

การสร้าง MyGPT สำหรับงานวิเคราะห์ทรัพยากรสารสนเทศ:
ประสบการณ์ของสำนักงานวิทยทรัพยากร จุฬาลงกรณ์มหาวิทยาลัย

รัฐธีร์ ปภัสสุริยโชติ*, อภิวัฒน์ แก้วหะวงษ์

สำนักงานวิทยทรัพยากร จุฬาลงกรณ์มหาวิทยาลัย

254 ถนนพญาไท แขวงวังใหม่ เขตปทุมวัน กรุงเทพมหานคร 10330

Creating MyGPT for Cataloging Tasks: The Experience of
the Office of Academic Resources, Chulalongkorn University

Rathtee Paphatsurichote*, Apiwat Kaewhawong

Office of Academic Resources, Chulalongkorn University

254 Phayathai Road, Wangmai, Pathumwan, Bangkok 10330

E-mail: rathtee.p@car.chula.ac.th

► รับบทความ 04 กุมภาพันธ์ 2568 ► แก้ไขบทความ 22 เมษายน 2568 ► ตอรับบทความ 28 เมษายน 2568

บทคัดย่อ

การศึกษานี้กล่าวถึงการทดลองของสำนักงานวิทยทรัพยากร จุฬาฯ ในการออกแบบ MyGPT สำหรับใช้สร้างระเบียบบรรณานุกรมหนังสือภาษาไทย วัตถุประสงค์ของการศึกษา คือ 1) ประเมินระดับความถูกต้องของระเบียบบรรณานุกรมหนังสือภาษาไทยที่สร้างโดย MyGPT 2) ศึกษารูปแบบความผิดพลาดที่สร้างโดย MyGPT ที่เกิดขึ้นซ้ำ 3) ศึกษาวิธีสร้างระเบียบบรรณานุกรมด้วย MyGPT ที่เหมาะสมในการปฏิบัติงาน วิธีการดำเนินการ คือ สร้าง MyGPT ด้วยคำสั่งที่ระบุขั้นตอนการทำงาน โครงสร้างข้อมูลที่ต้องการ และเขตข้อมูลที่ต้องการให้ปรากฏในระเบียบบรรณานุกรม จากนั้นทดสอบความสามารถในการสร้างระเบียบบรรณานุกรมของ MyGPT โดยใช้หนังสือภาษาไทย จำนวน 10 เล่ม เปรียบเทียบกับระเบียบบรรณานุกรมที่บรรณารักษ์สร้าง ผลการศึกษาพบว่า MyGPT ที่สร้างขึ้นยังไม่สามารถจัดทำระเบียบบรรณานุกรมที่นำไปใช้งานได้ทันทีโดยใช้รูปภาพส่วนนำและส่วนต่าง ๆ จากหนังสือ แต่เมื่อเปลี่ยนวิธีการเป็นการป้อนข้อมูลหนังสือเข้าสู่ MyGPT โดยตรงจะทำให้มีความถูกต้องมากขึ้น แสดงให้เห็นถึงศักยภาพของ MyGPT ที่พัฒนาต่อไปได้หากมีการทำงานร่วมกันระหว่างบรรณารักษ์และปัญญาประดิษฐ์ ด้านความถูกต้องของข้อมูลและโครงสร้างข้อมูล พบว่า เขตข้อมูลที่มีค่าคงที่มีคะแนนด้านความถูกต้องและโครงสร้างสูงกว่าเขตข้อมูลที่มีความซับซ้อนในการลงรายการ และความผิดพลาดที่เกิดขึ้นมากที่สุด คือ ความไม่สอดคล้องกับตามมาตรฐาน MARC, RDA และการกำหนดเลขผู้แต่ง

คำสำคัญ

ปัญญาประดิษฐ์, ปัญญาประดิษฐ์ในงานห้องสมุด, การวิเคราะห์ทรัพยากรสารสนเทศ

Abstract

The objectives of this study were to 1) Evaluate the level of accuracy of MyGPT-generated Thai bibliographic records. 2) Find the error patterns that occur repeatedly. 3) Examine the most appropriate method for creating bibliographic records using MyGPT. The method started with the creation of MyGPT

ผลงานนี้ได้รับการนำเสนอในการประชุมวิชาการระดับชาติ PULINET ครั้งที่ 15 วันที่ 15-17 มกราคม 2568

จัดโดยสำนักวิทยบริการ มหาวิทยาลัยมหาสารคาม และหน่วยงานห้องสมุดมหาวิทยาลัยส่วนภูมิภาค (PULINET)

with instructions about workflow, well-formed data structures, and the fields required in the records. Then generated ten Thai bibliographic records using MyGPT and compared the results with the records created by a librarian. The study found that MyGPT is still unable to generate fully accurate and instantly usable bibliographic records from images of front matter and other book sections, but entering books' data directly into MyGPT improved accuracy. The result showed the potential for enhancing MyGPT through the collaboration between librarians and artificial intelligence. For data accuracy and structure, constant fields received a higher score than those requiring complex recording methods. The most common mistakes were compliance with standards such as MARC, RDA, and cutter number assignments.

Keywords

Artificial Intelligence, Artificial Intelligence in Libraries Works, Cataloging

บทนำ (Introduction)

ปัจจุบันปัญญาประดิษฐ์ (Artificial Intelligence - AI) เป็นเครื่องมือที่เป็นที่รู้จักและถูกนำมาใช้ในงานด้านต่าง ๆ อย่างแพร่หลาย เช่น การช่วยสร้างเนื้อหาทั้งในรูปแบบตัวอักษร รูปภาพ ภาพเคลื่อนไหว และเสียง การตรวจและแก้ไขไวยากรณ์ภาษาอังกฤษ การวิเคราะห์ข้อมูล การสืบค้นข้อมูลเพื่อช่วยในการเขียนงานวิจัย และการช่วยเขียนโปรแกรมคอมพิวเตอร์ (Chan & Colloton, 2024, p. 13) ปัญญาประดิษฐ์ที่ทำให้การใช้ปัญญาประดิษฐ์เริ่มเป็นที่รู้จักและนิยมอย่างแพร่หลาย คือ ChatGPT ซึ่งเปิดตัวในช่วงปี พ.ศ. 2565 (OpenAI, 2022)

ChatGPT เป็นปัญญาประดิษฐ์ที่พัฒนาโดยบริษัท OpenAI จัดเป็นปัญญาประดิษฐ์แบบรู้สร้าง (Generative AI) จุดเด่นของปัญญาประดิษฐ์รูปแบบนี้คือ สร้างเนื้อหาใหม่ ๆ ได้ ไม่ว่าจะเป็นเนื้อหาประเภท ข้อความ ภาพเคลื่อนไหว รูปภาพ หรือเนื้อหาในรูปแบบอื่น ๆ โดยข้อมูลที่นำมาใช้ในการสร้างเนื้อหาใหม่มาจากแบบจำลองภาษาขนาดใหญ่ (Large Language Model - LLM) ที่ใช้เป็นข้อมูลในการสอนปัญญาประดิษฐ์ โดยปัญญาประดิษฐ์จะเรียนรู้รูปแบบ โครงสร้าง และลักษณะของข้อมูลจากแบบจำลองภาษาขนาดใหญ่เพื่อคาดเดาและสร้างเนื้อหาตามที่ใช้ต้องการ โดยการป้อนคำสั่งด้วยภาษาธรรมชาติ (Natural Language) ซึ่งมีลักษณะ คือ เป็นภาษาที่มนุษย์ใช้สื่อสารในชีวิตประจำวัน (Chan & Colloton, 2024, p. 45) โดยการป้อนคำสั่งแก่ปัญญาประดิษฐ์ในรูปแบบนี้เรียกว่า ข้อความพร้อมรับ (Prompt)

ในแวดวงห้องสมุดมีความสนใจในการนำปัญญาประดิษฐ์มาใช้งานห้องสมุดเช่นกัน จากงานวิจัยของ Mannheimer et al. (2024) พบว่า ลักษณะงานที่ห้องสมุดนำปัญญาประดิษฐ์มาใช้มากที่สุดสามอันดับแรก ได้แก่ งานด้านการสกัดเมตาดาตา (Metadata Extraction) การประมวลผลภาษาธรรมชาติ (Natural Language Processing - NLP) และการจดจำรูปภาพ (Image Recognition) โดยลักษณะข้อมูลที่นำมาใช้กับปัญญาประดิษฐ์มากที่สุดสามอันดับแรกคือ ข้อมูลตัวอักษรที่อยู่ในรูปแบบดิจิทัลตั้งแต่แรก (Text-Born Digital) รูปภาพที่ผ่านกระบวนการแปลงเป็นดิจิทัล (Images-Digitized) และข้อมูลตัวอักษรที่ผ่านกระบวนการแปลงเป็นดิจิทัล (Text-Digitized) โดยนำมาใช้เพื่อสร้างเมตาดาตา การสืบค้นทรัพยากรสารสนเทศของห้องสมุด จึงอาจกล่าวได้ว่า เป็นงานที่อยู่ในขอบข่ายของงานวิเคราะห์ทรัพยากรสารสนเทศ

สำนักงานวิทยทรัพยากร จุฬาลงกรณ์มหาวิทยาลัย เป็นหอสมุดกลางของมหาวิทยาลัย มีหน้าที่เป็นหน่วยงานที่จัดเตรียมระบบห้องสมุดอัตโนมัติ การอบรมด้านการลงรายการทรัพยากรสารสนเทศให้กับผู้ปฏิบัติงานในห้องสมุดคณะ วิทยาลัย และสถาบันวิจัยต่าง ๆ ของจุฬาลงกรณ์มหาวิทยาลัย รวมถึงดูแลข้อมูลระเบียบบรรณานุกรมที่อยู่ในระบบห้องสมุดอัตโนมัติ ปัญหาหนึ่งที่พบเกี่ยวกับงานวิเคราะห์ทรัพยากรสารสนเทศ คือ การขาดแคลนบุคลากรหรือกำลังคนที่ปฏิบัติงานด้านนี้ เช่น ห้องสมุดต่าง ๆ อาจมีบรรณารักษ์ประจำอยู่เพียงอัตราเดียว และอาจไม่มีพื้นฐานด้านการวิเคราะห์ทรัพยากรสารสนเทศ

แต่ห้องสมุดยังคงต้องจัดหาทรัพยากรสารสนเทศเพื่อให้บริการในห้องสมุด ทำให้ยังต้องมีการนำข้อมูลทรัพยากรสารสนเทศเข้าสู่ฐานข้อมูลห้องสมุด จากปัญหาดังกล่าว สำนักงานวิทยทรัพยากร จึงได้ริเริ่มโครงการศึกษาความเป็นไปได้ในการนำปัญญาประดิษฐ์มาใช้ในการวิเคราะห์ทรัพยากรสารสนเทศ โดยปัญญาประดิษฐ์ที่เลือกใช้ คือ ChatGPT ซึ่งมีฟังก์ชันการทำงานที่เรียกว่า MyGPT ที่สามารถใช้สร้างตัวแบบ ChatGPT ที่มีความถนัดเฉพาะงานได้ อีกทั้งยังสามารถแบ่งปันตัวแบบที่สร้างให้ผู้อื่นใช้งานได้ด้วย (OpenAI, 2024b) จึงสอดคล้องกับปัญหาที่พบ และอาจเป็นหนึ่งในองค์ประกอบของการแก้ไขปัญหาตามที่กล่าวมาข้างต้น

บทความนี้กล่าวถึงการทดลองการออกแบบ MyGPT เพื่อใช้ในการวิเคราะห์ทรัพยากรสารสนเทศ โดยสำนักงานวิทยทรัพยากร จุฬาลงกรณ์มหาวิทยาลัย และสิ่งที่พบจากการทดสอบความสามารถของ MyGPT ที่สร้างขึ้น มีวัตถุประสงค์ คือ ศึกษากระบวนการบรรณานุกรมที่สร้างขึ้นโดย MyGPT ในด้านความถูกต้อง และความผิดพลาดที่พบบ่อย และศึกษาวิธีที่เหมาะสมในการสร้างระเบียบบรรณานุกรมทรัพยากรสารสนเทศที่เป็นหนังสือภาษาไทยด้วย MyGPT

วัตถุประสงค์ (Objective)

1. ประเมินระดับความถูกต้องของระเบียบบรรณานุกรมหนังสือภาษาไทยที่สร้างโดย MyGPT ในด้านความถูกต้องของข้อมูลระเบียบบรรณานุกรมและด้านความถูกต้องของโครงสร้างระเบียบบรรณานุกรม
2. ศึกษารูปแบบความผิดพลาดของระเบียบบรรณานุกรมหนังสือภาษาไทยที่สร้างโดย MyGPT ที่เกิดขึ้นซ้ำ
3. ศึกษาวิธีที่เหมาะสมในการสร้างระเบียบบรรณานุกรมทรัพยากรสารสนเทศที่เป็นหนังสือภาษาไทยด้วย MyGPT ที่เหมาะสมกับสภาพแวดล้อมการปฏิบัติงานของผู้ปฏิบัติงานในห้องสมุด

วิธีการดำเนินการ (Methodology)

1. ศึกษาวรรณกรรมที่เกี่ยวข้อง
2. ร่างความต้องการและลักษณะการใช้งาน MyGPT โดยแนวคิดพื้นฐานของการออกแบบ คือ 1) มีความง่ายในการใช้งาน ผู้ใช้งานไม่ต้องใช้ข้อความพร้อมรับที่ซับซ้อนหรือใช้น้อยที่สุด และ 2) ผู้ใช้งานไม่จำเป็นต้องมีความรู้พื้นฐานด้านการวิเคราะห์ทรัพยากรสารสนเทศก็สามารถสร้างระเบียบบรรณานุกรมได้

จากแนวคิดพื้นฐานของการออกแบบ จึงเกิดเป็นแนวทางในการออกแบบ MyGPT สองลักษณะ คือ 1) นำเข้าข้อมูลที่เป็นรูปภาพส่วนประกอบที่มีข้อมูลสำคัญของหนังสือเข้าสู่ MyGPT เช่น ชื่อเรื่อง ชื่อผู้แต่ง สถานที่พิมพ์ สำนักพิมพ์ ปีพิมพ์ สารบัญ และเรื่องย่อหรือสาระสังเขปที่ปรากฏที่หน้าปกของหนังสือ

3. สร้าง MyGPT ตามร่างความต้องการ โดยข้อมูลที่ใช้ในการสอน MyGPT ให้สามารถสร้างระเบียบบรรณานุกรมได้ประกอบด้วย

3.1 คำสั่งเกี่ยวกับขั้นตอนการทำงานของ MyGPT และลักษณะของผลลัพธ์ที่ต้องการ โดยได้ระบุว่าเป็นบรรณานุกรมที่จัดทำ ต้องมีเขตข้อมูล MARC ขั้นต่ำ ดังนี้

- MARC Leader - ส่วนนำของชุดข้อมูลรูปแบบ MARC
- 008 - เขตข้อมูลระบุองค์ประกอบของข้อมูลแปรผันตามประเภททรัพยากรสารสนเทศ
- 020 - เลขมาตรฐานสากลประจำหนังสือ (ISBN) (กรณีที่มีข้อมูลปรากฏในรูปภาพ)
- 040 - เขตข้อมูลระบุหน่วยงาน ภาษา และมาตรฐานที่ใช้ในการลงรายการ กำหนดให้บันทึกเฉพาะ

มาตรฐานที่ใช้ในการลงรายการ ในเขตข้อมูลย่อย e (\$erda)

- 336, 337, 338 - รูปแบบของเนื้อหา อุปกรณ์ที่ต้องใช้ในการเข้าถึงเนื้อหา และสื่อที่เก็บเนื้อหา
- 082 - เลขหมู่ระบบจุดทศนิยมดิวอี้ (Dewey Decimal Classification - DDC) และเลขผู้แต่ง (Cutter

Number)

- 100 - ผู้แต่งหลัก (กรณีที่มีข้อมูลปรากฏในรูปภาพ)
- 245 - ชื่อเรื่องและข้อมูลระบุบุคคลหรือองค์กรที่มีความรับผิดชอบด้านเนื้อหา
- 264 - สถานที่พิมพ์ สำนักพิมพ์ และปีที่พิมพ์
- 505 - สารบัญ (กรณีที่มีข้อมูลปรากฏในรูปภาพ)
- 520 - เรื่องย่อหรือบทคัดย่อ (กรณีที่มีข้อมูลปรากฏในรูปภาพ)
- 6XX - หัวเรื่องหอสมุดรัฐสภาอเมริกัน (Library of Congress Subject Headings - LCSH)
- 7XX - ผู้แต่งร่วม หน่วยงาน หรือผู้รับผิดชอบเนื้อหาในด้านอื่น ๆ นอกจากผู้แต่งหลัก

สาเหตุที่ไม่รวมเขตข้อมูล 300 ซึ่งใช้บันทึกรายละเอียดด้านลักษณะทางกายภาพของทรัพยากรสารสนเทศลงในเขตข้อมูลที่ต้องปรากฏในขั้นต่ำ เพราะการนำเข้าข้อมูลโดยใช้รูปภาพไม่มีข้อมูลเลขหน้าที่ปรากฏในหน้าสุดท้ายของหนังสือ

ส่วนสาเหตุที่ระบุให้เขตข้อมูล 6XX ใช้เป็นหัวเรื่องหอสมุดรัฐสภาอเมริกัน เนื่องจากการเข้าถึงข้อมูลหัวเรื่องภาษาไทยออนไลน์โดยคณะทำงานกลุ่มวิเคราะห์ทรัพยากรสารสนเทศ ห้องสมุดสถาบันอุดมศึกษา (<https://webhost2.car.chula.ac.th/thaiccweb/>) ต้องมีการลงชื่อเข้าใช้ ดังนั้น MyGPT จึงไม่สามารถใช้งานข้อมูลจากเว็บไซต์ดังกล่าวได้ จึงได้ระบุให้ MyGPT สืบค้นหัวเรื่องจากเว็บไซต์ข้อมูลเชื่อมโยงของหอสมุดรัฐสภาอเมริกัน (Library of Congress Linked Data Service - <https://id.loc.gov/>) ซึ่งเป็นเว็บไซต์ที่เข้าถึงได้โดยไม่มีค่าใช้จ่ายและไม่ต้องการลงทะเบียนหรือลงชื่อเพื่อเข้าใช้

นอกจากนี้ยังได้ระบุในคำสั่งว่าหากไม่มีข้อมูลในส่วนใด ไม่ต้องบันทึกข้อมูลหรือแสดงเขตข้อมูลนั้น และ MyGPT ต้องทำตามคำสั่งที่ระบุไว้อย่างเคร่งครัด

3.2 คำสั่งระบุรายละเอียดรูปแบบระเบียบบรรณานุกรมที่ MyGPT ต้องจัดทำ ตัวอย่างเช่น รูปแบบการบันทึกข้อมูลชื่อผู้แต่ง ชื่อผู้รับผิดชอบด้านต่าง ๆ และปีที่พิมพ์ โดยมีรายละเอียด เช่น หากชื่อผู้แต่งเป็นชาวไทย ให้บันทึกข้อมูลในรูปแบบ [ชื่อ] [นามสกุล], \$ผู้แต่ง หรือปีที่พิมพ์ หากภาษาหลักของหนังสือเป็นภาษาไทย ให้บันทึกปีที่พิมพ์ในเขตข้อมูล 264 เขตข้อมูลย่อย c เป็นปี พ.ศ. เป็นต้น

3.3 คำสั่งระบุรายละเอียดโครงสร้างภายในเขตข้อมูล MARC ว่าในแต่ละเขตข้อมูลประกอบด้วยตัวบ่งชี้ (Indicators) เขตข้อมูลย่อย (Subfields) เครื่องหมายวรรคตอนต่าง ๆ และมีรูปแบบการบันทึกข้อมูลในแต่ละเขตข้อมูลอย่างไร เช่น เขตข้อมูล 245 ของหนังสือที่มีผู้แต่งหลักปรากฏในเขตข้อมูล 100 ตัวบ่งชี้แรกควรเป็น 1 และมีองค์ประกอบของเขตข้อมูล ดังนี้ 245 10\$a[ชื่อเรื่อง] :\$b[ชื่อเรื่องย่อหรือส่วนของชื่อเรื่องอื่น ๆ] /\$c[ผู้แต่ง หรือ ผู้รับผิดชอบอื่น ๆ] ทั้งนี้สามารถดูข้อมูลคำสั่งชุดที่ใช้ในการสร้าง MyGPT ในบทความนี้ได้ที่ส่วน “ข้อมูลเพิ่มเติม” ซึ่งปรากฏอยู่ท้ายบทความก่อนถึงส่วนรายการอ้างอิง

4. ทดสอบความสามารถในการสร้างระเบียบบรรณานุกรมหนังสือภาษาไทย ดังนี้

4.1 สุ่มเลือกหนังสือที่มีเนื้อหาเป็นภาษาไทย จำนวน 10 ชื่อเรื่อง เพื่อนำมาทดลองใช้สร้างระเบียบบรรณานุกรมด้วย MyGPT ที่สร้างขึ้น การทดสอบจะเริ่มจากการใช้วิธีการสร้างระเบียบบรรณานุกรมจากรูปภาพส่วนประกอบสำคัญของหนังสือเป็นอันดับแรก โดยทำการแปลงส่วนประกอบต่อไปนี้เป็นไฟล์รูปภาพแบบที่มีนามสกุลไฟล์เป็น .jpg ดังนี้ 1) หน้าปก 2) หน้าปกใน 3) หน้าลิขสิทธิ์ 4) สารบัญ และ 5) ปกหลัง (เฉพาะกรณีที่ปกหลังมีข้อความสรุปเนื้อหาของหนังสือเล่มนั้น ๆ) สาเหตุที่เลือกวิธีการนี้เป็นอันดับแรก เนื่องจากเป็นวิธีการที่ง่ายที่สุด ใช้เพียงอุปกรณ์ที่ใช้งานในชีวิตประจำวัน เช่น โทรศัพท์มือถือที่มีกล้องก็สามารถทำได้ และผู้ใช้งานไม่จำเป็นต้องใช้ข้อความพร้อมรับหรือมีปฏิสัมพันธ์กับ MyGPT ก็สามารถสร้างระเบียบบรรณานุกรมได้ ทั้งนี้ ทำการทดสอบ 2 รูปแบบ คือ 1) สร้างระเบียบบรรณานุกรมของหนังสือภาษาไทยทั้ง 10 ชื่อเรื่อง ภายในการสนทนาเดียว (การทดสอบแบบ A) และ 2) สร้างระเบียบบรรณานุกรมของหนังสือภาษาไทยโดยเปิดการสนทนาใหม่ทุกครั้งในการสร้างระเบียบบรรณานุกรมของแต่ละชื่อเรื่อง (การทดสอบแบบ B)

เนื่องจาก จากประสบการณ์ของผู้เขียน พบว่า ผลลัพธ์ที่ได้จาก MyGPT ที่ใช้เพียงหนึ่งการสนทนามีผลแตกต่างกับการเริ่มต้นการสนทนาใหม่ทุกครั้งที่ใช้งาน โดยช่วงเวลาที่ดำเนินการทดสอบ คือ ระหว่างวันที่ 17-24 พฤศจิกายน 2567 และเวอร์ชัน ChatGPT ที่ใช้คือ ChatGPT Team ซึ่งใช้ตัวแบบ GPT-4o

ทั้งนี้ หากพิจารณาแล้วพบว่า ระดับความถูกต้องของระเบียบบรรณานุกรมที่สร้างโดย MyGPT ในด้านความแม่นยำและด้านความถูกต้องของโครงสร้างระเบียบบรรณานุกรม ซึ่งจะดำเนินการต่อในขั้นตอนที่ 5 ไม่เป็นที่น่าพึงพอใจ จะดำเนินการทดสอบโดยใช้รูปแบบที่สอง คือ การสร้างระเบียบบรรณานุกรมจากการป้อนข้อมูลหนังสือเข้าสู่ MyGPT โดยตรง โดยใช้โครงร่างสำหรับกรอกข้อมูลที่ MyGPT สร้างขึ้นเพื่อให้ผู้ใช้งานป้อนข้อมูลที่จำเป็นในการสร้างระเบียบบรรณานุกรม เช่น ชื่อเรื่อง ชื่อผู้แต่ง สถานที่พิมพ์ ปีพิมพ์ สำนักพิมพ์ จำนวนหน้าเรื่องย่อหรือสาระสังเขป เป็นต้น และเปรียบเทียบว่าได้ผลของการสร้างระเบียบบรรณานุกรมที่ดีขึ้นหรือไม่ เพื่อใช้เป็นข้อมูลประกอบการคัดเลือกวิธีการที่เหมาะสมกว่าในการสร้างระเบียบบรรณานุกรมด้วย MyGPT

4.2 ตรวจสอบความถูกต้องของระเบียบบรรณานุกรมที่สร้างโดย MyGPT โดยเปรียบเทียบกับระเบียบที่บรรณารักษ์วิเคราะห์ทรัพยากรสารสนเทศเป็นผู้สร้าง

5. ประเมินระดับความถูกต้องของระเบียบบรรณานุกรมที่สร้างโดย MyGPT ในด้านความแม่นยำและด้านความถูกต้องของโครงสร้างระเบียบบรรณานุกรม และหารูปแบบความผิดพลาดที่เกิดขึ้นซ้ำ ๆ โดยเกณฑ์การให้คะแนน มีดังนี้

5.1 ความถูกต้อง แบ่งการลงรหัสเป็นสามระดับ คือ

“2” หมายถึง ข้อมูลที่ MyGPT บันทึกมีความถูกต้องและตรงตามที่บรรณารักษ์บันทึกข้อมูล

“1” หมายถึง ข้อมูลที่ MyGPT บันทึกมีความถูกต้องแต่ไม่ตรงกับที่บรรณารักษ์บันทึกข้อมูล

“0” หมายถึง ข้อมูลที่ MyGPT บันทึกไม่ถูกต้อง เช่น สะกดคำผิด เป็นข้อมูลที่ไม่มีปรากฏในตัวเล่มหรือเป็นข้อมูลที่ไม่มีอยู่จริง

5.2 โครงสร้างข้อมูล แบ่งการลงรหัสเป็นสองระดับ คือ

“1” หมายถึง ข้อมูลที่ MyGPT บันทึกมีการใช้ตัวบ่งชี้ เขตข้อมูลย่อย และเครื่องหมายวรรคตอนที่ใช้ในเขตข้อมูลถูกต้องตามรูปแบบ MARC และ RDA

“0” หมายถึง ข้อมูลที่ MyGPT บันทึกมีการใช้ตัวบ่งชี้ เขตข้อมูลย่อย และเครื่องหมายวรรคตอนที่ใช้ในเขตข้อมูลผิดพลาด หรือไม่ปฏิบัติตามรูปแบบ MARC และ RDA

5.3 วิธีคำนวณคะแนนด้านความถูกต้องและด้านโครงสร้างระเบียบบรรณานุกรม มีดังนี้

5.3.1 ด้านความถูกต้อง นับความถี่ของรหัสความถูกต้อง จากนั้นคูณด้วยค่ารหัสความถูกต้อง ตัวอย่างเช่น หากเขตข้อมูล 020 มีความถี่ของรหัสความถูกต้อง “2” จำนวน 9 ครั้ง และความถี่ของรหัสความถูกต้อง “1” จำนวน 1 ครั้ง สามารถคำนวณคะแนนความถูกต้องโดยนำความถี่ของรหัสความถูกต้องคูณด้วยค่ารหัสความถูกต้อง ดังนั้น คะแนนด้านความถูกต้องของเขตข้อมูล 020 จะคำนวณได้โดย $(9 \times 2) + (1 \times 1) = 19$ คะแนน

5.3.2 ด้านโครงสร้างระเบียบบรรณานุกรม นับความถี่ของรหัสระดับความถูกต้องของโครงสร้างระเบียบบรรณานุกรม และใช้ความถี่ดังกล่าวเป็นคะแนนด้านความถูกต้องของโครงสร้างระเบียบบรรณานุกรม ตัวอย่างเช่น หากเขตข้อมูล 020 มีความถี่ของรหัสความถูกต้องของโครงสร้างระเบียบบรรณานุกรม “1” จำนวน 9 ครั้ง คะแนนด้านโครงสร้างระเบียบบรรณานุกรม จะเท่ากับ $9 \times 1 = 9$ คะแนน

6. สรุปผลและสิ่งที่ค้นพบ

ผลการดำเนินการและอภิปรายผล (Result and Discussion)

เมื่อทดสอบและตรวจสอบความถูกต้องของระเบียบบรรณานุกรมโดยเปรียบเทียบกับระเบียบที่บรรณารักษ์วิเคราะห์ทรัพยากรสารสนเทศเป็นผู้สร้าง ได้ผลตามรายละเอียดตามตารางที่ 1 ดังนี้

คะแนนความถูกต้อง

เมื่อพิจารณาคะแนนความถูกต้องตามตารางที่ 1 อาจสรุปได้ว่า เขตข้อมูลที่ได้คะแนนความถูกต้องในภาพรวมร้อยละ 50 ขึ้นไป ทั้งสองการทดสอบ คือ เขตข้อมูล 020, 040, 336, 337, 338, 082 และ 100 ซึ่งเมื่อพิจารณาลักษณะของเขตข้อมูล พบว่า ส่วนใหญ่เป็นเขตข้อมูลที่มีค่าคงที่ เช่น 040 และ 336, 337, 338 ส่วนเขตข้อมูลที่มีค่าแปรผันไปตามข้อมูลหนังสือแต่ละเล่ม ได้แก่ เขตข้อมูล 020, 082 และ 100 ซึ่งแม้จะเป็นเขตข้อมูลที่มีค่าผันเปลี่ยนแต่เป็นเขตข้อมูลที่มีรูปแบบการลงรายการที่ไม่ซับซ้อน สามารถสกัดข้อมูลออกจากรูปภาพได้ง่าย และมีรูปแบบการลงรายการที่ค่อนข้างชัดเจน เช่น เขตข้อมูล 020 ซึ่งเป็นเลขมาตรฐานสากลประจำหนังสือ (ISBN) ที่ลงรายการตามข้อมูลที่ปรากฏในหนังสือและเป็นเลขที่มีจำนวนหลักที่แน่นอน คือ 10 หรือ 13 หลัก หรือในเขตข้อมูล 082 และ 100 ก็มีรูปแบบการลงรายการไม่ซับซ้อน เช่น เขตข้อมูล 082 ซึ่งเป็นเขตข้อมูลบันทึกเลขหมู่ตามการจัดหมู่หนังสือแบบทศนิยมดิวอี้ (Dewey Decimal Classification) มีรูปแบบการลงรายการที่แน่นอน คือ บันทึกเลขหมู่ตามหมวดหมู่ของหนังสือลงในเขตข้อมูลย่อย a ส่วนเขตข้อมูล 100 ซึ่งเป็นชื่อผู้แต่ง ลงรายการโดยใช้ชื่อและนามสกุลของผู้แต่งหลักซึ่งมักเป็นชื่อที่ปรากฏเด่นชัดที่สุดที่หน้าปกหรือหน้าปกในของหนังสือ ตามด้วยเขตข้อมูลย่อย e ระบุว่าเป็นผู้แต่ง

ส่วนเขตข้อมูลที่ได้คะแนนความถูกต้องในภาพรวมน้อยที่สุดทั้งสองการทดสอบ คือ เขตข้อมูล 008, เลขผู้แต่ง, 650 และ 7XX ซึ่งเมื่อพิจารณาแล้วจะเห็นว่าเป็นเขตข้อมูลที่มีความซับซ้อนในการลงรายการ หรือเป็นการลงรายการที่ต้องใช้ค่าตามที่ปรากฏในคู่มือต่าง ๆ

เขตข้อมูล 008 เป็นเขตข้อมูลที่สรุปลักษณะของระเบียบบรรณานุกรมระเบียบนั้นเพื่อใช้ประโยชน์ในการสืบค้น และมีการลงรายละเอียดในเชิงลึก เช่น ช่วงปีที่พิมพ์ มีภาพประกอบหรือเนื้อหาหลักพิเศษอย่างไร มีบรรณานุกรมและบรรณานุกรมหรือไม่ เป็นหนังสือที่ระลึกในโอกาสพิเศษ หรือกลุ่มผู้อ่านคือช่วงอายุใด เป็นต้น จึงอาจเป็นเหตุผลที่ทำให้การสกัดข้อมูลเพื่อลงรายการในเขตข้อมูลนี้มีคะแนนความถูกต้องในภาพรวมในระดับที่ไม่สูงมาก

เขตข้อมูลที่ได้คะแนนต่ำที่สุด (ร้อยละ 0) มี 2 เขตข้อมูล คือ เลขผู้แต่ง และ 650

ข้อมูลเลขผู้แต่ง เป็นเขตข้อมูลที่แปรผันตามชื่อผู้แต่งทรัพยากรสารสนเทศ ซึ่งหากเป็นผู้แต่งชาวไทยหรือเขียนชื่อด้วยภาษาไทย ตารางเลขผู้แต่งที่สำนักงานวิทยทรัพยากร จุฬาฯ ใช้คือ ตารางเลขผู้แต่งหนังสือภาษาไทยสำเร็จรูปของหอสมุดแห่งชาติ ซึ่งแม้ผู้เขียนจะได้จัดเตรียมไฟล์ในรูปแบบ PDF ของตารางเลขผู้แต่งดังกล่าวไว้เป็นข้อมูลเพื่อให้ MyGPT ใช้ และกำหนดคำสั่งให้ MyGPT เลือกใช้เลขผู้แต่งจากไฟล์ดังกล่าวแล้ว แต่พบว่า เลขผู้แต่งที่ MyGPT กำหนดไม่ตรงกับตารางเลขผู้แต่ง สาเหตุอาจเกิดจากไฟล์ PDF ที่ผู้เขียนเตรียมไว้นั้นอยู่ในรูปแบบที่ MyGPT ไม่สามารถทำความเข้าใจได้ อาจต้องลองเปลี่ยนรูปแบบจากไฟล์ PDF เป็นไฟล์ลักษณะอื่นที่ MyGPT สามารถทำความเข้าใจได้ เช่น ไฟล์ในรูปแบบตาราง เป็นต้น

เขตข้อมูล 650 หรือเขตข้อมูลสำหรับบันทึกข้อมูลหัวเรื่อง เป็นเขตข้อมูลที่ได้คะแนนต่ำอีกหนึ่งเขตข้อมูล เป็นเพราะการกำหนดหัวเรื่องให้กับทรัพยากรสารสนเทศนั้นจำเป็นต้องใช้หัวเรื่องหรือคำที่ปรากฏอยู่ในคู่มือ หรือฐานข้อมูล

หัวเรื่อง เช่น หัวเรื่องหอสมุดรัฐสภาอเมริกัน และหัวเรื่องภาษาไทยออนไลน์ โดยคณะทำงานกลุ่มวิเคราะห์ทรัพยากรสารสนเทศ หอสมุดสถาบันอุดมศึกษา ซึ่งแม้จะระบุในคำสั่งของ MyGPT ให้สืบค้นหัวเรื่องตามเว็บไซต์ที่เป็นแหล่งของหัวเรื่องแล้ว แต่ก็ยังไม่สามารถกำหนดหัวเรื่องได้ตรงตามที่บรรณารักษ์กำหนดได้หรือเป็นหัวเรื่องที่ใช้งานได้

ตารางที่ 1 คะแนนด้านความถูกต้องและด้านโครงสร้างระเบียบบรรณานุกรมที่สร้างโดย MyGPT ด้วยวิธีการอัปโหลดรูปภาพ

เชตข้อมูล MARC	Leader		008		020		040		336, 337, 338		082		เลขผู้แต่ง		100		245		264		505		520		650		7XX		ผลรวม	
คะแนน	A*	B**	A*	B**	A*	B**	A*	B**	A*	B**	A*	B**	A*	B**	A*	B**	A*	B**	A*	B**	A*	B**	A*	B**	A*	B**	A*	B**	A*	B**
จำนวนครั้งที่ถูกต้องระดับ 2	0	1	0	1	9	10	5	8	8	10	5	5	0	1	5	8	4	3	4	3	2	3	0	0	0	0	1	2	43	55
จำนวนครั้งที่ถูกต้องระดับ 1	6	7	3	0	1	0	5	2	0	0	4	1	0	0	2	0	2	1	2	1	0	0	6	9	2	0	1	0	34	21
คะแนนรวมด้านความถูกต้อง	6	9	3	2	19	20	15	18	16	20	14	11	0	2	12	16	10	7	10	7	4	6	6	9	2	0	3	4	120	131
	(30)	(45)	(15)	(10)	(95)	(100)	(75)	(90)	(80)	(100)	(70)	(55)	(0)	(10)	(60)	(80)	(50)	(35)	(50)	(35)	(20)	(30)	(30)	(45)	(10)	(0)	(15)	(20)	(43)	(47)
คะแนนด้านโครงสร้าง	9	10	10	9	9	9	9	9	8	10	8	7	5	5	5	8	8	8	1	0	0	0	9	10	5	2	5	8	91	95
	(90)	(100)	(100)	(90)	(90)	(90)	(90)	(90)	(80)	(100)	(80)	(70)	(50)	(50)	(50)	(80)	(80)	(80)	(10)	(0)	(0)	(0)	(90)	(100)	(50)	(20)	(50)	(80)	(65)	(68)
คะแนนรวมทั้งหมด***	15	19	13	11	28	29	24	27	24	30	22	18	5	7	17	24	18	15	11	7	4	6	15	19	7	2	8	12	211	226
	(50)	(64)	(44)	(37)	(94)	(97)	(80)	(90)	(80)	(100)	(74)	(60)	(17)	(24)	(57)	(80)	(60)	(50)	(37)	(24)	(14)	(20)	(50)	(64)	(24)	(7)	(27)	(40)	(51)	(54)
เปรียบเทียบร้อยละ A และ B	เพิ่มขึ้น 14		ลดลง 7		เพิ่มขึ้น 3		เพิ่มขึ้น 10		เพิ่มขึ้น 20		ลดลง 14		เพิ่มขึ้น 7		เพิ่มขึ้น 23		ลดลง 10		ลดลง 13		เพิ่มขึ้น 6		เพิ่มขึ้น 14		ลดลง 17		เพิ่มขึ้น 13		เพิ่มขึ้น 3	

ตารางที่ 2 เปรียบเทียบคะแนนด้านความถูกต้องและด้านโครงสร้างระเบียบบรรณานุกรมที่สร้างโดย MyGPT ด้วยวิธีการอัปโหลดรูปภาพกับการป้อนข้อมูลโดยตรง

เชตข้อมูล MARC	Leader		008		020		040		336, 337, 338		082		เลขผู้แต่ง		100		245		264		505		520		650		7XX		ผลรวม	
คะแนน	B**	C*	B**	C*	B**	C*	B**	C*	B**	C*	B**	C*	B**	C*	B**	C*	B**	C*	B**	C*	B**	C*	B**	C*	B**	C*	B**	C*	B**	C*
จำนวนครั้งที่ถูกต้องระดับ 2	1	0	1	2	10	10	8	10	10	10	5	3	1	0	8	8	3	10	3	0	3	10	0	10	0	1	2	9	55	83
จำนวนครั้งที่ถูกต้องระดับ 1	7	10	0	0	0	0	2	0	0	0	1	6	0	0	0	2	1	0	1	0	0	0	9	0	0	4	0	1	21	23
คะแนนรวมด้านความถูกต้อง	9	10	2	4	20	20	18	20	20	20	11	12	2	0	16	18	7	20	7	0	6	20	9	20	0	6	4	19	131	189
คะแนนด้านโครงสร้าง	10	10	9	10	9	10	9	10	10	10	7	10	5	10	8	10	8	3	0	10	0	10	10	10	2	1	8	10	95	124
คะแนนรวมทั้งหมด***	19	20	11	14	29	30	27	30	30	30	18	22	7	10	24	28	15	23	7	10	6	30	19	30	2	7	12	29	226	313
ร้อยละ	64	67	37	47	97	100	90	100	100	100	60	74	24	34	80	94	50	77	24	34	20	100	64	100	7	24	40	97	54	75
เปรียบเทียบร้อยละ A และ B	เพิ่มขึ้น 3		เพิ่มขึ้น 10		เพิ่มขึ้น 3		เพิ่มขึ้น 10		เท่าเดิม -		เพิ่มขึ้น 14		เพิ่มขึ้น 10		เพิ่มขึ้น 14		เพิ่มขึ้น 27		เพิ่มขึ้น 10		เพิ่มขึ้น 80		เพิ่มขึ้น 36		เพิ่มขึ้น 17		เพิ่มขึ้น 57		เพิ่มขึ้น 21	

*A คือ สร้างระเบียบบรรณานุกรมทั้งหมดภายในการสนทนาเดียว

**B คือ เป็การสนทนาใหม่ทุกครั้งทีสร้างระเบียบบรรณานุกรมของหนังสือแต่ละชื่อเรื่อง (เลือกใช้เป็นตัวเปรียบเทียบในตารางที่ 2 เนื่องจากได้คะแนนความถูกต้องและโครงสร้างที่สูงที่สุด)

*C = ป้อนข้อมูลโดยตรง

*** คะแนนเต็ม เท่ากับ 30 คะแนน ส่วนคะแนนเต็มของผลรวมทั้งหมด คือ 420 คะแนน

ตัวเลขในวงเล็บ คือ ร้อยละ

ผลการศึกษานี้ สอดคล้องกับผลการศึกษาของ Noruzi (2023) and Chow et al. (2024) ที่พบว่า หัวเรื่องที่ถูกกำหนดโดย ChatGPT มีความไม่แน่นอนและบางส่วนไม่ปรากฏอยู่ในหัวเรื่องหอสมุดรัฐสภาอเมริกัน นอกจากนี้ การกำหนดหัวเรื่องยังเป็นลักษณะหัวเรื่องทั่วไป ไม่มีความลึกซึ้งหรือเฉพาะเจาะจงได้เท่ากับที่บรรณารักษ์เป็นผู้กำหนด สาเหตุอาจเกิดจากสามกรณี คือ 1) ลักษณะคำสั่งของ MyGPT ที่ผู้เขียนใช้งานไม่สามารถทำให้ MyGPT เข้าถึงแหล่งข้อมูลหัวเรื่องที่เป็นเว็บไซต์ได้ อาจจะนำวิธีการเขียนโปรแกรมหรือเรียกข้อมูลผ่านโปรโตคอลต่าง ๆ มาใช้ประกอบ 2) ChatGPT ยังไม่มีความเข้าใจอย่างถ่องแท้เท่ามนุษย์ว่าหัวเรื่องแต่ละคำมีความหมายว่าอย่างไร และ 3) ข้อมูลในตัวแบบภาษาขนาดใหญ่ที่ใช้สอน ChatGPT ยังมีข้อมูลที่เป็นระเบียบบรรณานุกรมของห้องสมุดโดยเฉพาะข้อมูลที่เป็นภาษาไทยไม่มากเพียงพอที่จะทำให้ ChatGPT คาดเดาหรือระบุลักษณะของข้อมูลประเภทหัวเรื่องได้อย่างแม่นยำ

เขตข้อมูล 7XX ใช้บันทึกข้อมูลผู้แต่งและผู้รับผิดชอบเนื้อหาต่าง ๆ เพิ่มเติม พบข้อผิดพลาด เช่น การสะกดชื่อผิด การระบุชื่อผู้รับผิดชอบอื่น ๆ ที่ไม่ควรบันทึกในเขตข้อมูลนี้ เช่น สำนักพิมพ์ สาเหตุอาจเกิดจากลักษณะคำสั่งของ MyGPT ที่ผู้เขียนใช้ไม่ได้ระบุชัดเจนว่าควรนำชื่อและหน่วยงานผู้รับผิดชอบด้านเนื้อหาเพิ่มเติมลักษณะใดบ้างมาบันทึกไว้ในเขตข้อมูลนี้

คะแนนด้านโครงสร้างระเบียบบรรณานุกรม

เมื่อพิจารณาคะแนนด้านโครงสร้างระเบียบบรรณานุกรมในภาพรวมจากตารางที่ 1 พบว่า เขตข้อมูลส่วนใหญ่ได้รับคะแนนตั้งแต่ร้อยละ 50 เป็นต้นไป ดังนั้น ในส่วนนี้จึงขอเน้นเขตข้อมูลที่ได้รับคะแนนน้อยกว่าร้อยละ 50 ซึ่งได้แก่ เขตข้อมูล 264, 505 และ 650

พิจารณาจากเขตข้อมูลที่ได้คะแนนด้านโครงสร้างน้อยกว่าร้อยละ 50 มีลักษณะร่วมกัน คือ มีรายละเอียดโครงสร้างที่ค่อนข้างซับซ้อน หรือมีการเปลี่ยนแปลงและแปรผันตามทรัพยากรสารสนเทศที่เปลี่ยนแปลงไป ตัวอย่างเช่น เขตข้อมูล 264 เขตข้อมูลย่อย b ซึ่งเป็นชื่อสำนักพิมพ์มีกฎเกณฑ์ในการบันทึกข้อมูล คือ ไม่บันทึกคำว่าสำนักพิมพ์ บริษัท หรือคำอื่น ๆ นำหน้าชื่อสำนักพิมพ์ ส่วนปีพิมพ์ควรระบุไว้ในเขตข้อมูลย่อย c ทุกครั้ง หากมีปีพิมพ์มากกว่าหนึ่งปี ส่วนเขตข้อมูล 505 ซึ่งเป็นข้อมูลสารบัญ มีการกำหนดตัวบ่งชี้ตำแหน่งที่ 1 เพื่อบ่งชี้ว่าข้อมูลสารบัญที่บันทึกไว้นั้นเป็นข้อมูลสารบัญทั้งหมดหรือเป็นข้อมูลสารบัญเพียงบางส่วนจากที่ปรากฏในทรัพยากรสารสนเทศ และสุดท้าย คือ เขตข้อมูล 650 ที่ใช้ในการบันทึกข้อมูลหัวเรื่อง มีการใช้ตัวบ่งชี้ตำแหน่งที่ 2 แตกต่างไปตามแหล่งหรือคู่มือที่ใช้กำหนดหัวเรื่อง ส่วนเขตข้อมูลย่อยของหัวเรื่องย่อย (Subdivision) ก็มีการใช้งานที่หลากหลาย เช่น หากเป็นหัวเรื่องย่อยทั่วไป เขตข้อมูลย่อยที่ใช้ คือ เขตข้อมูลย่อย x (\$x) หากเป็นหัวเรื่องย่อยที่เป็นชื่อทางภูมิศาสตร์ จะใช้เขตข้อมูลย่อย z (\$z) หรือหัวเรื่องย่อยที่เป็นช่วงเวลา จะใช้เขตข้อมูลย่อย y (\$y) เป็นต้น

ผู้เขียนได้จัดเตรียมข้อมูลลักษณะตัวอย่างโครงสร้างของเขตข้อมูลข้างต้นแก่ MyGPT ในชุดคำสั่งที่ใช้สร้างและให้ข้อมูลอ้างอิงการกำหนดหัวเรื่องจากเว็บไซต์ข้อมูลเชื่อมโยงของหอสมุดรัฐสภาอเมริกัน (<https://id.loc.gov/>) แต่ดูเหมือนว่าผลลัพธ์ที่ได้ยังไม่เป็นไปตามคำสั่งทุกครั้ง อาจเป็นเพราะการสอนให้ MyGPT สามารถบันทึกเขตข้อมูลที่มีลักษณะซับซ้อนต้องใช้วิธีการเตรียมตัวอย่างของแต่ละกรณีเป็นจำนวนมากในลักษณะของตัวแบบภาษาขนาดใหญ่เพื่อให้ MyGPT เรียนรู้และคาดเดารูปแบบของข้อมูลได้ ซึ่งผลการศึกษาในส่วนนี้สอดคล้องกับข้อสังเกตของ Taniguchi (2024) ที่พบว่า ChatGPT มีปัญหาด้านการสร้างระเบียบบรรณานุกรมที่มีความซับซ้อน ซึ่งการจัดเตรียมข้อมูลที่มีคุณภาพสูงเพื่อใช้สอนปัญญาประดิษฐ์ให้สามารถจัดทำเมทาตาที่มีความถูกต้องจัดเป็นความท้าทายในการนำปัญญาประดิษฐ์มาใช้งาน

วิเคราะห์ทรัพยากรสารสนเทศอีกด้วย เนื่องจากห้องสมุดต้องดำเนินการเตรียมข้อมูลโดยทำความสะอาดข้อมูล (Data Cleaning) ซึ่งใช้เวลาและทรัพยากรอื่น ๆ เป็นจำนวนมาก (Mahmud, 2024)

เปรียบเทียบความแตกต่างของผลลัพธ์ที่ได้ระหว่างการทดสอบแบบ A และ B

เมื่อพิจารณาข้อมูลจากตารางที่ 1 พบว่า การสร้างระเบียบโดยเริ่มการสนทนากับ MyGPT ใหม่ทุกครั้ง (การทดสอบแบบ B) ทำให้ผลของการสร้างระเบียบบรรณานุกรมมีคะแนนเพิ่มขึ้น จำนวน 9 เขตข้อมูล จากทั้งหมด 14 เขตข้อมูล (ร้อยละ 64.29) ได้แก่ Leader, 020, 040, 336, 337, 338, เลขผู้แต่ง, 100, 505, 520, และ 7XX เนื่องจาก การสร้างระเบียบบรรณานุกรมของหนังสือทุกเล่มภายในการสนทนาเดียวจะเกิดการปะปนกันของข้อมูล ทำให้ MyGPT นำข้อมูลจากหนังสือเล่มหนึ่งมาปะปนกับหนังสือเล่มอื่น ๆ ทำให้ความถูกต้องของการใช้การสนทนาเดียว (การทดสอบแบบ A) น้อยกว่าแบบการเริ่มการสนทนาใหม่ทุกครั้ง (การทดสอบแบบ B)

อย่างไรก็ตาม หากพิจารณาคะแนนความถูกต้องและคะแนนด้านโครงสร้างระเบียบบรรณานุกรม ในภาพรวมจะเห็นว่าไม่จำเป็นที่จะเป็นการทดสอบด้วยวิธีใด คะแนนในภาพรวมยังอยู่ในช่วงร้อยละ 51 ถึง 54 ซึ่งอาจแปลความได้ว่า มีความผิดพลาดเกิดขึ้นประมาณครึ่งหนึ่ง ดังนั้น อาจกล่าวได้ว่า MyGPT ที่สร้างขึ้นนี้ยังไม่สามารถนำมาใช้ในการสร้าง ระเบียบบรรณานุกรมที่ถูกต้องได้ ซึ่งสอดคล้องกับการศึกษาของ Taniguchi (2024) ที่ทำการทดสอบและประเมินผล การสร้างระเบียบบรรณานุกรมในรูปแบบ MARC โดยใช้ ChatGPT และพบว่า มีข้อผิดพลาดเกิดขึ้นหลายจุด

อย่างไรก็ตามอาจนำ MyGPT มาใช้ช่วยในการสกัดและบันทึกเขตข้อมูลบางเขตข้อมูลได้ เนื่องจากเมื่อ พิจารณาคะแนนความถูกต้องแยกตามเขตข้อมูลโดยไม่ดูวิธีการทดสอบ และไม่รวมเขตข้อมูล 020, 040, 336, 337, 338 ซึ่งมี คะแนนความถูกต้องสูงอยู่แล้ว จะเห็นว่ามีหลายเขตข้อมูลที่มีคะแนนความถูกต้องอยู่ในเกณฑ์สูงจึงมีศักยภาพในการพัฒนา ให้ MyGPT สามารถสกัดและบันทึกข้อมูลได้ถูกต้องมากขึ้น เช่น เขตข้อมูล Leader, 082, 100, 245 และ 520 ซึ่งมีคะแนน ความถูกต้องตั้งแต่ร้อยละ 50 เป็นต้นไป

เปรียบเทียบความแตกต่างของผลลัพธ์ที่ได้ระหว่างสกัดข้อมูลจากรูปภาพกับการป้อนข้อมูลโดยตรง

ผลจากการสร้างระเบียบบรรณานุกรมด้วยวิธีการสกัดข้อมูลจากรูปภาพได้คะแนนรวมด้านความถูกต้อง และด้านโครงสร้างเพียงร้อยละ 51 ถึง 54 ผู้เขียนจึงได้ดำเนินการทดสอบเพิ่มเติมโดยใช้วิธีการสร้างระเบียบบรรณานุกรม จากการป้อนข้อมูลเกี่ยวกับหนังสือเข้าสู่ MyGPT โดยตรง ตามที่ได้กล่าวถึงในวิธีการดำเนินการว่า หากระดับความถูกต้อง ของระเบียบบรรณานุกรมที่สร้างโดย MyGPT ในด้านความแม่นยำและด้านความถูกต้องของโครงสร้างระเบียบบรรณานุกรม ไม่เป็นที่น่าพึงพอใจ จะดำเนินการทดสอบด้วยวิธีดังกล่าวเพิ่มเติม ซึ่งผลการทดสอบเป็นไปตามตารางที่ 2 และ 3 โดยได้ ทำการทดสอบเพิ่มเติมนี้ระหว่างวันที่ 19-20 ธันวาคม 2567 เวอร์ชัน ChatGPT ที่ใช้คือ ChatGPT Team ซึ่งใช้ตัวแบบ GPT-4o

จากข้อมูลในตารางที่ 2 พบว่า เมื่อเปรียบเทียบคะแนนจากการทดสอบการสร้างระเบียบบรรณานุกรม ด้วยวิธีการสกัดข้อมูลจากรูปภาพแบบเริ่มต้นการสนทนาใหม่ทุกครั้ง (การทดสอบแบบ B) ซึ่งมีคะแนนในภาพรวมสูงที่สุด ของวิธีการสกัดข้อมูลจากรูปภาพกับการป้อนข้อมูลหนังสือให้กับ MyGPT โดยตรง (การทดสอบแบบ C) คะแนนรวม ด้านความถูกต้องและด้านโครงสร้างเพิ่มขึ้นจาก ร้อยละ 54 เป็นร้อยละ 75 โดยทุกเขตข้อมูลมีคะแนนรวมเพิ่มขึ้น จึงอาจสรุป เบื้องต้นได้ว่า ณ ปัจจุบัน การสร้างระเบียบบรรณานุกรมของหนังสือภาษาไทยด้วยวิธีการสกัดข้อมูลจากรูปภาพยังได้ผลไม่ด้นัก ทั้งนี้ อาจเนื่องจากความสามารถของ ChatGPT ในการทำความเข้าใจภาษาไทยที่อยู่ในรูปภาพยังไม่ได้ถูกพัฒนา ให้เทียบเท่ากับการทำความเข้าใจภาษาอังกฤษ (OpenAI, 2024a) ทั้งนี้ อาจต้องทดสอบด้วยการพัฒนารูปแบบและวิธีการ

เตรียมข้อมูลหรือนำเข้าข้อมูลเพื่อให้ได้ผลลัพธ์ที่ดีขึ้น เช่น การปรับวิธีการเตรียมไฟล์รูปภาพ หรือไฟล์ดิจิทัลให้อยู่ในรูปแบบที่ MyGPT สามารถอ่านได้โดยการฝังตัวอักษรหรือเมตาดาตาที่บ่งชี้โครงสร้างของตัวอักษรในไฟล์ดิจิทัลว่าส่วนใดเป็นชื่อเรื่อง หัวข้อ เนื้อหา และมีวิธีการเรียงเนื้อหาหรือตัวอักษรอย่างไร และการใช้ไฟล์จากหนังสือทั้งเล่มแทนการใช้เฉพาะส่วนนำของหนังสือ

ตารางที่ 3 เปรียบเทียบคะแนนด้านความถูกต้องและด้านโครงสร้างระเบียบบรรณานุกรมที่สร้างโดย MyGPT ด้วยการทำซ้ำวิธีการป้อนข้อมูลเอง

เขตข้อมูล MARC	Leader			008			020			040			336, 337, 338			082			เลขผู้แต่ง		
คะแนนการทดสอบ C	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3
จำนวนครั้งที่ถูกต้องระดับ 2	0	9	0	2	0	0	10	10	10	10	10	10	10	10	10	3	5	5	0	0	0
จำนวนครั้งที่ถูกต้องระดับ 1	10	1	10	0	0	0	0	0	0	0	0	0	0	0	0	6	5	5	0	0	0
คะแนนรวมด้านความถูกต้อง	10	19	10	4	0	0	20	20	20	20	20	20	20	20	20	12	15	15	0	0	0
คะแนนด้านโครงสร้าง	10	10	10	10	10	10	10	10	10	10	10	10	10	10	10	10	10	10	10	9	10
คะแนนรวมทั้งหมด***	20	29	20	14	10	10	30	30	30	30	30	30	30	30	30	22	25	25	10	9	10
ร้อยละ	67	97	67	47	34	34	100	100	100	100	100	100	100	100	100	74	84	84	34	30	34

เขตข้อมูล MARC	100			245			264			505			520			650			7XX			ผลรวม		
คะแนนการทดสอบ C	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3
จำนวนครั้งที่ถูกต้องระดับ 2	8	7	8	10	10	10	0	0	0	10	10	10	10	10	10	1	0	1	9	8	0	83	89	74
จำนวนครั้งที่ถูกต้องระดับ 1	2	3	2	0	0	0	0	5	10	0	0	0	0	0	0	4	0	1	1	2	0	23	16	28
คะแนนรวมด้านความถูกต้อง	18	17	18	20	20	20	0	5	10	20	20	20	20	20	20	6	0	3	19	18	0	189	194	176
คะแนนด้านโครงสร้าง	10	10	10	3	3	4	10	5	0	10	10	10	10	10	10	1	0	9	10	10	10	124	117	123
คะแนนรวมทั้งหมด***	28	27	28	23	23	24	10	10	10	30	30	30	30	30	30	7	0	12	29	28	10	313	311	299
ร้อยละ	94	90	94	77	77	80	34	34	34	100	100	100	100	100	100	24	0	40	97	94	34	75	75	72

*** คะแนนเต็ม เท่ากับ 30 คะแนน ส่วนคะแนนเต็มของผลรวมทั้งหมด คือ 420 คะแนน

นอกจากนี้ เมื่อพิจารณาข้อมูลตามตารางที่ 3 ซึ่งเป็นการทำซ้ำวิธีการสร้างระเบียบบรรณานุกรมจากการป้อนข้อมูลเกี่ยวกับหนังสือเข้าสู่ MyGPT โดยตรง พบว่า แม้จะใช้ข้อมูลชุดเดียวกัน และใช้ MyGPT เดียวกัน แต่ดำเนินการต่างเวลากัน ผลของระเบียบบรรณานุกรมที่ได้นั้นมีความไม่คงที่ โดยตัวอย่างที่เห็นได้ชัด คือ เขตข้อมูล 650 และ 7XX ที่มีร้อยละของคะแนนรวมไม่เท่ากันในการทดสอบทั้งสามครั้ง ซึ่งอาจเป็นเพราะลักษณะของ ChatGPT คือ เป็นปัญญาประดิษฐ์แบบรู้สร้าง ซึ่งมีจุดเด่น คือ ความสามารถในการสร้างเนื้อหาใหม่ ๆ ผลลัพธ์จึงมีการเปลี่ยนแปลงไปได้ทุกครั้งที่มีการตอบกลับผู้ใช้งาน

ความผิดพลาดของระเบียบบรรณานุกรมที่พบ

จากข้อมูลในตารางที่ 4 และ 5 ซึ่งแสดงรายละเอียดข้อผิดพลาดที่พบจากระเบียบบรรณานุกรมที่สร้างโดย MyGPT ทั้งวิธีการสกัดข้อมูลจากรูปภาพและการป้อนข้อมูลหนังสือเข้าสู่ MyGPT โดยตรง สามารถจัดกลุ่มลักษณะของข้อมูลที่ผิดพลาดได้ตามตารางที่ 6 ซึ่งสรุปได้ว่า กลุ่มความผิดพลาดที่เกิดขึ้นมากที่สุด คือ ไม่เป็นไปตามมาตรฐาน MARC, RDA และการกำหนดเลขผู้แต่ง (วิธีอัปโหลดรูปภาพ (ร้อยละ 45.24), วิธีการป้อนข้อมูล (ร้อยละ 50.00)) นอกจากนี้ การใช้วิธีการป้อนข้อมูลหนังสือเข้าสู่ MyGPT โดยตรงยังทำให้จำนวนข้อผิดพลาดลดลงจาก 42 ข้อ เหลือ 14 ข้อ (ลดลงร้อยละ 66.67)

ข้อผิดพลาดที่เกิดขึ้นซ้ำจากทั้งสองวิธีการ ตามตารางที่ 4 และ 5 ได้แก่ ข้อผิดพลาดในเขตข้อมูล Leader, 008, เลขผู้แต่ง, 100, 245, 264, 650 และ 7XX ซึ่งโดยส่วนใหญ่เป็นเขตข้อมูลที่มีการแปรผันตามทรัพยากรสารสนเทศ ยกเว้น เขตข้อมูล Leader และ 008 ซึ่งลักษณะข้อมูลคงที่ มีจำนวนอักขระที่แน่นอน แต่การบันทึกข้อมูลให้แม่นยำตรงต้องดูลักษณะโดยรวมของทรัพยากรสารสนเทศ โดยเขตข้อมูลที่มีการทดสอบของผู้เขียนและการศึกษาของ Taniguchi (2024) มีความสอดคล้องกันว่ามีความถูกต้องในระดับต่ำ คือ เขตข้อมูล 100, 245, 264, และ 7XX นอกจากนี้ ยังอาจเป็นเพราะคำสั่งที่ใช้ในการอธิบายรายละเอียดการสร้างระเบียบบรรณานุกรมที่ผู้เขียนใช้ในการสร้าง MyGPT ยังมีความไม่รัดกุมหรือมีรายละเอียดไม่เพียงพอ จึงทำให้เกิดข้อผิดพลาดซ้ำในเขตข้อมูลดังกล่าว ผู้เขียนจึงดำเนินการปรับปรุงคำสั่งและรายละเอียดดังกล่าวใน MyGPT เพื่อทดสอบว่าสามารถลดข้อผิดพลาดที่เกิดขึ้นได้หรือไม่ ซึ่งจะกล่าวในส่วนต่อไป

การทดสอบการลดข้อผิดพลาดเบื้องต้นโดยการปรับปรุงชุดคำสั่งที่ใช้สร้าง MyGPT

ผู้เขียนได้ทดสอบการลดข้อผิดพลาดของการสร้างระเบียบบรรณานุกรมโดย MyGPT ดังนี้

1. ปรับปรุงคำสั่งที่ใช้อธิบาย MyGPT ในการสร้างระเบียบบรรณานุกรมให้มีความละเอียดเพิ่มมากขึ้น ในส่วนของ Leader, 008, เลขผู้แต่ง, 100, 245, 264, 650 และ 7XX เช่น กำหนดรูปแบบที่แน่นอนของ Leader การระบุตำแหน่งเขตข้อมูลต่าง ๆ ที่ใช้ในการลงรายการในเขตข้อมูล 008 การกำหนดวิธีการบันทึกข้อมูลผู้แต่งและเลือกเลขผู้แต่งชาวไทยอย่างละเอียด รูปแบบและลำดับการบันทึกข้อมูลในเขตข้อมูล 245 และ 264 อย่างละเอียด และการเลือกใช้หัวเรื่องและเขตข้อมูลย่อยของเขตข้อมูล 650 เป็นต้น

2. ทดสอบการแปลงตารางเลขผู้แต่งเป็นรูปแบบไฟล์ที่ MyGPT น่าจะสามารถทำความเข้าใจได้ โดยนำตารางเลขผู้แต่งบางส่วนมาปรับให้อยู่ในรูปแบบของตารางและบันทึกด้วยไฟล์ .xlsx

จากนั้นลองทดสอบการสร้างระเบียบบรรณานุกรมอีกครั้ง ด้วยวิธีการป้อนข้อมูลเข้าสู่ MyGPT โดยตรง พบว่า ยังพบข้อผิดพลาดจุดเดิม เช่น การบันทึกข้อมูลในเขตข้อมูล 008 ไม่สอดคล้องกับข้อมูลของหนังสือเล่มนั้น การเลือกใช้เขตข้อมูลย่อยในเขตข้อมูล 650 ไม่ถูกต้อง หรือในบางครั้งเขตข้อมูลถูกตัดออกถึงแม้ว่าจะกำหนดไว้ในคำสั่งแล้วว่าเป็นเขตข้อมูลที่ต้องปรากฏเมื่อสร้างระเบียบบรรณานุกรม อย่างไรก็ตาม การระบุคำสั่งโดยอธิบายอย่างละเอียดช่วยลดข้อผิดพลาดได้ในหลายกรณี เช่น การบันทึกข้อมูลชื่อผู้แต่งในเขตข้อมูล 100 และ 7XX หากเป็นชาวต่างชาติจะบันทึก

ข้อมูลโดยกลับนามสกุลขึ้นมาก่อนชื่อ หรือการบันทึกข้อมูลสำนักพิมพ์ ในเขตข้อมูล 264 มีการตัดคำว่า สำนักพิมพ์ บริษัท และจำกัด ออกก่อนการบันทึกข้อมูล เป็นต้น

ตารางที่ 4 ข้อผิดพลาดที่พบจากระเบียนบรรณานุกรมที่ MyGPT สร้างด้วยวิธีการอัปโหลดรูปภาพและความถี่ที่พบ*

Leader	008	020	040	336, 337, 338	082	เลขผู้แต่ง	100	245	264	505	520	650	7XX
ตำแหน่งที่ 9 ไม่บันทึกเป็น "a" เพื่อระบุว่ามีกรเข้ารหัสตัวอักษรแบบยูนิโคด (Unicode) (13)	บันทึกข้อมูลวันที่สร้างระเบียบไม่ถูกต้อง (12)	มีการบันทึกข้อมูล "Sqpaperback" ซึ่งเป็นข้อมูลที่บันทึกเพิ่มขึ้นมาออกเหนือจากข้อมูลในรูปภาพ (5)	มีการบันทึกเขตข้อมูลย่อย a, b และ c ซึ่งเป็นข้อมูลที่บันทึกเพิ่มขึ้นมาออกเหนือจากที่กำหนดไว้ (6)	ไม่มีข้อมูลปรากฏ (3)	ไม่มีข้อมูลปรากฏ (5)	ไม่มีข้อมูลปรากฏ (10)	สะกดชื่อผู้แต่งผิด (4)	สะกดชื่อผู้แต่งหรือชื่อผู้รับผิดชอบอื่นผิด (10)	มีการเพิ่มบันทึกที่เป็น ค.ศ. ขึ้นมา และได้บันทึกไว้ในเขตข้อมูลย่อย c (18)	ตัวบ่งชี้ตำแหน่งที่ 1 ไม่สอดคล้องกับข้อมูลสารบัญที่บันทึกไว้ (10)	แปลเป็นภาษาอังกฤษ (4)	หัวเรื่องไม่ปรากฏอยู่ในคู่มือ (15)	มีการบันทึกข้อมูลลงในเขตข้อมูล 700 หรือ 710 ทั้งที่หนังสือรายการนั้นไม่มีข้อมูลผู้รับผิดชอบเพิ่มเติม (14)
ตำแหน่งที่ 18 ไม่บันทึกเป็น "i" เพื่อระบุว่าจะเรียงรายการนั้น ๆ ใช้การลงรายการตามมาตรฐาน RDA (3)	บันทึกฉบับที่บันทึกเป็นปี พ.ศ. แทนที่ ค.ศ. (6)	มีเครื่องหมาย "-" คั่นระหว่างตัวเลข (1)	การบันทึกคำในเขตข้อมูลย่อย e ใช้ตัวอักษรใหญ่ทั้งหมด (1)		เลขผู้แต่งไม่เป็นไปตามคู่มือ (7)	ตัวบ่งชี้ตำแหน่งที่ 1 ไม่สอดคล้องกับการเรียงลำดับชื่อ-นามสกุลที่บันทึกในเขตข้อมูล (4)	ตัวบ่งชี้ตำแหน่งที่ 1 ไม่สอดคล้องกับการเรียงลำดับชื่อ-นามสกุลที่บันทึกในเขตข้อมูล (2)	ระบุสถานที่พิมพ์ผิดพลาด (5)	ไม่มีข้อมูลปรากฏ (8)	ไม่มีข้อมูลปรากฏ (1)	ตัวบ่งชี้ตำแหน่งที่ 2 ไม่เป็น "4" ในรายการที่ใช้หัวเรื่องภาษาไทย (11)	สะกดชื่อผู้รับผิดชอบเพิ่มเติมผิด (7)	
ตำแหน่งที่ 39 ไม่บันทึกเป็น "d" เพื่อแสดงว่าแหล่งของที่มาของระเบียบไม่ใช่หอสมุดแห่งชาติหรือความร่วมมือด้านการลงรายการ (6)					ตัวอักษรแรกไม่ได้มาจากชื่อผู้แต่ง (3)	ในเขตข้อมูลย่อย e ใช้คำว่า "ผู้เขียน" และ "author" แทน "ผู้แต่ง" (3)	ตัวบ่งชี้ตำแหน่งที่ 1 ไม่สอดคล้องกับการมีผู้แต่งหลักในเขตข้อมูล 100 (2)	การสะกดชื่อสำนักพิมพ์ผิดพลาด (4)	สก็ตข้อมูลออกมาไม่ถูกต้อง เช่น ขีอบผิด หรือเกิน มีการสะกดคำผิด (5)		ใช้หัวเรื่องย่อยไม่ถูกต้อง ตามคู่มือ (5)	ผู้รับผิดชอบเป็นชาวไทยแต่มีการบันทึกโดยนำนามสกุลมาขึ้นต้นก่อนแล้วตามด้วยชื่อ (4)	
มีการบันทึกข้อมูลกลุ่มผู้ใช้งานทรัพยากรสารสนเทศเป็น "g" ซึ่งเป็นข้อมูลที่บันทึกเพิ่มขึ้นมาออกเหนือจากข้อมูลในรูปภาพ (4)						ก่อนเขตข้อมูลย่อย e ไม่ปรากฏเครื่องหมาย ", " นำหน้า (3)		บันทึกชื่อสำนักพิมพ์โดยมีคำนำหน้า เช่น คำว่า บริษัท (3)	แปลสารบัญเป็นภาษาอังกฤษ (1)			ตัวบ่งชี้ตำแหน่งที่ 1 ไม่สอดคล้องกับชื่อผู้รับผิดชอบที่เป็นชาวไทย (2)	
							มีการบันทึกเขตข้อมูล 100 มากกว่า 1 ครั้งในระเบียน (2)		บันทึกปีที่พิมพ์ผิดพลาด (2)				
							แปลงชื่อผู้แต่งเป็นภาษาอังกฤษ (1)						

* ตัวเลขในวงเล็บ คือ ความถี่ที่พบ

ตารางที่ 5 ข้อผิดพลาดที่พบจากระเบียนบรรณานุกรมที่ MyGPT สร้างด้วยวิธีการป้อนข้อมูลเองและความถี่ที่พบ*

Leader	008	020	040	336, 337, 338	082	เลขผู้แต่ง	100	245	264	505	520	650	7XX
ตำแหน่งที่ 9 ไม่บันทึกเป็น "a" เพื่อระบุว่ามีหรือไม่มีตัวอักษรแบบยูนิโคด (Unicode) (11)	การลงข้อมูลไม่สอดคล้องกับเขตข้อมูลอื่น เช่น เขตข้อมูล 264 หรือ 300 เขตข้อมูลย่อย b (27)	-	-	-	-	เลขผู้แต่งไม่เป็นไปตามคู่มือ (20)	MyGPT ไม่รู้ว่าชื่อนั้นต้องกลับเอานามสกุลขึ้นก่อนหรือไม่ (5)	ไม่ใช่เครื่องหมาย = เพื่อแสดงชื่อเรื่องเทียบเคียงหรือชื่อเรื่องในภาษาอื่น (13)	มีการเพิ่มปีพิมพ์ที่เป็นปี ค.ศ. ขึ้นมา และไม่ได้บันทึกไว้ในเขตข้อมูลย่อย c (20)	-	-	ไม่มีข้อมูลปรากฏ (10)	มีการบันทึกข้อมูลลงในเขตข้อมูล 700 หรือ 710 ทั้งที่หนังสือรายการนั้นไม่มีข้อมูลผู้รับผิดชอบเพิ่มเติม (10)
						ตัวอักษรแรกไม่ได้มาจากชื่อผู้แต่ง (1)	แปลงชื่อผู้แต่งเป็นภาษาอังกฤษ (4)	ใช้เครื่องหมายคั่นระหว่างผู้แต่งหรือผู้รับผิดชอบที่มีมากกว่า 1 คน ผิดพลาด (1)	บันทึกชื่อสำนักพิมพ์โดยมีคำนำหน้า เช่น คำว่า บริษัท (12)			หัวเรื่องไม่ปรากฏอยู่ในคู่มือ (8)	MyGPT ไม่รู้ว่าชื่อนั้นต้องกลับเอานามสกุลขึ้นก่อนหรือไม่ (1)
												ใช้หัวเรื่องย่อย ไม่ถูกต้องตามคู่มือ (3)	แปลงชื่อผู้รับผิดชอบเพิ่มเติมเป็นภาษาอังกฤษ (1)

* ตัวเลขในวงเล็บ คือ ความถี่ที่พบ

ตารางที่ 6 การจัดกลุ่มความผิดพลาดที่พบจากระเบียนบรรณานุกรมที่ MyGPT สร้าง เปรียบเทียบจากวิธีที่ใช้สร้างข้อมูล

กลุ่มความผิดพลาด	วิธีอัปโหลดรูปภาพ		วิธีป้อนข้อมูลโดยตรง*	
	จำนวนข้อผิดพลาด	ร้อยละ	จำนวนข้อผิดพลาด	ร้อยละ
ไม่เป็นไปตามมาตรฐาน MARC, RDA และการกำหนดเลขผู้แต่ง	19	45.24	7	50.00
สกัดข้อมูลจากรูปภาพผิดพลาดหรือสะกดชื่อหรือคำผิด	10	23.81	2	14.29
เพิ่มเติมข้อมูลที่ไม่ปรากฏในรูปภาพหรือข้อมูลที่ป้อนสู่ MyGPT	8	19.05	4	28.57
ข้อมูลไม่ครบถ้วน	5	11.90	1	7.14
รวม	42	100.00	14	100.00

*เพื่อให้สามารถเปรียบเทียบค่ากับวิธีการอัปโหลดรูปภาพ ซึ่งมีการทดลอง 2 ครั้งได้ จึงคำนวณจากการสร้างระเบียนบรรณานุกรมโดยการป้อนข้อมูลให้ MyGPT โดยตรง

ในครั้งที่ 2 และ 3 นอกจากนี้ยังเป็นครั้งที่มียข้อผิดพลาดเกิดขึ้นมากด้วย จึงเป็นข้อมูลที่เหมาะสมเพื่อนำมาเปรียบเทียบกัน

ในส่วนนี้ อาจสรุปเกี่ยวกับวิธีที่เหมาะสมในการสร้างระเบียบบรรณานุกรมด้วย MyGPT ได้ คือ 1) คำสั่งที่ใช้ MyGPT ของผู้เขียนอาจมีปริมาณคำสั่งที่มากหรืออาจยังมีโครงสร้างไม่เหมาะสม จึงทำให้การสร้างระเบียบบรรณานุกรมมีข้อผิดพลาด 2) แม้ว่าจะระบุคำสั่งและคำอธิบายอย่างชัดเจน แต่หลายครั้ง MyGPT ไม่ปฏิบัติตามคำสั่งและคำอธิบาย ซึ่งวิธีการแก้ไขข้อผิดพลาด คือ เมื่อ MyGPT สร้างระเบียบบรรณานุกรมแล้ว หากผู้ใช้งานเห็นว่ามีส่วนใดที่ขาดหายไป หรือมีส่วนที่ไม่ถูกต้อง ผู้ใช้งานต้องใช้ข้อความพร้อมรับแจ้ง MyGPT ให้ทราบถึงจุดที่ต้องแก้ไข ผลลัพธ์ที่ได้จะมีความถูกต้องสมบูรณ์มากขึ้น แม้ว่าในปัจจุบันการสร้างระเบียบบรรณานุกรมโดย MyGPT หรือ ChatGPT ยังมีความผิดพลาดและไม่สมบูรณ์แต่ในบางส่วนก็ทำได้ถูกต้อง จึงอาจกล่าวได้ว่าปัญญาประดิษฐ์สามารถใช้เป็นจุดเริ่มต้นหรือเป็นผู้ช่วยในการสร้างระเบียบบรรณานุกรมได้ แต่ยังคงจำเป็นต้องมีการตรวจทาน แก้ไข และพัฒนาผลลัพธ์โดยผู้ที่มีความรู้ทางด้านการจัดทำเมทาตาตามมาตรฐานห้องสมุด ดังนั้น การทำงานร่วมกันระหว่างปัญญาประดิษฐ์กับมนุษย์จึงเป็นเรื่องที่สำคัญ ซึ่งจะทำให้ผลลัพธ์ที่ได้จากปัญญาประดิษฐ์มีความถูกต้องสมบูรณ์มากยิ่งขึ้น ซึ่งสอดคล้องกับบทความของ York et al. (2024) ที่กล่าวว่า การตรวจทานผลจากผู้เชี่ยวชาญด้านการลงรายการเป็นปัจจัยที่สำคัญมากในการพัฒนาผลลัพธ์ของการใช้งานปัญญาประดิษฐ์ นอกจากนี้ยังพบว่า ผลการศึกษาของผู้เขียนไม่สอดคล้องกับการศึกษาของ Brzustowicz (2023) ที่ทดสอบการสร้างระเบียบบรรณานุกรมด้วย ChatGPT และพบว่า ระเบียบที่ได้มีความถูกต้องเป็นไปตามมาตรฐานการจัดทำเมทาตาฯ ทั้งนี้ อาจเนื่องจากการสร้างระเบียบในงานวิจัยฉบับดังกล่าวใช้ทรัพยากรสารสนเทศที่มีเนื้อหาเป็นภาษาอังกฤษเป็นหลักและส่วนใหญ่มีปีพิมพ์ค่อนข้างเก่า ซึ่งมีโอกาสที่ข้อมูลของทรัพยากรสารสนเทศเหล่านั้นจะรวมถูกรวมอยู่ในตัวแบบภาษาขนาดใหญ่ที่ใช้ในการสอน ChatGPT ทำให้การสร้างระเบียบมีความถูกต้องแม่นยำมากกว่าหนังสือภาษาไทย และ 3) การเลือกเลขผู้แต่งภาษาไทยของ MyGPT ยังไม่ถูกต้อง แม้ว่าจะจัดเตรียมตารางเลขผู้แต่งในรูปแบบไฟล์ที่ MyGPT สามารถเข้าถึงและทำความเข้าใจได้ โดยปัญหาที่พบ คือ MyGPT ไม่สามารถเรียงลำดับสระและวรรณยุกต์ภาษาไทยได้ถูกต้อง จึงไม่สามารถเลือกเลขผู้แต่งที่ถูกต้องมาใช้งานได้

บทสรุป (Conclusion)

อาจสรุปข้อค้นพบจากการศึกษาครั้งนี้ได้ ดังนี้ 1) MyGPT ในการศึกษาครั้งนี้ยังไม่สามารถจัดทำระเบียบบรรณานุกรมหนังสือภาษาไทยได้ถูกต้องในระดับที่น่าไปใช้งานได้ทันที โดยเฉพาะหากนำเข้าสู่ข้อมูลด้วยรูปภาพส่วนนำของหนังสือ อาจเป็นเพราะตัวแบบของ ChatGPT สามารถสกัดและทำความเข้าใจข้อมูลที่เป็นภาษาอังกฤษได้ดีกว่าภาษาอื่น ๆ อย่างไรก็ตาม มีหลายส่วนที่ MyGPT สร้างข้อมูลที่มีความถูกต้องในระดับสูงได้ เช่น เขตข้อมูล Leader, 020, 040, 082, 100, 245, 336, 337, 338 และ 520 จึงอาจนำ MyGPT มาช่วยในการสกัดข้อมูลจากรูปภาพส่วนนำของหนังสือ หรือการป้อนข้อมูลหนังสือให้ MyGPT โดยตรงได้ จึงอาจกล่าวได้ว่า สามารถนำ MyGPT มาใช้เป็นจุดเริ่มต้นของการจัดทำระเบียบบรรณานุกรมหรือใช้เป็นผู้ช่วยในเบื้องต้นได้ 2) MyGPT มีศักยภาพในการพัฒนาต่อได้ ดังจะเห็นได้จากการเพิ่มขึ้นของคะแนนความถูกต้องของระเบียบบรรณานุกรมเมื่อมีการเปลี่ยนการนำเข้าข้อมูลด้วยรูปภาพเป็นการป้อนข้อมูลโดยตรง และการปรับปรุงชุดคำสั่งที่ใช้สร้าง MyGPT ให้รัดกุมและชัดเจนขึ้นเพื่อลดข้อผิดพลาด ดังนั้น หากปรับปรุงวิธีการเตรียมข้อมูล เช่น การเตรียมรูปภาพหรือไฟล์ดิจิทัลของหนังสือ หรือคู่มือและข้อมูลประกอบต่าง ๆ เช่น ข้อมูลเลขผู้แต่ง คู่มือการกำหนดหัวเรื่องและเลขหมู่ให้อยู่ในลักษณะที่ MyGPT สามารถอ่าน เข้าถึง และทำความเข้าใจโครงสร้างของเนื้อหาของหนังสือและคู่มือต่าง ๆ ได้น่าจะช่วยให้ผลลัพธ์ในการสร้างระเบียบบรรณานุกรมของ MyGPT พัฒนาขึ้นได้อย่างต่อเนื่อง โดยในการพัฒนานี้ต้องอาศัยการทำงานร่วมกันระหว่างมนุษย์หรือผู้เชี่ยวชาญด้านการลงรายการบรรณานุกรมกับปัญญาประดิษฐ์ (Human in the Loop) โดยผู้เชี่ยวชาญจะเป็นผู้ควบคุมคุณภาพและแก้ไข รวมถึงสอนปัญญาประดิษฐ์เมื่อพบข้อผิดพลาดเพื่อให้เกิดการพัฒนาของผลลัพธ์ 3) ด้านความถูกต้องและข้อผิดพลาดที่พบ พบว่า เขตข้อมูลที่มีค่าคงที่หรือมีรูปแบบการลงรายการไม่ซับซ้อน

มีแนวโน้มที่จะมีคะแนนความถูกต้องในระดับสูงกว่าเขตข้อมูลที่มีความผันแปรไปตามทรัพยากรสารสนเทศ 4) ผลลัพธ์จากการสร้างระเบียบบรรณานุกรมด้วย MyGPT มีการผันเปลี่ยนไป หากดำเนินการในช่วงเวลาที่ต่างกัน แม้ว่าจะใช้ชุดข้อมูลและชุดคำสั่งเดียวกันก็ตาม

ข้อจำกัดของการศึกษานี้ คือ 1) การใช้หนังสือภาษาไทยเป็นตัวอย่างในการสร้างระเบียบบรรณานุกรมเพียง 10 ชื่อเรื่อง เนื่องจากวิธีการที่ใช้สร้างระเบียบในช่วงแรกเป็นการอัปโหลดรูปภาพส่วนประกอบของหนังสือเข้าสู่ MyGPT ซึ่งมีข้อจำกัดในการใช้งาน คือ สามารถอัปโหลดได้เพียงครั้งละ 10 รูป และเมื่อถึงขีดจำกัดในการใช้งานต้องรอให้พ้นช่วงดังกล่าวจึงจะอัปโหลดรูปต่อไปได้ โดยช่วงเวลาสูงสุดที่ต้องรอหลังถึงขีดจำกัดใช้งานที่ผู้เขียนพบ คือ 4 ชั่วโมง ในการสร้างระเบียบบรรณานุกรมหนังสือ 10 เรื่อง จึงประสบปัญหาเรื่องระยะเวลาที่มีจำกัด 2) ใช้เพียงหนังสือภาษาไทย เนื่องจากในการปฏิบัติงานจริงนั้น หากเป็นหนังสือภาษาอังกฤษผู้ปฏิบัติงานจะสามารถค้นหาข้อมูลระเบียบบรรณานุกรมที่มีห้องสมุดอื่นจัดทำไว้แล้วได้ เช่น จากฐานข้อมูลหอสมุดรัฐสภาอเมริกัน และ WorldCat แต่หากเป็นหนังสือภาษาไทย การหาข้อมูลลักษณะดังกล่าวกระทำได้ยากกว่า นอกจากนี้ เริ่มมีผลการทดสอบการสร้างระเบียบบรรณานุกรมหนังสือภาษาอังกฤษด้วยปัญญาประดิษฐ์ออกมาบ้างแล้ว เช่น งานวิจัยของ Brzustowicz (2023), Chow et al. (2024), Taniguchi (2024), and York et al. (2024) แต่การสร้างระเบียบบรรณานุกรมหนังสือภาษาอื่น ๆ ยังมีจำกัด ผู้เขียนจึงเลือกใช้เฉพาะหนังสือภาษาไทยเพื่อเป็นการเติมเต็มช่องว่างของวรรณกรรมที่เกี่ยวข้อง 3) เนื่องจากยังไม่มีแนวทางการปฏิบัติที่ดีที่สุดในการสร้าง MyGPT คำสั่งที่ใช้ในการศึกษานี้เกิดจากประสบการณ์การลองผิดลองถูกของผู้เขียน ซึ่งอาจไม่ใช่วิธีการที่ดีที่สุด การศึกษานี้จึงเน้นการบอกเล่าประสบการณ์ที่พบและความคิดเห็นของผู้เขียนในการทดลองสร้าง MyGPT สำหรับใช้ในการสร้างระเบียบบรรณานุกรม

ข้อเสนอแนะสำหรับการศึกษาเพิ่มเติมต่อไป คือ 1) ทดลองปรับเปลี่ยนไฟล์รูปภาพ หรือไฟล์ดิจิทัลของหนังสือที่จะนำเข้าสู่ ChatGPT ให้มีลักษณะที่ ChatGPT เข้าถึงและทำความเข้าใจได้ เช่น การฝังตัวอักษร หรือมีเมทาตาทาบอกลักษณะโครงสร้างเนื้อหาที่อยู่ในไฟล์ดิจิทัล จากนั้น ทดสอบความสามารถในการสร้างรายการบรรณานุกรม 2) ทดลองการสร้าง MyGPT ด้วยชุดคำสั่งรูปแบบอื่น ๆ 3) ทดลองการสร้างระเบียบบรรณานุกรมโดยใช้ปัญญาประดิษฐ์อื่น ๆ นอกเหนือจาก ChatGPT หรือใช้ตัวแบบของ ChatGPT อื่น ๆ นอกเหนือจาก GPT-4o

ข้อมูลเพิ่มเติม (Availability of Shared Supplementary Data)

ผู้ที่สนใจนำคำสั่งที่ใช้สร้าง MyGPT ในบทความนี้ไปทดลองสร้าง MyGPT หรือนำไปพัฒนาต่อ รวมถึงต้องการศึกษาเพิ่มเติมในส่วนของตัวอย่างผลลัพธ์ที่ได้จากการสร้างระเบียบบรรณานุกรมจากการศึกษานี้ สามารถศึกษาได้จากชุดข้อมูลที่คุณเขียนได้นำฝากไว้ในคลังเก็บข้อมูลวิจัย Harvard Dataverse ชื่อ Replication data for: The instructions and files used for creating MyGPT for cataloging works in an AI pilot program of the Office of Academic Resources, Chulalongkorn University โดย Paphatsurichote (2024) ดังปรากฏในรายการอ้างอิงท้ายบทความ

รายการอ้างอิง (References)

- Brzustowicz, R. (2023). From ChatGPT to CatGPT: The implications of artificial intelligence on library cataloging. *Information Technology and Libraries*, 42(3), 1-22.
<https://doi.org/10.5860/ital.v42i3.16295>
- Chan, C. K. Y., & Colloton, T. (2024). *Generative AI in higher education: The ChatGPT effect*. Routledge.
<https://doi.org/10.4324/9781003459026>
- Chow, E. H. C., Kao, T. J., & Li, X. (2024). An experiment with the use of ChatGPT for LCSH subject assignment on electronic theses and dissertations. *Cataloging & Classification Quarterly*, 62(5), 574–588. <https://doi.org/10.1080/01639374.2024.2394516>

- Mahmud, M. R. (2024, July 30). AI in automating library cataloging and classification. *Library Hi Tech News*. <https://doi.org/10.1108/LHTN-07-2024-0114>
- Mannheimer, S., Bond, N., Young, S. W. H., Kettler, H. S., Marcus, A., Slipper, S. K., Clark, J. A., Shorish, Y., Rossmann, D., & Sheehey, B. (2024). Responsible AI practice in libraries and archives: A review of the literature. *Information Technology and Libraries*, 43(3), 1-29. <https://doi.org/10.5860/ital.v43i3.17245>
- Noruzi, A. (2023). The use of artificial intelligence in knowledge organization and subject indexing. *Informology*, 3(1), 1-8. <https://informology.org/2024/v3n1/editorial5.pdf>
- OpenAI. (2022, November 30). *Introducing ChatGPT*. <https://openai.com/index/chatgpt/>
- OpenAI. (2024a, November). *Image inputs for ChatGPT-FAQ*. <https://help.openai.com/en/articles/8400551-image-inputs-for-chatgpt-faq>
- OpenAI. (2024b, November). *Creating a GPT*. <https://help.openai.com/en/articles/8554397-creating-a-gpt>
- Paphatsurichote, R., & Kaewhawong, A. (2024). *Replication data for: The instructions and files used for creating MyGPT for cataloging works in an AI pilot program of the Office of Academic Resources, Chulalongkorn University (Version 3) [Data set]*. Harvard Dataverse. <https://doi.org/10.7910/DVN/TZ69WH>
- Taniguchi, S. (2024). Creating and evaluating MARC 21 bibliographic records using ChatGPT. *Cataloging & Classification Quarterly*, 62(5), 527–546. <https://doi.org/10.1080/01639374.2024.2394513>
- York, E., Hanegbi, D., & Ganor, T. (2024). Enriching bibliographic records using AI – a pilot by Ex Libris. *Internet Reference Services Quarterly*, 28(3), 287–291. <https://doi.org/10.1080/10875301.2024.2361871>