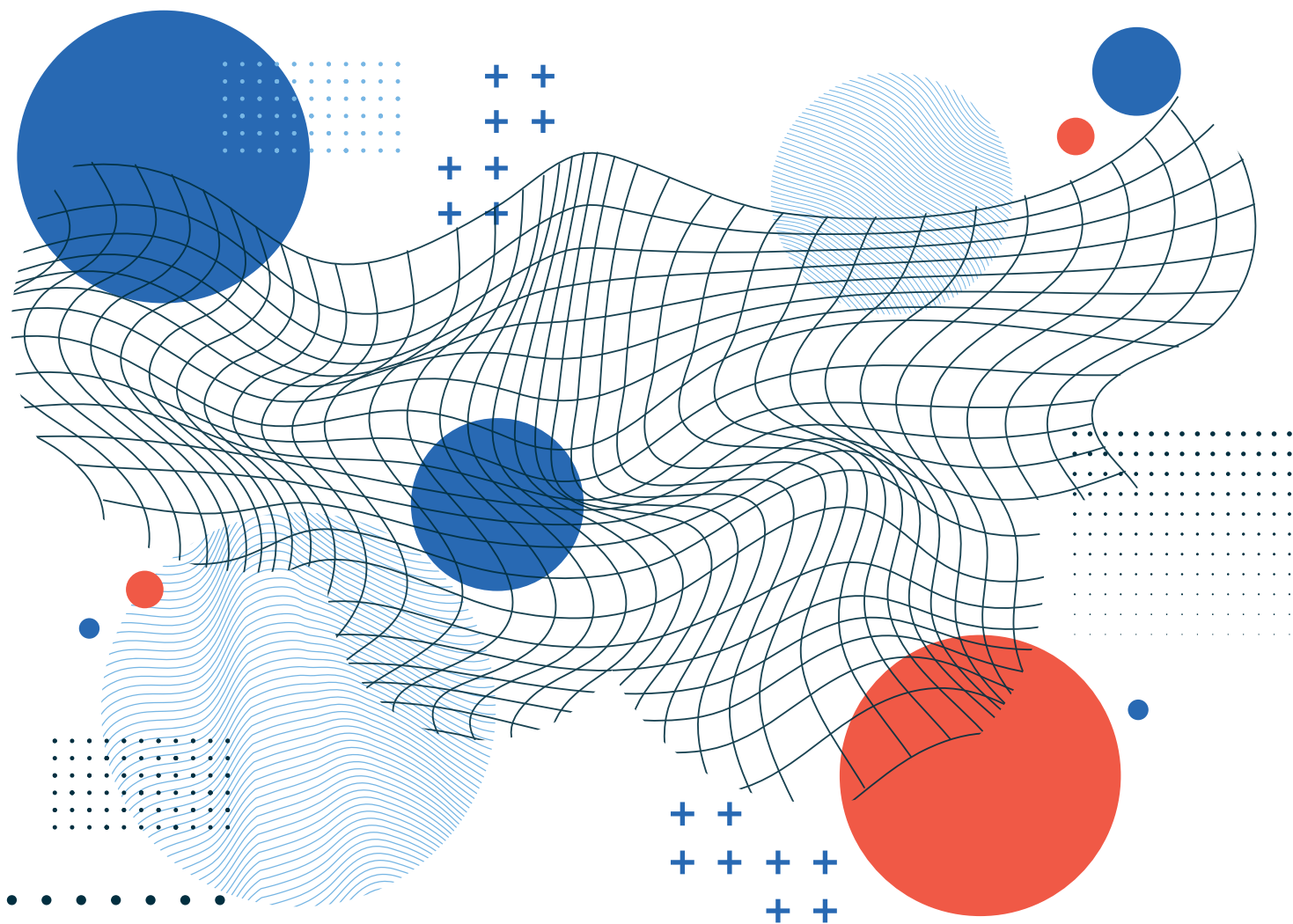


2026 ICEET

第十五屆 數位學習與教育科技國際研討會
International Conference on E-learning and Educational Technology

智慧教育×永續共榮

論文集 PROCEEDINGS



多個 AI 工具真的更安全嗎？生成式 AI 輔助資料分析中的自動化偏誤、態度-行為斷層與多工具悖論

When More AI Is Less Safe: Automation Bias, Attitude–Behavior Gaps, and the Multi-Tool Paradox in GenAI-Assisted Data Analysis

陳正和 (Jeng-Her Alex Chen)¹

¹ 助理教授，企業管理學群，元智大學 / Assistant Professor, Discipline of Business Management, Yuan Ze University

E-mail: alex@saturn.yzu.edu.tw

摘要

本研究探討大學商管學生在生成式 AI (GenAI) 輔助資料分析任務中的決策行為。以 161 名修習商業分析課程之大學生為研究對象 (來自 X、Y、Z 三班，資料來源為 Quiz01 與 W2 問卷之配對紀錄)，本研究提出三項研究問題：(1) 統計先備知識能否預測任務犯錯行為？(2) 自我報告之自動化偏誤量表 (ABI) 與批判思維指數 (CI) 能否預測實際犯錯？(3) 使用多個 AI 工具交叉驗證，是否能降低或反而提高犯錯率？

三項研究問題之答案如下：(1) 統計先備知識無法預測犯錯——Quiz01 分數在完整模型中 $OR = 1.215$ ($p = .556$)，天花板效應使測驗幾乎喪失區辨力；(2) ABI ($p = .858$) 與 CI ($p = .323$) 同樣無法預測犯錯，確認知識-行為斷層與態度-行為斷層並存；(3) 多工具驗證者犯錯率 (18.9%) 反而高於單一工具使用者 (6.5%， $OR = 2.839$ ， $p = .072$)，顯示「確認偏誤放大效應」。此外，Gemini 使用者犯錯機率为非 Gemini 使用者之 3.5 倍 ($OR = 3.515$ ， $p = .045$)，揭示工具特定精確度干擾。上述發現挑戰了現行 AI 素養教學的核心假設。本研究建議：AI 素養課程應從「比答案」升級為「比程序、比證據、比可重現性」；教師應引導學生以獨立計算方法 (而非多問幾個 AI) 進行交叉驗證，並針對不同任務類型評估各 GenAI 工具的精確度差異。本研究所提出的 TRV (可追溯性—可重現性—可驗證性) 框架，可直接轉化為 AI 輔助作業的評量標準，作為高等教育 AI 融入課程設計的操作指引。

關鍵詞：自動化偏誤、生成式 AI、知識-行為斷層、多工具驗證、AI 素養教育

When More AI Is Less Safe: Automation Bias, Attitude–Behavior Gaps, and the Multi-Tool Paradox in GenAI-Assisted Data Analysis

Jeng-Her Alex Chen

Assistant Professor, Discipline of Business Management, Yuan Ze University

Abstract

This study investigates decision-making behavior among undergraduate business students engaged in GenAI-assisted data analysis tasks. Using a matched dataset of 161 students (from three class sections) derived from Quiz01 knowledge assessments and W2 survey responses, the study addresses three research questions: (1) Can prior statistical knowledge predict task blunder behavior? (2) Can self-reported Automation Bias Index (ABI) and Critical Thinking Index (CI) predict actual task errors? (3) Does multi-tool cross-verification reduce or increase error rates?

Results address the three research questions as follows: (1) Prior statistical knowledge did not predict task blunders ($OR = 1.215$, $p = .556$), with ceiling effects rendering the knowledge test nearly non-discriminatory. (2) Neither ABI ($p = .858$) nor CI ($p = .323$) predicted task errors, confirming the co-existence of knowledge–behavior and attitude–behavior gaps. (3) Students using multiple AI tools for cross-verification exhibited a higher error rate (18.9%) than single-tool users (6.5%; $OR = 2.839$, $p = .072$), demonstrating a “confirmation bias amplification effect.” Additionally, Gemini users were 3.5 times more likely to commit errors than non-Gemini users ($OR = 3.515$, $p = .045$), revealing a tool-specific accuracy confound. These findings challenge core assumptions of current AI literacy education. The study recommends that AI literacy curricula shift from “comparing answers” to “comparing procedures, evidence, and reproducibility”; students should be trained to verify AI outputs using independent computational methods rather than querying multiple AI tools. The proposed TRV (Traceability–Reproducibility–Verifiability) framework offers a directly actionable evaluation rubric for AI-assisted assignments in higher education.

Keywords: automation bias, generative AI, knowledge–behavior gap, multi-tool verification, AI literacy education

