



WEB PHISHING HAKING

หลักการและเหตุผล

การจำลองฟิชชิ่งด้วย AI ช่วยตรวจจับภัยคุกคามจากฟิชชิ่งได้อย่างแม่นยำและรวดเร็ว โดยใช้ Machine Learning และ Natural Language Processing ในการวิเคราะห์อีเมลและเว็บไซต์ปลอม AI ลดข้อผิดพลาดจากการตรวจจับโดยมนุษย์และช่วยให้การป้องกันภัยคุกคามมีประสิทธิภาพมากขึ้น การสร้างแดชบอร์ดช่วยให้ผู้ใช้งานเข้าถึงข้อมูลที่จำเป็นในการตัดสินใจได้ดีขึ้น ทำให้การรักษาความปลอดภัยในโลกดิจิทัลมีความทันสมัยและมีประสิทธิภาพสูงขึ้น.

ขอบเขตการศึกษา

โปรเจกต์นี้ศึกษาการจำลองการทำฟิชชิ่งและใช้ AI ในการตรวจจับฟิชชิ่ง โดยการจำลองเว็บฟิชชิ่งและอีเมลปลอมเพื่อตรวจสอบความเสี่ยงจากฟิชชิ่งด้วย AI โมเดล เช่น BERT สำหรับข้อความและ Logistic Regression สำหรับ URL ผลการวิเคราะห์จะแสดงบนแดชบอร์ดที่พัฒนาโดยใช้ React.js และ Tailwind CSS โดยการใช้ Backend ในการส่งข้อมูลและประมวลผลความเสี่ยง ช่วยเพิ่มความปลอดภัยในโลกดิจิทัล.

โปรแกรมที่ใช้พัฒนา

NODE.JS
JAVASCRIPT
HUGGING FACE
REACT
TAILWINDCSS



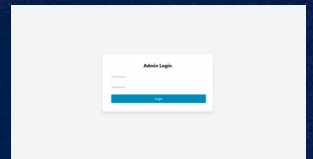
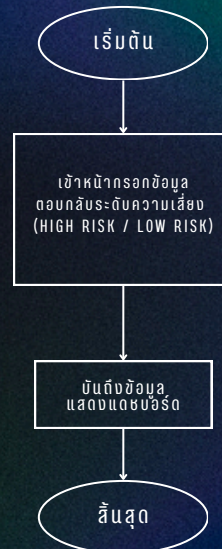
วัตถุประสงค์

- เพื่อจำลองการทำฟิชชิ่ง
- เพื่อพัฒนา AI ในการตรวจจับฟิชชิ่ง
- เพื่อพัฒนาแดชบอร์ด
- เพื่อเพิ่มความปลอดภัยในโลกดิจิทัล

อาจารย์ที่ปรึกษา

ศ.ดร.จักรชัย ใสอินทร์ (ศาสตราจารย์)
ผศ.ดร.สาริต กระเวนกิต

หลักการทำงาน



ผลการดำเนินงาน

การตรวจสอบการโจมตีจากการให้ AI ตรวจสอบเพื่อเพิ่มความแม่นยำและความปลอดภัยจากการประมวลผลธรรมชาติ เพื่อไม่ให้เกิดความเสี่ยงต่อการส่งข้อมูลให้มีฉ้อฉล

อ้างอิง

<https://huggingface.co/distilbert/distilbert-base-uncased-finetuned-sst-2-english#how-to-get-started-with-the-model>

ผู้จัดทำ

663380362-9 WORAKAN PRAPNOK
663380338-6 CHIRUS ARAYAPONGPAKDEE
663380342-5 CHAKHRIT CHARAWONG
663380578-6 APHITSARA NAKHONSUK
663380561-3 THEERATHAT MAHARIT
663380580-9 AMPHON KAEWPAKDEE
663380372-6 MATHAWEE SITTICHAINATE