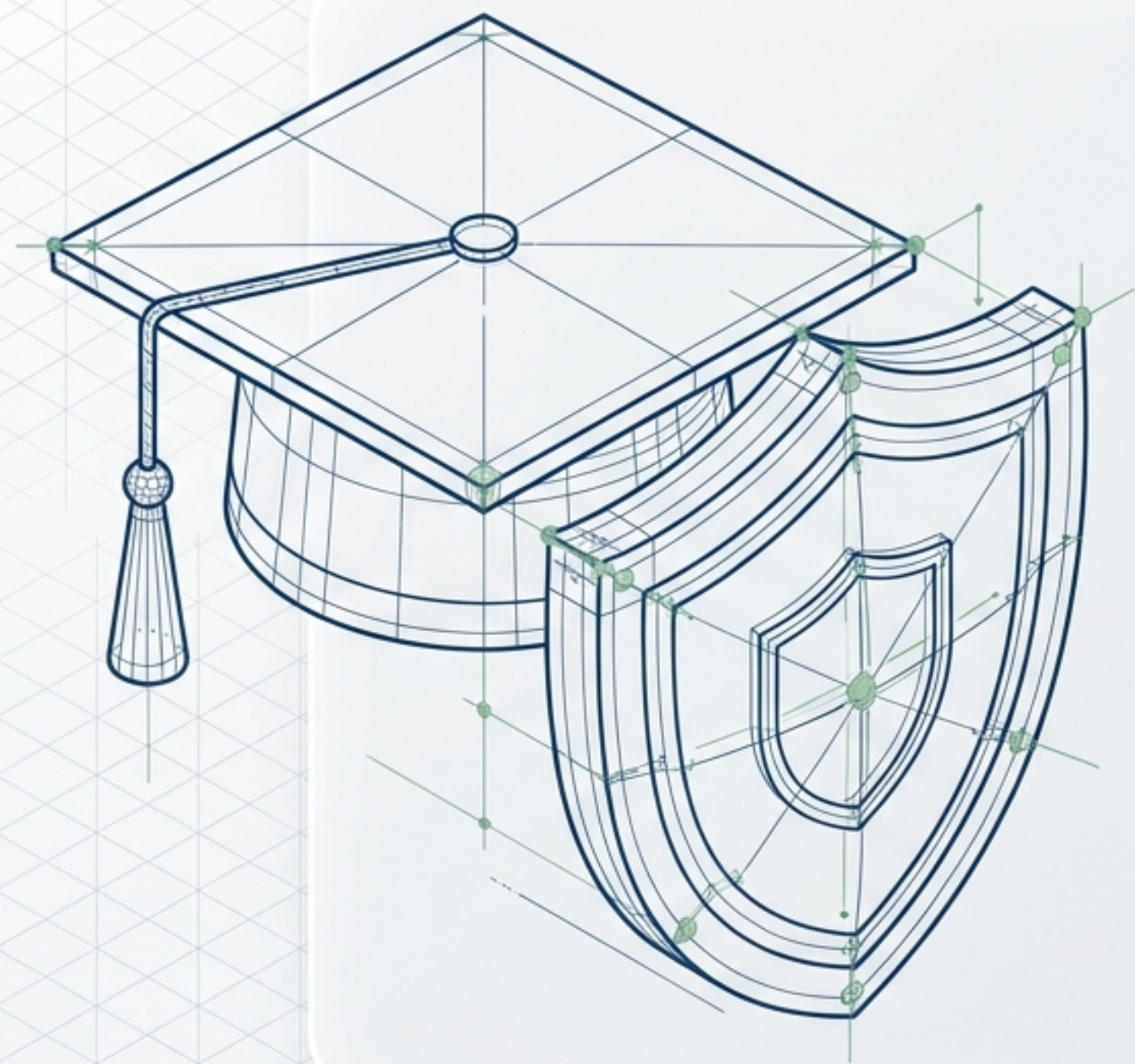




打造安全的 AI 智慧校園

學術與行政工作中的
AI 隱私防護指南

為什麼在校園中需要防護 AI 對話隱私？



- 預設情況下，許多公共 AI 模型會將使用者的對話內容用於未來的模型訓練。

- 在學術與行政環境中，若不慎輸入機密數據，將面臨嚴重的資料外洩風險。

第一道防線：關閉模型訓練（一般使用者必備）

目標：讓您的新對話不會被用來訓練 OpenAI 的模型。

四步設定：

1. 點選左下角帳號
2. 選擇 Settings（設定）
3. 進入 Data Controls（資料控制）
4. 關閉 Improve the model for everyone（使用我的內容改善模型）

註：關閉後仍可正常使用 ChatGPT，這是免費版使用者最重要的保護步驟。



帳號層級與資料保護差異

版本 (Version)	是否用於模型訓練 (Used for Training?)
免費版 (Free)	 可關閉
Plus	 可關閉
Team	 不用於訓練
Enterprise	 不用於訓練
Edu	 不用於訓練

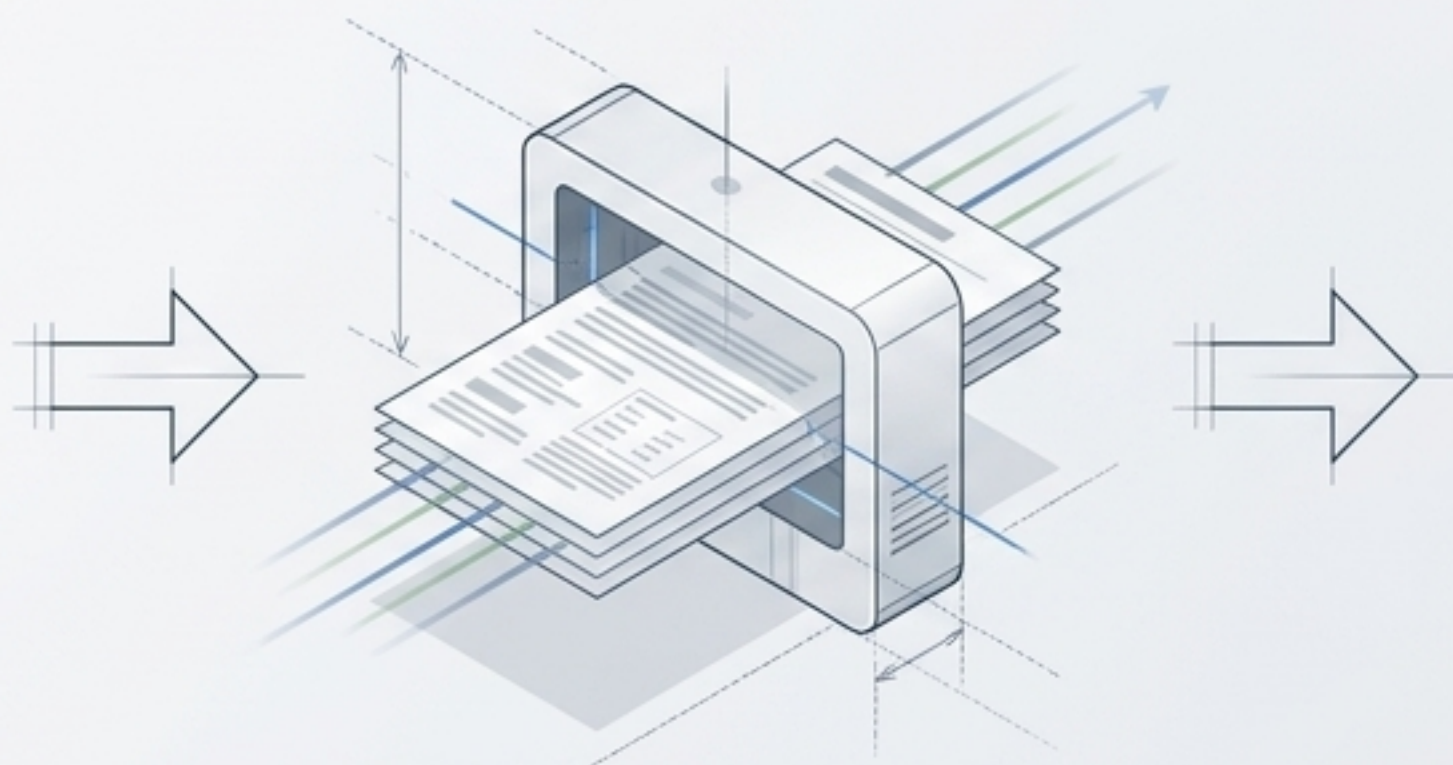
企業與教育版 (Edu) 提供最嚴格的資料保護政策，預設即阻斷訓練機制。

核心守則：敏感資料的去識別化

輸入 (Input)



過濾器 (Filter)



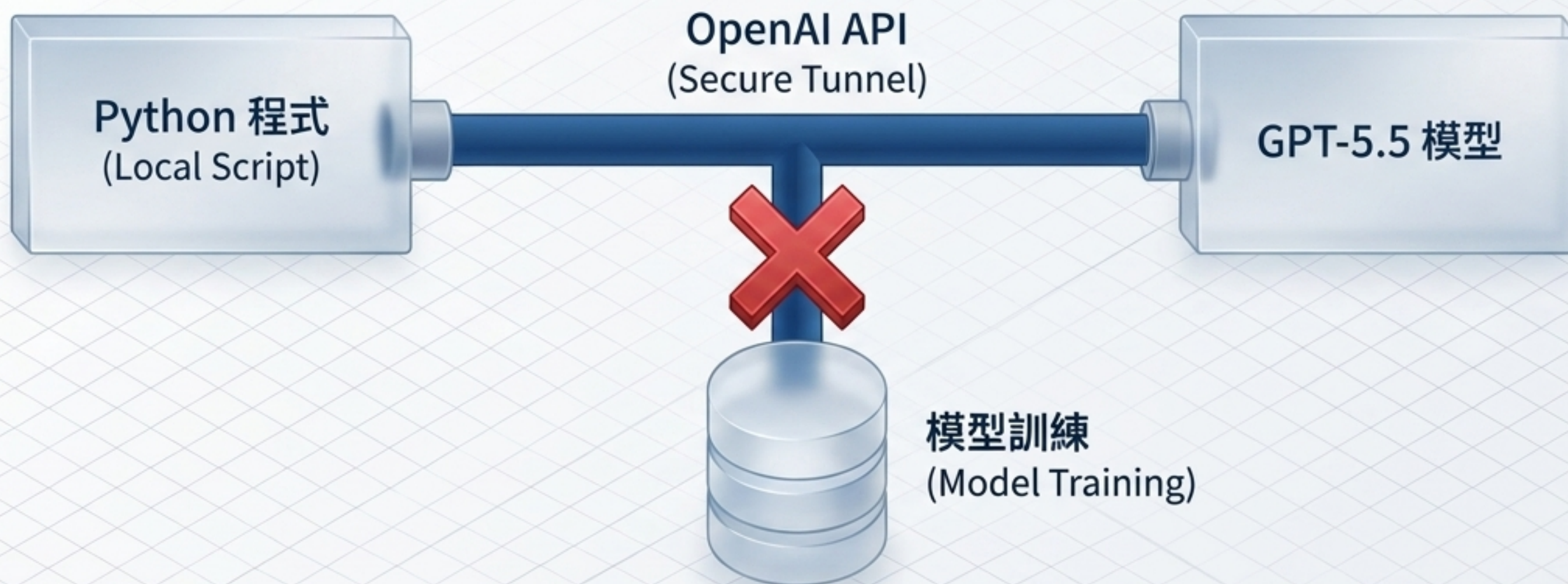
輸出 (Output)



黃金定律：即使已關閉模型訓練，也絕對避免直接輸入未公開研究或機密資料。

這是所有頂尖研究單位採用的標準資料衛生（Data Hygiene）流程。

開發者途徑：透過 API 建立安全通道



機制

自行開發程式，透過 OpenAI API 呼叫模型，而非使用網頁介面。

政策保障

依 OpenAI 現行政策，API 傳送的資料預設不會用於模型訓練。

適用情境

學校專案、研究計畫、醫療等需要較高資料控制，且需批次處理的情境。

終極物理防護：本地端執行 LLM



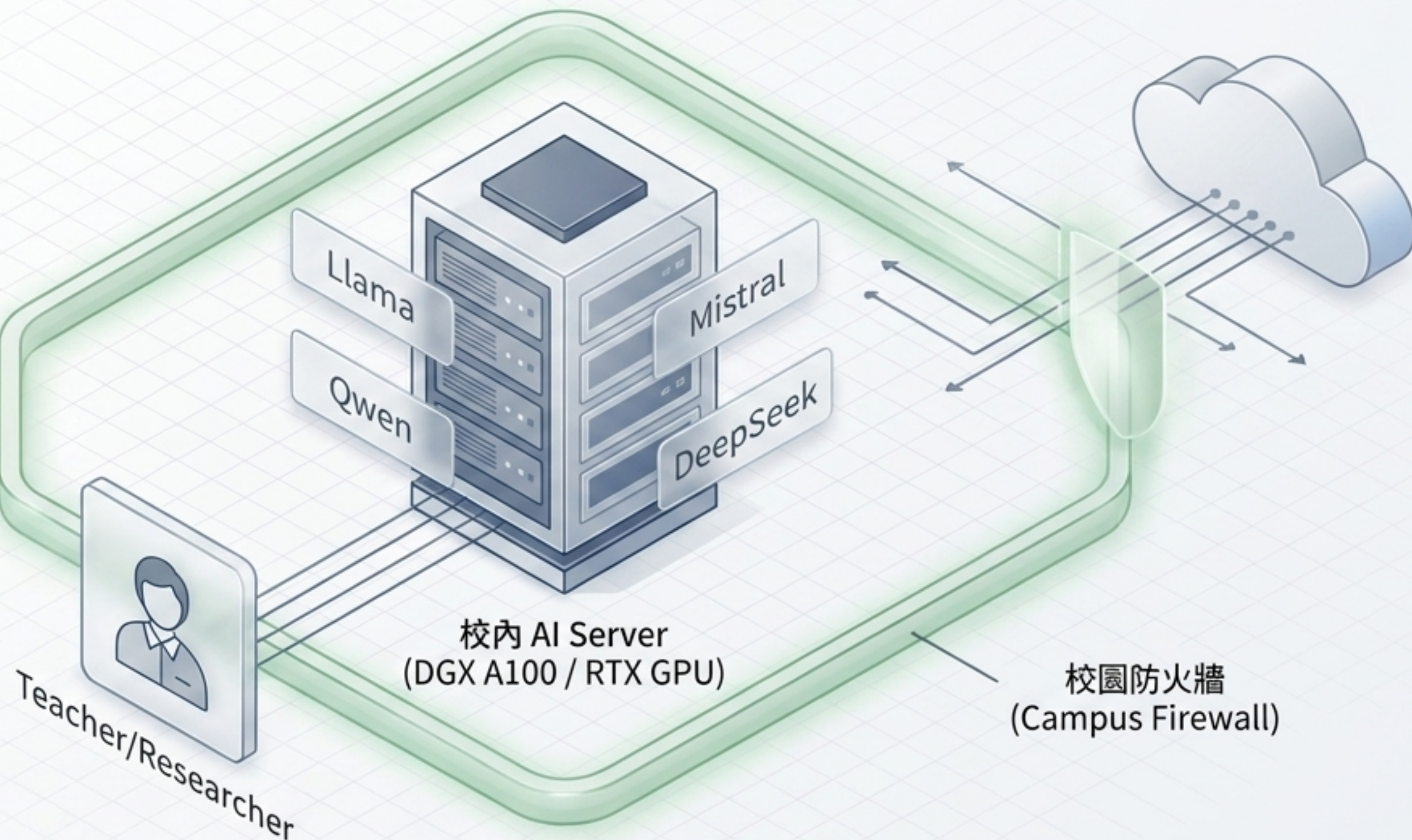
✓ 優點

- ✓ 資料完全留在本地，不經過任何雲端。
- ✓ 支援完全離線使用。

⚠ 缺點

- ⚠ 硬體門檻高（需要較好的 GPU）。
- ⚠ 模型能力通常不如最新的大型雲端模型。
- ⚠ 維護成本較高。

校園數位堡壘：智慧科技學院自架 AI



基礎設施：運用校內建置的 DGX A100 與 RTX GPU Server。

開源部署：自行架設強大開源模型。

絕對優勢：所有資料 100% 留在校內網路，完美契合醫療數據與研究機密。

決策藍圖：依資料敏感度選擇 AI 工具

Tier 3: 極高敏感度

數據醫療資料、個資、產學合作機密。

建議使用校內 DGX A100 (須付費)部署的本地 LLM，確保資料永不離開校園網路。

Tier 2: 中敏感度

數據類型：研究資料、未發表論文、校內文件。

優先使用 OpenAI API（無訓練政策）或校內自行部署的輕量模型。

Tier 1: 低敏感度

一般教學、行政、公開資訊。

使用 ChatGPT 網頁版，必須關閉「使用我的內容改善模型」。

打造安全 AI 環境的三大守則



1. 預設關閉訓練

在任何公共 AI 平台上，第一時間進入設定關閉資料共享。



2. 落實去識別化

不論平台多安全，永遠不要將真實身分證、病歷或機密直接輸入。



3. 善用校園資源

涉及機密與醫療數據時，請直接申請使用校園專屬的 DGX A100 封閉運算環境。