

สรุปประเด็นการฝึกอบรมที่มีการสอดแทรกสาระด้านจริยธรรม
ของผู้ปฏิบัติในผลงาน

เรื่อง "Applying Artificial Intelligence (AI)
for Researchers"

พศ. 2568

สถาบันวิจัยวิทยาศาสตร์สุขภาพ

สรุปประเด็นการฝึกอบรมที่มีการสอดแทรกสาระด้านจริยธรรมของผู้ปฏิบัติในสัมนา เรื่อง "Applying Artificial Intelligence (AI) for Researchers"

สถาบันวิจัยวิทยาศาสตร์สุขภาพ มหาวิทยาลัยเชียงใหม่ได้จัดอบรมแก่คณาจารย์และนักวิจัย เรื่อง "Applying Artificial Intelligence (AI) for Researchers" โดยวิทยากร Dr. José Guilherme Behrendorf Derraik แบบ Online training

Topic 1: Introduction to Artificial Intelligence: An Overview for Researchers

วันจันทร์ที่ 24 มีนาคม 2568 เวลา 9.00-10.30 น.

Zoom Meeting ID: 925 1774 9721

Invite Link: <https://cmu-th.zoom.us/j/92517749721>

ผู้เข้าร่วม จำนวน 33 คน

Topic 2: Practical Applications of AI in Research: Exploring Tools and Their Usage

วันจันทร์ที่ 31 มีนาคม 2568 เวลา 10.30-12.00 น.

Zoom Meeting ID: 913 1052 3743

Invite Link: <https://cmu-th.zoom.us/j/91310523743>

ผู้เข้าร่วม จำนวน 25 คน

ในการบรรยาย Topic 2 ผู้บรรยายได้อธิบายถึงการใช้ AI และได้ยกตัวอย่างจริยธรรมในบริบทของการใช้ AI ในประเด็นต่างๆ สอดแทรกในระหว่างการบรรยายซึ่งประกอบด้วย

1 ประเด็นจริยธรรมในการใช้ AI ในงานวิจัยและการตีพิมพ์ทางวิชาการ

ในช่วงสองปีที่ผ่านมาปัญหาที่พบบ่อยคือปัญหาการใช้ AI อย่างไม่ถูกต้องในงานวิจัย มีงานวิจัยจำนวนมากที่ถูกตีพิมพ์โดยไม่มีการตรวจสอบเนื้อหาอย่างละเอียด ซึ่งบางส่วนเกิดจากการใช้ AI ในทางที่ผิด เช่น การสร้างบทความปลอมหรือใช้ AI เขียนทั้งฉบับโดยไม่มีการกลั่นกรองทางวิชาการที่เหมาะสม ปัญหาความล้มเหลวของกระบวนการตรวจสอบคุณภาพ (Peer Review Process) ผู้บรรยายได้ชี้ให้เห็นว่า บทความที่มีข้อผิดพลาดร้ายแรงได้ผ่านกระบวนการตีพิมพ์โดยไม่มีการตรวจพบและได้มีการแจ้งถอดถอนบทความในภายหลัง ได้สะท้อนถึงปัญหาในระบบ และความเสี่ยงของการใช้ AI ในการสร้างเนื้อหาที่ขาดความถูกต้อง ความจำเป็นในการใช้ AI อย่างมีจริยธรรม ผู้บรรยายได้สร้างความตระหนักถึง การไม่ควรใช้ AI ในทางที่ผิด เช่น การสร้างงานวิจัยปลอม เพราะอาจส่งผลกระทบต่อความน่าเชื่อถือของวงการวิชาการและสำนักพิมพ์

2 จริยธรรมในเชิงปฏิบัติ:

ผู้บรรยายได้เสนอแนะให้มีการใช้ AI โดยเสริมศักยภาพในการทำงานแต่ไม่ควรให้ AI ทำงานแทนผู้ปฏิบัติงาน เช่น เสนอใช้ AI อย่างถูกต้องในการช่วยปรับปรุงข้อความในต้นฉบับแทนที่จะให้ AI เขียนทั้งหมด

2.1 ผู้บรรยายได้แนะนำแนวทางการเขียนตีพิมพ์โดยอาศัย AI เป็นเครื่องมือในการตรวจคำผิด การให้ข้อเสนอแนะเกี่ยวกับเนื้อหาทางวิทยาศาสตร์ การตรวจแก้/ปรับปรุง ร่าง manuscript ต้นฉบับและการตรวจทานก่อนส่งตีพิมพ์ โดยแนะนำ AI ที่ทำให้กระบวนการนี้มีความโปร่งใสเช่น ChatGPT และ Gemini (รายละเอียดใน Acknowledgment)

2.2 ผู้บรรยายได้แนะนำ AI ที่มีความปลอดภัยและสามารถรักษาความลับได้ดี เช่นในการทำงานด้าน health information ควรใช้ Perplexity Enterprise ในงาน medical research ที่เกี่ยวข้องกับข้อมูลของผู้ป่วยสามารถรักษาความลับผู้ป่วยด้วย AI ตัวอื่นเช่น GPT Enterprise และ Claude Enterprise เป็นต้น

2.3 ประเด็นจริยธรรมอื่นๆ

2.3.1 ประเด็น AI กำลังเปลี่ยนแปลงเศรษฐกิจและตลาดแรงงานอย่างรวดเร็ว ปัญหาที่เกิดขึ้นคืออาจนำไปสู่การว่างงานสูงถึง 30% ซึ่งเป็นความเปลี่ยนแปลงในระดับเดียวกับการปฏิวัติอุตสาหกรรม แต่เกิดขึ้นในเวลาอย่างรวดเร็วมากใช้เวลาเพียงไม่กี่ปีเท่านั้น ผู้บรรยายได้ชี้ให้เห็นถึงความจำเป็นในการใช้ AI อย่างมีจริยธรรม เนื่องจากอนาคตของตลาดแรงงานไม่แน่นอน การเรียนรู้วิธีใช้ AI อย่างมีจริยธรรมจะช่วยให้ผู้คนสามารถเก็บงานของตนไว้ได้นานที่สุด นอกจากนี้ได้อธิบายถึงความเร่งด่วนในการเรียนรู้ AI และการทำความเข้าใจ AI รวมถึงการใช้งานอย่างมีความรับผิดชอบเพื่อให้เกิดประโยชน์สูงสุดต่อตัวบุคคลและสังคม

2.3.2 ประเด็นข้อจำกัดของ AI ด้านจริยธรรมและการเซ็นเซอร์ข้อมูล

ผู้บรรยายได้ยกตัวอย่างถึงโมเดล AI บางตัว (เช่น โมเดลจาก Google) ที่มักปฏิเสธการตอบคำถามเกี่ยวกับบางหัวข้อโดยให้เหตุผลว่า "ไม่เหมาะสม" หรือ "ขัดกับจริยธรรม" ซึ่งบางครั้งอาจเป็นอุปสรรคต่อการทำงานวิจัย โดยเฉพาะเมื่อเกี่ยวข้องกับหัวข้อที่มีความละเอียดอ่อน เช่น ความแตกต่างทางเพศในการตอบสนองต่อการรักษาทางการแพทย์ โมเดล AI ทางเลือกที่มีการเซ็นเซอร์น้อยกว่า ได้อ้างถึง AI อย่าง Grok ที่มีแนวโน้มจะ "ไม่ถูกเซ็นเซอร์" และสามารถตอบคำถามที่ AI บางตัวหลีกเลี่ยง ซึ่งสะท้อนถึงการถกเถียงเกี่ยวกับ ขอบเขตของจริยธรรมในการควบคุม AI ระหว่างการปกป้องผู้ใช้จากเนื้อหาที่อาจเป็นอันตราย กับการเปิดโอกาสให้ AI สนับสนุนการวิจัยที่ต้องการข้อมูลเฉพาะทาง การใช้ AI ในการสร้างสไลด์และเอกสารนำเสนอ แม้จะไม่เกี่ยวข้องโดยตรงกับจริยธรรม แต่ได้เน้นว่า AI สามารถช่วยจัดทำเนื้อหาหรือออกแบบสไลด์ได้ ซึ่งเป็นตัวอย่างของการใช้ AI อย่างมีประสิทธิภาพ แต่ก็ต้องพิจารณาถึงข้อจำกัดด้วย

3 สรุปอภิปรายผลและข้อเสนอแนะ:

ในการจัดอบรมครั้งนี้ บุคลากรที่เข้าร่วมการอบรมได้เกิดความตระหนักในการใช้ AI อย่างมีจริยธรรม การใช้ AI ในงานวิจัยและการตีพิมพ์ทางวิชาการจำเป็นต้องยึดหลักจริยธรรมอย่างเคร่งครัด แม้ว่าในช่วงสองปีที่ผ่านมาปัญหาการใช้ AI อย่างไม่เหมาะสมในงานวิจัยและการตีพิมพ์ทางวิชาการได้กลายเป็นประเด็นสำคัญ ทั้งในแง่ของการสร้างบทความปลอม ความล้มเหลวของกระบวนการตรวจสอบ และการขาดความโปร่งใสในการใช้ AI ผู้บรรยายได้เสนอแนวทางการใช้ AI อย่างมีจริยธรรม โดยเน้นให้ AI เป็นเครื่องมือเสริมไม่ใช่ตัวแทนมนุษย์ พร้อมแนะนำเครื่องมือที่ปลอดภัยสำหรับงานด้านสุขภาพ เช่น GPT Enterprise ฯลฯ และควรมีการเปิดเผยการใช้ AI ด้วยนอกจากนี้ยังได้ชี้ให้เห็นข้อจำกัดของ AI ต่อการเซ็นเซอร์ข้อมูลในโมเดล AI บางตัว ซึ่งอาจเป็นอุปสรรคต่อการวิจัย จึงเน้นย้ำถึงความจำเป็นของบุคลากรภายในองค์กรในการเรียนรู้และใช้งาน AI อย่างมีความรับผิดชอบและโปร่งใสในทุกบริบท

ตามข้อบังคับมหาวิทยาลัยเชียงใหม่ว่าด้วยประมวลจริยธรรมผู้บริหารและปฏิบัติงานในมหาวิทยาลัยเชียงใหม่ พ.ศ. 2566 หมวดที่3 และ 4 ได้ระบุถึงจริยธรรมของผู้ปฏิบัติงานสายวิชาการ และจริยธรรม

ของผู้ปฏิบัติงานในมหาวิทยาลัย เพื่อให้สอดคล้องกับข้อบังคับดังกล่าว บุคลากรทั้งคณาจารย์และนักวิจัยจึงควรดำเนินการเปิดเผยการใช้ AI ด้วย ซึ่งในปัจจุบันได้มีแนวทางการเปิดเผยการใช้ AI ตามแนวทางของสำนักพิมพ์ต่างๆ เช่น Elsevier, JAMA และ NEJM ฯลฯ เป็นต้น ตัวอย่างการอ้างอิงมาตรฐานการใช้ AI ที่มีการระบุใน Elsevier publisher เพื่อความโปร่งใสมีดังนี้

1 The use of AI and AI-assisted writing technologies in scientific writing

Elsevier ไม่อนุญาตให้ AI หรือเครื่องมือช่วยเขียนด้วย AI เป็นผู้เขียนงานตีพิมพ์ เนื่องจากการเป็นผู้เขียนต้องมีความรับผิดชอบที่ทำได้โดยมนุษย์ เช่น การรับรองความถูกต้อง ความเป็นต้นฉบับ และการอนุมัติฉบับสุดท้าย

2 การ acknowledge AI: ผู้เขียนต้องเปิดเผยการใช้ AI ช่วยเขียนไว้ท้ายต้นฉบับ ก่อนเอกสารอ้างอิง โดยตั้งชื่อหัวข้อว่า Declaration of generative AI and AI-assisted technologies in the writing process

ให้ระบุชื่อเครื่องมือ เหตุผลในการใช้ และยืนยันว่าผู้เขียนได้ตรวจสอบและรับผิดชอบต่อเนื้อหาทั้งหมดดังนี้

“During the preparation of this work the author(s) used [NAME Tool/SERVICE] in order to [REASON]. After using this tool/service, the author(s) reviewed and edited the content as needed and take(s) full responsibility for the content of the publication”

3 นโยบายสำหรับauthors และเนื้อหาที่ได้รับมอบหมาย: นโยบายนี้มีเป้าหมายเพื่อสร้างความโปร่งใสและให้คำแนะนำแก่ผู้เขียน ผู้อ่าน ผู้ประเมินต้นฉบับ และบรรณาธิการ เกี่ยวกับการใช้ AI โดย นโยบายนี้เน้นความโปร่งใสในการใช้ AI ช่วยเขียน โดยครอบคลุมเฉพาะขั้นตอนการเขียน โดยไม่รวมการใช้ AI วิเคราะห์ข้อมูลในการวิจัย การปฏิบัติตามแนวทางข้างต้นจะช่วยส่งเสริมความน่าเชื่อถือ โปร่งใส และจริยธรรมในการใช้เทคโนโลยีปัญญาประดิษฐ์ ภายในองค์กรฯ

Acknowledgments

- The artificial intelligence models ChatGPT o3-mini-high (OpenAI, San Francisco, CA, USA) and Gemini 2.0 Flash Thinking Experimental (Google LLC, Mountain View, CA, USA) assistance with brainstorming improvements to text structure and clarity, and feedback on scientific content.
- JGB Derraiik used ChatGPT4o (Anthropic, San Francisco, CA, USA) and Gemini 2.0 Flash Thinking Experimental (Google LLC, Mountain View, CA, USA) for critical editorial review and revision of the draft manuscript, and Grammarly for Microsoft Office v6.8 (Grammarly Inc., San Francisco, CA, USA) for final proofreading.

References

- 1 <https://www.elsevier.com/about/policies-and-standards/the-use-of-generative-ai-and-ai-assisted-technologies-in-writing-for-elsevier>
 - The use of AI and AI-assisted writing technologies in scientific writing, Frequently asked questions: Why has Elsevier decided that AI and AI-assisted tools cannot be credited as an author on published work?
 - Elsevier endorsed the addition of a new section entitled “Declaration of Generative AI and AI-assisted technologies in the writing process”
 - Policy for book and commissioned content authors
- 2 A. Flanagan, J. Kendall-Taylor and K. Bibbins-Domingo. Guidance for Authors, Peer Reviewers, and Editors on Use of AI, Language Models, and Chatbots, JAMA 2023 Vol. 330 Issue 8 Pages 702-703, DOI: 10.1001/jama.2023.12500
- 3 A. Flanagan, R. Pirracchio, R. Khera, M. Berkwits, Y. Hsuen and K. Bibbins-Domingo, Reporting Use of AI in Research and Scholarly Publication—JAMA Network Guidance, JAMA 2024 Vol. 331 Issue 13 Pages 1096-1098, DOI: 10.1001/jama.2024.3471.
- 4 Koller D., Beam A., Manrai A., Ashley E., Liu X., Gichoya J., Holmes C., Zou J., Dagan N., Wong TY., Blumenthal D., and Kohane I. Why We Support and Encourage the Use of Large Language Models in NEJM AI Submissions. NEJM AI 2024;1(1),DOI:10.1056/Aie2300128, VOL. 1 NO. 1
- 5 ข้อบังคับมหาวิทยาลัยเชียงใหม่ว่าด้วยประมวลจริยธรรมผู้บริหารและปฏิบัติงานในมหาวิทยาลัยเชียงใหม่ พ.ศ. 2566

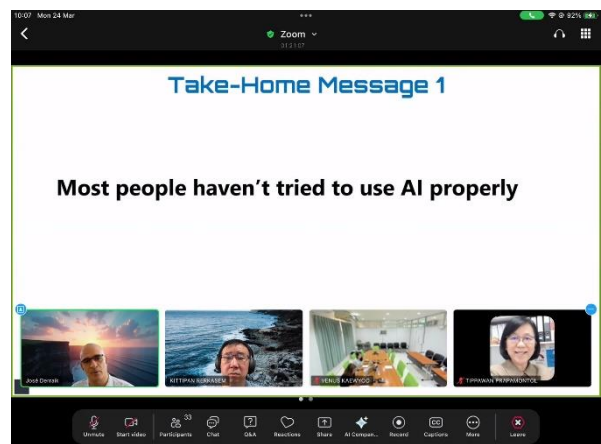
ภาพการจัดอบรมออนไลน์ เรื่อง “Applying Artificial Intelligence (AI) for Researchers”

Lecturer: Dr José Guilherme Behrendorf Derraik

วันจันทร์ที่ 24 มีนาคม 2568 เวลา 9.00-10.30 น. (15:00-16.30 in New Zealand)

Topic 1: Introduction to Artificial Intelligence: An Overview for Researchers

.....



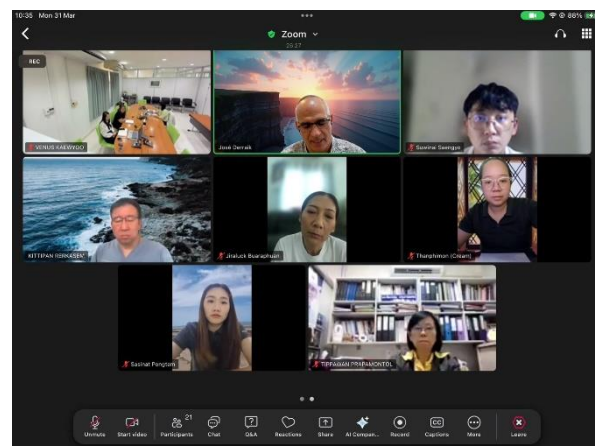
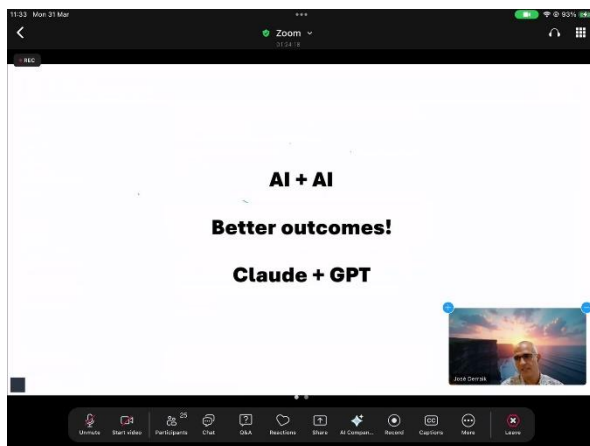
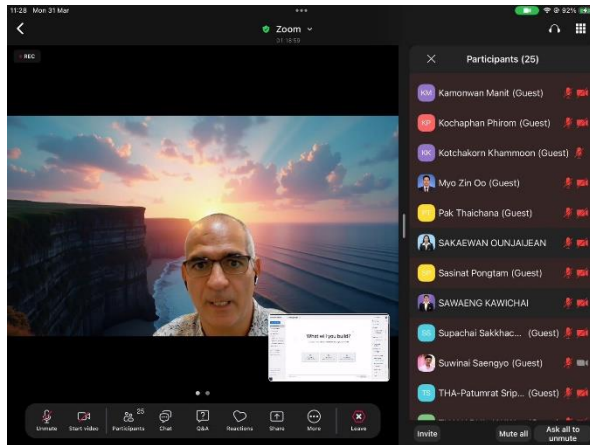
ภาพการจัดอบรมออนไลน์ เรื่อง “Applying Artificial Intelligence (AI) for Researchers”

Lecturer: Dr José Guilherme Behrendorf Derraik

วันจันทร์ที่ 31 มีนาคม 2568 เวลา 10.30-12.00 น. (16:30-18.00 in New Zealand)

Topic 2: Practical Applications of AI in Research: Exploring Tools and Their Usage

.....



Fraudulent Use of AI

- **Key (and frightening) indication of appalling editorial process**

- Authors prior to MS submission
- Editorial staff MS screening
- Possible Editor MS screening prior to peer-review
- Reviewer(s)
- Possibly Editor/Editorial staff again
- Authors during revision
- Editorial staff MS screening
- Possibly Editor/Editorial staff again
- Possibly Reviewer(s) again
- Production staff after acceptance
- Authors during proofreading



AI vs Ethics vs Publishers



Online Training

Applying AI for Researchers



**Topic 1: Introduction to
Artificial Intelligence:
An Overview for Researchers**

24 March
9:00-10:30 AM
(Thailand Time)
Zoom ID: 925 1774 9721





Dr. José Guilherme
Behrendorf Derraik



Open for registration
until 19 March 2025
Limited to only 30 participants

**Topic 2: Practical Applications
of AI in Research: Exploring
Tools and Their Usage**

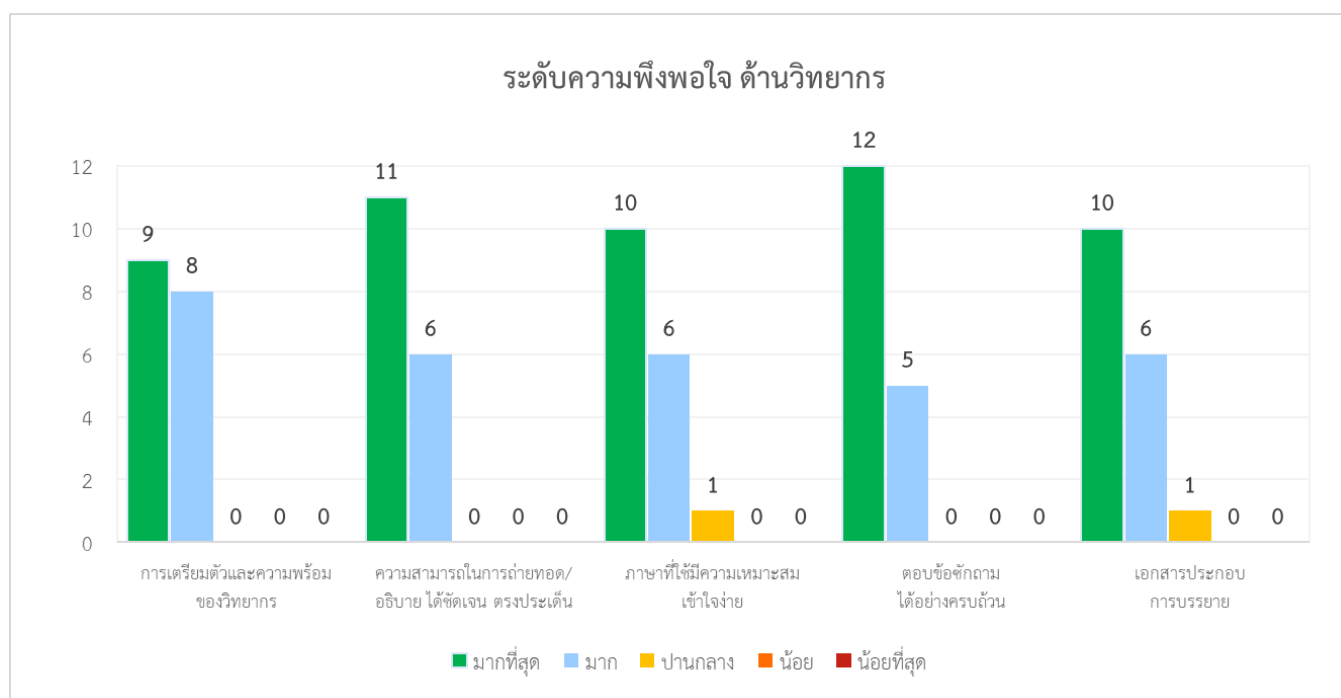
31 March
10:30-12:00 AM
(Thailand Time)
Zoom ID: 913 1052 3743

สรุปผลการประเมินการจัดอบรม

สถาบันวิจัยวิทยาศาสตร์สุขภาพได้จัดทำแบบประเมินกิจกรรมการอบรม (AI course training activity evaluation form) แบบออนไลน์ผ่าน google form โดยสำรวจความเห็นจากผู้เข้าร่วมการอบรม พบว่ามีผู้ตอบแบบประเมิน จำนวน 17 คน สรุปผลได้ดังนี้

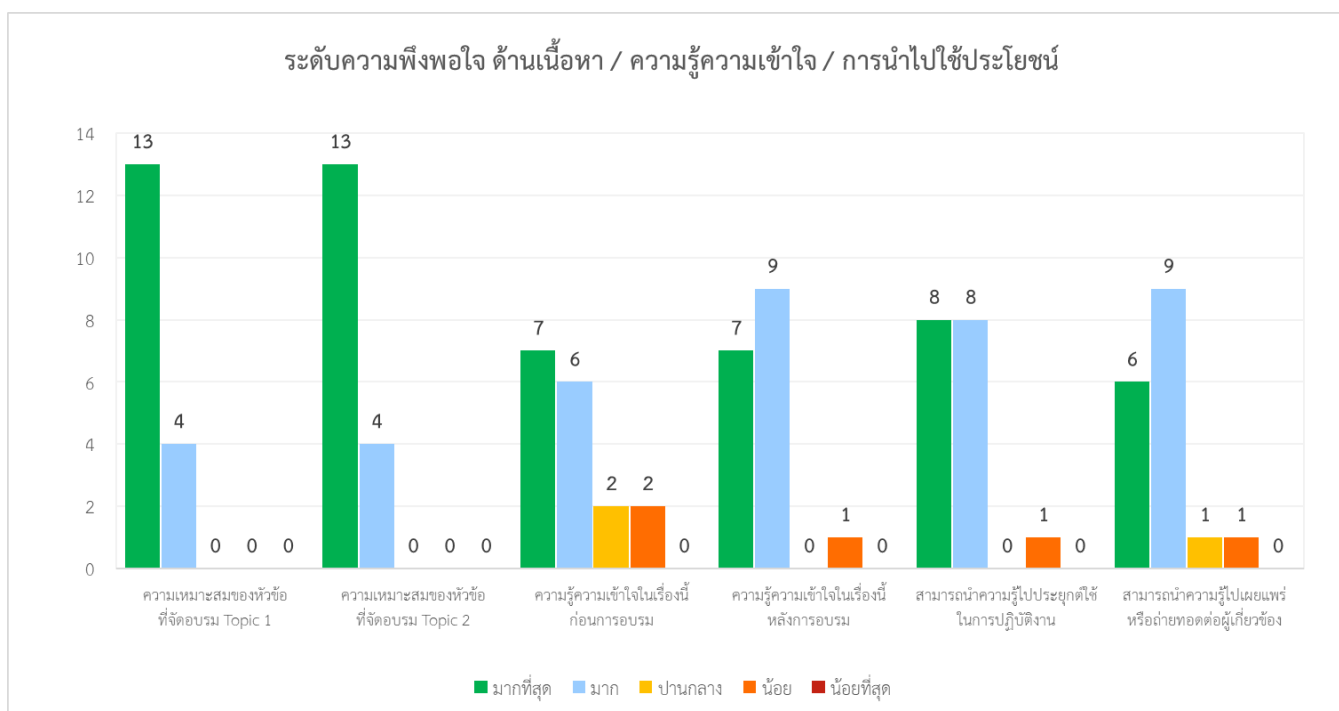
1. ความพึงพอใจ ด้านวิทยากร

การประเมินความพึงพอใจ แบ่งเป็น 5 ระดับ ได้แก่ พอใจมากที่สุด พอใจมาก พอใจปานกลาง พอใจน้อย และพอใจน้อยที่สุด โดยความพึงพอใจ ด้านวิทยากร แบ่งเป็น 5 หัวข้อ ดังนี้ 1) การเตรียมตัวและความพร้อมของวิทยากร ผลประเมินอยู่ในระดับพอใจมากที่สุด จำนวน 9 คน, พอใจมาก จำนวน 8 คน 2) ความสามารถในการถ่ายทอด/อธิบาย ได้ชัดเจน ตรงประเด็น ผลประเมินอยู่ในระดับพอใจมากที่สุด จำนวน 11 คน, พอใจมาก จำนวน 6 คน 3) ภาษาที่ใช้มีความเหมาะสม เข้าใจง่าย ผลประเมินอยู่ในระดับพอใจมากที่สุด จำนวน 10 คน, พอใจมาก จำนวน 6 คน, พอใจปานกลาง จำนวน 1 คน 4) ตอบข้อซักถามได้อย่างครบถ้วน ผลประเมินอยู่ในระดับพอใจมากที่สุด จำนวน 12 คน, พอใจมาก จำนวน 5 คน และ 5) เอกสารประกอบการบรรยาย ผลประเมินอยู่ในระดับพอใจมากที่สุด จำนวน 10 คน, พอใจมาก จำนวน 6 คน, พอใจปานกลาง จำนวน 1 คน



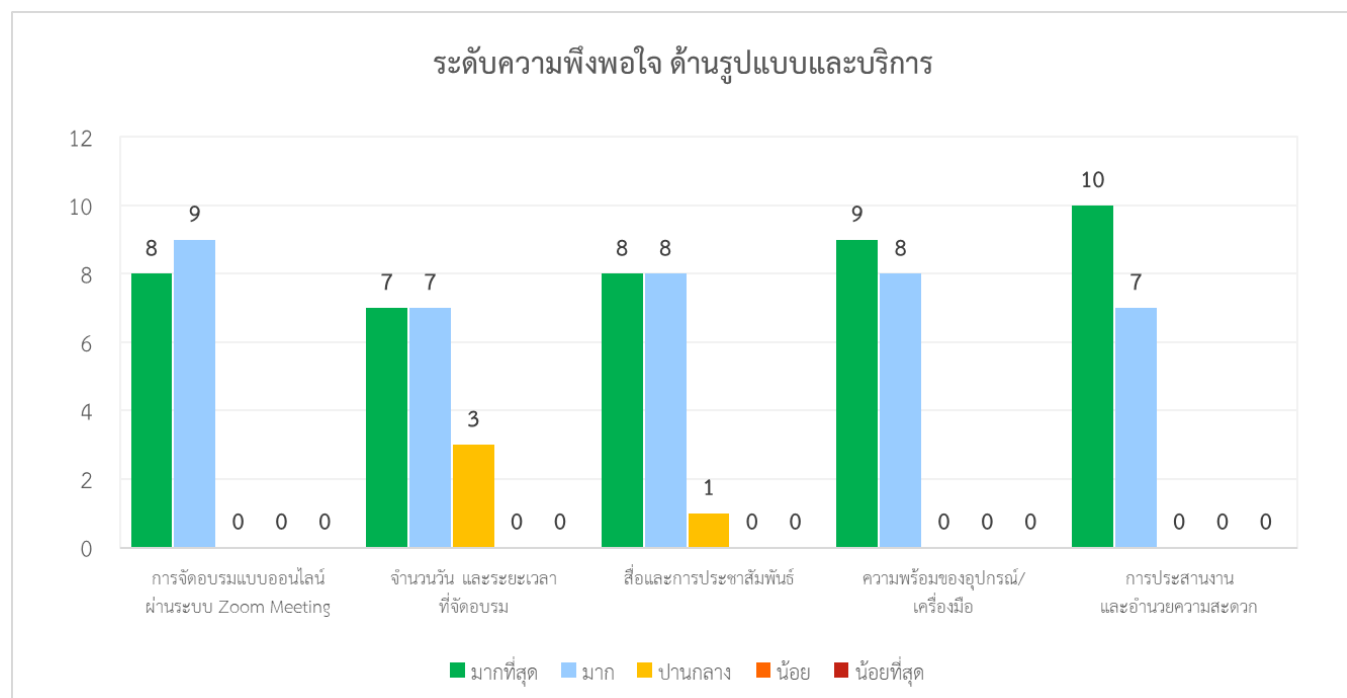
2. ความพึงพอใจ ด้านเนื้อหา / ความรู้ความเข้าใจ / การนำไปใช้ประโยชน์

ผลการประเมินความพึงพอใจ ด้านเนื้อหา / ความรู้ความเข้าใจ / การนำไปใช้ประโยชน์ แบ่งเป็น 6 หัวข้อ ดังนี้ 1) ความเหมาะสมของหัวข้อที่จัดอบรม Topic 1 : Introduction to Artificial Intelligence: An Overview for Researchers ผลประเมินอยู่ในระดับพอใจมากที่สุด จำนวน 13 คน, พอใจมาก จำนวน 4 คน 2) ความเหมาะสมของหัวข้อที่จัดอบรม Topic 2 : Practical Applications of AI in Research: Exploring Tools and Their Usage ผลประเมินอยู่ในระดับพอใจมากที่สุด จำนวน 13 คน, พอใจมาก จำนวน 4 คน 3) ความรู้ความเข้าใจในเรื่องนี้ ก่อนการอบรม มีความรู้ความเข้าใจมากที่สุด จำนวน 7 คน, มีความรู้ความเข้าใจมาก จำนวน 6 คน, มีความรู้ความเข้าใจปานกลาง จำนวน 2 คน, มีความรู้ความเข้าใจน้อย จำนวน 1 คน 4) ความรู้ความเข้าใจในเรื่องนี้หลังการอบรม มีความรู้ความเข้าใจมากที่สุด จำนวน 7 คน, มีความรู้ความเข้าใจมาก จำนวน 9 คน, มีความรู้ความเข้าใจน้อย จำนวน 1 คน แสดงให้เห็นถึงระดับความรู้ความเข้าใจที่เพิ่มขึ้นของผู้ที่เข้าร่วมการอบรม 5) สามารถนำความรู้ไปประยุกต์ใช้ในการปฏิบัติงาน ระดับมากที่สุด จำนวน 8 คน, ระดับมาก จำนวน 8 คน, ระดับน้อย จำนวน 1 คน และหัวข้อที่ 6) สามารถนำความรู้ไปเผยแพร่หรือถ่ายทอดต่อผู้เกี่ยวข้อง อยู่ในระดับมากที่สุด จำนวน 6 คน, ระดับมาก จำนวน 9 คน, ระดับปานกลาง จำนวน 1 คน และระดับน้อย จำนวน 1 คน



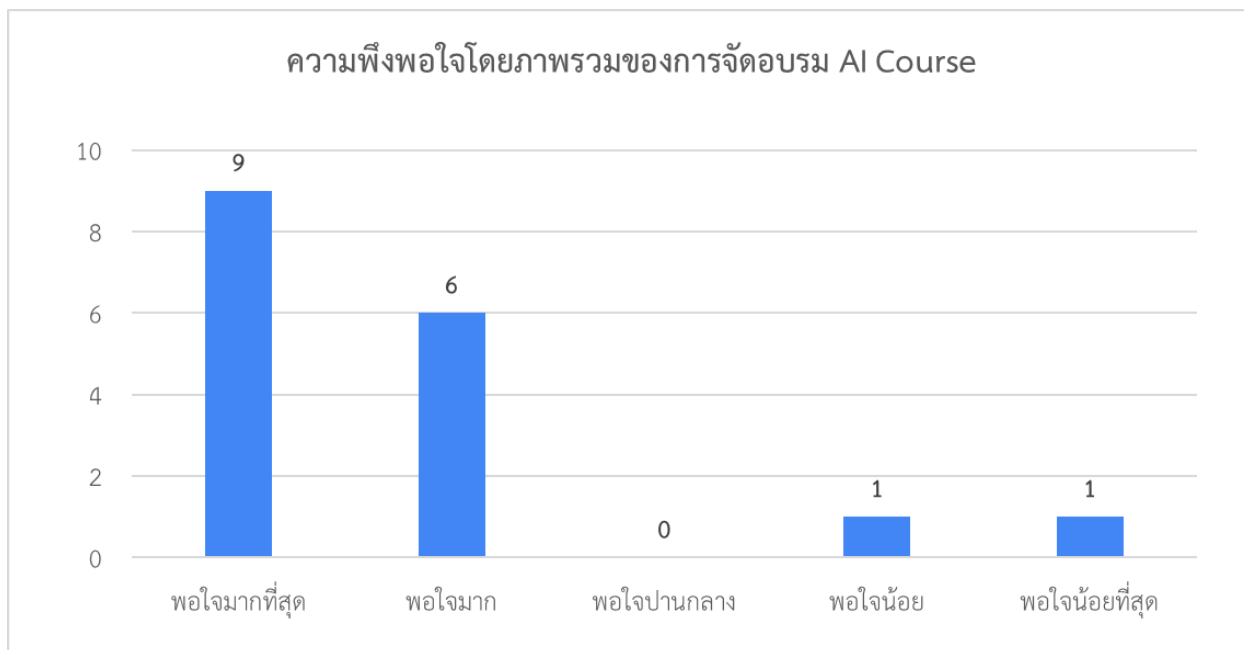
3. ความพึงพอใจ ด้านรูปแบบและบริการ

ผลการประเมินความพึงพอใจ ด้านรูปแบบและบริการ แบ่งเป็น 5 หัวข้อ ดังนี้ 1) การจัดอบรมแบบออนไลน์ผ่านระบบ Zoom Meeting ผลประเมินอยู่ในระดับพอใจมากที่สุด จำนวน 8 คน, พอใจมาก จำนวน 9 คน 2) จำนวนวัน และระยะเวลาที่จัดอบรม ผลประเมินอยู่ในระดับพอใจมากที่สุด จำนวน 7 คน, พอใจมาก จำนวน 7 คน, พอใจปานกลาง จำนวน 3 คน 3) สื่อและการประชาสัมพันธ์ ผลประเมินอยู่ในระดับพอใจมากที่สุด จำนวน 8 คน, พอใจมาก จำนวน 8 คน, พอใจปานกลาง จำนวน 1 คน 4) ความพร้อมของอุปกรณ์/เครื่องมือ ผลประเมินอยู่ในระดับพอใจมากที่สุด จำนวน 9 คน, พอใจมาก จำนวน 8 คน และ 5) การประสานงานและอำนวยความสะดวก ผลประเมินอยู่ในระดับพอใจมากที่สุด จำนวน 10 คน, พอใจมาก จำนวน 7 คน



4. ความพึงพอใจโดยภาพรวมของการจัดอบรม AI Course

การประเมินความพึงพอใจโดยภาพรวมของการจัดอบรม เรื่อง "Applying Artificial Intelligence (AI) for Researchers" แบ่งเป็น 5 ระดับ ได้แก่ พอใจมากที่สุด พอใจมาก พอใจปานกลาง พอใจน้อย และพอใจน้อยที่สุด ผลการประเมินของผู้เข้าร่วมการอบรม จำนวน 17 คน พบว่า มีความพอใจระดับมากที่สุด จำนวน 9 คน, พอใจมาก จำนวน 6 คน, พอใจน้อย จำนวน 1 คน และพอใจน้อยที่สุด จำนวน 1 คน



ความคิดเห็น ข้อเสนอแนะเพิ่มเติม

1. ควรมีเป็นระยะ ๆ
2. ประสงค์อยากจะอบรมแบบมีการฝึกทำไปด้วย แต่ระยะเวลาสั้นเกินไป

หัวข้อการอบรม/สัมมนา ที่ท่านอยากให้สถาบันวิจัยวิทยาศาสตร์สุขภาพจัดในครั้งต่อไป

1. การขอเขียนทุน
2. ตัวอย่างเพิ่มเติมการใช้ AI in research
3. การนำ AI มาช่วยงานอัตโนมัติในองค์กร