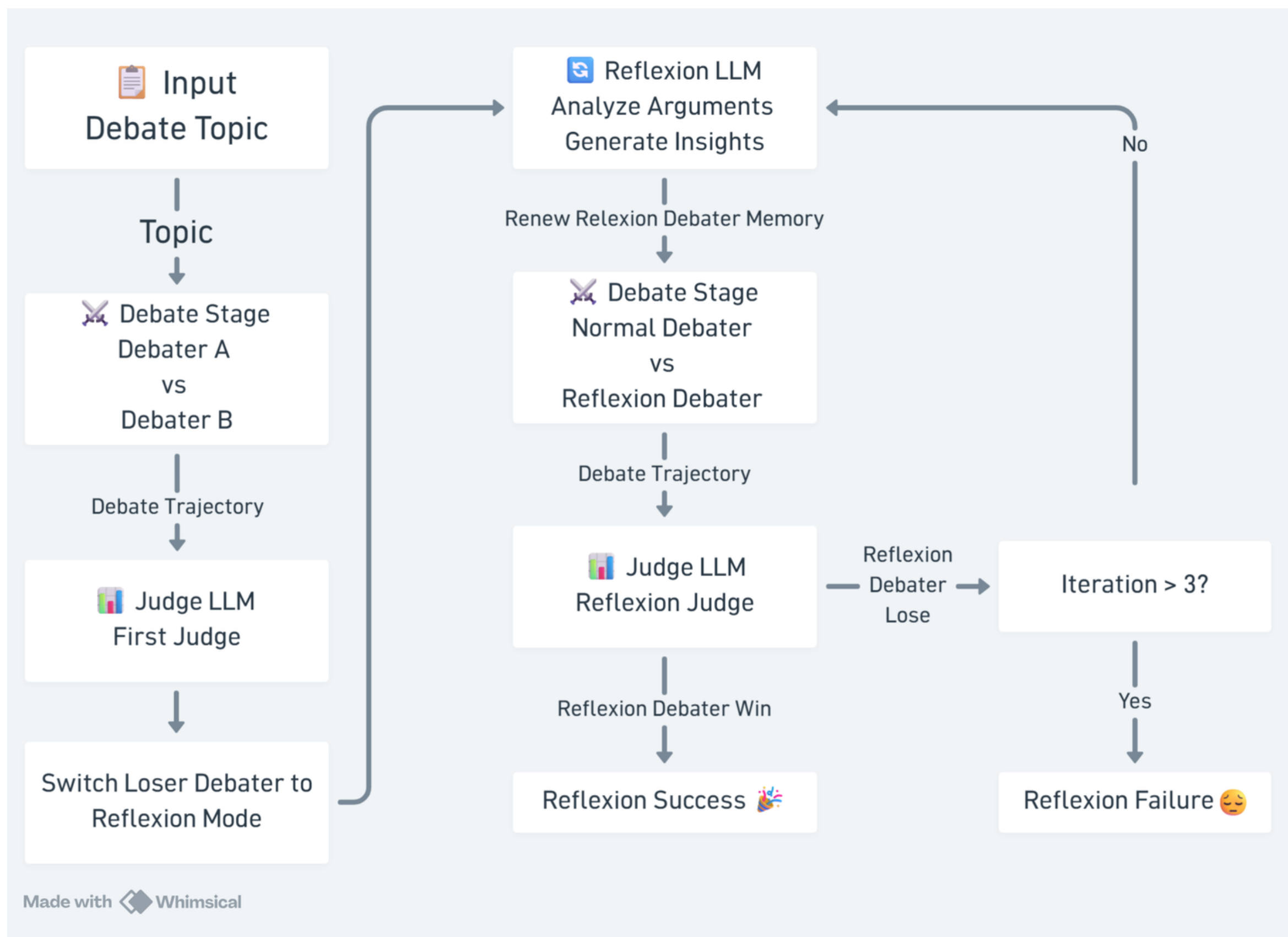


反思機制（Reflexion）

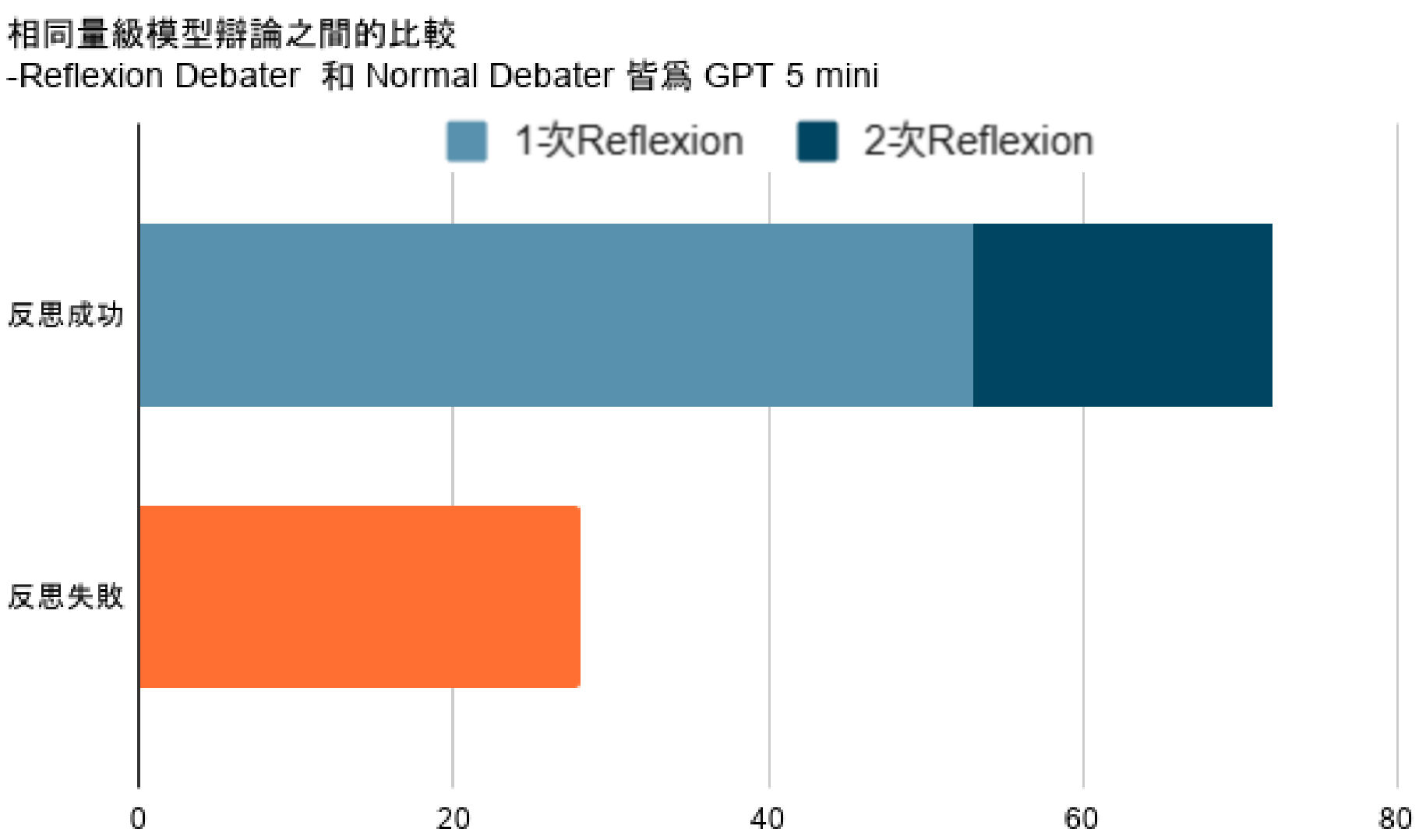
於人工智慧代理辯論任務中的應用與成效分析

系統流程圖

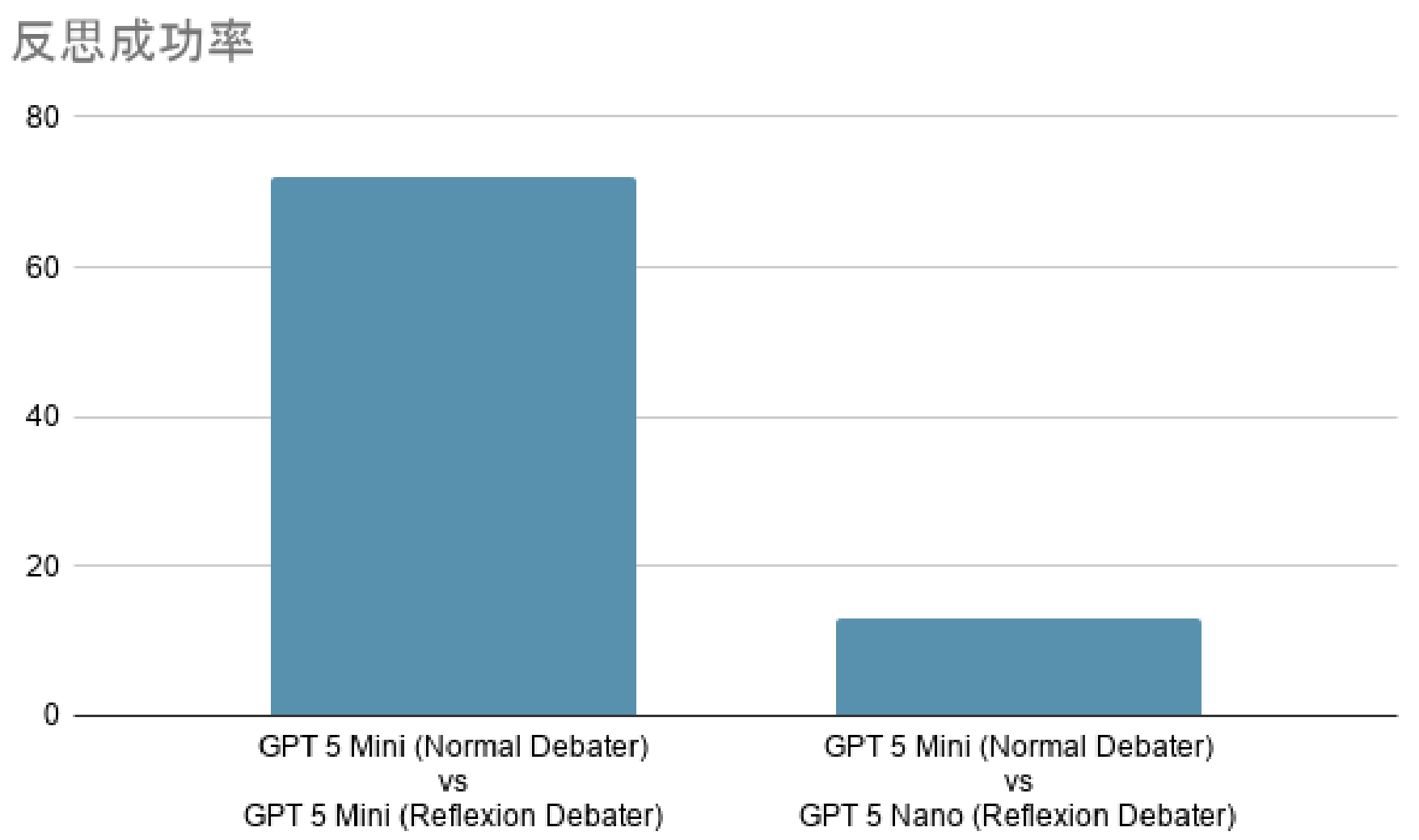


主要發現1

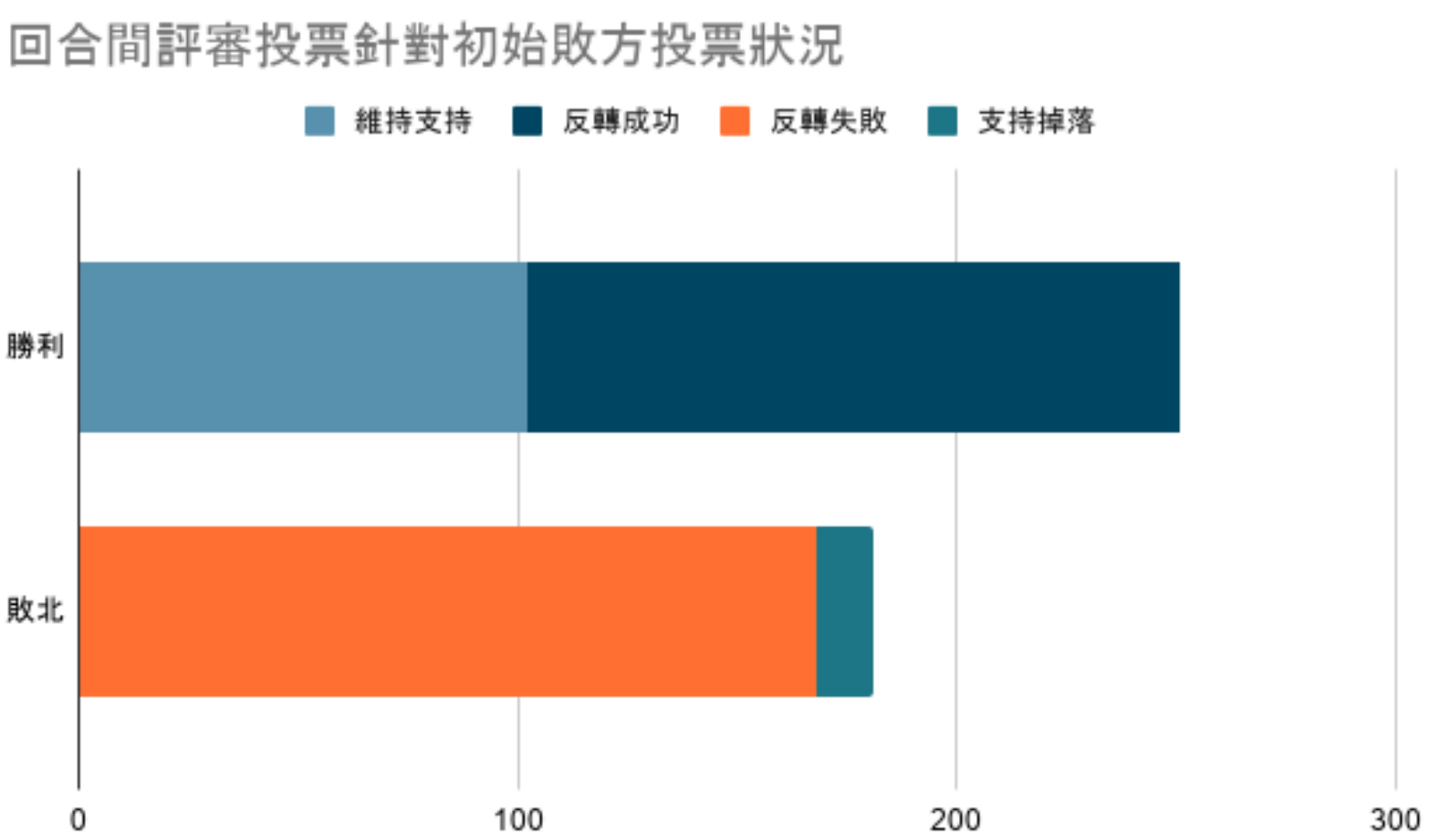
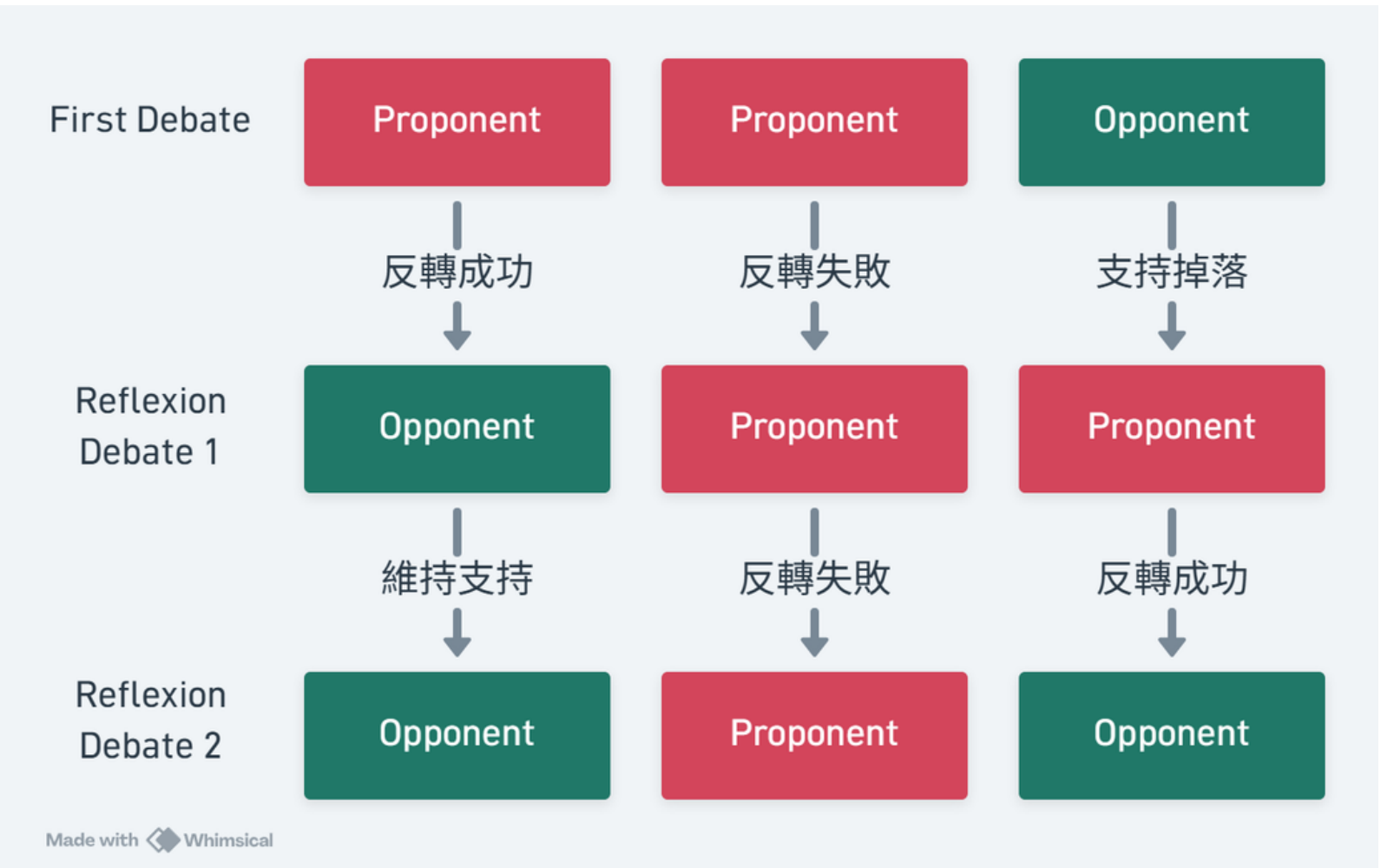
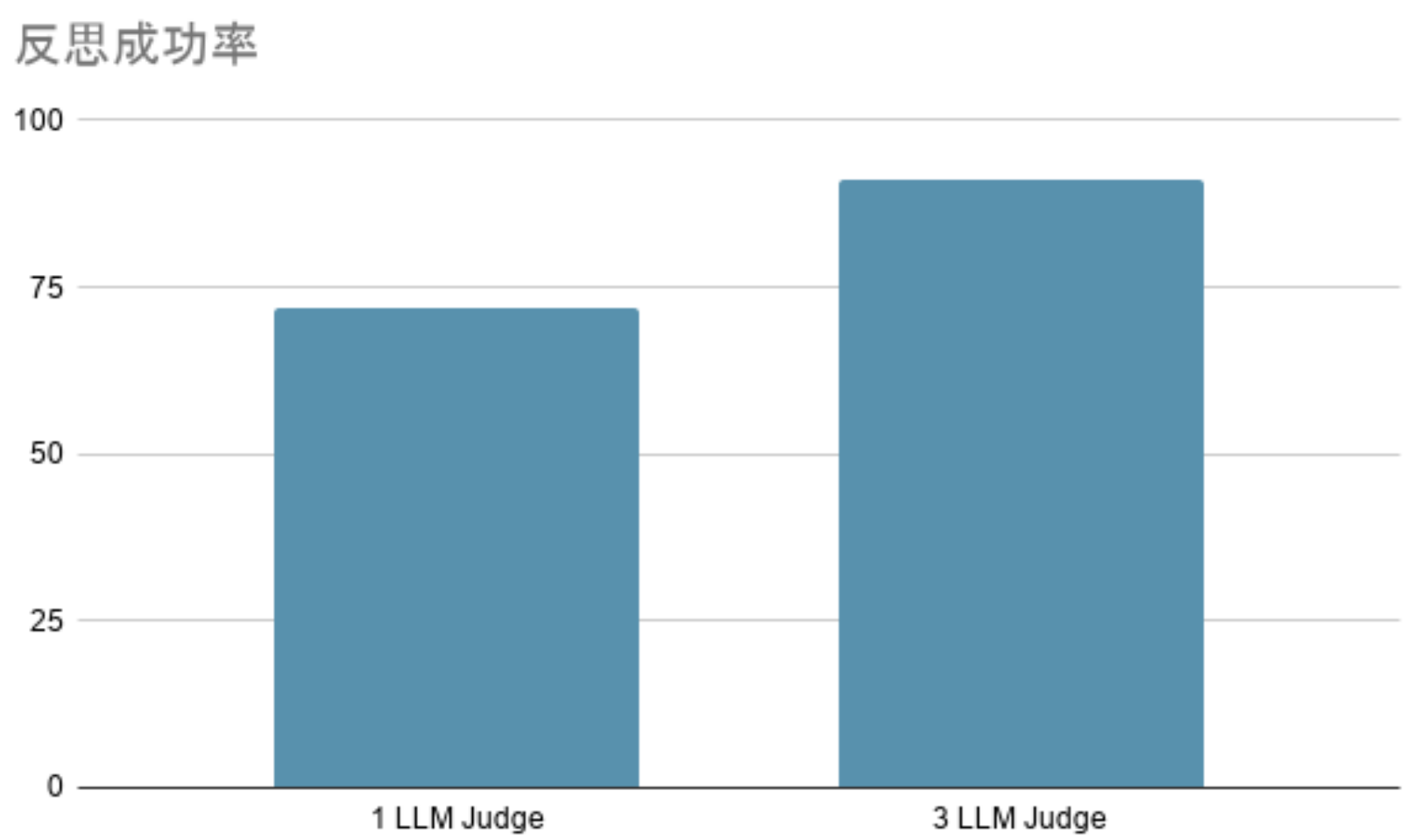
指導教授：朱威達
組員姓名：陳敬研



主要發現2



主要發現3



研究動機與目的

近年來，大型語言模型在多跳問答、程式碼生成等推理任務中等具有明確正解的任務當中表現良好。相較之下，對於缺乏標準答案的推理能力相關研究比較少。

本研究期望補充現有反思機制，聚焦於具明確正解任務之研究缺口，並評估辯論作為分析語言模型推理能力與策略演化之研究工具的可行性與限制。

什麼是REFLEXION

