

「數據的重量」刪除垃圾郵件也能為地球降溫？

看不見，不代表無重量

在2026年，人們逐漸習慣了雲端的輕盈。但是，你知道嗎？其實，數據是有「物理重量」的。

全球數位科技的碳排放量目前已達總量的3%~4%，正式超越全球航空業。每一位元（Bit）的存儲，背後都連動著真實世界的發電機組與冷卻系統。

數據中心的「大胃口」與AI代價

數據存儲在24小時運作的數據中心，2026年全球數據中心年用電量預計突破500TWh（占全球總量2%）。

- **AI耗能**：一次生成式AI（如ChatGPT或Gemini）查詢的能耗是傳統搜尋的10倍以上。
- **AI水足跡**：與AI對話20~50次，其散熱過程會蒸發約500毫升淡（等於一支礦泉水）。

一封郵件的連鎖反應

一封普通電郵從發送至儲存，平均產生約4克二氧化碳；若帶大型附件群發，最高可達50克。

- **驚人對比**：一名上班族一年收發郵件的排碳量約131千克，相當於開汽油車行駛500公里。
- **退訂的力量**：當你收到一封沒用的廣告郵件並將其留在收件箱時，它並非安靜地躺在那裡。雲端服務供應商會將這封郵件同步備份到全球多個伺服器上。每一天，這些伺服器都要耗電來維持這封郵件的「生命」。因此，「退訂（Unsubscribe）」比單純的「刪除」更有威力，因為它從源頭切斷了數據的二次產生與備份。

實踐「數位斷捨離」的四個方案

- 1 **清理「數位贅肉」**：在雲端儲存1GB的數據，每年大約會產生0.04千克的二氧化碳。建議定期搜尋信箱中的關鍵字「Unsubscribe」或「退訂」，取消那些你從來不看的電子報。同時清理雲端空間中重複、模糊的相片，這不僅能幫你省下月費，更能直接減少伺服器24小時運作與多重備份的壓力。
- 2 **優化影音習慣**：數據傳輸同樣耗能，從雲端下載或串流1GB的數據，耗電量約0.066kWh。在手機的小螢幕上

看影片時，適度將解析度從4K調低至1080p，肉眼分別不大但能大幅減少能耗。此外，將常用的音樂或影片下載到本地端播放，比每次重複透過網絡串流更環保。

- 1 **擁抱深色模式（Dark Mode）**：如果你的手機或電腦使用OLED螢幕，請切換到深色模式。這不僅能減輕眼睛疲勞，更因為OLED在顯示黑色時像素點是不發光的，在正常使用下，這一個小動作能幫你的設備節省高達30%的電量，變相延長電池壽命，減少電子廢物。
- 2 **智慧使用AI**：精準提問減少算力利用AI提升效率是2026年的必修課，但「精準提問」是關鍵。試著一次性給出清晰的背景與指令，減少無謂的往返對話。

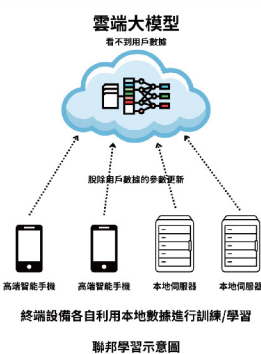
科技應以綠色為底色

一個擁有1TB雲端備份的家庭，每年維持數據存在的耗電量可能佔家庭總用電量的5%。下一次按下「儲存」前，請停一秒思考：這份數據真的有必要存在嗎？讓我們從刪除十張重覆照片開始，為地球「減重」。

2026年AI關鍵賽道——AI在私隱保障與輕量化部署上的趨勢

1.聯邦學習（Federated Learning）與私隱計算：

隨著各國對AI監管越來越嚴格，傳統將數據集中訓練的模式面臨挑戰。聯邦學習讓AI模型能在不接觸原始數據的前提下，於機構內部完成訓練，僅上傳加密後的參數。這種技術，正成為金融、醫療等敏感行業傾向使用的模式。模型在設備本地進行微調，僅回傳加密後的參數更新。



2.小型語言模型（Small Language Model）：

以往的大語言模型，如果要性能良好、準確度高，又有廣泛領域的知識，需要千億參數（100B）或以上。而SLM則是專注於特定領域，參數量可縮減至數十億（如：3B、7B、8B等）甚至更低。這些小模型，有條件可在高端手機、邊緣設備上離線運行，既降低運算成本，也大幅提升數據私隱性。一些參數量在數十億範圍的小模型，已證明在特定任務上表現優異，其核心技術為：模型量化（Quantization）及知識蒸餾（Knowledge Distillation）等。

3.端側AI（On-Device AI）普及化：

新一代旗艦手機晶片內置專用NPU（神經網絡處理器），使AI算力在終端設備上已有明顯的優勢，部份高端智能手機無需聯網即可實現即時語音翻譯、圖像生成或會議摘要等工作。

核心改變：

解決了AI應用的一大痛點——延遲（Latency）與流量（Token）消耗。

施少傑 工程師

區志煒 高級工程師

聯絡資料

資訊系統推廣室

地址：澳門新口岸上海街中華商會大廈六樓
電話：(853) 8898 0829
傳真：(853) 2878 8233

澳門生產力暨科技轉移中心總辦事處

地址：澳門新口岸上海街中華商會大廈七樓
電話：(853) 2878 1313
傳真：(853) 2878 8233
郵箱：cpttm@cpttm.org.mo
網址：www.cpttm.org.mo

數碼匯點

地址：澳門馬統領街廠商會大廈三樓
電話：(853) 8898 0601
傳真：(853) 2873 3085

時尚匯點

地址：澳門漁翁街海洋工業中心第二期十樓
電話：(853) 8898 0701
傳真：(853) 2831 2079



什麼是瀏覽器指紋

當我們登入某些網站的時候，網站會生成了一個用戶ID，但即使我們沒有登入，甚至打開了瀏覽器的「無痕模式」，網站依然能生成一個用戶ID，而且這個ID與我們登入時產生的完全相同——這就是瀏覽器指紋，一種能夠識別用戶身份的技術。

在我們瀏覽網站時，某些網站會悄悄收集與瀏覽器相關的各種細節，例如IP地址、作業系統版本、用戶時區、瀏覽器安裝的插件等。當這些數據集合在一起時，便形成了獨一無二的瀏覽器指紋。即使我們清除了瀏覽器Cookie或使用了隱私模式，網站仍然能夠識別出每個獨立的訪客。



瀏覽器指紋的主要用途

- **廣告追蹤與精準行銷**：這是最常見的用途。廣告商可以不依賴第三方Cookie追蹤你在不同網站的行為，進而向你投放個人化廣告。
- **分析與統計**：網站可以了解訪客的設備和瀏覽器分佈情況，以優化網站設計。
- **防範詐騙與增強安全**：銀行或支付平台可以利用指紋來判斷登入的裝置是否為你常用裝置。如果指紋突然改變，可能就代表帳號被盜，系統便會觸發額外的身份驗證。
- **網站防刷量**：社交平台都希望確保用戶的真實性，通常不樂見單一使用者操控大量帳號。透過瀏覽器指紋技術，網站便能識別出這些帳號是否由同一人操作，進而啟動相應的安全政策。

如何保護自己不被瀏覽器指紋追蹤？

雖然無法完全杜絕瀏覽器指紋的收集，但可以採取以下措施來降低被追蹤的風險：

使用反指紋追蹤的瀏覽器擴充功能：

例如 Privacy Badger、uBlock Origin（需在設定中啟用反指紋功能）或 Canvas Blocker，這些工具能阻擋或混淆常見的指紋讀取腳本。

使用指紋瀏覽器：

指紋瀏覽器會讓所有使用者的瀏覽器特徵幾乎一模一樣，使你的指紋「淹沒在人群之中」，讓網站難以區分個別用戶。

定期更新瀏覽器：

新版瀏覽器通常會加入更多隱私保護機制，減少指紋可以被讀取的資訊量。



Q | ispu@cpttm.org.mo

澳門生產力暨科技轉移中心希望透過出版「IT快訊」，收集並發放各地有關電子政府、電子商務、實用IT知識、軟件開發技巧及資訊科技教育等消息，讓社會各界人士作為參考（但恕本中心不會承擔任何責任），希望您覺得它有用，如對本刊物有任何意見或建議，歡迎發電郵到 ispu@cpttm.org.mo 與我們聯絡。



在這個人工智能（AI）大爆發的時代，AI模型已經成為我們工作與生活中的得力助手。然而，隨著使用頻率增加，許多企業主、開發者乃至個人使用者開始面臨兩個核心痛點：昂貴的訂閱費與Token費用，以及敏感資料的隱私風險。

目前的雲端AI服務大多採用「按量計費」或「月費訂閱模式」。對於需要處理大量文本或進行自動化流程的用戶來說，長期累積的Token費用相當驚人。更重要的是，在處理涉及個人隱私、法律文件或企業核心技術的資料時，雲端傳輸始終存在安全隱患。那麼「本地部署（Local Deployment）」就是一個不錯的解決方案。今天，要向大家介紹這款開源工具：Ollama

什麼是 Ollama？



Ollama 是一個開源框架，旨在簡化大型語言模型（LLM）在本地運行的流程。它將複雜的模型權重、配置與運行環境封裝在一起，讓你只需一行指令，就能在 macOS、Windows 或 Linux 上跑起 AI 模型。它支援目前市面上強大的開源模型，包括 GLM、Minimax、Gemma、Llama3 等。

備注：目前由 OpenAI 推出的 ChatGPT、Anthropic 推出的 Claude 並沒有開源。

你的電腦跑得動嗎？

想要在本地順暢運行AI，硬體效能，特別是記憶體及顯示卡是關鍵的要素。AI模型的運算高度依賴 VRAM（顯存）或統一記憶體（Unified Memory）。以下是針對不同規模模型的硬體建議：

基礎級（適合 1B~4B 參數模型，如 Phi-3，Gemma 2B）

- **應用場景：**回答速度快，適合簡單的文字摘要、翻譯。
- **記憶體／顯存：**至少 8GB
- **硬體設備：**一般辦公筆電、內顯設備或舊款顯示卡

進階級（適合 7B~30B 參數模型，如 GLM4.7，Gemma-9B）

- **應用場景：**這是目前主流的配置，邏輯推演與創意寫作能力有顯著提升，反應速度流暢。
- **記憶體／顯存：**建議 16GB 以上（VRAM或Mac的統一記憶體）
- **硬體設備：**搭載 NVIDIA-RTX3060/4060 以上的 PC，或 Apple M1/M2/M3（16GB RAM版）以上的 Mac

專業級（適合30B以上參數大模型，如 Llama 3-70B，Qwen3-128B）

- **應用場景：**具備強大的推理能力，適合處理複雜的編程問題或深度學術分析。
- **記憶體／顯存：**建議32GB至64GB以上
- **硬體設備：**NVIDIA RTX3090/4090 或多張顯卡並聯，或 Apple Mac Studio / MacBook Pro（Max 晶片、64GB RAM 以上版本）

快速部署指南

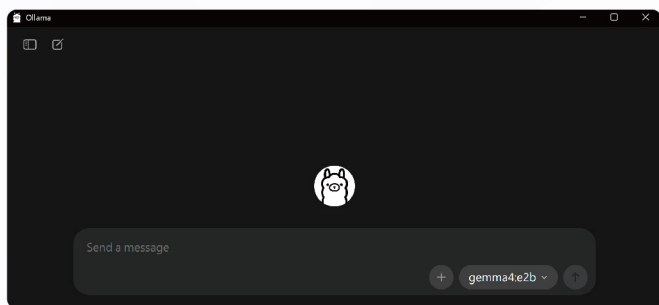
部署Ollama的過程比你想像中還要簡單，幾乎不需要程式背景：

下載與安裝：前往 Ollama 官網（<https://ollama.com/>），下載對應作業系統的 APP 安裝檔，及完成安裝。

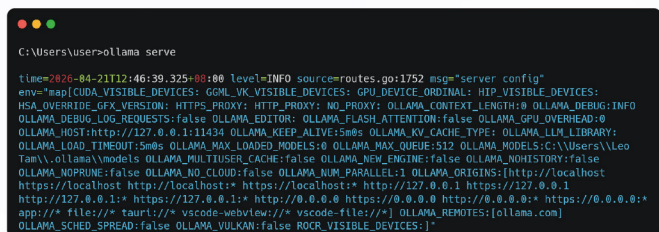
安裝好 Ollama 後，可以在 App 界面選擇想要的 LLM 模型，例如我選擇：gemma4-e2b

系統會自動下載 gemma4-e2b 模型（約7GB）。下載完成後，你就能直接與AI對話！

（Ollama APP 界面）



如果有 APIC all 的需要，也可以透過輸入終端命令 ollama serve 啟用 本機的伺服器，就可以 API 方式使用已安裝的模型。



（Ollama終端伺服器介面）

其他細節及流程，可以查詢官方文件：
<https://docs.ollama.com/>

優點與缺點：

優點：

- **數據主權：**資料完全儲存在你的硬碟中，無需上傳伺服器，從根源上杜絕洩密。
- **零成本運算：**除了初始硬體投資與電費，不再有Token費用，模型可24/7不間斷運行。

缺點：

- **硬體門檻較高：**這是最直接的障礙。如果顯卡或記憶體不夠，跑起來會很慢，甚至連跑起來都很吃力，舊電腦可能無緣體驗。
- **模型能力有上限：**雖然本地模型進步很快，但目前還是雲端的大模型（如 ChatGPT、Claude）比較聰明。本地版適合日常寫作、摘要，遇到特別複雜的邏輯問題，可能還是得求助雲端。

總結來說，Ollama 讓個人 AI 助手不再遙不可及。只要硬體達標，就能在確保隱私與零成本下享受模型便利。

什麼是 OpenClaw？近期在科技圈經常提及的「龍蝦」—— OpenClaw，其實是一個 AI Agent（人工智慧智能體）框架。它主打的概念是「住在你電腦裡的數字管家」。它不只是一個單純問答的聊天機器人，而是一個 24 小時在線、有記憶、能主動幫你執行任務的超級助理。無論是從堆積如山的文件中查閱資料、管理雜亂的桌面硬碟，還是自動撰寫工作文檔，只要你動動嘴下達指令，它就能在後台默默幫你搞定。

2026 年初，隨著各大科技廠商陸續佈局 AI Agent 領域，OpenClaw 憑藉著「開源、本地部署、高度定制」的優勢迅速走紅。然而，許多市民在躍躍欲試時，卻忽略了系統安全性的重要。我們必須有意識地為這隻「龍蝦」配置一個合適的「魚缸」——也就是說，將它部署在不同的環境，會直接影響你的隱私安全與電腦壽命。

三大部署方式的深度抉擇及起步指引

一、本機安裝

- **特點分析：**選擇在本機電腦直接安裝，優點在於所有資料始終存放在自己的物理裝置內，且反應速度最快。由於 OpenClaw 能直接讀取硬碟文件，它的靈活性極高。但缺點也顯而易見：AI 的權限極大，一旦因為理解錯誤而下達刪除指令，可能導致系統崩潰或重要文檔丟失。
- **操作指引：**第一步是先檢查電腦是否具備足夠的圖形處理器（GPU）效能。建議初學者前往 OpenClaw 官網下載「環境自檢工具」，若確認要安裝，請優先選擇一台「非工作用」的舊電腦作為測試機來部署，避免 AI 誤操作影響日常辦公。

二、雲端伺服器

- **特點分析：**目前許多互聯網大廠都推出了「一鍵部署」服務。其最大優勢是穩定性極高且不佔用自家電腦的硬體資源，無論你在公司、家中還是咖啡廳，只要有網路就能存取你的 AI 管家。然而，這意味著你的所有對話紀錄與文件都存放在服務商的伺服器上，你必須完全信任該廠商的資料保護政策與保安技術。
- **操作指引：**第一步是選擇信譽良好的平台，例如選擇阿里雲、騰訊雲、華為雲等大廠的應用。進入官網後一般都會找到類似「OpenClaw 一鍵部署」之類的功能，選擇最基礎的入門配置（通常每月僅需數十元），接下來都是比較傻瓜化的操作，系統會自動為你完成所有代碼安裝，接下來，你選擇連接 Whatsapp，Telegram，飛書等常見社交工具，便可以像和朋友聊天一樣向“龍蝦”發指令了。



三、本機虛擬機（VM）：

這是在電腦中劃分出一塊「獨立區域」來運行 OpenClaw。這種方式兼顧了私隱與安全性——即便AI在虛擬機內發生誤操作，也不會傷及你電腦的主系統。你可以像開關軟體一樣管理它，重新開啟即可恢復初始狀態。對於希望在私人電腦上體驗AI，又擔心系統受損的人來說，這是不錯的選擇。

- **操作指引：**第一步是安裝虛擬化軟件，例如「VMware Workstation Player」或「VirtualBox」等。安裝好軟體後，在裡面創建一個獨立的操作系統，再將 OpenClaw 安裝在裏面。這就像是在你的大房子裡蓋了一間專屬的「龍蝦魚缸」，將風險完全隔離在「魚缸」之外。腦作為測試機來部署，避免 AI 誤操作影響日常辦公。

在安裝好後，大家使用時也要注意遵守各種安全原則：

- **保護個人資訊：**嚴禁輸入身份證號碼、銀行帳戶、密碼等高度敏感數據，避免 AI 在學習過程中洩露隱私。
- **最小權限授權：**在安裝插件時，若 AI 請求存取你的相機或全硬碟權限，請先核實其功能必要性，非必要應予拒絕。
- **數據定期審計：**上傳檔案前先行檢查內容；並定期檢視 AI 工具連接的第三方服務，及時撤銷不再使用的權限。

結語：

技術如何日新月異，OpenClaw 是強大的助手，只是安裝好，能正常與「龍蝦」對話只是第一步，要讓「龍蝦」能真正協助你工作，便要花時間去「養龍蝦」。往後，筆者將繼續為大家介紹更多有關 OpenClaw 的進階應用技巧。

