

中製「生成式 AI 語言模型」檢測

語言模型 檢測項目		DeepSeek	豆包	文心一言	通義千問	騰訊元寶
		不合格項目（檢出不合格項目以✕標記）				
一、應用程式檢測（5類15項）	蒐集個資					
	蒐集位置	✕	✕	✕	✕	✕
	蒐集通訊錄				✕	✕
	蒐集剪貼簿	✕	✕		✕	✕
	蒐集截圖	✕	✕	✕	✕	✕
	讀取裝置上儲存空間	✕	✕		✕	✕
	逾越使用權限					
	過度填寫個資	✕	✕	✕		
	過度要求權限		✕		✕	✕
	強迫同意不合理隱私條款	✕	✕	✕	✕	✕
	未充分保障個資權利			✕	✕	✕
	數據回傳分享					
	未啟動時上傳非必要個資					
	逕向第3方SDK共享個資	✕	✕	✕		
	封包有無導向惡意連線位址				✕	✕
	擷取系統資訊					
	蒐集程式清單			✕	✕	
	蒐集設備參數	✕	✕	✕	✕	✕
	掌握生物特徵					
	蒐集臉部資訊		✕	✕		
二、生成內容檢測（10項）	安全性				✕	
	可解釋性	✕	✕	✕	✕	✕
	韌性	✕	✕	✕		✕
	公平性	✕	✕	✕	✕	
	準確性	✕	✕	✕	✕	
	透明性	✕	✕	✕		✕
	當責性					
	可靠性	✕	✕	✕	✕	
	隱私					
	資安	✕	✕	✕	✕	✕
總計		15	17	16	17	14

資料來源：國安局綜整