



แนวทางการสร้างสมดุลระหว่างเสรีภาพในการพูดกับมาตรฐานความน่าเชื่อถือและมาตรฐาน
ความปลอดภัยในโลกออนไลน์ด้วยการกลั่นกรองเนื้อหาที่มีประสิทธิภาพ

Seeking balance between free speech and online trust and safety
through an effective online moderation

โดย
บริษัท เนโอ โมเมนตัม จำกัด

ธันวาคม 2567

รายงานฉบับสมบูรณ์

แนวทางการสร้างสมดุลระหว่างเสรีภาพในการพูดกับมาตรฐานความน่าเชื่อถือและมาตรฐาน
ความปลอดภัยในโลกออนไลน์ด้วยการกลั่นกรองเนื้อหาที่มีประสิทธิภาพ

Seeking balance between free speech and online trust and safety
through an effective online moderation

โดย

นางสาวกุลชาดา ชัยพิพัฒน์
นายธนภณ เรามานะชัย
นางสาวกรรณกาญจน์ การเนตร
นายอภิเดช เตปิน
นายศุภฤกษ์ เพ็ชรพล

ธันวาคม 2567

Acknowledgement and Disclaimer

โครงการแนวทางการสร้างสมดุระหว่างเสรีภาพในการพูดและมาตรฐานความปลอดภัยในโลกออนไลน์ ด้วยการกลั่นกรองเนื้อหาที่มีประสิทธิภาพ ทีมวิจัยขอแสดงความขอบคุณเป็นอย่างยิ่งต่อโครงการติดตามและวิเคราะห์กระบวนการเลือกตั้งในโซเชียลมีเดีย (DEAL) และภาคีเครือข่าย ที่ให้การสนับสนุนชุดข้อมูลช่วงเลือกตั้ง เพื่อวิเคราะห์เปรียบเทียบกับช่วงหลังเลือกตั้ง ทำให้เห็นพัฒนาการการใช้โซเชียลมีเดียในประเด็นที่เกี่ยวข้องกับการเมืองของประเทศไทย ซึ่งเป็นประโยชน์อย่างยิ่งต่อการดำเนินงานของโครงการนี้

อย่างไรก็ดี เนื่องจากข้อจำกัดทั้งด้านทรัพยากร เครื่องมือ และบุคลากร จึงทำให้งานศึกษาครั้งนี้มีข้อจำกัดหลายประการ ดังนั้นเพื่อให้เกิดความเข้าใจที่ตรงกันของผู้อ่านและป้องกันการตีความที่อาจจะคลาดเคลื่อนไปจากวัตถุประสงค์การศึกษาและโครงสร้างของข้อมูล โครงการฯ จึงได้สรุปข้อจำกัดของงานศึกษาไว้ดังนี้

1. ความแตกต่างของวิธีการเก็บข้อมูลระหว่างสองช่วงเวลา

ในช่วงการเลือกตั้ง การเข้าถึงข้อมูลดำเนินการผ่าน API ของแพลตฟอร์มอย่างเป็นทางการ เช่น Twitter (X) และ Facebook (Crowdtangle) ทำให้ได้ข้อมูลที่เป็นโฆษณา (Ads) ปะปนน้อยหรือไม่มีเลย อย่างไรก็ตาม หลังการเลือกตั้ง การยกเลิกการใช้งาน API ของแพลตฟอร์ม ทำให้ต้องพึ่งพา Third-Party Platform ในการเก็บข้อมูล ส่งผลให้มีข้อมูลที่เป็น Ads และ ข้อมูลที่ไม่เกี่ยวข้อง (Noise) เข้ามาปะปนในชุดข้อมูล นอกจากนี้ ยังมีข้อจำกัดของทรัพยากรทำให้โครงการต้องจำกัดปริมาณข้อความที่เก็บให้ไม่เกิน 1.2 ล้านข้อความ

2. ข้อจำกัดของ Third-Party Platform ในการเก็บรวบรวมข้อมูล

โครงสร้างของข้อมูลที่เก็บรวบรวมจาก Third-Party Platform แตกต่างจากการใช้ API หรือ Crowdtangle โดยเฉพาะการเก็บทุกอย่างที่ปรากฏบนหน้าโปรไฟล์ ส่งผลให้มี Noise และ Ads เป็นจำนวนมาก รวมถึงข้อความภาษาอังกฤษที่มาจากบัญชีที่โพสต์เพียงครั้งเดียว นอกจากนี้ การกรองข้อมูลต้องพึ่งพาเกณฑ์การโพสต์มากกว่า 2 ครั้งขึ้นไป ซึ่งอาจมีผลต่อความครอบคลุมของข้อมูลที่แท้จริง นอกจากนั้นต้องกรองบัญชีที่เกี่ยวข้องกับข่าวหรือการเมืองโดยใช้ Amplification Scores และบัญชีที่ผู้เชี่ยวชาญระบุไว้ล่วงหน้า

3. ความถูกต้องและแม่นยำของ GPT Model ในการป้ายประเภทข้อมูล (Labeling)

ความแม่นยำในการ Label ข้อมูลโดยใช้ GPT Model มีความสม่ำเสมอมากกว่าการ Label โดยมนุษย์ เนื่องจากโมเดลทำงานตามเกณฑ์ที่กำหนดใน Prompt และตัวอย่างที่ได้รับ อย่างไรก็ตาม โมเดลอาจมีข้อผิดพลาดในการ Label ข้อมูลที่มีบริบทซับซ้อน เช่น การแยก Hate Speech ตามระดับ ซึ่งต้องอาศัยความเข้าใจที่ลึกซึ้งและการตีความบริบท นอกจากนี้ การ Label โดย GPT ในครั้งนี้เป็นการใช้ Prompt โดยตรง ไม่ใช้การ Train หรือ Fine-Tune โมเดล ทำให้มีความแม่นยำที่ยังสามารถพัฒนาได้ ในทางกลับกัน การ Label ข้อมูลโดยมนุษย์สามารถเข้าใจบริบทที่ซับซ้อนได้ดีกว่า แต่มีแนวโน้มเกิดความไม่สม่ำเสมอและ Bias ได้ง่าย รวมถึงใช้เวลามากกว่าการใช้โมเดล AI ในกระบวนการนี้

แม้จะมีข้อจำกัดดังกล่าว การศึกษานี้ยังคงเป็นความพยายามที่จะนำเสนอข้อมูลเชิงลึกเกี่ยวกับพลวัตของการสื่อสารออนไลน์ในบริบทการเมือง โดยเฉพาะอย่างยิ่งในช่วงที่ประเทศไทยเผชิญความท้าทายเกี่ยวกับเสรีภาพในการแสดงออกและความปลอดภัยในพื้นที่ดิจิทัล การระบุข้อจำกัดอย่างโปร่งใสจึงเป็นส่วนหนึ่งของความตั้งใจในการพัฒนาระบบการศึกษาในอนาคตให้มีประสิทธิภาพและครอบคลุมมากยิ่งขึ้น

Executive Summary

A surge of harmful content in Thailand's post-election period poses a threat to democratic processes and values, potentially exacerbating societal polarization in the fragile political transition.

This research project, undertaken by Neo Momentum in partnership with the Digital Election Analytic Lab (DEAL) and civil society organizations, delves into the implementation of community standards by online platforms. It focuses on online trust, safety, internet governance policies, and problematic online discourse during Thailand's 2023 general elections and the subsequent period.

The research aimed to address two primary objectives:

- First, it sought to identify gaps in the implementation of community standards by major online platforms, such as X (formerly Twitter), Meta (Facebook), and Google (YouTube). The study emphasized evaluating the impact of these standards on online trust and safety, as well as their alignment with broader internet governance policies.
- Second, the research explored actionable and inclusive strategies for online moderation and intervention. These approaches aim to promote fact-based, non-violent political communication while safeguarding the fundamental right to freedom of expression.

The research methodology employed a mixed-methods approach, incorporating both qualitative and quantitative methods. This included thematic analysis of malinformation, hate speech, and AI Deepfakes, alongside user behavior analysis using Social Network Analysis. The research also involved a desk review of community standards for major platforms that dominate Thailand's digital landscape.

Data collection encompassed both the pre-election period (April 1 - June 15, 2023) and the post-election period (January 1 - August 31, 2024). Data were gathered using licensed third-party platforms and leveraged techniques such as Human-AI content labeling, topic modeling, and behavioral analysis. Furthermore, the research included valuable insights from stakeholder consultations. This practice should serve as a *modus operandi* in processing massive datasets with speed to ensure lesser bias and higher accuracy and social relevance in data collection and analysis.

The research uncovered significant challenges following.

- **Community standards enforcement varied across platforms, with measures often insufficient to curb harmful content:** Platforms like Facebook and YouTube used labeling and fact-checking collaborations to tackle misinformation. However, these mechanisms lacked effectiveness, allowing harmful content to spread rapidly, especially during sensitive political events. X's policy of algorithmic adjustments instead of content removal further exposed inconsistencies in handling problematic discourse.
- **Delays and inconsistencies in content review processes allowed hate speech and malinformation to thrive during critical periods:** Slow response times in reviewing flagged content and uneven application of standards led to the unchecked proliferation of harmful narratives. This gap was particularly pronounced during Thailand's 2023 general elections, where such content influenced public perception and deepened societal divisions.
- **The evolution of harmful content dynamics highlighted a shift from election-focused narratives to more targeted attacks in the post-election period:** The study observed an



increase in personal attacks, reputational damage, and conspiracy theories aimed at discrediting political figures and institutions. These narratives often leveraged pre-existing divisions, further intensifying societal polarization and undermining democratic trust.

- **Limited use of AI-generated images and Deepfake technology in Thailand's political context:** The research found minimal use of AI Deepfake technology in Thai political discourse during the election and post-election periods. Instead, AI-generated images were observed more frequently, primarily used in election campaigns to create visually appealing materials and attract public attention. These findings suggest a less widespread application of AI Deepfake technology in the political landscape compared to other global contexts.

The findings of this research provide actionable recommendations for enhancing online moderation and intervention practices. These include fostering closer collaboration between platforms, policymakers, and civil society to ensure more robust enforcement of community standards. The study advocates for a balanced approach that prioritizes both the protection of free speech and the maintenance of a safe and trustworthy online environment. By addressing the gaps identified in this research, stakeholders can contribute to a healthier digital ecosystem that supports democratic resilience during political events, social crises, and future elections in Thailand.

Kulachada Chaipipat

Advisor to Cofact Thailand

15th December 2024

คำนำ

รายงานฉบับนี้จัดทำขึ้นเพื่อศึกษาและนำเสนอแนวทางที่เหมาะสมในการสร้างสมดุลระหว่างเสรีภาพในการแสดงออกกับการรักษาความเชื่อมั่นและความปลอดภัยออนไลน์ในประเทศไทย ภายใต้บริบทการเลือกตั้งสมาชิกสภาผู้แทนราษฎรไทยเป็นการทั่วไป พ.ศ. 2566 และผลกระทบที่ตามมา งานศึกษานี้เน้นไปที่ประเด็นข้อมูลลวงที่อยู่บนพื้นฐานของความจริง (Malinformation) และ ประชวาทา (Hate Speech) ซึ่งส่งผลกระทบอย่างมีนัยสำคัญต่อบรรยากาศทางสังคม ความเชื่อมั่นในกระบวนการประชาธิปไตย และความไว้วางใจในพื้นที่ดิจิทัล

ด้วยการผสมผสานระเบียบวิธีการวิจัยเชิงคุณภาพและเชิงปริมาณ งานศึกษานี้ได้รวบรวมและวิเคราะห์ข้อมูลจากแพลตฟอร์มสื่อสังคมออนไลน์ที่สำคัญ เพื่อศึกษาโครงสร้างเครือข่ายสังคม (Social Network Analysis: SNA) ลักษณะการกระจายตัวของข้อมูล และพฤติกรรมของผู้ใช้งานในช่วงก่อน ระหว่าง และหลังการเลือกตั้ง ผลการศึกษาชี้ให้เห็นถึงช่องว่างและข้อจำกัดในมาตรการจัดการเนื้อหาออนไลน์ของแพลตฟอร์มต่าง ๆ ตลอดจนความซับซ้อนที่เกี่ยวข้องกับวัฒนธรรมย่อยในสังคมไทย

รายงานฉบับนี้ยังนำเสนอข้อเสนอแนะที่สำคัญต่อองค์กรสื่อ แพลตฟอร์มดิจิทัล และภาคประชาสังคม โดยยึดมั่นในหลักการการกำกับดูแลร่วมกัน (Collaborative Governance) มากกว่าการให้อำนาจภาครัฐแทรกแซง ซึ่งอาจกระทบต่อหลักการเสรีภาพในการแสดงออก ข้อเสนอแนะดังกล่าวมุ่งหวังที่จะสร้างความสมดุลระหว่างการปกป้องสิทธิของประชาชนและการรักษาสภาพแวดล้อมออนไลน์ที่ปลอดภัย

หวังเป็นอย่างยิ่งว่ารายงานฉบับนี้จะประโยชน์ต่อผู้กำหนดนโยบาย นักวิชาการ แพลตฟอร์มดิจิทัล และภาคประชาสังคม ในการร่วมกันพัฒนากลยุทธ์ที่มีประสิทธิภาพและยั่งยืน เพื่อรับมือกับความท้าทายในโลกออนไลน์ที่เปลี่ยนแปลงอยู่ตลอดเวลา

ทีมวิจัย
บริษัท ซีโอ โมเมนตัม จำกัด
15 ธันวาคม 2567

สารบัญ

<i>Acknowledgement and Disclaimer</i>	i
<i>Executive Summary</i>	A
คำนำ.....	ก
สารบัญ	ข
สารบัญรูป.....	ง
สารบัญตาราง.....	จ
1. ภาพรวมของโครงการ	1
1.1 หลักการและเหตุผล	1
1.2 วัตถุประสงค์.....	1
1.3 ผลที่คาดว่าจะได้รับ.....	1
1.4 ขอบเขตการศึกษา	2
1.4.1 ขอบเขตด้านเวลา	2
1.4.2 ขอบเขตด้านพื้นที่.....	2
1.4.3 ขอบเขตด้านเนื้อหา	2
1.5 การดำเนินงาน	2
2. ระเบียบวิธีวิจัย (Methodology)	5
2.1 ข้อมูลที่ใช้วิเคราะห์	5
2.2 การเก็บรวบรวมข้อมูล	5
2.2.1 การเก็บรวบรวมข้อมูลเลือกตั้ง	6
2.2.2 การเก็บรวบรวมข้อมูลหลังเลือกตั้ง.....	6
2.3 การทำความสะอาดข้อมูล	6
2.4 การทดสอบความถูกต้องและแม่นยำของข้อมูล (Accuracy)	7
2.4.1 ผลการทดสอบความถูกต้องและแม่นยำของข้อมูลเลือกตั้ง	8
2.4.2 ผลการทดสอบความถูกต้องและแม่นยำของข้อมูลหลังเลือกตั้ง.....	10
2.5 การวิเคราะห์ข้อมูล	11
2.6 กรอบการศึกษา (Framework)	12
3. ผลการดำเนินงาน	13
3.1 กิจกรรมที่ 1 ทบทวนเอกสารนโยบายด้านความปลอดภัยของแพลตฟอร์ม	13
3.1.1 นโยบายและมาตรฐานความปลอดภัยของแพลตฟอร์ม	13
3.1.2 การวิเคราะห์กลไกการดำเนินงานและเปรียบเทียบวิธีการของแพลตฟอร์มในการจัดการกับปัญหา.....	13
3.1.3 กรณีศึกษาที่สำคัญที่เกิดขึ้นช่วงการเลือกตั้ง และหลังการเลือกตั้ง	15
3.2 กิจกรรมที่ 2 การวิเคราะห์โซเชียลมีเดีย (Social Media Analysis).....	22



3.2.1 การจำแนกประเภทข้อความตามแก่นสาร.....	22
3.2.2 การวิเคราะห์หัวข้อด้วยอัลกอริทึมการจำลองหัวข้อ.....	26
3.2.3 การวิเคราะห์ AI Deepfake	33
3.2.4 การวิเคราะห์เครือข่ายทางสังคม (Social Network Analysis: SNA)	35
3.3 กิจกรรมที่ 3 จัดวงเสวนาผู้มีส่วนได้ส่วนเสียและผู้ที่มีส่วนเกี่ยวข้อง (Stakeholder Discussion)	38
4. <i>อภิปรายผลการศึกษาและข้อเสนอแนะ</i>	40
4.1 ช่องว่างในการนำมาตรฐานชุมชนของแพลตฟอร์มที่เกี่ยวข้องกับความเชื่อมั่นและความปลอดภัยออนไลน์ไปปฏิบัติ.....	40
4.2 แนวทางการกำกับดูแลเนื้อหาออนไลน์หรือการแทรกแซงที่เป็นไปได้	41
4.3 ข้อเสนอแนะจากทีมวิจัย	42
4.4 ข้อเสนอแนะ	45
5. <i>บรรณานุกรม</i>	47
6. <i>ภาคผนวก</i>	48
ก. ตารางเปรียบเทียบแนวปฏิบัติของแพลตฟอร์มในการกลั่นกรองเนื้อหาที่เป็นอันตรายและข้อมูลบิดเบือน	48
ข. คำคัดออก (Excluded Keywords)	72

สารบัญรูป

รูป 1 สรุปจำนวนข้อมูลและบัญชีที่ใช้ในการวิเคราะห์.....	5
รูป 2 สัดส่วนข้อมูลดิบกับข้อมูลที่ทำความสะอาดแล้ว	7
รูป 3 โครงสร้างของ Confusion Matrix	8
รูป 4 กรณศึกษาที่ 1 ช่วงเลือกตั้ง	15
รูป 5 กรณศึกษาที่ 2 ช่วงเลือกตั้ง	16
รูป 6 กรณศึกษาที่ 3 ช่วงเลือกตั้ง	16
รูป 7 กรณศึกษาที่ 4 ช่วงเลือกตั้ง	17
รูป 8 กรณศึกษาที่ 5 ช่วงเลือกตั้ง	18
รูป 9 กรณศึกษาที่ 1 ช่วงหลังเลือกตั้ง	19
รูป 10 กรณศึกษาที่ 2 ช่วงหลังเลือกตั้ง	20
รูป 11 กรณศึกษาที่ 3 ช่วงหลังเลือกตั้ง	21
รูป 12 กรณศึกษาที่ 4 ช่วงหลังเลือกตั้ง	22
รูป 13 ผลการจำแนกประเภทข้อความตามแก่นสาระในภาพรวม	23
รูป 14 ผลการจำแนกประเภทข้อความตามแก่นสาระของข้อมูลช่วงเลือกตั้ง	24
รูป 15 ผลการจำแนกประเภทข้อความตามแก่นสาระของข้อมูลช่วงหลังเลือกตั้ง	25
รูป 16 เหตุการณ์ทางการเมืองที่สำคัญที่มียอดการมีส่วนร่วมสูงในโซเชียลมีเดีย	26
รูป 17 การใช้เทคโนโลยี AI Generated Images เพื่อผลิตภาพประกอบแคมเปญหาเสียง	34
รูป 18 AI Generated Images กับการ Demonization	34
รูป 19 การ์ตูนล้อการเมือง หรือ Editorial Cartoons	34
รูป 20 เครือข่ายบัญชีผู้ใช้งาน Facebook	35
รูป 21 เครือข่ายบัญชีผู้ใช้งาน X	37

สารบัญตาราง

ตาราง 1 การจำแนกสัดส่วนเนื้อหาที่เป็น Hate Speech โดยผู้เชี่ยวชาญของข้อมูลช่วงเลือกตั้ง	8
ตาราง 2 การจำแนกสัดส่วนเนื้อหาที่เป็น Malinformation โดยผู้เชี่ยวชาญของข้อมูลช่วงเลือกตั้ง	9
ตาราง 3 การวิเคราะห์ความถูกต้องและแม่นยำโดย Confusion Matrix ข้อมูลช่วงเลือกตั้ง	9
ตาราง 4 การจำแนกสัดส่วนเนื้อหาที่เป็น Hate Speech โดยผู้เชี่ยวชาญของข้อมูลช่วงหลังเลือกตั้ง	10
ตาราง 5 การจำแนกสัดส่วนเนื้อหาที่เป็น Malinformation โดยผู้เชี่ยวชาญของข้อมูลช่วงหลังเลือกตั้ง	10
ตาราง 6 การวิเคราะห์ความถูกต้องและแม่นยำโดย Confusion Matrix ข้อมูลช่วงหลังเลือกตั้ง	11
ตาราง 7 นโยบายและมาตรฐานความปลอดภัยของแพลตฟอร์มที่เกี่ยวข้องกับการเลือกตั้งและสถานการณ์ที่มีความ อ่อนไหวทางการเมือง	13
ตาราง 8 การจำลองหัวข้อที่ปรากฏขึ้นหลังการเลือกตั้งในช่วงที่ 1 ของแพลตฟอร์ม X	27
ตาราง 9 การจำลองหัวข้อที่ปรากฏขึ้นหลังการเลือกตั้งในช่วงที่ 2 ของแพลตฟอร์ม X	28
ตาราง 10 การจำลองหัวข้อที่ปรากฏขึ้นหลังการเลือกตั้งในช่วงที่ 3 ของแพลตฟอร์ม X	29
ตาราง 11 การจำลองหัวข้อที่ปรากฏขึ้นหลังการเลือกตั้งในช่วงที่ 1 ของแพลตฟอร์ม Facebook	30
ตาราง 12 การจำลองหัวข้อที่ปรากฏขึ้นหลังการเลือกตั้งในช่วงที่ 2 ของแพลตฟอร์ม Facebook	31
ตาราง 13 การจำลองหัวข้อที่ปรากฏขึ้นหลังการเลือกตั้งในช่วงที่ 3 ของแพลตฟอร์ม Facebook	32

1. ภาพรวมของโครงการ

1.1 หลักการและเหตุผล

แพลตฟอร์มดิจิทัลเป็นพื้นที่เชิงกลยุทธ์ที่มีความสำคัญต่อการสนทนาทางการเมืองในระบอบประชาธิปไตยของประเทศไทยนับตั้งแต่การตื่นตัวทางการเมืองของเยาวชน นิสิต นักศึกษา ตั้งแต่ปี 2562 และปรากฏเด่นชัดขึ้นในช่วงการเลือกตั้งทั่วไป ปี 2566 อย่างไรก็ตามเมื่อพ้นจากช่วงการเลือกตั้ง สื่อและแพลตฟอร์มหันเหความสนใจเรื่องความเชื่อมั่นและความปลอดภัยทางออนไลน์ไปจากประเด็นการเมืองและการเลือกตั้ง ข้อมูลที่เชื่อถือได้ ซึ่งเป็นพื้นฐานสำคัญสำหรับการอภิปรายทางการเมืองอย่างรู้เท่าทันและการหาทางออกทางการเมืองกลับตกอยู่ในความเสี่ยง กล่าวคือ พื้นที่ออนไลน์เต็มไปด้วยข้อมูลบิดเบือนและแคมเปญความเกลียดชังที่เพิ่มขึ้นอย่างต่อเนื่อง ซึ่งถูกจัดตั้งขึ้นเพื่อโจมตีคู่แข่งทางการเมือง ก่อให้เกิดความเสี่ยงต่อการเปลี่ยนผ่านทางการเมืองที่ปรารถนาของประเทศ ซึ่งมีความรุนแรงมากยิ่งขึ้นกว่าช่วงการเลือกตั้งทั่วไปปี 2566

ในช่วงการเลือกตั้งทั่วไปเมื่อปี 2566 โครงการการติดตามกระบวนการเลือกตั้งในสื่อสังคมออนไลน์ (Digital Election Analytic Lab: DEAL) ได้ติดตามและวิเคราะห์เครือข่ายสังคม พบว่า นอกจากกองทัพและรัฐบาลที่ได้มีใช้ปฏิบัติการข้อมูลข่าวสาร (Information Operation: IO) ที่ถูกอภิปรายในสภาผู้แทนราษฎร เมื่อปี 2563 (ไทยรัฐ, 2563) แล้วพรรคการเมืองใหญ่ทุกพรรค เช่น รวมไทยสร้างชาติ เพื่อไทย และก้าวไกลต่างก็ใช้ *ปฏิบัติการปั่นข้อมูลข่าวสารในโลกออนไลน์ (Online Influence Operations)* ทั้งแบบโปร่งใส (overt) และไม่โปร่งใส (covert) ในระดับที่ต่างกัน เพื่อโน้มน้าวผู้มีสิทธิเลือกตั้ง รวมทั้งโจมตีคู่แข่งทางการเมือง

ภายหลังการเลือกตั้ง ปฏิบัติการปั่นข้อมูลข่าวสารในโลกออนไลน์ก็ยังเข้มข้น และเปลี่ยนแปลงรูปแบบไปจากช่วงการเลือกตั้งที่อาจจะมุ่งเน้นแค่โน้มน้าวผู้มีสิทธิเลือกตั้ง มาสู่การโจมตีบุคคล การสร้างภาพลักษณ์ที่ไม่ดี และการทำลายความเชื่อถือของพรรคการเมืองและนักการเมือง เช่น การเปิดเผยข้อมูลส่วนบุคคล (doxing) และการออกมาแฉ (call-out tactics) ในหลายกรณี กลวิธีเหล่านี้ได้พัฒนาไปสู่การบิดเบือนข้อเท็จจริงในลักษณะ malinformation รวมถึงการใช้ทฤษฎีสมคบคิด เพื่อหลบเลี่ยงการตรวจสอบจากแพลตฟอร์ม และทำให้กระบวนการตรวจสอบข้อเท็จจริงทำได้ยากมากยิ่งขึ้น

ดังนั้น หากไม่มีกำกับดูแลที่เหมาะสม หรือการแทรกแซงที่มีประสิทธิภาพ บทสนทนาทางการเมืองออนไลน์ที่เป็นอันตรายเหล่านี้อาจทำลายความเชื่อมั่นในสถาบันประชาธิปไตยและคุณค่าของระบอบประชาธิปไตยมากขึ้น และผลักดันสังคมที่มีความแตกแยกมายาวนานให้เข้าสู่ความสุดโต่งและสภาวะที่ไม่สามารถปรองดองกันได้

1.2 วัตถุประสงค์

1.2.1 เพื่อระบุช่องว่างในการนำมาตรฐานชุมชนของแพลตฟอร์มที่เกี่ยวข้องกับความเชื่อมั่นและความปลอดภัยออนไลน์ไปปฏิบัติ รวมถึงนโยบายการกำกับดูแลอินเทอร์เน็ตในช่วงการเลือกตั้งและหลังการเลือกตั้ง ซึ่งนำไปสู่การเพิ่มขึ้นของข้อมูลที่มีแนวโน้มเป็นอันตรายและข้อมูลบิดเบือนในช่วงหลังการเลือกตั้งของไทยจนถึงปัจจุบัน

1.2.2 เพื่อสำรวจแนวทางที่เป็นไปได้ในการกำกับดูแลเนื้อหาออนไลน์ที่เหมาะสมหรือการแทรกแซงที่มีประสิทธิภาพจากแพลตฟอร์ม รวมถึงบทบาทของสื่อในการส่งเสริมการสื่อสารออนไลน์ที่อิงข้อเท็จจริงและปราศจากความรุนแรง ในขณะเดียวกันก็รักษาสมดุลของเสรีภาพในการแสดงออก

1.3 ผลที่คาดว่าจะได้รับ

1.3.1 รายงานสรุปมาตรฐานชุมชนของแพลตฟอร์ม

1.3.2 รายงานการวิเคราะห์สื่อสังคมออนไลน์

1.3.3 รายงานการประชุม รวมถึงรายชื่อผู้เข้าร่วม (ผู้ให้สัมภาษณ์ ถ้ามี) บันทึกการประชุม และประเด็นสำคัญที่ได้จากการอภิปราย

1.3.4 รายงานฉบับสมบูรณ์ ในประเด็นสถานะของข้อมูลบิดเบือนหลังการเลือกตั้งในประเทศไทย รวมถึงแนวทางที่เป็นไปได้ในการกำกับดูแลเนื้อหาออนไลน์ที่เหมาะสมหรือการแทรกแซงที่มีประสิทธิภาพจากแพลตฟอร์ม

1.4 ขอบเขตการศึกษา

1.4.1 ขอบเขตด้านเวลา

โครงการฯ จะศึกษาและวิเคราะห์บทสนทนาในสื่อสังคมออนไลน์ 2 ช่วง ได้แก่

- 1) ช่วงการเลือกตั้งสมาชิกสภาผู้แทนราษฎรไทยเป็นการทั่วไป พ.ศ. 2566 ของประเทศไทย ตั้งแต่ วันที่ 1 เมษายน 2566 ถึงวันที่ 15 มิถุนายน 2566
- 2) ช่วงหลังเลือกตั้งสมาชิกสภาผู้แทนราษฎรไทยเป็นการทั่วไป พ.ศ. 2566 ของประเทศไทย ตั้งแต่ วันที่ 1 มกราคม 2567 ถึงวันที่ 31 สิงหาคม 2567

1.4.2 ขอบเขตด้านพื้นที่

โครงการฯ จะศึกษาและวิเคราะห์เอกสารมาตรฐานชุมชน รวมทั้งบทสนทนาในสื่อสังคมออนไลน์ ใน 2 แพลตฟอร์ม ได้แก่ X (หรือ Twitter ซึ่งต่อไปนี้จะเรียกว่า X) และ Facebook เท่านั้น

1.4.3 ขอบเขตด้านเนื้อหา

โครงการฯ มุ่งศึกษาและวิเคราะห์สถานะของข้อมูลบิดเบือนหลังการเลือกตั้งในประเทศไทย รวมถึงแนวทางที่เป็นไปได้ในการกำกับดูแลเนื้อหาออนไลน์ที่เหมาะสมหรือการแทรกแซงที่มีประสิทธิภาพจากแพลตฟอร์ม ดังนั้นจะให้ความสำคัญกับการวิเคราะห์เนื้อหาใน 3 ประเด็นหลัก ได้แก่

- 1) **Hate Speech** คำพูดที่แสดงความเกลียดชัง แบ่งออกเป็น 5 ระดับ คือ
 - ระดับที่ 1 การด่าด้วยถ้อยคำหยาบคาย การดูหมิ่นสติปัญญา การดูถูกเหยียดหยาม และการสร้างความเป็นอื่น
 - ระดับที่ 2 การคุกคาม การข่มขู่ และการเปิดเผยข้อมูลส่วนตัว การเรียกร้องให้มีการบังคับใช้กฎหมายเพื่อบั่นทอนเสรีภาพในการแสดงออก
 - ระดับที่ 3 การเหมารวม การปะปน การลดทอนคุณค่า การดูถูกเหยียดหยาม การทำให้เป็นปศุสัตว์
 - ระดับที่ 4 การแบ่งแยก กีดกัน การปฏิเสธที่จะอยู่ร่วมกันในสังคม
 - ระดับที่ 5 การยุยง การสนับสนุน การเรียกร้องให้เกิดความรุนแรง การยกย่องผู้ก่อความรุนแรง การให้ความชอบธรรมกับการก่อความรุนแรง
- 2) **Malinformation** หมายถึง การบิดเบือนข้อมูลที่อยู่บนพื้นฐานของความจริง แต่ถูกนำมาใช้ในบริบทที่ผิดหรือแฝงเจตนาที่ทำให้ผู้รับสารเกิดความเข้าใจผิด มีเจตนาเพื่อก่อให้เกิดอันตรายทำลายชื่อเสียง โจมตีความคิด บุคคล องค์กร หรือประเทศ ฯลฯ¹
- 3) **AI Deepfake** หมายถึง การใช้เทคโนโลยี AI (Artificial Intelligence) สร้างเนื้อหาสื่อ เช่น ข้อความ รูปภาพ คลิปเสียง วิดีโอปลอมของบุคคล โดยมีเจตนาเพื่อหลอกลวง^{2,3}

1.5 การดำเนินงาน

ระยะเวลาการดำเนินโครงการ จำนวน 120 วัน (15 สิงหาคม 2567 – 15 พฤศจิกายน 2567 และขยายเวลาเพิ่ม 30 วัน สิ้นสุด 15 ธันวาคม 2567)

¹ Wardle, C. & Dias, D. (2017). *Information disorder: Toward an interdisciplinary framework for research and policy making*. Council of Europe. สืบค้นเมื่อ 15 ตุลาคม 2567 จาก <https://coilink.org/20.500.12592/321s46>

² State of Illinois. (m.d.). *United States Deepfake bill: 10 ILCS 5/29-21*. สืบค้นเมื่อ 15 ตุลาคม 2567 จาก <https://www.ilga.gov/legislation/fulltext.asp?DocName=&SessionId=108&GA=101&DocTypeId=HB&DocNum=5321&GAID=15&LegID=125899&SpecSess=&Session=>

³ Northwestern Buffett Institute for Global Affairs. (2023). *The Rise of Artificial Intelligence and Deepfakes* (BUFFETT BRIEF • JULY 2023). สืบค้นเมื่อ 15 ตุลาคม 2567 จาก https://buffett.northwestern.edu/documents/buffett-brief_the-rise-of-ai-and-deepfake-technology.pdf

วัตถุประสงค์	กิจกรรม	ระยะเวลา	ผลผลิต
1. การทบทวนเอกสารมาตรฐานชุมชนของแพลตฟอร์ม และการวิเคราะห์กรณีศึกษา	รวบรวมและทบทวนเอกสารมาตรฐานชุมชนที่จัดทำโดย Facebook, X และ YouTube วิเคราะห์กลไกการดำเนินงาน เช่น การแจ้งเตือนและการลบเนื้อหา ขั้นตอนการอุทธรณ์ นโยบายการเซ็นเซอร์ และระบบการรายงาน คัดเลือกกรณีศึกษาที่สำคัญที่เกิดขึ้นช่วงการเลือกตั้ง และหลังการเลือกตั้ง เพื่อเปรียบเทียบการดำเนินงานของแพลตฟอร์มในการจัดการกับเนื้อหาอันตรายและข้อมูลบิดเบือน เปรียบเทียบวิธีการของแต่ละแพลตฟอร์มในการจัดการกับปัญหาต่างๆ ได้แก่ malinformation, hate speech, และ AI deepfake	15 ส.ค. – 15 ก.ย. 2567	1. ตารางเปรียบเทียบแนวปฏิบัติของแพลตฟอร์มในการกลั่นกรองเนื้อหาที่เป็นอันตรายและข้อมูลบิดเบือน 2. ผลการวิเคราะห์กรณีศึกษาของการดำเนินการของแพลตฟอร์มต่อข้อมูลที่เป็น malinformation, hate speech, และ AI deepfake รวมถึงผลกระทบต่อความถูกต้องของข้อมูล การเซ็นเซอร์ และเสรีภาพในการแสดงออก
2. การวิเคราะห์ข้อมูลสื่อสังคมออนไลน์ที่เกี่ยวข้องกับการเลือกตั้ง	- วิเคราะห์ข้อมูลสื่อสังคมออนไลน์ที่เกี่ยวข้องกับการเลือกตั้งย้อนหลัง เพื่อทำความเข้าใจภูมิทัศน์ของบทสนทนาออนไลน์ในช่วงการเลือกตั้ง - หากจำเป็น ดำเนินการเก็บข้อมูลเพิ่มเติมและวิเคราะห์สภาพแวดล้อมออนไลน์ในปัจจุบัน โดยเน้นบัญชีผู้ใช้หรือเพจที่เผยแพร่ข้อมูลบิดเบือนทางการเมือง - ใช้เทคนิคการวิเคราะห์เครือข่ายเพื่อระบุรูปแบบการแพร่กระจายและตัวแสดงหลักที่มีส่วนในการเผยแพร่ข้อมูลที่เป็น malinformation, hate speech, และ AI deepfake - วิเคราะห์การดำเนินการของแพลตฟอร์มและการตอบสนองของสื่อเพื่อระบุช่องว่างและแนวทางในการปรับปรุงการกลั่นกรองเนื้อหาและกลยุทธการแทรกแซงออนไลน์		ข้อมูลเชิงลึกเกี่ยวกับการเพิ่มขึ้นของข้อมูลที่มีแนวโน้มเป็นอันตรายและข้อมูลบิดเบือนในช่วงหลังการเลือกตั้งของไทย จนถึงปัจจุบัน โดยเน้นประเด็นที่ต้องการการแทรกแซงและการกลั่นกรองเนื้อหา
3. วิเคราะห์ข้อมูลและจัดทำร่างรายงาน	จัดประชุมภายในทีมวิจัยเพื่อรวบรวมผลการศึกษาเบื้องต้นและจัดทำร่างรายงานฉบับแรก	15 ก.ย. – 15 ต.ค. 2567	รายงานความก้าวหน้า

วัตถุประสงค์	กิจกรรม	ระยะเวลา	ผลผลิต
4. จัดประชุมเชิงปฏิบัติการกับผู้มีส่วนได้ส่วนเสีย	- จัดประชุมสนทนาแบบปิด หรือ จัดกิจกรรม Tabletop Exercise ร่วมกับผู้มีส่วนได้ส่วนเสียหลัก รวมถึงตัวแทนแพลตฟอร์ม ผู้กำหนดนโยบาย องค์กรภาคประชาสังคม และสื่อมวลชน - อำนวยความสะดวกในการอภิปรายเพื่อเชื่อมช่องว่างและผลักดันการพัฒนามาตรฐานชุมชนที่บังคับใช้ได้ โดยพิจารณาจากผลการทบทวนเอกสาร กรณีศึกษา และการวิเคราะห์ข้อมูล - หากจำเป็น สัมภาษณ์ผู้มีส่วนได้ส่วนเสียในโลกออนไลน์ รวมถึงผู้ใช้ เพื่อยืนยันผลการศึกษาและรวบรวมข้อมูลเชิงลึกเพิ่มเติมเกี่ยวกับมุมมองต่อการดำเนินการของแพลตฟอร์มและการตอบสนองของสื่อ	สัปดาห์สุดท้ายของเดือนตุลาคม 2567	ข้อมูลเชิงลึกเพื่อยืนยันผลการศึกษาและข้อเสนอแนะจากผู้มีส่วนได้ส่วนเสียในการปรับปรุงการกลั่นกรองเนื้อหาออนไลน์ เพื่อส่งเสริมความถูกต้องของข้อมูล และเพิ่มความเชื่อมั่นและความปลอดภัยในโลกออนไลน์
5. สรุปผลการดำเนินงาน	ประชุมภายในทีมวิจัยเพื่อสรุปผลการดำเนินงานและปรับปรุงร่างรายงานฉบับสมบูรณ์	สัปดาห์แรกของเดือนพฤศจิกายน 2564	
6. จัดทำรายงานฉบับสมบูรณ์	ประชุมภายในทีมวิจัยเพื่อจัดทำรายงานฉบับสมบูรณ์	15 พ.ย. 2567	รายงานฉบับสมบูรณ์

2. ระเบียบวิธีวิจัย (Methodology)

โครงการศึกษาแนวทางการสร้างสมดุลระหว่างเสรีภาพในการพูดและมาตรฐานความปลอดภัยในโลกออนไลน์ ด้วยการกลั่นกรองเนื้อหาที่มีประสิทธิภาพ ได้มีการประยุกต์ใช้ระเบียบวิธีวิจัยทางสังคมศาสตร์มาเป็นกรอบการศึกษา (framework) หลักในการดำเนินงาน กล่าวคือ โครงการฯ ประยุกต์ใช้ระเบียบวิธีวิจัยแบบผสมผสานระหว่างวิธีวิจัยเชิงคุณภาพ (qualitative methods) และเชิงปริมาณ (quantitative methods) วิเคราะห์ข้อมูลโดยใช้แนวทางการวิเคราะห์แก่นสาระ (thematic analysis) ร่วมกับการวิเคราะห์พฤติกรรมของบัญชีผู้ใช้ โดยมีรายละเอียดดังนี้

2.1 ข้อมูลที่ใช้วิเคราะห์

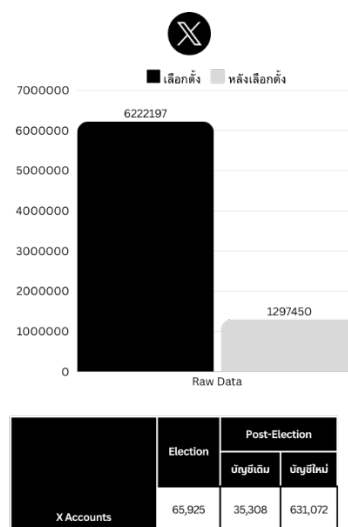
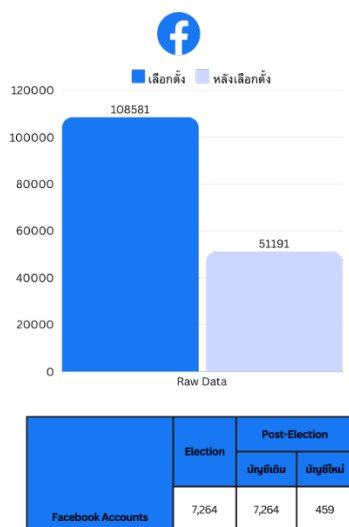
ข้อมูลที่ใช้โครงการฯ ใช้วิเคราะห์คือ ข้อมูลโซเชียลมีเดียในบริบทการเลือกตั้งทั่วไป พ.ศ. 2566 ของประเทศไทย โดยใช้ข้อมูล 2 ส่วนใหญ่ในการวิเคราะห์ คือ 1) ข้อมูลช่วงการเลือกตั้ง (“ข้อมูลเลือกตั้ง”) และ 2) ข้อมูลสื่อสังคมออนไลน์ในปัจจุบัน (“ข้อมูลหลังเลือกตั้ง”) ดังนี้

ข้อมูลเลือกตั้ง จะเป็นข้อมูลที่ครอบคลุมระยะเวลาตั้งแต่ 1 เมษายน 2566 – 15 มิถุนายน 2566 จำนวน 6,333,778 ข้อความ โดยมาจาก 73,189 บัญชี โดยแบ่งออกตามแพลตฟอร์มได้แก่

- X จำนวน 6,225,197 ข้อความ จาก 65,925 บัญชี
- Facebook จำนวน 108,581 ข้อความ จาก 7,264 บัญชี

ข้อมูลหลังเลือกตั้ง จะเป็นข้อมูลที่ครอบคลุมระยะเวลาตั้งแต่ 1 มกราคม 2567 – 31 สิงหาคม 2567 จำนวน 1,348,641 ข้อความ โดยมาจาก 674,103 บัญชี โดยแบ่งออกตามแพลตฟอร์มได้แก่

- X จำนวน 1,297,450 ข้อความ จาก 666,380 บัญชี
- Facebook จำนวน 51,191 ข้อความ จาก 7,723 บัญชี



รูป 1 สรุปจำนวนข้อมูลและบัญชีที่ใช้ในการวิเคราะห์

2.2 การเก็บรวบรวมข้อมูล

โครงการฯ ใช้วิธีการเก็บข้อมูลแบบบัญชีเป็นฐาน (actor-based) กล่าวคือ เป็นกระบวนการเก็บข้อมูลโดยเก็บข้อมูลจากบัญชีผู้ใช้งานบนแพลตฟอร์มหนึ่งๆ ซึ่งวิธีการนี้มุ่งเน้นการเก็บข้อมูลจากทุกกิจกรรมของบัญชีผู้ใช้นั้น ซึ่งรวมถึงโพสต์ ความคิดเห็น และการมีปฏิสัมพันธ์ต่างๆ (interaction) โดยโครงการฯ ต้องเตรียมบัญชีรายชื่อ Username หรือ User ID ของผู้ใช้งานที่ต้องการเก็บข้อมูลสำหรับแต่ละแพลตฟอร์ม และกำหนดกรอบเวลาที่ชัดเจนเพื่อให้ได้ข้อมูลในช่วงที่สนใจหรือตามที่กำหนดไว้ แทนการใช้คำสำคัญ (keywords) ในการค้นหาที่มักจะพบกับปัญหาเรื่องอคติของนักวิจัยในการเลือกคำสำคัญ

ในการศึกษาครั้งนี้ ใช้ข้อมูลโซเชียลมีเดีย 2 ช่วง ที่มีรูปแบบและเครื่องมือการเก็บรวบรวมข้อมูลที่แตกต่างกัน โดยมีรายละเอียดดังนี้

2.2.1 การเก็บรวบรวมข้อมูลเลือกตั้ง

การเก็บข้อมูลโซเชียลมีเดียในช่วงการเลือกตั้ง ระหว่างวันที่ 1 เมษายน 2566 – 15 มิถุนายน 2566 นั้น เครื่องมือที่ใช้ในการเก็บข้อมูลในแต่ละแพลตฟอร์มก็มีความแตกต่างกันออกไปตามนโยบายเรื่องการบริการข้อมูลของแพลตฟอร์ม โดยใน X โครงการฯ ใช้ Twitter API (Academic Tier Credentials) ในการเก็บข้อมูล และใช้ Crowdtangle เพื่อเก็บข้อมูลจาก Facebook

การเก็บรวบรวมข้อมูลแบบบัญชีเป็นฐานช่วงการเลือกตั้ง แบ่งการเก็บข้อมูลออกเป็น 11 ครั้ง โดยในการเก็บครั้งที่ 1 ใช้การสุ่มกลุ่มตัวอย่างแบบไม่ใช้หลักความน่าจะเป็น (nonprobability sampling) โดยใช้เทคนิคการสุ่มตัวอย่างเฉพาะเจาะจง (purposive sampling) ได้แก่ จากบัญชีที่เป็นทางการ (official account) ของพรรคการเมือง 67 พรรค ผู้ลงรับสมัครเลือกตั้งทั้งแบบแบ่งเขตและแบบบัญชีรายชื่อของแต่ละพรรค 6,324 คน บุคคลที่พรรคการเมืองจะเสนอให้สภาผู้แทนราษฎรพิจารณาให้ความเห็นชอบแต่งตั้งเป็นนายกรัฐมนตรี 62 คน จากนั้นในการเก็บข้อมูลครั้งที่ 2 -11 ใช้เทคนิคการสุ่มตัวอย่างแบบสโนว์บอล (snowball) โดยเลือกกลุ่มตัวอย่างจากบัญชีผู้ใช้ที่เข้ามามีปฏิสัมพันธ์ ได้แก่ การ รีทวีต (Re-tweeting) การชื่นชอบ การโคตทวีต (Co-tweeting) และการกล่าวถึง (Mention) กับบัญชีในการเก็บข้อมูลครั้งก่อนหน้า เพื่อพัฒนาเครือข่ายของบัญชีจนถึงจำนวนที่เราตั้งเกณฑ์เอาไว้ (Threshold)

2.2.2 การเก็บรวบรวมข้อมูลหลังเลือกตั้ง

สำหรับการเก็บข้อมูลโซเชียลมีเดียช่วงหลังการเลือกตั้งนั้น โครงการฯ เก็บข้อมูลเพียง 1 ครั้ง ระหว่างวันที่ 1 มกราคม 2567 – 31 สิงหาคม 2567 โดยใช้เครือข่ายบัญชีเดิมจากข้อมูลในช่วงการเลือกตั้งเป็นฐาน

นอกจากนั้นช่วงระหว่างปี 2566 – 2567 แพลตฟอร์มได้มีการเปลี่ยนแปลงนโยบายเรื่องการบริการข้อมูลของแพลตฟอร์ม ทำให้โครงการฯ ไม่สามารถเข้าถึงเครื่องมือ API ที่เป็นทางการ (Twitter API และ Crowdtangle) ของทั้ง X และ Facebook ได้ โครงการฯ จึงได้เปลี่ยนเครื่องมือการเก็บข้อมูลหลังการเลือกตั้งเป็น Apify ซึ่งเป็นผู้ให้บริการบุคคลที่สามที่ได้รับอนุญาต (licensed third-party)

2.3 การทำความสะอาดข้อมูล

การทำความสะอาดข้อมูล หมายถึง การจัดการกับข้อมูลที่เก็บมาได้ให้ตอบโจทยกับวัตถุประสงค์ของโครงการ กล่าวคือ วัตถุประสงค์ของโครงการคือข้อมูลที่เกี่ยวข้องกับการเดินทางการเมืองโดยเฉพาะอย่างยิ่งการเลือกตั้ง และเมื่อเก็บข้อมูลโดยวิธีการใช้บัญชีเป็นฐาน ทำให้ข้อมูลที่ได้มาจะเป็นทุกกิจกรรมของบัญชีนั้น ซึ่งจะมีข้อมูลที่ไม่เกี่ยวข้องกับการเลือกตั้ง ดังนั้น โครงการฯ จึงต้องทำความสะอาดข้อมูลโดยการคัดกรองข้อความที่ไม่เกี่ยวข้องกับการเลือกตั้งออก โดยใช้วิธีการกำหนดคำคัดออก (excluded keywords) โดยผู้เชี่ยวชาญ ดังนี้

ขั้นตอนที่ 1 ทีม Data Scientist สุ่มตัวอย่างจากฐานข้อมูลมา 800 แถว เพื่อส่งให้ผู้เชี่ยวชาญกำหนดคำคัดออก

ขั้นตอนที่ 2 ทีมผู้เชี่ยวชาญด้านสังคมศาสตร์ จำนวน 2 คน กำหนดคำคัดออกที่ไม่เกี่ยวข้องกับการเมือง (ดูคำคัดออกในภาคผนวก ข) และส่งกลับให้ทีม Data Scientist เพื่อทำการตัดข้อความที่มีคำคัดออกเป็นส่วนประกอบหนึ่งในประโยคออก

ขั้นตอนที่ 3 ทำซ้ำตั้งแต่ขั้นตอนที่ 1 รวม 6 ครั้ง

ผลการทำความสะอาดข้อมูลพบว่า ในช่วงเลือกตั้งเราพบข้อความที่ไม่เกี่ยวข้องกับการเลือกตั้งน้อยมาก โดยใน Facebook จากข้อมูลดิบเมื่อทำความสะอาดข้อมูลแล้ว มีข้อความที่ไม่เกี่ยวข้องกับการเลือกตั้ง คิดเป็นร้อยละ 1.24 เช่นเดียวกับ X ที่มีเพียงร้อยละ 2.49 ที่ไม่เกี่ยวข้องการเมือง ในขณะที่หลังเลือกตั้ง ใน X ที่กว่าร้อยละ 98 ของข้อมูลที่เก็บมาได้ไม่เกี่ยวข้องกับการเมือง และ Facebook อยู่ที่ประมาณร้อยละ 49%

ยิ่งไปกว่านั้นเมื่อพิจารณาเปรียบเทียบกับจำนวนบัญชีที่พบ จะเห็นว่าใน Facebook บัญชีที่เคยตื่นตัวช่วงเลือกตั้ง เมื่อผ่านช่วงเลือกตั้งมาแล้วพบว่ายังอยู่ครบทั้ง 7,264 บัญชี และยังพบบัญชีใหม่ที่มีปฏิสัมพันธ์ใน

เครือข่ายเพิ่มอีก 459 บัญชี รวมเป็น 7723 บัญชี ในขณะที่จำนวนบัญชีเพิ่มขึ้น แต่สังเกตเห็นว่าจำนวนข้อความที่เก็บรวบรวมได้หายไปกว่าครึ่ง กลับกันเมื่อพิจารณาแพลตฟอร์ม X คือ จากจำนวนบัญชีที่เคยตื่นตัวช่วงเลือกตั้ง 65,925 บัญชี พอเก็บข้อมูลใหม่หลังเลือกตั้ง โครงการฯ ได้คัดเลือกบัญชีที่เกี่ยวข้องกับการเมือง จำนวน 35,308 บัญชี โดยพิจารณาจากคะแนน score ที่มีนัยสำคัญต่อ network แต่ในขณะที่แม้ว่าโครงการฯ จะคัดเลือกบัญชีออกกว่า 50% แต่กลับพบบัญชีใหม่ที่มีปฏิสัมพันธ์ใน network เพิ่มขึ้นถึง 6.3 แสนกว่าบัญชี โดยพบว่า เป็นข้อความโฆษณาจากแพลตฟอร์ม กว่า 4 แสน บัญชี (วิเคราะห์จากการโพสต์เพียง 1 โพสต์ ในรอบ 8 เดือน) (รูป 2)



รูป 2 สัดส่วนข้อมูลดิบกับข้อมูลที่ทำความสะอาดแล้ว

2.4 การทดสอบความถูกต้องและแม่นยำของข้อมูล (Accuracy)

การวิเคราะห์เนื้อหา Hate Speech และ Malinformation ของการศึกษานี้ โครงการฯ ได้ใช้ GPT Model มาช่วยในการจำแนกประเภทข้อมูลที่โครงการสนใจทั้ง กล่าวคือ การนำโมเดลภาษาขนาดใหญ่ (Large Language Model หรือ LLM) มาช่วยในการจัดหมวดหมู่หรือวิเคราะห์ข้อมูลตามเงื่อนไขหรือคำสั่งที่เรากำหนด โดย LLM นั้นมีข้อได้เปรียบเนื่องจากสถาปัตยกรรมของโมเดลสามารถเข้าใจบริบทและความหมายเชิงลึกของข้อความได้ดี เนื่องจากได้รับการฝึกด้วยข้อมูลจำนวนมาก ครอบคลุมหลายภาษาและหลายบริบท ทำให้ GPT Model สามารถจัดการกับข้อความที่ซับซ้อนและมีภาษาสลับกัน เช่น ภาษาไทยผสมอังกฤษ หรือคำแสลงได้เป็นอย่างดี โดยประโยชน์ที่เราได้จากการนำ GPT Model มาใช้ในการวิเคราะห์ครั้งนี้ไม่ใช่เพียงแต่การนำมาช่วยในการจำแนกประเภทข้อมูลเท่านั้น แต่ยังช่วยลด Human Error และ ประหยัดเวลาในการ Label ข้อความจำนวนมากได้

การทำงานของ GPT Model ครั้งนี้ โครงการฯ ได้เขียน คำสั่ง(Instruction) หรือ ตัวอย่างอ้างอิง (Criteria) ที่ชัดเจน เพื่อเป็นการบอกรายละเอียดให้โมเดลเข้าใจบริบท หรือขอบเขตที่โมเดลจะต้องปฏิบัติตาม รวมไปถึงตัวอย่างของข้อความ คำอธิบายของ Criteria แต่ละอย่าง เพื่อให้โมเดลเห็นตัวอย่างในการนำไปใช้งาน และ ส่งข้อความที่ยังไม่มีการ Label ไปหาโมเดลเพื่อให้โมเดลประเมินและ Label ข้อความกลับมา โดยผลลัพธ์หรือคำตอบ (Response) ที่ได้ สามารถนำข้อความที่ได้รับการจัดหมวดหมู่จากโมเดลกลับไป Mapping กับข้อความ และนำไปตรวจสอบความถูกต้องในขั้นตอนถัดไป

การทดสอบความถูกต้องและแม่นยำของการจำแนกประเภทข้อมูล โครงการฯ มีขั้นตอนดังนี้

ขั้นตอนที่ 1 ทีม Data Scientist สุ่มตัวอย่างจากฐานข้อมูลที่ทำความสะอาดแล้วมา 400 แถว เพื่อส่งให้ผู้เชี่ยวชาญป้ายประเภทข้อมูล (labeling)

ขั้นตอนที่ 2 ทีมผู้เชี่ยวชาญด้านสังคมศาสตร์ จำนวน 2 คน ป้ายข้อมูลประเภทต่างๆ ได้แก่ เข้าข่ายเป็น Malinformation หรือไม่ และเข้าข่ายเป็น Hate Speech หรือไม่ ถ้าใช่ เป็น Hate Speech ในระดับใด โดย

ผู้เชี่ยวชาญ 2 คนทำงานแยกกันแต่ใช้ข้อมูลชุดเดียวกัน เมื่อแต่ละคนประเมินเสร็จแล้ว ผู้เชี่ยวชาญ 2 คนจะทำการถกเถียงในประเด็นที่ทั้ง 2 ประเมินแตกต่างกัน ก่อนที่จะสรุปเป็นข้อมูล 1 ชุดและส่งกลับให้ทีม Data Scientist

ขั้นตอนที่ 3 ทีม Data Scientist นำกลุ่มตัวอย่างชุดเดียวกับที่ส่งให้ผู้เชี่ยวชาญ นำไปให้ GPT Model ในการประเมินประเภทข้อมูล

ขั้นตอนที่ 4 ทีม Data Scientist นำข้อมูลที่ประเมินแล้วจากทั้งผู้เชี่ยวชาญ และจาก GPT Model มาวิเคราะห์หาความถูกต้องและความแม่นยำของ GPT Model ในการจำแนกประเภทข้อมูล โดยใช้ วิธี Confusion Matrix ซึ่งเป็นเครื่องมือที่ใช้ในการประเมินประสิทธิภาพของโมเดลการจำแนกประเภท (Classification Model) ด้วยเครื่องมือนี้สามารถเปรียบเทียบผลลัพธ์ที่โมเดลคาดการณ์กับค่าที่เป็นจริง (True Label) โดยโครงสร้างของ Confusion Matrix ประกอบไปด้วย ตารางที่มีค่าดังต่อไปนี้

True Positive (TP) คือ ข้อความที่ทั้งผู้เชี่ยวชาญและโมเดลประเมินตรงกันในหมวดหมู่ที่ใช้

True Negative (TN) คือ ข้อความที่ทั้งผู้เชี่ยวชาญและโมเดลประเมินตรงกันในหมวดหมู่ที่ไม่ใช่

False Positive (FP) คือ ข้อความที่โมเดลประเมินว่าใช่หมวดหมู่นั้นแต่ผู้เชี่ยวชาญตัดสินใจว่าไม่ใช่

False Negative (FN) คือ ข้อความที่โมเดลประเมินว่าไม่ใช่หมวดหมู่นั้นแต่ผู้เชี่ยวชาญตัดสินใจว่าใช่

		EXPERT	
		ใช่	ไม่ใช่
MODEL	ใช่	TP	FP
	ไม่ใช่	FN	TN

รูป 3 โครงสร้างของ Confusion Matrix

2.4.1 ผลการทดสอบความถูกต้องและความแม่นยำของข้อมูลเลือกตั้ง

ผลการตรวจสอบโดยผู้เชี่ยวชาญ ทำโดยใช้นักวิจัยทางสังคมศาสตร์ของโครงการ จำนวน 2 คน เพื่อจำแนกประเภทตัวอย่างข้อความที่ถูกสุ่มเลือกมา จำนวน 400 ข้อความ ด้วยเกณฑ์เดียวกันที่เป็นเงื่อนไขในการป้อนให้กับ GPT Model ผลการตรวจสอบพบว่า

ในเนื้อหาที่เป็น Hate Speech จากข้อความจำนวน 400 ข้อความ พบว่า มีเนื้อหาที่เกี่ยวข้องกับการเลือกตั้งจำนวน 269 ข้อความ เข้าข่ายเป็นเนื้อหา Hate Speech จำนวน 62 ข้อความ คิดเป็นร้อยละ 23.05 (ตาราง 1)

ตาราง 1 การจำแนกสัดส่วนเนื้อหาที่เป็น Hate Speech โดยผู้เชี่ยวชาญของข้อมูลช่วงเลือกตั้ง

Hate Speech	จำนวน	สัดส่วน
Level 1	33	12.27
Level 2	3	1.12
Level 3	19	7.06
Level 4	7	2.60
Level 5	0	0.00
ไม่เป็น Hate speech	207	76.95
รวม	269	100.00

ในเนื้อหาที่เป็น Malinformation พบว่า จากข้อความที่เกี่ยวข้องกับการเลือกตั้ง จำนวน 269 ข้อความ เข้าข่ายเป็นเนื้อหา Malinformation จำนวน 104 ข้อความ คิดเป็นร้อยละ 38.66 (ตาราง 2)

ตาราง 2 การจำแนกสัดส่วนเนื้อหาที่เป็น Malinformation โดยผู้เชี่ยวชาญของข้อมูลช่วงเลือกตั้ง

Malinformation	จำนวน	สัดส่วน
เข้าข่าย	104	38.66
ไม่เข้าข่าย	165	61.34
รวม	269	100.00

เมื่อเทียบผลการประเมินโดยใช้ Confusion Matrix แล้วมีรายละเอียดดังนี้ (ตาราง 3)

ข้อความประเภท Both คือข้อความที่ผู้เชี่ยวชาญประเมินว่าเป็นทั้ง Hate Speech และ Malinformation พบว่า โมเดลไม่สามารถทำนายได้เลย (TP = 0%) แต่แยกแยะข้อมูลที่ไม่ใช่ "Both" ได้ดี (TN = 93.64%)

Hate Speech Level 1 โมเดลพอเข้าใจบ้าง (TP ≈ 44.44%) และแยกข้อมูลที่ไม่ใช่คลาสนี้ได้ดี (TN = 93.64%)

Hate Speech Level 2 โมเดลล้มเหลวในการทำนาย (TP = 0%) แต่แยกแยะข้อมูลที่ไม่ใช่คลาสนี้ได้ดีมาก (TN = 96.50%)

Hate Speech Level 3 โมเดลเริ่มเข้าใจบ้าง (TP = 20%) และแยกข้อมูลที่ไม่ใช่คลาสนี้ได้ยอดเยี่ยม (TN = 98.75%)

Hate Speech Level 4 และ 5 โมเดลไม่ได้ทำนายคลาสนี้เลยเพราะไม่มีข้อมูลจริง (TP = 0%, TN = 100.00%)

Malinformation โมเดลทำนายได้ปานกลาง (TP ≈ 22.81%) และแยกแยะข้อมูลที่ไม่ใช่คลาสนี้ได้ในระดับปานกลาง (TN = 84.46%)

Normal คือข้อความที่เกี่ยวข้องกับการเมืองแต่ไม่ได้มีการใช้ข้อความอันตราย โมเดลทำงานได้ดีมากในคลาสนี้ (TP ≈ 89.32%) แต่ยังมีข้อผิดพลาดบางส่วนในการแยกข้อมูลที่ไม่ใช่ "Normal" (TN = 72.27%)

ตาราง 3 การวิเคราะห์ความถูกต้องและแม่นยำโดย Confusion Matrix ข้อมูลช่วงเลือกตั้ง

Class	TP		TN	
	N	%	N	%
Both	0	0.00	352	93.64
Hate Speech LV.1	4	44.44	362	93.64
Hate Speech LV.2	0	0.00	385	96.50
Hate Speech LV.3	1	20	391	98.75
Hate Speech LV.4	0	0.00	400	100
Hate Speech LV.5	0	0.00	0	0.00
Malinformation	13	22.81	298	84.46
Normal	251	89.32	86	72.27

การอ่านค่า

TP สูง + TN สูง คือ โมเดลเข้าใจคลาสเป้าหมายและแยกข้อมูลอื่นได้ดี

TP ต่ำ + TN สูง คือ โมเดลมีปัญหาในการเข้าใจคลาสเป้าหมาย

TP สูง + TN ต่ำ คือ โมเดลเข้าใจคลาสเป้าหมาย แต่สับสนกับข้อมูลคลาสอื่น

2.4.2 ผลการทดสอบความถูกต้องและแม่นยำของข้อมูลหลังเลือกตั้ง

ผลการตรวจสอบโดยผู้เชี่ยวชาญ ทำโดยใช้นักวิจัยทางสังคมศาสตร์ของโครงการ จำนวน 2 คน เพื่อจำแนกประเภทตัวอย่างข้อความที่ถูกสุ่มเลือกมา จำนวน 400 ข้อความ ด้วยเกณฑ์เดียวกันที่เป็นเงื่อนไขในการป้อนให้กับ GPT Model ผลการตรวจสอบพบว่า

ในเนื้อหาที่เป็น Hate Speech จากข้อความจำนวน 400 ข้อความ พบว่า มีเนื้อหาที่เกี่ยวข้องกับการเลือกตั้ง จำนวน 192 ข้อความ เข้าข่ายเป็นเนื้อหา Hate Speech จำนวน 122 ข้อความ คิดเป็นร้อยละ 63.54 (ตาราง 4)

ตาราง 4 การจำแนกสัดส่วนเนื้อหาที่เป็น Hate Speech โดยผู้เชี่ยวชาญของข้อมูลช่วงหลังเลือกตั้ง

Hate Speech	จำนวน	สัดส่วน
Level 1	57	29.69
Level 2	6	3.13
Level 3	41	21.35
Level 4	8	4.17
Level 5	10	5.21
ไม่เป็น Hate speech	70	36.46
รวม	192	100.00

ในเนื้อหาที่เป็น Malinformation พบว่า จากข้อความที่เกี่ยวข้องกับการเลือกตั้ง จำนวน 192 ข้อความ เข้าข่ายเป็นเนื้อหา Malinformation จำนวน 87 ข้อความ คิดเป็นร้อยละ 45.31 (ตาราง 5)

ตาราง 5 การจำแนกสัดส่วนเนื้อหาที่เป็น Malinformation โดยผู้เชี่ยวชาญของข้อมูลช่วงหลังเลือกตั้ง

Malinformation	จำนวน	สัดส่วน
เข้าข่าย	87	45.31
ไม่เข้าข่าย	105	54.69
รวม	192	100.00

เมื่อเทียบผลการปะป้ายโดยใช้ Confusion Matrix แล้วมีรายละเอียดดังนี้ (ตาราง 6)

ข้อความประเภท Both คือข้อความที่ผู้เชี่ยวชาญปะป้ายว่าเป็นทั้ง Hate Speech และ Malinformation พบว่า โมเดลไม่สามารถทำนายได้เลย (TP = 0%) แต่สามารถแยกแยะข้อมูลที่ไม่ใช่ "Both" ได้ดีมาก (TN ≈ 98.74%)

Hate Speech Level 1 โมเดลเข้าใจได้ค่อนข้างดี (TP ≈ 72.22%) และสามารถแยกแยะข้อมูลที่ไม่ใช่ "Level 1" ได้ดี (TN ≈ 93.83%)

Hate Speech Level 2 โมเดลล้มเหลวในการทำนาย (TP = 0%) แต่สามารถแยกข้อมูลที่ไม่ใช่ "Level 2" ได้ดีเยี่ยม (TN ≈ 99.75%)

Hate Speech Level 3 โมเดลเข้าใจในระดับหนึ่ง (TP ≈ 33.33%) และแยกแยะข้อมูลที่ไม่ใช่ "Level 3" ได้ดีมาก (TN ≈ 98.74%)

Hate Speech Level 4 และ 5 โมเดลไม่ได้ทำนายคลาสนี้เลยเพราะไม่มีข้อมูลจริง (TP = 0%, TN = 100.00%)

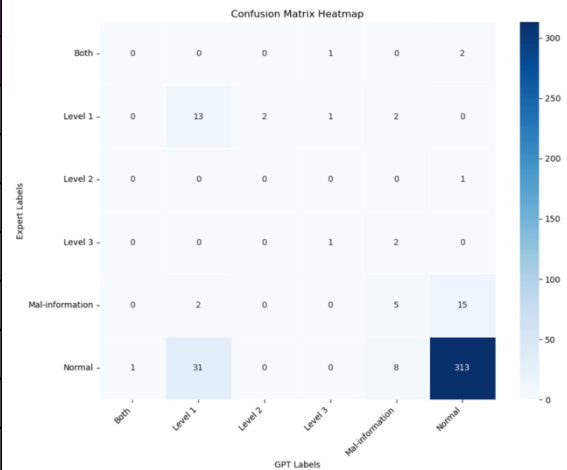
Malinformation โมเดลสามารถทำนายได้บางส่วน (TP ≈ 22.73%) และแยกแยะข้อมูลที่ไม่ใช่ "Malinformation" ได้ในระดับดี (TN ≈ 93.69%)

Normal คือข้อความที่เกี่ยวข้องกับการเมืองแต่ไม่มีการใช้ข้อความอันตราย โมเดลทำงานได้ดีมากในการทำนาย (TP ≈ 88.67%) แต่ยังมีปัญหาในการแยกข้อมูลที่ไม่ใช่ "Normal" (TN ≈ 60.38%)

อย่างไรก็ตามการที่ Accuracy สูงในรอบนี้ไม่ได้หมายความว่าโมเดลดี แต่แสดงให้เห็นว่าโมเดล Bias ต่อคลาสใหญ่ (Normal) อย่างชัดเจนเนื่องจากข้อมูลส่วนใหญ่จะตกไปทาง Normal การปรับสมดุลข้อมูลให้มีปริมาณเท่ากันจะช่วยให้ปัญหานี้และทำให้โมเดลมีความยุติธรรมและแม่นยำมากขึ้นสำหรับทุกคลาส

ตาราง 6 การวิเคราะห์ความถูกต้องและแม่นยำโดย Confusion Matrix ข้อมูลช่วงหลังเลือกตั้ง

Class	TP		TN	
	N	%	N	%
Both	0	0.00	394	98.74
Hate Speech LV.1	13	72.22	349	93.83
Hate Speech LV.2	0	0.00	398	99.75
Hate Speech LV.3	1	33.33	394	98.74
Hate Speech LV.4	0	0.00	0	0.00
Hate Speech LV.5	0	0.00	0	0.00
Malinformation	5	22.73	358	93.69
Normal	313	88.67	32	60.38



การอ่านค่า

TP สูง + TN สูง คือ โมเดลเข้าใจคลาสเป้าหมายและแยกข้อมูลอื่นได้ดี

TP ต่ำ + TN สูง คือ โมเดลมีปัญหาในการเข้าใจคลาสเป้าหมาย

TP สูง + TN ต่ำ คือ โมเดลเข้าใจคลาสเป้าหมาย แต่สับสนกับข้อมูลคลาสอื่น

2.5 การวิเคราะห์ข้อมูล

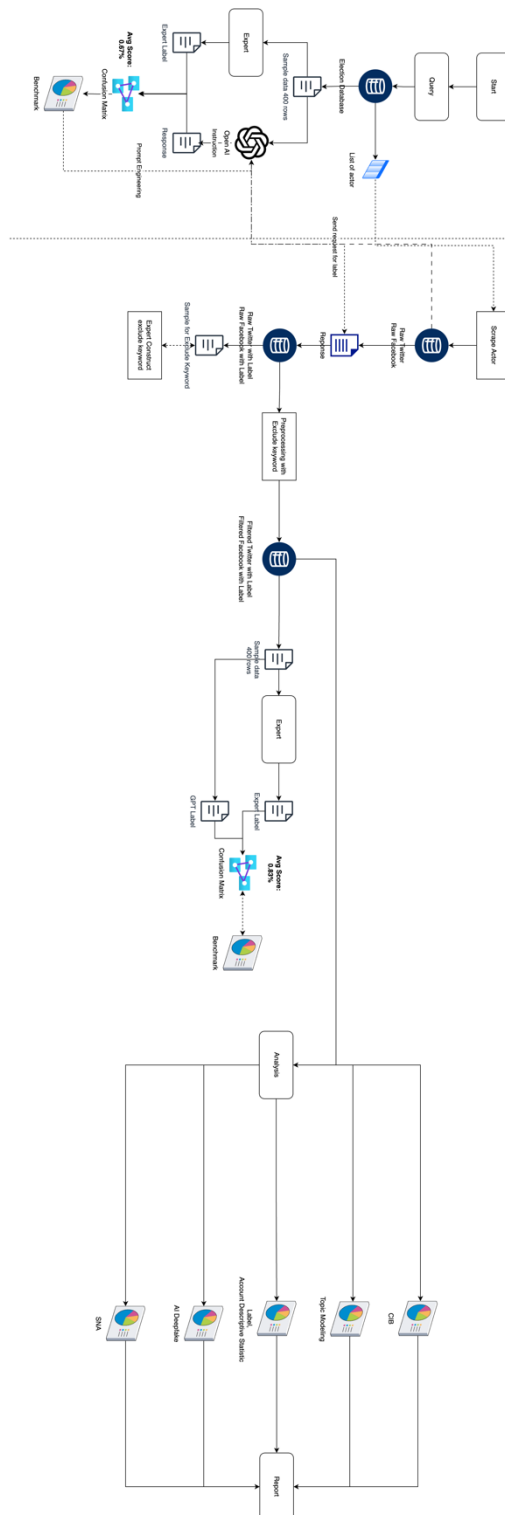
การวิเคราะห์ข้อมูลการศึกษาค้นคว้าครั้งนี้ โครงการฯ ใช้แนวทางการวิเคราะห์แก่นสาระ (thematic analysis) จากข้อมูลความคิดเห็นในสื่อสังคมออนไลน์ในช่วงการเลือกตั้งและหลังการเลือกตั้ง โดยแบ่งเป็น 2 ประเด็นใหญ่ ได้แก่ 1) Malinformation และ 2) Hate Speech ทั้ง 5 ระดับ

นอกจากนั้น โครงการฯ ได้ใช้การวิเคราะห์หัวข้อด้วยอัลกอริทึมการจำลองหัวข้อ (Topic Modelling Algorithms) เพื่อวิเคราะห์หัวข้อที่ปรากฏขึ้น (Emerging Topics) ซึ่งเป็นหนึ่งในเทคนิค Natural Language Processing (NLP) ที่ใช้ค้นหาและจัดกลุ่ม "Topic" ที่ซ่อนอยู่ในข้อความจำนวนมากโดยอัตโนมัติ เทคนิคนี้ทำงานโดยแปลงข้อความให้อยู่ในรูปแบบตัวเลข และใช้โมเดลทางสถิติเพื่อวิเคราะห์ความสัมพันธ์ระหว่างคำและหัวข้อในเอกสาร โมเดลที่ใช้วิเคราะห์เพื่อทำ Topic Modeling ในโครงการฯ นี้ ได้แก่ GSDMM (Gibbs Sampling Dirichlet Mixture Model) ซึ่งเหมาะสำหรับข้อความสั้น เช่น Tweets หรือข้อความโซเชียลมีเดีย

GSDMM ทำงานโดยใช้วิธี Gibbs Sampling ซึ่งเป็นกระบวนการสุ่มที่ปรับปรุงความน่าจะเป็นในการจัดกลุ่มข้อความ โดยสมมติว่าแต่ละข้อความเกี่ยวข้องกับหัวข้อเดียว โมเดลจะคำนวณว่าข้อความแต่ละข้อความควรถูกจัดกลุ่มในหัวข้อใดโดยพิจารณาจากการเพิ่มความคล้ายคลึงภายในกลุ่ม (intra-cluster similarity) และลดความแตกต่างระหว่างกลุ่ม (inter-cluster dissimilarity) กระบวนการนี้จะทำซ้ำจนกว่าการจัดกลุ่มจะมีเสถียรภาพ (convergence) อย่างไรก็ตามการจัดกลุ่มคำเป็นผลจากการวิเคราะห์ความสัมพันธ์ซึ่งกันและกันที่สำคัญไม่ใช่การสุ่ม

สำหรับการวิเคราะห์ AI Deepfake มีหลักการคล้ายกับการใช้ Model GPT ในการวิเคราะห์ข้อความ แต่ในกรณีนี้จะวิเคราะห์ข้อมูลในรูปแบบรูปภาพ โดยการนำภาพเข้าสู่ Vision Model (Transformer Model) ซึ่งเป็นโมเดลที่ออกแบบมาเพื่อจัดการกับข้อมูลภาพโดยเฉพาะ สำหรับโมเดลที่เลือกใช้ในครั้งนี้ ได้รับการฝึกฝนมาเพื่อจำแนกภาพที่ถูกสร้างขึ้นหรือ แก้ไขโดย AI (AI Generated/AI Manipulated) หรือภาพ AI Deepfake โดยหลักการทำงานของ Vision Model คือ โมเดลจะทำการเรียนรู้และเข้าใจคุณลักษณะ (Attributes) ของภาพที่มีพบใน Deepfake เช่น ความผิดปกติของแสงและเงา ขอบใบหน้า หรือรายละเอียดเล็กน้อยที่อาจไม่สอดคล้องกับความเป็นจริง ซึ่งเป็นการตรวจจับลักษณะเฉพาะของภาพที่สร้างโดย AI เมื่อเปรียบเทียบกับภาพจริง

2.6 กรอบการศึกษา (Framework)



3. ผลการดำเนินงาน

3.1 กิจกรรมที่ 1 ทบทวนเอกสารนโยบายด้านความปลอดภัยของแพลตฟอร์ม

โครงการฯ ได้รวบรวมและทบทวนเอกสารมาตรฐานชุมชนที่จัดทำโดย Meta (Facebook), X และ Google (YouTube) ดังรายละเอียดต่อไปนี้

3.1.1 นโยบายและมาตรฐานความปลอดภัยของแพลตฟอร์ม

โครงการฯ ได้รวบรวมและทบทวนเอกสารมาตรฐานชุมชนของแต่ละแพลตฟอร์มที่เกี่ยวข้องกับการเลือกตั้งและสถานการณ์ที่มีความอ่อนไหวทางการเมือง บริบทแวดล้อมที่อาจเกี่ยวข้อง นิยามความหมายและสิ่งที่แพลตฟอร์มจะดำเนินการเมื่อเกิดการละเมิดนโยบายและมาตรฐานความปลอดภัย ในหัวข้อต่อไปนี้

ตาราง 7 นโยบายและมาตรฐานความปลอดภัยของแพลตฟอร์มที่เกี่ยวข้องกับการเลือกตั้งและสถานการณ์ที่มีความอ่อนไหวทางการเมือง

Meta (Facebook และ Instagram)	X	Google (YouTube)
• Fact-Check label • Partner with third-party fact checkers • Identify false news • Combating Misinformation* • Related Community Standards* - Coordinating harm and promoting crime - Dangerous organizations and individuals - Fraud, scams and deceptive practices - Violence and incitement - Bullying and harassment - Hate speech - Privacy violations	• Overall measures • Misleading information about how to participate • Suppression • Intimidation • False or misleading affiliation • Exception	• Voter suppression • Candidate eligibility • Incitement to interfere with democratic processes • Election integrity

* มีเฉพาะ Facebook

ดูตารางเปรียบเทียบแนวปฏิบัติของแพลตฟอร์มในการกลั่นกรองเนื้อหาที่เป็นอันตรายและข้อมูลบิดเบือน⁴

3.1.2 การวิเคราะห์กลไกการดำเนินงานและเปรียบเทียบวิธีการของแพลตฟอร์มในการจัดการกับปัญหา

โครงการฯ ได้รวบรวมและทบทวนเอกสารมาตรฐานชุมชนที่จัดทำโดย Meta X และ Google แล้ววิเคราะห์กลไกการดำเนินงาน เช่น การแจ้งเตือนและการลบเนื้อหา ขั้นตอนการอุทธรณ์ นโยบายการเซ็นเซอร์ และระบบการรายงานพบว่า กฎชุมชน หรือ Community Standards คือกฎที่ผู้พัฒนาแพลตฟอร์มโซเชียลต่างๆ บังคับให้ผู้ใช้ปฏิบัติตามเพื่อความเป็นระเบียบเรียบร้อย โดยแต่ละแพลตฟอร์มจะมีกฎชุมชนที่แตกต่างกัน แต่ที่เหมือนกันคือมุ่งไปที่การให้สิทธิและเสรีภาพในการแสดงออกทางความคิดเห็นตามหลักประชาธิปไตย โดยไม่ละเมิดสิทธิของผู้อื่น แต่ความเข้มข้นในรายละเอียดและการบังคับใช้อาจจะแตกต่างกันไปตามนโยบายของบริษัทเจ้าของแพลตฟอร์มนั้นๆ

⁴ ดูรายละเอียดในภาคผนวก

ตารางรายละเอียดกฎหมาย (ภาคผนวก ก) ออกแบบมาเพื่อให้ทีมผู้วิจัยในโครงการนี้ใช้ในการเปรียบเทียบและอ้างอิงว่า เนื้อหาที่นำมาใช้ในการศึกษาเข้าข่ายผิดตามกฎหมายชุมชนข้อใดบ้าง แต่ละแพลตฟอร์มมีการลงโทษเจ้าของโพสต์ด้วยวิธีใด มีลักษณะการลงโทษที่เข้มข้นแตกต่างกันอย่างไร เพื่อที่ผู้วิจัยจะได้ใช้ในการวิเคราะห์และพิจารณาว่า แพลตฟอร์มต่างๆ ให้ความสำคัญในประเด็นที่มีความประปรายในสังคมไทยมากน้อยเพียงใด

โครงการฯ แบ่งรายละเอียดไปตามแพลตฟอร์มต่างๆ ได้แก่ Meta X และ Google โดยหยิบเฉพาะกฎหมายที่เกี่ยวกับข่าวลวงและ Hate Speech ในช่วงเลือกตั้งปี 2566 และหลังการเลือกตั้ง ซึ่งพบว่า Meta และ Google มีกฎหมายที่เกี่ยวกับการบิดเบือนเนื้อหาที่เกี่ยวกับการเลือกตั้งชัดเจนมากที่สุด โดยเฉพาะ YouTube มีการระบุอย่างชัดเจนถึงมาตรการลงโทษผู้ที่โพสต์เนื้อหาบิดเบือน สร้างความเข้าใจผิดเกี่ยวกับการเลือกตั้ง หากผิดตามกฎหมายจริงจะมีมาตรการลงโทษเริ่มจากการเตือน ปิดช่องชั่วคราว ไปจนถึงการปิดช่องถาวร

ส่วน Meta ผู้ที่เป็นเจ้าของแพลตฟอร์ม Facebook และ Instagram มีกฎหมายในลักษณะคล้ายคลึงกัน แต่แนวทางการปฏิบัติจะแตกต่างกันออกไป ขึ้นอยู่กับประเภทของเนื้อหา

กฎหมายของ Meta อัปเดตล่าสุดเมื่อวันที่ 18 กรกฎาคม 2024 แบ่งมาตรการการจัดการกับเนื้อหาบิดเบือนเกี่ยวกับการเลือกตั้งออกเป็น 3 รูปแบบ

- 1) การติดฉลากกับเนื้อหาทางการเมืองที่มีความบิดเบือน หากผู้เชี่ยวชาญตรวจสอบพบว่ามีเนื้อหาบิดเบือนเนื้อหา เช่น ข่าวลวง หรือการใส่ร้ายป้ายสี ซึ่งตรวจสอบโดยนักตรวจสอบข้อมูล (fact-checker) จะมีการติดฉลากเพื่อเตือนผู้อ่านว่ามีเนื้อหาบิดเบือน แต่จะไม่ลบเนื้อหานั้นออกจากระบบโดยทันที หากไม่ถึงขั้นผิดกฎหมายอย่างร้ายแรง
- 2) Facebook และ Instagram จับมือกับพันธมิตร ประกอบด้วยสำนักข่าวและหน่วยงานที่มีความเชี่ยวชาญในการตรวจสอบข่าวลวงกว่า 100 แห่งทั่วโลก ในการช่วยตรวจสอบข้อมูลข่าวลวงต่างๆ โดยพวกเขาสามารถแจ้งแพลตฟอร์มให้มีการติดฉลากกับเนื้อหาเหล่านั้นหากตรวจพบว่ามีเนื้อหาเสนอข้อมูลเท็จ แต่จะไม่สามารถลบเนื้อหาได้ทันที ซึ่งการจะลบเนื้อหาใดๆ ก็ตาม จะขึ้นอยู่กับดุลพินิจของผู้ตรวจสอบเนื้อหาภายในของ Meta
- 3) เนื้อหาเกี่ยวกับการเมืองที่มีความผิดร้ายแรง เมื่อพบจะถูกลบออกจากระบบ จะต้องเป็นเนื้อหาที่นำไปสู่ความรุนแรง เช่น 1. การสนับสนุนให้ทำร้ายร่างกาย หรือปลุกระดมให้ผู้อ่านไปทำร้ายร่างกายหรือทรัพย์สินของบุคคลใดบุคคลหนึ่ง 2. เนื้อหาข่าวลวงที่เกี่ยวข้องกับสุขภาพที่อาจนำไปสู่การบริโภคสินค้า หรือผลิตภัณฑ์ที่ส่งผลเสียต่อร่างกาย เช่น การรณรงค์ไม่ให้ไปฉีดวัคซีนทั้งๆ ที่มีผลการศึกษายืนยันชัดเจนว่าวัคซีนสามารถช่วยป้องกันการแพร่ระบาดของเชื้อโรคได้ และ 3. การนำเสนอข้อมูลบิดเบือนเกี่ยวกับวัน เวลา และสถานที่ตั้งของคูหาเลือกตั้ง รวมถึงการบิดเบือนข้อมูลของผู้สมัคร เช่น หมายเลขประจำตัวผู้สมัคร หรือการกล่าวหาว่าผู้สมัครผิดไปจากกฎหมายของคณะกรรมการการเลือกตั้ง เป็นต้น

โครงการฯ ตั้งข้อสังเกตว่ามาตรการของ Meta จะให้ความสำคัญกับเนื้อหาที่อาจนำไปสู่ความรุนแรงต่อร่างกาย เช่น การข่มขู่ทำร้ายร่างกาย การข่มขู่เอาชีวิต หรือเนื้อหาหลอกลวงที่นำไปสู่ความอันตรายต่อชีวิตมาเป็นอันดับแรก เนื้อหาเหล่านี้จะถูกลบออกโดยทันที แต่เนื้อหาที่เข้าข่ายการโฆษณาชวนเชื่อทางการเมือง หรือเนื้อหา Hate Speech ที่มีความคิดเห็นทางการเมืองสุดแทรกอยู่ อาจจะมีการติดฉลากว่าเป็นข้อมูลบิดเบือน แต่จะไม่ถูกลบโดยทันที ผู้ถูกกล่าวหาสามารถใช้ช่องทางรายงานเพื่อให้ทีมงานของ Meta ตรวจสอบว่าผิดกฎหมายข้อใดบ้าง ซึ่งบางครั้งอาจกินเวลานาน ทำให้ข้อความหรือเนื้อหาบิดเบือนดังกล่าวถูกแพร่กระจายไปเป็นวงกว้าง ยากต่อการควบคุม

ด้าน X มีการกำหนดมาตรการบังคับใช้กับข้อมูลบิดเบือนทางการเมือง ไม่ว่าจะเป็นการบิดเบือนวันเลือกตั้ง ข้อมูลเกี่ยวกับผู้สมัคร รวมถึงเนื้อหาที่อาจนำไปสู่ความรุนแรงเช่นกัน แต่จะเป็นการติดฉลากกับเนื้อหาหรือปรับอัลกอริทึมให้ค้นหาเนื้อหาเหล่านั้นได้ยากขึ้นแทนที่จะเป็นมาตรการลบเนื้อหานั้นออกจากระบบ

เป็นที่น่าสังเกตว่า นับตั้งแต่อัลลอน มัสก์ เข้ามาซื้อกิจการและขึ้นเป็นผู้บริหารของ X มีการปรับนโยบาย โดยเฉพาะกฎหมาย โดยเน้นเสรีภาพด้านการแสดงออกทางความเห็น ลดการบล็อกเนื้อหา โดยเปลี่ยนเป็นการปรับอัลกอริทึมให้แสดงเนื้อหาบิดเบือนยากขึ้น แต่ไม่ได้หมายความว่าเนื้อหาเหล่านั้นหายไป ดังนั้นหากผู้ใช้งานตั้งใจค้นหา

เนื้อหาเหล่านั้น หรือมีความสนใจในข่าวการเมืองขั้วใดเป็นพิเศษ ก็อาจจะทำให้เนื้อหาบิดเบือนปรากฏขึ้นบนหน้าฟีดได้ถึงแม้เนื้อหาเหล่านั้นอาจจะถูกตัดสินจากว่าเป็นเนื้อหาบิดเบือนก็ตาม

จากการเปรียบเทียบกฎชุมชนของทั้ง 3 แพลตฟอร์ม จะเห็นว่า Meta และ Google ให้ความสำคัญเกี่ยวกับการควบคุมและปิดกั้นเนื้อหาข่าวลวงและเนื้อหาบิดเบือนเกี่ยวกับการเลือกตั้งมากที่สุด โดยมีมาตรการจากเขาไปหาหนักขึ้นอยู่กับความรุนแรงของเนื้อหา ส่วน X มีมาตรการที่อ่อนที่สุด โดยใช้การตัดสินจากกับเนื้อหาบิดเบือนหรือการปรับอัลกอริทึมให้แสดงเนื้อหาเหล่านี้ยากขึ้น แทนที่จะเป็นการลบเนื้อหาเหล่านั้นออกไป

นอกจากนี้ทุกแพลตฟอร์มที่กล่าวมา ล้วนมีช่องทางร้องเรียนในกรณีที่ผู้โพสต์เนื้อหาสามารถยืนยันได้ว่าข้อมูลที่โพสต์นั้นถูกต้อง ไม่ผิดกฎหมาย หรือข้อมูลบิดเบือนที่อาจจะเคยถูกลบไปก่อนหน้านี้ แต่เมื่อพ้นช่วงที่มีการเลือกตั้งไปแล้ว ก็มีความเป็นไปได้ที่เนื้อหาที่อาจจะถูกลบหรือลือกมาก่อนอาจจะกลับเข้าสู่ระบบอีกครั้ง

อย่างไรก็ตาม กฎชุมชนทั้ง 3 แพลตฟอร์มที่กล่าวมา ล้วนเป็นกฎชุมชนที่ใช้กับผู้ใช้งานทั่วโลกและไม่เกี่ยวข้องกันเนื้อหาที่อาจเข้าข่ายกระทำผิดทางกฎหมายในแต่ละประเทศ

3.1.3 กรณีศึกษาที่สำคัญที่เกิดขึ้นช่วงการเลือกตั้ง และหลังการเลือกตั้ง

ช่วงเลือกตั้ง⁵

- 1) กรณีศึกษาที่ 1 Malinformation โจมตีคณะกรรมการการเลือกตั้ง (กกต.) พบว่า มีการเชื่อมโยงข้อมูลแล้วตีความใหม่โดยสื่อเจตนาให้เกิดความเข้าใจผิด ทำลายความน่าเชื่อถือในการจัดกระบวนการการเลือกตั้งของ กกต. อีกทั้งยังมีการสร้างข่าวลือ (rumor) เผยแพร่ทฤษฎีสมคบคิดกล่าวหาว่ากระบวนการทำงานของ กกต. ถูกต่างชาติเข้าแทรกแซง

จากข้อมูลกรณีศึกษาที่ยกมานี้ พบว่า โพสต์ 1 จากแพลตฟอร์ม Facebook และโพสต์ที่ 3 จาก X ยังสามารถค้นหาได้อยู่โดยไม่ได้อัปเดตจากใดๆ อาจเนื่องมาจากแพลตฟอร์มประเมินว่าข้อความดังกล่าวไม่ได้ละเมิดมาตรฐานชุมชน ในขณะที่โพสต์ที่ 2 จาก X ไม่สามารถค้นหาได้เนื่องมาจากผู้ใช้งานตั้งคำบัญชีเป็นส่วนตัว หรือลบเนื้อหา หรือลบบัญชีผู้ใช้งาน



รูป 4 กรณีศึกษาที่ 1 ช่วงเลือกตั้ง

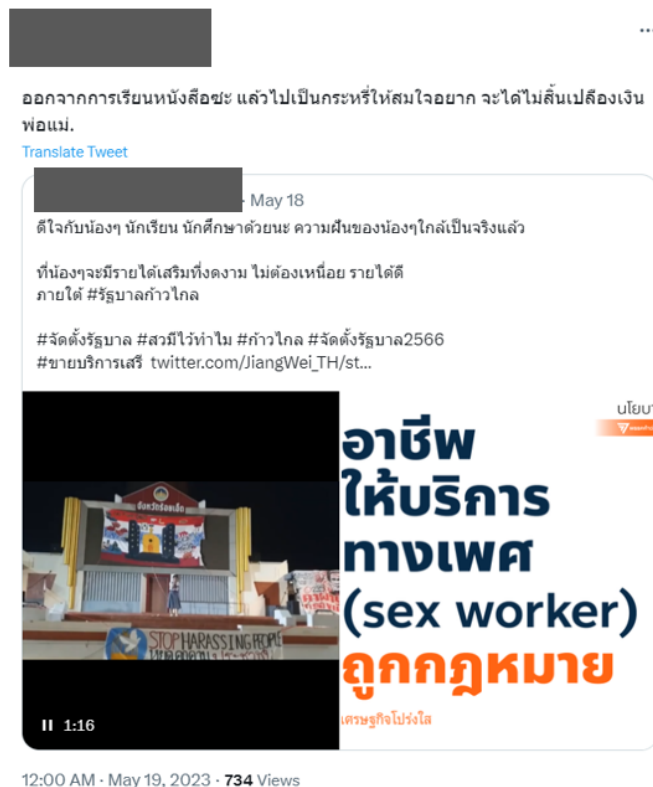
⁵ ข้อมูลจากรายงาน “ถอดรหัสปฏิบัติการบนข้อมูลข่าวสารในไทย #เลือกตั้ง2566” และฐานข้อมูลของโครงการติดตามและวิเคราะห์กระบวนการเลือกตั้งในโซเชียลมีเดีย (Digital Election Analytic Lab: DEAL)

- 2) กรณีศึกษาที่ 2 ข้อมูลที่บอกหมายเลขผู้สมัครแคนดิเดตนายกรัฐมนตรีและพรรคการเมืองผิด โดยมีการใช้รูป ตัดต่อข้อความและเบอร์พรรคการเมือง พบว่า ยังสามารถค้นหารูปที่ 1 ได้ใน X



รูป 5 กรณีศึกษาที่ 2 ช่วงเลือกตั้ง

- 3) กรณีศึกษาที่ 3 Malinformation โจมตีนโยบายของพรรคการเมือง ด้วยการโควท (quote) หรือ re-post ข้อมูลนโยบายของพรรคการเมืองซึ่งเป็นข้อมูลจริง แล้วนำไปผลิตเนื้อหาที่ตีความใหม่เพื่อสร้างความเข้าใจผิดในนโยบายของพรรคการเมือง ดังเช่นกรณีที่มีการตีความว่าการทำอาชีพให้บริการทางเพศ (sex worker) ถูกกฎหมายเท่ากับการส่งเสริมให้คนทำอาชีพนี้ เป็นต้น จากข้อมูลกรณีศึกษาจาก X นี้ พบว่ายังคงค้นหาได้อยู่ และไม่มีการติดฉลากใด



รูป 6 กรณีศึกษาที่ 3 ช่วงเลือกตั้ง

- 4) กรณีศึกษาที่ 4 Hate speech และ Malinformation โจมตีการหาเสียงของพรรคการเมือง ทั้งโพสต์ที่ 1 และ 2 จากผู้ใช้งาน X คนเดียวกัน มีการใช้ hate speech ในหลายระดับ ทั้งระดับที่ 1 การด่าด้วยถ้อยคำหยาบคาย ดูหมิ่นสติปัญญา “หน้าโง่” “อับริยจัญไร” ระดับที่ 3 การเหมารวมแปะป้าย “คนบุรีรัมย์” “คนใต้” เป็นต้น รวมถึงนำเอาบางช่วงบางตอนของคลิปวิดีโอลงพื้นที่หาเสียงของนักการเมืองมาตัดให้สั้นลง แล้วโพสต์โดยใส่ความคิดเห็นเจตนาทำลายชื่อเสียงของบุคคล โจมตีพรรคการเมือง ทั้งนี้ ปัจจุบันยังสามารถค้นหาโพสต์ที่ 1 เจอได้ใน X (ไม่มีการติดฉลาก) แต่ไม่พบโพสต์ที่ 2 อาจเนื่องจากผู้ใช้งานลบเนื้อหาออกด้วยตัวเอง



รูป 7 กรณีศึกษาที่ 4 ช่วงเลือกตั้ง

- 5) กรณีศึกษาที่ 5 Hate speech โจมตีผู้เห็นต่างทางการเมืองจาก X โพสต์ที่ 1-3 มีการใช้ hate speech ในหลายระดับ ทั้งระดับที่ 1 ด่าด้วยถ้อยคำหยาบคาย ดูหมิ่นสติปัญญา “โง่เหมือนควาย” ถูกเหยียดหยาม “หมารับใช้” ระดับที่ 3 การทำให้เป็นปศุสัตว์ “สัตว์นรก” การแปะป้ายเหมารวม “แกว่นคร” “ตังส้ม” “ควายแดง” โพสต์ที่ 1 มีข้อความอันตรายระดับที่ 4 แสดงการปฏิเสธที่จะอยู่ร่วมกันในสังคม “ขอให้...จับหายตายโหง...ถ้าไม่ตายก็ขอมันให้พิการไปตลอดชีวิต” ส่วนโพสต์ที่ 5 เป็นข้อความอันตรายระดับที่ 2 และ 5 คือ การข่มขู่ เรียกร้องให้เปิดเผยข้อมูลส่วนตัว และสนับสนุนให้เกิดความรุนแรง “เอาให้ตาย รบกว่นหมา” เมื่อตรวจสอบพบว่า มีเพียงโพสต์ที่ 1 ที่ยังสามารถค้นหาได้อยู่ แต่ไม่มีการติดฉลากใด ส่วนโพสต์ที่ 2 3 และ 5 ค้นหาไม่พบ อาจจะถูกลบไปโดยผู้ใช้งาน หรือผู้ใช้งานตั้งค่าเป็นบัญชีส่วนตัว หรือผู้ใช้งานลบบัญชี อย่างไรก็ตาม มีข้อสังเกตว่าโพสต์ที่ 5 ถูกลบไปตั้งแต่ช่วงเลือกตั้ง อาจเนื่องมาจากการลบโดยระบบเพราะละเมิดมาตรฐานชุมชน เนื่องด้วยช่วงนั้นยังเป็นช่วงที่ X มีมาตรฐานชุมชนที่เข้มข้นกว่าปัจจุบัน



1 ...

ดีกว่าสัตว์รบกวนนครฯ หารับใช้รับจ้างป่วน ชวนคนเผาเมืองครับ ดีแต่ตลก
กะโหลกทะลุบาตสองบาท ขอให้หมาตัวนั้น ถีบหาย ดายโหงในเร็ววัน จะตายแบบ
ไหนก็ได้ ถ้าไม่ตายก็ให้มันพิการไปตลอดชีวิต ให้มันสมกับความชั่วที่มันทำไว้กับ
ประเทศไทย คุณเห็นด้วยไหม ในฐานะผมคนสุราษฎร์ ใกล้เคียง กัน

12:28 PM · Apr 21, 2023 · 612 Views



2 ...

ตั้งส้มโง่เหมือนควายจะเล่นประเพณีพรรคการเมืองทำโพลได้หรือ อีดอกเค้าก็ทำกัน
ทุกพรรคแต่ไม่ได้ทำในลักษณะเจาะใจ หรือเอาผลโพลมาแสดงให้ดูเป็นการขึ้นำ พรรค
มึงก็มีค่า
แค่ทำกันเองโป๊ๆเบ๊ๆ

1:55 PM · Apr 21, 2023 · 3,963 Views



3 ...

ควายแดง โง่เหมือนเดิม เหมือนตอนไปเชียร์ฮิปโป ขึ้นราคา ก๊าซ lpg

12:26 PM · Apr 20, 2023 · 445 Views

[เอาให้ถึงตาย รบกวณหมา](#)

4

รูป 8 กรณีศึกษาที่ 5 ช่วงเลือกตั้ง

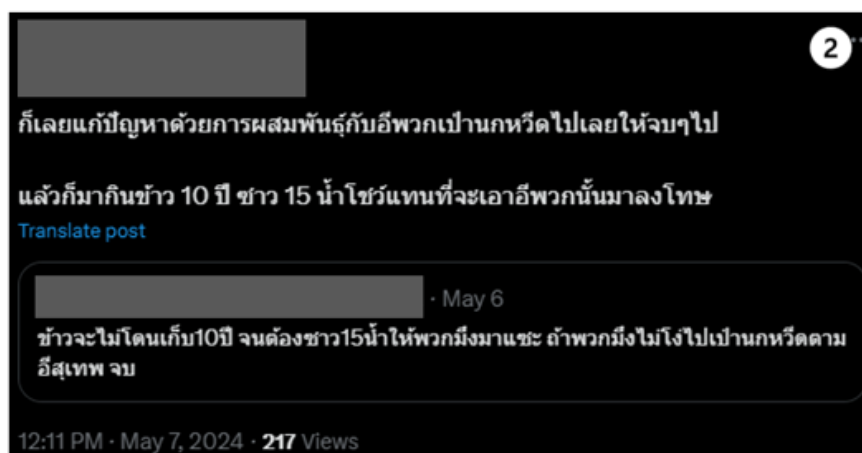
หลังเลือกตั้ง

- กรณีศึกษาที่ 1 Hate speech และ Malinformation โจมตีการจัดตั้งรัฐบาล ข้อมูลจาก X โพสต์ที่ 1-3 มีการใช้ข้อความอันตราย ระดับที่ 1 ดำด้วยถ้อยคำหยาบคาย ถูกเหยียดหยาม เช่น “ตอแหล” “เสียตัก” “พสมพันธุ์” “ตระกูลสัตว์” และข้อความทั้งหมดเป็น malinformation ใช้ข้อมูลจริงมาบิดเบือนโดยมีเจตนาโจมตีการจัดตั้งรัฐบาลของพรรคเพื่อไทย กล่าวหาว่าการจัดตั้งรัฐบาลของพรรคเพื่อไทยเป็นการบั่นทอนประชาธิปไตยของประเทศ หรือกล่าวหาว่าพรรคเพื่อไทยไม่มีความชอบธรรมในการเป็นแกนนำจัดตั้งรัฐบาล ข้อมูลทั้งหมดนี้ยังสามารถค้นหาได้ในแพลตฟอร์ม และไม่มีการติดฉลากใดๆ อาจเนื่องมาจากแพลตฟอร์มมองว่าเป็นเสรีภาพในการแสดงออก ไม่ได้ละเมิดกฎของชุมชน



รูป 9 กรณีศึกษาที่ 1 ช่วงหลังเลือกตั้ง

- 2) กรณีศึกษาที่ 2 Hate speech และ Malinformation โจมตีนโยบายของรัฐบาลใน X โดยโพสต์ที่ 1 และ 2 มีการใช้ข้อความอันตรายระดับที่ 1 การดูถูกเหยียดหยาม ดูหมิ่นสติปัญญา “ความรู้ต่ำ” “งูงมูก” “ผสมพันธุ์” มีการใช้ malinformation โจมตีนโยบายขึ้นค่าแรง (โพสต์ที่ 1) และกระกรว่งพาณิขย์เปิดประมูลข้าวในโครงการรับจำนำข้าว (โพสต์ที่ 2) ซึ่งกรณีการเปิดประมูลข้าวๆ นี้กลายเป็นวิวาทะที่ร้อนแรงระหว่างฝ่ายผู้สนับสนุนและฝ่ายผู้เห็นต่างจากรัฐบาลทั้งในโลกออนไลน์และออฟไลน์ ภายใต้วลี “ข้าว/เมา” รวมถึงโจมตีการผลักดันเมกะโปรเจกต์แลนด์บริดจ์ของรัฐบาล (โพสต์ที่ 3) ข้อมูลทั้งหมดนี้ยังสามารถค้นหาได้ในแพลตฟอร์ม X และไม่มีการติดฉลากใดๆ



รูป 10 กรณีศึกษาที่ 2 ช่วงหลังเลือกตั้ง

- 3) กรณีศึกษาที่ 3 Hate speech และ malinformation จาก X โจมตีนายเศรษฐา ทวีสิน อดีตนายกรัฐมนตรี และนางสาวแพทองธาร ชินวัตร นายกรัฐมนตรี จากพรรคเพื่อไทย โดยโพสต์ที่ 1 และ 3 เป็น malinformation โจมตีคุณสมบัติและความน่าเชื่อถือของนายกรัฐมนตรี ส่วนโพสต์ที่ 2 มีการใช้ข้อความอันตราย ระดับที่ 3 ลดทอนคุณค่านักการเมืองว่าเป็นคนเลว ชั่ว ระดับที่ 4 มีการปฏิเสธที่จะอยู่ร่วมกัน “ต้องรอไอ้เหี้ย ชั้น 14 มันทาย” เพื่อโจมตีนางสาวแพทองธาร ชินวัตร ซึ่งข้อความทั้งหมดยังสามารถค้นหาได้ และไม่มีการติดฉลากใดๆ



รูป 11 กรณีศึกษาที่ 3 ช่วงหลังเลือกตั้ง

- 4) กรณีศึกษาที่ 4 Hate speech และ malinformation เกี่ยวข้องกับนายทักษิณ ชินวัตร อดีตนายกรัฐมนตรีนี เนื่องจากช่วงเวลาเก็บข้อมูลหลังเลือกตั้งมีสถานการณ์ที่เกี่ยวข้องกับนายทักษิณ ชินวัตร เป็นระยะ ได้แก่ การพักโทษ และได้รับอภัยโทษ และอื่นๆ จึงทำให้พบบทสนทนาที่เกี่ยวข้องเป็นจำนวนมาก โดยโพสต์ที่พบใน X มีทั้งที่เป็นข้อความอันตราย ระดับที่ 1 (โพสต์ที่ 3) 3 (โพสต์ที่ 2) และ 5 (โพสต์ที่ 1) ผสมกับข้อความที่เป็น malinformation ดังเช่นโพสต์ที่ 4 ทั้งนี้ ยังคงสามารถค้นหาข้อความได้ และไม่พบว่ามี การติดฉลากใด



รูป 12 กรณีศึกษาที่ 4 ช่วงหลังเลือกตั้ง

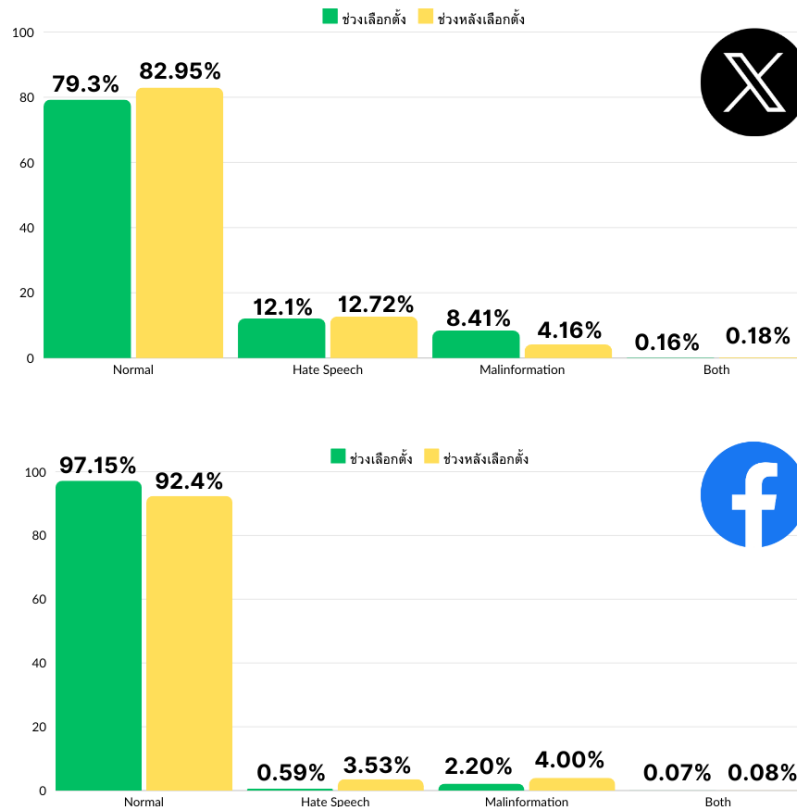
3.2 กิจกรรมที่ 2 การวิเคราะห์โซเชียลมีเดีย (Social Media Analysis)

ในกิจกรรมที่ 2 โครงการฯ ได้ดำเนินการวิเคราะห์ข้อมูลสื่อสังคมออนไลน์ที่เกี่ยวข้องกับการเลือกตั้ง *ย้อนหลัง* เพื่อทำความเข้าใจภูมิทัศน์ของบทสนทนาออนไลน์ในช่วงการเลือกตั้ง (เมษายน - มิถุนายน 2566) รวมทั้งเก็บข้อมูลเพิ่มเติมและวิเคราะห์สภาพแวดล้อมออนไลน์ในปัจจุบัน (มกราคม – สิงหาคม 2567) เมื่อวิเคราะห์ข้อมูลตามแนวทางการวิเคราะห์แก่นสาระ (thematic analysis) และใช้การวิเคราะห์หัวข้อด้วยอัลกอริทึมการจำลองหัวข้อ (Topic Modelling Algorithms) เพื่อวิเคราะห์หัวข้อที่ปรากฏขึ้น (Emerging Topics) ผลการดำเนินงานพบว่า

3.2.1 การจำแนกประเภทข้อความตามแก่นสาระ

โครงการฯ ได้กำหนดแก่นสาระของงานศึกษารังนี้ แบ่งเป็น 2 ประเด็นใหญ่ ได้แก่ 1) Malinformation และ 2) Hate Speech ทั้ง 5 ระดับ ได้แก่ ระดับที่ 1 การด่าด้วยถ้อยคำหยาบคาย การดูหมิ่นสติปัญญา การดูถูกเหยียดหยาม และการสร้างความเป็นอื่น ระดับที่ 2 การคุกคาม การข่มขู่ และการเปิดเผยข้อมูลส่วนตัว การเรียกร้องให้มีการบังคับใช้กฎหมายเพื่อบั่นทอนเสรีภาพในการแสดงออก ระดับที่ 3 การเหมารวม การปะปน การลดทอนคุณค่า การดูถูกเหยียดหยาม การทำให้เป็นปศุสัตว์ ระดับที่ 4 การแบ่งแยก กีดกัน การปฏิเสธที่จะอยู่ร่วมกันในสังคม ระดับที่ 5 การยุยง การสนับสนุน การเรียกร้องให้เกิดความรุนแรง การยกย่องผู้ก่อความรุนแรง การให้ความชอบธรรมกับการก่อความรุนแรง ผลการศึกษาพบว่า

1) ผลการวิเคราะห์ข้อมูลในภาพรวม



รูป 13 ผลการจำแนกประเภทข้อความตามเกณฑ์สาระในภาพรวม

ข้อมูลจาก Facebook เมื่อเปรียบเทียบข้อมูลก่อนและหลังการเลือกตั้ง พบแนวโน้มที่น่าสนใจ ดังนี้

- **ข้อความปกติ** มีสัดส่วนลดลงอย่างชัดเจนจาก 97.15% ในช่วงก่อนการเลือกตั้ง เหลือเพียง 92.4% หลังการเลือกตั้ง
- **ข้อความที่ถูกระบุว่าเป็น Hate Speech** เพิ่มขึ้นมากกว่า 5 เท่า จาก 0.59% เป็น 3.53%
- **ข้อความที่ถูกระบุว่าเป็น Malinformation** เพิ่มขึ้นเกือบสองเท่า จาก 2.20% เป็น 4.00%

การเพิ่มขึ้นของข้อความ Hate Speech และ Malinformation สะท้อนถึงการเปลี่ยนแปลงของเนื้อหาบนแพลตฟอร์ม Facebook หลังการเลือกตั้ง โดยข้อความในช่วงก่อนการเลือกตั้งส่วนใหญ่เป็นการโพสต์ที่เกี่ยวข้องกับนโยบายหรือการหาเสียง ซึ่งทำให้สัดส่วนข้อความปกติสูงกว่าช่วงหลังการเลือกตั้ง เมื่อเริ่มมีความคิดเห็น ข่าวดัง และข่าววงในมากขึ้น ส่งผลให้ข้อความที่มีปัญหาเพิ่มขึ้นอย่างเด่นชัด

ข้อมูลจาก X แสดงให้เห็นพฤติกรรมผู้ใช้งานที่แตกต่างจาก Facebook ดังนี้

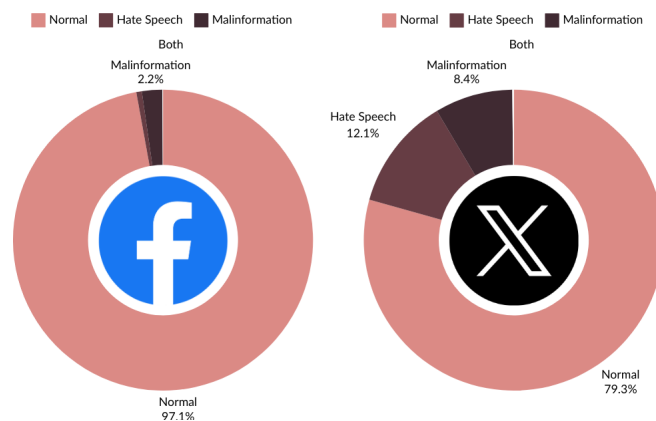
- **ข้อความปกติ** มีสัดส่วนเพิ่มขึ้นเล็กน้อย จาก 79.3% ในช่วงก่อนการเลือกตั้ง เป็น 82.95% หลังการเลือกตั้ง
- **ข้อความที่ถูกระบุว่าเป็น Hate Speech** เพิ่มขึ้นเล็กน้อย จาก 12.1% เป็น 12.72% ซึ่งยังคงมีสัดส่วนสูงกว่า Hate Speech บน Facebook
- **ข้อความที่ถูกระบุว่าเป็น Malinformation** ลดลงเกือบครึ่ง จาก 8.41% เป็น 4.16%

การเปลี่ยนแปลงนี้อาจสะท้อนถึงความสนใจและรูปแบบพฤติกรรมของผู้ใช้งานหลังการเลือกตั้ง โดยเฉพาะในช่วงที่มีประเด็นสำคัญทางการเมือง เช่น การจัดตั้งรัฐบาล การกลับมาของอดีตนายกรัฐมนตรี การยุบพรรคก้าวไกล และข่าวดังต่างๆ ซึ่งมีแนวโน้มส่งผลให้ Hate Speech บน X ยังคงอยู่ในระดับสูง ในขณะเดียวกัน การลดลงของ Malinformation อาจเป็นผลมาจากการเปลี่ยนแปลงเนื้อหาและความสนใจของผู้ใช้งาน

ข้อสังเกตเชิงเปรียบเทียบ

1. ข้อมูลจาก Facebook สะท้อนถึงลักษณะของแพลตฟอร์มที่เน้นการเผยแพร่เนื้อหาที่เป็นนโยบายหรือการหาเสียงในช่วงก่อนการเลือกตั้ง ขณะที่ช่วงหลังการเลือกตั้ง เนื้อหาที่เกี่ยวข้องกับความคิดเห็นส่วนบุคคลและข่าวลือเพิ่มขึ้น ส่งผลให้ข้อความปกติลดลงและข้อความที่มีปัญหาเพิ่มขึ้น
2. ข้อมูลจาก X ชี้ให้เห็นถึงธรรมชาติของแพลตฟอร์มที่มีการแสดงความคิดเห็นอย่างเปิดเผย โดยเฉพาะในประเด็นการเมือง ทำให้ Hate Speech มีสัดส่วนสูงกว่า Facebook อย่างต่อเนื่อง ขณะที่การลดลงของ Malinformation อาจสะท้อนถึงการเปลี่ยนแปลงโครงสร้างของข้อมูลที่เก็บรวบรวมหรือการปรับตัวของผู้ใช้งานหลังการเลือกตั้ง
3. ผลการศึกษานี้ชี้ให้เห็นถึงความแตกต่างในพฤติกรรมผู้ใช้งานและการจัดการเนื้อหาบนแพลตฟอร์ม Facebook และ X ซึ่งมีผลต่อการแพร่กระจายของ Hate Speech และ Malinformation ในบริบทการเมืองไทย

2) ผลการวิเคราะห์ข้อมูลช่วงเลือกตั้ง



รูป 14 ผลการจำแนกประเภทข้อความตามแก่นสาระของข้อมูลช่วงเลือกตั้ง

ในช่วงการเลือกตั้ง พบว่าแพลตฟอร์ม **Facebook** มีลักษณะการใช้งานที่เน้นเนื้อหาเชิงบวกและสร้างสรรค์มากกว่า โดยสัดส่วนข้อความแสดงถึงความเป็นปกติของเนื้อหาสูงถึง **97.15%** ซึ่งบ่งบอกว่าผู้ใช้งานส่วนใหญ่โพสต์เนื้อหาที่เกี่ยวข้องกับการสื่อสารนโยบายและกิจกรรมการหาเสียงมากกว่าการแสดงความคิดเห็นที่ก่อให้เกิดความขัดแย้ง

สำหรับข้อความในหมวดที่เป็นปัญหา

- **Hate Speech** มีสัดส่วนเพียง **0.59%** โดยไม่มีข้อความในระดับรุนแรงที่สุด (Level 5)
- **Malinformation** พบในสัดส่วน **2.20%**
- ข้อความที่อยู่ในทั้งสองหมวดพร้อมกัน มีเพียง **0.07%**

ข้อมูลนี้สะท้อนว่าแพลตฟอร์ม Facebook มีการเผยแพร่เนื้อหาที่เน้นการสื่อสารนโยบายและแคมเปญการเลือกตั้งเป็นหลักในช่วงเวลาดังกล่าว ส่งผลให้ข้อความที่เป็นปัญหา เช่น Hate Speech และ Malinformation มีสัดส่วนต่ำ

ในทางกลับกัน ข้อมูลจาก **X** แสดงลักษณะการใช้งานที่แตกต่าง โดยข้อความปกติมีสัดส่วนเพียง **79.3%** ต่ำกว่า Facebook อย่างชัดเจน ซึ่งสะท้อนถึงธรรมชาติของแพลตฟอร์มที่เปิดให้แสดงความคิดเห็นและวิพากษ์วิจารณ์มากกว่า พบ

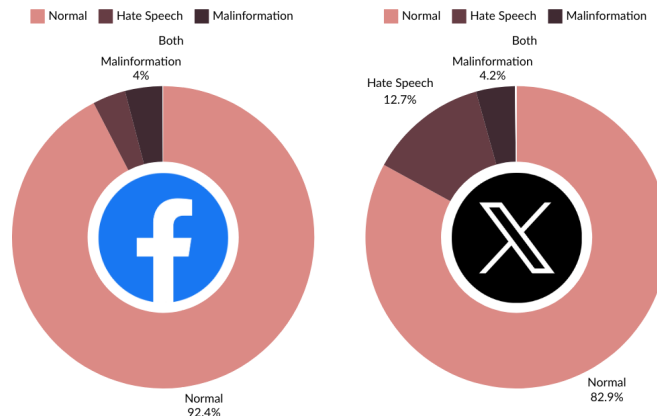
- **Hate Speech** มีสัดส่วนสูงถึง **12.1%** โดยส่วนใหญ่อยู่ใน **ระดับ 1 (10.42%)** และมีข้อความในระดับรุนแรง (Level 4-5) รวมกัน **0.15%**
- **Malinformation** พบในสัดส่วน **8.41%**

- ข้อความที่อยู่ในทั้งสองหมวดพร้อมกันมี **0.16%**

ข้อสังเกตเชิงเปรียบเทียบ

1. X มีข้อความที่เป็น Hate Speech และ Malinformation สูงกว่า Facebook อย่างมีนัยสำคัญ โดยเฉพาะ Hate Speech ในระดับรุนแรง (Level 4-5) ซึ่งไม่มีใน Facebook
2. Facebook มีข้อความปกติในสัดส่วนที่สูงกว่ามาก (97.15% เทียบกับ 79.3%)

3) ผลการวิเคราะห์ข้อมูลช่วงหลังเลือกตั้ง



รูป 15 ผลการจำแนกประเภทข้อความตามแก่นสารของข้อมูลช่วงหลังเลือกตั้ง

ในช่วงหลังการเลือกตั้ง พบว่า **92.4%** ของข้อความบนแพลตฟอร์ม **Facebook** เป็นข้อความปกติ โดยข้อความในหมวดที่เป็นปัญหา ได้แก่ **Malinformation** และ **Hate Speech** มีสัดส่วนรวมกันประมาณ **7.53%** โดย **Hate Speech** ส่วนใหญ่อยู่ในระดับเบา (Level 1) แต่ยังมีพบข้อความในระดับรุนแรง (Level 4-5) ส่วน **Malinformation** มีการกระจายตัวในระดับที่ไม่สูงมากเมื่อเทียบกับ X

ลักษณะดังกล่าวสะท้อนถึงการใช้งาน Facebook ที่ยังคงเน้นการแลกเปลี่ยนความคิดเห็นในเชิงสร้างสรรค์ แม้จะมีการเพิ่มขึ้นของข้อความที่เป็นปัญหาล็กน้อยหลังการเลือกตั้ง

ในขณะที่ข้อมูลจาก X พบว่า **82.95%** ของข้อความเป็นข้อความปกติ ขณะที่ข้อความในหมวด **Malinformation** และ **Hate Speech** รวมกันมีสัดส่วนสูงถึง **16.88%** โดย **Hate Speech** มีสัดส่วนถึง **12.72%** โดยส่วนใหญ่อยู่ในระดับเบา (Level 1) แต่ยังมีพบข้อความในระดับรุนแรง (Level 4) **Malinformation** ยังคงมีสัดส่วนที่น่าสังเกตในช่วงหลังการเลือกตั้ง ข้อความในหมวด Both พบในสัดส่วนที่สูงกว่า Facebook ซึ่งชี้ถึงความซับซ้อนของปัญหาบนแพลตฟอร์ม

ข้อสังเกตเชิงเปรียบเทียบ

1. **ข้อความปกติ** Facebook มีสัดส่วนข้อความปกติสูงกว่า X อย่างชัดเจน (92.4% เทียบกับ 82.95%)
2. **ข้อความที่เป็นปัญหา** X มีสัดส่วน Hate Speech และ Malinformation สูงกว่า โดยเฉพาะ Hate Speech ในระดับเบา (Level 1) และยังมีพบข้อความระดับรุนแรง (Level 4) มากกว่า Facebook
3. **ข้อความในหมวด Both** X มีข้อความในหมวดนี้สูงกว่า ซึ่งสะท้อนถึงปัญหาข้อความที่มีความซับซ้อนและการจัดการเนื้อหาที่อาจมีข้อจำกัดมากกว่า Facebook

การเปรียบเทียบข้อมูลระหว่างแพลตฟอร์มในช่วงหลังการเลือกตั้งชี้ให้เห็นว่า Facebook มีลักษณะการใช้งานที่ยังคงเน้นเนื้อหาปกติในสัดส่วนสูง แม้จะมีการเพิ่มขึ้นของข้อความที่เป็นปัญหา ขณะที่ X แสดงถึงลักษณะการใช้งานที่เปิดกว้างสำหรับความคิดเห็นและวิพากษ์วิจารณ์ ซึ่งนำไปสู่การเพิ่มขึ้นของ Hate Speech และ Malinformation ในระดับที่สูงกว่า ทั้งนี้ยังสะท้อนถึงความแตกต่างในธรรมชาติของผู้ใช้งานและความท้าทายในการจัดการเนื้อหาระหว่างสองแพลตฟอร์ม

3.2.2 การวิเคราะห์หัวข้อด้วยอัลกอริทึมการจำลองหัวข้อ

ในการวิเคราะห์หัวข้อด้วยอัลกอริทึมการจำลองหัวข้อ หรือ Topic Modeling ส่วนนี้จะขอนำเสนอแค่ผลการวิเคราะห์หัวข้อในข้อมูลหลังการเลือกตั้งเท่านั้น⁶

ในชุดข้อมูลหลังการเลือกตั้ง โครงการเก็บรวบรวมข้อมูลตั้งแต่วันที่ 1 มกราคม 2567 – 31 สิงหาคม 2567 โดยแบ่งข้อมูลเพื่อวิเคราะห์หัวข้อมาเป็น 3 ช่วงเวลา และเพื่อความแม่นยำในการตีความหัวข้อที่ปรากฏขึ้น โครงการฯ ได้รวบรวมเหตุการณ์สำคัญทางการเมืองไว้ ดังนี้

ช่วงที่ 1: 1 มกราคม – 31 มีนาคม 2567 เหตุการณ์ทางการเมืองที่สำคัญที่มียอดการมีส่วนร่วมสูงในโซเชียลมีเดียคือ การอภิปรายงบประมาณรายจ่ายประจำปี 2567 ในสภาผู้แทนราษฎร รัฐบาลพลัดถิ่นเมกะโปรเจกต์แลนด์บริดจ์ ศาลรัฐธรรมนูญนัดอ่านคำวินิจฉัยคดีล้มล้างการปกครองของพรรคก้าวไกล การพิทักษ์ทักษิณ ชินวัตร นิตยสาร TIME พาดหัวบทสัมภาษณ์นายกฯ เศรษฐาว่าเป็น “The salesman Thai Prime Minister” และการปรับ ครม. “เศรษฐา 1/1”

ช่วงที่ 2: 1 เมษายน – 30 มิถุนายน 2567 เหตุการณ์ทางการเมืองที่สำคัญที่มียอดการมีส่วนร่วมสูงในโซเชียลมีเดียคือ การปรับ ครม. “เศรษฐา 2” การเลือกตั้งท้องถิ่นจังหวัดบุรีรัมย์ การประชุมใหญ่สามัญประจำปี 2567 ของพรรคภูมิใจไทย เทศกาลสงกรานต์ ททท.จัดงานสงกรานต์สนามหลวงมีวงหมอลำ เนติพร (นุ้) เสียชีวิต กระทรวงพาณิชย์เปิดประมูลข้าวในโครงการรับจำนำข้าว เปิดรับสมัคร สว. และการคัดเลือก สว. ระดับประเทศ

ช่วงที่ 3: 1 กรกฎาคม – 31 สิงหาคม 2567 เหตุการณ์ทางการเมืองที่สำคัญที่มียอดการมีส่วนร่วมสูงในโซเชียลมีเดียคือ ประกาศผลการคัดเลือก สว. เปิดลงทะเบียนโครงการดิจิทัลวอลเล็ต ศาลรัฐธรรมนูญตัดสินยุบพรรคก้าวไกล เศรษฐา ทวีสิน ถูกปลดจากตำแหน่งนายกรัฐมนตรี แพทองธาร ชินวัตร ได้รับเลือกเป็นนายกรัฐมนตรี ทักษิณ ชินวัตร ได้รับอภิทักษิณ และวาทกรรม “รัฐบาลผสมข้ามชั่ว” ในกรณีพรรคประชาธิปัตย์ร่วมรัฐบาล



รูป 16 เหตุการณ์ทางการเมืองที่สำคัญที่มียอดการมีส่วนร่วมสูงในโซเชียลมีเดีย

อย่างไรก็ตาม นอกจากเหตุการณ์สำคัญในแต่ละช่วงแล้ว ยังพบประเด็นที่ผู้ใช้งานโซเชียลมีเดียมีประเด็นถกเถียงกันตลอดทั้งปี เช่น การเลือกตั้งท้องถิ่นที่ทยอยจัดขึ้นตลอดทั้งปี 2567 ประเด็น การเหยียดสตรีเพศ (misogyny) ที่พุ่งเป้าไปที่นักการเมืองหรือนักกิจกรรมผู้หญิง และประเด็น Foreign Information Manipulation and Interference (FIMI) กับวาทกรรม “รับทุนต่างชาติมาบ่อนทำลายไทย” เป็นต้น

ในตารางการจำลองหัวข้อจะประกอบด้วยหัวตาราง 4 คอลัมน์ โดยแต่ละคอลัมน์มีความหมายดังนี้

Cluster หมายถึง กลุ่มข้อความที่ถูกจัดกลุ่มเข้าด้วยกันตามความคล้ายคลึง หรือ ความเชื่อมโยงกันในเชิงเนื้อหา

Top Words หมายถึง คำสำคัญที่ปรากฏบ่อยที่สุดในแต่ละกลุ่มข้อความ สะท้อนธีมหรือเนื้อหาหลัก เช่น ใน Cluster 4 มีคำว่า “เสีย,” “แหม่ง,” “ควาย” ที่แสดงถึงการใช้ภาษาก้าวร้าวหรือเสียดสี

⁶ หากสนใจดูผลการวิเคราะห์หัวข้อของข้อมูลช่วงการเลือกตั้งสามารถดูเพิ่มเติมจาก รายงาน “ถอดรหัสปฏิบัติการบนข้อมูลข่าวสารในไทย #เลือกตั้ง2566” และฐานข้อมูลของโครงการติดตามและวิเคราะห์กระบวนการเลือกตั้งในโซเชียลมีเดีย (Digital Election Analytic Lab: DEAL)

ช่วงที่ 2 พบว่าหัวข้อที่ปรากฏขึ้นจากบทสนทนาออนไลน์นั้นสอดคล้องกับเหตุการณ์ทางการเมืองที่สำคัญที่มีขบวนการมีส่วนร่วมสูงในโซเชียลมีเดีย ทั้งประเด็น กระทรวงพาณิชย์เปิดประมูลข้าวในโครงการรับจำนำข้าว การเปิดรับสมัคร สว. และการคัดเลือก สว. การบริหารงานของอนุทิน ชาญวีรกูล สมัยเป็น รมต.สธ. ในการบริหารสถานการณ์โควิด (คาดว่ามาจากการเลือกตั้งท้องถิ่นจังหวัดบุรีรัมย์ การประชุมใหญ่สามัญประจำปี 2567 ของพรรคภูมิใจไทย) โดยพบว่าในช่วงนี้จะถกเถียงกันในเรื่องการกระทรวงพาณิชย์เปิดประมูลข้าวในโครงการรับจำนำข้าวมากที่สุด ทั้งเรื่องข้าวเก่ามีสารอะฟลาท็อกซินหรือไม่ การผสมข้าวเก่ามาหลอกลขายประชาชน การจัดอีเวนต์หุงข้าวโฮว้ รวมไปถึงการทำงานของสื่อมวลชนที่นำเสนอข้อมูลบิดเบือนในเรื่องการประมูลข้าว (ตาราง 9)

ตาราง 9 การจำลองหัวข้อที่ปรากฏขึ้นหลังการเลือกตั้งในช่วงที่ 2 ของแพลตฟอร์ม X

Cluster	Top Words	Document Count	Topic Gloss
6	ข้าว, บุด, แดก, ดี, ไม่ได้, คนอื่น, ควาย, สลัม, โท, โคตร, อาหาร, เฝี้ย, หน้า, ซอน, หมอ, ลัทธิ, จั, กบฏ, โง่, ครุ, หุงข้าว, แนน่า, เนอะ, เพยแพร่, อะฟลาท็อกซิน, ดี, โลก, แม่, สมอง, ศาล	152	กระทรวงพาณิชย์เปิดประมูลข้าวในโครงการรับจำนำข้าว
2	ไม่ได้, เพ้อไทย, ส้ม, เลือก, นางแบบ, นโยบาย, สว, แก, ต่อผล, รัฐบาล, ก้าวไกล, โง่, ดี, ประชาชน, เลือกตั้ง, สลัม, เพด็จการ, โคตร, รอบ, คาเสียง, ประเทศ, คนอื่น, ตระบัดสัตย์, ควาย, สส, พลังงาน, เฝี้ย, โง่, ระบบ, แคม	148	เปิดรับสมัคร สว. และการคัดเลือก สว.
1	หมา, ก็น, เฝี้ย, แม่, แม่่ง, ควาย, ตาย, ดี, เอ๊ย, อม, หี, เจิน, ขาย, หลอก, ย, บัญชี, ข่า, ตอนที, เลี้ยง, เกลียด, บ้า, ซื่อ, เมว, ยา, แดก, สลัม, ม้า, รวด, ติดคุก, ด	147	ประเด็นปัญหาสังคม เช่น ยาเสพติด ภัยอาชญากรรมแก๊งคอลเซ็นเตอร์ มิจฉาชีพ
4	เฝี้ย, ไม่ได้, เจิน, ดี, ภู, ประเทศ, นางแบบ, ล้าน, ขนาด, รัฐบาล, เลือก, ดี, คนไทย, ซื่อ, ติด, ก็น, เพื่อน, ประชาชน, ปัญญา, ควาย, โง่, แรงงาน, ตาย, พรรคเพื่อไทย, จบประมาณ, จบ, เรียน, ชินหาย, ก็ต้อง, แม่่ง	145	การค้า ด้อยค่ากันของกลุ่มผู้สนับสนุนพรรคการเมือง
5	ประเทศ, เฝี้ย, รัฐบาล, สลัม, ทุเรศ, รัก, สนับสนุน, ไม่ได้, คนอื่น, ตาย, ทำร้าย, โง่, ทักษณ, ด้อยค่า, ดีกว่า, เสือแดง, รสน, บ้าน, รีปล่า, แม่่ง, แดก, กฎหมาย, อวย, ร้าง, สุดท้าย, สอน, มือบ, ซาติ, ล, จัน	110	การค้ากลุ่มอนุรักษ์นิยม/establishment
7	ไม่ได้, คนไทย, ทำงาน, หมอ, คนอื่น, ตรวจ, รัฐบาล, เจิน, บังคับ, คนใน, ตัด, ขาย, คนไข้, มีปัญหา, ปัญหา, หลอก, สังคม, เรียน, เฝี้ยด, สลัม, IO, เก่ง, พัง, นางแบบ, สมอง, เหมือนกัน, ติด, ถูกเงิน, อวยวะ, ครอบครัว	93	ประเด็นปัญหาสังคม เช่น เวลาการทำงานของบุคลากรทางการแพทย์
3	เฟค, มะเร็ง, สาร, ก่อ, นิว, ไม่ได้, คนอื่น, ความคิดเห็น, เหมือนกัน, หลักฐาน, อะฟลาท็อกซิน, นิวส์, เลิก, เจอ, ตาย, ปลอ่ย, ดี, อา, แยก, โรงงาน, ซื่อ, ติด, แพน, เก้า, สนับสนุน, บางคน, ผู้หญิง, ถูกใจ, ลอ, จิ	85	กระทรวงพาณิชย์เปิดประมูลข้าวในโครงการรับจำนำข้าว และการถกเถียงว่าข้าว 10 ปีมีสารอะฟลาท็อกซินเป็นสารก่อมะเร็ง
0	ข้าว, เอาแต่, ขาย, รัฐบาล, แม่, ดี, คสช, ลิด, พกา, โครงการ, เพา, สภาพ, ไม่ได้, ขาดทุน, แยก, ความจริง, ประยุทธ์, ตาย, จำนำ, จี, ปู, อาหารสัตว์, นม, ตรงไหน, บิดเบือน, สือ, เจลย, ตรวจ, ข้าว, ดี	81	กระทรวงพาณิชย์เปิดประมูลข้าวในโครงการรับจำนำข้าว
9	วัคซีน, ตาย, ดี, ข้าว, หมอ, คนไทย, ติด, รัฐบาล, โควิท, จง, พ่อ, ข้อมูล, ดลก, SinoVac, ภัยชา, ติด, mRNA, โคตร, นิง, อาการ, Antivaxx, พลข้างเคียง, AstraZeneca, ปั่น, สมอง, ส้ม, โทษ, ๆๆ, J, ติดวัคซีน	71	การบริหารงานของอนุทิน ชาญวีรกูล สมัยเป็น รมต.สธ. ในการบริหารสถานการณ์โควิด (การเลือกตั้งท้องถิ่นจังหวัดบุรีรัมย์ การประชุมใหญ่สามัญประจำปี 2567 ของพรรคภูมิใจไทย)
8	อิสราเอล, อาฆา, ยิว, ปาเลสไตน์, ข่า, สงคราม, ไม่ได้, เซียร์, ไล, ประเทศ, กลายเป็น, บ้าน, รัฐ, ทะเล, fromrivertosea, สามิ, แม่่ง, ตัวประกัน, โนส, ทีมงาน, fromrivertothesea, เรียงร้อย, โลก, ทุน, จบ, เลือก, ขนาด, เหตุการณ์, หยุด, ด้วยซ้ำ	57	ประเด็นระดับนานาชาติ เช่น ความสัมพันธ์กับประเทศตะวันออกกลาง เหตุการณ์รัสเซีย-ยูเครน

หมายเหตุ เรียงตามยอด Document Count

ช่วงที่ 3 พบว่าหัวข้อที่ปรากฏขึ้นจากบทสนทนาออนไลน์นั้นสอดคล้องกับเหตุการณ์ทางการเมืองที่สำคัญที่มียอดการมีส่วนร่วมสูงในโซเชียลมีเดีย ทั้งประเด็น เศรษฐา ทวีสิน ถูกปลดจากตำแหน่งนายกรัฐมนตรี วาทกรรม "รัฐบาลผสมข้ามขั้ว" จากกรณีดังพรรคประชาธิปัตย์ร่วมรัฐบาล พบว่าในช่วงที่ 3 นี้มีประเด็นที่เด่นขึ้นมาคือ การยกเลิกโครงการบริหารงานของรัฐบาล เช่น การแก้กฎหมายให้ต่างชาติเช่าที่ดิน 99 ปี อย่างไรก็ตาม ใน Cluster 1 และ 2 Top Words ค่อนข้างกระจายกระจายจนโครงการฯ ไม่สามารถระบุ Topic Gloss ที่มีนัยสำคัญได้ (ตาราง 10)

ตาราง 10 การจำลองหัวข้อที่ปรากฏขึ้นหลังการเลือกตั้งในช่วงที่ 3 ของแพลตฟอร์ม X

Cluster	Top Words	Document Count	Topic Gloss
0	เหี้ย, ควาย, ไม่ได้, โง่, เสือก, แม่ง, หี, ปาก, ควย, ส้ม, เก่ง, ดี, ตา, บั๊, ซิบ, หาย, สัส, กระบอก, เอ๊ย, บิน, ร้องไห้, เกียง, ฤษ, ประชาชน, แม่, สันดาน, ประเทศ, ตาย, อีเหี้ย, วะ, แดก	230	การด่า ด้วยคำก่นของกลุ่มผู้สนับสนุนพรรคการเมือง
6	เหี้ย, ด, ตึง, โคตร, แม่ง, บ, บ้า, แอค, แวง, ตาย, ะ, อ, ง, ะ, ย, งง, ไม่ได้, ขนาด, แม่, เปน, เ, ุ้, โอ้ย, เอ้, ชั๊, เต, อด, ำ, คุม, ำ, ำ	169	การใช้ถ้อยคำที่รุนแรง ดูถูก และการแสดงความเห็นอย่างชัดเจน
5	ต่อแหละ, นางแบบ, เพ้อไทย, หลีกฐาน, ส้ม, สส, ไม่ได้, อู๊, แลนด์สไลด์, ขนาด, ลั่น, โหวต, แดก, รัฐบาล, เสือก, แม่นอน, ส้ม, โง่, รอบ, กปปส, คนอื่น, แพ้, กลัว, รัฐประหาร, ฉลาด, เก่ง, เสือกตึง, ข้าม, ขั้ว, โครงการ	126	วาทกรรม "รัฐบาลผสมข้ามขั้ว"
9	สส้ม, ชาดี, คนอื่น, ชาตินิยม, ส้ม, ประเทศ, ไม่ได้, นางแบบ, ไล้, สถาบัน, คนไทย, รัก, เสือกตึง, ดี, อวย, ชอบ, เพ้อไทย, ก้าวไกล, เลิก, ประชาชน, รัฐประหาร, บางคน, เป้, กฎหมาย, แก่, แฉ, อ้าง, ชัง, บท, การปกครอง	123	การด่า ด้วยคำก่นของกลุ่มผู้สนับสนุนพรรคการเมือง และการด่ากลุ่มอนุรักษนิยม/establishment
4	ไม่ได้, นางแบบ, ประเทศ, สมน้ำหน้า, เจ๊, ดี, พัง, เมิง, สอบ, PISA, ตก, ส้ม, ตอบ, เป็นได้, รัฐบาล, เกณฑ์, ศาลรัฐธรรมนูญ, ศาล, แจก, อาทิตย์, ซี้ซี้, ุ้, สะใจ, เศรษฐา, พรรคเพื่อไทย, ชาดี, คำถาม, แบงค์, เกล็ด, นายก	91	เศรษฐา ทวีสิน ถูกปลดจากตำแหน่งนายกรัฐมนตรี
8	ดี, ตาย, ไม่ได้, พช, ชี, คนอื่น, บ้าน, ลอง, ประเทศ, ซีวิต, ซื่อ, โคตร, หยุต, เหี้ย, จั, ลบ, ำ, แรง, จิน, ทวิต, อธิบาย, หายใจ, หมอ, ห่วย, พญ, คุมสติ, เต็ม, ยุโรป, อเมริกา, วง	85	การแก้กฎหมายให้ต่างชาติเช่าที่ดิน 99 ปี
3	มหาสิ, แม่, งานวิจัย, ไม่ได้, ก็น, จบ, ยัย, จริยธรรม, ซื่อ, บุด, ฟ้อง, มาตรฐาน, ดำว่า, ข้าว, พุงข้าว, แม่ง, หยุต, ส้ม, ASTV, ข่า, สว, ตัดสิทธิ์, STOPTHEGENOCIDE, ล้าง, เผาพันธุ์, WorldstandswithPALESTINEfromBangkoktoday, นัก, มหิดล, ขาย, ข่า	82	กระทรวงพาณิชย์เปิดประมูลข้าวในโครงการรับจำนำข้าว และ ประเด็นปัญหาสังคม เช่น การซื้อขายงานวิจัย/วุฒิการศึกษา รวมทั้ง ประเด็นระดับนานาชาติ เช่น ความสัมพันธ์กับประเทศตะวันออกกลาง
7	ป้าย, หลีกฐาน, เป้, ไม่ได้, มั่นใจ, สือ, กบ, นางแบบ, ลูกค้า, เกียง, เกล็ด, สรูป, ควาย, โฆษณา, ข้อนู, ป้ายโฆษณา, โปร, อัสราเอล, เจ้, ปลอม, บี, คนอื่น, โชว์, คำถาม, กฎหมาย, สันดาน, เหี้ย, หน้า, ครุ, จบ	79	การด่า ด้วยคำก่นของกลุ่มผู้สนับสนุนพรรคการเมือง
1	จิด, เจ๊, ตา, โฆษณา, หมอ, บริษั, เคส, สรูป, มะ, มาตรฐาน, กระเป๋า, สอง, แม่ง, คนอื่น, ยักยอก, แยม, กรณิ, ตื่น, เข้าข่าย, รอบ, เหี้ย, Filler, ผู้ถือ, หุ่น, นิ่ง, วิชาชีพ, คนไข้, ดี, โป, ผู้บริหาร, แขนง	67	Undefined
2	ข่าว, จิน, แก่, ดี, เลิก, ทีมงาน, กาว, บ้า, ตรงไหน, รัฐบาล, เล่า, ฝาก, BannedfromBroadcast, สารคดี, พื้นที่, หลึง, ฟู, อุดสาหกรรม, ทีม, สื่อสาร, รัก, กลายเป็น, ปลอ่ย, พลิก, ต่อสู้, พ่อแม่, โหด, กระต๊, มนุษย์, สา	48	Undefined

หมายเหตุ เรียงตามยอด Document Count

2) หัวข้อที่ปรากฏขึ้นใน Facebook

ช่วงที่ 1 พบว่าหัวข้อที่ปรากฏขึ้นจากบทสนทนาออนไลน์นั้นส่วนใหญ่เป็นประเด็นที่ไม่ได้สอดคล้องกับเหตุการณ์ทางการเมืองที่สำคัญที่โครงการได้ศึกษาไว้ก่อนหน้านี้ ทั้งนี้พบประเด็นที่น่าสนใจคือ การวิพากษ์วิจารณ์นโยบายรัฐบาล การวิพากษ์การจัดสรรทรัพยากร ในประเด็นต่างๆ เช่น โครงการ Digital Wallet การแก้ไขปัญหาเศรษฐกิจ ปัญหาหนี้สิน หนี้ครัวเรือน การผลักดันเมกะโปรเจก Entertainment complex เป็นต้น

เป็นที่น่าสนใจว่าประเด็นที่ได้รับการมีส่วนร่วมสูงอย่างศาลรัฐธรรมนูญนัดอ่านคำวินิจฉัยคดีล้มล้างการปกครองของพรรคก้าวไกล การพักโทษทักษิณ ชินวัตร หรือการอภิปรายงบประมาณรายจ่ายประจำปี 2567 ในสภาผู้แทนราษฎรอย่างที่พบใน X กลับไม่มากพอที่ โมเดลจะสามารถจัดกลุ่มได้ใน Facebook

อย่างไรก็ตามยังพบ การด่า ด้อยค่ากันของกลุ่มผู้สนับสนุนพรรคการเมือง เช่นเดียวกับ X แต่ที่น่าสนใจคือพอจะจัดกลุ่มได้ว่าการด่ากันเหล่านั้นกำลังใช้วาทกรรมอะไรในการโจมตีกัน เช่น วาทกรรม “ล้มเจ้า” “โกง” เป็นต้น (ตาราง 11)

ตาราง 11 การจำลองหัวข้อที่ปรากฏขึ้นหลังการเลือกตั้งในช่วงที่ 1 ของแพลตฟอร์ม Facebook

Cluster	Top Words	Document Count	Topic Gloss
7	บ้าน, อ่าง, ประชาชน, อุทสาหกรรม, โรงงาน, จิบหาย, ตาย, ฝาก, แผลง, วัคซีน, คนป่วย, เช็ดๆๆๆๆ, เข้, คบ, เรื่องมาก, หนู, จ้อ, ลาก, เตี้ยง, เหย้, ตอนแรก, อาศัย, ยึด, เอาคืน, โน่น, สารพัด, สุดท้าย, ชาวโลก, สมควร, ว่าง	11	การวิพากษ์วิจารณ์นโยบายรัฐบาล การวิพากษ์การจัดสรรทรัพยากร
4	ดี, หมอสำ, พัฒนา, ประเทศ, ชันต่ำ, พัก, มั่ง, หาเสียง, ยุติ, รัฐ, ราชการ, ศูนย์, กระจายอำนาจ, เต็ม, รูปแบบ, จังหวัด, เลือกดั้ง, ผู้ว่า, เจอ, ข้อมห้าม, แบ่งแยกดินแดน, แบน่อน, กว้านี้, หลง, ใจมา, ประเทศกำลังพัฒนา, ใหน, หว่า, ขยายตัว, สัตว์	11	การเลือกตั้งท้องถิ่น และการกระจายอำนาจ
3	แจก, ดิจิตอล, ประเทศ, รัฐบาล, วอลเล็ท, เงิน, ใช้นี้, 5, ล้าน, จำนำ, ข้าว, วิกฤต, ภู, ข้าง, เทวดา, กรอก, เลข, แบน, ฝากเงิน, ยุ่งยาก, บัตร, เบอร์, โน่น, ปลอ่ย, ดุด, บัญชี, ตัน, ปากกล้า, ขา, สัน	11	การวิพากษ์วิจารณ์นโยบายรัฐบาล การวิพากษ์การจัดสรรทรัพยากร เช่น วาทกรรม “โครงการ Digital Wallet จะเป็นหนี้เหมือนจำนำข้าว”
2	ประชาชน, 9, ห้า, ดำรวจ, ยุค, จำย, มาก, เมือง, ปฎิรูป, ลุง, ดูป, หากิน, เกิดขึ้น, ข้าราชการ, ปัญหา, จิต, ตำ, เทียม, ไอ้โง่, ประธาน, ชุมชน, จ้าง, หมอสำ, จบ, เงินสด, กอง, มะ, เจ้าอาวาส, เอี้ยว, บุตุ้ย	11	ประเด็นปัญหาสังคม เช่น การทุจริตในวงการตำรวจ
1	หลอก, รัก, มีงาช้าง, โลก, สังคม, ยุค, โซ, อาซา, มีลูก, ทิ้ง, สัตว์เลี้ยง, ศักดินา, สกานัน, ระวัจ, คุณหญิง, แจ่มใส, ศัลป์, ภริยา, บรรหาร, ศิป, นายกรัฐมนตรี, 21, แท้ง, โทร, รัฐมนตรี, กระทรวง, พัฒนา, ความมั่นคง, มนุษย์, ถ่าน	8	กัญจนา ศิลปอาชา โกชิวัด พังดัมมี
8	คุก, ลุง, ฐิ, ข่า, ไม่ได้, นอน, ฮือ, อา, กากิ, เงิน, สะกด, มารยาท, เป็นได้, ชน, ชันต่ำ, หมอสำ, สาวก, ดูป, ล้มเจ้า, บล๊อค, ความจริง, โง่, ด้อมส้ม, พุดถึง, คลุม, ซ้อม, ยัดคดี, ห, นึ่ง, เจอ, ไม้	7	การด่า ด้อยค่ากันของกลุ่มผู้สนับสนุนพรรคการเมือง ด้วยวาทกรรม “ล้มเจ้า”
6	โกง, ประเทศ, 1, อำนาจ, เก้าอี้, กฎหมาย, 2, เปิดทาง, ผลักดัน, ย, งานใหญ่, รับสมัคร, รางวัล, คู่มือ, เคล็ดลับ, ติดคุก, เครื่องตัดหญ้า, เครื่องรับ, พังงา, หม้าย, นึ่ง, มรั้ง, ซื่อ, กระชาก, วิญญาณ, รัฐประหาร, รัฐธรรมนูญ, ลดทอน, ใจหลวง, คัด	7	การด่า ด้อยค่ากันของกลุ่มผู้สนับสนุนพรรคการเมือง ด้วยวาทกรรม “โกง”
5	ความเห็น, ข้อมูล, ของปลอม, ปลอม, หมู่บ้าน, ชาวบ้าน, กระ, ย, รุนแรง, ปลา, ญี่ปุ่น, ข่า, ฟู, แยะ, ดีแต่, โซเซี่ยล, ผู้คน, กรม, ศัลป์, สัมภาษณ์, สิบ, คนใน, ยอมให้, อ้างอิง, วิทยานิพนธ์, คำบอกเล่า, เป็นเรื่อง, พุดโกหก, ติดใจ, คนภายนอก	6	Undefined
9	กด, อันตราย, กลางแดด, ทำงาน, หนี้สิน, ระมัดระวัง, 7, ความจริง, ประมาท, เมา, เมีย, ลอง, ชัก, 10, นาทิ, ดี, นึ่ง, ห้อง, แอร์, เหมือนกัน, ระวัจ, หลวมตัว, ลิงค์, สั่ง, พุง, เศรษฐกิจ, เพาจริง, ไม่ได้, เพาหลอก, โอกาส	5	การวิพากษ์วิจารณ์นโยบายรัฐบาล การวิพากษ์การจัดสรรทรัพยากร เช่น การแก้ไขปัญหาเศรษฐกิจ หนี้สิน
0	ภาชี, คาสโน, บาป, จบ, สรุป, กฎหมาย, รัฐบาล, อบายมุข, ปากทาง, ความ, จิบหาย, สุขภาพ, ไร้สาระ, มาตรา, 8, พรบ, จบประมาณ, มั่นคง, มั่งคั่ง,	5	รัฐบาลผลักดันเมกะโปรเจก Entertainment complex

Cluster	Top Words	Document Count	Topic Gloss
	ยังยืน, บทละคร, ว่าด้วย, อ้าง, ความมั่นคง, ประเทศชาติ, กระทรวงกลาโหม, ชئون, ยุทธโศภณ, สืบ, เอกสารสืบ, นำพิศวง		

***หมายเหตุ* เรียงตามยอด Document Count**

ช่วงที่ 2 พบว่าหัวข้อที่ปรากฏขึ้นจากบทสนทนาออนไลน์นั้นส่วนใหญ่เป็นประเด็นที่ไม่ได้สอดคล้องกับเหตุการณ์ทางการเมืองที่สำคัญที่โครงการได้ศึกษาไว้ก่อนหน้านี้ และประเด็นการพูดคุยใน Facebook ค่อนข้างกระจัดกระจาย และเฉพาะเรื่องลงไปในเรื่องประเด็นมาก เช่น สนธิ ลิ้มทองกุล จัดรายการ sonthitalk วิเคราะห์ทางออกของทักษิณ ชินวัตร ในคดี 112 สุดารัตน์ เกยุราพันธุ์ เผยข่าวลือ “เพื่อไทยจะแก้ รธน. ให้ไม่มีระบบบัญชีรายชื่อ” เกี่ยวกับ SQ321 สิงคโปร์แอร์ไลน์เผชิญเหตุฉุกเฉินที่ท่าอากาศยานสุวรรณภูมิ ประเทศไทย เป็นต้น

อย่างไรก็ตาม ยังพบการถกเถียงของผู้ใช้งานโซเชียลมีเดียในประเด็นต่างๆ เช่น กระทรวงพาณิชย์ เปิดประมูลข้าวในโครงการรับจำนำข้าว นโยบายรัฐบาลศึกษาเรื่อง Carbon Tax การเพิกถอนพื้นที่ “อุทยานแห่งชาติกับลาน” จนเกิดกระแส 2 กระแสที่ขัดแย้งกัน คือ #SAVEกับลาน กับ #SAVEชาวบ้านกับลาน รวมไปถึงการอภิปรายงบประมาณรายจ่ายประจำปี 2567 ในสภาผู้แทนราษฎร

ที่น่าสนใจคือ ใน Facebook ช่วงที่ 2 นี้ การใช้ถ้อยคำที่รุนแรง ถูก และการแสดงความเป็นศัตรูอย่างชัดเจน ทั้งในกลุ่มผู้สนับสนุนพรรคการเมืองต่างๆ การโจมตีกลุ่มอนุรักษ์นิยม หรือ establishment รวมไปถึงกลุ่มนักกิจกรรมทางการเมือง พบไม่มากพอที่จะวิเคราะห์สามารถจัดกลุ่มได้ (ตาราง 12)

ตาราง 12 การจำลองหัวข้อที่ปรากฏขึ้นหลังการเลือกตั้งในช่วงที่ 2 ของแพลตฟอร์ม Facebook

Cluster	Top Words	Document Count	Topic Gloss
4	บ่, ข้าว, 10, ขาย, ดับ, หมา, คดี, บัญชา, คดี, ข้อย, สี, กาก, เหย, เมีย, จำ, สุ, ไฟ, ไฟฟ้า, รื้อ, ศาล, เจ้า, ฟ้า, ตัว, รอง, 1, ดบ, ดี, 2, ตร, บัก	27	กระทรวงพาณิชย์เปิดประมูลข้าวในโครงการรับจำนำข้าว
7	ตำรวจ, ล้มเหลว, ประชาชน, ลาออก, , ก้าวไกล, ระบบ, การเมือง, ย, ดับ, 3, สังคม, ก่อนกำหนด, จงใจ, ทำลาย, ยุติธรรม, นายก, ห่วย, หลักฐาน, รัฐ, เลว, ด, นายก, ลุง, ขายชาติ, เจ้ง, 5555, ประเทศไทย, คนไทย, ดี	26	ประเด็นปัญหาสังคม เช่น การทุจริตในวงการตำรวจ จากกรณี ความขัดแย้งระหว่าง พล.ต.อ. ต่อศักดิ์ สุขวิมล (พ.ต.ธ.) และ พล.ต.อ.สุรเชษฐ์ หักพาล (รอง พ.ต.ธ.)
9	ไม่ได้, คนอื่น, คน, สมอง, ตาย, อ้าง, รัก, เอาเรื่อง, ยบ, สอน, เหมือนกัน, ลด, ชายแดน, โรง, ดูออก, สาม, ดี, แผลง, หมอสำ, แม่, กทม, เค, บวช, พม่า, ยืน, ดี, โรงฉาย, 150, 48, ฉาย	25	Undefined
6	ทักษิณ, สนธิ, ญาณ, คลิป, รายการ, ข้าว, ปัญหา, อานนท์, สะท้อน, กพ, รัฐบาล, 112, คิน, บ112, คินนี้, แอ, ศาลฎีกา, อัยการสูงสุด, วัง, คดี, มีย, 2030, สัง, ฟ้อง, ค, กบ, กลัว, เพื่อไทย, พรบ, เศรษฐา	25	สนธิ ลิ้มทองกุล จัดรายการ sonthitalk วิเคราะห์ทางออกของ ทักษิณ ชินวัตร ในคดี 112
0	สว, เลือก, อ้ว, โหวต, กกด, กระบวนการ, เลือกตั้ง, ไม่ได้, ย้ำ, องค์การอิสระ, อำนาจ, ตรวจสอบ, แต่งตั้ง, 67, 2567, ที่มา, การเลือกตั้ง, ช่วยกัน, เบา, าส, บาก, ผู้สมัคร, สภา, แอ้ง, ทุจริต, ระบบ, คุณหญิง, สุดารัตน์, เปิดทาง, กุณ, หนา	23	สุดารัตน์ เกยุราพันธุ์ เผยข่าวลือ “เพื่อไทยจะแก้ รธน. ให้ไม่มีระบบบัญชีรายชื่อ”
2	บน, กฎหมาย, แม่, ขาย, ยี่ห้อ, เครื่องบิน, โฆษณา, ไม่ได้, โล, ไกลเคียง, เทสียด, แม่ค้า, ดก, จลาจล, หร, โคตร, ก็น, บริษัท, ทีม, บ้อ, ท่า, เสียชีวิต, , ลุงตุ๋, 1, คินนี้, แนนอน, สิงคโปร์, แอ, หลุมอากาศ, ทัก	19	เกี่ยวกับ SQ321 สิงคโปร์แอร์ไลน์เผชิญเหตุฉุกเฉินที่ท่าอากาศยานสุวรรณภูมิ ประเทศไทย
1	เซียร์, ช้อง, คาร์บอน, ปลอ่ย, ป่า, รัฐบาล, ดักว่า, ชาวบ้าน, บ่, พ่อ, ทีม, เนียน, รับเคราะห์, โลก, ผลงาน, ไม่ได้, พิธีกร, คนไทย, บริการ, พัน, เวีย, ชีวิต, ข้าง, เหยียน, เครดิต, คนสวน, ปลอม, ดี, โรงบาล, ๆแล้ว	16	นโยบายรัฐบาลศึกษาเรื่อง Carbon Tax
8	คอร์รัปชัน, ต่อต้าน, ที่ดิน, โทง, ทุจริต, ขอบ, สังคม, ร้อน, องค์การ, สร้าง, เอกสารสิทธิ์, เจ้าหน้าที่, จังหวัด, สำเร็จ, วากรรม, ด้อยค่า, คนอื่น, อาภาศ, ขวน, เบา, าส, ข้าราชการ, จิตสำนึก, เพา, คนไทย, จง, วงจร, พื้นที่, ตรวจสอบ, ศาสนา, เจอ	15	การเพิกถอนพื้นที่ “อุทยานแห่งชาติกับลาน” #SAVEกับลาน กับ #SAVEชาวบ้านกับลาน
5	จน, ตอน, 68, บ้าน, ไทยสร้างไทย, แพง, อภิปราย, โจทย์, ี่, ประชาชน, , รอบ, รัฐบาล, หนี้, บัก, ย้อน, ฟัง, เหตุใด, ลูกขึ้น, อัด, ผิดฝาผิดตัว,	15	การอภิปรายงบประมาณรายจ่ายประจำปี 2567 ในสภาผู้แทนราษฎร

Cluster	Top Words	Document Count	Topic Gloss
	เปิดช่อง, ทุจริต, เพียบ, โดยเฉพาะ, กัฟ, ฟ้า, ชื่อ, วิทยุสื่อสาร, กัฟเรือ, ประเคน		
3	ติดสินบน, สะท้อน, บาป, สาปแช่ง, ทหาร, รับสินบน, เลว, ดี, ข้าราชการ, ความคิด, ท่าน, นิ่ง, ยึดอำนาจ, คอย, สถานะ, สนับสนุน, ประเทศ, ประชาชน, กลายเป็น, คอ, การณ์, ผู้นำ, นา, สังคม, สัก, สามานย์, กุณ, ราชวงศ์, ผู้พิพากษา, อีหม่าม	15	ประเด็นปัญหาสังคม เช่น การทุจริต การรับสินบน

***หมายเหตุ* เรียงตามยอด Document Count**

ช่วงที่ 3 พบว่าหัวข้อที่ปรากฏขึ้นจากบทสนทนาออนไลน์นั้นสอดคล้องกับเหตุการณ์ทางการเมืองที่สำคัญที่มีขบวนการมีส่วนร่วมสูงในโซเชียลมีเดีย ทั้งประเด็น ศาลรัฐธรรมนูญตัดสินยุบพรรคก้าวไกล และ เศรษฐา ทวีสิน ถูกปลดจากตำแหน่งนายกรัฐมนตรี แพทองธาร ชินวัตร ได้รับเลือกเป็นนายกรัฐมนตรี รวมไปถึงทักษิณ ชินวัตร ได้รับอภัยโทษ และเฉพาะเจาะจงลงไปก็กรณี เสรีพิศุทธิ์ เตมียาเวส ออกมาให้สัมภาษณ์ว่าได้เข้าพบกับทักษิณ ชินวัตร ที่ชั้น 14 โรงพยาบาลตำรวจ (ตาราง 13)

ตาราง 13 การจำลองหัวข้อที่ปรากฏขึ้นหลังการเลือกตั้งในช่วงที่ 3 ของแพลตฟอร์ม Facebook

Cluster	Top Words	Document Count	Topic Gloss
3	ทักษิณ, เมิง, นายก, พ่อ, เสรีพิศุทธิ์, แม่, ยา, ไม่ได้, 2, ลุง, ชัน, โคตร, ฟัง, ตาย, 14, นักโทษ, เน่า, ตา, 3, ลั่น, ก็น, สัตว์, เจอ, ชีวีต, เจ็น, เสีย, แดง, ตำรวจ, ข่าว, ดี	229	ทักษิณ ชินวัตร ได้รับอภัยโทษ และ เสรีพิศุทธิ์ เตมียาเวส ออกมาให้สัมภาษณ์ว่าได้เข้าพบกับทักษิณ ชินวัตรที่ชั้น 14 โรงพยาบาลตำรวจ
4	ประชาชน, อำนาจ, ดี, สร้าง, ขนาด, ยุบ, ไม่ได้, ไร่, เชื้อน, ประเทศ, สื่อ, หมอ, อ่าง, ชื่อ, ประชาธิปไตย, ต้า, ตัด, การเมือง, คนดี, สังคม, ปาก, น้ำ, น้ำท่วม, พัฒนา, นักการเมือง, ก้าวไกล, ความชอบธรรม, ปลา, พรรคการเมือง, นัก	189	ศาลรัฐธรรมนูญตัดสินยุบพรรคก้าวไกล
8	ทักษิณ, ชาดี, รัฐบาล, เพื่อไทย, โทง, เจ็น, 2, จันมือ, เสือแดง, ตาย, 1, ประชาชน, อำนาจ, ปชป, ประชาธิปัตย์, เขมร, ประเทศไทย, ข่าว, ผลประโยชน์, หยุต, ท่าน, ปชช, ทำลาย, เพจ, โขว, ทะเล, โลก, ช่วยกัน, กต, โคตร	185	วาทกรรม “รัฐบาลผสมข้ามชั่ว” ในกรณีพรรคประชาธิปัตย์ร่วมรัฐบาล
1	ประชาชน, สว, เสือ, รัฐบาล, ทักษิณ, ดี, พรรคการเมือง, ประเทศ, ยุบ, บ้าน, ชาวบ้าน, ไม่ได้, ข่าว, สส, กกต, โจ้, น้ำท่วม, หน้า, นายกฯ, เสีย, แดง, หลอกหลวง, ด้อมส้ม, เจ็น, ชอบ, บ่อม, ลุงบ่อม, บิ๊ก, เพื่อไทย, ลอน	184	ศาลรัฐธรรมนูญตัดสินยุบพรรคก้าวไกล และ เศรษฐา ทวีสิน ถูกปลดจากตำแหน่งนายกรัฐมนตรี แพทองธาร ชินวัตร ได้รับเลือกเป็นนายกรัฐมนตรี
7	ข่าว, 3, หนุม, ศพ, แม่, ตาย, สาว, คอมเมนต์, รก, แฟน, ดับ, ช้อง, ลั่น, เครื่องบิน, ไม่ได้, ข่า, ตัด, ออนไลน์, ทักษิณ, พ่อ, เจ็น, คา, กุณ, นัก, ชื่อ, จีน, ยิง, อุโมงค์, ชับ, อนุทิน	137	ประเด็นข่าวปัญหาสังคม
0	ข่าว, การเมือง, ข่าวการเมือง, ลั่น, 3, ธรรมา, ประเทศไทย, ทิศทาง, บาท, เจ็น, หลอก, รัสเซีย, จริยธรรม, แท้, รัฐบาล, กฎหมาย, สม, นายก, สถาบัน, ข่าวสด, ข่าวสาร, ทักษิณ, จีน, ประเทศ, ยูเครน, นายกุน, แท้, ทหาร, ตำรวจ, 2	103	ประเด็นข่าวปัญหาสังคม
5	เตรียม, คินนี้, 2030, รายการ, คลิป, เพจ, ทหาร, หนู, สิงค์, หน้า, คดี, ปลอม, ดี, ตาย, นิ, โฟ, ชอบ, กลัว, กี้, 112, 3, ลี้ภัย, ม112, เจอ, ฟ้อง, 1, บันเทิง, หนี, รัสเซีย, ชาย	91	ประเด็นข่าวปัญหาสังคม
6	เขมร, กัฟ, พม่า, 2, ชาวบ้าน, 1, ลาว, จีน, ยาดอง, ผู้ก่อการร้าย, เกล, คนไทย, พื้นที่, เจ้าหน้าที่, สื่อ, มรณะ, ปืน, ศพ, แด, คลัสเตอร์, ต่างชาติ, ต, พ่อ, อ, ทุเรียน, ตรวจ, ยิง, เสียชีวิต, ประเทศ, หนุม	90	ประเด็นข่าวปัญหาสังคม
9	เจ็น, พื้นที่, ไม่ได้, ลั่น, ประชาชน, จ, ต, คอร์รัปชัน, ท่าน, สนับสนุน, จังหวัด, 2567, รัฐมนตรี, แม่, ล้าง, บ้าน, เมือง, โครงการ, ผู้ว่าราชการ	60	Undefined

Cluster	Top Words	Document Count	Topic Gloss
	จังหวัด, นา, รัฐบาล, บาท, ระบบ, จำนวน, กราบ, เครือข่าย, ลงทะเบียน, สร้าง, จุด, สิ่งแวดล้อม		
2	คัน, สูญเสีย, รัสเซีย, ระบบ, ทหาร, รถถัง, ปืนใหญ่, ลำ, ยานเกราะ, ๑, รถ, โดรน, กระบอก, สิงหาคม, ขนส่ง, ชีปนาวุธ, ยูเครน, ข่าวด, อิสราเอล, กันภัย, ทางอากาศ, บ้อง, ตังมา, หนอนคาย, อาฆาต, จรวด, ล้างล้าง, ๒, ปาเลสไตน์, โจมตี	34	ประเด็นระดับนานาชาติ เช่น ความสัมพันธ์กับประเทศตะวันออกกลาง เหตุการณ์รัสเซีย-ยูเครน

หมายเหตุ เรียงตามยอด Document Count

กล่าวโดยสรุป การวิเคราะห์บทสนทนาออนไลน์ในช่วงหลังการเลือกตั้ง โดยใช้เทคนิค Topic Modeling เผยให้เห็นประเด็นสำคัญที่สะท้อนถึงความสนใจและความคิดเห็นของผู้ใช้งานในบริบททางการเมืองและสังคม ประเด็นที่พบได้แก่ การด่า ด้อยค่ากันของกลุ่มผู้สนับสนุนพรรคการเมือง ซึ่งเป็นการแสดงออกถึงความเป็นปฏิกิริยาและความแตกแยกในระดับกลุ่ม นอกจากนี้ยังมีการวิพากษ์วิจารณ์กลุ่มกลุ่มอนุรักษ์นิยม หรือ establishment ที่ตั้งคำถามเกี่ยวกับบทบาทและแนวคิดทางการเมืองของกลุ่มดังกล่าว การเผยแพร่ข้อมูลบิดเบือนเพื่อโจมตีฝ่ายตรงข้ามก็เป็นอีกหนึ่งประเด็นสำคัญที่พบในบทสนทนา โดยเฉพาะในรูปแบบของ Malinformation ซึ่งสะท้อนถึงความซับซ้อนของการสื่อสารในพื้นที่ออนไลน์

อีกประเด็นที่สำคัญคือการโจมตีผู้นำทางการเมืองและนักการเมือง ซึ่งครอบคลุมทุกพรรคการเมือง รวมถึงการตั้งคำถามต่อสถาบันทางการเมือง เช่น คณะกรรมการการเลือกตั้ง (กกต.) และศาลรัฐธรรมนูญ ประเด็นเหล่านี้แสดงให้เห็นถึงความไม่ไว้วางใจในระบบการเมืองในปัจจุบัน การวิพากษ์วิจารณ์นโยบายรัฐบาลและการจัดสรรทรัพยากรเป็นอีกหนึ่งหัวข้อที่ได้รับความสนใจ โดยเฉพาะในเรื่องการบริหารจัดการงบประมาณและการกระจายทรัพยากร ทั้งนี้บางหัวข้อยังมีความเชื่อมโยงกับเหตุการณ์ออฟไลน์ที่สำคัญ เช่น การแก้ไขกฎหมายมาตรา 112 การพักโทษอดีตนายกรัฐมนตรีทักษิณ ชินวัตร และการอภิปรายงบประมาณรายจ่ายปี 2567 เหตุการณ์เหล่านี้ชี้ให้เห็นถึงการสะท้อนประเด็นทางการเมืองในโลกออนไลน์ซึ่งเชื่อมโยงกับสภาพการณ์ในโลกจริง

นอกจากนี้ ยังมีประเด็นที่เกี่ยวข้องกับเหตุการณ์ระดับนานาชาติ เช่น ความสัมพันธ์ระหว่างไทยกับประเทศในตะวันออกกลาง ความขัดแย้งระหว่างรัสเซียและยูเครน และอิทธิพลของจีนในภูมิภาค เหล่านี้เป็นตัวอย่างที่สะท้อนถึงการรับรู้และถกเถียงในประเด็นที่กว้างไกลกว่าเพียงแค่การเมืองภายในประเทศ

เมื่อเปรียบเทียบเนื้อหาและลักษณะของการสนทนาระหว่างแพลตฟอร์ม Facebook และ X พบความแตกต่างที่ชัดเจน Facebook มีลักษณะการสนทนาที่เน้นการให้ข้อมูล (informative) การอภิปรายเนื้อหาในประเด็นต่าง โดยเนื้อหาส่วนใหญ่เกี่ยวข้องกับการอภิปรายในประเด็นการบริหารงานของรัฐบาล ความล้มเหลวของระบบราชการ และปัญหาสังคมและสิ่งแวดล้อม ลักษณะดังกล่าวสะท้อนให้เห็นว่าผู้ใช้งาน Facebook มักให้ความสำคัญกับการสนทนาเชิงข้อมูลมากกว่า ในทางกลับกัน X มีลักษณะการสนทนาที่แสดงออกถึงอารมณ์ความรู้สึกและความคิดเห็นส่วนบุคคลอย่างเข้มข้น (expressive/emotive) โดยมีการใช้ถ้อยคำรุนแรง การแสดงออกถึงความเป็นศัตรู และการโจมตีเชิงดูถูกอย่างชัดเจน ข้อความที่ปรากฏบน X มักสะท้อนถึงการเผชิญหน้าระหว่างผู้ใช้งานและกลุ่มการเมืองต่าง ๆ มากกว่าการแลกเปลี่ยนข้อมูลอย่างสร้างสรรค์

3.2.3 การวิเคราะห์ AI Deepfake

การวิเคราะห์ AI Deepfake โครงการใช้ Vision Model (Transformer Model) มาจำแนกภาพที่ถูกสร้างขึ้นหรือ แก้ไขโดย AI (AI Generated/AI Manipulated) หรือภาพ AI Deepfake ผลการศึกษาพบว่า

การใช้เทคโนโลยี AI Deep Fake ในบริบททางการเมืองไทย ทั้งในช่วงการเลือกตั้งและหลังการเลือกตั้ง พบว่ายังไม่มีมีการนำเทคโนโลยีดังกล่าวมาใช้อย่างแพร่หลายในแวดวงการเมือง ภาพหรือเนื้อหาที่เกี่ยวข้องกับ AI Deep Fake ที่พบมีจำนวนจำกัด โดยส่วนใหญ่มักเป็นภาพที่สร้างขึ้นจากเทคโนโลยี AI Generated Images ซึ่งปรากฏในรูปแบบที่หลากหลายมากกว่า Deep Fake ในเชิงดิจิทัล ภาพที่สร้างขึ้นจาก AI เหล่านี้มีบทบาทในช่วงการเลือกตั้ง โดยเฉพาะการนำมาใช้เพื่อผลิตภาพประกอบแคมเปญหาเสียงเพื่อดึงดูดความสนใจของประชาชน (รูป 17)



รูป 17 การใช้เทคโนโลยี AI Generated Images เพื่อผลิตภาพประกอบแคมเปญหาเสียง

แม้ว่าการใช้ AI Generated Images ในบริบทการเมืองไทยจะถูกนำมาใช้ในลักษณะสร้างสรรค์ เช่น การผลิตภาพประกอบแคมเปญหาเสียง แต่ก็พบว่ามีการใช้งานในด้านลบเช่นกัน โดยเฉพาะในลักษณะของ hate speech ตัวอย่างหนึ่งคือการสร้างภาพที่เปรียบเทียบ *มีชัย ฤชุพันธุ์* กับสัตว์ประหลาดหรือสิ่งที่ไม่ใช่มนุษย์ (รูป 18) ซึ่งเป็นการสื่อถึง demonization หรือการทำให้ผู้ถูกพาดพิงดูด้อยค่าหรือไม่น่าเชื่อถือในเชิงสัญลักษณ์ แต่อย่างไรก็ตาม การใช้งานในลักษณะนี้ยังคงพบในจำนวนน้อยมาก โดยจำนวนภาพที่ตรวจพบมีไม่ถึง 10 ภาพในช่วงที่ศึกษา



รูป 18 AI Generated Images กับการ Demonization

นอกจากนี้ อีกหนึ่งรูปแบบที่พบได้บ่อยคือการถูกล้อการเมือง หรือ Editorial Cartoons (รูป 19) ซึ่งใช้เพื่อสะท้อนความคิดเห็นและวิพากษ์วิจารณ์ประเด็นทางการเมืองในลักษณะที่สร้างสรรค์และเข้าถึงง่าย สื่อในรูปแบบนี้ยังคงมีบทบาทสำคัญมากกว่าเนื้อหาที่ใช้เทคโนโลยี AI Deep Fake โดยตรง การใช้งาน Deep Fake ที่ยังไม่แพร่หลายในบริบทนี้ อาจสะท้อนถึงความพร้อมทางเทคโนโลยีของผู้ใช้งาน รวมถึงความตระหนักถึงผลกระทบเชิงลบที่อาจเกิดขึ้นจากการใช้งาน Deep Fake ในมิติของการเมืองและสังคมไทย



รูป 19 การถูกล้อการเมือง หรือ Editorial Cartoons

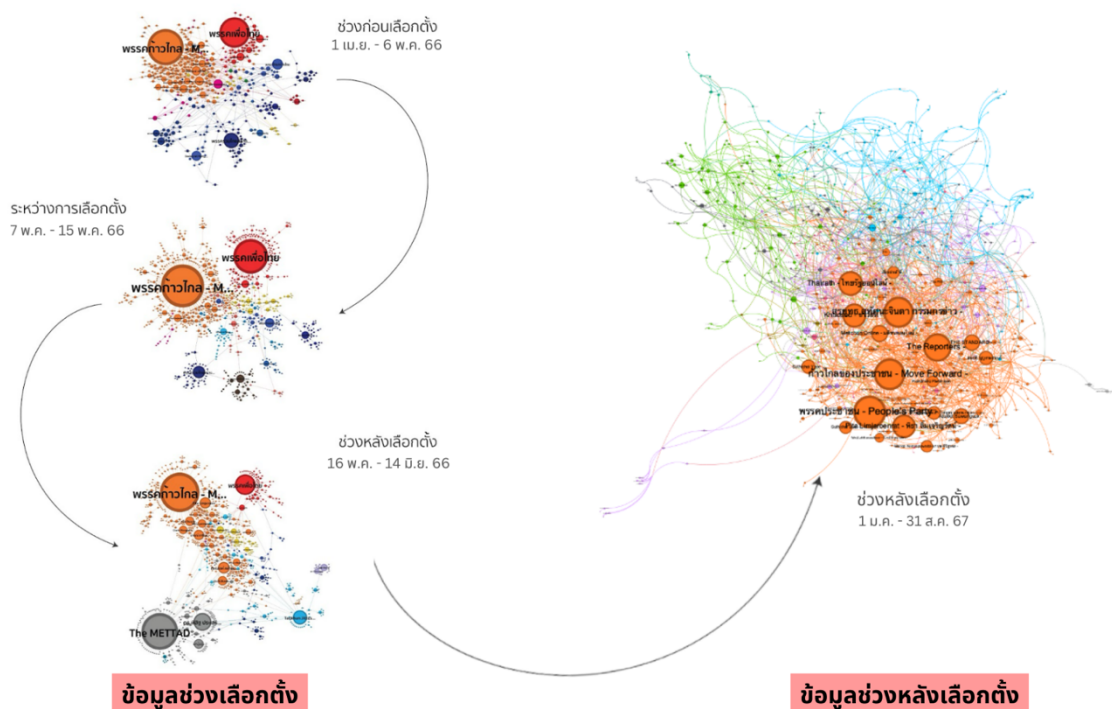
3.2.4 การวิเคราะห์เครือข่ายทางสังคม (Social Network Analysis: SNA)

แผนภาพเครือข่ายทางสังคม ประกอบด้วย **จุดวงกลม (node)** ซึ่งเป็นตัวแทนของแต่ละบัญชีผู้ใช้ ขณะที่ **วงกลม** หมายถึง ความสามารถในการเผยแพร่ข้อความของบัญชีนั้น ๆ ไปยังบัญชีอื่น โดยวงกลมที่มีขนาดใหญ่แปลว่าบัญชีนั้นมีความสามารถส่งต่อข้อความของตัวเองไปยังบัญชีอื่น ๆ ได้สูง หรืออาจกล่าวได้ว่าเป็นบัญชีที่มีอิทธิพลในการเผยแพร่ข้อความสูง และเส้นที่โยงเชื่อมกันระหว่างแต่ละวงกลมแสดงว่าบัญชีเหล่านั้นมีปฏิสัมพันธ์ระหว่างกันอย่างต่อเนื่องในช่วงเวลาที่เราศึกษา เช่น มีการแชร์ข้อความระหว่างกัน หรือตอบกลับข้อความกันอยู่หลายครั้ง โดยหากเส้นที่เชื่อมโยงกันมีระยะสั้นแปลว่าบัญชีคู่นั้นมีปฏิสัมพันธ์ระหว่างกันค่อนข้างถี่ แต่หากมีระยะห่างกันยาวแปลว่าไม่ค่อยมีปฏิสัมพันธ์ต่อกันมากนัก

เมื่อบัญชีหลายบัญชีมีปฏิสัมพันธ์ต่อกันหลายครั้งและเป็นไปอย่างต่อเนื่อง บัญชีเหล่านี้สามารถถูกจัดได้ว่าอยู่ใน **ชุมชน** หรือ **คลัสเตอร์** เดียวกัน เรามีข้อสังเกตว่าในแต่ละคลัสเตอร์หลัก บัญชีที่มีอิทธิพลในการเผยแพร่ข้อความมากที่สุด (วงกลมใหญ่ที่สุด) มักเป็นบัญชีทางการของพรรคการเมือง ทำให้เราจัดแบ่งคลัสเตอร์ตามพรรคการเมืองเป็นหลัก ซึ่งมีทั้งคลัสเตอร์พรรคก้าวไกล คลัสเตอร์พรรคเพื่อไทย คลัสเตอร์พรรครวมไทยสร้างชาติ คลัสเตอร์พรรคประชาธิปัตย์ ฯลฯ และในแต่ละคลัสเตอร์ เราพบได้ว่ามีบัญชีหลากหลายรูปแบบ ทั้งบัญชีของพรรคการเมือง นักการเมือง สำนักข่าว Influencer ทางการเมือง ผู้แสดงตัวสนับสนุนพรรคการเมืองหรือนักการเมือง ไปจนถึงบัญชีคนทั่วไป

บัญชีที่อยู่ในคลัสเตอร์ของพรรคการเมืองหนึ่ง ๆ แสดงพฤติกรรมได้สองลักษณะ ประการแรกซึ่งเป็นรูปแบบส่วนใหญ่ คือมักมีพฤติกรรมการเผยแพร่และส่งต่อข้อความในเชิงสนับสนุนพรรคการเมืองนั้น ๆ ประการที่สองคือบัญชีนั้นมีปฏิสัมพันธ์กับบัญชีอื่น ๆ ที่อยู่ในคลัสเตอร์พรรคการเมืองนั้นมากกว่าบัญชีในคลัสเตอร์อื่น ๆ แต่ไม่ได้หมายความว่าบัญชีนั้นชื่นชอบพรรคการเมืองที่ตนอยู่ในคลัสเตอร์ เช่น บัญชีหนึ่งอาจชอบเผยแพร่ข้อความด่าทอพรรคก้าวไกล แต่ถูกจัดอยู่ในคลัสเตอร์พรรคก้าวไกล เพราะพบว่ามันมีปฏิสัมพันธ์กับบัญชีอื่นในคลัสเตอร์พรรคก้าวไกลอยู่มาก โดยอาจเป็นลักษณะของการปฏิสัมพันธ์แบบการด่าทอหรือคุกคามต่อกัน เป็นต้น

1) เครือข่ายบัญชีผู้ใช้บน Facebook



รูป 20 เครือข่ายบัญชีผู้ใช้บน Facebook

ที่มา ดัดแปลงจาก DEAL. (2567). **ถอดรหัสปฏิบัติการบนข้อมูลข่าวสารในไทย #เลือกตั้ง2566**. The Digital Election Analytic Lab: DEAL.

จากแผนภาพแสดงเครือข่ายบัญชีผู้ใช้งาน Facebook (รูป 20) เห็นได้ว่า **ในข้อมูลช่วงเลือกตั้ง** คลัสเตอร์พรรคก้าวไกล (สีส้ม) คือคลัสเตอร์ที่ใหญ่ที่สุด ทั้งในช่วงก่อน ระหว่าง และหลังการเลือกตั้ง สะท้อนว่าประเด็นการสนทนาในช่วงการเลือกตั้งมักเกี่ยวข้องกับพรรคก้าวไกลเป็นส่วนใหญ่⁷ ถัดจากพรรคก้าวไกล คลัสเตอร์ที่ใหญ่รองลงมาในอันดับสองและสาม คือคลัสเตอร์พรรคเพื่อไทย (สีแดง) และรวมไทยสร้างชาติ (สีน้ำเงิน) ตามลำดับ บัญชีทางการเมืองของแต่ละพรรคมีอิทธิพลในการเผยแพร่ข้อความสูงสุดในคลัสเตอร์รอบตัวเอง โดยแฟนเพจพรรคก้าวไกลกับพรรคเพื่อไทยสามารถรักษาอิทธิพลในการเผยแพร่ข้อความได้อย่างต่อเนื่องตั้งแต่ช่วงก่อนการเลือกตั้ง ระหว่างการเลือกตั้ง ไปจนถึงหลังการเลือกตั้ง โดยเคียงคู่กันในฐานะแฟนเพจพรรคการเมืองที่มีอิทธิพลสูงสุดสองอันดับแรกในทั้งสามช่วงเวลา ซึ่งอาจมองได้ว่าเป็นผลจากความสำเร็จที่ทั้งสองพรรคทุ่มเทหาเสียงทางสื่อโซเชียลและผู้ใช้งาน Facebook อาจสนใจทั้งสองพรรคนี้เป็นทุนเดิม ขณะที่แฟนเพจพรรครวมไทยสร้างชาติมีอิทธิพลในการเผยแพร่ข้อความตามมาเป็นอันดับสาม

น่าสนใจว่าอิทธิพลในการเผยแพร่ข้อความของแต่ละพรรคการเมือง รวมถึงขนาดของคลัสเตอร์ไล่เรียงลำดับไปในทิศทางเดียวกันกับผลการเลือกตั้ง ส.ส. ในรูปแบบบัญชีรายชื่อ คือพรรคก้าวไกล พรรคเพื่อไทย และรวมไทยสร้างชาติ ได้คะแนนเสียงเป็นอันดับหนึ่ง สอง และสาม ตามลำดับ อย่างไรก็ตาม พบว่าอิทธิพลในการเผยแพร่ข้อความของแฟนเพจพรรครวมไทยสร้างชาตินั้น ค่อย ๆ ลดลงอย่างเห็นได้ชัดในช่วงหลังการเลือกตั้ง

เมื่อเก็บข้อมูลใหม่**ช่วงหลังเลือกตั้ง** พบว่า โครงสร้างโดยรวมของเครือข่ายยังคงมีลักษณะที่คล้ายคลึงกับช่วงก่อนการเลือกตั้ง อย่างไรก็ตาม จำนวนบัญชีผู้ใช้งานที่มีส่วนร่วมในเครือข่ายลดลงอย่างมีนัยสำคัญ ส่งผลให้ขนาดโดยรวมของเครือข่ายเล็กลง ปรากฏการณ์นี้สะท้อนถึงการลดลงของความเคลื่อนไหวและการมีส่วนร่วมของผู้ใช้งานในประเด็นทางการเมืองหลังผลการเลือกตั้งสิ้นสุด

ปัจจัยสำคัญที่อาจอธิบายการลดลงของความเคลื่อนไหวในเครือข่ายรวมถึงการเปลี่ยนแปลงของบริบททางการเมืองที่ทำให้ผู้ใช้งานบางส่วนหยุดแสดงความคิดเห็น การขาดประเด็นที่กระตุ้นการมีส่วนร่วมอย่างเข้มข้นในช่วงเวลานั้น และการกระจายตัวของเครือข่ายที่มากขึ้นในช่วงหลังการเลือกตั้ง เมื่อไม่มีประเด็นการพูดคุยที่ชัดเจนและเข้มข้นเทียบเท่ากับช่วงเลือกตั้ง

นอกจากนี้ Node ขนาดใหญ่ เช่น บัญชีนักการเมือง พรรคการเมือง หรือ Influencer ที่เคยมีบทบาทสำคัญในช่วงก่อนและระหว่างการเลือกตั้ง มีขนาดเล็กลงอย่างชัดเจน ซึ่งอาจเกี่ยวข้องกับการลดลงของการสื่อสารเชิงนโยบายหลังการเลือกตั้ง ในขณะเดียวกัน พบการเกิดขึ้นของ Node ใหม่ในเครือข่าย ซึ่งสามารถแบ่งได้เป็นสองประเภทได้แก่

1. บัญชีเก่าที่เปลี่ยนชื่อ การเปลี่ยนชื่อบัญชีอาจเป็นไปเพื่อเหตุผลส่วนตัว หรือในกรณีของบัญชีที่เกี่ยวข้องกับการเมือง อาจสะท้อนถึงการปรับเปลี่ยนกลยุทธ์การสื่อสารหรือการ Re-branding

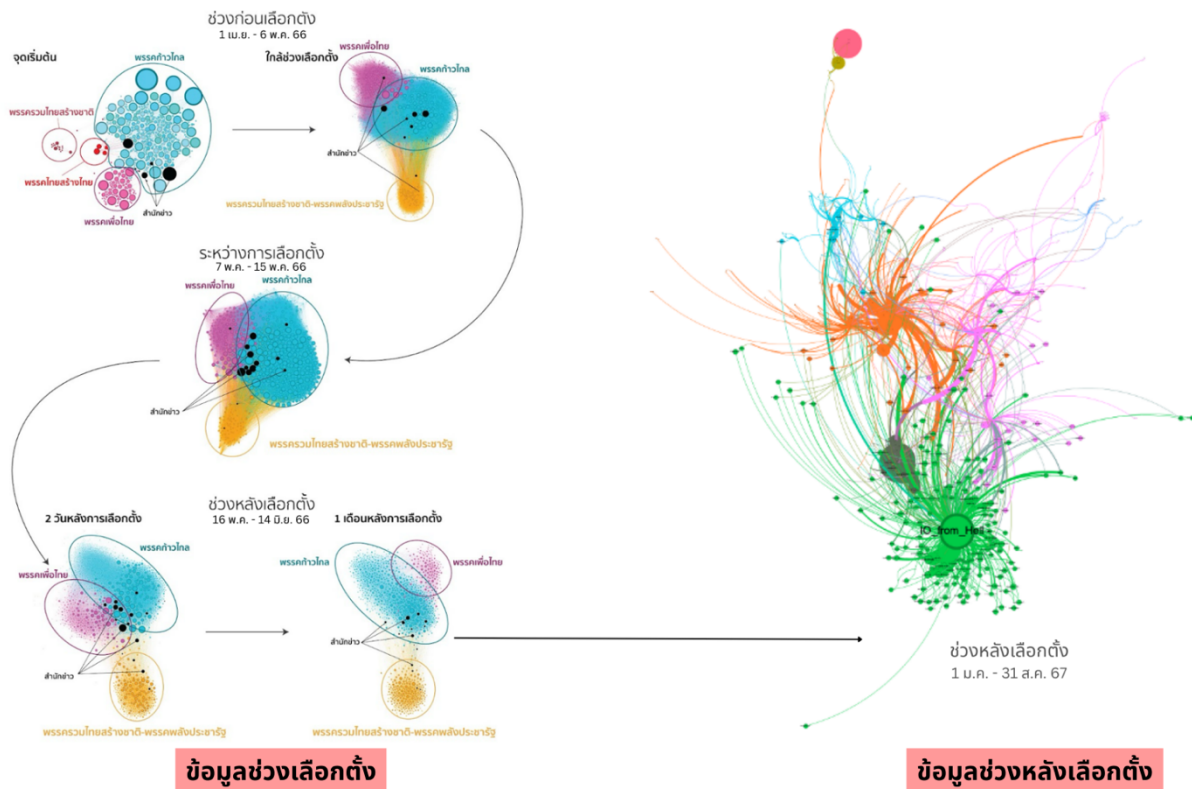
2. บัญชีใหม่ที่สร้างขึ้น บัญชีเหล่านี้อาจเป็นบัญชีที่สร้างขึ้นมาเพื่อพูดถึงประเด็นทางสังคมเฉพาะ หรือมีจุดประสงค์ในการแสดงความคิดเห็นเกี่ยวกับประเด็นนั้น ๆ โดยเฉพาะ

ปรากฏการณ์ที่น่าสังเกตอีกประการหนึ่งคือการกระจาย (fragmentation) ของเครือข่าย ซึ่งเกิดจากการลดลงของ Node สำคัญที่เคยทำหน้าที่เชื่อมโยงบัญชีผู้ใช้งานเข้าด้วยกัน การลดบทบาทของบัญชีที่เคยมีอิทธิพลสูงทำให้เครือข่ายมีลักษณะเป็นกลุ่มเล็ก ๆ มากขึ้น

อย่างไรก็ตาม โครงสร้างของเครือข่ายที่ยังคงรูปแบบเดิมแสดงให้เห็นว่ากลุ่มหรือชุมชนที่เคยมีบทบาทสำคัญในเครือข่ายยังคงรักษาสถานะของตนไว้ในระบบ แม้กิจกรรมจะลดลง การเปลี่ยนแปลงเหล่านี้สะท้อนให้เห็นถึงพลวัตของเครือข่ายทางสังคมในบริบทการเมืองหลังการเลือกตั้ง ซึ่งยังคงเป็นพื้นที่สำคัญในการแสดงออกและแลกเปลี่ยนความคิดเห็นในยุคดิจิทัล

⁷ เนื่องจากโครงการฯ ใช้ Actor-Based Scraping การที่บัญชีของคลัสเตอร์พรรคก้าวไกลมีมากที่สุด ส่วนหนึ่งเป็นเพราะบัญชีที่โครงการเก็บในช่วงแรก พบว่า พรรคก้าวไกลเป็นพรรคที่ใช้ Social Media มากที่สุด โดยมีจำนวนบัญชีทางการเมืองเยอะที่สุด ทั้ง X และ Facebook

2) เครือข่ายบัญชีผู้ใช้บน X



รูป 21 เครือข่ายบัญชีผู้ใช้บน X

ที่มา ดัดแปลงจาก DEAL. (2567). **ถอดรหัสปฏิบัติการบนข้อมูลข่าวสารในไทย #เลือกตั้ง2566**. The Digital Election Analytic Lab: DEAL.

จากแผนภาพแสดงเครือข่ายบัญชีผู้ใช้บน X (รูป 21) พบว่า เครือข่ายผู้ใช้งานบนแพลตฟอร์ม X และ Facebook ในช่วงเลือกตั้งพบความคล้ายคลึงในหลายประเด็น โดยโครงสร้างของคลัสเตอร์หลักยังคงมีลักษณะใกล้เคียงกัน ในทั้งสองแพลตฟอร์ม คลัสเตอร์ของพรรคก้าวไกลและพรรคเพื่อไทยเป็นกลุ่มใหญ่ที่สุด ตามด้วยคลัสเตอร์ที่สนับสนุนพรรครวมไทยสร้างชาติและพรรคพลังประชารัฐ ซึ่งมีการเชื่อมโยงระหว่างกันจนแยกไม่ออกเป็นคลัสเตอร์ย่อย ขนาดของคลัสเตอร์เหล่านี้มีความชัดเจนมากขึ้นเมื่อเข้าใกล้วันเลือกตั้ง เนื่องจากการมีส่วนร่วมของบัญชีผู้ใช้งานในประเด็นการเมืองที่เข้มข้นขึ้น

อย่างไรก็ตาม มีความแตกต่างที่สำคัญระหว่างสองแพลตฟอร์มในด้านอิทธิพลของบัญชีผู้ใช้งาน ใน Facebook บัญชีทางการเมืองของพรรคการเมืองมีบทบาทโดดเด่นในการเผยแพร่ข้อความ ส่วนใน X บัญชีที่มีอิทธิพลไม่ได้จำกัดเฉพาะบัญชีพรรคการเมือง แต่รวมถึงบัญชีนักการเมืองรายบุคคล สำนักข่าว และ Influencer ซึ่งได้รับการส่งเสริมโดยกลไกของแพลตฟอร์มที่มีลักษณะกระจายศูนย์มากกว่า Facebook ความแตกต่างนี้อาจมาจากลักษณะของข้อมูลที่เก็บรวบรวม โดยข้อมูลบน Facebook ถูกจำกัดให้เก็บจากบัญชีสาธารณะผ่าน CrowdTangle ขณะที่ X อนุญาตให้มีการรวบรวมข้อมูลที่หลากหลายกว่า

หลังการเลือกตั้ง การเก็บข้อมูลใหม่พบว่าโครงสร้างเครือข่ายบน X ยังคงรักษารูปแบบเดิม แต่ขนาดโดยรวมลดลงอย่างเห็นได้ชัด เนื่องจากจำนวนบัญชีผู้ใช้งานและความถี่ของกิจกรรมลดลงในช่วงหลังการเลือกตั้ง การลดลงนี้พบได้ทั้งในทั้งสองแพลตฟอร์ม แต่บน X มีความเด่นชัดกว่าในแง่ของการปรากฏตัวของบัญชีใหม่จำนวนมาก โดยเฉพาะบัญชีที่ไม่เกี่ยวข้องกับการเมืองโดยตรง เช่น บัญชีที่สร้างขึ้นเพื่อโฆษณา การลดลงของกิจกรรมและการกระจายของเครือข่ายยังชี้ให้เห็นถึงการลดลงของประเด็นการพูดคุยที่ชัดเจนและเข้มข้น รวมทั้งการเปลี่ยนชื่อบัญชีที่เกี่ยวข้องกับการเมือง อาจสะท้อนถึงการปรับเปลี่ยนกลยุทธ์การสื่อสารหรือการ Re-branding ก็ยังพบเจอใน X เช่น บัญชี

“suloveboss” ที่มีบทบาทมากในเรื่องพฤติกรรมการประสานงานกัน (Coordinated Inauthentic Behavior: CIB) ในช่วงการเลือกตั้ง เมื่อผ่านมาช่วงหลังเลือกตั้งโครงการได้ติดตามและพบว่าการเปลี่ยนชื่อ (username) ใหม่เป็น “Jeabpum” ซึ่งหลังจากเปลี่ยนชื่อแล้ว ลักษณะพฤติกรรมการโพสต์เนื้อหา ก็เปลี่ยนไปเล็กน้อย และ active น้อยลง เป็นต้น แม้ว่าความเคลื่อนไหวในเครือข่ายจะลดลง แต่คัสเตอร์หลักในทั้งสองแพลตฟอร์ม เช่น พรรคก้าวไกล พรรคเพื่อไทย และกลุ่มสนับสนุนพรรครวมไทยสร้างชาติ-พลังประชาชน ยังคงมีความต่อเนื่องในแง่ของการรักษาสถานะของกลุ่มความคิดเห็นในพื้นที่ออนไลน์ ข้อมูลนี้สะท้อนถึงความยืดหยุ่นและพลวัตของเครือข่ายทางสังคมในบริบทการเมืองดิจิทัล ที่เปลี่ยนแปลงไปตามสถานการณ์และบริบททางสังคมในแต่ละช่วงเวลา

3.3 กิจกรรมที่ 3 จัดวงเสวนาผู้มีส่วนได้ส่วนเสียและผู้ที่เกี่ยวข้อง (Stakeholder Discussion)

โครงการฯ ดำเนินการเก็บรวบรวมความคิดเห็นจากผู้มีส่วนได้ส่วนเสียและผู้ที่เกี่ยวข้อง โดยเข้าร่วมนำเสนอผลงานและร่วมเสวนา ในหัวข้อ Safe Social Media for Everyone ที่จัดขึ้นในวันที่ 29 พฤศจิกายน พ.ศ.2567 ณ โรงแรม Grand Centre Point ราชดำริ กรุงเทพฯ ซึ่งจัดขึ้นโดยมีวัตถุประสงค์เพื่อเผยแพร่ผลการวิจัยเกี่ยวกับมาตรฐานชุมชน กระบวนการจัดการเนื้อหา และปัญหาการแพร่กระจายคำพูดแสดงความเกลียดชังและข้อมูลผิดพลาดบนโซเชียลมีเดีย รวมถึงการสร้างความรู้ความตระหนักและการเรียกร้องให้แพลตฟอร์มรับผิดชอบต่อเนื้อหาที่เผยแพร่

ผู้เข้าร่วมหลักมาจากสองภาคส่วน คือ ภาคประชาสังคม และสื่อมวลชน ซึ่งได้ให้ความคิดเห็นและข้อเสนอแนะเกี่ยวกับมาตรฐานชุมชน กระบวนการจัดการเนื้อหา และปัญหบนโซเชียลมีเดีย และการกระตุ้นให้เกิดการสนับสนุนนโยบายการจัดการเนื้อหาที่โปร่งใสและมีประสิทธิภาพบนแพลตฟอร์มโซเชียลมีเดีย ดังนี้

ความคิดเห็น

1. โซเชียลมีเดียทำให้เกิดการแพร่กระจายของข้อมูลที่ผิดพลาด และการจัดการข้อมูลที่ไม่ทันเวลาจะส่งผลกระทบต่อการใช้และเข้าใจของประชาชนในวงกว้าง
2. ความแตกต่างในบริบททางสังคม วัฒนธรรม และภาษาอาจทำให้มาตรฐานเดียวกันไม่เพียงพอในการจัดการกับปัญหาของแต่ละประเทศ/พื้นที่ ที่มีบริบทเฉพาะ
3. ผู้สูงอายุและคนในชุมชนที่ขาดความรู้ความเข้าใจเกี่ยวกับการใช้อินเทอร์เน็ตและโซเชียลมีเดีย มักตกเป็นเหยื่อของมิจฉาชีพและการหลอกลวงออนไลน์ สะท้อนถึงช่องว่างทางความรู้ที่ทำให้ผู้สูงอายุมีความเสี่ยงสูงต่อการถูกหลอกลวง เช่น การหลอกให้โอนเงินหรือเสียทรัพย์สิน เป็นต้น รวมถึงปัญหาการเข้าถึงข้อมูลที่ถูกต้องและเครื่องมือในการป้องกันตัวเองจากภัยออนไลน์

ข้อเสนอแนะ

1. แพลตฟอร์มโซเชียลมีเดียควรมีบทบาทเชิงรุกมากขึ้นในการป้องกันและจัดการปัญหาเกี่ยวกับข้อมูลผิดพลาดและคำพูดแสดงความเกลียดชัง และให้ความสำคัญของการใช้ข้อมูลเชิงลึกที่สามารถปรับใช้ได้จริง เช่น ใช้สำหรับการพัฒนามาตรการใหม่ๆ หรือการเพิ่มความโปร่งใสในนโยบายของแพลตฟอร์ม
2. แพลตฟอร์มควรปรับปรุงระบบอัลกอริทึมให้สอดคล้องกับบริบทเฉพาะของแต่ละพื้นที่หรือวัฒนธรรม เพื่อให้สามารถจัดการข้อมูลและปัญหาที่เกิดขึ้นในแต่ละภูมิภาคได้อย่างมีประสิทธิภาพ
3. การแก้ไขปัญหาข้อมูลผิดพลาดและคำพูดแสดงความเกลียดชังบนโซเชียลมีเดียไม่สามารถทำได้เพียงฝ่ายเดียว จำเป็นต้องมีความร่วมมือและพัฒนาความเป็นเครือข่ายระหว่างภาคส่วนต่าง ๆ เช่น หน่วยงานรัฐ องค์กรประชาสังคม ภาควิชาการ และแพลตฟอร์มโซเชียลมีเดีย เพื่อพัฒนาแนวทางที่ครอบคลุมและมีผลกระทบเชิงบวกในระยะยาว
4. การส่งเสริมให้ประชาชนรู้เท่าทันสื่อ ต้องมุ่งเน้นการสร้างความรู้และความเข้าใจในการใช้สื่อออนไลน์อย่างปลอดภัยต่อตัวเองและไม่เป็นอันตรายต่อผู้อื่น โดยเฉพาะการเข้าถึงกลุ่มเปราะบางในสังคม เช่น ผู้สูงอายุ และคนในชุมชนที่อาจขาดความรู้เพียงพอในการป้องกันตนเองจากข้อมูลผิดพลาดหรือบิดเบือน ทั้งนี้ บทบาทของสื่อมวลชนมีความสำคัญอย่างยิ่งในการขยายข้อมูลที่ถูกต้องสู่สาธารณะอย่างกว้างขวาง และการพัฒนาแนวทางตรวจสอบข้อเท็จจริง (fact-checking) ที่มีประสิทธิภาพก็เป็นสิ่งจำเป็นในการสร้างความน่าเชื่อถือในข้อมูลที่เผยแพร่ นอกจากนี้ ยังควรสนับสนุนแรงจูงใจและการปรับตัวของ

สื่อมวลชนให้สามารถรับมือกับการเปลี่ยนแปลงของเทคโนโลยี พร้อมทั้งใช้สื่อเป็นเครื่องมือในการให้ความรู้ในระดับชุมชน เพื่อเสริมสร้างความรู้ความเข้าใจในกลุ่มประชาชนและลดการตกเป็นเหยื่อของข้อมูลที่ผิดพลาดในอนาคต

5. ในการส่งเสริมให้โซเชียลมีเดียเป็นพื้นที่ที่ส่งเสริมเสรีภาพในการแสดงออก ควรเรียกร้องให้มีกฎหมายป้องกันการฟ้องปิดปาก (SLAPP หรือ Strategic Lawsuits Against Public Participation) และการจัดการคดีที่มีเจตนากลั่นแกล้ง ผู้เข้าร่วมเสนอให้มีกฎหมายป้องกันการใช้ SLAPP เพื่อไม่ให้รัฐบาลหรือหน่วยงานอื่นๆ ใช้อำนาจทางกฎหมายในการปิดกั้นเสียงของผู้วิจารณ์หรือผู้ที่มีความคิดเห็นต่างออกไป รวมถึงการมีกลไกที่ช่วยป้องกันหรือกรองคดีที่มีเจตนากลั่นแกล้งโดยไม่สุจริต เช่น การฟ้องเพื่อกดดันหรือข่มขู่ เป็นต้น
6. การออกแบบระบบ(การสื่อสารทางอินเทอร์เน็ตและโซเชียลมีเดียหรือมาตรฐานชุมชน?)ที่ประเมินความเสี่ยงอย่างครอบคลุม โดยเฉพาะการปรับให้เหมาะสมกับอายุและบริบทของผู้ใช้ เช่น วัยรุ่นที่มักตกเป็นกลุ่มเป้าหมายของข้อมูลผิดพลาดหรือภัยออนไลน์ ระบบเหล่านี้ควรสามารถระบุความเสี่ยงและตอบสนองได้อย่างรวดเร็ว เพื่อป้องกันผลกระทบเชิงลบ นอกจากนี้ การสื่อสารควรปรับให้เหมาะสมกับกลุ่มเป้าหมายโดยใช้ภาษาที่เข้าถึงง่ายและเหมาะกับวัยรุ่นหรือชุมชนเฉพาะ เพื่อลดความเข้าใจผิดและเสริมสร้างความต้านทานต่อคำพูดแสดงความเกลียดชัง รวมถึงอาจมี internet expertise ในชุมชน เพื่อให้ความรู้และเพิ่มประสิทธิภาพในการสื่อสารการใช้งานอินเทอร์เน็ตและโซเชียลมีเดียอย่างปลอดภัย

4. อภิปรายผลการศึกษาและข้อเสนอแนะ

ในบทนี้ โครงการฯ จะอภิปรายผลการศึกษาเพื่อตอบวัตถุประสงค์ทั้ง 2 ข้อ คือ 1) เพื่อระบุช่องว่างในการนำมาตรฐานชุมชนของแพลตฟอร์มที่เกี่ยวข้องกับความเชื่อมั่นและความปลอดภัยออนไลน์ไปปฏิบัติ รวมถึงนโยบายการกำกับดูแลอินเทอร์เน็ตในช่วงการเลือกตั้งและหลังการเลือกตั้ง ซึ่งนำไปสู่การเพิ่มขึ้นของข้อมูลที่มีแนวโน้มเป็นอันตราย และข้อมูลบิดเบือนในช่วงหลังการเลือกตั้งของไทยจนถึงปัจจุบัน และ 2) เพื่อสำรวจแนวทางที่เป็นไปได้ในการกำกับดูแลเนื้อหาออนไลน์ที่เหมาะสมหรือการแทรกแซงที่มีประสิทธิภาพจากแพลตฟอร์ม รวมถึงบทบาทของสื่อในการส่งเสริมการสื่อสารออนไลน์ที่อิงข้อเท็จจริงและปราศจากความรุนแรง ในขณะเดียวกันก็ยังรักษาสมดุลของเสรีภาพในการแสดงออก นอกจากนี้ยังจะนำเสนอข้อสังเกตจากทีมวิจัย และข้อเสนอและต่อองค์การต่างๆ โดยมีรายละเอียดดังต่อไปนี้

4.1 ช่องว่างในการนำมาตรฐานชุมชนของแพลตฟอร์มที่เกี่ยวข้องกับความเชื่อมั่นและความปลอดภัยออนไลน์ไปปฏิบัติ

จากการวิเคราะห์ข้อมูลเชิงบรรยาย (descriptive analysis) เกี่ยวกับข้อความที่มีแนวโน้มเป็นอันตรายและข้อมูลบิดเบือนในช่วงการเลือกตั้งและหลังการเลือกตั้งในบริบทการเมืองไทย พบว่าแพลตฟอร์มต่าง ๆ เช่น Facebook (Meta), X, และ YouTube (Google) มีแนวทางปฏิบัติที่แตกต่างกันในการจัดการเนื้อหา ซึ่งมีช่องว่างสำคัญที่ส่งผลกระทบต่อความเชื่อมั่นและความปลอดภัยในพื้นที่ออนไลน์ โดยจากผลการรวบรวมนโยบายและมาตรฐานความปลอดภัยของแต่ละแพลตฟอร์ม โครงการฯ ได้สรุปช่องว่างในการนำมาตรฐานชุมชนของแพลตฟอร์มที่เกี่ยวข้องกับความเชื่อมั่นและความปลอดภัยออนไลน์ไปปฏิบัติ รวมถึงนโยบายการกำกับดูแลอินเทอร์เน็ตในช่วงการเลือกตั้งและหลังการเลือกตั้ง ซึ่งนำไปสู่การเพิ่มขึ้นของข้อมูลที่มีแนวโน้มเป็นอันตรายและข้อมูลบิดเบือนในช่วงหลังการเลือกตั้งของไทยจนถึงปัจจุบัน ที่สถานการณ์ยังมีความอ่อนไหวทางการเมือง ดังนี้

หนึ่งในช่องว่างที่ชัดเจนคือ การใช้ **การติดฉลาก (labeling) แทนการลบเนื้อหาโดยทันที** บน Facebook และ X เนื้อหาที่ถูกระบุว่าเป็นข้อมูลบิดเบือนหรือมีแนวโน้มเป็นอันตรายมักถูกติดฉลากเตือนผู้ใช้งานแทนที่จะถูกลบออก เช่น บน Facebook เนื้อหาจะถูกตรวจสอบเพิ่มเติมโดยบุคคลที่เป็นพันธมิตรที่เชื่อถือได้ (trusted partner) หรือผู้ตรวจสอบข้อเท็จจริง (fact-checkers) และยังคงเผยแพร่ต่อไปได้แม้จะมีคำเตือน ในขณะที่ X ใช้การติดฉลากเตือนเช่นกัน แต่ยังมี การปรับลดการมองเห็นของเนื้อหาผ่านกลไกการลดความสามารถในการค้นพบ (shadow banning) อย่างไรก็ตาม ทั้งสองแพลตฟอร์มยังคงอนุญาตให้เนื้อหาอยู่ในระบบ ซึ่งแตกต่างจาก YouTube ที่ใช้การลบเนื้อหาที่ตรวจพบว่าจะผิดมาตรฐานชุมชนทันที

ประเด็นต่อมาคือ **ความล่าช้าในกระบวนการตรวจสอบเนื้อหา** ซึ่งเป็นผลมาจากการพึ่งพาบุคคลในการพิจารณาความถูกต้องของเนื้อหา บน Facebook และ X ผู้ใช้งานที่โพสต์เนื้อหาที่ถูกติดฉลากหรือถูกลบสามารถยื่นอุทธรณ์ได้ ซึ่งกระบวนการตรวจสอบเหล่านี้มักใช้เวลานาน ส่งผลให้เนื้อหาที่มีแนวโน้มเป็นอันตรายสามารถแพร่กระจายในวงกว้างในระหว่างที่รอผลการตรวจสอบ โดยเฉพาะในช่วงเวลาสำคัญ เช่น การเลือกตั้ง ความล่าช้านี้อาจทำให้เกิดความเสียหายร้ายแรงต่อความเชื่อมั่นในข้อมูล

ยิ่งไปกว่านั้น **การปรับอัลกอริทึมให้การเข้าถึงยากขึ้น (Algorithm Adjustment) แทนการลบเนื้อหา** X ใช้ อัลกอริทึมในการลดการแสดงผลเนื้อหาบิดเบือน โดยเนื้อหาเหล่านี้จะปรากฏยากขึ้นในการค้นหา แต่ยังคงอยู่ในระบบและสามารถค้นพบได้หากผู้ใช้ตั้งใจค้นหา การดำเนินการเช่นนี้ช่วยลดการมองเห็นข้อมูลบิดเบือน แต่ไม่ได้กำจัดเนื้อหาออกจากแพลตฟอร์ม

นอกจากนี้ ในบริบทของภาษาไทยและการเมืองไทย พบว่า **มาตรฐานชุมชนของแพลตฟอร์มอาจไม่ครอบคลุมบริบทเฉพาะ** เช่น ข้อความที่แฝงนัยสำคัญในเชิงลบต่อความไว้วางใจในสังคม หรือข้อความที่ส่งผลกระทบต่อกระบวนการประชาธิปไตย บางครั้งข้อความเหล่านี้อาจไม่ถูกระบุว่าเป็น Hate Speech ตามมาตรฐานสากล แต่กลับมีอิทธิพลสำคัญในบริบทการเปลี่ยนผ่านทางการเมืองของไทย เช่น การเผยแพร่เรื่องเล่าที่บั่นทอนความเชื่อมั่นในกระบวนการทางกฎหมายและสถาบันการเมือง

ประเด็นสุดท้าย พบว่า แพลตฟอร์ม ยังไม่มีนโยบายหรือมาตรการที่ชัดเจนเกี่ยวกับการจัดการ AI Deepfake โดยเฉพาะบริบทที่เกี่ยวข้องกับการเมืองและการเลือกตั้ง

ดังนั้นเพื่อเชื่อมโยงเสรีภาพในการแสดงออก (freedom of expression) และความปลอดภัยในโลกออนไลน์ (online trust and safety) จำเป็นต้องมีการปรับปรุงมาตรฐานชุมชนและแนวทางปฏิบัติของแพลตฟอร์มให้โปร่งใสและตอบสนองต่อบริบทเฉพาะ โดยการพัฒนามาตรฐานชุมชนควรได้รับความร่วมมือจากทุกภาคส่วน รวมถึงการส่งเสริมการรู้เท่าทันสื่อ (media literacy) เพื่อให้ผู้ใช้งานสามารถประเมินและจัดการข้อมูลได้อย่างเหมาะสม นอกจากนี้ ระบบการจัดการเนื้อหาของแพลตฟอร์มควรมีความยืดหยุ่นเพียงพอที่จะรองรับความซับซ้อนทางสังคมและการเมืองในแต่ละบริบท เพื่อป้องกันการแพร่กระจายของเนื้อหาที่เป็นอันตรายและส่งเสริมสภาพแวดล้อมออนไลน์ที่ปลอดภัยและน่าเชื่อถือมากยิ่งขึ้น

4.2 แนวทางการกำกับดูแลเนื้อหาออนไลน์หรือการแทรกแซงที่เป็นไปได้

ในหัวข้อนี้ โครงการฯ มุ่งอภิปรายผลในประเด็นการสำรวจแนวทางที่เป็นไปได้ในการกำกับดูแลเนื้อหาออนไลน์ที่เหมาะสมหรือการแทรกแซงที่มีประสิทธิภาพจากแพลตฟอร์ม รวมถึงบทบาทของสื่อในการส่งเสริมการสื่อสารออนไลน์ที่อิงข้อเท็จจริงและปราศจากความรุนแรง ในขณะเดียวกันก็รักษาสมดุลของเสรีภาพในการแสดงออก จากกระบวนการศึกษาและผลการศึกษาในครั้งนี้ พบว่า

การกำกับดูแลเนื้อหาออนไลน์ในยุคปัจจุบันเผชิญความท้าทายที่สำคัญในการรักษาสมดุลระหว่างการป้องกันการแพร่กระจายของเนื้อหาที่เป็นอันตรายและการคงไว้ซึ่งเสรีภาพในการแสดงออก (freedom of expression) ในการอภิปรายนี้ ชื่อนำเสนอแนวทางที่เป็นไปได้สำหรับแพลตฟอร์มออนไลน์ รวมถึงบทบาทของสื่อในการสนับสนุนการสื่อสารที่อิงข้อเท็จจริงและปราศจากความรุนแรง

ประการแรก การใช้เทคโนโลยีควบคู่กับการมีส่วนร่วมของมนุษย์ (Human-AI Collaboration) เป็นแนวทางสำคัญในการกำกับดูแลเนื้อหา AI สามารถช่วยวิเคราะห์และประมวลผลข้อมูลจำนวนมากได้อย่างรวดเร็ว เช่น การระบุคำสำคัญหรือรูปแบบเนื้อหาที่เกี่ยวข้องกับ Hate Speech หรือ Malinformation ซึ่งจะช่วยให้เพิ่มประสิทธิภาพในกระบวนการตรวจสอบ อย่างไรก็ตาม การใช้ AI เพียงอย่างเดียวอาจไม่เพียงพอ เนื่องจาก AI มักขาดความเข้าใจในบริบททางสังคมและวัฒนธรรมที่ซับซ้อน การให้มนุษย์ทำหน้าที่ตรวจสอบขั้นสุดท้ายจะช่วยลดข้อผิดพลาดและเพิ่มความแม่นยำของกระบวนการตรวจจับและจัดการเนื้อหา

ประการที่สอง การสร้างกลไกการเฝ้าระวัง (Monitoring Mechanism) ที่สามารถแจ้งเตือนและตรวจจับการแพร่กระจายของเนื้อหาที่อาจเป็นอันตรายแบบเรียลไทม์ เป็นสิ่งจำเป็น โดยเฉพาะในช่วงเหตุการณ์สำคัญ เช่น การเลือกตั้งหรือวิกฤตทางสังคม กลไกดังกล่าวควรได้รับการพัฒนาโดยอาศัยความร่วมมือจากหลายภาคส่วน รวมถึงตัวแทนจากภาคประชาชน กลุ่มการเมือง สื่อมวลชน และผู้กำกับดูแลแพลตฟอร์มออนไลน์ เพื่อให้มั่นใจว่ามุมมองและความคิดเห็นที่หลากหลายได้รับการพิจารณา การมีส่วนร่วมนี้จะช่วยสร้างความไว้วางใจและความโปร่งใสในกระบวนการกำกับดูแล มากกว่าการพึ่งพากลไกที่เน้นการควบคุมหรือกำกับดูแลเพียงฝ่ายเดียว

ประการสุดท้าย บทบาทของสื่อในการส่งเสริมการสื่อสารออนไลน์ที่อิงข้อเท็จจริงและปราศจากความรุนแรงก็มีความสำคัญ สื่อสามารถทำหน้าที่เป็นตัวกลางในการถ่ายทอดข้อมูลที่ถูกต้อง และช่วยสนับสนุนการสร้างภูมิคุ้มกันทางความคิด (cognitive resilience) ให้กับประชาชนผ่านการส่งเสริมการรู้เท่าทันสื่อ (media literacy) ทั้งนี้ การสร้างสมดุลระหว่างการปกป้องเสรีภาพในการแสดงออกและการรับรองความปลอดภัยในพื้นที่ออนไลน์จำเป็นต้องมีการพัฒนานโยบายที่ยืดหยุ่นและตอบสนองต่อบริบทที่เปลี่ยนแปลงไป

ดังนั้น การกำกับดูแลเนื้อหาออนไลน์ที่มีประสิทธิภาพต้องอาศัยการบูรณาการระหว่างเทคโนโลยี มนุษย์ และความร่วมมือของหลายภาคส่วน การออกแบบระบบที่โปร่งใส มีความหลากหลาย และเคารพต่อบริบทเฉพาะจะช่วยสร้างสมดุลระหว่างการควบคุมเนื้อหาที่เป็นอันตรายและการรักษาสិทธิพื้นฐานของประชาชนในพื้นที่ออนไลน์ได้อย่างยั่งยืน

4.3 ข้อสังเกตจากทีมวิจัย

4.3.1 การเพิ่มขึ้นของสัดส่วนข้อความที่เป็นอันตรายทั้ง Hate Speech และ Malinformation ในช่วงหลังการเลือกตั้ง

การศึกษาพบว่าการเพิ่มขึ้นของสัดส่วนข้อความที่เป็นอันตรายทั้ง Hate Speech และ Malinformation บนแพลตฟอร์ม Facebook และ X ในช่วงหลังการเลือกตั้ง โดยมีข้อสังเกตสำคัญดังนี้

ในกรณีของ Facebook สัดส่วนข้อความที่เป็นอันตรายเพิ่มขึ้นอย่างมีนัยสำคัญจาก 2.85% เป็น 7.6% ซึ่งเพิ่มขึ้นเกือบ 3 เท่า การเพิ่มขึ้นนี้อาจสะท้อนถึงข้อจำกัดของนโยบายและกลไกการกลั่นกรองเนื้อหาของแพลตฟอร์ม เช่น ระบบตรวจจับอัตโนมัติที่ยังไม่สามารถจัดการกับข้อความที่ดัดแปลงได้ เช่น การใช้ตัวอักษรภาษาอังกฤษ ตัวเลข สัญลักษณ์อีโมจิ การเว้นวรรค หรือการขึ้นบรรทัดใหม่เพื่อหลีกเลี่ยงการตรวจจับ นอกจากนี้ แนวโน้มดังกล่าวอาจสะท้อนถึงความตึงเครียดทางการเมืองที่ยังคงมีอยู่หลังการเลือกตั้ง ซึ่งกระตุ้นให้เกิดการเผยแพร่ข้อความโจมตีคู่แข่งทางการเมืองอย่างต่อเนื่อง การเพิ่มขึ้นของข้อความลักษณะนี้อาจส่งผลกระทบต่อบรรยากาศทางสังคม ความสัมพันธ์ระหว่างกลุ่มการเมือง และความเชื่อมั่นในระบอบประชาธิปไตย

ในทางกลับกัน บนแพลตฟอร์ม X พบว่ามีสัดส่วนข้อความที่เป็นอันตรายลดลงจาก 20.67% ในช่วงการเลือกตั้ง เหลือ 17.06% ในช่วงหลังการเลือกตั้ง แม้ว่าจะลดลง แต่สัดส่วนดังกล่าวยังคงสูงกว่า Facebook อย่างมาก (17.06% เทียบกับ 7.6%) ข้อมูลนี้สะท้อนว่าแพลตฟอร์ม X อาจยังคงเผชิญกับข้อจำกัดในการจัดการเนื้อหาที่เป็นอันตราย เช่น ความไม่เข้มงวดของมาตรฐานชุมชน หรือความซับซ้อนของข้อความที่ถูกเผยแพร่โดยผู้ใช้งานที่อาศัยแพลตฟอร์มนี้เป็นพื้นที่ในการผลิตข้อความที่มีลักษณะเป็นอันตรายอย่างต่อเนื่อง

เมื่อพิจารณาแนวโน้มโดยรวม พบว่าหลังการเลือกตั้ง มีสัดส่วนข้อความที่เป็นอันตรายสูงขึ้น แม้ว่าจำนวนบัญชีที่แอกทีฟจะลดลงก็ตาม ข้อมูลจากการทำ Topic Modeling ชี้ว่า ในช่วงการเลือกตั้ง ผู้ใช้งานมุ่งเน้นการผลิตข้อความเพื่อโน้มน้าวผู้สืบทอดเลือกตั้งและโจมตีคู่แข่งทางการเมือง แต่ในช่วงหลังการเลือกตั้ง การใช้ข้อความที่เป็นอันตรายเปลี่ยนเป้าหมายไปที่การบั่นทอนความน่าเชื่อถือของสถาบันพรรคการเมืองและบุคคล รวมถึงการสร้างความขัดแย้งในสังคม การลดลงของบัญชีที่แอกทีฟอาจทำให้สัดส่วนของข้อความที่เป็นอันตรายเพิ่มขึ้น เนื่องจากบัญชีที่ยังคงอยู่มักเป็นบัญชีที่มุ่งเน้นการเผยแพร่ข้อความลักษณะนี้อย่างต่อเนื่อง

ข้อสังเกตนี้ชี้ให้เห็นถึงความจำเป็นในการพัฒนากลไกการจัดการเนื้อหาออนไลน์ที่มีประสิทธิภาพและตอบสนองต่อพฤติกรรมผู้ใช้งานที่เปลี่ยนแปลงไปในช่วงเวลาต่าง ๆ รวมถึงการเสริมสร้างความเข้มแข็งของมาตรฐานชุมชนที่สามารถรองรับความซับซ้อนของข้อความและบริบทเฉพาะในแต่ละแพลตฟอร์มและสถานการณ์ได้อย่างยั่งยืน

4.3.2 ข้อความ Hate Speech ที่ตรวจพบในบริบทการเมืองไทยกับ “วัฒนธรรมชาวเน็ตไทย”

จากการตรวจสอบข้อความที่ถูกระบุว่าเป็น Hate Speech โดยโมเดล AI พบข้อสังเกตหลายประการที่สะท้อนถึงวัฒนธรรมออนไลน์ของผู้ใช้งานชาวไทย และความท้าทายในการจัดการเนื้อหาในภาษาที่มีลักษณะซับซ้อนอย่างภาษาไทย

ในกรณีของ Hate Speech ระดับ 5 ซึ่งควรถูกระบุว่าเป็นข้อความที่มีลักษณะยุยงหรือสนับสนุนความรุนแรงอย่างชัดเจน การวิเคราะห์โดยโมเดลกลับพบการตรวจจับคำหรือวลีที่ไม่เข้าข่าย hate speech จริง เช่น ชื่อคน วลีเปรียบเปรย หรือคำอุทาน ตัวอย่างเช่น “เฒ่าหัว” “ยังได้ยัง” หรือ “คมแบบปาดคอคนได้” สะท้อนถึงความท้าทายของ AI ในการทำความเข้าใจบริบทเฉพาะของภาษาไทย รวมถึงวัฒนธรรมย่อย (subculture) และการใช้คำในเชิงล้อเลียนหรือประชดประชันของชาวเน็ตไทย การตรวจจับลักษณะนี้ชี้ให้เห็นว่าภาษาไทยมีความซับซ้อน ทั้งในด้านโครงสร้างไวยากรณ์และบริบทการใช้งาน โดยเฉพาะอย่างยิ่งเมื่อผู้ใช้งานออนไลน์มักสร้าง “มิม” หรือวลีที่ได้รับความนิยมในกลุ่มสังคมออนไลน์ ซึ่งอาจดูรุนแรงในความหมายผิวเผิน แต่ไม่ได้มีเจตนาร้าย

อีกประเด็นสำคัญคือการพบ Hate Speech ระดับ 1 ซึ่งประกอบด้วยข้อความหยาบคายและการดูถูกสติปัญญาในอัตราที่สูง ทั้งบน Facebook และ X โดยเฉพาะใน X ซึ่งมีสัดส่วนถึง 89.81% ของข้อความที่ถูกระบุว่าเป็น Hate Speech การใช้ถ้อยคำลักษณะนี้สะท้อนถึงวัฒนธรรมชาวเน็ตไทยที่ถ้อยคำหยาบคายและการวิพากษ์วิจารณ์อย่างตรงไปตรงมาเป็นเรื่องปกติในบทสนทนาออนไลน์ นอกจากนี้ ลักษณะของแพลตฟอร์ม X ซึ่งสนับสนุนการ

แสดงออกเชิงอารมณ์อย่างรวดเร็วและเข้มข้น อาจส่งเสริมพฤติกรรมการเลียนแบบในกลุ่มผู้ใช้งานและการแสดงความคิดเห็นที่ไม่ผ่านการกลั่นกรอง

ในทางกลับกัน การเพิ่มขึ้นของ Hate Speech ระดับ 3 บน Facebook หลังการเลือกตั้ง (15.85%) อาจสะท้อนถึงการเปลี่ยนแปลงของลักษณะการโจมตีทางวาจา จากการใช้ถ้อยคำหยาบคายทั่วไปในระดับ 1 ไปสู่คำพูดที่มีลักษณะลดทอนคุณค่าและเหมารวมกลุ่มเป้าหมาย ซึ่งแสดงถึงเจตนาการโจมตีที่มีความซับซ้อนและลึกซึ้งมากขึ้น ลักษณะนี้อาจเชื่อมโยงกับบริบทของความขัดแย้งทางการเมืองในสังคมไทย ซึ่งการแบ่งแยกทางอุดมการณ์และความตึงเครียดในโลกออนไลน์ถูกสะท้อนออกมาในบทสนทนาออนไลน์

การตรวจพบ Hate Speech ระดับ 4 ในจำนวนที่น้อยสะท้อนให้เห็นว่าการแสดงออกถึงการยุยงหรือสนับสนุนความรุนแรงในระดับสูงมีอยู่ แต่ระบบกลั่นกรองของแพลตฟอร์มอาจสามารถตรวจจับและจัดการข้อความประเภทนี้ได้ในระดับหนึ่ง อย่างไรก็ตาม ความท้าทายที่สำคัญคือการจัดการกับ Hate Speech ในระดับที่ต่ำกว่า ซึ่งอาจไม่ถูกมองว่าเป็นภัยร้ายแรงในทันที แต่สามารถสะสมผลกระทบในเชิงลบต่อบรรยากาศของสังคมออนไลน์ในระยะยาว

ข้อสังเกตนี้ชี้ให้เห็นถึงความจำเป็นในการพัฒนากลไกที่สามารถจัดการกับ Hate Speech อย่างมีประสิทธิภาพ โดยคำนึงถึงบริบทเฉพาะของภาษาไทยและวัฒนธรรมออนไลน์ เช่น การสร้างเครื่องมือที่ช่วยให้ผู้ใช้งานตระหนักถึงผลกระทบของคำพูด และการสนับสนุนการแสดงออกที่สร้างสรรค์มากกว่าการควบคุมหรือจำกัดเสรีภาพโดยตรง แพลตฟอร์มควรปรับปรุงมาตรฐานชุมชนให้ครอบคลุมความซับซ้อนของภาษาไทย และส่งเสริมความร่วมมือระหว่างเทคโนโลยีและมนุษย์เพื่อเพิ่มความแม่นยำในการจัดการเนื้อหาในระยะยาว

4.3.3 ข้อความ Malinformation กับผลกระทบต่อความไว้วางใจในสังคม

การศึกษาพบแนวโน้มที่แตกต่างกันในสัดส่วนของข้อความที่จัดเป็น Malinformation ระหว่างแพลตฟอร์ม X และ Facebook ในช่วงหลังการเลือกตั้ง โดยใน X สัดส่วนของ Malinformation ลดลงจาก 10.60% เหลือ 5.01% ขณะที่ใน Facebook สัดส่วนกลับเพิ่มขึ้นจาก 2.26% เป็น 4.32%

การลดลงของ Malinformation ใน X อาจสะท้อนถึงการลดการผลิตหรือเผยแพร่เนื้อหาบิดเบือนโดยผู้ใช้งานเอง ซึ่งอาจเกิดจากความสนใจในประเด็นทางการเมืองที่ลดลงหลังการเลือกตั้ง นอกจากนี้ การลดลงดังกล่าวอาจเกี่ยวข้องกับการปรับปรุงกลไกการกลั่นกรองของแพลตฟอร์ม แม้ว่าการกลั่นกรองยังคงมีข้อจำกัด โดยเฉพาะในกรณีที่ข้อความบิดเบือนมีลักษณะเชิงบริบท เช่น การนำข้อเท็จจริงบางส่วนมาใช้อย่างผิดบริบท ซึ่งอาจหลีกเลี่ยงการถูกระบุว่าเป็นข้อมูลบิดเบือนได้

ในทางกลับกัน การเพิ่มขึ้นของ Malinformation ใน Facebook อาจสะท้อนถึงพฤติกรรมของผู้ใช้งานที่ใช้แพลตฟอร์มนี้เป็นพื้นที่เผยแพร่ข้อมูลบิดเบือนในลักษณะที่ซับซ้อนมากขึ้น เช่น การบิดเบือนเชิงบริบทหรือการแฝงข้อความที่ทำให้เกิดความเข้าใจผิดในประเด็นทางการเมือง การเพิ่มขึ้นนี้อาจบ่งชี้ถึงความแตกต่างในวิธีการกลั่นกรองของแพลตฟอร์ม Facebook และ X ซึ่งอาจนำไปสู่ความสามารถในการจัดการกับเนื้อหาบิดเบือนที่ไม่เท่าเทียมกัน

แนวโน้มที่แตกต่างนี้อาจเกี่ยวข้องกับความท้าทายหลายประการ เช่น กลยุทธ์และมาตรการของแต่ละแพลตฟอร์มในการตรวจสอบและจัดการเนื้อหา Malinformation ใน X ซึ่งใช้วิธีการลดการมองเห็น (shadow banning) อาจช่วยลดการแพร่กระจายของข้อมูลบิดเบือนบางส่วนได้ ในขณะที่ Facebook ซึ่งยังคงอนุญาตให้เนื้อหาหลายประเภทอยู่ในระบบ อาจทำให้เนื้อหาบิดเบือนมีโอกาสร่วงหลายมากขึ้น

ตัวอย่าง Malinformation ที่ตรวจพบในการศึกษา ได้แก่

1. **ประเด็นเชิงนโยบายรัฐ** เช่น โครงการแลนด์บริดจ์ นโยบายดิจิทัลวอลเล็ต การขายข้าวในโครงการจำนำข้าวของกระทรวงพาณิชย์ และโครงการ entertainment complex
2. **การแก้ไขรัฐธรรมนูญ**
3. **ประเด็นทางเศรษฐกิจ** เช่น สัดส่วนหนี้สาธารณะต่อ GDP ซึ่งถูกเชื่อมโยงกับวาทกรรมเกี่ยวกับ “หนี้ชั่วลูกชั่วหลาน”
4. **การเมืองและการเลือกตั้ง** เช่น ระบบทักซิธ ข่าวดูเกี่ยวกับดิลลัส การคัดเลือก ส.ว. การซื้อสิทธิขายเสียง และการแทรกแซงของต่างชาติ เป็นต้น

5. ประเด็นความขัดแย้งในสังคม เช่น กรณี Save กับ ลาน กับ Save ชาวบ้าน กับ ลาน

การถกเถียงในประเด็นเหล่านี้มักมาพร้อมกับการปกป้องหรือการโจมตีข้อมูลจากทั้งสองฝั่งที่เกี่ยวข้อง ทำให้ผู้ใช้งานทั่วไปที่อาจเข้าถึงข้อมูลเพียงด้านเดียวมีแนวโน้มที่จะเชื่อข้อมูลที่ได้รับโดยไม่ตั้งคำถามต่อความครบถ้วนหรือบริบทที่ซับซ้อนของข้อมูล นี่เป็นปัจจัยสำคัญที่สามารถสร้าง **ความแบ่งแยก (polarization)** และบั่นทอน **ความไว้วางใจ (trust)** ในระดับสังคม

ข้อสังเกตนี้ชี้ให้เห็นถึงผลกระทบของ Malinformation ที่มีต่อความไว้วางใจและการอยู่ร่วมกันในสังคม การจัดการกับปัญหานี้จำเป็นต้องพัฒนากลไกการตรวจสอบที่สามารถเข้าใจบริบทของเนื้อหา รวมถึงส่งเสริมการรู้เท่าทันสื่อในระดับประชาชน เพื่อเพิ่มความสามารถในการแยกแยะและประเมินข้อมูลอย่างรอบด้าน นอกจากนี้ แพลตฟอร์มควรเสริมสร้างมาตรฐานชุมชนและกลไกการกลั่นกรองที่สามารถตอบสนองต่อความซับซ้อนของข้อมูลและพฤติกรรมการใช้งานของผู้คนในแต่ละแพลตฟอร์มได้อย่างมีประสิทธิภาพและเหมาะสมกับบริบทเฉพาะของแต่ละประเทศ

4.3.4 “บอกละ” และการ Monetization บน X

จากการวิเคราะห์เครือข่ายบัญชีผู้ใช้งานบนแพลตฟอร์ม X หลังการเลือกตั้ง พบปรากฏการณ์ที่น่าสนใจเกี่ยวกับการเพิ่มขึ้นของจำนวนบัญชีใหม่ในเครือข่าย ขณะที่จำนวนบัญชีที่เคยแอคทีฟในช่วงการเลือกตั้งกลับมีการลดลงอย่างมีนัยสำคัญ กล่าวคือ แพลตฟอร์ม X พบความเปลี่ยนแปลงที่น่าสนใจ โดยจากบัญชีเดิม 65,925 บัญชีที่เคยมีบทบาทในช่วงเลือกตั้ง หลังการเลือกตั้งโครงการได้คัดเลือกบัญชีที่เกี่ยวข้องกับการเมืองเหลือ 35,308 บัญชี อย่างไรก็ตาม ในการเก็บข้อมูลรอบใหม่ พบว่ามีบัญชีใหม่เข้ามาในเครือข่ายเพิ่มขึ้นถึง **631,072 บัญชี** เมื่อตรวจสอบเบื้องต้น พบว่ากว่า 400,000 บัญชีเป็นบัญชีที่ใช้สำหรับโฆษณา (Ads accounts) ซึ่งโพสต์ข้อความเพียงหนึ่งครั้งในรอบ 8 เดือน ส่วนที่เหลือเป็นบัญชีที่ชาวเน็ตเรียกว่า “บอกละ”

บัญชี “บอกละ” เหล่านี้ไม่ได้โพสต์เนื้อหาเกี่ยวกับการเมือง แต่บางบัญชีมีการใช้โปรแกรมแปลภาษาเพื่อโพสต์ข้อความในภาษาไทยหรือภาษาอื่น ๆ เช่น อังกฤษ ญี่ปุ่น เกาหลี อินเดีย และสันสกฤต เป็นต้น เมื่อทำความเข้าใจข้อมูลโดยการตัดบัญชีบอกละและบัญชีโฆษณาออก เหลือเพียง 17,902 ข้อความที่เกี่ยวข้องกับการเมืองจาก 1,759 บัญชีในเครือข่าย

ดังนั้นการพบเจอบัญชีบอกละจำนวนมากหลังการเลือกตั้ง สามารถเป็นไปได้ใน 3 กรณี คือ

1. การลดลงของกิจกรรมทางการเมืองในบัญชีเดิม บัญชีที่เคยโพสต์เนื้อหาการเมืองในช่วงการเลือกตั้งส่วนใหญ่หยุดโพสต์ข้อความที่เกี่ยวข้องกับการเมืองหลังการเลือกตั้ง สะท้อนถึงความสนใจที่ลดลงต่อประเด็นทางการเมืองในกลุ่มผู้ใช้งานเดิม

2. การยืนยันพฤติกรรมบัญชีจัดตั้ง (Organized Account) บัญชีทางการเมืองเดิมในช่วงการเลือกตั้งได้เปลี่ยนรูปแบบกิจกรรมของตัวเองมาเป็นบัญชีบอกละเอง สะท้อนว่าบัญชีที่เคลื่อนไหวในช่วงการเลือกตั้งถูกจัดตั้งขึ้นโดยองค์กรหนึ่ง เช่น บริษัทเอเจนซีทางการตลาดและโฆษณา เพื่อวัตถุประสงค์บางอย่าง และพอสิ้นสุดการเลือกตั้ง องค์กรนั้นก็ใช้เครือข่ายบัญชีเดิมของตัวเองเพื่อสร้างรายได้ของผู้ใช้งานจากแพลตฟอร์ม

3. การเพิ่มขึ้นของบัญชีใหม่ที่ไม่เกี่ยวข้องกับการเมือง เป็นบัญชีที่ถูกสร้างขึ้นใหม่หลังการเลือกตั้งเพื่อวัตถุประสงค์การเป็นบัญชีบอกละโดยเฉพาะ

อย่างไรก็ตาม ทั้ง 3 สมมติฐานข้างต้น โครงการไม่สามารถพิสูจน์ได้ว่า กว่า 6 แสนบัญชีที่เข้ามามีส่วนร่วมในเครือข่ายบัญชีช่วงเลือกตั้งจัดอยู่ในกลุ่มไหน ในสัดส่วนเท่าไรบ้าง เนื่องจากข้อจำกัดหลายด้านทั้งทรัพยากร ทุน เวลา และเครื่องมือการเข้าถึงข้อมูลอันเนื่องมาจากการเปลี่ยนแปลงนโยบายของแพลตฟอร์ม ทำให้ข้อมูลใหม่ที่โครงการฯ เก็บมาได้ั้น โครงการฯ ไม่ทราบถึงวันที่สร้างบัญชี รวมทั้งการเปลี่ยนชื่อ username ของบัญชี ก็ยากที่จะตรวจสอบย้อนกลับ ดังนั้นโครงการฯ จึงขอเสนอเป็นสมมติฐาน 3 scenario ที่น่าจะเป็นไปได้ดังที่กล่าวไปข้างต้น

แต่สิ่งที่โครงการฯ ยืนยันได้คือ การเพิ่มขึ้นของบัญชีใหม่ในเครือข่าย X ส่วนใหญ่เป็นบัญชีที่ไม่ได้มีส่วนเกี่ยวข้องกับการเมืองหรือเนื้อหาที่เป็นอันตราย โดยเฉพาะกลุ่มบัญชี “บอกละ” ที่คาดว่าเข้ามาเพื่อสร้าง engagement เพื่อเพิ่มการมองเห็นหรือดึงดูดความสนใจบนแพลตฟอร์มกับเครือข่ายในลักษณะเชิงพาณิชย์ (monetization) มากกว่าที่จะมีวัตถุประสงค์ทางการเมือง

ปรากฏการณ์ “บอทแฮก” บ่งชี้ถึงการเปลี่ยนแปลงในพฤติกรรมการใช้งานและกลไกของแพลตฟอร์ม X หลังการเลือกตั้ง แม้ว่าบัญชีเหล่านี้จะไม่ได้เผยแพร่ข้อความที่เป็นอันตราย แต่การปรากฏตัวจำนวนมากของพวกเขาอาจส่งผลต่อความสมดุลของเครือข่ายในแง่ของการสร้างข้อมูลหรือการเปลี่ยนทิศทางของบทสนทนาในเครือข่ายการศึกษาเพิ่มเติมเกี่ยวกับบทบาทและผลกระทบของบัญชีเหล่านี้ต่อระบบข้อมูลของ X จะช่วยสร้างความเข้าใจที่ลึกซึ้งยิ่งขึ้นเกี่ยวกับพลวัตของเครือข่ายในยุคดิจิทัล

4.4 ข้อเสนอแนะ

จากการศึกษาการจัดการเนื้อหาออนไลน์ในช่วงการเลือกตั้งและหลังการเลือกตั้งในประเทศไทย พบว่าปัญหา (Malinformation) และ Hate Speech ยังคงเป็นความท้าทายสำคัญต่อการสร้างความไว้วางใจในพื้นที่ออนไลน์ แม้ว่าแพลตฟอร์มและผู้มีส่วนเกี่ยวข้องต่าง ๆ จะมีมาตรการจัดการปัญหาเหล่านี้ แต่ผลการวิเคราะห์ชี้ให้เห็นถึงช่องว่างทั้งในเชิงกลไกและการบังคับใช้ที่ต้องได้รับการปรับปรุงเพื่อเพิ่มประสิทธิภาพและความโปร่งใส การลดผลกระทบของข้อมูลที่เป็นอันตรายและส่งเสริมความปลอดภัยในสภาพแวดล้อมดิจิทัลจึงต้องอาศัยความร่วมมือจากหลายภาคส่วน ได้แก่ องค์กรสื่อ แพลตฟอร์มโซเชียลมีเดีย และภาคประชาสังคม

อย่างไรก็ตาม ในข้อเสนอแนะนี้ไม่ได้รวมถึงบทบาทของภาครัฐ เนื่องจากโครงการฯ ยึดมั่นในจุดยืนที่เชื่อว่าการกำกับดูแลที่มีประสิทธิภาพควรเกิดจากความร่วมมือระหว่างภาคประชาชน สื่อ และแพลตฟอร์ม มากกว่าการให้อำนาจรัฐเข้ามาแทรกแซง ซึ่งอาจนำไปสู่การบ่อนทำลายหลักการเสรีภาพในการแสดงออก (freedom of expression) การมุ่งเน้นการกำกับดูแลร่วมกันของภาคส่วนที่หลากหลายจะช่วยรักษาสมดุลระหว่างเสรีภาพของผู้ใช้งานและความปลอดภัยในสภาพแวดล้อมออนไลน์ได้อย่างยั่งยืน

ข้อเสนอแนะในส่วนนี้มุ่งเน้นถึงบทบาทและความรับผิดชอบของแต่ละภาคส่วน โดยเน้นการสร้างสมดุลระหว่างเสรีภาพในการแสดงออกและการปกป้องผู้ใช้งานจากผลกระทบของ Malinformation และ Hate Speech ข้อเสนอแนะดังกล่าวไม่เพียงแต่ตอบสนองต่อสถานการณ์ปัจจุบัน แต่ยังวางรากฐานสำหรับการพัฒนาแนวทางที่ยั่งยืนในการจัดการเนื้อหาออนไลน์ในระยะยาว โดยมีรายละเอียดต่อไปนี้

1. ข้อเสนอต่อองค์กรสื่อ

การจัดการเนื้อหาออนไลน์ในบริบทของข้อมูลผิดพลาดและคำพูดแสดงความเกลียดชังเป็นประเด็นสำคัญที่ต้องการความร่วมมือระหว่างองค์กรสื่อ แพลตฟอร์ม และภาคประชาสังคม เพื่อสร้างความสมดุลระหว่างเสรีภาพในการแสดงออกและความปลอดภัยในพื้นที่ออนไลน์ องค์กรสื่อควรทำหน้าที่เป็นด่านแรกในการป้องกันข้อมูลบิดเบือน โดยเน้นการตรวจสอบข้อเท็จจริงก่อนเผยแพร่ หลีกเลี่ยงการนำเสนอข่าวลือหรือข่าวที่ไม่สามารถตรวจสอบแหล่งที่มาได้ นอกจากนี้ สื่อควรทำหน้าที่ส่งเสริมการรู้เท่าทันสื่อในสังคมไทย เพื่อสร้างความตระหนักรู้เกี่ยวกับวิธีการบริโภคข้อมูลอย่างมีวิจารณญาณ และสนับสนุนการพัฒนาวัฒนธรรมประชาธิปไตยที่มีรากฐานมั่นคง ดังนี้

- **ตรวจสอบความถูกต้องของข้อมูล** สื่อควรเป็นด่านแรกในการป้องกันข้อมูลบิดเบือน โดยหลีกเลี่ยงการนำเสนอข่าวเท็จจริง ข่าวลือ ข่าววงใน หรือข่าวที่ไม่สามารถตรวจสอบแหล่งที่มาได้ พร้อมทั้งยกระดับการตรวจสอบข้อเท็จจริง (fact-checking) ก่อนเผยแพร่
- **ส่งเสริมการรู้เท่าทันสื่อ** รับผิดชอบให้ความรู้ประชาชนเพื่อสร้างความเข้าใจในการบริโภคและตรวจสอบข้อมูลอย่างมีวิจารณญาณ รวมถึงสนับสนุนการสร้างวัฒนธรรมประชาธิปไตยในสังคมไทย
- **บทบาทในการขยายข้อมูลที่ถูกต้อง** สื่อมวลชนควรทำหน้าที่ขยายข้อมูลที่ถูกต้องและให้ความรู้เกี่ยวกับเทคโนโลยีใหม่ ๆ ในระดับชุมชน โดยเฉพาะกลุ่มเปราะบาง เช่น ผู้สูงอายุและผู้ที่ยากต่อการเข้าถึงข้อมูลเพียงพอ และทำงานร่วมกับองค์กรภาคประชาสังคมที่ทำหน้าที่ตรวจสอบข่าวลวงเพื่อการตรวจสอบข่าวลวงในประเด็นสาธารณะที่ส่งผลกระทบต่อสังคมอย่างสูง และมีการเผยแพร่ข้อมูลที่ถูกต้องออกไปอย่างเป็นระบบและสม่ำเสมอ
- **ปรับตัวตามเทคโนโลยี** สนับสนุนการปรับตัวของสื่อให้ทันต่อเทคโนโลยี พร้อมใช้สื่อเป็นเครื่องมือในการให้ความรู้ที่เข้าถึงได้ง่ายและเหมาะสมกับกลุ่มเป้าหมาย เช่น เด็กและเยาวชนหรือชุมชนเฉพาะ

2. ข้อเสนอต่อแพลตฟอร์ม

บทบาทของแพลตฟอร์มออนไลน์มีความสำคัญในการพัฒนาเครื่องมือที่ช่วยลดผลกระทบของข้อมูลบิดเบือน และคำพูดแสดงความเกลียดชัง การปรับปรุงระบบอัลกอริทึมเป็นสิ่งจำเป็น เพื่อให้สามารถกระจายเนื้อหาได้อย่างสมดุล และลดการมองเห็นเนื้อหาที่สร้างความแตกแยก แพลตฟอร์มควรให้ผู้ใช้งานมีทางเลือกในการควบคุมเนื้อหาที่ต้องการเห็น เช่น ระบบกรองเนื้อหาและกลไกการรายงาน นอกจากนี้ การทำงานร่วมกับหน่วยงานภายนอก เช่น ภาครัฐ สังคมและนักวิชาการ สามารถช่วยพัฒนากลไกการกลั่นกรองเนื้อหาให้เหมาะสมกับบริบทเฉพาะ เช่น การเปิดให้เข้าถึงข้อมูลผ่าน API สำหรับการวิจัยเพื่อพัฒนามาตรการที่มีประสิทธิภาพ แพลตฟอร์มยังควรเสริมสร้างความโปร่งใสในกระบวนการจัดการเนื้อหาออนไลน์ เพื่อสร้างความไว้วางใจจากสังคมดังนี้

- **ปรับปรุงระบบอัลกอริทึม** ลดการมองเห็นเนื้อหาที่สร้างความแบ่งแยกหรือกระตุ้นอารมณ์ที่รุนแรง พร้อมกระจายเนื้อหาที่มีความสมดุลและสอดคล้องกับบริบทเฉพาะของแต่ละพื้นที่หรือวัฒนธรรม
- **เพิ่มการควบคุมโดยผู้ใช้** โดยพัฒนาโลกที่ผู้ใช้สามารถควบคุมเนื้อหาที่พวกเขาเห็นได้สะดวกและใช้งานง่ายยิ่งขึ้น เช่น ระบบกรองเนื้อหา (content filtering) และกลไกการรายงาน (reporting mechanisms)
- **ส่งเสริมความร่วมมือ** แพลตฟอร์มควรทำงานร่วมกับหน่วยงานอื่น เช่น ภาครัฐ สังคมและนักวิชาการ เพื่อแบ่งปันข้อมูลและพัฒนากลไกการกลั่นกรองเนื้อหาที่มีประสิทธิภาพ นอกจากนี้ ควรสนับสนุนการวิจัยโดยการเปิดให้เข้าถึงข้อมูลผ่าน API แบบ Academic Tier Credentials
- **โปร่งใสและมีความรับผิดชอบ** เพิ่มความโปร่งใสในกระบวนการจัดการเนื้อหาออนไลน์และสร้างกลไกตอบสนองต่อภัยคุกคามอย่างรวดเร็ว
- **สร้างกลไกป้องกันกลุ่มเปราะบาง** ออกแบบระบบที่ประเมินความเสี่ยงได้อย่างครอบคลุม เช่น การป้องกันกลุ่มเปราะบางจากข้อมูลบิดเบือนและภัยออนไลน์ พร้อมปรับวิธีการสื่อสารให้เหมาะสมกับกลุ่มเป้าหมาย

3. ข้อเสนอต่อภาคประชาสังคม

ในส่วนของภาคประชาสังคม ควรทำหน้าที่เป็นผู้สนับสนุนและเชื่อมโยงความร่วมมือระหว่างภาคส่วนต่าง ๆ โดยมีบทบาทสำคัญในการเรียกร้องความรับผิดชอบจากแพลตฟอร์ม ติดตามการปฏิบัติตามมาตรฐานชุมชน และพัฒนานโยบายที่สร้างสมดุลระหว่างเสรีภาพและความปลอดภัย นอกจากนี้ ภาคประชาสังคมควรจัดหาทรัพยากรและพื้นที่ปลอดภัยสำหรับผู้ที่ได้รับผลกระทบจากการละเมิดออนไลน์ รวมถึงสนับสนุนการรณรงค์ให้มีการออกกฎหมายที่ป้องกันการฟ้องร้องเชิงกลั่นแกล้ง (SLAPP) เพื่อคุ้มครองสิทธิในการแสดงความคิดเห็นของประชาชน ดังนี้

- **เรียกร้องความรับผิดชอบจากแพลตฟอร์ม** ทำหน้าที่ติดตามและตรวจสอบการปฏิบัติตามมาตรฐานชุมชนของแพลตฟอร์ม รวมถึงสนับสนุนการปรับปรุงนโยบายที่สร้างสมดุลระหว่างเสรีภาพและความปลอดภัย
- **ประสานความร่วมมือระหว่างภาคส่วน** ส่งเสริมความร่วมมือระหว่างหน่วยงานรัฐ ภาคประชาสังคม นักวิชาการ และแพลตฟอร์ม เพื่อพัฒนาแนวทางแก้ไขปัญหาคือครอบคลุมและยั่งยืน
- **จัดหาทรัพยากรและพื้นที่ปลอดภัย** สนับสนุนการจัดหาพื้นที่ปลอดภัยและทรัพยากรให้กับผู้ที่ได้รับผลกระทบจากการละเมิดออนไลน์ รวมถึงสนับสนุนการฟื้นฟูและการเสริมสร้างความรู้เท่าทันสื่อ
- **สนับสนุนกฎหมายป้องกันการฟ้องปิดปาก (SLAPP)** เรียกร้องให้มีการออกกฎหมายป้องกันการฟ้องร้องเชิงกลั่นแกล้งที่มุ่งปิดกั้นการแสดงความคิดเห็น เพื่อส่งเสริมเสรีภาพในการแสดงออกในพื้นที่ออนไลน์

5. บรรณานุกรม

- DEAL. (2567). กอตรหัสปฏิบัติการบันทึกข้อมูลข่าวสารในไทย #เลือกตั้ง2566. The Digital Election Analytic Lab: DEAL.
- ไทยรัฐ. (26 กุมภาพันธ์ 2563). “วิโรจน์ ลักขณาอดิศร” แอปพลิเคชันโอโอ โยง “บิ๊กตู่” จะเร่งสร้างแตกแยก. เข้าถึงได้จาก <https://www.thairath.co.th/news/politic/1780653>
- พิรงรอง งามสุด. (2558). **ประชุมจากกับโลกออนไลน์**. กรุงเทพมหานคร: มูลนิธิเพื่อการศึกษาประชาธิปไตยและการพัฒนา (โครงการจัดพิมพ์คบไฟ).
- Adam Geitgey. (18 July 2018). Natural Language Processing is Fun! How computers understand Human Language. Retrieved from <https://medium.com/@ageitgey/natural-language-processing-is-fun-9a0bff37854e>
- Jacob Murel and Eda Kavlakoglu. (19 January 2024). What is a confusion matrix?. Retrieved from <https://www.ibm.com/topics/confusion-matrix>
- Northwestern Buffett Institute for Global Affairs. (2023). **The Rise of Artificial Intelligence and Deepfakes (BUFFETT BRIEF • JULY 2023)**. Retrieved from https://buffett.northwestern.edu/documents/buffett-brief_the-rise-of-ai-and-deepfake-technology.pdf
- Omar Abid. (7 November 2024). What Are Large Vision Models and How Do They Work?. Retrieved from <https://www.phdata.io/blog/what-are-large-vision-models-and-how-do-they-work/>
- Seungjun (Josh) Kim. (8 November 2022). **Let us Extract some Topics from Text Data — Part II: Gibbs Sampling Dirichlet Multinomial Mixture (GSDMM) Learn how to use GSDMM for Topic Modeling and how it compares to LDA**. Geek Culture. Retrieved from <https://medium.com/geekculture/let-us-extract-some-topics-from-text-data-part-ii-gibbs-sampling-dirichlet-multinomial-mixture-9e82d51b0fab>
- Social Science Research Council. (n.d.). **Classifying and identifying the intensity of hate speech**. Retrieved from <https://items.ssrc.org/disinformation-democracy-and-conflict-prevention/classifying-and-identifying-the-intensity-of-hate-speech/>
- State of Illinois. (n.d.). **United States Deepfake bill: 10 ILCS 5/29-21**. Retrieved from <https://www.ilga.gov/legislation/fulltext.asp?DocName=&SessionId=108&GA=101&DocType=LD=HB&DocNum=5321 &GAID=15&LegID=125899&SpecSess=&Session=>
- United Nations Human Rights Office of the High Commissioner. (2024). **Hate speech and incitement to hatred in the electoral context**. Retrieved from <https://www.ohchr.org/sites/default/files/2024-05/information-note-hate-speech-incitement-hatred-in-electoral-context.pdf>
- Wardle, C. & Dias, D. (2017). **Information disorder: Toward an interdisciplinary framework for research and policy making**. Council of Europe. Retrieved from <https://coilink.org/20.500.12592/321s46>
- Weber, A. (2009). **Manual on Hate Speech**. France: Council of Europe.

6. ภาคผนวก

ก. ตารางเปรียบเทียบแนวปฏิบัติของแพลตฟอร์มในการกลั่นกรองเนื้อหาที่เป็นอันตรายและข้อมูลบิดเบือน

Measures or Activities	Description	Action (from Platforms)
Facebook & Instagram		
Fact-Check label	Keep an eye out for warning labels to help you identify content that's been debunked by fact-checkers, so you can decide what to read, trust and share.	Labeling, censoring and can be taken down by moderator if violate the community standards
Partner with third-party fact checkers	Facebook and Instagram partner with nearly 100 independent fact-checking organizations globally who review and rate content.	Fact-check partners can label the content as misinformation. Then the moderator can take it down if it violates the community standards.
Identify false news	Identify potential misinformation using signals, like feedback from people on Facebook, Instagram, and Threads, and surface the content to fact-checkers. Fact-checkers may also identify content to review on their own.	Labeling, censoring and can be taken down by moderator if violate the community standards
Facebook		
Combating Misinformation	Our policies articulate different categories of misinformation and try to provide clear guidance about how we treat that speech when we see it. For each category, our approach reflects our attempt to balance our values of expression, safety, dignity, authenticity and privacy. We remove misinformation where it is likely to directly contribute to the risk of imminent physical harm. We also remove content that is likely to directly contribute to interference with the functioning of political processes. In determining what constitutes misinformation in these categories, we partner with independent experts who possess	We REMOVE the following types of misinformation: I. Physical harm or violence We remove misinformation or unverifiable rumours that expert partners have determined are likely to directly contribute to a risk of imminent violence or physical harm to people. We define misinformation as content with a claim that is determined to be false by an authoritative third party. We define an unverifiable rumour as a claim whose source expert partners confirm is extremely hard or impossible to trace, for which authoritative sources are absent, where there is not enough specificity for the claim to be debunked, or where the claim is too incredulous or too irrational to be believed. We know that sometimes misinformation that might appear benign could, in a specific context, contribute to a risk of offline harm, including threats of violence that could contribute to a



Measures or Activities	Description	Action (from Platforms)
	knowledge and expertise to assess the truth of the content and whether it is likely to directly contribute to the risk of imminent harm. This includes, for instance, partnering with human rights organisations with a presence on the ground in a country to determine the truth of a rumour about civil conflict. For all other misinformation, we focus on reducing its prevalence or creating an environment that fosters a productive dialogue. We know that people often use misinformation in harmless ways, such as to exaggerate a point ("This team has the worst record in the history of the sport!") or in humour or satire ("My husband just won Husband of the Year.") They also may share their experience through stories that contain inaccuracies. In some cases, people share deeply held personal opinions that others consider false or share information that they believe to be true but others consider incomplete or misleading. Recognising how common such speech is, we focus on slowing the spread of hoaxes and viral misinformation, and directing users to authoritative information. As part of that effort, we partner with third-party fact-checking organisations to review and rate the accuracy of the most viral content on our platforms (see here to learn more about how our fact-checking programme works). We also provide resources to increase media and digital literacy so that people can decide what to read, trust and share themselves. We require people to disclose, using our AI disclosure tool, whenever they post organic	heightened risk of death, serious injury or other physical harm. We work with a global network of non-governmental organisations (NGOs), not-for-profit organisations, humanitarian organisations and international organisations that have expertise in these local dynamics. In countries experiencing a heightened risk of societal violence, we work proactively with local partners to understand which false claims may directly contribute to a risk of imminent physical harm. We then work to identify and remove content making those claims on our platform. For example, in consultation with local experts, we may remove out-of-context media falsely claiming to depict acts of violence, victims or perpetrators of violence, weapons or military hardware. II. Harmful health misinformation We consult with leading health organisations to identify health misinformation likely to directly contribute to imminent harm to public health and safety. The harmful health misinformation that we remove includes the following: Misinformation about vaccines. We remove misinformation primarily about vaccines when public health authorities conclude that the information is false and likely to directly contribute to imminent vaccine refusals. III. Voter or census interference In an effort to promote election and census integrity, we remove misinformation that is likely to directly contribute to a risk of interference with people's ability to participate in those processes. This includes the following: Misinformation about the dates, locations, times and methods for voting, voter registration or census participation. Misinformation about who can vote, qualifications for voting, whether a vote will be counted and what information or materials must be provided in order to vote. Misinformation about whether a candidate is running or not.



Measures or Activities	Description	Action (from Platforms)
	content with photorealistic video or realistic-sounding audio that was digitally created or altered, and we may apply penalties if they fail to do so. We may also add a label to certain digitally created or altered content that creates a particularly high risk of misleading people on a matter of public importance. Finally, we prohibit content and behaviour in other areas that often overlap with the spread of misinformation. For example, our Community Standards prohibit fake accounts, fraud and coordinated inauthentic behaviour. As online and offline environments change and evolve, we will continue to evolve these policies. Accounts that repeatedly share the misinformation listed below may, in addition to having their content enforced on in accordance with this policy, receive decreased distribution, limitations on their ability to advertise, or be removed from our platforms.	Misinformation about who can participate in the census and what information or materials must be provided in order to participate. Misinformation about government involvement in the census, including, where applicable, that an individual's census information will be shared with another (non-census) government agency. False or unverified claims that the U.S. Immigration and Customs Enforcement (ICE) is at a voting location. Explicit false claims that people will be infected by COVID-19 (or another communicable disease) if they participate in the voting process. False claims about current conditions at a US voting location that would make it impossible to vote, as verified by an election authority. We have additional policies intended to cover calls for violence, the promotion of illegal participation and calls for coordinated interference in elections (see violence and incitement section).
Related Community Standards		
Coordinating harm and promoting crime	Prohibit people from facilitating, organising, promoting or admitting to certain criminal or harmful activities targeted at people, businesses, property or animals. We allow people to debate and advocate for the legality of criminal and harmful activities, as well as draw attention to harmful or criminal activity that they may witness or experience as long as they do not advocate for or coordinate harm.	Remove: Voter and/or census fraud - Offers to buy or sell votes with cash, gifts, services or other material goods, except if shared in condemning, awareness raising, news reporting, or humorous or satirical contexts. - Advocating, providing instructions for or demonstrating explicit intent to illegally participate in a voting or census process, except if shared in condemning, awareness raising, news reporting, or humorous or satirical contexts. For the following Community Standards, we require additional information and/or context to enforce: - Imagery that is likely to deceive the public as to its origin if:



Measures or Activities	Description	Action (from Platforms)
		- The entity depicted, or an authorised representative, objects to the imagery, and - The imagery has the potential to cause harm to members of the public. - Voter or census interference, including: - Calls for coordinated interference that would affect an individual's ability to participate in an official election or census. - Claims that voting or census participation may or will result in law enforcement consequences (for example, arrest, deportation or imprisonment). - Threats to go to an election site to monitor or watch voters or election officials' activities if combined with a reference to intimidation (e.g. "Let's show them who's boss!", "They want a war? We'll give them a war."). - Threats to go to a post-election activity site if combined with a reference to intimidation (e.g. "Let's show them who's boss!", "They want a war? We'll give them a war.").
Dangerous organisations and individuals	We do not allow organisations or individuals that proclaim a violent mission or are engaged in violence to have a presence on our platforms. We assess these entities based on their behaviour both online and offline – most significantly, their ties to violence. Under this policy, we designate individuals, organisations and networks of people. These designations are divided into two tiers that indicate the level of content enforcement, with Tier 1 resulting in the most extensive enforcement because we believe that these entities have the most direct ties to offline harm.	Remove Glorification, Support and Representation of various dangerous organisations and individuals. These concepts apply to the organisations themselves, their activities and their members. These concepts do not proscribe peaceful advocacy for particular political outcomes. Remove: Glorification, defined as any of the below: - Legitimising or defending the violent or hateful acts of a designated entity by claiming that those acts have a moral, political, logical or other justification that makes them acceptable or reasonable. E.g. "Hitler did nothing wrong." - Characterising or celebrating the violence or hate of a designated entity as an achievement or accomplishment; E.g. "Hizbul Mujahideen is winning the war for a free and independent Kashmir" - An aspirational statement of membership or statement that you would like to be a designated entity or the perpetrator of a violating violent event. E.g. "I wish I could join ISIS and be part of the Khilafah" Support, defined as any of the below: Material Support - Any act which improves the financial status of a designated entity – including funnelling money towards or away from a designated entity;



Measures or Activities	Description	Action (from Platforms)
		E.g. "Donate to the KKK!" - Any act which provides material aid to a designated entity or event; E.g. "If you want to send care packages to the Sinaloa Cartel, use this address:" - Recruiting on behalf of a designated entity or event; E.g. "If you want to fight for the Caliphate, DM me" - Other support Channelling information or resources, including official communications, on behalf of a designated entity or event E.g. Directly quoting a designated entity without caption that condemns, neutrally discusses or is a part of news reporting. Putting out a call to action on behalf of a designated entity or event; E.g. "Contact the Atomwaffen Division – (XXX) XXX-XXXX" Representation, defined as any of the below: -Stating that you are a member of a designated entity, or are a designated entity; E.g. "I am a grand dragon of the KKK." - Creating a Page, profile, event, group or other Facebook entity that is or purports to be owned by a designated entity or run on their behalf, or is or purports to be a designated event. E.g. A Page named "American Nazi Party".
Fraud, scams and deceptive practices	We aim to protect users and businesses from being deceived out of their money, property or personal information. We achieve this by removing content and combatting behaviour that purposefully employs deceptive means – such as wilful misrepresentation, stolen information and exaggerated claims – to either scam or defraud users and businesses, or to drive engagement. This includes content that seeks to coordinate or promote those activities using our platform. We allow people to raise awareness and educate others as well as condemn these activities.	Do Not Allow: Loan fraud and scams Content that: - Offers loans requiring the user to pay an advance fee to obtain a loan. - Offers loans with guarantee or near-guarantee of approval, either explicitly stated or implicitly understood based on context (such as claims to approve loan without asking for financial information). Note: We also look for other signals to determine if an entity is posting legitimate, non-fraudulent content, such as when it is a verified entity and a bank or financial institution. Gambling fraud and scams



Measures or Activities	Description	Action (from Platforms)
		Gambling fraud and scams Content that: - Offers real money gambling services ("Real money" is real-world currency that can be used to buy goods or services in the real world, including national currencies such as US Dollars and virtual currencies such as Bitcoin): - with a guarantee of winning. - implying or admitting to have rigged the outcome of a game or match. - soliciting people to enable match fixing or looking for help or tips on how to fix a match or game. Investment or financial fraud and scams Investment opportunities. Content that: Offers investment opportunities where returns on investment are guaranteed or risk-free. Offers investment opportunities where returns on investment or compensation is partly or fully based on recruitment of others to participate in the scheme. Offers investment opportunities where the opportunity is of a "get-rich-quick" nature and/or claims that a small investment can be turned into a large amount. Money/cash flip. Content that: Offers to turn a certain sum of money into a larger one through flipping or trick or strategy involving explicit mentions of "cash flip", "money flip" or similar terminology. Money muling and laundering fraud and scams Money muling. Content that: Offers or asks for money muling (causing victims to be unknowing participants in money laundering by offering money or share of profits in exchange for allowing others to use their bank accounts or transferring money on behalf of others). Offers or asks for money muling by offering employment to accept and transfer money to third parties using the victim's bank account.



Measures or Activities	Description	Action (from Platforms)
		Money laundering. Content that: Requests, solicits or offers to facilitate money laundering, which is an attempt to make illegally obtained money appear legitimate by disguising the origin of the money through a complex sequence of financial transactions, including through any of the following means: Seeking transfer of funds through SWIFT (Society for Worldwide Interbank Financial Telecommunications) or similar methods, Seeking or offering details on types of bank accounts available to support receipt or transfer of cash. Inauthentic identity fraud and scams Content that: Attempts to scam or defraud users by misrepresenting the identity of the poster or nature of a request: Charity fraud and scam, which are fraudulent requests for money or donations for charitable causes together with claims that the donation is urgent and includes information, such as bank accounts, where money can be sent. Romance fraud and scam, which are fraudulent attempts to establish online romantic relationships by seeking non-sexual companionship or relationship and offering or asking for money or its equivalent in exchange. Established business/entity fraud and scams, which involve falsely claiming to represent, or speak in the voice of, an established business or entity, in an attempt to scam or defraud. Product or reward fraud and scams - Government grant fraud and scam - Tangible, spiritual or illuminati fraud and scam - Insurance fraud and scams - Job fraud and scams - Debt relief and credit repair fraud and scam - Giveaway fraud and scam



Measures or Activities	Description	Action (from Platforms)
		- Advance fee fraud and scam Fake documents fraud and scams - Fake or forged documents - Fake or counterfeit currency - Fake or counterfeit vouchers - Fake or forged educational and professional certificates Stolen information, goods or services fraud and scam - Carding fraud and scam - PII fraud and scam - Fake review fraud and scam - Subscription fraud and scam - Cheating fraud and scam. Content that: Involves sharing, selling, trading or buying of: future exam papers or answer sheets. Products or services that enable cheating in exams. Products or services that enable passing drug tests in an unauthorised manner Unauthorised use of devices fraud and scam - Device manipulation fraud and scam. Content that Calls for buying, selling, trading or sharing of any manipulated, altered or fake measurement devices. Admits to, promotes or solicits use of physical manipulation of devices to achieve inaccurate pricing. - Digital content fraud and scam. Content that Offers or asks for products that facilitate or encourage access to digital content in an unauthorised manner. These include, but are not limited to: augmented set top boxes, fully loaded/KODI installed boxes and KODI services. Deceptive and misleading practices



Measures or Activities	Description	Action (from Platforms)
		- Misleading health practices. Content that Promotes false or misleading health claims or guarantees in a weight loss context by employing clickbait tactics, such as the use of sensational language that make exaggerated or extreme claims
Violence and incitement	While we understand that people commonly express disdain or disagreement by threatening or calling for violence in non-serious and casual ways, we remove language that incites or facilitates violence and credible threats to public or personal safety. This includes violent speech targeting a person or group of people on the basis of their protected characteristic(s) or immigration status. We remove content, disable accounts and work with law enforcement when we believe that there is a genuine risk of physical harm or direct threats to public safety. We also try to consider the language and context in order to distinguish casual or awareness-raising statements from content that constitutes a credible threat to public or personal safety. In determining whether a threat is credible, we may also consider additional information such as a person's public visibility and the risks to their physical safety.	We remove: We remove threats of violence against various targets. Threats of violence are statements or visuals representing an intention, aspiration, or call for violence against a target, and threats can be expressed in various types of statements, such as statements of intent, calls for action, advocacy, expressions of hope, aspirational statements and conditional statements. We do not prohibit threats when shared in awareness-raising or condemning context, when less severe threats are made in the context of contact sports, or certain threats against violent actors, such as terrorist groups. Universal protections for everyone Everyone is protected from the following threats: Threats of violence that could lead to death (or other forms of high-severity violence) Threats of violence that could lead to serious injury (mid-severity violence). We remove such threats against public figures and groups not based on protected characteristics when credible, and we remove them against any other targets (including groups based on protected characteristics) regardless of credibility Admissions to high-severity or mid-severity violence (in written or verbal form, or visually depicted by the perpetrator or an associate), except when shared in a context of redemption, self-defence, contact sports (mid-severity or less), or when committed by law enforcement, military or state security personnel Threats or depictions of kidnappings or abductions, unless it is clear that the content is being shared by a victim or their family as a plea for help, or shared for informational, condemnation or awareness-raising purposes Additional protections for private adults, all children, high-risk persons and persons or groups based on their protected characteristics:



Measures or Activities	Description	Action (from Platforms)
		In addition to the universal protections for everyone, all private adults (when self-reported), children and persons or groups of people targeted on the basis of their protected characteristic(s), are protected from threats of low-severity violence. Other violence In addition to all of the protections listed above, we remove the following: Content that asks for, offers or admits to offering services of high-severity violence (for example, hitmen, mercenaries, assassins, female genital mutilation) or advocates for the use of these services Instructions on how to make or use weapons where there is language explicitly stating the goal to seriously injure or kill people, or imagery that shows or simulates the end result, unless with context that the content is for a non-violent purpose such as educational self-defence (for example, combat training, martial arts) and military training Instructions on how to make or use explosives, unless with context that the content is for a non-violent purpose such as recreational uses (for example, fireworks and commercial video games, fishing) Threats to take up weapons or to bring weapons to a location or forcibly enter a location (including but not limited to places of worship, educational facilities, polling places or locations used to count votes or administer an election), or locations where there are temporary signals of a heightened risk of violence. Threats of violence related to voting, voter registration, or the administration or outcome of an election, even if there is no target. Glorification of gender-based violence that is either intimate partner violence or honour-based violence
Bullying and harassment	We distinguish between public figures and private individuals because we want to allow discussion, which often includes critical commentary of people who are featured in the news or who have a large public audience. For public figures, we remove attacks that are severe, as well as certain attacks where the public figure is directly tagged in the post	Measures are prioritized into 4 tiers: Tier 1: Universal protections for everyone: Everyone is protected from: Unwanted contact that is: Repeated, OR



Measures or Activities	Description	Action (from Platforms)
	or comment. We define public figures as state- and national-level government officials, political candidates for those offices, people with over one million fans or followers on social media and people who receive substantial news coverage. For private individuals, our protection goes further: We remove content that's meant to degrade or shame, including, for example, claims about someone's sexual activity. We recognise that bullying and harassment can have more of an emotional impact on minors, which is why our policies provide heightened protection for anyone under the age of 18, regardless of user status.	Sexually harassing, OR Is directed at a large number of individuals with no prior solicitation. Calls for self-injury or suicide of a specific person, or group of individuals. Attacks based on their experience of sexual assault, sexual exploitation, sexual harassment or domestic abuse. Statements of intent to engage in a sexual activity or advocating to engage in a sexual activity. Severe sexualised commentary. Derogatory sexualised photoshop or drawings Attacks through derogatory terms related to sexual activity (for example: whore, slut). Claims that a violent tragedy did not occur. Claims that individuals are lying about being a victim of a violent tragedy or terrorist attack, including claims that they are: Acting or pretending to be a victim of a specific event, or Paid or employed to mislead people about their role in the event. Threats to release an individual's private phone number, residential address, email address or medical records (as defined in the Privacy Violations policy). Calls for, or statements of intent to engage in, bullying and/or harassment. Content that degrades or expresses disgust towards individuals who are depicted in the process of, or straight after, menstruating, urinating, vomiting or defecating Everyone is protected from the following, but for adult public figures, they must be purposefully exposed to: Calls for death and statements in favour of contracting or developing a medical condition. Celebration or mocking of death or medical condition. Claims about sexually transmitted infections. Derogatory terms related to female gendered cursing. Statements of inferiority about physical appearance. Tier 2: Additional protections for all minors, private adults and limited scope public figures (for example, individuals whose primary fame is limited to their activism or journalism, or those who become famous through involuntary means):



Measures or Activities	Description	Action (from Platforms)
		In addition to the universal protections for everyone, all minors (private individuals and public figures), private adults and limited scope public figures are protected from: Claims about sexual activity, except in the context of criminal allegations against adults (non-consensual sexual touching). Content sexualising another adult (sexualisation of minors is covered in the Child Sexual Exploitation, Abuse and Nudity policy). All minors (private individuals and public figures), private adults and limited scope public figures) are protected from the following, but for minor public figures, they must be purposefully exposed to: Dehumanising comparisons (in written or visual form) to or about: Animals and insects, including subhuman creatures, that are culturally perceived as inferior. Bacteria, viruses, microbes and diseases. Inanimate objects, including rubbish, filth, faeces. Content manipulated to highlight, circle or otherwise negatively draw attention to specific physical characteristics (nose, ear and so on). Content that ranks them based on physical appearance or character traits. Content that degrades individuals who are depicted being physically bullied (except in fight-sport contexts). Tier 3: Additional protections for private minors, private adults and minor involuntary public figures: In addition to all of the protections listed above, all private minors, private adults (who must self-report) and minor involuntary public figures are protected from: Targeted cursing. Claims about romantic involvement, sexual orientation or gender identity. Calls for action, statements of intent, aspirational or conditional statements, or statements advocating or supporting exclusion. Negative character or ability claims, except in the context of criminal allegations and business reviews against adults. Expressions of contempt or disgust, or content rejecting the existence of an individual, except in the context of criminal allegations against adults.



Measures or Activities	Description	Action (from Platforms)
		When self-reported, private minors, private adults and minor involuntary public figures are protected from the following: - First-person voice bullying. - Unwanted manipulated imagery. - Comparison to other public, fictional or private individuals on the basis of physical appearance. - Claims about religious identity or blasphemy - Comparisons to animals or insects that are not culturally perceived as intellectually or physically inferior ("tiger", "lion"). - Neutral or positive physical descriptions. - Non-negative character or ability claims. - Attacks through derogatory terms related to a lack of sexual activity. Tier 4: Additional protections for private minors only: Minors get the most protection under our policy. In addition to all of the protections listed above, private minors are also protected from: - Allegations about criminal or illegal behaviour. - Videos of physical bullying against minors, shared in any context. - Bullying and harassment through pages, groups, events and messages The protections of Tiers 1 to 4 are also enforced on Pages, groups, events and messages.
Hate speech	We define hate speech as a direct attack against people – rather than concepts or institutions – on the basis of what we call protected characteristics: race, ethnicity, national origin, disability, religious affiliation, caste, sexual orientation, sex, gender identity and serious disease. We define attacks as violent or dehumanising speech, harmful stereotypes, statements of inferiority, expressions of contempt, disgust or dismissal, cursing and calls for exclusion or segregation. We also prohibit the use of harmful stereotypes, which we define as dehumanising	Measures are prioritized into 3 tiers: Tier 1 Content targeting a person or group of people (including all groups except those who are considered non-protected groups described as having carried out violent crimes or sexual offences or representing less than half of a group) on the basis of their aforementioned protected characteristic(s) or immigration status with: Violent speech or support in written or visual form



Measures or Activities	Description	Action (from Platforms)
	comparisons that have historically been used to attack, intimidate or exclude specific groups, and that are often linked with offline violence. We consider age a protected characteristic when referenced along with another protected characteristic. We also protect refugees, migrants, immigrants and asylum seekers from the most severe attacks, though we do allow commentary and criticism of immigration policies. Similarly, we provide some protections for characteristics such as occupation, when they're referenced along with a protected characteristic. Sometimes, based on local nuance, we consider certain words or phrases as frequently used proxies for PC groups. We also prohibit the usage of slurs that are used to attack people on the basis of their protected characteristics. However, we recognise that people sometimes share content that includes slurs or someone else's hate speech to condemn it or raise awareness. In other cases, speech, including slurs, that might otherwise violate our standards can be used self-referentially or in an empowering way. Our policies are designed to allow room for these types of speech, but we require people to clearly indicate their intent. If the intention is unclear, we may remove content.	Dehumanising speech or imagery in the form of comparisons, generalisations or unqualified behavioural statements (in written or visual form) to or about: Insects (including but not limited to: cockroaches, locusts) Animals in general or specific types of animals that are culturally perceived as intellectually or physically inferior (including but not limited to: Black people and apes or ape-like creatures; Jewish people and rats; Muslim people and pigs; Mexican people and worms) Filth (including but not limited to: dirt, grime) Bacteria, viruses or microbes Disease (including but not limited to: cancer, sexually transmitted diseases) Faeces (including but not limited to: shit, crap) Subhumanity (including but not limited to: savages, devils, monsters, primitives) Sexual predators (including but not limited to: Muslim people having sex with goats or pigs) Violent criminals (including but not limited to: terrorists, murderers, members of hate or criminal organisations) Other criminals (including, but not limited to, "thieves", "bank robbers" or saying "All [protected characteristic or quasi-protected characteristic] are 'criminals'"). Certain objects (women as household objects, property or objects in general; Black people as farm equipment; transgender or non-binary people as "it") Statements denying existence (including but not limited to: "[protected characteristic(s) or quasi-protected characteristic] do not exist", "no such thing as [protected characteristic(s) or quasi-protected characteristic]") Harmful stereotypes historically linked to intimidation, exclusion or violence on the basis of a protected characteristic, such as Blackface; Holocaust denial; claims that Jewish people control financial, political or media institutions; and references to Dalits as menial labourers Mocking the concept, events or victims of hate crimes even if no real person is depicted in an image. Tier 2 Content targeting a person or group of people on the basis of their protected characteristic(s) with:



Measures or Activities	Description	Action (from Platforms)
		Generalisations that state inferiority (in written or visual form) in the following ways: Physical deficiencies are defined as those about: Hygiene, including, but not limited to: filthy, dirty, smelly. Physical appearance, including, but not limited to: ugly, hideous. Mental deficiencies are defined as those about: Intellectual capacity, including, but not limited to: dumb, stupid, idiots. Education, including, but not limited to: illiterate, uneducated. Mental health, including, but not limited to: mentally ill, retarded, crazy, insane. Moral deficiencies are defined as those about: Character traits culturally perceived as negative, including, but not limited to: coward, liar, arrogant, ignorant. Derogatory terms related to sexual activity, including, but not limited to: whore, slut, perverts. Other statements of inferiority, which we define as: Expressions about being less than adequate, including, but not limited to: worthless, useless. Expressions about being better/worse than another protected characteristic, including, but not limited to: "I believe that males are superior to females." Expressions about deviating from the norm, including, but not limited to: freaks, abnormal. Expressions of contempt (in written or visual form), which we define as: Self-admission to intolerance on the basis of protected characteristics, including, but not limited to: homophobic, islamophobic, racist. Expressions that a protected characteristic shouldn't exist. Expressions of hate, including, but not limited to: despise, hate. Expressions of dismissal, including, but not limited to: don't respect, don't like, don't care for Expressions of disgust (in written or visual form), which we define as: Expressions suggesting that the target causes sickness, including, but not limited to: vomit, throw up. Expressions of repulsion or distaste, including, but not limited to: vile, disgusting, yuck. Cursing, except certain gender-based cursing in a romantic break-up context, defined as: Referring to the target as genitalia or anus, including, but not limited to: cunt, dick, asshole. Profane terms or phrases with the intent to insult, including, but not limited to: fuck, bitch, motherfucker.



Measures or Activities	Description	Action (from Platforms)
		Terms or phrases calling for engagement in sexual activity, or contact with genitalia, anus, faeces or urine, including, but not limited to: suck my dick, kiss my ass, eat shit. Tier 3 Content targeting a person or group of people on the basis of their protected characteristic(s) with any of the following: Segregation in the form of calls for action, statements of intent, aspirational or conditional statements, or statements advocating or supporting segregation. Exclusion in the form of calls for action, statements of intent, aspirational or conditional statements, or statements advocating or supporting, defined as Explicit exclusion, which means things such as expelling certain groups or saying that they are not allowed. Political exclusion, which means denying the right to political participation. Economic exclusion, which means denying access to economic entitlements and limiting participation in the labour market. Social exclusion, which means things such as denying access to spaces (physical and online) and social services, except for gender-based exclusion in health and positive support groups. Content that describes or negatively targets people with slurs, where slurs are defined as words that inherently create an atmosphere of exclusion and intimidation against people on the basis of a protected characteristic, often because these words are tied to historical discrimination, oppression and violence. They do this even when targeting someone who is not a member of the PC group that the slur inherently targets.
Privacy violations	We remove content that shares, offers, or solicits personally identifiable information or other private information that could lead to physical or financial harm, including financial, residential and medical information, as well as private information obtained from illegal sources. We recognise that private information may become publicly available through news coverage, court filings, press releases or other	We do not allow: Content that shares or asks for private information, either on Facebook or through external links, as follows: Personally identifiable information (PII) Content sharing personally identifiable information (information that uniquely identifies an individual) of the poster or others. This includes:



Measures or Activities	Description	Action (from Platforms)
	sources. When that happens, we may allow the information to be posted.	National identification numbers such as social security numbers (SSN), passport numbers or individual taxpayer identification numbers (ITIN). Government IDs of law enforcement, military or security personnel. Records or official documentation of civil registry information such as marriage, birth, death, name change or gender recognition documents Immigration and work status documents such as green cards, work permits or visas Content asking for personally identifiable information of others Personal contact information Content sharing personal contact information of others, except when made public by the individual or when shared or solicited to promote charitable causes, facilitate finding missing people, animals or owners of missing objects, or contact a business or service provider (unless it is established that the personal contact information is shared without the consent of the individual). Financial information. Content sharing or asking for personal financial information about the poster or others, defined as information about a person's individual finances including: Non-public financial records or statements. Bank account numbers with security or PIN codes. Digital payment method information with login details, security or PIN codes. Credit or debit card information with validity dates or security PINs or codes. Content sharing or asking for non-public financial information of a business or organisation, defined as information about a business' or organisation's finances, except when originally shared by the organisation itself (including subsequent shares with the original context intact) or shared through public reporting requirements (for example as required by stock exchanges or regulatory agencies), including: Non-public financial records or statements Bank account numbers accompanied by security or PIN codes. Digital payment method information accompanied by login details, security or PIN codes. Residential information



Measures or Activities	Description	Action (from Platforms)
		Content sharing full private residential addresses of others, including building name or pins on a map identifying the address (even if the pins are in an off-platform link), except in the following contexts: - when shared to promote charitable causes, facilitate finding missing people, animals or owners of missing objects, or contact a business or service providers - when the residence is an official residence or embassy provided to a high-ranking public official Content sharing partial private residential addresses of others (except when the residence is an official residence or embassy provided to a high-ranking public official): When shared in the context of organising protests or surveillance of the resident and the location of the residence is identified by any one of the following: - Street - Town/city or neighbourhood (only for towns/cities with fewer than 50,000 residents) - Postal code GPS pins or pins on a map identifying any of these (even if the pins are in an off-platform link) Imagery that displays the external view of private residences if all of the following conditions apply: - The residence is a single-family home, or the resident's unit number/building name is identified in the image/caption. - The location of the residence is identified by any one of the following: - Street - Town/city or neighbourhood (only for towns/cities with fewer than 50,000 residents) - Postal code GPS pins or pins on a map identifying any of these (even if the pins are in an off-platform link) The content identifies the resident(s). Either that resident objects to the exposure of their private residence, or there is context of organising protests against the resident. The imagery of the residence is not being shared because the residence is the focus of a news story (except when shared in the context of organising protests against the resident) Content asking for private residential information of others (except when the residence is an official residence or embassy provided to a high-ranking public official)



Measures or Activities	Description	Action (from Platforms)
		Content asking for location of safe houses or exposes information about safe houses by sharing any of the below (except when the safe house is actively promoting information about their facility) Actual address (note: "Post Box only" is allowed) Images of the safe house Identifiable city/neighbourhood of the safe house Information exposing the identity of the safe house residents
X		
Overall measures	You may not use X's services for the purpose of manipulating or interfering in elections or other civic processes, such as posting or sharing content that may suppress participation, mislead people about when, where, or how to participate in a civic process, or lead to offline violence during an election.	Labeling: Any attempt to undermine the integrity of civic participation undermines our core tenets of freedom of expression and as a result, we will apply labels to violative posts informing users that the content is misleading.
Misleading information about how to participate	You may not advance verifiably false or misleading information about how to participate in an election or other civic process. This includes but is not limited to: - misleading information about procedures to participate in a civic process (for example, that you can vote by Post, text message, email, or phone call in jurisdictions where these are not a possibility); - misleading information about requirements for participation, including identification or citizenship requirements; - misleading claims that cause confusion about the established laws, regulations, procedures, and methods of a civic process, or about the actions of officials or entities executing those civic processes; and - misleading statements or information about the official, announced date or time of a civic process.	Posts that are enforced under this policy will have their reach restricted on X by: - Excluding the post from search results, trends, and recommended notifications - Removing the post from the For you and Following timelines - Restricting the post's discoverability to the author's profile - Restricting Likes, replies, reposts, quotes, bookmarks, share, pin to profile, or Edit post - Downranking the Post in replies Posts enforced under this policy will receive a label informing both Post authors and viewers that we have limited the Post's visibility. Post authors are able to submit an appeal on the label if they think we incorrectly limited their Post's visibility.



Measures or Activities	Description	Action (from Platforms)
Suppression	You may not advance verifiably false or misleading information about the circumstances surrounding a civic process intended to intimidate or dissuade people from participating in an election or other civic process. This includes but is not limited to: - misleading claims that polling places are closed, that polling has ended, or other misleading information relating to votes not being counted; - misleading claims about police or law enforcement activity related to voting in an election, polling places, or collecting census information; - misleading claims about long lines, equipment problems, or other disruptions at voting locations during election periods;	
Intimidation	You may not engage in or promote behaviors that may coerce others to refrain from participating in a civic process. This includes, but is not limited to: - inciting or promoting violent behaviors intentionally near a location where an electoral process is being conducted, including polling stations and vote counting locations; - inciting the disruption or destruction of procedures, infrastructure, or election equipment that is necessary for someone to participate in a civic process; - inciting others to harass voters or poll workers;	



Measures or Activities	Description	Action (from Platforms)
	- promoting the brandishing of firearms near polling locations to intimidate voters and election workers; - threats regarding voting locations or other key places or events	
False or misleading affiliation	You may not create fake accounts which misrepresent their affiliation, or share content that falsely represents its affiliation, to a candidate, elected official, political party, electoral authority, or government entity	
Exception	Not all false or untrue information about politics or civic processes constitutes manipulation or interference. In the absence of other policy violations, the following are generally not in violation of this policy: - inaccurate statements about an elected or appointed official, candidate, or political party; - organic content that is polarizing, biased, hyperpartisan, or contains controversial viewpoints expressed about elections or politics; - discussion of public polling information; voting and audience participation for competitions, game shows, or other entertainment purposes; - using X pseudonymously or as a parody, commentary, or fan account to discuss elections or politics.	No action will be taken
Youtube		



Measures or Activities	Description	Action (from Platforms)
Voter suppression	Content aiming to mislead voters about the time, place, means, or eligibility requirements for voting, or false claims that could materially discourage voting. <u>Example:</u> - Telling viewers they can vote through inaccurate methods like texting their vote to a particular number. - Giving made up voter eligibility requirements like saying that a particular election is only open to voters over 50 years old. - Telling viewers an incorrect voting date. - Claiming that a voter's political party affiliation is visible on a vote-by-mail envelope. - False claims that non-citizen voting has determined the outcome of past elections. - False claims that Brazilian electronic voting machines have been hacked in the past to change an individual's vote.	If your content violates this policy, we will remove the content and send you an email to let you know. If we can't verify that a link you post is safe, we may remove the link. Note that violative URLs posted within the video itself or in the video's metadata may result in the video being removed. If this is your first time violating our Community Guidelines, you'll likely get a warning with no penalty to your channel. You will have the chance to take a policy training to allow the warning to expire after 90 days. The 90 day period starts from when the training is completed, not when the warning is issued. However, if the same policy is violated within that 90 day window, the warning will not expire and your channel will be given a strike. If you violate a different policy after completing the training, you will get another warning. If you get 3 strikes within 90 days, your channel will be terminated. We may terminate your channel or account for repeated violations of the Community Guidelines or Terms of Service. We may also terminate your channel or account after a single case of severe abuse, or when the channel is dedicated to a policy violation. We may prevent repeat offenders from taking policy trainings in the future.
Candidate eligibility	Content that advances false claims related to the technical eligibility requirements for current political candidates and sitting elected government officials to serve in office. Eligibility requirements considered are based on applicable national law, and include age, citizenship, or vital status. <u>Example:</u> - Claims that a candidate or sitting government official is not eligible to hold office based on false info about the age required to hold office in that country/region.	



Measures or Activities	Description	Action (from Platforms)
	- Claims that a candidate or sitting government official is not eligible to hold office based on false info about citizenship status requirements to hold office in that country/region. - Claims that a candidate or sitting government official is ineligible for office based on false claims that they're deceased, not old enough or otherwise do not meet eligibility requirements.	
Incitement to interfere with democratic processes	Content encouraging others to interfere with democratic processes. This includes obstructing or interrupting voting procedures. <u>Example:</u> - Telling viewers to create long voting lines with the purpose of making it harder for others to vote. - Telling viewers to hack government websites to delay the release of elections results. - Telling viewers to incite physical conflict with election officials, voters, candidates, or other individuals at polling locations to deter voting.	
Election integrity	Content advancing false claims that widespread fraud, errors, or glitches occurred in certain past elections to determine heads of government. Or, content that claims that the certified results of those elections were false. <u>Example:</u> - Content advancing false claims that widespread fraud, error, or glitches changed the outcome of the German parliamentary (Bundestag) elections, delegitimizes the formation of the new government or	



Measures or Activities	Description	Action (from Platforms)
	the election and appointment of the next German Chancellor. - False claims that widespread fraud, error, or glitches changed the outcome of the 2018 Brazilian presidential election.	



ข. คำคัดออก (Excluded Keywords)

ครั้งที่ 1 จำนวน 469 คำ ได้แก่

จุงอาเชน, JoongArchen, เด็กหงส์, เด็วด, อดุณสวัสดิ์, เกี้ยวก๊วง, มิวศุภคิษฏ, MewSuppasit, โก๋ย่าง, หนูเล็ก, กอดหัวใจ, YouTube, NEW YEAR, HNY, สวัสดิ์ปีใหม่, TENLEE, เดนลี่, จุงดิง, dunknatachai, guysivakorn, ใพนจำของกาย, ปอนดักวิรินทร์, pparavit, phuwintang, pondphuwin, #WeAreSeries, จินโจะ, HBD, เลี้ยงปลา, ทางนุกยุง, octotul, Tu_Pakorn, ดุลย์ภากร, mewtul, The33rdSMA, ZeeNuNew, SMAinBKK, นูนิว, NuNew, NanaNu, CwrNew, สุขสันต์วัน, สวัสดิ์วัน, แมนสรวง, ManSuang, NetflixTH, MilePhakphum, Nnattawin, โบกอม, นื่องกอม, TayNew, Tawan_V, Newwiew, เตมิว, TayTawan, วิชญ์ภาส, BibleWichapas, bsumone, ขอWS, LastTwilightSeries, jimmyujp, sea_tawinan, JimmySea, จิมมีซี, Zhangzhehan, ZhangSanJian, จางจื่อฮั่น, OceanZhang511, psyfe_member, PE7PE7POWPOW, SPICE_PCF, PSYCHICFEVER, FHERO, BearKnuckle, HighCloudEnt, ไม้เลื้อยดอก, มอนิ่ง, birthday, ChaoTalay, ชาวทะเล, วายซีฟ, ที่ดิน+ขาย, ที่ดิน+ซื้อ, กินกันกับเตมิว, บริฟเหี้ย, บริฟเหี้ย, PerthTanapon, Perthppe, GMMTV, KDPPE, ตลาดนัดgmm, JISUNG, NCTDREAM, TOP100KPOPMKNAES, หาวโรอ้อย, Sumettikul, Chimon, PerthChimon, เพ็ริชชีมอน, chimonac, ChimonWachirawit, ซีม่อนโจงะไครหลั, DOYOUNGTEN, DOYOUNG, DOYOUNG+TEN, เปรียวกาก, โมนเมน, เลิกงาน, sunseangnim, mixxiw, wixxiws, เอิร์ทมิทซ์, EarthMix, HappyAllyDay, ปฏิกรณ์, reactor, อย่ำมาเล่นกับไฟ, calumscott, BOGUMMY, ParkBoGum, พักโบกอม, ชวนชม, *Kaowoat*, PlayboyTheSeries, จิตอง, สละโสด, อร์อัย, TOTY, ZeePruk, zee_pruk, มั่งเฮียก็หาว่าซื้อ, BamBam1A, SHINYU, โดคยอม, ซินยู, SEVENTEEN, ทักทาย, MilanFashionWeek, จุงเอ็ง, AFGvIRE, MichaelKors, APPLE MUSIC, NightWalker, พร้อเอเดอร์กาหลั, GMMTVHappyWeekend, Cutie Pie, เป็กพลิตโชค, PeckPaLitChoke, peckpalit, คุณได้ไปต่อ, mhokday, BloodfromLove, กอดบัตร, กอด+บัตร, บัตร+เงิน, Mewlions, มิวเลี่ยนส์, ndcodeine, ssummersnoww, PinGPinZ, สุมุนวิท11, Youngjae, ยองแจ, Linkin Park, Lying from You, Shadowban, ดิดอง, phuwintom, 5 Tag, 5 Food, แจมิน, JAEMIN, SMTOWNGLOBAL, ROSÉ, BLACKPINK, TOP200FACESKPOP, ปิ่นภักดี, EatWithBirb, ApoNattawin, ขำพ้าเคียงเธอ, นุ่นแพรว, WangYibo, หวังอี้บ้อ, นุ่นเปรม, BounPrem, prem_space, รักเธอทั้งหมดของหัวใจ, Janhae, itdthairadex2024, Earthpirapat, Earth_Pirapat, jinyoung, GOT7, Gmmtlivehouse, กาลเวลาพิสูจน์รักแท้, แจจ, TheLoyalPinGalaPremiere, ซาอี, มินารีนา, LISA, ROCKSTAR, BounConcert, BUSbecauseofyouishine, Pokemon, Anime, sooneclipse, งานบอล, งานบอล+ช่อง3, 54ปี3Miracles, PEEMWASU, NEXnattakit, PHUTATCHAI, COPPERdechawat, THAichayanon, ตลาดนัดBUS, BOC, Fuaiz_thanawat, Zunshine, Congratulations, Congrast, อ๊วง, kaewoic, NANON, B612, mynameisnnon, nanon_korapat, ดาวเด็ม, giveaway, กอว, บุญมาพาก, SaintLaurent, เจ็ดสี่เอมืองสิงห์, TheLegendarySeven, netsiraphop, Nanning, NetJames, เน็ตเจมส์, JamesSu, ohmtpk, Tpkkrubs, 22ndOHMTPKDAY, Boeing 737, david_sommers, ฝนดี, ยืมได้ตัวเอง, #พิชซ่าชานม, พร็อกาส, SWEATSHIRT, free shipping, ส่งฟรี, สอบถาม DM, ซิงหุขุขุ, ssingss, SIngHarit, Minmin_SoraSora, Bb0un_India, bb0un, คินน์พอร์ช, มายอาโบ, Mile1stSoloConcert, AuthenticMileconcert, รักแรกของเพิร์ล, First_Chalongrat, m1ke, MookWorranit, RIIZE, concert, ticket, มิวดลย์, TheSignกลางสิงห์, idolfactoryTH, TheSign7Hunters, babiibabe, GroomWithDCxGP, DCxGiantPiccolo, DrChoicePet, Pramote, TorSaksit, ใจซ่อนรัก, TheSecretOfUs, Fuaiz, fitoxyคู่ซี้เดย์พิชชีนดี้, Saint_sup, MingEr, PIRAMYST, fitoxy, ฅมอนอิง, yhglaags, ใจพิชชี, เซ็นต์คูกพงษ์, ดวงใจทองพรม, Jaipisut, BIGHIT_MUSIC, bts_bighit, UMG, GeffenRecords, Best_SoraSora, MinminSoraSora, BestSoraSora, SoraSora, Catsolute, RalphLauren, HommesThailand, cactus, KFC, สนใจ ib, สอบถามเพิ่มเติม, gentlewoman, gentlewomen, Saintsup, เซ็นต์ซู, thconcert, fan+benefit, ส่งต่อ, Hutiaoxia Gorge, สวัสดิ์ยาม, YOASOBI, gigi goode, สวอนน์, Jamessu_w, GMMTV2024, cherriri_s, sansun_zzz, KIM NAMJOON, #BTS, เรอฟอร์แมช, LYKN, bebec, ฟรีน, สโรชา, ตู้หิ้ว, Top Albums, Spotify, Top Artists, GMMTV2024PART2, GMMTV2024PART1, ต้าหู้จื่ออู่ฟโรด, ต้าหู้จื่อ, biblesumett, SPACELESS, ไม้ต้องมีที่ให้อันอยู่แต่ขอ, แท้มีอันอยู่ก็พอ, 2024BTSFESTA, BTS11thAnniversary, ฅนกรอบ, DreamGL, sanrio, ซานริโอ, กล้องลุ่ม, arttoy, คุโรมิ, kuromi, ฟีมาย, ส่งต่อ, เสื้อผ้า, เสื้อผ้ามือสอง, เสื้อผ้ากาหลั, โล้ดูเสื้อผ้า, ส่งต่อcintage, MEWBR, jin, bts_twt, หลิงหลิงคอง, นันจุน, NuNew3rdSingle, Alyana, DAOU, oueijia, SourandSweet, BamBam, พี่ดา, พี่โก, mmarkpkk, MarkPakinStandUpBDParty, หอมพอนคลาย, Maison, Emma, Dragon, FriendofKFTThailand, สสส, GreenyRose, คอนเสิร์ต, mile, innsarin, ValentinoSGxGreatInn, markohm, SuperMARKday26th, NheeTamGaliTeo, กิซอโปไชย, winmetawin, ออฟโรด, CHUANGAsia, uniqlo, LOLFANFEST2024, ฅมชดล็กอวอร์ด, KOMCHADLUEK, JIMIN, TOP200KPOPVOCALISTS, ฅารามาธิกมิเนโอะ, เมษ, พฤษก, เมณู, กรกฏ, สิงห์, กันย์, ดุลย์, พิจก, ธนู, มังกร, ฅมร์, มีน, Aries, Taurus, Gemini, Cancer, Leo, Virgo, Libra, Scorpio, Sagittarius, Capricorn, Aquarius, Pisces, ฅูแตก, BibleJes, JesBible, jesjpp, JENNIE, KomChadLuekAward, Kom Chad Luek Award, สหายท้อ, riBBon, perses_official, perses, ฅวมส่ง, fernnnfernnfernn, BighugChimon, คอน+เกา, คอน+ไทย, ฅนั๊ก, JAMFILM, NATARAJA, ฅารุราชครั้งที่15, ฅารุราช, film_tnp20, Firstkpp, FirstKanaphan, เพ็ริสขัวดิง, คินดะอิจิ, ReadWithBirb, GMMTVOuting2024OUTING, ฅาร์คโอม, off_tumcial, AtthaphanP, KobarahTrueLove, OffGun, ออฟกัน, สติทเกอร์+ไลน์, สติทเกอร์+line, Bammie, สตรีน, Streaming, Stream, ฅอกจิน, TheAstronaut, ฅอชอกจินกลับบ้าน, WaitingForTheAstronaut, ฅิปมินจู, สวสค์, ปร+ถนพ

ครั้งที่ 2 จำนวน 274 คำ ได้แก่

โอดะ, ทำอาหาร, จุนกุ, JUNKYU, TREASURE, Liverpool, #LFC, Khatauri, Bharat, เรื่องสยอง, ฮิโตะบาชิระ, ตงตง, พระคุ้มครอง, หาวโรทาน, วัฟ, สลัส, หอมกลิ่นความรัก, จูดี, โรสสัด, ทานอะโร, นาฬิกาปลุก, กองหลัง, กองหน้า, ทำบุญ, บัตรคิว, ฟอสนุค, EnchanteSeries, kasibook, fforce_, ฅรขจีน, Indie, India, Indian, EMSPIRE, จ้าวลู่ซือ, OnThisDay, LiverpoolFC, iPhone, fforcejs, อดอดอร์, ขอนญญาดี, ขอนญญาต, อลิส, โบอา, Mexicans, Spanish, Africa, น้ำตก, AIS, PradaSS24, JAEHYUN, Prada, LiuTao, หลิวทา, ขนฟู, กิต้าร์, กิตาร์, ม้องจิน, baseball, Lucifer, พี่สาว, มัดจำ, ตลาดนัดบ๊งกัน, บังกัน, FERRAGAMO, JENO, Beckham, Malaysian, Malaysia, tnbau1, besevboom, อู๋บูม, AouBoom, aou_tnbknr, HiddenAgendaSeries, ฅะมุตะมิ, anunkamol, อนันกมล, lookkaewkamollak, andaanunta, MLTR, Elton John, Beatles, chimpanzees, sunsun, sun_sun817, Sneketoshi, ZMNation, WWF, ลารีนาร์, work from, เลขเด็ด, ผัวมัน, ผัวแท้, ผัวผสม, มีลเคิลฟ,

MilkLove, 23point5, loverrick, panigyy, tawattannn, Google, cryptomarket, คัลเลน, คัลเลน+จวง, นักแสดง, WPL2024, TATAWPLFinal, TATAWPL2024, BarackObama, Harvard University, Good morning, แม่มด, WaveToEarthLiveBangkok2024, tchaikovsky, อายากิน, ChenRcj, nutcracker, จัดส่ง, ตลาดนัดdenhyphen, ตลาดนัด, alan_pasawee, ALANpasawee, introvert, extrovert, Dr.PONG, พัก, น่องจิ, JFKairport, airbus, light-years, Milky Way, Bitoey, metronome, MAESTRO, 17IsRightHere, หมอน้อง, หมอพง, JustmineNika, Mdundomusic, TwendeMbeleNaNewEra, ุบิ, forduk, FordMustang, Roadtrip, ซีรีส์, ซีรีส์, โดม่อน, YIBO, UNIQ, TOP100MostHandsomeFace, จักย, เซกิโระ, หงอกง, เอสเซนส์, สกินแคร์, ฟีบรีด, earthquake, สายเทา, Chelsea, ยาสลิม, BirthBoomDay26th, RachataJam, แอมเบด, คอนเสิร์ต, olympics2024, herooftheday, พีสัก, แฮวอน, gardening, hiking, HealthierNJ, รั้ว, EXO, ผัว+สัว, อีนัส, review+mask, แพนเค้ก, วาฬเพชฌ, เคส, ตัวต่อ, เพลง+วง, บัตร+คิว, นางเอก, เมคอัพ, Esports, Interoperable, พิน, Debut, TheTraineeSeries, ผีรัก, กันธรรพพันธ์, GunAtthaphan, forcebook, Bridgerton, Jenny, ฝันของมงกุฎ, jirayu_jj, jamesjigallery, jamesjirayu, เจนส์จิรายุ, เจนส์จิ, กวิน, กวินนิค, Winpww, winpw, กันสมาย, Inter fans, Inter fan, Interfan, Perfect10Liners, ซ้อปLazada44 กับGF, Lazada, Pre-order, พี่, ประจักษ์หลั่งลือ, PeipeiKrit, PhanP, shein, kristtps, KristPerawat, crossbody, สุลสันต์, quotedreplies, กุมาร, HappyMothersDay, Apple, ฟ่อน 0%, from this+to this, AHGASE_10YearsAndStillSoaring, AHGASE, เด็กผัว, กาทิ, โคดี้, เน็ตฟลิก, น้าค่อม, ปลาญ, บุรี, มุนิก, ปาร์ตี, โน้สอุดม, GawinCaskey, GawinKrist, sea tawinan, jimmy, wallpaper, LingOrm, RomandMultyxLingOrm, บำรุงราษฎร์, พทยา, Pattaya, SoiLengKee, #LoveBullyรักให้ร้ายตอนจบ, LoveBullyรักให้ร้าย, ไก่จิก, ทอดมัน, บอลโลก2026, บอลไทย, WorldCup2026, WorldCup, Singapore, doi, loathe, ทิคาเมซ, สุมเมเรียน, MATCHA, matcha_mosimann, whitefoxgmm, Autism, MiniBrains, Neurodiversity, Neurobiology, StemCellScience

ครั้งที่ 3 จำนวน 107 คำ ได้แก่

จิจิสุฟุซซา, จิจิ, อิกคิว, patrickananda, เรื่องของจิจิ, กัดบัต, แคมบอยอน, อาร์มี, แอ้ว, ยุกัน, อิกจิ, บิว, เฮจิ, จงอิน, โดยอง, ซิปอย, มาย+ภาคภูมิ, เจนนี่, มวยไทย, TrueID, Fourthnattawat, เจมีโน้, WayV, ช่อง8, SM+TEN, คบ+จ้อ, แฟน+เล็ก, แมมี้อตาล, เต๊ส, แอม, โซยานโด, ดวงจันทร์+ลงทักษ์, ออฟ, ซาล็อต, วอลดีน, ซิปเปอร์, คาร์ลมา, บัตรแพง, จ่าน, หลาน+จ่าน, หยาง+ด, ซาแซง, บัตร์โกลั, บัตรดอย, บัตร์หลุม, ฮัฟ, จูออน, มินา+น้ำผึ้ง, แคลเรอีน, เดปป์, แอมเบอร์, โควทาม, โจ้, แกยอน, ไสวอน, โซวอน, My Mate Nate, เจโน่, เจนเวน, จองกุก, ตึกฟ้า, มัมหมี่, เจมีนาลี, คอสตุน, คสล, #beckysangels, #srchafreen, 28บีวอยู่bocจบ, ANAN+WAR, BarcodeTin, beoncloud_th, BoycottBOC, BuildJakapan, BuildProvenInnocent, charlotte, Davikah, DomundiTV, EWaraha, FilmThanapat, FreenBecky, freeten, GulfKanawut, HYBE, ikea, jeffsatur, JusticeForBuildJakapan, Maverick, MGI, MGT, ohmnanon, QingQing, Top Gun, VogueThailand, warwanarat, yinyin_anw, กล้าม, กันดั้ม, เก็นยะ, โทโฮมิวจักรพันธ์, คัม, คัมแบค, คุโระกาคิ, คุโละพัมกิ้น, เจสัน, แจ็คสัน, ซิซัว

ครั้งที่ 4 จำนวน 138 คำ ได้แก่

ChristineMae, R4UTheSeries, Reverse4YouTheSeries, ยองพิลา, เสอเกิล, เมเจอร์, จอมแก่น, คลึงจัง, LINGLING, โชซี่, เบ๊น, อีซิง, เลย์จาง, FILA, LAYxFLA, AUGAW, ซลัยปรีน, ตัวท่อน, องครักษ์แห่งดินแดนน้ำแข็ง, PorSuppakarn, ที่พื้นฟ้า, Paris2024, จันทรใจเบ๊น, OlympicGames, กว๊น ย่า, โชทาร์โ, น้องดรีม, คู่ซี้, เจิ้งเนีย, พุดดิ้ง, เทรเซอร์, โอซาก้า, BBNaaja, ฮอนด้า, Honda, โม่กี, ฤๅเดช, ค่ายฮั่น, แม่ออนไลน์, น้องอ้ง, เอ็ม+กาด, Ybookfair8, Winny, Satang, วินนี่สตางค์, เก่งน้ำปิง, KengNamping, วอลเลย์บอลชาย, คลินิก+ชิด, AlwaysWithDawon, Hashimoto, Hayato, Yoshiwara, Last+Idol, TrailsNAles, TrailHead, SoKno, slowknorunningclub, 90DayFiance, HEESEUNG, HEECoverpt2_SweetMelodies, thechi, Eurovision, M+Countdown+Stage, Ruan, Gen1es, วิกตอเรีย, พิมพ์กรกนก, Cumulus, แอกลาส, Lightitup, ATLAS, learnprintmaking, ฮันมิน, เปิดตัวThePackageTH, MSuppasit, รูปสิงหา, รูป+สิง, Bethesda, quadcities, ฮีโร่ตรก+สัตว์เลี้ยง, Basketball, Toronto, NBA, GetWellSoonFreen, FREEN, ฟรีน, Sugarcoat, พี่ต๋อง, tongthk, ยูจี, โกะโง, รางดาว, น้องโย, แมคลาสu, XiaoZhan, HereWeGo , nfl, NASDAQ100, ซาฮีนอู, CHAEUNWOO, MysteryElevator, JUSTONE10MINUTE, NJwx, ไชเรน, ScottishIndependence, VoteSNP, พ่อเฒ่าหล่อคอนั้น30ยังแจ๋ว, TheWhisperCafe, วิลวิล, กวังอิล, LUCY, นิด+ซิง+บอล, เบนดริ์กทุกทุ่ง, อรุณสวัสดิ์, ญานญา, นายท่านกร ชิด, GaoQingchen, vivoV30, JobbyNoo, ohmpawat, โอมกวิต, OneLegacy, เหวล่าตาย, Nodted, NodtNutthasid, กระดุมขมถั่ว, Tongthk, ปลาน้อยขของต๋อง, TongAquarium, เก่ง+น้ำปิง, Kenq+Namping, KenqHarit, nampingnapat, harit_keng, nampingster, 4minute

ครั้งที่ 5 จำนวน 95 คำ ได้แก่

ชิวาวา, น้องวอจ, เน้นย่นหม้อย่น-Teaser, TEASER, ชาวเฉ็ก, ชาวด+เช็ก, THIRD_RentItFromYou, RentItFromYou, มั่งขำว, from_alienbunny, พึดนิ่ม, ลูกค้ำ+อุดหนุน, ซิป+ไอเฟน, DICEforPRIZE, สดเตนไคซ์, PRIZE, Flex1045xTHIRD, เติร์ด+MV, Binance, ชีออนบง, นิซคุณ, Nickkhun, ASICSTH, CelebrationofMovement, GQThailand, มั่งงะ, โยชี, Guangjie, กวงเจี๋ย, เป็นไอเฟน+กัน, เป็นไอเฟน+กัน, อัสฟ้า, คริสสิงโต, พิโมน, ไทเนอร, มาการอง, ฮอกทอย, KANGDANIEL, KBS+Ent, Juncao, ดุจฮัส, ไอ้ดั่งสิงโ, ENGLot, LoveBully, วันนี3กุ่ม15มีรักให้ร้าย, ClubFridayTheSeries, ชาลื้อดออสติน, CharlotteAustin, itscharlotty, DestroyLonely, ตัวไอเทน+ชาย, seakeen, OnlyBooSeries, Miniso, ซากาโมโด้, คานเย่, เซปเซป, ซง+พูน, RedRoseTogether, เฌอแถม, JakeB4rever, InfinityxJJChengFM, BUILD+LOVE, หมวดล๊อก, Metal+Alchemist, ดงยอน, Elden+Ring, ComicSquare, RaareeReview, CQ7, LineShop, babyshower, juniorcricket, เหวารินทร์คงคา, AniEx, ARProject, sgsupatt, BoycottNike, BoycottEnglandGames, มียากะ, นาคีร์, หลัอี้จ๊ว, CPop, LuoYizhou, SIAMSCAPE, พระจันทร์ของป๋อป๋อ, PorporEuphonie, EuphonieTH, ArcjewelThailand, คิรส์กี้, นิยายยริ, บ้านผะขี้พWL, ยนะชีพยริ, 1001OurRendezvous

ครั้งที่ 6 จำนวน 60 คำ ได้แก่

ทุน่มะลา, NCT_TEN_TEN, TEN_TEN, TENzenter, DylanWang, WangHedi, วังเอ๋อดี, พัสด+Review, wsmลืขิต, VOTEFORKHUNCHA!, ซามุไรฟูลูกอ่อน, Dewnitikorn, ดิวนิติกอร์, ป๋อนรักอะอ้าเต็มคำไห้บาเบ, หนึ่งปีท้อะอ้า, หลิวเต๋อ, ลิฟท้ออย, LiftOil, ลิฟท์ลู่พอน, liftaha, liftsupoj, ฝน้ำตาล, ดันเมชิ, กามิกาชะ, โทะระ, ริวซ, เกคิตสึ, แบทเกิลบอย, bbillkin, KAZZMAGAZINE, พอร์ตพิท, LoveSeaTheSeries, FortFTS, Peatwasu,



ComeFortZon, CaptainPeat, FortPeat, แมนู, วอลเลย์บอลหญิง, ออยรนา, oilthana, Reactionไม่กล้าฟัน, หลินมาลิน, Travelokaล่าเหรียญรุ่น
มีดPerth, จีซอน, พี่เก้อ, ชนชุง, ลีนจัน, ดลลี้, Phuwin, เหมื่อนวิฑาก, ฝดาษ, cashflow , diseases, Catholic, ลีน+ฟัน, จัน, วิมานขนาน, ส่วนลด