

การประยุกต์ใช้ปัญญาประดิษฐ์แบบรู้สร้างทั้งในด้านการวิเคราะห์ และการวิจัยในสังคมศาสตร์

Generative AI's Role in Social Science Both in Research Assistance and Analysis

กานต์ ยืนยง

Kan Yuenyong

นักศึกษาปริญญาเอก คณะรัฐประศาสนศาสตร์ สถาบันบัณฑิตพัฒนบริหารศาสตร์

Ph.D. Candidate, Graduate School of Public Administration,

National Institute of Development Administration

E-mail: sikkha@gmail.com

This article is presented in the 17th National Conference on Public Administration of
Public Administration Association of Thailand at Ubon Ratchathani University, November 17, 2023

Received: October 12, 2023; Revised: December 2, 2023; Accepted: December 15, 2023

บทคัดย่อ

ในยุคที่เทคโนโลยีปัญญาประดิษฐ์แบบรู้สร้างเติบโตอย่างต่อเนื่อง โดยเฉพาะโมเดลภาษาขนาดใหญ่ที่กำลังได้รับความสนใจเป็นพิเศษในการประยุกต์ใช้เพื่อช่วยงานวิจัยในด้านต่าง ๆ บทความนี้วิเคราะห์วิวัฒนาการการเรียนรู้ของเครื่องจักรตลอดจนการถือกำเนิดของโมเดลภาษาขนาดใหญ่ โดยจะเน้นการวิเคราะห์ไปที่โครงการของ DARPA ในการผสมผสานวิจัยคุณภาพและวิธีการคำนวณที่ขั้นสูงโดยเฉพาะผ่านโครงการ Next Generation Social Science (NGS2) และ Ground Truth (GT) ต่อมาโครงการเหล่านี้จะมีพัฒนาการเปลี่ยนเป็นการใช้คอมพิวเตอร์ในสังคมศาสตร์ (Computational Social Science) และกระบวนการวิจัยแบบผสมผสานซึ่งจะรวมวิธีวิจัยเชิงคุณภาพ (QUAL) และวิธีวิจัยเชิงปริมาณ (QUAN) ซึ่งจะทำให้งานวิจัยมีความครอบคลุมหลากหลายมากขึ้น แต่แม้ว่าปัญญาประดิษฐ์แบบรู้สร้างจะมีศักยภาพมาก แต่ยังมีโอกาสที่จะผิดพลาดหรือที่เรียกกันว่าการเกิด “การหลอน” (Hallucination) ในบทความนี้จะพูดถึงมาตรการในการเพิ่มความเที่ยงตรงเรื่องนี้ นอกจากนี้บทความนี้ยังกล่าวถึงตัวอย่างการนำปัญญาประดิษฐ์แบบรู้สร้างมาใช้งานจริงเพื่อตรวจจับแนวโน้มและการสร้างฉากทัศน์แบบอัตโนมัติ

คำสำคัญ: การวิจัยแบบผสมผสาน; ปัญญาประดิษฐ์แบบรู้สร้าง; โมเดลภาษาขนาดใหญ่;
การใช้คอมพิวเตอร์ในสังคมศาสตร์

Abstract

In an era where generative artificial intelligence technology is continuously growing, especially large language models that are receiving significant attention for their application in assisting various research endeavors, this paper analyzes the evolution of machine learning up to the inception of large language models. The discussion highlights DARPA's projects in integrating qualitative research with advanced computational methods, specifically through the Next Generation Social Science (NGS2) and Ground Truth (GT) projects. Subsequently, these initiatives have evolved into Computational Social Science and mixed method research, which combines qualitative (QUAL) and quantitative (QUAN) research methods, resulting in more diverse and comprehensive research outcomes. However, while generative artificial intelligence holds great potential, there are chances of errors, commonly referred to as "hallucination". This paper addresses measures to enhance accuracy concerning this aspect. Furthermore, the paper also presents practical examples of utilizing generative artificial intelligence for trend detection and automated scenario creation.

Keywords: Mixed Method Research; Generative Artificial Intelligence; Large Language Model; Computational Social Science

บทนำ

ความจริงฐาน (Ground Truth) หรืออาจจะนิยามว่า “ข้อมูลโลกจริง” ก็ได้ ซึ่งต่อไปนี้นับวันจะกลายเป็นคำสามัญที่เราต้องทำความเข้าใจมากขึ้นทุกที ในพจนานุกรมออกซ์ฟอร์ดได้ชี้ให้เห็นว่าคำศัพท์คำนี้มีที่มาศตวรรษก่อนนี้แล้วในบทกวีของเฮนรี เอลลิสัน (Henry Ellison, pp. 1811-1880) เรื่อง “นิทานนเรนเทศไซบีเรีย” ตั้งแต่ปี ค.ศ. 1833 แต่ถูกนำมาใช้ในแวดวงภาพเทคโนโลยีสำรวจข้อมูลระยะไกล (remote sensing) ที่เกี่ยวข้องกับแวดวงการทำแผนที่ (cartography) อุตุนิยมวิทยา (meteorology) ภาพถ่ายทางอากาศและภาพถ่ายดาวเทียม เพราะจำเป็นจะต้องมีการสอบวัดความถูกต้องและความสมนัยระหว่าง สิ่งที่เกิดขึ้นในภาพถ่ายและภาพวาดที่ได้จากอุปกรณ์ กับความจริงบนภาคพื้นสนามข้างนอกนั้น (ground truth) พจนานุกรมออกซ์ฟอร์ดยังชี้เพิ่มเติมอีกด้วยว่ารากของคำว่าความจริงฐานในภาษาอังกฤษอาจมีที่มาจากคำในภาษาเยอรมันว่า “Grundwahrheit” ซึ่งมีความหมายว่า “ความจริงขั้นมูลฐาน” ตั้งแต่คริสต์ศตวรรษ 1650 มาก่อนหน้านี้แล้ว (Woodhouse, 2021, pp. 041501-2-3)

ในไม่ช้าคำว่าความจริงฐานนี้จะถูกนำไปใช้กับแวดวงการวิจัยด้านวิศวกรรมชีวภาพ เพราะต้องสอบวัดความถูกต้องระหว่างผลที่ได้จากแบบจำลองในคอมพิวเตอร์ (in silico) หรือในเนื้อเยื่อเพาะเลี้ยง (in vitro) ซึ่งทั้งคู่ต้องจัดให้เป็นผลที่ได้จากสิ่งที่ยังไม่เกิดขึ้นจริงหรือเกิดขึ้นในห้องทดลองหรือแบบจำลอง (in situ) จัดเป็นเรื่องราวตรงข้ามกับความจริงฐานที่ได้จากผลที่เกิดขึ้นในสิ่งมีชีวิตจริงๆ (in vivo) ซึ่งเกิดขึ้นนอกห้องทดลองหรือนอกแบบจำลอง (out situ)

ความจริงเรื่องนโยบายสาธารณะ อย่างเช่นเงินดิจิทัล 10,000 บาท นี้ก็มีส่วนเกี่ยวข้องกับความจริงฐานที่ว่าน้อยอยู่นั่นเอง เพราะการอภิปรายสาธารณะมักจะเกิดขึ้นจากต่างฐานความรู้ ต่างประสบการณ์ กระทั่ง

ต่างวิธีวิทยา แล้วอะไรจะเป็นความจริงฐานที่จะบอกได้ว่า สิ่งที่เราตั้งสมมติฐานจากหลายฝ่ายนั้น ใช่ หรือ ไม่ใช่ อะไรที่เป็นนิยาม ผลสำเร็จที่เกิดขึ้น อะไรเป็นข้อจำกัดที่ยอมรับได้ อะไรเป็นผลกระทบทั้งโดยตั้งใจ และผลกระทบโดยไม่ตั้งใจเชิงนโยบายที่อาจส่งผลกระทบต่อสาธารณะในวงกว้างทั้งด้านบวกและด้านลบ เราจะใช้อะไร เป็นมาตรวัดประสิทธิภาพและประสิทธิผลของนโยบายสาธารณะนั้นๆ แต่หากมองข้ามข้อถกเถียงทั่วไปให้ลงลึกไปยังปมหลักของข้อขัดแย้งระดับรากฐาน ทั้งนี้ให้สมมติว่าผู้อภิปรายมีความรู้พื้นฐาน ประสบการณ์ และความสามารถผลิตงานวิจัยสนับสนุนไม่ต่างกัน ซึ่งแน่นอนว่าข้อสมมติฐานนี้ไม่เป็นจริง แต่เอาเป็นว่าเราสมมติให้ตัวแปรนี้ไม่มีผลเอาไว้ก่อน

ปัญหาที่การอภิปรายในระดับนี้คือมักจะมีการอภิปรายจากองค์ความรู้ที่แตกต่างกัน จากผู้ที่ยืนอยู่บนรากฐานงานวิจัยเชิงปริมาณ และผู้ที่ยืนอยู่บนรากฐานงานวิจัยเชิงคุณภาพ กล่าวในแง่ที่ลึกที่สุดคือเรื่องญาณวิทยา (epistemology) หรือการสร้างความรู้จากสองพวกนี้มีรากที่ต่างกันสุดขั้ว ฝ่ายหนึ่ง คือฝ่ายปริมาณ เชื่อว่า “ความจริงนั้นมีอยู่” และความจริงแยกอยู่ต่างหากจึงต้องหาวิธีกันอคติของผู้สังเกตการณ์ให้มากที่สุด เราสามารถสกัดความจริงออกมาได้หากมีระเบียบวิธีวิจัยที่เป็นวิทยาศาสตร์และรัดกุมพอ อาทิ การกำหนดระเบียบวิธีทางสถิติ หรือการกำหนดระเบียบวิธีการทดลอง ส่วนอีกฝ่ายหนึ่งคือฝ่ายคุณภาพได้ทำการปฏิเสธสัจพจน์ (axiom) อ้างอิงของฝ่ายปริมาณทั้งหมด แล้วกำหนดสัจพจน์ของฝั่งตนบ้างว่า “ความจริงนั้นไม่มี” เพราะความจริงมีหลายมิติความจริงไม่สามารถแยกออกจากตัวผู้สังเกตการณ์ ดังนั้นจึงควรอนุมานได้ว่าการมีอคตินั้นเป็นเรื่องธรรมดา สิ่งที่เราเพียงทำได้คือการสกัดประสบการณ์ของผู้ร่วมสังเกตการณ์ ความจริงจึงย่อมแปรเปลี่ยนไปได้ตามการรับรู้และประสบการณ์ของผู้สังเกตการณ์ อนึ่งในเชิงปรัชญา เรื่องนี้มีการแยกนิกาย (schism) ของกระบวนวิธีทั้งสองแบบในยุคของอิมมานูเอล คานท์ (Immanuel Kant, pp. 1724-1804)

แต่แล้วในภายหลังก็มีผู้คิดได้ว่า ในเมื่อต่างฝ่ายต่างก็มีจุดอ่อน ในเมื่อต่างฝ่ายต่างมีจุดแข็ง ถ้าเช่นนั้นทำไมไม่นำมารวมกันเสียละ ฝ่ายนี้คือกลุ่มที่ใช้กระบวนกรวิจัยแบบผสมผสาน (Mixed-Method Methodology) ซึ่งนับวันเริ่มได้รับความนิยมในงานวิจัยจากวิทยาศาสตร์สุขภาพมากขึ้นทุกที เพราะเป็นศาสตร์สาขาที่ต้องปะทะกับทั้งวิทยาศาสตร์กายภาพและวิทยาศาสตร์สังคมไปพร้อมกัน เนื่องจากการแพทย์ในแง่หนึ่งก็เป็นเรื่องเทคนิคทางชีววิทยา ในอีกแง่หนึ่งก็เป็นปรากฏการณ์สังคมไปด้วย

แต่ก็นั่นแหละการวิจัยผสมผสาน จะผสมและผสานอย่างไร เอาวิจัยเชิงคุณภาพมาก่อนเพื่อสร้างทฤษฎี แล้วเอาวิจัยเชิงปริมาณตามมาเพื่อสอบทฤษฎี หรือกลับกัน หรือจะทำวิจัยทั้งสองแบบคู่ขนานกันไปเพื่อยืนยันกันและกัน (triangular research) การผสมแบบนี้ถ้าคิดให้ดีมีได้ไม่จำกัด ทั้งลำดับและทั้งจำนวน ขั้นตอนการวิจัยของแต่ละประเภท ว่ากันที่จริงก็มีลักษณะตามอำเภอใจอยู่ไม่น้อย นี่ทำให้แวดวงกรวิจัยแบบผสมผสานในระยะหลังจึงเริ่มหา (1) วิธีสร้างมาตรฐานให้ยอมรับกันได้, (2) เริ่มหาปรัชญาที่มาสร้างความชอบธรรมให้กระบวนวิธีวิจัยอย่างนี้ ซึ่งในไม่ช้าก็ไปค้นเจอกันที่ Pragmatism philosophy ของพวกปรัชญาอเมริกัน และ (3) ในเมื่อกระบวนกรวิจัยผสมผสานนี้มาจากต่างฝ่ายต่างเผ่าพันธุ์ ดังนั้นก็จึงต้องหาพื้นที่กลางที่พอยอมรับกันได้จากทั้งสองฝ่าย ตรงนี้แหละที่ “ความจริงฐาน” จะเข้ามามีบทบาท

ส่วนในแวดวงเอไอหรือระบบปัญญาประดิษฐ์ โดยเฉพาะพวกที่ทำวิจัยเรื่องเครือข่ายประสาทเทียม (artificial neuron network) ไม่ใช่พวกสายระบบผู้เชี่ยวชาญ (expert system) ในสมัยก่อนที่มักจะเป็นอัลกอริทึมตรรกะแบบตายตัว ซึ่งในหนังสืออย่าง Levine, Drang and Edelson (1991) ก็พูดถึงแนวคิดเรื่องเครือข่ายประสาทสมัยใหม่อยู่บ้างแต่ทั้งเทคนิคในการอิมพลีเมนต์และเทคโนโลยีฮาร์ดแวร์ในขณะนั้นยังไม่มี

ประสิทธิภาพเทียบเท่าในปัจจุบัน คนที่แยกสองเรื่องนี้ไม่ออกจึงต้องถือว่าไม่รู้อะไรเกี่ยวกับเรื่องเอไอสมัยใหม่เลย คนพวกนี้มักจะเชื่อว่า “เอไอใช้ไม่ได้จริง” ในทางตรงข้ามระบบเครือข่ายใยประสาทเทียมจะมีกลไกการเรียนรู้ทั้งไปข้างหน้าและย้อนกลับ และปรับให้เที่ยงตรงมากขึ้น เที่ยงตรงกับอะไร? เที่ยงตรงกับโจทย์ที่เราต้องการให้เอไอนั้นแก้โจทย์ นั่นคือเราต้องกำหนดโจทย์เป้าหมายเบื้องต้นซึ่งก็คือความจริงฐานนี้มาให้เอไอตอบสนองให้สามารถยืนยันความถูกต้องให้ได้ก่อน ทั้งนี้โดยนัยยะคือการตรวจสอบประสิทธิภาพของโมเดลนั้นกับข้อมูลทั้งที่มีอยู่ และข้อมูลที่จะป้อนต่อไปในอนาคตไปด้วยในตัว

ผู้ทำการทดลองหรือฝึกสอนเอไอนั้น เทียบสิ่งที่เอไอตอบกับโจทย์ที่ตั้ง (ความจริงฐาน) แล้วเทียบค่าที่ได้ ค่าความผิดพลาดที่เกิดขึ้นเรียกว่า ค่าต้นทุน (Cost) ค่าสูญเสีย (Loss) ค่าผิดพลาด (Error) หรือค่าเหลือ (Residual) เป็นต้น อุดมคติที่เราต้องการคือค่าผิดพลาดจะต้องน้อยที่สุด แต่ก็ไม่ใช่ให้น้อยจนเหลือศูนย์ทีเดียว เพราะไม่เช่นนั้นจะเกิดปัญหาที่เรียกว่า โมเดลเข้ากันได้กับข้อมูลมากเกินไป (overfitting) เพราะหากเป็นเช่นนั้นเวลาโมเดลไปเจอข้อมูลที่แตกต่างจากแบบแผนไปมากก็จะไม่สามารถทำนายหรือใช้ประโยชน์ได้ในทางกลับกันโมเดลที่หลวมเกินไป (underfitting) ก็จะล้มเหลวในการทำนายข้อมูลส่วนใหญ่ การสร้างสมดุลระหว่างสองประเภทต้องใช้ทั้งศาสตร์และศิลป์

นักวิจัยและผู้ฝึกสอนปัญญาประดิษฐ์จะต้องเชี่ยวชาญทั้งการหาความจริงฐานที่ว่านี้ เพื่อพิจารณาค่าผิดพลาด และหาช่วงจุดเปลี่ยน (inflection point) ที่ทำให้ได้ดุลยภาพระหว่างโมเดลที่ไม่เข้ากันได้กับข้อมูลมากเกินไป และไม่หลวมมากเกินไป โดยต้องมีค่าใช้จ่ายในการฝึกสอนเอไอนั้นให้น้อยที่สุดอีกด้วย

ในตอนนี้อาจเรื่องความจริงฐานนี้ทำให้เราพอจะเห็นภาพความเชื่อมโยงขององค์ความรู้ด้านต่างๆ ในบทความนี้แล้ว บทนำในเรื่องความจริงฐานนี้ถือเป็นแกนกลางที่ร้อยรัดเนื้อหาที่เหลือทั้งหมดของบทความเข้าด้วยกัน ซึ่งเราจะเริ่มที่โครงการของ DARPA ที่ใช้ชื่อว่าความจริงฐานนี้เช่นกันแต่จะมีความหมายที่ต่างกันออกไปและไปจบลงที่โครงการที่ชื่อว่า “สังคมศาสตร์ยุคถัดไป” (Next Generation Social Science หรือ NGS2) ลำดับจากนั้นเราจะเปลี่ยนไปพูดถึงจุดเริ่มต้นของการสร้างองค์ความรู้เรื่องการวิจัยแบบผสมผสานซึ่งในที่สุดจะมีส่วนเชื่อมโยงกับเรื่องการสร้างแบบจำลองในคอมพิวเตอร์ในสังคมศาสตร์ และปิดท้ายด้วยการกล่าวถึงการนำปัญญาประดิษฐ์แบบรู้สร้างมาใช้ในการวิเคราะห์และช่วยในการทำงานวิจัยซึ่งถือว่าเป็นอีกแขนงหนึ่งของคอมพิวเตอร์ในสังคมศาสตร์นี้เช่นเดียวกัน

ความจริงฐานของสำนักงานวิจัยโครงการขั้นสูงด้านกลาโหมสหรัฐอเมริกา

สำนักงานวิจัยโครงการขั้นสูงด้านกลาโหมสหรัฐอเมริกา (The Defense Advanced Research Projects Agency: DARPA) ถูกก่อตั้งขึ้นมาโดยมติของรัฐบาลสหรัฐฯ เมื่อปี ค.ศ. 1958 ในฐานะศูนย์กลางเพื่อการวิจัยและพัฒนาของกระทรวงกลาโหม โดยมีงบประมาณประจำปีราว 3 พันล้านเหรียญสหรัฐ เพื่อเป็นการตอบโต้ต่อความตื่นตระหนกเมื่อสหภาพโซเวียตได้ประสบความสำเร็จในการปล่อยดาวเทียมสปุตนิกขึ้นสู่อวกาศเมื่อวันที่ 4 ตุลาคม ปี ค.ศ. 1957 โดยในระยะเริ่มแรกเป็นข้อเสนอของนีล แม็คเอลรอย (Neil H. McElroy, 1904-1972) รัฐมนตรีว่าการกระทรวงกลาโหมภายใต้รัฐบาลของประธานาธิบดีไอเซนฮาวร์ เมื่อวันที่ 20 พฤศจิกายนในปีเดียวกันนั้นเอง โดยข้อเสนอเริ่มแรกเขาทำในชื่อของสำนักงานวิจัยโครงการขั้นสูง (Advanced Research Project Agency: ARPA) เพียงเท่านั้น (Jacobsen, 2015, pp. 46-51)

โครงการสำคัญชิ้นหนึ่งของสำนักงานฯ ได้แก่ระบบอินเทอร์เน็ต ซึ่งถือเป็นโครงการที่มีต้นกำเนิดมาจากวิสัยทัศน์ของลิกไลเดอร์ (Joseph Carl Robnett Licklider, 1915-1990) โดยแท้ ในช่วงปี ค.ศ. 1962 ลิกไลเดอร์ได้รับการติดต่อจากแจ๊ค รุยนา (Jack Ruina, 1923-2015) ผู้อำนวยการสำนักงานฯ ในขณะนั้นเพื่อเข้ามาทำงานวิจัยเรื่องระบบควบคุมและสั่งการ (Command and Control: C2) สถานการณ์ในช่วงที่ลิกไลเดอร์เข้ามาร่วมงานในตำแหน่งหัวหน้าสำนักเทคนิคประมวลผลข้อมูล (Information Processing Techniques Office: IPTO) ต้องถือว่าเป็นวิกฤตอย่างยิ่งของสหรัฐฯ และทั้งโลกในขณะนั้นด้วย เนื่องจากเหตุการณ์วิกฤตชิปนาอูทิวคิวกาที่เครื่องบินสอดแนม U2 ของสหรัฐฯ สามารถตรวจจับและถ่ายภาพทางอากาศของที่ตั้งชิปนาอูทิวคิวกาที่สหภาพโซเวียตมาติดตั้งที่คิวบาเอาไว้ได้ วิกฤตการณ์ในช่วงนั้นกินเวลาสิบสามวัน ตั้งแต่วันที่ 16-28 ตุลาคมในปีเดียวกัน ประธานาธิบดีเคนเนดีถึงกับประกาศความพร้อมสงครามนิวเคลียร์ในรหัส DEFCON 2 ซึ่งถือเป็นขั้นตอนก่อนสงครามนิวเคลียร์เต็มรูปแบบ (DEFCON 1) เป็นครั้งแรกและครั้งเดียวในประวัติศาสตร์ นอกจากนั้นสิ่งที่เป็นที่รับรู้กันน้อยในสาธารณะกว่านั้นคือในขณะที่กำลังมีการประกาศความพร้อมในขั้นสูงสุดที่เคยมีมานั้นเองสหรัฐฯ ทำการทดลองระเบิดอาวุธนิวเคลียร์ในชั้นบรรยากาศระดับสูงสองลูกภายใต้รหัสเชคเมต (Checkmate) และบลูจิลทริปเปิลไพรม์ (Bluegill Triple Prime) เมื่อวันที่ 20 และ 26 เดือนเดียวกันนั้นเอง การทดลองดังกล่าวในด้านหนึ่งก็เพื่อส่งสัญญาณป้องปรามสหภาพโซเวียตและในอีกด้านหนึ่งก็เพื่อเก็บข้อมูลปฏิกิริยาคริสโตฟิลอส (Christofilos effect) ซึ่งเป็นปฏิกิริยาที่แถบคลื่นแม่เหล็กของโลกทำตัวเป็นตาข่ายดักจับอนุภาคที่มีประจุ (charged particle radiation) จำนวนมหาศาลที่เกิดจากการระเบิดของอาวุธนิวเคลียร์ในชั้นบรรยากาศ และทำให้อนุภาคที่มีประจุเหล่านั้นวิ่งไปนามแนวสนามแม่เหล็กโลกในลักษณะพุ่งปลาวีงทวนกระแสและสร้างแถบคลื่นรังสีเทียม (artificial radiation belt) ซึ่งอาจทำอันตรายให้กับดาวเทียมในวงโคจรได้ อนึ่งโครงการดังกล่าวดำเนินการภายใต้โครงการร่มที่เรียกว่า “ปฏิบัติการอ่างปลา” (Operation Fishbowl) ฝ่ายสหภาพโซเวียตตอบโต้ด้วยการทดลองระเบิดนิวเคลียร์ในชั้นบรรยากาศจำนวนสองลูกเช่นกันในวันที่ 22 และ 28 เดือนเดียวกัน ต่อมาภายหลังวิกฤตนิวเคลียร์คิวบาคคลี่คลายลงจะมีการลงนามระหว่างสองฝ่ายในข้อตกลงสนธิสัญญาว่าด้วยการใช้อวกาศอย่างสันติ (The Outer Space Treaty) ในปีรุ่งขึ้นเพื่อห้ามการทดลองระเบิดนิวเคลียร์ในชั้นบรรยากาศอีก ภายใต้บรรยากาศตึงเครียดของการแข่งขันระหว่างมหาอำนาจเช่นนี้เองลิกไลเดอร์ได้ชี้แนะให้สำนักงานฯ พยายามคิดค้นวิธีการใช้งานคอมพิวเตอร์ที่มากไปกว่างานคำนวณอย่างงานด้านระบบบัญชีหรือระบบเงินเดือนอย่างที่นิยมกันในช่วงนั้น เขาถึงกับเขียนบันทึก “เครือข่ายคอมพิวเตอร์อินเตอร์กาแล็คติก” (Intergalactic Computer Network) เพื่อวางแนวทางในการเชื่อมโยงเครือข่ายคอมพิวเตอร์ระหว่างกัน คำว่า “อินเตอร์กาแล็คติก” ตั้งใจให้หมายถึงการเชื่อมต่อคอมพิวเตอร์ต่างชนิดต่างรูปแบบกัน (2015, pp. 147-151)

แม้ลิกไลเดอร์จะลาออกจากสำนักงานฯ ในปี ค.ศ. 1965 (เขาจะกลับมาทำงานที่สำนักงานฯ ในตำแหน่งเดิมอีกครั้งในระยะเวลาสั้นๆ ในช่วงปี ค.ศ. 1974-1975) แต่วิสัยทัศน์เรื่องระบบเครือข่ายคอมพิวเตอร์เริ่มแจ่มชัดลงไปแล้วแนวคิดระบบควบคุมและสั่งการ (C2) ในไม่ช้าก็เพิ่ม “การสื่อสาร” (Communication) เข้ามาอีกหนึ่งขากลายเป็น C ที่ 3 หรือ C3 (ปัจจุบันมีการปรับเป็น C6ISR คือ Command, Control, Communications, Computers, Cyber-defense, Combat systems, Intelligence, Surveillance, และ Reconnaissance) ในช่วงเวลาใกล้เคียงกันนั้นเองโครงการระบบปฏิบัติการคอมพิวเตอร์อย่างยูนิกซ์ (Unix) ซึ่งได้รับอิทธิพลมาจากมัลติคส์ (Multics) ก็ได้ถือกำเนิดขึ้น

ในปี ค.ศ. 1970 โดยเคน ทอมป์สัน (Ken Thompson, 1943) และเดนนิส ริชชี (Dennis Ritchie, 1941-2011) นักวิจัยของเบลล์แล็บหลังจากที่เบลล์ถอนตัวออกจากการพัฒนามัลติทาสก์ ในช่วงแรกโครงข่าย ARPA จะใช้โปรโตคอล NCP (Network Control Protocol) เพื่อเชื่อมต่อระบบคอมพิวเตอร์ของมหาวิทยาลัยแคลิฟอร์เนีย ลอสแอนเจลิส (UCLA) และสถาบันวิจัยสแตนฟอร์ด (Stanford Research Institute: SRI) เข้าด้วยกัน ในช่วงแรกโปรโตคอล NCP จะมีเป้าหมายเพื่อให้สามารถสื่อสารระหว่างเครือข่ายคอมพิวเตอร์ได้เท่านั้น แต่ในเวลาถัดมาจะมีการพัฒนาโปรโตคอล TCP/IP ขึ้นเพื่อให้มีความสามารถในการเชื่อมต่อแบบกระจายศูนย์ มีการเชื่อมโยงกับเครือข่ายคอมพิวเตอร์ต่างมหาวิทยาลัยมากขึ้นและสร้างผลพลอยได้ในการทำให้โครงข่ายมีความยืดหยุ่นทนทานต่อกรณีระบบเครือข่ายเสียหายจากการโจมตีจากระเบิดนิวเคลียร์ได้ด้วย ระบบเครือข่าย ARPA นี้เองภายหลังจะวิวัฒนาการมาเป็นระบบเครือข่ายอินเทอร์เน็ตที่เราใช้กันในปัจจุบัน และภายหลังจะมีการแยกเครือข่ายคอมพิวเตอร์ของกองทัพออกไปต่างหากในปี ค.ศ. 1983 ในชื่อ MILNET (2015, 244-247)

นอกจากโครงการโครงข่าย ARPA ที่กลายมาเป็นอินเทอร์เน็ตที่รู้จักกันดีแล้ว สำนักงานฯ ยังมีโครงการอื่นอีกมากที่สำคัญที่บทความนี้ให้ความสนใจคือโครงการสังคมศาสตร์ยุคถัดไป (Next Generation Social Science: NGS2) ที่มีเป้าหมายในการกำหนดตัวแปรที่เกี่ยวข้องและกลไกเชื่อมโยงการเป็นเหตุเป็นผลเพื่อจะอธิบายและคาดการณ์การผุดบังเกิดของอัตลักษณ์ร่วม (collective identities) (Asif, 2019, p. 9) โครงการนี้สามารถสืบย้อนกลับไปได้ยังความสนใจของลิลเลเดอร์ในส่วนที่เกี่ยวข้องกับการประยุกต์โครงการพฤติกรรมศาสตร์ (Behavioral Sciences Program) ในปีที่เขาเริ่มทำงานกับสำนักงานฯ นั่นเอง โดยในเมื่อไทยมีการสำรวจข้อมูลของทหารบก ทหารเรือ และทหารอากาศจำนวน 2,950 นาย การเก็บข้อมูลมีมากกว่าเรื่องข้อมูลกายภาพส่วนตัวแต่ยังรวมไปถึงความเชื่อทางศาสนา ข้อมูลเชื้อสายบรรพบุรุษ ไปจนถึงความคิดความเชื่อต่อสถาบันพระมหากษัตริย์ภายใต้โครงการ “การสำรวจมานุษยมิติของกองทัพไทย” (Anthropometric Survey of The Royal Thai Armed Forces) ในกรณีหากเมืองไทยตกอยู่ในสภาวะสงคราม สำนักงานฯ จะมีข้อมูลของทหารไทยพร้อมและสามารถใช้คอมพิวเตอร์คาดการณ์ได้ว่าใครจะขึ้นมาเป็นผู้บังคับบัญชาเป็นต้น (Jacobsen, 2015, pp. 153-155) โครงการ NGS2 อยู่ภายใต้การนำของอดัมรัสเซลล์ (Adam Russell) นอกจากโครงการดังกล่าวเขายังริเริ่มโครงการ Forensic Social Science Supercollider (FS3) เพื่อทดสอบความเที่ยงตรงของการอนุมานเกี่ยวกับพฤติกรรมด้านสังคมของมนุษย์อย่างไม่เคยทำได้มาก่อนเลย และโครงการความจริงฐาน (Ground Truth: GT) ซึ่งมีส่วนคาบเกี่ยวกับทั้ง NGS2 และ FS3 แต่โครงการนี้นอกจากจะทำความเข้าใจเกี่ยวกับปัจจัยต่างๆ ที่เกี่ยวข้องกับพฤติกรรมของมนุษย์ยังสามารถสร้างแบบจำลองเพื่อจำลองสถานการณ์ได้อีกด้วย (González, 2022)

6 | DARPA Ground Truth: TAI Simulation Requirements

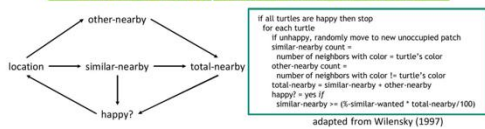
1. **Simulation accessibility:** Can the simulations handle social science data collection methods?

Data Collection Methods

Observational data	Event journals
Interviews	Passive data collection
Surveys	Randomized trial
Ethnographic observations	Experiments
Laboratory experiments	Proxy experiments...

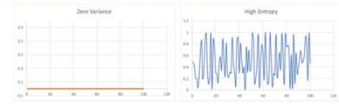
2. **Verifiability of ground truth:** Does the ground truth accurately represent the simulation?

Ground Truth Represents Causal Structure



3. **Plausibility:** Is the simulation a self sustaining virtual world?

Simulation-Driven "Interesting" Behavior



4. **Complexity:** How complex is the simulation?

Multiple Dimensions of Complexity



5. **Flexibility:** Can the TAI team manipulate complexity?

Ground Truth Represents Causal Structure



ภาพที่ 1 ข้อมูลโครงการ Ground Truth ของ DARPA

ที่สำคัญคือบทบาทในการลงทุนของภาครัฐที่เกี่ยวข้องกับโครงการพัฒนาและวิจัยโดยเฉพาะอย่างยิ่งการลงทุนในการพัฒนาและวิจัยจากสำนักงานฯ ทั้งในเรื่องของ ARPANET ที่ภายหลังวิวัฒนาการมาเป็นระบบอินเทอร์เน็ตก็กลายมาเป็นฐานให้กับผลิตภัณฑ์ที่ต้องถือว่าเปลี่ยนแปลงโครงสร้างสนามแข่งขันตลาดโทรศัพท์มือถืออย่างไอโฟน หรือระบบนำร่อง GPS ในทุกวันนี้ก็มีที่มาจากโครงการนำร่องทางทหารที่ชื่อ NAVSTAR ส่วนระบบจอสัมผัสของไอโฟนก็มีรากฐานมาจากโครงการที่มีส่วนเกี่ยวข้องกับงานวิจัยที่ได้รับทุนสนับสนุนจากมูลนิธิวิทยาศาสตร์แห่งชาติและสำนักงานข่าวกรองกลาง และเช่นเดียวกับระบบผู้เชี่ยวชาญอัจฉริยะ SIRI ที่ใช้งานในไอโฟนครั้งหนึ่งก็เคยเป็นโครงการระบบปัญญาประดิษฐ์ที่ถูกพัฒนาในสำนักงานฯ มาก่อนเช่นเดียวกัน (Mazzucato, 2015, p. 6) ดังนั้นจึงอาจกล่าวได้ว่าปัจจัยการลงทุนในการพัฒนาและวิจัยจากภาครัฐโดยเฉพาะอย่างยิ่งด้านกลาโหมดังที่ได้อภิปรายความเป็นมาของโครงการ ARPANET จะมีส่วนอย่างมากในการส่งเสริมการสร้างนวัตกรรมในภาคเอกชนในภายหลังและนี่เป็นปัจจัยขับเคลื่อนหลักอันหนึ่งที่สร้างความเปลี่ยนแปลงในสังคม เราอาจจะมองได้ว่าโครงการที่ทางสำนักงานฯ ได้ลงทุนอย่างต่อเนื่องและยาวนานมักมีโอกาสในการถูกขับเคลื่อนไปสู่นวัตกรรมภาคเอกชนและกลายเป็นกระแสหลักในระยะต่อมา ในส่วนถัดไปของบทความเราจะพุดถึงปัจจัยขับเคลื่อนจากแวดวงที่สำคัญอีกส่วนหนึ่งคือแวดวงวิชาการ

กระบวนการวิจัยแบบผสมผสาน

ในขณะที่ฟากกระทรวงกลาโหมอย่างสำนักงานวิจัยโครงการขั้นสูงด้านกลาโหมสหรัฐอเมริกาซึ่งรับผิดชอบโครงการวิจัยและพัฒนาขั้นแนวหน้าหลายโครงการ และต่อมานำมาสู่พัฒนาการที่ส่งผลกระทบต่อการศึกษาและการสร้างแบบจำลองทางสังคมศาสตร์ซึ่งยืนอยู่บนพื้นฐานของพฤติกรรมศาสตร์นั่นเอง แต่ต้องถือว่าโครงการและเทคนิคเหล่านี้ยังคงอยู่ในระเบียบวิจัยเชิงปริมาณและมีแนวโน้มในการใช้คอมพิวเตอร์มาช่วย

สนับสนุนการวิจัยตามวิสัยทัศน์ของลิลิเดออร์นับแต่ต้น ทางภาคพลเรือนในแวดวงวิชาการได้เริ่มมีข้อถกเถียงถึงระเบียบวิธีวิจัยที่ใช้ในศาสตร์ด้านสังคมศาสตร์โดยเฉพาะอย่างยิ่งระเบียบวิจัยเชิงคุณภาพ ในหนังสือที่ชื่อว่า “การออกแบบงานวิจัยด้านสังคม: การอนุมานที่เป็นวิทยาศาสตร์ในงานวิจัยเชิงคุณภาพ” (Designing Social Inquiry: Scientific Inference in Qualitative Research) ของแกรี่ คิง (Gary King, 1958) โรเบิร์ต คีโฮฮาน (Robert Keohane, 1941) และซิดนีย์ เวอร์บา (Sidney Verba, 1932-2019) หนังสือเล่มนี้มักถูกอ้างอิงตามอักษรตัวแรกของนามสกุลของผู้แต่งทั้งสามว่า “KKV” ในหนังสือเล่มนี้ผู้แต่งทั้งสามได้เสนอว่างานวิจัยที่นำเสนอทฤษฎีที่มีความเป็นเหตุเป็นผลที่ดีจำเป็นจะต้องสามารถถูกหักล้างได้ มีความคงเส้นคงวาของตรรกะที่ใช้ภายในตัวงานโดยไม่สร้างสมมติฐานที่ขัดแย้งกันเอง ตัวแปรตามที่ต้องการอธิบายได้รับอิทธิพลจากกลุ่มตัวแปรอิสระซึ่งเป็นตัวแปรภายในระบบที่สังเกต (endogenous) มากกว่าจะเป็นตัวแปรได้รับอิทธิพลจากภายนอก (exogenous) กรอบทฤษฎีจะต้องมีแนวคิดที่ชัดเจนและต้องสามารถวิเคราะห์แยกแยะได้ งานวิจัยที่สร้างทฤษฎีขึ้นมานั้นจะต้องมีอำนาจอธิบายที่แข็งแกร่ง (leverage) กล่าวคือทฤษฎีนั้นมีความเรียบง่ายแต่ภายใต้ความเรียบง่ายนั้นสามารถใช้อธิบายปรากฏการณ์ได้อย่างกว้างขวาง ซึ่งผู้วิจัยทั้งสามเห็นว่างานวิจัยเชิงคุณภาพในช่วงเวลานั้นไม่สามารถให้การอนุมานเพื่อหาสาเหตุ (causal inference) ที่แข็งแกร่งได้ เนื่องจากนักวิจัยเชิงคุณภาพต้องประสบปัญหาในการพิจารณาว่าตัวแปรใดกันแน่ที่มีผลต่อตัวแปรตามที่กำลังพิจารณาอยู่อย่างแท้จริง ทั้งนี้เป็นเพราะส่วนใหญ่่งานวิจัยเชิงคุณภาพขาดจำนวนตัวอย่างที่ใช้สังเกตการณ์มากเพียงพอที่จะใช้ประเมินตัวแปรอิสระที่กำลังศึกษา แต่จากการสร้างโมเดลสมการถดถอยโดยทั่วไปเราสามารถอนุมานได้ว่าการเพิ่มจำนวนตัวแปรที่ศึกษาในขณะที่จำกัดตัวอย่างในการสังเกตการณ์จะทำให้ระดับแห่งความเป็นอิสระ (Degrees of Freedom: df) ลดลงไปตามสมการ $df = \text{จำนวนตัวอย่าง} - \text{จำนวนตัวแปรที่ศึกษา} - 1$ ทำให้กรอบทฤษฎีที่ได้เข้ากับตัวอย่างมากเกินไป (overfitting) ขาดการใช้ได้เป็นการทั่วไปและดังนั้นจึงมีความน่าเชื่อถือน้อย (King, Keohane & Verba, 1994)

งาน KKV นี้สร้างแรงกระเพื่อมในแวดวงวิจัยเชิงคุณภาพอย่างยิ่งจนกระทั่งก่อให้เกิดการโต้กลับอย่างกว้างขวาง งานชิ้นสำคัญเช่น “หลัง KKV: วิธีวิทยาแบบใหม่ในการทำวิจัยเชิงคุณภาพ” (After KKV: The New Methodology of Qualitative Research) ความจริงต้องบอกว่างาน “หลัง KKV” นี้เป็นการทบทวนวรรณกรรมที่โต้กลับงาน KKV จะตรงกว่า อาทิ ได้มีการยกตัวอย่างจากเฮนรี แบริดจ์ (Henry E Brady) ลาร์รี บาร์เทลส์ (Larry Bartels) และเดวิด ฟรีดแมน (David Freedman) ว่า KKV พุดเกินเลยไปมากเรื่องความแข็งแกร่งของงานวิจัยเชิงปริมาณ (จะเห็นว่ามุมมองวิจารณ์ของ KKV มองออกมาจากจุดยืนของงานวิจัยเชิงปริมาณ) และลดทอนคุณค่าและประโยชน์ที่ได้รับจากงานวิจัยเชิงคุณภาพ นอกจากนี้ “หลัง KKV” ยังได้ชี้ให้เห็นว่าฝ่ายวิจัยเชิงคุณภาพได้ข้ามไปพ้นจากสิ่งที่ KKV ได้เสนอมาไปสู่แนวทางใหม่ อาทิ กระบวนการติดตามและการสังเกตกระบวนการความเป็นเหตุเป็นผล วิธีวิทยาเรื่องทฤษฎีเช็ดและตรรกะ รวมไปถึงกลยุทธ์ที่ใช้ในการประสานกันระหว่างการวิจัยเชิงคุณภาพและการวิจัยเชิงปริมาณเข้าด้วยกัน (Mahoney, 2010, pp. 122-123) “หลัง KKV” ชี้ให้เห็นว่าการวิจัยเชิงปริมาณแท้จริงมีลักษณะเป็นการสังเกตที่มาจากชุดข้อมูล (Data-Set Observations: DSOs) ในขณะที่การวิจัยเชิงคุณภาพมีลักษณะเป็นการสังเกตกระบวนการที่เป็นเหตุเป็นผล (Causal-Process Observations: CPOs) กระบวนการอย่างแรกจะนำไปสู่การทดสอบสมมติฐานหรือทดสอบทฤษฎี ในขณะที่กระบวนการอย่างหลังจะนำไปสู่ทั้งการสร้างทฤษฎี (Theory Development) และทดสอบทฤษฎี (Theory Testing) ในคราวเดียวกันด้วย ด้านการทดสอบทฤษฎีนั้นทำได้ทั้งผ่านการ

ทดสอบการสังเกตกระบวนการที่เป็นเหตุเป็นผลจากตัวแปรตาม (Independent Variable CPOs) หรือจากการสังเกตกลไก (Mechanism CPOs) หรือแม้กระทั่งผลข้างเคียง (Auxiliary Outcome CPOs) (2010, pp. 124-128)

ในกรณีการการสังเกตกระบวนการที่เป็นเหตุเป็นผลจากตัวแปรตามใช้ในงานศึกษาเรื่องความสัมพันธ์ต่างประเทศที่ไม่มีการใช้อาวุธนิวเคลียร์เพราะเป็นเรื่องต้องห้าม (taboo issue) ของนินา แทนเนนวัลด์ (Nina Tannenwald) ฝ่าย DSOs อาจมองว่าจำนวนตัวอย่างที่ใช้ในการศึกษาน้อยไป คือ 4 กรณีศึกษา แต่ฝ่าย CPOs จะโต้แย้งว่าปริมาณข้อโต้แย้งในเชิงนโยบายที่เกิดขึ้นทั้งหมดในห้วงเวลาดังกล่าวอ้างถึง “เรื่องต้องห้าม” ในการใช้อาวุธนิวเคลียร์นั้นเพียงพอแล้วในการพิจารณาว่า “เรื่องต้องห้ามอาวุธนิวเคลียร์” ดังกล่าวเป็นตัวแปรอิสระที่ไปกำหนดตัวแปรตาม คือการกำหนดนโยบายต่างประเทศว่าจะใช้อาวุธนิวเคลียร์หรือไม่ อย่างไรก็ตามวิธีการใช้ CPOs ในกรณีตัวแปรตามนี้ ต้องระมัดระวังเรื่องการอคติในการคัดเลือกตัวแปรอิสระด้วย (2010, pp. 125-128)

ส่วนกรณีที่สองคือการสังเกตกระบวนการที่เป็นเหตุเป็นผลจากกลไกนั้น ฝ่าย CPOs จะใช้เทคนิคการวิเคราะห์แบบเบย์ (Bayesian analysis) เข้ามาจับในกรณีที่มีผลลัพธ์ที่มีลักษณะสุดขั้วเกิดขึ้นในการสังเกตการณ์ แม้จะมีเพียงจำนวนน้อยแต่อาจนำไปสู่การสร้างแบบแผนทางทฤษฎีในรูปแบบใหม่ได้ กรณีนี้ในฝ่าย DSOs มักจะนิยมคัดทิ้งเพราะถือว่าเป็นค่ารบกวนมากกว่าจะมีนัยยะที่สำคัญ ข้อควรระวังของฝ่าย CPOs ที่ใช้วิธีก็เช่นเดียวกับกรณีแรกคือต้องระมัดระวังอคติที่กำหนดแบบแผนทางทฤษฎีให้ไปขึ้นกับตัวแปรน้อยชุดที่มีลักษณะสุดขั้วจนเกินไป (2010, pp. 128-129)

สำหรับกรณีสุดท้ายคือการสังเกตกระบวนการที่เป็นเหตุเป็นผลจากผลข้างเคียง ในกรณีนี้ฝ่าย CPOs จะใช้ผลที่เกิดขึ้นข้างเคียงมาเป็นเหตุสนับสนุนทฤษฎีแม้ว่าส่วนของผลที่เกิดขึ้นไม่ได้ชัดเจนลงไปด้านหนึ่งด้านใดก็ตาม ทั้งนี้ฝ่าย CPOs ให้เหตุผลว่าผลข้างเคียงที่เกิดขึ้นอาจเกิดจากกลไกที่เราสงสัยค้นได้ไม่แน่ชัด (2010, pp. 129-131)

อย่างไรก็ตาม KKV อ้างตนว่าเป็นความพยายามในการเชื่อมโยงวิธีวิทยาในการวิจัยเชิงปริมาณและเชิงคุณภาพเข้าด้วยกันโดยการสร้างตรรกะผนวกรวมเพื่อช่วยให้อนุมานได้ทั้งสองวิธีวิทยา (Mahoney, 2010, p. 139; King, Keohane & Verba, 1994, p. ix) แต่ “หลัง KKV” ก็ยังเน้นให้เห็นว่าข้อเสนอของ KKV มาจากจุดยืนด้านวิธีวิทยาในการวิจัยเชิงปริมาณโดยเฉพาะอย่างยิ่งกรอบวิเคราะห์แบบสมการถดถอย (regression-oriented analysis) ดังนั้นข้อเสนอของ KKV แท้จริงแล้วเป็นการทำให้วิธีวิทยาการวิจัยเชิงคุณภาพต้องเดิมตามกรอบการวิจัยเชิงปริมาณโดยขาดข้อเสนอที่จะรวมความแตกต่างของวิธีวิทยาทั้งสองเข้าด้วยกันอย่างแท้จริง “หลัง KKV” ไม่ได้เสนอข้อสังเคราะห์ที่ลึกกว่า KKV แต่เสนอการเชื่อมวิธีวิทยาในการวิจัยที่เป็นกลางกว่า KKV ไม่ว่าจะเป็นการใช้วิธีวิทยาเชิงปริมาณเป็นส่วนเสริมของวิธีวิทยาเชิงคุณภาพ หรือกลับกันคือใช้วิธีวิทยาเชิงคุณภาพเป็นส่วนเสริมวิธีวิทยาเชิงปริมาณก็สามารถทำได้เช่นกัน (2010, pp. 139-143)

ข้อเสนอที่รวมวิธีวิทยาที่แตกต่างกันทั้งฟากการวิจัยเชิงปริมาณและเชิงคุณภาพนั้นต้องทำในระดับปรัชญา และเป็นข้อเสนอของเดวิด มอร์แกน (David Morgan) ในวารสารการวิจัยแบบผสมผสานโดยมอร์แกนได้เสนอว่าเราควรมองเรื่องวิธีวิทยาการวิจัยในสังคมศาสตร์ให้มากกว่าเฉพาะเทคนิควิธีจากมุมมองเรื่องกระบวนการทัศน์ของโทมัส คูห์น โดยเฉพาะนิยามเรื่อง “ความเชื่อร่วมกัน” (shared beliefs) และเขาได้เสนอว่าด้วยกรอบนี้กระบวนการทัศน์ทางสังคมศาสตร์ได้ย้ายจากกระบวนการทัศน์แบบปฏิฐานนิยม (Positivism) ซึ่งเป็น

รากฐานของการวิจัยเชิงปริมาณไปสู่กระบวนทัศน์แบบอภิปรัชญา (Metaphysics) ซึ่งวางฐานอยู่บนการขับเคลื่อนทางภววิทยาที่เน้นเรื่องประสบการณ์ซึ่งเป็นรากฐานของการวิจัยเชิงคุณภาพ แต่นี้ไม่ใช่ว่าการทิ้งกระบวนทัศน์แบบปฏิฐานนิยมเสียทีเดียว กระบวนทัศน์แบบอภิปรัชญาเพียงเสนอว่ากระบวนทัศน์แบบปฏิฐานนิยมเป็นเพียงทางเลือกหนึ่งจากหลายทางเลือกที่เป็นไปได้เท่านั้น (Morgan, 2007, p. 64) มอร์แกนได้เสนอปรัชญาแฟร์กมาติซึมเพื่อท้าทายโดยตรงต่อกระบวนทัศน์แบบอภิปรัชญาโดยตรง โดยเน้นน้ำหนักไปที่แนวคิด “แนวทางการกระทำ” (line of actions) ของ วิลเลียม เจมส์ (William James) และจอร์จ เฮอร์เบิร์ต มีด (George Herbert Mead) รวมไปถึงแนวคิด “การยืนยันที่มีเหตุผล” (warranted assertions) ของจอห์น ดิวอี้ (John Dewey) ไปจนถึงแนวคิดเรื่อง “การใช้งานได้จริง” (workability) ของทั้งเจมส์และดิวอี้ ทั้งนี้มอร์แกนเห็นว่านิยามเรื่อง “ความเชื่อร่วมกัน” (shared beliefs) ของชุมชนวิชาการในแนวคิดเรื่องกระบวนทัศน์ของคูห์นนั้นทำให้เขาเชื่อมั่นว่าจะสามารถไปให้พ้นจากจุดยืนเรื่อง “ญาณวิทยา” หรือทฤษฎีความรู้ (epistemology) ที่ส่งผลในที่สุดทำให้เกิดความเห็นที่ไม่เป็นอันหนึ่งอันเดียวกัน (incommensurability) ซึ่งเป็นแกนหลักของกระบวนทัศน์แบบอภิปรัชญาได้ ทั้งนี้ตามความเห็นของปรัชญาแฟร์กมาติซึมเห็นว่าพหุมีจุดร่วมกันได้ระหว่างนักวิจัยหรือปรัชญาต่างสาขากันที่ “ความหมายร่วม” (joint action) และ “การกระทำร่วมกัน” (joint action) ดังนั้นจึงเป็นการย้ายจุดศูนย์กลางจากเรื่อง “ญาณวิทยา” มาสู่ “วิธีวิทยา” (methodology) เพราะสุดท้ายปัญหาทางสังคมศาสตร์ไม่ได้มุ่งจุดเน้นที่ข้อถกเถียงทางปรัชญาที่ตามธรรมชาติแล้วยากจะลงรอยกันแบบเดียวกับนักปรัชญา ทำให้เกิดข้อแตกต่างของวิธีวิทยาวิจัยทั้งสามแบบคือ ในขณะที่วิธีวิทยาวิจัยเชิงคุณภาพเชื่อในเรื่อง “การอนุมาน” (induction) เน้นหนักเรื่องอัตวิสัยและบริบท ส่วนวิธีวิทยาวิจัยเชิงปริมาณจะเชื่อเรื่อง “การนิรนัย” (deduction) เน้นหนักเรื่องภววิสัยและกฎเกณฑ์ทั่วไป แต่วิธีวิทยาแบบผสมผสานจะเชื่อเรื่อง “การสืบสวนหาสาเหตุ” (abduction) เน้นหนักเรื่องอัตวิสัยที่เชื่อมโยงต่อกัน (intersubjectivity) และความสามารถในการถ่ายโอนผลการวิจัยข้ามไปมาได้ (transferability) (2007, pp. 64-73)

อนึ่ง ในระยะหลังเมื่อวิธีวิทยาการวิจัยแบบผสมผสานได้รับความนิยมแพร่หลายมากขึ้นโดยเฉพาะวิทยาศาสตร์สุขภาพ จึงพยายามจัดทำมาตรฐานโดยเฉพาะเครื่องมือแบบประเมินวิธีวิทยาการวิจัยแบบผสมผสาน (Mixed Methods Appraisal Tool: MMAT) และมีการกำหนดแนวทางการออกแบบการวิจัยแบบผสมผสานเอาไว้ด้วย (Hong et al., 2018, p. 6) เราจะเห็นได้ว่าในแวดวงวิชาการด้านสังคมศาสตร์จะเริ่มมีการเปิดกว้างในการวิจัยจากช่วงเริ่มแรกที่เน้นหนักไปที่วิธีวิทยาการวิจัยเชิงปริมาณซึ่งมีฐานกระบวนทัศน์แบบปฏิฐานนิยมแล้วเริ่มคลี่คลายมาเป็นการวิจัยเชิงคุณภาพที่มีฐานกระบวนทัศน์อภิปรัชญาซึ่งก็เท่ากับเปิดกว้างให้กับแนวทางการวิจัยที่หลากหลายมากขึ้น ทั้งนี้การเน้นความสำคัญไปยังเรื่องญาณวิทยาจึงยังมีข้อจำกัดที่ไม่สามารถหาข้อตกลงร่วมระหว่างกระบวนทัศน์ที่ขัดแย้งกันได้ มาจนถึงการถือกำเนิดขึ้นของการวิจัยแบบผสมผสานที่มีฐานกระบวนทัศน์อยู่บนปรัชญาแฟร์กมาติซึมซึ่งสามารถหาพื้นที่กลางระหว่างกระบวนทัศน์หรือความเชื่อแบบต่างๆ ได้ (ในกรณีนี้ผู้เขียนเสนอว่านี่คืออีกความหมายหนึ่งของ “ความจริงฐาน” ที่เราได้พูดถึงในตอนก่อนหน้านี้) ไปจนถึงมาตรฐานของกระบวนการวิจัยแบบผสมผสานที่เริ่มก่อตัวขึ้น นอกจากนี้ในระยะหลังจากการแพร่หลายของการใช้คอมพิวเตอร์และระบบอินเทอร์เน็ตซึ่งสามารถผลิตข้อมูลที่เกิดจากพฤติกรรมออนไลน์ผ่านสื่อสังคมของชุมชนอินเทอร์เน็ตปริมาณมากอย่างต่อเนื่องซึ่งถือว่าเป็นข้อมูลขนาดใหญ่ (Big Data) จึงเริ่มมีการใช้คอมพิวเตอร์ในทางสังคมศาสตร์ (Computational Social Science: CSS)

การประยุกต์ใช้คอมพิวเตอร์นี้เป็นไปมากกว่าการเก็บข้อมูลและสังเคราะห์ข้อมูลเหมือนในยุคแรก แต่มีการประยุกต์ใช้ได้หลากหลายมากขึ้น ในกรณีนี้มีการเสนอให้การใช้คอมพิวเตอร์ในทางสังคมศาสตร์เป็นส่วนขยายเพิ่มเติมของวิธีวิจัยแบบผสมผสานด้วยเช่นกัน (Rodriguez and Storer, 2019, p. 5) ด้านเคลาดิโอ ซิโอฟฟี-เรวิลลา (Claudio Cioffi-Revilla) ได้กำหนดคอมพิวเตอร์ในทางสังคมศาสตร์เอาไว้ว่าประกอบไปด้วยลักษณะต่างๆ 5 ภาคใหญ่ กล่าวคือ ภาคที่หนึ่งเป็นองค์ประกอบต่างๆ ทางกายภาพเช่น ตัวคอมพิวเตอร์ หน่วยความจำ ฮาร์ดแวร์ และอื่นๆ ภาคที่สองเป็นระบบหรืออัลกอริทึมสกัดข้อมูลอัตโนมัติ ภาคที่สามเป็นระบบเครือข่ายสังคม ภาคที่สี่เป็นความซับซ้อนทางสังคม และภาคที่ห้าเป็นการจำลองทางสังคมศาสตร์ด้วยคอมพิวเตอร์ซึ่งในส่วนนี้เราจะเห็นว่าละม้ายกับโครงการ NGS2 ที่เราได้กล่าวถึงไปข้างต้นแล้ว ทั้งนี้จะเห็นว่าสองภาคนี้เป็นส่วนที่ผูกองค์ความรู้ด้านวิทยาการคอมพิวเตอร์เข้ากับองค์ความรู้ทางสังคมศาสตร์ (Cioffi-Revilla, 2022, pp. 22-25) บทความของซิโอฟฟี-เรวิลลาถูกตีพิมพ์ในหนังสือ “คู่มือคอมพิวเตอร์ในสังคมศาสตร์” (Handbook of Computational Social Science) เล่มแรก จึงต้องถือว่าเป็นหมุดหมายสำคัญที่ทำให้รากฐานเรื่องคอมพิวเตอร์ในสังคมศาสตร์ในโลกตะวันตกได้ลงหลักปักฐานในสังคมศาสตร์และวิทยาการคอมพิวเตอร์ไปพร้อมกันอย่างมั่นคงจากที่ก่อนหน้านี้ยังมีการกำหนดอัตลักษณ์ของตนเองไม่แจ่มชัดนัก ในส่วนสุดท้ายของบทความนี้จะกล่าวถึง การประยุกต์ใช้ “ปัญญาประดิษฐ์แบบรู้สร้าง” (Generative AI: Gen AI) โดยเฉพาะโมเดลภาษาขนาดใหญ่ (Large Language Model: LLM) ซึ่งเป็นเทคโนโลยีคอมพิวเตอร์ที่มีการพัฒนาต่อยอดจากเทคโนโลยีการเรียนรู้ของเครื่องจักร (Machine Learning: ML) การเรียนรู้เชิงลึกของเครื่องจักร (Deep Learning: DL) และระบบเครือข่ายประสาทเทียม (Artificial Neuron Network: ANN) ซึ่งทั้งสามเรื่องนี้มีการกล่าวถึงในหนังสือ “คู่มือคอมพิวเตอร์ในสังคมศาสตร์” เล่มที่สองอยู่ด้วยเช่นกัน (โปรดดูใน De Veaux and Eck, 2022) แต่ยังไม่ได้ลงลึกถึงเรื่องปัญญาประดิษฐ์แบบรู้สร้างเนื่องจากเป็นเทคโนโลยีที่เพิ่งมีการเปิดตัวเมื่อปลายปี ค.ศ. 2022 ซึ่งเป็นช่วงหลังจากที่หนังสือทั้งสองเล่มตีพิมพ์ออกมาแล้ว แต่อย่างไรก็ตามเนื่องจากการประสบความสำเร็จอย่างมากของปัญญาประดิษฐ์แบบรู้สร้างและความสามารถในการสร้างข้อความและรูปภาพได้ในระดับที่แทบไม่แตกต่างจากคนทั่วไปดังนั้นโดยตัวมันเองก็มีผลต่อการเปลี่ยนแปลงเชิงสังคม และมีศักยภาพในการเข้ามาช่วยงานด้านงานวิเคราะห์และวิจัยด้านสังคมศาสตร์ด้วยเช่นกัน การเปลี่ยนแปลงอย่างถึงรากถึงโคนซึ่งเป็นผลที่สะสมมายาวนานเหล่านี้ จะมารวมศูนย์อยู่ที่การเปิดตัวปัญญาประดิษฐ์แบบรู้สร้างของภาคเอกชนในตอนถัดไป

ปัญญาประดิษฐ์แบบรู้สร้าง

เรากำลังอยู่ในยุคที่น่าตื่นตาตื่นใจในประวัติศาสตร์ เมื่อบริษัทโอเพนเอไอ (OpenAI) ได้เปิดตัวแชทบอทระดับโลกที่มีชื่อว่า “แชทจีพีที” (ChatGPT) เมื่อวันที่ 30 พฤศจิกายน ค.ศ. 2022 แล้วนับจากวันนั้นเป็นต้นมา ปริมาณผู้ใช้งานแชทจีพีทีได้เพิ่มขึ้นถึง 1 ล้านคนโดยใช้เวลาเพียง 5 วันเท่านั้นถือว่าเป็นเวลาที่รวดเร็วมากกันอย่างมากมายเมื่อลองเทียบสถิติในการสร้างปริมาณผู้ใช้งานในระดับนี้กับเฟสบุ๊คซึ่งใช้เวลา 10 เดือน ทวิตเตอร์ใช้เวลา 2 ปี และเน็ตฟลิกซ์ใช้เวลาถึง 3 ปีครึ่งกว่าจะบรรลุยอดผู้ใช้เท่าแชทจีพีที ที่น่าสนใจกว่านั้นคือแชทจีพีทียังบรรลุยอดผู้ใช้งาน 100 ล้านคนในหนึ่งเดือน และมีผู้เข้ามาเยี่ยมชมเว็บไซต์แชทจีพีทีถึง 1 พันล้านครั้งภายในสองเดือนเท่านั้น ด้วยความสำเร็จอย่างถล่มทลายแบบไม่มีใครคาดคิดมาก่อนแบบนี้ทำให้โอเพนเอไอเร่งเปิดตัว “เครื่องยนต์” ของแชทจีพีทีรุ่นที่ 4 (GPT-4) ในเดือนถัดมา เรื่องนี้ยังส่งผลให้บริษัทไมโครซอฟท์ซึ่ง

เป็นยักษ์ใหญ่ในวงการไอทีที่เพิ่มการลงทุนในโอเพนเอไอในระดับหนึ่งหมื่นล้านเหรียญสหรัฐ โดยจะขยายการลงทุนอย่างต่อเนื่องเป็นเวลาหลายปี ไม่เพียงเท่านั้นสตาร์ทอัพ นาเดลลา ซีอีโอของไมโครซอฟท์ยังข่มขู่แข่งอย่างกูเกิลว่าไมโครซอฟท์จะติดตั้งแชทจีพีทีที่เข้ากับเครื่องมือการค้นหา (Search Engine) ของตนเองที่ชื่อว่า “บิง” (Bing) และพร้อมจะติดตั้งแชทจีพีทีที่เข้ากับผลิตภัณฑ์หลักของตนทั้งหมดไม่ว่าจะเป็น “โคไพโลต” (Copilot) เพื่อช่วยให้โปรแกรมเมอร์เขียนโค้ดได้อย่างสะดวกง่ายดายขึ้น หรือจะเป็นการผสานงานร่วมกันระหว่างแชทจีพีทีและออฟฟิศ 365 ซึ่งสตาร์ทอัพยังระบุว่าการแข่งขันครั้งนี้ไมโครซอฟท์จะได้เปรียบเพราะต่อไปเมื่อบิงของไมโครซอฟท์มีผู้ใช้งานมากขึ้นมันจะเบียดตลาดเครื่องมือการค้นหาของกูเกิล เพราะบิงไม่ต้องพึ่งพาการหารายได้จากการโฆษณาซึ่งเป็นรายได้หลักของกูเกิล เรื่องนี้ส่งผลอย่างสำคัญต่อยุทธศาสตร์การแข่งขันเพราะสตาร์ทอัพกำลังบอกว่าจะให้อูเกิลปรับแผนมาใช้เอไอติดตั้งกับเครื่องมือการค้นหาเหมือนกับไมโครซอฟท์แต่คนก็จะไม่สนใจไปหาข้อมูลในเว็บไซต์อีกต่อไป เพราะได้รับข้อมูลจากเอไอไปเรียบร้อยแล้วนั่นเอง จึงถือเป็นการทำลายแหล่งรายได้หลักของกูเกิลลงไปด้วย ปัจจุบันมีการเปิดตัวโครงการปัญญาประดิษฐ์สร้างแบบโมเดลภาษาขนาดใหญ่ขึ้นมาแข่งขันกันมากมาย งานของ Zhao et al. (2023) ต้องถือว่าการสำรวจโมเดลภาษาขนาดใหญ่ได้อย่างกว้างขวางที่สุดและมีการปรับปรุงข้อมูลของวรรณกรรมที่เกี่ยวข้องลงบนเว็บไซต์ GitHub อีกด้วยจึงถือว่ามีประโยชน์ต่อการติดตามสถานการณ์ภาพรวมในอุตสาหกรรมอย่างยิ่ง

ในการให้สัมภาษณ์รายการ จีส์ตีมินิสต์ (หรือ 60 นาที) ซึ่งต้องถือว่าเป็นรายการที่มีชื่อเสียงมากในสหรัฐฯ กูเกิลใช้โอกาสนี้ในการประชาสัมพันธ์แผนการเอไอของกูเกิลอย่างรอบด้านและครบเครื่อง เรื่องที่น่าสนใจคือซุนดรา พิชัย (Sundar Pichai, 1972) และทีมบริหารของเขายังได้พูดถึงศัพท์สำคัญสองคำคือคำว่า “การผุดบังเกิด” (emergence) หมายถึงความสามารถที่ผุดขึ้นมาใหม่ๆ โดยที่เราไม่ได้คาดคิดมาก่อนของเอไอที่มีพารามิเตอร์ขนาดใหญ่เมื่อมันถูกฝึกฝนด้วยข้อมูลมหาศาล และคำว่า “การหลอน” (hallucination) ของเอไอ เมื่อบางครั้งมันอาจให้คำตอบที่ไม่ถูกต้องเพียงตรงแต่สามารถให้คำตอบกับผู้ใช้ได้อย่างมั่นใจ ซึ่งพิชัยยังบอกว่าในวงการยังไม่พบวิธีการแก้ปัญหาเรื่องนี้ได้อย่างเบ็ดเสร็จ

ประวัติของเอไอและแชทจีพีที

จุดกำเนิดขององค์ความรู้ด้านปัญญาประดิษฐ์ นอกจากจะถือกันว่ายู่ที่บทความ “Computing Machinery and Intelligence” ของอลัน ทัวริง (Alan Turing, 1912-1954) ที่ตีพิมพ์ในปี ค.ศ. 1950 ซึ่งบทความนี้นำเสนอแนวคิดที่มีชื่อเสียงของเขาในการทดสอบ (เรียกว่าการทดสอบของทัวริง) เพื่อพิสูจน์ว่าเครื่องจักรนั้นจะมีความชาญฉลาดเท่ามนุษย์หรือไม่ (Turing, 1950) นอกจากนั้นยังมีบทความ “A Logical Calculus of the Ideas Immanent in Nervous Activity” ของวอร์เรน แม็คคูลอค (Warren McCulloch, 1898-1969) และวอลเตอร์ พิตต์ส (Walter Pitts, 1923-1969) ซึ่งได้นำเสนอแนวคิดที่ว่าเมื่อจัดวางใยประสาทเทียมได้อย่างเหมาะสม มันจะทำให้คอมพิวเตอร์มีความสามารถในการใช้เหตุผลเหมือนสมองของมนุษย์ขึ้นมาได้ (McCulloch and Pitts, 1943) แนวคิดเรื่องนี้ของแม็คคูลอคและพิตต์สในภายหลังจะสร้างแรงบันดาลใจให้แฟรงค์ โรเซนบลัทส์สร้างต้นแบบของระบบโครงข่ายใยประสาทเทียมขึ้นมาจริงๆ โดยตั้งชื่อว่าเพอร์เซ็ปตรอน (Perceptron) แต่ในช่วงนั้นเพอร์เซ็ปตรอนถูกพัฒนาขึ้นมาจากฮาร์ดแวร์และระบบอิเล็กทรอนิกส์ยังไม่ได้เป็นระบบคอมพิวเตอร์ที่เป็นเครือข่ายที่ซับซ้อนและใช้คู่กับซอฟต์แวร์สมัยใหม่ดังปัจจุบัน อย่างไรก็ตามอย่างไรก็ดีเพอร์เซ็ปตรอนแสดงให้เห็นว่าแนวคิดของแม็คคูลอคและพิตต์สมีความเป็นไปได้

(Rosenblatt, 1958; กอบเกียรติ, 2565, หน้า 71-76; ปริญา, 2562, หน้า 105-112) ภายหลังเดวิด รูเมลฮาร์ต (David Rumelhart) จอห์น ฮินตัน (Geoffrey Hinton) และโรนัลด์ วิลเลียมส์ (Ronald Williams) จะเสนอแนวคิดเรื่องโครงข่ายประสาทเทียมแบบแพร่ย้อนกลับ (Backpropagation) เพื่อใช้เรียนรู้ข้อมูลเพิ่มเติมในโครงข่ายประสาทเทียมแบบหลายชั้น (Multilayer Feedforward network) จนทำให้โมเดลมีความเที่ยงตรงมากขึ้น ด้วยแนวคิดนี้เองถือเป็นการปฏิวัติวงการปัญญาประดิษฐ์ขนานใหญ่และผ่านช่วง “ฤดูหนาวเอไอ” (AI Winter) ที่ปราศจากความก้าวหน้าในการค้นคว้าวิจัยอยู่หลายปี (Rumelhart, Hinton, & Williams, 1986)

ทั้ง ยาน เลอคุน (Yann André LeCun) (ใน Lecun, Bottou, Bengio and Haffner, 1998) โยชัว เบงจิโอ (Yoshua Bengio) (ใน Bengio, Lamblin, Popovici and Larochelle, 2006) และจอห์น ฮินตัน (ใน Hinton, Osindero, and The, 2006) ต่างคนก็พัฒนาองค์ความรู้เรื่องโครงข่ายประสาทเทียม (Artificial Neuron Network) ต่อมาอีกมาก จนกระทั่งทั้งสามคนได้รับรางวัล ACM A.M. Turing Award ซึ่งถือกันว่าย เทียบเท่ากับรางวัลโนเบลในวงการคอมพิวเตอร์ ในปี ค.ศ. 2018 งานของทั้งสามคนนี้ ผสมกับงานวิจัยที่ต้อง ถูกรื้อสร้างจุดเปลี่ยนครั้งสำคัญทั้งในองค์ความรู้เรื่องโครงข่ายประสาทเทียมและการประมวลผลภาษาธรรมชาติ (Natural Language Processing: NLP) อย่างงานของทีมีวิจัยจากบริษัทกูเกิลเรื่อง “ตัวแปลง” (Transformer) แนวคิดที่ถือว่าการปฏิบัติของ “ตัวแปลง” ที่เปลี่ยนแปลงไปจากเครือข่ายประสาทที่ใช้ สถาปัตยกรรมแบบวิเคราะห์การเกิดซ้ำไม่ว่าจะเป็นเรื่อง ข้อมูลอนุกรมเวลา ราคาหุ้น และภาษาอย่าง Recurrent Neuron Networks (RNN) (ดูข้อมูล RNN ได้ที่ Géron, 2017, pp. 383-415) คือจะมีกลไกหลัก เพิ่มเข้ามาคือ Attention และ Self-attention เพื่อจดจำคำยืมไว้ในรูปแบบเครือข่ายที่เกี่ยวข้อว่าข้อมูลที่ ป้อนเข้ามาส่วนใดมีความสำคัญ (Vaswani et al., 2017; กอบเกียรติ, 2565, หน้า 572-574) ทั้งหมดนี้ต่อมา จะกลายเป็นจุดกำเนิดของปัญญาประดิษฐ์แบบรู้สร้างที่เป็นโมเดลภาษาขนาดใหญ่ (Large Language Model: LLM) อย่างแซทจีพีทีในที่สุด

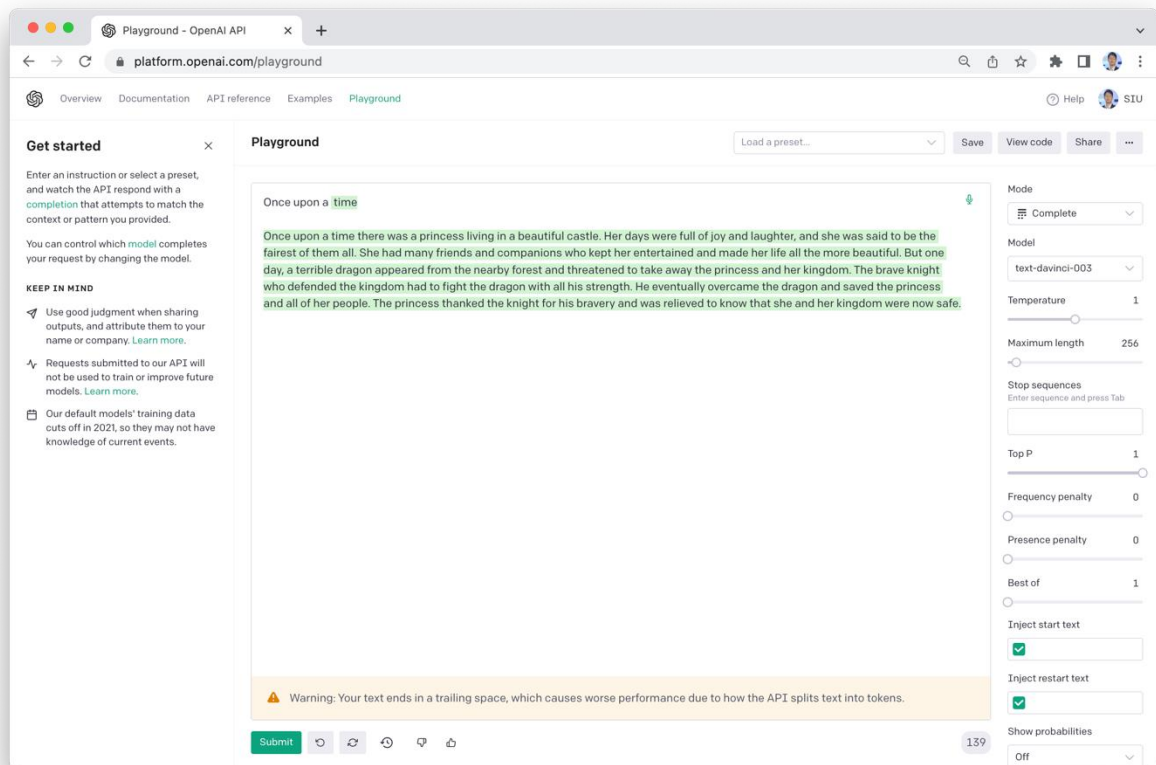
จีพีทีเวอร์ชันแรก (GPT-1) ปรากฏตัวเมื่อเดือนมิถุนายน ค.ศ. 2018 โดยถูกพัฒนาขึ้นมาจาก สถาปัตยกรรม Transformer แล้วนำมาปรับแต่ง (fine-tuned) ให้ตัวเอไอมีความเข้าใจภาษาธรรมชาติ เพิ่มเติมมากขึ้น ซึ่งแม้จะไม่ได้ถูกปรับแต่งดังกล่าวตัวจีพีทีเวอร์ชันแรกก็มีความสามารถประมวลผล ภาษาธรรมชาติเป็นอย่างดีอยู่แล้วใช้ชุดข้อมูล “คลังข้อมูลหนังสืออิเล็กทรอนิกส์” (Book Corpus Data Set) จากเว็บไซต์สมชเวิร์ด (Smashwords) รวบรวมเล่ม ตัวจีพีทีเวอร์ชันนี้จะมีพารามิเตอร์จำนวน 117 ล้าน ชุด ความสามารถของมันแม้จะเป็นเพียงเวอร์ชันแรกแต่ก็เป็นที่น่าประทับใจ กล่าวคือมันสามารถตอบคำถาม ได้แม้จะไม่ได้มีการป้อนตัวอย่างให้ก่อน (เรียกว่าการป้อนข้อมูลแบบ “ศูนย์ตัวอย่าง” หรือ zero-shot learning ในกรณีที่มีการยิงคำถามโดยป้อนตัวอย่างหนึ่งชุดจะเรียกว่า one-shot learning ถ้าต้องป้อน ตัวอย่างมากกว่าหนึ่งชุดขึ้นไปจะเรียกว่า few-shot learning)

เมื่อปีรุ่งขึ้นในเดือนกุมภาพันธ์ ค.ศ. 2019 โอเพ่นเอไอก็ได้เปิดตัวจีพีทีรุ่นที่สอง (GPT-2) ซึ่งมีขนาด ใหญ่กว่าจีพีทีรุ่นแรก แต่สถาปัตยกรรมเบื้องหลังนั้นเหมือนเดิมทุกอย่าง จีพีทีรุ่นนี้สามารถตอบสนองคำสั่งให้ ทำงานได้หลากหลายแบบมากขึ้นโดยไม่จำเป็นต้องมีการฝึกสอนใดๆ ก่อนเลย ความสำคัญของจีพีทีรุ่นที่สองนี้ คือมันพิสูจน์ให้เห็นว่ายิ่งเราขยายขนาดของเอไอประเภทโมเดลภาษาขนาดใหญ่ (Large Language Model: LLM) นี้มากขึ้นเท่าใด มันก็ยิ่งมีความสามารถเพิ่มเป็นทวีคูณเท่านั้นแม้จะไม่ได้มีการปรับแต่งใดๆ ให้มันเลย

ทีมงานโอเพ่นเอไอใช้ชุดข้อมูลที่สร้างขึ้นจากการใช้โปรแกรมเข้าไปอ่านเอกสารจำนวน 8 ล้านชุดในเว็บเรดดิต (Reddit) ทำให้ได้ฐานข้อมูลขนาด 40 GB แล้วนำมาเก็บไว้เป็นฐานข้อมูลชื่อว่า WebText อนึ่งแชทจีพีทีเวอร์ชันสองจะมีจำนวนพารามิเตอร์ 1.5 พันล้านชุดซึ่งต้องถือว่ามีความใหญ่กว่าแชทจีพีทีเวอร์ชันแรกถึงสิบเท่าเลยทีเดียว

ส่วนจีพีทีรุ่นที่สาม (GPT-3) มีขนาดใหญ่กว่าจีพีทีรุ่นที่สองในระดับทวีคูณถึงสองเท่าคือมีจำนวนพารามิเตอร์ถึง 1 แสน 7 หมื่น 5 พันล้านชุด และถูกฝึกสอนด้วยชุดข้อมูลที่เป็นตัวอักษรที่แตกต่างกันถึงห้าชุด โดยมีจำนวนชุดข้อมูลมากกว่าชุดที่เคยถูกป้อนให้กับจีพีทีรุ่นที่สอง ที่สำคัญในจีพีทีรุ่นที่สามนี้ทางโอเพ่นเอไอได้เปิดให้เข้าถึงด้วยการใช้ช่องทางการเชื่อมต่อผ่านทาง API หรือ Application Programming Interface

เมื่อมีการเปิดตัวจีพีทีรุ่นที่สามนี้ออกไป ก็มีคำวิจารณ์เชิงบวกจากหลากหลายทัศนะ อาทิ จีพีทีรุ่นที่สามแสดงศักยภาพในการเป็นหนึ่งในสิบเทคโนโลยีแห่งปี ค.ศ. 2021 หรือแม้กระทั่งอาร์ราม ซาเบตกีทวิตถึงว่าพอได้ใช้งานจีพีทีรุ่นที่สามเหมือนกำลังเห็นอนาคต และเขาก็ได้เล่าว่าเขาใช้มันเพื่อเขียนเพลง แต่งนิทาน ร่างแถลงข่าว แต่งโน้ตกีตาร์ ร่างคำสัมภาษณ์ ร่างบทความ แต่งคู่มือเทคนิค ฯลฯ (Kublikand Saboo, 2022, pp. 9-14)



ภาพที่ 2 ภาพหน้าจอ Playground ของโอเพ่นเอไอ

ในช่วงแรกทางโอเพ่นเอไอได้เตรียม “เครื่องยนต์” ของจีพีทีรุ่นที่สามไว้ให้เราสืบแบบด้วยกันซึ่งเครื่องยนต์ทั้งสี่แบบนี้ต่างก็ถูกตั้งชื่อตามชื่อนักวิทยาศาสตร์เรียงลำดับตามตัวอักษรภาษาอังกฤษ โดยเริ่มตั้งแต่ เออด้า (Ada ตั้งชื่อตาม Ada Lovelace), แบบเบจ (Babbage ตั้งชื่อตาม Charles Babbage),

คูรี (Curie ตั้งชื่อตาม Madame Marie Curie) และดาร์วินชี (Davinci ตั้งชื่อตาม Leonardo da Vinci) เครื่องยนต์ทั้งสองแบบนี้มีเป้าหมายและวัตถุประสงค์การใช้งานที่ต่างกันไป หากเราไม่ระบุตัวไหน (โปรดดูจากภาพ Playground ด้านบน) โอเพ่นเอไอจะกำหนดค่าปริยายเป็นดาร์วินชีซึ่งเครื่องยนต์ตัวนี้จะให้ประสิทธิภาพที่ดีที่สุดแต่ก็มีค่าใช้จ่ายในการใช้งาน API ที่สูงที่สุดเช่นกัน สำหรับคูรีนั้นทางโอเพ่นเอไอพยายามสร้างสมดุลระหว่างประสิทธิภาพและความเร็วในการใช้งาน ส่วนแบบเบจแม้ว่าจะมีความเร็วกว่าคูรีแต่ก็จะมีข้อจำกัดในการเข้าใจตรรกะและบริบทของข้อมูลที่ป้อนเข้าไป ด้านเอได้นั้นจะมีความเร็วในการใช้งานสูงที่สุดและมีค่าใช้จ่ายถูกที่สุดด้วยแต่ก็ต้องแลกกับการที่มันไม่สามารถใช้งานที่อาศัยตรรกะที่ซับซ้อนได้ โดยส่วนใหญ่เอได้อาจจะถูกใช้ประมวลผลสำหรับงานง่าย ๆ ทั่วไป หลังจากพัฒนาเครื่องยนต์ทั้งสองแบบนี้แล้วทางโอเพ่นเอไอก็ได้พัฒนาตัวแบบเอไอใหม่ที่ชื่อว่า อินสตรัคจีพีที (InstructGPT) ซึ่งถูกฝึกสอนให้มีความเข้าใจคำสั่งและปฏิบัติตามคำสั่งได้ดีกว่าเดิม และยังให้ข้อความที่มีปัญหา (toxic) ที่น้อยกว่าและให้ข้อมูลที่ถูกต้องตรงความเป็นจริงมากกว่าจีพีทีรุ่นที่สามจากโมเดลทั้งสองแบบที่เราได้พูดถึงไปแล้ว เจ้าอินสตรัคจีพีทีนี้ถูกฝึกสอนด้วยมนุษย์ ซึ่งเป็นการป้อนกลับการเรียนรู้ของปัญญาประดิษฐ์ด้วยวงจรป้อนกลับโดยมนุษย์ (Reinforcement Learning by Human Feedback: RLHF) ตัวอินสตรัคจีพีทีตัวนี้เองที่เป็นตัวต้นแบบของแชทจีพีทีอันโด่งดังในปัจจุบันนั่นเอง (2022, pp. 9-14)

```
In [3]: from transformers import GPT2LMHeadModel, GPT2Tokenizer
```

```
# Load pre-trained GPT-2 model and tokenizer
model_name = "gpt2-medium"
model = GPT2LMHeadModel.from_pretrained(model_name)
tokenizer = GPT2Tokenizer.from_pretrained(model_name)

# Function to generate a response using GPT-2
def generate_response(prompt):
    inputs = tokenizer.encode(prompt, return_tensors="pt")
    outputs = model.generate(
        inputs,
        max_length=150,
        num_return_sequences=1,
        temperature=0.8, # Adjust the temperature to control randomness
        top_p=0.9, # Use nucleus sampling to reduce repetition
        repetition_penalty=1.2, # Apply a repetition penalty to discourage repeating phrases
    )
    response = tokenizer.decode(outputs[0], skip_special_tokens=True)
    return response

# Interactive chat loop
print("Chatbot is ready to talk! (type 'quit' to exit)")
while True:
    user_input = input("You: ")
    if user_input.lower() == 'quit':
        break

    response = generate_response(user_input)
    print("Chatbot: " + response)
```

```
Chatbot is ready to talk! (type 'quit' to exit)
You: Hello my name is Peter
```

```
The attention mask and the pad token id were not set. As a consequence, you may observe unexpected behavior. Please pass your input's 'attention_mask' to obtain reliable results.
Setting 'pad_token_id' to 'eos_token_id':50256 for open-end generation.
```

```
Chatbot: Hello my name is Peter and I am a professional photographer.
```

```
I have been shooting for over 20 years, but it was only recently that the passion of photography really took hold in me. It's something you can't teach or learn from experience; your own personal journey to becoming an artist has always fascinated us.
```

```
(photo by: David Hockney)
You: quit
```

ภาพที่ 3 หน้าจอแสดงการเรียกใช้จีพีทีเวอร์ชันสองผ่านเอพีไอ

เมื่อทำการเรียกใช้จีพีทีเวอร์ชันสองผ่านเอพีไอจะเห็นว่าข้อความที่ถูกสร้างขึ้นผ่านจีพีทีเวอร์ชันนี้ซึ่งยังไม่ได้ถูกปรับแต่งใดๆ จะให้ข้อความที่ไม่ได้มีความเกี่ยวข้องกับข้อมูลที่เรป้อนเข้าไปเลย (ในที่นี้เราเขียนโปรแกรมภาษาไพธอนให้เป็นรูปแบบแชทบอท) และยังไม่มีความสามารถในการรับทราบคำตอบที่ได้ที่เราคาดหวังการตอบรับกลับมา

```
In [1]: import openai
import re
import os

# Initialize the OpenAI API key
openai.api_key = os.environ["OPENAI_API_KEY"]

response = openai.Completion.create(
    # engine="text-davinci-002",
    engine="text-davinci-002",
    prompt="write a scenario regarding geopolitical situation between China and the United States?",
    temperature=0.7,
    max_tokens=150,
    top_p=1,
    frequency_penalty=0,
    presence_penalty=0,
)

generated_text = response.choices[0].text.strip()
print(generated_text)
```

The economic relationship between the United States and China has always been a contentious one, with the two countries regularly sparring over trade practices and currency manipulation. In recent years, the geopolitical situation between the two countries has become increasingly strained, as China has become more assertive in the South China Sea and the United States has responded by stationing military assets in the region. The situation has led to a number of tense standoffs between the two countries, and many experts believe that it is only a matter of time before the two nations come into direct conflict with one another.

ภาพที่ 4 หน้าจอแสดงการเรียกใช้จีพีทีเวอร์ชันสามผ่านเอพีไอ

ในทางตรงข้าม ข้อความที่ได้รับจากการที่จีพีทีเวอร์ชันสามที่สร้างขึ้นจากดา วินชี (Text-Davinci-002) กลับมีลักษณะตอบรับคำสั่งที่เราป้อนถามได้ดีกว่า มีความหมายมากกว่า และมีข้อมูลที่ถูกต้องเที่ยงตรงมากกว่าจีพีทีเวอร์ชันสองก่อนหน้านี้อย่างมาก

การประยุกต์ใช้เอไอ

ปกติเอไอถูกใช้ทั่วไปกับระบบคอมพิวเตอร์อยู่แล้วอย่างแพร่หลายในชีวิตประจำวัน เช่น ระบบแนะนำสินค้าในเว็บไซต์ค้าปลีกออนไลน์ ระบบค้นหาของ Google เครือข่ายสังคมออนไลน์ เป็นต้น เอไอช่วยให้เว็บไซต์เหล่านี้สามารถให้บริการได้อย่างมีประสิทธิภาพมากขึ้น

ในอนาคตอันใกล้ เอไอจะถูกพัฒนาให้สามารถทำงานร่วมกับมนุษย์ได้ดียิ่งขึ้น โดยเฉพาะอย่างยิ่งในงานบริการ การแพทย์ และการศึกษา เช่น แชทบอทสำหรับการซักประวัติผู้ป่วย ระบบวิเคราะห์ข้อมูลการเรียนรู้ของนักเรียน เป็นต้น อย่างไรก็ตาม การพัฒนาเอไออย่างมีจริยธรรมและการควบคุมการใช้งานที่เหมาะสมจะเป็นเรื่องสำคัญเพื่อให้เอไอสามารถอำนวยความสะดวกในการดำเนินชีวิต โดยไม่สร้างผลกระทบในทางลบต่อสังคม โดยทั่วไปเรื่องหลักจริยธรรมมีข้อถกเถียงกันมาก แต่หลักการโดยทั่วไปที่ระบบปัญญาประดิษฐ์ควรจะต้องบรรจुर่องหลักจริยธรรมก็ได้แก่ การให้น้ำหนักเรื่องความปลอดภัย การให้น้ำหนักเรื่องความยุติธรรม เคารพสิทธิส่วนบุคคล สนับสนุนการสร้างความร่วมมือ ต้องโปร่งใส ลดการใช้งานในทางที่ก่อให้เกิดอันตราย ยึดหลักการตรวจสอบได้ ยึดมั่นคุณค่าสิทธิมนุษยชน เสริมสร้างหลักความหลากหลาย

และการมีส่วนร่วม หลีกเลี่ยงการรวมศูนย์อำนาจ เคารพกฎหมายและนโยบายของทางการ ต้องคำนึงถึงผลกระทบด้านปัญหาการว่างงานที่อาจเกิดขึ้นจากเอไอด้วย ฯลฯ (Russell and Norvig, 2022, pp. 1037-1041; Jobin, Lenca, and Vayena, 2019) การพัฒนา AI แบบอัตโนมัติจะมีบทบาทสำคัญมากขึ้นในอนาคต เพราะช่วยให้สามารถสร้าง AI ได้รวดเร็วและมีประสิทธิภาพมากขึ้น

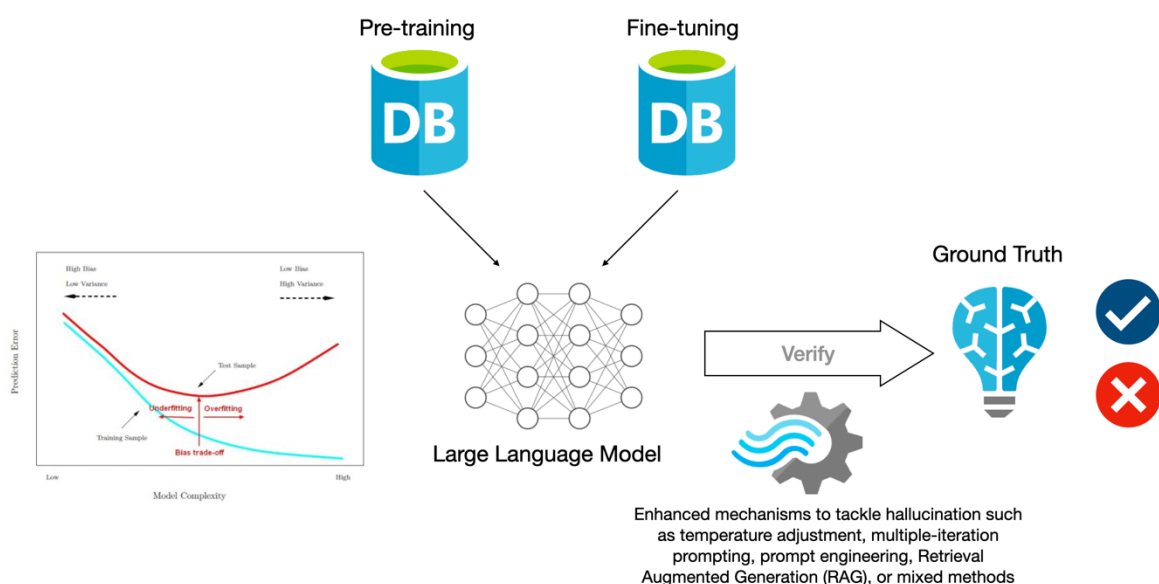
ด้าน Machine Learning Operations (MLOps) ก็จะมีการพัฒนาเครื่องมือที่ช่วยออกแบบ ML Pipeline, จัดการกระบวนการฝึกสอนและประเมินผลโมเดล, ติดตามและตรวจจับปัญหาของโมเดลที่นำไปใช้จริง, รวมถึงปรับโมเดลให้ Up-to-date อยู่เสมอโดยอัตโนมัติ

เทคโนโลยีเหล่านี้จะช่วยให้สามารถพัฒนาและประยุกต์ใช้เอไอได้อย่างรวดเร็ว มีประสิทธิภาพ และคุ้มค่ามากยิ่งขึ้น ซึ่งจะทำให้เอไอแพร่หลายในวงกว้างมากขึ้นในอนาคตอันใกล้นี้ จากการสำรวจของ EdX เกือบครึ่งหนึ่งของซีโอโอเชื่อว่าเอไอสามารถและควรรับภาระหน้าที่ส่วนใหญ่แทนตัวเอง ซึ่งผู้บริหารเหล่านี้เชื่อว่าการที่เอไอเข้ามาช่วยทำงานอัตโนมัติในการแบ่งเบาภาระของตนจะทำให้มีเวลามากขึ้นสำหรับการเป็นผู้นำเชิงกลยุทธ์ แต่แม้ว่า 79% จะกังวลว่าตนอาจล้าหลังหากไม่ปรับตัวให้เข้ากับเอไอ ผู้เชี่ยวชาญกล่าวว่าเครื่องจักรยังไม่สามารถเลียนแบบการคิดเชิงกลยุทธ์ของมนุษย์ได้ สรุปได้ว่า การสำรวจเปิดเผยว่าผู้บริหารเห็นเอไอไม่ใช่แค่ทางเลือกแต่เป็นความจำเป็น แม้ว่าคุณสมบัติแท้จริงของผู้นำจะยังไม่สามารถถูกแทนที่ด้วยเอไอได้ ส่วนการสำรวจเมื่อเร็วๆ นี้ของ KPMG เปิดเผยว่า ซีโอโอชาวอเมริกันมองโลกในเชิงบวกต่อการลงทุนในเอไอ โดยมีถึง 72% ที่ระบุปัญหาประดิษฐ์แบบรู้สร้างเป็นความจำเป็นลำดับต้นๆ และคาดว่าจะได้ผลตอบแทนจากการลงทุนภายใน 3 ถึง 5 ปี ความกระตือรือร้นนี้ ตามมาหลังความก้าวหน้าอย่างมากของแชทจีพีทีซึ่งแสดงถึงศักยภาพที่สร้างความเปลี่ยนแปลงอย่างใหญ่หลวงของเอไอต่ออุตสาหกรรมต่างๆ แม้ว่าเทคโนโลยีนี้จะถูกมองในแง่ดีว่าก่อให้เกิดการเติบโตของรายได้และประสิทธิภาพการดำเนินงาน แต่ก็ยังมีการอภิปรายอย่างต่อเนื่องในด้านลบที่เกิดผลกระทบต่อการจ้างงาน นอกจากนี้ซีโอโอยังเริ่มให้ความสำคัญกับการกลับไปทำงานที่สำนักงานมากขึ้น โดยมีถึง 62% ที่คาดว่าจะมีงานจะทำงานที่สำนักงานอย่างถาวรภายใน 3 ปี ซึ่งเพิ่มขึ้นอย่างมากจากปีที่แล้ว แม้รูปแบบการทำงานและภูมิทัศน์เทคโนโลยีจะเปลี่ยนแปลงอยู่ตลอดเวลา แต่ซีโอโอดูเหมือนจะมุ่งมั่นที่จะผสมผสานวัฒนธรรมสำนักงานแบบดั้งเดิมเข้ากับการพัฒนาเอไอขึ้นแนวหน้า

ด้านจากกลุ่มนักวิจัยข้ามสาขาวิชา โดยเฉพาะอย่างยิ่งจากมหาวิทยาลัยสแตนฟอร์ดและมหาวิทยาลัยนอร์ทเวสเทิร์น ในบรรดาสถาบันอื่นๆ สำคัญของการศึกษานี้ไม่เพียงแต่ทันต่อเหตุการณ์เท่านั้น แต่ยังสะท้อนประเด็นที่ชุมชนวิทยาศาสตร์กำลังพยายามแก้ไข นั่นคือคอขวดของการประเมินโดยผู้ประเมินที่มีคุณภาพ งานวิจัยชิ้นนี้นำ GPT-4 มาใช้ในกระบวนการอัตโนมัติเพื่อแสดงความคิดเห็นต่อบทความวิชาการแล้วเปรียบเทียบความคิดเห็นเหล่านั้นกับความคิดเห็นของผู้ทบทวนโดยมนุษย์ ซึ่งน่าประหลาดใจในแง่ที่ประเด็นซึ่งเอไอยกขึ้นมานั้นมักจะมีคาบเกี่ยวใกล้เคียงกับประเด็นที่ผู้วิจารณ์ที่เป็นมนุษย์สองคนยกขึ้นมายังเป็นบทความที่อ่อนด้อยคุณภาพมากเท่าใด ก็ยังแสดงลักษณะนี้อย่างเด่นชัด นอกจากนี้ มากกว่าครึ่งหนึ่งของนักวิจัย 308 คนที่สำรวจ รู้สึกว่าข้อเสนอแนะที่สร้างโดยเอไอมีประโยชน์ ในขณะที่ 82.4% มองว่าข้อเสนอแนะของเอไอมีประโยชน์มากกว่าข้อเสนอแนะบางอย่างที่ได้รับจากมนุษย์ อย่างไรก็ตาม GPT-4 ก็ยังมีข้อจำกัดในการวิจารณ์การออกแบบวิธีการวิจัย สรุปได้ว่า หลักฐานแสดงให้เห็นอย่างชัดเจนว่า เรากำลังอยู่เพียงจุดเริ่มต้นของภูเขาน้ำแข็งเท่านั้น LLMs เช่น GPT-4 กำลังเข้าสู่บริเวณที่พวกมันจะกลายเป็นเครื่องมือสำคัญเสริมการทำงานของมนุษย์ แม้ว่าพวกมันยังไม่สามารถแทนที่คุณค่าลึกซึ้งของข้อเสนอแนะจากมนุษย์ได้

แต่ประโยชน์ของพวกมัน โดยเฉพาะอย่างยิ่งในสถานการณ์ที่ขาดแคลนทรัพยากรหรือต้องแข่งกับเวลา ก็ไม่ควรมองข้าม ยุคของการสืบค้นทางวิทยาศาสตร์ด้วยความช่วยเหลือของเอไอได้มาถึงแล้ว มิใช่ที่กำลังจะมา (Liang et al., 2023)

อย่างไรก็ตามการใช้งานปัญญาประดิษฐ์แบบรู้สร้างโดยเฉพาะโมเดลภาษาขนาดใหญ่มีปัญหาในเรื่องที่คาบเกี่ยวกับปัญหาความหลอน (hallucination) คือเอไออาจให้คำตอบที่จินตนาการไปเองโดยไม่เกี่ยวกับความจริง ดังที่ได้กล่าวไว้แต่ต้นแล้ว ผู้เขียนได้กำหนดมาตรการที่จะแก้ปัญหาเหล่านี้ขึ้นตั้งแต่ (1) การปรับแต่งอุณหภูมิ (temperature setting) เพื่อให้เอไอลดการใช้จินตนาการลงและยึดกับข้อมูลที่เป็นจริงมากขึ้น ปัญหาในทางกลับของมาตรการนี้คือในบางกรณีเราอาจต้องกำหนดให้เอไอใช้จินตนาการเพื่อสร้างงานเชิงสร้างสรรค์ อาทิ การเขียนนิยายวิทยาศาสตร์ หรือการสร้างฉากทัศน์อนาคต เป็นต้น (2) การโต้ตอบกับเอไอหลายระลอกเพื่อค่อย ๆ ชี้นำให้เอไอทำงานตามเป้าหมายของงาน (multiple-iteration prompting) (3) การใช้เทคนิควิศวกรรมพรอมต์ (prompt engineering) เช่นการวางระบบสายโซ่เหตุผล (chain of thought) (4) การสร้างเนื้อหาโดยใช้ระบบสืบค้นเสริม (Retrieval Augmented Generation: RAG) โดย RAG จะรับอินพุตเข้ามาแล้วทำการเรียกค้นเอกสารที่เกี่ยวข้องหรือสนับสนุนอินพุตนั้นจากแหล่งข้อมูล เช่น วิกิพีเดีย เอกสารเหล่านั้นจะถูกนำมาต่อกับอินพุตเดิมแล้วส่งเข้าสู่โมเดลการสร้างข้อความ เพื่อสร้างเอาต์พุตขั้นสุดท้าย ทำให้ RAG สามารถปรับตัวเข้ากับสถานการณ์ที่ข้อเท็จจริงอาจเปลี่ยนแปลงไปตามเวลา ซึ่งมีประโยชน์มาก เนื่องจากความรู้ในโมเดลภาษาขนาดใหญ่มีลักษณะคงตัวไม่เปลี่ยนแปลงหากไม่ได้รับการฝึกฝนเพิ่มเติม แต่ RAG จะช่วยให้โมเดลภาษาสามารถเข้าถึงข้อมูลล่าสุดเพื่อสร้างเอาต์พุตที่น่าเชื่อถือได้โดยไม่ต้องฝึกฝนใหม่ (5) หรือใช้วิธีเหล่านี้ผสมผสานกันทั้งหมด ลักษณะเหล่านี้เป็นปัญหาพิเศษจากปัญญาประดิษฐ์แบบรู้สร้างที่พัฒนาอยู่บนระบบโครงข่ายประสาทหลายชั้นที่ซับซ้อนและมีขนาดใหญ่ ถูกฝึกฝนด้วยข้อมูลที่คัดสรรมาจํานวนมหาศาลที่ต่างออกไปจากระบบผู้เชี่ยวชาญแบบเดิมที่ปราศจากคุณสมบัติเหล่านี้ ทำให้มันมีความสามารถที่เรียกว่า “ระบบการแทนที่ความเป็นเหตุเป็นผลทางอ้อม” (Meta-Causal Representation) ซึ่งเกิดจากแบบจำลองนามธรรมที่จับความสัมพันธ์เชิงเหตุและผลที่อยู่เบื้องหลังข้อมูลได้ โดยไม่จำเป็นต้องสังเกตการณ์จากโลกจริงโดยตรง มุมมองนี้สอดคล้องกับอีเลีย ซัสคีเวอร์ (Ilya Sutskever) ที่เชื่อว่า โครงข่ายประสาทเทียมขนาดใหญ่นั้นเรียนรู้ “ผ่านตัวแทนโลกจริง” (learning representation) ด้วยการฝึกฝนจากข้อมูลที่ป้อนให้ สิ่งที่โมเดลทำได้จึงอาจไม่ใช่เพียงแค่การหาค่าความสัมพันธ์ทางสถิติ แต่เราจำเป็นต้องตระหนักว่าถึงแม้โมเดลเหล่านี้จะทำนายอย่างแม่นยำและซับซ้อนจากการแทนแบบนามธรรม แต่ “ความเข้าใจแบบจักรกล” นี้แตกต่างจากความเข้าใจแบบมนุษย์ แม้จะจับเหตุผลเชิงสาเหตุและความซับซ้อนของโลกได้บ้าง แต่โมเดลเหล่านี้ยังไร้ซึ่งสำนึกและความลึกซึ้งในการเข้าใจเทียบเท่ามนุษย์



ภาพที่ 5 ระบบการแทนที่ความเป็นเหตุเป็นผลทางอ้อม

ภาพที่ 5 แสดงรูปการเรียนรู้ของเครื่องจักรหรือ “ระบบการแทนที่ความเป็นเหตุเป็นผลทางอ้อม” (Meta-Causal Representation) เนื่องจากเหตุผลแบบเครื่องจักรเป็นการเรียนรู้แบบข้อมูลตัวแทน ดังนั้นจึงต้องมีการสอบเทียบกับ “ความจริงฐาน” เพื่อวัดความเที่ยงตรงของการตอบสนองของโมเดลอยู่อย่างสม่ำเสมอ วิธีการลดค่าความผิดพลาดนี้ หรืออาจเรียกอีกอย่างว่าลด “ความหลอน” ของโมเดลอาจทำได้วิธีอื่นประกอบกันด้วยก็ได้ นอกจากนี้การฝึกสอนโมเดลด้วยข้อมูลเฉพาะทาง (fine-tuning) ซ้ำก็ช่วยยั้งช่วยให้โมเดลตอบสนองคำสั่งอย่างเที่ยงตรงมากขึ้นไปอีก

อนึ่ง ผู้เขียนประสบความสำเร็จในการใช้ปัญญาประดิษฐ์แบบรู้สร้างมาสร้างกลไกการตรวจจับแนวโน้มทางสังคมอัตโนมัติจากข้อมูลผ่านโซเชียลมีเดียที่ชื่อว่า “PulsarWave” และได้ส่งผลงานเข้าประกวดในงานเสิ่คาคาตอนของ Google Vertex AI และได้รับการคัดเลือกให้อยู่ใน 6 อันดับแรก งานชิ้นนี้แสดงให้เห็นถึงมาตรการในการลด “ความหลอน” เพื่อให้โมเดลปัญญาประดิษฐ์แบบรู้สร้างสามารถทำงานได้อย่างอัตโนมัติ ได้รับการยอมรับในระดับสากล นอกจากนี้ผู้เขียนยังได้รับอนุมัติให้ทำโครงการร่วมกับศูนย์ทรัพยากรคอมพิวเตอร์เพื่อการคำนวณขั้นสูง (ThaiSC) เพื่อทำการฝึกสอน (fine-tuning) ให้กับปัญญาประดิษฐ์แบบรู้สร้างเพื่อเพิ่มความเชี่ยวชาญเฉพาะทางในด้านภูมิรัฐศาสตร์ภายใต้ชื่อโครงการ โมเดลภาษาขนาดใหญ่ (Large Language Model, LLM) สำหรับการวิเคราะห์ความเสี่ยงทางภูมิรัฐศาสตร์ผ่านซูเปอร์คอมพิวเตอร์ “LLMs for Geopolitical Risk Analysis via Supercomputing” เพื่อปรับปรุงประสิทธิภาพของตัว PulsarWave ในอนาคตอีกด้วย

สรุปและข้อเสนอแนะ

บทความนี้แสดงให้เห็นถึงแรงขับเคลื่อนหลักจากสามกระแสด้วยกันคือ (1) แรงขับเคลื่อนจากการลงทุนในการวิจัยและพัฒนาของภาครัฐโดยเฉพาะอย่างยิ่งจากสำนักงานวิจัยโครงการขั้นสูงด้านกลาโหมสหรัฐอเมริกาหรือ DARPA ด้วยวัตถุประสงค์ทางการด้านการป้องกันประเทศ แต่จากทั้งงานวิจัยและข้อมูลในประวัติศาสตร์จะเห็นว่าการริเริ่มจากสำนักงานจะส่งผลต่อการสร้างนวัตกรรมและการบริการทางด้านเทคโนโลยีในภาคเอกชนตามไปด้วยเช่นการเกิดขึ้นของระบบอินเทอร์เน็ต หรือระบบนำร่องผ่านดาวเทียม GPS เป็นต้น สำนักงานฯ ได้ลงทุนวิจัยและพัฒนาเพื่อนำคอมพิวเตอร์มาใช้ในการงานทางด้านสังคมศาสตร์หลายโครงการด้วยกันไม่ว่าจะเป็น NGS2 หรือ GT เป็นต้น ผลการลงทุนและพัฒนาเหล่านี้ย่อมส่งผลถึงภาคเอกชนในอนาคต

นอกจากนี้ในภาควิชาการโดยเดวิด มอร์แกนก็ชี้ให้เห็นว่าระเบียบวิธีวิจัยในทางสังคมศาสตร์เกิดการเปลี่ยนกระบวนทัศน์จากกระบวนทัศน์ปฏิฐานนิยม มาเป็นกระบวนทัศน์อภิปราย ซึ่งทำให้เกิดการเปิดกว้างต่อบรรทัดระเบียบวิธีวิจัยที่เปิดกว้างมากขึ้นไม่ว่าจะเป็นแนวทางการวิจัยเชิงปริมาณแบบเดิม หรือการวิจัยเชิงคุณภาพที่ได้รับความนิยมมากขึ้น แต่กระบวนทัศน์นี้ก็ยังมีปัญหาอยู่ในแง่ที่ว่าเกิดความขัดแย้งระหว่างกันจนไม่สามารถรวมขอมกันได้ มอร์แกนจึงเสนอว่าการเกิดขึ้นของระเบียบวิธีวิจัยแบบผสมผสานแท้ที่จริงอยู่ภายใต้การเปลี่ยนกระบวนทัศน์ครั้งใหม่คือกระบวนทัศน์ภายใต้ปรัชญาแฟร์กมาติซึม โดยเขาได้เสนอพื้นที่กลางที่เป็นที่ยอมรับของทุกฝ่าย และย้ายจุดศูนย์กลางในเรื่องญาณวิทยามาเป็นวิธีวิทยาโดยมุ่งเน้น “การสืบสวนหาสาเหตุ” (abduction) เน้นหนักเรื่องอัตวิสัยที่เชื่อมโยงต่อกัน (intersubjectivity) และความสามารถในการถ่ายโอนผลการวิจัยข้ามไปมาได้ (transferability) ทำให้สร้างความชอบธรรมต่อระเบียบวิธีวิจัยแบบผสมผสานที่กำลังได้รับความนิยมมากขึ้นทุกขณะ นอกจากนี้การเกิดขึ้นของกระบวนทัศน์ภายใต้ปรัชญาแฟร์กมาติซึมยังเปิดกว้างต่อ “การใช้คอมพิวเตอร์ในสังคมศาสตร์” (Computational Social Science) ซึ่งได้รับความนิยมมากขึ้นเช่นกันหลังพัฒนาการอย่างมากของคอมพิวเตอร์และการผลิตข้อมูลผ่านสื่อทางสังคมอย่างมหาศาลทำให้การวิเคราะห์พฤติกรรมของมนุษย์ผ่านเทคโนโลยีเช่นการเรียนรู้ของเครื่องจักรที่จำเป็นต้องอาศัยข้อมูลจำนวนมากนี้เป็นไปได้ สาขา “การใช้คอมพิวเตอร์ในสังคมศาสตร์” เริ่มลงหลักปักฐานในแวดวงวิชาการอย่างเป็นทางการและมีความมั่นคงแสดงให้เห็นจากการตีพิมพ์หนังสือ “คู่มือคอมพิวเตอร์ในสังคมศาสตร์” (Handbook of Computational Social Science) ออกมาเป็นจำนวนสองเล่มเมื่อปี ค.ศ. 2022 ที่ผ่านมานี้เอง

ต่อมาหลังการระดมองค์ความรู้ในสาขาปัญญาประดิษฐ์มาเป็นเวลายาวนานผ่านช่วง “ฤดูหนาวเอไอ” ในที่สุดเมื่อวันที่ 30 เดือนพฤศจิกายน ปี 2022 บริษัทโอเพ่นเอไอได้เปิดตัวให้บริการปัญญาประดิษฐ์สร้างแบบโมเดลภาษาขนาดใหญ่และประสบความสำเร็จอย่างสูง ปัญญาประดิษฐ์สร้างนี้มีความสามารถมากกว่าเทคโนโลยีการเรียนรู้ของเครื่องจักรดังที่เคยเป็นมา เนื่องจากมันมีความสามารถ “ระบบการแทนที่ความเป็นเหตุเป็นผลทางอ้อม” (Meta-Causal Representation) จนสามารถมีความเข้าใจแบบที่เรียกว่า “ความเข้าใจแบบจักรกล” ซึ่งผู้เขียนตีความว่าไม่ใช่แค่การสร้างโมเดลจากแบบจำลองทางสถิติธรรมดาทั่วไป แต่ความเข้าใจแบบจักรกล กับความเข้าใจของมนุษย์ ยังมีช่องว่างห่างระหว่างกันอยู่มาก หลายครั้งที่ปัญญาประดิษฐ์สร้างแสดง “ความหลอน” (hallucination) ออกมา ซึ่งผู้นำปัญญาประดิษฐ์สร้างไปอิมพลีเม้นต์ให้เป็นระบบอัตโนมัติจะต้องหมั่นตรวจสอบความเที่ยงตรงของผลลัพธ์ที่ได้อยู่เสมอ และจำเป็น

จะต้องเข้าใจมาตรการต่างๆ ในการลดปัญหาที่เกิดจากความหลอนของปัญญาประดิษฐ์แบบรูปร่างนี้ไปพร้อมกันด้วย ผู้เขียนได้แสดงให้เห็นผ่านการทดลองสร้างระบบที่ชื่อว่า “PulsarWave” และประสบความสำเร็จในระดับนานาชาติแสดงให้เห็นว่ามาตรการที่ผู้เขียนนำเสนอเป็นที่ยอมรับในระดับสากลในเวลา

กิตติกรรมประกาศ

ผู้เขียนขอขอบคุณศูนย์ทรัพยากรคอมพิวเตอร์เพื่อการคำนวณขั้นสูง (ThaiSC) สวทช. สำหรับการจัดหาทรัพยากรคอมพิวเตอร์ในงานวิจัยประกอบบทความชิ้นนี้

เอกสารอ้างอิง

- กอบเกียรติ สระอุบล. (2565). *เรียนรู้ AI: Deep Learning ด้วย Python*. สำนักพิมพ์ อินเทอร์เน็ตมีเดีย.
- ปริญญา สงวนสัตย์. (2562). *Artificial Intelligence with Machine Learning, AI สร้างได้ด้วยแมชชีนเลิร์นนิง*. โอดีซี พรีเมียร์.
- Asif, M. A. (2019). Technologies of Power - From Area Studies to Data Science. *Journal for Digital Cultures. Spectres of AI*, Nr. 5, S. 2019: 1–13. <https://doi.org/10.25969/mediarep/13498>.
- Bengio, Y., Lamblin, P., Popovici, D., and Larochelle, H.. (2006). Greedy layer-wise training of deep networks. In *Proceedings of the 19th International Conference on Neural Information Processing Systems (NIPS'06)*. MIT Press, Cambridge, MA, USA, 153-160.
- Cioffi-Revilla, C. (2022). *The Scope of Computational Social Science*. In Handbook of Computational Social Science: Theory, Case Studies and Ethics, Volume 1, edited by Uwe Engel, Anabel Quan-Haase, Sunny Xun Liu and Lars Lyberg, pp. 17-32. London: Routledge.
- De Veaux, R. D. and Eck A. (2022). *Machine Learning Methods for Computational Social Science*. In Handbook of Computational Social Science: Theory, Case Studies and Ethics, Volume 2, edited by Uwe Engel, Anabel Quan-Haase, Sunny Xun Liu and Lars Lyberg, pp. 291-321. Routledge.
- Géron, A. (2017). *Hands-On Machine Learning with Scikit-Learn and TensorFlow: Concepts, Tools, and Techniques to Build Intelligence Systems*. O'Reilly Media, Inc.
- González, R. J. (2022). *War Virtually: The Quest to Automate Conflict, Militarize Data, and Predict the Future*. University of California Press.
- Hinton, G. E., Osindero, S., and Teh Y. W. (2006). A fast learning algorithm for deep belief nets. *Neural Comput*, 18(7), 1527-54. <http://www.doi.org/10.1162/neco.2006.18.7.1527>. PMID: 16764513.

- Hong QN, Pluye P, Fàbregues S, Bartlett G, Boardman F, Cargo M, Dagenais P, Gagnon M-P, Griffiths F, Nicolau B, O’Cathain A, Rousseau M-C, Vedel I. (2018). *Mixed Methods Appraisal Tool (MMAT), version 2018*. Registration of Copyright (#1148552), Canadian Intellectual Property Office, Industry Canada.
- Jacobsen, A. (2015). *The Pentagon’s Brain: An Uncensored History of DARPA, America’s Top Secret Military Research Agency*. Little, Brown and Company.
- Jobin, A., Lenca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature machine intelligence*, 1(9), 389-399.
- King G., Keohane, R. O., Verba, S. (1994). *Designing Social Inquiry: Scientific Inference in Qualitative Research*. Princeton University Press.
- Kublik, S., and Saboo, S. (2022). *GPT-3: Building Innovative NLP Products Using Large Language Model*. O’Reilly Media, Inc.
- Lecun, Y., Bottou, L., Bengio Y., and Haffner P. (1998). Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11), 2278-2324, <http://www.doi.org/10.1109/5.726791>.
- Levine, R., Drang, D. E., and Edelson, B. (1991). *AI and Expert Systems: A Comprehensive Guide* (Second Edition). McGraw-Hill Book Co-Singapore.
- Liang, W., Zhang, Y., Cao, H., Wang, B., Ding, D. Y., Yang, X., ... & Zou, J. (2023). *Can large language models provide useful feedback on research papers? A large-scale empirical analysis*. arXiv preprint arXiv: 2310.01783.
- Mahoney, J. (2010). After KKV: The New Methodology of Qualitative Research. *World Politics*, 62, 120-147. <http://doi.org/10.1017/S0043887109990220>
- Mazzucato, M. (2015). *The Entrepreneurial State: Debunking Public vs Private Sector Myths*. PublicAffairs.
- McCulloch, W.S. and Pitts, W. A. (1943). logical calculus of the ideas immanent in nervous activity. *Bulletin of Mathematical Biophysics*, 5, 115–13. <https://doi.org/10.1007/BF02478259>
- Morgan, D. L. (2007). Paradigm Lost and Pragmatism Regained: Methodological Implications of Combining Qualitative and Quantitative Methods. *Journal of Mixed Methods Research*, 1, 48. <https://doi.org/10.1177/2345678906292462>
- Rodriguez, M. Y. and Storer, H. (2019). A computational social science perspective on qualitative data exploration: Using topic models for the descriptive analysis of social media data, *Journal of Technology in Human Services*, <http://doi.org/10.1080/15228835.2019.1616350>

- Rosenblatt, F. (1958). The perceptron: A probabilistic model for information storage and organization in the brain. *Psychological Review*, 65(6), 386-408.
<https://doi.org/10.1037/h0042519>
- Rumelhart, D., Hinton, G. & Williams, R. (1986). Learning representations by back-propagating errors. *Nature*, 323, 533-536. <https://doi.org/10.1038/323533a0>
- Russel, S., and Norvig, P. (2022). *Artificial Intelligence: A Modern Approach* (Fourth Edition) Global Edition. Pearson Education Limited.
- Turing, A. (1950). Computing Machinery and Intelligence. *Mind, New Series*, 59(236), (Oct., 1950), 433-460.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., ... & Polosukhin, I. (2017). Attention is all you need. *Advances in neural information processing systems*, 30.
- Woodhouse, I. H. (2021). On ‘ground’ truth and why we should abandon the term. *J. Appl. Rem. Sens.* 15(4), 041501. <https://doi.org/10.1117/1.JRS.15.041501>.
- Zhao, W. X., Zhou, K., Li, J., Tang, T., Wang, X., Hou, Y., ... & Wen, J. R. (2023). *A survey of large language models*. arXiv preprint arXiv: 2303.18223.