

繁體中文 AI 書齋版

AI 解讀古代文字

AI 復原的古代記錄

金京鎮律師

kimkj.com • AI 書齋 • 2026

AI 解讀古代文字

AI 復原的古代記錄

金京鎮律師

被火燒成焦炭的卷軸、深埋泥土中的木片、破碎散落的泥板，長久以來都是無法解讀的物品。本書追尋打破這份沉默的十二起事件。從深埋於火山灰中的赫庫蘭尼姆卷軸無需展開便能讀取的虛擬展開、以X-ray重現的恩戈地《利未記》、依筆跡辨別抄寫員人數的以賽亞卷軸鑑定、拼合數萬塊碎片的數位黏合劑、《吉爾伽美什史詩》的復原、阿卡德語帳簿的翻譯、甲骨文的三維虛擬綴合、新羅木簡的草書判讀、處理線形文字A與印度河文字及朗戈文字的計算語言學、填補殘損希臘銘文的語言模型，一路延伸至即使原件消失亦能留存的多光譜成像檔案庫。各章依序由事件檔案、科學調查隊的追蹤、揭曉的真相所構成。為了讓沒有背景知識的讀者也能輕鬆閱讀，書中以淺顯易懂的方式解說專門術語，並在每個子章節附上經查證的參考資料。附錄中則彙整了讀者在家中即可直接查閱的公開資料管道。本書係根據 NotebookLM 筆記「基於AI之古代語言與古文獻解讀研究典藏 (2024-2026)」中所蒐集的 251 件資料撰寫而成。

公開研究資料庫

本書係透過人工智慧整理製作而成，蒐集並綜合了以全球主要語言撰寫的資料。無論資料是以何種語言寫成，NotebookLM 均不會限制提問所使用的語言。讀者可以使用自己的語言，向籌備本書時所使用的同一個研究資料庫提問。即使部分資料是以韓文、英文或其他語言寫成，情況亦然。只要在 NotebookLM 中開啟設定，將輸出語言指定為偏好的語言，回答與生成結果便會以該語言呈現。

籌備本書時所使用的公開 NotebookLM 資料庫（基於AI之古代語言與古文獻解讀研究典藏 2024-2026，資料 251 件）：<https://notebook.google.com/notebook/cf32a6a1-66cc-473d-99d8-ed5a51596eb3>

本書的韓文原版位於 kimkj.com AI書齋：<https://kimkj.com/ai-library/?mod=document&uid=12641>

目錄

點擊各章標題即可前往該章正文。若想從頭開始閱讀，請選擇第1章。

- 第1章. 火山灰吞噬的哲學家書房：赫庫蘭尼姆的碳化卷軸 • 4
- 第2章. 焦黑聖櫃中綻放的話語：隱基底封印的《利未記》 • 11
- 第3章. 隱藏在以賽亞卷軸中的兩道身影：抄寫員鑑定 • 17
- 第4章. 拼合數萬片泥板與羊皮紙拼圖的數位黏著劑 • 23
- 第5章. 拼合泥板的碎裂斷面：吉爾伽美什史詩的計算學重建 • 30
- 第6章. 古代美索不達米亞收據與商人的私語：AI翻譯 • 37
- 第7章. 沉睡在龜腹甲中的帝國未來：甲骨文與3D虛擬綴合 • 44
- 第8章. 從泥坑中歷劫生還的毛筆字：新羅木簡與草書解讀人工智慧 • 51
- 第9章. 愛琴海失落的謎團：線形文字A與幾何網格 • 58
- 第10章. 破解太陽週期與動物印章的密碼：印度河文字與朗戈朗戈文字的聲音 • 65
- 第11章. 人工智慧，重讀殘破的希臘銘文：基礎模型與 Ithaca • 73
- 第12章. 永不焚毀的虛擬圖書館：超高解析度多光譜檔案庫與未來 • 80
- 附錄. 專屬我的AI考古漫遊記：為大眾打造的全球探索入口網站 • 87

第1章. 火山灰吞噬的哲學家書房：赫庫蘭尼姆的碳化卷軸

1.1. 案件檔案

火山掩埋書房之日

西元79年，義大利南部的維蘇威火山爆發。火山一口氣噴出大量高溫氣體與火山灰。山腳下的海濱有一座名為赫庫蘭尼姆的城市。那是羅馬富豪消暑的度假城市。這座城市也在短短一天內被掩埋於地下[1]。

城外矗立著一座大型宅邸。如今學者們稱此處為紙草莊園。屋主推測為盧基烏斯·卡爾普爾尼烏斯·皮索·凱索尼努斯。他是尤利烏斯·凱撒的岳父。希臘哲學家菲洛德穆斯經常出入這所宅邸。菲洛德穆斯是伊比鳩魯學派的學者。宅邸一隅有一間收藏卷軸的書房。由於從中出土了特別多他的著作，因此被視為與他關係密切的藏書庫[1]。

莎草紙是古代地中海世界的紙張。它是用生長在埃及尼羅河畔的植物製成的。將植物莖部的髓心切成薄片，橫向一字排開。在其上方再縱向鋪上一層。噴水加壓後，植物的汁液便發揮黏合劑的作用。乾燥後便成為一張紙。因此莎草紙表面會留下縱橫交錯的紋理。這些紋理日後成為重要的線索[7]。

那一天，高溫氣體湧入書房。卷軸在沒有明火的情況下直接被高溫烘烤。扮演紙張角色的莎草紙變成了焦黑的炭。像這樣燃燒成炭的過程稱為碳化。卷軸萎縮，表面隆起[2]。其上方堆積了超過20公尺的火山沉積物。火山沉積物是火山灰與岩石碎屑堆積凝固而成的地層[1]。

深埋地下後便隔絕了空氣。接觸不到空氣，植物纖維就不易腐爛。這場災難反而發揮了保護文物的封蓋作用。卷軸就這樣挺過了2,000年[2]。

地下隧道中出土的黑色棒狀物

1738年，拿坡里王室開始開掘這片區域。1750年，打井的工人發現了色彩斑斕的馬賽克地板。宅邸的位置就此曝光。發掘工作由卡爾·韋伯負責。他是來自瑞士的軍事工程師[1]。

韋伯在堅硬的火山沉積物中垂直向下開鑿坑道。這種坑道稱為垂直豎井。隨後又如迷宮般開鑿出狹窄的隧道。隧道內昏暗且空氣渾濁。壁面坍塌的危險也始終存在。即便在這種條件下，韋伯仍仔細繪製了宅邸的平面圖。多虧了那張地圖，後人才能知曉哪個房間出土了什麼物品[1]。

1752年起，工人們開始取出宛如黑色木炭棒的物品。起初以為那是木炭塊。查明後才發現那是莎草紙卷軸。據悉從這座宅邸出土的卷軸約有1,100卷。這是古代世界藏書庫完整留存至今的唯一已知案例[1]。

這處藏書庫中菲洛德穆斯的著作格外豐富。定稿、抄本與草稿擺放在一起。學者們認為他曾在此宅邸中著書立說並整理文稿[1]。費盡千辛萬苦展開的卷軸中，記載著關於哲學、修辭學與倫理學的文章。修辭學是研究如何以言語和文字說服他人的學問。至於未能展開的卷軸中記載了什麼，當時無人知曉[2]。

強行展開而損毀的卷軸

學者們渴望閱讀其中的文字。問題在於卷軸一經觸碰便會碎裂[2]。1753年，安東尼奧·皮亞喬神父製造了一台機器。那是一具利用絲線與鉛錘一點一點拉開卷軸的裝置[3]。這台機器雖然搶救了幾卷，但也撕裂了內層[2]。

到了19世紀，英國化學家漢弗里·戴維出手嘗試。他在卷軸上使用了氯氣、碘蒸氣與酸。這是試圖藉由化學藥品將紙層逐一分離的方法。結果並不理想。藥品使莎草紙軟化崩解[3]。

以物理方式強行展開的嘗試大多損壞了文物。寫有文字的碎屑散落一地甚至化為烏有。最終學界停止了開啟的嘗試。時至今日，仍有數百卷未曾打開的卷軸留存[2]。這些文物被分別保存在拿坡里、巴黎與牛津的典藏機構中[4]。

擋在X-ray面前的碳之壁

科學家們開始尋找不打開卷軸也能閱讀的方法。首先想到的是電腦斷層掃描。這就是通常所稱的CT技術。CT是一種無需切開物體即可測量內部密度差異的裝置。密度較高的物質會大幅阻擋X-ray。因此在斷面影像中會呈現明亮色調[5]。

含有金屬成分的墨水在CT影像中會顯現白色光亮。在中世紀歐洲，含有鐵成分的墨水被廣泛使用。以色列隱基底出土的碳化卷軸也類似。這件文物是古代以色列的希伯來語抄本。由於墨水中似乎混有金屬，因此在影像上留下了痕跡。於是在2016年，它在未被展開的情況下被成功判讀[5]。

赫庫蘭尼姆抄寫員的情況則不同。抄寫員是以手工抄寫文字的人。他們將燃燒木材產生的煙灰與水和樹脂混合製成墨水。墨水的主要成分是碳。然而，被火山高溫烘烤的莎草紙也變成了碳。紙張與字跡變成了同一種物質[6]。

若密度相同，X-ray穿過這兩個部位的程度是一樣的。斷面影像上不會出現分界線。學者們將這道障礙稱為「碳墨難題」。2015年，義大利與法國的研究團隊嘗試了另一種方法。這是一種名為相位差成像的技術。它利用了X-ray在穿過物質邊界時行進方向會略微偏折的特性。即使密度相同，界面處的訊號也會有所不同。研究團隊利用這種方法，在捲曲的卷軸內部捕捉到了幾個字母的輪廓。但還達不到能夠閱讀句子的程度[6]。長期以來，人們都相信這道難關無法跨越[2]。

參考資料

- [1] Herculaneum Society. The Villa of the Papyri. University of Oxford. <https://www.herculaneum.ox.ac.uk/index.php/the-villa-of-the-papyri/> (2026年9月3日檢索)
- [2] Vesuvius Challenge. FAQ. scrollprize.org. <https://scrollprize.org/faq> (2026年9月3日檢索)
- [3] Villa of the Papyri. Wikipedia. https://en.wikipedia.org/wiki/Villa_of_the_Papyri (2026年9月3日檢索)
- [4] University of Kentucky. (2026, 6月 25日). The day the Herculaneum scrolls began speaking again. UKNow. <https://uknow.uky.edu/research/day-herculaneum-scrolls-began-speaking-again> (2026年9月3日檢索)
- [5] Seales, W. B., Parker, C. S., Segal, M., Tov, E., Shor, P., & Porath, Y. (2016). From damage to discovery via virtual unwrapping: Reading the scroll from En-Gedi. Science Advances, 2(9), e1601247. <https://doi.org/10.1126/sciadv.1601247>
- [6] Mocella, V., Brun, E., Ferrero, C., & Delattre, D. (2015). Revealing letters in rolled Herculaneum papyri by X-ray phase-contrast imaging. Nature Communications, 6, 5895. <https://doi.org/10.1038/ncomms6895>
- [7] Papyrus. Encyclopaedia Britannica. <https://www.britannica.com/topic/papyrus-writing-material> (2026年9月3日檢索)

1.2. 科學調查隊的追蹤 1：剝開虛擬洋蔥皮

不打開卷軸也能閱讀的構想

美國肯塔基大學的布倫特·席爾斯教授選擇了另一條路。他的構想是不去打開文物實體，而是在電腦內部將其展開。保留實物原樣，僅獲取三維影像。接著在螢幕上一層一層剝離紙層。這是一種宛如剝洋蔥皮般逐步深入內部的方式[1]。

研究團隊將這種方法命名為虛擬展開。英文稱為 Virtual Unwrapping。2016年，該方法首次獲得成功。研究對象是前一節提到的隱基底卷軸。該文物於1970年在古猶太會堂遺址出土。外觀是一團焦黑的炭塊。研究團隊在螢幕上將其展開後，希伯來文字跡赫然顯現。那是舊約《利未記》的開篇[1]。

虛擬展開由五個步驟組成：獲取三維影像、分層、賦予墨水數值、展平、拼接。在隱基底抄本中，這一順序完全行得通。因為墨水在影像中已經清晰可見[1]。

赫庫蘭尼姆卷軸的條件比隱基底抄本更為惡劣。受到火山高溫擠壓，其截面扭曲變形。原本的圓筒形狀坍塌，層與層之間緊緊沾黏在一起。加上墨水在影像中完全看不到。在五個步驟中，第二步與第三步同時成為掣肘[2]。

因此在赫庫蘭尼姆卷軸上改變了順序。先進行分層並將其展平。尋找墨水則在此之後單獨進行。1.3節便是講述尋找墨水的故事[2]。

粒子加速器拍攝的三維影像

第一步是完整拍攝卷軸內部[1]。醫院的CT通常只能分辨1毫米左右的細節。而莎草紙的層次遠比這更薄。單層的厚度甚至不及頭髮絲的粗細。使用醫院設備會導致層與層之間模糊重疊[3]。

研究團隊求助於粒子加速器。最初造訪的是英國的鑽石光源[3]。後來也一併使用了法國格勒諾勃的歐洲同步輻射研究設施[4]。這些設施能產生極為明亮且沿直線行進的X-ray。光線越強，越能在短時間內獲得清晰的影像[3]。

研究團隊將卷軸立在圓筒內牢牢固定。因為在拍攝過程中稍有晃動，影像就會變模糊。由此獲得的資料解析度約為單一像素8微米。像素是指連深度數值也一併具備的立體畫素。可以將其視為把小畫家中的一個點換成了微小的正立方體[3]。

資料量極其龐大。僅拍攝一卷卷軸就會累積約300TB。這是足以填滿數千支智慧型手機儲存空間的容量。處理這些資料本身就成了一項新課題[4]。

法國格勒諾勃的設施被稱為第四代光源。2020年新設備投入運作。歐洲21個國家投入1.5億歐元提升其性能。該設施最初是為了探究物質與生命現象而建立的。起初誰也沒料到它會被用來解讀古代卷軸[4]。

剝離紙層的體積分割

獲得三維影像後，下一步是分層階段。這項作業稱為體積分割。也就是在立體影像中將紙層逐一分離的工作。電腦會追蹤每一張斷面影像中紙張表面走過的路徑。將這些線條拼接在一起，便會產生單張紙的立體模型。這個模型稱為網格[1]。

最初是由人工用滑鼠一點一點標記邊界[2]。一卷卷軸所包含的紙張長度超過1公尺[4]。單靠人工追蹤實在太長且太過密集。加上紙層相互壓合黏結，使得邊界十分模糊[2]。

隨著參賽者們推出自動化工具，處理速度大幅提升。瑞士的朱利安·施利格開發的程式便是代表。該程式順著相鄰斷面影像的紋理連接紙層。將人工需要數天才能完成的工作縮短至幾小時[2]。

這項工作規模之大，絕非單一實驗室所能獨力承擔。主辦方直接將掃描資料公開在網路上。參賽者們也公開了自己開發的程式原始碼。後人在他人打造的工具之上進一步疊加新功能。判讀出的文本與程式碼至今仍開放所有人下載[5]。

展平並拼接

分層之後，紙張模型依然呈捲曲狀。維持原樣的話，字跡看起來會是扭曲的。因此需要進行將模型展平的運算。這一步驟稱為平坦化[1]。

電腦將模型的頂點轉換為微小的珠子。並假設珠子與珠子之間由彈簧相連。接著反覆計算將這張網鋪平在地面上的過程。使彈簧伸長或收縮的幅度降到最低。如此一來，文字的形狀便不會嚴重變形[1]。

最後是拼接。卷軸無法一次全部展開。電腦將其分成數個區域分別展平。必須把展平的各個區域重新拼合為一張。此時會利用莎草紙特有的網狀纖維紋理。若兩個區域的纖維紋理能夠吻合，即代表它們原本是相鄰的。完成這項比對作業後，虛擬展開便大功告成[1]。

這五個步驟無法一次完成。任何一個步驟出錯，後續步驟都會全部偏差。若紙層劃分錯誤，其他頁面的文字就會混雜進來。因此研究團隊會對同一卷卷軸反覆進行處理。用肉眼檢視結果並調整設定後重新執行[5]。

參考資料

[1] Seales, W. B., Parker, C. S., Segal, M., Tov, E., Shor, P., & Porath, Y. (2016). From damage to discovery via virtual unwrapping: Reading the scroll from En-Gedi. *Science Advances*, 2(9), e1601247.

<https://doi.org/10.1126/sciadv.1601247>

[2] Vesuvius Challenge. (2024, 2월 5일). Vesuvius Challenge 2023 Grand Prize awarded: we can read the scrolls! scrollprize.org. <https://scrollprize.org/grandprize> (2026年9月3日檢索)

[3] Vesuvius Challenge. FAQ. scrollprize.org. <https://scrollprize.org/faq> (2026年9月3日檢索)

[4] University of Kentucky. (2026, 6월 25일). The day the Herculaneum scrolls began speaking again. UKNow. <https://uknow.uky.edu/research/day-herculaneum-scrolls-began-speaking-again> (2026年9月3日檢索)

[5] Vesuvius Challenge. (2026, 6월 25일). An entire Herculaneum scroll has been read for the first time. scrollprize.org. <https://scrollprize.org/firstscroll> (2026年9月3日檢索)

1.3. 科學調查隊的追蹤 2：尋找墨水的微觀指紋

寫在碳上的碳

即便虛擬展開完成，螢幕上依然是一片空白。紙張雖然展平了，卻看不到文字[1]。原因正是前面提到的碳難題。莎草紙是碳，墨水也是碳。X-ray無法區分兩者[2]。

席爾斯教授的研究團隊提出了一項假說。那就是墨水絕不可能完全沒有留下任何痕跡[1]。墨水以液體狀態接觸紙張。液體會滲透至纖維之間。並在乾燥過程中略微拉扯表面[3]。

如此一來，紙張表面的形狀就會產生微小的變化。即使密度相同，起伏凹凸也會有所不同。研究團隊決定找出這種起伏[3]。由於那是肉眼不可見的微小尺度，因此只能交由機器處理[1]。

此處使用的工具是機器學習。這並非由人工逐一編寫規則的方式。而是向其展示大量附有答案的範例，讓它自行找出規則。這與向其展示數千張貓的照片後，它便能認出貓是同樣的道理。在尋找墨水時，則是向其展示標記了有墨水處與無墨水處的樣本[1]。

名為Crackle的線索

2023年3月，一場名為維蘇威挑戰賽的競賽拉開帷幕[1]。由席爾斯教授與納特·弗里德曼、丹尼爾·格羅斯共同發起[4]。主辦方將掃描資料免費公開在網路上。設立獎金號召全球程式設計師參與[1]。

最初幾個月裡，沒有人找到墨水。那年夏天，一位名叫凱西·漢德默的參賽者找到了線索。他只是用肉眼長時間凝視著展開後的表面資料。隨後他發現了一種特殊的紋理。那是細密龜裂的痕跡[1]。

想像乾旱時的稻田地面即可。乾涸的泥土會像網狀般裂開。沾有墨水的地方也出現了類似的細紋。參賽者們將這種紋理稱為龜裂紋（crackle）。漢德默憑藉這項發現獲得了首個墨水獎[1]。

漢德默在發現後立刻在線上聊天室公布。這次分享改變了比賽的走向。其他參賽者隨即將這種紋理當作訓練材料[1]。

只透過小視窗檢視

美國大學生盧克·法里托搶先行動。他將龜裂紋理教給機器學習模型。模型以機率回答何處有墨水。隨著訓練進行，答案變得越來越準確[1]。

問題在於訓練資料太少。如果資料太少，機器就會把答案死記下來。這被稱為過度擬合。這就像是不理解考題、只死記答案的學生一樣[1]。

參賽者們藉由旋轉與翻轉資料來增加數量。即使旋轉圖片，細紋的性質依然不變。機器也因此只掌握了紋理本身的特徵[1]。

他們也調整了輸入視窗的大小。起初是一次展示較寬廣的區域。範圍一寬，整個字母的形狀就會盡收眼底。於是機器開始憑空捏造字母。在沒有墨水的地方畫上看似合理的筆畫。這被稱為人工智慧的幻覺[1]。

研究人員將視窗縮小至64乘64畫素。在那樣的範圍內，字母的形狀並不會顯現。機器只能根據細紋的質感來進行判斷。憑空捏造的情況大幅減少。取而代之的是，真實的筆觸如實呈現在螢幕上[1]。

負責這項計算的模型是由兩種架構串聯而成。其中之一是3D卷積神經網路。這種結構能掃描微小立方體內的紋理並擷取特徵。另一個則是Transformer。這種結構用來分析相互分離的訊號彼此有多大關聯。兩者結合使用，就不會遺漏薄層中模糊的紋理[1]。

機器產生的畫面並不會立刻被視為正確答案。古文字學專家會另外解讀該畫面。檢視希臘語單字的排列是否符合文法。也確認前後相連的句子語意是否通順。若機器與人工的解讀不一致，該對象就會重新計算[1]。

僅憑表面起伏

此後，尋找墨水的方法仍持續進步。參賽者優素福·納德爾採用了預先學習的方法。這種方式是先向機器展示大量沒有標記答案的卷軸資料。機器會自行掌握紙草纖維的外觀形狀。接著再用少量的標記答案僅教它尋找墨水[1]。

2026年3月，一項探討起伏本身的研究發表了。該研究由喬治·安傑洛蒂、費代麗卡·尼科拉爾迪、保羅·亨德森與布倫特·席爾斯共同撰寫。研究團隊使用光學儀器測量了已展開的紙草樣本表面。將僅包含高低起伏的資料讓機器學習。並未放入亮度或密度資訊[3]。

即便如此，機器依然辨別出了有墨水的地方。這意味著僅憑表面的高低起伏就能找到墨水。研究團隊還一同計算出測量需要達到多細密的程度。這個數值成為下次拍攝的解析度基準。該論文於同年刊登在學術期刊《Scientific Reports》上[3]。

儘管如此，尋找墨水仍非成熟完善的技術。墨水殘留較淡的區段依然顯得模糊。在層疊沾黏處，拍攝資料本身就一片模糊。有時兩張不同的紙甚至會黏合在一起。在某個卷軸上表現良好的模型，到了另一個卷軸上卻可能徒勞無功。因此，每個環節都需要人工以肉眼確認並動手調整。研究至今仍在一步步向前邁進[5]。

參考資料

- [1] Vesuvius Challenge. (2024, 2월 5일). Vesuvius Challenge 2023 Grand Prize awarded: we can read the scrolls! scrollprize.org. <https://scrollprize.org/grandprize> (2026年9月3日檢索)
- [2] Mocella, V., Brun, E., Ferrero, C., & Delattre, D. (2015). Revealing letters in rolled Herculaneum papyri by X-ray phase-contrast imaging. Nature Communications, 6, 5895. <https://doi.org/10.1038/ncomms6895>
- [3] Angelotti, G., Nicolardi, F., Henderson, P., & Seales, W. B. (2026). Ink detection from surface topography of the Herculaneum papyri. Scientific Reports. <https://doi.org/10.1038/s41598-026-58467-1>
- [4] University of Kentucky. (2026, 6월 25일). The day the Herculaneum scrolls began speaking again. UKNow. <https://uknow.uky.edu/research/day-herculaneum-scrolls-began-speaking-again> (2026年9月3日檢索)
- [5] Vesuvius Challenge. (2026, 7월 10일). Open problems: why reading every Herculaneum scroll is still a challenge. scrollprize.org. https://scrollprize.org/2026_open_problems (2026年9月3日檢索)

1.4. 揭曉的真相

2023年找回的第一個單字

最初的成果是一個單字。2023年10月12日公布了這項消息。主角是美國的大學生盧克·法里托。當時他年僅二十一歲。他讀出的單字是波菲拉斯（porphyras）。這是代表紫色的希臘語。這在古代文獻中並非尋常詞彙。法里托因獲得首字獎而得到4萬美元。第二名的優素福·納德爾獲得了1萬美元。這是首次在未展開的卷軸中讀出單字的案例[7]。

2024年2月5日大會公布了比賽結果。大獎由三人共同獲得。分別是埃及的優素福·納德爾、美國的盧克·法里托，以及瑞士的尤利安·席利格。獎金為70萬美元[1]。

他們所解讀的文物是法蘭西學會收藏的卷軸。名稱為PHerc. Paris. 4。三人從這卷卷軸中復原了2,000多個希臘文字母。那是十五欄的文字。欄是指卷軸上縱向分段書寫的文字區塊。相當於整個卷軸約5%的份量[1]。

評審標準相當嚴格。必須提交四處各140字的段落。每個段落都必須解讀出85%以上的文字才算通過。古文字學專家在隱去參賽者姓名的情況下進行了審查[1]。

內容是關於享樂的探討。文中探討稀少之物是否必然比常見之物帶來更多快樂。正文中有一段以食物為例的文句。其主旨是，並不認為稀有之物就必然會讓人更感快樂。這段文字充分體現了伊比鳩魯學派的思想[1]。

找回標題的卷軸

下一項成果則是標題。2025年5月5日頒發了首座標題獎。得獎者是馬塞爾·洛特與米哈·諾瓦克。獎金為6萬美元[2]。

目標文物是牛津大學博德利圖書館收藏的PHerc. 172。螢幕上顯現的文字是菲洛德穆斯的《論惡德》。經研讀為第一卷。這一點由拿坡里費德里克二世大學的費代麗卡·尼科拉爾迪研究團隊確認。這是首次在不展開的情況下讀出捲起卷軸標題的案例[2]。

在這卷卷軸中，隨後又解讀出了70多欄內容。這是探討貪婪、傲慢與諂媚等性格缺陷的文章。如今人們已能從頭梳理其論述脈絡[3]。

首次完整讀完的一卷

2026年6月25日傳來了更重大的消息。那就是從頭到尾完整讀完了一整卷卷軸的發表。文物是拿坡里國家圖書館收藏的PHerc. 1667。比賽參賽者們稱之為4號卷軸[4]。

拍攝是在法國的歐洲同步輻射設施進行。之後依序經過了虛擬展開與尋找墨水。復原文字的長度約1.4公尺。分成大約二十二欄。完整讀完一整卷赫庫蘭尼姆卷軸，這還是第一次[4]。

內容是探討倫理與人格特質的希臘語論著[3]。正文中提及了斯多葛學派哲學家克律西波斯的侄子。其名為阿里斯托克瑞翁[4]。成書年代推測為西元前3世紀末至2世紀左右。作者身分目前尚未確定[3]。

研究結果也以論文形式公開。包括喬治·安傑洛蒂、史蒂芬·帕森斯與費代麗卡·尼科拉爾迪在內，共有27人共同署名[5]。資料、程式碼與解讀文本亦一併公開[4]。

研究團隊同時公布了驗證流程。他們以更高的解析度重新拍攝了2023年解讀的PHerc. Paris. 4。由新資料得出的解讀文本與過去的結果完全吻合。這證明了文字並非機器憑空捏造。這也意味著相同的方法可以應用到其他卷軸上[4]。

仍剩下的600卷

同一份發表中也包含了另一項消息。編號為PHerc. 139的卷軸標題被解讀了出來。那是菲洛德穆斯的《論眾神》第八卷。在此之前，這部著作僅為人所知有第一卷。這是首次揭露該著作篇幅超過八卷的事實[3]。

尚未解讀的卷軸仍有600多卷[3]。大會設立了2027年大獎。總獎金為100萬美元。第一名可獲得80萬美元。競賽目標為指定的十三卷卷軸。各欄必須解讀出殘留文字的70%以上。截止日期為2027年6月25日[6]。

這段故事中仍有許多空白。即使讀出了文字，領會其含義仍是另一項工作。光是從學術角度整理一欄解讀文本，也需要耗費漫長的時間。此外也還有作者尚未查明的卷軸[3]。

儘管如此，方向已經明確。兩千年來無人能開啟的典籍，在螢幕上被展開。原以為已被燒毀殆盡的藏書，正在一點一滴重見天日[4]。

成就這項事業的並非單獨一人。有負責拍攝的物理學家，也有負責分層的程式設計師。還有為大家解讀古希臘語的古文字學家。參加比賽的學生與上班族也出了力。歷經兩千年後典籍得以重見天日，正是這些人共同創造的成果。剩餘的卷軸也可以透過相同的方式被開啟[4]。

參考資料

[1] Vesuvius Challenge. (2024, 2월 5일). Vesuvius Challenge 2023 Grand Prize awarded: we can read the scrolls! scrollprize.org. <https://scrollprize.org/grandprize> (2026年9月3日檢索)

[2] Vesuvius Challenge. (2025, 5월 5일). \$60,000 First Title Prize awarded! scrollprize.substack.com. <https://scrollprize.substack.com/p/60000-first-title-prize-awarded> (2026年9月3日檢索)

[3] University of Kentucky. (2026, 6월 25일). The day the Herculeaneum scrolls began speaking again. UKNow. <https://uknow.uky.edu/research/day-herculeaneum-scrolls-began-speaking-again> (2026年9月3日檢索)

[4] Vesuvius Challenge. (2026, 6월 25일). An entire Herculeaneum scroll has been read for the first time. scrollprize.org. <https://scrollprize.org/firstscroll> (2026年9月3日檢索)

- [5] Angelotti, G., Parsons, S., Nicolardi, F., Nader, Y., Johnson, S., Josey, D., et al. (2026). Complete virtual unwrapping and reading of a rolled Herculaneum papyrus. arXiv:2606.29085. <https://arxiv.org/abs/2606.29085> (2026年9月3日檢索)
- [6] Vesuvius Challenge. Open Prizes. scrollprize.org. <https://scrollprize.org/prizes> (2026年9月3日檢索)
- [7] Vesuvius Challenge. (2023, 10월 12일). First word discovered in unopened Herculaneum scroll by 21yo computer science student. scrollprize.org. <https://scrollprize.org/firstletters> (2026年9月3日檢索)

第2章. 焦黑聖櫃中綻放的話語：隱基底封印的《利未記》

2.1. 案件檔案

死海西側的綠洲

隱基底是位於死海西岸的綠洲，有時也寫作恩戈地。周圍雖是乾涸的沙漠，但這裡有泉水湧出。有了水，人們便聚集於此。數千年來，人們在這片水邊建立聚落生活。舊約聖經《撒母耳記上》中也出現了這片土地的名字，記載著大衛為了躲避掃羅王而藏身於隱基底曠野的故事。

根據遺跡調查，這座丘陵從西元前8世紀後半葉開始就有人居住。後來在西元600年左右的一場大火中，村落化為廢墟[1]。這是一片將一千三百多年的歲月層層堆積在泥土中的土地。

死海周圍降雨稀少，空氣乾燥。這種氣候有助於古物長久保存。木材、皮革和布料比其他地方更不易腐爛。死海古卷出自該地區，原因也在此。

1970年，這裡展開了大規模發掘。發掘工作由丹·巴拉格（Dan Barag）和埃胡德·內策爾（Ehud Netzer）主持，兩人當時任職於耶路撒冷希伯來大學考古研究所。約瑟夫·波拉特（Yosef Porath）於1970年5月5日挖掘出了這卷卷軸[1]。發掘團隊找到的是古猶太會堂的遺址。會堂是猶太社群聚集祈禱與研讀聖經的建築，地面上鋪設著以不同顏色小石子鑲嵌而成的馬賽克拼貼。

會堂內有擺放聖櫃的空間。聖櫃是存放神聖聖經卷軸的櫃子。猶太社群將這個櫃子視為會堂中最神聖之處。發掘團隊從該聖櫃中清除了燃燒後的殘骸，並在殘骸中發現了一塊黑色物體[1]。

出自聖櫃這一事實本身就是一條重要線索。會堂的聖櫃裡不會隨便放入任何文件，只會存放社群研讀並恪守的聖經卷軸。這為推斷該黑色物體的真實身份提供了依據。

古代聖經抄本本身就是珍貴的資料。在沒有印刷機的時代，全靠人工手抄。每次謄抄都有文字產生變動的可能。一旦有古老抄本面世，就能確認經文是如何傳承流變的。這就是學者們對出自聖櫃的黑色物體孜孜以求的原因。

大火留下的黑色塊狀物

火是抹除記錄的力量。無論是紙張還是皮革，在烈火中都會化為灰燼。但在隱基底，情況有些不同。被殘骸壓住的卷軸在缺氧狀態下被緩慢烘烤，無法散成灰燼，反而像木炭一樣凝固硬化。這種如木炭般硬化的轉變稱為碳化。在空氣不足的火勢中燃燒時，就會發生這種現象。

這與炭窯中木材轉化為木炭的過程與原理相同。若直接燃燒木材，只會剩下灰燼；若用泥土覆蓋隔絕空氣再行烘烤，就會留下黑色的木炭。碳化卷軸也是這樣形成的。

碳化的塊狀物被掩埋在倒塌建築的灰燼與泥土之下。由於接觸不到空氣，腐朽的過程也隨之停滯。皮革本是極易腐爛的材質，若直接暴露在地表，恐怕早已蕩然無存。摧毀文物的火，同時也是讓文物得以留存的力量。

發掘團隊取出的物體乍看之下像是一塊焦炭。然而在斷裂的側面上，仍殘留著一圈圈捲起的層次。考古學家判定這是一卷捲起的皮革卷軸[1]。這件文物以其發現地點命名，「隱基底卷軸」之名便由此而來。

這卷卷軸是由動物皮革鞣製而成的羊皮紙[1]。而義大利赫庫蘭尼姆出土的卷軸則是由莎草紙製成。莎草紙是將植物莖條壓平黏合而成的植物性紙張。這兩件文物的底材截然不同，這項差異日後成為區分閱讀方法的分水嶺。

羊皮紙是一種製作繁瑣的材料。必須先去除動物皮革上的毛髮與油脂並加以洗淨，再固定於框架上緊繃乾燥。在如此打磨出的薄片上書寫文字。它比紙張更具韌性且保存更久，但造價昂貴。

其出土狀態也非比尋常。古老卷軸多半以碎裂的殘片出土，1947年起在昆蘭洞穴發現的死海古卷也多為碎片。而隱基底的文物出土時，卻完整維持著捲曲的形態，這意味著內部還保留著數層結構。

無法觸碰的文物

要閱讀文字，就必須展開卷軸。然而，這件文物卻無法展開。因為在高溫烘烤下，皮革收縮並緊緊黏在一起。外表堅硬如石，內部卻脆弱易碎，處於指尖一碰就會化為齏粉的狀態[1]。

並非完全沒有辦法。可以用刀剖開層疊，也可以塗抹化學藥劑使皮革軟化。但這兩種方法都會毀壞文物。卷軸一旦破碎，便再也無法復原。

僅憑觀察外觀，甚至無法得知是否存在文字。碳化的表面一片漆黑且毫無紋樣。墨水是黑的，皮革也是黑的，肉眼根本無法分辨。內部寫有文字這一點，也只是從層疊捲曲的形態中所作的推測。

文物工作者始終將這一點放在首位：文物在世上絕無僅有，無法承擔任何一次失敗。因此在文物保護科學中，存在著「不觸碰而進行調查」的原則。這種調查方法稱為非侵入式調查。「侵入」意味著深入內部，而「非侵入」則正好相反。隱基底卷軸正是一件僅允許進行非侵入式調查的文物。

45年的沉默

卷軸被送進了以色列的文物典藏庫。日後由以色列文物局負責保管這件文物[1]。在箱子中等待，是當時唯一的保存方法。典藏庫是一個維持恆溫恆濕的空間，能隔絕空氣、光線與蟲害，防止文物進一步受損。即便是一件無法解讀的文物，維持其現狀的工作仍持續不斷。

負責死海古卷保存工作的普尼娜·肖爾（Pnina Shor）接手了管理工作。日後成功解讀這卷卷軸的研究中，她也是共同作者之一[1]。

這份等待自1970年一直延續至2015年[1][2]。在長達45年的時間裡，沒有人見過其中的文字。由於表面漆黑一片，甚至難以確認是否有字。學者之間甚至流傳著這是一件永遠無法解讀的文物的說法。

等待是有原因的。以1970年的技術，根本無法處理這個塊狀物。將物體內部細緻分層拍攝的成像設備並不普及；即便能夠拍攝，也沒有能夠運算如此龐大數據的電腦。只能任文物維持原狀，靜待技術發展。

在此期間，世界各地也面臨著類似的困擾。義大利的赫庫蘭尼姆卷軸同樣受阻於這道障礙。過去試圖以物理方式展開的嘗試，只不過是摧毀了文物。這兩件文物的出土地與材質雖有不同，但問題卻是一致的：必須找到一種無需觸碰即可閱讀內部的方法。

打破沉默的並非某種精湛的手工技巧，而是在不切開物體的情況下透視內部的成像技術與電腦運算。席爾斯研究團隊長期致力於研究不觸碰文物而閱讀古文字的方法。這項研究獲得了美國國家科學基金會與Google的資助[2]。他們需要一件能夠檢驗長年累積之方法的文物，而隱基底卷軸正成為了這個對象。

2015年，美國肯塔基大學布倫特·席爾斯（Brent Seales）率領的研究團隊接下了這項工作。研究成果於2016年9月21日刊登在學術期刊《科學進展》（Science Advances）上[1][2]。

現在剩下的疑問只有一個：在不觸碰的情況下，究竟是如何讀出內部文字的？答案就在成像設備與運算流程之中。

參考資料

[1] Seales, W. B., Parker, C. S., Segal, M., Tov, E., Shor, P., & Porath, Y. (2016). From damage to discovery via virtual unwrapping: Reading the scroll from En-Gedi. *Science Advances*, 2(9), e1601247.
<https://doi.org/10.1126/sciadv.1601247>

[2] University of Kentucky Research. (2016, 9월 22일). Virtually Unwrapping the En-Gedi Scroll. [research.uky.edu](https://research.uky.edu/news/virtually-unwrapping-en-gedi-scroll).
<https://research.uky.edu/news/virtually-unwrapping-en-gedi-scroll> (2026年9月3日檢索)

2.2. 科學鑑識團隊的追查：X光數位淨化

不切開即可透視內部的掃描成像

研究團隊選用的第一個工具是顯微CT（micro-CT）。CT指的是電腦斷層掃描，原理與醫院用於檢查體內狀況的掃描相同。它從多個角度發射X光以測量物體內部的密度，電腦將這些數值匯總並建立立體地圖。顯微CT則是以極微小的刻度來進行這項工作的設備。

掃描隱基底卷軸使用的是一台名為Bruker SkyScan 1176的設備。負責掃描的是以色列的默克爾技術公司（Merkel Technologies），該公司將掃描數據捐贈給了以色列文物局。掃描刻度為每格18微米[1]。1微米是1毫米的千分之一，相當於頭髮直徑的幾分之一。

這種刻度大小是有原因的。若要區分層疊，測量刻度必須遠小於單層皮革的厚度。如果刻度大於單層厚度，兩層就會模糊地混成一團。掃描刻度是決定能夠辨識之層數極限的關鍵條件。

掃描生成的數據並非單張照片，而是一個由微小立方體顆粒密集堆疊而成的立體塊。這些顆粒被稱為體素（voxel）。如果說照片的像素是平面顆粒，那麼體素就是立體顆粒。每個顆粒中都包含著表示該處密度的數值。

在這個階段，文物在無需觸碰的情況下被轉移到了電腦中。自此之後，要處理的對象不再是皮革，而是數字。數字無論怎樣操作都不會碎裂；若計算有誤，從頭再來即可。這就是非侵入式調查的威力所在。

混有金屬的墨水留下的亮點

獲得了立體數據，並不代表文字就能看見。要顯現文字，墨水與底材的密度必須有所差異。而在赫庫蘭尼姆卷軸中，這項條件並不成立。當地的墨水是由木炭粉製成，而碳化的莎草紙本身也是碳塊。兩者密度相近，因此在X光下難以區分。

隱基底卷軸的條件則較為理想。研究團隊在掃描數據中確認了有文字處的密度高於周圍。因此推測這種墨水中可能含有鐵或鉛等金屬。只不過由於沒有裸露的外表面，無法直接分析其成分。研究團隊在論文中坦承了這項局限[1]。

金屬會大量吸收X光。吸收量多的地方在畫面上呈現為亮部；碳化皮革吸收X光較少，因此顯得暗淡。文字與底材的明暗對比便由此產生。這項差異成為電腦找出文字的第一條線索。隱基底卷軸之所以比赫庫蘭尼姆卷軸更早被解讀，原因也正在於此。

不過，這種墨水的成分至今尚未確定。研究團隊提出的僅是密度較高的觀察結果及相應的推測[1]。古代抄寫員究竟摻入了什麼，必須取得樣本才能知曉。不將未經證實之事當作定論，是該領域的原則。

虛擬展開的四個步驟與斑痕消除

找到了亮點，並不代表就能閱讀文字。文字分散在層層捲曲的曲面上，必須在電腦內部將這些曲面展平。這項工作稱為「虛擬展開」。席爾斯研究團隊將虛擬展開劃分為四個步驟[1]。

第一步是找出層疊的表面。電腦將立體數據切成薄片，逐一標定皮革層的坐標。將收集到的這些點連接起來，建立出薄片狀的曲面。如果各層互相黏連，在這一階段就會迷失方向；若將兩層當作一層來讀取，文字就會相互混淆。

第二步是為該曲面賦予亮度值。讀取曲面所經坐標的掃描值，並轉移繪製到曲面上。密度較高的墨水在此時會呈現為明亮的筆畫。

第三步是將皺縮的曲面展平。將立體壓成平面時，形狀會被拉伸或扭曲。這與將地球儀展平為紙地圖時兩極地區會被放大是同樣的問題。演算法會盡可能減少這種形變來展平曲面。

第四步是將分別展開的多個曲面拼接在一起。對齊曲面與曲面之間的邊界，融合成一張完整的大圖像[1]。

在四個步驟中，第一步最為繁難。一旦層位判定錯誤，後續的三個步驟都會產生偏差。錯誤展開的圖像甚至會出現原本不存在的字跡。這就是研究人員必須手動校對這一階段的原因。

經過這四個步驟後，黑色的圓柱體就化為一幅平面圖。研究團隊以此方式展開了五層皮革。展開後的面積估計約為94平方公分[1]，比筆記本的一頁稍小一些。

展開後的圖像依然存在問題。皮革各處厚薄不一，皺褶依然存在，掃描過程中產生的雜訊也夾雜其中。因此有的地方過亮，有的地方過暗。文字若與斑痕混在一起，便難以辨讀。

研究團隊透過演算法消除了這些斑痕。首先測量周圍像素的平均亮度，繪製出底材的亮度圖。這張圖僅記錄廣泛分布的明暗趨勢。接著將原始圖像除以這張底材圖。相除之後，大範圍的亮度差異消失了，只清晰留下集中在狹小範圍內的墨水信號。經過此過程的圖像中，希伯來文字終於顯現出來[1]。

延續至2026年的課題

虛擬展開此後也被應用於其他文物。2026年6月25日，維蘇威挑戰賽發布了消息。內容是他們已從頭到尾完整解讀了一卷赫庫蘭尼姆卷軸。文物編號為 PHerc. 1667。拍攝工作是在法國格勒諾勃的歐洲同步輻射設施進行的。展開出的希臘語文本長約1.4公尺，共二十二欄。經確認為西元前2世紀的倫理學著作[2]。

同步加速器是利用磁場使高速運行的電子發生偏轉的裝置。電子偏轉時會產生極強的X光。用這種光進行拍攝，即使是微小的密度差異也能顯現出來。這是實驗室拍攝儀器難以企及的清晰度。它是比在隱基底所使用的設備更高階的工具。

這項成果得益於拍攝方法的改進。在新的拍攝法中，碳墨可以直接在資料內部看見。它捕捉到了原本因密度差異太小而看不見的墨水。它透過拍攝技術，製造出了在隱基底是自然具備的亮度差異。據悉，剩餘的600多卷卷軸也可以採用相同的方法[3]。

課題依然存在。2026年7月10日公開的整理歸納指出了剩餘的六個問題。在層疊擠壓的區段，X光會發生散射，導致畫面模糊。尋找層次的運算會出現穿孔或將兩層黏合在一起的情況。體素資料中並未包含哪一層與哪一層相連的資訊。人工製作的訓練標記在顆粒尺度上不夠精確。從一卷卷軸上學得的模型在其他卷軸上效果不佳。墨水究竟位於層的哪個深度也不夠明確[4]。

解決方案也一併被提出。如果標註的標準答案不足，可以讓模型直接從資料本身學習規則。自主學習方式的訓練模型被視為候選方案。在尋找層次的運算中，也開始採用新的神經網路架構[4]。

新的線索也隨之出現。2026年3月29日公開的一項研究，著眼於墨水會改變表面形狀這一特點。研究團隊精確測量了已展開紙草卷的表面起伏，並將其傳授給機器學習演算法。結果僅憑表面起伏就能分辨出文字所在的位置。該研究還探討了要讀取這些起伏需要達到何種程度的拍攝解析度[5]。這是一項將在隱基底憑藉密度獲得的線索，從另一種物理性質中重新找回的研究。

參考資料

- [1] Seales, W. B., Parker, C. S., Segal, M., Tov, E., Shor, P., & Porath, Y. (2016). From damage to discovery via virtual unwrapping: Reading the scroll from En-Gedi. *Science Advances*, 2(9), e1601247. <https://doi.org/10.1126/sciadv.1601247>
- [2] Vesuvius Challenge. (2026, 6月 25日). An entire Herculaneum scroll has been read for the first time. [scrollprize.org](https://scrollprize.org/firstscroll). <https://scrollprize.org/firstscroll> (2026年9月3日檢索)
- [3] ASNT. (2026, 8月 13日). X-Ray Imaging and AI Recover First Complete Text from Sealed Ancient Scroll. *Materials Evaluation*. <https://www.asnt.org/me/26/8/x-ray-ai-recover-text-sealed-ancient-scroll> (2026年9月3日檢索)
- [4] Vesuvius Challenge. (2026, 7月 10日). Open Problems: Why Reading Every Herculaneum Scroll Is Still a Challenge. [scrollprize.org](https://scrollprize.org/2026_open_problems). https://scrollprize.org/2026_open_problems (2026年9月3日檢索)
- [5] Angelotti, G., Nicolardi, F., Henderson, P., & Seales, W. B. (2026, 3月 29日). Ink Detection from Surface Topography of the Herculaneum Papyri. *arXiv:2603.27698*. <https://arxiv.org/abs/2603.27698> (2026年9月3日檢索)

2.3. 揭開的真相

螢幕上浮現的《利未記》

電腦螢幕上浮現的字跡是希伯來文字。研究團隊展開了五層羊皮卷，從中顯現出兩欄文字。經文是舊約聖經的第三卷書《利未記》。解讀出的部分是《利未記》第1章與第2章[1]。

合成的圖像中包含了卷軸單幅頁面上的三十五行文字。其中有十八行保留了文字。其餘十七行則是學者們復原的。復原並不是由電腦憑空捏造文字。而是將殘留的筆畫與已知的經文進行比對，從而填補缺損部分的工作。每行所包含的字母與字間空格約為三十三個到三十四個[1]。

論文將殘存的行與復原的行分別加以說明，這一點至關重要。因為這樣讀者才能區分哪些是已確認的事實。如果把殘存的內容與填補的內容混在一起發表，後人就無法進行驗證。多虧於此，其他學者才能使用相同的資料重新進行審視與考證。

《利未記》第1章探討的是燔祭的規矩。燔祭是將祭物在祭壇上完全焚燒獻上的祭禮。第2章探討的是以穀物獻祭的規矩。在細麵中調油並加上乳香的方法就記載在此。就這樣，在被火燒毀的聖櫃中，出現了用火焚燒獻祭的規矩。

復原的文字中也包含了《利未記》的第一節經文[1]。即耶和華在會幕中召喚摩西並對他說話的句子。書卷的開篇得以留存，意味著這卷卷軸正是《利未記》整卷書的開端。

「兩欄」這個說法也有其含義。古代卷軸是將文字劃分在縱向修長的一欄一欄中書寫的。寫滿一欄後便移到下一欄。出現兩欄，意味著這樣的兩欄文字相連顯現了出來。

解讀出來的不僅僅是文字。連抄寫員下筆施力書寫時的筆畫粗細也留在了畫面上。這是由人工親手抄寫的手抄本的證據。

填補千年的空白

研究團隊也測定了卷軸的年代。透過放射性碳年代測定，得出了西元3世紀至4世紀的數值[1]。這種方法透過測量殘留在文物中的碳變化來判定其年代。這個數值填補了古代聖經抄本研究中長久以來的空白。

我們先來看看這個空白究竟是什麼。在昆蘭洞穴出土的死海古卷聖經抄本，年代介於西元前3世紀到西元1世紀之間。相對地，以完整形態留存至今的標準希伯來語聖經抄本，則是西元1000年左右的產物。這兩批文獻之間存在著將近一千年的空白。在此期間，能說明經文是如何傳抄繼承的抄本極為罕見。

空白越大，疑慮也隨之加深。因為無法確認中世紀的抄本是否如實抄錄了古代的經文。學者之間甚至曾有觀點認為，中世紀的抄寫員可能潤飾改動了經文。

隱基底卷軸正是從這段空白的正核心出土的。而解讀出來的經文，與中世紀希伯來語聖經的子音文本完全一致[1]。希伯來語只書寫子音字母，而不標記母音。這種僅記載子音的經文便被稱為「子音文本」。這個中世紀文本被稱為「馬索拉文本」。這是猶太教抄寫員訂定嚴格規則並恪守至今的標準文本。甚至連劃分段落的位置都完全相同[2]。這一結果表明，經文在漫長的歲月中被原封不動地傳抄了下來。

據悉，猶太教抄寫員在抄寫聖經時極其嚴格地遵守各項規則。逐字清點以進行核對的方法也是這些規則之一。隱基底卷軸成為證明這些規則確實得到切實執行的實物證據。一份抄本為長久以來的爭論增添了一項實證材料。

這份抄本還增添了一項紀錄。它是迄今為止在聖櫃中發現的《摩西五經》抄本中年代最久遠的[1]。《摩西五經》是指舊約聖經前五卷書的統稱。

已確立的事實與遺留的課題

已經查明的事實與仍在討論中的問題，必須分開來看。

已查明的事實如下。解讀出的經文是《利未記》第1章與第2章。碳定年測得的數值指向西元3世紀至4世紀。子音文本與中世紀的標準文本一致。墨水的密度高於皮革，推測其中混有金屬成分[1]。

遺留的問題同樣很明確。墨水的成分未能進行直接分析[1]。因為卷軸沒有外露的表面，無法取得樣本。也有以字體風格為依據，對年代提出不同看法的觀點。古文字學家阿達·雅爾德尼認為其年代在西元1世紀後半期或2世紀初。古文字學家是透過比對古代字體風格來斷定年代的學者。對此，德魯·隆埃克則反駁認為放射性碳定年更具可信度[2]。這兩種觀點目前尚未達成定論。

保存時間也是尚未解開的謎題。卷軸製成於西元3世紀至4世紀。猶太會堂約在西元600年左右被燒毀[1]。若是如此，這份抄本便在聖櫃中放置了兩百多年。該社群為何長期保留這份古老的抄本，目前仍不得而知。僅有推測認為他們可能十分珍視代代相傳的抄本。

解讀出的範圍也僅僅是卷軸的一小部分。成功展開的只有外側的五層[1]。內側深處的層疊被擠壓得更為嚴重。在受擠壓的區段無法剝離分層的問題，直到2026年依然存在。這正是研究赫庫蘭尼姆卷軸的學者們至今仍在全力攻克的課題[3]。

要解讀內側，必須在兩個方面同時取得進展。需要能將層次劃分得更加細緻的拍攝解析度。將沾黏在一起的層次剝離分解的運算法也必須進一步提升。一旦這兩個條件齊備，隱基底卷軸便能再次登上拍攝台。因為文物依然完好留存，這一切才成為可能。

輔助年代測定本身的方法也在不斷增加。2025年，一項利用機器學習掌握古希伯來語抄本字體風格以推估年代的模型發表了。該模型同時學習了放射性碳定年數值與字體資料。這種方法甚至能測量出人眼難以察覺的筆畫曲折特徵[4]。隨著此類工具的不斷積累，圍繞隱基底卷軸的年代爭論也能重新得到檢驗。

虛擬展開所開闢的道路並不僅限於聖經抄本。在摺疊密封的古代信件上，也採用了同樣的方法。脆弱易碎的古籍與文獻同樣引入了拍攝與運算技術。「無需觸碰即可閱讀」的原則，正逐漸成為處理文獻資料的基本準則。

電腦呈現的圖像本身並不是最終的答案。在2016年的研究中，電腦所做的工作僅止於展開各層並生成圖像。從該圖像中辨讀文字並賦予其意義的則是人類。因此，古文字學家必須逐字重新進行審視。這起事件留下的並不僅僅是一篇經文。它開創了一種處理觸碰即碎的文物的方法。其流程是：透過拍攝獲取數據，透過運算展開層次，最後由人工進行審查。這一流程也同樣應用於後續將提及的多件文物之中。

在隱基底解讀出的《利未記》也是歷經了此番審查，才成為學術界的正式紀錄。打破這長達四十五年沉默的，正是機器與人類的攜手合作。拍攝設備測量了密度，運算法展開了層次。隨後由人類辨讀出文字並賦予其意義。在下一章中，我們將繼續探討考察文字字形本身的故事。

參考資料

- [1] Seales, W. B., Parker, C. S., Segal, M., Tov, E., Shor, P., & Porath, Y. (2016). From damage to discovery via virtual unwrapping: Reading the scroll from En-Gedi. *Science Advances*, 2(9), e1601247. <https://doi.org/10.1126/sciadv.1601247>
- [2] En-Gedi Scroll. (2026, 9월 3일). Wikipedia. https://en.wikipedia.org/wiki/En-Gedi_Scroll (2026年9月3日檢索)
- [3] Vesuvius Challenge. (2026, 7월 10일). Open Problems: Why Reading Every Herculaneum Scroll Is Still a Challenge. scrollprize.org. https://scrollprize.org/2026_open_problems (2026年9月3日檢索)
- [4] Popović, M., Dhali, M. A., Schomaker, L., van der Plicht, J., Rasmussen, K. L., La Nasa, J., Degano, I., Colombini, M. P., & Tigchelaar, E. (2025). Dating ancient manuscripts using radiocarbon and AI-based writing style analysis. *PLOS ONE*, 20(6), e0323185. <https://doi.org/10.1371/journal.pone.0323185>

第3章. 隱藏在以賽亞卷軸中的兩道身影：抄寫員鑑定

3.1. 案件檔案

從洞穴出土的皮革卷軸

1947年，在死海西側昆蘭地區的洞穴中發現了古老的皮革卷軸。放羊的貝都因人在懸崖洞穴內找到了陶罐與卷軸。這些文物被稱為死海古卷。此後，多個洞穴陸續展開了發掘工作。至今收集到的死海古卷已整理為約930件抄本。用手抄寫的書籍稱為抄本。若以單頁計算，皮革與紙草殘片超過數千張[1]。記錄的年代被認為介於西元前3世紀至西元1世紀之間[2]。

死海是一座高鹽度的湖泊，周圍是乾旱的沙漠。這裡幾乎不下雨，空氣乾燥。洞穴內的溫度也沒有太大起伏。這正是皮革與紙草得以不腐爛而留存下來的原因。能留下如此龐大規模古代文獻群的地方實屬罕見。

在這些文物中，有一件卷軸吸引了學者們的目光。那就是以賽亞大卷軸。學界使用的符號為1QIsaa。意思是出自昆蘭第一洞穴的以賽亞書第一號抄本。這是一件西元前2世紀以希伯來語寫在獸皮上的抄本。長度約7.3公尺。幾乎完整收錄了舊約聖經以賽亞書的內容。內文分為54個垂直長條欄位[2]。在學術術語中稱這些欄位為欄（column）。這是死海洞穴出土的聖經抄本中保存狀態最好的一件。

卷軸是由多張皮革用線縫合拼接而成。在一張皮革上寫下幾欄後，再接著寫下一張。寫完的卷軸捲起來存放。閱讀時則一邊展開一側，一邊捲起另一側來閱讀。「卷軸」之名便由此而來。

這件抄本的價值來自於它的年代。死海古卷中包含了現存最古老的希伯來語聖經抄本[3]。以賽亞大卷軸也是其中之一。這讓人們擁有了可以對照聖經經文在漫長歲月中如何傳承的珍貴資料。然而，關於製作這件卷軸的人，卻始終無人知曉。

圍繞著究竟由誰抄寫而分歧的學界

解讀卷軸與查明是誰寫的，是兩件截然不同的事。雖然文字都能讀懂，但這件抄本上卻沒有抄寫者的名字。負責抄寫文字的人被稱為抄寫員。這件卷軸上既沒有抄寫員的簽名，也沒有註明日期[2]。因此，學者們只能單憑字體形態來推斷人數。

面對這個問題，學界站在了兩個分岔路口。一方認為是由一名抄寫員從頭到尾獨自抄寫。其依據是整件卷軸的字形十分勻稱。另一方則認為是由兩人以上分工抄寫。其依據是過了中間部分後，筆畫形態開始略有變化。

兩種主張都面臨著一個難題。即使是同一個人寫的字，開頭和結尾也不會完全相同。手疲勞時筆畫會變鬆散，換了筆後筆畫粗細也會改變。若不知道這種變化的幅度有多大，就無法判斷是否換了人。兩種主張都仰賴肉眼觀察的直觀印象。當時缺乏能衡量哪一方正確的客觀標準。

學界認為昆蘭曾存在一個製作並守護卷軸的社群。人們將卷軸放入陶罐並藏匿在洞穴中。這似乎是為了在戰火逼近時保護這些文獻。將它們藏起來的人也沒有留下名字。尋找製作卷軸之手，也是一條走近那些人的途徑。

這場爭論並不僅止於字跡的討論。如果抄本是由多人分工製作，便意味著存在著抄寫員共同工作的組織體系。這關係到查明聖經經文究竟經過哪些人的手才流傳至今[2]。這正是死海古卷研究者長期執著於這個問題的原因。

肉眼鑑定的局限

研究古代字跡的學問稱為古文字學。傳統方式是學者拿著放大鏡仔細觀察字形。測量字體的高度、傾斜度與筆畫粗細，並描摹記錄下來。有時也會將不同抄本的字跡按年代順序排列，以掌握其演變趨勢。長期以來，這種方法是唯一的鑑定工具。

這種方法存在著弱點。它在很大程度上依賴學者的肉眼與經驗。即使看著同一個字，不同學者也經常給出不同的評價[2]。人的肉眼無法精確測量微小的數值差異。壓下蘆葦筆的力道強弱或墨水暈開的厚度，很難單憑肉眼來衡量。記憶力也有限度。要將相距數公尺遠的兩個字在腦海中並列比較絕非易事。結果只能不斷重複意見的對立。

另一個障礙是那個時代的教育方式。立志成為抄寫員的人，都必須接受完全臨摹老師字跡的訓練。他們被教育不要展現個人風格。結果，出自同一所學校的抄寫員，其筆跡外觀幾乎一模一樣。對人類學者而言，這無異於一場濃霧。根本無法分辨卷軸前後的微小差異，究竟是同一個人的狀態變化，還是換了人抄寫[2]。

一旦鑑定出現偏差，隨之而來的判斷也會跟著失誤。如果抄寫員只有一人，代表一份抄本凝聚了一個人的時間；如果是兩人，則意味著存在分工合作的組織。製作抄本之機構的規模與性質也將隨之不同。僅憑一項筆跡鑑定，就牽動著諸多結論。

想想在教室裡照著同一個字帖練習聽寫的孩子們，情況也十分相似。大家努力寫得端正的作業本，表面上很難區分。即便如此，轉折筆畫的習慣在每個孩子身上依然有些許不同。問題在於，當時缺乏一把能夠測量這種差異的尺。

情況在抄本照片數位化公開後發生了轉變。實體卷軸對光線和濕氣十分敏感，無法頻繁展開。有了照片，便能在不觸碰文物的狀況下隨意查看無數次。以色列文物局與Google合作建立了死海古卷數位資料庫。他們利用名為MegaVision的設備，照射多種波長的光線來拍攝文物。甚至捕捉到了人眼看不見的近紅外光，讓模糊的文字重現。其目標是達到與親眼端詳實物相近的畫質[1]。

隨著電腦可讀取的高解析度照片不斷累積，引進新工具的道路被打開了。照片是一組數字的集合。構成畫面的微小點被稱為像素。每個像素都附有亮度值。這意味著筆畫的粗細與角度可以用數字來測量。將肉眼鑑定轉化為計算鑑定的準備工作已經就緒。自此，AI偵探正式登場。

參考資料

- [1] Israel Antiquities Authority. (n.d.). The Digital Library. The Leon Levy Dead Sea Scrolls Digital Library. <https://www.deadseascrolls.org.il/about-the-project/the-digital-library> (2026年9月3日檢索)
- [2] Popović, M., Dhali, M. A., & Schomaker, L. (2021). Artificial intelligence based writer identification generates new evidence for the unknown scribes of the Dead Sea Scrolls exemplified by the Great Isaiah Scroll (1QIsaa). PLOS ONE, 16(4), e0249769. <https://doi.org/10.1371/journal.pone.0249769>
- [3] University of Groningen. (2021, 4月 21日). Cracking the code of the Dead Sea Scrolls. rug.nl. <https://www.rug.nl/news/2021/04/cracking-the-code-of-the-dead-sea-scrolls?lang=en> (2026年9月3日檢索)

3.2. 科學鑑識團隊的追蹤：留在字體裡的生物特徵指紋

調查團隊正式成軍

荷蘭格羅寧根大學研究團隊接下了這起案件。古文字學家姆拉登·波波維奇（Mladen Popović）率領團隊。AI研究員蘭伯特·肖馬克（Lambert Schomaker）與馬魯夫·達利（Maruf Dhali）一同參與[1]。他們於2017年率先發表了利用電腦測量死海古卷筆跡的方法[2]。進一步改良該方法並應用於以賽亞大卷軸的研究成果，於2021年發表[1, 5]。

研究團隊設定的目標非常明確：將人類的主觀印象評估轉化為數字。首先測量同一個人書寫時所產生的波動幅度；接著比較卷軸前後的差異是否大於該波動幅度。如果差異更大，就意味著是由不同人所寫。

2017年建立的特徵綜算法曾用於比較不同的抄本[2]。研究團隊將該綜算法集中重新應用於單一的以賽亞大卷軸上。在同一件抄本內部找出更換抄寫員的位置，遠比較不同抄本更加困難。因為字跡之間的差異微小得多。

這種方法有一個優點：只要記錄下計算程序，任何人都能使用相同的照片得出相同的答案。無須仰賴學者的個人眼力。如果結果有誤，也能回溯追查是在哪一個步驟出錯。即便對鑑定結果再次產生爭議，爭論的依據也留存為客觀的數字。

清除污漬的第一步

第一項工作是從照片中僅保留墨水痕跡。這項處理被稱為二值化。就是將所有像素劃分為墨水或背景其中之一的處理過程，亦即將照片分為黑色字體與明亮背景的工作。

這一步驟至關重要。卷軸在洞穴中歷經了約兩千年。皮革表面出現了污漬、黴菌斑痕與細微裂紋。在電腦看來，這些痕跡與墨水筆畫都屬於暗色像素。若保留背景，機器學習到的將是皮革紋理而非字跡。如此一來，機器區分的將是皮革種類而非抄寫員，導致鑑定失準[1]。

研究團隊專門打造了負責此項工作的神經網路，命名為BiNet。神經網路是模仿人腦神經元連結所構建的計算架構。BiNet透過觀察人工標註墨水與背景的樣本，學習判斷標準。完成學習的BiNet便能從新照片中剔除污漬，僅留下筆畫。

在此必須遵守一項規則：不能更動原始筆畫的粗細與邊緣形狀。因為抄寫員的書寫習慣恰恰體現在那些邊緣之中。若為了去除雜質而將筆畫修飾得過於平滑，證據就會隨之消失。BiNet的設計宗旨就是盡可能完整保留墨水痕跡[1]。這項技術於2019年先以公開預印本形式發布[3]。重新整理其架構與性能的論文，則刊登於2026年的國際文件識別學術期刊上[4]。

衡量整體筆跡習慣的尺規

從清理乾淨的字跡中擷取兩種特徵。第一種是紋理特徵，第二種是形狀特徵。

紋理特徵被稱為Hinge。Hinge是指門上的鉸鏈。在筆畫轉折處測量兩個方向的角度，並統計該角度對出現的頻率製成表格。這張表格以數字記錄了書寫時手腕旋轉的習慣。無須逐字辨識文字即可進行計算[1]。

形狀特徵則被稱為Fraglet。Fraglet是指將筆畫切成短小的片段。將切下的筆畫片段彙整成清單。研究團隊利用從多件抄本中擷取的片段建立了標準清單。當輸入新的字跡時，便會計算該片段與清單中哪一個項目相似。人們在學習文字時掌握的筆畫基本形態便會在此顯現[1]。

這兩種特徵有一個共同點：即使不知道文字的涵義也能進行計算。即使電腦讀不懂希伯來語，也能測量筆畫的角度與彎曲形狀。即便夾雜著受損的文字，也能從殘存的筆畫中提取數值。這在處理古代文獻時是一項極為實用的特性。

兩種特徵著眼於不同的層面。Hinge觀察的是手部的動作，而Fraglet觀察的則是習得的字體形態。研究團隊分別使用了這兩項特徵，同時也將兩者合併使用。這是為了確認不同的衡量標準是否能得出相同的答案[1]。

殘留在單一阿列夫字母上的痕跡

前述的兩種特徵是綜合衡量整體字跡。這一次，輪到將單一字母大幅放大來仔細觀察了。

研究團隊單獨挑選了一個希伯來字母進行精密檢驗。那就是阿列夫。阿列夫是希伯來語的第一個字母，也是卷軸中最常出現的字母。它的形狀是由一條對角線筆畫加上兩條短筆畫組成。

選擇阿列夫是有原因的。樣本越多，統計就越可靠。只有頻繁出現的字母，才能在前後區域平均收集足夠的數量。整件卷軸中共有5,011個阿列夫。研究團隊利用電腦從中篩選出758個用於分析[1]。僅挑選了筆畫未中斷且保存完整的字母。

挑選出的阿列夫被分為兩組：出自卷軸前半部的與出自後半部的。將每組的影像疊加在同一位置，產生平均形狀。這就像將多張半透明紙疊在一起從上方俯視一樣。經常重疊的筆畫會變深，而因人而異的筆畫則會變淡。將這兩張平均影像相減，剩下的就只有差異。透過顏色便能直觀看出這些差異集中在哪個筆畫上[1]。

分裂成兩群的點陣

轉化為數字的特徵是一份包含數百個數值的長清單。人類無法單憑肉眼比對這種清單。因此，研究團隊將清單縮減為兩個數值，並在平面上繪製成點。數值相近的字跡會聚集在一起，差異較大的字跡則會拉開距離。這是透過反覆計算並微調點的位置所繪製出的圖表。

在54欄之中，照片完好的欄位全都呈現在這個平面上。第16欄與第46欄因缺乏可用的照片而被排除。當點繪製完成後，呈現出了兩個群落。前半部的欄位形成一個群落，後半部的欄位則形成了另一個群落[1]。這兩個群落之間存在著肉眼可見的明顯間隙。

研究團隊驗證了這種分歧是否純屬偶然。他們使用了卡方檢定與t檢定等統計方法。卡方檢定是判斷差異是否出於偶然的計算方式。t檢定則是比較兩組平均數差異的計算方式。兩種方法都是計算偶然出現這種樣貌的機率。若該機率極小，便不視為偶然。光憑同一人筆跡中產生的波動，並不會出現如此巨大的差異[1]。

多種評估標準得出相同答案這一點也很重要。無論是用鉸鏈（Hinge）測量，還是用筆畫片段（Fraglet）測量，都在同一個位置劃出了分界。在Aleph的平均圖像中，也在同一處出現了分歧。界線落在了第27欄與第28欄之間[1]。若不同方法皆指向同一處，該界線就不是因照片瑕疵而產生的。而是機器替人眼畫出了曾被忽略的界線。

研究團隊也探討了相反的可能性。這可能是因照片亮度或前後皮革狀況不同所導致的結果。然而，在透過二值化去除背景之後，界線依然存在[1]。這意味著原因不在於皮革，而在於字跡本身。

參考資料

- [1] Popović, M., Dhali, M. A., & Schomaker, L. (2021). Artificial intelligence based writer identification generates new evidence for the unknown scribes of the Dead Sea Scrolls exemplified by the Great Isaiah Scroll (1QIsaa). PLOS ONE, 16(4), e0249769. <https://doi.org/10.1371/journal.pone.0249769>
- [2] Dhali, M. A., He, S., Popović, M., Tigchelaar, E., & Schomaker, L. (2017). A digital palaeographic approach towards writer identification in the Dead Sea Scrolls. Proceedings of the 6th International Conference on Pattern Recognition Applications and Methods, 693-702. <https://doi.org/10.5220/0006249706930702>
- [3] Dhali, M. A., de Wit, J. W., & Schomaker, L. (2019). BiNet: Degraded-manuscript binarization in diverse document textures and layouts using deep encoder-decoder networks. arXiv:1911.07930. <https://arxiv.org/abs/1911.07930> (2026年9月3日檢索)
- [4] Dhali, M. A., de Wit, J. W., & Schomaker, L. (2026). BiNet: A deep encoder-decoder network for binarizing degraded ancient manuscripts. In Document Analysis and Recognition - ICDAR 2025 Workshops (Lecture Notes in Computer Science, pp. 166-193). Springer Nature Switzerland. https://doi.org/10.1007/978-3-032-09371-4_11
- [5] University of Groningen. (2021, 4月 21日). Cracking the code of the Dead Sea Scrolls. rug.nl. <https://www.rug.nl/news/2021/04/cracking-the-code-of-the-dead-sea-scrolls?lang=en> (2026年9月3日檢索)

3.3. 揭曉的真相

在第27欄換手

分析結果支持了多位抄寫員的說法。《以賽亞大卷軸》是由兩人分工抄寫的。第一位抄寫員抄寫了第1欄到第27欄。第二位抄寫員則接著抄寫了第28欄到第54欄[1]。

這條界線也與其他證據相互吻合。第27欄下方留出了相當於三行高度的空白。第28欄則是從新拼接的皮革張開始。該處也是《以賽亞書》新章節開始的地方[1]。這是交接工作極為合適的結尾。電腦找到的統計界線與卷軸的物理界線恰好落在同一個位置。

在第27欄之後更換了抄寫員的推測，在此之前便已存在。這是根據皮革拼接處與正文劃分處所提出的見解。只不過當時未能形成學界共識。AI分析為這項推測賦予了具體數據[2]。

這項發現意義重大。一份聖經抄本是由多人接力製作的依據，終於以數據的形式呈現[1]。兩位抄寫員既未留下姓名，也未留下簽名。留下來的，只有手部動作的習慣。兩千年後的機器解讀出了這種習慣。波波維奇表示，雖然終究無法得知他們的名字，但透過字跡，彷彿正在與他們握手[2]。

在同一所學校學習的兩人

兩位抄寫員的字跡依然彼此相似。即使經由機器測量，差異也相當微小。研究團隊將這種相似性本身解讀為另一項證據。這意味著兩人的書寫風格非常相近[1]。研究團隊認為，兩位抄寫員可能在同一地點受訓，或是師出同門[2]。

由此可以想見古代抄寫員學校的面貌。受訓學員以筆畫為單位熟習公認的字體。唯有壓抑個人個性，才能成為優秀的抄寫員。多虧了這種紀律，單一卷軸內的字跡才能保持均勻一致。

即便如此，身體的習慣依然無法抹去。握蘆葦筆的角度因人而異。運筆的速度與下筆的力道也各不相同。這些差異會留在筆畫轉折的角度與筆畫邊緣的形狀上。人眼無法綜合這些微小差異來做出判斷。電腦則能彙整數千個數值，計算出平均值與離散程度。被紀律掩蓋的個人痕跡，就這樣顯露了出來[1]。

這項結果顯示，製作抄本是一項協同合作的工作。要完成一卷長卷軸需要耗費數月。沒有理由非得由一個人從頭到尾獨立承擔。分工合作並交接工作的可能性由此浮現[1]。聖經經文正是經由這樣一雙雙手的傳遞，流傳至今。

機器目前尚無法做到的事

這項成果並不代表AI是萬能的。同一研究團隊開發的下一項工具展現了其局限性。

研究團隊於2025年推出了一款名為Enoch的年代推估模型。Enoch是一款透過字跡形貌來推斷抄本撰寫年代的模型。放射性碳定年法是一種利用古物中殘留的碳來測量年代的方法。Enoch學習了以該方法測得年代的24個樣本。研究團隊也同時公布了將其應用於135件未知年代抄本的結果。預測誤差平均介於27.9年至30.7年之間。古文字學家在檢視這項預測後，評估有79%的結果是合理的[3]。

其餘的21%則引發了爭議。古文字學家克里斯多福·羅爾斯頓（Christopher Rollston）舉出一個案例進行反駁。Enoch將包含《但以理書》的4Q114抄本判定為公元前230年至前160年之間。羅爾斯頓認為，《但以理書》第7章至第12章指向了某一歷史事件。即公元前167年安條克四世玷污耶路撒冷聖殿的事件。該經文不可能寫於事件發生之前。然而Enoch給出的時間範圍大多早於前167年。羅爾斯頓指出Enoch推估的年代過早。他表示，Enoch應當只是工具箱中眾多工具的其中之一[4]。

這場爭論帶來了啟示。機器只看字形。無法解讀文字記載的歷史事件或文句含義。因此，機器的答案不是定論，而是候選方案。在縮小候選範圍後，必須由歷史學家比對語境進行判定。決定測量哪些字母的人也是古文字學家。挑選Aleph的判斷，源自長期觀察古代字跡的經驗。

訓練資料的數量也是一大障礙。Enoch僅透過24個已知年代的樣本進行學習[3]。古代文物很難增加樣本數量。因為若要測量年代，就必須取下一小部分文物進行燃燒。依據少數樣本訓練的模型，一旦超出學習範圍就會產生動搖。抄寫員鑑定的情況也十分相似。頻繁出現的字母樣本充足，但罕見的字母卻只有寥寥數個。在樣本匱乏之處，機器的判斷力便會減弱。

因此，研究團隊並未將AI視為唯一的裁判。皮革拼接處、正文劃分處，以及學者長期累積的見解都會一併納入考量。唯有多項證據指向同一處時，結論才能成立。《以賽亞大卷軸》得出解答也是遵循這種方式。機器只是以數據回答了人類提出的問題。

下一個調查目標

對《以賽亞大卷軸》的研究引申出了一項更宏大的計畫。歐洲研究委員會於2026年6月23日決定授予波波維奇高級研究資助（Advanced Grant）。研究經費為250萬歐元，研究期限為5年[5]。新計畫的名稱為「追蹤抄寫員與卷軸」（Tracing Scribes and Scrolls）[6]。意指追尋抄寫員與卷軸的足跡。

以色列文物局於2026年6月30日宣布參與該計畫。比薩大學、拿坡里腓特烈二世大學、南丹麥大學與魯汶大學也共同參與。柏林與杜林的埃及博物館亦名列其中[6]。

這一次的問題從「由誰抄寫」轉向了「在何處製作」。研究團隊將調查約250件死海古卷樣本。測量皮革、紙草與墨水的化學成分。並將其成分與古埃及紙草的成分進行比對。若得知材料來源，便能縮小抄本製作地點的範圍。AI負責整理這些化學資料，並找出字跡中隱藏的規律[7]。波波維奇表示，結合實驗室分析與古文字研究，將能探討過去無法觸及的問題[6]。

這種方法不只適用於《以賽亞大卷軸》。死海古卷中仍有許多抄寫者未明的抄本。將出自同一手筆的抄本歸納在一起，就能建立單一抄寫員的作品目錄。隨著這類目錄的累積，便能描繪出抄寫員彼此之間如何相互關聯[1]。找出兩位抄寫員的這項研究，正是這條探索之路的起點。

這項研究留下的另一項成果是方法論本身。去除照片瑕疵、將筆畫轉化為數據，並以統計進行檢定的流程得以確立。相同的流程亦可應用於其他語言的古代文獻。在本書後續章節中將介紹的泥板與木簡研究，也循著相似的路徑前進。

源自《以賽亞大卷軸》的調查尚未結束。找尋到兩人身影的工具，如今啟程探尋卷軸的製作地點。

研讀古代文獻並不止於辨識文字。還必須追問是誰在何處寫下了這些字。《以賽亞大卷軸》是以數據回答「究竟有幾人抄寫」這一問題的首例。卷軸在兩千年間始終保持沉默。如今，人類從字跡的顫動中讀出了答案。

參考資料

- [1] Popović, M., Dhali, M. A., & Schomaker, L. (2021). Artificial intelligence based writer identification generates new evidence for the unknown scribes of the Dead Sea Scrolls exemplified by the Great Isaiah Scroll (1QIsaa). PLOS ONE, 16(4), e0249769. <https://doi.org/10.1371/journal.pone.0249769>
- [2] University of Groningen. (2021, 4월 21일). Cracking the code of the Dead Sea Scrolls. rug.nl. <https://www.rug.nl/news/2021/04/cracking-the-code-of-the-dead-sea-scrolls?lang=en> (2026年9月3日檢索)
- [3] Popović, M., Dhali, M. A., Schomaker, L., van der Plicht, J., Rasmussen, K. L., La Nasa, J., Degano, I., Colombini, M. P., & Tigchelaar, E. (2025). Dating ancient manuscripts using radiocarbon and AI-based writing style analysis. PLOS ONE, 20(6), e0323185. <https://doi.org/10.1371/journal.pone.0323185>
- [4] Steinmeyer, N. (2025, 6월 20일). Can AI date the Dead Sea Scrolls? Bible History Daily, Biblical Archaeology Society. <https://www.biblicalarchaeology.org/daily/biblical-artifacts/dead-sea-scrolls/can-ai-date-the-dead-sea-scrolls/> (2026年9月3日檢索)
- [5] University of Groningen. (2026, 6월 23일). Mladen Popović awarded ERC Advanced Grant to trace the origins of the Dead Sea Scrolls. rug.nl. <https://www.rug.nl/about-ug/latest-news/news/archief2026/nieuwsberichten/mladen-popovi-awarded-erc-advanced-grant-to-trace-the-origins-of-the-dead-sea-scrolls?lang=en> (2026年9月3日檢索)
- [6] The Jerusalem Post. (2026, 6월 30일). New AI-powered research project aims to uncover the origins of the Dead Sea Scrolls. jpost.com. <https://www.jpost.com/archaeology/article-900938> (2026年9月3日檢索)
- [7] Archaeology News Online Magazine. (2026, 6월). Where the Dead Sea Scrolls come from: AI and chemical analysis trace origins. archaeologymag.com. <https://archaeologymag.com/2026/06/dead-sea-scrolls-origins-studied-with-ai-and-chemical-analysis/> (2026年9月3日檢索)

第4章. 拼合數萬片泥板與羊皮紙拼圖的數位黏著劑

4.1. 案件檔案

尋找山羊時發現的陶罐

1947年春天，死海西側的懸崖上發生了一件事。幾名貝都因少年正在尋找逃跑的山羊。少年們在岩石之間看見了狹窄的洞穴入口。洞穴內立著幾個高大的陶罐。陶罐裡裝著在皮革上書寫的卷軸[1]。

這些卷軸就是死海古卷。死海古卷是大約西元前3世紀至西元1世紀之間寫成的猶太文獻。隨後九年間，附近又陸續發現了十處洞穴。出土卷軸的洞穴總共有十一處。這裡出土的文獻超過900篇。這些文獻是以希伯來語、亞蘭語和希臘語寫成的[1]。大部分寫在將綿羊或山羊皮削薄製成的羊皮紙上。也有寫在名為紙草的植物莖部紙張上的抄本[7]。

卷軸中包含了希伯來聖經抄本中年份極為久遠的文獻。同時也發現了記錄社群規範與曆法的文章[1]。乾燥荒野中的洞穴將這些文字守護了2,000年。

洞穴附近有一處名為庫姆蘭的古代聚落遺址。多數學者認為是曾居住在此地的猶太社群留下了這些卷軸。他們抄寫聖經，並記錄下自己的規範。據推測，他們是在羅馬軍隊逼近時，將卷軸藏進了洞穴。

第4洞穴的碎片堆

各洞穴中文物的保存狀態大不相同。第1洞穴的卷軸在陶罐中保存得相當完好。第4洞穴的情況則截然不同。該洞穴的卷軸並未裝在陶罐中，而是直接堆放在地面上。光是這一個地方，就出土了約500篇卷軸、多達數千件的碎片[1]。

洞穴內的濕氣逐漸侵蝕了皮革。蟲蟻與野獸啃食了皮革。歷經漫長歲月，卷軸碎裂成了細屑。最後留下的，只是一堆指甲般大小的碎片。

以色列文物局所保存的碎片超過25,000件[2]。大部分碎片上並未標明究竟是從哪一個卷軸脫落的。單一碎片上僅殘留兩三個字的情況也屢見不鮮。學者們將這些碎片陳列在大型玻璃板上，全憑肉眼進行比對拼合。

拼合的線索僅有寥寥數種。學者檢視字形是否相似，觀察皮革顏色是否相近，並比對斷裂邊緣是否契合。用這種方法，一次只能比對兩三件碎片。當碎片數量增加到以萬為單位時，需要比對的配對組合數量便會失控暴增。

以手工拼圖耗時漫長。人類的肉眼無法同時比對數萬塊碎片的邊緣。再加上碎片的邊緣早已磨損且蜷曲變形。即便學者們耗費數十年心血，依然留下了許多未能找到原本位置的碎片。

當初只有少數學者擁有閱覽這些文獻的特權，這也是一個問題。外界的研究人員甚至連照片都難以見到[1]。1991年9月，美國亨廷頓圖書館打破了這一門檻。該圖書館無條件公開了所藏的底片膠卷[3]。隨著照片的公開，利用電腦拼合碎片的嘗試也隨之展開。

文物分散在多個國家也是一大阻礙。在發現初期，部分卷軸經由伯利恆的古董商被轉手賣出[1]。資助發掘費用的國家機構也分到了碎片。因此，源自同一卷軸的碎片被分散存放在不同的城市。若想確認兩塊碎片能否接合，研究人員甚至必須親自搭乘飛機前往。

市場上出現的偽造碎片

隨著死海古卷碎片的價格上漲，出現了另一個問題。大約從2002年開始，古董市場上出現了號稱死海古卷碎片的物品。市場上出現了大約70件寫有聖經章節的殘片[4]。學者們將這些物品稱為「2002年後碎片」。這些物品缺乏記錄其何時出土於何處洞穴的文獻記錄[6]。

美國華盛頓的聖經博物館也收購並展出了16件這類碎片。多數學者懷疑其字形顯得生硬不自然。該博物館委託藝術品防偽調查公司進行精密鑑定。主導調查的是調查員柯萊特·羅爾[4]。

2020年3月，調查結果出爐。16件全數為現代製造的贗品[4][5]。調查團隊使用3D顯微鏡觀察其表面。材質並非真正死海古卷那種薄羊皮紙。至少有15件是厚重且紋理粗糙的皮革。墨水甚至沉積在早已龜裂的縫隙之中[4]。這正是皮革裂開之後才書寫文字的明證。

這起事件震撼了整個聖經考古學界。喬治華盛頓大學的克里斯多福·羅斯頓教授表示，如今市場上仍不斷出現贗品。他透露自己經常遇到堅信自己藏品是真品的人，但幾乎全都是假貨。留有發掘記錄的文物與在市場上購得的文物之間的差異，在此暴露無遺[6]。

拼合破碎殘片與篩除贗品，本質如出一轍。兩者皆須解讀人眼無法捕捉的細微物理痕跡。AI介入了這項工作。2026年6月，歐洲研究委員會撥發了新的研究經費。他們為一項探討死海古卷出處的5年期研究資助了250萬歐元。研究名稱為「追蹤抄寫員與卷軸」。該研究結合了材料成分分析、AI以及古代文字研究[7]。

本章將追蹤兩起事件。其一是重新拼接散落的碎片。其二是篩除贗品。這兩項工作面對相同的資料，提出了不同的疑問。前者的疑問是這兩者原本是否為一體；後者的疑問則是這件物品原本是否為古物。

參考資料

- [1] Israel Antiquities Authority. (n.d.). Discovery and publication. The Leon Levy Dead Sea Scrolls Digital Library. <https://www.deadseascrolls.org.il/learn-about-the-scrolls/discovery-and-publication> (2026年9月3日檢索)
- [2] Israel Antiquities Authority. (n.d.). The digital library. The Leon Levy Dead Sea Scrolls Digital Library. <https://www.deadseascrolls.org.il/about-the-project/the-digital-library> (2026年9月3日檢索)
- [3] Deseret News. (1991, 9月 22日). Research library will give scholars access to Dead Sea Scroll photos. Deseret News. <https://www.deseret.com/1991/9/22/18942605/research-library-will-give-scholars-access-to-dead-sea-scroll-photos/> (2026年9月3日檢索)
- [4] Greshko, M. (2020, 3月 13日). "Dead Sea Scrolls" at the Museum of the Bible are all forgeries. National Geographic. <https://www.nationalgeographic.com/history/article/museum-of-the-bible-dead-sea-scrolls-forgeries> (2026年9月3日檢索)
- [5] Solly, M. (2020, 3月 16日). All of the Museum of the Bible's Dead Sea Scrolls are fake, report finds. Smithsonian Magazine. <https://www.smithsonianmag.com/smart-news/all-museum-bibles-dead-sea-scrolls-are-fake-report-finds-180974425/> (2026年9月3日檢索)
- [6] Govier, G. (2026, 4月 8日). New Dead Sea Scrolls exhibit is the real deal. Christianity Today. <https://www.christianitytoday.com/2026/04/new-dead-sea-scrolls-exhibit-is-the-real-deal/> (2026年9月3日檢索)
- [7] The Jerusalem Post. (2026, 6月 30日). New AI-powered research project aims to uncover the origins of the Dead Sea Scrolls. The Jerusalem Post. <https://www.jpost.com/archaeology/article-900938> (2026年9月3日檢索)

4.2. 科學鑑識團隊的追蹤 1：幾何邊緣親和度與物理指紋

將撕裂的邊緣轉換為數字

拼合破碎殘片的第一個線索是邊緣形狀。用紙張徒手撕開時，會產生如鋸齒般參差不齊的線條。相鄰兩片殘片的線條彼此契合。人類憑目測比對這些線條，電腦則將這些線條轉換為數字。

電腦沿著邊緣測量彎曲程度。將彎曲程度依序排列，便形成波浪狀的數列。當兩塊碎片的數列相互翻轉重疊且吻合良好時，便成為拼接候選對象。這種方法被稱為邊緣親和度計算，用我們的話來說就是邊緣契合分數。

這個分數用我們的話來說，可以稱為邊緣契合分數。這是用來表示兩塊碎片的邊緣有多吻合的數值。

僅觀察邊緣的方法有著明顯的局限。古代碎片的邊緣早已磨損。因乾燥蜷曲而失去原有形狀的部分也不在少數。因此，研究人員同時著眼於邊緣之外的線索。

材料本身留下的指紋

光靠邊緣是不夠的。碎片的邊緣往往早已磨損。因此研究人員轉而觀察材料本身。

紙草是將植物莖部剖成細條，交錯縱橫疊放製成。因此表面會留下如紋理般的莖部紋路。原本同屬一張的兩塊碎片，其紋理是彼此連貫的[1]。

以色列海法大學的羅伊·阿維特布爾研究團隊解讀了這種紋理。研究團隊將碎片邊緣附近裁切成小方塊。神經網路藉由觀察兩個方塊來評估紋理是否連續並打出分數。再整合這些分數來判定整個碎片的配對[1]。

測試對象是一組保存狀態不佳的死海古卷碎片。由於這組碎片並無已知答案，因此也出現了許多錯誤的候選組合。即便如此，仍成功排除了98%不可能配對的組合[1]。使人工審查的工作量縮減至五十分之一。

這個結果的意義十分明確。電腦並非確定了配對，而是替人清除了無需查看的組合。在碎片拼合工作中，AI所扮演的角色大多是這類篩選過濾工作。

法國波爾多的安東萬·皮隆研究團隊在此採用了孿生神經網路。孿生神經網路將兩張影像並置輸入，判定兩者是否來自同一出處。研究團隊將此模型命名為 PapieSNet。該模型以約500件紙草碎片進行學習。不看文字，僅依據紙張紋理來挑選配對。在實際資料中，配對成功率達79%[2]。

羊皮紙上也有類似的指紋。皮革上會殘留生長毛髮的毛孔痕跡。取自同隻動物皮革的碎片，這些痕跡的排列會相互連貫。蛀蟲蛀穿卷軸留下的蟲洞也是線索。將捲起的卷軸攤平後，蟲洞會排列成規律的重複圖樣。卷軸內層周長較短，外層周長較長。因此，蟲洞之間的間距越往外層會逐漸變寬。只要測量這個間距，便能推估碎片原本處於第幾層。

要處理這類物理指紋，必須先整理照片。以色列文物局的死海古卷照片中，同時拍攝了比例尺、色卡與號碼牌。背景為黑色，容易與墨水混淆。布朗森·布朗德沃斯特研究團隊開發了一套分為四個步驟的預處理方法，先行清除這些干擾物。該團隊同時也公開了附有標準答案的資料集[3]。

公開附有標準答案的資料集意義深遠。其他研究人員得以利用相同的資料測試自己的方法。也能以數據量化比對兩種方法孰優孰劣。古文書研究正以此種方式轉變為彼此驗證的學術領域。

擴展至泥板、甲骨與紙張的技術

相同的思維也延伸應用至其他文物。德國慕尼黑路德維希-馬克西米利安大學的恩里克·希門尼斯教授領導著電子巴比倫圖書館（eBL）計畫。該計畫所使用的碎片拼合工具名為 Fragmentarium[4]。

美索不達米亞人使用蘆葦筆在濕潤的泥土上壓印楔形文字。這些泥板極易破碎。研究團隊在大英博物館與伊拉克國家博物館拍攝了泥板照片。自2018年起，他們將大約22,000件碎片納入資料處理[5]。電腦相互比對碎片上殘留文字所轉寫出的字串。此時借鑑了生物學中比對基因序列的方法。這是因為古代抄寫員在各個地區使用的拼寫方式不盡相同[4]。

成果直接反映在數字上。該計畫在短短5年間找到了約1,500件新的拼接組合。而在過去的150年間，學者們找到的拼接組合總共約為5,000件[4]。

泥板上有一項羊皮紙所沒有的額外線索：斷裂面的立體形狀。泥板大多經日曬乾燥，僅有少數經過火烤。在硬化過程中，其內部堅固程度並不均勻。因此，每個斷裂面都會產生不同的凹凸起伏。只要將這些斷裂面測量為立體形狀，便能找到彼此契合的配對。

在中國，相同的技術也被應用於甲骨碎片上。甲骨是商代人們占卜時刻在龜腹甲與牛骨上的記錄。河南大學的張重生研究團隊建立了甲骨綴合資料集。該團隊的綴合演算法取得了前10名準確率98.39%的成果[6]。張展研究團隊也利用甲骨碎片照片訓練了 Deep Rejoining 模型。該研究團隊同樣公開了資料集[7]。

「前10名準確率」這項說法有必要加以解釋。機器並非只給出單一標準答案，而是依序給出十個可能性最高的候選對象。正解落在這十個候選對象之中的機率為98.39%[6]。最後的判斷仍由人類做出。

整合資料的工作也同步進行。安陽師範學院營運名為「殷契文淵」的甲骨資料平台。截至2023年4月，已上傳了23萬餘張甲骨照片。同時也匯集了152部甲骨文相關著作與3,000餘篇論文[8]。

敦煌莫高窟出土的佛經碎片同樣散落於各個國家。陳明坤研究團隊於2026年提出了一種新方法。這種方法即使沒有人類提供標準答案，也能自主學習。此種學習方式稱為自監督學習。研究團隊收集了完整的抄本以及已知配對的碎片，同時也將尚未找到配對的孤立碎片納入資料中。電腦從這些資料中自行生成練習題目。研究團隊設計了一種同時包含邊緣方向的新編碼方式，隨後利用孿生神經網路對兩塊碎片的相似程度進行評分[9]。

在死海古卷方面，歐洲與以色列的研究人員共同打造了一個線上工作平台，名稱為 Scripta Qumranica Electronica。在這個工作平台上，研究人員可以在螢幕上移動並旋轉碎片照片來進行排列。排列的結果可以與其他研究人員分享[10]。碎片拼合工作從過去單憑個人大腦思索的作業，轉變為由眾人共同檢驗的協作歷程。

參考資料

- [1] Abitbol, R., Shimshoni, I., & Ben-Dov, J. (2021). Machine learning based assembly of fragments of ancient papyrus. *Journal on Computing and Cultural Heritage*, 14(3), 1-21. <https://doi.org/10.1145/3460961>
- [2] Pirrone, A., Beurton-Aimar, M., & Journet, N. (2019). Papy-S-Net: A Siamese network to match papyrus fragments. *Proceedings of the 5th International Workshop on Historical Document Imaging and Processing*, 78-83. <https://doi.org/10.1145/3352631.3352646>
- [3] Brown-deVost, B., Kurar-Barakat, B., & Dershowitz, N. (2024). Segmenting Dead Sea Scroll fragments for a scientific image set. *arXiv*. <https://arxiv.org/abs/2406.15692>
- [4] Jiménez, E. (2023, 5월 25일). AI is helping us read ancient Mesopotamian literature. *The Conversation*. <https://theconversation.com/ai-is-helping-us-read-ancient-mesopotamian-literature-205091> (2026年9月3日檢索)
- [5] Phys.org. (2023, 2월 2일). Researcher uses AI to make texts that are thousands of years old readable. *Phys.org*. <https://phys.org/news/2023-02-ai-texts-thousands-years-readable.html> (2026年9月3日檢索)
- [6] Zhang, C., Zong, R., Cao, S., Men, Y., & Mo, B. (2020). AI-powered oracle bone inscriptions recognition and fragments rejoining. *Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence*, 5309-5311. <https://doi.org/10.24963/ijcai.2020/779>
- [7] Zhang, Z., Zhang, H., Guo, A., Jiao, Q., Liu, Y., Li, B., & Gao, F. (2025). Deep rejoining model and dataset of oracle bone fragment images. *npj Heritage Science*, 13(1), 66. <https://doi.org/10.1038/s40494-025-01651-9>
- [8] 王者. (2023, 4월 8일). "殷契文淵"数据平台为甲骨文研究提供助力. *人民日报*. https://news.hangzhou.com.cn/gnxw/content/2023-04/08/content_8509027.htm (2026年9月3日檢索)
- [9] Chen, M., Xiang, Y., & Zheng, Y. (2026). A self-supervised learning approach for pairwise matching of ancient Dunhuang fragments. *Journal on Computing and Cultural Heritage*, 19(1), 1-19. <https://doi.org/10.1145/3776642>
- [10] Scripta Qumranica Electronica. (n.d.). Scripta Qumranica Electronica. University of Haifa. <https://megillot.haifa.ac.il/index.php/en/research/scripta-qumranica-electronica> (2026年9月3日檢索)

4.3. 科學鑑識團隊的追蹤 2：偽作偵探 AI

揪出16件贗品的方法

聖經博物館買下的16件死海古卷殘片全都是假貨[1][2]。調查團隊不只是用肉眼看文字，還用物理和化學方法仔細剖析了材料與墨水。這項鑑定並未使用人工智慧，而是由人透過顯微鏡與成分分析查了出來[1]。

聖經博物館的鑑定並非一次完成。2018年10月，德國聯邦材料研究與檢驗所率先檢驗了其中5件。這5件因不符合古代文物的特徵而撤出展覽[5]。隨後將剩餘殘片全數檢驗的結果於2020年出爐[1][2]。

第一個證據是底材。真正的死海古卷寫在削薄的羊皮紙上。聖經博物館的殘片中，至少有15件是皮革。這種皮革比羊皮紙更厚，紋理也更粗糙[1]。

第二個證據是墨水著墨的位置。調查團隊以3D顯微鏡檢視表面。裂縫深處積有墨水，破損的邊緣也能看到墨水向下流淌的痕跡[1]。

只要思考順序，就能明白為什麼這個痕跡具有決定性。真正的抄本是在平整的羊皮紙上寫字，經過漫長歲月後，皮革乾燥而裂開。因此真品的文字筆劃會在裂縫處斷開。假貨的順序正好相反。偽造者找來已經裂開的古代皮革，在上面塗上墨水，墨水便流進了裂開的縫隙深處[1]。

第三個證據是塗抹在表面的物質。調查團隊發現了殘片曾浸泡在琥珀色液體中的痕跡。這種液體推測是用獸皮製成的動物膠。偽造者企圖藉此模仿真正死海古卷表面光滑的質感[1]。

這三種證據從不同方向指向了同一個答案。單看一項或許純屬巧合，但三者吻合就絕非偶然。文物鑑定中之所以要同時匯集不同種類的證據，原因就在這裡。

手留下的痕跡

製造假貨的人會看著真品照片照樣描摹字跡。這種描摹會留下痕跡。

人在書寫熟悉的字跡時，運筆流暢而不停頓。臨摹陌生的字形時則截然不同。筆劃寫到一半會停下來確認下一筆的位置，筆尖速度變慢，線條也會產生微小的顫抖。這種差異在放大檢視時，會顯現在筆劃的外輪廓上。

電腦能將這種差異以數值捕捉下來。2026年6月，英國布拉德福大學的哈桑·烏加爾研究團隊發表了相關研究。研究團隊開發出判定古代素描作品是否為真品的AI。該模型只檢視從畫作中提取的5個數值：測量圖樣重複的程度與亮度值分散的程度，也測量亮處與暗處的差異，並加上質地均勻的程度與線條如細枝般分叉的程度[3]。

研究團隊以10位畫家的真品素描來訓練模型。作品年代跨越1510年左右至1917年。資料蒐集自6座釋出公開資料的美術館與博物館，大都會藝術博物館與阿什莫林博物館均包含在內[3]。

論文刊登了針對900件判定進行測試的結果。平衡準確率為87.6%。將真品判定為真品的比率為77.8%，將他人作品誤判為該畫家之作的比率為2.6%[3]。

這項研究還有一個值得注意的成果。廣泛使用的大型深度學習模型在這項任務中成績反而較差。大型模型經常將真品誤判為贗品。當資料量較少時，針對問題量身設計的小型模型表現更好[3]。

這個結果對古文獻研究同樣深具參考價值。現存的古代抄本數量稀少，沒有足夠的資料可以餵飽大型模型。這意味著預先指定觀察重點的設計，效果反而更好。

測量筆跡的機器

相同的技術在研究真正的文物時也有所應用。荷蘭格羅寧根大學的姆拉登·波波維奇研究團隊以AI分析了死海古卷的筆跡。研究團隊利用神經網路單獨提取出墨水留下的痕跡，隨後將文字的彎曲程度與傾斜度轉化為數值[4]。

研究團隊使用24件已完成碳定年的抄本來訓練模型。該模型的名字是以諾（Enoch）。以諾估算了135件年代不詳抄本的年齡。平均誤差約在28年至31年之間。結果顯示，有多件抄本的年代比至今已知的更為古老[4]。

這項技術對真偽判定也有所助益。真正抄寫員的字跡在各個時代都具有一定的統計特徵。假貨則脫離了這種分布。人眼雖然覺得兩種字跡相似，機器卻能以數字呈現出差異。

不過，這種幫助僅止於多獲得一份參考資料。以諾的研究有需要審慎看待之處。該模型是為了預測年代而受訓，並非為了分辨真假而訓練。不能僅因年代估算出現異常，就斷定它是假貨。機器的運算結果唯有在其所學習的任務範圍內才具有意義。

機器無法取代的事物

機器的判定同樣存在極限。最終判定聖經博物館殘片真偽的，並非單靠AI。顯微鏡觀察、化學檢驗與文獻學知識均共同參與其中[1][2]。

來源記錄的分量依然不減。來源記錄是記載文物出處的記錄。羅爾斯頓教授建議將擁有發掘記錄的文物與從市場購買的文物區別看待。經由發掘出土的卷軸，留有是在哪座洞穴、何時出土的記錄；流落於市場上的殘片則沒有這種記錄[5]。

機器為人類解讀人眼遺漏的微小訊號，人則負責判斷該訊號代表何種意義。辨別偽作的工作，必須兩者兼備才能水落石出。

參考資料

[1] Greshko, M. (2020, 3월 13일). "Dead Sea Scrolls" at the Museum of the Bible are all forgeries. National Geographic. <https://www.nationalgeographic.com/history/article/museum-of-the-bible-dead-sea-scrolls-forgeries> (2026年9月3日檢索)

[2] Solly, M. (2020, 3월 16일). All of the Museum of the Bible's Dead Sea Scrolls are fake, report finds. Smithsonian Magazine. <https://www.smithsonianmag.com/smart-news/all-museum-bibles-dead-sea-scrolls-are-fake-report-finds-180974425/> (2026年9月3日檢索)

- [3] Ugail, H., Ritch-Frel, J., Matuzava, I., & Stork, D. G. (2026). Verification of historical sketches via one-class learning on compact feature representations. PLOS One, 21(6), e0344796. <https://doi.org/10.1371/journal.pone.0344796>
- [4] Popović, M., Dhali, M. A., Schomaker, L., van der Plicht, J., Lund Rasmussen, K., La Nasa, J., Degano, I., Perla Colombini, M., & Tigchelaar, E. (2025). Dating ancient manuscripts using radiocarbon and AI-based writing style analysis. PLOS One, 20(6), e0323185. <https://doi.org/10.1371/journal.pone.0323185>
- [5] Museum of the Bible. (2018, 10월 22일). Museum of the Bible releases research findings on fragments in its Dead Sea Scrolls collection. Museum of the Bible. <https://www.museumofthebible.org/newsroom/museum-of-the-bible-releases-research-findings-on-fragments-in-its-dead-sea-scrolls-collection> (2026年9月3日檢索)

4.4. 揭開的真相

劃分確定拼合與候選拼合的界線

即使電腦指出兩塊殘片可以拼合，事情也並未就此結束。機器給出的只不過是機率。研究人員嚴格區分確定的拼合與候選的拼合。

唯有多種證據同時吻合時，確定的拼合才會獲得認可。首先，斷裂面的形狀必須嚴絲合縫地契合。其次，材質的紋理必須連貫。紙草的莖部紋路必須相連[1]，羊皮紙的毛孔痕跡也必須延續。接著，跨越邊界的文字筆劃必須自然相接。最後，拼合後的文句必須符合該語言的語法。

候選拼合則是僅符合其中部分條件的情況。依據蟲蛀蛀孔的間距推測殘片位置的結果便屬於此類；墨水成分相同而歸為同一卷軸群組的結果亦屬候選。這類成果雖能成為研究的起點，卻無法成為結論。其流程是由電腦篩選出候選對象，再由人工進行確認[1]。

若不嚴守此一劃分將會非常危險。若僅因機率高就將兩塊殘片拼在一起，便可能產生原本不存在的句子。一旦該句子被引用，錯誤就會蔓延到後續研究中。因此在資料記錄上，確定與候選都會分開記載。

機器給出的數值也是以此區分為前提。篩選率高並不代表拼合就此確定[1]，正確答案落在排名靠前的候選名單中的機率也是如此[2]。無論如何，最後的判定都必須由人來做。

重獲新生的文字

藉由這種方式，真實的文獻得以重獲新生。Fragmentarium自2018年起處理了約22,000塊泥板殘片。該程式尋獲了多處新的拼合與新的抄本。研究團隊於2023年2月將Fragmentarium向大眾公開，並同時推出了《吉爾伽美什史詩》的數位版本[3]。

在重獲新生的文本中，包含了《吉爾伽美什史詩》。《吉爾伽美什史詩》是美索不達米亞的英雄故事。2022年11月，程式辨認出一塊殘片。這塊殘片出自西元前130年左右抄寫的版本。在目前確認的吉爾伽美什抄本中，其年代相對較晚，這意味著直到那個時代，人們仍持續抄寫著這個故事[3]。

楔形文字在不同地區有著相異的拼法，這是因為抄寫員會根據自己的方言改變用字。人工以肉眼查找時，這種差異是一道巨大的高牆。研究團隊借用比對基因序列的方法跨越了這道障礙，使彼此略有差異的字串也能進行比對[4]。

死海古卷方面的資料基礎也更加寬廣。以色列文物局將超過25,000件殘片文物拍攝成高解析度照片並公開[5]。研究人員能在這些照片上移動並排列殘片[6]。自動清除比例尺與色卡等干擾物的影像處理方法也已公開發布[7]。

2026年開啟的下一扇門

2026年6月展開的「追蹤抄寫員與卷軸」研究開啟了下一扇門。該研究提出的疑問比拼接殘片更進一步，探問死海古卷是在何處、由何人所寫[8]。

研究團隊檢驗了約250件羊皮紙、紙草與墨水樣本。所謂樣本，是指用於檢驗的微小樣品。研究並結合了AI筆跡分析與古文字研究。以色列文物局、比薩大學與那不勒斯大學共同合作，南丹麥大學與比利時魯汶大學也參與其中，並與柏林及杜林博物館所收藏的埃及紙草卷進行比對[8]。

若樣本成分相同，出自同一工坊的可能性便大幅提升。這項資訊亦能應用於殘片拼合。若先將材料出處相同的殘片歸類在一起，需要比對的配對數量就會大幅減少。

剩餘的殘片與未完的工作

儘管成果豐碩，但至今仍有許多殘片尚未找到自己的歸宿。文物局的收納盒中，依然留有不知出自哪卷卷軸的殘片[8]。泥板方面也有許多尚未被發現的拼合[4]。

設備的差異同樣是個棘手問題。若要使用3D掃描器測量文物，就必須將文物移至設備前。但許多博物館難以將文物移出館外。因此有些文物僅有照片資料，缺乏3D立體資料。這類文物便無法採用比對斷裂面的方法。

技術也帶來了待解的課題。殘片照片在各國是由不同的設備拍攝而成。當照明與背景不同時，即便相同的演算法也會給出不同的答案。因此，研究人員建立了附有標籤解答的公開資料集，以便能用相同的標準來比對彼此的成果[7]。敦煌殘片的研究也朝向即使沒有人工標註答案也能自我學習的方向發展[9]。

這項技術並非要取代人類。機器負責縮減需要比對的配對範圍，人則負責判斷文意。拼合重現的句子是否符合那個時代的語氣與文風，唯有經由人類研讀才能得知。正是因為兩者的角色互不重疊，才得以相輔相成、共同協作。

開放程度也已不同以往。過去僅有少數學者能看到原始照片，如今泥板資料與死海古卷照片都在網路上公開[4][5]，連學生都能在螢幕前嘗試移動並排列殘片[6]。

數位黏著劑並非以物理方式黏合殘片。文物依然原封不動地留在盒子裡，卷軸僅在電腦螢幕上重新接合。因此一旦判定有誤，隨時都可以復原回溯。散落長達2,000年的文字，正以這種方式一點一滴地重回完整的篇章。

參考資料

- [1] Abitbol, R., Shimshoni, I., & Ben-Dov, J. (2021). Machine learning based assembly of fragments of ancient papyrus. *Journal on Computing and Cultural Heritage*, 14(3), 1-21. <https://doi.org/10.1145/3460961>
- [2] Zhang, C., Zong, R., Cao, S., Men, Y., & Mo, B. (2020). AI-powered oracle bone inscriptions recognition and fragments rejoining. *Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence*, 5309-5311. <https://doi.org/10.24963/ijcai.2020/779>
- [3] Phys.org. (2023, 2월 2일). Researcher uses AI to make texts that are thousands of years old readable. *Phys.org*. <https://phys.org/news/2023-02-ai-texts-thousands-years-readable.html> (2026年9月3日檢索)
- [4] Jiménez, E. (2023, 5월 25일). AI is helping us read ancient Mesopotamian literature. *The Conversation*. <https://theconversation.com/ai-is-helping-us-read-ancient-mesopotamian-literature-205091> (2026年9月3日檢索)
- [5] Israel Antiquities Authority. (n.d.). The digital library. The Leon Levy Dead Sea Scrolls Digital Library. <https://www.deadseascrolls.org.il/about-the-project/the-digital-library> (2026年9月3日檢索)
- [6] Scripta Qumranica Electronica. (n.d.). Scripta Qumranica Electronica. University of Haifa. <https://megillot.haifa.ac.il/index.php/en/research/scripta-qumranica-electronica> (2026年9月3日檢索)
- [7] Brown-deVost, B., Kurar-Barakat, B., & Dershowitz, N. (2024). Segmenting Dead Sea Scroll fragments for a scientific image set. *arXiv*. <https://arxiv.org/abs/2406.15692>
- [8] The Jerusalem Post. (2026, 6월 30일). New AI-powered research project aims to uncover the origins of the Dead Sea Scrolls. *The Jerusalem Post*. <https://www.jpost.com/archaeology/article-900938> (2026年9月3日檢索)
- [9] Chen, M., Xiang, Y., & Zheng, Y. (2026). A self-supervised learning approach for pairwise matching of ancient Dunhuang fragments. *Journal on Computing and Cultural Heritage*, 19(1), 1-19. <https://doi.org/10.1145/3776642>

第5章. 拼合泥板的碎裂斷面：吉爾伽美什史詩的計算學重建

5.1. 案件檔案

跨越4千年的故事

吉爾伽美什史詩在現存的敘事文學中屬於年代相當古老的一部。比這更早的蘇美歌謠也有流傳至今。主角是統治美索不達米亞城市烏魯克的國王吉爾伽美什[1]。他與好友恩奇都一同前往遙遠的雪松森林。兩人與守護該森林的怪物胡姆巴巴展開戰鬥[2]。後來恩奇都死去，吉爾伽美什便啟程尋找不死之道[1]。

記錄這段故事的載體不是紙張，而是泥板。古代的書吏將蘆葦斜切削尖，製成尖筆。用尖筆壓印在濕潤的泥板上，就會留下楔形的痕跡。這不是劃線寫出，而是按壓印出的文字。泥土乾燥後，那些痕跡便隨之硬化固定。因此這種文字被稱為楔形文字。立志成為書吏的學徒們，以抄寫這部史詩來練習書寫。多虧如此，抄本殘片在現今中東的多處遺址中出土。甚至在埃及的阿瑪納遺址中，也出土了這部史詩的殘片。這意味著該故事曾傳播到遠離美索不達米亞的地方[1]。

流傳至今的版本主要有兩種。其一是西元前第二千紀初寫成的古巴比倫版本。另一種是後來整理完成的標準巴比倫版本。標準版本由十二塊泥板構成。其中十一塊是正篇。第十二塊泥板被視為後來增補的附錄。今天我們閱讀的譯本，就是以這個標準版本為基礎。殘缺空白處則以古巴比倫版本和其他殘片來填補。整理標準版本的人，據悉是一位名叫辛-勒基-烏寧尼的書吏。他將散落各處的古老歌謠編纂成一部長篇作品。他決定了每塊泥板要容納的內容，並調整了行數。然而，即使是這個標準版本，殘缺之處也極為可觀。從頭到尾完整留存的抄本，連一份都沒有[1]。

故事的內容即使在今天讀來也不令人感到陌生。吉爾伽美什是一位力大無窮卻欺壓百姓的國王。眾神創造了恩奇都並派他前去與之抗衡。兩人大打出手後結為摯友。他們合力擊敗了胡姆巴巴，也打倒了天之公牛。作為代價，恩奇都染病身亡。看見摯友死亡的吉爾伽美什，領悟到自己終究也會死去。第11塊泥板中，收錄了他去拜訪在大洪水中倖存的人並聆聽其故事的場景[1]。

這段故事曾被遺忘超過兩千年。這是因為能識讀楔形文字的人消失了。直到19世紀，學者們重新解讀楔形文字，這部史詩才重返人世。這一切的開端，是一名研究人員在翻檢從尼尼微運來的泥板箱子時，辨認出了洪水的故事。此後的150年間，學者們蒐集散落的抄本，一點一滴地填補了正文[1]。

尼尼微與西元前612年的大火

標準版本的抄本大多出自一座古代圖書館。那是西元前7世紀亞述國王亞述巴尼拔在尼尼微王宮建立的藏書庫。這座圖書館廣泛蒐集了占卜書、醫學書、辭典和文學作品。國王下令從帝國各地徵集泥板，並致信臣子命其尋訪缺藏的書籍。亞述巴尼拔是一位能親自識讀楔形文字的國王[3]。

西元前612年左右，尼尼微陷入了一場大火。王宮傾塌，圖書館也一同焚燬。這場大火對文物同時帶來了兩種截然不同的影響。原本只在陰涼處風乾的泥板，在烈火中被燒得堅硬[4]。燒過的泥板如同窯中出產的磚塊，獲得了能夠長久保存的體魄。如果身處同一場大火中的是莎草紙或皮革，早已化為灰燼消失了。

然而，倒塌建築物的重量將那些堅硬的泥板砸得粉碎。如今大英博物館保管的尼尼微出土泥板與殘片超過3萬件[3][4]。這是19世紀中葉發掘團從土堆中收集並裝船運回的文物[3]。同一個故事，是以散落成數十塊殘片的狀態呈現在我們面前。

殘片的大小各不相同。有的如手掌般大，有的僅如大拇指指甲般小。一塊殘片上殘存的文字往往只有寥寥數行。沒有人能預先告知這幾行文字究竟屬於史詩十二塊泥板中的哪一部分。識讀出文字內容，與為其找到原本的位置，是截然不同的兩項工作。

當一塊殘片找回它原本的位置，其含意甚至可能完全改變。只有前後脈絡相連，才能知曉那是誰說的話、屬於哪一個場景。分離時只是一連串不知所云的單詞排列，此時才成為完整的句子。一次成功的拼合，甚至能讓故事中的一整幕場景完整重現。

無法靠雙手拼合35萬行

拼合破碎殘片的工作，長期以來只能仰賴人類的雙手與記憶。學者在庫房中拿出一塊殘片，接著回想其他殘片的照片與手摹圖，試圖尋找配對。當判定兩塊殘片出自同一塊泥板時，這項成果便被稱為拼合。

這種方式受限於人類的大腦記憶容量。在亞述學誕生後的150年間，如此拼合成功的案例約有5,000件[5]。然而遺留下來的殘片數量遠不止於此。光是尼尼微出土的泥板，大英博物館就藏有3萬多件[3]。若加上其他收藏，數量更是大幅增加。其中絕大多數自古代以來就從未被閱讀過[5]。全世界博物館中留存的楔形文字泥板總數，估計約有50萬件[6]。

這項工作之所以困難，還有另一個原因。殘片分散在多個國家。大英博物館、柏林、耶魯大學、費城的賓夕法尼亞大學博物館各自擁有收藏品[7]。若要親手比對兩塊殘片，學者必須橫跨兩大洲奔波。

楔形文字本身也極為棘手。同一個詞彙，不同的書吏寫法略有差異。拼寫會隨著地域與時代而有所不同[5]。即便用肉眼比較兩塊殘片上的文字，也很難看出它們是否出自同一部作品。

再加上一個楔形符號往往有多種讀音。同一個印記在不同語境下具有不同的音值。面對數十萬行這樣的文字來尋找配對，絕非人類肉眼所能完成。

因此，德國慕尼黑路德維希-馬克西米利安大學的恩里克·希門尼斯教授選擇了另一條道路。2018年，他啟動了一項盡可能完整復原巴比倫文獻的計畫。該計畫旨在將散落於世界各地的泥板釋讀文本匯集一處[5]。這項計畫的名稱為電子巴比倫文獻圖書館，簡稱 eBL。eBL 所建立的殘片拼合作業平台名為 Fragmentarium。截至2024年，該工作平台已累積了約2萬5千件泥板的釋讀文本。若以行數計算，已超過35萬行[7]。

這項工作之所以迫切，原因在於時間不等人。泥板縱使經過火燒，也並非能永恆留存。鹽分與濕氣會使表面膨脹剝落，抹去上頭的文字。在博物館庫房中等待識讀的泥板，至今仍有數十萬件[5]。其中蘊藏了多少故事，無人知曉。

識讀楔形文字的人數也是個問題。全球的亞述學研究者並不多。在留存的泥板中，已被翻譯的僅占一部分[6]。一個人窮盡一生所能閱讀的泥板數量終究有限。這正是向機器尋求協助的原因。

參考資料

- [1] George, A. R. (2003). *The Babylonian Gilgamesh Epic: Introduction, Critical Edition and Cuneiform Texts*. Oxford University Press.
- [2] Al-Rawi, F. N. H., & George, A. R. (2014). Back to the Cedar Forest: The beginning and end of Tablet V of the Standard Babylonian Epic of Gilgamesh. *Journal of Cuneiform Studies*, 66, 69-90.
<https://doi.org/10.5615/jcunestud.66.2014.0069>
- [3] The Ashurbanipal Library Project. (n.d.). What is the Library? Oracc, University of Pennsylvania.
<https://oracc.museum.upenn.edu/asbp/whatisthelibrary/index.html> (2026年9月3日檢索)
- [4] British Museum. (n.d.). What was Ashurbanipal's Library? britishmuseum.org.
<https://www.britishmuseum.org/research/projects/what-was-ashurbanipals-library> (2026年9月3日檢索)
- [5] Jiménez, E. (2023, 5월 25일). AI is helping us read ancient Mesopotamian literature. *The Conversation*.
<https://theconversation.com/ai-is-helping-us-read-ancient-mesopotamian-literature-205091> (2026年9月3日檢索)
- [6] Wang, Z. (2026). TabletCraft: Bridging a 4,000-year cultural gap with bidirectional Akkadian NMT and cuneiform rendering. arXiv:2608.02609. <https://arxiv.org/abs/2608.02609> (2026年9月3日檢索)
- [7] Cobanoglu, Y., Laasonen, J., Simonjetz, F., Khait, I., Cohen, S., Földi, Z., Härtinen, A., Heinrich, A., Mitto, T., Rozzi, G., Sáenz, L., & Jiménez, E. (2024). Transliterated cuneiform tablets of the electronic Babylonian Library platform. *Journal of Open Humanities Data*, 10, 19. <https://doi.org/10.5334/johd.148>

5.2. 科學鑑識隊的追蹤：泥板邊緣數位融合與 Fragmentarium

將文字轉寫為字母的第一步

Fragmentarium 的出發點在於文字。研究人員看著泥板照片，將楔形文字轉寫為羅馬字[1]。這項工作稱為音譯。這類似於在韓語中將漢字以韓文字音標記出來一樣。也就是將一個個楔形符號按照其音值轉換為字母。對於破損而無法辨讀的部分，則以空白符號標註。

如此建立的音譯文本，成為電腦可以處理的字串。eBL 從2018年5月至2023年8月持續累積這些音譯文本。其成果是讓約2萬5千件泥板、超過35萬行的文字轉化為可供檢索的資料。eBL 管理的泥板記錄目錄超過26萬件。大英博物館、耶魯巴比倫收藏、賓夕法尼亞大學博物館的藏品均涵蓋其中[1]。

音譯工作由人工進行。這是因為機器尚無法從照片中完美擷取出文字。eBL 研究團隊持續進行這項音譯工作超過五年。其成果已作為任何人皆可下載的公開資料釋出[1]。若想讓機器工作，首先必須由人類為其建立資料。

宛如基因檢測的字串比對

核心演算法來自一個意想不到的領域。那就是生物學家比對基因序列時所使用的字串對齊方法[2]。基因是由 A、T、G、C 四個字母長串連接而成的鏈條。比對兩種生物的基因時，即使字母略有差異，也會將它們對齊在相同位置進行觀察。

楔形文字也面臨同樣的問題。因為每位書吏書寫同一個詞的方式都不盡相同。希門尼斯教授認為，這種書寫上的變異與基因突變十分相似。因此，他將生物學的對齊演算法加以修改，使其適用於楔形文字[2]。

電腦實際計算的是簡短的文字片段。將音譯文本按幾個字母切分出來的片段被稱為 N-gram。系統會計算兩塊泥板所共享的 N-gram 數量。共享較長片段的配對組合會獲得更高的評分。評分較高的配對組合將作為候選名單呈現在人類面前。

Fragmentarium 透過網頁介面開啟。世界上任何國家的研究人員都可以連線查看相同的資料。倫敦的殘片與慕尼黑的殘片可以並排在同一個螢幕上進行比對[8]。再也不需要收拾行囊搭乘飛機了。

這種方式即使遇到拼寫不一致也能應對。即使某個單詞的拼法不同，只要周圍的片段重疊，依然能獲得分數。eBL 已將這個配對程式上傳至公開儲存庫。其他研究人員也可以下載並使用相同的工具[6]。

電腦在文字識讀本身也提供協助。一個楔形符號往往有多種發音。eBL 訓練的模型在從多種意涵中選取正確選項的問題上，達到了98%的準確率[2]。

舉例來說便很容易理解這項工作為何不可或缺。同一個楔形印記在某些句子中只讀作一個音。在其他句子中則被當作代表整個詞彙的符號來識讀。如果識讀錯誤，配對分數就會失真。糾正識讀結果，正是提高拼合準確度的關鍵。

成果以數字展現出來。亞述學在150年間以人工拼接的拼合成果約為5,000件。而 eBL 僅用了五年時間，就額外找到了約1,500件拼合[2]。光是在建立音譯資料集的過程中，就發現了約1,250件拼合[1]。此外，雖然斷面沒有直接相接、但確認屬於同一部作品的殘片也有數千件[2]。

填補殘缺處的語言模型

即使將殘片拼接在一起，文本各處依然存在空白。這種空白處被稱為缺文。從這裡開始，便由語言模型登場接棒。

2023年，特拉維夫大學與阿里埃勒大學的研究團隊推出了阿卡德語翻譯模型。該模型無論輸入楔形文字符號，還是輸入字母音譯文本，都能將其翻譯成英語。它同時處理了這兩種輸入形式。但該模型並未處理將英語反向翻譯回阿卡德語的功能。用於訓練的資料是已經出版的泥板釋讀文本[3]。

2026年，出現了連反方向也能處理的工具。這是一款名為 TabletCraft 的開源軟體。該工具能夠進行阿卡德語至英語、英語至阿卡德語的雙向翻譯。訓練資料為11萬6千對句子。楔形符號轉換表收錄了超過1萬4千組對應關係。開發者表示，這項研究旨在打造能識讀剩餘50萬件泥板的工具[4]。

這類模型會根據空格的前後語境，以機率推算並提出可能填入的詞彙。巴比倫詩歌中有一種將結構相同的兩行並列放置的規則。模型也會將這種形式視為線索。模型所給出的並非確定的答案，而是候選清單。

連字體都能識讀的2026年工具

2026年5月18日，德國維爾茨堡大學公開了一項新工具。該工具由美茵茨學術院與多特蒙德工業大學共同開發。名稱為 Paleographicum。該工具從泥板照片中學習每一個楔形符號的具體形態。隨後便能在整個藏品庫中，搜尋出以相同手法按壓書寫的符號[5]。

這套工具處理的照片達7萬張，包含的符號超過500萬個。以往比較五塊殘片的筆跡需要花費三天。使用這套工具，只需5分鐘即可完成。如同每個人筆跡各異，每位抄寫員壓印楔形文字的角度與間距也各不相同。以這些差異為線索，便能篩選出源自同一塊泥板的殘片[5]。

這套工具處理的西臺楔形文字共有375個符號。目前已知的西臺泥板殘片達3萬件。主導該研究的丹尼爾·施韋默教授表示，藉由這套工具可省下數千小時[5]。

研究團隊的下一個目標是辨識出個別抄寫員。這意味著要找出同一人書寫的所有泥板。一旦實現，便能描繪出某位抄寫員曾在何處工作、後來又遷徙至何地。3,500年前人們的職業生涯便能透過文獻資料重見天日[5]。

拼合碎裂泥板邊緣的另一條途徑

透過形狀而非文字來拼合殘片的研究也持續了許久。碎裂的泥板會留下斷面。這個面稱為斷裂面。若兩塊殘片的斷裂面能夠相互契合，即代表它們來自同一塊泥板。

走在這條道路上的代表性專案，就是虛擬楔形文字泥板復原專案。英國基爾大學的珊卓·伍利主持該專案，提姆·柯林斯與埃爾蘭德·格爾肯共同參與。研究團隊使用3D掃描器掃描殘片，並讓電腦判斷兩塊殘片的斷裂面是否能夠相互契合。年代久遠的殘片邊緣往往已經磨損，研究中甚至將這種磨損程度納入了運算考量[9]。

這套方法的成果之一出現在大洪水故事中，那就是名為《阿特拉哈西斯》的洪水史詩第三塊泥板殘片。其中一塊殘片在倫敦，另一塊在日內瓦。兩地相距約1,000公里。由於無法將實物拼對在一起，過去一直無法確認其拼接關係。當研究人員將3D資料進行疊合比對後，兩塊殘片成功接合了起來[9]。該方法於2014年以論文形式發表[10]。

處理3D資料的研究至今仍在持續。2026年3月發表的一篇論文中，探討了利用泥板的3D點雲資料來分類文物資訊的神經網路。點雲是將物體表面記錄為無數個點的座標之資料。用3D掃描器掃描泥板時便會產生這種資料，使表面的起伏與厚度轉化為數字記錄下來。這項研究取得了優於先前基於 Transformer 模型的效果[7]。

然而，目前 Fragmentarium 實際使用的主要利器並非3D形貌，而是照片與轉寫文本，也就是文字與圖像[1][2]。斷裂面拼合必須逐一掃描殘片，需要耗費大量人力工夫。這兩條途徑並非相互取代，而是並行不悖。

從照片中辨識符號的工作也進展迅速。2026年6月，eBL 研究團隊發表了新論文，提出一種能在泥板照片中自動定位楔形符號位置的方法。作為分析對象的殘片照片超過8萬7千張，辨識出的符號達290萬個，成效較先前的方法大幅提升[11]。

參考資料

- [1] Cobanoglu, Y., Laasonen, J., Simonjetz, F., Khait, I., Cohen, S., Földi, Z., Härtinen, A., Heinrich, A., Mitto, T., Rozzi, G., Sáenz, L., & Jiménez, E. (2024). Transliterated cuneiform tablets of the electronic Babylonian Library platform. *Journal of Open Humanities Data*, 10, 19. <https://doi.org/10.5334/johd.148>
- [2] Jiménez, E. (2023, 5月 25日). AI is helping us read ancient Mesopotamian literature. *The Conversation*. <https://theconversation.com/ai-is-helping-us-read-ancient-mesopotamian-literature-205091> (2026年9月3日檢索)
- [3] Gutherz, G., Gordin, S., Sáenz, L., Levy, O., & Berant, J. (2023). Translating Akkadian to English with neural machine translation. *PNAS Nexus*, 2(5), pgad096. <https://doi.org/10.1093/pnasnexus/pgad096>
- [4] Wang, Z. (2026). TabletCraft: Bridging a 4,000-year cultural gap with bidirectional Akkadian NMT and cuneiform rendering. *arXiv:2608.02609*. <https://arxiv.org/abs/2608.02609> (2026年9月3日檢索)
- [5] Julius-Maximilians-Universität Würzburg. (2026, 5月 18日). AI reads cuneiform: A milestone for Ancient Near Eastern Studies. *EurekAlert!*. <https://www.eurekalert.org/news-releases/1128630> (2026年9月3日檢索)

- [6] Electronic Babylonian Literature. (n.d.). ngram-matcher [소프트웨어]. GitHub. <https://github.com/ElectronicBabylonianLiterature/ngram-matcher> (2026年9月3日檢索)
- [7] Hagelskjær, F. (2026). A novel network for classification of cuneiform tablet metadata. arXiv:2603.03892. <https://arxiv.org/abs/2603.03892> (2026年9月3日檢索)
- [8] electronic Babylonian Library. (n.d.). Fragmentarium. ebl.lmu.de. <https://www.ebl.lmu.de/fragmentarium> (2026年9月3日檢索)
- [9] International Association for Assyriology. (2022, 10월 22일). In the spotlight: The "Virtual Cuneiform Tablet Reconstruction Project". iaassyriology.com. <https://iaassyriology.com/in-the-spotlight-the-vctr-project/> (2026年9月3日檢索)
- [10] Collins, T., Woolley, S. I., Hernandez Munoz, L., Lewis, A., Ch'ng, E., & Gehlken, E. (2014). Computer-assisted reconstruction of virtual fragmented cuneiform tablets. In Proceedings of the 2014 International Conference on Virtual Systems and Multimedia (pp. 70-77). IEEE. <https://doi.org/10.1109/VSMM.2014.7136691>
- [11] Che, W., Garcés Arias, E., Niaz, A., Bender, A., & Jiménez, E. (2026). Automated sign detection across the Electronic Babylonian Library: A large-scale dataset and end-to-end cuneiform OCR pipeline. arXiv:2606.22608. <https://arxiv.org/abs/2606.22608> (2026年9月3日檢索)

5.3. 揭曉的真相

重返《吉爾伽美什史詩》的詩行

Fragmentarium 為《吉爾伽美什史詩》新找到了20件殘片。多虧於此，超過100行的正文得以重歸原位，好幾個曾消失的場景也重獲新生。恩奇都勸阻吉爾伽美什不要殺死洪巴巴的場面便是其中之一。兩人在殺死洪巴巴後前往尼普爾的場景也重新串連了起來。尼普爾是供奉神祇恩利爾的宗教中心城市[1]。

2025年出現了一個展現這套方法有多麼迅速的案例。巴格達大學的安馬爾·法迪爾與希梅內斯教授復原了一首讚美巴比倫的長詩。這首詩已被遺忘超過1,000年。研究團隊透過 eBL 平台找到了屬於同一首詩的30件抄本。若以人工處理，這項工作原本需要數十年。將散落的抄本拼集之後，長達250行的詩篇絕大部分重新接續了起來[3]。

這首詩描繪了巴比倫這座城市的一天，呈現了河流、城門與人們的日常勞作，其中也包含了關於女祭司司職責的全新記載。能發現多達30件抄本，意味著這首詩在當時廣為流傳，極有可能是當時學校學生用來抄寫練習的文本[3]。

這項工作所耗費的時間展現了該方法的威力。若要以人工尋找30件抄本，必須逐一翻查世界各地博物館的藏品目錄；但在 eBL 平台上，同樣的工作只需一次搜尋即可完成[3]。

這套方法並非僅用於《吉爾伽美什史詩》。巴比倫尼亞的占卜書、醫學典籍、詞典與祈禱文也同樣適用此方法。eBL 所找到的拼接涵蓋了多種類型的文獻[5]。

由人手率先解開的雪松森林

並非所有謎團都由機器解開。《吉爾伽美什史詩》中廣為人知的新發現，其實出自人手。

2011年，伊拉克蘇萊曼尼亞博物館收購了一塊泥板。法魯克·阿爾拉維（Al-Rawi）與安德魯·喬治解讀了這塊泥板，並於2014年發表了研究結果。該泥板記載了標準版第五塊泥板的前半部與結尾部分，從中發掘出了約20行全新正文[4]。

這段全新正文以聲音生動描繪了雪松森林。猴群喧鬧啼叫，蟬鳴如合唱，百鳥齊鳴。這些聲音每天取悅著守護森林的洪巴巴。文中還出現了兩人初見森林時驚訝駐足的場景，亦有吉爾伽美什在殺死洪巴巴後對砍伐樹木感到後悔的詩句[4]。美索不達米亞平原並沒有棲息猴群的森林，因此這段描寫究竟源自何處，成了學者們持續探討的謎題。

這塊泥板的來歷並不光明磊落。博物館是向文物商收購而來的，這是為了防止戰後在伊拉克被盜掘的文物遭到流散而採取的措施。該文物究竟出土自哪處遺址並無任何記錄[4]。此類文物在尋找相互匹配的殘片時，所擁有的線索往往更少。

這項發現即使沒有人工智慧也得以實現。然而，今後能否頻繁獲得同類發現，則取決於工具。蘇萊曼尼亞泥板為世人所知，需要機緣巧合與人類的慧眼。eBL 正試圖將這種偶然轉變為搜尋[5]。

圓形方舟的故事

大洪水的故事也是類似的情況。如果說《吉爾伽美什史詩》第五塊泥板描寫了雪松森林，那麼第十一塊泥板記載的便是大洪水[6]。以下將探討的泥板並非《吉爾伽美什史詩》的抄本，而是記載了相同洪水傳承、年代更為古老的泥板[8]。

大英博物館的歐文·芬克爾博士於2009年檢視了一塊泥板。這是一位私人收藏家帶到博物館的泥板，他研判該泥板約書寫於西元前1750年左右。尺寸約為寬11.5公分、長6公分，上面刻有60行文字[8]。內容講述神指示一個人建造一艘船的故事，其中也出現了讓動物兩兩登船的指示[7]。

引人注目的是船的形狀。這艘船不是狹長的箱型，而是一艘圓形的船，外形就如同在河流上使用的圓形皮舟（coracle）[7]。泥板上記錄了繩索與木架結構，還提到了防止滲水所塗抹的瀝青用量。瀝青是一種黑而黏稠的天然油脂。芬克爾博士於2014年將這一解讀成果集結成書出版[8]。

這塊泥板並非由人工智慧解讀，而是由人類逐行辨讀出來的。人工智慧的價值體現在其他層面：將散落於世界各地的洪水故事抄本匯集在一處，並找出彼此重疊的詩行[5]。

機器做的事與人做的事

該領域將成果分為三個層次來處理。

第一層是物理拼接，即兩塊殘片的斷裂面能夠實際相互吻合的情況。此時，兩塊殘片便被確認出自同一塊泥板。

第二層是文獻拼接，即兩塊殘片雖然沒有直接相連，但被證實屬於同一部作品的不同部分。Fragmentarium 所找到的大多數案例皆屬此類[5]。屬於同一部作品的殘片便以此方式確認了數千件[2]。

這兩個層次的性質截然不同。物理拼接一旦經由肉眼確認，便成為既成事實；而文獻拼接則依賴機率推估，因此一旦出現新殘片就有可能發生變更。eBL 會將這兩個層次區分開來進行記錄[5]。

第三層是詮釋，即探討復原後的正文所蘊含的意義。雪松森林中的猴叫聲究竟源自何種記憶，機器無從知曉；圓形船究竟如何演變成聖經中方形的方舟，機器也無法回答。這些問題屬於研究古代文獻的學者與歷史學家的職責所在。

機器所提出的候選配對，必須經過人類的檢視才能成為正式記錄。學者打開候選清單，並列檢視兩塊殘片的照片。在仔細審視文字字形、行距乃至泥板厚度之後，才最終確認拼接[5]。機器所做的工作，僅止於從無數配對中挑選出少數值得一探的選項[2]。

2026年6月發表的一篇綜述文章也劃下了相同的界線。現今的人工智慧無法自行解讀楔形文字，它只是一項協助人類研究的工具[9]。維爾茨堡大學的 Paleographikum 也是如此，它僅僅是將耗時三天的字跡比對縮短至5分鐘而已。決定比對內容並判斷比對結果的，依舊是人類[10]。

改變的唯有速度。在過去150年間僅有5,000件的拼接，如今在五年間就增加了1,500件[2]。吉爾伽美什所苦苦追尋的永恆生命，最終由他的故事所獲得。在泥土中破碎著等待了2,600年的語句，如今正一一重歸原位。

未來尚待完成的工作依然繁多。全世界50萬件泥板中，已被解讀的僅占一部分[11]。收錄進 eBL 轉寫資料庫的泥板有2萬5千件[5]。3萬件西臺殘片才剛完成照片整理[10]。拼合斷裂面的3D作業才剛剛起步。下一代的學子將會繼承這份清單，屆時《吉爾伽美什史詩》的留白處必將少於現今。尼尼微大火所留下的殘片尚未被全數解讀，而能夠將其逐一辨讀的工具正在逐年增加。

參考資料

[1] Ofgang, E. (2024, 8월 12일). Piecing together an ancient epic was slow work. Until AI got involved. The New York Times. https://edition.pagesuite.com/tribune/article_popover.aspx?guid=235bb893-7f2d-41cd-9857-139c7d1e8355 (2026年9月3日檢索)

- [2] Jiménez, E. (2023, 5월 25일). AI is helping us read ancient Mesopotamian literature. The Conversation. <https://theconversation.com/ai-is-helping-us-read-ancient-mesopotamian-literature-205091> (2026年9月3日檢索)
- [3] Fadhil, A. A., & Jiménez, E. (2025). Literary texts from the Sippar Library V: A hymn in praise of Babylon and the Babylonians. *Iraq*, 1-58. <https://doi.org/10.1017/irq.2024.23>
- [4] Al-Rawi, F. N. H., & George, A. R. (2014). Back to the Cedar Forest: The beginning and end of Tablet V of the Standard Babylonian Epic of Gilgameš. *Journal of Cuneiform Studies*, 66, 69-90. <https://doi.org/10.5615/jcunestud.66.2014.0069>
- [5] Cobanoglu, Y., Laasonen, J., Simonjetz, F., Khait, I., Cohen, S., Földi, Z., Härtinen, A., Heinrich, A., Mitto, T., Rozzi, G., Sáenz, L., & Jiménez, E. (2024). Transliterated cuneiform tablets of the electronic Babylonian Library platform. *Journal of Open Humanities Data*, 10, 19. <https://doi.org/10.5334/johd.148>
- [6] George, A. R. (2003). *The Babylonian Gilgamesh Epic: Introduction, Critical Edition and Cuneiform Texts*. Oxford University Press.
- [7] PBS NewsHour. (2014, 1월 24일). New evidence suggests Noah's ark may have been round. pbs.org. <https://www.pbs.org/newshour/world/new-evidence-suggests-noahs-ark-may-have-been-round> (2026年9月3日檢索)
- [8] Finkel, I. (2014). *The Ark Before Noah: Decoding the Story of the Flood*. Hodder & Stoughton.
- [9] Centre for the Study of Manuscript Cultures. (2026, 6월 14일). Can artificial intelligence decipher cuneiform tablets? Universität Hamburg. <https://www.csmc.uni-hamburg.de/publications/mesopotamia/2026-06-14.html> (2026年9月3日檢索)
- [10] Julius-Maximilians-Universität Würzburg. (2026, 5월 18일). AI reads cuneiform: A milestone for Ancient Near Eastern Studies. EurekAlert!. <https://www.eurekalert.org/news-releases/1128630> (2026年9月3日檢索)
- [11] Wang, Z. (2026). TabletCraft: Bridging a 4,000-year cultural gap with bidirectional Akkadian NMT and cuneiform rendering. arXiv:2608.02609. <https://arxiv.org/abs/2608.02609> (2026年9月3日檢索)

第6章. 古代美索不達米亞收據與商人的私語：AI翻譯

6.1. 案件檔案

刻在泥板上的收據

古代美索不達米亞人在潮濕的泥板上寫字。他們將蘆葦末端削出稜角做成鐵筆。鐵筆是不蘸墨水、只留下壓痕的書寫工具。他們將這支鐵筆壓入泥中後隨即提起[1]。壓出的痕跡成了形似楔子的三角形。這種文字的名字叫楔形文字。學界也同時使用「楔形文字」這個名稱。楔形就是楔子形狀的意思。

泥土乾燥後會像石頭一樣堅硬。與紙張或皮革不同，即使歲月流逝也不會腐爛。即使發生火災也不會化為灰燼，灼熱的高溫反而會把泥土像磚塊一樣燒得十分堅硬。當城市因戰爭被焚毀時，泥板反而變得更加堅固[3]。多虧於此，4千年前的帳簿得以留存至今。

泥板多數只有手掌大小，是可以放在口袋裡隨身攜帶的尺寸。抄寫員揉捏濕泥做成泥板。抄寫員是代人書寫和誦讀的人。寫完字後，便放在陽光下曬乾。特意放進窯裡燒製的情況很少見。如今博物館進行燒製是為了保存而採取的措施[2]。若有需要修改之處，只要抹平表面擦掉即可。

據估計，迄今發掘出的泥板約有50萬塊[4]。地下還埋藏著更多尚未出土的泥板。上面記錄的文字不只有國王的誇耀。有收到幾袋穀物的收據，有借出白銀並約定利息的契約，也有妻子寫給遠行丈夫的信件[6]。

一張收據的篇幅很短。物品名稱與數量排在最前面，接著是付款人與收款人的名字，最後寫上日期與證人的名字。這與今天在商店拿到的收據順序十分相似。這樣的一張文件內容有限，但當數萬張聚集在一起時，情況就完全不同了。哪一年的糧價上漲得以顯現，哪座城市買賣了什麼也看得一清二楚。

絕大多數未被翻譯

問題在於這些文字幾乎無法被閱讀。在出土的泥板中，被翻譯成現代語言的僅是一小部分[5]。其餘的都沉睡在博物館的箱子裡。

能寫字的人並非人人可當。要成為抄寫員，必須在學校學習很長時間。學生們在泥板上反覆書寫符號。寫錯的泥板就加水重新揉捏製作。留存下來的習字板也作為文物被發掘出來。這些習字板向我們展現了那個時代的課堂情景。

識字的人很少，大多數人都是請抄寫員代筆。收到信後，也是由抄寫員代為讀出。信件語句之所以具備鄭重的格式，原因正在於此：說話的人與記錄的人並非同一人。

解讀一塊泥板也需要經過多道步驟。首先要辨認一個個符號；接下來確定符號代表什麼讀音；然後劃分詞彙並查找含義；最後串聯句子進行翻譯。如果是熟悉的格式，很快就能完成；但若是陌生的文件，就必須耗費漫長的時間細細推敲。

原因在於能夠解讀的人力不足。研究楔形文字的學問稱為亞述學。培養一名學者需要耗費數十年，因為必須同時掌握古代語言與古代歷史。因此，全世界能讀懂這種文字的人極其稀少[5]。

新出土泥板的數量遠遠超過了學者解讀的速度。研究人力沒有增加，裝著泥板的箱子只管不斷堆積。2026年6月，英國約克大學啟動了一項新的研究專案。專案名稱為「邁向楔形文字遺產的普及化接觸」。這項計畫打算將楔形文字資料也翻譯成阿拉伯文，讓任何人都能查閱。該專案將持續3年至2029年[7]。

將資料彙整於一處的工作也在同步推展。「楔形文字數位圖書館計畫」正是這樣的專案。它匯集散落於全球的泥板資訊並在網路上公開，將照片、音譯與翻譯並列呈現。所謂音譯，是用拉丁字母記錄楔形符號讀音的文字。至今已有超過40萬件文物被收錄進目錄[4]。唯有資料匯集起來，機器才能獲得學習的素材。

泥板目前分散在多個國家。伊拉克和土耳其的博物館藏有大量泥板；倫敦、巴黎和柏林的博物館也占有相當大的比重；美國和日本的大學裡也有收藏。如果一份文件的正反面分處不同國家，就很難拼合對齊。只有把照片同時上傳到網路上，才能相互對照拼合。

被國王記錄排擠的帳簿

日常文件被排到後面，有著由來已久的原因。19世紀和20世紀初的發掘隊最先挖掘的是王宮與神廟。巨大的石碑與黃金飾品吸引了人們的目光。學者們也優先解讀王室銘文與法典，把時間先花在翻譯神話與史詩上。

在此期間，商人的收據和信件被排到了後面。糧食配給清單被視為瑣碎的差事記錄。單獨拆開來看，內容顯得無足輕重。學者們便將親手抄錄這些枯燥文件的工作往後推延。就這樣，普通人的聲音被長期埋沒了。

這種偏頗讓我們失去了許多東西：物價如何波動變得難以知曉，女性從事何種工作被掩蓋了，孩童與僕從的名字僅殘留在清單之中。國王的銘文記錄的是誇耀，收據卻不會誇耀，它只記錄交易的事實。因此，收據往往成了更誠實的史料。

一個楔形符號具有多種含義

即便想借助機器來解讀，楔形文字也絕非易事。因為單一符號往往身兼數職[8]。在某些位置上，它是表音的音節符號；在另一些位置上，它直接代表整個詞彙的意義；而在某些位置上，它甚至不發音，僅作為標記，用來提示後面的詞彙是樹木還是神祇。這種標記被稱為限定詞。

楔形文字並非只記錄單一語言的文字。蘇美語與阿卡德語共用相同的文字；後來西臺語與埃蘭語也借用了這種文字。同一個符號在不同語言中發音各異。閱讀者首先必須確定該文件使用的是何種語言。

代表樹木的蘇美語符號「GiŠ」（吉什）就是一個很好的例子。若附在家具名稱前面，它是表示材料為木頭的標記；若單獨出現，它本身就是代表樹木的詞彙；而在其他位置上，它僅表示「giš」這個讀音。閱讀者必須觀察前後文句來判斷其扮演的角色。

數字的記錄方式也與現今不同。美索不達米亞人以60為一個計數單位。時鐘上的60分與60秒正是源於此處。收據上的數量也是以這種方式記錄的。穀物與白銀使用不同的單位；即使是同一個符號，在代表穀物與代表白銀時數值也有所不同。機器必須把這種單位的差異也一併學會。

連寫詞彙也是一大障礙。古代抄寫員在詞與詞之間不空格，電腦很難判斷詞語從何處斷開[8]。泥板破碎磨損也是問題；楔形壓痕一旦變淺，即使是人眼也很難分辨。阿卡德語詞形變化繁複也是一大難題：一個詞根後面可以附加多個詞尾；根據是主詞還是受詞，字尾讀音會有所不同；隨著時態與人稱的不同，又會再次變化。必須先查明在字典中檢索的基本型。

泥板破損的原因也有很多：在地下受到擠壓並受潮；發掘時邊緣可能會剝落；早期的文物交易導致破碎的泥板散落各國；出土後鹽分滲出甚至會遮蓋字跡[2]。文字消失的空白處被學者們稱為缺損處。

正因這層層壁壘，數十萬張帳簿一直以來始終保持著沉默。

參考資料

- [1] Truman State University. (2026). How it works. Truman State University Cuneiforms. <https://exhibits.truman.edu/cuneiform/how-it-works/> (2026年9月3日檢索)
- [2] Guinan, A., Oller, G., & Ormsby, D. (1976). Nippur rebaked: The conservation of cuneiform tablets. Expedition, 18(3). <https://www.penn.museum/sites/expedition/nippur-rebaked/> (2026年9月3日檢索)
- [3] Bouscaren, D. (2023). Assyrian women of letters. Archaeology, 76(6). <https://archaeology.org/issues/november-december-2023/features/assyrian-women-of-letters/> (2026年9月3日檢索)
- [4] Cuneiform Digital Library Initiative. (2026). About the CDLI. cdli.earth. <https://cdli.earth/about> (2026年9月3日檢索)
- [5] Gutherz, G., Gordin, S., Sáenz, L., Levy, O., & Berant, J. (2023). Translating Akkadian to English with neural machine translation. PNAS Nexus, 2(5), pgad096. <https://doi.org/10.1093/pnasnexus/pgad096>
- [6] Michel, C. (2020). Women of Assur and Kanesh: Texts from the Archives of Assyrian Merchants (Writings from the Ancient World 42). SBL Press.
- [7] Cuneiform Digital Library Initiative. (2026). Research Programme "Towards Universal Access to Cuneiform Heritage" 2026-2029. cdli.earth. <https://cdli.earth/postings/239> (2026年9月3日檢索)

[8] Gordin, S., Gutherz, G., Elazary, A., Romach, A., Jiménez, E., Berant, J., & Cohen, Y. (2020). Reading Akkadian cuneiform using natural language processing. PLoS ONE, 15(10), e0240511. <https://doi.org/10.1371/journal.pone.0240511>

6.2. 科學搜查隊的追蹤：阿卡德語AI翻譯

開闢兩條翻譯途徑

2023年，特拉維夫大學的研究團隊率先開拓了道路。蓋伊·古特爾茨（Gai Gutherz）與其同事研發了一款將阿卡德語翻譯為英語的翻譯器[1]。阿卡德語是美索不達米亞長期使用的語言，以楔形文字書寫，如今已無人使用。

研究團隊開闢了兩條翻譯途徑。第一條途徑名為C2E，是將轉寫為電腦文字的楔形符號輸入以獲取英語的途徑。這並不是直接辨識泥板照片的方式；由於不需要經過學者的音譯，步驟較為精簡。第二條途徑名為T2E，是將學者以拉丁字母轉寫楔形符號讀音的文本輸入的途徑。這種以拉丁字母轉寫的文本稱為音譯。在音譯階段，符號的含義已被梳理過一次，因此T2E方面的翻譯更加精準[1]。

製作音譯文本並非易事。學者看著符號確定讀音；若有多種讀音，則根據上下文來選擇。這種選擇本身就是一種判斷。經過一次判斷處理後的文本，機器更容易處理。這正是T2E得分較高的原因所在。

翻譯器的內部結構分為兩部分：前端將原文轉化為數字陣列，後端則根據這些數字逐一挑選英語單詞。前端稱為編碼器，後端稱為解碼器。編碼器同時容納詞彙的含義與順序；解碼器則在參考先前挑選詞彙的同時決定下一個詞彙。

直接從照片中辨識符號的研究也在另行推進。2026年6月，電子巴比倫圖書館的研究團隊發表了研究成果。機器掃描了8萬7千多張破損泥板的照片，識別出的符號痕跡接近290萬個。檢測表現比以往的方法最高提升了37%。但對於破損嚴重以及行排列紊亂的泥板，目前依然較為薄弱[6]。

以8千份文獻學習的機器

機器藉由觀察例句來學習翻譯。研究團隊從名為ORACC的公開資料庫中蒐集文獻。ORACC是一個整理並公開楔形文字文獻的網路資料庫。團隊從中蒐集的文獻共計8,056篇，句子數量超過5萬句[1]。

蒐集而來的文獻種類繁多：刻有國王功績的銘文約有3千篇，行政信件約有2千篇，此外還包含審判記錄與占星報告[1]。每份文獻旁都附有學者製作的英語譯文。機器反覆研讀這些成對的文本數萬次，自主找出了其中的規律。

機器學習的過程與準備考試十分相似：首先同時看題目與答案；接著遮住答案自己試著解答；若答錯了，便反思究竟錯在哪裡。如此反覆檢視數百萬次。與人類不同之處，只在於速度與耐力。

機器並非將詞彙作為整體來處理，而是將詞彙拆解為更小的片段來處理。這些片段被稱為詞元（Token）。即使是罕見的詞彙，只要拆解成片段就能處理。對於附加大量詞尾的阿卡德語而言，這種方法非常合適。

翻譯品質是透過名為BLEU的評分標準來衡量的。這是一種計算機器翻譯與人工翻譯重疊程度的方法，以重疊詞組的比例來計分，介於0分到100分之間。一般認為超過30分就代表語意通順，但這並非絕對標準，其解讀會隨文獻類型與句子長度而有所不同。

研究團隊選擇了名為卷積神經網路的架構。這是一種將相鄰資訊結合在一起學習的架構，因為在資料量較少時學習速度快且穩定。驗證結果顯示，C2E獲得了36.52分；T2E得分更高，達37.47分。以往使用的自動檢索方式僅獲得27.09分與23.51分。在句子短小且格式固定的文獻中，該模型的表現尤為優異[1]。

劃分資料的方法也至關重要。必須將用於學習的部分與用於測試的部分分開留存。若事先給機器看測試文獻，所得的分數便毫無意義，這就像人類考試前預先知道了考題一樣。研究團隊嚴格遵守了這種劃分原則來評定分數。

即便是同一篇原文，不同人的翻譯也會有所差異。因此不能單憑分數來評定優劣，除了分數之外，還需結合人工評估。將譯文展示給學者並讓其評定等級，分別詢問語意是否通順，以及是否存在錯誤之處。學者認為堪用的翻譯，在T2E中約占44%，在C2E中約占38%[1]。當評分與人工評估出現分歧時，會以人類的判斷為準。

阿卡德語翻譯屬於低資源任務。所謂低資源，意味著可供學習的例句極少。英語和法語翻譯器可以學習數千萬個句子，而阿卡德語擁有的對照句對甚至不及那個數目的千分之一。當資料匱乏時，大型模型反而容易迷失方向。2025年還舉辦了多個研究團隊同台競技的共同任務。該任務是尋找詞彙的原型並推測缺失的文字。資料取自電子巴比倫圖書館與ARCHIBAB。利用相同的資料評定成績，哪種方法更優便一目了然[2]。

詞形原型還原是翻譯的基礎前置作業。學界將這項工作稱為詞元還原（Lemmatization）。它是將因附加詞尾而變化的詞彙還原至字典收錄的原型。唯有完成這項前置工作，按詞彙進行檢索與統計才成為可能。

識破看似合理的謊言

機器翻譯有一個危險的壞毛病。那就是在毫無根據的情況下，編造出看似合理的句子。這種毛病被稱為幻覺。研究團隊將幻覺分為兩種情況進行探討。一種是看著原文卻錯誤拼湊文意的情況。另一種則是無視原文，編造其他故事的情況[1]。

殘破的泥板特別容易產生幻覺。這是因為文字消失的空白處，機器會用想像力去填補。因此，翻譯成果務必由學者重新審閱。機器產出的句子只是初稿，而非最終判定。

減少幻覺的方法有很多種。將破損之處明確標示為空格。讓機器在該處回答「不知道」。也會嘗試將同一個句子翻譯多次。若每次的答案不同，便不予採信。學者看著這些標記，決定該重新檢查何處。

將楔形符號轉寫為字母並進行斷詞，是先於翻譯的工作。在同一研究脈絡下，這項工作也試著交由機器處理。名為雙向LSTM的神經網路成績最為突出。準確率最高達到了97%[3]。

雙向開啟的大門

迄今為止的翻譯器都只是單向流動。只負責將古代文獻翻譯成現代英語。2026年5月，美國南加州大學的王兆輝開啟了相反的方向。他所打造的工具名為TabletCraft[4]。

TabletCraft也能將英語翻譯成阿卡德語。它會將翻譯結果繪製成楔形符號，甚至生成泥板圖像。訓練時使用了11萬6千對雙向句子。符號對應表則具備1萬4千多個。測試成績方面，阿卡德語翻英語為49.1分。英語翻阿卡德語則拿到了48.5分[4]。

這款工具已被應用於博物館體驗設備與大學課堂中。開發者也明確指出了其局限。絕不能將機器編造的阿卡德語當作真正的古代文獻來引用。該工具擅長處理的文本是新亞述時期的王室碑文[5]。

雙向翻譯具有實用價值。學生可以嘗試將自己的句子轉換為楔形符號。親手描摹符號，便能看清文字的結構。在博物館裡，參觀者可以試著刻下自己的名字。古代文字遂從單純的展品變成了可以親手操作的工具。

上述研究都將資料與程式一同公開[1][5]。公開後，其他團隊便能重複進行相同的測試。若有錯誤之處也能迅速顯露出來。古代語言研究的參與人數較少。公開與否極大程度地左右了發展速度。

使用翻譯器的方法也正在逐步確立。學者將機器的初稿呈現在螢幕上。將原文與譯文並排比對。在存疑的詞彙上做標記。累積的標記便成為下一次學習的資料。

參考資料

- [1] Gutherz, G., Gordin, S., Sáenz, L., Levy, O., & Berant, J. (2023). Translating Akkadian to English with neural machine translation. PNAS Nexus, 2(5), pgad096. <https://doi.org/10.1093/pnasnexus/pgad096>
- [2] Gordin, S., Sahala, A., Spencer, S., & Klein, S. (2025). EvaCun 2025 shared task: Lemmatization and token prediction in Akkadian and Sumerian using LLMs. In Proceedings of the Second Workshop on Ancient Language Processing (pp. 242-250). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2025.alp-1.33>
- [3] Gordin, S., Gutherz, G., Elazary, A., Romach, A., Jiménez, E., Berant, J., & Cohen, Y. (2020). Reading Akkadian cuneiform using natural language processing. PLoS ONE, 15(10), e0240511. <https://doi.org/10.1371/journal.pone.0240511>
- [4] Wang, Z. (2026, 5월 14일). TabletCraft: Bridging a 4,000-year cultural gap with bidirectional Akkadian NMT and cuneiform rendering. arXiv:2608.02609 (2026년 5월 14일 제출, 2026년 8월 공개). <https://arxiv.org/abs/2608.02609> (2026년 9월 3일檢索)

- [5] Buyukyildirim, K. (2026, 8월 11일). TabletCraft opens a two-way AI translation channel for archaeologists. Arkeonews. <https://arkeonews.net/tabletcraft-opens-a-two-way-ai-translation-channel-for-archaeologists/> (2026年9月3日檢索)
- [6] Che, W., Garcés Arias, E., Niaz, A., Bender, A., & Jimenez, E. (2026, 6월 21일). Automated sign detection across the Electronic Babylonian Library: A large-scale dataset and end-to-end cuneiform OCR pipeline. arXiv:2606.22608. <https://arxiv.org/abs/2606.22608> (2026年9月3日檢索)

6.3. 揭曉的真相

卡尼什商人的信件

土耳其中部開塞利附近有一座名為屈爾特佩的山丘。4千年前這裡的名字叫卡尼什。亞述商人在此建立了貿易據點。商人們的據點社群被稱為卡魯姆。卡魯姆原本是指河畔碼頭的詞彙。卸貨的地點變成了商人們的聚落區域。

卡尼什的卡魯姆是亞述商人的自治組織。他們自行制定規則並裁決爭端。向當地國王繳納稅款以獲得安全保障。貨物馱在驢背上運送。從阿淑爾到卡尼什是需要耗費數週的漫長路程。

自1948年展開的發掘中，大批泥板相繼出土。迄今出土的文獻約有2萬3千5百塊。這是主持發掘工作的菲克里·庫拉科烏盧所公佈的數字[1]。內容並非國王的自誇功績。而是借據與貨物託運單。託運單即所裝運貨物的清單目錄。此外還有審判紀錄與家書。

信件中原原本本地保留了催促與抱怨。詢問何時匯款。責備貨物送達太遲。附上保重身體的問候。也叮囑兒子要好好用功。4千年前人們的牽掛與擔憂，就這樣如實呈现在字裡行間。

這些文獻揭示了古代經濟的輪廓架構。商人們從阿淑爾運來紡織品與錫。在安納托利亞將其換成白銀與黃金。其中也記載了該如何繳納稅款與通行費。發生紛爭時如何判決亦有記錄留存。

錫是製造青銅不可或缺的金屬。青銅是由銅與錫混合製成的。在安納托利亞很難獲取錫。錫必須從遙遠的東方運來[8]。這種匱乏正是開拓漫長貿易路線的驅動力。

留守在阿淑爾的女性所承擔的職責也得以明朗。商人的妻子與女兒在家中織造布料並寄出[2][8]。她們調節產量並討價還價。在寄給丈夫的信中還會計較市價行情。賽西兒·米歇爾收集了300多篇此類文獻並翻譯成英語。米歇爾任職於法國國家科學研究中心與漢堡大學[2]。

不拆開信封閱讀

古代商人會將重要的泥板裝進泥信封中密封。若要閱讀內部內容，就必須打破信封。然而信封上也蓋有印章印記。若打碎信封，該資訊就會隨之消失。

信封是防止偽造的裝置。先在內側的泥板上寫下契約內容。再於其上覆蓋一層泥土製成信封。接著在信封表面重寫相同的內容。見證人們滾動印章加蓋印記。若外表與內部不一致，篡改塗改的事實便會曝光。

人們需要一種在完整保留信封的情況下讀取內部內容的方法。2026年6月，開塞利考古博物館給出了答案。他們在不拆開信封的情況下，成功讀取了一塊泥板的內容。這是與法國國家科學研究中心及漢堡大學合作完成的工作。他們採用了不切開物體即可透視內部的CT掃描等方法。裡面記載的內容是小麥與大麥的交易約定。庫拉科烏盧表示，這項攝影技術連一顆大麥粒都能清晰呈現[1]。如此取得的文字圖像成為翻譯器的輸入資料。能讀取的文獻增加了，機器也才能獲得學習的素材。

獎金5萬美元的翻譯大賽

翻譯大賽也應運而生。名為Deep Past Initiative的機構籌辦了這場比賽[3]。大賽在人工智慧競賽網站Kaggle上舉行[4]。任務是將卡尼什的商業文獻翻譯成英語[3]。比賽於2025年12月16日開跑。參賽截止日為2026年3月23日。獎金為5萬美元。營運費用則由一家金融公司贊助[7]。

大賽資料向大眾公開，供任何人下載[4]。公開資料在比賽結束後依然留存。後續研究者得以在這些資料的基礎上繼續研究。這形成了一場比賽留下一份珍貴資料的結構。

屈爾特佩文獻中也有許多尚未翻譯的篇章[3]。發掘工作至今仍每年持續進行[1]。每當新的泥板出土，便會被增補到清單中。若能解讀的人手沒有增加，清單只會越來越長。這正是需要機器協助的原因所在。

相同的情況並不僅止於單一遺址。世界各地的博物館中堆積著大量類似的文獻。那些地方同樣堆滿了尚未解讀的藏箱。格式固定的文獻隨處可見。穀物帳簿與稅務登記冊便是如此。在卡尼什行得通的方法，也能應用於其他文獻庫。

大賽簡介重點介紹了文獻中出現的真實人物。像是阿淑爾納達與阿淑爾伊迪等商人的名字。伊什塔爾拉馬西與恩南阿淑爾的名字也名列其中[3]。4千年前一個商人家族的名字成為了大賽的主題。

競賽模式自有其好處。不再僅限於少數研究室，全世界的開發者都能投身於同一個問題。用同一把尺衡量成績，哪種方法更優秀便一目了然。獎金也成為聚集人才的動力。

這批文獻資料非常適合機器翻譯。格式固定且句子簡短。相同的表達方式在多篇文獻中反覆出現。機器很擅長找出重複的規則。學者先前已經翻譯好的文獻也不在少數。那些譯文便成了機器的教材。

機器做不到的事

儘管成果豐碩，眼前的高牆依然分明。第一道高牆是資料的匱乏。現代語言的翻譯器是汲取數億個句子成長的。阿卡德語成對的對照句子卻僅有數萬個[5]。資料一旦不足，在罕見的地名與人名上就很容易出錯。

第二道高牆是符號的多重含義。同一個符號究竟代表表音、單詞還是標記，令人混淆[6]。前後上下文若太短，機器也會無所適從。

第三道高牆是機器的判斷模式。機器只會計算出現哪個詞彙的機率較高。它無法判斷該句子在歷史上是否合情合理。因此，學者的檢驗在最後階段必不可少[5]。當機器與人類結為搭檔時，翻譯才真正值得信賴。

分工合作的模式正在扎根。機器迅速瀏覽數萬塊泥板並梳理出大致意思。學者從中挑選出重要的文獻。只針對挑出的文獻親手細緻地重新研讀。這樣一來，得以閱讀的文獻數量便大幅增加。

譯文需要標籤。必須註記哪一段是由機器產出。由人工修正之處也要予以標明。如此一來，後人才能追溯依據。歷史研究正是始於追溯依據。

在翻譯中，專有名稱是個棘手的部分。人名與地名無法照字面意譯。必須依照讀音如實音譯。然而同一個名字在不同文獻中的寫法可能各異。機器很容易將初次見到的名字置換為風馬牛不相及的單詞。因此需要另外整理專有名稱清單供機器參考。

儘管如此，方向依然明確。以往一名學者畢生只能閱讀數百塊泥板。如今機器先產出初稿翻譯。學者則將時間花在潤飾初稿上。在木箱中沉睡了4千年的收據正等待著輪候。不再只是國王，普通百姓的心聲終於開始被世人聽見。

卡尼什的商人們未曾想過自己的信件能長久留存。他們只不過是記錄了那一天的算帳帳目。那些計算如今卻呈現了一整個時代的經濟面貌。最終流傳下來的，比起炫耀自誇，更是切切實實的記錄。

參考資料

- [1] Türkiye Today. (2026, 6월 13일). 4,000-year-old clay tablet sealed in mud envelope goes on display in central Türkiye. Türkiye Today.
<https://www.turkiyetoday.com/culture/4000-year-old-clay-tablet-sealed-in-mud-envelope-goes-on-display-in-central-turkiye-3221841> (2026年9月3日檢索)
- [2] Michel, C. (2020). Women of Assur and Kanesh: Texts from the Archives of Assyrian Merchants (Writings from the Ancient World 42). SBL Press.
- [3] Deep Past Initiative. (2026). Deep Past Challenge: Introduction. deeppast.org.
<https://www.deeppast.org/challenge/intro/> (2026年9月3日檢索)
- [4] Kaggle. (2026). Deep Past Challenge: Translate Akkadian to English. kaggle.com.
<https://www.kaggle.com/competitions/deep-past-initiative-machine-translation> (2026年9月3日檢索)

- [5] Gutherz, G., Gordin, S., Sáenz, L., Levy, O., & Berant, J. (2023). Translating Akkadian to English with neural machine translation. PNAS Nexus, 2(5), pgad096. <https://doi.org/10.1093/pnasnexus/pgad096>
- [6] Gordin, S., Gutherz, G., Elazary, A., Romach, A., Jiménez, E., Berant, J., & Cohen, Y. (2020). Reading Akkadian cuneiform using natural language processing. PLoS ONE, 15(10), e0240511. <https://doi.org/10.1371/journal.pone.0240511>
- [7] Deep Past Initiative. (2025, 12월 16일). Deep Past Initiative partners with Kaggle on prize challenge to translate 4,000-year-old tablets. EIN Presswire. <https://www.einpresswire.com/article/875885564/deep-past-initiative-partners-with-kaggle-on-prize-challenge-to-translate-4-000-year-old-tablets> (2026年9月3日檢索)
- [8] Bouscaren, D. (2023). Assyrian women of letters. Archaeology, 76(6). <https://archaeology.org/issues/november-december-2023/features/assyrian-women-of-letters/> (2026年9月3日檢索)

第7章. 沉睡在龜腹甲中的帝國未來：甲骨文與3D虛擬綴合

7.1. 案件檔案

在藥房發現的三千年前文字

1899年，清朝學者王懿榮認出了一種古老文字。關於這件事，流傳著一個由來已久的故事。說的是王懿榮生了病，配得了一味名為龍骨的藥材。龍骨是將古老獸骨磨碎製成的藥粉。古人相信這種骨頭是龍的骨頭。據說他拿到的骨頭當時尚未磨成粉。他在那表面看見了陌生的紋樣[1]。

這個故事並未被證實為史實。中國學界認為沒有支持這一說法的紀錄。古文字學家李學勤在2007年的文章中否認了這個故事。理由是作為藥材出售的龍骨早已被搗碎，不可能看見文字。實際上是古董商收購並運送了出自安陽的骨頭。學界認為是那些骨頭流轉到金石學家王懿榮手中並接受了鑑定[2]。

王懿榮於隔年離世，未能找出骨頭的出處。他的朋友劉鶚在1903年出版了第一本收錄這些文字的書籍。1908年，學者羅振玉查明了骨頭的出處[1]。原來那些骨頭一直出自河南省安陽的田地。

安陽的那片土地正是殷墟。殷墟是商朝最後的都城遺址。商朝自西元前13世紀至西元前11世紀統治此地[3]。正式的發掘工作始於1928年。在此之前，農民挖出的骨頭都被當作藥材賣出。甚至有人會刮掉刻有文字的一面再行出售[1]。歷經三千年留存的紀錄差點化為藥材而消失殆盡。

君王向神明問卜的國度

刻在殷墟骨頭上的文字被稱為甲骨文。在東亞，目前尚未發現比甲骨文更早的文字紀錄[3]。

商朝人在面臨重大事情時會叩問神明。他們會問是否會下雨，會問農事是否會順利，會問是否可以出兵征戰，也會問王妃分娩的日期。像這樣叩問神意的儀式稱為占卜。占卜是王室直接管理的國家事務[3]。

商朝的王妃婦好也是一位率領軍隊的司令官。婦好的名字在甲骨文紀錄中出現過多次。1976年，殷墟完整發掘了婦好墓[4]。墓中的青銅器與玉器展現了她的地位。甲骨文就是這樣與陵墓一同揭示了古人的生活。

甲骨文上記錄的問卜並不僅限於國家大事。其中也有為君王的牙痛擔憂的問卜，也有詢問是否可以出門打獵的提問。因此，這些骨頭成了呈現商朝人日常生活的紀錄[3]。

以火灼烤的骨頭，以刀刻下的文字

商朝人在占卜中使用兩種材料。一種是烏龜扁平的腹甲，另一種是牛寬闊的肩胛骨。這兩種材料合稱為甲骨[3]。

準備過程非常嚴謹。史官們將骨頭表面打磨平滑，接著在背面排成一列鑿出小凹槽。隨後將燒得熾熱的棒子接觸鑿出的凹槽。受熱的骨頭表面便會劈啪裂開。君王與占卜官觀察裂紋的方向與角度，並將那裂紋解讀為神明的回答[3]。

解讀結束後，史官們便拿起刀。他們將占卜的過程一筆一劃刻在骨頭表面。這些刻下的文字正是甲骨文。

記錄遵循著固定的順序。首先記錄占卜的日期與負責人姓名，接著記錄向神明提出的問題，然後記錄君王作出的吉凶判斷，最後記錄實際發生的事情。這是依序填寫四個部分的方式。只不過四個部分都完整保留下來的紀錄並不多見。有時因骨頭碎裂而缺失了一兩處，有時則是史官一開始就留白[3]。即便如此，記錄的順序本身是一致的。

這種規律性在日後提供了極大的幫助。因為格式固定時，便能縮小缺漏處的範圍。前面的欄位若有日期，後面的欄位就會是問題。電腦也能學習這種順序來推測缺失的文字[3]。

記錄日期的方法也是固定的。商朝將十個與十二個符號配對來計日。這種配對每六十天循環一次。骨頭上即使只剩下兩個日期的符號，也能推算其順序。縮減殘斷紀錄所屬年代的線索便由此而來[3]。

碎裂成15萬片的紀錄

商朝滅亡後，殷墟被埋入泥土之中。在漫長的歲月裡，骨頭不可能完好無損。泥土的重量壓碎了脆弱的腹甲，水分與微生物侵蝕著骨質紋理。原本完整的一大片腹甲分裂成了數十塊碎片。

迄今為止為世人所知的甲骨殘片約有15萬片[5]。這些文物並未聚集在一處。在19世紀末被發現後，相當多的甲骨流向了古董市場[1]。因此，源自同一塊骨頭的殘片被分散保存在不同國家的博物館中[5]。有的文物位於中國安陽與北京，也有文物由歐洲與韓國的機構收藏[6]。

將散落的碎片拼接復原為原本完整狀態的工作稱為綴合。長久以來，綴合僅能依賴人類的肉眼。學者們拿著放大鏡仔細審視數萬張黑白拓本。拓本是將紙敷在骨頭上，用墨輕捶拓印出來的複本。這就像在硬幣上覆蓋紙張，用鉛筆塗抹拓出紋樣一樣。學者們比對骨頭的紋理與斷裂邊緣的曲度，並逐一觀察殘斷筆畫的末端是否能銜接相連[5]。

這種方式緩慢而艱辛。有時為了找到一對碎片就要耗費數年時間。未能找到配對的碎片只能原封不動地留在箱子裡。如果沒有配對，文句也只能殘缺不全。只留下問題而結果消失的紀錄堆積如山。

若估算這項工作的規模，其難度便更加顯而易見。若有15萬片殘片，從中選取兩片進行比對的組合數就超過100億對。即使一個人一天比對一百對，終其一生也看不完。更何況這些殘片還分散在多個國家[5]。很多時候甚至連將兩片殘片並列對照的機會都沒有。

3千字至今仍沉默不語

即使將殘片全部拼合，問題依然存在。因為文字本身無法被解讀。在所有的甲骨文中，已確認的不同字形約有4,500字。然而學界對含義達成共識的字僅約1,500字左右。剩下的約3,000字至今仍不知其意[7]。

難以解讀的原因有很多。刀痕磨損導致輪廓模糊；骨頭碎裂使得文字的一半消失；每位史官握刀的習慣不同，刻出同一個字的方式也各異。因此產生了許多意義相同但形體各異的字。這種字被稱為異體字[3]。現代的漢字字典無法解釋這些三千年前的形體。

這片沉默也在商朝歷史中留下了空白。若讀不出祭祀的名稱，就無法計算祭祀的種類；若讀不出地名，就無法描繪國家的疆界；若讀不出人名，就無從知曉是誰負責了何事。一個文字的沉默，便連結著歷史上一幕場景的沉默。

不知其意的3千個字，以及失去配對的數萬件殘片依然留存。為打破這片沉默，人工智慧走進了殷墟。

參考資料

- [1] Mark, E. (2016, 2월 26일). Oracle bones. World History Encyclopedia. https://www.worldhistory.org/Oracle_Bones/ (2026年9月3日檢索)
- [2] 姚小鸥. (2020, 7월 16일). 谁是甲骨文的最早发现者. 澎湃新闻. https://www.thepaper.cn/newsDetail_forward_8297008 (2026年9月3日檢索)
- [3] Chen, Z., Hua, W., Li, J., Zhu, Y., Zhi, X., Liu, Z., Chen, T., Zhang, W., & Zhai, G. (2026). Oracle bone inscriptions information processing: a comprehensive survey. npj Heritage Science, 14, 220. <https://doi.org/10.1038/s40494-026-02511-w> (2026年9月3日檢索)
- [4] Smarthistory. War and sacrifice: The tomb of Fu Hao. Smarthistory. <https://smarthistory.org/tomb-of-fu-hao/> (2026年9月3日檢索)
- [5] Zhang, C., Wang, B., Chen, K., Zong, R., Mo, B., Men, Y., Almpandis, G., Chen, S., & Zhang, X. (2022). Data-driven oracle bone rejoining: A dataset and practical self-supervised learning scheme. In Proceedings of the 28th ACM SIGKDD Conference on Knowledge Discovery and Data Mining (pp. 4482-4492). ACM. <https://doi.org/10.1145/3534678.3539050>
- [6] 梁文艳. (2026, 5월 26일). 以科技赋能，让甲骨文"活"起来. 中国青年报. https://zqb.cyol.com/pc/content/202605/26/content_426168.html (2026年9月3日檢索)
- [7] Liu, Y., Guan, H., Wang, P., Wang, X., Wan, J., Zhang, K., Zheng, H., Liu, X., Kuang, Z., Yang, H., Li, B., Liu, Y., Jin, L., & Bai, X. (2026). AlphaOracle: Oracle bone script decipherment via human-workflow-inspired deep learning. The Innovation, 7(11), 101462. <https://doi.org/10.1016/j.xinn.2026.101462>

7.2. 鑑識調查隊的追蹤：甲骨專用多模態模型與3D虛擬綴合

將散落的資料匯集於同一畫面

調查始於蒐集資料，甲骨研究亦是如此。安陽師範學院與中國社會科學院的研究團隊於2019年推出了名為「殷契文淵」的資料平台。殷契文淵將甲骨拓本、文字資訊與學術論文彙整於一處，並免費公開[1]。

該平台不僅僅上傳照片，還將甲骨殘片進行高解析度拍攝，並製作成3D模型上傳。3D模型是從多個方向掃描物件，在電腦中以立體形式重現的資料。旋轉螢幕畫面，即可檢視骨頭的正面、背面以及斷裂的側面。研究團隊以公釐級的精度構建這些模型，即使是淺淺刻下的刀痕深度也記錄在這些資料中[1]。

找回流散至海外之甲骨的工作也與該平台接軌。研究團隊選擇引進3D資料來代替實物。他們與德國、韓國、法國等地的收藏機構合作，蒐集了900多件資料[1]。實物仍留在原收藏地，只有其樣貌回到了安陽。

2025年10月，該平台上新增了名為「殷契行知」的工具。這項工具能辨識文字、比對拓本並檢索相關文獻[1]。原本堆積資料的倉庫，搖身變成了能回答問題的助手。

資料累積了多少，也能從數據中得到證實。2026年發表的一篇綜述論文介紹了36種與甲骨相關的公開資料集。其中既有用於文字辨識的資料集，也有用於綴合研究的資料集[2]。

同時使用多種資料的方式稱為多模態。在甲骨研究中，這意味著要同時檢視三種資料：一是墨拓出的黑白拓本，二是色彩與質感鮮明的彩色照片，最後一種則是包含深度資訊的3D掃描資料。拓本筆畫清晰但缺乏深度，3D資料雖包含深度但筆畫顯得模糊。將這三者疊合比對，便能填補彼此的不足[2]。

遵循人類解讀文字的順序

2026年問世的AlphaOracle是在這些資料上運作的人工智慧。該模型由華中科技大學與安陽師範學院等單位共同開發。其設計原則是完全遵循人類學者的工作順序[3]。

人類學者會經歷四個步驟：首先審視拓本並整理筆畫；接著仔細分析字形結構；然後尋找出現同一個字的其他文句；最後運用古代文獻知識推敲含義是否合適。AlphaOracle也相應地分拆為分別負責這四個步驟的組件[3]。

這種架構是有原因的。若機器只是直接給出答案，學者很難予以信任。如果將步驟分開，就能看清答案是在哪一個階段出現了分歧。學者便能退回該階段，與機器對照彼此的見解[3]。

該模型所閱讀的資料量超越了一個人終其一生所能閱覽的規模。它處理了73,883件拓本，同時也處理了45,364件掃描資料與23,755篇研究論文。研究團隊讓86名專家試用該工具，結果分析單一文字所需的時間減少了64%。79%的參與者回答該工具非常有用[3]。

將文字拆解為部件來審視

其他研究團隊則選擇了不將文字整體死記硬背的方法。甲骨文字是由反覆出現的圖形部件組成的[4]。諸如手形、足形、器皿形等部件會分佈在多個不同的字中。只要掌握部件的含義，即使是初次見到的字也能推測其意。

基於這種構想，於2026年4月發表的便是OB-Radix資料集。OB-Radix收錄了934個不同的字，構成這些字的部件圖像達1,853幅，部件種類共有478種。每個部件都附有經過專家核實的說明[4]。

人工智慧遇到生疏的文字時，會先將其拆解為部件，接著推敲這些部件是以何種關係排列配置。兩個部件是上下相疊還是並排排列，往往決定了含義的不同。模型甚至會考量這種結構關係來提出字義建議[4]。

同年6月，名為OracleAnalyzer的模型也隨之問世。該模型以甲骨文資料重新訓練了一個30億參數規模的小型語言模型。研究團隊全新構建並加入了甲骨文推理資料與偏好資料。偏好資料是彙集了人類評選為更佳答案的資料。結果顯示，小型模型的表現遠優於大型模型[5]。這意味著只要精選資料，不必單純以模型規模硬拼。

復原並重新描繪模糊筆畫的模型

也有在解讀之前先使其清晰顯現的研究。2024年發表的OBSD正是這樣的工具。該工具使用了能重新繪製圖像的擴散模型[6]。擴散模型是一種從模糊斑點中逐步生成清晰圖像的方法。

OBSD接收磨損模糊的甲骨文字圖像作為輸入，然後以圖像形式提出該字對應至何種現代漢字的建議。這項研究在2024年國際學術會議上被評選為最佳論文[6]。將文字解讀轉化為圖像生成問題的構想獲得了高度認可。

由電腦拼合碎裂殘片的斷面

與文字解讀並行，殘片拼合的研究也在持續推進。2022年公開的OB-Rejoin是用於綴合研究的資料集，其中收錄了998張真實甲骨拓本。同一個研究團隊還一併推出了名為S3-Net的拼合神經網路[7]。

這個神經網路會計算碎片邊緣的曲折彼此有多相似。若兩個邊緣吻合，得分就會提高。機器會依得分由高到低的順序排列候選對象。同一個研究團隊在2020年也發表過先前的研究。當時前十個候選對象中包含正確答案的比例為98.39%[11]。

在拼拼圖時，人類通常會先看碎片的形狀，而非圖案。機器也採用相同的方法。只不過，機器一天就能比對數十萬次。

機器給出的不是單一正確答案，而是候選名單。名單中最頂端的項目有時也會出錯。因此，研究團隊將前十個候選中是否包含正確答案作為性能評估基準[7]。人類只需確認十個項目的情況，與必須確認十萬個項目的情況截然不同。

首都師範大學甲骨文研究中心與河南大學合作開發了綴合程式。其名為「綴多多」，自2020年起投入使用。該程式能一次比對18萬張拓本來尋找匹配的碎片。研究中心以這種方式整理了約1,500筆甲骨碎片紀錄。過去要找到一片碎片的配對，有時需要耗費八到十年。如今只需幾天就能得出候選結果[8]。

邁向句子單位的考驗

辨識單一文字與理解整段句子是兩回事。2026年6月底公開的S-OBI，正是用來衡量這項差異的測試基準。研究團隊從專家確認過的95件原始拓本中抽取題目。題目共有695道，且設計為以句子為單位作答。接著，他們讓多個大型人工智慧解答相同的題目。測試結果顯示，各模型之間的成績差異相當明顯。同時也顯示出，要掌握整個句子的脈絡目前仍然十分困難[9]。

這項測試還揭示了另一項事實：一般人工智慧並不擅長解答甲骨文問題[9]。即使是廣泛學習了世上各種文字的模型，在面對三千年前的文句時也會動搖。這是因為那些文字與語法是它們從未學過的。因此，針對甲骨文專門重新訓練的研究正在持續開展[5]。

突破破碎斷面以連結上下文的研究也隨之出現。有一支研究團隊提出了一項新任務，旨在判定兩個句子原本是否刻在同一塊骨頭上。他們公開了名為OBID-ACR的資料集，並提出了結合字形資訊與句意資訊的方法[10]。這為在實際拼合碎片之前先連結文意脈絡開闢了道路。

參考資料

- [1] 梁文艳. (2026, 5월 26일). 以科技赋能，让甲骨文"活"起来. 中国青年报. https://zqb.cyol.com/pc/content/202605/26/content_426168.html (2026年9月3日檢索)
- [2] Chen, Z., Hua, W., Li, J., Zhu, Y., Zhi, X., Liu, Z., Chen, T., Zhang, W., & Zhai, G. (2026). Oracle bone inscriptions information processing: a comprehensive survey. npj Heritage Science, 14, 220. <https://doi.org/10.1038/s40494-026-02511-w> (2026年9月3日檢索)
- [3] Liu, Y., Guan, H., Wang, P., Wang, X., Wan, J., Zhang, K., Zheng, H., Liu, X., Kuang, Z., Yang, H., Li, B., Liu, Y., Jin, L., & Bai, X. (2026). AlphaOracle: Oracle bone script decipherment via human-workflow-inspired deep learning. The Innovation, 7(11), 101462. <https://doi.org/10.1016/j.xinn.2026.101462>
- [4] Zhang, J., Li, R., Pang, H., Xia, D., Zhu, Z., Zhang, Q., Li, C., & Yang, X. (2026, 4월 8일). Specializing large models for oracle bone script interpretation via component-grounded multimodal knowledge augmentation. arXiv:2604.06711. <https://arxiv.org/abs/2604.06711> (2026年9月3日檢索)
- [5] Song, Z., Wang, Y., Ma, Z., Yu, Z., Wang, T., Zhang, J., Wang, T., & Yu, K. (2026, 6월 24일). OracleAnalyser: Analysing implicit semantics of oracle bone scripts through MLLMs with post-training. arXiv:2606.25906. <https://arxiv.org/abs/2606.25906> (2026年9月3日檢索)

- [6] Guan, H., Yang, H., Wang, X., Han, S., Liu, Y., Jin, L., Bai, X., & Liu, Y. (2024). Deciphering oracle bone language with diffusion models. In Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers). Association for Computational Linguistics. <https://aclanthology.org/2024.acl-long.831/> (2026年9月3日檢索)
- [7] Zhang, C., Wang, B., Chen, K., Zong, R., Mo, B., Men, Y., Almpandis, G., Chen, S., & Zhang, X. (2022). Data-driven oracle bone rejoining: A dataset and practical self-supervised learning scheme. In Proceedings of the 28th ACM SIGKDD Conference on Knowledge Discovery and Data Mining (pp. 4482-4492). ACM. <https://doi.org/10.1145/3534678.3539050>
- [8] 北京日報. (2026, 8月 28日). 高校團隊"喚醒"三千年前文明碎片. 新浪教育. <https://edu.sina.com.cn/l/2026-08-28/doc-inipwhrw0020843.shtml> (2026年9月3日檢索)
- [9] Li, Z., Chen, Z., Chen, T., & Zhai, G. (2026, 6月 30日). Beyond single character: Evaluating MLLMs for sentence-level oracle bone inscription understanding. arXiv:2606.31169. <https://arxiv.org/abs/2606.31169> (2026年9月3日檢索)
- [10] Zhang, H., Wang, T., Zhang, Z., Li, B., Sun, H., Hou, C., Wang, N., Yu, Y., Jiao, Q., Xiong, J., & Liu, Y. (2026). A multi-modal dataset and method for bone-level association prediction in oracle bone inscriptions. npj Heritage Science, 14, 19. <https://doi.org/10.1038/s40494-025-02282-w> (2026年9月3日檢索)
- [11] Zhang, C., Zong, R., Cao, S., Men, Y., & Mo, B. (2020). AI-powered oracle bone inscriptions recognition and fragments rejoining. In Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence (IJCAI-20), 5309-5311. <https://doi.org/10.24963/ijcai.2020/779>

7.3. 揭曉的真相

改變一個字，地圖就跟著改變

AlphaOracle提出的第一項成果涉及一個單字。研究團隊將一個爭議已久的單字輸入該系統中。該字即為後來演變為漢字「勞」的字[1]。系統整理了筆畫形狀與構字部件，接著找出並列出該字出現的其他句子。隨後，再對照古代文獻知識縮小了字義範圍。

系統提出的解答是，該字屬於地名或氏族名稱。氏族是指源自同一祖先的親族家族。這一結果不同於以往將該字視為修飾語的解讀。研究團隊表示，這項解讀已通過古文字學家的檢驗[1]。

一個字的意義若改變，包含該字的所有句子也都會隨之改變。若它是人名，該句子講述的就是個人的故事；若它是地名，相同的句子就變成了地方行政的紀錄。商朝如何統治該地區的歷史也必須重新描繪。研究團隊報告指出，這一結果促使人們重新審視對商朝行政與社會的認知[1]。

這項研究中值得關注的並非答案本身，而是系統同時給出了依據。AlphaOracle展示了它是根據哪些句子與構字部件才選出該答案的[1]。學者可以沿著這些依據追溯，決定是否同意。以人類能夠審查的形式提出答案，正是這項研究的核心所在。

耗時數年的配對縮短至數天

在綴合領域也取得了成果。這正是前面提到的「綴多多」程式所達成的成就。該程式找到了學者們長久以來忽略的300多組配對[2]。這是在一百多年來人工尋找的配對成果之上所增添的數字。

由人類承擔的工作依然存在。學者會逐一確認機器給出的名單，並最終以實物比對來完成確認。

安陽師範大學的研究室也發表了自家的成果。截至2026年5月，新拼合的甲骨已達155件。同一個研究室辨識出重複紀錄的案例則接近7,000件[3]。過去曾有一塊骨頭分別收錄在不同圖錄中的情況，這會導致同一塊骨頭被當作兩件不同的文物來計算。比起新綴合的數量，辨識出重複的數量要多得多。這意味著機器所做的工作不僅在於新發現，也在於檢驗既有的文獻紀錄。

碎片一旦拼合，句子就會變長。很多情況下一側只留下提問，而占卜結果卻在另一側。只有當兩塊碎片合在一起時，占卜的日期與實際發生的事情才能連成一貫。如此一來，它便不再是數條零碎的簡短紀錄，而是成為單一事件的完整紀錄。

這項差異直接影響了歷史敘述。若只看殘斷的紀錄，事件顯得支離破碎；若看銜接完整的紀錄，前後順序便一目了然。商王於何時提問、何時得到結果，都能得以確立。

辨識出7,000筆重複紀錄同樣具有重要意義[3]。若將同一塊骨頭算作兩件，從資料數量開始就會出錯，甚至會把同一個句子當作兩筆不同的紀錄來引用。機器透過大量重新比對拓本，過濾掉了這些重複現象。這雖然不是引人注目的重大發現，卻是端正基礎資料的關鍵工作。

散落海外的甲骨在螢幕上相聚

阻礙綴合的壁壘曾是距離。從同一塊骨頭上脫落的碎片若分散在不同國家，便無法並列檢視。即使互通照片，也無從得知厚度與曲折程度。而三維資料降低了這道門檻。

前述的900餘件海外資料便是突破這道壁壘的具體案例[3]。實物依然留在原處，被帶過來的則是表面的形態資訊。

安陽的研究人員可以在螢幕上旋轉檢視該碎片，甚至可以將北京的碎片與歐洲的碎片放在同一螢幕上，比對兩者的斷面是否吻合。

公開的方式也隨之改變。過去必須親自造訪收藏機構才能見到文物，光是取得許可往往就要耗時數月。如今，即使是學生也能在家中查閱甲骨資料[3]。接觸資料的人數增加，負責檢驗的眼睛也隨之增加；檢驗的眼睛變多，篩除錯誤綴合的機會自然也隨之增加。

機器目前仍無法做到的事

在取得成果的同時，局限性也變得一清二楚。2026年6月底公開的S-OBI測試以具體數字呈現了這項局限。多個大型人工智慧解答了相同的甲骨文句子題目。相較於辨識單一文字的題目，理解整個句子的題目成績明顯下滑。原因也隨之浮現：模型從保存完好的文字就開始辨識錯誤，接著帶著這個錯誤繼續推論，導致最終答案出錯[4]。要掌握三千年前文句的脈絡，目前仍高度依賴人類的專業。

更需要警惕的問題是憑空捏造。人工智慧是依機率來挑選答案的。在甲骨文這類訓練資料匱乏的領域中，它可能會捏造出看似合理但實際上並不存在的字。商代的語言也沒有足夠的相鄰語言資料可供比對[5]，處於難以翻查字典來驗證答案的境地。

資料分布不均的問題依然存在。迄今收集的甲骨資料高度集中在特定時期與特定發掘地點。以偏差資料訓練出的機器，一旦遇到陌生區域的甲骨便容易出錯。2026年發表的綜述論文也將資料的品質與均衡性列為今後的挑戰[5]。

從數字來看，前方的路依然漫長。目前仍有大約3,000個意義不明的文字。AlphaOracle提出依據所破解的字，僅僅是其中之一[1]。考慮到破解一個字所投入的精力，剩餘的工作量十分龐大。不過，一旦方法確立成熟，破解下一個字的速度就會稍有提升。

人類與機器分工合作

因此，目前的研究模式已確立為人機協作。機器廣泛提出候選對象並整理依據，人類則透過文獻與實物來驗證這些候選。

前面提到的分析時間縮短，也是這種分工的結果。並非機器取代了判斷，而是預先為人類篩選出需要檢視的資料。專家們之所以認為這項工具實用，原因亦在於此[1]。

在綴合方面也確立了相同的模式。過去需要耗費數年的比對工作，如今程式幾天之內就能完成[2]。學者帶著那些候選結果走進庫房，將實際的骨片拼合在一起，確認斷面是否吻合。最終的判斷依然落在人類的指尖上。

這項分工附帶了一個條件：機器給出答案的依據必須能夠呈現給人類。若看不到依據，學者便無從查證；無法查證的答案，便無法納入歷史敘述之中。這也是AlphaOracle將流程劃分為多個階段的原因[1]。

三千年前，商朝的國王看著骨頭上的裂紋探詢未來。今天的研究人員看著同一塊骨頭的破裂斷面探尋過去。提問的方向雖截然相反，注視的痕跡卻依然相同。

參考資料

- [1] Liu, Y., Guan, H., Wang, P., Wang, X., Wan, J., Zhang, K., Zheng, H., Liu, X., Kuang, Z., Yang, H., Li, B., Liu, Y., Jin, L., & Bai, X. (2026). AlphaOracle: Oracle bone script decipherment via human-workflow-inspired deep learning. *The Innovation*, 7(11), 101462. <https://doi.org/10.1016/j.xinn.2026.101462>
- [2] 北京日报. (2026, 8월 28일). 高校团队"唤醒"三千年前文明碎片. 新浪教育. <https://edu.sina.com.cn/l/2026-08-28/doc-inipwhrw0020843.shtml> (2026年9月3日檢索)
- [3] 梁文艳. (2026, 5월 26일). 以科技赋能，让甲骨文"活"起来. 中国青年报. https://zqb.cyol.com/pc/content/202605/26/content_426168.html (2026年9月3日檢索)
- [4] Li, Z., Chen, Z., Chen, T., & Zhai, G. (2026, 7월 1일). Beyond single character: Evaluating MLLMs for sentence-level oracle bone inscription understanding. arXiv:2606.31169. <https://arxiv.org/abs/2606.31169> (2026年9月3日檢索)
- [5] Chen, Z., Hua, W., Li, J., Zhu, Y., Zhi, X., Liu, Z., Chen, T., Zhang, W., & Zhai, G. (2026). Oracle bone inscriptions information processing: a comprehensive survey. *npj Heritage Science*, 14, 220. <https://doi.org/10.1038/s40494-026-02511-w> (2026年9月3日檢索)

第8章. 從泥坑中歷劫生還的毛筆字： 新羅木簡與草書解讀人工智慧

8.1. 案件檔案

在木頭上書寫的時代

曾經有一段紙張非常珍貴的時期。古代東亞人會將木頭削薄打磨，並在上面寫字。寫有文字的這種木板被稱為木簡。木簡的寬度通常相當於兩三根手指併攏。長度則大多在一拃左右。也有將木頭削製成柱狀的。這種木簡的四個面都能書寫文字。因為是在多個面上寫字的木板，所以被稱為多面木簡。

木簡的用途十分多樣。其一是繫在物品上的標籤。從地方運送穀物或魚醬到王京時，會用繩子將木簡綁在物品上。上面記錄了寄送地的郡邑名稱、物品名稱與數量。另一種是官府的公文。出土的木簡中，便有記錄官員姓名與職級、接收物品及日期的文物[1]。也有為了練字而反覆書寫相同漢字的練習用木簡。

木頭比紙張堅固，因此可以重複使用。用完的木簡，可以用刀將表面薄薄地削去一層。露出新的表面後，便能在上面重新寫字。這時削下來的薄木屑也會在遺跡中一同出土。這些碎屑被稱為削屑。意思是用來抹去字跡而削下的木頭碎屑。從扶餘官北里遺跡出土的 329 件木簡類文物中便可看出這一點。其中有 247 件是削屑。這代表附近曾有製作與修改公文的機構設施[1]。

歷史學家珍視木簡是有原因的。史書通常是在事件結束之後才撰寫的。其中會夾雜著撰寫者的個人判斷與王朝的現實考量。木簡則接近於當天事務當天記錄下來的檔案。它並不是為了展示給後人看而刻意修飾的文章。誰在何時寄送了幾石糧食，都被原原本本地保留了下來。在事件發生當下產生的這類記錄，被稱為第一手史料。木簡上也有寫錯的字和修改過的痕跡。因此研究人員會將多件木簡並列對照比較。

朝鮮半島與日本列島都有木簡出土。這兩個地方過去都是以漢字製作公文的社會。但與此同時，他們也都在苦思如何記錄自己的母語。新羅人結合漢字的意與音，來記錄韓語的語序。這種表記法被稱為吏讀與鄉札。日本人也借用漢字的讀音來記錄自己的語言。木簡正是說明這種嘗試始於何時的實物證據。因為紙質文書幾乎沒有留存下來。這也是兩國木簡並列對照研究得以持續開展的原因。

泥土守護的千年歲月

木簡出土的地點通常很固定。如宮殿旁的池塘底部、城牆內側積水的低窪濕地，以及填塞水溝的淤泥層。人們用過即丟的木頭沉入水中，就這樣被掩埋了起來。

水下的淤泥幾乎無法透入空氣。使木頭腐爛的，大多是呼吸空氣生存的好氧微生物。一旦空氣被阻絕，這些微生物便無法發揮作用。雖然仍有在無氧環境下生存的微生物存在，但木頭腐朽的速度會大幅減緩。因此木頭才能在歷經千年之後依然保持原貌。即使在皮革與紙張消失得無影無蹤的土壤裡，木簡依然留存了下來。不是墓穴，而是垃圾坑守護了這些記錄。

但取而代之的是發生了另一種變化。水分滲入木材內部，填滿了細胞壁。土壤中的鐵質與各種礦物質也一同滲入。木材變得飽含水分。這種木材被稱為吸水木材，即「飽水木材」。

飽水木材取出後絕不能直接晾乾。隨著水分流失，它會迅速收縮並乾裂。這就是在發掘現場必須立刻將木簡浸入水桶的原因。在保存修復室中，專業人員會用藥劑慢慢置換出木材內部的水分。單是這個過程就需要耗費數月甚至數年。判讀工作必須在此之前完成。因為在處理過程中，木材表面的顏色可能會發生改變。

黑色木頭上的黑色字跡

滲入礦物質的木材顏色會變深。許多木材會由褐色漸漸轉為漆黑。木材變黑的這種現象被稱為黑化。問題在於寫字所用的材料。古人將燃燒松木所得的松煙調和動物膠，製成了墨。墨的主要成分是碳，顏色也是黑的。

黑化的木頭上寫著黑色的墨字。無論用多明亮的光線照射，兩者都難以分辨。再加上在水中度過漫長歲月，大量的墨粒已被沖刷流失。木紋深色的縱向紋理與毛筆筆觸交雜在一起。泥土與腐殖質填滿了毛筆劃過的凹痕。光憑肉眼，甚至很難看出上面是否存在文字。

判讀若出現困難，學說也會產生分歧。面對同一件木簡，不同學者可能會解讀出不同的漢字。只要改變一個字，句意就會改變，歷史解釋也會隨之不同。因此，木簡研究人員長久以來都在尋找一雙新的眼睛。他們需要能夠捕捉人眼所遺漏細節的工具。

有待判讀的木簡持續在增加。慶南咸安城山山城不斷出土新羅木簡。2025 年 8 月又新確認了兩件帶有銘文的木簡[2]。銘文是指文物上所書寫的文字。忠南扶餘官北里遺跡出土了大量百濟泗沘時期的木簡類文物[1]。京畿楊州大母山城則出土了推測為 5 世紀的木簡。若屬實，這將成為目前已知百濟文字資料中年代極早的實物[3]。

日本的情況也是如此。奈良文化財研究所蒐集各地出土的木簡資訊並對外公開。這個資料庫的名稱是「木簡庫」[4]。在該資料庫中，不僅能依照木簡上的文字進行檢索，還能根據用途、年代、材質與出土地區進行查詢[5]。等待被閱讀的木頭就這樣堆積如山。

參考資料

- [1] 파이낸셜뉴스. (2026, 2월 5일). 부여 관북리 '백제 피리' 첫 확인..목간 329점 무더기 발견. 파이낸셜뉴스. <https://www.fnnews.com/news/202602051448564101> (2026年9月3日檢索)
- [2] 경남일보. (2025, 8월 7일). 함안 성산산성서 새로운 목간 2점 명문 확인. 경남일보. <https://www.gnnews.co.kr/news/articleView.html?idxno=616028> (2026年9月3日檢索)
- [3] 뉴시스. (2025, 11월 20일). 1500년 전 '백제 문자' 깨어났다...양주 대모산성서 5세기 목간 출토. 뉴시스. https://www.newsis.com/view/NISX20251120_0003410107 (2026年9月3日檢索)
- [4] 奈良文化財研究所. 木簡庫(MOKKAN-KO). <https://mokkanko.nabunken.go.jp/ja/> (2026年9月3日檢索)
- [5] カレントアウェアネス・ポータル. (2018). 奈良文化財研究所、木簡に関する情報を検索するシステム「木簡庫」(MOKKAN-KO)を公開. 国立国会図書館. <https://current.ndl.go.jp/car/35714> (2026年9月3日檢索)

8.2. 科學搜查隊的追查 1：數位墨水放大鏡

使用肉眼看不見的光

人類的肉眼只能看見狹窄範圍內的光。在彩虹的紅色之外，也存在著光。這種光被稱為紅外線。因為波長較長，肉眼無法捕捉，但可以用相機拍攝。遙控器發送給電視的訊號也是紅外線。

在木簡判讀中使用紅外線，是因為材料本身的特性。墨汁中的碳微粒能良好地吸收紅外線。黑化的木紋則會大量反射紅外線。在人眼看來同樣漆黑的兩種物質，在紅外線相機前卻涇渭分明。吸收光線的一方拍出來較暗，反射光線的一方拍出來則較亮。原以為已經消失的毛筆筆觸，就這樣在照片上浮現出來。

主要使用的光是波長約 800 奈米的近紅外線。奈米是將 1 公尺分成 10 億等分的長度。人眼能看見的紅光大約在 700 奈米附近結束。這意味著使用的是剛好落在可見光之外的光。拍攝是在遮蔽外部光線的暗室中進行。照射時會讓光線均勻散佈，避免偏向某一側。因為若只從單側照射，木紋的每一道凹槽都會產生長長的陰影。

紅外線拍攝已成為木簡研究的基本流程。發掘報告中刊載的木簡照片旁，通常都會並列附上紅外線照片。奈良文化財研究所很早就開發了利用紅外線相機讓木簡文字影像更加清晰的設備[1]。

從紅外線到高光譜

並非僅憑一張紅外線照片就能讓所有文字都重現。根據污漬的位置與木紋的狀態，能看清文字的波長各不相同。因此，高光譜影像技術應運而生。

普通照片只儲存紅、綠、藍三種數值。高光譜影像則不同。它將同一個畫面細分成數十到數百個狹窄的波長區間進行拍攝。可以想像成將光線細分成彩虹般的各種顏色後分別拍攝。畫面中的每一個點，都會依序記錄各波長下的亮度值。因此產生的不是單張照片，而是一整疊厚厚的照片集。

透過比對這些數值，便能區分殘留墨跡之處與泥土污漬沉積之處。因為墨與泥土在各個波長下亮度變化的規律截然不同。有些文字在 900 奈米附近清晰可見。有些文字則在其他波長區間重獲清晰。研究人員藉由挑選不同的波長區間，尋找最容易判讀的最佳組合。電腦也會自動比對多個區間來推敲候選結果。原本需要人工逐張用肉眼確認的工作，透過運算大幅減少了。

國立伽耶文化遺產研究所於 2025 年首次採用了這種方法。也就是將高光譜影像技術導入咸安城山山城木簡的判讀工作。該研究所表示，這種方法比既有的紅外線拍攝更為精密。據其說明，即使是肉眼或一般拍攝難以辨識的文字，也能藉此重現[2]。

按字形檢索文字

即便照片變得清晰，下一個問題依然存在。必須判定模糊的毛筆筆觸究竟是哪一個漢字。以草書隨性連筆書寫的字，形體變化極大。許多字甚至只剩下寥寥幾筆筆劃。由於漢字種類多達數萬個，光靠人工翻查字典也有其極限。

奈良文化財研究所與東京大學史料編纂所開發了一套協助此項工作的系統。其名為 Mojizo。使用者裁切並上傳想要判讀的文字圖像。系統便會在從木簡與古文書中蒐集的文字圖像資料庫中進行檢索。並從中找出外形相似的字展示出來[3]。由人工審視候選字後做出判斷。它不是直接給出答案的機器，而是調出相似案例的搜尋工具。

這類工具改變的是工作流程。過去學者必須依賴記憶與字典來推想候選字。如今機器會先從數萬件資料中篩選出相似的文字。人工則在縮小範圍後的候選字中結合上下文來斟酌挑選。判讀所需的時間大幅縮短。

辨別機器繪製的筆劃

如今的生成式人工智慧能讓模糊的圖像變得清晰。它能連結斷裂的線條，並填補空白處。將這項技術應用於木簡照片時，效果看起來確實很好。但這之中存在著陷阱。

機器補上的筆劃並非真正殘留在木頭上的墨跡。而是從學習過的數萬件字跡平均值中推導出的猜測。若木紋恰好呈現圓弧狀，機器就可能將其判讀為圓弧筆劃。如此一來，原本不存在的字便似模似樣地產生了。史料遭到污染，歷史解釋也會隨之走樣。

因此在研究現場有一套明確規則。原始照片與校正結果必須並列展示。實際確認到的墨跡與機器生成的像素必須明確區分標示。機器的答案不作為定論，而僅視為候選清單。最終判讀由人工把關。在發表判讀方案時，也必須同時說明是使用何種設備、在何種條件下拍攝的。

政府層面的籌備也在進行中。國家科技諮詢會議於 2026 年 4 月 16 日審議通過了國家遺產廳的施行計畫。該計畫預計在 2026 年至 2030 年間，將人工智慧、機器人與感測器導入國家遺產的調查與修復中。今年的投資規模為 140 億韓元。其中 44 億韓元被編列用於智慧型尖端保存技術的研發[4]。國立金海博物館於 2026 年 6 月 23 日舉辦了學術研討會。主題為「伽耶遺產研究的科學視角」。現場進行了七項發表，探討非破壞性分析、數位修復以及 3D X-ray 顯微鏡技術[5]。以肉眼無法窺見的視角透視文物內部的方法，正從木簡逐步拓展至更廣闊的領域。

參考資料

- [1] 奈良文化財研究所. 赤外線カメラを利用した木簡文字画像鮮明化システムの開発. 学術機関リポジトリデータベース. <https://irdb.nii.ac.jp/01144/0004944317> (2026年9月3日檢索)
- [2] 경남일보. (2025, 8월 7일). 함안 성산산성서 새로운 목간 2점 명문 확인. 경남일보. <https://www.gnnews.co.kr/news/articleView.html?idxno=616028> (2026年9月3日檢索)
- [3] 奈良文化財研究所, 東京大学史料編纂所. 木簡・くずし字解読システム MOJIZO. <https://aimojizo.nabunken.go.jp/> (2026年9月3日檢索)
- [4] 뉴시스. (2026, 4월 16일). "AI, 블록체인으로 국가유산 보존관리"…과기자문회의, 유산청 시행계획 의결. 뉴시스. https://www.newsis.com/view/NISX20260416_0003594216 (2026年9月3日檢索)
- [5] 경남일보. (2026, 6월 17일). 국립김해박물관, 23일 '가야학술제전' 심포지엄 개최. 경남일보. <https://www.gnnews.co.kr/news/articleView.html?idxno=639314> (2026年9月3日檢索)

8.3. 科學搜查隊的追查 2：KuroNet 與 Miwo

連筆書寫的字跡無法分割

日本留存了數百萬冊在江戶時代以前撰寫的古文書。這些文書是以名為「崩字（Kuzushiji）」的草體字寫成的。Kuzushiji 是自 8 世紀起使用了千餘年的書寫體。1900 年日本改革了文字制度。此後學校不再教授這種字體。因此現今絕大多數的日本人連 150 年前的書籍都無法閱讀[1]。

難以閱讀的原因並不僅在於字形本身。有一種毛筆不離紙面、將多個字連續書寫的書法技法。前一個字的末筆與後一個字的首筆連成一體。如此一來，字與字之間便不留任何分界線。整行字看起來就像是一條長長的線。

電腦的文字辨識長久以來都是按照相同的流程運作。首先將每個字用方形外框逐一裁切分割。接著判斷被裁切的外框內的圖像屬於哪一個字。這種方式必須有明確的分界線才能運作。面對連筆書寫的草字，在第一階段就宣告崩潰。由於前面出錯，後面也會跟著錯，因此根本無法提高準確率。

不分割文字的 KuroNet

日本人文學開放資料共同利用中心的研究團隊開闢了另一條道路。Tarin Clanuwat、Alex Lamb 與北本朝展共同打造了這套模型。名稱為 KuroNet。該模型於 2019 年的文件分析與辨識國際研討會上首次發表[1]。2020 年進一步優化了架構，並將成果重新刊載於學術期刊上[2]。

KuroNet不預先切割字元，而是將整頁古文書完整輸入。模型的核心結構是U-Net。U-Net是一種先逐漸縮小圖像以擷取特徵，接著再放大回原本尺寸的神經網路。縮小端與放大端相互對稱，因而呈現英文字母U的形狀。在縮小的過程中學習筆畫的粗細與方向，在放大的過程中則將該資訊還原到原始位置。

這個神經網路能同時預測整頁中的兩件事：一是字元的中心在哪裡，另一則是該位置的字元是什麼。結果以機率地圖呈現，是將畫面各處存在字元的可能性標示出來的圖像。只要沿著機率高的地方追蹤，就能同時確定字元的位置與名稱。由於沒有切割步驟，即使面對連筆草書也不會失效。

訓練時使用了日本國文學研究資料館製作的資料，名稱為「日本古典籍Kuzushiji資料集」，內含超過100萬字的字元樣本[2]。同一個研究團隊為了讓機器學習研究人員便於使用，也公開了縮小版，模仿手寫數字資料集製作的Kuzushiji-MNIST便是其一。他們也一併推出了增加字數的版本以及僅收錄漢字的版本[3]。多虧於此，即使不懂日本古典文學的研究人員也能投入這個課題。

KuroNet於2022年10月交接了任務。辨識模型的名稱變更為Ruri[4]。Ruri是將物件偵測技術調整為適合草書的辨識模型。即使文字背後的底色或花紋複雜，它也能比以往辨識得更好。該模型利用補強後的相同資料集進行訓練，能夠解讀出以往無法辨識的字元[5]。

走入掌中的Miwo

將原本在伺服器上運行的模型放入智慧型手機，是下一個課題。就這樣誕生的應用程式便是Miwo。該應用程式於2021年8月30日首次公開[5]。日文標記為「みを」，英文標記為「Miwo」[6]。

使用方法很簡單。只要用智慧型手機相機拍下古文書，片刻後畫面上就會出現轉換為現代文字的結果。使用者可以透過滑桿將原本與結果重疊比對，也會標示出原本是什麼變體假名。變體假名是現今不再使用的古代平假名字形。它還具備將結果匯出為文字檔案的功能。拍攝的相片不會保留在伺服器上。在Android與iPhone上皆可免費使用[6]。

自2022年10月26日推出的1.1版起，辨識模型更換為Ruri[5]。同一個月，該應用程式獲得了2022年度日本優良設計獎（Good Design Award）[7]。使用紀錄持續累積，自公開以來，截至2026年9月3日為止辨識的影像已達3,523,373件[6]。這意味著人們正在圖書館、寺院以及地方檔案館拍攝古文書。

大眾共同解讀

機器先辨識、人類再修改的方式也應用在其他地方。有一個名為「大家的翻刻（Minna de Honkoku）」的線上翻刻專案。翻刻是指將古代文字轉錄為現代文字的工作。其日文標記為「みんなで翻刻」，意思是大家一起轉錄。該專案於2017年由京都大學古地震研究會發起，現由國立歷史民俗博物館與東京大學地震研究所共同營運[8]。

參與方式十分開放。只要有網路瀏覽器，任何人都能轉錄古文書。畫面上會顯示人工智慧預先辨識的初稿結果，人類只需修正錯誤的字元即可。該網站引進了人文學開放資料共同利用中心的Kuzushiji辨識服務。截至目前為止，已有1萬800多名參與者轉錄了超過4,800萬個字[8]。

轉錄出來的文字不會只是白白堆積。人類修正後的標準答案會再次成為訓練資料。訓練資料增加，模型就會進步；模型進步，初稿結果就更準確，需要人工介入的地方也更少。一旦這個循環開始運轉，速度就會越來越快。

解讀出來的文書用途也很廣泛。這項專案最初是由地震研究人員發起的。古代的日記與官府文書中記載了地震、海嘯與荒年的日期及災情。原本因草書而無法著手的紀錄一旦轉化為資料，便能推算過去的災難週期。古老的毛筆字就這樣成了當今的防災資料。

到這裡為止是日本紙本古文書的故事。KuroNet、Ruri與Miwo都是解讀寫在紙上的Kuzushiji的模型。目前尚未確認有這些模型自動解讀泥坑中木簡的案例。木簡上的字數少、載體為木材，且墨跡大多脫落模糊。用於訓練的標準答案資料也不如紙本古文書累積得那麼多。如果韓日兩國の木簡資料進一步匯聚，情況可能會有所改變。目前這仍是一份期待，而非既有的成果。

參考資料

- [1] Clanuwat, T., Lamb, A., & Kitamoto, A. (2019). KuroNet: Pre-modern Japanese Kuzushiji character recognition with deep learning. 2019 International Conference on Document Analysis and Recognition (ICDAR), 607-614. <https://doi.org/10.1109/ICDAR.2019.00103>
- [2] Lamb, A., Clanuwat, T., & Kitamoto, A. (2020). KuroNet: Regularized residual U-Nets for end-to-end Kuzushiji character recognition. SN Computer Science, 1, 177. <https://doi.org/10.1007/s42979-020-00186-z>
- [3] Clanuwat, T., Bober-Irizar, M., Kitamoto, A., Lamb, A., Yamamoto, K., & Ha, D. (2018). Deep learning for classical Japanese literature. arXiv. <https://arxiv.org/abs/1812.01718>
- [4] ROIS-DS 人文学オープンデータ共同利用センター. KuroNetくずし字認識サービス(AI OCR). CODH. <https://codh.rois.ac.jp/kuronet/> (2026年9月3日検索)
- [5] ROIS-DS 人文学オープンデータ共同利用センター. みを(miwo)とは？ CODH. <https://codh.rois.ac.jp/miwo/about/> (2026年9月3日検索)
- [6] ROIS-DS 人文学オープンデータ共同利用センター. みを(miwo)：AIくずし字認識アプリ. CODH. <https://codh.rois.ac.jp/miwo/> (2026年9月3日検索)
- [7] 国立情報学研究所. (2022, 10월 26일). 세계初のAIくずし字認識アプリ「みを(miwo)」が2022年度グッドデザイン賞を受賞. 国立情報学研究所. <https://www.nii.ac.jp/news/release/2022/1026.html> (2026年9月3日検索)
- [8] みんなで翻刻. みんなで翻刻について. <https://honkoku.org/> (2026年9月3日検索)

8.4. 揭開的真相

月池木簡揭示的王室生活

慶州東宮與月池是統一新羅王室的離宮與池塘。1975年月池發掘時首次出土了木簡，此後隨著調查持續進行，出土數量不斷增加。國立慶州博物館於月池館展出此處出土的文物[1]。

月池木簡展現了王室的一天。有的木簡是懸掛在陶甕上的標籤。標籤木簡的末端左右兩側刻有凹槽，用繩子繞過凹槽繫在物品上，用途與今日的宅配貨單相似。月池出土的木簡上寫有「鹿甲」與「甕」等字。「鹿甲」是用鹿肉製成的肉醬，「甕」則是陶甕。這代表新羅王室記錄了肉醬的種類與醃製時間並加以管理[2]。

這樣的紀錄並未記載於史書中。《三國史記》記載了國王、戰爭與制度，卻未曾記錄廚房裡送進了什麼。木簡填補了這個空白。泥坑中的一片木板，成了宮廷的起居帳簿。

韓國木簡學會與國立慶州博物館於2026年4月3日和4日舉行了國際學術研討會。主題為「新羅王京慶州的空間與文字生活」，在慶州華白國際會議中心進行了兩天。會議將月城護城河以及東宮與月池出土的資料一併探討，試圖透過木簡、碑石與墨書重塑新羅王京的面貌。韓日兩國的研究人員共同發表了成果[3]。

城山山城木簡的處罰報告

咸安城山山城是新羅於6世紀修築的山城。城牆內側的低濕地中不斷出土木簡。截至2025年8月為止，此處出土的墨書木簡共有247件，佔韓國出土古代木簡總數的一半以上[4]。

在第18次發掘調查中新解讀的兩件木簡是由松木類製成。一件是四面皆書寫文字的多面木簡，另一件則是雙面木簡[4]。多面木簡的四面中有三面載有行政內容，是針對個人進行處罰的紀錄[5]。製作年代推測為6世紀中葉[4]。

這次解讀是藉由前面介紹過的超光譜影像所獲得的成果。無論肉眼或一般照片都無法辨認的文字得以重現生機[4]。沉睡在泥坑中超過一千年的新羅下級官吏報告書，化為具體文句重新復甦。城山山城木簡上也大量記載了運送穀物的鄉邑名稱，是展現6世紀新羅如何劃分並治理地方的珍貴資料，亦能藉此推敲稅賦是循哪條路徑運抵山城的。

百濟王宮的大批文書

扶餘官北里遺址是被認定為百濟泗泚時期王宮遺址的地方。國立扶餘文化遺產研究所於2024年與2025年進行了第16次發掘調查，此處共出土了329件木簡類文物。其中寫有文字的木簡為82件，削屑為247件[6]。這是韓國國內單一遺址的最大數量，且屬於泗泚時期木簡中較早期的產物[7]。

木簡中夾雜著記錄有年代的物件，出現了「庚申年」與「癸亥年」的干支紀年，分別推斷為540年與543年。這意味著這些是百濟從熊津遷都至泗泚之後不久的紀錄。官吏的人事紀錄、記載國家財政的帳簿以及官職名稱一同得到了確認[7]。竹製的百濟管樂器實物也首次出土[8]。

2026年7月22日和23日舉行了國際學術研討會，由國立扶餘文化遺產研究所、韓國木簡學會與百濟學會共同籌辦。主題為「扶餘官北里遺址樂器與木簡」，於釜山海雲臺區舉行。研討會以忠南大學朴淳發榮譽教授的主旨演講拉開序幕，在三個分組會議中發表了十二篇論文。韓、中、日三國的研究人員共同探討了官北里木簡的意涵[8]。

在碑石前需審慎

在木簡上行得通的方法，並非適用於所有古代文字。一旦載體換成石頭，情況就大不相同了。廣開土大王碑就是一個代表案例。這座立於414年的碑石刻有高句麗的歷史。由於碑面長年經受風雨侵蝕，文字已磨損不堪。在這裡，我們必須讀取的不是墨跡，而是鑿刻的凹痕。

因此研究人員透過比對拓本來進行研究。原石拓本是在碑面上塗抹石灰之前拓印的拓本，直接將紙貼在原始石碑上搗打拓印而成，至今確認的僅有十幾份。到了19世紀後半期，出現了另一種拓本，即為了讓文字看起來更清晰而在碑面塗抹石灰後拓印的「石灰拓本」。此外，也製作許多所謂的「雙鉤加墨本」。這是將紙鋪在碑文上，沿著字元輪廓勾勒線條，再用濃墨填滿其餘空白，使其看起來像是拓本。這並非捶打拓印的拓本，而是摹繪的複製品，在製作過程中字形可能被抄寫錯誤[9]。因此，相同的文字在不同拓本中可能會呈現出不同模樣。

爭論的核心是辛卯年條文句。該文句的詮釋與古代韓日關係史交織在一起，引發了長期的爭端，文字是否遭到人為竄改曾是爭論的焦點[10]。若想查明原本的文字，必須將早期的原石拓本與多種石灰拓本進行比對。建立能夠甄別各拓本的標準是首要任務[9]。

坊間流傳著三維量測與人工智慧為這場爭論提供了答案的說法。但若仔細核實，往往找不到被引為依據的具體研究。人們期待在木簡上行得通的方法在碑石上也能奏效，但實際的研究結果並非如此。前面介紹的技術能讓原本看不見的事物顯現出來，要確定顯現之物的意涵，則是下一步的事。

因此，結論具有雙重層面。紅外線與超光譜影像讓新羅與百濟木簡上消失的墨跡重現生機；無分割辨識模型解讀了日本紙本古文書中難以拆分的連筆草書。這兩項技術處理的對象雖不同，方向卻是一致的：機器率先指出人眼

遺漏的筆畫。與此同時，如果缺乏區分機器描繪的筆畫與原本殘留筆畫的規範，同樣的技術便可能製造出原本不存在的歷史。喚醒毛筆字的是機器，而為那些文字負責的則是人類。

參考資料

- [1] 국립경주박물관. 월지관. 국립경주박물관. <https://gyeongju.museum.go.kr/kor/html/sub02/0201.html> (2026年9月3日檢索)
- [2] 국사편찬위원회. 목간. 우리역사넷. https://contents.history.go.kr/mobile/tt/view.do?levelId=tt_b36 (2026年9月3日檢索)
- [3] 경북일보. (2026, 3월 31일). 1500년 전 신라의 목소리 깨어난다...경주서 국제 목간 학술대회. 경북일보. <https://www.kyongbuk.co.kr/news/articleView.html?idxno=4068614> (2026年9月3日檢索)
- [4] 경남일보. (2025, 8월 7일). 함안 성산산성서 새로운 목간 2점 명문 확인. 경남일보. <https://www.gnnews.co.kr/news/articleView.html?idxno=616028> (2026年9月3日檢索)
- [5] 경남도민일보. (2025, 8월 7일). 누군가를 처벌한 신라 행정 보고, 함안 성산산성 목간서 확인. 경남도민일보. <https://www.idomin.com/news/articleView.html?idxno=943461> (2026年9月3日檢索)
- [6] 파이낸셜뉴스. (2026, 2월 5일). 부여 관북리 '백제 피리' 첫 확인..목간 329점 무더기 발견. 파이낸셜뉴스. <https://www.fnnews.com/news/202602051448564101> (2026年9月3日檢索)
- [7] 경향신문. (2026, 2월 17일). 부여 관북리서 발견된 목간 329점, '백제 피리'만큼이나 중요한 이유. 경향신문. <https://www.khan.co.kr/article/202602171427001> (2026年9月3日檢索)
- [8] 이데일리. (2026, 7월 21일). 백제 왕궁터 출토 목간 329점 · 획적 발굴 성과, 국제학술대회서 조명. 이데일리. <https://edaily.co.kr/News/Read?mediaCodeNo=257&newsId=03276726645516816> (2026年9月3日檢索)
- [9] 국립중앙박물관. 광개토대왕비 탁본. 국립중앙박물관 큐레이터 추천 소장품. <https://www.museum.go.kr/MUSEUM/contents/M0501000000.do?schM=view&relicRecommendId=254462> (2026年9月3日檢索)
- [10] 국사편찬위원회. 광개토대왕릉비. 우리역사넷. https://contents.history.go.kr/mobile/kc/view.do?levelId=kc_r100300 (2026年9月3日檢索)

第9章. 愛琴海失落的謎團：線形文字A與幾何網格

9.1. 案件檔案

克里特島宮殿出土的三種文字

愛琴海南端坐落著克里特島。青銅時代，米諾斯文明在這座島嶼興起。米諾斯人造船往來於地中海。他們運送橄欖油、葡萄酒與穀物。宮殿收集這些物資，再分發給人民。

1900年前後，英國考古學家亞瑟·埃文斯發掘了克諾索斯宮殿遺址。他在地下發現了刻有文字的泥板。這些文字並非只有一種。埃文斯將發現的文字分為三類：克里特象形文字、線形文字A和線形文字B[1]。

米諾斯文明延續了一千多年。到了青銅時代晚期，它經歷了巨大的變遷。北方聖托里尼島的火山爆發便是其中之一。火山爆發的時間推測在西元前17世紀至前16世紀之間。確切年代至今仍有爭議[7]。據信火山灰飄散到了克里特島，海嘯襲擊了海岸。西元前1450年左右，邁錫尼希臘人佔領了這座島嶼[1]。火山爆發與邁錫尼人的統治之間相隔了一百多年。很難將這兩起事件直接歸結為因果關係。從那之後，米諾斯人的語言便從歷史紀錄中消失了。

這三種文字使用的年代各不相同。線形文字A的使用時期約在西元前1750年至西元前1450年之間。線形文字B出現的時間則較晚。它是邁錫尼希臘人佔領克里特島之後所使用的文字。這兩種文字共享了許多形狀相似的符號[1]。

泥板能夠留存下來純屬偶然。米諾斯人並不燒製泥板。他們在濕泥上刻寫文字晾乾，使用一年左右便丟棄。然而宮殿發生火災，使泥板如同陶瓷般被燒透。唯有經大火燒過的泥板耐住了三千五百年的歲月。今日留存的資料，正是大火留下的遺產。

在埃文斯分類的三種文字中，克里特象形文字最早出現。這種文字生動地保留了人頭或器皿的形狀。線形文字A則更接近將那些圖畫簡化為線條的樣貌。這兩種文字在近兩百年的時間裡同時並存使用。並非一種消失之後才出現另一種。有觀點認為線形文字A是模仿克里特象形文字創造出來的。然而，無法斷定這三種文字呈現筆直的單一演變脈絡。它們共享若干符號是明確的事實，但彼此銜接的軌跡至今尚未釐清[1]。

線形文字B已被解讀

線形文字B在1950年代被解讀出來。這是由英國建築師麥可·文特里斯完成的。他證實了線形文字B所記錄的語言是希臘語的早期形態。古典學家約翰·查德威克協助了他的研究[1]。

文特里斯從繪製一張表格著手。這張表格橫行放置子音，縱列放置母音[8]。這類表格被稱為音節網格。它是預先為每個音節劃定填入格位的圖表。文特里斯透過統計方法，逐步縮小各個符號應歸屬格位的範圍[1]。

他將克諾索斯和斐斯托斯等地名套入已縮小範圍的格位中。隨著地名契合無誤，其餘符號也隨之相繼解開。學者們嘗試將相同的方法套用於線形文字A。也就是將線形文字B的音值直接帶入形狀相似的符號中[3]。發音雖然讀得出來，但沒有人聽得懂那是何種語言[1]。

帶入音值後卻解不出意思，這種情況相當罕見。因為「識讀文字」與「理解語言」是兩回事。懂韓文的人，即使陌生的外語用韓文標音也能唸得出來。發音雖然讀得出來，意思卻不得而知。面對線形文字A的學者們，處境正是如此。

沒有以雙語刻寫的石碑

解讀古代文字時，學者通常會尋找相互對照的文獻。即以兩種語言並列記載相同內容的文獻。這種文獻被稱為雙語銘文。埃及象形文字便是拜羅塞塔石碑所賜而得以破解。羅塞塔石碑上以三種文字刻寫了同一道法令。其中之一便是學者早已能夠識讀的古希臘語。

美索不達米亞楔形文字也走過了類似的歷程。伊朗貝希斯敦懸崖上留有以三種語言鐫刻的銘文。這篇銘文的三種語言起初都無法識讀。其中古波斯語率先被破譯。因為它的符號較少且詞與詞之間有所區隔。學者們以率先解開的古波斯語為踏板，進而讀懂了其餘文字。

在克里特島上卻未曾出土此類文獻。沒有任何文物在線形文字A旁附有希臘語對譯。米諾斯人留下的文字僅以米諾斯人自身的語言寫成。缺乏可資比對的對照文本，確認詞義的途徑便徹底受阻[1]。

泥板上也見不到任何足以作為線索的標記。破解埃及象形文字的商博良擁有王名框（cartouche）可循。那是環繞國王名字的長橢圓形邊框。邊框內的文字事先便確定是人名。商博良從中獲得了幾個音值，進而解開了其餘符號。線形文字A泥板上則沒有這種飾框[1]。

問題還不止於此。關鍵在於我們不知曉線形文字A所記錄之語言的根源。學者將這種語言稱為米諾斯語。米諾斯語屬於哪一個語系至今仍未明朗。有人主張其屬於印歐語系，也有人主張屬於西亞系。其中被認為依據較充分的是閃語系與呂基亞語系。呂基亞語是昔日通行於今日土耳其西南部的古老語言[1]。但無論哪一種說法，都尚未獲得學界的一致認同。

線形文字A並非只出土於克里特島。在愛琴海的諸多島嶼及希臘本土也都有所發現[1]。就連被火山掩埋的聖托里尼島遺址中亦有出土。文字沿著米諾斯人往來的航路傳播開來。這些資料展現了青銅時代地中海世界緊密連結的廣闊網絡。

學者們之所以執著於線形文字A，原因也正在於此。線形文字A名列歐洲土地上出土的最古老文字之一。這是米諾斯人以自身語言留下的唯一紀錄。若線形文字A被破譯，他們信仰什麼、買賣什麼都將大白於世。希臘神話中留存的米諾陶洛斯傳說，其根源亦深植於此。

1,400件簡短的帳簿

現存資料的數量也相當匱乏。銘文是指刻寫或繪製在文物上的文字。迄今確認的線形文字A銘文約有1,370件。上面所記載的符號全部加總也僅有7,400個左右[1]。這與埃及或美索不達米亞浩繁的文獻規模完全無法相提並論。

文字篇幅也極其簡短。泥板絕大多數是倉庫帳簿。記錄了何物入庫幾件、何人支領多少。通篇只見人名、品名與數字整齊羅列。堪稱完整句子的語句少之又少。缺乏句子，探求語法結構的途徑也就隨之減少。

計算符號數量時需要分層歸類。表記音節的符號有97個。完整代表物品名稱的表意符號約有50個。記數符號有4個，分數符號有17個。由兩個符號疊加形成的複合符號已知約有150個[6]。將這些全數彙整的標準符號清單達到300個[2]。其中頻繁出現的符號並不多。其餘多是僅出現幾次便不再現身的符號。若要交由電腦進行統計分析，相同的符號必須多次出現。僅出現幾次的符號，統計方法根本無法捕捉其規律。

因此，研究人員決定重新建構資料本身。這是一項將手摹圖樣轉換為電腦可處理形式的工作[2]。2026年，甚至出現了利用光線重新辨讀刻痕的方法[4]。這種方法能顯現肉眼容易遺漏的細淺凹痕[5]。喚醒愛琴海沉默的工作，就這樣在重新打磨資料中再度啟程。

參考資料

- [1] Braović, M., Krstinić, D., Štula, M., & Ivanda, A. (2024). A systematic review of computational approaches to deciphering Bronze Age Aegean and Cypriot scripts. *Computational Linguistics*, 50(2), 725-779.
https://doi.org/10.1162/coli_a_00514
- [2] Salgarella, E., & Castellan, S. (2021). SigLA: The Signs of Linear A. A palaeographical database. In *Grapholinguistics in the 21st Century 2020: Proceedings (Grapholinguistics and Its Applications, Vol. 5, pp. 945-962)*. Fluxus Editions.
<https://sigla.phis.me/> (2026年9月3日檢索)
- [3] Younger, J. G. (n.d.). Linear A texts and inscriptions in phonetic transcription. University of Kansas.
<http://people.ku.edu/~jyounger/LinearA/> (2026年9月3日檢索)
- [4] Giorgi, L., & Lopez, S. (2026). Aegean Bronze Age writing systems through RTI: New potential and perspectives. *Digital Applications in Archaeology and Cultural Heritage*, 41, e00508.
<https://www.sciencedirect.com/science/article/abs/pii/S2212054826000159> (2026年9月3日檢索)
- [5] GreekReporter. (2026, 3월 14일). New technology reveals hidden secrets of 3,500-year-old Bronze Age Aegean scripts. *greekreporter.com*.
<https://greekreporter.com/2026/03/14/new-technology-reveals-hidden-secrets-aegean-bronze-age-scripts/> (2026年9月3日檢索)

- [6] Del Frio, M. (2022). Linear A. Mnamon: Ancient writing systems in the Mediterranean. Scuola Normale Superiore. <https://doi.org/10.25429/sns.it/lettere/mnamon022> (2026年9月3日檢索)
- [7] Pearson, C. L., Brewer, P. W., Brown, D., Heaton, T. J., Hodgins, G. W. L., Jull, A. J. T., Lange, T., & Salzer, M. W. (2018). Annual radiocarbon record indicates 16th century BCE date for the Thera eruption. Science Advances, 4(8), eaar8241. <https://doi.org/10.1126/sciadv.aar8241>
- [8] Stavrakakis, E. (2024). Ventris' Grids. Bulletin of the Institute of Classical Studies, 67(1), 86-103. <https://doi.org/10.1093/bics/qbae033>

9.2. 科學鑑識調查：非監督音韻網格配對

沒有字典也能尋找配對的機器

打造翻譯機通常需要對譯句對。將韓文句子與英文句子並列排整，展示給機器學習。這類資料被稱為平行語料庫。線形文字A並不存在平行語料庫[3]。

因此，研究人員選擇了一種無需標準答案的學習方法。這種方式不預先給予標準答案，僅憑統計數據尋找規律。這被稱為非監督式學習。機器會統計符號以何種順序相連出現，進而與其他文字系統的統計數據進行比對，尋找可能的配對。

2019年，美國麻省理工學院研究團隊提出了這種方法。這是由駱嘉明（Jiaming Luo）、曹袁（Yuan Cao）與蕾吉娜·巴齊萊（Regina Barzilay）共同開發的模型。他們將連結兩個文字系統的任務轉化為最小費用流問題。具體做法是在符號與符號之間建立渠道，並賦予水流通過時的代價。耗費代價最低的水道，其為正確配對的機率最高[1]。

研究團隊將此模型應用於烏加里特文字。結果比先前的最佳成績高出5個百分點。他們也將其應用於線形文字B。在與希臘語同源的詞彙中，有67.3%被正確對譯[1]。不過，線形文字B原本就已有既定答案。研究團隊事先已知該文字記錄的是希臘語。

研究團隊選擇線形文字B作為測試基準是有原因的。唯有在已知答案的問題上，才能衡量模型的效能[1]。若直接應用於沒有答案的問題，根本無從得知結果是否正確。因此必須先以已讀懂的文字衡量效能，再逐步邁向未解讀文字[3]。

模型輸出的並非完整的翻譯文。而是各符號可能對應的候選音值清單[1]。學者接手這份清單後，與文物逐一比對並加以篩選。電腦承擔的是為人類縮減審閱工作量的角色[3]。

將語音特徵轉化為坐標

線形文字A的情況則截然不同。連應該與何種語言比對都尚未查明[3]。2021年，同一個實驗室提出了應對這種條件的模型。這是由駱嘉明、弗雷德里克·哈特曼（Frederik Hartmann）、恩里克·桑圖斯（Enrico Santus）與巴齊萊、曹袁共同開發的[2]。

該模型同時處理了兩大難題。第一個難題是缺乏詞界。古代文獻往往沒有分詞空格。第二個難題是親屬語言尚未確定[2]。

研究團隊將語音特徵轉化為數值坐標代入其中。他們以國際音標（IPA）為基礎，將子音與母音的性質置於坐標上。雙唇音之間的坐標彼此接近。與舌根音之間的坐標則相距甚遠。有了這些坐標，機器就能減少產生離譜荒謬的配對[2]。

研究團隊先以哥德語和烏加里特文字測試了該模型。接著將其應用於尚未解讀含義的伊比利語。伊比利文字的符號音值已大致為人所知。尚未解開的是那些語音所記錄的語言。研究團隊並未發現伊比利語與巴斯克語同屬親屬語言的證據。此一結果與學界的主流觀點相符。同時也展示了比對各候選語言以判斷誰可能為親屬語言的程序[2]。同樣的方法亦可直接應用於線形文字A的研究[3]。

空格揭示的規律

機器比對的並非僅有符號的外形。符號與相鄰符號搭配的慣性亦納入運算。有些符號只出現在詞首。有些符號只出現在詞尾。甚至有些符號彼此絕不相鄰並列[3]。

這類統計反映了共同出現的頻率程度，稱為共現統計。彙整共現統計後，表格中的空格規律便會顯現。位於同一橫行的符號，具有相同子音的機率很高。位於同一縱列的符號，具有相同母音的機率很高[8]。文特里斯當年徒手進行的工作，電腦如今能以快得多的速度反覆進行[3]。

電腦完成這項運算的速度遠比人類快得多。若有300個符號，配對組合的數量便達數萬種之多。這是人類在紙上一筆一劃推算難以負荷的龐大量體。文特里斯歷經數年才完成了這項工作[3]。

到此為止，僅是描繪出語音的結構。即使勾勒出結構，也無法由此得知那些語音所代表的含義。填滿網格與理解語言本就是兩碼事[3]。

將字形轉化為資料

若要让機器進行統計分析，資料必須具備機器可讀的形式。單憑紙本書籍中刊載的圖畫，根本無法讓電腦進行運算。承擔這項任務的正是名為SigLA的資料平台。這是由劍橋大學的艾絲特·薩爾加拉（Ester Salgarella）與雷恩大學的西蒙·卡斯特蘭（Simon Castellan）共同建置的[4]。

SigLA中收錄了300個標準符號與400多件手繪摹寫的銘文。在銘文中可以找到3,000個以上的個別符號。資料與圖形皆已公開供所有人使用。甚至可以單獨拆出單一符號的筆畫來進行比對[4]。

2026年3月發表的研究進一步拓展了資料的維度。拉維尼亞·喬爾吉與薩拉·洛佩斯將反射轉換成像應用於愛琴海文字。這是一種從多個方向對同一件文物打光並拍攝多張照片的方法。將這些照片合成後，表面的細微起伏便會顯現出來。甚至可以推斷出刻刀先刻下了哪一筆[7]。

這類資料累積起來後，對其他文字也能發揮作用。2022年波隆那大學的研究團隊重新審視了賽普勒斯-米諾斯文字。這項研究由米歇爾·科拉查、法比奧·坦布里尼、米格爾·瓦萊里奧與西爾維婭·費拉拉共同進行。他們利用未加入先驗知識的神經網路比對符號形狀。結果得出了必須重新檢討以往分為三組之既有分類的結論。從統計學來看，視為單一文字體系更為合適[5]。

電腦產出的結果不會直接成為論文。它必須經過考古學家與文獻學家重新審視的過程[3]。他們會核對哪塊泥板出自哪間房間，比對運筆習慣以確認是否出自同一名書吏之手[7]。只有通過這項審查後，才能納入學界的資料[3]。

同時探討三種愛琴海文字的研究也陸續展開。這種方法將克里特象形文字、線形文字A與賽普勒斯-米諾斯文字放在一起比對。將三種文字相互對照時，單看一種文字時看不見的規律便會浮現。電腦在此類大規模比對中正能發揮優勢[3]。

機器無法代勞之事

這些方法並未能揭開線形文字A的意義。眼前仍有三道必須跨越的難關。

第一是比對的對象。非監督式對齊只有在未解讀文字與已知文字之間存在親緣關係時才能發揮效用[1]。米諾斯語並沒有存活至今的親屬語言[3]。座標雖然可以描繪出來，卻沒有能固定該座標的基準點。

第二是資料量。1,370件銘文用於機器學習顯得太少[3]。資料不足時，大型神經網路會編造出看似合理的虛假內容。這種錯誤被稱為幻覺。

第三是驗證。2026年7月刊登於The Conversation的一篇文章探討了這個問題。該文由都柏林城市大學的珍·阿德金斯撰寫。她指出，透過統計比對模式無法憑空創造出不存在的含義。必須有已知語言或對譯文本作為立足點，才能區分偶然巧合與真正的規律。機器只能縮減候選範圍，甄別正誤依然是人類的責任[6]。

參考資料

- [1] Luo, J., Cao, Y., & Barzilay, R. (2019). Neural decipherment via minimum-cost flow: From Ugaritic to Linear B. In Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics. Association for Computational Linguistics. <https://aclanthology.org/P19-1303/> (2026年9月3日檢索)
- [2] Luo, J., Hartmann, F., Santus, E., Barzilay, R., & Cao, Y. (2021). Deciphering undersegmented ancient scripts using phonetic prior. Transactions of the Association for Computational Linguistics, 9, 69-81. https://doi.org/10.1162/tacl_a_00354
- [3] Braović, M., Krstinić, D., Štula, M., & Ivanda, A. (2024). A systematic review of computational approaches to deciphering Bronze Age Aegean and Cypriot scripts. Computational Linguistics, 50(2), 725-779. https://doi.org/10.1162/coli_a_00514
- [4] Salgarella, E., & Castellan, S. (2021). SigLA: The Signs of Linear A. A palaeographical database. In Grapholinguistics in the 21st Century 2020: Proceedings (Grapholinguistics and Its Applications, Vol. 5, pp. 945-962). Fluxus Editions. <https://sigla.phis.me/> (2026年9月3日檢索)
- [5] Corazza, M., Tamburini, F., Valério, M., & Ferrara, S. (2022). Unsupervised deep learning supports reclassification of Bronze age cypriot writing system. PLOS ONE, 17(7), e0269544. <https://doi.org/10.1371/journal.pone.0269544>
- [6] Adkins, J. (2026, 7월 27일). Cracking the code: can AI help us decipher ancient languages? The Conversation. <https://theconversation.com/cracking-the-code-can-ai-help-us-decipher-ancient-languages-288238> (2026年9月3日檢索)
- [7] Giorgi, L., & Lopez, S. (2026). Aegean Bronze Age writing systems through RTI: New potential and perspectives. Digital Applications in Archaeology and Cultural Heritage, 41, e00508. <https://www.sciencedirect.com/science/article/abs/pii/S2212054826000159> (2026年9月3日檢索)
- [8] Stavrakakis, E. (2024). Ventris' Grids. Bulletin of the Institute of Classical Studies, 67(1), 86-103. <https://doi.org/10.1093/bics/qbae033>

9.3. 揭曉的真相

數字與分數已獲破解

在線形文字A中，最先被破解的是計數方式。米諾斯人使用的是十進位制。一條豎線代表1，一條橫線代表10。圓圈代表100，放射光芒的圓圈代表1000[3]。這個規則光看圖像就能推測出來。

困難的是分數符號。泥板上出現了17個分數符號[8]。若計入將兩個符號疊寫的情況，則接近30個[9]。究竟哪個符號代表多少數值，長期以來一直未獲解答。

2021年波隆那大學的研究團隊著手研究了這個問題。這項研究由米歇爾·科拉查、西爾維婭·費拉拉、芭芭拉·蒙泰基、法比奧·坦布里尼與米格爾·瓦萊里奧共同進行。他們結合了觀察符號外形的傳統方法與電腦運算[1]。

研究團隊首先梳理了泥板上符號遵循的規則。他們統計了哪些符號會相鄰出現，哪些符號絕不相鄰。接著利用電腦檢驗數值的組合。候選組合接近400萬種。他們透過統計檢定以及與其他文明分數體系的比較，逐步縮小了候選範圍[1]。

最終得出的答案如下：米諾斯分數體系中最小的數值是60分之1。藉由這個體系，絕大多數以60為分母的數值都能記錄下來。代表10分之1的符號也直接沿用至線形文字B。在線形文字B中，該符號被用作計量乾穀的單位。兩種文字的記錄之間有著約250年的時間差。即使跨越了這段時差，計數法則依然留存了下來[1]。

計算法則被解開後，帳簿的結構也開始清晰。泥板遵循著將品項符號與數量配對記錄的方式。相同的格式在多處遺址中反覆出現。這意味著克里特島的各座宮殿共用著相同的會計規則[3]。即使不知道具體含義，文書的性質也能以此方式顯露出來。

地名雖能讀出，含義卻未解開

在語音方面也有所收穫。美國堪薩斯大學的約翰·揚格持續公開將線形文字A銘文轉寫為語音的清單。這種轉寫是將線形文字B中形狀相同的符號代入相同的音值[2]。

在這份轉寫中出現了許多疑似克里特島古地名的詞彙。例如克諾索斯與費斯托斯等流傳至後世的地名[2]。地名能夠對應吻合，意味著音值假說並非全盤皆錯[3]。

音值假說的局限也十分明顯。即便兩種文字中字形看似相同，也不能保證發音也相同。創造線形文字B的人借用了別人的文字來記錄自己的語言。在借用的過程中，發音有可能發生了轉變[3]。因此，轉寫文本僅作為暫定讀法使用[2]。

然而進展也就到此為止。除去地名之後，其餘詞彙的意義無法串連。即使能照著讀音讀出來，也不知道是什麼意思。語音與意義之間尚未架起橋樑[3]。

2026年的騷動

2026年夏天出現了宣稱已解讀線形文字A的發表。這是AI工程師湯姆·迪米諾提出的主張。他並非受過正式語言學訓練的人[4]。

迪米諾表示自己編寫了Python腳本來檢索兩個資料庫：收集手繪摹寫資料的GORILA，以及前述的SigLA。在最初的發表中，他宣稱為40個符號賦予了音值，並整理出包含408個詞彙的清單[4]。到了8月，他發表了一篇41頁的文章，將清單擴充至508項，同時附上了42個符號的讀法以及443件銘文的翻譯。他主張米諾斯語屬於包含希伯來語與亞拉姆語在內的閃語族[5]。

學界並未接受這項主張。因為該研究從未發表過經過同儕審查的論文[4]。據傳羅格斯大學與劍橋大學正在進行審查，但目前尚未得出結論[5]。宣稱解讀線形文字A的主張在過去百年間屢見不鮮，因此研究人員對新主張皆採取謹慎態度[6]。

同在8月，也出現了另一種詮釋。那是關於由尤克塔斯山頂峰聖所出土的石桌銘文。這是一張向神明奠酒奉獻的奠酒石桌。該觀點提出刻在石頭上的兩個名字與後來的希臘神祇有關聯。這項見解同樣停留在假說階段，因為線形文字A尚未被確切解讀[7]。

這段插曲對整個AI研究也留下了啟示。機器擅長在存在規則的地方找出規則。在如分數符號這般必須合乎邏輯的資料中取得了成果[1]。然而，有些領域是光靠規則無法觸及的。哪種語音對應哪種意義，純屬約定俗成的問題。當知曉這項約定的人全部消失時，要藉由運算重新復原這條路就會變得十分狹窄[6]。

將留存資料保持公開

研究方式也改變了。過去銘文圖像只收錄於昂貴的圖錄中，唯有取得幾本專著的人才能進行統計。如今資料已在網路上公開，即便是學生也能下載並清點符號[3]。

資料一旦公開，驗證也會變得迅速。一旦有人提出新的讀法，其他人便能利用相同的資料進行查證[6]。在2026年夏天的騷動中，驗證過程也正是如此運作[4]。提出主張的一方與審查的一方都在看著同一份資料[5]。

得益於資料公開，研究的門檻也降低了。不隸屬於大學的人也能進行符號統計。隨著新的視角增多，過去被忽略的規則有時也會被發現。相對地，未經檢驗的主張也隨之增加。公開與驗證必須相輔相成[6]。

尚未解讀的文字

AI在線形文字A中實現的成果有四項。機器將分散的資料匯集於一處；將符號的形狀與筆畫轉換為數值以便相互比對[3]；分數符號的數值則是藉由電腦運算協助人類判讀而縮小了範圍，把擁有數百萬個候選組合的問題縮減到人類得以審查的規模[1]。

機器未能達成的事同樣顯而易見。米諾斯語究竟是何種語言至今仍未揭曉。能扮演羅塞塔石碑角色的文獻至今仍在地下或根本不存在[3]。即便轉寫出了讀音，也無法得知該詞彙究竟指涉何物[6]。

賽普勒斯-米諾斯文字的案例清楚展現了這種差異。神經網路得出了既有分類產生動搖的結果[10]。即便如此，並不代表已解讀出該文字。機器所梳理的是文字的結構，而非語言的含義。

看來剩下的路有兩條。其一是出土新文物，只要有一件用兩種語言書寫的文獻就能徹底扭轉局勢。其二是更精細地整理資料，若能將每一筆畫都記錄下來，後續研究便能在此基礎上展開[3]。

線形文字A至今仍是未解讀的文字。機器並非在其面前開啟了大門，而是描繪出了鎖的形狀。開啟大門依然取決於新出土的文物與人類的判斷[6]。愛琴海的泥板至今仍未開口訴說自己的語言。

參考資料

- [1] Corazza, M., Ferrara, S., Montecchi, B., Tamburini, F., & Valério, M. (2021). The mathematical values of fraction signs in the Linear A script: A computational, statistical and typological approach. *Journal of Archaeological Science*, 125, 105214. <https://doi.org/10.1016/j.jas.2020.105214>
- [2] Younger, J. G. (n.d.). Linear A texts and inscriptions in phonetic transcription. University of Kansas. <http://people.ku.edu/~jyounger/LinearA/> (2026年9月3日檢索)
- [3] Braović, M., Krstinić, D., Štula, M., & Ivanda, A. (2024). A systematic review of computational approaches to deciphering Bronze Age Aegean and Cypriot scripts. *Computational Linguistics*, 50(2), 725-779. https://doi.org/10.1162/coli_a_00514
- [4] GreekReporter. (2026, 7월 30일). AI engineer claims to have deciphered Linear A. [greekreporter.com](https://greekreporter.com/2026/07/30/ai-engineer-claim-decipher-linear-a/). <https://greekreporter.com/2026/07/30/ai-engineer-claim-decipher-linear-a/> (2026年9月3日檢索)
- [5] GreekReporter. (2026, 8월 13일). AI engineer claims to have deciphered Linear A, identifies Minoan language as Semitic. [greekreporter.com](https://greekreporter.com/2026/08/13/ai-engineer-claim-linear-a-minoans-semitic-language/). <https://greekreporter.com/2026/08/13/ai-engineer-claim-linear-a-minoans-semitic-language/> (2026年9月3日檢索)
- [6] Adkins, J. (2026, 7월 27일). Cracking the code: can AI help us decipher ancient languages? *The Conversation*. <https://theconversation.com/cracking-the-code-can-ai-help-us-decipher-ancient-languages-288238> (2026年9月3日檢索)
- [7] GreekReporter. (2026, 8월 14일). Linear A inscription points to Minoan origins of Greek gods Dionysos and Demeter. [greekreporter.com](https://greekreporter.com/2026/08/14/minoan-linear-a-inscription-greek-gods-dionysos-demeter/). <https://greekreporter.com/2026/08/14/minoan-linear-a-inscription-greek-gods-dionysos-demeter/> (2026年9月3日檢索)
- [8] Del Freo, M. (2022). Linear A. Mnamon: Ancient writing systems in the Mediterranean. *Scuola Normale Superiore*. <https://doi.org/10.25429/sns.it/lettere/mnamon022> (2026年9月3日檢索)
- [9] Salgarella, E., & Castellan, S. (2021). SigLA: The Signs of Linear A. A palaeographical database. In *Grapholinguistics in the 21st Century 2020: Proceedings (Grapholinguistics and Its Applications, Vol. 5, pp. 945-962)*. Fluxus Editions. <https://sigla.phis.me/> (2026年9月3日檢索)
- [10] Corazza, M., Tamburini, F., Valerio, M., & Ferrara, S. (2022). Unsupervised deep learning supports reclassification of Bronze age cypriot writing system. *PLOS ONE*, 17(7), e0269544. <https://doi.org/10.1371/journal.pone.0269544>

第10章. 破解太陽週期與動物印章的密碼： 印度河文字與朗戈朗戈文字的聲音

10.1. 案件檔案

刻在印章上的四個字

印度河文明從約西元前2600年延續至約西元前1900年[1]。其發源地為今日的巴基斯坦與印度西北部。從城市的規模與數量來看，是青銅時代數一數二的文明。青銅時代是使用青銅多於鐵的古代時期。據悉，城市與村落分布的土地面積超過100萬平方公里[2]。

1920年代，哈拉帕與摩亨佐達羅展開了大規模發掘[1]。統籌這項工作的人是英國考古學家約翰·馬歇爾。他是印度考古調查局的負責人。哈拉帕遺址於1921年由達亞·拉姆·薩尼負責發掘。摩亨佐達羅遺址則於1922年由拉卡爾達斯·班納吉負責。馬歇爾整合了這些發掘成果，向世人宣告了一個新文明的存在[3]。

地下出土了筆直的磚路與排水設施。此外，還出土了大量比手掌還小的方形印章[1]。

單從城市的樣貌就能得知許多事情。道路如棋盤般筆直延伸，家家戶戶都連接著下水道。磚塊按統一尺寸燒製。測量重量的砝碼也有規格。這意味著這是一個制定計畫並遵守規則的社會。然而，他們留下的文字至今仍無法解讀。

印章是用名為滑石的軟石製成的。滑石是用指甲劃過也會留下痕跡的石頭。印章下方雕刻著動物圖案。例如獨角獸、隆背牛、大象等圖案。圖案上方的空白處則排成一系列符號[1]。

印章是用來按壓蓋印的物品。按壓在柔軟的黏土上，圖案與符號就會原樣印上去。在捆綁貨物的繩子上敷上黏土並蓋上印章，就成了封泥。在印度河各城市中，蓋有印章的泥封也一同出土[1]。

學者們將這一群符號稱為印度河文字或哈拉帕文字[1]。雖然名稱中帶有文字二字，但至今仍無法解讀。既不知其讀音，也不知其意義。被發現超過百年以來一直是如此[2]。

「無法解讀」這句話包含兩層涵義。一層是不知道符號的讀音，另一層是不知道符號的意義。目前處於兩者皆不知的狀態。

問題在於長度。單一印章上雕刻的符號通常只有四、五個[1]。即使是一面上刻下的最長記錄也不超過十七個字。就算把分刻在三面上的全部加起來，也只有二十六個字[4]。這是連課堂上用的筆記本一行都填不滿的篇幅。

刻在印章或器皿上的這類文字稱為銘文。至今收集到的印度河銘文約有3千件[5]。件數雖然不算少，但每一件都太短了。光靠收集短文很難查明句子的結構。因為根本沒有機會觀察主詞與述語是如何排列的。

跨越大海的木板符號

南太平洋中央有一座名為拉帕努伊的島嶼。我們更熟悉的名字是復活節島。距離有人居住的鄰近陸地超過2000公里。這裡也是以巨大的石人雕像摩艾而聞名的地方。這座島上也出土了刻有未知符號的木板。這批記錄在今天被稱為朗戈朗戈文字[6]。

這些符號不是用毛筆寫的字。而是用鋒利的黑曜石碎片或骨製工具在木頭上雕刻而成的。刻出的形狀類似人、鳥、魚[6]。舉起手腳的人形經常出現。

閱讀方向也很陌生。閱讀者從木板下方的左端開始。讀完一行後，將木板旋轉180度。在上下顛倒的狀態下閱讀下一行。這種每行顛倒木板的方式稱為反向牛耕體[6]。簡單來說就是倒著讀的方式。古希臘也有每行改變方向的書寫方式。那種方式不旋轉木板，只改變字的方向。

這種方式對研究人員是一大考驗。必須先確定哪一行是第一行，才能理清順序[6]。如果順序抓錯，符號相連的順序就會全盤顛倒。這是在將資料輸入電腦之前，必須由人先做出判斷的部分。

殘存的文物並不多。散落於世界各地博物館與私人收藏家手中的真品僅有26件。全部都在拉帕努伊島之外[6]。上面雕刻的符號全部加起來也只有約1萬5千字[7]。這連一本書的篇幅都不到。

文物如此稀少是有原因的。19世紀島上人口大幅銳減，木板也大量消失。殘存的物品流散至歐洲的博物館與修道院。其中有四件保存在羅馬的一處修會中[8]。

這些木板是何時製作的，也是長期爭論不休的問題。曾有人懷疑這是歐洲人踏上該島後模仿製造的。2024年，西爾維婭·費拉拉研究團隊測定了羅馬四件木板的年代[8]。他們採用測量木材中殘留碳含量來測定年代的方法。這種方法稱為放射性碳定年法。

四件當中有三件是19世紀的木材。其餘一件則檢測出是年代久遠得多的木材。比歐洲船隻抵達該島的1720年代早了200多年[8]。拉帕努伊人自行創造文字的可能性大為增加。

研究團隊附帶了一個前提條件。砍伐樹木的時間並不等於雕刻符號的時間。樹木的年齡只能說明雕刻是發生在其後。因為也有可能是後來拿老木頭削製使用的。此外，測定出較早年代的木板也僅有一件而已[8]。

這究竟算不算文字

圍繞著這兩種符號體系，學界爭論了許久。爭論並非始於解讀。而是首先追究這究竟是不是記錄語言的文字[1][9]。

2004年，史蒂夫·法默、理查·斯普羅特與麥可·威策爾共同發表了論文。他們主張印度河符號並非文字。其根據在於記錄過於簡短，且沒有任何長篇文獻。也就是說，這些符號更接近家族標記、物品的所有權標記或護身符[9]。

反方則試圖從統計中尋找答案。他們主張計算符號出現的頻率以及前後相連的規則[5]。人類的語言帶有類似習慣口吻的統計慣性。這種方法試圖測量符號中是否也存在這種習慣。

這種方法的優點在於即使不知道意義也能使用。只需將符號轉換為編號並計算其順序即可。這也是電腦擅長的工作。

阻礙解讀的三道高牆

兩種符號體系至今皆尚未被解讀。阻擋在前的高牆有三道。

沒有可供對照的文本。埃及聖書體多虧了羅塞塔石碑才得以破解。羅塞塔石碑上也以希臘文書寫了相同的內容。法國學者商博良將已知文字與未知文字並列進行比對。將相同內容以兩種語言並列書寫的文本稱為對照本。在印度河遺址與拉帕努伊島上，並未出土過這樣的記錄[2][6]。

也不知道記錄的是何種語言。符號所記錄的實際語言稱為底層語言。關於印度河文字的底層語言，達羅毗荼語系說與印度-雅利安語系說互相對立。也有主張認為它不屬於任何已知語族[2]。

最後一道高牆是資料的數量。今日的AI靠著龐大的文本茁壯成長。然而這裡僅有的，只有數千件印度河印章與26件木板[1][6]。資料不足時，AI便會編造出看似合理的謊言。這種現象稱為幻覺。

這三道高牆互相牽制。若知道底層語言，就能比對讀音數值。若有對照本，就能查明底層語言。若資料豐富，就能透過統計掌握線索。只要其中任何一個無法解決，其餘的也只能停滯不前。

2025年1月5日，印度坦米爾納德邦設立了懸賞獎金。這是在欽奈舉行的印度河文明發現100週年國際學術研討會開幕式上發布的公告[10]。內容指出，將向能夠破解印度河文字並獲得考古學專家認可者提供100萬美元獎金[10][11]。截至2026年9月，尚無已知的獲獎者。

參考資料

[1] Rao, R. P. N., Yadav, N., Vahia, M. N., Joglekar, H., Adhikari, R., & Mahadevan, I. (2009). Entropic evidence for linguistic structure in the Indus script. *Science*, 324(5931), 1165. <https://doi.org/10.1126/science.1170391>

- [2] Khan, M. A., Khurshid, S. K., & Aslam, M. (2025). Modern algorithms for ancient scripts: A review of AI-based techniques in Indus civilization research. *Cultural and Historical Heritage: Preservation, Presentation, Digitalization*, 11(2), 11-21. <https://doi.org/10.55630/KINJ.2025.110201>
- [3] Grand Valley State University. (n.d.). Archaeologist of the month: John Marshall. gvsu.edu. <https://www.gvsu.edu/archaeology/module-spotlight-view.htm?entryId=A748C17B-A353-F70E-D7422D41F7889433> (2026年9月3日檢索)
- [4] Farmer, S. (n.d.). Claims concerning the longest Indus "inscription". safarmer.com. <https://safarmer.com/indus-longestinscription/> (2026年9月3日檢索)
- [5] Yadav, N., Joglekar, H., Rao, R. P. N., Vahia, M. N., Adhikari, R., & Mahadevan, I. (2010). Statistical analysis of the Indus script using n-grams. *PLoS ONE*, 5(3), e9506. <https://doi.org/10.1371/journal.pone.0009506>
- [6] Australian Museum. (2026, 5月 28日). Kohau Rongorongo: Talking tablet. australian.museum. <https://australian.museum/learn/cultures/pasifika-collections/rapa-nui-collections/kohau-rongorongo-talking-tablet/> (2026年9月3日檢索)
- [7] Ross, L. (2026, 3月 31日). Reading Rongorongo (Version 2). Zenodo. <https://doi.org/10.5281/zenodo.19362491> (2026年9月3日檢索)
- [8] Ferrara, S., Tassoni, L., Kromer, B., Wacker, L., Friedrich, M., Tonini, F., Lastilla, L., Ravanelli, R., & Talamo, S. (2024). The invention of writing on Rapa Nui (Easter Island): New radiocarbon dates on the Rongorongo script. *Scientific Reports*, 14, 2794. <https://doi.org/10.1038/s41598-024-53063-7>
- [9] Farmer, S., Sproat, R., & Witzel, M. (2004). The collapse of the Indus-script thesis: The myth of a literate Harappan civilization. *Electronic Journal of Vedic Studies*, 11(2), 19-57. <https://doi.org/10.11588/ejvs.2004.2.620>
- [10] The Tribune. (2025, 1月 6日). Stalin announces \$1 mn prize for deciphering Indus Valley script. [tribuneindia.com. https://www.tribuneindia.com/news/india/stalin-announces-1-mn-prize-for-deciphering-indus-valley-script/](https://www.tribuneindia.com/news/india/stalin-announces-1-mn-prize-for-deciphering-indus-valley-script/) (2026年9月3日檢索)
- [11] Smithsonian Magazine. (2025, 1月 30日). Officials are offering \$1 million to anyone who can decode this ancient script. [smithsonianmag.com. https://www.smithsonianmag.com/smart-news/officials-are-offering-1-million-to-anyone-who-can-decode-this-ancient-script-180985922/](https://www.smithsonianmag.com/smart-news/officials-are-offering-1-million-to-anyone-who-can-decode-this-ancient-script-180985922/) (2026年9月3日檢索)

10.2. 科學搜查隊的追蹤：資訊熵測量儀

測量文字指紋的尺

印度河印章與拉帕努伊的木板從未相遇過。然而兩地的研究人員卻拿起了同一種工具。因為面對意義不明的一連串符號，他們能做的事情十分相似。

面對無法解讀的文字，研究人員選擇了不同的提問。他們不再追問是什麼意思，而是詢問是如何排列的。因為即使不知道意思，順序還是可以計算的。

此時使用的工具是資訊熵。熵是測量接下來會出現什麼有多難以知曉的數值。數值越大，代表接下來的內容越難預測。數值越小，代表接下來的內容幾乎已經確定。

這個概念源自通訊工程。是為了測量透過電話線發送訊號時乘載了多少資訊而製造的尺規。它不考慮意義，純粹依據順序來計算。這正是它可以套用在無法解讀的文字上的原因。

想像一下教室的座位就很容易理解。如果抽籤決定座位，就不知道身旁會坐誰。這時熵就很高。如果按座號順序入座，身旁會坐誰就已經確定了。這時熵就很低。

人類的語言介於這兩個極端之間。雖然語法約束了順序，但要說的話卻可以自由選擇。因此人類使用的語言，其熵落在中間地帶。以這個中間地帶為尺規，正是這項研究的出發點。

從韓語中也能確認這一點。在「去學校 (hakgyo-e)」之後可以接的詞有很多。「去 (ganda)」、「去了 (gatda)」、「去了一趟 (danyeowatda)」皆有可能。但並非任何詞都能接在後面。雖然有選項但已被收窄的這種狀態，正是語言的特徵。

測量印度河符號

2009年，包含拉傑什·拉奧與妮莎·亞達夫在內的研究團隊將這把尺套用在印度河符號上。他們計算了在已知前一個符號的情況下，下一個符號的不確定性[1]。在已知前一個符號時，測量下一個符號有多可預測的這個數值稱為條件熵。

他們也同時準備了對照組。一邊放置了英語、梵語、蘇美語等人類語言。另一邊則放置了非人類語言的事物。例如生物的基因序列、電腦程式碼以及隨機生成的符號序列[1]。

結果顯示，印度河符號更偏向人類語言那一邊。它比隨機符號序列更具規律，又比電腦程式碼更自由。研究團隊將這項結果發表於《Science》期刊[1]。

同一個研究團隊還提出了一種機率模型。這種方式是將前一個符號確定後出現下一個符號的機率製成表格。像這樣僅根據前一個符號來計算下一個符號的模型稱為馬可夫模型。若仔細檢視這張表，便能看出哪些符號容易出現在開頭，也能看出哪些符號容易出現在結尾[2]。

同一個研究團隊在隔年也調查了符號的頻率。他們收集了約3千件銘文，並從中挑選出保存狀態良好的1,548件。少數符號佔據了絕大部分的使用率，其餘符號則極少出現。這種形式是人類語言中常見的分佈規律[3]。

這種分佈有一個名稱，稱為齊夫定律。這是一項指稱少數常用詞佔據了文章絕大部分篇幅的規則。韓語也是如此，像「eun、neun、i、ga」這樣的詞佔據了前列。相反地，僅出現過一次的詞則非常多。印度河符號也呈現出相同的模式[3]。

在排列順序上也展現出規律。23個符號佔據了文本末尾的80%。相反地，要填滿文本開頭的80%則需要82個符號[3]。這意味著出現在末尾的符號要有限得多。這與韓語中助詞或語尾集中在句末的現象非常相似。

研究團隊還建立了根據前一個符號預測下一個符號的模型[2]。該模型以約75%的機率還原了因磨損而難以辨識的符號[3]。這是必須存在規律才能達到的成果。

不能誤解「還原」這個詞。模型填補的僅僅是符號的形狀。那個符號代表什麼聲音，依然無從得知。這只是補上了缺失的一塊拼圖，並不是讀懂了畫面。

反駁與再度反駁

這個模型也有其局限性。因為它是僅根據前面一兩個符號來預測下一個符號的方式[2]。人類的語句往往是相隔更遠的詞彙之間相互呼應。如果資料太短，就無法學習那種長距離的呼應關係。

針對這一結果，隨即出現了反駁。理查·斯普羅特於2010年發表了批評文章。他指出，僅憑熵值相似這一事實，並不能斷定其就是文字。因為非人類語言的符號體系也能產生類似的數值[4]。

拉奧研究團隊於同年在同一本學術期刊上做出了回應。他們表明自己並非證明了印度河符號就是文字。其解釋是，他們只是透過統計數據佐證了其屬於文字的可能性而已[5]。

史普勞特在2014年將批判又推進了一步。他收集了多套非人類語言的符號資料。其中包括美索不達米亞的神明象徵、北美洲的圖騰柱，以及穀倉牆上繪製的星形圖案。他甚至將天氣預報中使用的圖示也納入其中。他以相同的標準將這些資料與印度河符號進行了比對[6]。

他所設計的判定程序，將印度河符號歸類為非人類語言的一方。史普勞特寫道，單憑統計數據是遠遠不夠的。他主張必須成功完成解讀，或者另有考古證據證明該文化確實使用文字[6]。這場爭論至今尚未結束。

2026年4月發表的最新分析也面臨著同樣的疑問。阿席希·奈爾針對1,916件印度河銘文重新設計了檢驗方法。他利用電腦生成了模仿非語言符號系統的假資料。並將真實的印度河資料與假資料的統計結果並列進行了比對[7]。

印度河資料與任何一種假資料都沒有完全吻合。它落在了兩種假資料之間的中間地帶。現實中存在的任何非語言符號系統，都無法完全再現印度河的統計特徵。奈爾寫道，目前應暫緩判定這究竟是不是文字。這篇文章是在通過學術期刊審查前，先行上傳到網路上的預印本[7]。

辨認符號的眼睛

在進行統計之前，有一件事必須先做。那就是將刻在石頭上的模糊痕跡轉寫為符號。長久以來，這項工作一直依賴人類的雙眼與雙手。

轉寫比想像中要困難。印章歷經了4千年的磨損。即便同屬一個符號，外觀也會因刻印者的不同而有些許差異。有時對於該看作兩個符號還是一個符號，判斷也會產生分歧。一旦這項判斷有所改變，後續的統計結果也會隨之改變。

2017年，薩提什·帕拉尼亞潘與羅諾喬伊·阿迪卡里將這項工作交給了電腦。他們開發了一套只要輸入印章照片就能輸出符號清單的處理流程[8]。

2025年，瓦什納維·迪克西特研究團隊承接了這一脈絡。他們分別建立了從印章照片中辨識符號的神經網路，以及辨識動物圖畫的神經網路。他們將資料分成五等份進行輪流測試。兩個神經網路都展現了平均85%左右的準確率[9]。

若要減少這個問題，就必須公開判定標準並相互比對。2025年，穆扎法爾·阿里·汗研究團隊對該項工作進行了統整。他們蒐集了2009年至2025年間發表的印度河文字相關人工智慧研究，並按類型進行分類。同時也記錄了經常重複出現的錯誤[10]。

2025年，阿圖爾·夏爾馬與蘇巴吉特·羅伊·喬杜里公開了一款工具。名稱為AI-EPIGRAPHY。使用者只要輸入符號順序，它就會回傳機率較高的解讀候選項。這是一款結合了頻率分析與連鎖機率計算的工具[11]。

2026年5月，有人發表了一項重構資料庫基礎本身的嘗試。貝亞塔·梅傑西研究團隊建立了一個將罕見且未解讀文字彙整在一起的資料庫。他們將高畫質照片、轉寫符號、註釋與出處資訊整合進單一架構中。這是基於一種判斷：蒐集分散的資料是一項優先於解讀的功課[12]。

參考資料

- [1] Rao, R. P. N., Yadav, N., Vahia, M. N., Joglekar, H., Adhikari, R., & Mahadevan, I. (2009). Entropic evidence for linguistic structure in the Indus script. *Science*, 324(5931), 1165. <https://doi.org/10.1126/science.1170391>
- [2] Rao, R. P. N., Yadav, N., Vahia, M. N., Joglekar, H., Adhikari, R., & Mahadevan, I. (2009). A Markov model of the Indus script. *Proceedings of the National Academy of Sciences*, 106(33), 13685-13690. <https://doi.org/10.1073/pnas.0906237106>
- [3] Yadav, N., Joglekar, H., Rao, R. P. N., Vahia, M. N., Adhikari, R., & Mahadevan, I. (2010). Statistical analysis of the Indus script using n-grams. *PLoS ONE*, 5(3), e9506. <https://doi.org/10.1371/journal.pone.0009506>
- [4] Sproat, R. (2010). Ancient symbols, computational linguistics, and the reviewing practices of the general science journals. *Computational Linguistics*, 36(3), 585-594. https://doi.org/10.1162/coli_a_00011
- [5] Rao, R. P. N., Yadav, N., Vahia, M. N., Joglekar, H., Adhikari, R., & Mahadevan, I. (2010). Entropy, the Indus script, and language: A reply to R. Sproat. *Computational Linguistics*, 36(4), 795-805. https://doi.org/10.1162/coli_c_00030
- [6] Sproat, R. (2014). A statistical comparison of written language and nonlinguistic symbol systems. *Language*, 90(2), 457-481. <https://doi.org/10.1353/lan.2014.0031>
- [7] Nair, A. (2026, 4월 20일). How non-linguistic is the Indus sign system? A synthetic-baseline scorecard. arXiv:2604.17828. <https://arxiv.org/abs/2604.17828> (2026年9月3日檢索)
- [8] Palaniappan, S., & Adhikari, R. (2017). Deep learning the Indus script. arXiv:1702.00523. <https://arxiv.org/abs/1702.00523>
- [9] Dixit, V., Hussain, N., Basak, S., Atturu, D., Mitra, D., & Bhattacharya, U. (2025). Deep learning in archiving Indus script and motif information. *Journal of Computer Applications in Archaeology*, 8(1), 156-169. <https://doi.org/10.5334/jcaa.175>
- [10] Khan, M. A., Khurshid, S. K., & Aslam, M. (2025). Modern algorithms for ancient scripts: A review of AI-based techniques in Indus civilization research. *Cultural and Historical Heritage: Preservation, Presentation, Digitalization*, 11(2), 11-21. <https://doi.org/10.55630/KINJ.2025.110201>
- [11] Sharma, A., & Roy Chowdhury, S. (2025). AI-EPIGRAPHY: An interactive tool for computational decipherment of the Indus Valley script. *Proceedings of IndiaHCI '25*. <https://doi.org/10.1145/3768633.3770145>

[12] Megyesi, B., Rattenborg, R., Lang, B., Waldispuhl, M., & Heder, M. (2026). Building a corpus and database for rare and undeciphered scripts. Proceedings of the Fourth Workshop on Language Technologies for Historical and Ancient Languages (LT4HALA 2026) @ LREC 2026, 184-196. ELRA Language Resources Association. <https://aclanthology.org/2026.lt4hala-1.17/> (2026年9月3日檢索)

10.3. 揭曉的真相

計算月亮的木板

在拉帕努伊的木板中，有一塊名為瑪瑪麗（Mamari）的木板。這塊木板上有個外觀與其他部位截然不同的區段。一連串新月形狀的符號接連出現[1]。

德國學者托馬斯·巴特爾於20世紀中葉注意到了這個區段。他認為這是記錄太陰日期的紀錄[1]。1990年，雅克·吉重新整理了資料，提出了相同的結論[2]。

2011年，保羅·霍利進一步鞏固了這一解讀。瑪瑪麗木板上的曆法由三十個夜晚組成。這個數字與實際的月相週期吻合。也與波里尼西亞諸島的曆法一致[1]。

霍利還將大溪地和紐西蘭的曆法放在一起進行了比對。每個夜晚所取的名稱在各島嶼之間都很相似。拉帕努伊的曆法比其他島嶼落後了大約三天。這被認為是在島嶼長期彼此相距遙遠期間所產生的差異[1]。

符號的排列也遵循著規則。在新月符號之間夾雜著其他符號，將其劃分為八組。最後的兩個新月則代表看不見月亮的夜晚[1]。

同樣的排列也出現在木板之外。島上的岩畫中，有同時雕刻著海龜圖案與新月的作品。在那裡，最後兩個新月同樣位於海龜的背甲之內。這與木板將無月之夜單獨分開繪製的方式如出一轍[1]。木板與岩石共享了相同的計算法。

這項解讀在朗戈朗戈文字研究中獲得了廣泛認可。即便如此，解讀出來的也只是日期的順序。拉帕努伊人究竟如何稱呼那些夜晚，至今仍不得而知。

人工智慧給出的答案及其份量

計算月亮與計算太陽是息息相關的。日食發生在太陽與月亮重疊之時。因此，主張解讀出曆法，往往很快就會演變成主張解讀了整個天空。

2026年3月，羅根·羅斯發表了一篇預印本論文。他在個人電腦上運行語言模型來分析朗戈朗戈文字。他建立了一部由89個符號組成的字典，並聲稱以此涵蓋了近97%的殘存紀錄。他還主張瑪瑪麗木板是一張預測日食與月食的對照表[3]。

這項主張尚未經過學術界的檢驗。羅斯的文章並非通過學術期刊審查的論文。而是先行上傳到網路典藏庫的預印本[3]。預印本是任何人都能上傳的。其內容是否正確，還需要其他研究者後續進行審視與驗證。

面對這樣的主張時，必須同時檢視其檢驗方法。殘存的朗戈朗戈文字紀錄僅有26件、約1萬5千字[3]。在資料如此匱乏的情況下，很容易產生純屬巧合的吻合結果。若隨意為89個符號分配讀音，幾乎可以拼湊出任何句子。因此，學術界首先會質疑其他人以相同的方法是否也能得出相同的結果。

若是真正的解讀，就必須通過一項考驗。那就是拿出一塊在制定解讀方案時未使用過的其他木板來嘗試閱讀。在那塊木板上也必須解讀出語意通順的文句。記錄邁錫尼希臘語的線形文字B便是通過這項考驗的例子[4]。朗戈朗戈文字與印度河文字至今仍未跨過這道門檻。

機器能做的事與不能做的事

到目前為止，人工智慧在這兩種文字上所取得的成果是明確的。

機器能從模糊的痕跡中辨識符號[5]。人工用肉眼比對需要耗費數月的工作，機器能在幾小時內完成。它還能藉由前後文脈絡，復原磨損而難以辨認的符號[6]。更能將散落各處的文物照片與記錄彙整於一處以供檢索[7]。

蒐集資料的工作並不起眼。然而若沒有這項工作，後續的運算也無從談起。若照片的明暗與角度各不相同，便無法比對符號。若每位研究者的轉寫規則不同，統計結果也會產生偏差。2026年問世的資料庫，正是試圖將這一標準統一起來的嘗試[7]。

在統計方面也有所斬獲。多項研究已證實印度河符號並非隨機排列[6][8]。排列存在規則的結論已一再被重現。但這並不代表就能得出它是文字的結論。存在規則與記錄口語是兩碼子事。

然而，機器未能做到的事情依然巨大。印度河符號中每一個符號的讀音仍然無人知曉[9]。朗戈朗戈符號的讀音同樣不得而知[3]。這兩種文字目前都沒有任何一個確立字義的詞彙。印度河印章上所記錄的究竟是人名還是物品名稱，也尚未定論[9]。

原因在於資料本身的性質。若有對照文本，解讀就會容易得多。但在印度河遺跡與拉帕努伊島上都沒有這樣的文本[9]。並非沒有在缺乏對照文本的情況下被解開的先例。線形文字B就是這樣被解開的。但當時具備推測底層語言的線索，且資料也相當充裕[4]。而印度河文字甚至連底層語言是什麼都尚未確定[9]。朗戈朗戈文字則是殘存的文字量太少[3]。

2026年4月的分析再次以數據展現了這種處境。1,916件印度河銘文的統計特徵，既不符合人類語言的統計規律，也不符合非語言符號的統計規律[10]。這意味著延續了一百多年的疑問至今仍未有定論。

站在等待的一方

研究未解讀文字需要具備兩種態度。其一是將新工具運用到極致的態度；其二則是對產生的結果抱持懷疑的態度。

人工智慧擅長前者。它負責計數、比對並羅列候選項。後者則是人類的責任。人類必須判斷在這些候選項中，哪一個才合乎那個時代的常理。該判斷需要綜合考量遺址的位置、墓葬出土的隨葬品以及留存下來的口頭傳說。

拉帕努伊的曆法解讀就是一個例子。計算新月符號是運算；而將那三十個夜晚與波里尼西亞諸島的曆法進行比對則是考證[1]。當這兩項工作緊密結合時，解讀才具備了說服力。

印度河印章上的動物圖畫也是此類資料。那隻獨角獸象徵著什麼，我們至今仍不得而知。瑪瑪麗木板上的三十個夜晚亦是如此。雖然找出了順序，但我們並不知道人們在那些夜晚究竟做了些什麼。

承認「不知道」本身也有其價值。唯有劃清何處是已查明的事實、何處是推測的界線，後續的研究才能展開。若模糊了這條界線，看似合理的故事就會偽裝成事實。在研究未解讀文字的過程中，劃定這條界線與得出研究成果同等重要。

本章探討的兩種文字彼此從未有過交集。兩者在時代上相差三千年以上，居住地也分處地球的兩端。然而，研究人員面臨的困境卻幾乎完全相同：短碎的資料、缺乏對照文本、未知的底層語言。這正是同一套工具何以同時應用於兩者的原因。

坦米爾納德邦祭出的100萬美元獎金也反映了這一現況[11]。自設立獎金以來已逾一年，卻仍無人能領取。計算工具的速度雖然變快了，但解開謎團的鑰匙至今仍未出現。

參考資料

- [1] Horley, P. (2011). Lunar calendar in rongorongo texts and rock art of Easter Island. *Journal de la Societe des Oceanistes*, 132, 17-38. <https://doi.org/10.4000/jso.6314>
- [2] Guy, J. B. M. (1990). The lunar calendar of Tablet Mamari. *Journal de la Societe des Oceanistes*, 91(2), 135-149. <https://doi.org/10.3406/jso.1990.2882>
- [3] Ross, L. (2026, 3월 31일). Reading Rongorongo (Version 2). Zenodo. <https://doi.org/10.5281/zenodo.19362491> (2026年9月3日檢索)
- [4] Chadwick, J. (1958). *The decipherment of Linear B*. Cambridge University Press.
- [5] Palaniappan, S., & Adhikari, R. (2017). Deep learning the Indus script. arXiv:1702.00523. <https://arxiv.org/abs/1702.00523>

- [6] Yadav, N., Joglekar, H., Rao, R. P. N., Vahia, M. N., Adhikari, R., & Mahadevan, I. (2010). Statistical analysis of the Indus script using n-grams. PLoS ONE, 5(3), e9506. <https://doi.org/10.1371/journal.pone.0009506>
- [7] Megyesi, B., Rattenborg, R., Lang, B., Waldispühl, M., & Heder, M. (2026). Building a corpus and database for rare and undeciphered scripts. Proceedings of the Fourth Workshop on Language Technologies for Historical and Ancient Languages (LT4HALA 2026) @ LREC 2026, 184-196. ELRA Language Resources Association. <https://aclanthology.org/2026.lt4hala-1.17/> (2026年9月3日檢索)
- [8] Rao, R. P. N., Yadav, N., Vahia, M. N., Joglekar, H., Adhikari, R., & Mahadevan, I. (2009). Entropic evidence for linguistic structure in the Indus script. Science, 324(5931), 1165. <https://doi.org/10.1126/science.1170391>
- [9] Khan, M. A., Khurshid, S. K., & Aslam, M. (2025). Modern algorithms for ancient scripts: A review of AI-based techniques in Indus civilization research. Cultural and Historical Heritage: Preservation, Presentation, Digitalization, 11(2), 11-21. <https://doi.org/10.55630/KINJ.2025.110201>
- [10] Nair, A. (2026, 4월 20일). How non-linguistic is the Indus sign system? A synthetic-baseline scorecard. arXiv:2604.17828. <https://arxiv.org/abs/2604.17828> (2026年9月3日檢索)
- [11] The Tribune. (2025, 1월 6일). Stalin announces \$1 mn prize for deciphering Indus Valley script. [tribuneindia.com. https://www.tribuneindia.com/news/india/stalin-announces-1-mn-prize-for-deciphering-indus-valley-script/](https://www.tribuneindia.com/news/india/stalin-announces-1-mn-prize-for-deciphering-indus-valley-script/) (2026年9月3日檢索)

第11章. 人工智慧，重讀殘破的希臘銘文：基礎模型與 Ithaca

11.1. 案件檔案

刻在石頭上的國家記錄

古希臘人將重要的事情刻在石頭上。制定新法律時，便將其內容刻在石柱上立於廣場。徵收稅金的帳簿也刻在石頭上留存。供奉給神殿的物品清單同樣留在石頭上。刻在石頭或青銅上的這類文字稱為銘文。蒐集並研讀銘文的學問是金石學。意思是研究刻在石頭或金屬上的文字之學問[4]。

銘文大致上是那個時代的人直接留下的記錄。極少有後人抄寫與修改的痕跡。因此對歷史學家而言是珍貴的資料。然而，銘文並不總是保持最初的面貌。有些石碑在日後被抹去字跡或重新刻寫。甚至還留有磨平石頭後重新書寫的案例。雅典曾與哪座城市、在何時攜手結盟，銘文也能告訴我們[4]。

銘文是為了讓大眾觀看而製作的文字。立在廣場或神殿前，供路過的人閱讀。其用途與貼在學校走廊上的公告相似。因此，銘文上會清楚端正地刻上當天的官員姓名與日期。歷史學家便將這些姓名與日期視為線索[4]。

磨滅的文字，缺損部

問題在於時間。數千年來，石碑歷經風吹雨打。石材表面逐漸磨損，字跡的刻痕也變淺了。埋在地底的石碑則受到濕氣與微生物侵蝕。有些石碑甚至碎裂，只剩下一部分[1]。

結果，銘文上便出現了字跡消失的空白處。學者們稱這種空白為缺損部，拉丁語為 lacuna[1]。缺損部往往貫穿句子的正中央。若人名消失，就無從得知是誰所為；若數字消失，就無法得知繳納了多少；若城市名稱消失，便弄不清是與哪個國家訂立的盟約。

學者們很早以前就嘗試填補這些空白。他們尋找相同時期、內容相似的其他銘文進行比對。希臘銘文有固定的句式結構。公布法律的公文幾乎都依相同的順序排列句子。因此，只要前後文還保留著，就能推測出消失的詞彙[1]。

石碑過於沉重，無法搬進研究室。因此學者們將濕紙敷在石碑上按壓。紙張乾燥後，字跡刻痕便清楚浮現。這種紙質副本稱為拓本。學者們將拓本帶回，在書桌前解讀文字。古老的印刷圖錄中，往往只收錄抄錄下來的文字而非照片。因此後人無法看到石頭的外觀，只能單純閱讀文字[4]。

推測是有規則的。首先要計算空白處原本有幾個字。測量石頭上殘存文字之間的間距，就能大致推算出來。接著，只將字數相符的詞彙列為候選。如此填入的文字必須寫在方括號內。對於字形殘缺難以辨認的文字，則在其下方標上小點。這項約定稱為萊頓標記法。這是為了區分石頭上殘存的文字與人類的推測而制定的規則[1, 4]。

圍繞三筆劃 Sigma 的爭端

要解讀銘文，也必須知道文字是何時雕刻的。傳統的方法是透過字形來判定年代。因為字形會隨著時代產生微妙的變化[2]。

在雅典銘文中經常出現的例子是 Sigma。Sigma 是希臘字母的其中一個字母。早期雕刻的是三筆劃的 Sigma。後來四筆劃的 Sigma 得到廣泛使用。長期以來，學者們將三筆劃的 Sigma 視為基準。若發現此類文字，便認定為西元前 446 年左右之前的銘文[2]。

然而，有些學者開始懷疑這項基準。因為字體演變的速度因地區而異。有些地區長期沿用舊字體。即便在同一座城市內，每位工匠雕刻的習慣也不盡相同[2]。

因此，圍繞雅典與其他城市所簽訂的條約銘文年代，爭論持續不斷。一方主張其早於西元前 446 年左右；另一方則認為是刻於西元前 420 年代的文字[2]。西元前 420 年代正值伯羅奔尼撒戰爭如火如荼之際。這場戰爭始於西元前 431 年。年代若相差數十年，歷史的脈絡也會隨之改變。因為雅典以武力壓制同盟城市的時間點將整個位移。

填補空白伴隨著風險。因為一位學者填補的詞彙，到了下一代可能會被當作原文來閱讀。只要印本漏掉了方括號，推測就可能演變成事實[1]。因此，金石學家總會在自己的推測處留下標記。

同時也必須查明銘文最初立於何處。希臘各城邦使用各自不同的方言。即使是同一個詞，各城市的拼法也不盡相同[2]。學者們透過這些拼寫差異與官職名稱來推測出土地點[4]。出土地點即是該石碑出土的地方。

單憑人力無法讀盡

有待解讀的銘文不斷堆積。帕卡德人文學會蒐集的希臘語銘文達 178,551 件[2]。羅馬時代的拉丁語銘文資料也超過 176,000 件[3]。相較之下，全球能夠解讀銘文的學者卻寥寥無幾[4]。

單是整理數據就絕非易事。將這些資料整理成機器可讀取的格式後，僅剩下 78,608 件。超過一半的資料因過於簡短或保存狀況太差而被剔除。光是挑選提供給機器的訓練資料，就是一項龐大的工程[2]。

填補空白對人類而言同樣困難。有一項測驗是讓受過專業訓練的金石學家填補文字遭到磨滅的銘文。人類給出的第一個答案正確率僅有 25%[2]。這意味著四次中有三次是錯誤的。在更早之前的實驗中，人類填補文字的錯誤率則高達 57.3%[1]。

莎草紙的情況更為嚴峻。紙莎草是生長在埃及水邊的一種植物名稱。古人將這種草的莖髓切成薄片，橫豎交錯重疊後敲打並壓平曬乾。這樣製成的書寫材料也稱為莎草紙。在埃及乾燥的沙地中出土了數十萬張。至今尚未有人解讀的希臘語莎草紙據估計達一百萬件[5]。這些紙上寫著稅收收據、書信與契約書，是完整記錄古人日常生活的珍貴記錄。

解讀工作並非打開箱子就能完成。因為紙張破碎且墨跡已經磨滅。必須先進行保存維護並拼湊散落的碎片，之後才能判讀文字。然而，能勝任這一切工作的人力嚴重匱乏[4]。

資料的讀取方式也是個問題。長期以來，銘文資料僅以書籍中的文字清單形式累積，並未同時收錄石頭的照片或材質資訊[4]。這意味著能夠輸入機器的資訊僅有文字。

發掘工作至今仍在持續。每年都有新的石碑與新的莎草紙出土。需要解讀的資料不斷增加，但能解讀的人手卻毫無增長[4]。兩者之間的差距逐年擴大。

因此，許多記錄未能被解讀，依然留在庫房的箱子中。石頭與紙張俱在，卻缺少能重現其上文字的人。為了解開這一瓶頸，研究人員引進了機器[4]。

參考資料

- [1] Assael, Y., Sommerschild, T., & Prag, J. (2019). Restoring ancient text using deep learning: a case study on Greek epigraphy. In Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP) (pp. 6368-6375). Association for Computational Linguistics. <https://doi.org/10.18653/v1/D19-1668>
- [2] Assael, Y., Sommerschild, T., Shillingford, B., Bordbar, M., Pavlopoulos, J., Chatzipanagiotou, M., Androutsopoulos, I., Prag, J., & de Freitas, N. (2022). Restoring and attributing ancient texts using deep neural networks. *Nature*, 603(7900), 280-283. <https://doi.org/10.1038/s41586-022-04448-z>
- [3] Assael, Y., Sommerschild, T., Cooley, A., Pavlopoulos, J., Shillingford, B., Herms, B., Suresh, P., Maynard, B., Grayston, J., Wulgaert, R., Prag, J., Mullen, A., & Mohamed, S. (2025). Contextualizing ancient texts with generative neural networks. *Nature*, 645(8079), 141-147. <https://doi.org/10.1038/s41586-025-09292-5>
- [4] Sommerschild, T., Assael, Y., Pavlopoulos, J., Stefanak, V., Senior, A., Dyer, C., Bodel, J., Prag, J., Androutsopoulos, I., & de Freitas, N. (2023). Machine learning for ancient languages: A survey. *Computational Linguistics*, 49(3), 703-747. https://doi.org/10.1162/coli_a_00481
- [5] Reply. (2026, 3월 19일). Austrian Academy of Sciences is developing the Ancient Greek AI "Apollo" with Mistral AI and Reply. <https://www.reply.com/en/newsroom/news/austrian-academy-of-sciences-is-developing-the-ancient-greek-ai-apollo-with-mistral-ai-and-reply> (2026年9月3日檢索)

11.2. 科學鑑識隊的追蹤：古代文本自動完成裝置 Ithaca 與 Apollo

藉由填空題學習的機器

人工智慧解讀古文字的方法，與語文測驗中的填空題非常相似。首先讓機器看大量完整的銘文，接著遮蔽其中幾個字，讓它猜測該填入什麼內容。如果答案錯誤，便微調機器內部的數值[4]。反覆進行這個過程數百萬次後，機器就能掌握希臘語句子的行文習慣。

透過這種方式學習的機器並不會死記詞彙，而是以機率記住某個詞後面常接續哪些詞[4]。因此，即使面對從未見過的銘文，也能推導出適合填入空白處的候選詞彙。

從 Pythia 到 Ithaca

2019 年，研究人員推出了名為 Pythia 的模型。該模型由 Yannis Assael、Thea Sommerschield 與 Jonathan Prag 共同打造。Pythia 是一個用於預測希臘語銘文中受損文字的神經網路[1]。神經網路是一種透過觀察資料自主學習規則的計算架構，模擬了人類大腦神經元之間的連結。

Pythia 同時從字元與詞彙層級進行讀取。因為要處理詞彙毀損了一半的資料，需要雙管齊下的視角。當人類填補文字的錯誤率為 57.3% 時，Pythia 的錯誤率僅為 30.1%。而正確答案落在 Pythia 前 20 個候選名單中的比例高達 73.5%[1]。

2022 年，Ithaca 問世。它由 Google DeepMind 與多所大學共同研發，包括牛津大學、威尼斯卡福斯卡里大學及雅典經濟與商業大學。研究成果於 2022 年 3 月發表於學術期刊《自然》（Nature）[2]。其名稱取自史詩《奧德賽》中的伊薩卡島（Ithaca），蘊含著歷經漫長漂泊後重返家園之意。

Ithaca 如何進行判斷

Ithaca 採用稱為 Transformer 的架構。Transformer 是一種能衡量句子內部各詞彙之間關聯程度的神經網路。文字磨損處會被轉換為代表空白的符號輸入，模型則計算替換該符號之文字的機率[2]。

Ithaca 主要執行三項任務：填補缺失文字、推斷銘文原本位於何處，以及推測銘文刻寫的時間。這三項任務並非分開學習，而是一併訓練。方言與拼寫習慣是辨識地點的線索；用詞習慣與句式結構則是判斷時代的線索。當單一神經網路同時處理這些線索時，彼此便能相輔相成[2]。

Ithaca 並不觀看石頭的照片，而是僅閱讀轉錄下來的文字進行判斷[2]。這與過去藉由筆劃形態來判定年代的方法截然不同。

它輸出年代的方式也與眾不同。Ithaca 不會直接指定單一具體年份，而是將西元前 800 年至西元 800 年以每十年為單位劃分，並賦予每個區間一個機率，以圖表呈現[2]。學者則能根據圖表中的高峰，判斷哪一個時期最具可能性。

用於訓練的資料來自帕卡德人文學會的銘文典藏。前面提及的 78,608 件即為該批資料。單是出現十次以上的詞彙，就高達 35,884 個[2]。

研究團隊在開發 Ithaca 時設定了三大目標：與學者協同合作、輔助學者判斷，以及展示判斷過程。他們的目標並非取代人類的機器，而是打造能夠融入學者既有工作流程的工具[2]。

Ithaca 不會只給出單一答案，而是會連同機率列出多個填補空白的候選答案。此外，它還會在螢幕上標示出是根據哪些文字做出此判斷。學者能親眼確認機器是依據什麼給出答案。最終的判斷權仍保留給人類。研究團隊亦公開了 Ithaca 的程式碼與資料[2]。

二十個候選答案的意涵

機器給出的答案並非只有一個。它會依機率高低將候選答案排列排序。首位候選答案正確的比率稱為 Top-1 準確率；而正確答案落在前二十個候選名單內的比率則是 Top-20 準確率[1]。

這項區分至關重要。學者需要的往往不是單一正確答案，而是一份值得審視的候選清單。只要前二十個名單中包含正確答案，學者就能從中做出抉擇[1]。有時甚至能從清單中發現自己未曾想到的詞彙[2]。

跨足拉丁語的 Aeneas

同一個研究團隊將研究方向拓展至拉丁語銘文。2025 年 7 月，他們在《自然》期刊上發表了名為 Aeneas 的模型[3]。該名稱取自羅馬建國傳奇的主角埃涅阿斯（Aeneas）[5]。

Aeneas 透過超過 176,000 件拉丁語銘文資料進行學習。這是整合了羅馬、海德堡與克勞斯·斯拉比（Clauss-Slaby）三大銘文資料庫的資料。在缺損十個字元以內的情況下，正確答案落在前 20 個候選名單內的比例達 73%；即便在不知缺損字數的情況下，也達到了 58%[3]。在不知缺損長度的情況下推斷答案，對人類而言同樣是一大難題。

它還能準確推斷出銘文出自羅馬的哪一個行省。行省是羅馬統治下的地方行政區。在 62 個選項中，其命中率達到 72%[5]。至於雕刻年代，則將歷史學家劃定的範圍進一步縮小至平均 13 年之內[3, 5]。

Aeneas 具備一項新功能：能夠結合銘文的照片進行綜合判斷[3]。這意味著它不僅檢視轉錄文字，還同時觀察石頭的外觀樣貌。它還能搜尋具備相似文句的其他銘文並列呈現，在短短數秒內找出學者過去需耗費數日才能查得的相似案例[5]。

研究團隊將這項工具試用於羅馬皇帝奧古斯都的記錄。Aeneas 針對該文銘刻的時期同時給出了兩個區間。其中一個是西元前10年至西元前1年之間。另一個則是西元10年至西元20年之間，且機率更高[5]。學術界原本分歧的兩種主張，就這樣直接呈現在機率圖表上。

研究團隊讓23位歷史學家試用這項功能。學者們回答表示，在每十件碑文中就有九件獲得了新的研究靈感[3, 5]。該模型已公開發布，供所有人使用[5]。

完整學習希臘語的 Apollo

2026年3月，奧地利科學院公布了一款新模型。名稱為 Apollo。由人工智慧企業 Mistral AI 與資訊科技企業 Sail Reply 共同打造[7]。如果說 Ithaca 是處理單件碑文的工具，那麼 Apollo 的目標則是學習整門希臘語的模型[6]。

這類模型被稱為基礎模型。意思是為了廣泛應用於多種任務而經過全面學習的基礎模型。它並非只為執行單一任務而製作的工具[4]。無論是碑文修復、文獻檢索還是字跡辨讀，都使用相同的底層基礎[6]。

Apollo 學習的資料約為6億字詞規模的希臘語文獻。此外還加入了數萬件已出版的碑文與紙草文獻[6, 7]。研究團隊表示，這在迄今所蒐集的希臘語資料中規模名列前茅[7]。

奧地利國家圖書館收藏了大量紙草文獻。僅該圖書館就保存了數萬件[7]。其中夾雜著許多尚未被解讀的文件。

領導這項研究的是奧地利考古研究所的安娜·多爾加諾夫（Anna Dolganov）。她表示，全球還有上百萬件尚未被任何人解讀過的希臘語紙草文獻[7]。Apollo 負責從這些資料中修復文字並檢索內容。解讀手寫字跡也包含在目標之中[6, 7]。研究團隊表示，這項工作可能只需數小時而非數年即可完成[7]。

參考資料

- [1] Assael, Y., Sommerschield, T., & Prag, J. (2019). Restoring ancient text using deep learning: a case study on Greek epigraphy. In Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP) (pp. 6368-6375). Association for Computational Linguistics. <https://doi.org/10.18653/v1/D19-1668>
- [2] Assael, Y., Sommerschield, T., Shillingford, B., Bordbar, M., Pavlopoulos, J., Chatzipanagiotou, M., Androutsopoulos, I., Prag, J., & de Freitas, N. (2022). Restoring and attributing ancient texts using deep neural networks. Nature, 603(7900), 280-283. <https://doi.org/10.1038/s41586-022-04448-z>
- [3] Assael, Y., Sommerschield, T., Cooley, A., Pavlopoulos, J., Shillingford, B., Herms, B., Suresh, P., Maynard, B., Grayston, J., Wulgaert, R., Prag, J., Mullen, A., & Mohamed, S. (2025). Contextualizing ancient texts with generative neural networks. Nature, 645(8079), 141-147. <https://doi.org/10.1038/s41586-025-09292-5>

- [4] Sommerschild, T., Assael, Y., Pavlopoulos, J., Stefanak, V., Senior, A., Dyer, C., Bodel, J., Prag, J., Androutsopoulos, I., & de Freitas, N. (2023). Machine learning for ancient languages: A survey. Computational Linguistics, 49(3), 703-747. https://doi.org/10.1162/coli_a_00481
- [5] Google DeepMind. (2025, 7월 23일). Aeneas transforms how historians connect the past. deepmind.google. <https://deepmind.google/blog/aeneas-transforms-how-historians-connect-the-past/> (2026年9月3日檢索)
- [6] Mistral AI. (2026). Austrian Academy of Sciences. mistral.ai. <https://mistral.ai/customers/austrian-academy-of-sciences/> (2026年9月3日檢索)
- [7] Reply. (2026, 3월 19일). Austrian Academy of Sciences is developing the Ancient Greek AI "Apollo" with Mistral AI and Reply. reply.com. <https://www.reply.com/en/newsroom/news/austrian-academy-of-sciences-is-developing-the-ancient-greek-ai-apollo-with-mistral-ai-and-reply> (2026年9月3日檢索)

11.3. 揭開的真相

62%與72%

研究團隊以具體數字驗證了 Ithaca 的實力。這是一項故意遮蔽碑文字符並要求進行填補的測試。他們也讓受過專業訓練的金石學家接受了同樣的測試[1]。

人類獨自解題時，首選答案正確的比例為25%。而 Ithaca 獨自作答時則達到了62%。這是超過兩倍的差距[1]。

這項測試還有第三種條件。即由人類參考 Ithaca 給出的候選名單共同解題。結果準確率提升到了72%[1]。這一數值既高於人類獨自作答，也高於機器獨自運作的結果。

這個數字正是這項研究的核心所在。機器憑藉讀完78,608件碑文的記憶來計算機率。人類則清楚該碑文處於何種歷史脈絡之中。當機器列出20個候選選項時，人類能以歷史知識加以篩選。當兩者結合時，準確率達到了最高[1, 2]。

研究也驗證了推斷地點與時期的能力。Ithaca 將希臘世界劃分為84個區域進行學習。它能以71%的準確率推斷出土地不明碑文的原始出處。至於銘刻年代，它給出的數值與歷史學家判定的範圍平均僅相差29.3年。過去單憑肉眼判斷字體形態的方法，誤差有時甚至長達一百年[1]。

西元前421年的答案

Ithaca 也為長年以來的爭論給出了答案。這就是前面提到的雅典條約碑文的年代問題。這些碑文過去因帶有三筆劃的 Sigma，而被認為早於西元前446年左右[1]。

研究團隊將這些碑文輸入 Ithaca。Ithaca 並非依據字形，而是藉由文本本身來進行判斷。它會分析哪些詞彙常與哪些詞彙搭配使用。也會考量哪些官職名稱會與特定表述同時出現[1]。

Ithaca 給出的平均年代為西元前421年。這比過去的標準晚了一個世代。這一數值與近期出現的其他證據趨向一致。與近代學者重新斷代的年代平均僅相差5年。而按舊標準判定的年代則相差了27年之久[1]。

機器並未獨自改寫歷史。而是為長期分歧的兩種主張之一增添了份量[1, 2]。學者們也因此重新審視僅憑單一字體就敲定年代的慣例。單單一項工具，就改變了長達百年的論爭走向。

機器編造出看似合理的句子

儘管取得了這些成果，依然存在遺留的問題。那就是學習資料存在偏頗。希臘語碑文資料中，雅典的文獻佔了多數。年代也多集中於古典時期。資料較少的邊陲地區碑文，其準確度也相應較低[2]。

另一個問題是難以得知機器為何做出這樣的回答。神經網路內部包含數百萬個數值。雖然它能顯示更關注哪些字，卻無法用語言解釋選擇該答案的理由[2]。

當缺漏空白變長時也會產生問題。填補幾個字尚能應付自如，但若整段多個詞彙完全消失，難度就極高。在 Aeneas 的測試中，若不知道缺失長度，成績也會從73%驟降至58%[3]。此時模型可能會生成文法合規、看似合

理但實際上根本不存在的句子。這種憑空捏造的答案被稱為「幻覺」。意思是機器像真的一樣編造出來的錯誤答案。在處理古代文獻時，幻覺尤其危險。因為原本不存在的歷史可能會被當成事實而固定下來[2]。

2026年8月發表的一項研究正面探討了這項局限。這是在多種模型與條件下修復受損古文件的實驗。研究團隊得出結論：無法將修復工作完全交由自動化處理。根據填補文件部位的不同，成績差異極大。固定套話反覆出現的文件能填補得很好。但每個人用詞各異的句子則難以處理。是否知曉缺失長度，結果也有所不同。研究團隊指出，應當將機器打造成協助古文字學家的工具[4]。

研究人員為了降低這種風險而制定了規則。機器填補的文字也加上方括號。機率較低的候選選項則由學者進行篩選。傳統金石學家恪守的標註約定，同樣直接套用於機器上[2]。

因此，Ithaca 給出的答案並非歷史事實。而是值得檢驗的假說清單。唯有經學者與其他資料比對確認後，才能真正成為史料[2]。史料是闡明歷史的根據資料。

七年間發生的事

這個領域的變化速度可以用數字來衡量。2019年，Pythia 將人類57.3%的錯誤率降至30.1%[5]。2022年，Ithaca 將首選答案的準確率提升至62%[1]。2025年，Aeneas 跨足拉丁語，將年代誤差縮小至平均13年[3]。2026年，學習整門希臘語的 Apollo 正式啟動[6]。

七年之內，涵蓋的目標語言增加了，準確率也提升了。處理的資料也從文字拓展到了圖像[3]。在這段期間，唯一一件事未曾改變。那就是由人類做出最終判斷的原則[2]。

圖像與語境，下一步

研究領域仍在持續拓展。2026年發表的一篇論文中，研究人員使用希臘語資料重新訓練了廣泛應用的大型語言模型。這項實驗採用了 Meta 公開的 Llama 3.1 模型。在紙草文獻修復中，首選答案正確的比例為73.5%。在碑文修復中則達到了63.7%。正確答案落在前20個候選選項中的比例分別為86.0%與83.0%。這意味著即便不是專為碑文設計的模式，也能勝任這項工作[7]。

還出現了還原文字圖像的研究。2026年4月發布的 Mesa 是一種填補碑文照片中磨損筆劃的方法。它以同塊石碑上保存完好的文字為範本，復原磨損處的痕跡[8]。填補文句的模式與填補圖像的模式正並駕齊驅地發展。

失落的悲劇作品也確實重現於世。2022年，在埃及費拉德爾菲亞遺址出土了一張紙草紙。這是一張推測書寫於西元3世紀的紙草。2024年，學者們解讀了其內容並對外發表。那是歐里庇得斯所作悲劇《波呂伊多斯》與《伊諾》中的97行對話。直到那時為止，這兩部作品僅能透過簡短的引文與故事梗概為人所知[9]。

這項修復是由人類的手和眼睛完成的，並非機器的功勞。全球還有上百萬件這樣的紙草文獻。光憑人類的肉眼無法讀完如此堆積如山的紙草。像 Apollo 這樣的模型所鎖定的目標正是此處[6]。

隨著工具的公開，研究方式也發生了改變。Ithaca 與 Aeneas 的程式碼和資料皆已公開[1, 10]。即使身處大學研究室之外，任何人都能在網頁上試用。只要上傳一張碑文照片，系統便會給出候選名單、地圖以及年代圖表[10]。在過去，這往往是需要耗費數年的比對工作[3]。

尚待完成的工作依然繁多。光是希臘語紙草就有上百萬件尚未被解讀[6]。碑文資料在地區與年代上的分佈並不均勻[2]。機器難以妥善處理的長篇殘缺部分也依然存在[3]。這些問題必須透過更優質的資料與更多的驗證來逐一減少[4]。

機器給出候選選項，由人類進行挑選。將 Ithaca 從62%提升至72%的方法就是解答[1]。要讓沉默的石碑與紙草重新開口說話，兩者必須並肩攜手。

參考資料

- [1] Assael, Y., Sommerschield, T., Shillingford, B., Bordbar, M., Pavlopoulos, J., Chatzipanagiotou, M., Androutsopoulos, I., Prag, J., & de Freitas, N. (2022). Restoring and attributing ancient texts using deep neural networks. *Nature*, 603(7900), 280-283. <https://doi.org/10.1038/s41586-022-04448-z>

- [2] Sommerschield, T., Assael, Y., Pavlopoulos, J., Stefanak, V., Senior, A., Dyer, C., Bodel, J., Prag, J., Androutsopoulos, I., & de Freitas, N. (2023). Machine learning for ancient languages: A survey. *Computational Linguistics*, 49(3), 703-747. https://doi.org/10.1162/coli_a_00481
- [3] Assael, Y., Sommerschield, T., Cooley, A., Pavlopoulos, J., Shillingford, B., Herms, B., Suresh, P., Maynard, B., Grayston, J., Wulgaert, R., Prag, J., Mullen, A., & Mohamed, S. (2025). Contextualizing ancient texts with generative neural networks. *Nature*, 645(8079), 141-147. <https://doi.org/10.1038/s41586-025-09292-5>
- [4] Zhang, S., Caraffa, E., Zuffrano, A., Modesti, M., & Colavizza, G. (2026, 8월 28일). Text restoration of ancient documents with language models (arXiv:2608.28170). arXiv. <https://arxiv.org/abs/2608.28170> (2026년 9월 3일 검색)
- [5] Assael, Y., Sommerschield, T., & Prag, J. (2019). Restoring ancient text using deep learning: a case study on Greek epigraphy. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)* (pp. 6368-6375). Association for Computational Linguistics. <https://doi.org/10.18653/v1/D19-1668>
- [6] Reply. (2026, 3월 19일). Austrian Academy of Sciences is developing the Ancient Greek AI "Apollo" with Mistral AI and Reply. reply.com. <https://www.reply.com/en/newsroom/news/austrian-academy-of-sciences-is-developing-the-ancient-greek-ai-apollo-with-mistral-ai-and-reply> (2026년 9월 3일 검색)
- [7] Cullhed, E. (2026). Instruction-tuning pretrained causal language models to restore ancient Greek papyri and inscriptions. *Digital Scholarship in the Humanities*, 41(1), 69-76. <https://doi.org/10.1093/dllc/fqaf131>
- [8] Toulatzis, V., Theodoridou, S., & Fudos, I. (2026, 4월 19일). MESA: A training-free multi-exemplar deep framework for restoring ancient inscription textures (arXiv:2604.17390). arXiv. <https://arxiv.org/abs/2604.17390> (2026년 9월 3일 검색)
- [9] O'Grady, E. (2024, 10월 11일). Unearthed papyrus contains lost scenes from Euripides' plays. *The Harvard Gazette*. <https://news.harvard.edu/gazette/story/2024/10/unearthed-papyrus-contains-lost-scenes-from-euripides-plays/> (2026년 9월 3일 검색)
- [10] Google DeepMind. (2025, 7월 23일). Aeneas transforms how historians connect the past. deepmind.google. <https://deepmind.google/blog/aeneas-transforms-how-historians-connect-the-past/> (2026년 9월 3일 검색)

第12章. 永不焚毀的虛擬圖書館： 超高解析度多光譜檔案庫與未來

12.1. 案件檔案

亞歷山大圖書館消失的方式

埃及亞歷山大港曾有古代世界最大的圖書館。托勒密王朝於西元前3世紀左右建立。國王們試圖將世上的書籍聚集於一處。進港船隻上的書籍會被扣留並加以抄寫。抄本歸還給原主，原本則由圖書館收藏。古老文獻記載當時的卷軸多達數十萬卷 [1]。

許多人認為這座圖書館是因一場大火而消失。然而研究人員認為它是逐漸衰落的。西元前145年，國王驅逐了外國學者。隨著學者離去，維護書籍的人手也一同離開。到了羅馬時代，建築多次遭到破壞。購買新書的資金也中斷了 [1]。

古代的文字本就難以流傳至今。當時的抄寫員比起創作新文章，更常做的是抄寫現有的書籍。常被翻閱的紙草卷軸過了大約一百年便會破損難讀。在乾燥的埃及地下有時能保存得更久。因此必須有人及時重新抄寫，文字才能延續下去。一旦停止抄寫，那本書便會從世上消失。

我們今天閱讀的古代文獻只是那個時代文字的一小部分。僅留下書名而正文已佚失的書籍要多得多。古希臘悲劇作家所撰寫的作品大多也是這樣消失的。留存至今的書籍，都是運氣好而在多處散布有抄本的作品。

破壞至今仍在持續

摧毀紀錄的力量並非只存在於過去。戰爭至今仍不斷摧毀圖書館與檔案館。聯合國教科文組織正在逐一核實並統計烏克蘭境內遭到破壞的文化設施。截至2026年7月29日，已確認的受損地點達545處。其中包括43處博物館、22處圖書館、5處檔案館。宗教設施與古老建築也有數百處一同受損 [2]。

戰爭不僅僅是摧毀文物。在缺乏管理的文物庫房中，文物也可能被偷偷轉移盜走。失去出土紀錄的文物會大幅喪失研究價值。因為一塊泥板出土於哪個房間、與哪些泥板存放在一起，這本身就是重要的情報資訊。

氣候同樣威脅著紀錄。海邊的遺跡受海浪侵蝕而逐漸崩壞。地下的文物若遇降雨增加或地下水上升，損壞速度會更快。甚至有些地方的消逝速度比發掘速度還要快。

受損的文物並不意味著完全消失。只要殘留有形態，就有辦法加以重建復原。因此，在意外發生前預先完成測量至關重要。在持續發生戰爭的地區，正進行著以3D方式測量建築與文物的作業。

每種材質的壽命各不相同

承載紀錄的材質各自以不同的方式損壞。美索不達米亞的泥板是透過壓印黏土刻上文字的板子。宮殿焚毀時產生的高溫將泥板如陶瓷般堅硬地燒製成型。未曾遇火的泥板一旦受水浸泡，便會重新化為泥土。

泥板還有另一個局限。抄寫員會將用過的泥板浸水抹平表面後重複使用。因此年代久遠的泥板大多重新變回了泥土。這也是為何現存記錄往往只剩下王國滅亡前最後幾年的原因。因為只有宮殿被焚毀時留在書庫中的最後一批泥板才被燒得堅硬。

紙草、皮革與紙張則更加脆弱。這些材料是源自植物與動物的有機物。在濕熱環境中往往撐不過幾百年便會腐爛。羅馬士兵所使用的蠟板上，蜂蠟早已蕩然無存。留下來的僅有鐵筆劃透至木板表面所留下的淺淺刮痕。

堅硬的材料也未必安全。刻在龜腹甲與牛肩胛骨上的甲骨文，在地下也因受壓而碎裂。今天留存的甲骨多半是比手掌還小的碎片。甚至有一句話分散在三塊碎片上，分別被存放在不同庫房的情況 [3]。

大火所拯救的圖書館

紀錄存留下來的的方式有時出人意料。亞述國王亞述巴尼拔在尼尼微建了一座大型書庫。他自西元前669年左右開始統治國家。書庫中收集了神話、占卜、醫學與行政文書 [4]。

約在西元前612年，尼尼微遭到焚毀。如果是紙質書籍，當時就會化為烏有。然而泥板經火燒烤後，反而變得更加堅硬。這些泥板在1840年代至1930年代期間歷經多次發掘。如今大英博物館收藏了3萬多件此類碎片 [4]。

然而，出土時的狀態並非完好無損。當城市陷落時，泥板遭到蓄意砸碎並丟棄在宮殿各處。發掘人員見到的不是井然有序的書架，而是散落一地的碎片堆。同一部作品的前後部分出現在不同房間是常有的事 [4]。重新拼合這些碎片，是當今計算考古學的一大課題。計算考古學是利用電腦研究古代文物的領域。

已經重獲新生的文字

進行數位化也有其瓶頸。世上的文物遠遠多過人力所能應付的數量。前述的尼尼微碎片堆就是一個例子 [4]。逐件拍攝並附上解讀出的文字需要漫長的時間。因此需要機器來提升人類的處理速度。

拯救逐漸消逝的文物主要有兩個方向。一是長久保護物品本身的保存工作；另一種則是在不觸碰物件的情況下解讀內部並轉移到電腦中。後者的方法已經多次取得成果。

在以色列恩戈地出土的卷軸被火燒成了炭塊。如果試圖用手展開，它就會直接碎裂。研究團隊利用不需切開物體即可檢視內部的 CT 掃描透視內部。電腦透過數學方法追蹤捲曲的皮革層並將其展開為平面。在上面顯現出含有金屬成分的墨水痕跡，並成功解讀出《利未記》的開篇 [5]。

被維蘇威火山灰掩埋的赫庫蘭尼姆卷軸也迎來了突破。2026年6月25日，研究團隊宣布已從頭到尾完整解讀了一卷卷軸。該卷軸的編號為 PHerc. 1667。其內容包含了二十二欄希臘語篇幅的倫理論述 [6]。作為挑戰任務指定的十三卷密封卷軸也設立了新的獎金。截止日期為2027年6月25日 [7]。

虛擬圖書館並非只是堆滿照片的倉庫。文物的形貌、墨水成分以及解讀出的文字必須整合繫結在一起。如此一來，日後出現新方法時，才能對過去的資料重新進行運算分析。即使原件消失，若要讓驗證得以延續，這種整合資料庫是不可或缺的。

拍攝與判讀需要耗費資金與人力。決定優先拍攝哪些文物也並非易事。位於危險地區的文物光是拍攝本身就很困難。因此各國都會排定優先順序來進行這項工作。

這些事件留下了一個疑問：即便物品最終損毀，其中的文字是否仍能留存？而尋找這個答案的場所，正是虛擬圖書館。

參考資料

- [1] Mark, J. J. (2023, 7월 25일). Library of Alexandria. World History Encyclopedia. https://www.worldhistory.org/Library_of_Alexandria/ (2026年9月3日檢索)
- [2] UNESCO. (2026). Damaged cultural sites in Ukraine verified by UNESCO. UNESCO. <https://www.unesco.org/en/ukraine-war/damaged-cultural-sites> (2026年9月3日檢索)
- [3] Ma, Y. R. (2026). Towards Computational Chinese Paleography. arXiv:2601.06753. <https://arxiv.org/abs/2601.06753> (2026年9月3日檢索)
- [4] The British Museum. What was Ashurbanipal's Library? The British Museum. <https://www.britishmuseum.org/research/projects/what-was-ashurbanipals-library> (2026年9月3日檢索)
- [5] Seales, W. B., Parker, C. S., Segal, M., Tov, E., Shor, P., & Porath, Y. (2016). From damage to discovery via virtual unwrapping: Reading the scroll from En-Gedi. Science Advances, 2(9), e1601247. <https://doi.org/10.1126/sciadv.1601247>
- [6] Vesuvius Challenge. (2026, 6월 25일). An entire Herculaneum scroll has been read for the first time. scrollprize.org. <https://scrollprize.org/firstscroll> (2026年9月3日檢索)
- [7] Vesuvius Challenge. (2026). Open Prizes. scrollprize.org. <https://scrollprize.org/prizes> (2026年9月3日檢索)

12.2. 科學鑑識的追蹤：知識圖譜與光譜分析

串聯文物的龐大地圖

在過去的研究室中，文物的形狀與文字是分開處理的。形狀由文物修復科學家測量，文字則由文獻學家解讀。兩者的成果被放進了不同的抽屜裡。因此即使對同一件文物進行了兩次研究，也無法將結果整合在一起。

計算考古學將這兩者呈現在同一張地圖上。那張地圖的名字就叫知識圖譜。知識圖譜是用節點與連線繪製資訊的圖象。節點代表事物，連線則代表事物之間的關係。

在這張地圖上，一塊泥板就是一個節點。書寫該泥板的抄寫員是一個節點，泥板出土的城市也是一個節點。墨水成分與3D外形資訊也各自成為節點。泥板上記錄的人名與物品名稱也被提取為節點。電腦則尋找連接這些節點的連線並加以學習。

一旦地圖建立，提問的方式便截然不同。過去是拿著一塊泥板詢問上面寫了什麼字。如今則可以詢問使用這種墨水的其他文書在哪裡。甚至能一次查出同一位抄寫員書寫的泥板出土於哪些城市。

2026年2月，波蘭亞捷隆大學研究團隊公開了這樣的一張地圖。這張地圖中包含了大約320萬條人物、事件與物品的關係。研究團隊讓機器推薦關係，並能進一步解釋其原因。專家們閱讀這些解釋並驗證其是否實用 [1]。

測量形狀的層級

地圖的第一層是文物的外部形態。古老的石碑與泥板因表面磨損，字跡凹槽會變淺。僅憑肉眼很難看出是否殘留凹痕。因此研究人員使用3D掃描器掃描其表面。

掃描器將表面轉換為數百萬個點儲存起來。電腦將這些點連成三角形網格。在該網格上，可以從多個方向投射虛擬光線。只要改變光線方向，淺淺的凹槽也會產生陰影，從而使文字顯現出來。

機器還能學習凹槽的斷面形狀。人使用工具刻出的痕跡，其斷面均勻且方向一致。自然形成的裂紋則並非如此。學習了這種差異的模型便能區分文字與刮痕瑕疵。

以相同的方法也能測量碎裂的斷面。機器會比對兩塊碎片的斷面曲度是否相互契合。這項運算替代了人工親手拿著碎片拼合比對的工作。甚至能與遠在其他國家博物館、無法親手觸及的碎片進行比對拼合。

捲曲的卷軸則採用了另一種運算。在恩戈地卷軸的研究中，研究團隊在電腦內將皮革層一層層分離開來。隨後將每一層展開為平面，使文字重現出來 [2]。

以光線解讀墨水的層級

地圖的第二層是墨水。不同材質的墨水對光線的吸收與反射各有不同。多光譜成像是將可見光以及紅外線、紫外線分開拍攝的方法。在某些波長下，背景會變暗而墨水則明亮顯現。

得益於這種差異，被抹去的文字也能重見天日。古人為了節省昂貴的羊皮紙，常將原有的文字擦拭後重新書寫。擦拭過的表面仍會殘留少許墨水成分。利用多種波長拍攝並疊加後，底層的痕跡便會浮現出來。這種經過刮拭重寫的古書卷被稱為重寫本（Palimpsest）。

埃及西奈山的聖凱瑟琳修道院中藏有許多此類文獻。美國與希臘的研究團隊在此利用多光譜成像拍攝了74冊文獻。拍攝頁數達6,800頁。從底層重新顯現出305篇被抹去的文獻。這些文獻是以包含希臘語和敘利亞語在內的十種語言書寫而成 [3]。

此外也有直接測量成分本身的方法。拉曼光譜法是向墨水照射光線並讀取散射出來的光譜圖樣。圖樣因物質而異，宛如指紋一般。人工智慧透過辨識這種圖樣，區分出是用炭黑製成的墨，還是含有鐵成分的墨水。由於不同工坊的墨水配方略有差異，這便成為推測其產地的線索。

在赫庫蘭尼姆卷軸上則使用了更強的光。研究團隊利用能產生極強 X-ray 的同步輻射加速器拍攝卷軸內部。碳化紙草與碳墨的密度相近，單憑常規方式無法分辨。深度學習模型學習了紋理的微細差異，找出了原本有文字的地方 [4]。

串聯文獻與文獻的層次

地圖的第三層是文獻之間的關係。古代書籍是透過抄寫流傳的。每次抄寫都會引入些許微小錯誤。學者們順著這些錯誤追蹤，勾勒出抄本的譜系。

電腦在這項工作中借鑑了生物學所使用的演算法。這與比對基因序列來繪製生物系譜圖的方式十分相似。透過計算文獻之間的差異，排列出哪份抄本源自哪份抄本。這樣繪製出的譜系，讓人們得以推測遺失的原版樣貌。

字體筆跡也是串聯關係的線索。每位抄寫員運筆的力道與筆畫角度各有不同。人工智慧將這種書寫習慣轉化為數值加以儲存。將數值相近的文獻歸類在一起，便能呈現出抄寫員的關聯網絡。

已在運作的典藏庫

整合這些層級並實際運作的檔案庫已有好幾個。Trismegistos 是一項為古代世界文獻編號的計畫。該資料庫將超過100種語言、約96萬件文獻資料納入單一體系中。各國對同一件文物稱呼不一的問題，藉由這些編號得到了解決 [5]。

Pelagios 網絡是一個將分散資料相互連結的社群。他們制定了串聯地點、人物與時間的共通方式。即便古地名隨時代變更，也能將其歸納至同一座標 [6]。還有由慕尼黑營運的電子巴比倫圖書館（Electronic Babylonian Library）。這裡將世界各地的楔形文字碎片匯整至單一資料庫中。學者們在此搜尋碎片，找回遺失的語句 [7]。

共通編號看似枯燥，力量卻很強大。編號一致，不同機構的資料才能串連成一線 [5]。照片在博物館、解讀出的文字在大學、發掘記錄在研究所的情況屢見不鮮。當編號將這三者繫在一起，分散的資訊便凝聚成單一文物的完整故事。

公開資料的方式也改變了。以往要等到論文發表後，才能看見照片與解讀出的文字。如今的研究團隊則會同時公開掃描資料與程式。如此一來，缺乏設備之國家的研究人員也能利用相同資料重新進行計算 [8]。

尚待解決的課題

維蘇威挑戰賽於2026年7月10日整理並公布了尚存的問題。在受壓變形之處，卷軸的層與層無法相互分離。電腦沿著層面追蹤時，有時會錯誤跳躍至相鄰層。墨水由表面滲入多深，至今也尚未明確。在單一卷軸上訓練的模型，應用到其他卷軸時便失去效力。資料量龐大，普通電腦難以處理。研究團隊設定的下一個目標時程是2027年6月 [8]。

處理古代文字的其他領域也面臨同樣的壁壘。統整中國古文字研究的2026年論文，將資料匱乏列為重大難題。並且建議與其取代人工判斷，不如朝輔助人類的方向發展 [9]。

參考資料

- [1] Kutt, K., Sroka, E., Ishchuk, O., & do Valle Miranda, L. (2026). A three-stage neuro-symbolic recommendation pipeline for cultural heritage knowledge graphs. arXiv:2602.19711. <https://arxiv.org/abs/2602.19711> (2026年9月3日檢索)
- [2] Seales, W. B., Parker, C. S., Segal, M., Tov, E., Shor, P., & Porath, Y. (2016). From damage to discovery via virtual unwrapping: Reading the scroll from En-Gedi. Science Advances, 2(9), e1601247. <https://doi.org/10.1126/sciadv.1601247>
- [3] Early Manuscripts Electronic Library. Sinai Palimpsests Project. EMEL. <https://www.emelibrary.org/sinai-palimpsests-project/> (2026年9月3日檢索)
- [4] Vesuvius Challenge. (2026, 6月 25日). An entire Herculaneum scroll has been read for the first time. scrollprize.org. <https://scrollprize.org/firstscroll> (2026年9月3日檢索)
- [5] KU Leuven. Trismegistos portal. Trismegistos+. <https://tmplus.kuleuven.be/trismegistos-portal> (2026年9月3日檢索)
- [6] Pelagios Network. Pelagios Network. <https://pelagios.org/> (2026年9月3日檢索)
- [7] Ludwig-Maximilians-Universität München & Bayerische Akademie der Wissenschaften. electronic Babylonian Library. <https://www.ebl.badw.de/> (2026年9月3日檢索)
- [8] Vesuvius Challenge. (2026, 7月 10日). Open problems: Why reading every Herculaneum scroll is still a challenge. scrollprize.org. https://scrollprize.org/2026_open_problems (2026年9月3日檢索)
- [9] Ma, Y. R. (2026). Towards Computational Chinese Paleography. arXiv:2601.06753. <https://arxiv.org/abs/2601.06753> (2026年9月3日檢索)

12.3. 後記

回顧十一個事件

這本書跟隨了十一個事件的腳步。起點是埋在火山灰下的赫庫蘭尼姆紙草卷。一碰即碎的炭塊，電腦對其進行了虛擬展開。碳化紙草與碳質墨水在X-ray下看起來幾乎毫無二致。機器學會了辨識殘留在表面上的微細紋理差異，從而找出文字 [1]。

恩戈地的《利未記》卷軸也透過虛擬展開重現了首段經文。在此處，含有金屬成分的墨水在X-ray下清晰顯現 [2]。在《以賽亞書》卷軸中，字跡則是關鍵線索。人工智慧量測筆劃的微小習氣，證實抄寫員並非僅有一人。那曾是肉眼看來完全一致的筆跡 [3]。

從第四個事件開始，則是接連不斷的碎片拼合。散落的泥板與羊皮紙碎片根據斷面形狀重新拼接。這套計算方法同樣用於辨別偽造文物。《吉爾伽美什史詩》中缺失的詩行，也透過拼合碎片得以補全。

美索不達米亞的收據藉由機器翻譯得以開口說話。數千年前的糧食帳簿與信件被轉譯成現代語言。龜腹甲上的甲骨文透過3D掃描與專用模型拼合相連 [4]。新羅的木簡也由文字識別模型解讀出來。日本的草書體文獻同樣以這種方式破解 [5]。

我們也看到了尚未開啟的大門。線形文字A、印度河文字與朗戈朗戈文字依然沒有解答。機器衡量符號排列的規則性，藉此推敲那是否為語言。解讀其意涵的工作，則留給了下一個世代。

最後一個事件是古代碑文自動填補。Ithaca是用於填補希臘語碑文殘缺文字的模型。該模型在填補文字的同時，也能推斷碑文的製作地點與年代 [6]。隨後推出的Aeneas則處理拉丁語碑文。Aeneas能同時參照轉錄文字與碑文照片 [7]。

這十一個事件有一個共同點。沒有任何一個問題是由機器獨自解決的。物理學家負責成像拍攝，計算機科學家負責運算。解讀古代語言的學者確認結果並賦予意義。唯有這三個領域圍繞著同一件文物攜手合作，歷史的一行字句才得以重獲新生。

機器目前仍做不到的事

在走到這一步的過程中，機器的弱點也一同暴露出來。第一道障礙是資料匱乏。線形文字A留存的符號全部加起來也不過7,400個左右。單憑這樣的數量，機器難以學習其規則 [8]。當今的語言模型是閱讀數十億字詞成長的。已消亡古語言的資料規模遠遠無法與之相比。

在資料匱乏的情況下勉強給出答案，機器便會捏造出原本不存在的字詞。缺漏的空白愈長，答案的可信度便愈低。這可能產生看似合乎語境、歷史上卻從未存在過的句子。一旦這樣的句子刊載於論文中，後續的研究便會將其視為既成事實。

訓練資料一旦失衡傾斜，答案也會產生偏差。若特定地區的碑文偏多，其他地區的碑文也會被套用該地區的風格來解讀。機器往往無法自行解釋為何給出該答案。因此，近來的研究正轉向要求同時提供推論依據 [9]。

卷軸解讀研究團隊也公開了一份殘留問題清單。在擠壓損毀之處，層與層無法相互分離。在一卷卷軸上訓練的模型，應用於其他卷軸時有時也會失去效力。釐清目前無法做到什麼，才能邁出下一步 [10]。

處理機器所產出之文字的規則同樣不可或缺。明確解讀出的文字與推測出的文字，標示必須有所區隔。同時也必須記錄是由哪個模型根據何種資料做出如此解讀。唯有如此，日後其他研究者才能重新進行相同的計算驗證。

原件始終是原件

數位修復再精細，也無法取代原件。學界將掃描資料視為數位分身，而非原件本身。拍攝過程中融入了人為的抉擇。因為是由人來決定使用何種光源、以何種角度拍攝。這些抉擇一旦改變，最終呈現的影像也會隨之不同。

實物中留存著掃描器未能捕捉的資訊。泥板表面的土粒構造便是一例。龜腹甲上微小氣孔的排列，也只有在實物上才能得見。墨水滲入紙草纖維深處的程度，同樣唯有在實物上才能測量 [10]。

新的測量技術未來仍會層出不窮。若此刻損毀遺棄文物，將來便無法對其運用新技術。因此，應當避免切割文物或強行將其攤平。虛擬圖書館的目的絕非捨棄原件，而是為了在盡量減少接觸原件的前提下實現解讀 [2]。觸摸文物的次數減少，原件的壽命便得以延長。

攜手解讀的未來

人工智慧並非排擠歷史學家的工具。Ithaca的研究以數據印證了這一點。當Ithaca獨自復原缺失文字時，準確率為62%。歷史學家獨自進行時為25%。而當雙方攜手合作時，準確率則攀升至72% [6]。

這項數據同時揭示了兩件事：在這項實驗中，機器獨自表現優於人類；而當與人類攜手時，成果更加卓越。單憑任何一方，都無法走到今天這一步。

機器能迅速提供大量候選字。人類則推敲該候選字在該時代是否合乎常理。學界將這種人機交替驗證的方式稱為「人在迴路中的協同合作」。若要使這個迴路持續運轉，必須持續有具備古文字識讀能力的人才。若無人能讀，機器所提取的文字也只會淪為另一道謎題。

教育傳承與學術研究同等重要。解讀古代語言的方法，向來是由人傳授給人。若這條鏈條中斷，便無人能驗證機器所提取的文字。新世代必須同時學習古代文字與電腦運算。

這個迴路並非僅屬於專家。日本已推出協助解讀草書體的免費應用程式 miWo。只要拍下照片，螢幕上便會顯示現代文字。初次接觸古文獻的人也能讀懂第一行 [5]。在博物館玻璃展櫃前舉起智慧型手機，有時便成了研究的開端。

未來尚有許多工作要做。赫庫蘭尼姆至今仍留有封存未啟的卷軸 [11]。在烏克蘭，如今圖書館與檔案館仍遭受毀壞 [12]。世上的古代文字，消逝的速度依然更快。

讀過這本書的人也有力所能及之事。公開的掃描資料任何人都可以下載查閱 [10]。拼合碎片的競賽也有學生參與其中 [11]。許多地方亦有民眾共同參與轉錄古代文獻。即使一人只負責一行，資料也能聚沙成塔。

守護記錄並非僅是陳年往事。今天我們所建立的資料，有朝一日也會成為古代文獻。檔案格式一旦改變，現今的文件也將無法被讀取。因此，各圖書館持續將資料轉移至新格式。建立虛擬圖書館的工程中，始終離不開維護管理的雙手。

儘管如此，方向已經確立。兩千年前埋於火山灰之下的卷軸得以重見天日並被解讀 [1]。從一卷卷軸中獲得的方法，成為下一卷卷軸研究的起點 [11]。公開資料與程式碼，其他國家的研究人員也能循著相同的道路前行 [10]。文物終有一天會化作塵土，但其承載的字句將永世長存。永不焚毀的圖書館，便是在這樣的一行行文字中逐步建立起來。

參考資料

- [1] Vesuvius Challenge. (2026, 6월 25일). An entire Herculaneum scroll has been read for the first time. scrollprize.org. <https://scrollprize.org/firstscroll> (2026年9月3日檢索)
- [2] Seales, W. B., Parker, C. S., Segal, M., Tov, E., Shor, P., & Porath, Y. (2016). From damage to discovery via virtual unwrapping: Reading the scroll from En-Gedi. Science Advances, 2(9), e1601247. <https://doi.org/10.1126/sciadv.1601247>
- [3] Popović, M., Dhali, M. A., & Schomaker, L. (2021). Artificial intelligence based writer identification generates new evidence for the unknown scribes of the Dead Sea Scrolls exemplified by the Great Isaiah Scroll (1QIsaa). PLOS ONE, 16(4), e0249769. <https://doi.org/10.1371/journal.pone.0249769>
- [4] Ma, Y. R. (2026). Towards Computational Chinese Paleography. arXiv:2601.06753. <https://arxiv.org/abs/2601.06753> (2026年9月3日檢索)
- [5] ROIS-DS Center for Open Data in the Humanities. 미을(miwo): AIକୁずし字認識アプリ. CODH. <https://codh.rois.ac.jp/miwo/> (2026年9月3日檢索)
- [6] Assael, Y., Sommerschield, T., Shillingford, B., Bordbar, M., Pavlopoulos, J., Chatzipanagiotou, M., Androutsopoulos, I., Prag, J., & de Freitas, N. (2022). Restoring and attributing ancient texts using deep neural networks. Nature, 603(7900), 280-283. <https://doi.org/10.1038/s41586-022-04448-z>

- [7] Assael, Y., Sommerschield, T., Cooley, A. E., Shillingford, B., Pavlopoulos, J., Suresh, P., Herms, B., Grayston, J., Maynard, B., Dietrich, N., Wulgaert, R., Prag, J., Mullen, A., & Mohamed, S. (2025). Contextualizing ancient texts with generative neural networks. *Nature*. <https://doi.org/10.1038/s41586-025-09292-5>
- [8] Braović, M., Krstinić, D., Štula, M., & Ivanda, A. (2024). A systematic review of computational approaches to deciphering Bronze Age Aegean and Cypriot scripts. *Computational Linguistics*, 50(2), 725-779. https://doi.org/10.1162/coli_a_00514
- [9] Kutt, K., Sroka, E., Ishchuk, O., & do Valle Miranda, L. (2026). A three-stage neuro-symbolic recommendation pipeline for cultural heritage knowledge graphs. arXiv:2602.19711. <https://arxiv.org/abs/2602.19711> (2026年9月3日檢索)
- [10] Vesuvius Challenge. (2026, 7월 10일). Open problems: Why reading every Herculaneum scroll is still a challenge. scrollprize.org. https://scrollprize.org/2026_open_problems (2026年9月3日檢索)
- [11] Vesuvius Challenge. (2026). Open Prizes. scrollprize.org. <https://scrollprize.org/prizes> (2026年9月3日檢索)
- [12] UNESCO. (2026). Damaged cultural sites in Ukraine verified by UNESCO. UNESCO. <https://www.unesco.org/en/ukraine-war/damaged-cultural-sites> (2026年9月3日檢索)

附錄. 專屬我的AI考古漫遊記：為大眾打造的全球探索入口網站

A.1. Vesuvius Challenge 官方資料典藏庫

能看到什麼

維蘇威挑戰賽是一項國際競賽。競賽內容是以電腦解讀碳化卷軸。卷軸出土的地點是赫庫蘭尼姆。那是位於維蘇威火山旁的古羅馬城市。卷軸是以紙草製成。紙草是壓平草莖製成的古代紙張。官方網址是 scrollprize.org[1]。該網站將競賽所用的資料向所有人公開[2]。無需註冊會員，也完全免費[2][5]。

首頁上方排列著選單。Prizes、Problems、Tutorials & In-depth、Data 排在前面。其後接著 Milestones、Community、Press、Donate。截至2026年9月，已有45件卷軸與紙草殘片完成掃描。其中一件已經從頭到尾被完整解讀[1]。

2026年6月25日，競賽團隊公布了該件卷軸的成果。這件卷軸被稱為 PHerc. 1667[3]。這是在不展開卷軸的情況下，將內容從頭到尾完整解讀出來的首例[3][4]。展開後的紙草長度約為1.4公尺。古代書籍是將文字劃分成垂直欄位來書寫的。這一個欄位稱為「欄」。在這件卷軸中，出現了大約二十二個以希臘字母書寫的欄。內容是探討斯多噶學派倫理學的文章。講述了人的本性、衝動以及道德成長。被認為是西元前2世紀左右的作品[3]。

並非整件卷軸都完整保存了下來。如今留存的是卷軸內側的核心部分。外層在19世紀與1980年代嘗試以手工展開時破碎損毀。表面剝落之處，文字也隨之消失。因此解讀出的文章處處有缺漏[3]。

掃描是在法國格勒諾勃進行的。地點在歐洲同步輻射研究中心 ESRF。那是一座能產生極亮 X-ray 的大型研究設施。其中使用了 BM18 光束線[4]。光束線是將 X-ray 引出並照射到物體上的實驗通道。

在家跟著做的步驟

首先進入 scrollprize.org，點擊 Get Started 按鈕。這個畫面為初次造訪者指引了三個步驟[5]。

第一步是尋找墨跡。在畫面上塗繪標記肉眼可見的墨跡。接著模型就會學習這些標記，並預測其餘部分[5]。

第二步是虛擬展開。也就是用電腦將卷軸展開。在 CT 斷面影像中順著單層紙草追蹤。CT 斷面影像就像切成薄片拍攝內部構造一樣的照片。將該層展平後確認成果[5]。

第三步是掌握整體流程。以小規模連續嘗試掃描、展開與解讀[5]。

必須注意體驗步驟與實際修復順序不同。體驗是從較簡單的尋找墨跡開始。實際修復則是掃描在先。接著以虛擬方式展開卷軸。解讀墨跡則是放在最後一步[5]。

若想直接下載資料，可前往 scrollprize.org/data。資料存放在公開儲存庫 [s3://vesuvius-challenge-open-data/](https://s3.amazonaws.com/scrollprize-open-data) 中。也可以透過網頁瀏覽器開啟瀏覽。說明使用方法的範例檔案則存放在 GitHub 上[2]。

競賽也為挑戰者設立了獎金。截至2026年9月，剩餘獎金為214萬美元。2027年大獎規模達100萬美元。第一名可獲得80萬美元。截止日期為2027年6月25日。尚未找出文字的卷軸也設有獎金。只要在4平方公分的面積內解讀出10個文字即可。此時每件卷軸提供5萬美元獎金。此外還有每月頒發的進度獎。當月評選出的最佳提交成果可獲得2萬美元[6]。

注意事項

資料以 CC BY-NC 4.0 條款公開。不得用於營利目的。使用時必須註明出處。針對早期資料，另有指定的引用論文。卷軸1號至4號、殘片1號至6號便是如此[2]。此時需引用 Parsons 等人於2023年發表的論文[2][7]。

也必須事先了解檔案容量非常龐大。3D 掃描資料整體達到數 TB 之多[2]。若要執行深度學習，需要效能優良的顯示卡。如果家用電腦效能較弱，可以先從瀏覽資料開始。

挑戰獎金時附帶一定條件。必須加入競賽的 Discord 頻道。程式碼必須以如 MIT 等自由開放的條款公開。在公布獲獎前，不得向外界透露成果[6]。

參考資料

- [1] Vesuvius Challenge. (2026). Vesuvius Challenge: Reading the Herculaneum Scrolls with machine learning. scrollprize.org. <https://scrollprize.org/> (2026年9月3日檢索)
- [2] Vesuvius Challenge. (2026). Data. scrollprize.org. <https://scrollprize.org/data> (2026年9月3日檢索)
- [3] Vesuvius Challenge. (2026, 6월 25일). An entire Herculaneum scroll has been read for the first time. scrollprize.org. <https://scrollprize.org/firstscroll> (2026年9月3日檢索)
- [4] Angelotti, G., Parsons, S., Nicolardi, F., et al. (2026). Complete virtual unwrapping and reading of a rolled Herculaneum papyrus. arXiv:2606.29085. <https://arxiv.org/abs/2606.29085> (2026年9月3日檢索)
- [5] Vesuvius Challenge. (2026). Get Started. scrollprize.org. https://scrollprize.org/get_started (2026年9月3日檢索)
- [6] Vesuvius Challenge. (2026). Open Prizes. scrollprize.org. <https://scrollprize.org/prizes> (2026年9月3日檢索)
- [7] Parsons, S., Parker, C. S., Chapman, C., Hayashida, M., & Seales, W. B. (2023). EduceLab-Scrolls: Verifiable recovery of text from Herculaneum papyri using X-ray CT. arXiv:2304.02084. <https://doi.org/10.48550/arXiv.2304.02084>

A.2. eBL 電子巴比倫文學入口網站

入口網站的架構

eBL 是 electronic Babylonian Library 的縮寫。中文意為電子巴比倫圖書館。古代美索不達米亞人使用楔形文字書寫。楔形文字是在濕潤的泥板上壓印出如釘子般形狀的文字。該入口網站將那些泥板殘片匯集於一處。收錄對象是記錄巴比倫文學的殘片。這裡並非收錄世上所有楔形文字資料的場所[1]。

該入口網站設於德國慕尼黑大學（路德維希-馬克西米利安大學）。由恩里克·希門尼斯教授主持。計畫於2018年4月啟動。由亞歷山大·馮·洪堡基金會的索菲亞·科瓦列夫斯卡婭獎資助款營運[2]。

網址是 www.ebl.lmu.de。無需登入即可進入，且不收費。畫面左側有八個項目。eBL Project、Library、Corpus、Signs 排在前面。其後是 Akkadian Dictionary 和 Bibliography。最後則是 News 與 Archaeology[1]。

Library 是收集各個殘片的地方。過去曾被稱為 Fragmentarium。eBL 的研究人員在此處找到了大約1,200件殘片拼接的案例[1]。Corpus 則以作品為單位展示正文。《吉爾伽美什史詩》整理為約3,150行。創世神話《埃努瑪·埃利什》則為1,096行[3]。2024年發表的資料論文呈現了其規模。該入口網站收錄了約25,000件泥板。轉寫文本超過35萬行[4]。轉寫文本是將楔形文字轉換為拉丁字母書寫的文字。

在畫面上能做些什麼

點擊 Library 就會出現搜尋框。可依名稱、年份、標題進行搜尋。亦提供單字搜尋與進階搜尋。若不知道要找什麼，可按下 I'm feeling lucky 按鈕。這樣就會隨機開啟任一塊殘片[5]。

開啟一個殘片後，畫面會分成數個區塊。上方記載著收藏機構、尺寸及出土地點。字體與文獻類型也會一併顯示。例如 K.5743 是大英博物館的庫雲吉克藏品。庫雲吉克是古城尼尼微所在的山丘。這是指出土於該山丘的文物統稱。尺寸為3.81公分乘2.54公分。字體為新亞述式[6]。即新亞述帝國時代所使用的文字形體。

中間部分顯示帶有行號的轉寫文本。下方則附有泥板照片。同時也會註明與哪一塊殘片拼合在一起[6]。

網站也說明了引用文章時所使用的格式。寫上殘片編號與 eBL edition。其後附上網址與造訪日期[1]。

最新消息與注意事項

在 News 項目中可以閱讀通訊簡訊。2026年5月28日發布的第22期為最新一期。在此期中，入口網站推出了全新畫面配置。收藏機構也增加了數處。位於仁川的國立世界文字博物館亦為其中之一。新增上傳了大英博物館3,100張以上的照片。耶魯皮博迪博物館的照片也新增了1,000張以上。此外還新增了可直接在辭典中登錄新單字的功能[7]。

在使用前有幾點需要注意。Library 的轉寫文本並非定稿，而是暫定版本。研究團隊至今仍持續在進行修改。Corpus 中整理的作品正文亦是如此。這是將迄今為止蒐集的殘片拼合而成的階段性成果。若拼接上新的殘片，行數與內容都會發生變化。因此在用於學校作業或文章寫作時，應一併記下造訪日期[1]。

轉寫文本以 CC BY-NC-SA 4.0 條款公開。照片的授權條款則有所不同。若要再次使用照片，必須另行取得授權[1]。是由收藏泥板的博物館負責給予許可[1][4]。必須先閱讀畫面下方照片說明所附的指引[1]。

參考資料

- [1] electronic Babylonian Library. (2026). About: Library. ebl.lmu.de. <https://www.ebl.lmu.de/about/library> (2026年9月3日檢索)
- [2] Jiménez, E. (2023, 5월 25일). AI is helping us read ancient Mesopotamian literature. The Conversation. <https://theconversation.com/ai-is-helping-us-read-ancient-mesopotamian-literature-205091> (2026年9月3日檢索)
- [3] electronic Babylonian Library. (2026). Corpus. ebl.lmu.de. <https://www.ebl.lmu.de/corpus> (2026年9月3日檢索)
- [4] Cobanoglu, Y., Laasonen, J., Simonjetz, F., Khait, I., Cohen, S., Földi, Z., Härtinen, A., Heinrich, A., Mitto, T., Rozzi, G., Sáenz, L., & Jiménez, E. (2024). Transliterated cuneiform tablets of the electronic Babylonian Library platform. Journal of Open Humanities Data, 10(1), 19. <https://doi.org/10.5334/johd.148>
- [5] electronic Babylonian Library. (2026). Library. ebl.lmu.de. <https://www.ebl.lmu.de/library> (2026年9月3日檢索)
- [6] electronic Babylonian Library. (2026). K.5743, eBL edition. ebl.lmu.de. <https://www.ebl.lmu.de/library/K.5743> (2026年9月3日檢索)
- [7] electronic Babylonian Library. (2026, 5월 28일). eBL Newsletter 22. ebl.lmu.de. <https://www.ebl.lmu.de/news> (2026年9月3日檢索)

A.3. Miwo 使用指南

這是款什麼樣的 App

日本的古籍與古文書中大量使用了草書體。這種字體被稱為草書字（くずし字，Kuzushiji）。在江戶時代的印刷本與古文書中經常可見[3]。如今大多數的日本人都無法閱讀這種字體。

Miwo 是一款拍攝這種草書字即可將其轉換為現代文字的智慧型手機 App。由日本的人文學開放數據共同利用中心 CODH 開發[1]。費用全免[2]。App 名稱取自《源氏物語》第十四帖〈澁標〉（Miotsukushi）。「澁標」是指引航道的標記。寓意著要在古文書的汪洋中成為引路指南[3]。

該 App 於2021年8月30日首次問世。2022年10月26日升級至 1.1 版。此時加入了名為「瑠璃（Ruri）」的新辨識模型。即使文字背景帶有花紋或顏色，瑠璃也能精準找出[1]。模型訓練採用了「日本古典籍草書字資料集（日本古典籍くずし字データセット）」。這是專為機器學習打造的古文字圖像集合。包含超過一百萬字的資料。該 App 於2022年榮獲 Good Design Award[3]。截至2026年9月，該 App 辨識的照片數量已超過350萬張[1]。

安裝與使用步驟

首先開啟應用程式商店。iPhone 和 iPad 使用 App Store。Android 裝置使用 Google Play。在搜尋欄中輸入「みを」或「miwo」[1]。以 App Store 為準，App 大小約為 52 MB。iPhone 和 iPad 只要是 iOS 11.0 以上版本即可[2]。

開啟 App 後會出現拍照畫面。將古文書置於畫面中並按下拍攝按鈕。也可以匯入已拍好的照片。接著按下辨識按鈕，AI 就會解讀文字[1]。

結果畫面具備多種功能。畫面上設有一條調節拉桿。左右滑動此拉桿即可讓兩個畫面重疊顯示。能夠將原始照片與解讀出的文字並列比對。亦有點選日語假名文字的功能。它會顯示該字最初源自哪一個漢字。解讀出的文字可透過複製按鈕轉移至其他 App[1]。也可以將結果儲存起來，日後再次調出查看[1][3]。

必須注意需要保持網路連線。照片會傳送至 CODH 伺服器進行辨識。傳送的照片不會留存在伺服器中[1]。

能順利解讀與難以解讀的資料

使用 App 前若先了解其局限，就不易感到失望。AI 主要是以江戶時代的印刷本為核心進行訓練學習。該時期的書籍多為木版印刷本。是在木板上刻上文字並印於紙上的書籍。因此對於該時期的印本能順利解讀。而其他時代的資料或手寫文書，準確度則會下降[3]。

紙張狀況也會改變結果。有污漬或蟲蛀之處難以解讀。紙張的紋理、光線、陰影也會造成干擾[3]。最好在明亮且均勻的光線下拍攝。

也有無法解讀的對象。目前不支援刻在石碑或招牌上的草書體文字[3]。

開發團隊表明辨識結果可能存在錯誤。並強調 AI 草書辨識無法達到百分之百完美。因此重要資料仍需由人工再次確認[3]。

將古文字轉寫為現代文字的工作稱為翻刻。另外有一個平台是由大眾分工進行這項工作。那就是名為「大家來翻刻」的市民參與型翻刻網站。中文意為「大家的翻刻」。該網站由與 Miwo 開發團隊不同的機構負責營運。其性質相當於由人工來修飾潤色 App 所解讀出的文字。網址為 honkoku.org。迄今已有超過10,800人參與。轉寫文字已超過4,800萬字。2025年8月進行了第二次改版。新版本目前正在試營運中。日本國立歷史民俗博物館參與其中。東京大學地震研究所與京都大學研究團隊亦共同維護營運[4]。

參考資料

[1] ROIS-DS 人文学オープンデータ共同利用センター. (2026). AIくずし字認識アプリ「みを（miwo）」. CODH.
<https://codh.rois.ac.jp/miwo/> (2026年9月3日檢索)

[2] Clanuwat, T. (2022, 10월 27일). miwo (Version 1.1) [Mobile app]. App Store.
<https://apps.apple.com/jp/app/miwo/id1581794085> (2026年9月3日檢索)

[3] ROIS-DS 人文学オープンデータ共同利用センター. (2026). みを（miwo）とは？. CODH.
<https://codh.rois.ac.jp/miwo/about/> (2026年9月3日檢索)

[4] みんなで翻刻. (2026). みんなで翻刻 Minna de Honkoku. honkoku.org. [https://honkoku.org/](https://honkoku.org) (2026年9月3日檢索)

支持這本 AI 書齋著作

如果這本書對您有幫助，歡迎透過 PayPal 支持後續的翻譯、公開版本以及開放的人工智慧知識計畫。

PayPal

<https://paypal.me/kimkjkj>



請掃描 QR 碼，或直接開啟上方連結。

金京鎮律師