



แบบฟอร์มเสนอ

เค้าโครง ดุษฎีนิพนธ์

ปรัชญาดุษฎีบัณฑิต (ปร.ด.) เทคโนโลยีดิจิทัลมีเดีย

คณะวิทยาศาสตร์และเทคโนโลยี มหาวิทยาลัยเทคโนโลยีราชมงคลสุวรรณภูมิ

ชื่อเรื่อง (ภาษาไทย) กรอบการจัดกลุ่มภัยคุกคามด้วยปัญญาประดิษฐ์ เพื่อการลดระยะเวลาในการตอบสนองต่อเหตุการณ์และผลกระทบทางธุรกิจ ในสภาพแวดล้อมความมั่นคงปลอดภัยไซเบอร์ กรณีศึกษาในประเทศไทย

ชื่อเรื่อง (ภาษาอังกฤษ) AI-driven Threat Clustering Framework for Reducing Incident Response Time and Business Impact in Cybersecurity Environments: A Case Study in Thailand

ผู้เสนอ

อนรรักษ์ ต้นตราธิปไตย

รหัสประจำตัว 168490432006

สาขาวิชาเทคโนโลยีดิจิทัลมีเดีย

1. ความเป็นมาและความสำคัญของปัญหา

1.1 บทนำ

ในยุคดิจิทัลโครงสร้างพื้นฐานเทคโนโลยีสารสนเทศและบริการดิจิทัลได้กลายเป็นองค์ประกอบสำคัญของการดำเนินงานทั้งในระดับองค์กรและระดับประเทศ ในเวลาเดียวกันภัยคุกคามทางไซเบอร์ก็กลับมีความซับซ้อนรุนแรงขึ้นอย่างต่อเนื่อง กระจายกันโจมตีจากหลายแหล่งพร้อมกัน และโจมตีโดยการแทรกซึมมาในระบบของเราที่ละเล็กละน้อย สะสมมาโดยใช้เวลานาน ค่อยๆ แทรกแซงเข้ามา บางครั้งภัยคุกคามที่เข้ามาเป็นภัยคุกคามใหม่ที่เราไม่เคยเจอมาก่อน ยังไม่รู้ถึงวิธีการแก้ไข ทำให้กลไกป้องกันแบบดั้งเดิมที่อาศัย signature-based detection หรือ static rules ไม่สามารถรับมือได้อย่างมีประสิทธิภาพเท่าที่ควร [1] [2]

งานทบทวนล่าสุดด้าน AI-driven cybersecurity ระบุว่าเทคนิคปัญญาประดิษฐ์ โดยเฉพาะ machine learning และ deep learning สามารถช่วยเพิ่มความสามารถในการตรวจจับภัยคุกคามที่ซับซ้อน วิเคราะห์รูปแบบการโจมตีจากข้อมูลขนาดใหญ่ และลดระยะเวลาในการตอบสนองต่อเหตุการณ์ได้อย่างมีนัยสำคัญ เมื่อเปรียบเทียบกับระบบตรวจจับแบบอาศัยกฎหรือ signature แบบเดิม [1] งานสำรวจด้าน intrusion detection บน benchmark datasets เช่น CICIDS2017 และ CSE-CIC-IDS2018 ก็ชี้ให้เห็นว่าอัลกอริทึม ML/DL หลายประเภทมีศักยภาพสูงในการยกระดับ accuracy, recall และ F1-score ของระบบตรวจจับเมื่อเทียบกับแนวทางดั้งเดิม [3] แม้ AI จะช่วยยกระดับการตรวจจับได้ดีขึ้น แต่ช่องว่างสำคัญของงานประยุกต์ในปัจจุบัน ยังขาดสภาพแวดล้อมจำลองที่สมจริงและปลอดภัย สำหรับใช้ทดสอบ วิเคราะห์ และประเมินภัยคุกคามโดยไม่กระทบต่อระบบผลิตจริง งานทบทวนด้าน cyber attack simulation ระบุว่าชุดเครื่องมือจำลองจำนวนมากยังมีข้อจำกัดในด้านความสมจริงของพฤติกรรมผู้ใช้ ความหลากหลายของสถานการณ์โจมตี และการเชื่อมต่อกับกลไกตอบสนองด้านความมั่นคงปลอดภัยจริง ทำให้การทดลองจำนวนมากยังยากต่อการเปรียบเทียบและการนำไปใช้เชิงปฏิบัติ [4] [5] การเชื่อมผลการวิเคราะห์ภัยคุกคามเข้ากับ business impact และ incident response time งานด้าน incident response ล่าสุดของ NIST เน้นว่าการจัดการเหตุการณ์ทางไซเบอร์ควรถูกบูรณาการเข้ากับการบริหารความเสี่ยงทั้งระบบ โดยมีเป้าหมายไม่เพียงแต่ตรวจจับให้ได้ แต่ต้องลดจำนวนและผลกระทบของเหตุการณ์ และช่วยให้การตอบสนองและการกู้คืนระบบมีประสิทธิภาพมากขึ้น [6] NIST Cybersecurity Framework 2.0 ยังชี้ว่าการจัดการ cyber risk ที่ดีต้องเชื่อมโยงไปถึง govern , detect , respond , recover อย่างครบวงจร [7]

ด้วยเหตุนี้งานวิจัยนี้จึงเสนอแนวทาง AI-driven Threat Clustering Framework โดยยกระดับจากการตรวจจับภัยคุกคามรายเหตุการณ์ไปสู่การจัดกลุ่มภัยคุกคาม ตามลักษณะพฤติกรรม ความรุนแรง และบริบทของระบบ เพื่อสนับสนุนการจัดลำดับความสำคัญในการตอบสนองต่อเหตุการณ์และลดผลกระทบต่อกิจกรรมทางธุรกิจ แนวคิดดังกล่าวมีความสำคัญเชิงวิชาการและเชิงปฏิบัติ เพราะทำให้ AI ไม่ได้ทำหน้าที่เพียง classification หรือ anomaly detection เท่านั้น แต่กลายเป็นกลไกสนับสนุนการตัดสินใจ ที่ช่วยให้ผู้ดูแลระบบเลือกแนวทางตอบสนองที่เหมาะสมต่อภัยแต่ละกลุ่มได้รวดเร็วขึ้น [1] [5] นอกจากนี้หากเชื่อมผลการจัดกลุ่มดังกล่าวเข้ากับกรอบมาตรฐาน เช่น NIST CSF 2.0 และ ISO/IEC 27001 ก็จะช่วยเพิ่มผลลัพธ์ของระบบสามารถใช้ได้ทั้งในเชิงเทคนิค เชิงการกำกับดูแล และเชิงการเตรียมพร้อมด้าน compliance [7] [8]

การออกแบบระบบ AI สำหรับงานด้าน cybersecurity จำเป็นต้องคำนึงถึงความท้าทายใหม่ เช่น adversarial machine learning, evasion, model extraction และการโจมตีต่อโมเดล AI เอง ซึ่งอาจทำให้ระบบเรียนรู้ผิดพลาดหรือตอบสนองไม่สอดคล้องกับสถานการณ์จริง [9] ดังนั้นกรอบการวิจัยในงานนี้จึงมีได้มุ่งเพียงสร้างโมเดลที่แม่นยำ แต่ยังมุ่งสร้างสภาพแวดล้อมจำลองที่เปิดโอกาสให้เกิดการประเมินซ้ำ ปรับปรุงโมเดลอย่างต่อเนื่อง และสร้างต้นแบบที่นำไปใช้จริงเพื่อ ลดเวลาความเสียหายและช่วยให้การกู้คืนระบบทำได้อย่างรวดเร็วขึ้น เมื่อเกิดการโจมตีทางไซเบอร์ในสภาพแวดล้อมระดับองค์กรหรือระดับประเทศ [6] [9]

โดยสรุปงานวิจัยนี้เพื่อพัฒนากรอบการจัดกลุ่มภัยคุกคามด้วยปัญญาประดิษฐ์ ที่สามารถเชื่อมข้อมูลภัยคุกคาม การวิเคราะห์รูปแบบการโจมตี การประเมินผลกระทบทางธุรกิจ และการตอบสนองต่อเหตุการณ์แบบอัตโนมัติ ภายใต้สภาพแวดล้อมจำลองที่สมจริงและตรวจสอบได้อย่างเป็นระบบ อันจะนำไปสู่การยกระดับประสิทธิภาพการตรวจจับ การตอบสนอง และการกู้คืนระบบในงานความมั่นคงปลอดภัยไซเบอร์ระดับประเทศ [4] [6] [7]

2. วัตถุประสงค์

2.1 พัฒนารอบแนวคิดการจัดกลุ่มภัยคุกคามทางไซเบอร์ด้วยปัญญาประดิษฐ์ เพื่อจำแนกรูปแบบภัยคุกคามและสนับสนุนการตัดสินใจเชิงรุกอย่างเป็นระบบ

2.2 ประเมินประสิทธิภาพของกรอบแนวคิดที่พัฒนาขึ้น ในการลดระยะเวลาในการตรวจจับ วิเคราะห์ และตอบสนองต่อเหตุการณ์ความมั่นคงปลอดภัยไซเบอร์

2.3 ศึกษาความสามารถของกรอบแนวคิดในการลดผลกระทบทางธุรกิจ โดยเชื่อมโยงผลการจัดกลุ่มภัยคุกคาม เข้ากับการจัดลำดับความสำคัญของการตอบสนองต่อเหตุการณ์

3. สมมติฐานของการวิจัย/คำถามการวิจัย

3.1 กรอบแนวคิด AI-driven Threat Clustering สามารถจัดกลุ่มภัยคุกคามทางไซเบอร์ตามลักษณะพฤติกรรมและระดับความรุนแรงได้อย่างมีประสิทธิภาพ และเหนือกว่าวิธีการตรวจจับแบบดั้งเดิมหรือไม่

3.2 การใช้กรอบแนวคิดร่วมกับระบบตอบสนองอัตโนมัติ สามารถลดระยะเวลาในการตรวจจับและตอบสนองต่อเหตุการณ์ ได้อย่างมีนัยสำคัญหรือไม่

3.3 การประยุกต์ใช้ผลการจัดกลุ่มภัยคุกคามร่วมกับการประเมินผลกระทบทางธุรกิจ สามารถลดความเสียหาย และช่วยให้การกู้คืนระบบ ทำได้รวดเร็วขึ้นหรือไม่ ภายใต้กรอบมาตรฐานด้านความมั่นคงปลอดภัยไซเบอร์

4. ขอบเขตของงานวิจัย

4.1 ประชากร ได้แก่ สภาพแวดล้อมโครงสร้างพื้นฐานด้านเทคโนโลยีสารสนเทศและระบบความมั่นคงปลอดภัยไซเบอร์ขององค์กรในประเทศไทย ทั้งภาครัฐและเอกชน ซึ่งประกอบด้วยระบบเครือข่าย เซิร์ฟเวอร์ ฐานข้อมูล และกลไกป้องกัน เช่น Firewall, IDS/IPS และ SIEM โดยมุ่งเน้นบริบทขององค์กรที่มีความเสี่ยงต่อภัย เช่น Malware, Phishing, DDoS, APT และเหตุการณ์ผิดปกติที่มีลักษณะคล้าย Zero-day

4.2 กลุ่มตัวอย่าง ได้แก่ สภาพแวดล้อมจำลองโครงสร้างพื้นฐานไอทีอัจฉริยะ ที่ผู้วิจัยออกแบบขึ้นโดยอ้างอิงสถาปัตยกรรมขององค์กรทั่วไปในประเทศไทย และผสมผสาน benchmark intrusion datasets เช่น CICIDS2017 และ CSE-CIC-IDS2018 ร่วมกับสถานการณ์โจมตีสังเคราะห์ใน cyber range เพื่อรองรับการจัดกลุ่มภัยคุกคาม การประเมินผลกระทบ และการทดลอง response playbooks

4.3 ตัวแปรในการวิจัย

การวิจัยครั้งนี้กำหนดตัวแปรที่ใช้ศึกษาออกเป็น 2 ประเภท ดังนี้

4.3.1 ตัวแปรอิสระ ได้แก่

1. ประเภทของภัยคุกคามที่ถูกจำลอง (เช่น Malware, Phishing, DDoS, Zero-day-like anomaly)
2. เทคนิคและโมเดล AI/ML ที่ใช้ใน threat representation learning และ clustering (เช่น Random Forest, LSTM, Autoencoder, HDBSCAN/K-Means)
3. กลไกการตอบสนองอัตโนมัติ (เช่น Alert, Block, Isolate, Quarantine, Log)

4.3.2 ตัวแปรตาม ได้แก่

1. ความแม่นยำในการตรวจจับและจัดกลุ่มภัยคุกคาม
2. ระยะเวลาในการตอบสนองต่อเหตุการณ์
3. อัตราความผิดพลาดของระบบ (False Positive / False Negative)
4. คะแนนผลกระทบทางธุรกิจ (Business Impact Score)
5. ระดับความสอดคล้องกับกรอบ NIST CSF 2.0 และ ISO/IEC 27001

4.3.3 ระยะเวลาในการวิจัย

ระยะเวลาในการวิจัย ตั้งแต่เดือนกันยายน พ.ศ. 2569 ถึง สิงหาคม พ.ศ. 2570

5. วิธีดำเนินการวิจัย

การวิจัยครั้งนี้เป็นการวิจัยและพัฒนา (Research and Development: R&D) มีวัตถุประสงค์เพื่อศึกษาออกแบบ พัฒนา และประเมินกรอบ AI-driven Threat Clustering Framework สำหรับลดระยะเวลาในการตอบสนองต่อเหตุการณ์และลดผลกระทบทางธุรกิจจากภัยคุกคามไซเบอร์ โดยมุ่งเน้นให้ระบบสามารถทำหน้าที่วิเคราะห์ จัดกลุ่มภัยคุกคาม สนับสนุนการตัดสินใจ และช่วยผู้ดูแลระบบตอบสนองต่อเหตุการณ์ได้อย่างมีประสิทธิภาพ การดำเนินการวิจัยแบ่งออกเป็น 3 ระยะหลัก ดังนี้

5.1 ระยะที่1 ศึกษาและออกแบบกรอบแนวคิดการวิจัย

ในระยะนี้เป็นการศึกษาองค์ความรู้ วิเคราะห์ปัญหา และออกแบบกรอบแนวคิดของระบบ

5.1.1 ศึกษาเอกสาร แนวคิด ทฤษฎี และงานวิจัยที่เกี่ยวข้องกับ Artificial Intelligence , Machine Learning in Cybersecurity , Threat Detection , Threat Clustering , Incident Response Framework , Cybersecurity Standards (NIST CSF , ISO/IEC 27001)

5.1.2 ศึกษาข้อมูลและแหล่งข้อมูลที่ใช้ในการวิจัย เช่น ชุดข้อมูลมาตรฐาน (CICIDS2017, CSE-CIC-IDS2018) รูปแบบ log และ network traffic และ ประเภทภัยคุกคามไซเบอร์

5.1.3 วิเคราะห์ปัญหาและข้อจำกัดของระบบเดิม เช่น การตรวจจับแบบ signature-based

ความล่าช้าในการตอบสนอง การขาดการเชื่อมโยงกับผลกระทบทางธุรกิจ

5.1.4 สังเคราะห์องค์ประกอบสำคัญของระบบ ได้แก่ Input (Threat Data, Infrastructure Data) , Process (AI Learning, Clustering, Analysis) , Output (Response, Impact Reduction)

5.1.5 ออกแบบกรอบแนวคิดการวิจัย (Conceptual Framework) โดยกำหนดตัวแปรความสัมพันธ์และกระบวนการทำงานของระบบ

5.1.6 นำกรอบแนวคิดที่ออกแบบเสนอผู้เชี่ยวชาญเพื่อตรวจสอบความเหมาะสม ความถูกต้อง และความครอบคลุม ก่อนปรับปรุงให้สมบูรณ์

5.2 ระยะที่ 2 การพัฒนาแบบจำลองและระบบการทดลอง

ระยะนี้เป็นการพัฒนาระบบต้นแบบและดำเนินการทดลอง

5.2.1 ออกแบบสถาปัตยกรรมของระบบ ประกอบด้วย Simulation Environment , AI Model, Response Engine , Dashboard

5.2.2 ออกแบบและเตรียมข้อมูล ประกอบด้วย Cleaning , Preprocessing , Feature Engineering , Data Normalization

5.2.3 พัฒนาโมเดล AI สำหรับการวิเคราะห์ภัยคุกคาม ได้แก่ Supervised Learning(Random Forest) และ Deep Learning (LSTM, Autoencoder)

5.2.4 พัฒนาโมเดลการจัดกลุ่มภัยคุกคาม โดยใช้เทคนิค เช่น K-Means , DBSCAN , HDBSCAN

5.2.5 พัฒนาแบบจำลองการประเมินผลกระทบทางธุรกิจ เพื่อใช้ในการวิเคราะห์ระดับความเสียหาย และจัดลำดับความสำคัญของเหตุการณ์

5.2.6 พัฒนากลไกตอบสนองอัตโนมัติ เช่น การแจ้งเตือน การบล็อก การแยกระบบ และการบันทึกเหตุการณ์

5.2.7 พัฒนาแดชบอร์ดและส่วนแสดงผล เพื่อแสดงผลการวิเคราะห์ภัยคุกคาม การจัดกลุ่ม และผลกระทบทางธุรกิจ

5.2.8 ทดสอบการทำงานเบื้องต้นของระบบ และปรับปรุงระบบให้พร้อมสำหรับการทดลองจริง

5.3 ระยะที่ 3 การทดลองและประเมินผลระบบ

ระยะนี้เป็นการประเมินประสิทธิภาพของระบบที่พัฒนา

5.3.1 ออกแบบการทดลองโดยเปรียบเทียบระหว่าง ระบบเดิม (Traditional / Rule-based) และ ระบบที่พัฒนาขึ้น (AI-driven Framework)

5.3.2 ดำเนินการทดลองตามสถานการณ์จำลอง เช่น Malware Attack , Phishing Attack , DDoS Attack , Zero-day Scenario

5.3.3 เก็บรวบรวมข้อมูลผลการทดลอง เช่น Detection Result, Response Time, Error Rate

5.3.4 ประเมินผลประสิทธิภาพของระบบ โดยใช้ตัวชี้วัด ได้แก่ Accuracy, Precision, Recall, F1-score , Response Time, MTTR , False Positive / False Negative , Business Impact Score

5.3.4 ประเมินผลประสิทธิภาพของระบบ โดยใช้ตัวชี้วัด ได้แก่ Accuracy, Precision, Recall, F1-score , Response Time, MTTR , False Positive / False Negative , Business Impact Score

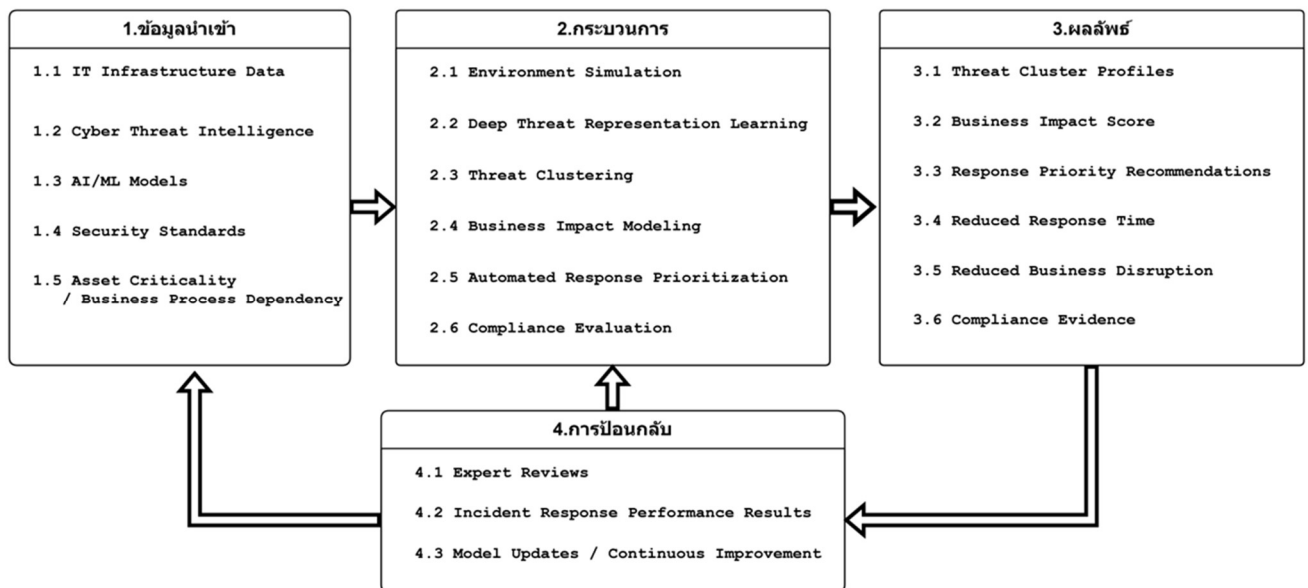
5.3.5 วิเคราะห์ผลและเปรียบเทียบกับระบบเดิม เพื่อประเมินว่าระบบที่พัฒนาขึ้นสามารถลดระยะเวลา

การตอบสนอง และลดผลกระทบทางธุรกิจได้อย่างมีนัยสำคัญหรือไม่

5.3.6 ประเมินความสอดคล้องกับมาตรฐานความมั่นคงปลอดภัยไซเบอร์
เช่น NIST CSF และ ISO/IEC 27001

5.3.7 สรุปผลการวิจัย อภิปรายผล และเสนอแนวทางปรับปรุงระบบในอนาคต

6. กรอบแนวคิดการวิจัย



ภาพที่ 1 กรอบแนวคิดการวิจัย AI-driven Threat Clustering Framework สำหรับการลดระยะเวลาในการตอบสนองและผลกระทบทางธุรกิจในสภาพแวดล้อมความมั่นคงปลอดภัยไซเบอร์ระดับประเทศ

จากภาพที่1 แสดงกรอบแนวคิดการวิจัยในครั้งนี้ เป็นการกำหนดโครงสร้างของระบบในลักษณะ IPOF (Input / Process / Output / Feedback) โดยมีองค์ประกอบในแต่ละส่วนดังนี้

1. ข้อมูลนำเข้า (Input)

ข้อมูลนี้ถูกนำมาใช้เพื่อสะท้อนบริบทของระบบจริง และเป็นข้อมูลพื้นฐานเพื่อใช้ในการวิเคราะห์รูปแบบของภัยคุกคามไซเบอร์ในเชิงลึก เป็นองค์ประกอบในส่วนของข้อมูลนำเข้าเป็นปัจจัยพื้นฐานที่ใช้ในกระบวนการวิเคราะห์ ประกอบด้วย

1.1 IT Infrastructure Data

เป็นข้อมูลโครงสร้างพื้นฐานด้านเทคโนโลยีสารสนเทศ เช่น ระบบเครือข่าย เซิร์ฟเวอร์ และระบบรักษาความปลอดภัย

1.2 Cyber Threat Intelligence

เป็นข้อมูลเกี่ยวกับภัยคุกคามไซเบอร์ รูปแบบการโจมตี และพฤติกรรมของผู้โจมตี

1.3 AI/ML Models

เป็นโมเดลปัญญาประดิษฐ์และการเรียนรู้ของเครื่องที่ใช้ในการวิเคราะห์ข้อมูล

1.4 Security Standards

เป็นมาตรฐานด้านความมั่นคงปลอดภัย เช่น NIST Cybersecurity Framework และ ISO/IEC 27001

1.5 Asset Criticality / Business Process Dependency

เป็นข้อมูลเกี่ยวกับความสำคัญของทรัพยากรและกระบวนการทางธุรกิจ

2. กระบวนการ (Process)

กระบวนการเป็นแกนหลักของการวิจัย โดยประกอบด้วย 6 ขั้นตอนสำคัญ ได้แก่

2.1 Environment Simulation

เป็นการจำลองสภาพแวดล้อมระบบไอทีเพื่อสร้างสถานการณ์ทดสอบที่ใกล้เคียงกับระบบจริง

2.2 Deep Threat Representation Learning

เป็นการใช้เทคนิค Deep Learning เพื่อเรียนรู้รูปแบบข้อมูลภัยคุกคามจาก log และ network traffic

2.3 Threat Clustering (แกนหลักของงานวิจัย)

เป็นการจัดกลุ่มภัยคุกคามโดยใช้เทคนิค Unsupervised Learning เพื่อแยกประเภทภัยคุกคามตามพฤติกรรม ช่วยให้สามารถตรวจจับภัยคุกคามแบบใหม่หรือ Zero-day ได้ดีขึ้น

2.4 Business Impact Modeling

เป็นการประเมินผลกระทบทางธุรกิจ โดยพิจารณาจากความรุนแรงของภัยและความสำคัญของระบบ

2.5 Automated Response Prioritization

เป็นการจัดลำดับความสำคัญของเหตุการณ์และดำเนินการตอบสนองแบบอัตโนมัติ

2.6 Compliance Evaluation

เป็นการประเมินความสอดคล้องของระบบกับมาตรฐานความมั่นคงปลอดภัย

3. ผลลัพธ์ (Output)

ผลลัพธ์ของกรอบแนวคิดครอบคลุมทั้งด้านเทคนิค การดำเนินงาน และการกำกับดูแล ประกอบด้วย

3.1 Threat Cluster Profiles

เป็นกลุ่มของภัยคุกคามที่ถูกจัดหมวดหมู่ตามพฤติกรรม

3.2 Business Impact Score

เป็นคะแนนประเมินผลกระทบทางธุรกิจ

3.3 Response Priority Recommendations

เป็นคำแนะนำในการตอบสนองตามลำดับความสำคัญ

3.4 Reduced Response Time

เป็นการลดระยะเวลาในการตรวจจับและตอบสนอง

3.5 Reduced Business Disruption

เป็นการลดผลกระทบต่อการดำเนินงานขององค์กร

3.6 Compliance Evidence

เป็นหลักฐานที่ใช้ในการประเมินความสอดคล้องกับมาตรฐาน

4. การป้อนกลับ (Feedback Loop)

ผลลัพธ์ดังกล่าวครอบคลุมทั้งด้านเทคนิค การดำเนินงาน และการกำกับดูแล กลไกนี้ช่วยให้ระบบสามารถเรียนรู้และปรับตัว (Adaptive System) ได้ตามสถานการณ์ที่เปลี่ยนแปลง

4.1 Expert Reviews

การประเมินและข้อเสนอแนะจากผู้เชี่ยวชาญ

4.2 Incident Response Performance Results

ผลลัพธ์ด้านประสิทธิภาพของการตอบสนอง

4.3 Model Updates / Continuous Improvement

การปรับปรุงโมเดล AI และระบบอย่างต่อเนื่อง

กรอบแนวคิดนี้แสดงให้เห็นว่า การบูรณาการข้อมูลโครงสร้างพื้นฐาน ขาวกรองภัยคุกคาม เทคโนโลยี AI/ML และการประเมินผลกระทบทางธุรกิจ ผ่านกระบวนการ Threat Clustering และการตอบสนองอัตโนมัติ จะช่วยเพิ่มประสิทธิภาพในการตรวจจับ ลดระยะเวลาในการตอบสนอง และลดผลกระทบทางธุรกิจได้อย่างมีนัยสำคัญ พร้อมทั้งสนับสนุนการพัฒนาระบบความมั่นคงปลอดภัยไซเบอร์เชิงรุกในระดับองค์กรและระดับประเทศ

7. นิยามศัพท์

7.1 กรอบการจัดกลุ่มภัยคุกคามด้วยปัญญาประดิษฐ์ (AI-driven Threat Clustering Framework) หมายถึง กรอบการวิจัยที่ผู้วิจัยพัฒนาขึ้นเพื่อใช้วิเคราะห์ข้อมูลภัยคุกคามทางไซเบอร์จากหลายแหล่ง โดยประยุกต์ใช้ปัญญาประดิษฐ์และการเรียนรู้ของเครื่องในการสร้างตัวแทนข้อมูล และจัดกลุ่มภัยคุกคามตามลักษณะพฤติกรรม รูปแบบการโจมตี และระดับความเสี่ยง เพื่อสนับสนุนการตัดสินใจในการตอบสนองต่อเหตุการณ์และลดผลกระทบทางธุรกิจอย่างเป็นระบบ

7.2 สภาพแวดล้อมจำลองอัจฉริยะ (Intelligent Simulation Environment) หมายถึง สภาพแวดล้อมจำลองด้านโครงสร้างพื้นฐานเทคโนโลยีสารสนเทศและความมั่นคงปลอดภัยไซเบอร์ที่ถูกออกแบบให้สามารถสร้างสถานการณ์โจมตี วิเคราะห์ข้อมูล และทดสอบการตอบสนองต่อเหตุการณ์ได้ภายใต้เงื่อนไขที่ควบคุมได้ ทำซ้ำได้ และไม่กระทบต่อระบบจริง โดยมีลักษณะใกล้เคียงกับสภาพแวดล้อมการปฏิบัติงานจริงขององค์กรหรือโครงสร้างพื้นฐานสำคัญ

7.3 การจัดกลุ่มภัยคุกคาม (Threat Clustering) หมายถึง กระบวนการใช้เทคนิค Unsupervised Learning หรือเทคนิคการเรียนรู้เชิงลึก เพื่อจัดกลุ่มเหตุการณ์หรือพฤติกรรมทางไซเบอร์ที่มีรูปแบบใกล้เคียงกันให้อยู่ในกลุ่มเดียวกัน โดยมีเป้าหมายเพื่อค้นหารูปแบบภัยคุกคามที่ซ่อนอยู่ ตรวจสอบภัยที่ไม่เคยปรากฏมาก่อน และลดความซับซ้อนของข้อมูลเหตุการณ์จำนวนมากให้อยู่ในรูปที่สามารถนำไปใช้ตัดสินใจเชิงปฏิบัติการได้ง่ายขึ้น

7.4 ขาวกรองภัยคุกคามไซเบอร์ (Cyber Threat Intelligence) หมายถึง ข้อมูล ความรู้ หรือข้อบ่งชี้ที่เกี่ยวข้องกับภัยคุกคามทางไซเบอร์ เช่น รูปแบบการโจมตี ตัวบ่งชี้การคุกคาม วิธีการและเทคนิคของผู้โจมตี

ตลอดจนบริบทของเหตุการณ์ ซึ่งนำมาใช้เพื่อเพิ่มความสามารถในการวิเคราะห์ ตรวจสอบ และตอบสนองต่อภัยคุกคามได้อย่างแม่นยำและทันที่

7.5 ระยะเวลาในการตอบสนองต่อเหตุการณ์ (Incident Response Time) หมายถึง ระยะเวลาตั้งแต่ระบบสามารถตรวจพบหรือระบุเหตุการณ์ที่เป็นภัยคุกคามทางไซเบอร์ได้ ไปจนถึงการเริ่มต้นหรือดำเนินการตอบสนองที่กำหนดไว้แล้วเสร็จในระดับที่สามารถควบคุมผลกระทบของเหตุการณ์ได้ ในงานวิจัยนี้ใช้เป็นตัวชี้วัดสำคัญของประสิทธิภาพเชิงปฏิบัติการของระบบ เนื่องจากแนวทาง incident response สมัยใหม่มุ่งลดทั้งจำนวนและผลกระทบของเหตุการณ์ผ่านการตรวจจับ ตอบสนอง และกู้คืนที่มีประสิทธิภาพมากขึ้น

7.6 ผลกระทบทางธุรกิจ (Business Impact) หมายถึง ระดับความเสียหายหรือผลกระทบที่เหตุการณ์ทางไซเบอร์ก่อให้เกิดต่อการดำเนินงานขององค์กร เช่น การหยุดชะงักของบริการ ความล่าช้าของกระบวนการสำคัญ การสูญเสียข้อมูล ความเสียหายทางการเงิน หรือผลกระทบต่อความเชื่อมั่นของผู้มีส่วนได้ส่วนเสีย โดยในงานวิจัยนี้ใช้เป็นแนวคิดสำคัญในการเชื่อมผลการวิเคราะห์ทางเทคนิคไปสู่การจัดลำดับความสำคัญของการตอบสนองและการกู้คืนระบบ

7.7 คะแนนผลกระทบทางธุรกิจ (Business Impact Score: BIS) หมายถึง ค่าดัชนีเชิงผสมที่ผู้วิจัยกำหนดขึ้นเพื่อใช้ประเมินระดับความสำคัญของเหตุการณ์ทางไซเบอร์ โดยอิงจากปัจจัย เช่น ความรุนแรงของภัยคุกคาม ความสำคัญของทรัพยากรสารสนเทศ ความเชื่อมโยงกับกระบวนการทางธุรกิจ และความไวต่อการหยุดชะงักของบริการ ทั้งนี้ BIS ถูกใช้เพื่อสนับสนุนการจัดลำดับความสำคัญของการตอบสนองและเพื่อประเมินว่าระบบที่พัฒนาขึ้นสามารถลดผลกระทบทางธุรกิจได้มากเพียงใด

7.8 การจัดลำดับความสำคัญของการตอบสนองอัตโนมัติ (Automated Response Prioritization) หมายถึง กลไกที่ใช้ผลจากการจัดกลุ่มภัยคุกคามและการประเมินผลกระทบทางธุรกิจในการกำหนดลำดับความเร่งด่วนของการตอบสนองต่อเหตุการณ์ เช่น การแจ้งเตือน การบล็อก การกักกันระบบ หรือการยกระดับเหตุการณ์ไปยังผู้เชี่ยวชาญ โดยมีจุดมุ่งหมายเพื่อให้การตอบสนองเป็นไปอย่างเหมาะสมกับระดับความเสี่ยงและลดช่วงเวลาความเสียหายขององค์กรให้มากที่สุด

7.9 มาตรฐาน NIST Cybersecurity Framework 2.0 (NIST CSF 2.0) หมายถึง กรอบมาตรฐานด้านความมั่นคงปลอดภัยไซเบอร์ที่จัดทำโดย National Institute of Standards and Technology (NIST) เพื่อใช้เป็นแนวทางในการบริหารความเสี่ยงไซเบอร์ขององค์กร โดยประกอบด้วยฟังก์ชันหลัก ได้แก่ Govern, Identify, Protect, Detect, Respond และ Recover ซึ่งในงานวิจัยนี้ใช้เป็นกรอบอ้างอิงในการประเมินความพร้อมและความสอดคล้องของระบบที่พัฒนาขึ้น

7.10 มาตรฐาน ISO/IEC 27001:2022 หมายถึง มาตรฐานสากลด้านระบบบริหารจัดการความมั่นคงปลอดภัยสารสนเทศ ซึ่งกำหนดข้อกำหนดสำหรับการจัดตั้ง การดำเนินการ การคงไว้ และการปรับปรุงระบบบริหารจัดการความมั่นคงปลอดภัยสารสนเทศอย่างต่อเนื่อง ในงานวิจัยนี้ใช้เป็นกรอบอ้างอิงในการประเมินความพร้อมและความสอดคล้องของระบบกับข้อกำหนดด้านการบริหารความเสี่ยง การควบคุม และความยืดหยุ่นขององค์กรต่อภัยคุกคามทางไซเบอร์

7.11 การเรียนรู้เชิงลึกตัวแทนภัยคุกคาม (Deep Threat Representation Learning)

หมายถึง กระบวนการใช้โมเดล Deep Learning ในการเรียนรู้คุณลักษณะเชิงลึกของข้อมูล log, traffic หรือเหตุการณ์ทางไซเบอร์ เพื่อแปลงข้อมูลดิบที่มีมิติสูงให้เป็นตัวแทนข้อมูลที่เหมาะสมต่อการจัดกลุ่มและการวิเคราะห์ภัยคุกคาม โดยมีเป้าหมายเพื่อเพิ่มความสามารถของโมเดลในการตรวจจับรูปแบบภัยที่ซับซ้อนหรือเปลี่ยนแปลงตลอดเวลา

7.12 การโจมตีต่อโมเดลปัญญาประดิษฐ์ (Adversarial Machine Learning) หมายถึง รูปแบบการโจมตีที่มุ่งแทรกแซงหรือบิดเบือนกระบวนการเรียนรู้และการตัดสินใจของระบบ AI/ML เช่น การโจมตีแบบ evasion, poisoning และ privacy attacks ซึ่งอาจทำให้โมเดลตรวจจับภัยคุกคามผิดพลาดหรือมีความน่าเชื่อถือลดลง ในงานวิจัยนี้ใช้เป็นแนวคิดประกอบเพื่อพิจารณาความทนทานและข้อจำกัดของระบบที่พัฒนาขึ้นในสภาพแวดล้อมภัยคุกคามยุคใหม่

8. ประโยชน์ที่คาดว่าจะได้รับ

8.1 ได้องค์ความรู้เชิงลึกในการพัฒนา กรอบการจัดกลุ่มภัยคุกคามด้วยปัญญาประดิษฐ์ ที่สามารถบูรณาการเทคนิคการเรียนรู้ของเครื่อง การวิเคราะห์ข้อมูลภัยคุกคาม และการประเมินผลกระทบทางธุรกิจ เพื่อยกระดับประสิทธิภาพของระบบความมั่นคงปลอดภัยไซเบอร์ในเชิงองค์กรรวม

8.2 ได้กรอบแนวคิดและแนวทางในการออกแบบระบบวิเคราะห์ภัยคุกคามที่เชื่อมโยง การตรวจจับ การตอบสนอง และการกู้คืน เข้าด้วยกันอย่างเป็นระบบ ซึ่งสามารถใช้เป็นต้นแบบสำหรับการพัฒนาระบบ Cybersecurity เชิงรุก ในระดับองค์กรและระดับประเทศ

8.3 ได้ระบบต้นแบบที่สามารถจัดกลุ่มภัยคุกคาม วิเคราะห์รูปแบบการโจมตี และประเมินผลกระทบทางธุรกิจ เพื่อสนับสนุนการตัดสินใจและการจัดลำดับความสำคัญในการตอบสนองต่อเหตุการณ์ได้อย่างมีประสิทธิภาพและเป็นระบบ

8.4 ช่วยลดระยะเวลาในการตรวจจับและตอบสนองต่อเหตุการณ์ รวมถึงลดช่วงเวลาความเสียหายของระบบ และเพิ่มความสามารถในการกู้คืนระบบ ซึ่งเป็นตัวชี้วัดสำคัญของประสิทธิภาพเชิงปฏิบัติการด้านความมั่นคงปลอดภัยไซเบอร์

8.5 ช่วยลดผลกระทบทางธุรกิจ จากเหตุการณ์ภัยคุกคามไซเบอร์ โดยผ่านการประเมินระดับความเสี่ยง และการจัดลำดับความสำคัญของการตอบสนอง ซึ่งส่งผลต่อความต่อเนื่องทางธุรกิจ และความยืดหยุ่นขององค์กร

8.6 ได้แนวทางและกลไกในการประเมินความสอดคล้องของระบบความมั่นคงปลอดภัยไซเบอร์กับมาตรฐานสากล เช่น NIST Cybersecurity Framework (CSF 2.0) และ ISO/IEC 27001 ซึ่งสามารถใช้เป็นหลักฐานเชิงประจักษ์ในการตรวจประเมินและยกระดับการกำกับดูแลด้านความมั่นคงปลอดภัย

8.7 สร้างองค์ความรู้และต้นแบบที่สามารถนำไปต่อยอดในงานวิจัยและการใช้งานจริง เช่น ระบบ Security Operations Center (SOC) ระบบ SOAR (Security Orchestration, Automation, and Response) ระบบวิเคราะห์ภัยคุกคามเชิงอัจฉริยะ

8.8 ส่งเสริมการพัฒนาแนวคิดและเทคโนโลยีด้าน AI for Cybersecurity โดยเฉพาะในบริบทของการจัดการภัยคุกคามเชิงซับซ้อนและการตัดสินใจอัตโนมัติ ซึ่งสามารถนำไปประยุกต์ใช้กับองค์กรภาครัฐ ภาคเอกชน และโครงสร้างพื้นฐานสำคัญของประเทศในอนาคต

9. เอกสารอ้างอิง

- [1] A. H. Salem, S. M. Azzam, O. E. Emam, and A. A. Abohany, “Advancing cybersecurity: a comprehensive review of AI-driven detection techniques,” *J. Big Data*, vol. 11, no. 1, Dec. 2024, doi: 10.1186/s40537-024-00957-y.
- [2] A. Vassilev, A. Oprea, A. Fordyce, H. Anderson, X. Davies, and M. Hamin, “NIST Trustworthy and Responsible AI NIST AI 100-2e2025 Adversarial Machine Learning A Taxonomy and Terminology of Attacks and Mitigations”, doi: 10.6028/NIST.AI.100-2e2025.
- [3] A. Kumar and D. Pandey, “Enhancing intrusion detection with ML and deep learning: A survey of CICIDS 2017 and CSE-CIC-IDS2018 datasets,” *AIP Conf. Proc.*, vol. 3168, no. 1, Jul. 2024, doi: 10.1063/5.0222131/3300765.
- [4] A. Jaber and L. Fritsch, “Towards AI-powered Cybersecurity Attack Modeling with Simulation Tools: Review of Attack Simulators,” *Lecture Notes in Networks and Systems*, vol. 571 LNNS, pp. 249–257, 2023, doi: 10.1007/978-3-031-19945-5_25.
- [5] A. I. Mallick and R. Nath, “Simulating Cyber Threats: A Review of AI-powered Attack Simulators for Enhanced Cybersecurity,” *World Sci. News*, 2024.
- [6] A. Nelson, S. Rekhi, M. Souppaya, and K. Scarfone, “NIST Special Publication 800 NIST SP 800-61r3 Incident Response Recommendations and Considerations for Cybersecurity Risk Management A CSF 2.0 Community Profile”, doi: 10.6028/NIST.SP.800-61r3.
- [7] C. Pascoe, S. Quinn, K. Scarfone, C. Pascoe, S. Quinn, and K. Scarfone, “The NIST Cybersecurity Framework (CSF) 2.0,” Feb. 2024, doi: 10.6028/NIST.CSWP.29.
- [8] “ISO/IEC 27001:2022 - Information security management systems.” Accessed: Jun. 08, 2026. [Online]. Available: <https://www.iso.org/standard/27001>
- [9] A. Vassilev, A. Oprea, A. Fordyce, H. Anderson, X. Davies, and M. Hamin, “Adversarial Machine Learning: A Taxonomy and Terminology of Attacks and Mitigations,” Mar. 2025, doi: 10.6028/NIST.AI.100-2E2025.