

114 年 AI 應用規劃師(初級、中級)

講義與仿真試題

補充講義

1. 名詞釋義(新增 11 月)

2. 樣板題目解析

五南出版社

Copyright ©黃朝健版權所有

目錄

1. 名詞釋義

表一：生成式 AI 常見名詞 P3

表二：生成式 AI 常見名詞 P5

2. 樣板題目解析

1 月初級樣題 P8

1 月中級樣題 P9

4 月初級樣題 P11

9 月初級樣題 P17

9 月中級樣題 P25

1. 名詞釋義

◎表一：生成式 AI 常見名詞

縮寫	中文名稱	英文名稱 / 常用寫法	一句話簡單說明
RAG	檢索增強生成	Retrieval-Augmented Generation	先從知識庫找相關資料，再把找到的內容餵給模型生成更正確的回答。
LLM	大型語言模型	Large Language Model	以大量文本訓練、能理解與產生自然語言的通用 AI 模型。
Agentic AI	代理式 AI / 多重代理 LLM	Agentic AI	會自行規劃步驟、呼叫工具/API、與多代理協作來完成任務的 LLM 系統。
CoT	思維鏈	Chain of Thought	讓模型逐步寫出推理過程，以提高解題與可解釋性。
ToT	樹狀思維	Tree of Thought	讓模型同時探索多條推理分支，選最佳路徑解題。
Graph Prompting	圖結構提示	Graph Prompting	用節點與邊的圖結構引導模型理解關係與流程。
Auto Prompting	自動提示生成	Auto Prompting	由系統自動產生或優化提示詞，讓模型表現更好。
Few-shot	少樣本提示	Few-shot Prompting	在提示中給少量範例，讓模型依樣學習完成任務。
Zero-shot	零樣本提示	Zero-shot Prompting	不給範例，只靠描述任務讓模型直接完成。
ICL	上下文提示	In-context Learning	模型只靠提示裡的範例「臨時學會」任務規則。
Temperature	溫度參數	Temperature	調整輸出隨機性：高溫更發散，低溫更保守。
LoRA	低秩適應	Low-Rank Adaptation	只訓練小量參數增量，以低成本微調大模型。
KD	知識蒸餾	Knowledge Distillation	用大模型教小模型，讓小模型更精簡還保留表現。
CF	災難性遺忘	Catastrophic Forgetting	模型學新任務時把舊知識「忘掉」的現象。
Text Generation	文字生成	Text Generation	依輸入條件自動產生自然語言文字內容。
Seq2Seq	序列 to 序列	Sequence-to-Sequence	把一段序列轉成另一段序列的模型（如翻譯）。
Multi-Vector	多向量表示	Multi-vector Representation	為同一段內容存多個向量以捕捉不同語意面向。
Chunking	資料塊切分	Chunking	把長文本切成可檢索/可處理的小片段。

Fine-tuning	微調	Fine-tuning	以任務資料持續訓練現有模型以貼近特定需求。
MLOps	機器學習運維	Machine Learning Operations	管理模型從資料、訓練到部署監控的全流程實務。
MCP	微服務通訊平台	MCP (see note)	一套讓服務/工具用標準協議互通的介面層，便於代理調用外部資源。
Solution Graph	解決方案圖譜	Solution Graph	以圖/節點方式列出解題步驟與依賴，幫助規劃與追蹤。
Benchmark	基準測試集	Benchmark	用來客觀比較模型表現的標準題庫與評測方式。
MMLU	大規模多任務語言理解	Massive Multitask Language Understanding	測多領域知識與推理能力的大型綜合測試集。
GSM8K	小學數學 8K	Grade School Math 8K	小學數學應用題集，用來評估基本數學推理。
MATH	高階數學推理	MATH (Benchmark)	難度較高的數學競賽/證明題基準，用於高階推理評測。
C-Eval	中文 MMLU	C-Eval	以中文學科考題評估模型知識與推理的基準。
Copilot	GitHub Copilot	GitHub Copilot	在程式編輯器中即時補全與建議程式碼的 AI 助手。
Codex	OpenAI Codex 技術	OpenAI Codex	早期能把自然語言轉成程式碼的 OpenAI 模型技術。
API-calling	API 呼叫	API Calling	讓模型根據語意自動選擇並呼叫外部工具/服務。
Task-planner	任務規劃器	Task Planner	把目標拆成可執行步驟與順序的系統元件。
Vector Retrieval	向量檢索	Vector Retrieval (Dense Retrieval)	用嵌入向量比對語意相近內容來找資料。
Auto Gen	LLM 代理庫	AutoGen (Agent Library)	支援多代理對話協作與工具調用的 LLM 代理框架。
Task Weaver	LLM 代理庫	TaskWeaver (Agent Library)	以「任務/工具」為中心的代理框架，可編排複雜流程。
Flowise	LLM 應用構建平台	Flowise	以拖拉節點方式快速搭建與部署 LLM 應用的平台。
LLaMA Factory	LLM 訓練框架	LLaMA-Factory	一站式微調 LLaMA 等模型的開源訓練工具箱。
Axolotl	LLM 訓練框架	Axolotl	支援各種微調策略 (SFT、LoRA 等) 的開源訓練框架。
Domain Shift	領域遷移	Domain Shift	部署資料分佈與訓練時不同而導致效能下滑。
No Code	無程式碼	No-code	用圖形介面就能做應用，不用寫程式。

Low Code	低程式碼	Low-code	主要用配置與少量程式碼就能快速開發。
----------	------	----------	--------------------

◎表二：生成式 AI 常見名詞

縮寫	中文名稱	英文名稱 / 常用寫法	一句話簡單說明
AI	人工智慧	Artificial Intelligence	讓電腦用資料學習，做出像人類思考/決策的行為。
XAI	可解釋性 AI	Explainable AI	讓模型的決策理由可被人理解與檢查的方法。
LIME	局部可解釋模型	Local Interpretable Model-agnostic Explanations	用「附近的擬合」來解釋單一預測為何這樣判斷。
Sandbox	監理沙盒	Regulatory Sandbox	在有限風險與監管下試跑新型服務/技術的制度。
ETL	資料擷取、轉換、載入	Extract, Transform, Load	取資料→轉格式/清理→載入到資料庫或倉儲。
OHE	獨熱編碼	One-Hot Encoding	把類別欄位轉成 0/1 向量，利於模型使用。
Cross Enc.	交叉編碼	Cross-Encoder	把兩段輸入合併後一併編碼，學到更細的匹配關係。
Rel. DB	關聯式資料庫	Relational Database	用表格/關聯來管理結構化資料（如 SQL）。
Non-Rel. DB	非關聯式資料庫	Non-relational / NoSQL Database	非表格型資料庫，擅長大規模與彈性結構資料。
SFS	監督式特徵選擇	Supervised Feature Selection	用標記資料挑出對預測最有幫助的特徵。
SL	監督式學習	Supervised Learning	有標答案的資料來訓練模型做預測/分類。
UL	非監督式學習	Unsupervised Learning	沒標答案，從資料中找結構或群組。
Clustering	叢集（分群）	Clustering	依相似度把資料自動分成幾群。
KNN	K-近鄰演算法	K-Nearest Neighbors	看最近的 K 個鄰居多數是什麼就跟著判。
LL	邏輯迴歸	Logistic Regression	用機率方式做二元分類的線性模型。
BC	貝氏分類器	(Naive) Bayes Classifier	用條件機率與先驗知識來做分類。
MSE	均方誤差	Mean Squared Error	預測與真值差的平方平均，越小越好。
CE	交叉熵損失	Cross-Entropy Loss	衡量分類機率分佈的差異，越小越準。

Reg.	正規化	Regularization	加入「懲罰」避免模型過度貼合訓練資料。
L1 / Lasso	L1 正規化	L1 Regularization / Lasso	偏好稀疏權重、可做特徵選擇的正規化。
L2 / Ridge	L2 正規化	L2 Regularization / Ridge	平滑權重、降低過擬合的正規化。
B-V	偏差變異權衡	Bias - Variance Tradeoff	簡單模型易有偏差，高度彈性易有變異，需取平衡。
RNN	循環神經網路	Recurrent Neural Network	透過循環連結處理序列資料（如文字、時間序列）。
Transformer	Transformer 模型	Transformer	用注意力機制並行處理序列的主流架構。
VAE	變分自編碼器	Variational Autoencoder	讓編碼器學會「分佈」並能生成樣本的自編碼器。
Disc. AI	鑑別式 AI	Discriminative AI	專注於分界/條件機率，做分類或回歸。
Gen AI	生成式 AI 概念	Generative AI	學資料分佈並生成新內容（文、圖、音等）。
ReLU	ReLU 激活函數	Rectified Linear Unit (ReLU)	大於 0 直接通過，小於等於 0 置 0 的啟動函數。
VG	梯度消失	Vanishing Gradient	反向傳播時梯度變很小，導致學不動。
EG	梯度爆炸	Exploding Gradient	梯度過大造成訓練不穩或發散。
Pruning	剪枝	Pruning	移除不重要權重/節點，讓模型更小更快。
QL	Q 學習	Q-learning	不用環境模型，學狀態 - 行動價值表的強化學習。
DQL	深度 Q 學習	Deep Q-learning	用神經網路近似 Q 值，解決大型狀態空間。
H-o-t-L	人在迴圈上	Human-on-the-loop	人類在外圍監督/覆核 AI 決策並可介入修正。
FL	聯邦學習	Federated Learning	資料不離開裝置，只傳梯度/參數來共同訓練。
Data Drift	數據漂移 / 領域飄移	Data/Domain Drift	資料分佈隨時間改變，模型表現跟著下滑。
Inf.	即時推論	Real-time Inference	在低延遲要求下即刻給出模型結果。
Batch Inf.	批次推論	Batch Inference	把很多資料累積後一次離線推論。
On-prem	地端部署	On-premise Deployment	在自家機房/自管硬體上部署模型與服務。

Fine-tuning	微調（模型部署技術）	Fine-tuning	以少量任務資料續訓，讓基礎模型貼合情境。
-------------	------------	-------------	----------------------

2. 樣板題目解析

▼1 月初級樣題

科目一：人工智慧基礎概論

題號	題目（精簡）	選項	正確解答	解題重點
1	關於 AI，下列何者正確？	(A) 只處理結構化 (B) 涵蓋多領域 (C) 只學術用 (D) 不能用於金融	B	AI 是跨領域技術與方法的集合。
2	何者非大數據特性？	(A) 量大 (B) 速度快 (C) 多樣性 (D) 儲存位置固定	D	大數據強調 3V/5V，與固定儲存位置無關。
3	銀行建置聊天機器人需哪些領域？	(A) 資料庫 (B) 機器學習+NLP (C) 網頁開發 (D) 資安	B	對話系統核心是 ML/NLP。
4	生成式 AI 最可能用於創建新圖像/影片？	(A) 品質檢測 (B) 醫影分析 (C) 監控 (D) 虛擬現實圖像	D	生成式模型擅長新媒體內容合成。

科目二：生成式 AI 應用與規劃

5	AI 治理中國際合作的重要性？	(A) 統一標準 (B) 避免濫用 (C) 促進轉移 (D) 以上皆是	D	國際協作同時涵蓋標準、倫理與技轉。
6	No/Low-Code 平台主要特色？	(A) 大量寫碼 (B) 模板快速建置 (C) 僅供專業 (D) 只能靜態網站	B	以視覺化/模板加速開發、降門檻。
7	關於 No-Code AI 工具正確敘述？	(A) 已完全取代傳統 (B) 只適大型企業 (C) 降低複雜與成本的新方法 (D) 功能完全相同	C	核心價值是降複雜、降成本、提速。
8	生成內容時的品質措施？	(A) 直接用於學術 (B) 適當標注來源 (C) 減少審查 (D) 排除所有生成資料	B	引用/出處標注是最基本治理要件。
9	教師如何引導學生使用生成式 AI？	(A) 不用 (B) 無限制 (C) 訂立清楚規範並說明 (D) 只鼓勵用來交作業	C	訂規範+示範與倫理教育並行。
10	生成式 AI 風險管理	(A) 內容偏見/歧視 (B) 系統中斷 (C) 儲存成本	A	偏見、歧視、幻覺

	中的倫理風險？	(D) 培訓成本		等屬倫理風險。
--	---------	----------	--	---------

▼1 月中級樣題

科目一：人工智慧技術應用與規劃

題號	題目（精簡）	選項	正確解答	解題重點
1	何者未用 AI/ML？	(A) 用預定義規則下棋 (B) DNN 語音辨識 (C) 自駕 + 定義規則 (D) NLP 聊天機器人	A	純規則系統非 ML。
2	文本處理中「將文本切成詞彙單位」	(A) 詞形還原 (B) 停用詞移除 (C) 斷詞 / Tokenization (D) TF-IDF	C	Tokenization 是語料前處理首步。
3	NLP 在 ML 應用中的主要用途	(A) 情緒分析 (B) 影像識別 (C) 預測維護 (D) 供應鏈優化	A	情緒/意見探勘為典型 NLP 任務。
4	關於深度學習何者不正確？	(A) CNN 適合影像 (B) ReLU 可減緩梯度問題 (C) RNN 適序列 (D) Elman 網路適影像	D	Elman (RNN 變體) 非影像主力。
5	業界部署趨勢	(A) 採用 AutoML 增加 (B) 轉向更簡單演算法 (C) 雲端使用下降 (D) 依賴手調超參	A	AutoML/雲端/MLOps 為主流。

科目二：大數據處理分析與應用

題號	題目（精簡）	選項	正確解答	解題重點
6	身高常態分佈的檢定（多年級）何者不正確？	(A) 一年級 vs 二年級 → t 檢定 (B) 一/二/三年級三組 → t 檢定 (C) 二/三/四年級 → F 檢定 (D) 一年級均數=170 → 卡方檢定	D	檢定平均數用 t/z；卡方多用變異/列聯。題目答案標為 D。
7	關於 ROC 正確為何？	(A) 繪真陽性率 vs 假陽性率 (B) 評估準確率 (C) AUC 一定=1 (D) 只適二元	A	ROC 核心座標為 TPR-FPR。
8	何者不屬於特徵工程？	(A) 轉換 (B) 萃取 (C) 挑選 (D) 預測	D	預測屬於建模/推論非 FE。
9	用過往銷售資料預測下季銷量	(A) 決策樹分類 (B) K-means (C) PCA (D) 線性迴歸	D	目標是連續量 → 迴歸。
10	低結構化文本/影像的 FE 方法	(A) 改善 (B) 建構 (C) 特徵學習 (D) 選擇	C	由模型自學特徵（如深度學習）。

科目三：機器學習技術與應用

題號	題目（精簡）	選項	正確解答	解題重點
11	MapReduce： Map/Reduce 在做什麼？	(A) Map：映射；Reduce：統合歸納 (B) Map：地圖式搜尋 (C) Reduce：過濾 (D) Reduce：生成更多資料	A	Map 做對映，Reduce 聚合彙總。
12	「加寬不加深」的 CNN 架構	(A) R-CNN (B) Inception (C) ResNet (D) VGG19	B	Inception 以並聯/寬化著稱。
13	訓練誤差低、測試誤差高代表？	(A) 泛化強 (B) 過度擬合 (C) 欠擬合 (D) 無相關	B	高方差情形。
14	迴歸模型常用評估指標	(A) R^2 (B) F1 (C) AUC (D) Precision	A	決定係數量化解釋力。
15	深度學習中的數據潛在問題	(A) 標註品質直接影響性能 (B) 品質完美 (C) 無類不平衡 (D) 不需領域知識	A	Data quality/annotation 是關鍵瓶頸。

▼4 月初級樣題

科目一：人工智慧基礎概論

題號	題目（精簡）	選項	正確解答	解題重點
1	AI 敘述何者正確	(A) 僅能處理結構化數據 (B) 涵蓋多種專業領域與技術 (C) 只能在學術研究中應用 (D) 無法應用於金融領域	B	AI 是一個跨學科的領域，涵蓋多種專業技術。
2	適合訓練電腦下圍棋、自動駕駛等動態重複地互動的問題	(A) 監督式學習 (B) 非監督式學習 (C) 半監督式學習 (D) 強化學習	D	針對動態重複互動、追求最大化累積獎勵的問題，最適合使用強化學習 (Reinforcement Learning)。
3	深度學習模型中，降低過擬合 (Overfitting) 的方法	(A) 增加訓練數據量 (B) 增加模型的複雜度 (C) 增加學習率 (D) 增加正則化項	D	增加正則化項 (Regularization) 是在深度學習中常用來限制模型複雜度，從而降低過擬合問題的方法。
4	AI 治理中，國際合作的重要性	(A) 統一 AI 發展標準 (B) 避免 AI 技術的濫用 (C) 促進 AI 技術的轉移 (D) 以上皆是	D	AI 治理中的國際合作重要性包括統一標準、避免技術濫用、以及促進技術轉移，因此以上皆是。
5	設計辨識並過濾垃圾郵件的系統應選擇何種演算法	(A) 監督式學習 (B) 非監督式學習 (C) 半監督式學習 (D) 強化學習	A	垃圾郵件過濾是典型的分類問題，屬於有標籤數據的學習，應選擇監督式學習 (Supervised Learning)。
6	歐盟《人工智慧法》(AIA) 中，社會信用評分（基於年齡、種族）與公眾場所遠程生物辨識的風險等級	(A) 不可接受風險 (B) 高風險 (C) 有限風險 (D) 小或低風險	A	根據歐盟 AIA，這類可能侵犯基本人權或導致歧視的應用（如基於敏感特徵的社會信用評分、廣泛的生物辨識執法），屬於不可接受風險。
7	何者非大數據時代資料的特性	(A) 資料量大 (B) 資料變動速度快 (C) 資料多樣性 (D) 資料存儲位置固定	D	大數據的特性通常包含資料量大 (Volume)、變動速度快 (Velocity) 和多样性 (Variety)，資料存儲位置固定並非其特性。
8	關於 K 平均法 (K-means)，何者「不」正確	(A) 找出 k 個互不交集的群集 (B) 不同的起始群集中心，可能造成不同分群結果 (C) 容易受雜訊與離群值影響其群集中心 (D) 可以處理類別型資料	D	K-means 是一種基於距離的聚類方法，主要適用於數值型資料，而非類別型資料。

9	哪一種資料類型不屬於非結構化資料	(A) X 光醫學影像 (B) 監控錄影畫面 (C) 客服電話錄音 (D) 組織內部的關係型資料庫記錄	D	關係型資料庫記錄是依照預定結構儲存的，因此屬於結構化資料，而非非結構化資料。
10	最常用來儲存員工年齡、年資、貨品銷售量等資料的型態	(A) 文字型 (B) 數值型 (C) 日期型 (D) 布林型	B	年齡、年資和銷售量都是可以量化計算的數值，應使用數值型 (Numeric) 資料。
11	在品質管理中，標準差顯著偏大通常意味著什麼	(A) 資料點高度集中，產品質量穩定 (B) 生產過程波動大，產品品質不穩定 (C) 資料無法反映產品實際狀況 (D) 中位數數值高，品質良率較高	B	標準差是衡量數據分散程度的指標。標準差偏大意味著數據變異大，反映生產過程波動大，產品品質不穩定。
12	何者不屬於資料集中趨勢衡量的方法	(A) 平均數 (B) 中位數 (C) 眾數 (D) 標準差	D	平均數、中位數和眾數是用來衡量資料的集中趨勢，而標準差 (Standard Deviation) 是用來衡量變異程度 (分散程度) 的方法。
13	醫院研究心血管疾病成因，收集屬性變數，建立模型	(A) 決策樹 (B) 線性迴歸 (C) DBSCAN (D) K-means 聚類	A	此研究目的是根據屬性變數將個體劃分為「病患」或「正常人」兩種類別，屬於分類問題，決策樹 (Decision Tree) 是適合的分類模型。
14	銀行建立聊天機器人可透過哪一種領域技術達成	(A) 資料庫管理技術 (B) 機器學習與自然語言處理 (C) 網頁開發技術 (D) 網路安全技術	B	聊天機器人需要理解人類的文本或語音，這仰賴於自然語言處理 (NLP) 技術，而 NLP 屬於機器學習的應用範疇。
15	線性迴歸模型最適合解決哪種類型的問題	(A) 圖像分類 (B) 銷售額預測 (C) 聚類分析 (D) 遊戲策略學習	B	線性迴歸 (Linear Regression) 模型是用來預測連續數值輸出的，最適合解決如銷售額預測等迴歸問題。
16	何者不是常見的特徵選取技術或方法	(A) 皮爾森積差相關分析 (B) 主成分分析 (C) 迴歸分析 (D) 隨機森林	C	迴歸分析 (Regression Analysis) 是一種統計建模方法，而皮爾森相關分析、PCA 和隨機森林都可以作為特徵選取或降維的技術。
17	交叉驗證的主要目的	(A) 提高模型的訓練速度 (B) 驗證數據是否線性可分 (C) 減少模型的過擬合風險 (D) 測試模型的容錯能力	C	交叉驗證 (Cross-validation) 透過將數據分為多個子集進行訓練和驗證，以評估模型在未知數據上的表現，從而減少過擬合風險。
18	神經網路與傳統機器學習模型的主要區別	(A) 無法處理非線性數據 (B) 透過多層結構學習複雜特徵 (C) 只適用於迴歸問題 (D) 不需要大量數據支持	B	神經網路 (尤其是深度學習) 的核心優勢在於其多層結構，使其能夠自動地從原始數據中學習和提取複雜的階層式特徵。
19	關於生成對抗網路 (GAN) 的描述何者正確	(A) GAN 由生成器和鑑別器組成 (B) 僅用於分類問題 (C) 結果始終高度可解釋 (D) 不能生成高品質的數據	A	生成對抗網路 (GAN) 的基本結構由生成器 (Generator) 和鑑別器 (Discriminator) 兩個對抗的網路組成。

20	生成式 AI 最可能被使用來創建新的圖像或影片內容的應用領域	(A) 產品品質檢測 (B) 醫學影像分析 (C) 監控系統 (D) 虛擬現實圖像	D	生成式 AI 的主要功能是創造新內容，因此最適用於如虛擬現實圖像的內容生成。
21	目前生成式 AI 的主要應用，不包括下列哪一項	(A) 創建合成數據樣本 (B) 模擬數據分佈 (C) 分類醫學影像 (D) 生成文本	C	分類醫學影像屬於鑑別式 AI (Discriminative AI) 的任務，而非生成式 AI 的主要應用。
22	生成式 AI 支援鑑別式 AI 的典型案列	(A) 模擬交通場景以訓練自動駕駛模型 (B) 使用 CNN 對腫瘤分類 (C) 使用 SVM 分析風險 (D) 創建更好的分類演算法	A	生成式 AI 可用於創建合成數據（如模擬交通場景），這些數據隨後可用於訓練鑑別式模型（如自動駕駛模型），這是一種典型的支援案列。
23	關於自然語言處理 (NLP) 核心技術，何者不正確	(A) 語音識別技術將語音轉換為文本 (B) 自然語言生成技術可以生成自然流暢的文本 (C) 語意分析技術主要用於語音識別和機器翻譯 (D) 機器翻譯技術自動翻譯文本	C	語意分析的目標是理解文本的含義，雖然它對機器翻譯等有幫助，但將其應用範圍限制於語音識別和機器翻譯是不正確的，語意分析是許多 NLP 應用（如問答、摘要）的基礎。
24	關於「負責任的 AI」，下列敘述何者較為正確	(A) AI 系統的開發者對 AI 系統的行為負責 (B) AI 系統的使用者對 AI 系統的結果負責 (C) AI 系統本身對其行為負責 (D) 政府對 AI 系統的發展負責	A	在負責任的 AI 框架下，AI 系統的開發者（或組織）通常被視為對系統的設計、行為和潛在後果負主要責任。
25	關於生成式 AI 的基本原理，下列敘述何者較正確	(A) 通過分析大量數據來生成新數據，模擬數據分佈以創造與訓練數據相似的結果 (B) 主要通過預定義的規則來進行數據處理和分類 (C) 專注於數據分類和迴歸預測 (D) 通過自動化的方式清洗數據	A	生成式 AI 的核心原理是學習訓練數據的底層分佈，並利用此分佈生成新的、相似的數據樣本。

科目二：生成式 AI 應用與規劃

題號	題目（精簡）	選項	正確解答	解題重點
1	No Code / Low Code 平台的主要特色	(A) 需撰寫大量程式碼 (B) 運用模板快速建立應用程式 (C) 僅供專業開發人員使用 (D) 只能製作靜態網站	B	No Code / Low Code 平台的主要特色是利用視覺化介面和模板，快速建立應用程式，大幅減少手動編程的需求。
2	關於 No Code AI 工具，何者最為準確	(A) 完全取代傳統的 AI 開發模式 (B) 只適用於大型企業 (C) 是一種降低 AI 技術複雜性和開發成本的新興方法 (D) 工具都具有完全相同的功能和性能	C	No Code AI 工具的目的是民主化 AI 技術，透過降低技術門檻，來減少開發複雜性和成本。
3	哪種情況，選擇 Low Code 平台比 No Code 平台更適合	(A) 非技術人員快速開發 (B) 應用需求簡單，無需自訂 (C) 需較複雜的業務邏輯並使用自訂整合功能 (D) 預算和時間極度有限	C	Low Code 平台雖然仍使用視覺化工具，但允許開發者加入自訂程式碼和複雜的整合功能，更適合處理複雜的業務邏輯。
4	生成式 AI 與 No Code / Low Code 平台最不適合的應用	(A) 自動生成程式碼 (B) 自動化生成行銷文案 (C) 快速開發個人化 App (D) 自動化生成法律判決	D	自動化生成法律判決涉及高風險、倫理與社會責任，需要高度的人類判斷和法律合規性，因此最不适合交由自動化或 No Code/Low Code 系統執行。
5	關於 No Code / Low Code 平台，下列敘述何者較正確	(A) 兩者完全相同 (B) Low Code 平台不需要任何程式設計知識 (C) Low Code 平台更適合開發靈活且可擴展的解決方案 (D) No Code 平台可以無限客製化	C	Low Code 平台由於允許自訂程式碼和整合，因此比 No Code 平台提供更大的靈活性和可擴展性。
6	使用 Low-Code 平台進行開發時，企業應特別留意哪一項潛在風險	(A) 可能造成企業內部敏感資料的洩露 (B) 難以進行大規模的應用擴展和維護 (C) 開發成本將大幅增加 (D) 可能有未經 IT 部門管理的應用程式擴散	D	Low-Code 平台易於使用，可能導致非 IT 部門快速創建應用程式，造成**「影子 IT」(Shadow IT)**，即未經正式管理的應用程式擴散，帶來潛在風險。
7	適用於改善客戶體驗的技術方案	(A) 智慧排程系統 (B) 消費行為洞察模型 (C) 預測性維護工具 (D) 自然語言處理 (NLP) 和生成式回應模組	D	改善客戶體驗通常需要即時且個性化的互動，NLP 和生成式回應模組能夠提供智慧化的客戶服務和問答。

8	下列哪一項技術是生成式 AI 的基礎	(A) 決策樹模型 (B) 聚類演算法 (C) 生成對抗網路 (D) 隨機森林技術	C	生成對抗網路 (GAN) 是當前許多生成式 AI 應用 (如圖像生成) 的基礎或核心技術之一。
9	能使用 DALL·E-2 生成各式逼真的圖片，最關鍵的應用技術為何	(A) 卷積神經網路 (CNN) (B) 生成對抗網路 (GAN) (C) 擴散模型 (Diffusion Model) (D) 自然語言處理 (NLP)	C	如 DALL·E-2 和 Midjourney 等先進的圖像生成模型，其核心技術是擴散模型 (Diffusion Model)。
10	何者不是生成式 AI 核心技術	(A) Variational Autoencoders (B) Generative Adversarial Networks (C) Visual Geometry Group (VGG) (D) Autoregressive Models	C	VAE、GAN 和 Autoregressive Models 都是用於生成數據或學習數據分佈的架構；VGG 是用於圖像分類的卷積神經網路 (CNN) 架構，屬於鑑別式模型。
11	使用生成式 AI 技術或工具生成內容時，應採取哪一項措施以確保內容品質	(A) 使用內容直接進行學術報告 (B) 適當標注引用來源 (C) 減少人工參與的審查過程 (D) 排除所有生成的資料	B	為了維護學術和專業誠信，並確保內容品質的可追溯性，應當適當標注引用來源。
12	哪一項不是生成式 AI 工具在使用體驗方面的優化方向	(A) 提供更直觀的操作設計 (B) 支援自然語言指令 (C) 提供智慧化的參數調整建議 (D) 限制使用者自訂生成內容	D	生成式 AI 的優化旨在提升用戶的效率和創造力，限制使用者自訂生成內容不屬於優化方向。
13	學校教師如何引導學生正確使用生成式 AI 工具	(A) 不應使用 AI 工具於教學場域 (B) 無限制地使用 AI 工具 (C) 訂立清晰的使用規範並進行說明 (D) 僅鼓勵學生利用 AI 完成課堂作業	C	引導學生正確且負責任地使用 AI，最重要的方法是訂立清晰的使用規範並向學生說明倫理和學術要求。
14	企業有效支援生成式 AI 的運行，內部 IT 環境最需要具備何種條件	(A) 提供更多的辦公設備 (B) 精簡企業內部流程 (C) 擁有高效能運算資源與彈性儲存空間 (D) 增加部門之間的交流機會	C	運行和訓練生成式 AI 模型需要大量的計算能力和存儲容量，因此高效能運算資源 (HPC) 與彈性儲存空間是必要的 IT 條件。
15	企業希望透過 AI 理解文本訊息，並調用圖片和影片進行回覆，應選擇哪一種模型	(A) 強化學習模型 (B) 多模態模型 (C) 圖像分類模型 (D) 單模態大語言模型	B	該需求涉及處理和整合多種數據類型 (文本、圖片、影片)，需要多模態模型 (Multimodal Model) 來理解並在不同模態間進行轉換。
16	哪一種情境最能展現提示工程 (Prompt Engineering) 的價值	(A) 輸入模糊問題，AI 給出確定答案 (B) 輸入具體問題，AI 給出相關但不完全符合的答案 (C) 輸入精準有架構的問題，AI 生成符合架構的答案 (D) 輸入簡單問題，AI 給出複雜答案	C	提示工程的價值在於優化輸入指令，使 AI 能夠理解需求並生成精準且符合預期結構的答案。
17	生成式 AI 在圖像生成領域的發展趨勢	(A) 僅能生成低解析度圖像 (B) 逼真度顯著提升，並可處理更複雜的圖像	B	隨著技術進步，生成式 AI 在圖像生成方面趨向於更高的逼真度、更豐富

		生成任務 (C) 只能生成靜態圖像 (D) 發展主要集中於圖像風格轉換		的風格多樣性，以及處理更複雜的生成任務。
18	關於生成式 AI，敘述哪些正確 (A. 不帶有偏見; B. 具有高度準確性; C. 準確性建議需經過人類審核; D. 每次生成內容可能不同; E. 具有高度安全性)	(A) C、D (B) A、C、D (C) A、B、D、E (D) A、C、D、E	A	生成式 AI 由於訓練數據的影響可能帶有偏見，且輸出不保證高度準確，因此需要人類審核 (C)；同時，其生成過程具有一定的隨機性，導致每次生成的內容可能不同 (D)。
19	財務部門導入生成式 AI 結合自動化系統進行報告產出，最可能實現的效益	(A) 增加財務收益 (B) 顯著提升報告準確性並減少人工錯誤 (C) 加速市場推廣活動 (D) 提高品牌曝光度	B	導入 AI 自動化報告產出，最直接且顯著的效益是減少重複性的人工錯誤，從而提升報告的準確性。
20	在生成式 AI 導入過程中，資料安全與隱私保護的哪一方面是最重要的考量	(A) 設定目標優先級 (B) 增強客服回饋能力 (C) 資料視覺化能力 (D) 權限控管與合規要求	D	在處理敏感數據的 AI 系統中，權限控管 (確保只有授權人能存取) 和合規要求 (遵守隱私法規) 是資料安全與隱私保護中最重要考量。
21	企業將資料安全管理外包給第三方服務供應商，屬於哪種風險應對策略	(A) 風險緩解 (B) 風險轉移 (C) 風險接受 (D) 風險規避	B	將風險的處理責任轉移給第三方 (如保險或外包服務商)，屬於風險轉移策略。
22	在生成式 AI 的風險管理中，下列哪一項屬於倫理風險	(A) AI 生成的內容可能帶有偏見或歧視 (B) 系統運行中斷可能導致企業業務受到影響 (C) 因資料需求增加而引起的存儲成本上升 (D) 員工培訓成本增加	A	偏見或歧視屬於社會公平和公正範疇，是 AI 系統運行中典型的倫理風險。
23	在企業導入 AI 的實施/營運階段，為持續發揮價值，下列步驟的正確排序應為何 (A. AI 價值擴散; B. 上線部署; C. 模型監控與優化)	(A) ACB (B) ABC (C) BAC (D) BCA	D	實施/營運階段的邏輯順序為：首先上線部署 (B)，然後對模型進行監控與優化 (C)，最後在確認效果後進行價值擴散 (A)，故正確排序為 BCA。
24	管理生成式 AI 系統的隱私風險時，哪一種技術最能確保數據使用的安全性	(A) 強化學習 (B) 深度學習 (C) 零信任架構 (D) 注意力機制	C	零信任架構 (Zero Trust Architecture, ZTA) 是一種安全策略，要求對所有嘗試存取系統資源的用戶和設備進行嚴格的驗證，即使是內部網路，能有效確保數據使用的安全性。
25	驗證概念驗證 (POC) 時，若企業希望確保模型生成的公平性，最適合採用哪種評估策略	(A) 壓力測試 (B) 對抗性測試 (C) 偏差檢測 (D) 延遲測試	C	為了確保模型的公平性 (Fairness)，必須使用偏差檢測 (Bias Detection) 策略來識別和量化模型輸出中可能存在的偏見或歧視。

▼9 月初級樣題

科目一：人工智慧基礎概論

題號	題目（精簡）	選項	正確解答	解題重點
1	關於 AI，下列敘述何者正確？	(A) AI 僅能處理結構化數據的分析 (B) AI 涵蓋多種專業領域與技術 (C) AI 系統只能在學術研究中應用 (D) AI 無法應用於金融領域	B	人工智慧 (AI) **是一個廣泛的領域，涵蓋多種專業領域與技術。
2	最適合訓練電腦下圍棋、自動駕駛等動態重複地互動的問題？	(A) 監督式學習 (B) 非監督式學習 (C) 半監督式學習 (D) 強化學習	D	強化學習 (Reinforcement Learning) 透過與環境的動態互動來學習最佳策略，最適合這類序列決策問題。
3	深度學習模型中，通常用來降低過擬合 (Overfitting) 問題？	(A) 增加訓練數據量 (B) 增加模型的複雜度 (C) 增加學習率 (D) 增加正則化項	D	**增加正則化項 (Regularization) **是常用的技術，用來懲罰複雜模型，從而降低過擬合風險。
4	哪一項敘述符合 AI 治理的核心原則？	(A) AI 系統的風險等級應在設計階段預先評估，並於部署後持續監控 (B) 只要 AI 模型輸出結果準確，是否公平或透明並非治理重點 (C) AI 治理應聚焦在企業內部自行負責，不需外部規範參與 (D) 對於低風險 AI 系統皆可免除透明性與問責機制之要求	A	AI 治理的核心原則包括在設計階段評估風險，並在部署後持續監控。
5	設計辨識並過濾垃圾郵件的系統，應該選擇哪一種機器學習演算法？	(A) 監督式學習 (B) 非監督式學習 (C) 半監督式學習 (D) 強化式學習	A	垃圾郵件過濾是分類任務，且需要已標記的數據（是/否垃圾郵件），因此屬於監督式學習。
6	歐盟《人工智慧法》(AIA) 的風險分級中，基於敏感特徵的信用評分及遠程生物辨識系統屬於哪一風險等級？	(A) 不可接受風險 (B) 高風險 (C) 有限風險 (D) 小或低風險	A	這種基於敏感特徵進行評分或在公共場所使用遠程生物辨識的應用，被視為對個人權利和自由構成威脅的不可接受風險。
7	語音輸入生成圖像的任務分類？	(A) 語音生成圖片 (B) 圖像生成語音 (C) 語音翻譯 (D) 圖像標註	A	根據描述，模型接收語音內容後產生圖像，屬於語音生成圖片的任務。
8	下列何者非大數據時代資料的特性？	(A) 資料量大 (B) 資料變動速度快 (C) 資料多樣性 (D) 資料存儲位置固定	D	大數據通常具備 3V（或更多 V），包括資料量大、變動速度

				快、多樣性，而其存儲位置通常是分散或彈性調整的，非固定。
9	關於 K 平均法 (K-means)，何者「不」正確？	(A) 希望找出 k 個互不交集的群集 (B) 不同的起始群集中心，可能會造成不同的分群結果 (C) 容易受雜訊與離群值 (Outlier) 影響其群集中心 (D) 可以處理類別型資料	D	K-means 依賴歐氏距離計算，無法直接處理類別型資料。
10	判斷公司財務資料中 1,000,000 元支出是否為異常值，最合適的處理方式？	(A) 以 Z-score 方法量化異常程度 (B) 以資料中的眾數作為參考基準 (C) 先透過 PCA 降低維度 (D) 直接將該筆金額替換為資料中位數	A	Z-score 方法能有效量化數據點偏離平均值的程度，適合用於判斷極端值或異常值。
11	哪一種資料類型不屬於非結構化資料？	(A) X 光醫學影像 (B) 監控錄影畫面 (C) 客服電話錄音 (D) 組織內部的關係型資料庫記錄	D	關係型資料庫記錄通常以表格形式儲存，屬於結構化資料。
12	員工年齡、員工年資、貨品銷售量等資料，最常用哪種資料型態儲存？	(A) 文字型 (B) 數值型 (C) 日期型 (D) 布林型	B	這些欄位的值是連續或離散的數字，屬於**數值型 (Numeric) **資料。
13	欲預測新的單一客戶是否可能完成購買行為 (已知購買結果作為標籤)，最合適的學習方式與資料搭配？	(A) 使用非監督式學習，分析點擊路徑 (B) 使用非監督式學習，將客戶分群後預測轉換率 (C) 使用監督式學習，針對未標記資料直接預測 (D) 使用監督式學習，以已知購買結果作為標籤進行訓練	D	預測購買行為是二元分類，且擁有已知標籤 (是否購買)，因此最適合採用監督式學習。
14	產品的生產過程中標準差顯著偏大，通常意味著什麼？	(A) 資料點高度集中 (B) 生產過程波動大，產品品質不穩定 (C) 資料無法反映產品實際狀況 (D) 中位數數值高	B	標準差衡量資料的分散程度，標準差大意味著數據波動大，即生產過程波動大，產品品質不穩定。
15	專案目標是識別高價值客戶的行為模式，對於明顯高於其他資料點的數值 (離群值)，哪一種處理方式最為合適？	(A) 立即刪除離群值 (B) 視為錯誤值並全部替換為平均值 (C) 保留離群值並標註為高價值異常點，納入後續模型訓練考量 (D) 將離群值全數轉換為中位數	C	當離群值可能代表有意義的現象 (如高價值行為) 時，應保留並標註，納入模型訓練考量。
16	下列何者不屬於資料集中趨勢衡量的方法？	(A) 平均數 (Mean) (B) 中位數 (Median) (C) 眾數 (Mode) (D) 標準差 (Standard Deviation)	D	標準差 (Standard Deviation) 是衡量資料變異程度 (分散程度) 的方法，不屬於集中趨勢。
17	收集病患 (50 名) 與正常人 (150 名) 的年齡、血壓等屬性變數，適合使用下列哪一種機器學習模型來建立？	(A) 決策樹 (B) 線性迴歸 (C) DBSCAN (D) K-means 聚類	A	這是有標籤的分類問題 (病患/正常人)，決策樹是一種適合的監督式分類模型。
18	客戶姓名在不同系統中拼寫不一致，導致資料無法正確	(A) 資料轉換 (Data Transformation) (B) 資料擷取 (Data Extraction) (C) 資料載入	A	處理資料清洗、格式標準化和修正不一致等操作，屬於 **ETL

	對應，此類資料品質問題應該在 ETL 哪一個流程步驟中進行處理？	(Data Loading) (D) 型態轉換 (Type Conversion)		流程中的資料轉換 (Transformation) **步驟。
19	銀行想建立聊天機器人，可透過下列哪一種領域技術來達成？	(A) 資料庫管理技術 (B) 機器學習與自然語言處理 (C) 網頁開發技術 (D) 網路安全技術	B	聊天機器人涉及理解和生成人類語言，核心技術是機器學習與自然語言處理 (NLP)。
20	線性迴歸模型最適合解決哪種類型的問題？	(A) 圖像分類 (B) 銷售額預測 (C) 聚類分析 (D) 遊戲策略學習	B	線性迴歸用於預測連續的數值結果，如銷售額預測。
21	AI 模型遭客戶抱怨不公平，不同族群預測結果存在顯著差異，根據金管會指引，最適當的處置方式？	(A) 重新訓練新的模型，並重新部署以修正結果 (B) 啟動人工覆核與調整機制，並持續監控族群間預測效果 (C) 記錄模型預測結果並提交備查 (D) 由風控或合規單位先行審視其公平性影響，再決定是否啟動調整機制	B	《金融業運用人工智慧 (AI) 指引》要求在出現不公平情況時，應啟動人工覆核與調整機制，並持續監控。
22	請問下列何者不是常見的特徵選取技術或方法？	(A) 皮爾森積差相關分析 (B) 主成分分析 (PCA) (C) 迴歸分析 (D) 隨機森林	C	迴歸分析是一種建模方法，而其他選項 (相關分析、PCA、基於模型的特徵重要性如隨機森林) 是常見的特徵選取或降維技術。
23	交叉驗證的主要目的是什麼？	(A) 提高模型的訓練速度 (B) 驗證數據是否線性可分 (C) 減少模型的過擬合風險 (D) 測試模型的容錯能力	C	交叉驗證旨在評估模型在未見數據上的泛化能力，從而減少模型的過擬合風險。
24	哪項指標最適合用來衡量模型在偵測異常停機時的「漏報率」(即未能正確偵測出異常事件的比例)？	(A) 準確率 (B) 召回率 (C) F1 分數 (D) 假陽性率	B	召回率 (Recall) 衡量的是所有實際為陽性 (異常停機) 的事件中，模型正確偵測出的比例。召回率低即表示漏報率高。
25	神經網路與傳統機器學習模型的主要區別是什麼？	(A) 神經網路無法處理非線性數據 (B) 神經網路透過多層結構學習複雜特徵 (C) 神經網路只適用於迴歸問題 (D) 神經網路不需要大量數據支持	B	神經網路的主要特點是其透過多層結構 (深層學習) 來自動學習複雜的特徵表示。
26	下列關於生成對抗網路 (GAN) 的描述正確的是哪一項？	(A) GAN 由生成器和鑑別器組成 (B) GAN 僅用於分類問題 (C) GAN 的結果始終高度可解釋 (D) GAN 不能生成高品質的數據	A	GAN 的核心是兩個網路：**生成器 (Generator) 和鑑別器 (Discriminator) **相互對抗學習。
27	希望減少輸入特徵的維度，以提升模型運算效率，並觀察變數間潛在的整體結構關係，最適合的方法？	(A) 套用主成分分析 (PCA) 以擷取主要變異方向並轉換新變數 (B) 利用離散化方法將連續變數轉為分類型欄位 (C) 使用標準化方法將所有特徵縮放至相同數值區間 (D) 以 ETL 技術移除空值欄位	A	主成分分析 (PCA) 是最適合用於降維同時最大化保留數據變異性 (資訊量) 的方法。

28	關於目前生成式 AI 的主要應用，不包括下列哪一項？	(A) 創建合成數據樣本 (B) 模擬數據分佈 (C) 分類醫學影像 (D) 生成文本	C	分類醫學影像屬於**鑑別式 AI (Discriminative AI)** 的任務，非生成式 AI 的主要應用。
29	哪項是生成式 AI 支援鑑別式 AI 的典型案例？	(A) 模擬交通場景以訓練自動駕駛模型 (B) 使用 CNN 對腫瘤分類 (C) 使用 SVM 分析風險 (D) 創建更好的分類演算法	A	模擬交通場景 (生成合成數據) 是為了提供給自動駕駛模型 (鑑別式模型) 訓練，屬於生成式 AI 支援鑑別式 AI 的案例。
30	關於自然語言處理 (NLP) 核心技術，下列敘述何者不正確？	(A) 語音識別技術將語音轉換為文本 (B) 自然語言生成技術可以生成自然流暢的文本 (C) 語意分析技術理解文本的語意，並主要用於語音識別和機器翻譯 (D) 機器翻譯技術自動翻譯文本	C	語意分析理解文本的語意，應用廣泛，不主要用於語音識別和機器翻譯 (這些技術有各自的核心模塊)。
31	關於「負責任的 AI」，下列敘述何者較為正確？	(A) AI 系統的開發者對 AI 系統的行為負責 (B) AI 系統的使用者對 AI 系統的結果負責 (C) AI 系統本身對其行為負責 (D) 政府對 AI 系統的發展負責	A	「負責任的 AI」的核心原則之一是開發者應對其創建的 AI 系統的行為負責。
32	關於生成式 AI 的基本原理，下列敘述何者較正確？	(A) 生成式 AI 通過分析大量數據來生成新數據，模擬數據分佈以創造與訓練數據相似的結果 (B) 生成式 AI 主要通過預定義的規則來進行數據處理和分類 (C) 生成式 AI 專注於數據分類和迴歸預測 (D) 生成式 AI 通過自動化的方式清洗數據	A	生成式 AI 的基本原理是分析大量數據，學習其潛在分佈，進而生成新的、與原始數據相似的數據。
33	下列模型中，何者並非以「產生新資料」為主要設計目的？	(A) 支援向量機 (Support Vector Machine) (B) 變分自編碼器 (C) 自迴歸模型 (D) 擴散模型	A	支援向量機 (SVM) 是用於分類或迴歸的鑑別式模型，而非生成新資料的模型。
34	某份資料中出現多個欄位儲存相同的成績資訊，造成資料結構重複與使用混淆，此種情形屬於下列哪一種資料品質問題？	(A) 重複資料 (B) 冗餘資料 (C) 格式錯誤資料 (D) 缺失資料	B	**冗餘資料 (Redundant Data)** 指多個欄位重複儲存相同或相似的資訊，導致資料結構不必要地重複。
35	預測商品的退貨機率，模型的輸出為是否退貨 (是/否)。在此模型中，「是否退貨」應歸類為下列哪一類變數？	(A) 特徵 (B) 標籤 (Label / Target) (C) 超參數 (D) 正則化係數	B	「是否退貨」是模型希望預測的結果，因此屬於標籤 (Label) 或目標變數 (Target)。

科目二：生成式 AI 應用與規劃

題號	題目（精簡）	選項	正確解答	解題重點
1	何者最能表達 No Code / Low Code 平台的主要特色？	(A) 需要撰寫大量程式碼 (B) 運用模板快速建立應用程式 (C) 僅供專業開發人員使用 (D) 只能製作靜態網站	B	No Code / Low Code 平台的主要優勢是通過圖形介面或模板，讓使用者運用模板快速建立應用程式。
2	下列何者不是推理模型 (Reasoning Model) 的主要特點？	(A) 具備多步驟邏輯推理能力 (B) 具備良好的可解釋性與邏輯一致性 (C) 採用強化式學習的訓練方式 (D) 回應內容結構清晰、推理脈絡完整	C	推理模型通常透過預訓練、微調或提示工程進行訓練和引導，不主要採用強化式學習的訓練方式。
3	哪種情況，選擇 Low Code 平台可能比 No Code 平台更為適合？	(A) 需要非技術人員快速進行開發與應用 (B) 應用需求簡單，無需自訂功能 (C) 需要較複雜的業務邏輯並使用自訂整合功能 (D) 預算和時間極度有限	C	Low Code 平台透過允許開發者加入少量程式碼，更適合處理較複雜的業務邏輯並使用自訂整合功能。
4	關於生成式 AI 與 No Code / Low Code 平台的應用，下列何者最不適合？	(A) 自動生成程式碼 (B) 自動化生成行銷文案 (C) 快速開發個人化 App (D) 自動化生成法律判決	D	自動化生成法律判決涉及高度倫理、法規與社會責任，是目前生成式 AI 最不適合自動化的應用。
5	關於 No Code / Low Code 平台，下列敘述何者較正確？	(A) 兩者完全相同 (B) Low Code 平台不需要任何程式設計知識 (C) Low Code 平台更適合開發靈活且可擴展的解決方案 (D) No Code 平台可以無限客製化	C	Low Code 平台由於允許客製化程式碼的加入，因此更適合開發靈活且可擴展的解決方案。
6	若希望引導模型逐步推理以提升答案的準確性與可解釋性 (Chain-of-Thought)，應設計哪一組提示語？	(A) 請回答這個數學問題，並直接給出答案 (B) 幫我算出 125×12 是多少？ (C) 請詳細列出每一個思考步驟，最後再給出答案 (D) 請用一句話簡要回答問題	C	提示語要求模型「詳細列出每一個思考步驟」，是啟用思維鏈 (Chain-of-Thought) 推理特性的關鍵。
7	使用 Low-Code 平台進行開發時，企業應特別留意哪一項潛在風險？	(A) 可能造成企業內部敏感資料的洩露 (B) 難以進行大規模的應用擴展和維護 (C) 開發成本將大幅增加 (D) 可能有未經 IT 部門管理的應用程式擴散	D	Low-Code 門檻低，可能導致非 IT 專業人員開發應用，造成未經 IT 部門管理的應用程式擴散 (Shadow IT)。
8	下列哪一種技術方案適用於改善客戶體驗？	(A) 智慧排程系統 (B) 消費行為洞察模型 (C) 預測性維護工具 (D) 自然語言處理 (NLP) 和生成式回應模組	D	自然語言處理 (NLP) 和生成式回應模組用於提升客戶互動品質 (如聊天機器人)，適用於改善客戶體驗。

9	依《個資法》規定，哪一種情況下，中央目的事業主管機關有權對國際資料傳輸行為進行限制？	(A) 公司規模未達主管機關要求標準 (B) 公司曾發生個人資料外洩事件紀錄 (C) 傳輸資料包含當事人已公開資訊 (D) 接收國家之個人資料保護法規尚未完善，可能損害當事人權益	D	依據《個資法》，當接收國家之個人資料保護法規尚未完善，可能損害當事人權益時，主管機關有權限制國際傳輸。
10	下列哪一項技術是生成式 AI 的基礎？	(A) 決策樹模型 (B) 聚類演算法 (C) 生成對抗網路 (D) 隨機森林技術	C	**生成對抗網路 (GAN) **是生成式 AI 的重要基礎技術之一。
11	能使用 DALL·E-2 生成各式逼真的圖片，最關鍵的應用技術為何？	(A) 卷積神經網路 (B) 生成對抗網路 (C) 擴散模型 (Diffusion Model) (D) 自然語言處理	C	DALL·E-2 和 Midjourney 等先進的圖像生成模型，其核心技術是擴散模型 (Diffusion Model)。
12	在企業級數據管道 (ETL) 中，No Code / Low Code 平台的主要角色為何？	(A) 可作為前端數據可視化工具 (B) 取代部分傳統 ETL 解決方案，但可能無法處理過於客製化的邏輯 (C) 僅適用於小型數據應用 (D) 無法應用於任何數據處理場景	B	No Code / Low Code 平台具備一定資料處理能力，可以取代部分傳統 ETL 解決方案，但客製化程度受限。
13	下列何者不是生成式 AI 核心技術？	(A) Variational Autoencoders (VAE) (B) Generative Adversarial Networks (GAN) (C) Visual Geometry Group (VGG) (D) Autoregressive Models (AR Model)	C	VGG 是一種卷積神經網路 (CNN) 架構，主要用於圖像分類，屬於鑑別式模型，不是生成式 AI 的核心技術。
14	使用生成式 AI 技術或工具生成內容時，應採取哪一項措施以確保內容品質？	(A) 使用內容直接進行學術報告 (B) 適當標注引用來源 (C) 減少人工參與的審查過程 (D) 排除所有生成的資料	B	確保生成內容品質和負責任使用，應適當標注引用來源。
15	小莉老師製作個人化學習系統的操作步驟邏輯順序排列？	(A) d → a → c → b (B) a → d → b → c (C) d → a → b → c (D) a → b → c → d	C	合理的生成式 AI 學習系統設計順序：d (建立知識結構) → a (輸入學生特徵) → b (生成練習題與回饋) → c (生成適應性學習計畫)。
16	下列哪一項不是生成式 AI 工具在使用體驗方面的優化方向？	(A) 提供更直觀的操作設計 (B) 支援自然語言指令 (C) 提供智慧化的參數調整建議 (D) 限制使用者自訂生成內容	D	優化方向應是增強使用者的自由度，因此限制使用者自訂生成內容不是優化方向。
17	學校教師如何引導學生正確使用生成式 AI 工具？	(A) 不應使用 AI 工具於教學場域 (B) 無限制地使用 AI 工具 (C) 訂立清晰的使用規範並進行說明 (D) 僅鼓勵學生利用 AI 完成課堂作業	C	教師應訂立清晰的使用規範並進行說明，引導學生負責地使用 AI 工具。
18	企業使用第三方 AI API 處理內部資料，為避免敏感資訊遭記憶或進一步分析，應優先採取哪一種作法？	(A) 啟用 Zero-Retention 模式或採用具企業資安保障的 API 服務 (B) 使用虛擬機建立隔離環境 (C) 將原始數據進行視	A	優先作法是從服務供應商端確保資料不被保留或用於模型訓練，即啟用 Zero-Retention 模式或採用具企業資安保障的 API 服務。

		覺化轉換 (D) 將輸入資料轉換為向量嵌入		
19	企業若要有效支援生成式 AI 的運行，內部 IT 環境最需要具備下列何種條件？	(A) 提供更多的辦公設備 (B) 精簡企業內部流程 (C) 擁有高效能運算資源與彈性儲存空間 (D) 增加部門之間的交流機會	C	生成式 AI 的運行（特別是訓練與推理）對硬體要求高，需要高效能運算資源（如 GPU）與彈性儲存空間。
20	企業希望 AI 理解客戶文本訊息，並根據文本內容調用相對應的圖片和影片進行回覆，應選擇哪一種模型？	(A) 強化學習模型 (B) 多模態模型 (C) 圖像分類模型 (D) 單模態大語言模型	B	處理多種模態的輸入（文本）和輸出（圖片/影片），應選擇多模態模型。
21	哪一種情境最能展現提示工程（Prompt Engineering）的價值？	(A) 使用者輸入一個模糊的問題，AI 給出一個非常確定的答案 (B) 使用者輸入一個非常具體的問題，AI 給出一個相關但不完全符合的答案 (C) 使用者輸入一個精準有架構的問題，AI 生成符合架構的答案 (D) 使用者輸入一個簡單的問題，AI 給出一個非常複雜的答案	C	提示工程的價值在於優化輸入，使得 AI 能夠生成符合精準架構和要求的答案。
22	下列敘述何者最能反映生成式 AI 在圖像生成領域的發展趨勢？	(A) 生成式 AI 目前僅能生成低解析度圖像 (B) 隨著生成式 AI 技術的進步，圖像不僅風格多樣，且逼真度顯著提升，並可處理更複雜的圖像生成任務 (C) 生成式 AI 無法生成動態圖像 (D) 發展主要集中於圖像風格轉換	B	圖像生成領域的趨勢是逼真度顯著提升，且能處理複雜的圖像生成任務。
23	關於生成式 AI，下列 A-E 敘述哪些正確？ (A. 生成內容不會帶有偏見 B. 具有高度準確性 C. 內容準確性建議需經過人類審核 D. 每次生成的內容都可能不同 E. 具有高度安全性)	(A) C、D (B) A、C、D (C) A、B、D、E (D) A、C、D、E	A	生成式 AI 可能帶有偏見 (A 錯誤)、可能產生虛假信息 (B 錯誤)、安全性並非絕對 (E 錯誤)，但準確性需要人類審核 (C 正確) 且生成內容具隨機性 (D 正確)。
24	財務部門導入生成式 AI 結合自動化系統進行報告產出，最有可能實現的效益為何？	(A) 增加財務收益 (B) 顯著提升報告準確性並減少人工錯誤 (C) 加速市場推廣活動 (D) 提高品牌曝光度	B	在報告產出方面，導入 AI 自動化最直接的效益是提升報告的準確性並減少人工錯誤。
25	訓練瑕疵偵測模型時，使用高解析度影像儲存於 NAS，GPU 利用率偏低，最可能原因？	(A) 模型架構過於複雜 (B) 訓練資料品質不穩定 (C) 高解析影像從 NAS 載入產生 I/O (Input/Output) 瓶頸，導致 GPU 等待 (D) 批次大小設定不當	C	當 GPU 使用率遠低於飽和時，常見原因是大容量數據傳輸緩慢，即高解析影像從 NAS 載入產生 I/O 瓶頸。
26	在生成式 AI 導入過程中，資料安全與隱私保護的哪一方面是最重要的考量？	(A) 設定目標優先級 (B) 增強客服回饋（反饋）能力 (C) 資料視覺化能力 (D) 權限控管與合規要求	D	在資料安全與隱私保護方面，最重要的是確保符合法規和企業政策的權限控管與合規要求。

27	若企業將資料安全管理外包給第三方服務供應商，屬於哪種風險應對策略？	(A) 風險緩解 (B) 風險轉移 (C) 風險接受 (D) 風險規避	B	將管理職責和相關風險轉移給外部服務供應商，屬於風險轉移策略。
28	在生成式 AI 的風險管理中，下列哪一項屬於倫理風險？	(A) AI 生成的內容可能帶有偏見或歧視 (B) 系統運行中斷可能導致企業業務受到影響 (C) 因資料需求增加而引起的存儲成本上升 (D) 員工培訓成本增加	A	AI 生成的內容可能帶有偏見或歧視，直接涉及社會公平性，因此屬於倫理風險。
29	在企業導入 AI 的實施/營運階段，為持續發揮導入 AI 的價值，下列步驟的正確排序？ (A. AI 價值擴散 B. 上線部署 C. 模型監控與優化)	(A) ACB (B) ABC (C) BAC (D) BCA	D	實施/營運階段的正確順序為：B. 上線部署 → C. 模型監控與優化 → A. AI 價值擴散。
30	AI 生成的報告中，出現虛構的法規條文與企業名稱，最可能的原因為下列何者？	(A) 模型語料涵蓋不足 (B) 提示語過度簡略 (C) 模型缺乏檢索能力 (D) 模型產生幻覺內容，虛構與真實資訊混雜出現	D	這種聽起來合理但事實錯誤的現象，是大型語言模型 (LLM) 的典型問題，稱為模型產生幻覺內容 (Hallucination)。
31	在管理生成式 AI 系統的隱私風險時，下列哪一種技術最能確保數據使用的安全性？	(A) 強化學習 (B) 深度學習 (C) 零信任架構 (Zero Trust Architecture) (D) 注意力機制	C	**零信任架構 (Zero Trust Architecture) **假設所有內部和外部的實體都不可信任，從而提供最高等級的數據安全保障。
32	在驗證生成式 AI 應用的概念驗證 (POC) 時，若企業希望確保模型生成的公平性，最適合採用哪種評估策略？	(A) 壓力測試 (B) 對抗性測試 (C) 偏差檢測 (Bias Detection) (D) 延遲測試	C	為了確保模型生成的公平性，最適合的評估策略是進行偏差檢測 (Bias Detection)。
33	產出需要邏輯清晰、高度準確性與一致性的文本時，哪一種溫度 (Temperature) 設定最能有助於降低回答的不確定性？	(A) 將溫度設定在 0.1 至 0.3，使輸出結果趨於穩定可預測 (B) 將溫度設定在 0.5 至 0.7 (C) 將溫度設定在 1.0 至 1.2 (D) 維持預設溫度不變	A	**較低的溫度設定 (例如 0.1 至 0.3) **會使模型輸出趨向於最有可能的結果，降低隨機性，從而提高答案的一致性與準確性。
34	專門用於評估大型語言模型 (LLM) 在「台灣本土特有知識」方面表現的基準測試資料集？	(A) HumanEval (B) Table Understanding (C) MBPP (D) TTQA	D	TTQA (Taiwanese Text Question Answering) 是專門針對台灣本土特有知識設計的基準測試資料集。
35	某中型零售企業打算導入 AI 客服系統，哪一個步驟最適合作為接下來的第一階段任務？	(A) 設計符合品牌風格的 AI 回應邏輯與對話範本 (B) 選擇合適的 AI 技術供應商或開源方案 (C) 蒐集過往客服對話記錄，進行資料清洗與前處理 (D) 明確定義導入目標與關鍵績效指標	D	任何 AI 導入流程的第一步都是明確定義導入目標與關鍵績效指標 (KPI)，作為後續設計、評估和依據。

▼9 月中級樣題

科目一：人工智慧技術應用與規劃

題號	題目（精簡）	選項	正確解答	解題重點
1	何者未用 AI/ML 技術？	(A) 預定義規則的象棋遊戲 (B) 深度神經網路的語音識別 (C) 預定義導航規則的自動駕駛汽車 (D) NLP 的聊天機器人	A	人工智慧 (AI) 或機器學習 (ML) 技術涉及數據學習或推論，而使用預定義規則的系統不屬於 AI/ML 技術。
2	將文本轉換為詞彙單位的處理方法？	(A) 詞形還原 (B) 停用詞移除 (C) 斷詞 (D) TF-IDF	C	**斷詞 (Tokenization)** 是在文本資料處理中，將連續的文本轉換為詞彙單位的過程。
3	NLP 在 ML 應用中的主要用途？	(A) 情緒分析 (B) 圖像識別 (C) 預測性維護 (D) 供應鏈優化	A	情緒分析是自然語言處理 (NLP) 在機器學習應用中的一個主要且典型的用途。
4	關於深度學習模型，何者不正確？	(A) CNN 適合影像辨識 (B) ReLU 減緩梯度爆炸與消失 (C) RNN 適合序列資料 (D) Elman 神經網路適合影像辨識	D	Elman 神經網路屬於遞迴神經網路 (RNN) 的變體，主要用於處理序列相關資料，不適合影像辨識。
5	機器學習模型在業界部署的主要趨勢？	(A) 採用 AutoML 技術 (B) 轉向更簡單算法 (C) 雲平台使用率下降 (D) 依賴手動超參數調整	A	業界部署趨勢是越來越多地採用自動化機器學習 (AutoML) 技術，以提高效率。
6	何者「最適合」使用迴歸模型進行預測？	(A) 預測購物車商品類別 (B) 判斷疾病類別 (C) 推薦電影類型 (D) 預測明年某產品的總銷售額	D	迴歸 (Regression) 模型用於預測連續的數值，總銷售額是連續數值預測。
7	電腦視覺辨識技術配對圖中 (a)~(d)？	(A) (a)實例分割，(b)語義分割，(c)物件偵測，(d)圖像分類 (B) (a)圖像分類，(b)語義分割，(c)物件偵測，(d)實例分割 (C) (a)圖像分類，(b)物件偵測，(c)語義分割，(d)實例分割 (D) (a)圖像分類，(b)物件偵測，(c)實例分割，(d)語義分割	B	正確配對為：(a)圖像分類、(b)語義分割、(c)物件偵測、(d)實例分割。
8	企業導入生成式 AI，可選擇何種模型壓縮技術減少記憶體？	(A) 參數剪枝 (B) 增加訓練數據量 (C) 增加模型層數 (D) 使用更高維數據	A	**參數剪枝 (Parameter Pruning)** 是一種模型壓縮技術，透過移除不重要的權重來減少模型大小和記憶體使用。

9	供應鏈攻擊如何影響 AI 系統安全？	(A) 員工因釣魚郵件被盜帳號 (B) 第三方開源模型或數據被植入惡意內容 (C) 只會受到 DDoS 影響 (D) 僅影響硬體層級	B	供應鏈攻擊的風險在於 AI 系統依賴的第三方開源模型或數據可能被植入惡意內容。
10	NLP 中將文字轉換為向量的詞嵌入技術？	(A) TF-IDF (B) Word2Vec (C) Stop Words (D) Bag-of-Words	B	Word2Vec 是一種詞嵌入 (Word Embedding) 技術，能將文字轉換為向量，以利機器學習處理。
11	多模態學習中，早期融合 (Early Fusion) 的主要特徵？	(A) 輸出結果進行決策合併 (B) 在模型輸入或特徵提取階段整合 (C) 僅處理單一來源 (D) 利用注意力機制在深層隱藏層融合	B	早期融合方法的主要特徵是在模型輸入階段或特徵提取階段整合不同模態資料。
12	哪項技術最有助於強化醫療多模態 AI 系統在處理影像與文本數據時的整合能力？	(A) 利用預先定義的規則 (B) 僅使用 CNN 架構 (C) 利用單一模態資料 (D) 採用 Transformer 架構	D	採用 Transformer 架構因其強大的序列處理和注意力機制，有助於整合醫療影像與臨床文本數據。
13	線上音樂平台希望將用戶劃分為不同類型 (事前未定義)，哪種模型最適合？	(A) 邏輯迴歸 (B) 決策樹 (C) DBSCAN (D) 線性迴歸	C	基於密度之含噪空間聚類法 (DBSCAN) 是非監督式學習的聚類模型，適用於事前沒有定義類型的分群任務。
14	特徵尺度差異極大時，最能有效解決此問題的預處理方式？	(A) 移除尺度較小的欄位 (B) 對所有特徵進行 Z-score 標準化 (C) 進行 Min-Max 正規化 (D) 加上常數	B	**Z-score 標準化 (Standardization) **能將數據轉換為平均值 0、標準差 1 的分佈，有效處理尺度差異極大的問題。
15	AI 模型部署至營運系統，進入系統整合測試階段，哪項驗證最應優先執行？	(A) 確認模型在開發機上訓練集表現是否達標 (B) 驗證模型服務與資料平台、前後端介面協同是否正常 (C) 檢查程式碼提交規範 (D) 審閱模型報告與 API 文件的格式	B	在系統整合測試階段，最優先的是驗證模型服務與資料平台、前後端介面的協同運作是否正常和資料格式流程的一致性。

科目二：大數據處理分析與應用

題號	題目（精簡）	選項	正確解答	解題重點
1	關於常態分佈身高檢測，何者不正確？	(A) 檢測一年級和二年級平均差異用 t 檢定 (B) 檢測一年級、二年級、三年級平均差異用 t 檢定 (C) 檢測二年級、三年級、四年級平均差異用 F 檢定 (D) 檢測一年級平均是否等於 170 公分，可以利用卡方檢定	D	卡方檢定主要用於類別資料的獨立性檢定，不適合用於檢測單一數值平均數是否等於特定數值（應使用 Z 檢定或 t 檢定）。
2	關於接受者操作特徵（ROC）曲線，何者正確？	(A) 繪製了真陽性率與假陽性率的關係 (B) 用於評估模型的準確性 (C) 曲線下的面積（AUC-ROC）始終等於 1 (D) 只適用於二元分類問題	A	ROC 曲線繪製的是模型在不同分類閾值下，真陽性率（True Positive Rate）與假陽性率（False Positive Rate）之間的關係。
3	何者不屬於特徵工程（Feature Engineering）？	(A) 轉換 (B) 萃取 (C) 挑選 (D) 預測	D	特徵工程包括轉換（Transformation）、萃取（Extraction）和挑選（Selection），但不包括「預測」。
4	預測下一季產品銷售數量，適合哪個類型的模型？	(A) 決策樹分類器 (B) K-means 分群 (C) 主成份分析 (D) 線性迴歸	D	產品銷售數量是連續變數，應使用迴歸模型，其中線性迴歸適合此類預測目標。
5	對於低結構化的文本或圖像資料，哪一種特徵工程方法最為適用？	(A) 特徵改善 (B) 特徵建構 (C) 特徵學習 (D) 特徵選擇	C	對於低結構化的複雜資料（如圖像、文本），通常採用深度學習模型進行特徵學習（Feature Learning）。
6	哪種方法屬於非監督式學習中的降維技術？	(A) K-均值聚類 (B) 隨機森林 (C) 支持向量機 (D) 主成分分析（PCA）	D	主成分分析（PCA）是常見的非監督式學習降維技術。
7	IQR 法異常值偵測，哪個範圍外的數值會被視為異常？	(A) 超過 $Q2 \pm 2 \times IQR$ (B) 高於 $Q1 - IQR$ 或低於 $Q3 + IQR$ (C) 高於 $Q1 - 2 \times IQR$ 或低於 $Q3 + 2 \times IQR$ (D) 低於 $Q1 - 1.5 \times IQR$ 或高於 $Q3 + 1.5 \times IQR$	D	IQR 法中，離群值（異常值）定義為低於 $Q1 - 1.5 \times IQR$ 或高於 $Q3 + 1.5 \times IQR$ 的數值。
8	差分隱私（Differential Privacy）在數據應用中的目的？	(A) 添加隨機噪聲掩蓋個體訊息 (B) 加密所有數據 (C) 移除多數類樣本 (D) 提高數據完整性	A	差分隱私的目的在於添加隨機噪聲來掩蓋個體訊息，從而保護數據隱私。

9	PCA 將資料降維至兩個主成分表示哪種情況？	(A) 僅保留兩筆樣本資料 (B) 將原始高維資料轉換為兩個主成分表示 (C) 每筆資料僅保留兩個原始欄位的數值 (D) 模型將限制只能用於分類任務	B	這表示將原始高維資料轉換為由兩個主成分表示的新資料空間。
10	常態分配：平均數 70，標準差 10，成績 90 的 Z 分數？	(A) 0 (B) 1 (C) 2 (D) 3	C	Z 分數計算為 $(X - \mu) / \sigma$ ；即 $(90 - 70) / 10 = 2$ 。
11	在 AI 倫理治理的背景 下，「透明性」是指什麼？	(A) AI 系統的運算速度公開 (B) AI 系統僅使用公開數據 (C) AI 系統決策流程清晰且可解釋 (D) 所有 AI 模型都是開源的	C	「透明性」通常是指 AI 系統的決策流程清晰且可解釋。
12	企業希望即時監控交易異常，應選擇哪一類數據處理架構？	(A) 批次式資料倉儲 (B) 離線式資料湖 (C) 串流處理系統 (D) 冷資料備援架構	C	處理即時監控的需求，應選擇具備低延遲處理能力的串流處理系統。
13	分析社群網路使用者之間的互動結構，應使用哪種分析方法？	(A) 文字探勘 (B) 主成分分析 (C) 圖論分析 (D) 分群分析	C	分析社群網路中的節點（使用者）與邊（互動）的結構，應使用圖論分析（Graph Theory Analysis）。
14	為了加速大數據環境下的 AI 模型訓練，以下哪一項為常見技術？	(A) 早期停止 (B) 批次分群 (C) 混合精度訓練 (D) 主成分分析	C	**混合精度訓練（Mixed Precision Training）**是利用低精度浮點數來加速大型模型訓練的常見技術。
15	訓練集(x_train)資料集個數為 10000 筆？ (CIFAR-10 資料集)	(A) 訓練集(x_train)資料集個數為 100000 筆 (B) 測試集(x_test)資料集個數為 10000 筆 (C) 訓練集(x_train)是 Pandas 資料框物件 (D) 像素值轉換為的程式碼	B	CIFAR-10 數據集通常包含 50,000 筆訓練集和 10,000 筆測試集。

科目三：機器學習技術與應用

題號	題目（精簡）	選項	正確解答	解題重點
1	MapReduce 計算框架中，Map 和 Reduce 所負責處理資料的問題？	(A) Map：一組資料映射成另一組資料；Reduce：統合與歸納資料 (B) Map：地圖式的搜索資料；Reduce：統合與歸納資料 (C) Map：一組資料映射成另一組資料；Reduce：過濾不符合的資料 (D) Map：一組資料映射成另一組資料；Reduce：生成更多的資料	A	在 MapReduce 中，Map 階段負責將一組資料轉換（映射）成另一組資料，而 Reduce 階段負責對這些資料進行統合與歸納。

2	何種 CNN 是將卷積層加寬而非加深？	(A) R-CNN (B) Inception (C) ResNet (D) VGG19	B	Inception 架構（例如 GoogLeNet）的特點是使用多種卷積核並行處理，達到加寬卷積層的效果。
3	模型的訓練誤差低、但測試誤差很大，是什麼情況？	(A) 模型的泛化能力強 (B) 模型出現過度擬合 (Overfitting) (C) 模型出現欠擬合 (Underfitting) (D) 訓練資料和測試資料之間沒有相關性	B	訓練集表現優秀，但在未見過的測試集上表現差，表示模型對訓練數據學習得太過細節，發生過度擬合 (Overfitting)。
4	哪一種指標通常用於評估迴歸模型的效能？	(A) R^2 (B) F1-分數 (C) AUC (D) Precision	A	R^2 （決定係數）是用於評估迴歸模型效能的常用指標。
5	深度學習研究與應用中，數據可能存在的潛在問題？	(A) 數據標註品質鮮少被討論，但它卻直接影響模型性能 (B) 數據品質是完美可信賴的 (C) 大部分情況下，數據不存在類別不平衡問題 (D) 數據不需要領域知識的輔助	A	數據的標註品質是影響模型性能的關鍵因素，但常被忽視。
6	在分類任務中，深度學習模型通常搭配哪一種輸出函數？	(A) Tanh (B) ReLU (C) Sigmoid 或 Softmax (D) Dropout	C	分類任務通常使用 Sigmoid（二元分類）或 Softmax（多元分類）作為輸出層的激活函數。
7	哪一種學習任務不適合使用監督式學習方法處理？	(A) 客戶信用風險分類 (B) 預測未來銷售額 (C) 找出資料中的潛在群集 (D) 判斷影像是否為貓或狗	C	找出資料中的潛在群集（聚類）屬於非監督式學習任務，因為事先沒有標籤。
8	神經網路中，前向傳播主要依賴哪一種數學操作？	(A) 機率積分 (B) 矩陣乘法與向量內積 (C) 對數變換 (D) 條件機率推論	B	前向傳播過程主要涉及矩陣乘法與向量內積的線性代數運算。
9	特徵縮放中，標準化 (Standardization) 的主要作用？	(A) 將數據範圍限制在 0 到 1 且標準差-1 (B) 使數據平均值為 0 且標準差為 1 (C) 移除數據中的異常值 (D) 增加特徵間的相關性	B	標準化 (Z-score normalization) 的主要作用是使數據平均值為 0 且標準差為 1。
10	關於準確率 (Accuracy) 的計算方式，下列何者正確？	(A) $TP / (TP+FP)$ (B) $(TP+TN) / (TP+FP+TN+FN)$ (C) $TN / (TN+FP)$ (D) $TP / (TP+FN)$	B	準確率是所有正確預測 (TP, True Positive + TN, True Negative) 佔總樣本數的比率。
11	關於損失函數 (Loss Function) 的主要功能？	(A) 記錄模型預測歷史 (B) 計算模型結構複雜度 (C) 控制模型的學習率 (D) 衡量模型預測與真實值之間的誤差	D	損失函數的主要功能是衡量模型預測與真實值之間的誤差。
12	歐盟 GDPR 中，被遺忘權 (Right to be Forgotten) 賦予資料主體哪項權利？	(A) 要求平台永久備份其個資 (B) 在符合條件下請求刪除其個人資料 (C) 將個人資料轉換為匿名格式保存 (D) 限制企業將資料輸出至境外伺服器	B	被遺忘權主要賦予資料主體在符合條件下請求刪除其個人資料的權利。

13	調整資料中不同群體（例如分類標籤）之間樣本數量不平衡的情況，屬於哪一種技術？	(A) 資料重抽樣 (B) 特徵選擇 (C) 模型正則化 (D) 超參數調整	A	調整樣本比例以平衡類別的作法屬於**資料重抽樣 (Data Resampling)**技術。
14	優化器中，哪個方法會自動調整每個參數的學習率，特別適用於稀疏資料？	(A) Momentum (B) Adagrad (C) Adam (D) SGD	B	Adagrad 能夠自動調整每個參數的學習率，並累積歷史梯度信息，對稀疏資料特別有效。
15	程式碼片段依正確順序排序，以完成模型的建立與預測？	(A) $c \rightarrow a \rightarrow b \rightarrow d$ (B) $a \rightarrow c \rightarrow b \rightarrow d$ (C) $c \rightarrow b \rightarrow a \rightarrow d$ (D) $b \rightarrow a \rightarrow c \rightarrow d$	A	機器學習模型建立與預測的標準流程是：資料處理 (c) $\rightarrow$ 模型訓練 (a) $\rightarrow$ 模型評估/預測 (b) $\rightarrow$ 結果應用 (d)。