

HUMAN HACKER VS AI HACKER

ความคิดต่างกัน แต่มีเป้าหมายเดียวกัน
องค์กรของคุณพร้อมรับมือหรือยัง?

Jedsada Thongkanluang

Managing Director

T-NET IT Solution Co., Ltd.



VAPT

Vulnerability Assessment
& Penetration Testing



SOC & SIEM

Security Monitoring
& Threat Detection



INCIDENT RESPONSE

Detect - Respond
Recover



SECURITY CONSULTING

Risk Assessment
& Security Strategy

About us



T-NET IT Solution Company Limited was established in Thailand and registered in Thailand on October 3, 2019 by a team of researchers and practitioners specializing in computer security in Thailand (Thai Computer Emergency Response Team, ThaiCERT).

“We are proud to provide All-encompassing cybersecurity services which complies with information security management standards”



Location



T-NET IT SOLUTION Co., Ltd.

131 Moo 9, Thailand Science Park,
Phaholyothin Road, Khlong Nueng Subdistrict,
Khlong Luang District, Pathum Thani Province
Zip code 12120 Thailand

Our Business



**“Provide All-encompassing Cybersecurity”
services**

➤ **IT Security Consultant** ◀

➤ **Designed IT Architecture. Development software** ◀

➤ **Implementation Cybersecurity Standard and Cybersecurity Act.** ◀

Company Services

Consultants

All of IT cycle.

Act., Practice, ISMS, SMS, BCMS



Training



Hacking [VA, Pentest] Network, SOC
ISMS, PDPA.... Cyber security, CompTIA

Vulnerability assessment & Penetration test

OS, Network, Website, Application
White box, Gray box, Black box



SOC and Log

Security Operation Center

Log compile Act.

SIEM



Software & Hardware

Develop software

IT Device Shop



Secure design

Network & Infrastructure

Data Center, System Integrator

Cloud solution



Audit



Act., Practice, ISMS, SMS, BCMS

Security baseline

OS [Windows, Linux]

Network, Application, Virtual Machine



Code review

Statics Application Security Testing

Dynamic application security testing



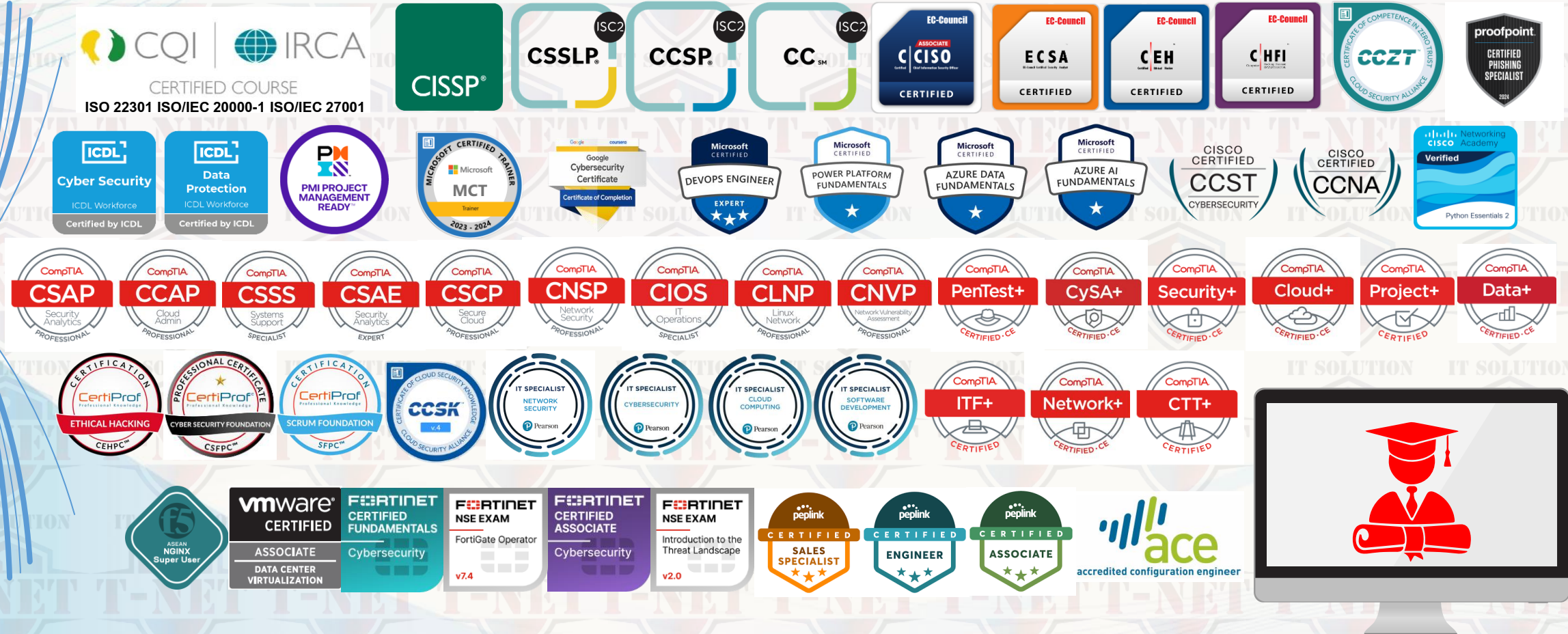
Digital forensics

Chain of Custody

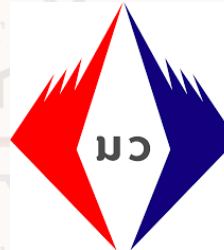
[Acquisition, Analysis, Reporting]



Our Certification



Company portfolio



MAHASARAKHAM
UNIVERSITY





นายเจษฎา ทองก้านเหลือง (ต่อม)

กรรมการผู้จัดการ (Managing Director)

| บริษัท ที-เน็ต ไอที โซลูชัน จำกัด



“ Cybersecurity ไม่ใช่แค่การป้องกัน แต่คือการเตรียมพร้อม รับมือ และพัฒนาอย่างต่อเนื่อง เพื่อให้องค์กรสามารถดำเนินธุรกิจได้อย่างมั่นคงและปลอดภัย ”

ประวัติย่อ

ผู้เชี่ยวชาญด้าน Cybersecurity, Digital Forensics, Network Security, Cloud Security และ Data Protection ที่มีประสบการณ์ในการให้คำปรึกษา ออกแบบระบบ และถ่ายทอดองค์ความรู้ด้านความมั่นคงปลอดภัยไซเบอร์ แก่หน่วยงานภาครัฐและเอกชน รวมถึงองค์กรด้านอุตสาหกรรม พลังงาน และโครงสร้างพื้นฐานสำคัญของประเทศ (Critical Infrastructure)

ปัจจุบันดำรงตำแหน่ง กรรมการผู้จัดการ บริษัท ที-เน็ต ไอที โซลูชัน จำกัด ผู้ให้บริการด้าน Cybersecurity, Network Infrastructure, OT Security, SOC Services, Digital Forensics และการพัฒนาบุคลากรด้าน ความมั่นคงปลอดภัยไซเบอร์

นอภาที่ยังเป็นผู้ดูแลชุมชนและกลุ่มความรู้ด้านการทดสอบเจาะระบบ (Ethical Hacking) และ Cybersecurity ภายใต้กลุ่ม “สอนแลกเปลี่ยน” ที่มุ่งเน้นการแบ่งปันองค์ความรู้และพัฒนาศักยภาพบุคลากรด้านความ ปลอดภัยไซเบอร์ในประเทศไทย

ผลงานด้านการอบรมและวิทยากร

- ✓ Cybersecurity Awareness
- ✓ Ethical Hacking Workshop
- ✓ Penetration Testing
- ✓ Incident Response
- ✓ Digital Forensics
- ✓ SOC Analyst Training
- ✓ OT/ICS Cybersecurity
- ✓ PDPA Compliance
- ✓ ISO/IEC 27001
- ✓ Cloud Security
- ✓ Zero Trust Architecture
- ✓ Network Security
- ✓ Security Monitoring & Threat Hunting

รูปแบบการอบรม



In-house Training



Public Training



Workshop & Hands-on Lab



Tabletop Exercise (TTX)



Cyber Drill Exercise



คุณวุฒิและใบรับรองวิชาชีพ

EC-Council Certifications

- Certified Ethical Hacker (CEH)
- Certified Hacking Forensic Investigator (CHFI)
- EC-Council Certified Security Analyst (ECSA)
- Certified Chief Information Security Officer (CISO)
- Certified Network Defender (CND)
- EC-Council Continuing Education (ECE)
- Certified Cybersecurity Educator Professional (CCEP)
- Certified Red Team Operations Management (CRTOM)

CompTIA Certifications

CompTIA

- CompTIA Security+
- CompTIA CySA+
- CompTIA Project+
- CompTIA Cloud+
- CompTIA Network+



Cloud & Zero Trust Certifications

- Certificate of Cloud Security Knowledge v5 (CCSK)
- Certificate of Competence in Zero Trust v1.0 (CCZT)



IT Specialist Certifications

- IT Specialist (ITS): Cybersecurity
- IT Specialist (ITS): Network Security

Vendor Certifications



- Peplink Certified Engineer (PCE)
- Peplink Sales Specialist (PSS)
- Fortinet NSE 1
- Fortinet NSE 2

Professional Certifications

CertiProf

- CertiProf: Cyber Security Foundation
- CertiProf: Remote Work and Virtual Collaboration
- CertiProf: Scrum Foundation Professional Certificate
- SKILLFRONT: ISO/IEC 27001 INFORMATION SECURITY ASSOCIATE

Professional Qualification Institute (TPQI)



- นักจัดการคุ้มครองข้อมูลส่วนบุคคล ระดับ 5 (สถาบันคุณวุฒิวิชาชีพ)
- นักบริหารจัดการระบบเครือข่ายคอมพิวเตอร์ระดับ 4 (สถาบันคุณวุฒิวิชาชีพ)



ช่องทางการติดต่อ



E-Mail : Jedsada@tnetitsolution.co.th



Facebook : ผู้ดูแลกลุ่มสอนแลกเปลี่ยนแบบหมู่ๆ



Website : www.tnetitsolution.co.th



ความเชี่ยวชาญ



Cybersecurity Management



Ethical Hacking & Penetration Testing



Incident Response & Digital Forensics



Security Operations Center (SOC)



Network Security & Network Defense



Cloud Security



OT Security / Industrial Cybersecurity



Data Protection & PDPA Compliance



Zero Trust Architecture



Cybersecurity Awareness & Professional Training



ETHICAL HACKING



NETWORK SECURITY



CLOUD SECURITY



DIGITAL FORENSICS



INCIDENT RESPONSE



DATA PROTECTION



ZERO TRUST



SOC SERVICES

ยุคของผู้โจมตี AI เริ่มต้นแล้ว

AI เร่งนวัตกรรมองค์กร — และเพิ่มศักยภาพผู้โจมตีอย่างก้าวกระโดด



องค์กรใช้ AI ขับเคลื่อนทุกมิติ



- ✓ AI ฝังอยู่ในระบบขององค์กร
- ✓ AI ช่วยให้งานได้รวดเร็ว อัตโนมัติ และแม่นยำ



เขียนโค้ด



วิเคราะห์ข้อมูล



จัดการเวิร์กโฟลว์



ตัดสินใจอัตโนมัติ

VS



ผู้โจมตีใช้ AI เพิ่มความเร็วและความรุนแรง

ปี 2025 ผู้โจมตีที่ใช้ AI
เพิ่มจำนวนการโจมตี

89%

เมื่อเทียบกับปีก่อนหน้า



สร้าง Phishing และ
ล่าข้อมูลอัตโนมัติ



ลดเวลาจาก “เจาะเข้า”
สู่ “ลงมือทำ”



ยกระดับผู้โจมตีมือใหม่
ให้เก่งขึ้น



เพิ่มศักยภาพผู้โจมตีขั้นสูง



AI สร้างมิตความเสี่ยงใหม่: เป้าหมายคือระบบ AI เอง



- AI ฝังในระบบองค์กร
เช่น Development Pipeline,
SaaS Platform และ
Workflow
- ระบบ AI กลายเป็น
ส่วนหนึ่งของพื้นผิว
การโจมตี



ผู้โจมตีใช้
Prompt อันตราย



สร้างคำสั่งที่
ไม่ได้รับอนุญาต



ความปลอดภัยต้องทำให้ทัน
ความเร็วของนวัตกรรม

- ป้องกันตลอดวงจรชีวิตของ AI
- ตรวจสอบ ตอบสนอง และเรียนรู้
ได้อย่างรวดเร็ว



ความเร็วคือความจริงใหม่ของการโจมตี

- ✓ “ความเร็ว” คือคุณลักษณะสำคัญของการบุกรุกยุคใหม่
- ✓ ผู้โจมตีหลบเลี่ยงการตรวจจับได้เร็วขึ้น
- ✓ องค์กรต้องตรวจจับ ป้องกัน และตอบสนองให้เร็วกว่าเดิม



ความมั่นคงปลอดภัยคือโครงสร้างพื้นฐานของยุค Agentic

ในยุค Agentic การโจมตีอัตโนมัติด้วย AI
คือโครงสร้างพื้นฐานสำคัญที่ต้องปกป้อง
AI และธุรกิจของคุณ

“ AI เร่งอนาคต
ความปลอดภัยคือสิ่งที่ปกป้องมัน
**Secure AI.
Secure Tomorrow.** ”



สรุปประเด็นสำคัญ

AI มอบประโยชน์มหาศาล — แต่ที่มาพร้อมความเสี่ยงใหม่ องค์กรต้องใช้ AI อย่างมั่นใจ
พร้อมป้องกันด้วยความเร็วและความชาญฉลาดที่เท่าทันผู้โจมตี



AI กลายเป็นหัวใจ
ขององค์กรยุคใหม่



ผู้โจมตีใช้ AI
เพิ่มความเร็ว



ความปลอดภัยต้อง
เร็วกว่าเสมอ



ปกป้องระบบ AI
= ปกป้องธุรกิจ



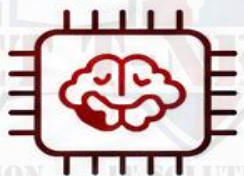
ลงทุนในความปลอดภัย
คือการลงทุนในอนาคต

สรุปสถานการณ์ภัยคุกคามไซเบอร์ ปี 2025-2026

ภัยคุกคามรุนแรงขึ้น รวดเร็วขึ้น และซับซ้อนกว่าเดิม

1

การโจมตีด้วย AI



89%

การโจมตีด้วย AI เพิ่มขึ้น
เมื่อเทียบกับปีก่อน

2

เวลายืดระบบลดลง



29 นาที

เวลาเฉลี่ยยืดระบบลดเหลือ
เร็วที่สุดเพียง 27 วินาที

3

การโจมตีไม่ใช้มัลแวร์



82%

ของการโจมตีไม่ใช้มัลแวร์
ใช้บัญชีผู้ใช้หรือเครื่องมือที่ถูกต้อง

4

กลุ่มผู้โจมตีใหม่เพิ่มขึ้น



24 กลุ่ม

พบกลุ่มผู้โจมตีใหม่
รวมติดตามแล้ว 281 กลุ่ม

5

กิจกรรมจากจีนเพิ่มขึ้น



38%

กิจกรรมจากกลุ่มที่เชื่อมโยงกับจีน
ทุกภาคส่วนเพิ่มขึ้น โดยโลจิสติกส์พุ่ง
85%

6

การใช้ช่องโหว่ Zero-day



42%

การใช้ช่องโหว่ Zero-day
ก่อนเปิดเผยสู่สาธารณะ เพิ่มขึ้น

7

การใช้บัญชีผู้ใช้จริง



35%

ของเหตุการณ์บน Cloud
เกิดจากการใช้บัญชีผู้ใช้จริงที่ถูกขโมย

8

การโจมตี Cloud เพิ่มขึ้น



37%

การโจมตี Cloud เพิ่มขึ้น
จากกลุ่มผู้โจมตีระดับรัฐเพิ่มขึ้น
266%



ประเด็นสำคัญ

ภัยคุกคามไซเบอร์ในปัจจุบัน รวดเร็ว ซับซ้อน และมุ่งเป้าไปที่การใช้เครื่องมือที่ถูกต้องและบัญชีผู้ใช้จริงมากขึ้น
องค์กรต้องยกระดับการตรวจจับภัยคุกคามเชิงรุก (Proactive Detection) และตอบสนองเหตุการณ์
ให้รวดเร็วกว่าเดิม เพื่อป้องกันความเสียหายอย่างมีประสิทธิภาพ

ภาพรวมการโจมตีแบบมีปฏิสัมพันธ์ (Interactive Intrusions) ทั่วโลก ปี 2025

มกราคม – ธันวาคม 2025

สรุปประเด็นสำคัญ



อเมริกาเหนือเผชิญการโจมตี
สูงสุดของโลก **55%**



เอเชียตะวันออกเฉียงและยุโรป
มีความเสี่ยงรองลงมา
5% และ **9%**



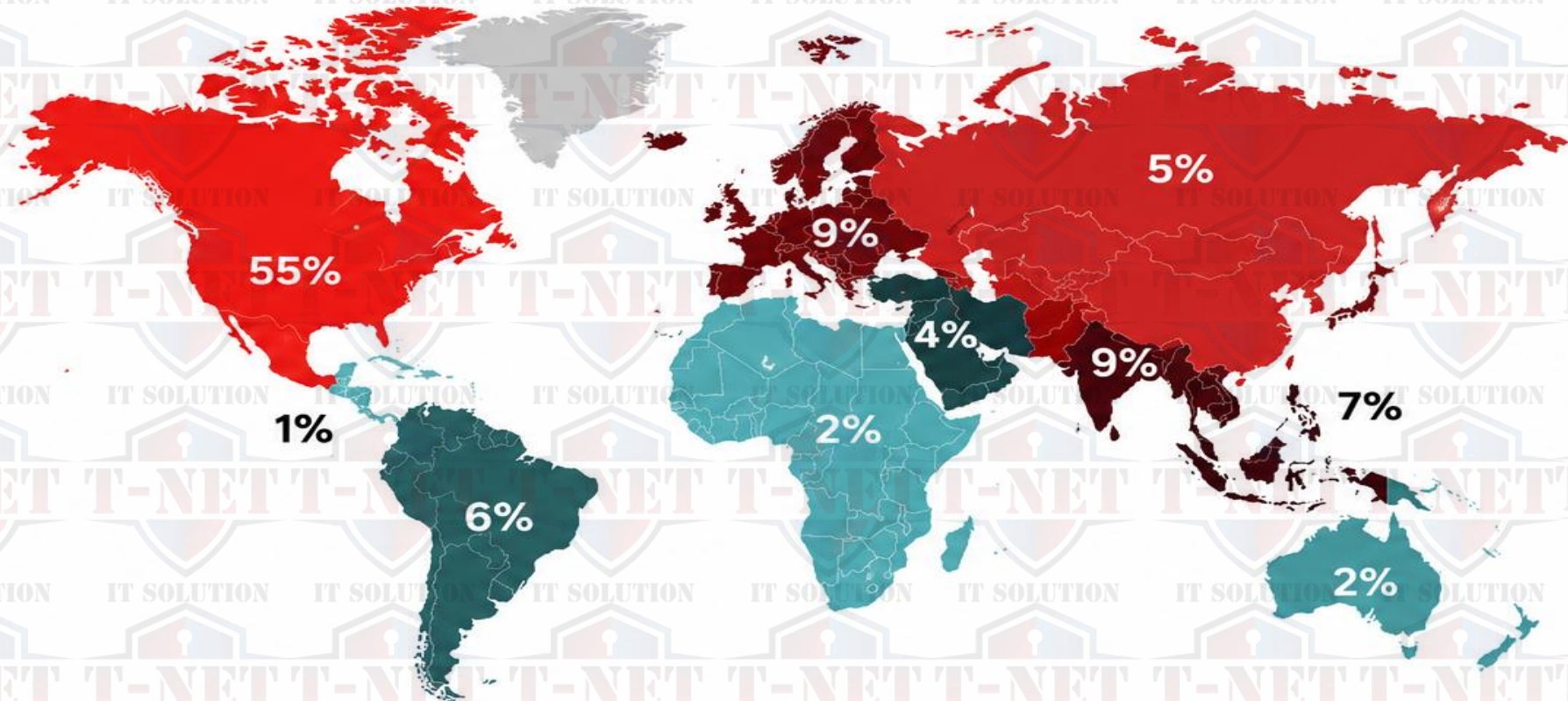
ภูมิภาคตะวันออกกลาง
มีการโจมตีที่ **4%**



อเมริกาใต้ **6%** และ
เอเชียตะวันออกเฉียงใต้ **7%**



แอฟริกา และโอเชียเนีย
มีส่วนการโจมตี **2%** เท่ากัน



อเมริกาเหนือ เอเชียตะวันออกเฉียง ยุโรป เอเชียใต้ เอเชียตะวันออกเฉียงใต้ ตะวันออกกลาง อเมริกาใต้ โอเชียเนีย แอฟริกา อเมริกากลาง



องค์กรทุกภูมิภาคควรเสริมความพร้อมด้านความปลอดภัยไซเบอร์
ตรวจจับภัยคุกคาม**เชิงรุก** และตอบสนองต่อเหตุการณ์ได้อย่าง**รวดเร็ว**



ตรวจจับ
เชิงรุก



ป้องกัน
อย่างมีประสิทธิภาพ



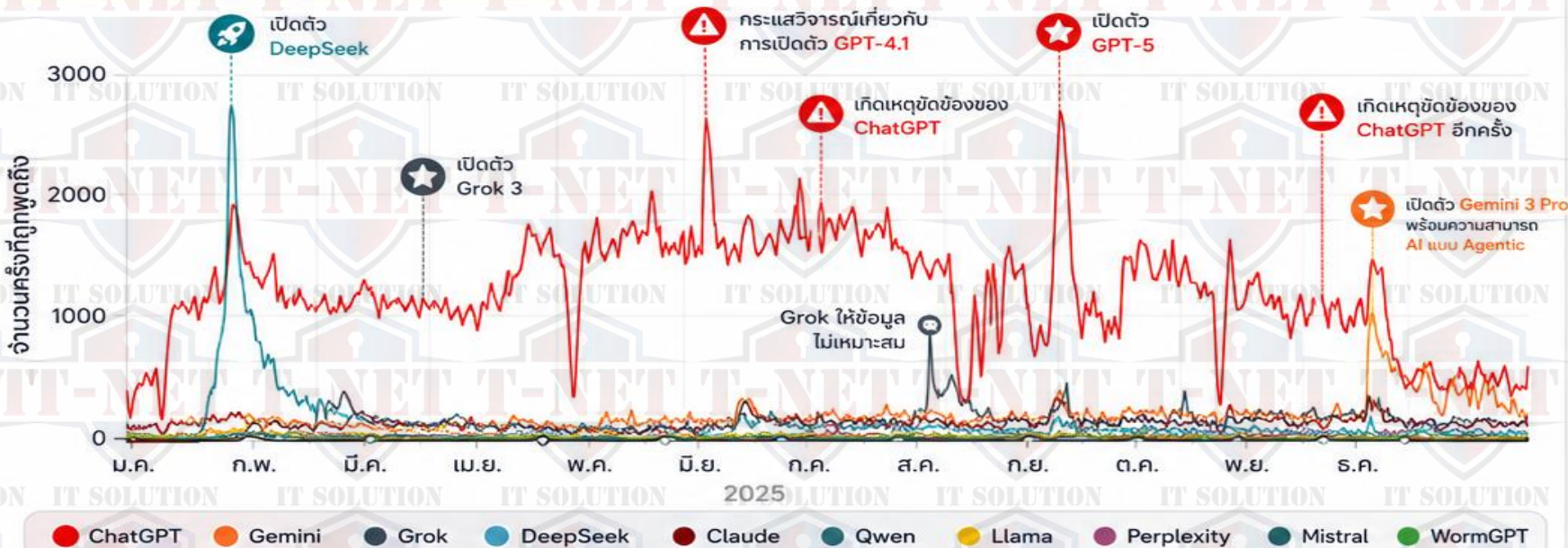
ตอบสนอง
รวดเร็ว

10 อันดับโมเดล AI

ที่ถูกพูดถึงมากที่สุดตลอดปี 2025

การติดตามการพูดถึงโมเดล AI จากแหล่งข้อมูลออนไลน์ทั่วโลก เพื่อวิเคราะห์แนวโน้มความนิยมและเหตุการณ์สำคัญตลอดปี 2025

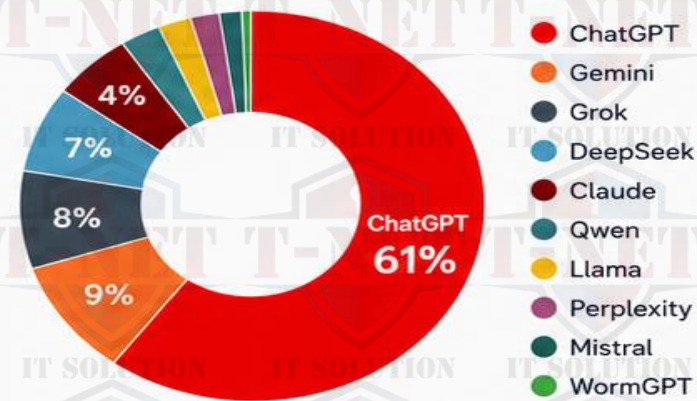
แนวโน้มการพูดถึงโมเดล AI ตลอดปี 2025



เหตุการณ์สำคัญตลอดปี 2025



สัดส่วนการพูดถึงตลอดปี 2025



สรุปประเด็นสำคัญ

- ChatGPT ครองอันดับ 1 ตลอดปี คิดเป็น 61% ของการพูดถึงทั้งหมด
- DeepSeek เปิดตัวช่วง ก.พ. และได้รับความนิยมสูงสุดในระยะสั้น
- เหตุขัดข้องของ ChatGPT เกิดขึ้น 2 ครั้ง ใน มี.ย. และ ธ.ค.
- Gemini 3 Pro เปิดตัวใน ธ.ค. พร้อมความสามารถ AI แบบ Agentic ได้รับความนิยมเพิ่มขึ้น



ข้อเสนอแนะสำหรับองค์กร

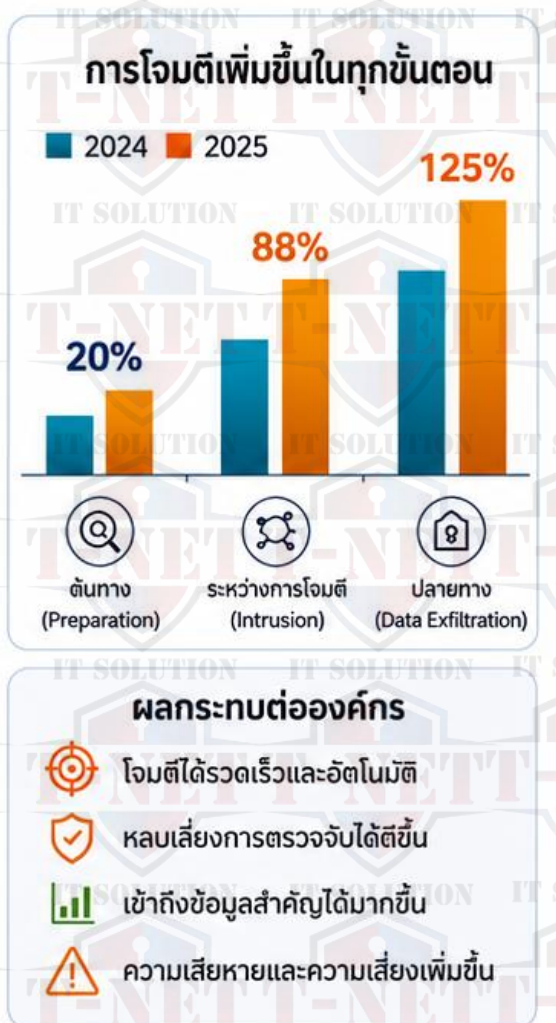
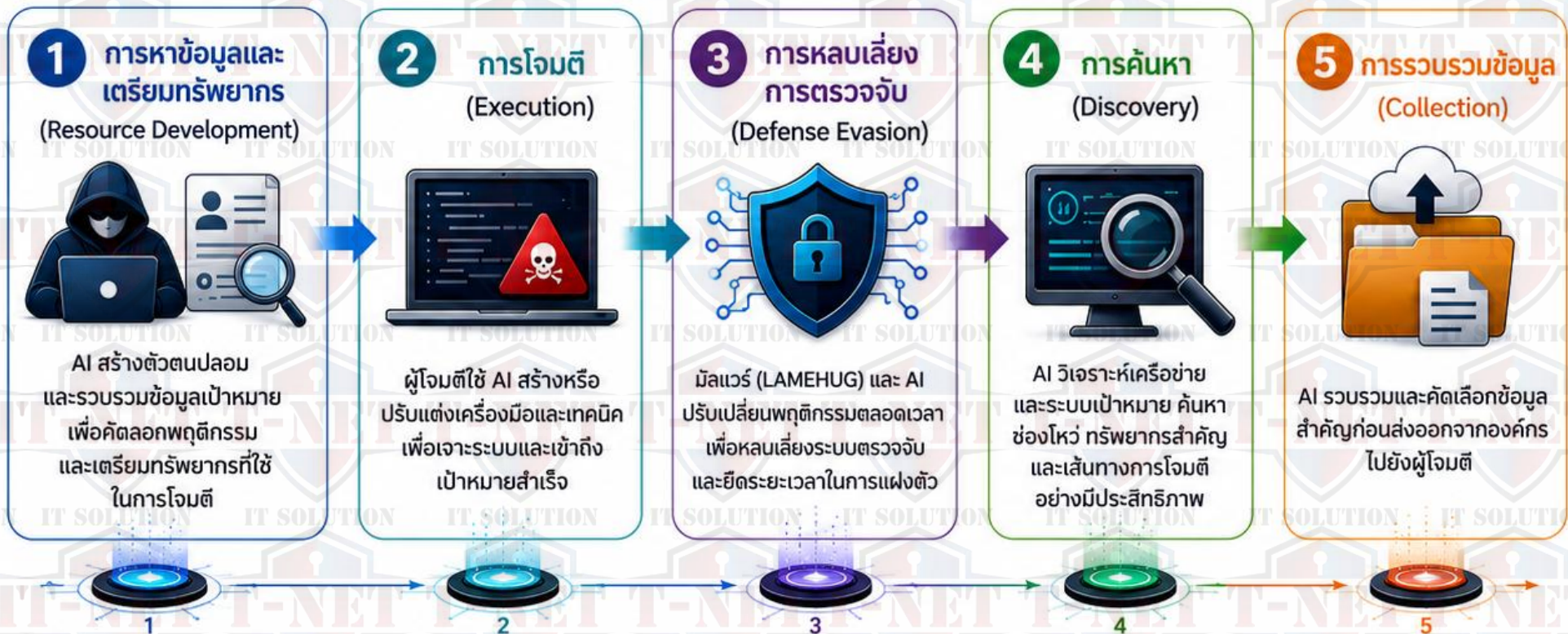
ติดตามแนวโน้มและเหตุการณ์สำคัญของโมเดล AI อย่างต่อเนื่อง เพื่อประเมินความเสี่ยงและโอกาสในการนำไปใช้งาน รวมถึงเตรียมแผนรับมือกับเหตุขัดข้องและประเด็นด้านความปลอดภัยของข้อมูล





ภัยคุกคามจาก AI: การวิเคราะห์เชิงกลยุทธ์

AI ถูกนำมาใช้ในทุกขั้นตอนของการโจมตีไซเบอร์ ทำให้รวดเร็ว แม่นยำ และตรวจจับยากขึ้น



แนวทางรับมือ

- ใช้ระบบตรวจจับและสอบสวนที่ขับเคลื่อนด้วย AI
- เฝ้าระวังข้อมูลและพฤติกรรมแบบเรียลไทม์
- ยกระดับการฝึกอบรมและความตระหนักรู้
- กำหนดนโยบายและควบคุมการใช้งาน AI ภายในองค์กร

สรุป AI ทำให้ภัยคุกคามไซเบอร์มีความซับซ้อนมากขึ้น องค์กรต้องพัฒนากลยุทธ์และเทคโนโลยีด้านความปลอดภัยให้ทันต่อการเปลี่ยนแปลง

แหล่งอ้างอิง

- CrowdStrike Global Threat Report 2025
- Verizon Data Breach Investigations Report 2025
- IBM Cost of a Data Breach Report 2024
- Microsoft Digital Defense Report 2024

■ 2024 ■ 2025

AI HACKING THREATS 2026

ภัยคุกคามทางไซเบอร์ที่ขับเคลื่อนด้วย AI

01



AI Agent Attacks

- การบุกรุกระบบอัตโนมัติ
- การสร้างโปรไฟล์เป้าหมาย
- การโจมตีแบบประสานงาน

02



Deepfake & Voice Cloning

- การปลอมแปลงใบหน้าและเสียง
- การสวมรอยตัวตน
- การหลอกลวงด้วยสื่อสังเคราะห์

03



AI Phishing & Social Engineering

- การสร้างอีเมลหลอกลวงอัตโนมัติ
- การโจมตีแบบเฉพาะบุคคล (Spear Phishing)
- การหลอกลวงข้อมูลสำคัญ

04



AI-powered Malware

- มัลแวร์ปรับเปลี่ยนพฤติกรรมได้
- หลบเลี่ยงการตรวจจับ
- แพร่กระจายและโจมตีอัตโนมัติ

05



Prompt Injection

- หลอก AI ด้วยคำสั่งอันตราย
- บังคับให้เปิดเผยข้อมูล
- เปลี่ยนพฤติกรรมของ AI

06



AI Supply Chain Attacks

- โจมตีโมเดลและไลบรารี AI
- แทรกโค้ดหรือข้อมูลอันตราย
- โจมตีห่วงโซ่อุปทานซอฟต์แวร์

07



AI Reconnaissance

- รวบรวมข้อมูลเป้าหมายอัตโนมัติ
- ค้นหาช่องโหว่ด้วย AI
- วิเคราะห์พื้นผิวการโจมตี (Attack Surface)

08



Credential Theft & Identity Attacks

- ขโมยข้อมูลรับรอง (Credentials)
- สวมรอยตัวตนผู้ใช้งาน
- ยึดครองบัญชีผู้ใช้ (Account Takeover)

09



AI Defense Evasion

- หลบเลี่ยง EDR/AV/SIEM
- เปลี่ยนพฤติกรรมเพื่อหลบการตรวจจับ
- ลดโอกาสถูกวิเคราะห์

10



Shadow AI & Data Leakage

- ใช้ AI ที่ไม่ได้รับอนุญาต
- ข้อมูลสำคัญรั่วไหลสู่ภายนอก
- เพิ่มความเสี่ยงด้านการกำกับดูแลข้อมูล



สรุป

AI ทำให้การโจมตีไซเบอร์มีความรวดเร็ว อัตโนมัติ แม่นยำ และซับซ้อนมากขึ้น
องค์กรต้องยกระดับการตรวจจับ การป้องกัน และการบริหารความเสี่ยงด้าน AI อย่างต่อเนื่อง



เฝ้าระวัง
อย่างต่อเนื่อง



ป้องกัน
เชิงรุก



พัฒนาคน
และกระบวนการ



ใช้ AI เพื่อ
ความปลอดภัย



HUMAN HACKER **V** AI HACKER

DIFFERENT MINDS. SAME TARGET.



HUMAN HACKER

ใช้ประสบการณ์ ความคิดสร้างสรรค์ และเทคนิคที่ซับซ้อน
เพื่อเจาะระบบและหาผลประโยชน์จากช่องโหว่

STRENGTHS

-  Experience
-  Creativity
-  Critical Thinking
-  Business Logic Abuse
-  Social Engineering

CHARACTERISTICS

-   **เข้าใจระบบธุรกิจ**
เข้าใจโครงสร้างระบบและกระบวนการทางธุรกิจ เพื่อนำไปใช้โจมตี
-   **ดัดแปลงเทคนิคใหม่**
คิดค้นและปรับเปลี่ยนเทคนิคเพื่อหลีกเลี่ยงการตรวจจับ
-   **วางแผนโจมตีหลายชั้น**
วางแผนการโจมตีอย่างเป็นขั้นตอนเพื่อใช้บรรลุเป้าหมาย
-   **ปรับเปลี่ยนตามสถานการณ์**
ปรับกลยุทธ์และเทคนิคให้เหมาะสมกับสถานการณ์ที่เปลี่ยนแปลง

-  **RECONNAISSANCE**
- WHOIS
 - DNS
 - OSINT

-  **EXPLOITATION**
- > exploit.sh
 - > target found
 - > executing...

-  **SOCIAL ENGINEERING**
- > person + speech bubble
 - > star

-  **PRIVILEGE ESCALATION**
- > exploit success
 - > access
 - > granted

- DATA EXFILTRATION**
- > collecting data
 - > exfiltrating...



คิดเป็น วางแผนเป็น
โจมตีอย่างชาญฉลาด



ค้นหาจุดอ่อน



วางแผน



ลงมือโจมตี



ปรับตัว



บรรลุเป้าหมาย



สรุป

Human Hacker อาศัยประสบการณ์ ความคิดสร้างสรรค์ และทักษะเฉพาะทาง
ในการวางแผนและดำเนินการโจมตีที่ซับซ้อน เพื่อให้บรรลุเป้าหมาย



เฝ้าระวัง
อย่างต่อเนื่อง



ป้องกันเชิงรุก
กันสมัย



พัฒนาคน
และกระบวนการ



ใช้ AI เพื่อ
เสริมการป้องกัน

วงจรการโจมตีทางไซเบอร์ (CYBER KILL CHAIN)



สรุป

วงจรการโจมตีทางไซเบอร์ประกอบด้วย 7 ขั้นตอน
ผู้โจมตีจะดำเนินการอย่างเป็นลำดับ เพื่อให้บรรลุเป้าหมาย
การเข้าใจวงจรนี้ช่วยให้องค์กรสามารถ **ป้องกัน ตรวจสอบ และตอบสนอง**
ได้อย่างมีประสิทธิภาพ



ป้องกัน



ตรวจสอบ



แจ้งเตือน



ตอบสนอง



AI HACKER

ความเร็ว อัตโนมัติ การวิเคราะห์ขนาดใหญ่
ขับเคลื่อนการโจมตีในระดับมหาศาล

STRENGTHS



Speed

ทำงานรวดเร็วมาก



Automation

ทำงานอัตโนมัติ 24/7



Large-scale Analysis

วิเคราะห์ข้อมูลจำนวนมาก



AI-assisted Exploitation

ช่วยค้นหาและใช้ช่องโหว่ได้แม่นยำ



Autonomous Reconnaissance

ค้นหาเป้าหมายได้ด้วยตนเอง

CHARACTERISTICS



Scan หลายล้าน IP

ค้นหาเป้าหมายจำนวนมาก
ภายในเวลาอันรวดเร็ว



วิเคราะห์ได้ลึกทันที

AI วิเคราะห์ช่องโหว่และบั๊กในโค้ด
ได้อย่างรวดเร็วและแม่นยำ



สร้าง Phishing Email

สร้างอีเมลฟิชชิ่งอัตโนมัติ
ปรับให้เหมาะกับเป้าหมายอัตโนมัติ



เขียน Malware Prototype

สร้างโค้ดมัลแวร์ต้นแบบได้ภายในไม่กี่วินาที
ปรับเปลี่ยนรูปแบบได้อัตโนมัติ



Generate Exploit

สร้างสคริปต์หรือโค้ด Exploit
สำหรับเจาะช่องโหว่โดยอัตโนมัติ

AUTONOMOUS RECONNAISSANCE



TARGET DISCOVERY

5,285,412 IPs

SCANNING...

CODE ANALYSIS

- Scan Source Code
- Detect Vulnerabilities
- Analyze Dependencies
- Identify Security Flaws
- Risk Assessment

98%

VULNERABILITY FOUND

PHISHING GENERATION

To: CEO
Subject: Urgent Action Required

REALISTIC | PERSONALIZED

HIGH SUCCESS RATE



MALWARE PROTOTYPE

- File Infection
- Persistence Mechanism
- Anti-Analysis
- Data Exfiltration
- Command & Control
- Evasion Technique



EXPLOIT GENERATION

- Auto Exploit
- Vulnerability Discovery
- Payload Generation
- Exploit Delivery
- Privilege Escalation



AI ขับเคลื่อนการโจมตีแบบอัตโนมัติ



1. RECONNAISSANCE

ค้นหาเป้าหมาย



2. SCANNING

สแกนช่องโหว่



3. ANALYSIS

วิเคราะห์ด้วย AI



4. PHISHING

สร้างฟิชชิ่ง



5. MALWARE

สร้างมัลแวร์



6. EXPLOIT

สร้าง Exploit



7. COMPROMISE

เจาะระบบสำเร็จ



สรุป

AI Hacker ใช้ความเร็ว การทำงานอัตโนมัติ และการวิเคราะห์ข้อมูลขนาดใหญ่
เพื่อค้นหาเป้าหมาย วิเคราะห์ช่องโหว่ สร้างมัลแวร์และ Exploit รวมถึงเจาะระบบได้อย่างมีประสิทธิภาพ



ป้องกัน

เตรียมพร้อม
ก่อนถูกโจมตี



ตรวจจับ

เฝ้าระวังและตรวจจับ
พฤติกรรมผิดปกติ



แจ้งเตือน

แจ้งเตือนเร็ว
ลดความเสียหาย



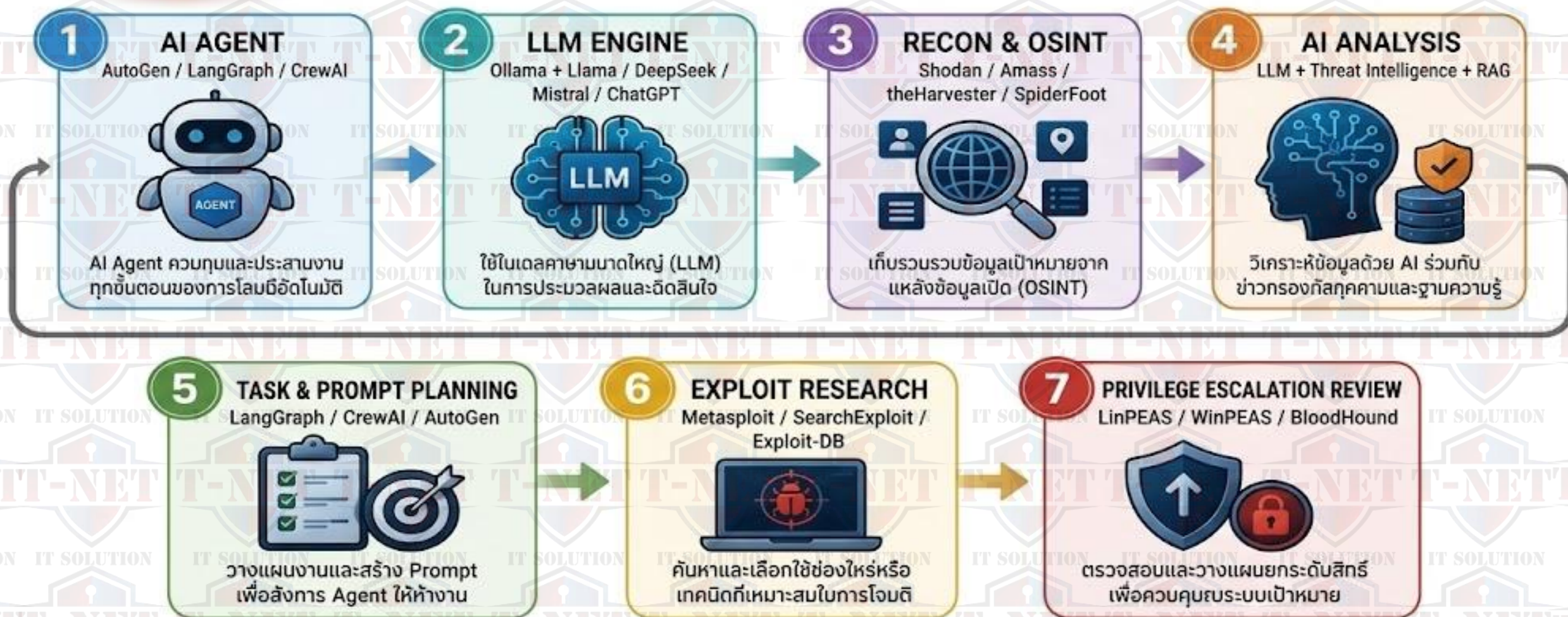
ตอบสนอง

ตอบสนองทันที
ควบคุมสถานการณ์



AI Agent อัตโนมัติ: ATTACK WORKFLOW

ขั้นตอนการโจมตีใช้เบอร์ด้วย AI Agent แบบอัตโนมัติ



AI Agent ทำงานแบบอัตโนมัติ ครอบคลุมตั้งแต่การรวบรวมข้อมูล วิเคราะห์ วางแผน ค้นหาช่องโหว่จนถึงการยกระดับสิทธิ์ เพื่อเพิ่มประสิทธิภาพและความรวดเร็วในการโจมตี



รวดเร็ว



แม่นยำ



อัตโนมัติ



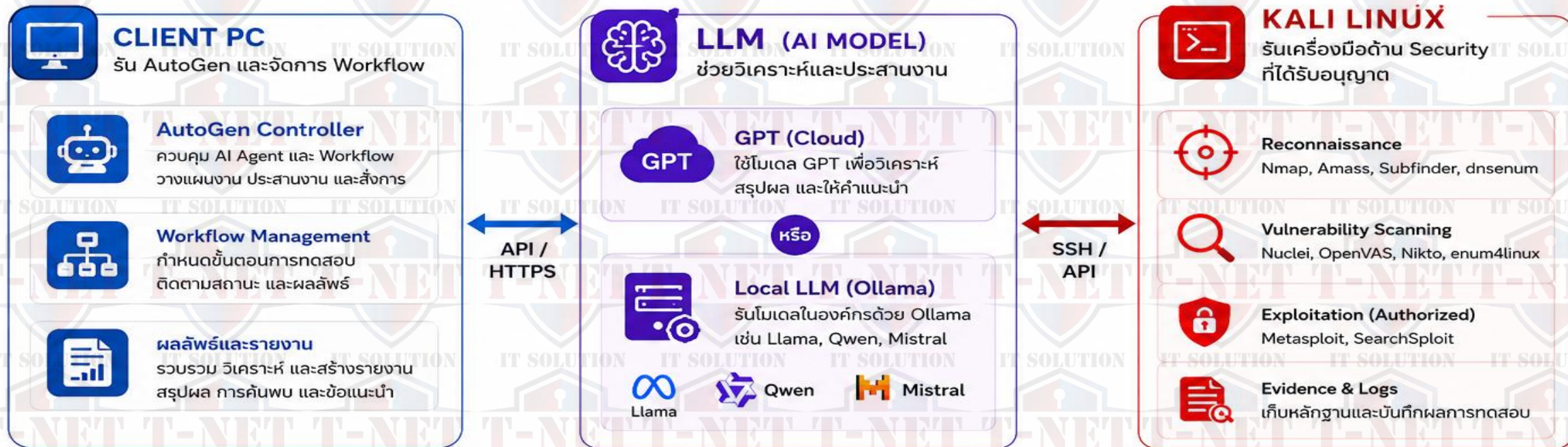
แบบเนียน

AI Agent Framework

AI Agent Framework	ผู้พัฒนา	จุดเด่น	เหมาะกับ Penetration Testing
LangGraph	LangChain	ออกแบบ Workflow และ State Machine สำหรับ AI Agent	Recon, Vulnerability Analysis, Pentest Workflow
OpenAI Agents SDK	OpenAI	สร้าง AI Agent ที่เรียกใช้ Tools, MCP และ Multi-Agent ได้	Tool Orchestration, Automation
CrewAI	CrewAI	AI Agent หลายบทบาททำงานร่วมกัน	Recon, Reporting, Incident Response
AutoGen	Microsoft	Multi-Agent Collaboration ให้ Agent สนทนา และแบ่งงานกัน	Research, Analysis, Code Review
Semantic Kernel	Microsoft	เชื่อม LLM กับ .NET/Python และระบบองค์กร	Enterprise Security Automation

สถาปัตยกรรมการทดสอบเจาะระบบด้วย AutoGen

ทำงานร่วมกันอย่างปลอดภัย ภายในองค์กรที่ได้รับการอนุญาต



**ปลอดภัย
ได้รับอนุญาต**

- ทดสอบเฉพาะระบบภายในองค์กรตามขอบเขตที่กำหนด
- ปฏิบัติตามกฎหมายและนโยบายความปลอดภัยขององค์กร
- เก็บรักษาข้อมูลอย่างปลอดภัย ไม่รั่วไหล

ขั้นตอนการทำงาน (Workflow)



หมายเหตุ: ใช้สำหรับการทดสอบเจาะระบบที่ได้รับอนุญาตเท่านั้น (Authorized Penetration Testing)

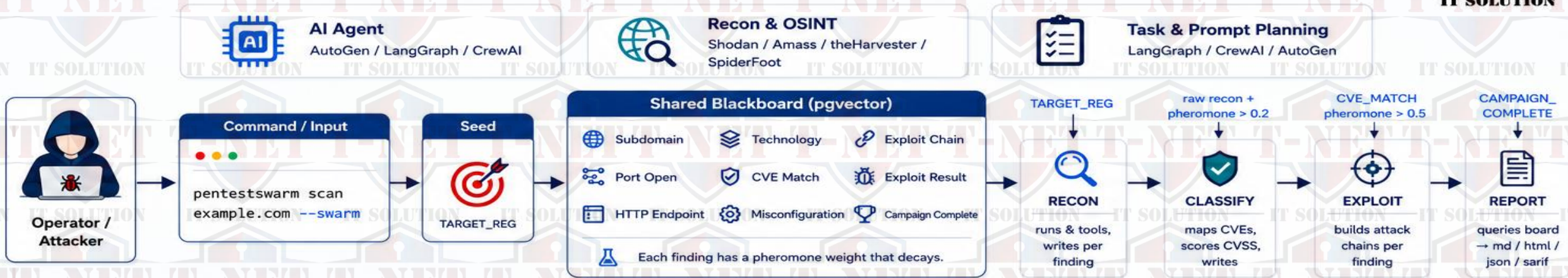
AI-powered Penetration Testing

เครื่องมือ	ประเภท	จุดเด่น
PentestGPT	AI Pentest Assistant	ใช้ LLM ช่วยวิเคราะห์ วางแผน และตีความผลการทดสอบเจาะระบบ
XBOW	Autonomous AI Pentesting Platform	ใช้ AI ทำ Vulnerability Discovery และ Penetration Testing แบบอัตโนมัติในขอบเขตที่กำหนด
PentestSwarm	Multi-Agent Pentesting Framework	ใช้ AI Agent หลายตัวแบ่งบทบาทและทำงานร่วมกันในกระบวนการ Pentest
HackerGPT	AI Security Assistant	ผู้ช่วยด้าน Security สำหรับการวิเคราะห์ การค้นคว้า และการอธิบายผลด้านความมั่นคงปลอดภัย
Archon	AI Security Agent Framework	เฟรมเวิร์กสำหรับสร้าง AI Agent เพื่อสนับสนุนงาน Security Automation และ Penetration Testing Workflow

DEMO: Pentest Swarm AI Diagram



AI-assisted autonomous penetration testing workflow, environment readiness, and live campaign results



LLM Engine
Ollama + Llama / DeepSeek / Mistral / ChatGPT

AI Analysis
LLM + Threat Intelligence + RAG

Exploit Research
Metasploit / SearchSploit / Exploit-DB

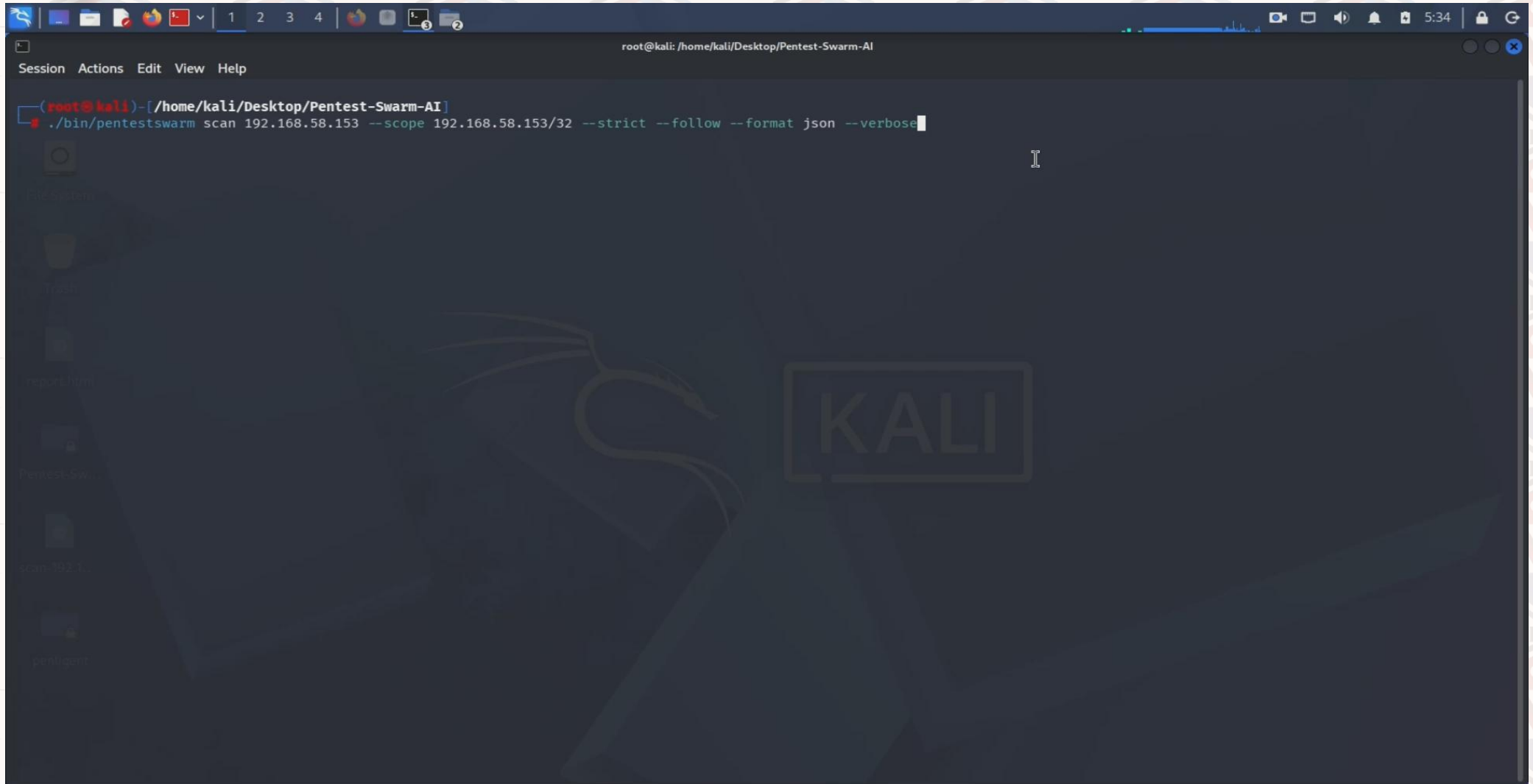
Privilege Escalation Review
LinPEAS / WinPEAS / BloodHound

1 Environment Readiness (Doctor Check)	
8/8 infra checks passed	
API server reachable	listening on :8080
PostgreSQL connection	listening on :5432
Redis connection	listening on :6379
Ollama running	listening on :11434
Docker daemon	v29.6.1
Go version	go1.26.4
Disk space (>10GB)	check passed
RAM (>8GB)	unable to determine

Security Tools	
Reconnaissance	subfinder, dnsx, httpx, naabu, katana, gau, nmap, amass
Vulnerability scanners	nuclei, sqlmap
Content discovery	ffuf, gobuster
Source / secret scanning	trufflehog, gitleaks
Evidence capture	gowitness
Missing	semgrep → "pip install semgrep"
15/16 tools present	

3 Live Campaign Findings	
Target	192.168.58.153
Status Flow	Initialized → Recon → Classifying → Planning → Executing
Recon Tools Used	naabu, httpx, nuclei
Recon Results	0 subdomains, 0 hosts, 6 endpoints
Classified Findings	15 total
Severity Summary	4 Critical, 0 High, 11 Medium
Built Attack Paths	3
Top Path	known_exploit → vsftpd 2.3.4 (50% estimated success)
Critical Findings	
vsftpd 2.3.4 — CVSS 10.0 — exploited via CVE-2011-2523 — success	
Samba 3.0.20-Debian — CVSS 10.0 — exploited via CVE-2007-2447 — success	
distccd (RCE vulnerable) — CVSS 9.8 — exploited via CVE-2004-2687	
Medium Highlights	
Multilidae, DVWA, DAV, phpMyAdmin, TWiki	
Apache 2.2.8, OpenSSH 4.7p1, PHP 5.2.4, PostgreSQL default creds, MySQL port open	
OS detected: Metasploitable2 (Ubuntu-based Linux)	
05:34:15 [*] Starting reconnaissance on 192.168.58.153	
05:35:02 [*] Recon tools (naabu, httpx, nuclei) running	
05:35:45 [✓] Found 0 subdomains, 0 hosts, 6 endpoints	
05:36:08 [✓] Classifying 15 raw findings	
05:37:10 [✓] Built 3 attack paths	
05:38:15 [✓] Step succeeded: vsftpd 2.3.4 (CVE-2011-2523)	
05:38:33 [✓] Step succeeded: Samba 3.0.20-Debian (CVE-2007-2447)	

AI Agent Attack Workflow



The screenshot shows a Kali Linux terminal window with the title bar "root@kali: /home/kali/Desktop/Pentest-Swarm-AI". The terminal content shows the command `./bin/pentestswarm scan 192.168.58.153 --scope 192.168.58.153/32 --strict --follow --format json --verbose` being entered. The terminal window has a menu bar with "Session", "Actions", "Edit", "View", and "Help". The background of the terminal window features a large, faint watermark of a dragon and the word "KALI".

```
root@kali: /home/kali/Desktop/Pentest-Swarm-AI
Session Actions Edit View Help
(root@kali)-[/home/kali/Desktop/Pentest-Swarm-AI]
└─$ ./bin/pentestswarm scan 192.168.58.153 --scope 192.168.58.153/32 --strict --follow --format json --verbose
```




PENLIGENT

The World's First Agentic AI Hacker

การทำงานของ Penlagent

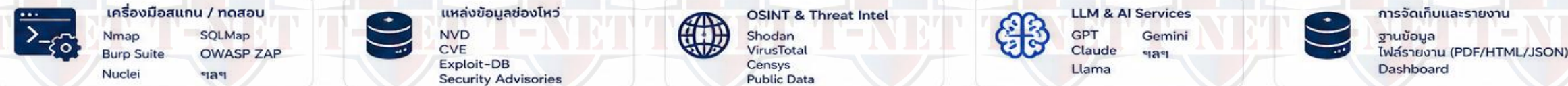
AI-Powered Penetration Testing Platform



AI Agent ทำงานร่วมกันแบบอัตโนมัติ

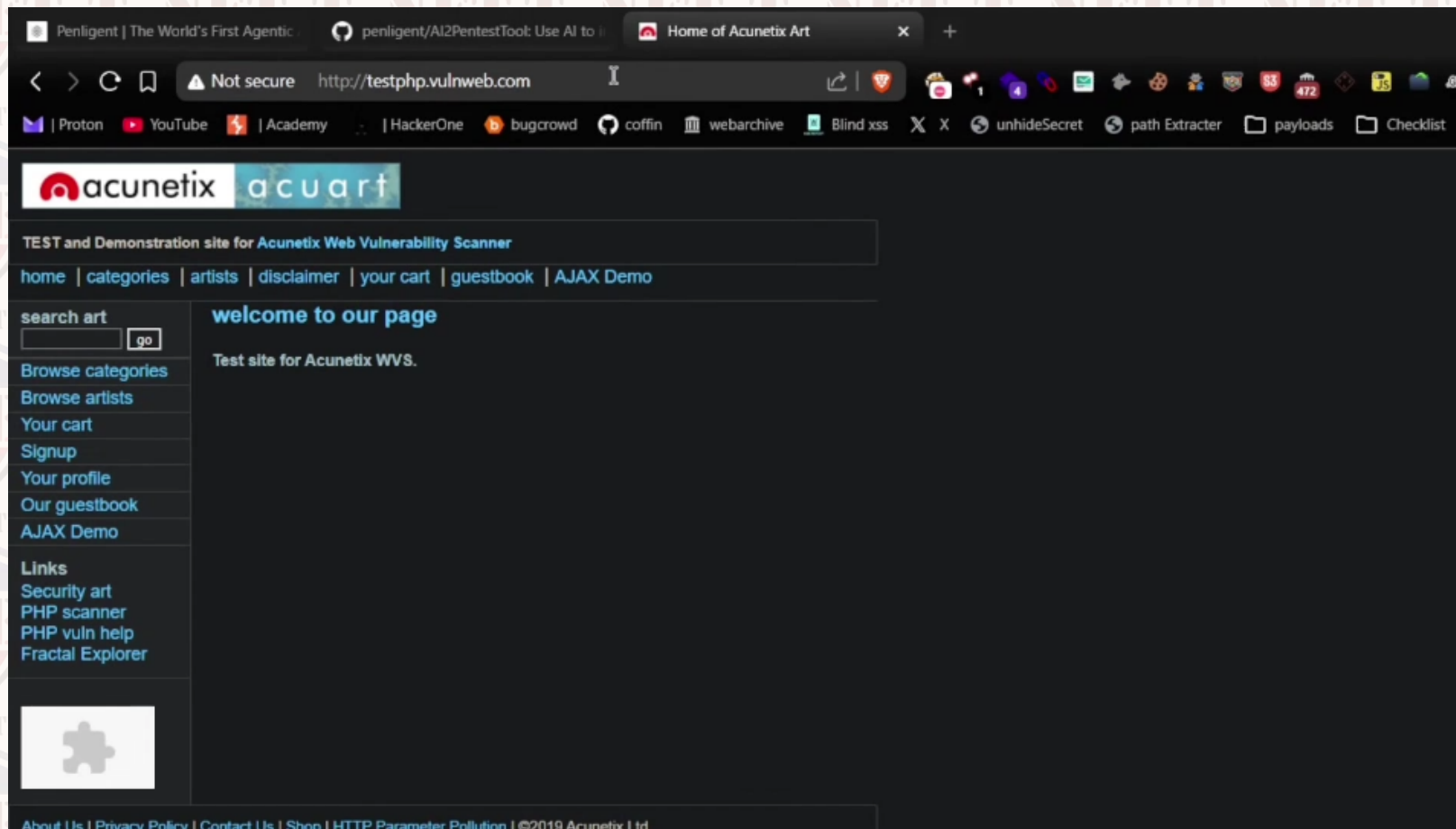


เครื่องมือและแหล่งข้อมูลที่เชื่อมต่อ



หมายเหตุ : Penlagent ทำงานภายใต้ขอบเขตและการอนุญาตที่ชัดเจนเท่านั้น ผลลัพธ์ที่ได้ควรได้รับการตรวจสอบและยืนยันโดยผู้เชี่ยวชาญก่อนนำไปใช้งาน

AI Agent Attack Workflow



ความแตกต่างระหว่าง การโจมตีแบบมนุษย์ vs การโจมตีอัตโนมัติด้วย AI



Human Attack



Manual Recon

มนุษย์ค้นหาข้อมูลเป้าหมายด้วยตนเอง
ใช้เวลาและประสบการณ์



Manual Analysis

วิเคราะห์ข้อมูลและหาจุดอ่อนด้วยตนเอง
อาศัยความเชี่ยวชาญเฉพาะบุคคล



Manual Exploit

ค้นหาและเลือกใช้ช่องโหว่ด้วยตนเอง
ทดสอบและปรับแต่งแบบ Manual



Manual Phishing

สร้างอีเมลหลอกลวงด้วยตนเอง
ใช้ภาษาและเทคนิคพื้นฐาน



Manual Lateral Movement

วางแผนและเคลื่อนที่ในระบบด้วยตนเอง
ตัดสินใจแบบทีละขั้นตอน



Human Speed

ความเร็วจำกัดตามศักยภาพของมนุษย์
ทำได้ครั้งละไม่กี่อย่าง

VS



AI-powered Autonomous Attack



AI Recon

AI เก็บรวบรวมข้อมูลอัตโนมัติจากหลายแหล่ง
รวดเร็วและครอบคลุมกว่า



LLM Analysis

ใช้ LLM วิเคราะห์ข้อมูลและหาจุดอ่อน
เข้าใจบริบทและสรุปได้อย่างแม่นยำ



AI-assisted Exploit Research

AI ค้นหาและแนะนำช่องโหว่ พร้อมโค้ดตัวอย่าง
ปรับแต่งได้เหมาะกับเป้าหมาย



AI-generated Phishing

AI สร้างอีเมลแบบเนียนเฉพาะบุคคล
หลอกลวงได้แบบเนียนและมีประสิทธิภาพสูง



AI-assisted Decision

AI ช่วยตัดสินใจและเลือกเส้นทางที่ดีที่สุด
ปรับตัวตามสถานการณ์แบบเรียลไทม์



Machine Speed

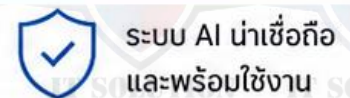
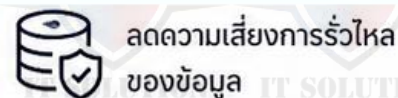
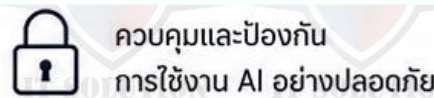
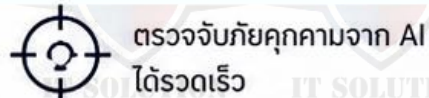
ทำงานได้หลายอย่างพร้อมกัน 24/7
รวดเร็ว แม่นยำ และต่อเนื่อง



AI ทำให้การโจมตีเร็วขึ้น อัตโนมัติมากขึ้น และตรวจจับได้ยากกว่าเดิม
องค์กรต้องยกระดับการป้องกันให้เท่าทันยุค AI

AI Agent Detection & Defense Workflow

ขั้นตอนการตรวจจับและป้องกันภัยคุกคามจาก AI Agent อย่างเป็นระบบ



AI Security Guidelines

- ☐ แนวปฏิบัติการใช้ปัญญาประดิษฐ์อย่างมั่นคงปลอดภัย
- ☐ จัดทำโดย : สำนักงานคณะกรรมการการรักษาความมั่นคงปลอดภัยไซเบอร์แห่งชาติ (สกมช. : NCSA Thailand)
- ☐ ดาวน์โหลดจากเว็บไซต์ สกมช.
 - https://www.ncsa.or.th/standards?utm_source=chatgpt.com
- ☐ เว็บไซต์หลัก NCSA Thailand
 - <https://www.ncsa.or.th>



มาตรฐานและแนวปฏิบัติสำคัญด้าน AI Security

แนวทางสากลสำหรับการพัฒนา ใช้งาน และบริหารจัดการระบบ AI อย่างมั่นคงปลอดภัย

1 NIST AI RMF 1.0 AI Risk Management Framework



- ✓ กรอบการบริหารความเสี่ยง AI แบบครบวงจร
- ✓ ครอบคลุมตลอดวงจรชีวิตของระบบ AI
- ✓ เน้นธรรมาภิบาล ความโปร่งใส และความน่าเชื่อถือ

2 OWASP Top 10 for LLM Applications



- ✓ จัดอันดับความเสี่ยง 10 อันดับแรกของการใช้งาน LLM
- ✓ ช่วยให้นักพัฒนาและองค์กรลดความเสี่ยงได้อย่างตรงจุด

3 OWASP GenAI Security Project



- ✓ ชุดทรัพยากรและแนวทางความปลอดภัยสำหรับ Generative AI
- ✓ Tools, Cheatsheets, Prompt Security, Testing Guide และอื่นๆ
- ✓ อัปเดตต่อเนื่องโดยชุมชนความปลอดภัย

4 MITRE ATLAS Adversarial Threat Landscape for AI Systems



- ✓ ฐานข้อมูลเทคนิคการโจมตีระบบ AI
- ✓ ช่วยวิเคราะห์ภัยคุกคามและวางแผนการป้องกันได้อย่างเป็นระบบ
- ✓ อ้างอิงจากการโจมตีจริงในโลกไซเบอร์

5 ISO/IEC 42001 AI Management System



- ✓ มาตรฐานระบบบริหารจัดการ AI
- ✓ กำหนดนโยบาย กระบวนการ และการควบคุมเพื่อการใช้งาน AI อย่างรับผิดชอบ
- ✓ บูรณาการเข้ากับ ISO/IEC 27001 ได้ง่าย

6 ISO/IEC 23894 AI Risk Management



- ✓ แนวทางการบริหารความเสี่ยงของระบบ AI
- ✓ ครอบคลุมการระบุ ประเมิน และจัดการความเสี่ยง
- ✓ ใช้ควบคู่กับ AI RMF และ ISO/IEC 42001

7 NIST SP 800-61 Rev.3 Incident Response



- ✓ แนวทางรับมือเหตุการณ์ความมั่นคงปลอดภัยเวอร์ชันล่าสุด
- ✓ ครอบคลุมการเตรียมพร้อม ตรวจสอบ ตอบสนอง กู้คืน และปรับปรุงหลังเกิดเหตุ
- ✓ ใช้ได้กับระบบ IT และ AI

8 NIST CSF 2.0 Cybersecurity Framework



- ✓ กรอบการบริหารความมั่นคงปลอดภัยไซเบอร์
- ✓ ยืดหยุ่น ปรับใช้ได้กับทุกขนาดองค์กร
- ✓ ใช้เป็นแนวทางหลักในการสร้าง Resilience



การบูรณาการมาตรฐานเหล่านี้ จะช่วยให้องค์กรพัฒนาและใช้งาน AI ได้อย่างปลอดภัย น่าเชื่อถือ และสอดคล้องตามแนวทางสากล



<https://www.tnetitsolution.co.th>



tnetitsolution



tnetitsolution



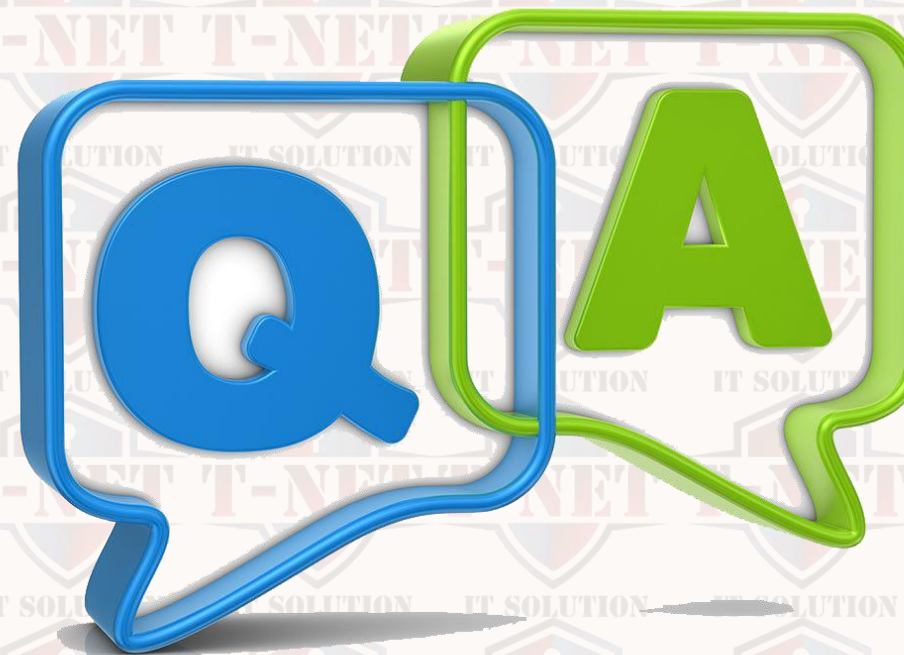
tnetitsolution



+66(0)-8-8874-4741



Info@tnetitsolution.co.th



Search T-NET IT Solution

