

การศึกษาปัญญาประดิษฐ์เพื่อการใช้งานอย่างสร้างสรรค์ THE STUDY ARTIFICIAL INTELLIGENCE FOR CREATIVE USE

ไพสันดี ศรีแปง¹ และธันวา ชัยแก้ว²

มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย^{1, 2}

Paisan Sripaeng¹ and Thanwa Chaikaw²

Mahachulalongkornrajavidyalaya University^{1, 2}

Email: paisan252119@gmail.com¹

(Received : October 28, 2025; Revised December 16, 2025; Accepted 29 December 29, 2025)



บทคัดย่อ

ในยุคปัจจุบัน ปัญญาประดิษฐ์ (Artificial Intelligence: AI) ได้เข้ามามีบทบาทสำคัญในการเปลี่ยนแปลงโครงสร้างทางเศรษฐกิจ สังคม และวิถีชีวิตของมนุษย์อย่างรวดเร็วและรอบด้าน AI ช่วยเพิ่มประสิทธิภาพการทำงานในหลายภาคส่วน ในด้านธุรกิจ การเสริมสร้างนวัตกรรมทางการแพทย์ การพัฒนาการศึกษาแบบใหม่ และการยกระดับคุณภาพชีวิตในหลายมิติ แต่การพัฒนาและการใช้งาน AI อย่างกว้างขวางย่อมมาพร้อมกับความเสี่ยงและผลกระทบที่ซับซ้อน ทั้งในด้านการสูญเสียงาน ความปลอดภัยไซเบอร์ การละเมิดความเป็นส่วนตัว และการสร้างข้อมูลเท็จที่อาจบั่นทอนเสถียรภาพทางสังคมและประชาธิปไตย บทความนี้มีวัตถุประสงค์เพื่อวิเคราะห์องค์ประกอบและการทำงานของ AI รวมถึงการสำรวจประโยชน์ อันตราย และข้อถกเถียงเชิงจริยธรรมที่เกี่ยวข้องกับเทคโนโลยีดังกล่าว พร้อมทั้งนำเสนอแนวทางการกำกับดูแลและข้อเสนอแนะเชิงนโยบายเพื่อสนับสนุนการพัฒนา AI อย่างมีความรับผิดชอบ โปร่งใส และเคารพต่อสิทธิมนุษยชน การสร้างสมดุลระหว่างการใช้ประโยชน์จากศักยภาพของ AI กับการป้องกันผลกระทบเชิงลบถือเป็นความท้าทายสำคัญที่ทุกภาคส่วนต้องร่วมกันกำหนดทิศทางและออกแบบกรอบการกำกับดูแลอย่างรอบคอบ

คำสำคัญ: ปัญญาประดิษฐ์, AI, จริยธรรม, ความปลอดภัยไซเบอร์, การกำกับดูแล AI

Abstract

In the present era, Artificial Intelligence (AI) plays a crucial role in rapidly and comprehensively transforming economic, social, and human lifestyles. AI enhances operational efficiency across various sectors, including business, the development of medical innovations, the advancement of new educational models, and the improvement of quality of life in many dimensions. However, the widespread development and use of AI also come with complex risks and impacts, including job loss, cybersecurity issues, privacy violations, and the

creation of false information that could undermine social stability and democracy. This article aims to analyze the components and workings of AI, explore its benefits, dangers, and ethical debates related to the technology, and propose regulatory frameworks and policy recommendations to support the responsible, transparent development of AI that respects human rights. Striking a balance between harnessing the potential of AI and mitigating its negative impacts is a critical challenge that all sectors must address by collaboratively setting directions and designing comprehensive and sustainable governance frameworks.

Keyword: Artificial Intelligence, AI, Ethics, Cybersecurity, AI Governance

บทนำ

ในยุคที่เทคโนโลยีเคลื่อนตัวอย่างรวดเร็วเหนือขอบเขตจินตนาการของมนุษย์ หนึ่งในปรากฏการณ์สำคัญที่สุดที่กำลังพลิกโฉมโลกคือ ปัญญาประดิษฐ์ (Artificial Intelligence: AI) จากเครื่องมือที่ถูกพัฒนาขึ้นเพื่อช่วยแก้ปัญหาทางเฉพาะทาง AI ได้พัฒนามาเป็น “ผู้ช่วยอัจฉริยะ” ที่สามารถคิด วิเคราะห์ และตัดสินใจได้ในระดับที่ท้าทายบทบาทของมนุษย์ในหลายมิติ เทคโนโลยีนี้ได้กลายเป็นพลังขับเคลื่อนที่ทรงอิทธิพลทั้งต่อระบบเศรษฐกิจ การเมือง สังคม และวัฒนธรรมโลกอย่างรวดเร็ว (Brynjolfsson & McAfee, 2014)

ปัจจุบัน AI ได้ถูกนำมาใช้ในวงกว้าง ตั้งแต่การพัฒนา ระบบแนะนำสินค้าบนแพลตฟอร์มออนไลน์ การประมวลผลภาพถ่ายทางการแพทย์ การบริหารจัดการทรัพยากรองค์กร ไปจนถึง รถยนต์ไร้คนขับ และ ระบบสนทนาอัตโนมัติ ที่สามารถเข้าใจและโต้ตอบภาษามนุษย์ได้อย่างแม่นยำ (Russell & Norvig, 2021) ทั้งนี้ ความก้าวหน้าในเทคนิค Machine Learning และ Deep Learning ช่วยให้ AI สามารถเรียนรู้จากข้อมูลขนาดใหญ่และปรับปรุงประสิทธิภาพของตนเองได้อย่างต่อเนื่อง (Russell & Norvig, 2021) วิวัฒนาการทางเทคนิคของ AI ก้าวล้ำอย่างรวดเร็ว โดยมีองค์ประกอบหลัก ได้แก่ Machine Learning, Deep Learning, Natural Language Processing (NLP), Computer Vision, Control and Decision Making, Data Processing, Neural Network, Algorithm, และ Big Data (Gebu et al., 2020) ทั้งนี้ การพัฒนา AI แบบ Generative AI ซึ่งสามารถ “สร้างข้อมูลใหม่” ได้อย่างทรงพลัง ไม่ว่าจะเป็น ข้อความใหม่, ภาพ, วิดีโอ, หรือแม้แต่ เสียงจำลอง ได้ยกระดับศักยภาพของ AI สู่ระดับใหม่อย่างที่ไม่เคยมีมาก่อน (Russell & Norvig, 2021) แม้ AI จะได้ผลตอบรับด้านโอกาสเชิงนวัตกรรม ที่กว้างขวาง ตั้งแต่ การแพทย์เฉพาะทาง ไปจนถึง การศึกษาแบบยืดหยุ่นไร้ขอบเขต และ บริการสาธารณะอัจฉริยะ แต่ อันตรายและความเสี่ยง ของ AI ก็ทวีความรุนแรงและซับซ้อนมากขึ้นเช่นกัน (Zuboff, 2019) ประเด็นที่น่ากังวลได้แก่ การสูญเสียงาน ในหลายอาชีพ, การละเมิดความเป็นส่วนตัว, การสร้างข้อมูลเท็จและข่าวปลอม ผ่าน Generative AI, การโจมตีทางไซเบอร์ เช่น Zero-Day attacks และ มัลแวร์อัจฉริยะ, รวมถึง ความเสี่ยงต่อความสัมพันธ์ของมนุษย์ และ ผลกระทบเชิงวัฒนธรรมจาก AI ที่สร้าง “ความจริงใหม่” ผ่าน Generative AI ที่ลดความน่าเชื่อถือของข้อมูล บนฐานที่

สังคมยังไม่มีกรอบกำกับดูแลที่เพียงพอ (Couldry & Mejias, 2019; SANS Institute, 2022) แต่ AI ก็ได้สร้างคำถามเชิงจริยธรรม และข้อถกเถียงทางสังคม อย่างกว้างขวาง ไม่ว่าจะเป็นประเด็นเรื่อง ความโปร่งใสในการตัดสินใจของ AI, ผลกระทบต่อสิทธิความเป็นส่วนตัว, การใช้งานด้วยหุ่นยนต์ทำงานแทนแรงงานคน ในระบบการผลิตในอุตสาหกรรมต่าง ๆ, หรือแม้แต่ความเสี่ยงที่ AI อาจถูกนำไปใช้ในทางที่ก่อให้เกิด ผลกระทบเชิงลบต่อความมั่นคงของสังคม ปรากฏการณ์เหล่านี้ชี้ให้เห็นว่า AI ไม่ได้เป็นเพียงเทคโนโลยีใหม่ แต่เป็น พลังทางสังคม ที่กำลังเปลี่ยนแปลงโครงสร้างของโลกอย่างมีนัยสำคัญ

บทความนี้มีจุดมุ่งหมายเพื่อศึกษาวิเคราะห์องค์ประกอบและการทำงานของ AI รวมถึงประโยชน์และอันตรายที่อาจเกิดขึ้นจากการพัฒนาและการใช้งาน AI พร้อมนำเสนอข้อถกเถียงเชิงวิชาการที่เกี่ยวข้อง และแนวทางเชิงนโยบายสำหรับการกำกับดูแล AI อย่างมีความรับผิดชอบ โดยคำนึงถึงจริยธรรมและสิทธิมนุษยชน การสร้างสมดุลระหว่างการขับเคลื่อนนวัตกรรมกับการปกป้องสังคมจึงเป็นภารกิจสำคัญที่ทุกภาคส่วนต้องร่วมกันกำหนดทิศทางในอนาคต ข้อเสนอแนะด้านการกำกับดูแล AI จะคำนึงถึงความซับซ้อนเชิงอำนาจของข้อมูลและความจำเป็นในการรักษาสมดุลระหว่างการขับเคลื่อนนวัตกรรมกับการคุ้มครองสิทธิและเสรีภาพของประชาชน เพื่อศึกษาข้อถกเถียงเชิงจริยธรรมที่เกี่ยวข้องกับการพัฒนาและการใช้งาน AI โดยเฉพาะประเด็นเรื่องความเป็นส่วนตัว ความเป็นธรรม ความโปร่งใส และความรับผิดชอบ

ความเป็นมาของปัญญาประดิษฐ์

ปัญญาประดิษฐ์ (Artificial Intelligence) หรือ AI คือ โปรแกรมที่แสดงออกเลียนแบบสติปัญญาของมนุษย์ผ่านการเรียนรู้อัลกอริทึม (กระบวนการและขั้นตอนที่ถูกตั้งไว้สำหรับการแก้ปัญหา) ให้เหตุผลและแก้ไขอัลกอริทึมนั้น ๆ เพื่อปรับปรุงผลลัพธ์ให้ออกมามีประสิทธิภาพมากที่สุด โดย AI จะมีโมเดลพื้นฐาน (Foundation Model) ซึ่งเปรียบเสมือนสมองที่ได้รับการฝึกฝนทดสอบลักษณะของข้อมูลและประมวลผลออกมา เช่น Large Language Model หรือที่รู้จักกันในนาม LLM เช่น ChatGPT, Claud และ Image Model ที่มี input เป็นภาพ ซึ่ง Generative AI ก็เป็นหนึ่งในประเภทของ AI ที่สามารถเรียนรู้ลักษณะของข้อมูล ให้เข้าใจว่าข้อมูลน่าจะถูกสร้างมาอย่างไร (Distribution of data) เพื่อสามารถสร้างข้อมูลใหม่ขึ้นเองได้ เช่น การสร้างข้อความใหม่ ๆ การสกัดข้อมูลจากข้อความ การจำแนกข้อความ และการสรุปข้อความ ในธุรกิจต่าง ๆ ณ ปัจจุบันจึงมีการนำ AI เข้ามาใช้งานมากมายซึ่งจะมีทั้งในรูปแบบ Predictive AI และ Generative AI โดยอาจจะถูกนำไปใช้ประโยชน์ที่หลากหลาย เช่น การทำ Process Automation, Cognitive Insight, และ Cognitive Engagement เป็นต้น

ปัญญาประดิษฐ์ (Artificial Intelligence : AI) ได้กลายเป็นหนึ่งในปรากฏการณ์ทางเทคโนโลยีที่มีอิทธิพลสูงสุดต่อโครงสร้างทางเศรษฐกิจ สังคม และวัฒนธรรมของโลกยุคใหม่ ในช่วงไม่กี่ทศวรรษที่ผ่านมา ความก้าวหน้าอย่างก้าวกระโดดของ AI ไม่เพียงเปลี่ยนวิถีชีวิตของคนในระดับปัจเจก หากยังมีนัยสำคัญต่อพลวัตของอำนาจ กระบวนการสร้างความรู้ และเสรีภาพของมนุษย์ในสังคมร่วมสมัย (Zuboff, 2019) งานวิจัยและเอกสารจำนวนมากในช่วงหลังจึงได้พยายามวิเคราะห์ผลกระทบเชิง

โครงสร้างของ AI รวมถึงข้อถกเถียงเชิงจริยธรรมที่ล้อมรอบการพัฒนาและการทำงานของเทคโนโลยีนี้ ในเชิงเทคนิค AI มีองค์ประกอบที่หลากหลาย โดยเริ่มตั้งแต่ Machine Learning (ML) ที่ช่วยให้ระบบสามารถเรียนรู้จากข้อมูลโดยไม่ต้องพึ่งพาการเขียนโปรแกรมแบบตายตัว ต่อมาพัฒนาสู่ Deep Learning ที่ใช้โครงข่ายประสาทเทียม (Neural Networks) หลายชั้น ทำให้สามารถประมวลผลข้อมูลที่ซับซ้อน เช่น ภาพ เสียง และภาษาธรรมชาติ ได้อย่างมีประสิทธิภาพ (Goodfellow et al., 2016) นอกจากนี้ Natural Language Processing (NLP) และ Computer Vision ก็เป็นองค์ประกอบสำคัญที่ทำให้ AI สามารถเข้าใจและสังเคราะห์ภาษา รวมถึงตีความและวิเคราะห์ข้อมูลภาพและวิดีโอได้อย่างแม่นยำ แนวโน้มสำคัญที่เกิดขึ้นในช่วงหลังคือการพัฒนา Generative AI ที่สามารถ "สร้างข้อมูลใหม่" ได้ เช่น ข้อความ ภาพ เสียง และวิดีโอ ส่งผลให้ AI ก้าวข้ามจากการเป็นเครื่องมือเชิงรับ มาเป็นผู้สร้างวาทกรรมและความหมายในสังคม (Couldry & Mejias, 2019)

การเติบโตอย่างรวดเร็วของ AI ได้จุดประกายคำถามเชิงจริยธรรมและผลกระทบเชิงสังคมที่ซับซ้อนมากขึ้น งานวิจัยของ Gebru et al. (2020) ได้วิพากษ์ให้เห็นว่า AI มิใช่ระบบที่เป็นกลาง หากแต่สะท้อนอคติของข้อมูลและกระบวนการออกแบบที่ฝังอยู่ในระบบเทคโนโลยี ข้อมูลที่ใช้ในการฝึก AI มักมีอคติเชิงเชื้อชาติ เพศ และชนชั้นอย่างไม่รู้ตัว ส่งผลให้การตัดสินใจของ AI ในหลากหลายบริบท เช่น การคัดกรองผู้สมัครงาน การอนุมัติสินเชื่อ หรือการประเมินความเสี่ยงทางอาชญากรรม นอกจากนี้ การใช้ AI ในการเก็บและวิเคราะห์ข้อมูลส่วนบุคคลยังละเมิดสิทธิความเป็นส่วนตัวในระดับลึก โดยเฉพาะเมื่อข้อมูลถูกใช้เพื่อสร้างโปรไฟล์ผู้บริโภค หรือเพื่อควบคุมพฤติกรรมของประชาชนในบริบททางการเมืองและเศรษฐกิจ (Zuboff, 2019)

ภัยคุกคามทางไซเบอร์ก็เป็นอีกมิติหนึ่งที่ AI เข้ามามีบทบาทอย่างน่ากังวล การใช้ AI ในการพัฒนา Zero-Day attacks หรือมัลแวร์ขั้นสูงที่สามารถปรับตัวได้แบบ real-time ทำให้ระบบความปลอดภัยไซเบอร์ดั้งเดิมไม่สามารถรับมือได้อย่างมีประสิทธิภาพ (SANS Institute, 2022) นอกจากนี้ การเกิดขึ้นของ Generative AI ยังเปิดช่องให้เกิดการผลิตข้อมูลเท็จ (Fake News) และ Deepfake ที่มีความสมจริงสูง ส่งผลกระทบต่อกระบวนการประชาธิปไตย และความน่าเชื่อถือของข้อมูลสาธารณะ ประเด็นเหล่านี้สะท้อนให้เห็นถึงข้อถกเถียงเชิงจริยธรรมที่สำคัญเกี่ยวกับ AI ได้แก่ ความเป็นธรรม (Fairness), ความโปร่งใส (Transparency), ความรับผิดชอบ (Accountability), และการเคารพสิทธิความเป็นส่วนตัว (Privacy) (Gebru et al., 2020) AI ที่ถูกออกแบบมาโดยขาดความตระหนัก Couldry & Mejias (2019) ชี้ว่า AI และ Big Data ได้กลายเป็นระบบความรู้ใหม่ ที่มีอิทธิพลต่อการกำหนดสิ่งที่สามารถรู้ได้ (epistemic boundaries) และสิ่งที่สามารถทำได้ทุกอย่าง (ontological possibilities) ในสังคมยุคดิจิทัล การที่ข้อมูลและอัลกอริทึมถูกควบคุมโดยกลุ่มทุนหรืออำนาจรัฐจึงสร้างความเสี่ยงเชิงโครงสร้างต่อประชาธิปไตย และเสรีภาพในการเข้าถึงความรู้

อันตรายของปัญญาประดิษฐ์ (AI)



แผนภาพที่ 1 ระวัง!! อาชญากรรมออนไลน์ ที่มาพร้อมกับ AI

ที่มา: ศูนย์ต่อต้านข่าวปลอม ประเทศไทย

แม้ปัญญาประดิษฐ์ (AI) จะถูกยกย่องว่าเป็นเทคโนโลยีปฏิวัติโลก แต่การเติบโตอย่างรวดเร็วของ AI ก็มาพร้อมกับความเสี่ยงที่สังคมต้องตระหนักรู้ การแทนที่งานมนุษย์และผลกระทบทางเศรษฐกิจ AI กำลังเปลี่ยนแปลงตลาดแรงงานอย่างรวดเร็ว โดยเฉพาะงานที่มีรูปแบบซ้ำๆ เช่น งานผลิต งานบริการลูกค้า และแม้แต่งานวิเคราะห์ข้อมูลบางประเภท การศึกษาของ Brynjolfsson และ McAfee (2014) ชี้ว่าเทคโนโลยีนี้จะทำให้งานหลายตำแหน่งหายไป ส่งผลกระทบต่อปัจเจกบุคคลและเศรษฐกิจในภาพใหญ่ การละเมิดความเป็นส่วนตัวในยุคข้อมูลคืออำนาจ กล่าวคือ AI สามารถรวบรวมและวิเคราะห์ข้อมูลส่วนบุคคลได้ละเอียดกว่ายุคใดๆ ที่ผ่านมา บริษัทและรัฐบาลใช้ AI ติดตามพฤติกรรมผู้คน เช่น การโฆษณาแบบกำหนดเป้าหมายที่ละเมิดความเป็นส่วนตัวการสอดส่องประชาชนโดยรัฐ (Zuboff, 2019) ปัญหาคือ ผู้ใช้ส่วนใหญ่ไม่รู้ว่าข้อมูลของตนถูกนำไปใช้อย่างไร หรือแม้แต่ให้ความยินยอมอย่างแท้จริงหรือไม่ ภัยคุกคามทางไซเบอร์ที่ทวีความรุนแรงขึ้น โดย AI ถูกใช้ทั้งป้องกันและโจมตีระบบดิจิทัล เช่น มัลแวร์ที่ปรับตัวได้เอง การโจมตีแบบ Zero-Day ที่หาช่องโหว่ได้รวดเร็ว บทหนึ่งคือจริยศาสตร์ที่โจมตีแบบอัตโนมัติ (SANS Institute, 2022) ในอนาคต AI อาจพัฒนาระบบโจมตีที่ทำงานโดยไม่ต้องพึ่งมนุษย์ ทำให้ภัยไซเบอร์อันตรายขึ้นมาก ตามมาด้วยข้อมูลเท็จและ Deepfake AI สร้างเนื้อหาปลอมที่สมจริง เช่น วิดีโอ Deepfake ปลอมตัวบุคคลสำคัญ ข่าวปลอมที่โน้มน้าวความคิดคนหมู่เสี่ยงเลียนแบบที่ใช้ในการหลอกลวงสิ่งนี้ทำให้สังคมแยกแยะความจริงจากเรื่องแต่งได้ยากขึ้น ส่งผลต่อประชาธิปไตยและความเชื่อมั่นในสื่อ (Couldry & Mejias, 2019) อีกทั้ง AI ขี่นำอคติทางสังคม ระบบ AI เรียนรู้จากข้อมูลมนุษย์ ซึ่งมักมีอคติแฝงอยู่ เช่น อคติทางเชื้อชาติ เพศ หรือชนชั้น ทำให้ระบบคัดเลือกงานอาจกีดกันกลุ่มบางกลุ่ม การอนุมัติสินเชื่อหรือการประเมินความเสี่ยงอาจไม่เป็นธรรม และ

กระบวนการยุติธรรมอาจได้รับอิทธิพลจากอัลกอริทึมที่มีอคติ (Gebru et al., 2020) และ AI สร้างความไม่เท่าเทียมทางอำนาจ ข้อมูลและ AI กลายเป็นเครื่องมือของกลุ่มอำนาจ ไม่ว่าจะเป็นบริษัทใหญ่หรือรัฐบาล หากขาดการตรวจสอบ เช่นประชาชนอาจถูกควบคุมผ่านอัลกอริทึม อำนาจกระจุกตัวอยู่ในมือคนกลุ่มเล็กๆ เสรีภาพและความเป็นประชาธิปไตยถูกคุกคาม

จริยธรรมกับการกำกับดูแลอย่างรับผิดชอบ

เพื่อรับมือกับความเสี่ยงและผลกระทบเหล่านี้ งานวิจัยหลายชิ้นได้เสนอแนวทางการกำกับดูแล AI ที่หลากหลาย เช่น การสร้างกรอบกฎหมายและมาตรฐานสากลที่สามารถปรับตัวได้ตามพลวัตของเทคโนโลยี (Zuboff, 2019), การกำหนดเกณฑ์จริยธรรมขั้นต่ำในการออกแบบและใช้งาน AI (Gebru et al., 2020), การพัฒนา Explainable AI ที่โปร่งใสและสามารถตรวจสอบได้ (Couldry & Mejias, 2019), การจัดตั้งกลไกตรวจสอบแบบอิสระที่มีส่วนร่วมจากภาคประชาชน และการเสริมสร้าง AI literacy ให้กับประชาชนเพื่อป้องกันการถูกชักนำโดยข้อมูลปลอม และการใช้งาน AI ที่ขาดความรับผิดชอบ จากข้อมูลที่ได้กล่าวมาผู้เขียนเสนอแนวคิดเพื่อลดความเสี่ยงเหล่านี้ โดยสังคมต้อง

1) การออกกฎหมายควบคุม AI ที่โปร่งใสและเป็นธรรม ปัจจุบันการพัฒนา AI ก้าวหน้าเร็วกว่ากรอบกฎหมาย ส่งผลให้เกิด ช่องว่างทางกฎหมาย ที่อาจถูกใช้ในทางที่ผิด (Zuboff, 2019) จึงเสนอให้ ออกกฎหมายเฉพาะ กำหนดหลักการพื้นฐาน 4 ประการ ความโปร่งใส (เปิดเผยแหล่งข้อมูลและหลักการทำงาน) ความเป็นธรรม (ป้องกันการเลือกปฏิบัติ) ความรับผิดชอบ (กำหนดผู้รับผิดชอบหากเกิดความเสียหาย) การคุ้มครองผู้เสียหาย (กลไกเยียวยาผู้ได้รับผลกระทบ) และจัดตั้งหน่วยงานกำกับดูแลอิสระ เพื่อตรวจสอบการใช้ AI ในภาคส่วนสำคัญ เช่น สาธารณสุข การเงิน และกระบวนการยุติธรรม

2) การตรวจสอบอคติในระบบ AI อย่างเป็นระบบเนื่องจาก AI มักสืบทอด อคติทางสังคม ซึ่งมักสะท้อน อคติทางเชื้อชาติ เพศ อายุ ชนชั้น และปัจจัยทางสังคมอื่นๆ จากข้อมูลฝึกสอน (Gebru et al., 2020)

3) ส่งเสริมการรู้เท่าทันดิจิทัล เพื่อให้ประชาชนรับมือกับข้อมูลเท็จ เสริมสร้างการรู้เท่าทันดิจิทัล เพื่อรับมือกับ ข้อมูลเท็จจาก AI (Couldry & Mejias, 2019) ของประชาชนจึงเป็นภูมิคุ้มกันที่สำคัญ ที่จะช่วยให้สังคมสามารถรับมือกับ ข้อมูลเท็จและการบิดเบือน ได้อย่างมีประสิทธิภาพ ภาครัฐและภาคการศึกษาควรจัดทำ หลักสูตรและกิจกรรม เพื่อเสริมสร้างความสามารถในการ วิเคราะห์วิจารณ์ และ ตรวจสอบความน่าเชื่อถือของข้อมูล ให้แก่ประชาชนทุกกลุ่มวัย โดยเฉพาะเยาวชน ซึ่งเป็นกลุ่มที่มีโอกาสเผชิญกับข้อมูลเท็จผ่าน โซเชียลมีเดีย ในชีวิตประจำวัน นอกจากนี้ควรสนับสนุน การวิจัยและพัฒนาเครื่องมือทางเทคโนโลยี ที่ช่วยให้ผู้ใช้งานสามารถตรวจจับและระบุเนื้อหาที่ถูกสร้างขึ้นโดย AI

4) สร้างความร่วมมือระหว่างประเทศ เพื่อป้องกันการใช้ AI ในทางที่ผิด AI เป็นเทคโนโลยีที่มี ลักษณะไร้พรมแดน (Borderless Technology) การควบคุมและกำกับดูแล AI อย่างมี

ประสิทธิภาพจึงไม่สามารถพึ่งพา กฎหมายภายในประเทศ เพียงอย่างเดียวได้ ความร่วมมือระหว่างประเทศเป็นสิ่งจำเป็นเพื่อสร้าง มาตรฐานสากล และ กรอบความร่วมมือระหว่างรัฐ ที่จะช่วยป้องกันการพัฒนาและการใช้งาน AI ในทางที่ผิด

สรุป

จากการศึกษาพบว่า การพัฒนาและประยุกต์ใช้ AI มีทั้งประโยชน์และความเสี่ยงที่ต้องพิจารณาอย่างรอบคอบ ในด้านบวก AI สามารถเพิ่มประสิทธิภาพการทำงาน สร้างนวัตกรรม และพัฒนาระบบบริการในหลายสาขา เช่น การแพทย์ การศึกษา และการบริหารภาครัฐอย่างไ้ก็ตาม ยังหลงเหลือประเด็นสำคัญที่ต้องตระหนัก ได้แก่

- 1) ผลกระทบต่อตลาดแรงงานจากการแทนที่งานบางประเภทด้วยระบบอัตโนมัติ
- 2) ความเสี่ยงด้านความปลอดภัยของข้อมูลส่วนบุคคล
- 3) การใช้เทคโนโลยีในทางที่ผิด เช่น การสร้างข้อมูลเท็จหรือการโจมตีทางไซเบอร์
- 4) การสะท้อนอคติทางสังคมผ่านระบบอัลกอริทึม

ยิ่งไปกว่านั้น การควบคุมข้อมูลและเทคโนโลยีโดยกลุ่มอำนาจบางกลุ่มอาจส่งผลกระทบต่อความมั่นคงของระบบประชาธิปไตยและสิทธิเสรีภาพของประชาชน เพื่อให้การพัฒนา AI เป็นไปอย่างมีความรับผิดชอบ จำเป็นต้องยึดหลักการสำคัญ 4 ประการ คือ

- 1) ความเป็นธรรม (Fairness)
- 2) ความโปร่งใส (Transparency)
- 4) ความรับผิดชอบ (Accountability)
- 4) การคุ้มครองความเป็นส่วนตัว (Privacy)

นอกหลักการตรวจสอบ 4 ประการดังกล่าวมา ต่อจากนี้ยังต้องมีรอบกำกับดูแลที่เหมาะสม โดยคำนึงถึงมิติทางสังคมและจริยธรรม พร้อมทั้งส่งเสริมการมีส่วนร่วมของทุกภาคส่วนในการกำหนดนโยบายที่เกี่ยวข้องเพื่อให้การพัฒนา AI ในอนาคตได้คำนึงถึงผลกระทบทางสังคมและจริยธรรมเป็นสำคัญ โดยมุ่งสร้างสมดุลระหว่างความก้าวหน้าทางเทคโนโลยีกับการรักษาคุณค่าพื้นฐานให้สมดุลกับชีวิตในทุกมิติ

เอกสารอ้างอิง

- Brynjolfsson, E., & McAfee, A. (2014). The second machine age: Work, progress, and prosperity in a time of brilliant technologies. W. W. Norton & Company.
- Couldry, N., & Mejias, U. A. (2019). The costs of connection: How data is colonizing human life and appropriating it for capitalism. Stanford University Press.

- Gebru, T., Morgenstern, J., Vecchione, B., Vaughan, J. W., Wallach, H., Daumé III, H., & Crawford, K. (2020). Datasheets for datasets. *Communications of the ACM*, 64(12), 86–92. <https://doi.org/10.1145/3458723>
- Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep learning*. MIT Press.
- Russell, S., & Norvig, P. (2021). *Artificial intelligence: A modern approach* (4th ed.). Pearson.
- SANS Institute. (2022). Cybersecurity trends and emerging threats. <https://www.sans.org/white-papers/40245/>
- Zuboff, S. (2019). *The age of surveillance capitalism: The fight for a human future at the new frontier of power*. PublicAffairs.

