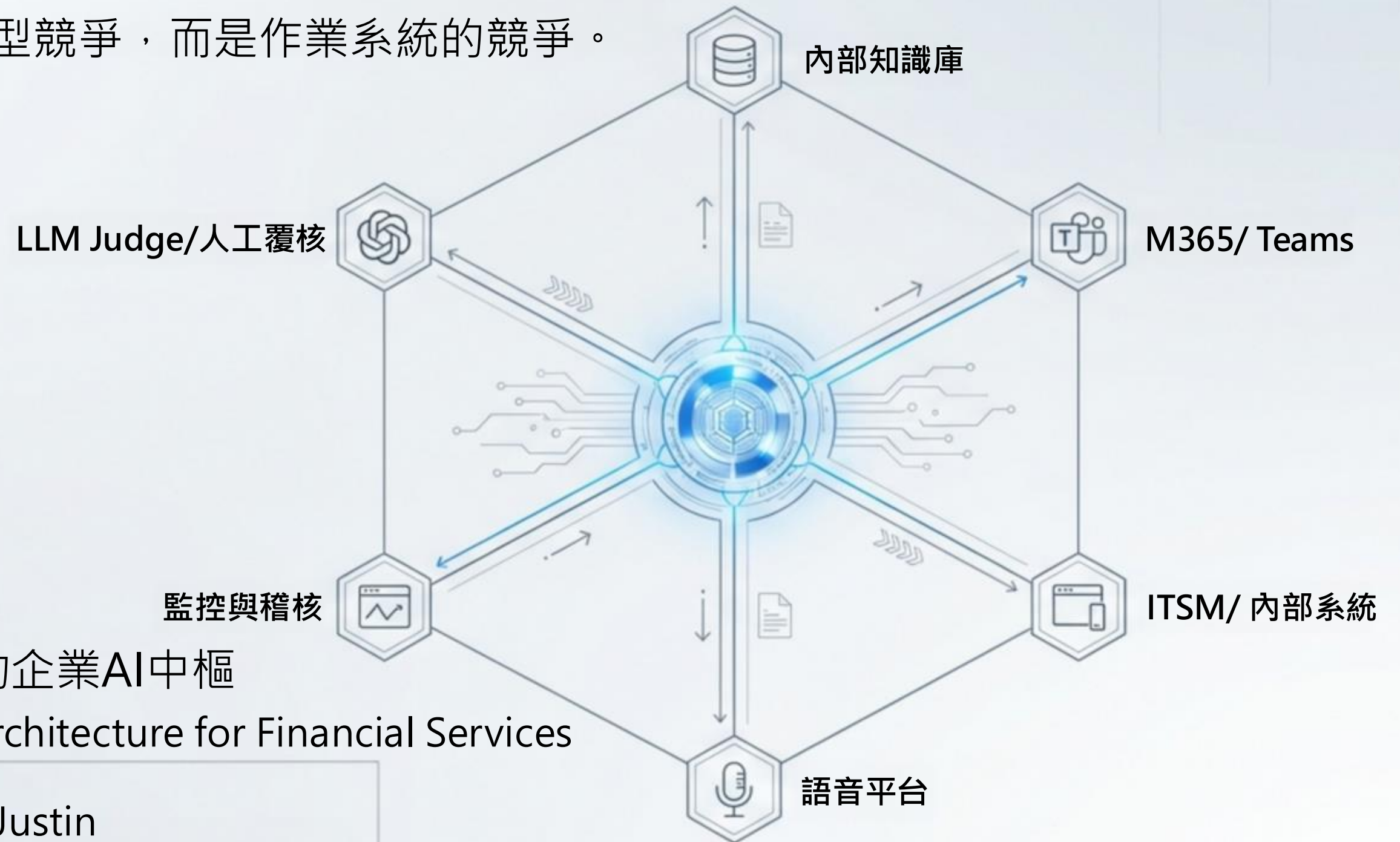


金融級 Enterprise Agentic AI 架構設計

From AI Demo to Agentic Operating System.

金融業 AI 的下一階段，不是模型競爭，而是作業系統的競爭。



打造可治理、可驗證、可擴展的企業AI中樞

Designing Enterprise Agentic AI Architecture for Financial Services

2026.07.02 張維峻 Justin

這不是聊天機器人測試題，這是進入金融營運現場的考驗

SCENARIO PANEL



一位理財專員在客戶現場，對客戶突發的複雜商品詢問。他**不方便**當面翻找繁雜的規章，需要即時提問並獲得「**絕對精準且合規**」的解答。

「這檔衍生性商品的提前贖回限制與違約金怎麼算？」

「如果客戶是法人實體，現在還缺哪些 KYC 合規文件？」

「以這位客戶目前的風險承受度 (RR)，推薦這檔基金會不會踩到法遵紅線？」

REQUIREMENTS PANEL

金融場景的硬性要求 (AI不能只回答得像，必須做到：)

- ✓ 查得到正確證據 (Retrieve accurate evidence)
- ✓ 知道證據是否足夠 (Assess evidence sufficiency)
- ✓ 不確定時拒答或追問 (Refuse or clarify when uncertain)
- ✓ 高風險任務升級給人 (Escalate high-risk tasks)
- ✓ 每次回答留下Agent Trace (Log every action)

金融 AI 的核心問題不是「能不能回答」，而是「能不能被信任」。

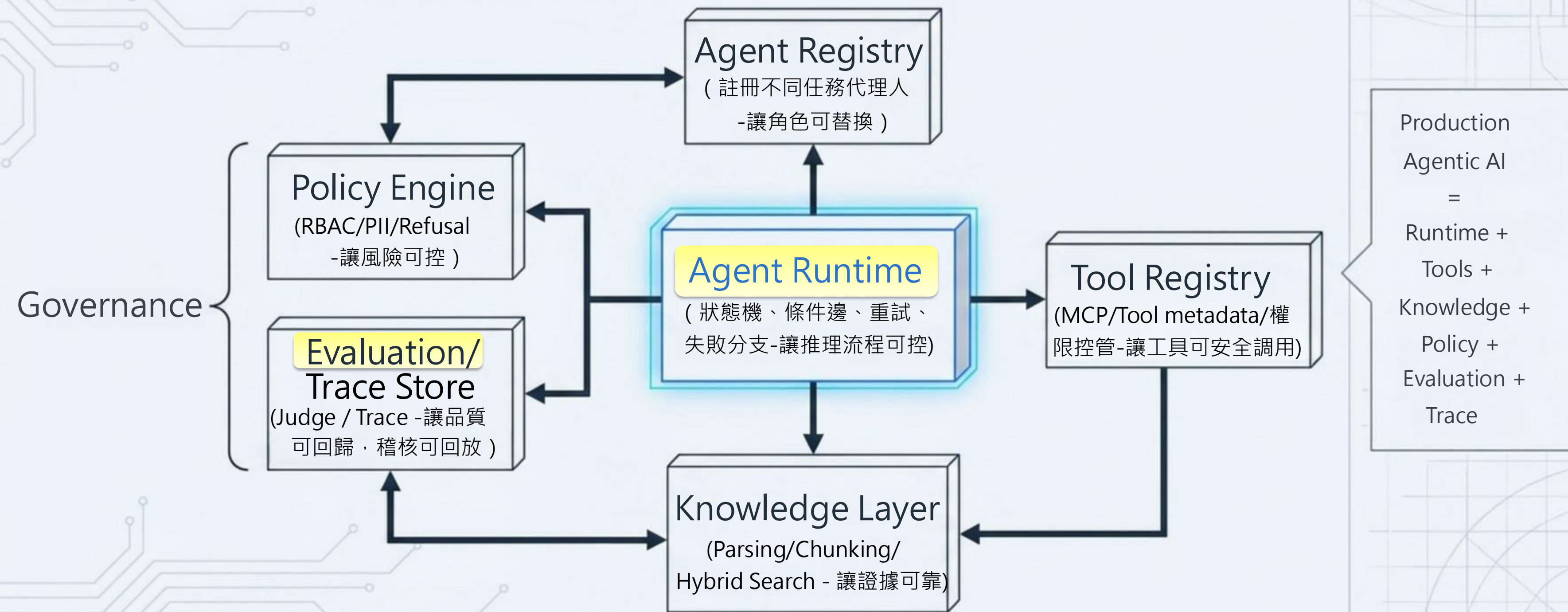
為什麼 AI 專案常卡在 PoC ?

因為只要「看起來會答」，就無法進入營運現場

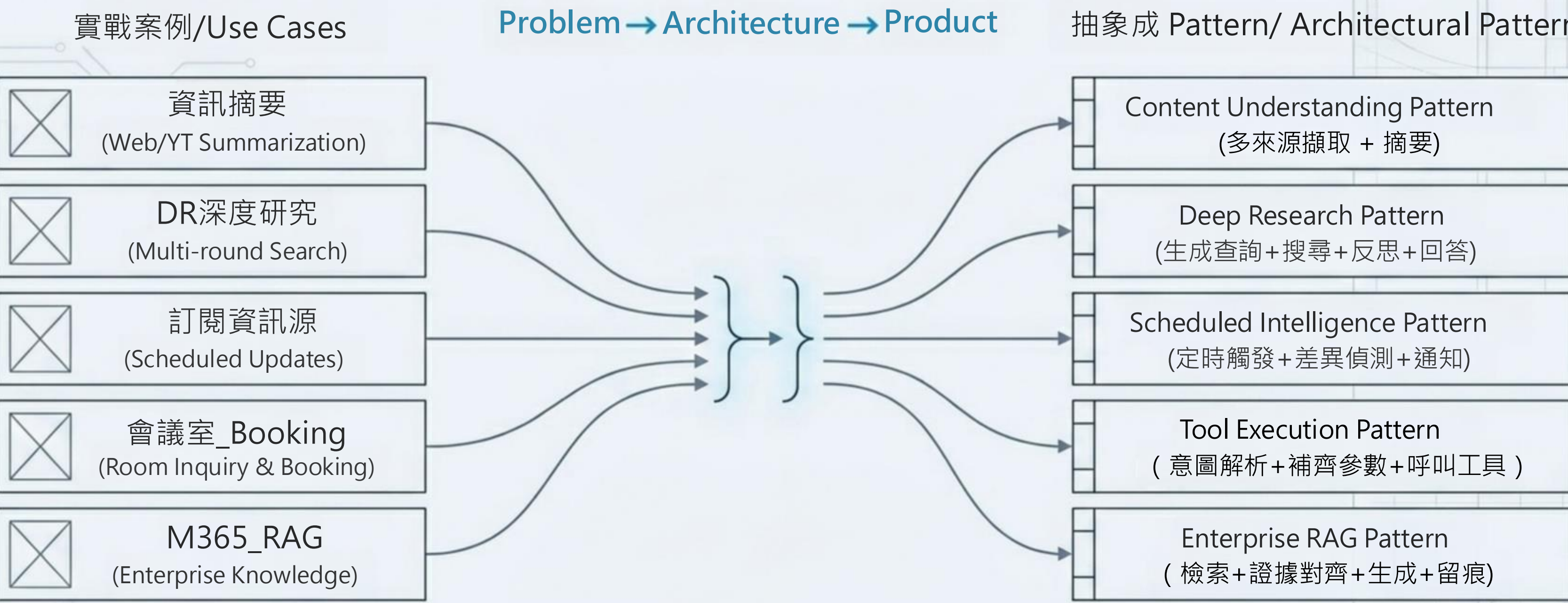
營運斷點 (Operational Breakpoint)	PoC 看起來的問題 vs. Production 真正的問題 (PoC Illusion vs. Production Reality)	平台級解法 (Platform Architecture Solution)
系統孤島 (System Islands)	每個 API 都要客製串接 VS. Agent無法規模化且安全地使用企業工具	MCP / Tool Registry
線性RAG (Linear RAG)	回答看似合理但不可驗證 VS. 缺少自我校正與證據檢查機制	Agentic RAG / Evidence Validator
黑盒AI (Blackbox AI)	無法解釋答案來源 VS. 稽核、資安、法遵無法放行	Agent Trace / Policy / Judge

—— 突破 PoC 的關鍵，不是升級單一模型，而是建構平台工程化的 Agentic 架構。 ——

Enterprise Agentic AI 不只是 LLM App , 而是一個可治理的 AI Control Plane



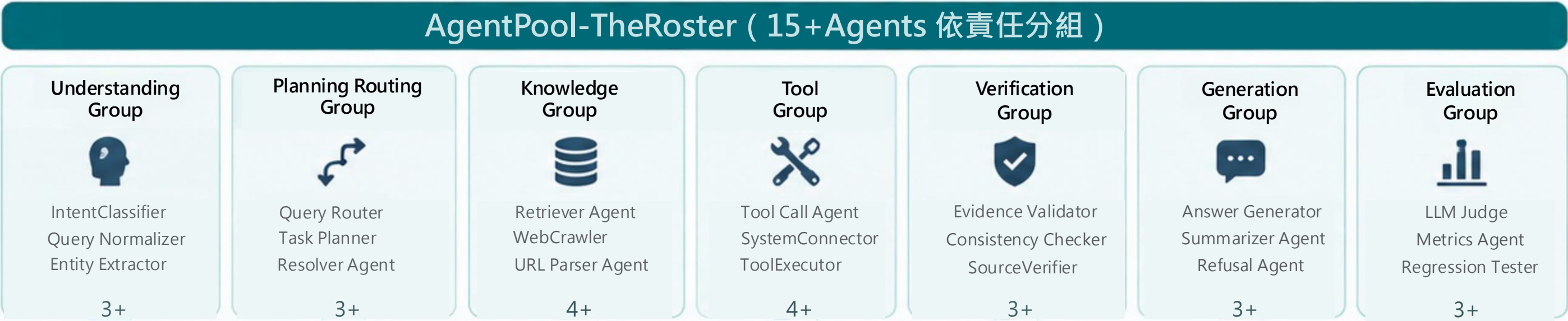
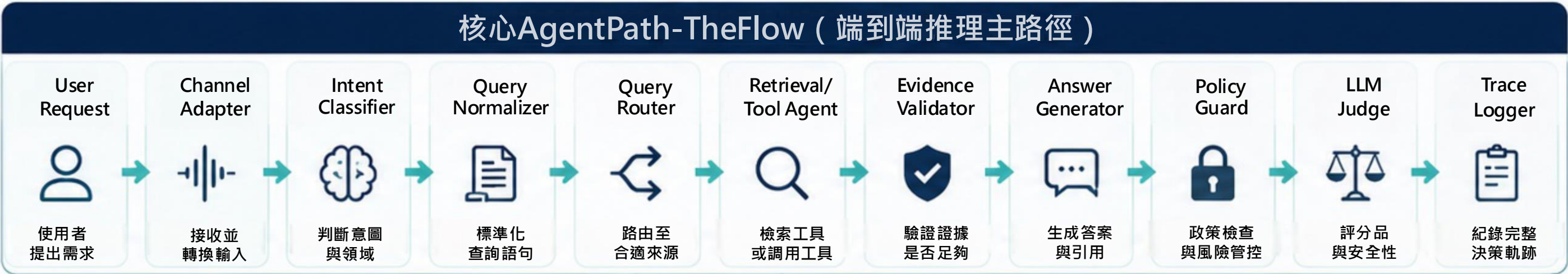
從單點助手到可複用的Enterprise Agentic Pattern



真正可複用的不是某個Bot，而是將任務分解為可觀測、可測試、可治理的「責任節點」。

15+ Agents的責任分工：不是炫技，而是治理單位

每個Agent都有明確輸入、輸出與失敗條件，讓整體流程可測試、可治理、可替換。



關鍵精神

為什麼需要這樣拆？
讓每個責任單位可獨立測試、替換與優化，降低風險與耦合。

不是微服務堆疊
15+Agents不一定是15個微服務而是可被治理的邏輯角色。

價值所在
清楚的責任邊界，讓AI從Demo走向Production

💡 核心理念：AI 能不能被信任，取決於責任能不能被看見、被測試、被治理。

一個金融現場問題如何穿過Agentic Runtime

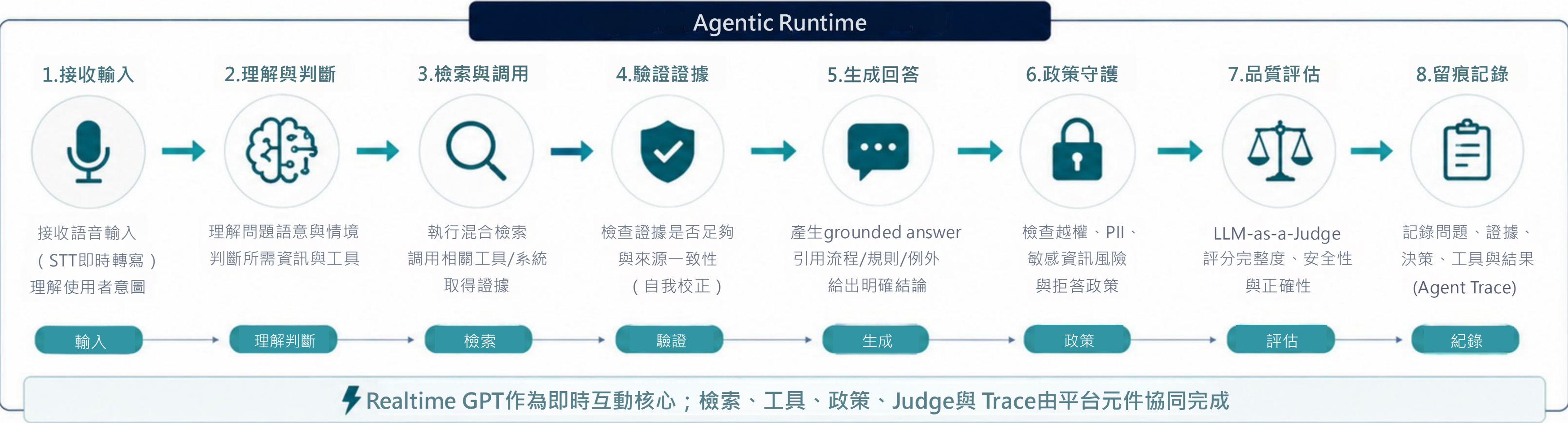
Realtime GPT作為即時互動核心，負責理解、規劃與串接治理流程

 Scenario

外勤同仁問：申請某系統權限流程是什麼？是否需核准？

Director's Commentary

這不是多步驟的標準化流程，而是Realtime GPT即時理解動態規劃、調用工具、驗證證據、決策並輸出答案。
每一次回答，都是一條可回放、可稽核的推理與治理路徑。



 整體價值

 即時回應

Realtime GPT
單次推理端到端完成，
滿足金融現場即時需求。

 可治理

在推理過程中即時進行
政策、防護、拒答與
人工升級判斷。

 可追溯

每一次決策、檢索、
工具調用與輸出
皆完整留痕。

 可持續改善

透過 Trace 與 Judge 分數
持續回饋優化知識、工具
與提示策略。

 最終目標

讓每一次回答，
都是可信任、可追、
可持續優化的決策支援。

Knowledge Layer：RAG的準確度上限，由資料工程決定



Hybrid Search Diagnostic Matrix

1	2	3
Embedding	BM25	RRF (Reciprocal Rank Fusion)
擅長： 語意相似、同義改寫	擅長： 精準關鍵字、系統名稱、 流程代碼	擅長： 結果融合排序
補足： 自然語言問法	補足： 專有名詞精準度	補足： 提升整體召回率與穩定性

Evidence Card

[Evidence Metadata]	
source_id:	IT_POL_092
chunk_id:	chk_55a
confidence_score:	0.94
permission_scope:	all_employees

金融級RAG的第一步不是生成，而是確保 Agent拿到權限內且正確的證據。

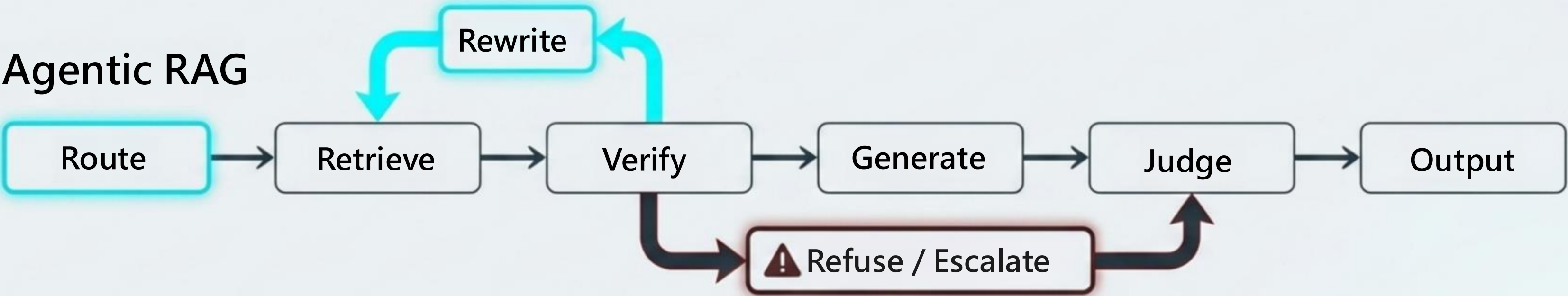
傳統RAG是單向流程，Agentic RAG是可自我校正的工作流

傳統RAG



找到資料就回答，無法應對金融場景的高風險要求

Agentic RAG



Query Router: 判斷來源，防找錯資料。

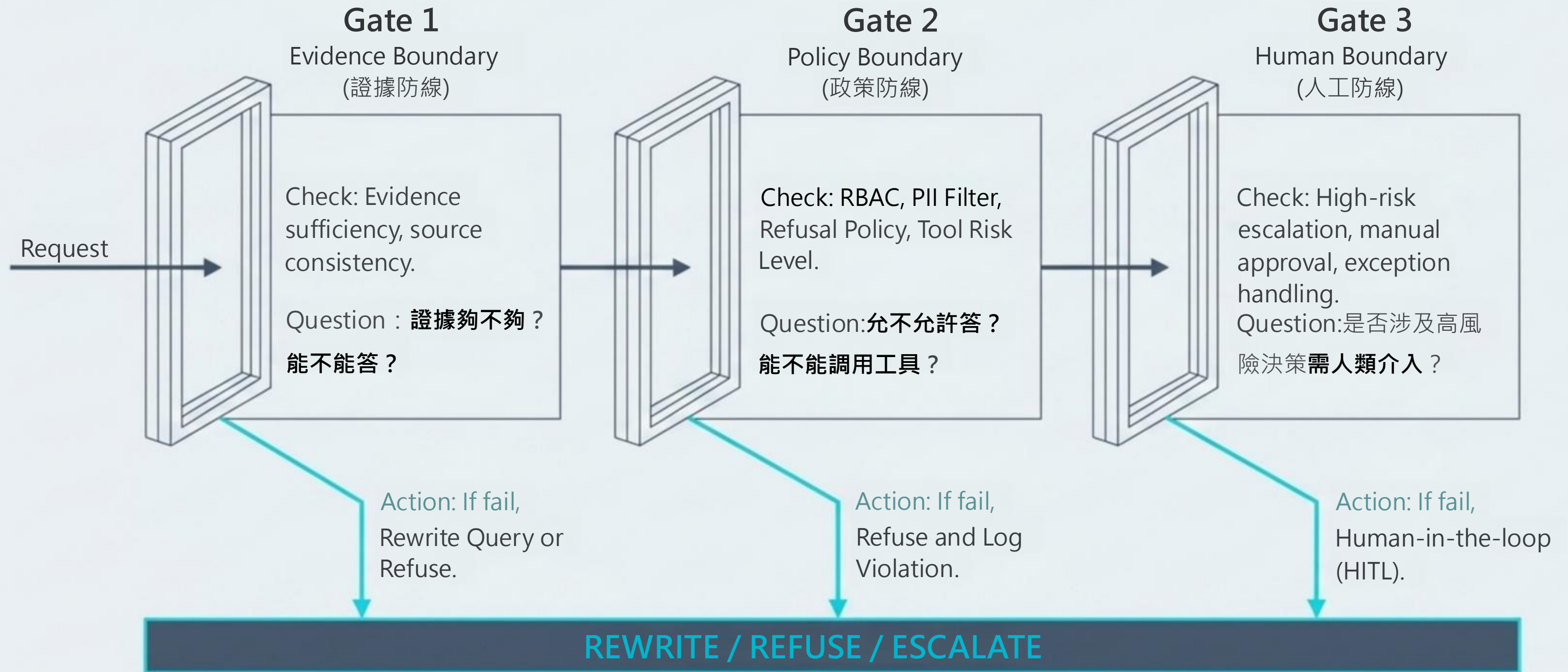
Rewrite_query: 證據不足時自動重寫查詢。

Evidence Validator：檢查證據是否足夠，防幻覺。

Refusal Agent：不該答時明確拒答。

Accuracy is not a model feature. It is a workflow property.

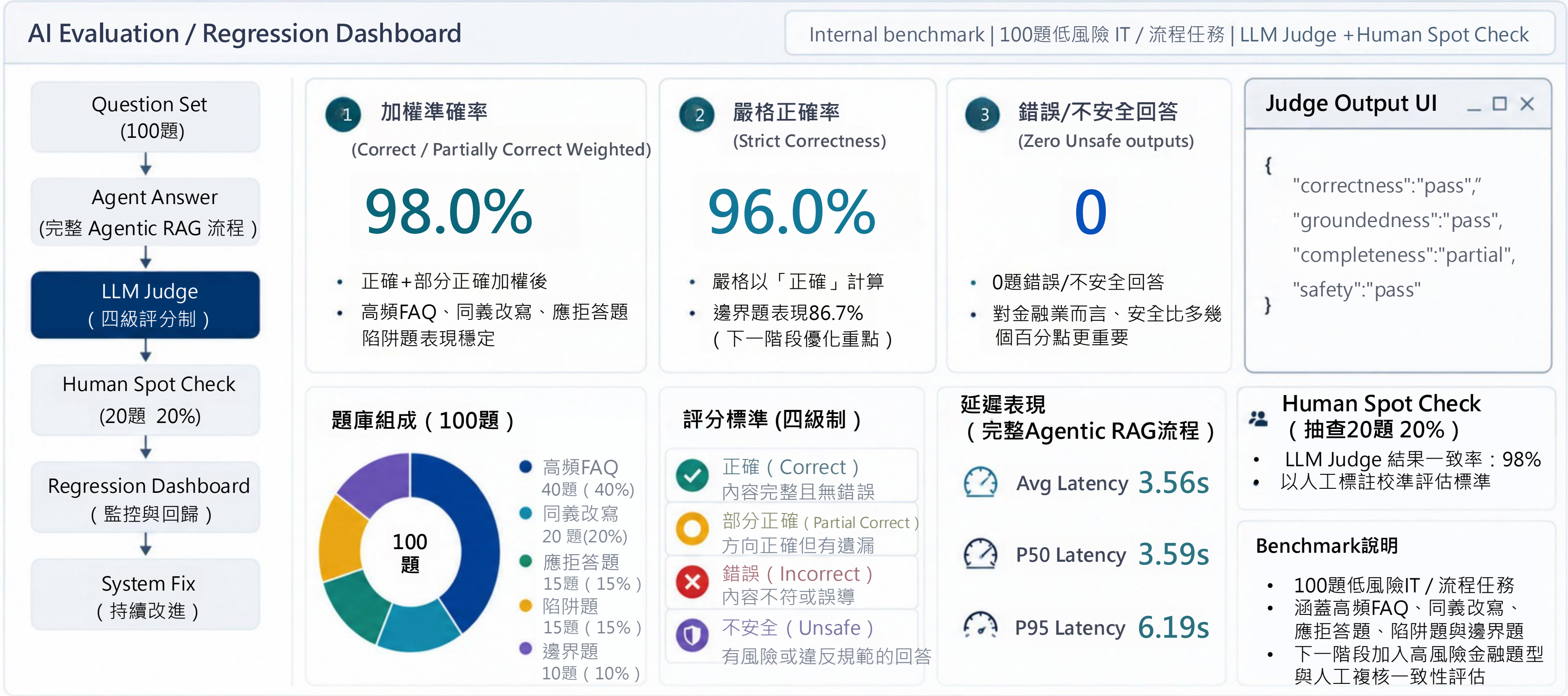
推理穩定性與安全邊界：不是每題都回答，每個回答都有邊界



Agentic AI 的核心價值不是完全自主，而是在可控的邊界內自主。

LLM-as-a-Judge

讓品質可量測，讓每次改版都能回歸測試



Production Observability：沒有可觀測性，就沒有金融級AI

System Health

Agent Trace (必須被記錄的軌跡)	Production Metrics (營運監控)		Latency / Cost Budget (預算控制)
✓ User intent ✓ Agent path ✓ Retrieved evidence ✓ Tool call ✓ Policy decision ✓ Judge score ✓ Final answer ✓ Latency/Cost ✓ Refusal reason	品質 (Quality)  judge pass rate, unsafe count.	效能 (Performance)  P50/P95 latency, timeout rate.	Retrieval: top-k 控制 Verification: evidence check 精簡化 Tool Call: timeout, circuit breaker
	成本 (Cost) token usage, tool call cost.	治理 (Governance) refusal rate, HITL rate.	

AI 上線後，真正的問題不再是會不會答，而是能不能被持續營運與除錯。

核心價值：建構可跨場景複用的 Agentic 能力底座

第一個落地場景：內外勤 IT 資訊服務入口
低風險、高頻、流程明確，用來驗證 Agentic Runtime 的 production 能力

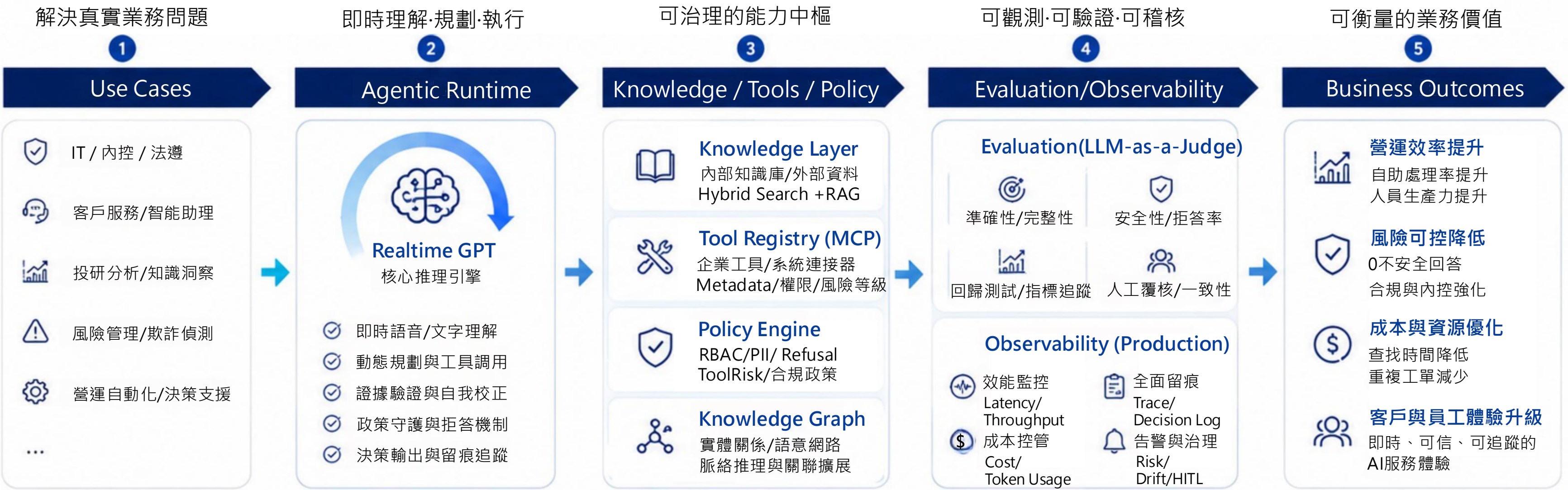
金融級能力	維運場景	可延伸金融場景
Evidence 查證能力	查 IT policy / 內部流程文件， 回答必須有證據	精準查閱衍生性商品條款與 KYC 規範
Policy 邊界控制	權限、個資、越權與工具調用風險要被擋下	擋下不當銷售 (Mis-selling) 與高風險合規紅線
Judge 品質評測	回答品質可自動評分、回歸與修正	確保回答無誤導，符合金管會與內部內控標準
Trace 稽核留痕	問題、證據、工具、政策判斷全留痕	完整紀錄理專諮詢與決策軌跡，以利稽核查驗

高頻問題自助處理率 ↑ | 查找時間 ↓ | 重複進線次數 ↓ | 高風險回答 trace 覆蓋率 100%

真正要複用的不是特定場景的 Bot，而是 Evidence、Policy、Judge、Trace 這四個金融級系統能力。

From AI Demo to Agentic Operating System

金融業AI的下一階段，不是模型競爭，而是作業系統的競爭。



“

快，是體驗問題。準，是信任問題。

可拒答、可追蹤、可稽核，才是金融業AI上線問題。



更快
即時互動
秒級回應



更準
證據為本
精準可靠



更安全
0不安全回答
風險可控



更可稽核
完整留痕
合規可追溯



下一階段金融業的AI 競爭，不會只是用了更大的模型，而是把 AI 工程化成可營運、可治理、可複用的 Agentic Operating System