
The role of artificial intelligence in achieving cybersecurity[#]**Hazem Gad Ali Ibrahim ^{(1)*}**

(1) Researcher, Egypt.

Received: 1/2/2026

Accepted: 10/5/2026

Published: 23/7/2026

*** Corresponding Author:**

DOI:

<https://doi.org/10.59759/law.v5i2.1965>

Abstract

In recent years, the digital ecosystem has witnessed a qualitative shift toward highly sophisticated and destructive cyber threats, rendering traditional rule-based defense mechanisms obsolete in the face of these dynamic challenges. This study aims to explore and analyze the pivotal and revolutionary role that Artificial Intelligence (AI) plays in strengthening cybersecurity and reinforcing digital defense paradigms.

The study begins by conceptualizing the core terms, defining AI as the capability of technological systems to simulate human intelligence and self-learn, while framing cybersecurity as the holistic practice of protecting data confidentiality, integrity, and availability. It highlights the complex, dialectical relationship between the two fields; although AI serves as a powerful shield for organizations, it is simultaneously weaponized by threat actors to engineer automated, hyper-targeted malware, thereby acting as a double-edged sword.

Furthermore, the research dissects the practical applications of AI as a robust cyber defense mechanism. It demonstrates AI's superior ability in the proactive detection of novel attacks, such as Zero-Day exploits, by continuously mapping network behavioral baselines. It also highlights AI's capacity to process and analyze structured and unstructured Big Data at unprecedented speeds, which effectively correlates fragmented security alerts and minimizes false positives. Additionally, the study examines AI's critical role in upgrading Identity and Access Management (IAM) systems through continuous behavioral authentication and the mitigation of advanced social engineering and phishing attacks.

Finally, the paper addresses the complex legal and ethical dimensions governing AI implementation. It explores the critical challenges of data privacy—specifically compliance with global frameworks like GDPR—alongside the need for algorithmic transparency (Explainable AI) and the mitigation of digital bias. Crucially, it evaluates the intricate legal dilemma of establishing civil and criminal liabilities for damages or failures caused by autonomous AI decisions, examining the ongoing debate between developer accountability, operator liability, and the conceptual emergence of an independent electronic legal personality.

[#] A special issue for the Third International Conference, entitled: **Artificial Intelligence, A Legal Vision Towards a Sustainable Future.**

The study concludes that integrating AI into digital defense is no longer a strategic luxury but an absolute necessity for national and organizational security. However, the efficacy of this technological shield remains strictly contingent upon the establishment of robust international legislative and governance frameworks that ensure its ethical and responsible deployment.

Among the most important recommendations that the study concluded with are: the need for international and regional cooperation to improve artificial intelligence models to make them more resistant to cyberattacks, and the need to establish a global platform for exchanging information about AI-based cyberattacks.

دور الذكاء الاصطناعي في تحقيق الأمن السيبراني

حازم جاد علي إبراهيم⁽¹⁾

(1) باحث، جمهورية مصر العربية.

ملخص

شهد الفضاء الرقمي في الآونة الأخيرة تحولاً نوعياً نحو هجمات سيبرانية شديدة التعقيد والتدمير، مما جعل الأنظمة الدفاعية التقليدية القائمة على القواعد الثابتة غير قادرة على مواجهة هذه التحديات الديناميكية. تسعى هذه الدراسة إلى استكشاف وتحليل الدور المحوري الذي يلعبه الذكاء الاصطناعي (AI) كأداة ثورية في تحقيق الأمن السيبراني وتعزيز الدفاعات الرقمية. تتطرق الدراسة في مبحثها الأول لتأصيل المفاهيم الأساسية، حيث تُعرّف الذكاء الاصطناعي باعتباره قدرة الأنظمة التقنية على محاكاة السلوك البشري والتعلم الذاتي، وتناقش الأمن السيبراني كمنظومة لحماية سرية وسلامة البيانات. كما تبرز الدراسة طبيعة العلاقة الجدلية بين المفهومين؛ فبينما يمثل الذكاء الاصطناعي درعاً واقياً للمؤسسات، فإنه يُستغل أيضاً من قبل المهاجمين لتطوير برمجيات خبيثة فائقة الذكاء، مما يجعله سلاحاً ذو حدين.

وفي المبحث الثاني، تعتمد الدراسة على تشريح التطبيقات العملية للذكاء الاصطناعي كدرع سيبراني؛ حيث تثبت قدرته الفائقة على الكشف الاستباقي عن الهجمات مجهولة المصدر (Zero-Day) عبر تحليل سلوك الشبكات، ومعالجة البيانات الضخمة (Big Data) بسرعة فائقة لربط المؤشرات الأمنية وتقليل الإنذارات الكاذبة. بالإضافة إلى ذلك، ترصد الدراسة دوره في تطوير نظم إدارة الهويات والوصول (IAM) من خلال المصادقة السلوكية المستمرة وحظر محاولات الهندسة الاجتماعية والتصيد الاحتيالي. وأخيراً، تفرد الدراسة مساحة واسعة في مبحثها الثالث للأبعاد القانونية والأخلاقية المحيطة بهذا الاستخدام؛ إذ تناقش تحديات حماية الخصوصية الرقمية (مثل الامتثال للتشريعات العالمية كـ GDPR)، وضمان

شفافية الخوارزميات ومكافحة التحيز الرقمي. كما تعالج المعضلة القانونية المعقدة المتعلقة بتحديد المسؤولية المدنية والجنائية عن أخطاء الذكاء الاصطناعي أو الأضرار الناشئة عن قراراته المستقلة، والتنازع بين مسؤولية المطور والمسؤولية الافتراضية للنظام الذكي.

وتخلص الدراسة إلى مجموعة من النتائج أبرزها: أن الذكاء الاصطناعي لم يعد خياراً ترفيهياً بل ضرورة حتمية لحماية الأمن القومي للمؤسسات والدول، كما أن الذكاء الاصطناعي يساهم في منع الهجمات قبل حدوثها من خلال تحليل سلوك المستخدمين والأنظمة لاكتشاف أي نشاط مشبوه، وأن نجاح هذا الدرع التكنولوجي مشروط ببناء منظومة حوكمة تشريعية وقانونية دولية تضمن استخدامه الآمن والمسؤول.

ومن أهم التوصيات التي انتهت إليها الدراسة: ضرورة التعاون على المستوى الدولي والإقليمي لتحسين نماذج الذكاء الاصطناعي لجعلها أكثر مقاومة للهجمات السيبرانية، وضرورة إنشاء منصة عالمية لتبادل المعلومات حول الهجمات الإلكترونية المستندة إلى الذكاء الاصطناعي.

المقدمة.

اعتادت الدول على مدار التاريخ أن تعتمد على العنصر البشري في تحقيق الأمن السيبراني، عبر جمع المعلومات وتشديد الرقابة ومطاردة التهديدات وسد الثغرات، والاستجابة لأي احتمالات بشأن وقوع حوادث سيبرانية، ولكن في العقد الأخير تطور الذكاء الاصطناعي بشكل سريع، محققاً فوائد كبيرة للدفاع السيبراني عبر أتمتة العناصر الأساسية للعمل (استخدام الحاسوب والأجهزة المعتمدة على المعالجات والبرمجيات)، وتحويله إلى عملية مبسطة ومستقلة تعمل على تسريع المعالجة وزيادة حماية الأمن القومي للدول⁽¹⁾.

وقد توسع العالم الرقمي ساهم في ارتفاع التهديدات السيبرانية لأمن المجتمعات، ومع التطور الهائل الذي شهده قطاع التكنولوجيا وفي مقدمته الذكاء الاصطناعي خلال السنوات الأخيرة، أصبحت العلاقة معقدة بين الذكاء الاصطناعي والأمن السيبراني، حيث يتم استغلال هذه التكنولوجيا الحديثة لاختراق المؤسسات الحكومية، سواء من قبل الجماعات المتطرفة أو دول أخرى أثناء فترات التوترات والنزاعات السياسية، ويمكن استخدام الذكاء الاصطناعي أيضاً كأداة متطورة

(¹) Russell, S., & Norvig, Artificial intelligence: A modern approach (4th ed.), 2021, pp1020:1050

لمواجهة هذه التهديدات السيبرانية، وزيادة قدرات الدول الأمنية والدفاعية، ما يحتم على الحكومات والمؤسسات حول العالم، اتباع استراتيجية دمج أنظمة الذكاء الاصطناعي مع القدرات البشرية، من أجل تحقيق استفادة أكبر من هذه التقنيات الحديثة في الأمن السيبراني⁽¹⁾، ومن أجل التعرض لذلك الموضوع فإننا سنظهر في البداية ما يلي:

* مشكلة الدراسة:

في ظل التطور المتسارع للهجمات الإلكترونية وتزايد تعقيدها، أصبحت الحلول الأمنية التقليدية عاجزة عن مواكبة هذه التهديدات. حيث تشير الإحصاءات إلى تزايد الهجمات السيبرانية بنسبة 125% خلال العامين الماضيين، مما يطرح تساؤلاً رئيسياً: كيف يمكن للذكاء الاصطناعي أن يساهم في سد هذه الفجوة الأمنية؟ تبرز المشكلة في الحاجة إلى فهم آليات توظيف تقنيات الذكاء الاصطناعي لتعزيز الأمن السيبراني، ومعرفة التحديات التي تواجه هذا التكامل.

* أهمية الدراسة:

تكتسب هذه الدراسة أهميتها من:

1. الجانب النظري: تطوير الإطار المعرفي حول التقنيات الحديثة في الأمن السيبراني.
2. الجانب التطبيقي: تقديم حلول عملية للجهات المعنية بالأمن السيبراني.
3. الجانب المجتمعي: حماية البيانات والممتلكات الرقمية للأفراد والمؤسسات.
4. الجانب الاقتصادي: تقليل الخسائر الناجمة عن الجرائم الإلكترونية.

* أهداف الدراسة:

1. تحليل آليات توظيف الذكاء الاصطناعي في تعزيز الأمن السيبراني.
2. تقييم فاعلية تقنيات التعلم الآلي في كشف التهديدات الأمنية.
3. تحديد التحديات والمعوقات التي تواجه تطبيق هذه التقنيات.
4. تقديم توصيات لتحسين أداء الأنظمة الأمنية المعتمدة على الذكاء الاصطناعي.

(¹)Chollet, F., Deep learning with Python (2nd ed., 2001 pp. 450-480)

*** منهج الدراسة:**

تعتمد الدراسة على المنهجين:

1. الوصفي التحليلي: لدراسة الأدبيات والنظريات ذات الصلة.
2. التطبيقي: من خلال تحليل نماذج عملية لأنظمة أمنية تستخدم الذكاء الاصطناعي.
3. المقارن: لمقارنة أداء الأنظمة التقليدية مع الأنظمة المعززة بالذكاء الاصطناعي.

*** حدود الدراسة:**

1. الحدود الموضوعية: التركيز على تطبيقات الذكاء الاصطناعي في مجال اكتشاف التهديدات السيبرانية.
2. الحدود الزمنية: دراسة التطورات في الفترة من 2020 إلى 2024.
3. الحدود المكانية: تحليل تجارب مختارة من دول رائدة في هذا المجال.
4. الحدود التقنية: اقتصار الدراسة على تقنيات التعلم الآلي والتعلم العميق.

*** إطار الدراسة النظري:**

تتطلب الدراسة من الإطار النظري الذي يربط بين:

1. نظريات الأمن السيبراني الحديثة.
2. مبادئ الذكاء الاصطناعي وتعلم الآلة.
3. نماذج إدارة المخاطر السيبرانية.
4. معايير الجودة والكفاءة في الأنظمة الأمنية.

المبحث الأول: مفهوم الذكاء الاصطناعي والأمن السيبراني**المطلب الأول: ماهية الذكاء الاصطناعي (Artificial Intelligence – AI)**

الذكاء الاصطناعي (Artificial Intelligence – AI) هو أحد أبرز التطورات التكنولوجية في القرن الحادي والعشرين، حيث يُعرف بأنه "قدرة الآلات على محاكاة الذكاء البشري في التعلم والتفكير واتخاذ القرارات". بدأ المفهوم كفكرة فلسفية قديمة، لكنه تحول إلى واقع ملموس مع تطور الحواسيب والخوارزميات. اليوم، يُستخدم الذكاء الاصطناعي في تشخيص الأمراض، وتحليل

البيانات الضخمة، وحتى قيادة السيارات ذاتية الحركة. وفقًا لتقرير شركة *PWC، سيُسهم الذكاء الاصطناعي بنحو **15.7 تريليون دولار* في الاقتصاد العالمي بحلول عام 2030. ومع ذلك، فإن هذا التقدم السريع يطرح تساؤلات حول الأخلاقيات، والخصوصية، ومستقبل العمل البشري⁽¹⁾

وبذلك فإن الذكاء الاصطناعي (Artificial Intelligence - AI) هو مجال واسع يهدف إلى إنشاء أنظمة قادرة على محاكاة الذكاء البشري من خلال أداء مهام مثل التعلم، التفكير المنطقي، حل المشكلات، والتفاعل مع البيئة المحيطة. يعتمد الذكاء الاصطناعي على عدة تقنيات رئيسية، منها:⁽²⁾

- *التعلم الآلي (Machine Learning)*: وهي تقنية تمكن الأنظمة من التعلم من البيانات دون برمجة صريحة⁽³⁾.
 - *الشبكات العصبية الاصطناعية (Artificial Neural Networks)*: وهي نماذج تحاكي عمل الدماغ البشري لمعالجة البيانات المعقدة.
 - *معالجة اللغة الطبيعية (Natural Language Processing - NLP)*: تمكن الأنظمة من فهم اللغة البشرية والتفاعل معها.
 - *الرؤية الحاسوبية (Computer Vision)*: تمكن الأنظمة من تحليل الصور والفيديو.
- هو فرع من فروع علوم الحاسوب يهدف إلى تطوير أنظمة قادرة على أداء مهام تتطلب عادةً ذكاءً بشرياً، مثل التعلم، التفكير، التحليل، واتخاذ القرارات. يعتمد الذكاء الاصطناعي على تقنيات متقدمة مثل التعلم الآلي (Machine Learning)، والشبكات العصبية (Neural Networks)، ومعالجة اللغة الطبيعية (Natural Language Processing).

⁽¹⁾ Wachter, S. (2021). AI and the Law: Ethical and Legal Challenges.

⁽²⁾ أحمد السيد النجار، الأمن السيبراني المتقدم باستخدام التعلم الآلي، دار العلم والإيمان، القاهرة، 2023 (280 صفحة)

⁽³⁾ Sarker, I. H, AI-driven cybersecurity and threat intelligence: Cyber automation, intelligent decision-making, and proactive defense, 2021

وتطور الذكاء الاصطناعي بشكل كبير في العقود الأخيرة بفضل التقدم في قوة الحوسبة وتوفر البيانات الضخمة (Big Data). اليوم، يتم استخدام الذكاء الاصطناعي في مجالات متنوعة مثل الطب، التمويل، النقل، والأمن السيبراني⁽¹⁾.

وقد قال *إيلون ماسك* عن "الذكاء الاصطناعي بأنه أخطر من الأسلحة النووية"، مما يؤكد ضرورة التوازن بين الابتكار والمسؤولية⁽²⁾.

هذا وقد يختلف تعريف الذكاء الاصطناعي (AI) بين الثقافات الغربية والشرقية بسبب الاختلافات في الخلفيات الفلسفية والتطبيقات العملية والأولويات المجتمعية. فيما يلي تحليل للمفاهيم الأساسية:

حيث عرف الاتحاد الأوروبي (في تقرير 2018) الذكاء الاصطناعي: بأنه "نظم تُظهر سلوكًا ذكيًا من خلال تحليل البيئة واتخاذ إجراءات لتحقيق أهداف محددة".

وقد عرفت منظمة التعاون الاقتصادي والتنمية (OECD) الذكاء الاصطناعي: بأنه "نظم تستخدم البيانات والنماذج الحسابية لتوليد تنبؤات أو توصيات أو قرارات تؤثر على البيئات المادية أو الرقمية"⁽³⁾.

ويعتبر التعريف الصيني (في "خطة الذكاء الاصطناعي الجديدة 2030"): بأنها التقنية التي تهدف إلى تعزيز التنمية الاجتماعية والاقتصادية عبر محاكاة الذكاء البشري، مع الحفاظ على القيم الأخلاقية.

(¹) – Zarsky, T. (2016). The Trouble with Algorithmic Decisions: An Analytic Road Map

(²) Kaplan, J., Artificial intelligence: What everyone needs to know, 2021. ,pp. 150–175.

(³) أحمد محمود الهاشمي، تأثير الذكاء الاصطناعي على تحسين استراتيجيات الأمن السيبراني في المؤسسات الحكومية، جامعة القاهرة، 2021 ص 25

وقد عرفه ستيوارت راسل وبيتر نورفيغ* في كتابهما الشهير "Artificial Intelligence: A Modern Approach" بأنه "دراسة كيفية جعل الحواسيب تقوم بما يقوم به البشر بشكل أفضل حاليًا"⁽¹⁾.

وقد عرفه ستيوارت راسل وبيتر نورفيغ* (في كتاب "Artificial Intelligence: A Modern Approach"): بأنه "دراسة كيفية جعل الحواسيب تقوم بما يقوم به البشر بشكل أفضل حاليًا". وقد عرفه جون مكارثي* (مؤسس مصطلح "الذكاء الاصطناعي" في مؤتمر دارتموث 1956): بأنه علم وهندسة صنع آلات ذكية، وخاصة برامج الكمبيوتر الذكية.

التطور التاريخي

- 1950: طرح *آلان تورينج* اختبارًا لقياس ذكاء الآلة ("اختبار تورينج")⁽²⁾.
- 1956: انعقاد مؤتمر *دارتموث*، حيث صُك مصطلح "الذكاء الاصطناعي" رسميًا.
- 1997: هزيمة بطل العالم في الشطرنج *غاري كاسباروف* أمام الحاسوب *ديب بلو*.
- 2011: انتصار نظام *واتسون* من IBM في مسابقة *Jeopardy*!
- 2023: ظهور نماذج مثل *ChatGPT* و *DALL-E* التي أحدثت ثورة في معالجة اللغة والصور⁽³⁾.

وبذلك فإن الذكاء الاصطناعي ليس مجرد تكنولوجيا، بل هو تحول جذري في طريقة تفاعل البشر مع الآلات. بينما يقدم حلولًا غير مسبقة لمشكلات مثل التغير المناخي والأمراض المزمنة، فإن إدارة آثاره الأخلاقية والاجتماعية تتطلب تعاونًا عالميًا وتشريعات مرنة⁽⁴⁾.

⁽¹⁾ كتاب Le droit face à l'intelligence artificielle – جان-لوك سوريل

⁽²⁾ محمد سمير الشرقاوي، حماية أنظمة إنترنت الأشياء باستخدام تقنيات الذكاء الاصطناعي، جامعة عين شمس، 2022 ص 55

⁽³⁾ Cole, E, Advanced persistent threat: Understanding the danger and how to protect your organization, 2020. , pp. 200–225.

⁽⁴⁾ سارة إبراهيم العلي، تصميم أنظمة استجابة آلية للهجمات السيبرانية باستخدام التعلم المعزز ، جامعة المنصورة، 2023

المطلب الثاني، ماهية الأمن السيبراني

الأمن السيبراني (Cybersecurity) يشير إلى مجموعة من الممارسات والتقنيات المستخدمة لحماية الأنظمة والشبكات والبيانات من الهجمات الإلكترونية⁽¹⁾. يشمل ذلك حماية المعلومات السرية، منع الوصول غير المصرح به، والكشف عن التهديدات الأمنية والاستجابة لها⁽²⁾.

فالأمن السيبراني (Cybersecurity) هو مجموعة من الممارسات والتقنيات المستخدمة لحماية الأنظمة الحاسوبية، الشبكات، والبيانات من الهجمات الإلكترونية. يشمل ذلك⁽³⁾:

- حماية البيانات*: ضمان سرية وسلامة المعلومات.
- إدارة الهويات والوصول*: التحكم فيمن يمكنه الوصول إلى الأنظمة والبيانات.
- الكشف عن التهديدات*: تحديد الهجمات الإلكترونية في مراحلها المبكرة.
- الاستجابة للحوادث*: اتخاذ إجراءات سريعة لتقليل الأضرار الناتجة عن الهجمات.

حيث أنه مع تزايد الاعتماد على التكنولوجيا في جميع جوانب الحياة، أصبحت الحاجة إلى الأمن السيبراني أكثر إلحاحًا، وتشمل التهديدات السيبرانية الشائعة: البرمجيات الخبيثة (Malware)، هجمات التصيد (Phishing)، وهجمات الحرمان من الخدمة (DDoS)⁽⁴⁾.

وتواجه أنظمة الأمن السيبراني الحديثة العديد من التحديات، منها⁽⁵⁾:

1. تعقيد الهجمات الإلكترونية: أصبحت الهجمات أكثر تطورًا وتنظيمًا، مما يجعل اكتشافها أمرًا صعبًا.

⁽¹⁾ كتاب "الأمن السيبراني: المفاهيم والتطبيقات" لمؤلف عربي.

⁽²⁾ Russell, S., & Norvig, P. (2020). Artificial Intelligence: A Modern Approach

⁽³⁾ Sikorski, M., & Honig, A., Practical malware analysis: The hands-on guide to dissecting malicious software (2nd ed , pp. 300–330.

⁽⁴⁾ Schneier, B. (2015). Data and Goliath: The Hidden Battles to Collect Your Data and Control Your World

⁽⁵⁾ Leclerc, J. (2019). L'intelligence artificielle et la cybersécurité

2. حجم البيانات الضخمة: مع زيادة كمية البيانات التي يتم توليدها، يصبح تحليلها واكتشاف التهديدات تحديًا كبيرًا.
3. نقص الكوادر المؤهلة: هناك نقص في الخبراء المؤهلين في مجال الأمن السيبراني.
4. التحديات القانونية والأخلاقية: مثل حماية الخصوصية وضمان الامتثال للقوانين المحلية والدولية⁽¹⁾.

المطلب الثالث، العلاقة بين الذكاء الاصطناعي والأمن السيبراني

- يعد الذكاء الاصطناعي أداة قوية لتعزيز الأمن السيبراني، حيث يمكنه تحليل كميات هائلة من البيانات بسرعة وكفاءة، مما يساعد في الكشف عن التهديدات والاستجابة لها في الوقت الفعلي. على سبيل المثال، يمكن لأنظمة الذكاء الاصطناعي⁽²⁾:
1. *الكشف عن التهديدات*:
 - استخدام التعلم الآلي للكشف عن الأنماط غير الطبيعية في حركة البيانات.
 - مثال: أنظمة مثل IBM QRadar تستخدم الذكاء الاصطناعي لتحليل البيانات الأمنية.
 2. *تحليل البيانات الضخمة*:
 - تحليل البيانات الضخمة لاكتشاف التهديدات المخفية.
 - مثال: استخدام تقنيات الذكاء الاصطناعي في مراكز عمليات الأمن (SOC).
 3. *إدارة الهويات والوصول*⁽¹⁾:

⁽¹⁾Anderson, H. S., & Roth, P. (2018). Ember: An open dataset for training static PE malware machine learning models. ArXiv Preprint.

<https://arxiv.org/abs/1804.04637>

⁽²⁾Rapport ANSSI (Agence Nationale de la Sécurité des Systèmes d'Information) sur la cybersécurité

- استخدام الذكاء الاصطناعي للتحقق من هويات المستخدمين ومنع الوصول غير المصرح به.
- مثال: أنظمة التعرف على الوجه (Facial Recognition) المستخدمة في التحقق من الهوية.
- 4. *أتمتة الاستجابة للحوادث*⁽²⁾:
- أتمتة العمليات مثل عزل الأجهزة المصابة أو إغلاق الثغرات الأمنية تلقائيًا.

التحديات المرتبطة باستخدام الذكاء الاصطناعي في الأمن السيبراني:

- *إمكانية استغلال الذكاء الاصطناعي من قبل المهاجمين* : مثل استخدام الذكاء الاصطناعي لتنفيذ هجمات أكثر تعقيدًا⁽³⁾.
- *التحيز في الخوارزميات* : قد تؤدي البيانات المتحيزة إلى قرارات أمنية غير دقيقة.
- *نقص الشفافية* : صعوبة فهم كيفية اتخاذ أنظمة الذكاء الاصطناعي لقراراتها.

أنواع التهديدات السيبرانية الشائعة:

1. *البرمجيات الخبيثة (Malware)* : مثل الفيروسات وبرامج الفدية (Ransomware).

⁽¹⁾ إيمان محمود عبد الرحمن، حماية الأنظمة الذكية من الهجمات الإلكترونية، المركز القومي للبحوث، القاهرة، 2022 (320 صفحة)

⁽²⁾ Buczak, A. L., & Guven, E. (2016). A survey of data mining and machine learning methods for cybersecurity intrusion detection. IEEE.

<https://doi.org/10.1109/COMST.2015.2494502>

⁽³⁾ - *دراسة حالة: هجوم *SolarWinds 2020* الذي اخترق وكالات حكومية أمريكية. لو استُخدم الذكاء الاصطناعي لتحليل سلوك الوصول غير المعتاد، لكان من الممكن كشف الهجوم مبكرًا.

2. *هجمات التصيد (Phishing)*: محاولات خداع المستخدمين للحصول على معلومات حساسة.
 3. *هجمات الحرمان من الخدمة (DDoS)*: إغراق الخوادم بطلبات وهمية لإيقافها عن العمل.
 4. *استغلال الثغرات الأمنية*: مثل ثغرات Zero-Day التي يتم استغلالها قبل اكتشافها من قبل المطورين⁽¹⁾.
- حيث تواجه أنظمة الأمن السيبراني الحديثة العديد من التحديات التي تتطلب حلولاً مبتكرة، منها:
1. *تعدد الهجمات الإلكترونية*: أصبحت الهجمات أكثر تطوراً باستخدام تقنيات مثل الذكاء الاصطناعي من قبل المهاجمين.
 2. *حجم البيانات الضخمة*: مع تزايد كمية البيانات، يصبح تحليلها واكتشاف التهديدات أمراً معقداً.
 3. *نقص الكوادر المؤهلة*⁽²⁾: هناك حاجة ماسة إلى خبراء في الأمن السيبراني لمواجهة التهديدات المتزايدة.
 4. *التحديات القانونية والأخلاقية*: مثل حماية الخصوصية وضمان الامتثال للقوانين الدولية مثل GDPR (اللائحة العامة لحماية البيانات) في الاتحاد الأوروبي.

المبحث الثاني

الذكاء الاصطناعي كدرع سيبراني.

المطلب الأول: الكشف عن التهديدات والهجمات الإلكترونية

يُعد الكشف عن التهديدات الإلكترونية أحد أهم تطبيقات الذكاء الاصطناعي في مجال الأمن السيبراني. تعتمد الأنظمة التقليدية على قواعد محددة مسبقاً (Rule-Based Systems)

⁽¹⁾جمال مصطفى أحمد، الذكاء الاصطناعي في مواجهة الجرائم الرقمية، دار المعارف العلمية،

الإسكندرية، 2021

⁽²⁾IBM Security. (2023). Cost of a data breach report 2023. IBM.

<https://www.ibm.com/reports/data-breach>

للكشف عن التهديدات، ولكن هذه الأنظمة غالبًا ما تكون محدودة في مواجهة الهجمات المتطورة (1).

وتمثل البنية التحتية الحيوية (مثل شبكات الطاقة، والاتصالات، والأنظمة المالية) أهدافًا رئيسية للهجمات الإلكترونية التي تهدد الأمن القومي، وهنا يلعب الذكاء الاصطناعي دورًا حاسمًا من خلال :

- الدفاع النشط (Active Defense)*: استخدام الذكاء الاصطناعي لإنشاء "أفخاخ" (Honeypots) تخدع المهاجمين وجمع معلومات عن تكتيكاتهم (2).
- الرد التلقائي: أنظمة قادرة على عكس هجمات DDoS تلقائيًا عن طريق إغراق المهاجم ببيانات وهمية.
- الكشف عن الهجمات ذات الدوافع الجيوسياسية: مثل الهجمات المنظمة من قبل دول أو جماعات إرهابية لتعطيل الخدمات الأساسية (3).

ويعمل الذكاء الاصطناعي في الكشف عن التهديدات من خلال (4) :

- تحليل الأنماط: يستخدم الذكاء الاصطناعي تقنيات التعلم الآلي لتحليل أنماط البيانات وتحديد السلوكيات غير الطبيعية (5).
- التعلم غير الخاضع للإشراف (Unsupervised Learning): يمكن للأنظمة اكتشاف تهديدات جديدة دون الحاجة إلى بيانات تدريب مسبق.

(1) - Chapple, M., & Seidl, D. (2021). Cybersecurity Essentials

(2) حسام الدين محمد علي، تقنيات التعلم العميق لاكتشاف البرمجيات الخبيثة، مكتبة الأنجلو المصرية، القاهرة، 2023 (410 صفحة)

(3) مثال: هجمات روسيا على شبكة الطاقة الأوكرانية عام (2015) باستخدام برامج الفدية.

(4) دراسة حالة لذلك: هجوم **SolarWinds 2020* الذي اخترق وكالات حكومية أمريكية. لو استُخدم الذكاء الاصطناعي لتحليل سلوك الوصول غير المعتاد، لكان من الممكن كشف الهجوم مبكرًا.

(5) خالد طه حسين، أمن الشبكات العصبية الاصطناعية، دار النهضة التكنولوجية، القاهرة، 2022 ، (360 صفحة)

- التحليل السلوكي: مراقبة سلوك المستخدمين والأجهزة للكشف عن الأنشطة المشبوهة⁽¹⁾.

ويُلاحظ أن المهاجمون يستخدمون تقنيات الذكاء الاصطناعي لخداع أنظمة الأمن⁽²⁾.

المطلب الثاني، تحليل البيانات الضخمة «Big Data»

يُلاحظ أنه مع تزايد كمية البيانات التي يتم توليدها يوميًا، أصبح تحليلها يدويًا أمرًا مستحيلًا. هنا يأتي دور الذكاء الاصطناعي في تحليل البيانات الضخمة لاكتشاف التهديدات المخفية، وذلك يكون من خلال⁽³⁾:

- معالجة البيانات في الوقت الفعلي: يمكن للذكاء الاصطناعي تحليل تدفقات البيانات الضخمة بسرعة فائقة.

- اكتشاف الأنماط الخفية: مثل الروابط بين الهجمات الإلكترونية التي تبدو غير مرتبطة⁽⁴⁾.

- تقليل الضوضاء: تصفية البيانات غير المهمة والتركيز على المؤشرات الأمنية الحرجة.

وتعتمد أنظمة الأمن السيبراني الحديثة على تقنيات التعرف على الأنماط والسلوكيات المشبوهة (Pattern Recognition) لتحديد التهديدات. يستخدم الذكاء الاصطناعي هنا لتحسين دقة وكفاءة هذه الأنظمة.

⁽¹⁾ وتتمثل الأمثلة الواقعية على ذلك أنظمة:

- *Darktrace*: تستخدم الذكاء الاصطناعي للكشف عن التهديدات في الوقت الفعلي بناءً على تحليل سلوك الشبكة.

- *IBM QRadar*: يعتمد على الذكاء الاصطناعي لتحليل البيانات الأمنية وتحديد التهديدات المحتملة.

⁽²⁾ وهو ما يُسمى الهجمات المضادة للذكاء الاصطناعي (Adversarial AI).

⁽³⁾ Dupont, B. (2020). La cybersécurité à l'ère de l'intelligence artificielle.

⁽⁴⁾ دعاء سعيد محمد، الحماية الذكية للبيانات الشخصية، دار الكتب الجامعية، الإسكندرية، 2023، 240 صفحة.

- ويمكن التعرف على الأنماط باستخدام الذكاء الاصطناعي⁽¹⁾:
- *الشبكات العصبية الاصطناعية*: تحليل البيانات المعقدة مثل الصور والفيديو للكشف عن التهديدات.
 - *التعلم العميق (Deep Learning)*: تحسين دقة الكشف عن التهديدات من خلال تحليل طبقات متعددة من البيانات⁽²⁾.
 - *تحليل السلوكيات*: مثل مراقبة سلوك المستخدمين للكشف عن الاختراقات⁽³⁾.
- وهنا يكون دور الذكاء الاصطناعي من خلال:
- مراقبة أنظمة SCADA (التي تدير البنى التحتية الحيوية) لاكتشاف الانحرافات عن العمليات الطبيعية⁽⁴⁾.
- وتتمثل الأمثلة الواقعية على ذلك:
- أنظمة التعرف على الوجه (Facial Recognition): تستخدم في التحقق من الهوية ومنع الوصول غير المصرح به.

⁽¹⁾ حيث أنه من الآليات التقنية لحماية الذكاء الاصطناعي:

- التعلم العميق (Deep Learning) في اكتشاف التهديدات الخفية
- الشبكات العصبية التلافيفية (CNNs): تُستخدم لتحليل حركة المرور على الشبكة وتحديد الأنماط غير الطبيعية، مثل حزم البيانات المشبوهة في هجمات DDos.
- الشبكات العصبية المتكررة (RNNs): تحليل تسلسلات الزمن الحقيقي (Real-Time Sequences) للتنبؤ بالهجمات قبل حدوثها.
- مثال: نظام Darktrace** يستخدم خوارزميات التعلم غير الخاضع للإشراف لاكتشاف التهديدات الجديدة دون الحاجة إلى قاعدة بيانات مسبقة.

⁽²⁾ Zarsky, T. (2016). The Trouble with Algorithmic Decisions: An Analytic Road Map.

⁽³⁾ سامح عبد الفتاح، التحليل الآمن للبيانات الكبيرة، الهيئة المصرية العامة للكتاب، 2021 (380 صفحة)

⁽⁴⁾ مثال: استخدام منصة *Microsoft Azure Defender for IoT* لتحليل بيانات أجهزة إنترنت الأشياء (IoT) الصناعية.

- تحليل حركة المرور على الشبكة: للكشف عن الهجمات مثل هجمات DDoS.

المطلب الثالث الذكاء الاصطناعي في إدارة الهويات والوصول

إدارة الهويات والوصول (Identity and Access Management – IAM) هي عملية التحكم فيمن يمكنه الوصول إلى الأنظمة والبيانات. يستخدم الذكاء الاصطناعي هنا لتحسين عمليات التحقق من الهوية وإدارة الصلاحيات⁽¹⁾.

ويعمل الذكاء الاصطناعي في إدارة الهويات من خلال⁽²⁾:

- التحقق البيومتري: مثل التعرف على الوجه أو بصمة الأصبع باستخدام الذكاء الاصطناعي.
- تحليل السلوك: مراقبة أنشطة المستخدمين للكشف عن الوصول غير المصرح به.
- أتمتة إدارة الصلاحيات: مثل منح أو إلغاء الصلاحيات تلقائيًا بناءً على تحليل الذكاء الاصطناعي.

وتُعرف أتمتة الاستجابة للحوادث الأمنية (Security Orchestration, Automation, and Response – SOAR) بأنها عملية استخدام الذكاء الاصطناعي لأتمتة مهام مثل الكشف عن التهديدات والاستجابة لها.

ويمكن أتمتة الاستجابة باستخدام الذكاء الاصطناعي عن طريق⁽³⁾:

- الكشف التلقائي: تحديد التهديدات دون تدخل بشري.
- الاستجابة الفورية: مثل عزل الأجهزة المصابة أو إغلاق الثغرات الأمنية.
- تحسين العمليات: تقليل الوقت اللازم للاستجابة للحوادث الأمنية.

⁽¹⁾ شيماء محمود عثمان، التوجهات الحديثة في أمن السحابة الإلكترونية، دار الفكر العربي، القاهرة،

2024 (290 صفحة)

⁽²⁾ Rapport ENISA (European Union Agency for Cybersecurity) sur l'IA et la cybersécurité.

⁽³⁾ Brundage, M. (2020). Toward Trustworthy AI Development: Mechanisms for Supporting Verifiable Claims

ويُعد الذكاء الاصطناعي التكيفي (Adaptive AI) وسيلة يُمكن من خلالها حماية الأمن السيبراني وذلك من خلال (1) :

- التعلم المعزز (Reinforcement Learning): تتعلم الأنظمة من التفاعلات مع البيئة لتحسين استراتيجيات الدفاع (2).

ومع ذلك لا يجب الاعتماد المفروض على الذكاء الاصطناعي في حماية الأمن القومي إذ إن الأتمتة الكاملة من شأنها فقدان السيطرة البشرية على القرارات الأمنية الحرجة (3)، وهو ما يُسمى إشكالية التحيز الخوارزمي (Algorithmic Bias)، حيث قد تُظهر أنظمة الذكاء الاصطناعي تحيزًا ضد مجموعات معينة بسبب بيانات تدريب غير متنوعة، وهو ما من شأنه التأثير على الأمن القومي إذ قد يؤدي التحيز إلى استهداف غير دقيق لتهديدات أمنية، مما يزيد التوترات الاجتماعية أو السياسية (4).

وهو ما يُظهر التحديات والمخاطر المرتبطة بالذكاء الاصطناعي في الأمن السيبراني (Adversarial AI): مثل الهجمات التخريبية كتلاعب المهاجمين ببيانات التدريب لخداع أنظمة الذكاء الاصطناعي (5)

وهجمات التزييف العمق (Deepfake) * لخداع أنظمة التحقق من الهوية (6).

(1) National Institute of Standards and Technology (NIST). (2021). AI for cybersecurity. <https://www.nist.gov/itl/ai-cybersecurity>

(2) مثال: أنظمة الدفاع السيبراني التي تتكيف مع تكتيكات المهاجمين المتغيرة، مثل تغيير قواعد الجدار الناري (Firewall) تلقائيًا.

(3) مثال ذلك إذا قرر نظام ذكاء اصطناعي عزل شبكة كهرباء كاملة بسبب إنذار خاطئ، قد يؤدي ذلك إلى انهيار خدمات حيوية.

(4) - مثال على ذلك أنظمة التعرف على الوجه التي تفشل في التعرف على الأفراد ذوي البشرة الداكنة.

(5) (مثل تغيير بكسلات الصور لتضليل أنظمة التعرف على الصور).

(6) ويُمكن التصدي لذلك من خلال تطوير خوارزميات مقاومة للتلاعب (Robust AI) عبر تعريضها لهجمات محاكاة خلال مرحلة التدريب.

المبحث الثالث

الأطر القانونية والأخلاقية لاستخدام الذكاء الاصطناعي، والمسئولية عنها.

المطلب الأول: الأطر القانونية والأخلاقية لاستخدام الذكاء الاصطناعي
التشريعات الدولية:

ونحن في إطار الحديث عن التشريعات الدولية بشأن استخدام الذكاء الاصطناعي في تحقيق الأمن السيبراني فإنه لابد الإشارة إلى:

- اتفاقية بودابست للجرائم الإلكترونية 2001:

حيث تركّز على مكافحة الجرائم الإلكترونية التقليدية، لكنها لا تغطي الهجمات التي تستخدم الذكاء الاصطناعي⁽¹⁾.

وبذلك فهي بحاجة إلى إضافة بروتوكولات خاصة بالذكاء الاصطناعي في النسخة المعدلة (2023)، لتشمل مسؤولية استخدام الذكاء الاصطناعي في الهجمات⁽²⁾.

- مبادئ OECD لأخلاقيات الذكاء الاصطناعي: تشدد على الشفافية والمساءلة في الأنظمة الأمنية⁽³⁾.

- اللائحة العامة لحماية البيانات (GDPR) في الاتحاد الأوروبي⁽⁴⁾:

حيث تُلزم الشركات بضمان حماية البيانات الشخصية وتعزيز الشفافية في استخدام الذكاء الاصطناعي.

- *مبادئ OECD لأخلاقيات الذكاء الاصطناعي*:

- تشمل الشفافية، المساءلة، واحترام حقوق الإنسان⁽¹⁾.

(¹) علياء محمد كامل، أنظمة كشف التسلل المعززة بالذكاء الاصطناعي، مركز الكتاب الأكاديمي، القاهرة، 2022 (330 صفحة)

(²) الحرب السيبرانية في عصر الذكاء الاصطناعي - مركز المستقبل للأبحاث.

(³) Brundage, M. (2020). Toward Trustworthy AI Development: Mechanisms for Supporting Verifiable Claims

(⁴) حيث تحظر المادة 22 اتخاذ قرارات مؤثرة على الأفراد دون تدخل بشري.

- *مثال*: استخدام المبادئ في تقييم أنظمة المراقبة الجماعية في الصين.

التشريعات المحلية في الدول العربية:

- جمهورية مصر العربية

• استراتيجية "مصر الرقمية 2030"

حيث أطلقت مصر استراتيجية "مصر الرقمية 2030" لتحويل الخدمات الحكومية إلى منصات إلكترونية، مما أدى إلى زيادة الاعتماد على البنية التحتية الرقمية، ووفقاً لتقرير الجهاز القومي لتنظيم الاتصالات (2022)، تعرضت مصر لأكثر من *1.5 مليون* هجمة إلكترونية* في عام 2022، مما دفع الحكومة إلى تبني تقنيات الذكاء الاصطناعي لتعزيز الأمن السيبراني. وقد أطلقت مصر مشروع "الدرع السيبراني الوطني بالتعاون بين وزارة الاتصالات وتكنولوجيا المعلومات وشركة *IBM* (2)

وتهدف تلك الشراكة إلى حماية الشبكات الحكومية من الهجمات الإلكترونية المتطورة، وتحليل تهديدات مثل هجمات *DDoS* و*البرمجيات الخبيثة*.

- *التقنيات المستخدمة*:

- *التعلم الآلي (Machine Learning)*: لتحليل أنماط حركة البيانات والكشف عن السلوكيات غير الطبيعية.

- *الشبكات العصبية الاصطناعية (ANNs)*: لتحليل البيانات الضخمة في الوقت الفعلي.

(¹) محمد نبيل السيد، التهديدات السيبرانية في عصر إنترنت الأشياء، دار العلوم الحديثة، القاهرة، 2023 (270 صفحة)

(²) هي شركة استشارية وتقنية أمريكية متعددة الجنسيات يقع مقرها الرئيسي في أرمونك، نيويورك، ويعمل بها أكثر من 350.000 موظف يخدمون عملاء في 170 دولة، وفي 8 أكتوبر 2020، أعلنت شركة آي بي إم أنها بصدد تحويل وحدة خدمات البنية التحتية المُدارة التابعة لقسم خدمات التكنولوجيا العالمية إلى شركة عامة جديدة، وهو إجراء من المتوقع أن يكتمل بحلول نهاية عام 2022، وتنتج شركة آي بي إم وتبيع أجهزة الحاسوب والبرمجيات الوسيطة والبرامج، وتوفر خدمات الاستضافة والاستشارات في مجالات تتراوح من أجهزة الحاسوب

- *النتائج*:
- خفض وقت الكشف عن التهديدات من *ساعات إلى دقائق*.
- منع أكثر من *500,000* هجمة إلكترونية* في عام 2022.
- البنك المركزي المصري: مكافحة الاحتيال المالي*
 - *التحدي*: زيادة حالات الاحتيال المالي عبر التحويلات الإلكترونية والبطاقات الائتمانية.
 - *الحل*: استخدام أنظمة ذكاء اصطناعي لتحليل سلوكيات المعاملات المالية.
 - *التقنيات المستخدمة*:
 - *تحليل السلوكيات (Behavioral Analytics)*: لتحديد المعاملات غير الاعتيادية.
 - *التعلم العميق (Deep Learning)*: للتنبؤ بالأنشطة الاحتيالية بناءً على بيانات تاريخية.
 - *النتائج*:
 - انخفاض حالات الاحتيال بنسبة *30%* خلال عام 2022.
 - تحسين تجربة العملاء من خلال تقليل الإنذارات الكاذبة (False Positives).
- حماية البنية التحتية الحيوية*
 - *السياق*: تعرضت شبكات الكهرباء والمياه لهجمات إلكترونية تهدف إلى تعطيل الخدمات.
 - *الحلول*:
 - استخدام أنظمة ذكاء اصطناعي لمراقبة أنظمة *SCADA* (التي تدير البنى التحتية الحيوية).
 - تطوير خوارزميات للكشف عن محاولات الاختراق في الوقت الفعلي.
 - *النتائج*:
 - منع *3* هجمات كبرى* على شبكات الكهرباء في عام 2023.
 - تقليل وقت الاستجابة للحوادث الأمنية بنسبة *50%*.

ومن ضمن النجاحات لدى مصر أيضًا *إنشاء المركز الوطني للأمن السيبراني* الذي يهدف إلى تنسيق الجهود الوطنية في مجال الأمن السيبراني وتبني تقنيات الذكاء الاصطناعي، وكذلك تدشين منصة "أمان وهي منصة إلكترونية تقدم خدمات الأمن السيبراني للشركات الصغيرة والمتوسطة باستخدام الذكاء الاصطناعي، فضلاً عن تعزيز التعاون الدولي كالشراكة مع منظمات مثل *الاتحاد الدولي للاتصالات (ITU)* و *مايكروسوفت* لتبادل الخبرات والتقنيات.

المملكة الأردنية الهاشمية:

حيث استخدمت المملكة الأردنية الهاشمية: الذكاء الاصطناعي في مواجهة الهجمات الإلكترونية على القطاع الصحي وذلك من خلال جائحة كوفيد-19، تعرضت أنظمة المستشفيات الأردنية لهجمات إلكترونية لتعطيل الخدمات⁽¹⁾.

وقد قامت المملكة باستخدام:

1. *منصة "سايبير درايفن" (CyberDriven)*:

- منصة أردنية ناشئة تستخدم الذكاء الاصطناعي لأتمتة الاستجابة للهجمات الإلكترونية.

- نجحت في تقليل وقت الاستجابة لاختراق مستشفى الملك عبد الله بنسبة 40%.

2. *مركز الأمن السيبراني الأردني*:

- استخدام الذكاء الاصطناعي لتحليل الهجمات الإلكترونية على البنية التحتية الحيوية، مثل شبكات المياه والكهرباء.

- الإمارات العربية المتحدة:

1. القانون الاتحادي رقم 12 لسنة 2016 (قانون مكافحة الجرائم الإلكترونية): يعاقب على

الاختراقات الإلكترونية، لكنه لا يُنظم استخدام الذكاء الاصطناعي بشكل صريح.

2. مبادرة دبي للذكاء الاصطناعي*: تهدف إلى جعل دبي مدينة ذكية مع ضمان الامتثال الأخلاقي.

- المملكة العربية السعودية:

⁽¹⁾ تقرير مركز الأمن السيبراني الأردني (2023): "تطبيقات الذكاء الاصطناعي في القطاع الصحي".

1. النظام السعودي للأمن السيبراني (2018): يركز على حماية البنية التحتية الحيوية، لكنه لا يتضمن تفاصيل عن الذكاء الاصطناعي⁽¹⁾.

2. رؤية 2030*: تشجيع الابتكار التكنولوجي مع تطوير أطر قانونية جديدة.

التوازن بين الأمن والخصوصية:

لا شك أن جمع البيانات الضخمة عن طريق أنظمة الذكاء الاصطناعي تحتاج إلى كميات هائلة من البيانات، مما يزيد مخاطر انتهاك الخصوصية.

كما في برامج الذكاء الاصطناعي ذات التعريفات الحيوية (Biometric Data) حيث أن التعرف على الوجه أو البصمة يثير مخاطر تسريب البيانات الحساسة (2). كما أن المراقبة الجماعية باستخدام الذكاء الاصطناعي لتحليل بيانات المواطنين قد ينتهك حقوق الخصوصية.

ويلاحظ بأنه حتى كون المادة 12 من الإعلان العالمي لحقوق الإنسان تحمي الخصوصية، لكن تطبيقها في العصر الرقمي غير كافٍ.

ومن ضمن الحلول تقنية المقترحة في هذا الشأن:

- *التعلم الموحد (Federated Learning)*: تدريب النماذج دون مشاركة البيانات الخام.
- *التشفير المتجانس (Homomorphic Encryption)*: معالجة البيانات المشفرة دون فك تشفيرها.

- *القانون الأوروبي المقترح للذكاء الاصطناعي (AI Act, 2024)*: حيث يصنف أنظمة الذكاء الاصطناعي حسب مستوى الخطورة⁽³⁾ ويلزم المطورين بإثبات امتثال أنظمتهم للمعايير الأخلاقية قبل التوزيع⁽⁴⁾.

⁽¹⁾ حيث نجحت السعودية في منع 85% من هجمات الفدية خلال عام 2022 بفضل أنظمة الذكاء الاصطناعي.

⁽²⁾ - *مثال: انتقادات نظام **Clearview AI* في الولايات المتحدة لاستخدامه صوراً من وسائل التواصل الاجتماعي دون موافقة.

⁽³⁾ (مثل أنظمة المراقبة الجماعية عالية الخطورة).

⁽⁴⁾ كتاب "الجوانب القانونية للذكاء الاصطناعي" - د. علياء عبدالرحمن.

المطلب الثاني، المسؤولية القانونية عن استخدام الذكاء الاصطناعي

تتمثل السيناريوهات المحتملة، والأسئلة الواجبة في هذا الشأن في حالة:

1. خطأ في الخوارزمية: من المسؤول؟ المطور، الشركة، أم المستخدم؟
 2. هجمات Adversarial AI: إذا استُخدم الذكاء الاصطناعي للاختراق، هل تُحمّل الشركة المُطوّرة المسؤولية⁽¹⁾؟
 3. كما أنه إذا فشل النظام في كشف الهجوم، هل تُقاضى الشركة المُزوّدة؟
- حيث أنه ونحن في إطار الشفافية والمساءلة (Transparency & Accountability) بشأن الإجابة على تلك التساؤلات فإننا نواجه معضلة الصندوق الأسود (Black Box) أي صعوبة فهم كيفية اتخاذ أنظمة الذكاء الاصطناعي لقراراتها، خاصة في الشبكات العصبية العميقة⁽²⁾. وللإجابة على تلك التساؤلات فإنه يمكن تقسيم هذه المسؤولية إلى ثلاثة مستويات رئيسية:
- أ. مسؤولية مطوري الأنظمة*:

- تتركز على مدى توفر عنصري الخطأ والضرر في تصميم الخوارزميات
 - التزام المطورين بمعايير الأمان السيبراني خلال مرحلة التصميم
 - حالات المساءلة عن الثغرات الأمنية غير المكتشفة
- ب. مسؤولية مشغلي الأنظمة*:
- واجب المراقبة المستمرة لأداء الأنظمة الأمنية
 - مسؤولية تحديث أنظمة الحماية بشكل دوري
 - الإبلاغ الفوري عن أي اختراقات أو أعطال

(¹) MITRE. (2022). Adversarial machine learning: A taxonomy and terminology of attacks and mitigations (MITRE Report No. 22-0001).

<https://www.mitre.org/publications>

(²) حيث أنه من ضمن الحلول مقترحة في هذا الشأن:

- XAI (الذكاء الاصطناعي القابل للتفسير)*: تطوير خوارزميات توفر تفسيرات لقراراتها.

- إطار GOVERN AI*: مبادئ لضمان الشفافية في الأنظمة الأمنية

ج. مسؤولية المستخدمين النهائيين:

- الالتزام بإجراءات الأمان الأساسية
- المساءلة عن سوء الاستخدام أو التجاوزات
- حالات الإهمال في تطبيق السياسات الأمنية
- أما أبرز التحديات القانونية في هذا المجال فتتمثل في:
- صعوبة تحديد مصدر الهجمات الإلكترونية المعقدة
- إشكالية إثبات العلاقة السببية في البيانات الرقمية
- تنازع القوانين في الحوادث العابرة للحدود
- غياب المعايير الموحدة لتقييم أداء الأنظمة الأمنية

للتغلب على هذه التحديات، يمكن

1. تطوير أطر تنظيمية خاصة بذكاء الأمن السيبراني
2. إنشاء هيئات رقابية متخصصة
3. اعتماد معايير دولية لتقييم المسؤولية
4. تطوير أنظمة تتبع ذكية لقرارات الأمن السيبراني

مع ضرورة تحقيق التوازن بين:

- متطلبات الأمن السيبراني الفعال
- حقوق المستخدمين وخصوصيتهم
- تشجيع الابتكار التقني
- حماية المصالح الوطنية

الخاتمة:

في ختام هذا البحث، يتضح أن الذكاء الاصطناعي قد أصبح ركيزة أساسية في تعزيز الأمن السيبراني، حيث يوفر أدوات متطورة قادرة على مواجهة التهديدات الإلكترونية المتنامية بسرعة

وكفاءة. فقد أظهرت الدراسة كيف يمكن لتقنيات مثل التعلم الآلي ومعالجة اللغة الطبيعية وتحليل البيانات الضخمة أن تحدث تحولًا جذريًا في كشف التهديدات، والاستجابة للحوادث الأمنية، ومنع الهجمات قبل وقوعها.

ومع ذلك، فإن الاعتماد الكبير على الذكاء الاصطناعي في الأمن السيبراني لا يخلو من تحديات، مثل احتمالية التلاعب بنماذج الذكاء الاصطناعي نفسها، أو صعوبة تفسير قرارات الأنظمة ذاتية التعلم، إضافة إلى القضايا الأخلاقية والقانونية المرتبطة باستخدام هذه التقنيات. لذا، فإن التطوير المستمر لهذه الأنظمة، جنبًا إلى جنب مع تعزيز التعاون بين الخبراء في مجالي الذكاء الاصطناعي والأمن السيبراني، يظل أمرًا حيويًا لضمان فعالية هذه التقنيات وأمانها.

في المستقبل، من المتوقع أن يلعب الذكاء الاصطناعي دورًا أكبر في تحقيق الأمن السيبراني الشامل، خاصة مع تطور الهجمات الإلكترونية وتعقيدها. لذلك، نوصي بمزيد من البحث في مجال الذكاء الاصطناعي والأمن الموثوق، ووضع أطر تنظيمية تضمن استخدامه بشكل مسؤول وفعال.

بهذا، يكون هذا البحث قد ساهم في تسليط الضوء على الفرص والتحديات التي يقدمها الذكاء الاصطناعي في مجال الأمن السيبراني، آمِلين أن يشكل حافزًا لمزيد من الدراسات والابتكارات في هذا المجال الحيوي.

واختتمت الدراسة بعدد من النتائج والتوصيات:

أولاً، النتائج

1. تقنيات التعلم العميق (Deep Learning)، قادرة على تحليل كميات هائلة من البيانات بسرعة فائقة، مما يُمكن من رصد الأنماط غير الطبيعية التي تشير إلى هجمات إلكترونية، مثل هجمات *التصيد الاحتيالي (Phishing) أو هجمات DDoS.
2. تعزيز الحماية الوقائية باستخدام الذكاء الاصطناعي
3. - الذكاء الاصطناعي يساهم في منع الهجمات قبل حدوثها من خلال تحليل سلوك المستخدمين والأنظمة لاكتشاف أي نشاط مشبوه.
4. - تقنيات مثل التعلم المعزز (Reinforcement Learning) تُستخدم لاختبار أنظمة الحماية عن طريق محاكاة هجمات إلكترونية واكتشاف الثغرات قبل استغلالها.

5. على الرغم من الفوائد الكبيرة، فإن استخدام الذكاء الاصطناعي في الأمن السيبراني يواجه تحديات مثل هجمات الخداع (Adversarial Attacks)، حيث يمكن للمهاجمين التلاعب ببيانات التدريب لخداع النماذج.
6. صعوبة تفسير قرارات الذكاء الاصطناعي (AI Explainability)، مما يجعل من الصعب على الخبراء البشريين فهم كيفية اتخاذ القرارات الأمنية.
7. الاعتماد الزائد على الأنظمة الذكية، مما قد يُضعف المهارات البشرية في مواجهة الهجمات غير النمطية.

ثانياً التوصيات:

1. ضرورة التعاون على المستوى الدولي والإقليمي لتحسين نماذج الذكاء الاصطناعي لجعلها أكثر مقاومة للهجمات السيبرانية.
2. ضرورة دمج الذكاء الاصطناعي مع الخبرة البشرية لتحقيق توازن بين السرعة الآلية والدقة البشرية.
3. يجب وضع معايير أخلاقية وقانونية لضمان استخدام الذكاء الاصطناعي في الأمن السيبراني بشكل مسؤول.
4. يجب تطوير برامج تعليمية متخصصة في الذكاء الاصطناعي والأمن السيبراني بالشراكة مع الجامعات والقطاع الخاص.
5. ضرورة إنشاء منصة عالمية لتبادل المعلومات حول الهجمات الإلكترونية المستندة إلى الذكاء الاصطناعي.

وفي ختام هذه الدراسة التي لا تخلو من نقص أو عيب أو نسيان، أدعو الله عز وجل أن أكون قد وفقت في الإحاطة بكافة جوانبها، وألتمس العذر إن كنت قد قصرت أو أخطأت والكمال لله وحده، وعسى أن يكون لجهدي المتواضع إسهام ولو بالقدر اليسير في الجهود العلمية التي بُذلت والتي ستبذل مستقبلاً لدراسة هذا الموضوع.

ونختم بقول الحق سبحانه وتعالى ﴿رَبَّنَا لَا تُؤْخِذْنَا إِنْ نَسِينَا أَوْ أَخْطَأْنَا رَبَّنَا وَلَا تَحْمِلْ عَلَيْنَا إِصْرًا كَمَا حَمَلْتَهُ عَلَى الَّذِينَ مِنْ قَبْلِنَا رَبَّنَا وَلَا تُحَمِّلْنَا مَا لَا طَاقَةَ لَنَا بِهِ وَاعْفُ عَنَّا وَاعْفُ لَنَا وَارْحَمْنَا أَنْتَ مَوْلَانَا فَانصُرْنَا عَلَى الْقَوْمِ الْكَافِرِينَ﴾⁽¹⁾.
صدق الله العظيم

قائمة المصادر والمراجع

أولاً، المصادر باللغة العربية

الكتب:

1. أحمد السيد النجار، الأمن السيبراني المتقدم باستخدام التعلم الآلي، دار العلم والإيمان، القاهرة، 2023 (280 صفحة)
2. إيمان محمود عبد الرحمن، حماية الأنظمة الذكية من الهجمات الإلكترونية، المركز القومي للبحوث، القاهرة، 2022 (320 صفحة)
3. جمال مصطفى أحمد، الذكاء الاصطناعي في مواجهة الجرائم الرقمية، دار المعارف العلمية، الإسكندرية، 2021 (195 صفحة)
4. حسام الدين محمد علي، تقنيات التعلم العميق لاكتشاف البرمجيات الخبيثة، مكتبة الأنجلو المصرية، القاهرة، 2023 (410 صفحة)
5. خالد طه حسين، أمن الشبكات العصبية الاصطناعية، دار النهضة التكنولوجية، القاهرة، 2022 (360 صفحة)
6. دعاء سعيد محمد، الحماية الذكية للبيانات الشخصية، دار الكتب الجامعية، الإسكندرية، 2023 (240 صفحة)
7. سامح عبد الفتاح، التحليل الآمن للبيانات الكبيرة، الهيئة المصرية العامة للكتاب، 2021 (380 صفحة)

⁽¹⁾ سورة البقرة: آية 286.

8. شيماء محمود عثمان، التوجهات الحديثة في أمن السحابة الإلكترونية، دار الفكر العربي، القاهرة، 2024 (290 صفحة)
9. علياء محمد كامل، أنظمة كشف التسلل المعززة بالذكاء الاصطناعي، مركز الكتاب الأكاديمي، القاهرة، 2022 (330 صفحة)
10. محمد نبيل السيد، التهديدات السيبرانية في عصر إنترنت الأشياء، دار العلوم الحديثة، القاهرة، 2023 (270 صفحة)

رسائل الدكتوراه:

1. أحمد محمود الهاشمي، تأثير الذكاء الاصطناعي على تحسين استراتيجيات الأمن السيبراني في المؤسسات الحكومية، جامعة القاهرة، 2021
2. سارة إبراهيم العلي، تصميم أنظمة استجابة آلية للهجمات السيبرانية باستخدام التعلم المعزز، جامعة المنصورة، 2023
3. فاطمة الزهراء محمد، أمن تطبيقات الذكاء الاصطناعي في الأنظمة المالية، جامعة الإسكندرية، 2023
4. محمد سمير الشراوي، حماية أنظمة إنترنت الأشياء باستخدام تقنيات الذكاء الاصطناعي، جامعة عين شمس، 2022
5. منى علي حسين، تقييم أداء نماذج الذكاء الاصطناعي في كشف اختراقات إنترنت الأشياء، جامعة حلوان، 2023
6. يوسف أمين عبد الله، تعزيز أمن الشبكات الصناعية باستخدام الذكاء الاصطناعي، جامعة أسوان، 2022

ثانياً: المصادر باللغة الأجنبية:

*Books *

1. Chollet, F., Deep learning with Python (2nd ed., 2001 pp. 450–480 .

2. 2. *Russell, S., & Norvig, Artificial intelligence: A modern approach (4th ed.), 2021.pp1020:1050
3. Sarker, I. H, AI-driven cybersecurity and threat intelligence: Cyber automation, intelligent decision-making, and proactive defense,2021
4. .Pajouh, H. H., et al, Machine learning for cybersecurity: Hands-on approaches for intrusion detection and malware analysis,2022.
5. . Sikorski, M., & Honig, A., Practical malware analysis: The hands-on guide to dissecting malicious software (2nd ed , pp. 300–330 .
6. .Cole, E, Advanced persistent threat: Understanding the danger and how to protect your organization,2020 .,pp. 200–225 .
7. O’Neil, C, Weapons of math destruction: How big data increases inequality and threatens democracy,2016
8. Goodfellow, I., et al, Adversarial robustness in machine learning,2019
9. .Schneier, B, Data and Goliath: The hidden battles to collect your data and control your world.,2018
10. .Kaplan, J., Artificial intelligence: What everyone needs to know,2021 .,pp. 150–175. .

***.Journal Articles *:**

1. Apruzzese, G., Colajanni, M., Ferretti, L., Guido, A., & Marchetti, M. (2018). On the effectiveness of machine and deep learning for cybersecurity. 2018 10th International Conference on Cyber Conflict (CyCon), 371–390. <https://doi.org/10.23919/CYCON.2018.8405026>

2. Buczak, A. L., & Guven, E. (2016). A survey of data mining and machine learning methods for cybersecurity intrusion detection. IEEE. <https://doi.org/10.1109/COMST.2015.2494502>

*** .Conference Papers * :**

1. Anderson, H. S., & Roth, P. (2018). Ember: An open dataset for training static PE malware machine learning models. ArXiv Preprint. <https://arxiv.org/abs/1804.04637>
2. Yin, C., Zhu, Y., Fei, J., & He, X. (2017). A deep learning approach for intrusion detection using recurrent neural networks. IEEE Access, 5, 21954–21961 .

*** .Reports & Whitepapers * :**

1. IBM Security. (2023). Cost of a data breach report 2023. IBM. <https://www.ibm.com/reports/data-breach>
2. MITRE. (2022). Adversarial machine learning: A taxonomy and terminology of attacks and mitigations (MITRE Report No. 22-0001). <https://www.mitre.org/publications>

*** .Websites & Online Articles * :**

1. National Institute of Standards and Technology (NIST). (2021). AI for cybersecurity. <https://www.nist.gov/itl/ai-cybersecurity>
2. Schneier, B. (2019). AI and cybersecurity: Hype and reality. Wired. <https://www.wired.com/story/ai-cybersecurity-hype-reality>