

讓CLD走出心智的陰影

資料驅動與AI協作的建模革新路徑

廖東山
元智大學管理學院

於
2025 中華系統動力學學會年會暨研討會



內容

CONTENTS

- 01 問題意識與時代挑戰
- 02 方法論反思：CLD的學術挑戰與知識重構
- 03 四階段AI協作CLD建構流程
- 04 實作案例分享
- 05 方法論貢獻與理論定位

讓CLD走出心智的陰影：資料驅動與AI協作的
建模革新路徑

問題意識

- Littlejohns et al. (2021)：針對2018至2021年間公共衛生領域中CLD應用進行系統性回顧，
 - 多數未明確說明其建構流程與參與角色
 - 缺乏模型形成的透明性與可複審性。
- Jalali與Beaulieu (2023) 則針對72篇CLD研究進行分析，指出
 - 僅有25%的研究清楚標示因果連結的資料來源
 - CLD在方法論透明度上仍存結構性缺口。

當CLD成為制度模擬的起點，它的主觀性與資料貧乏問題，卻讓我們一再失去對建模的信任（Stermann，2000；Richardson，1991）

CLD的學術挑戰

- CLD 的三大基本科學思辨的挑戰，均可歸因於SD研究者的“心智模式依賴”現象
 - **主觀偏誤**：心智模式常基於個人經驗與直觀理解，缺乏明確建模依據，導致CLD建構結果反映的是「認知投射」而非「系統結構」。
 - **資料貧乏**：模型來自腦中想像而非語料萃取或資料驗證，即便圖像看似合理，事實上缺乏經得起挑戰的知識基礎。
 - **過程不透明**：心智建模往往無正式記錄與版本控管，其他研究者難以重建其形成邏輯，使模型失去複現性與學術信任度。



CLD若僅是研究者的腦圖，它注定成為『美觀卻無法驗證』的敘事結構。

CLD的知識重構轉捩點

心智模式僅能視為一種接近「認識論的工具 (epistemic heuristic) 」而非「科學模型」論證的基礎。

- 模型即實踐 (modeling-as-practice)
 - 模型不應再被視為研究者個人心智圖像(靜態輸出)
 - 應理解為一種嵌入語境、結合資料與跨角色互動的動態知識建構

(Eidin et al. , 2024)
- 建構新命題：CLD模型是一種知識生產歷程，不只是結果圖像
 - 透過「建模實踐轉向」的視角，CLD不再是研究者的直覺產物
 - 結合多元資料、社會語境與人機協作共同生成的知識平台
 - CLD 本身承載著知識系統建構的角色
 - CLD的實踐即賦予政策對話媒介與跨領域理解工具的實用價值。

資料驅動結合AI協作的建模革新路徑

- 當代系統建模正處於一場知識革新的關鍵轉折。
 - 面對CLD長期以來「主觀偏誤、資料貧乏與過程不透明」三重挑戰
 - 資料驅動與AI協作的CLD建模路徑，不僅是一項技術創新，更是一種方法論思維的重構。
 - 建模應從個人心智的投射轉向語境資料的對話

將建模過程視為一場與資料、語義與文本的深度對話，才能讓CLD成為反映現實系統結構與動態邏輯的認識與詮釋平台。這種轉向不只是技術流程的調整，更是方法論立場的轉變



四階段AI協作CLD建構流程

我們不是要用AI取代思考，而是讓AI成為建模的語義橋梁與資料分析的協作者。

語境設定與資料準備

強調語境敏感與多文本並置（訪談、觀察、政策文本、新聞資料等）

語義萃取與因果語句分析

結合質性分析工具與生成式AI：

- 標記（編碼）因果語句
- 萃取變數
- 產生變數組合結構

結構關係與視覺建模

將碎片語句轉化為回饋迴路：

- 分析變數關係（命題建立）
- 建立CLD初稿
- 理論概化的最適原則

版本控管與來源註記

強調追溯性與邏輯透明性，使用版本註記表與出處對照表

- 透過質性分析工具
- 生成式AI+人工記錄

實作案例分享

ESG-S基礎之治理準平台構型：以制度模糊場域為例的CLD建構

- **階段1：語境設定與資料準備**

- RQ：在制度模糊與正式治理缺位的情境中，行動者如何建構穩定且具永續導向的治理秩序？
- 資料來源以深度訪談為主。
 - 訪談一（2024年6月）：聚焦其個人行動歷程、農場經營邏輯與合作經驗。
 - 訪談二（2025年1月）：了解合作規則如何形成、服務標準如何建立與關係如何維持。
 - 非正式田野觀察與歷史脈絡補充資訊
 - 輔以其過往受訪記錄與公開發表內容作為場域補充材料（惟不列為正式資料來源，僅協助語境理解與理論敏感性建立）。

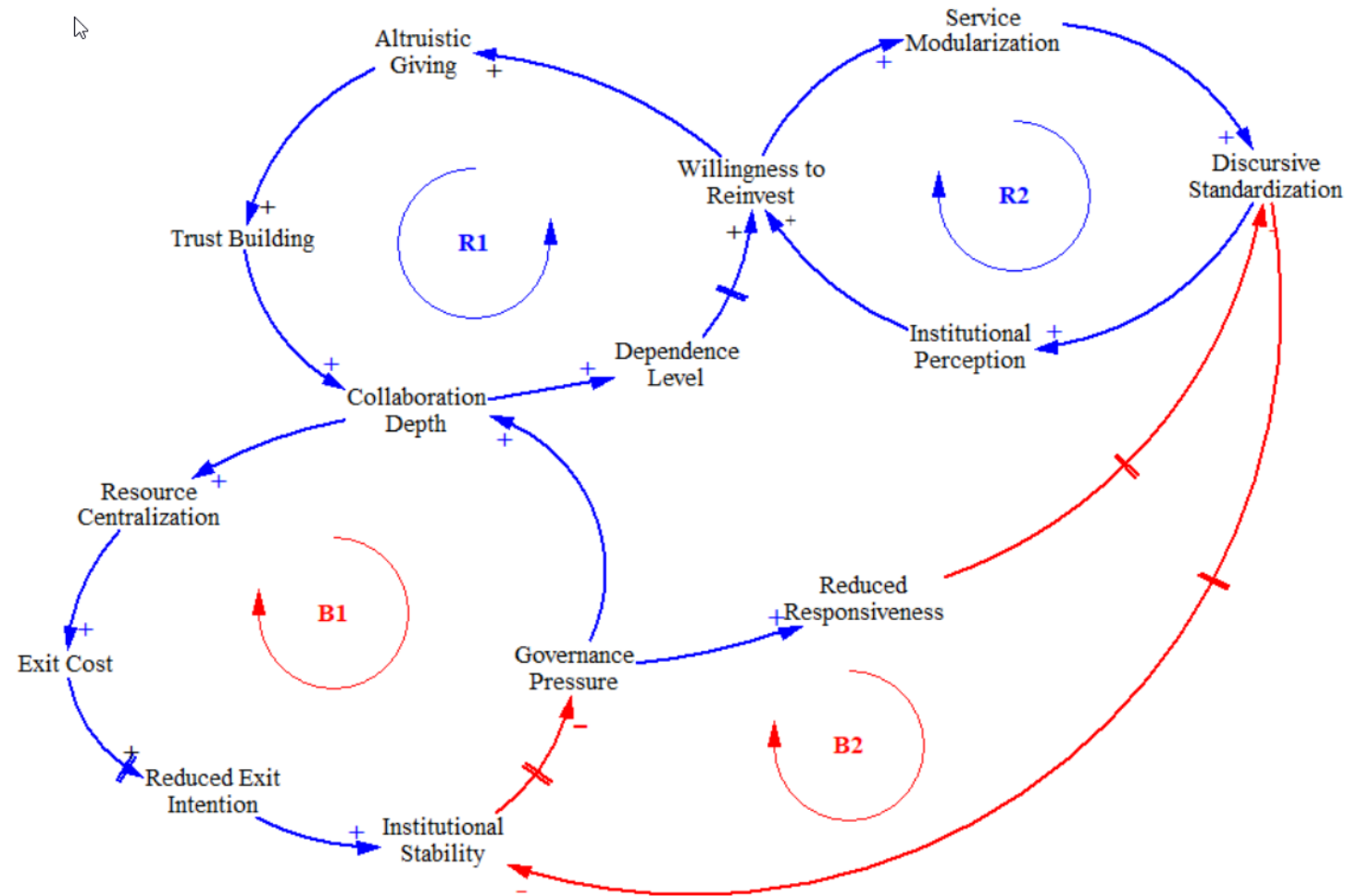
階段2：語義萃取與因果語句分析

- 依CGT進行資料分析，依序進行
 - 起始編碼
 - 聚焦編碼
 - 理論範疇建構
 - 輔以備忘錄書寫與反覆比較策略
- 工具
 - Maxqda + ChatGPT

理論範疇	聚焦編碼（核心節點變數）	起始編碼（關鍵語料）
利他治理	讓利	資材內攬、讓利固著、風險吸納、讓利擴大、去中介轉利、議價權共享
	去中介化	承接瓶頸、重建交付、整合通路、簡化決策、取代調度、內建耕作
	語用規則取代契約	內化流程、統一語意、語言劃界、語句制度化、預設結構、隱性約束、鎖定路徑
資源互賴	構築退出成本	製造相對損失感、封閉服務路徑、壟斷資源調度、設置排他節點、整併功能節點、擴張規模優勢
	服務模組化	作業模組串接、需求驅動模組發展、推動模組化複製擴散
	角色配置	預設角色範疇、前置行動節奏、階段性角色嵌入、去技能化角色限定、壓縮協作界面、語用角色檢核、制度性角色定型
語境平台	語境化準平台（協作）	模糊化平台角色、透明化協作授權、即時化利益回饋、價值導向協作正當化、非商業化平台擴散
	制度默契	內化作業語境、預設協議邊界、建立共享邏輯、公平性默契、可學習性制度穩定
	標準化形成	統一分配規格、包裝標準化、議價制度化、物流標準化、儲藏標準化

階段3：結構關係與視覺建模

- 利他治理的行動結構
 - 命題一：行動者透過讓利與去中介化方式建立信任與合作穩定性
 - 命題二：語用規則逐步取代契約，構成行動預期與治理默契
- 資源互賴與制度感知的生成
 - 命題三：高互賴結構形成退出成本，進而產生制度依附感
 - 命題四：制度穩定性由服務模組與角色配置的集成所建構
- 語境平台性的浮現與治理邏輯
 - 命題五：平台功能在語境中逐步浮現，非先行設計
 - 命題六：標準化與制度默契由共用歷程積累，而非技術規範建立



SQPGC:語境範疇與治理構形之中層理論暨CLD建構

結構關係與視覺建模

R1 = 命題1~3 的視覺呈現為例

- **Altruistic Giving**（利他讓利） 以「不賺中間價」、「我幫你先墊」、「大家有就好」等行動語句為表徵，展現紀元農莊將資源讓利視為治理啟動機制。
 - **Trust Building**（信任建構） 讓利形成「你不會害我」的行動記憶，降低對合作風險的感知，提升對組織正當性與長期承諾的信任基礎。
 - **Collaboration Depth**（合作深度） 信任提升使行動者願意交付更多工作面向（如配送、包裝、上架），形成模組化的合作流程。
 - **Dependence Level**（依賴強度） 合作項目增加帶來功能與效率上的結構性依賴，如「不走這套流程就會卡關」，角色難以替換，強化制度認同。
 - **Willingness to Reinvest**（再投入意願） 由於信任穩定與流程熟悉，行動者願意長期投入資源與行動，增強黏著性。
 - 回到 **Altruistic Giving** 再投入意願提升後，組織持續願意讓利與維持信任模式，正向循環完成。
- 命題一：利他讓利作為秩序生成的初始行動→ 對應 **Altruistic Giving** 啟動治理迴路本節指出，治理啟動並非來自制度規範，而是行動者自願讓利所建立的行動正當性（例如「你不賺我就信你」的語境），正是迴路的第一節點。
 - 命題二：讓利結構提升信任交換的穩定性→ 對應 **Trust Building → Collaboration Depth**行動者因讓利行動而建立起信任記憶，願意嘗試更多模組合作，如「你弄我就配合」，信任成為合作深化的邏輯中介。
 - 命題三：信任與合作經驗形成角色依賴與再投入→ 對應 **Dependence Level → Willingness to Reinvest**當合作變為默契、角色難以取代，便產生制度黏著性。參與者轉而內部化對流程與角色的依賴，表現為「我也不會走了」的語境語句。

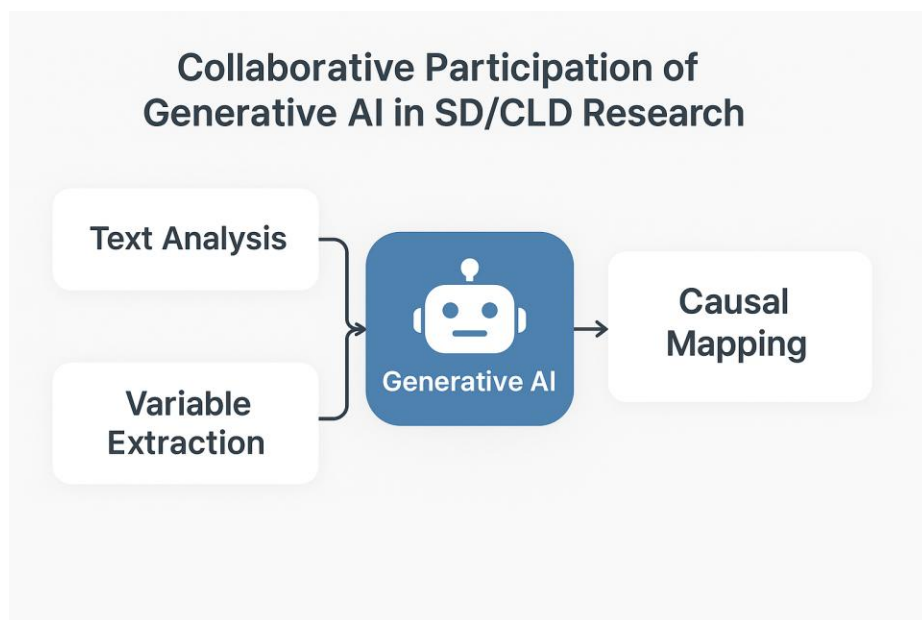
階段4：版本控管與來源註記

R1 迴路分析為例：利他信任驅動的治理生成機制

節點	邏輯功能	語境語料例示	對應命題
Altruistic Giving （利他讓利）	治理啟動，建立初始正當性與信任前提	「我不賺中間價」、 「你種我來幫你處理」	命題一：利他讓利作為秩序生成的初始行動
Trust Building （信任建構）	降低風險感知，提升行動黏著與合作期待	「我知道他不會害我」、 「信得過才敢交給他」	命題二：讓利結構提升信任交換的穩定性
Collaboration Depth （合作深度）	行動分工擴張，角色功能模組化	「冷藏、貼紙、通路我們都處理好」	命題二：讓利結構提升信任交換的穩定性
Dependence Level （依賴強度）	角色互賴提升、合作難以中斷，制度內嵌啟動	「這一套流程你不照就卡關」、「久了你就離不開」	命題三：信任與合作經驗形成角色依賴與再投入
Willingness to Reinvest （再投入意願）	行動者因信任與穩定體系持續投入資源與參與	「他們願意一直合作就是因為有信任」	命題三：信任與合作經驗形成角色依賴與再投入

版本控管：可隨語境語料的增減而修訂版本
來源註記：可隨語料之起始或開放編碼之來源而記錄之

對SD科學論證哲學與方法途徑的貢獻



- **由心智推論轉向語境實證：重新定位CLD的知識屬性**
 - 透過語言資料的結構性萃取與因果語句的來源註記，可讓CLD由心智映射的工具，轉化為具備資料語境支撐與版本歷程記錄的「證據性模型」。
- **將建模視為知識生產歷程，而非心智映射的圖形化的結果**
 - CLD建構歷程本身應該是一種可記錄、可複審、可對話的知識生產過程。
 - 回應了建模即實踐觀點 (Eidin, et.al., 2024; Martinez-Moyano & Richardson, 2013)，主張模型不僅是科學論述的結果，也是一種過程性、參與性的知識實作。
- **生成式AI在SD/CLD研究中不應被視為取代者，而是作為語意橋接與結構輔助的建模協作者**
 - 你應該不會認為當你使用統計軟體進行統計分析，就是統計工具幫你做研究吧！？

結語

必須正視CLD長期以來的知識正當性爭議

讓CLD走出心智的陰影，並不意味著否定個人思考的重要性，而是希望透過資料、語言與技術的結合，讓這些思考可以被追溯、被挑戰、被再建構。報告者展示的，是一種AI時代的CLD建模實踐途徑，也是一個系統動力學在跨界知識建構中重新定位的起點。

感謝聆聽

廖東山

元智大學管理學院

Email: valenliao@saturn.yzu.edu.tw

LineID: valenliao0214