

การใช้ปัญญาประดิษฐ์ในการบังคับใช้กฎหมาย ขอบเขตของจริยธรรม ผลกระทบต่อสิทธิมนุษยชน

The use of artificial intelligence in law enforcement: the scope of ethics and the impact on human rights

เศรษฐากรณ์ วงษ์อารยะสกุล^{1*} และ ดาริณี ปันกันสกุล²

Setthakorn Wongarayasakul¹ and Darinee Pungunsakul²

¹สถาบันนวัตกรรมสังคมเพื่อการพัฒนาที่ยั่งยืน

²วิทยาลัยเกษตรและเทคโนโลยีมหาสารคาม

¹Social Innovation for Sustainable Development Institute

²Mahasarakham College of Agriculture and Technology

* E-mail : Setthakorn824@gmail.com

วันที่รับบทความ : 21 ตุลาคม 2568, วันที่แก้ไขบทความ ครั้งที่ 1 : 25 พฤศจิกายน 2568

วันที่แก้ไขบทความ ครั้งที่ 2 : 17 ธันวาคม 2568

วันที่ตอบรับการตีพิมพ์ : 20 ธันวาคม 2568

บทคัดย่อ

การวิจัยครั้งนี้มีวัตถุประสงค์เพื่อศึกษา 1. ศึกษาขอบเขตของการใช้ปัญญาประดิษฐ์ในการบังคับใช้กฎหมาย 2. ศึกษาประเด็นด้านจริยธรรมและหลักความรับผิดชอบต่อการใช้งาน และ 3. ศึกษาผลกระทบที่อาจเกิดขึ้นต่อสิทธิมนุษยชนในระบบยุติธรรม เป็นงานวิจัยเชิงคุณภาพ ผู้ให้ข้อมูลหลัก คือ ผู้เชี่ยวชาญด้านกฎหมาย ปัญญาประดิษฐ์ และสิทธิมนุษยชน จำนวน 10 ท่าน เก็บข้อมูลโดยศึกษาเอกสาร สัมภาษณ์เชิงลึก และสังเกตภาคสนาม วิเคราะห์เนื้อหาและวิเคราะห์เชิงประเด็น ตรวจสอบคุณภาพข้อมูลด้วย Triangulation และ Member Check ภายใต้อำนาจจริยธรรมการวิจัย เพื่อสะท้อนมุมมองต่อการประยุกต์ใช้ AI ในกระบวนการยุติธรรม

ผลการวิจัย พบว่า การนำปัญญาประดิษฐ์ AI มาใช้บังคับกฎหมายมีขอบเขตกว้าง ครอบคลุมทั้งการเฝ้าระวัง การวิเคราะห์ข้อมูลอาชญากรรม และการพยากรณ์พื้นที่เสี่ยง ซึ่งช่วยเพิ่มประสิทธิภาพและความรวดเร็วในการปฏิบัติงานของเจ้าหน้าที่ อย่างไรก็ตาม มีการพบประเด็นด้านจริยธรรมที่สำคัญ เช่น ปัญหาอคติของอัลกอริทึม ความโปร่งใสของกระบวนการตัดสินใจ และความรับผิดชอบเมื่อเกิดเหตุระบบทำงานผิดพลาด ทำให้ส่งผลกระทบต่อสิทธิมนุษยชน โดยเฉพาะเรื่องสิทธิความเป็นส่วนตัว และสิทธิในการได้รับความเป็นธรรมในกระบวนการยุติธรรม ดังนั้น การใช้ AI ในภาครัฐควรที่จะอยู่ภายใต้หลักของจริยธรรม ความโปร่งใส และกลไกการกำกับดูแล เพื่อป้องกันการละเมิดสิทธิและสร้างความไว้วางใจจากประชาชน

คำสำคัญ : ปัญญาประดิษฐ์, การบังคับใช้กฎหมาย, จริยธรรม, สิทธิมนุษยชน, เทคโนโลยีดิจิทัล

Abstract

This research aims to study: 1. the scope of artificial intelligence use in law enforcement; 2. ethical issues and accountability principles in its use; and 3. the potential impact on human rights in the justice system. It is a qualitative study with a key informant of 10 experts in law, artificial intelligence, and human rights. Data were collected through document review, in-depth interviews, and field observations. The analysis involved content analysis and thematic analysis, with data quality verified through triangulation and member checks, all conducted under research ethics principles, to reflect perspectives on the application of AI in the justice process.

Research findings indicate that the use of artificial intelligence (AI) in law enforcement has a broad scope, covering surveillance, crime data analysis, and risk area prediction, which helps improve the efficiency and speed of officers' operations. However, significant ethical issues have been identified, such as algorithmic bias, transparency in decision-making processes, and accountability when system errors occur, affecting human rights, particularly privacy rights and the right to a fair justice process. Therefore, the use of AI in the public sector should be guided by principles of ethics, transparency, and oversight mechanisms to prevent rights violations and build public trust.

Keywords: Artificial Intelligence, Law Enforcement, Ethics, Human Rights, Digital Technology

บทนำ

ในยุคปัจจุบันปัญหาประดิษฐ์ AI ได้เข้ามามีบทบาทในหลายมิติของชีวิตประจำวัน ตั้งแต่เรื่องการสื่อสาร เรื่องการแพทย์ เรื่องการเงิน ไปจนถึงการพัฒนานวัตกรรมต่างๆ ในอุตสาหกรรม (Grimson E., 2025) ซึ่งชี้ให้เห็นถึงพลังของปัญหาประดิษฐ์ AI ในยุคที่มาจากการพัฒนาเทคโนโลยี Deep Learning และมีการประมวลผลข้อมูลขนาดใหญ่ (Big Data) โดยมีตัวอย่างที่น่าสนใจคือ ปัญหาประดิษฐ์ AI ที่ช่วยตรวจจับโรค โดยการถ่ายภาพทางการแพทย์ซึ่งมีความแม่นยำกว่าแพทย์ผู้เชี่ยวชาญในบางกรณี หรือการนำปัญหาประดิษฐ์ AI มาใช้ในการขนส่ง เช่น รถยนต์ไร้คนขับ หรือการวิเคราะห์ข้อมูล เพื่อช่วยในการตัดสินใจในธุรกิจให้มีประสิทธิภาพ ปัญหาประดิษฐ์ AI ยังเป็นแรงขับเคลื่อนในโลกของโซเชียลมีเดีย แอปพลิเคชันที่เป็นระบบผู้ช่วยอัจฉริยะ เช่น Siri , Alexa ทั้งยังมีการเน้นย้ำถึงความสำคัญของการศึกษาและการพัฒนาทักษะในด้านปัญหาประดิษฐ์ AI โดยเฉพาะในช่วงที่อุตสาหกรรม

กำลังเปลี่ยนแปลงไปอย่างรวดเร็ว ซึ่งเชื่อว่าปัญญาประดิษฐ์ AI จะไม่เพียงแต่จะเป็นเครื่องมือช่วยเหลือเท่านั้น แต่จะกลายเป็นพันธมิตรที่ทรงพลังของมนุษย์อีกด้วย (Russell, S., & Norvig, P. 2021)

ปัญญาประดิษฐ์ AI เริ่มต้นช่วงกลางศตวรรษที่ 20 บุคคลสำคัญอย่าง Alan Turing ผู้ซึ่งบุกเบิกแนวคิดของ เครื่องจักรที่สามารถคิดได้ ซึ่งได้นำไปสู่การสร้างเครื่องคอมพิวเตอร์รุ่นแรก และมีการพัฒนาอัลกอริทึมพื้นฐานขึ้นด้วย และนักวิจัยยังได้มีการเริ่มทดลองสร้างโปรแกรมที่สามารถแก้ปัญหา ซึ่งนับได้ว่าเป็นก้าวแรกที่สำคัญของปัญญาประดิษฐ์ AI ในยุคต่อมาปัญญาประดิษฐ์ AI ได้รับความสนใจมากขึ้น ซึ่งเป็นช่วงที่เกิดยุคความรู้ และการพัฒนาระบบผู้เชี่ยวชาญ ซึ่งสามารถนำไปใช้ในวงการแพทย์ วงการการเงิน และการผลิต แม้ในขณะนั้นปัญญาประดิษฐ์ AI ยังคงมีข้อจำกัดในด้านต่าง ๆ เช่น ด้านความสามารถและทรัพยากรคอมพิวเตอร์ แต่ได้สร้างรากฐานสำคัญให้กับการพัฒนาในปัจจุบันมากขึ้น เมื่อมองไปยังอนาคต ศักยภาพที่ไร้ขีดจำกัดของปัญญาประดิษฐ์ AI เห็นว่าในอนาคตปัญญาประดิษฐ์ AI อาจมีบทบาทในการแก้ไขปัญหาระดับโลก เช่น การเปลี่ยนแปลงสภาพภูมิอากาศ การพัฒนาทางด้านการแพทย์เฉพาะบุคคล และการพัฒนาเศรษฐกิจที่ยั่งยืน ซึ่งความท้าทายที่มาพร้อมกับการพัฒนาปัญญาประดิษฐ์ AI เช่น ประเด็นเรื่องจริยธรรม เรื่องความเป็นส่วนตัว และการกระจายผลประโยชน์ของ ปัญญาประดิษฐ์ AI อย่างเท่าเทียม จึงมีการเรียกร้องให้เกิดการกำกับดูแล และการทำงานร่วมกันระหว่างภาควิชาการ ภาคธุรกิจ และรัฐบาล เพื่อให้ปัญญาประดิษฐ์ AI เติบโตในทางที่เป็นประโยชน์ต่อทุกคน เนื่องจากปัญหาในปัจจุบันยังพบว่า การใช้เทคโนโลยีปัญญาประดิษฐ์ AI ด้านความมั่นคง และตำรวจบางแห่งในต่างประเทศได้มีการรายงานปัญหาอคติของอัลกอริทึม (Algorithmic Bias) การละเมิดสิทธิความเป็นส่วนตัว และการจับกุมผิดพลาดจากระบบจดจำใบหน้า (Angwin et al., 2016; Buolamwini & Gebru, 2018)

ในประเทศไทย หน่วยงานทางด้านกฎหมายได้เริ่มทดลองใช้เทคโนโลยีปัญญาประดิษฐ์ AI ในงานที่เกี่ยวข้องกับการตรวจสอบและวิเคราะห์ข้อมูลทางอาชญากรรม ซึ่งยังขาดกรอบกำกับดูแลในด้านจริยธรรม และด้านสิทธิมนุษยชน จึงทำให้เกิดความเสี่ยงในการละเมิดสิทธิของประชาชนโดยไม่ได้ตั้งใจ (สำนักงานพัฒนารัฐบาลดิจิทัล, 2566; กิตติคุณ ตันรักษ, 2564; TIJ, 2564; OHCHR, 2021) ซึ่งผู้วิจัยได้ศึกษาขอบเขตของการใช้ปัญญาประดิษฐ์ในการบังคับใช้กฎหมาย ในประเด็นทางด้านจริยธรรมและหลักความรับผิดชอบ ซึ่งรวมไปถึงผลกระทบที่อาจเกิดขึ้นต่อสิทธิมนุษยชน และเพื่อเสนอแนวทางการใช้เทคโนโลยีอย่างมีความรับผิดชอบ โปร่งใส และเคารพสิทธิมนุษยชน เพื่อประโยชน์ต่อการพัฒนาระบบยุติธรรมไทยให้ก้าวสู่ยุคดิจิทัลได้อย่างยั่งยืน ปัญญาประดิษฐ์ AI ไม่ได้เป็นเพียงนวัตกรรมทางเทคโนโลยีเท่านั้น แต่เป็นเครื่องมือที่เปลี่ยนวิถีที่เราดำรงชีวิตอยู่ในทุก ๆ วัน และไม่ว่าจะเป็นเรื่องการทำงาน ซึ่งปัญญาประดิษฐ์ AI มีความสามารถในการเปลี่ยนแปลงโลกอย่างมาก ตั้งแต่ในอดีตที่เป็นเพียงแนวคิด จนถึงปัจจุบันที่เป็นหัวใจของการปฏิวัติอุตสาหกรรมต่าง ๆ ไปจนถึงอนาคตที่เต็มไปด้วยความเป็นไปได้ หากเราใช้ปัญญาประดิษฐ์ AI อย่างมีความรับผิดชอบแล้ว และยังมีเป้าหมายเพื่อประโยชน์ส่วนรวม ปัญญาประดิษฐ์ AI จะไม่เพียงช่วยให้เราก้าวไปข้างหน้าเท่านั้น แต่ยังช่วยสร้างโลกที่ดีขึ้นสำหรับทุกคนได้อีกด้วย

วัตถุประสงค์ของการวิจัย

1. เพื่อศึกษาขอบเขตของการใช้ปัญญาประดิษฐ์ในการบังคับใช้กฎหมาย
2. ศึกษาประเด็นด้านจริยธรรมและความรับผิดชอบในการใช้เทคโนโลยีปัญญาประดิษฐ์ AI
3. เพื่อศึกษาผลกระทบต่อสิทธิมนุษยชนของระบบยุติธรรมการนำเทคโนโลยีปัญญาประดิษฐ์ AI มาใช้

การทบทวนวรรณกรรมและแนวคิด

แนวคิดปัญญาประดิษฐ์และวิวัฒนาการของเทคโนโลยีปัญญาประดิษฐ์

ความหมายของปัญญาประดิษฐ์

ในยุคแรกเครื่องคอมพิวเตอร์ส่วนใหญ่นั้นจะมีหน้าที่ในการคำนวณ และประมวลผลโปรแกรมของระบบงานต่าง ๆ ตามที่มนุษย์เกิดความต้องการ โดยมีมนุษย์เป็นผู้ควบคุม ป้อนคำสั่ง และรอผลลัพธ์จากเครื่องคอมพิวเตอร์ ซึ่งการตัดสินใจในสิ่งต่าง ๆ ยังคงอยู่บนพื้นฐานการสังเคราะห์ของมนุษย์ ซึ่งทำให้แนวคิดที่จะพัฒนาให้เครื่องคอมพิวเตอร์ให้มีประสิทธิภาพสูงขึ้น เพียงพอที่จะอำนวยความสะดวกและลดกำลังการสนับสนุนจากมนุษย์ลงให้น้อยที่สุด ดังนั้น จึงมีความจำเป็นต้องพัฒนาเครื่องคอมพิวเตอร์ให้มีความคิด มีหลักการที่สมเหตุสมผล เทียบเคียงกับสมองของมนุษย์ และเพื่อให้เครื่องคอมพิวเตอร์สามารถวิเคราะห์ข้อมูลเพื่อตัดสินใจว่าสิ่งใดถูกต้องตามจุดประสงค์ของการทำงาน และด้วยเหตุนี้จึงเป็นที่มาของการพัฒนาสมองของเครื่องคอมพิวเตอร์ ให้มีกระบวนการคิดที่สมบูรณ์แบบแทนการคิดของมนุษย์ ที่เรียกว่า ปัญญาประดิษฐ์ (Artificial Intelligence) (Russell & Norvig, 2021).

ปัญญาประดิษฐ์ หรือเรียกว่า AI (Artificial Intelligence) เป็นแนวทางการพัฒนาทางด้านคอมพิวเตอร์อีกรูปแบบหนึ่ง ซึ่งทำให้เครื่องคอมพิวเตอร์ได้คิดและตัดสินใจได้ซึ่งมีความใกล้เคียงกับมนุษย์ มีหลักการจากการศึกษาวิธีการคิด การตัดสินใจ หรือหลักของเหตุและผลจากมนุษย์ เพื่อนำไปใช้พัฒนาศักยภาพของเครื่องคอมพิวเตอร์ให้ตอบสนองการทำงานที่มากกว่าการเป็นเพียงแค่เครื่องจักรกลหรือโปรแกรมทั่ว ๆ ไป โดยเริ่มจากการนำแนวคิดมากำหนดเป็นขั้นตอนให้เครื่องคอมพิวเตอร์ได้ทำงาน ได้แก้ปัญหา ได้ตัดสินใจ และได้เรียนรู้ด้วยตนเองซึ่งจะส่งผลให้เครื่องคอมพิวเตอร์มีความฉลาดขึ้น สามารถทำงานในระบบที่มีความซับซ้อนได้อย่างมีประสิทธิภาพโดยไม่ต้องอาศัยแรงงานจากมนุษย์

สำนักงานตำรวจแห่งชาติ (2567) หน่วยงานยุติธรรม เช่น สำนักงานตำรวจแห่งชาติ และหน่วยงานความมั่นคง เริ่มทดลองใช้ปัญญาประดิษฐ์ AI ในการตรวจสอบข้อมูลอาชญากรรมและการวิเคราะห์ข้อมูลเชิงพฤติกรรม แต่ยังไม่มีการยอมรับจริยธรรมของรัฐที่รองรับการใช้ AI อย่างเป็นระบบ

ผลกระทบต่อสิทธิมนุษยชนจากการใช้ปัญญาประดิษฐ์ AI ความเสี่ยงของปัญญาประดิษฐ์ AI ต่อสิทธิมนุษยชน สิทธิในความเป็นส่วนตัว สิทธิในการไม่ถูกเลือกปฏิบัติ สิทธิในกระบวนการยุติธรรมที่เป็นธรรม สิทธิในเสรีภาพการแสดงออก

Angwin et al. (2016) ระบบ COMPAS ความเสี่ยงการกระทำผิดซ้ำในสหรัฐ มือกติสูงต่อคนผิวดำ Buolamwini & Gebru (2018) พบความคลาดเคลื่อนรุนแรงของระบบจดจำใบหน้าในการจำแนกเพศของผู้หญิงผิวสี สถาบันเพื่อการยุติธรรมแห่งประเทศไทย (2564) ประเมินความเสี่ยงของการใช้ปัญญาประดิษฐ์ AI ในกระบวนการยุติธรรมไทย ไทยยังไม่มี หลักสิทธิมนุษยชนดิจิทัล (Digital Human Rights) เป็นทางการ ขณะที่คณะกรรมการสิทธิมนุษยชนแห่งชาติเตือนว่าระบบเฝ้าระวังอัจฉริยะของรัฐอาจละเมิดสิทธิความเป็นส่วนตัวของประชาชนได้หากขาดกฎหมายควบคุม

แนวคิดในด้านสิทธิมนุษยชนในยุคเทคโนโลยีดิจิทัล คือ การคุ้มครองสิทธิขั้นพื้นฐานของมนุษย์ให้เท่าทันการเปลี่ยนแปลงของเทคโนโลยี โดยเน้นให้เทคโนโลยีดิจิทัลและเทคโนโลยีปัญญาประดิษฐ์ AI เป็นเครื่องมือที่ “ส่งเสริมสิทธิ” มากกว่าจะเป็นเครื่องมือที่ “ละเมิดสิทธิ” ของประชาชน ทั้งในมิติของความเป็นส่วนตัว ความยุติธรรม เสรีภาพ และศักดิ์ศรีความเป็นมนุษย์.

สหประชาชาติ (United Nations Human Rights Council. 2021) การเทคโนโลยีดิจิทัลต้องไม่ละเมิดสิทธิขั้นพื้นฐาน คือ สิทธิในเรื่องความเป็นส่วนตัว (Right to Privacy), สิทธิในความเสมอภาค (Equality before the Law), และสิทธิในกระบวนการยุติธรรม (Right to Due Process) การนำปัญญาประดิษฐ์ AI มาใช้ในกระบวนการบังคับใช้กฎหมายจึงต้องยึดหลักความจำเป็นและความได้สัดส่วน (Necessity & Proportionality Principle) Toronto Declaration (2018) การพัฒนาเทคโนโลยีปัญญาประดิษฐ์ AI ต้องไม่เลือกปฏิบัติและต้องสามารถตรวจสอบย้อนกลับได้เพื่อปกป้องศักดิ์ศรีความเป็นมนุษย์ (Human Dignity) แนวคิดนี้จึงเป็นรากฐานสำคัญในการกำหนดกรอบการใช้เทคโนโลยีปัญญาประดิษฐ์ AI อย่างมีจริยธรรมและเคารพสิทธิมนุษยชนในสังคมประชาธิปไตย

งานวิจัยที่เกี่ยวข้อง

สำนักงานพัฒนารัฐบาลดิจิทัล (องค์การมหาชน). (2566). ส่งเสริมการใช้ AI ในภาครัฐ ยังไม่มีกรอบสิทธิมนุษยชน ซึ่งได้นำเสนอนโยบายรัฐบาลดิจิทัล พ.ศ. 2565–2570 โดยเน้นให้หน่วยงานภาครัฐเร่งประยุกต์ใช้เทคโนโลยีปัญญาประดิษฐ์ (AI) เพื่อเพิ่มประสิทธิภาพการให้บริการรัฐ ลดความซ้ำซ้อน และยกระดับความโปร่งใสในการทำงานราชการ

กิตติคุณ ตันรักษ์ (2564) เสนอให้มีกฎหมายควบคุมความรับผิดจากปัญญาประดิษฐ์ AI มาใช้ตัดสินใจในบริบทของทางกฎหมาย หรือการตรวจสอบความผิดแทนเจ้าหน้าที่ที่เป็นมนุษย์ ซึ่งทำให้เกิดปัญหาใหม่ด้านความรับผิดทางกฎหมาย (liability) ออกกฎหมายควบคุมความรับผิดของปัญญาประดิษฐ์ AI โดยมีกำหนดขั้นตอนการตรวจสอบเรื่องความโปร่งใส เรื่องการบันทึกข้อมูลในการทำงาน และกลไกตรวจสอบในความผิดพลาด เพื่อคุ้มครองประชาชนในยุคที่ปัญญาประดิษฐ์ AI มีบทบาทเหนือกระบวนการยุติธรรมมากขึ้น

สถาบันเพื่อการยุติธรรมแห่งประเทศไทย (2564) วิเคราะห์การใช้เทคโนโลยีปัญญาประดิษฐ์ AI ในระบบ

ยุติธรรมไทย เช่น ระบบวิเคราะห์ข้อมูลอาชญากรรม มีการบันทึกพฤติกรรมผู้ต้องขัง หรือระบบประเมินความเสี่ยง การกระทำผิดซ้ำ โดยระบุประเด็นสำคัญว่าปัญญาประดิษฐ์ AI อาจจะช่วยเพิ่มประสิทธิภาพของกระบวนการยุติธรรม พร้อมกับกลไกการตรวจสอบอิสระ (independent oversight mechanism)

คณะกรรมการสิทธิมนุษยชนฯ เตือนเรื่องผลกระทบจากการเฝ้าระวังอัจฉริยะของรัฐ แม้เทคโนโลยีเหล่านี้ได้ช่วยเพิ่มประสิทธิภาพในด้านความมั่นคง แต่ถ้าหากไม่มีหลักเกณฑ์ด้านสิทธิมนุษยชนรองรับ ก็อาจจะนำไปสู่ปัญหาสำคัญได้ เช่น

1. เรื่องการละเมิดสิทธิในความเป็นส่วนตัว ในการเก็บข้อมูลประชาชนที่เกินความจำเป็น
2. เรื่องความเสี่ยงของการนำข้อมูลที่เก็บมาไปใช้อย่างผิดวัตถุประสงค์
3. เรื่องอคติของระบบจดจำใบหน้า ที่อาจทำให้เกิดการจับกุมผิดพลาด
4. เรื่องการขาดกลไกในการตรวจสอบการใช้งานของรัฐ ทำให้การใช้อำนาจรัฐไม่มีการถ่วงดุล

ระเบียบวิธีวิจัย

เป็นการวิจัยเชิงคุณภาพ (Qualitative Research) ศึกษาขอบเขตการใช้ปัญญาประดิษฐ์ในงานตำรวจ งานยุติธรรม ประเด็นด้านจริยธรรม และผลกระทบต่อสิทธิมนุษยชนในบริบทไทยและเทียบเคียงของต่างประเทศ

ผู้ให้ข้อมูลหลัก

- หน่วยงานภาครัฐที่เกี่ยวข้องกับเรื่องของการบังคับใช้กฎหมาย เช่น ตำรวจ, ปปง., กระทรวงยุติธรรม
- นักวิชาการในด้านกฎหมาย
- ผู้เชี่ยวชาญในด้าน AI และจริยธรรมเทคโนโลยี
- ผู้แทนในส่วนภาคประชาชนในด้านสิทธิมนุษยชน (Human Rights Advocates)

คัดเลือกแบบเฉพาะเจาะจง (Purposive Sampling) เพื่อให้ได้ข้อมูลจากผู้เชี่ยวชาญที่มีความรู้ตรงประเด็น จำนวน 10 ท่าน

เครื่องมือที่ใช้ในการวิจัย (Research Instruments)

1. แบบสัมภาษณ์กึ่งโครงสร้าง (Semi-structured Interview Guide)
2. แบบสังเกตภาคสนาม (Observation Record)
3. แบบบันทึกเอกสาร (Document Analysis Form)

การเก็บรวบรวมข้อมูล (Data Collection)

1. ศึกษาจากเอกสาร (Documentary Study)
2. เก็บจากการสัมภาษณ์เชิงลึก (In-depth Interview)
3. การสังเกต (Observation)

การวิเคราะห์ข้อมูล (Data Analysis) โดยใช้การวิเคราะห์เนื้อหา (Content Analysis) และใช้การวิเคราะห์เชิงประเด็น (Thematic Analysis)

ผลการวิจัย

ผลการวิจัยตามวัตถุประสงค์ข้อ 1. ศึกษาขอบเขตของการใช้ปัญญาประดิษฐ์ในการบังคับใช้กฎหมาย พบว่า เทคโนโลยีปัญญาประดิษฐ์ AI ได้ถูกนำมาใช้ในกระบวนการยุติธรรม ทั้งด้านการตรวจจับและการเฝ้าระวัง โดยผ่านระบบแบบจดจำใบหน้า มีการวิเคราะห์ภาพจากกล้องวงจรปิด มีการคาดการณ์พื้นที่เสี่ยงต่อคดีด้านอาชญากรรม (Predictive Policing) และได้มีการวิเคราะห์ข้อมูลที่มีขนาดใหญ่ (Big Data Analytics) และเพื่อช่วยในการสืบสวนสอบสวน หรือมีการระบุตัวของผู้ต้องสงสัย ซึ่งยังช่วยเพิ่มความรวดเร็วและช่วยลดข้อผิดพลาดที่อาจเกิดจากมนุษย์ ซึ่งในการใช้งานนั้นโดยส่วนใหญ่อยู่ในระดับทดลอง ซึ่งยังต้องอาศัยการกำกับหรือควบคุมจากเจ้าหน้าที่อยู่ เนื่องจากยังไม่มีกรอบกฎหมายรองรับที่มีชัดเจนเกี่ยวกับการนำเทคโนโลยีปัญญาประดิษฐ์ AI มาใช้สำหรับการบังคับใช้กฎหมาย

ผลการวิจัยตามวัตถุประสงค์ข้อ 2. ศึกษาประเด็นด้านจริยธรรมและความรับผิดชอบในการใช้เทคโนโลยีปัญญาประดิษฐ์ AI พบว่า ปัญหาสำคัญ คือ อคติของอัลกอริทึม (Algorithmic Bias) ซึ่งเกิดจากข้อมูลฝึกสอนระบบ อาจจะไม่ครอบคลุมหรือเกิดความเอนเอียง ทำให้ผลลัพธ์การพยากรณ์อาชญากรรมนั้นมีแนวโน้มการเลือกปฏิบัติต่อกลุ่มประชากรในบางกลุ่ม เช่น ชนกลุ่มน้อย หรือผู้มีรายได้น้อย และนอกจากนี้ยังพบปัญหาในด้านความโปร่งใส (Transparency) และความรับผิดชอบ (Accountability) เมื่อระบบเทคโนโลยีปัญญาประดิษฐ์ AI ได้มีการตัดสินใจที่เกิดความผิดพลาด เนื่องจากยังไม่มีแนวทางในการกำหนดผู้รับผิดชอบอย่างชัดเจน และผู้เชี่ยวชาญส่วนใหญ่ยังคงเห็นว่า การนำปัญญาประดิษฐ์ AI มาใช้ในกระบวนการยุติธรรมจะต้องอยู่ภายใต้หลัก ความเป็นธรรมและการควบคุมที่เกิดจากมนุษย์ (Human Oversight) เพื่อให้การตัดสินใจของปัญญาประดิษฐ์ AI ผ่านการตรวจสอบของเจ้าหน้าที่ผู้มีอำนาจ

ผลการวิจัยตามวัตถุประสงค์ข้อ 3. ศึกษาผลกระทบต่อสิทธิมนุษยชนในระบบยุติธรรม พบว่า ในการนำเทคโนโลยีปัญญาประดิษฐ์ AI มาใช้ที่ขาดการควบคุม อาจก่อให้เกิดผลกระทบต่อสิทธิมนุษยชนในหลายด้าน หลายมิติ โดยเฉพาะสิทธิในเรื่องความเป็นส่วนตัว (Right to Privacy) และจากการเก็บประมวลผลข้อมูลส่วนบุคคล สิทธิในความเสมอภาค (Equality before the Law) จากความเอนเอียงของระบบ และสิทธิในการได้รับความเป็นธรรมทางกระบวนการ (Due Process) ซึ่งหากประชาชนนั้นไม่สามารถตรวจสอบความเป็นเหตุและผลของตัดสินใจจากเทคโนโลยีปัญญาประดิษฐ์ AI ได้ ข้อมูลจากผู้เชี่ยวชาญจึงระบุว่า ประเทศไทยควรเร่งจัดทำ กรอบจริยธรรมด้านเทคโนโลยีปัญญาประดิษฐ์ AI เพื่อความยุติธรรม (Ethical AI for Justice Framework) สำหรับป้องกันการละเมิดสิทธิและสร้างความไว้วางใจต่อการใช้เทคโนโลยี AI ในภาครัฐ

อภิปรายผล

อภิปรายผลการวิจัยตามวัตถุประสงค์ข้อ 1 พบว่า การนำปัญญาประดิษฐ์ AI เข้ามาประยุกต์ใช้ในกระบวนการบังคับใช้กฎหมายซึ่งมีขอบเขตที่ขยายตัวกว้างขึ้น และครอบคลุมตั้งแต่เริ่มโดยการเฝ้าระวัง การตรวจจับ การวิเคราะห์ข้อมูล ไปจนถึงการพยากรณ์แนวโน้มอาชญากรรม ซึ่งสอดคล้องกับแนวคิดของ Perry et al. (2013) ที่อธิบายว่า Predictive Policing คือ การใช้ข้อมูลและอัลกอริทึมเพื่อคาดการณ์เหตุการณ์อาชญากรรมล่วงหน้า เพื่อจัดสรรทรัพยากรเจ้าหน้าที่อย่างมีประสิทธิภาพ ผลการศึกษานี้จึงยืนยันได้ว่าการใช้ปัญญาประดิษฐ์ AI จะช่วยเพิ่มประสิทธิภาพและความแม่นยำในการสืบสวน แต่ในขณะเดียวกันก็ยังคงต้องอาศัยการควบคุมจากมนุษย์ และกรอบกฎหมายที่ชัดเจน สอดคล้องกับ INTERPOL (2024) ที่ระบุว่าเทคโนโลยีปัญญาประดิษฐ์ AI สามารถเป็นเครื่องมือเสริมที่สำคัญในงานด้านตำรวจได้ เมื่อมีการออกแบบอย่างมีความรับผิดชอบและคำนึงถึงเรื่องสิทธิมนุษยชน

อภิปรายผลการวิจัยตามวัตถุประสงค์ข้อ 2 พบว่า จริยธรรมและความรับผิดชอบต่อการใช้เทคโนโลยีปัญญาประดิษฐ์ AI เป็นส่วนหนึ่งที่ได้รับ ความสนใจมากในงานวิจัยนี้ ซึ่งผลการศึกษาพบว่าอคติของอัลกอริทึม (Algorithmic Bias) และความไม่โปร่งใสของการตัดสินใจ ยังเป็นปัญหาหลักที่อาจส่งผลกระทบต่อระบบความยุติธรรม และยังสอดคล้องกับงานของ Buolamwini, J., & Gebru, T. (2018) ที่มีความเห็นว่า ระบบจดจำใบหน้าของเทคโนโลยีปัญญาประดิษฐ์ AI นั้น ได้มีอัตราความผิดพลาดสูงในกลุ่มคนผิวสีและเพศหญิง ซึ่งเป็นผลจากข้อมูลฝึกสอนที่ไม่มีความหลากหลาย และข้อค้นพบนี้ก็มีความสอดคล้องกับรายงานของ Angwin et al. (2016) ซึ่งวิจารณ์โปรแกรม COMPAS ที่ใช้ในศาลสหรัฐอเมริกาในการประเมินความเสี่ยงของผู้ต้องหา และยังเห็นว่ามีแนวโน้มให้คะแนนความเสี่ยงสูงกว่าความเป็นจริงในกลุ่มคนผิวดำ โดยทั้งหมดนี้ได้สะท้อนว่าหากไม่มีการควบคุมด้านจริยธรรมและการตรวจสอบความเป็นกลางของระบบปัญญาประดิษฐ์ AI นั้นอาจกลายเป็นเครื่องมือสร้างความเหลื่อมล้ำทางกระบวนการยุติธรรมแทนที่จะลดความเหลื่อมล้ำในกระบวนการยุติธรรม

อภิปรายผลการวิจัยตามวัตถุประสงค์ข้อ 3 ผลกระทบต่อสิทธิมนุษยชน ผลการวิจัยพบว่า การใช้เทคโนโลยีปัญญาประดิษฐ์ AI โดยขาดการกำกับและดูแลที่ชัดเจนนั้น อาจเกิดการละเมิดสิทธิความเป็นส่วนตัว และสิทธิในการได้รับความเป็นธรรม ซึ่งยังสอดคล้องกับแนวคิดของ Council of Europe, (2021) ที่กำหนดให้การใช้ปัญญาประดิษฐ์ AI ในกระบวนการยุติธรรม และต้องอยู่ภายใต้ของหลักสิทธิมนุษยชนสากล คือ ความจำเป็น ความได้สัดส่วน และความโปร่งใสของการดำเนินการ ซึ่งเช่นเดียวกับแนวทางของ Toronto Declaration, (2018) ที่ระบุว่า การพัฒนาและใช้เทคโนโลยีปัญญาประดิษฐ์ AI ต้องไม่ก่อให้เกิดการเลือกปฏิบัติ และยังต้องมีระบบการตรวจสอบที่สามารถตรวจสอบย้อนกลับได้ (Traceability) เพื่อคุ้มครองศักดิ์ศรีของความเป็นมนุษย์ ซึ่งข้อค้นพบนี้ยังมีความสอดคล้องกับงานของ Floridi et al., (2018) ที่เสนอหลัก AI for Good Society โดยได้เน้นย้ำว่าการพัฒนาเทคโนโลยีปัญญาประดิษฐ์ AI ที่ดีควรอยู่บนพื้นฐานของคุณค่าทางด้านศีลธรรมและผลประโยชน์ของมนุษย์เป็นสำคัญ

ซึ่งผลการวิจัยนี้ยังได้มีความสอดคล้องกับแนวทางของ European Commission, (2022) ที่เสนอแนวคิด Trustworthy AI โดยได้เน้นให้ทุกระบบของปัญญาประดิษฐ์ AI ต้องมีคุณสมบัติ 3 ประการ ด้วยกัน คือ 1. มีความชอบธรรมตามกฎหมาย (Lawfulness) 2. มีการเคารพหลักจริยธรรม (Ethical Compliance) และ 3. มีความมั่นคงปลอดภัยทางเทคนิค (Technical Robustness) ซึ่งแนวคิดนี้ยังสอดคล้องกับข้อเสนอของผู้เชี่ยวชาญที่ให้ความสำคัญกับการควบคุมโดยมนุษย์ (Human Oversight) เพื่อให้การตัดสินใจขั้นสุดท้ายยังคงอยู่ในอำนาจของเจ้าหน้าที่รัฐ มิใช่ปล่อยให้ระบบอัตโนมัติเป็นผู้ชี้ขาดเพียงลำพัง

การอภิปรายผล ในเรื่องปัญญาประดิษฐ์ AI จะมีศักยภาพในการเพิ่มประสิทธิภาพของกระบวนการบังคับใช้กฎหมาย แต่ด้วยการนำมาใช้โดยขาดหลักของจริยธรรมและขาดกลไกตรวจสอบที่เข้มแข็ง ย่อมมีความเสี่ยงต่อการละเมิดสิทธิมนุษยชนและเกิดการบั่นทอนความเชื่อมั่นของสาธารณชน ซึ่งสอดคล้องกับ Bryson (2019) ที่เสนอว่า ความสำเร็จของปัญญาประดิษฐ์ AI ในสังคมขึ้นอยู่กับความไว้วางใจทางสังคมมากกว่าความก้าวหน้าทางเทคนิค ดังนั้น ประเทศไทยควรมีการกำหนดกรอบจริยธรรมเทคโนโลยีปัญญาประดิษฐ์ AI แห่งชาติ สำหรับภาคยุติธรรม (National AI Ethics Framework for Justice Sector) เพื่อใช้เป็นแนวทางในการพัฒนาและใช้งานเทคโนโลยีที่มีความรับผิดชอบ เกิดความโปร่งใส และเคารพสิทธิมนุษยชนอย่างแท้จริง

ข้อเสนอแนะ

การใช้เทคโนโลยีปัญญาประดิษฐ์ (AI) ในกระบวนการบังคับใช้กฎหมาย แม้ว่าจะช่วยเพิ่มประสิทธิภาพและความรวดเร็วในการทำงานของเจ้าหน้าที่นั้น แต่หากขาดกรอบในการกำกับดูแลที่ชัดเจน อาจนำไปสู่การละเมิดสิทธิและความไม่เป็นธรรมในสังคม ดังนั้น ข้อเสนอแนะที่ได้จากการศึกษาค้นคว้าครั้งนี้ควรประกอบไปด้วย

1. ข้อเสนอเชิงนโยบาย (Policy Recommendations)

1.1 รัฐ ควรจัดทำกรอบจริยธรรมแห่งชาติไว้สำหรับการใช้ปัญญาประดิษฐ์ AI ในกระบวนการยุติธรรม (National AI Ethics Framework for Justice) เพื่อจะได้กำหนดมาตรฐานด้านความโปร่งใส ความรับผิดชอบ และสิทธิมนุษยชนในการใช้ปัญญาประดิษฐ์ AI ของหน่วยงานภาครัฐ

1.2 ควรจะมีกฎหมายหรือระเบียบที่สำหรับการควบคุม การเก็บ การใช้ และการเผยแพร่ข้อมูลส่วนบุคคลในระบบปัญญาประดิษฐ์ AI เพื่อให้สอดคล้องกับพระราชบัญญัติคุ้มครองข้อมูลส่วนบุคคล พ.ศ. 2562 (PDPA) และหลักของสิทธิมนุษยชนสากล

1.3 ควรจัดตั้งคณะกรรมการอิสระด้านการกำกับเทคโนโลยีปัญญาประดิษฐ์ AI เพื่อความยุติธรรม มีหน้าที่ตรวจสอบความเป็นธรรมของอัลกอริทึม (Algorithmic Fairness) และพิจารณาผลกระทบต่อสิทธิมนุษยชน

2. ข้อเสนอเชิงปฏิบัติ (Practical Recommendations)

2.1 หน่วยงานบังคับใช้กฎหมายควรพัฒนาระบบตรวจสอบอคติของเทคโนโลยีปัญญาประดิษฐ์ AI (Bias Audit System) เพื่อประเมินความเป็นกลางและลดความผิดพลาดของข้อมูล

2.2 ควรมีการจัดอบรมเจ้าหน้าที่ผู้ใช้ระบบเทคโนโลยีปัญญาประดิษฐ์ AI ด้านจริยธรรมดิจิทัลและด้านสิทธิมนุษยชน เพื่อให้เข้าใจขีดจำกัดของเทคโนโลยี และไม่ใช้ผลลัพธ์ของเทคโนโลยีปัญญาประดิษฐ์ AI โดยขาดดุลยพินิจของมนุษย์

2.3 ควรส่งเสริมให้หน่วยงานภาครัฐและภาคเอกชนร่วมมือกับสถาบันการศึกษาในการสร้าง ต้นแบบการใช้ปัญญาประดิษฐ์ AI เพื่อความยุติธรรมอย่างรับผิดชอบ (Responsible AI Sandbox) เพื่อทดสอบระบบในพื้นที่จำกัดก่อนนำไปใช้จริง

3. การนำผลการวิจัยไปใช้ประโยชน์ (Research Utilization)

3.1 การกำหนดนโยบายด้านเทคโนโลยีปัญญาประดิษฐ์ AI ของกระทรวงยุติธรรมและสำนักงานตำรวจแห่งชาติ

3.2 การพัฒนาแนวทางการอบรมเจ้าหน้าที่ภาครัฐและผู้บังคับใช้กฎหมายด้านเทคโนโลยีและสิทธิมนุษยชน

3.3 การสร้างหลักสูตรการศึกษาด้าน “AI และจริยธรรมทางกฎหมาย” ในระดับมหาวิทยาลัย เพื่อพัฒนาองค์ความรู้ใหม่ด้านนิติศาสตร์และเทคโนโลยี

เอกสารอ้างอิง

กิตติคุณ ต้นรักษ์. (2564). *กฎหมายกับปัญญาประดิษฐ์: กรอบคิดพื้นฐานและประเด็นทางกฎหมาย*. กรุงเทพฯ: สำนักพิมพ์วิญญูชน.

สถาบันเพื่อการยุติธรรมแห่งประเทศไทย (TIJ). (2564). *AI กับกระบวนการยุติธรรมไทย: ความเสี่ยง โอกาส และแนวทางเชิงนโยบาย*. กรุงเทพฯ: TIJ.

สำนักงานตำรวจแห่งชาติ. (2567). *รายงานสถานการณ์อาชญากรรมทางเทคโนโลยีในประเทศไทย*. กรุงเทพฯ: ศูนย์ปราบปรามอาชญากรรมทางเทคโนโลยีสารสนเทศ (ศปอส.ตร.).

สำนักงานพัฒนารัฐบาลดิจิทัล (องค์การมหาชน). (2566). *นโยบายและแนวทางรัฐบาลดิจิทัล พ.ศ. 2565–2570*. กรุงเทพฯ: สพร. (DGA).

Angwin, J., Larson, J., Mattu, S., & Kirchner, L. (2016). Machine bias: There's software used across the country to predict future criminals—and it's biased against blacks. ProPublica. Retrieved from <https://www.propublica.org/article/machine-bias-risk-assessments-in-criminal-sentencing>

- Bryson, J. J. (2019). The past decade and future of AI's impact on society. In Towards a New Enlightenment? BBVA OpenMind.
- Buolamwini, J., & Gebru, T. (2018). Gender shades: Intersectional accuracy disparities in commercial gender classification. *Proceedings of Machine Learning Research*, 81, 1–15.
- Council of Europe. (2021). *Guidelines on Artificial Intelligence and Human Rights*. Strasbourg: Council of Europe Publications.
- European Commission. (2022). *Ethics guidelines for trustworthy AI*. Brussels: European Union.
- Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., ... & Vayena, E. (2018). AI4People-An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28(4), 689–707. <https://doi.org/10.1007/s11023-018-9482-5>
- Grimson E. (2025) สรุปรูป AI อดีต ปัจจุบัน อนาคต สถาบันเทคโนโลยีแมสซาชูเซตส์ (MIT) สืบค้นจาก <https://techsauce.co/tech-and-biz/ai-past-present-future-eric-grimson-mit-keynote>
- Perry, W. L., McInnis, B., Price, C. C., Smith, S. C., & Hollywood, J. S. (2013). *Predictive policing: The role of crime forecasting in law enforcement operations*. RAND Corporation.
- Russell, S., & Norvig, P. (2021). *Artificial intelligence: A modern approach* (4th ed.). Harlow, England: Pearson Education.
- Toronto Declaration. (2018). *Protecting the right to equality and non-discrimination in machine learning systems*. Amnesty International and Access Now. Retrieved from <https://www.accessnow.org/the-toronto-declaration>
- United Nations Human Rights Council. (2021). *The right to privacy in the digital age: Report of the United Nations High Commissioner for Human Rights*. Geneva: United Nations Publications.