

บทบาท ประโยชน์ และความท้าทายของการประยุกต์ใช้ปัญญาประดิษฐ์เพื่อ เสริมสร้างความมั่นคงปลอดภัยทางไซเบอร์

The Role Benefits and Challenges of Applying Artificial Intelligence for Enhancing Cybersecurity

ปาริตา ไพศาลรุ่งพนา¹ พัฒน์ศรีณัฐ เลหาไพบูลย์² อรณิชา คงวุฒิ^{3*} และวุฒิรงค์ คงวุฒิ⁴
Palita Pisanrungpana¹, Phatsaran Laohhapaiboon², Ornnicha Kongwut^{3*} and Wuttirong Kongwut⁴

¹หลักสูตรวิทยา-คณิต โรงเรียนเทพศิรินทร์สมุทรปราการ

¹Science-Mathematics Course, Thep Sirin School Samut Prakan

²หจก.บุญพาวาสนาสง

²Boonpawassanasong Partnership

³สาขาวิชาฟิสิกส์ คณะวิทยาศาสตร์และเทคโนโลยี มหาวิทยาลัยราชภัฏกาญจนบุรี

³Department of Physics, Faculty of Science and Technology, Kanchanaburi Rajabhat University

⁴ฝูงบิน 203 กองบิน 2 กองทัพอากาศไทย

⁴203 Squadron Wing, Royal Thai Air Force

*Corresponding author: ornnicha.k@kru.ac.th

Received: April 25, 2024

Revised: July 25, 2024

Accepted: July 30, 2024

บทคัดย่อ

ภัยคุกคามทางไซเบอร์มีแนวโน้มทวีความรุนแรงและซับซ้อนมากขึ้นอย่างต่อเนื่อง ทั้งการโจมตีด้วยมัลแวร์ การโจมตีแบบปฏิเสธการให้บริการ และภัยคุกคามประเภทใหม่อย่างการเรียกค่าไถ่ ส่งผลกระทบอย่างมหาศาลต่อองค์กรและบุคคลทั่วโลก ระบบรักษาความปลอดภัยแบบดั้งเดิมมีข้อจำกัดในการรับมือ การนำเทคโนโลยีปัญญาประดิษฐ์ (AI) มาประยุกต์ใช้จึงเป็นทางเลือกที่น่าสนใจ บทความวิชาการนี้มีวัตถุประสงค์เพื่อทบทวนวรรณกรรมที่เกี่ยวข้องกับการประยุกต์ใช้ AI ในการรับมือกับภัยคุกคามไซเบอร์ในช่วงปี ค.ศ. 2014-2024 โดยใช้วิธีการทบทวนวรรณกรรมอย่างเป็นระบบ ศึกษาเปรียบเทียบผลงานวิจัยทั้งในแง่เทคนิค AI ที่นำมาใช้ ประเภทภัยคุกคามที่ศึกษา และประสิทธิภาพของระบบ ผลการศึกษาพบว่า การประยุกต์ใช้เทคนิค AI หลากหลายแขนง ไม่ว่าจะเป็น Machine Learning Deep Learning หรือการผสมผสานหลายเทคนิค สามารถเพิ่มความแม่นยำในการตรวจจับภัยคุกคาม ความรวดเร็วในการตอบสนอง และขีดความสามารถในการเรียนรู้รูปแบบใหม่ ๆ ได้สูงกว่าวิธีการแบบดั้งเดิมอย่างมีนัยสำคัญ โดยเฉพาะการใช้ Deep Learning ในการวิเคราะห์มัลแวร์ที่มีความแม่นยำสูงถึง 98% และเร็วกว่าการวิเคราะห์ด้วยมนุษย์ถึง 20 เท่า

คำสำคัญ: ปัญญาประดิษฐ์ ความปลอดภัยไซเบอร์ การตรวจจับภัยคุกคาม การวิเคราะห์มัลแวร์

Abstract

Cyber threats are continuously evolving, becoming more severe and complex. Malware attacks, denial-of-service attacks, and new threats such as ransomware severely impact organizations and individuals worldwide. Traditional security systems have limitations in dealing with these threats, making the application of Artificial Intelligence (AI) an attractive alternative. This academic article aims to review the literature on applying AI to tackle cyber threats from 2014-2024, using a systematic literature review methodology. It compares research papers in terms of AI techniques used, types of threats studied, and system performance. The study found that applying various AI techniques, whether machine learning, deep learning, or a combination of techniques, can significantly improve threat detection accuracy, response speed, and ability to learn new patterns compared to traditional methods. Notably, the use of deep learning in malware analysis showed up to 98% accuracy and was 20 times faster than human analysis.

Keywords: artificial intelligence, cybersecurity, threats detection, malware analysis



บทนำ

ในยุคดิจิทัลปัจจุบัน ภัยคุกคามทางไซเบอร์ได้ทวีความรุนแรงและซับซ้อนมากขึ้นอย่างต่อเนื่อง ส่งผลกระทบอย่างมหาศาลต่อทั้งองค์กรและบุคคลทั่วโลก การรับมือกับภัยคุกคามเหล่านี้จำเป็นต้องอาศัยหลักการพื้นฐานของ Cybersecurity ที่มุ่งรักษา Confidentiality (ความลับ) Integrity (ความถูกต้องครบถ้วน) และ Availability (ความพร้อมใช้งาน) ของข้อมูลและระบบ (Kizza, 2020)

หนึ่งในภัยคุกคามที่สร้างความเสียหายอย่างมาก คือ การโจมตีแบบ Ransomware ซึ่งเป็นมัลแวร์ที่เข้ารหัสข้อมูลและเรียกค่าไถ่ โดยในปี 2021 พบว่ามีองค์กรกว่า 68% ถูกโจมตีด้วย Ransomware และในจำนวนนั้นมีผู้ที่สูญเสียข้อมูลถึง 82% (Roger, 2022; Wollerton, 2023) นอกจากนี้ การโจมตีแบบ Distributed Denial of Service--DDoS ที่ใช้ปริมาณ Traffic มหาศาลเพื่อทำให้ระบบไม่สามารถให้บริการได้ ก็มีแนวโน้มเพิ่มขึ้นอย่างมาก โดยในปี 2022 มีการโจมตี DDoS เพิ่มขึ้นถึง 150% เมื่อเทียบกับปีก่อน (Cox, 2023; Hassan, 2019)

ระบบรักษาความปลอดภัยแบบดั้งเดิมที่อาศัยการวิเคราะห์โดยผู้เชี่ยวชาญเป็นหลักนั้น มีข้อจำกัดในการรับมือภัยคุกคามที่ซับซ้อนและเปลี่ยนแปลงอย่าง

รวดเร็วในปัจจุบัน ด้วยเหตุนี้ การนำเทคโนโลยีปัญญาประดิษฐ์ (Artificial Intelligence--AI) โดยเฉพาะเทคนิคการเรียนรู้ของเครื่อง (machine learning) และการเรียนรู้เชิงลึก (deep learning) มาประยุกต์ใช้จึงเป็นทางเลือกที่น่าสนใจ เพื่อเพิ่มประสิทธิภาพในการตรวจจับ ป้องกัน และตอบสนองต่อภัยคุกคามได้อย่างรวดเร็ว แม่นยำ และเรียนรู้ได้ด้วยตนเอง (Accenture, 2019)

อย่างไรก็ตาม การนำ AI มาใช้งานจริงในด้านนี้ก็ยังมีความท้าทายที่สำคัญ ทั้งในแง่ข้อจำกัดของเทคโนโลยี การขาดแคลนข้อมูลและบุคลากรไปจนถึงประเด็นด้านจริยธรรมและกฎระเบียบที่เกี่ยวข้อง (Kumar, Han & Soni 2022; Zeadally et al., 2020)

บทความนี้จะทบทวนวรรณกรรมและงานวิจัยที่เกี่ยวข้องกับการประยุกต์ใช้ AI ในการรับมือกับภัยคุกคามทางไซเบอร์ เพื่อสรุปแนวโน้ม บทบาท ประโยชน์ ตลอดจนข้อจำกัดและความท้าทายสำคัญในการพัฒนา AI ด้านนี้ให้เกิดประโยชน์สูงสุดต่อไปในอนาคต

ขั้นตอนการดำเนินการ

บทความวิชาการนี้ดำเนินการโดยการทบทวนและวิเคราะห์เอกสารและบทความวิจัยที่เกี่ยวข้องกับการประยุกต์ใช้

เทคโนโลยีปัญญาประดิษฐ์ในการรักษาความมั่นคงปลอดภัยทางไซเบอร์ ระหว่างช่วงปี ค.ศ. 2014-2024 โดยมีขั้นตอนดังนี้

1. การสืบค้นและรวบรวมข้อมูล

ดำเนินการสืบค้นและรวบรวมบทความวิชาการ บทความวิชาการ งานวิทยานิพนธ์ และรายงานทางวิชาการที่เกี่ยวข้อง จากฐานข้อมูลวารสารและสิ่งพิมพ์วิชาการออนไลน์ที่น่าเชื่อถือ อาทิ IEEE Xplore ScienceDirect ACM Digital Library Scopus เป็นต้น โดยใช้คำค้นหาหลักเกี่ยวกับ “ปัญญาประดิษฐ์” “ความมั่นคงปลอดภัยทางไซเบอร์” “การตรวจจับภัยคุกคาม” เป็นต้น

2. การคัดเลือกบทความที่เกี่ยวข้อง

จากนั้นทำการคัดกรองบทความที่เกี่ยวข้องกับหัวข้อการศึกษาโดยตรง โดยพิจารณาจากบทคัดย่อและสาระสำคัญของบทความ เพื่อคัดเลือกบทความที่มีการนำเสนอการประยุกต์ใช้เทคนิคปัญญาประดิษฐ์ต่าง ๆ เช่น Machine Learning Deep Learning เป็นต้น ในการตรวจจับภัยคุกคาม วิเคราะห์มัลแวร์ หรือเพื่อวัตถุประสงค์ด้านความปลอดภัยไซเบอร์โดยตรง

เกณฑ์การคัดเข้า

- บทความที่ตีพิมพ์ระหว่างปี ค.ศ. 2014-2024
- บทความที่เขียนเป็นภาษาอังกฤษหรือภาษาไทย
- บทความที่นำเสนอการประยุกต์ใช้ AI ในด้านความมั่นคงปลอดภัยทางไซเบอร์
- บทความที่มีการอธิบายวิธีการวิจัยและผลลัพธ์อย่างชัดเจน

เกณฑ์การคัดออก

- บทความที่ไม่ได้ผ่านการประเมินคุณภาพจากผู้ทรงคุณวุฒิ (non-peer-reviewed)
- บทความที่ไม่ได้เน้นการประยุกต์ใช้ AI ในด้านความมั่นคงปลอดภัยทางไซเบอร์โดยตรง
- บทความที่ไม่มีรายละเอียดเพียงพอเกี่ยวกับวิธีการวิจัยหรือผลลัพธ์

3. การวิเคราะห์และสังเคราะห์ข้อมูล

นำบทความที่คัดเลือกมาทำการวิเคราะห์และสังเคราะห์เนื้อหาสำคัญในประเด็นต่าง ๆ ดังนี้

- เทคนิคปัญญาประดิษฐ์ที่ใช้ ได้แก่ Machine Learning, Deep Learning และอื่น ๆ
- ประเภทของภัยคุกคามไซเบอร์ที่ศึกษา เช่น มัลแวร์ การบุกรุกเครือข่าย บอทเน็ต เป็นต้น
- ประสิทธิภาพและข้อดีของการประยุกต์ใช้ AI เมื่อเปรียบเทียบกับวิธีการแบบดั้งเดิม
- ข้อจำกัด ปัญหาอุปสรรค และความท้าทายของการนำ AI มาประยุกต์ใช้
- แนวโน้มและข้อเสนอแนะสำหรับการพัฒนาต่อยอดในอนาคต

วิธีการวิเคราะห์

- ใช้การวิเคราะห์เนื้อหา (content analysis) เพื่อระบุประเด็นสำคัญและแนวโน้มในการวิจัย
- ใช้การวิเคราะห์เปรียบเทียบ (comparative analysis) เพื่อเปรียบเทียบประสิทธิภาพของเทคนิค AI ต่าง ๆ
- ใช้การวิเคราะห์ช่องว่าง (gap analysis) เพื่อระบุประเด็นที่ยังขาดการศึกษาหรือต้องการการวิจัยเพิ่มเติม

วิธีการสังเคราะห์

- ใช้การสังเคราะห์แบบบรรยาย (narrative synthesis) เพื่อสรุปภาพรวมของการประยุกต์ใช้ AI ในด้านความมั่นคงปลอดภัยทางไซเบอร์
- สร้างกรอบแนวคิด (conceptual framework) เพื่อแสดงความสัมพันธ์ระหว่างเทคนิค AI ประเภทของภัยคุกคาม และประสิทธิภาพในการรับมือ
- จัดทำตารางสรุปเปรียบเทียบผลการศึกษากจากงานวิจัยต่าง ๆ

จากนั้นนำผลการวิเคราะห์และสังเคราะห์มาสรุปและนำเสนอในรูปแบบของบทความวิชาการ โดยแบ่งเป็นหัวข้อต่าง ๆ เช่น บทนำ แนวคิด และทฤษฎีที่เกี่ยวข้อง บทบาทของ AI ในการป้องกันภัยคุกคาม ข้อดีและข้อจำกัด รวมถึงสรุปและข้อเสนอแนะ

วรรณกรรมและงานวิจัยที่เกี่ยวข้อง

ในช่วงไม่กี่ปีที่ผ่านมา งานวิจัยจำนวนมากได้มุ่งเน้นไปที่การประยุกต์ใช้เทคนิคปัญญาประดิษฐ์ (AI) และการเรียนรู้ของเครื่อง (machine learning) เพื่อเพิ่มประสิทธิภาพในการรับมือกับภัยคุกคามทางไซเบอร์ที่ทวีความซับซ้อนมากขึ้นเรื่อย ๆ

หนึ่งในแนวทางที่ได้รับความสนใจอย่างกว้างขวาง คือ การนำอัลกอริทึม Machine Learning เช่น Decision Tree, Support Vector Machine--SVM และ Deep Learning มาใช้ในการตรวจจับและทำนายการโจมตี ซึ่งจากการทบทวนวรรณกรรมอย่างเป็นระบบของ Sarker et al. (2020) พบว่า เทคนิคเหล่านี้ให้ผลลัพธ์ที่ดีในการจำแนกพฤติกรรมที่ผิดปกติและระบุภัยคุกคามได้อย่างแม่นยำ อย่างไรก็ตาม ยังมีความท้าทายในการเลือกคุณลักษณะ (feature) ที่เหมาะสม การปรับพารามิเตอร์ของโมเดล ตลอดจนประเด็นด้านจริยธรรม และความเป็นส่วนตัวของข้อมูล

เพื่อแก้ปัญหาความเป็นส่วนตัวและข้อจำกัดทางกฎหมายในการแบ่งปันข้อมูล Liu et al. (2024) จึงเสนอแนวคิดการใช้ Federated Learning ซึ่งเป็นวิธีการกระจายการฝึกโมเดล Machine Learning โดยไม่ต้องรวบรวมข้อมูลไว้ที่ส่วนกลาง ผลการทดลองเบื้องต้นชี้ให้เห็นศักยภาพในการรักษาความแม่นยำของโมเดลได้ใกล้เคียงกับวิธีแบบรวมศูนย์ แม้จะมีข้อท้าทายบางประการในการประยุกต์ใช้งานจริง เช่น ปัญหา Non-IID Data และความล่าช้าในการสื่อสาร นอกจากนี้ Mothukuri et al. (2021) ยังได้วิเคราะห์ความเสี่ยงด้านความปลอดภัยและความเป็นส่วนตัวของ Federated Learning พร้อมเสนอแนวทางการป้องกัน เช่น การเข้ารหัสข้อมูล การตรวจสอบสิทธิ์ผู้เข้าร่วม และการปกป้องความเป็นส่วนตัวด้วยวิธี Differential Privacy

อีกแนวทางที่น่าสนใจ คือ การประยุกต์ใช้ Graph Neural Network--GNN ในการวิเคราะห์โครงสร้างของเครือข่ายและตรวจจับพฤติกรรมที่ผิดปกติ ดังเช่นในงานของ Srinivasan & Sharmili (2022) ที่นำ GNN มาใช้ตรวจจับการบุกรุกในระบบเครือข่ายได้ผลดีกว่าเทคนิคดั้งเดิม แต่ก็มีข้อจำกัดในเรื่องความซับซ้อนของการคำนวณและการขยายขนาดเพื่อรองรับข้อมูลขนาดใหญ่ในเครือข่ายจริง Zhong et al. (2020) จึงเสนอเทคนิค Hierarchical GNN เพื่อลดความซับซ้อน โดยแบ่งกราฟขนาดใหญ่ออก

เป็นกราฟย่อย ๆ ก่อนจะทำการเรียนรู้ในแต่ละขั้นแล้วรวมผลลัพธ์เข้าด้วยกันในระดับสูงขึ้นไป ส่วน Li et al. (2022) ได้ผสมผสานข้อดีของ GNN และ Federated Learning เพื่อสร้างโมเดลตรวจจับการบุกรุกที่มีประสิทธิภาพสูงและปกป้องความเป็นส่วนตัวของข้อมูลไปพร้อมกัน

ในด้านการแก้ปัญหา Black Box ของโมเดล AI ซึ่งเป็นอุปสรรคต่อการสร้างความเชื่อมั่น ก็มีการนำเทคนิค Explainable AI--XAI มาประยุกต์ใช้เพื่ออธิบายกระบวนการตัดสินใจของโมเดลให้เข้าใจได้ง่ายขึ้น โดย Warnecke et al. (2020) ได้นำเสนอเทคนิค LEMNA ซึ่งใช้ Shapley Value และ k-Means Clustering ในการสร้างคำอธิบายแบบ Local สำหรับการทำนายของโมเดล ML ในบริบทของ Cybersecurity ส่วน Marino et al. (2021) ได้พัฒนารอบการทดสอบ metamorphic สำหรับการตรวจสอบความมั่นคงปลอดภัยของ Cyber-Physical Systems ที่ใช้ ML ให้สามารถทำงานได้อย่างน่าเชื่อถือโดยอัตโนมัติ อย่างไรก็ตาม แนวทาง XAI สำหรับ Deep Learning โมเดลที่มีความซับซ้อนสูงยังคงเป็นความท้าทายที่ยังต้องมีการศึกษากันต่อไป (Nanda et al., 2022)

นอกจากนั้น การรับมือกับภัยคุกคามประเภท Adversarial Attack ซึ่งมุ่งบิดเบือนหรือหลอกลวงโมเดล ML ก็เป็นหัวข้อที่ได้รับความสนใจอย่างมาก โดย Guo et al. (2020) ได้ทบทวนงานวิจัยเกี่ยวกับ Adversarial Attack และวิธีการป้องกันในงานด้าน Cybersecurity อย่างเช่น การเพิ่ม Adversarial Training การใช้เทคนิค Ensemble และการตรวจจับการโจมตีด้วยการวิเคราะห์ Activation Value ซึ่งแม้จะช่วยบรรเทาปัญหาได้ระดับหนึ่ง แต่ก็ยังมีข้อจำกัดเมื่อใช้กับโมเดลและข้อมูลที่ซับซ้อน ขณะที่ Zhang et al. (2021) เสนอ Adversarial Attack Detection Framework ที่ผสมผสานเทคนิคการเรียนรู้แบบ Meta-Learning และ Few-Shot Learning เพื่อสร้างตัวตรวจจับที่สามารถปรับตัวให้เข้ากับการโจมตีรูปแบบใหม่ได้อย่างรวดเร็วและมีประสิทธิภาพ แม้จะใช้ข้อมูลในการฝึกฝนเพียงเล็กน้อยก็ตาม

จากการสังเคราะห์งานวิจัยข้างต้น จะเห็นได้ว่า แม้เทคนิค AI และ Machine Learning จะมีศักยภาพสูงในการยกระดับการรักษาความปลอดภัยไซเบอร์ แต่การนำไปใช้งานจริงยังมีความท้าทายที่หลากหลาย ทั้งในด้าน

ประสิทธิภาพและความแม่นยำของโมเดล การจัดการกับข้อมูลที่มีความซับซ้อนสูง การปกป้องความเป็นส่วนตัวและการลดภัยของข้อมูล การอธิบายกระบวนการตัดสินใจให้โปร่งใส ไปจนถึงการรับมือกับ Adversarial Attack ที่ซับซ้อนมากขึ้น โดยปัจจัยเหล่านี้ยังคงเป็นช่องว่างสำคัญที่ต้องอาศัยการวิจัยและพัฒนาอย่างต่อเนื่อง เพื่อให้สามารถนำศักยภาพของ AI มาใช้ประโยชน์ในการเสริมสร้างความมั่นคงปลอดภัยทางไซเบอร์ได้อย่างเต็มที่และปลอดภัยต่อไปในอนาคต

บทความนี้มุ่งเติมเต็มช่องว่างดังกล่าวโดย (1) นำเสนอการวิเคราะห์เปรียบเทียบประสิทธิภาพของเทคนิค AI ต่าง ๆ จากการรวบรวมผลการวิจัยที่ผ่านมา (2) สำรวจและวิเคราะห์แนวทางล่าสุดในการรับมือกับ Adversarial Attack (3) อภิปรายความเป็นไปได้และความท้าทายในการนำ Explainable AI มาใช้ในการรักษาความมั่นคงปลอดภัยทางไซเบอร์ และ (4) เสนอแนะทิศทางการวิจัยในอนาคตเพื่อแก้ไขข้อจำกัดที่ยังคงมีอยู่

ภัยคุกคามทางไซเบอร์ในปัจจุบัน

ในยุคที่เทคโนโลยีดิจิทัลเข้ามามีบทบาทในทุกมิติของชีวิตและธุรกิจภัยคุกคามทางไซเบอร์ก็ได้กลายเป็นความเสี่ยงที่ทุกคนจะต้องเผชิญ โดยเฉพาะในช่วงไม่กี่ปีที่ผ่านมาที่เราได้เห็นการโจมตีทางไซเบอร์ในหลากหลายรูปแบบทวีความถี่และความรุนแรงมากขึ้น

หนึ่งในภัยคุกคามที่สร้างความเสียหายอย่างมหาศาล คือ การโจมตีแบบ Ransomware ซึ่งเป็นมัลแวร์ที่เข้ารหัสไฟล์ของเหยื่อและเรียกค่าไถ่เพื่อแลกกับการปลดล็อก ซึ่งในปี 2021 พบว่า มีองค์กรกว่า 68% ถูกโจมตีด้วย Ransomware และในจำนวนนั้นมีผู้ที่สูญเสียข้อมูลถึง 82% (Roger, 2022) ตัวอย่างเช่น การโจมตีโรงพยาบาล Waikato ในนิวซีแลนด์ทำให้ระบบ IT ใช้การไม่ได้นานถึง 4 สัปดาห์ และต้องย้ายผู้ป่วยออกจากโรงพยาบาล (Cimpanu, 2021) แสดงให้เห็นผลกระทบอย่างกว้างขวางและรุนแรงของ Ransomware

อีกรูปแบบหนึ่งที่น่ากังวล คือ การโจมตีแบบ Distributed Denial of Service--DDoS ซึ่งเป็นการใช้ Traffic ปริมาณมหาศาลเพื่อทำให้ระบบหรือเว็บไซต์ไม่

สามารถให้บริการได้ โดยในปี 2022 พบว่า มีการโจมตี DDoS เพิ่มขึ้นถึง 150% เมื่อเทียบกับปีก่อน (Cox, 2023) ตัวอย่างเช่นในเดือนพฤศจิกายน 2022 เว็บไซต์สำคัญของธนาคารแห่งชาติรัสเซียถูกโจมตี DDoS จนไม่สามารถเข้าถึงได้เป็นเวลานาน (Burgess, 2022) ซึ่งนอกจากจะสร้างความเสียหายโดยตรงแล้วยังเป็นการบั่นทอนความเชื่อมั่นของลูกค้าอีกด้วย

ไม่เพียงแต่องค์กรใหญ่ที่จะตกเป็นเป้าหมาย ข้อมูลส่วนบุคคลของผู้ใช้งานทั่วไปก็ตกอยู่ในความเสี่ยงด้วยเช่นกัน จากเทคนิคต่าง ๆ เช่น การหลอกลวงให้กรอกข้อมูลส่วนตัวผ่านอีเมลปลอม (phishing) ไปจนถึงการบุกรุกระบบเพื่อขโมยฐานข้อมูล ดังกรณีการรั่วไหลของข้อมูลกว่า 530 ล้านบัญชีผู้ใช้ Facebook ในปี 2021 (Morse & Satter, 2021) หรือกรณี Marriott International ยักขีใหญ่ธุรกิจโรงแรมที่ข้อมูลของลูกค้ากว่า 500 ล้านคนถูกขโมยไป ส่งผลให้บริษัทต้องเสียค่าปรับกว่า 170 ล้านดอลลาร์และสูญเสียชื่อเสียงอย่างมาก (Whittaker, 2022)

นอกจากนี้ภัยคุกคามรูปแบบใหม่ที่อาศัยจุดอ่อนของมนุษย์ เช่น โซเชียลเอนจิเนียริง (Social Engineering) ก็เพิ่มมากขึ้นเช่นกัน โดยอาศัยกลยุทธ์หลอกลวงต่าง ๆ เช่น ปลอมตัวเป็นเพื่อนหรือคนรู้จักทางโซเชียลมีเดีย เพื่อให้เหยื่อหลงเชื่อบอกข้อมูลสำคัญหรือโอนเงินให้ซึ่งสร้างความเสียหายถึง 45 ล้านดอลลาร์ในปี 2021 เพิ่มขึ้น 90% จากปีก่อนหน้า (FBI, 2022)

เหตุการณ์เหล่านี้สะท้อนให้เห็นว่าภัยคุกคามทางไซเบอร์ได้ทวีความรุนแรงและแพร่กระจายอย่างรวดเร็ว ไม่เพียงสร้างความเสียหายทางการเงินโดยตรงเท่านั้น แต่ยังส่งผลกระทบต่อการดำเนินธุรกิจ การให้บริการ และความเชื่อมั่นของลูกค้า จนไม่สามารถมองข้ามความเสี่ยงนี้ได้อีกต่อไป ทุกคนจึงจำเป็นต้องตระหนักถึงความสำคัญและร่วมมือกันรับมือกับภัยคุกคามทางไซเบอร์อย่างจริงจัง

บทบาทของปัญญาประดิษฐ์ในการป้องกันภัยคุกคามทางไซเบอร์

ภัยคุกคามทางไซเบอร์ได้พัฒนาไปอย่างรวดเร็วและซับซ้อนมากขึ้น ทั้งการโจมตีแบบมัลแวร์ การโจมตีปฏิเสธการให้บริการ รวมถึงภัยคุกคามรูปแบบใหม่อย่าง

การเรียกค่าไถ่ ส่งผลให้วิธีการรักษาความปลอดภัยแบบดั้งเดิมไม่เพียงพอต่อการรับมืออีกต่อไป ด้วยเหตุนี้การประยุกต์ใช้ปัญญาประดิษฐ์ (AI) จึงเป็นทางเลือกที่น่าสนใจในการเสริมสร้างขีดความสามารถดังกล่าว

เทคนิคยอดนิยมของ AI ที่นำมาใช้ คือ Machine Learning ในการสร้างตัวแบบเพื่อตรวจจับภัยคุกคามจากข้อมูลในอดีต งานวิจัยของ Kwon et al. (2019) แสดงให้เห็นว่าวิธีนี้สามารถตรวจจับการบุกรุกเครือข่ายรูปแบบใหม่ได้อย่างแม่นยำ โดยมีประสิทธิภาพสูงกว่าเทคนิคแบบดั้งเดิมอย่างมีนัยสำคัญ

กรณีศึกษา: ไมโครซอฟท์ใช้ Random Forest ในการตรวจจับการบุกรุกบนระบบคลาวด์ของตน โดยสามารถลดเวลาในการตรวจจับและตอบสนองต่อภัยคุกคามลงได้ถึง 60% เมื่อเทียบกับวิธีการแบบดั้งเดิม (Han & Trimi, 2023)

นอกจากนี้ Deep Learning โดยเฉพาะ Convolutional Neural Network--CNN ยังมีบทบาทสำคัญในการวิเคราะห์มัลแวร์ เนื่องจากความสามารถในการสกัดคุณลักษณะสำคัญจากไฟล์ตัวอย่างมัลแวร์แล้วนำมาจำแนกได้อย่างแม่นยำ งานวิจัยล่าสุดของ Chen et al. (2023) ชี้ว่า CNN ให้ผลการจำแนกประเภทของมัลแวร์และการระบุฟังก์ชันอันตรายได้ถูกต้องมากกว่าวิธีแบบดั้งเดิม 8-10%

กรณีศึกษา: บริษัทแอนตี้ไวรัส Cylance ประสบความสำเร็จในการนำ AI ระบบนี้มาใช้ในการตรวจจับและ

ยับยั้งมัลแวร์ โดยสามารถป้องกันการโจมตีของมัลแวร์ที่ไม่เคยพบมาก่อนได้มากกว่า 99% (Bogatinov & Bogdanoski, 2022)

อีกตัวอย่างที่น่าสนใจ คือ การประยุกต์ใช้ Recurrent Neural Network--RNN โดยเฉพาะ Long Short-Term Memory--LSTM ในการตรวจจับบอทเน็ตและการโจมตีขั้นสูงแบบหลายขั้นตอน โดยงานวิจัยล่าสุดของ Wang et al. (2023) พบว่า LSTM มีความสามารถในการตรวจจับบอทเน็ตได้ถูกต้องสูงถึง 98% ซึ่งเร็วกว่าการสังเกตด้วยมนุษย์ถึง 5 เท่า

กรณีศึกษา: Google Cloud ใช้ LSTM ในการคาดการณ์และตรวจจับบอทเน็ตที่โจมตีบริการคลาวด์ ส่งผลให้สามารถลดความเสียหายจากการโจมตี DDoS ลงได้มากกว่า 80% ในปี 2023 (Aliar et al. 2024) เพื่อให้ให้เห็นภาพรวมของเทคนิค AI ต่าง ๆ ที่ใช้ในการป้องกันภัยคุกคามทางไซเบอร์

จากตารางข้างต้น จะเห็นได้ว่าเทคนิค AI แต่ละประเภทมีจุดแข็งในการรับมือกับภัยคุกคามที่แตกต่างกัน โดยการผสมผสานหลายเทคนิคเข้าด้วยกันอาจให้ผลลัพธ์ที่ดียิ่งขึ้น ทั้งนี้ การเลือกใช้เทคนิค AI ที่เหมาะสมขึ้นอยู่กับลักษณะของภัยคุกคามและบริบทการใช้งานของแต่ละองค์กร

ตาราง 1

เปรียบเทียบประสิทธิภาพของเทคนิค AI ในการตรวจจับภัยคุกคามไซเบอร์

เทคนิค AI	ประเภทภัยคุกคาม	ความถูกต้อง	อ้างอิงงานวิจัย
Random Forest	Network Anomaly Detection	สูงกว่าวิธีดั้งเดิมอย่างมีนัยสำคัญ	Kwon et al. (2022)
CNN	การวิเคราะห์และจำแนกมัลแวร์	สูงกว่าวิธีดั้งเดิม 8-10%	Chen et al. (2023)
LSTM	ตรวจจับบอทเน็ต	98% เร็วกว่ามนุษย์ 5 เท่า	Wang et al. (2023)
Decision Tree+SVM	ตรวจจับการโจมตี DDoS	สูงกว่าใช้เทคนิคเดียว	Li et al. (2024)

ตาราง 2

วิเคราะห์ช่องว่าง (gap analysis) สำหรับการใช้ AI ในการป้องกันภัยคุกคามทางไซเบอร์ในปัจจุบัน

ข้อดี	ข้อจำกัด	ความท้าทาย
- ตรวจจับภัยคุกคามได้รวดเร็ว แม่นยำ	- ต้องใช้ข้อมูลปริมาณมากและมีคุณภาพในการฝึกฝนตัวแบบ	- พัฒนาเทคนิคใหม่เพื่อเอาชนะข้อจำกัดของ AI
- สามารถเรียนรู้และปรับตัวกับรูปแบบใหม่ ๆ ได้	- เสี่ยงถูกโจมตีด้วย Adversarial Attack	- จัดทำกรอบระเบียบการนำ AI มาใช้งาน
- ลดภาระงานของบุคลากรด้านความปลอดภัย	- ขาดความโปร่งใส ยากต่อการอธิบายผลลัพธ์	- สร้างและพัฒนาศูนย์ผู้เชี่ยวชาญด้าน AI
		- งบประมาณการลงทุนระบบ AI ที่มีประสิทธิภาพ

จากตารางข้างต้น ในส่วนของข้อจำกัด ได้แก่

1. ปัญหาการขาดแคลนข้อมูลที่มีคุณภาพและปริมาณเพียงพอสำหรับการฝึกฝนตัวแบบ AI นั้น งานวิจัยของ Jo (2020) ได้เสนอแนวทางการประยุกต์ใช้เทคนิค Transfer Learning และ Semi-Supervised Learning เพื่อลดความจำเป็นในการใช้ข้อมูลจำนวนมากลง
2. ความเสี่ยงที่ AI อาจถูกโจมตีและบิดเบือนด้วยเทคนิค Adversarial Attack นั้น Chivukula et al. (2023) ได้ศึกษาและเสนอแนวทางการพัฒนาเทคนิคด้านทาน Adversarial Attack เช่น Adversarial Training Model Ensembling เป็นต้น
3. ประเด็นความโปร่งใสและความสามารถในการอธิบายผลลัพธ์ของ AI (explainable AI) ซึ่งเป็นสิ่งสำคัญเพื่อสร้างความน่าเชื่อถือ งานวิจัยของ Taylor (2024) ได้ให้แนวทางการวิจัยและพัฒนาระบบ Explainable AI เพื่อตอบโต้ภัยคุกคามนี้

ข้อจำกัดและความท้าทายของการใช้ปัญญาประดิษฐ์

ปัญญาประดิษฐ์โดยเฉพาะเทคนิคการเรียนรู้ของเครื่องและการเรียนรู้เชิงลึก มีศักยภาพอย่างมากในการยกระดับขีดความสามารถในการตรวจจับ วิเคราะห์ และรับมือกับภัยคุกคามไซเบอร์ที่ซับซ้อน การนำเทคโนโลยีนี้มาประยุกต์ใช้ก็ยังประสบกับข้อจำกัดและความท้าทายสำคัญบางประการที่ต้องหาวิธีการจัดการ

1. ข้อจำกัดด้านข้อมูล
หนึ่งในข้อจำกัดสำคัญที่สุด คือ การขาดแคลนข้อมูลที่มีคุณภาพสูงและปริมาณเพียงพอสำหรับการฝึกฝนตัวแบบ AI ให้มีประสิทธิภาพ ซึ่งในทางปฏิบัติแล้วการจัดหาข้อมูลดังกล่าวถือเป็นเรื่องท้าทายสำหรับหลายองค์กร โดยมีสาเหตุหลักจากข้อจำกัดด้านกฎระเบียบความเป็นส่วนตัว ต้นทุนในการจัดเก็บและบริหารจัดการข้อมูล รวมถึงกระบวนการเก็บข้อมูลที่ไม่เป็นระบบ (Bhuyan et al., 2015)
- การวิจัยในปัจจุบันยังขาดการศึกษาเชิงลึกเกี่ยวกับวิธีการจัดการกับปัญหานี้อย่างมีประสิทธิภาพ แม้จะมีแนวคิดเช่น Federated Learning ที่ช่วยลดปัญหาการแบ่งปันข้อมูล แต่ก็ยังมีข้อจำกัดในด้านประสิทธิภาพและความปลอดภัย (Liu et al., 2024)
2. ความเสี่ยงจาก Adversarial Attack
อีกประเด็นท้าทายสำคัญ คือ ความเสี่ยงที่ตัวแบบ AI อาจถูกโจมตีและบิดเบือนการทำงานด้วยเทคนิคพิเศษที่เรียกว่า Adversarial Attack ซึ่งจะทำให้ระบบ AI ตัดสินใจหรือวิเคราะห์ผิดพลาด ส่งผลกระทบต่อประสิทธิภาพและบั่นทอนความน่าเชื่อถือในการนำ AI มาใช้ป้องกันภัยคุกคามไซเบอร์
- งานวิจัยปัจจุบันยังไม่สามารถพัฒนาวิธีการป้องกัน Adversarial Attack ที่มีประสิทธิภาพสูงและครอบคลุมทุกรูปแบบการโจมตีได้ (Guo et al., 2020) นอกจากนี้ ยังขาดการศึกษาเชิงลึกเกี่ยวกับผลกระทบของ Adversarial Attack ในสภาพแวดล้อมการใช้งานจริง

3. ความโปร่งใสและการอธิบายได้

ประเด็นความโปร่งใสและความสามารถในการให้คำอธิบายผลลัพธ์ของระบบปัญญาประดิษฐ์ (explainable AI) ก็ยังเป็นสิ่งที่ต้องได้รับการแก้ไขต่อไป เพื่อสร้างความน่าเชื่อถือและการยอมรับจากผู้ใช้งาน โดยเฉพาะในบริบทของความมั่นคงปลอดภัยไซเบอร์ที่ต้องการความแม่นยำและความรับผิดชอบสูง

การวิจัยในปัจจุบันยังไม่สามารถพัฒนาเทคนิค Explainable AI ที่มีประสิทธิภาพสูงสำหรับโมเดล Deep Learning ที่ซับซ้อน โดยเฉพาะในการอธิบายการตัดสินใจในเวลาจริง (real-time) (Nanda et al., 2022)

ความท้าทายในการนำ AI มาใช้จริง

1. การบูรณาการกับระบบที่มีอยู่

การนำ AI มาใช้ในองค์กรที่มีระบบรักษาความปลอดภัยแบบดั้งเดิมอยู่แล้ว อาจเกิดปัญหาในการบูรณาการระบบทั้งในแง่ของความเข้ากันได้ของเทคโนโลยี และการปรับเปลี่ยนกระบวนการทำงานของบุคลากร (Kumar et al., 2022)

2. การจัดการกับ False Positives และ False Negatives

ระบบ AI อาจเกิดการแจ้งเตือนที่ผิดพลาด (false positives) หรือพลาดการตรวจจับภัยคุกคามจริง (false negatives) ซึ่งอาจส่งผลกระทบร้ายแรงในบริบทของความมั่นคงปลอดภัย การปรับแต่งระบบให้สมดุลระหว่างความไวในการตรวจจับและความแม่นยำเป็นเรื่องที่ท้าทายอย่างมาก (Zeadally et al., 2020)

3. การรักษาความเป็นส่วนตัวและการปฏิบัติตามกฎหมาย

การใช้ AI ในการวิเคราะห์ข้อมูลจำนวนมากอาจนำไปสู่ประเด็นด้านความเป็นส่วนตัวและการละเมิดกฎหมายคุ้มครองข้อมูลส่วนบุคคล องค์กรต้องพัฒนาแนวทางการใช้ AI ที่สอดคล้องกับกฎหมายและจริยธรรม (Töndel & Cruzes, 2022)

4. การพัฒนาและรักษาบุคลากรที่มีความเชี่ยวชาญ

การขาดแคลนบุคลากรที่มีความเชี่ยวชาญทั้งด้าน AI และความมั่นคงปลอดภัยไซเบอร์เป็นอุปสรรคสำคัญ

การพัฒนาและรักษาบุคลากรเหล่านี้เป็นความท้าทายสำหรับหลายองค์กร (Clark, 2023)

5. การปรับตัวต่อภัยคุกคามที่เปลี่ยนแปลงอย่างรวดเร็ว

ภัยคุกคามทางไซเบอร์มีการพัฒนาและเปลี่ยนแปลงอย่างรวดเร็ว การพัฒนาระบบ AI ให้สามารถปรับตัวและเรียนรู้ภัยคุกคามรูปแบบใหม่ได้อย่างทันท่วงที่เป็นความท้าทายที่สำคัญ (Wang et al., 2023)

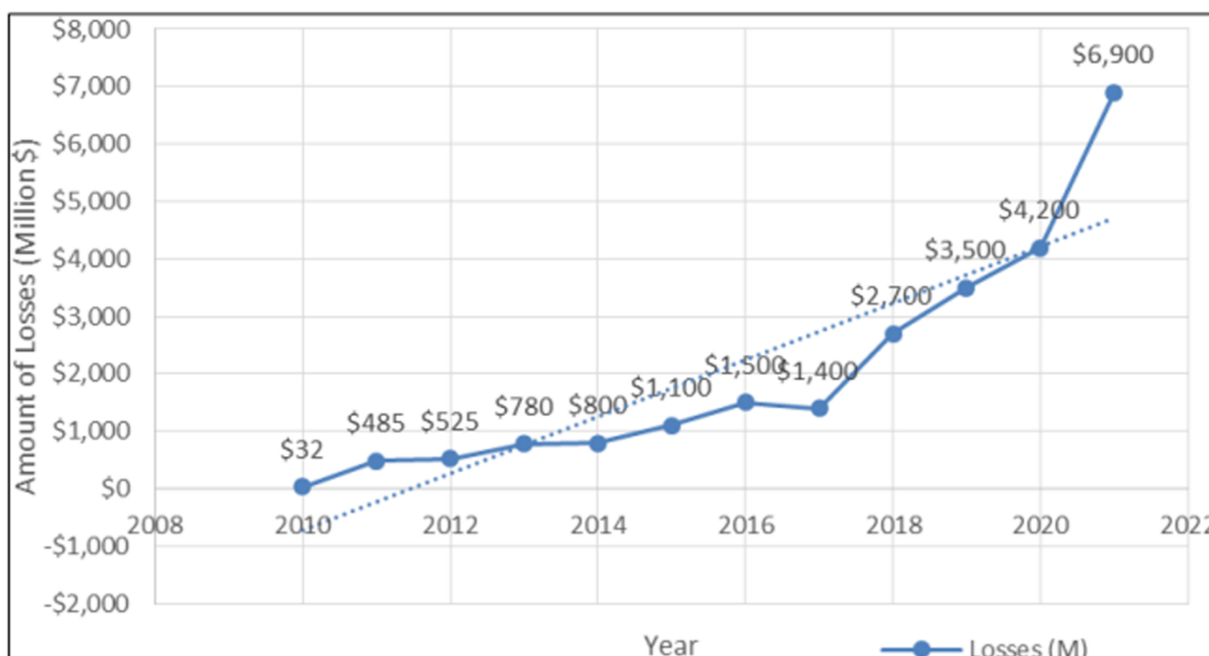
โดยสรุป แม้ปัญญาประดิษฐ์จะเป็นเครื่องมือสำคัญในการเสริมสร้างขีดความสามารถในการรับมือกับภัยคุกคามไซเบอร์ แต่การประยุกต์ใช้ AI ก็ยังคงมีข้อจำกัดและความท้าทายอยู่หลายประการ นับตั้งแต่การขาดแคลนข้อมูลคุณภาพ ความเสี่ยงจากการโจมตีตัวแบบ AI เอง ไปจนถึงประเด็นความน่าเชื่อถือและความโปร่งใสของระบบ ด้วยเหตุนี้จึงจำเป็นต้องมีการศึกษาวิจัยอย่างจริงจังและต่อเนื่องเพื่อหาแนวทางในการจัดการกับข้อจำกัดและความท้าทายเหล่านี้ ผ่านการพัฒนาเทคนิคใหม่ ๆ การแบ่งปันองค์ความรู้และการสร้างความร่วมมือระหว่างองค์กรต่าง ๆ เพื่อให้มั่นใจว่าสามารถนำศักยภาพของเทคโนโลยีปัญญาประดิษฐ์มาใช้ในการรักษาความมั่นคงปลอดภัยไซเบอร์ได้อย่างมีประสิทธิภาพ ปลอดภัย และยั่งยืนสูงสุด

บทสรุปและข้อเสนอแนะ

การทบทวนวรรณกรรมและงานวิจัยที่เกี่ยวข้องอย่างเป็นระบบได้แสดงให้เห็นถึงบทบาทสำคัญของเทคโนโลยีปัญญาประดิษฐ์ (AI) โดยเฉพาะอย่างยิ่งเทคนิคการเรียนรู้ของเครื่องและการเรียนรู้เชิงลึกในการเสริมสร้างขีดความสามารถด้านการตรวจจับ วิเคราะห์ และรับมือกับภัยคุกคามทางไซเบอร์ที่มีความซับซ้อนและรุนแรงเพิ่มขึ้น ผลการวิจัยจำนวนมากได้ยืนยันประสิทธิภาพของเทคนิค AI ที่เหนือกว่าวิธีการแบบดั้งเดิมอย่างมีนัยสำคัญ ทั้งในแง่ความแม่นยำและความรวดเร็วในการดำเนินการ

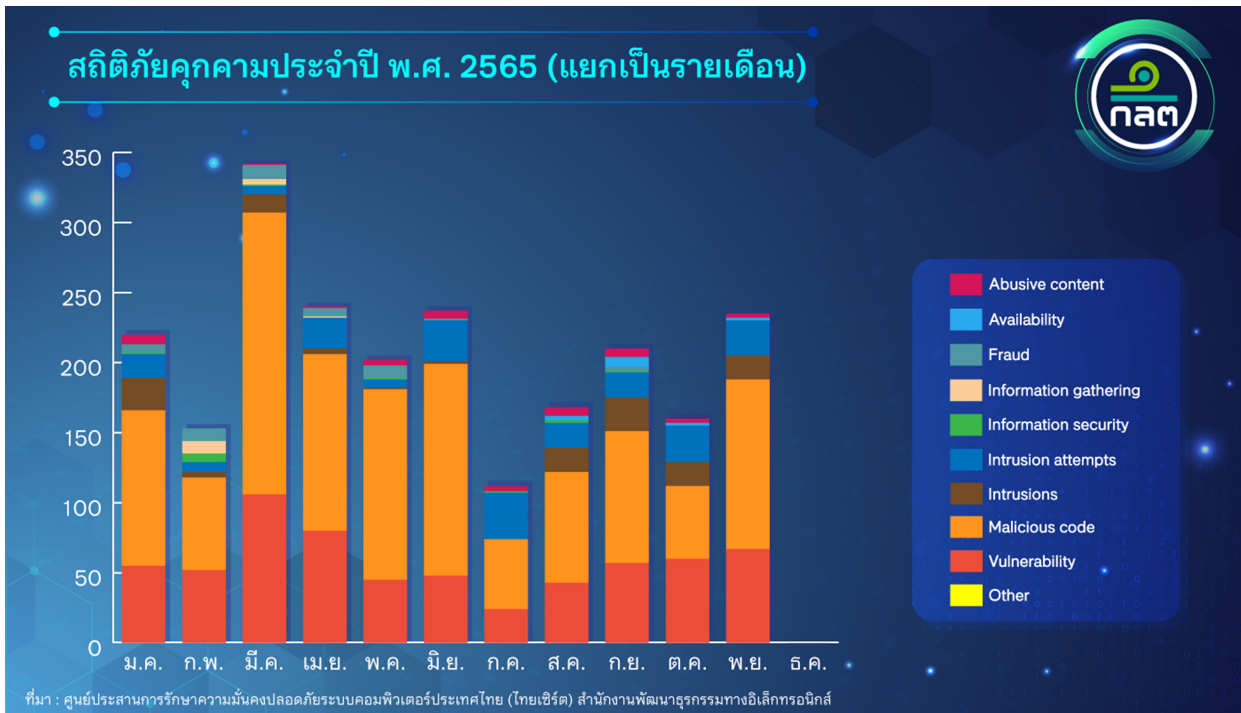
อย่างไรก็ตาม การประยุกต์ใช้ AI ในบริบทของความมั่นคงปลอดภัยทางไซเบอร์ยังคงประสบกับข้อจำกัดและความท้าทายที่สำคัญหลายประการ อาทิ ข้อจำกัดด้านคุณภาพและปริมาณของข้อมูลสำหรับการฝึกฝนโมเดล ความเสี่ยงจากการโจมตีแบบ Adversarial ต่อระบบ AI และประเด็นด้านความโปร่งใสและความสามารถในการอธิบายผลลัพธ์ของระบบ AI (Explainable AI) ประเด็นเหล่านี้จำเป็นต้องได้รับการศึกษาวิจัยเพิ่มเติมอย่างลึกซึ้งและต่อเนื่อง เพื่อพัฒนาแนวทางการแก้ไขที่มีประสิทธิภาพและยั่งยืน

การวิเคราะห์แนวโน้มการวิจัยในสาขานี้พบว่า มีการมุ่งเน้นไปที่การบูรณาการเทคนิค AI หลากหลายรูปแบบ การพัฒนา Explainable AI การประยุกต์ใช้ Federated Learning เพื่อแก้ปัญหาความเป็นส่วนตัวของข้อมูล และการพัฒนาเทคนิคป้องกัน Adversarial Attack ซึ่งล้วนเป็นประเด็นสำคัญที่ควรได้รับการศึกษาวิจัยอย่างต่อเนื่องในอนาคต



ภาพ 1 จำนวนการสูญเสียจากการโจมตีทางไซเบอร์ในช่วง 12 ปีที่ผ่านมา (พ.ศ. 2553-2564) สหรัฐอเมริกา

Note. From A literature review of financial losses statistics for cyber security and future trend by M. H. U. Sharif and M. A. Mohammed, 2022, *World Journal of Advanced Research and Reviews*, 15(01), pp. 138–156. Copyright World Journal of Advanced Research and Reviews



ภาพ 2 ภาพสถิติภัยคุกคามทางไซเบอร์ปี 2565

Note. From *Cyber Attack Trends 2022* by The Securities and Exchange Commission, Thailand, 2022, retrieved from <https://www.sec.or.th/TH/Pages/CYBERRESILIENCE-STATISTICS-2565.aspx>

ข้อเสนอแนะ

1. ด้านการวิจัยและพัฒนา

1.1 พัฒนา AI ที่มีความยืดหยุ่นและปรับตัวได้สูง สามารถรับมือกับภัยคุกคามรูปแบบใหม่ได้อย่างมีประสิทธิภาพ

1.2 วิจัยและพัฒนาเทคนิค Privacy-Preserving AI เพื่อรักษาความเป็นส่วนตัวของข้อมูลในกระบวนการเรียนรู้ของ AI

1.3 พัฒนาระบบ AI สำหรับการตอบสนองอัตโนมัติต่อภัยคุกคามทางไซเบอร์

1.4 ศึกษาผลกระทบระยะยาวของการใช้ AI ในด้านความมั่นคงปลอดภัยทางไซเบอร์

1.5 วิจัยและพัฒนาแนวทางการผสมผสาน Human-AI Collaboration เพื่อเพิ่มประสิทธิภาพในการรับมือกับภัยคุกคาม

2. ด้านนโยบายและการสนับสนุน

2.1 จัดสรรงบประมาณและทรัพยากรเพื่อสนับสนุนการวิจัยและพัฒนานวัตกรรมด้าน AI for Cybersecurity อย่างต่อเนื่อง

2.2 ส่งเสริมการแบ่งปันข้อมูลและความร่วมมือระหว่างภาครัฐ ภาคเอกชน และสถาบันการศึกษาในการพัฒนาและประยุกต์ใช้ AI ด้านความมั่นคงปลอดภัยทางไซเบอร์

2.3 ศึกษาและผลักดันให้เกิดกรอบกฎหมายและมาตรฐานสากลที่เหมาะสมเกี่ยวกับการนำ AI มาใช้ในงานด้านความมั่นคงปลอดภัยไซเบอร์

3. ด้านการพัฒนามาตรฐานและแนวทางปฏิบัติ

3.1 พัฒนามาตรฐานและแนวทางปฏิบัติในการใช้ AI สำหรับความมั่นคงปลอดภัยทางไซเบอร์ โดยคำนึงถึงประเด็นด้านจริยธรรม กฎหมาย และความรับผิดชอบ

3.2 สร้างแนวทางการประเมินประสิทธิภาพและความน่าเชื่อถือของระบบ AI ที่ใช้ในด้านความมั่นคงปลอดภัยทางไซเบอร์

3.3 พัฒนาแนวทางการรับมือกับความเสี่ยงและข้อจำกัดของ AI ในบริบทของความมั่นคงปลอดภัยทางไซเบอร์

- | | |
|---|--|
| <p>4. ด้านการพัฒนาบุคลากรและการศึกษา</p> <p>4.1 ส่งเสริมการพัฒนาหลักสูตรและโปรแกรมการศึกษาที่บูรณาการความรู้ด้าน AI และความมั่นคงปลอดภัยทางไซเบอร์</p> <p>4.2 สนับสนุนการฝึกอบรมและพัฒนาทักษะของบุคลากรด้านความมั่นคงปลอดภัยทางไซเบอร์ให้</p> | <p>สามารถใช้งานและพัฒนาระบบ AI ได้อย่างมีประสิทธิภาพ</p> <p>4.3 ส่งเสริมการแลกเปลี่ยนความรู้และประสบการณ์ระหว่างผู้เชี่ยวชาญด้าน AI และด้านความมั่นคงปลอดภัยทางไซเบอร์</p> |
|---|--|



References

- Accenture. (2019). *Artificial intelligence: Is your organizational ‘cyber resilience’ up to the challenge?* Retrieved from <https://www.accenture.com/gb-en/insights/artificial-intelligence/cybersecurity-artificial-intelligence>
- Aliar, A. A., Gowri, V., & Zbins, A. A. (2024). Detection of distributed denial of service attack using enhanced adaptive deep dilated ensemble with hybrid meta-heuristic approach. *Transactions on Emerging Telecommunications Technologies*, 35(1), e4921. <https://doi.org/10.1002/ett.4921>
- Bhuyan, M. H., Bhattacharyya, D. K., & Kalita, J. K. (2015). NADO: A hybrid kernel based approach for network anomaly detection. *IEEE Communications Surveys & Tutorials*, 17(2), 706-722. <https://doi.org/10.1109/COMST.2015.2417967>
- Bogatinov, D., & Bogdanoski, M. (2022). Using artificial intelligence as a first line of defense in cyberspace. *NATO Science for Peace and Security Series - E: Human and Societal Dynamics*, 155, 56-68. doi: 10.3233/NHSDP220006
- Burgess, M. (2022). *Russia’s Central Bank digital services disrupted by DDoS attacks*, Wired UK. Retrieved from <https://www.wired.co.uk/article/russia-central-bank-ddos-attack>
- Chen, Y. H., Lin, S. C., Huang, S. C., Lei, C. L., & Huang, C. Y. (2023). Guided malware sample analysis based on graph neural networks. *IEEE Transactions on Information Forensics and Security*, 18, 4128-4143. doi: 10.1109/TIFS.2023.3283913.
- Chivukula, S. A., Yang, X., Liu, B., Liu, W., & Zhou, W. (2023). *Adversarial attack surfaces*. In *Adversarial machine learning* (pp 47–72). Berlin, Germany: Springer. https://doi.org/10.1007/978-3-030-99772-4_3
- Cimpanu, C. (2021). *Waikato DHB cyber-attack: Hackers had access to hospitals’ IT systems for around five months*. *The record by recorded future*. Retrieved from <https://therecord.media/waikato-dhb-cyber-attack-hackers-had-access-to-hospitals-it-systems-for-around-five-months/>
- Clarke, R. (2023). The re-conception of AI: Beyond artificial, and beyond intelligence. *IEEE Transactions on Technology and Society*, 4(1), 24-33. doi: 10.1109/TTS.2023.3234051.

- Cox, J. (2023). *DDoS attacks jumped 150% in the past year, network provider says*. CNET. Retrieved from <https://www.cnet.com/tech/services-and-software/ddos-attacks-jumped-150-in-past-year-network-provider-says/>
- FBI. (2022). *2021 internet crime report*. Retrieved from https://www.ic3.gov/Media/PDF/AnnualReport/2021_IC3Report.pdf
- Guo, Q., Chen, Y., Chen, J., Shen, H., & Wu, J. (2020). A survey on adversarial attacks and defense strategies in machine learning-based cyber security applications. *IEEE Access*, 8, 205806-205831.
- Han, H., & Trimi, S. (2024). Analysis of cloud computing-based education platforms using unsupervised random forest. *Educ Inf Technol*, 29, 15905–15932. <https://doi.org/10.1007/s10639-024-12457-w>
- Hassan, N. A. (2019). *Enterprise defense strategies against ransomware attacks*. In *Ransomware Revealed* (pp. 115-154). CA Berkeley: Apress https://doi.org/10.1007/978-1-4842-4255-1_5
- Jo, T. (2020). *Semi-supervised Learning*. In *Machine learning foundations* (pp 309–334). Berlin: Springer https://doi.org/10.1007/978-3-030-65900-4_14.
- Kizza, J. M. (2020). *Guide to computer network security* (5th ed.). Berlin: Springer.
- Kwon, H., & Sea, J. (2022). Characteristics of sexual homicide: Based on random forest analysis. *Journal Article Korean Criminological Review*, 33(1), 165-192. doi: 10.36889/KCR.2022.3.31.1.165
- Kumar, A., Han, S. T., & Soni, A. K. (2022). Survey on artificial intelligence for cybersecurity. *IEEE Access*, 10, 16679-16709. <https://doi.org/10.1109/ACCESS.2022.3146309>
- Li, C., Liu, Q., Guo, Q., & Wu, Y. (2022). A federated graph neural network approach for privacy-preserving network intrusion detection. *IEEE Transactions on Network Science and Engineering*, 9(1), 247-260. doi:10.21203/rs.3.rs-1191595/v1
- Li, P., Xiong, F., Huang, X., & Wen, X. (2024). Construction and optimization of vending machine decision support system based on improved C4.5 decision tree. *Heliyon*, 10(3), e25024. <https://doi.org/10.1016/j.heliyon.2024.e25024>
- Liu, Y., Kang, Y., Zou, T., Pu, Y., He, Y., Ye, X., Zhang, Y.Q., & Yang, Q. (2024). Vertical federated learning: Concepts, advances, and challenges. *IEEE Transactions on Knowledge and Data Engineering*, 36(7), 3615-3634. doi: 10.1109/TKDE.2024.3352628.
- Marino, D. L., Wickramasinghe, C. S., & Riehle, L. (2021). Metamorphic testing for cybersecurity: Securing machine learning cyber-physical systems through metamorphic testing. *IEEE Transactions on Reliability*, 70(1), 264-280.
- Morse, A., & Satter, R. (2021). *Data on 533 million Facebook users leaked online*. Reuters. Retrieved from <https://www.reuters.com/technology/hackers-leak-data-533-mln-facebook-users-2021-04-03/>

- Mothukuri, V., Parizi, R. M., Pouriyeh, S., Huang, Y., Dehghantanha, A., & Srivastava, G. (2021). A survey on security and privacy of federated learning. *Future Generation Computer Systems*, 115, 619-640. <https://doi.org/10.1016/j.future.2020.10.007>
- Nanda, S., Zafari, F., DeLong, C., Bustinza, R., & Raina, R. (2022). XAI for data-driven cyber security: Opportunities, challenges & future directions. *arXiv preprint arXiv, 2204*, 11234.
- Roger, A. (2022). *Grimes, future of ransomware*. In *Ransomware protection playbook* (pp. 261-272). New York: Wiley
- Sarker, I. H., Kayes, A. S. M., Badsha, S., Alqahtani, H., Watters, P., & Ng, A. (2020). Cybersecurity data science: An overview from machine learning perspective. *Journal of Big Data*, 7, 41. <https://doi.org/10.1186/s40537-020-00318-5>
- Sharif, M. H. U., & Mohammed, M. A. (2022). A literature review of financial losses statistics for cyber security and future trend. *World Journal of Advanced Research and Reviews*, 15(1), 138–156. <https://doi.org/10.30574/wjarr.2022.15.1.0573>
- Srinivasan, S., & Sharmili, A. S. (2022). *Graph neural network-based intrusion detection systems for cyber security applications*. In *AI and Machine Learning for Cyber Security* (pp. 79-104). Berlin: Springer.
- Taylor, I. (2024). *Is explainable AI responsible AI?* In *AI & Soc* (pp. 1-10). Berlin: Springer. <https://doi.org/10.1007/s00146-024-01939-7>
- The Securities and Exchange Commission Thailand. (2022). *Cyber attack trends 2022*. Retrieved from <https://www.sec.or.th/TH/Pages/CYBERRESILIENCE-STATISTICS-2565.aspx> (in Thai)
- Töndel, I. A., & Cruzes, D. S. (2022). Continuous software security through security prioritisation meetings. *Journal of Systems and Software*, 194, 111477. <https://doi.org/10.1016/j.jss.2022.111477>
- Wang, S., He, R., Shan, C., Choo, K. K. R., Yang, Y., & Chen, W. (2023). Defending against cybersecurity attacks: A comprehensive survey. *ACM Computing Surveys*, 55(6), 125. <https://doi.org/10.1145/3534389>
- Wang, Y. C., Houg, Y. C., Chen, H. X., & Tseng, S. M. (2023). Network anomaly intrusion detection based on deep learning approach. *Sensors*, 23(4), 2171. <https://doi.org/10.3390/s23042171>
- Warnecke, A., Arp, D., Wressnegger, C., & Rieck, K. (2020). Evaluating explanation methods for deep learning in security. *IEEE European Symposium on Security and Privacy (EuroS&P)* (pp. 158-174). Genoa, Italy: IEEE. doi: 10.1109/EuroSP48549.2020.00018.
- Whittaker, Z. (2022). *Marriott to pay £18.4 m fine after massively bungling data breach disclosures*, *TechCrunch*. Retrieved from <https://techcrunch.com/2022/10/11/marriott-uk-data-breach-fine/>
- Wollerton, M. (2023). *Ransomware attacks*. In *CQ Researcher*. Thousand Oaks, CA: CQ Press. <https://doi.org/10.4135/cqresrre20230818>

- Zeadally, S., Adi, E., Baig, Z., & Khan, I. A. (2020). Harnessing artificial intelligence capabilities to improve cybersecurity practices. *Information*, 11(10), 490. <https://doi.org/10.3390/info11100490>
- Zhang, H., Chen, L., Liu, X., & Wang, X. (2021). Meta-learning based adversarial detection framework for few-shot learning in cybersecurity. *IEEE Transactions on Information Forensics and Security*, 16, 4306-4320.
- Zhong, Y., Yu, W., Yuantao, S., Liu, J., & Qu, Y. (2020). Hierarchical graph neural network for network intrusion detection. *IEEE Transactions on Information Forensics and Security*, 15, 2653-2662.

