子，研究者可能可接受ChatGPT回傳10篇文章，但若回傳的文章是1000篇，是否還可以接受？

如果是大量的文獻的時候，就需要請它幫忙列出具體的任務以及研究問題，更重要



的是要寫出具體的篩選標準。舉例來說，只挑選有理論結果且經過同儕審閱的文章。也就是列出篩選標準，請ChatGPT限縮目標。

接下來就可以篩選刊名、摘要等詳細的資訊，再請ChatGPT幫忙決定符合篩選標準的文章。這是在做初篩選時，ChatGPT可以幫忙的部分，經過篩選後，就可以從剩餘符合篩選標準的文章開始著手進行研究。

目前在學術界，LLM直接生成論文內容仍存在諸多爭議。許多研討會和期刊都明文禁止使用LLM直接生成論文內容，但允許使用LLM輔助論文撰寫，例如，ICML（International Conference on Machine Learning）允許使用LLM進行文獻回顧和論文修改潤飾。

三、ChatGPT協助發想論文標題

如果要寫論文前言的各個段落的話，可以很明確地講，每一個段落要包含什麼樣的觀點和內容。首先條列最重要的觀點，再給明確的指示應該如何構築語句或者段落。假設要寫一段關於強化學習技術用於訓練語言模型的介紹文，明確提供希望文章要包含以下幾個觀點，再請ChatGPT自行編排順序。基於人類回饋的RLHF這個方法背後的優勢是，給它明確要包含哪些觀點，當然觀點還是由研究者制定，然後指示如何編排順序，產生的第一版就能是蠻不錯且可讀性還不錯的文章。

四、ChatGPT探索可能解決的方案

如果遇到研究問題，試著用ChatGPT探索一些可能解決的方案。舉例來說，在數學領域有一個很有名的猜想「孿生質數猜想（Twin prime conjecture）」，猜想內容大致是問是不是存在無窮多組相差為二的質數，這個問題應該是屬於開放性問題，我們可以用下面的步驟詢問ChatGPT。

首先提供研究背景，想要證明「Twin prime conjecture」，然後請它提供五個可能的解決方案，再給它具體的指示要求每個解決方案應該要提供一篇相關的文章，並且要提供引用文獻，然後更重要的是所有的引用文獻都要經過事實查核。它回覆的內容，其中之一就提到一篇2014年很有突破性的論文〈Bounded gaps between primes⁵〉，這是由數學家張益唐提出跟孿生質數相關的文章，ChatGPT也提供了相對應的引用文獻。另外，它也提供

5 Zhang, Y. (2014). Bounded gaps between primes. *Annals of Mathematics*, 179(3) 1121-1174. doi: 10.4007/annals.2014.179.3.7



從不同切入角度的四篇文章，這些文章確實都有真實存在的一些解法。然後，各自都給了一篇引用文獻，也確實檢查過文章，所以提供ChatGPT明確的指示，讓它提供引用文獻，而且要經過事實查核，以此例來說回覆頗為精準。

五、ChatGPT協助回覆審稿者

ChatGPT可在投稿時回覆審稿者（reviewer）的意見，特別是有字數限制的需求。ChatGPT很擅長改寫句子，在努力地寫下大部分完整的回覆意見，排列優先順序，請它依字數限制改寫，就不需要人為的縮寫。

回覆審稿者的內容字數是有限制的。最困難的是審稿者可能問了20個問題，但只能用一頁A4版面就要回答這20個問題。通常先寫完一版兩三頁回覆，還需要花很多時間、字字珠璣把複雜的內容改寫成很短的句子，善用各種創意縮成一頁，這樣才夠完整回覆。

使用ChatGPT，可以先不管篇幅，先直接、詳細地寫下每一個問題的回覆，排列問題重要性的優先順序。然後，直接請ChatGPT把「文字縮到多少個字元，同時保留大部分的意思。」基本上它都可以完成指示，縮減成一頁或要求的字元數。

捌、結論

綜上所述，可發現從發想研究題目、文獻回顧、撰寫文章、到回覆審稿者，基本上大部分的研究階段都可以透過ChatGPT語言模型來輔助。

其中，最關鍵的概念是：善用類比（In-Context Learning）、做步驟拆解提示（Chain-of-Thought Prompting）把複雜的東西化成簡單的東西、給定明確的情境（Context），明確地告訴如何回覆及其他研究上具體需求，然後透過追問產生更精準的回覆。

本文為「知識經濟時代之圖書館服務系列專題演講（三十一）：生成式AI應用與挑戰」（113.07.16）之演講紀錄，由陳進發及林彥汝協助整理，並經主講者過目授權同意刊登。