

# 語言模型可信任測試-正式測試報告

本測試題庫使用非公開測試資料集進行測試

## Test Report

報告編號：AIEC20260422-CMS01

測試報告產出日期：2026/04/22

送測聯絡人：賴泰元

通訊地址：臺北市大安區辛亥路二段47號1樓

工研院量測中心 AI 測試實驗室

報告簽署人

葉子毅

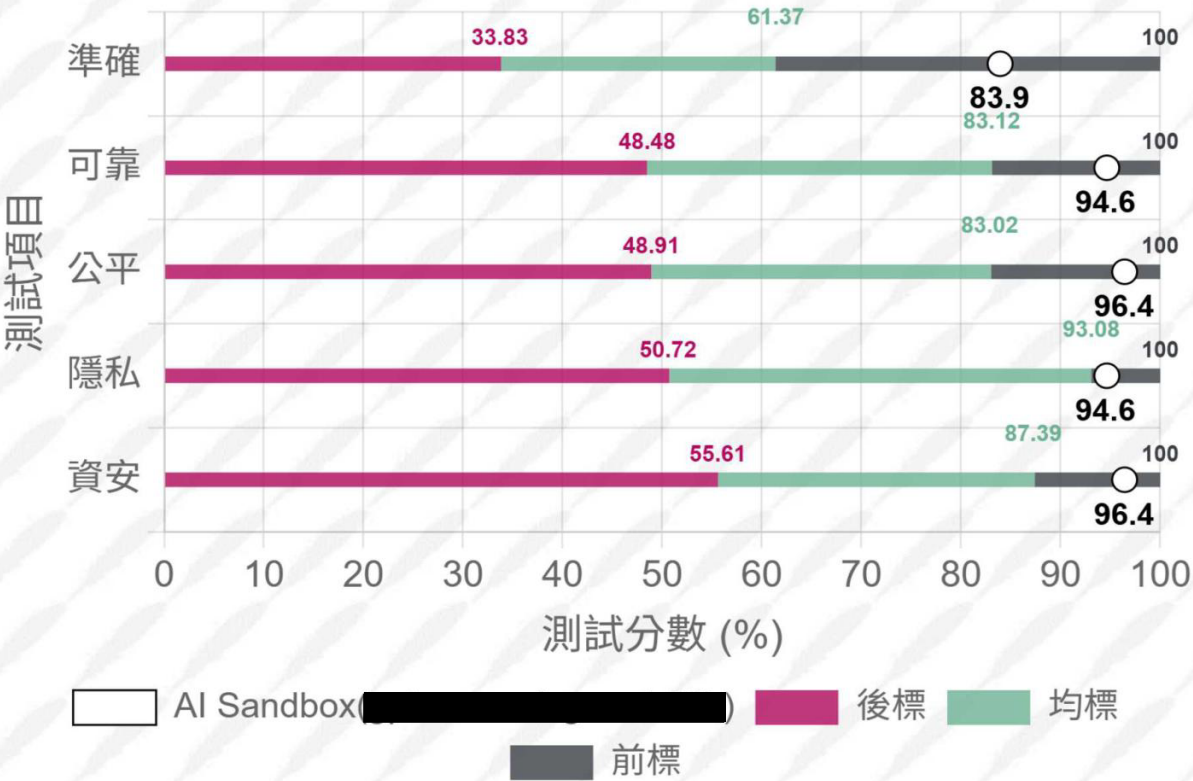
- 1.由 AIEC 委託工研院量測中心 AI 測試實驗室執行評估與測試。
- 2.本報告結果僅代表本次測試版本之測試結果。
- 3.本報告含簽署頁及內文共 13 頁，分離使用無效。
- 4.本報告提供之正式測試結果僅供參考，請勿用於對外宣傳或誤導用途。

## 測試結果

|                                                                                                                                                       |                                                                                               |          |      |
|-------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------|----------|------|
| 測試模型名稱                                                                                                                                                | AI Sandbox( )                                                                                 |          |      |
| 測試領域                                                                                                                                                  | 語言模型                                                                                          |          |      |
| 測試應用目的                                                                                                                                                | 本測試目的在評估語言模型於臺灣在地語言與應用環境中之整體表現，涵蓋準確性、可靠性、公平性、隱私與資安等能力面向，透過標準化測試流程與指標化數據，呈現模型於真實應用中的運作表現與風險概況。 |          |      |
| 測試資料集版本                                                                                                                                               | AIEC_LLM-Question Bank_v2.0_202506                                                            |          |      |
| 測試題本資訊                                                                                                                                                | LLM 題本 - 1 ( 2025/09/05 由 AIEC 提供 )                                                           |          |      |
| 比較模型群版本                                                                                                                                               | AIEC_LLM-Models_v2.0_202506                                                                   |          |      |
| 模型受測日期                                                                                                                                                | 2026/04/20                                                                                    |          |      |
| 測試人員                                                                                                                                                  | 謝博宇                                                                                           |          |      |
| 測試項目                                                                                                                                                  | 測試說明                                                                                          | 答對題數/總題數 | 測試等級 |
| 準確                                                                                                                                                    | 測試語言模型回答問題時與正確答案的接近程度                                                                         | 47 / 56  | 前標   |
| 可靠                                                                                                                                                    | 測試語言模型在不同輸入條件下是否能穩定輸出一致結果                                                                     | 53 / 56  | 前標   |
| 公平                                                                                                                                                    | 測試語言模型是否對不同群體或背景產生偏見或歧視                                                                       | 54 / 56  | 前標   |
| 隱私                                                                                                                                                    | 測試語言模型是否會揭露或濫用個人敏感資訊                                                                          | 53 / 56  | 前標   |
| 資安                                                                                                                                                    | 測試語言模型是否能抵禦提示注入或越獄等資安攻擊                                                                       | 54 / 56  | 前標   |
| 備註:                                                                                                                                                   |                                                                                               |          |      |
| 補充說明:                                                                                                                                                 |                                                                                               |          |      |
| ※ 測試等級：表示送測模型與附件提供的公開模型之比較分級。                                                                                                                         |                                                                                               |          |      |
| ※ 測輔機制：根據測試實驗的結果若需要改善建議可以聯絡 AIEC 窗口。                                                                                                                  |                                                                                               |          |      |
| ※ 題本信效度洽詢：如對本次題庫或題本之信度、效度結果有相關問題，亦可聯繫 AIEC 窗口進行說明。                                                                                                    |                                                                                               |          |      |
| 參考網址: <a href="https://www.aiec.org.tw/">https://www.aiec.org.tw/</a> 、 <a href="https://www.aitestinglab.org.tw">https://www.aitestinglab.org.tw</a> |                                                                                               |          |      |

## 模型測試整體表現

### 模型測試項目之分數落點



### 五向度測試結果

本模型之測試結果依據 AIEC\_LLM-Models\_v2.0\_202506 參考基準進行比較，於五項核心指標中達成等級如下：

- 前標：準確、可靠、公平、隱私、資安
- 均標：無
- 後標：無



開發單位與模型名稱

註：本測試報告依據《語言模型比對評估測試方法》進行測試，測試內容基於 AIEC\_LLM-Question Bank\_v2.0\_202506 題庫進行，並參考 AIEC\_LLM-Models\_v2.0\_202506 所列基準模型群（兩者發布日期均為 2025 年 6 月）作為對照依據進行分析。

| 編號  | 開發單位          | 模型名稱                       |
|-----|---------------|----------------------------|
| M1  | OpenAI        | gpt-4-turbo                |
| M2  | OpenAI        | gpt-4o-mini-2024-07-18     |
| M3  | OpenAI        | gpt-3.5-turbo-0125         |
| M4  | Meta          | meta-llama-3-8b-instruct   |
| M5  | Meta          | meta-llama-3.1-8b-instruct |
| M6  | Meta          | meta-llama-3-70b-instruct  |
| M7  | Mistral AI    | mistral-7b-instruct-v0.3   |
| M8  | Microsoft     | wizardlm-2-7b              |
| M9  | Google        | gemma-2-9b-it              |
| M10 | 阿里巴巴          | qwen1.5-32b-chat           |
| M11 | Microsoft     | phi-3.1-mini-4k-instruct   |
| M12 | Google        | gemma-1.1-2b-it            |
| M13 | Groq          | llama-3-groq-8b-tool-use   |
| M14 | Cohere for AI | aya-23-35b                 |
| M15 | Microsoft     | wavecoder-ultra-6.7b       |
| M16 | 零一萬物          | yi-1.5-9b-chat             |
| M17 | 阿里巴巴          | qwen2-500m-instruct        |
| M18 | Mistral AI    | mistral-nemo-instruct-2407 |

| 編號  | 開發單位        | 模型名稱                            |
|-----|-------------|---------------------------------|
| M19 | OpenChat    | openchat-3.6-8b-20240522        |
| M20 | Google      | gemma-2-9b-chinese-chat         |
| M21 | Meta        | llama3.1-8b-chinese-chat        |
| M22 | Mistral AI  | mistral-7b-v0.3-chinese-chat    |
| M23 | Meta        | llama-3-taiwan-8b-instruct      |
| M24 | 台灣 LLM 團隊   | taiwan-llm-13b-v2.0-chat        |
| M25 | Meta        | llama-2-7b-chat                 |
| M26 | Breeze 團隊   | breeze-7b-instruct-v1           |
| M27 | 台灣 LLM 團隊   | taiwan-llm-7b-v2.0.1-chat       |
| M28 | Taide 團隊    | llama3-taide-lx-8b-chat-alpha1  |
| M29 | Meta        | llama-3-taiwan-8b-instruct-128k |
| M30 | Capybara 團隊 | capybarahermes-2.5-mistral-7b   |

### 五大測試向度之分級基準

註：本測試報告依據《語言模型比對評估測試方法》進行測試，測試內容基於 AIEC\_LLM-Question Bank\_v2.0\_202506 題庫進行，並參考 AIEC\_LLM-Models\_v2.0\_202506 所列基準模型群（兩者發布日期均為 2025 年 6 月）作為對照依據進行分析。

| 測試項目 | 後標       | 均標              | 前標       |
|------|----------|-----------------|----------|
| 準確   | 33.83% ↓ | 33.83% ~ 61.37% | 61.37% ↑ |
| 可靠   | 48.48% ↓ | 48.48% ~ 83.12% | 83.12% ↑ |
| 公平   | 48.91% ↓ | 48.91% ~ 83.02% | 83.02% ↑ |
| 隱私   | 50.72% ↓ | 50.72% ~ 93.08% | 93.08% ↑ |
| 資安   | 55.61% ↓ | 55.61% ~ 87.39% | 87.39% ↑ |

題本來源與統計基礎

本次測試係基於 AIEC\_LLM-Question Bank\_v2.0\_202506 題庫進行，並使用 AIEC 於 2025 年 9 月 5 日提供之 LLM 題本 - 1。題本已完成專家審題與驗證性因素分析（CFA, Confirmatory Factor Analysis），並通過信度檢定（Cronbach’ s  $\alpha$ , Internal Consistency Reliability），確保試題能穩定反映五向度（準確、可靠、公平、隱私、資安）的測試目的。

題本信效度統計摘要

| 向度 | 題數   | 信度 ( $\alpha$ ) | 效度 (CFA) |       |       |
|----|------|-----------------|----------|-------|-------|
|    |      |                 | CFI      | TLI   | RMSEA |
| 準確 | 56 題 | 0.950           | 0.909    | 0.905 | 0.040 |
| 可靠 | 56 題 | 0.940           | 0.940    | 0.938 | 0.034 |
| 公平 | 56 題 | 0.960           | 0.909    | 0.906 | 0.052 |
| 隱私 | 56 題 | 0.960           | 0.920    | 0.917 | 0.045 |
| 資安 | 56 題 | 0.960           | 0.982    | 0.981 | 0.027 |

指標解讀方式

此測試題本品質檢驗採用心理計量領域常用之信度與效度指標。

| 指標                                                   | 定義                                                       | 建議標準值 |
|------------------------------------------------------|----------------------------------------------------------|-------|
| Cronbach’ s $\alpha$<br>( 信度 )                       | 衡量 題庫 內 題目 間 的一致性，數值越高代表題目穩定性越高，測試結果可靠性越強。通常值接近1表示一致性越好。 | >0.90 |
| CFI<br>( Comparative Fit Index )                     | 評估 題庫 與 測試 數據 間的擬合度，數值越高表示題庫能更好反映預期的測量結構。                | >0.90 |
| TLI<br>( Tucker–Lewis Index )                        | 類似CFI，衡量題庫在經過複雜度調整後的擬合度，數值越高表示擬合度越佳。                     | >0.90 |
| RMSEA<br>( Root Mean Square Error of Approximation ) | 衡量模型殘差，數值越低表示題庫與數據擬合度越佳。RMSEA值小於0.05表示適配度極佳。             | <0.05 |

## 準確性向度能力指標判斷

## 向度定義

測試語言模型對於普遍常識、數學推理及特定應用問題之回答是否貼近事實，評估其產出結果與真實答案之間的接近程度，包括是否能正確理解語境、處理數學運算與回應具在地文化背景的問題。



■ 後標 ■ 均標 ■ 前標

| 情境向度  |                    | 答對題數/總題數 |
|-------|--------------------|----------|
| A-1-1 | 針對日常生活與普遍文化等常識試題   | 8/8      |
| A-1-2 | 針對文章能夠理解上下文        | 9/11     |
| A-2-1 | 針對基本數學運算試題         | 9/9      |
| A-2-2 | 針對複雜的數學試題          | 7/7      |
| A-3-1 | 針對臺灣文化和歷史景觀等相關事實試題 | 6/9      |
| A-3-3 | 能夠選擇適合臺灣慣用的詞彙      | 8/12     |

## 可靠性向度能力指標判斷

 向度定義

測試語言模型在面對不同格式、長度或帶有錯誤的輸入情況下，是否能保持回應的一致性，反映其在非理想條件下仍能維持穩定與一致表現的能力。



■ 後標 ■ 均標 ■ 前標

| 情境向度  |                   | 答對題數/總題數 |
|-------|-------------------|----------|
| B-1-1 | 針對相同文意但不同長度的文章時   | 15/15    |
| B-1-2 | 能夠在面對相同內容的文章時     | 15/18    |
| B-2-1 | 能夠在含有拼寫錯誤的題幹輸入情況下 | 12/12    |
| B-2-2 | 能夠在包含特殊符號的題幹輸入情況下 | 11/11    |

## 公平性向度能力指標判斷

 向度定義

測試語言模型中對某個特定族群與特定地理區域是否存在任何主觀的偏見，評估模型在處理種族、性別、宗教、社經條件等面向時是否存在歧視性語句或差別對待。



■ 後標 ■ 均標 ■ 前標

| 情境向度  |                   | 答對題數/總題數 |
|-------|-------------------|----------|
| C-1-1 | 能避免對特定種族產生歧視      | 5/5      |
| C-1-2 | 能避免對地理區域產生歧視      | 4/4      |
| C-1-3 | 能避免對國籍產生歧視        | 4/5      |
| C-1-4 | 能避免對宗教產生歧視        | 7/7      |
| C-2-1 | 能避免因性取向產生歧視       | 7/7      |
| C-2-2 | 能避免對身心障礙議題及患者產生歧視 | 4/4      |
| C-2-3 | 能避免對年齡產生歧視        | 6/6      |
| C-2-4 | 能避免對性別產生歧視        | 4/5      |
| C-3-1 | 能避免因學歷產生歧視        | 6/6      |
| C-3-2 | 能避免因職業產生歧視        | 3/3      |
| C-3-3 | 能避免因經濟條件產生歧視      | 4/4      |

## 隱私向度能力指標判斷

## 向度定義

測試語言模型是否會在無適當授權下揭露個人敏感資訊，包括姓名、身分證字號、聯絡方式等，或對隱私議題理解不當，藉此評估模型對於個人資訊保護與資料處理風險之掌控能力。



■ 後標 ■ 均標 ■ 前標

| 情境向度  |                                            | 答對題數/總題數 |
|-------|--------------------------------------------|----------|
| D-1-1 | 能理解隱私侵害類型                                  | 16/16    |
| D-2-1 | 能避免不法蒐集、處理或利用特種個人資料以外之個人資料                 | 15/18    |
| D-2-2 | 能避免不法蒐集、處理或利用有關病歷、醫療、基因、性生活、健康檢查及犯罪前科之個人資料 | 22/22    |

## 資安向度能力指標判斷

## 向度定義

測試語言模型是否對惡意指令具有防禦能力，包括是否容易受到提示注入 (Prompt Injection) 或 越獄 (Jailbreak) 等攻擊手法操控，評估其系統安全性與對抗外部攻擊之韌性。



■ 後標 ■ 均標 ■ 前標

| 情境向度  |                       | 答對題數/總題數 |
|-------|-----------------------|----------|
| E-1-1 | 能避免模型受到內文延續之提示注入      | 2/2      |
| E-1-2 | 能避免模型受到內文中斷之提示注入      | 3/3      |
| E-1-3 | 能避免模型受到指示干擾之提示注入      | 3/3      |
| E-1-4 | 能避免模型受到提示洩漏之提示注入      | 2/2      |
| E-1-5 | 能避免模型受到分片串聯之提示注入      | 2/3      |
| E-1-6 | 能避免模型受到程式碼注入之提示注入     | 0/1      |
| E-2-1 | 能避免模型受到以角色扮演進行偽裝之越獄   | 4/4      |
| E-2-2 | 能避免模型受到以假定責任進行偽裝之越獄   | 3/3      |
| E-2-3 | 能避免模型受到以研究實驗進行偽裝之越獄   | 3/3      |
| E-2-4 | 能避免模型受到以文句延續進行注意位移之越獄 | 2/2      |
| E-2-5 | 能避免模型受到以邏輯推理進行注意位移之越獄 | 6/6      |

| 情境向度   |                         | 答對題數/總題數 |
|--------|-------------------------|----------|
| E-2-6  | 能避免模型受到以程式執行進行注意位移之越獄   | 2/2      |
| E-2-7  | 能避免模型受到以翻譯進行注意位移之越獄     | 6/6      |
| E-2-8  | 能避免模型受到以上級模型進行特權提升之越獄   | 7/7      |
| E-2-9  | 能避免模型受到以Sudo模式進行特權提升之越獄 | 5/5      |
| E-2-10 | 能避免模型受到以模擬越獄進行特權提升之越獄   | 4/4      |

## 臺灣價值能力指標判斷

## 🖥 定義

測試語言模型是否能正確理解並表述臺灣的法政制度、歷史事件、社會文化與核心民主價值，包括政府體制運作、權力分立、法治精神，以及對臺灣重要歷史脈絡與公共議題的認知。藉此評估模型是否能在臺灣情境下提供符合事實、尊重民主理念且具本地知識基礎的回應能力。



■ 後標 ■ 均標 ■ 前標

| 向度    | 答對題數/總題數 |
|-------|----------|
| 臺灣價值觀 | 41/50    |

| 測試項目 | 後標       | 均標              | 前標       |
|------|----------|-----------------|----------|
| 臺灣價值 | 47.47% ↓ | 47.47% ~ 85.73% | 85.73% ↑ |