

Arabian Gulf Journal of Humanities and Social Studies

ISSN: 3080-4086

الإصدار السادس - العدد السابع عشر || تاريخ الإصدار 2026-08-20



التوازن بين توظيف الذكاء الاصطناعي و متطلبات الامن السيبراني دراسة في التحديات القانونية

Balancing the Use of Artificial Intelligence with Cybersecurity Requirements: A Study of Legal Challenges

المدرس المساعد ميثم فلاح حسن

Maitham Falah Hasan

جامعة الفراهيدي – كلية القانون

المدرس المساعد عبد الكريم عدنان اشرف

Abdulkareem Adnan Ashraf

جامعة الفراهيدي – كلية القانون

DOI: <https://doi.org/10.64355/agjhss61737>

مجلة خليج العرب للدراسات الإنسانية والاجتماعية || هذه المقالة مفتوحة المصدر موزعة بموجب شروط وأحكام ترخيص مؤسسة المشاع الإبداعي (CC BY-NC-SA)

Clarivate | ProQuest

Ulrichsweb™



Crossref doi

ISSN INTERNATIONAL  
STANDARD  
SERIAL  
NUMBER  
INTERNATIONAL CENTRE



Google Scholar

معرفة  
e-Marefa



شبكة المعلومات العربية  
Arab Educational Information Network

AskZad



Scilit

INTERNATIONAL  
Scientific Indexing

CC creative commons

## الملخص:

يتناول هذا البحث طبيعة العلاقة بين الذكاء الاصطناعي والأمن السيبراني، في ظل التطور المتسارع الذي تشهده التقنيات الرقمية وتزايد الاعتماد على الأنظمة الذكية في مختلف المجالات وقد أدى هذا التطور إلى نشوء علاقة متبادلة بين الذكاء الاصطناعي والأمن السيبراني، إذ أصبح الذكاء الاصطناعي أداة مهمة في تعزيز القدرات الدفاعية للأنظمة الرقمية، من خلال تحليل البيانات الضخمة، واكتشاف الأنماط غير الاعتيادية، والتنبؤ بالهجمات السيبرانية والاستجابة لها بصورة أسرع وأكثر كفاءة. وفي المقابل، أفرز استخدام الذكاء الاصطناعي في البيئة السيبرانية مجموعة من المخاطر والتحديات الجديدة، ولا سيما إمكانية استغلال تقنيات الذكاء الاصطناعي من قبل المهاجمين في تطوير أساليب أكثر تعقيداً لتنفيذ الهجمات السيبرانية، فضلاً عن صعوبة اكتشاف بعض الهجمات التي تعتمد على تقنيات متقدمة وقابلة للتكيف. الأمر الذي يجعل العلاقة بين المجالين ذات طبيعة مزدوجة، تجمع بين كون الذكاء الاصطناعي وسيلة لتعزيز الأمن السيبراني، وكونه في الوقت ذاته مصدراً محتملاً لتهديده. ويهدف البحث إلى بيان طبيعة هذه العلاقة، وتحليل أهم أوجه توظيف الذكاء الاصطناعي في حماية الأنظمة والشبكات والمعلومات، فضلاً عن الوقوف على أبرز المخاطر والتحديات التي قد تنتج عن استخدامه في المجال السيبراني، مع بيان الحاجة إلى وضع أطر قانونية وتنظيمية وتقنية قادرة على تحقيق التوازن بين الاستفادة من تطبيقات الذكاء الاصطناعي والحد من مخاطره. وقد خلص البحث إلى أن تحقيق الأمن السيبراني في ظل التطور المتزايد للذكاء الاصطناعي يتطلب اعتماد رؤية متكاملة تجمع بين التطور التقني والتنظيم القانوني والرقابة المؤسسية، بما يضمن توظيف الذكاء الاصطناعي لتعزيز الحماية السيبرانية والحد من إمكانية استخدامه كوسيلة لارتكاب الهجمات والاعتداءات الرقمية. **الكلمات المفتاحية:** الذكاء الاصطناعي، الأمن السيبراني، الهجمات السيبرانية، المخاطر السيبرانية، الأمن الرقمي، الحماية السيبرانية.

## Abstract:

This research examines the nature of the relationship between Artificial Intelligence and Cybersecurity in light of the rapid development of digital technologies and the increasing reliance on intelligent systems across various fields. This technological development has created a reciprocal relationship between Artificial Intelligence and Cybersecurity, as Artificial Intelligence has become an important tool for enhancing

the defensive capabilities of digital systems through big data analysis, detection of abnormal patterns, prediction of cyberattacks, and faster and more efficient response mechanisms.

At the same time, the use of Artificial Intelligence in the cyber environment has created a number of emerging risks and challenges, particularly the possibility of exploiting Artificial Intelligence technologies by malicious actors to develop more sophisticated methods of conducting cyberattacks. Furthermore, some attacks based on advanced and adaptive technologies may become increasingly difficult to detect. Consequently, the relationship between the two fields has a dual nature: Artificial Intelligence can serve as an effective means of strengthening Cybersecurity, while simultaneously representing a potential source of cybersecurity threats.

The research aims to clarify the nature of this relationship and analyze the main applications of Artificial Intelligence in protecting systems, networks, and information. It also seeks to identify the most significant risks and challenges arising from its use in the cyber domain, while highlighting the need for legal, regulatory, and technical frameworks capable of achieving a balance between benefiting from Artificial Intelligence applications and mitigating their potential risks.

The research concludes that achieving Cybersecurity in the context of the continuous development of Artificial Intelligence requires an integrated approach combining technological advancement, legal regulation, and institutional oversight. Such an approach should ensure the effective use of Artificial Intelligence to strengthen cybersecurity protection while limiting its potential exploitation as a means of conducting cyberattacks and digital violations.

**Keywords:** Artificial Intelligence, Cybersecurity, Cyberattacks, Cyber Risks, Digital Security, Cyber Protection.

#### المقدمة

أدى التطور المتسارع في تقنيات المعلومات والاتصالات إلى إحداث تحولات جوهرية في مختلف مجالات الحياة، وأصبح الفضاء السيبراني يمثل بيئة أساسية لممارسة الأنشطة الاقتصادية والاجتماعية والإدارية والعلمية، الأمر الذي جعل حماية هذه البيئة والمحافظة على أمنها واستقرارها من أهم التحديات التي تواجه الدول والمؤسسات والأفراد في العصر الرقمي. وفي ظل هذا التطور برز الذكاء الاصطناعي بوصفه إحدى أهم التقنيات الحديثة القادرة على إحداث تغيير واسع في طبيعة استخدام الأنظمة الرقمية ووسائل حمايتها، فضلاً عن تأثيره المباشر في طبيعة المخاطر والتهديدات التي تواجه الأمن السيبراني ويتميز الذكاء الاصطناعي بقدرته على معالجة كميات كبيرة من البيانات وتحليلها واكتشاف الأنماط والعلاقات بينها، فضلاً عن قدرته على التعلم والتنبؤ واتخاذ بعض القرارات بصورة آلية أو شبه مستقلة. وقد جعلت هذه الخصائص منه أداة مهمة في تعزيز القدرات السيبرانية، من خلال اكتشاف الأنشطة غير الطبيعية، وتحليل السلوك الشبكي، والتنبؤ بالهجمات، والاستجابة السريعة للحوادث الأمنية. وفي المقابل، فإن القدرات ذاتها يمكن أن يستفيد منها المهاجمون في تطوير أساليب أكثر تعقيداً في ارتكاب الهجمات السيبرانية، الأمر الذي يجعل العلاقة بين الذكاء الاصطناعي والأمن السيبراني علاقة ذات طبيعة مزدوجة، تجمع بين الحماية والتهديد في آن واحد ولا تقتصر أهمية العلاقة بين الذكاء الاصطناعي والأمن السيبراني على الجانب التقني، وإنما تمتد إلى أبعاد قانونية وأمنية واقتصادية واجتماعية؛ إذ إن اختراق الأنظمة التي تعتمد على الذكاء الاصطناعي أو التلاعب ببياناتها قد يؤدي إلى نتائج تتجاوز الأضرار التقنية المباشرة، لتصل إلى المساس بالخصوصية وسلامة المعلومات واستمرارية الخدمات، فضلاً عن إمكانية التأثير في قطاعات حيوية تمس الأمن الوطني والاقتصاد والخدمات العامة وفي المقابل، يمثل توظيف الذكاء الاصطناعي في الأمن السيبراني وسيلة متقدمة للانتقال من الأساليب التقليدية في مواجهة الهجمات إلى أساليب أكثر قدرة على التنبؤ والكشف والاستجابة. فبدلاً من الاقتصار على اكتشاف الهجوم بعد وقوعه، يمكن للأنظمة المدعومة بالذكاء الاصطناعي تحليل السلوك والبيانات والمؤشرات المختلفة بهدف التعرف على الأنماط التي قد تشير إلى وجود تهديد قبل تطوره إلى هجوم فعلي، ومن ثم فإن الذكاء الاصطناعي يمكن أن يسهم في تعزيز عناصر السرعة والدقة والاستباقية في منظومة الأمن السيبراني.

#### أولاً: أهمية البحث

تتبع أهمية البحث من حادثة موضوع العلاقة بين الذكاء الاصطناعي والأمن السيبراني، ومن اتساع نطاق استخدام تقنيات الذكاء الاصطناعي في مختلف القطاعات الحيوية. كما تبرز أهميته في محاولة بيان الطبيعة المزدوجة لهذه العلاقة، إذ يمكن للذكاء الاصطناعي أن يكون وسيلة لتعزيز الأمن السيبراني، وفي الوقت ذاته أداة يمكن استغلالها لتطوير الهجمات والتهديدات السيبرانية.

### ثانيا : مشكلة البحث

تتمثل مشكلة البحث في بيان طبيعة العلاقة بين الذكاء الاصطناعي والأمن السيبراني، ومدى قدرة تقنيات الذكاء الاصطناعي على تعزيز حماية الأنظمة والشبكات، في مقابل المخاطر التي يمكن أن تنتج عن استخدامها أو استغلالها في تنفيذ الهجمات السيبرانية.

### ثالثا :منهج البحث

يعتمد البحث على المنهج التحليلي من خلال تحليل طبيعة الذكاء الاصطناعي والأمن السيبراني وبيان أوجه التفاعل بينهما، كما يعتمد على المنهج المقارن عند تناول بعض الاتجاهات الدولية والتنظيمية ذات الصلة، فضلاً عن الاستفادة من التقارير والدراسات الدولية المتخصصة في مجال الأمن السيبراني والذكاء الاصطناعي

### رابعا :خطة البحث

ينقسم البحث إلى مبحثين:

المبحث الأول: الذكاء الاصطناعي كأداة لتعزيز الأمن السيبراني

المبحث الثاني: مخاطر وتحديات الذكاء الاصطناعي في مجال الأمن السيبراني

### المبحث الأول

#### الذكاء الاصطناعي كأداة لتعزيز الأمن السيبراني

لقد برز الذكاء الاصطناعي وخاصة التعلم الآلي والتعلم العميق، كقوة تحويلية في مجال الأمن السيبراني، حيث يوفر قدرات غير مسبقة للكشف عن التهديدات، والاستجابة للحوادث، وإدارة الثغرات الأمنية. إن قدرة الذكاء الاصطناعي على معالجة وتحليل كميات هائلة من البيانات بسرعة ودقة تتجاوز القدرات البشرية تجعله حليفاً لا غنى عنه في مواجهة المشهد المتطور للتهديدات السيبرانية .

### المطلب الاول

#### فعالية الذكاء الاصطناعي في تعزيز الامن السيبراني

يلعب الذكاء الاصطناعي دورا بارزا في أي مجال يتصل به ويساهم في تطويره لو تم استخدامه في الطريقة الصحيحة والمناسبة فيساهم في تطوير المجال المتصل به ، وكذلك الحال بالنسبة لو اتصل الذكاء الاصطناعي مع الامن السيبراني وهذا ما سنتناوله في الفرع الاول دور الذكاء الاصطناعي الكشف عن التهديدات وتحليلها و الثاني الذكاء الاصطناعي و دوره في منع الجريمة السيبرانية

## الفرع الاول

### دور الذكاء الاصطناعي الكشف عن التهديدات وتحليلها

يُعد الكشف عن التهديدات وتحليلها أحد أبرز المجالات التي أحدث فيها الذكاء الاصطناعي ثورة في الأمن السيبراني. تعتمد الأنظمة التقليدية للكشف عن التهديدات على التوقعات المعروفة للبرمجيات الخبيثة أو أنماط الهجوم المحددة مسبقاً، مما يجعلها غير فعالة ضد التهديدات الجديدة وغير المعروفة. هنا يأتي دور الذكاء الاصطناعي من خلال (1)

**اولا : الكشف عن البرمجيات الخبيثة:** تستخدم خوارزميات التعلم الآلي مثل آلات المتجهات الداعمة والشبكات العصبية، لتحليل خصائص الملفات، وسلوك البرامج، وأنماط التعليمات البرمجية لتحديد ما إذا كانت خبيثة، حتى لو كانت متغيرة أو مشفرة. يمكن للتعلم العميق، على وجه الخصوص، اكتشاف الأنماط المعقدة في البيانات الثنائية أو تدفقات الشبكة التي قد تشير إلى وجود برمجيات خبيثة متطورة (2).

**ثانيا : الكشف عن التسلسل :** يمكن لنماذج الذكاء الاصطناعي مراقبة حركة مرور الشبكة وسجلات النظام وسلوك المستخدمين لتحديد الانحرافات عن الأنماط الطبيعية. على سبيل المثال، يمكن لخوارزميات الكشف عن الشذوذ تحديد محاولات الوصول غير المصرح به، أو الأنشطة المشبوهة داخل الشبكة (3)، أو هجمات حجب الخدمة الموزعة من خلال تحليل تدفقات البيانات في الوقت الفعلي.

**ثالثا : تحليل السلوك :** من خلال بناء نماذج سلوكية للمستخدمين والأجهزة والتطبيقات، يمكن للذكاء الاصطناعي تحديد أي سلوك غير عادي قد يشير إلى اختراق أو تهديد داخلي على سبيل المثال، إذا بدأ حساب مستخدم في الوصول إلى ملفات حساسة في أوقات غير معتادة أو من مواقع غير مألوفة، يمكن لنظام مدعوم بالذكاء الاصطناعي أن يطلق تنبيهاً (4).

**رابعا : الاستجابة للحوادث والتحقيق الجنائي الرقمي :** لا يقتصر دور الذكاء الاصطناعي على الكشف عن التهديدات فحسب، بل يمتد ليشمل تعزيز قدرات الاستجابة للحوادث والتحقيق الجنائي الرقمي، مما يقلل من وقت الاستجابة ويحد من الأضرار المحتملة (5).

1. **الاستجابة التلقائية للحوادث:** يمكن لأنظمة الأمن المعتمدة على الذكاء الاصطناعي (مثل منصات التشغيل الآلي للأمن والاستجابة) أن تتخذ إجراءات تلقائية فور اكتشاف تهديد. قد تشمل هذه الإجراءات عزل الأجهزة المصابة، أو حظر عناوين IP الضارة، أو تطبيق قواعد جدار الحماية الجديدة، مما يقلل بشكل كبير من وقت الاستجابة للمهاجمين داخل الشبكة (6).

2. **التحقيق الجنائي الرقمي المعزز بالذكاء الاصطناعي:** في مرحلة ما بعد الحادث، يمكن للذكاء الاصطناعي تسريع عملية التحقيق من خلال تحليل كميات هائلة من البيانات اللوغاريتمية وسجلات الشبكة، وصور الذاكرة لتحديد السبب الجذري للهجوم، ومسار الاختراق، والبيانات المتأثرة. يمكن لخوارزميات التعلم العميق تحديد الأنماط الخفية في البيانات التي قد تفوت المحققين البشريين، مما يوفر رؤية قيمة لإعادة بناء سيناريو الهجوم (7).

3. **تحليل التهديدات :** يمكن للذكاء الاصطناعي تجميع وتحليل كميات هائلة من معلومات التهديدات من مصادر متعددة (تقارير، منتديات، شبكات اجتماعية) لتحديد التهديدات الناشئة، وتكتيكات المهاجمين، وتقنياتهم مما يساعد المؤسسات على الاستعداد بشكل أفضل للهجمات المستقبلية (8).

4. إدارة الثغرات ونقاط الضعف: تُعد إدارة الثغرات الأمنية عملية معقدة وتستغرق وقتاً طويلاً. يمكن للذكاء الاصطناعي

تبسيط هذه العملية وجعلها أكثر فعالية من خلال تحديد الثغرات المحتملة وتحديد أولوياتها(9).

5. مسح الثغرات الأمنية: يمكن لأنظمة الذكاء الاصطناعي تحليل التعليمات البرمجية المصدر والتطبيقات والشبكات لتحديد

الثغرات الأمنية المعروفة وغير المعروفة بشكل أكثر كفاءة من الماسحات الضوئية التقليدية. يمكنها أيضاً التنبؤ بالثغرات

المحتملة بناءً على أنماط الكود أو التكوينات الخاطئة(10).

6. الأمن التنبؤي والوقائي: يمثل الأمن التنبؤي والوقائي قمة استخدام الذكاء الاصطناعي في الأمن السيبراني، حيث ينتقل

من الاستجابة إلى التهديدات إلى التنبؤ بها ومنعها قبل وقوعها(11).

التنبؤ بالتهديدات يمكن لنماذج التعلم الآلي تحليل البيانات التاريخية للتهديدات، والاتجاهات العالمية، وحتى العوامل الجيوسياسية للتنبؤ بالأنواع المحتملة للهجمات التي قد تستهدف مؤسسة معينة في المستقبل القريب. هذا يسمح للمؤسسات باتخاذ تدابير وقائية استباقية.

الصيد الاستباقي للتهديدات بدلاً من انتظار التنبيهات، يمكن للذكاء الاصطناعي مساعدة فرق الصيد الاستباقي للتهديدات في البحث عن مؤشرات الاختراق المخفية داخل الشبكة. يمكنه تحديد الأنماط الشاذة أو السلوكيات الخفية التي قد تشير إلى وجود مهاجم لم يتم اكتشافه بعد(12).

يمكن للذكاء الاصطناعي إنشاء بيانات محاكاة لاختبار فعالية الدفاعات الحالية ضد أنواع مختلفة من الهجمات، بما في ذلك الهجمات المعززة بالذكاء الاصطناعي. هذا يساعد على تحديد نقاط الضعف في البنية التحتية الأمنية قبل أن يستغلها المهاجمون(13).

## الفرع الثاني

الذكاء الاصطناعي و دوره في منع الجريمة السيبرانية

كل تقدم علمي يشهده العالم يكون له جانبين سلبي و ايجابي فأن ظهور الذكاء الاصطناعي يمكن ان يكون له جانب سلبي ويتمثل كأداة او وسيلة لتسهيل ارتكاب الجريمة لكن في نفس الوقت له جانب ايجابي وهو امكانية استخدامه كحارس رقمي لمنع ارتكاب الجرائم عامة ومنع ارتكاب بعض الجرائم بصورة خاصة ويمكن تحديد الدور الذي يلعبه الذكاء الاصطناعي كأداة تمنع ارتكاب الجريمة من خلال بعض النقاط الاساسية وهي :

### 1. تحليل البيانات والتنبؤ بالجريمة:

من خلال تحليل الأنماط الإجرامية حيث يمكن للذكاء الاصطناعي تحليل البيانات التاريخية للجرائم لتحديد الأنماط والاتجاهات الإجرامية. على سبيل المثال، يمكن للخوارزميات التنبؤ بمناطق الجريمة المحتملة بناءً على عوامل مثل الوقت والمكان والظروف الاجتماعية(14).

2. الوقاية من الجريمة:  
يمكن استخدام أنظمة الذكاء الاصطناعي للتنبؤ بالجرائم قبل حدوثها. هذه الأنظمة تعتمد على تحليل البيانات في الوقت الفعلي لتحديد السلوكيات المشبوهة(15).
3. المراقبة والأمن:  
من خلال أنظمة المراقبة الذكية يمكن للذكاء الاصطناعي تحليل لقطات الكاميرات الأمنية للكشف عن السلوكيات المشبوهة، مثل السرقة أو التحرش. هذه الأنظمة قادرة على التعرف على الوجوه وتحليل الحركات غير الطبيعية(16).
4. التعرف على الأصوات(17): يمكن للذكاء الاصطناعي تحليل الأصوات للكشف عن التهديدات المحتملة، مثل الصراخ أو الكلمات العدوانية.
5. التحقيق الجنائي:  
يمكن تحليل الأدلة الرقمية: من خلال قدرة للذكاء الاصطناعي تحليل الأدلة الرقمية، مثل رسائل البريد الإلكتروني وملفات الوسائط الاجتماعية، للكشف عن الأدلة التي قد تفيد في التحقيقات(18).
6. التعرف على الصور والفيديوهات المزيفة: يمكن للذكاء الاصطناعي الكشف عن الصور والفيديوهات المزيفة (Deepfake) التي قد تستخدم في الجرائم الإلكترونية(19).
7. المراقبة الإلكترونية الذكية والتعرف على الوجوه حيث من خلال أنظمة الكاميرات الذكية المدعومة بالذكاء الاصطناعي يمكنها التعرف على المشتبه بهم في الوقت الفعلي ومطابقتهم بقواعد البيانات الأمنية ، ويمكن تحليل لغة الجسد والسلوك المشبوه في الأماكن العامة لاكتشاف النشاط الإجرامي المحتمل(20).

## المطلب الثاني

تكامل الذكاء الاصطناعي مع منصات الامن السيبراني واثره بتعزيز الحماية السيبرانية

فرض التحول الرقمي المتسارع تسارعاً كبيراً في تعقيد الهجمات السيبرانية وحجمها، مما أدى إلى قصور الآليات الدفاعية التقليدية القائمة على التوقيعات الرقمية والقواعد الثابتة. يركز هذا المبحث على دراسة المنصات الدفاعية الحديثة التي تدمج تقنيات الذكاء الاصطناعي (AI)، والتعلم الآلي، والذكاء الاصطناعي التوليدي لتحقيق الحماية الاستباقية والأتمتة في البيئات الرقمية المعقدة، وسنتكلم في هذا المطلب عن انواع المنصات الامن السيبراني و دور الذكاء الاصطناعي في توظيف تقنياته

## الفرع الاول

التصنيف الوظيفي للمنصات الامنية المعتمدة على الذكاء الاصطناعي

يمكن تصنيف المنصات الأمنية المعتمدة على الذكاء الاصطناعي وفقاً لطبقة الدفاع المستهدفة والآلية الخوارزمية المستخدمة إلى ثلاثة أطر رئيسية:(21)

1. أنظمة حماية الأجهزة الطرفية والاستجابة المتقدمة: (AI-Powered EDR/XDR)
  - CrowdStrike Falcon: تعتمد على معمارية سحابية هجينة تدمج بين التعلم الآلي الموجه وغير الموجه في محرك التهديدات، إضافة إلى المساعد التوليدي (Charlotte AI) لترجمة بيانات الأحداث المعقدة إلى استجابات فورية .
  - SentinelOne Singularity: تتخصص في تقديم آليات استجابة آلية نسيجية تشمل الاستجابة المحلية على الأجهزة (On-device AI Engine) ومكافحة برمجيات الفدية عبر الترجيع الآلي للبيانات (Autonomous Rollback) مدعومة بمساعد الذكاء الاصطناعي (Purple AI).

## 2. أنظمة المراقبة الذاتية وتحليل سلوك الشبكات: (Autonomous Network & Cloud Defense)

- **Darktrace:** تطبيق نموذج "التعلم الذاتي (Self-Learning AI)" القائم على التعلم غير الموجه (Unsupervised Learning) لإنشاء نمط سلوكي قياسي (Pattern of Life) لكل جهاز ومستخدم داخل المؤسسة، مما يتيح اكتشاف الشذوذ دون الحاجة لمؤشرات اختراق سابقة. (IoCs) (22)
- **Wiz:** تركز على أمن البنية التحتية السحابية (CSPM/CNAPP) من خلال تحليل الرسوم البيانية الأمنية (Security Graph) المدعومة بالذكاء الاصطناعي لتحديد مسارات الهجمات الحرجة وترتيب أولويات المخاطر عبر البيانات متعددة السحابات.

## 3. أنظمة التنسيق والتحليل للعمليات الأمنية: (AI-Driven SOC & Security Operations)

- **Microsoft Security Copilot:** يعتمد على نماذج اللغات الكبيرة (LLMs) لمعالجة اللغات الطبيعية وترجمتها إلى استعلامات أمنية (Security Queries)، وتلخيص الحوادث الأمنية المعقدة لتقليل زمن التحقيق الإجمالي.
- **Palo Alto Networks (Cortex XSIAM):** تعتمد على معمارية مراكز العمليات المستقلة (Autonomous SOC) بدمج إدارة السجلات والبيانات الضخمة مع التنسيق الآلي (SOAR) لتقليل من قدر الإنذارات الخاطئة وأتمتة معالجة الحوادث. (23)

### الفرع الثاني

#### توظيف تقنيات الذكاء الاصطناعي في منصات الأمن السيبراني وانعكاساته

لم يعد توظيف الذكاء الاصطناعي في مجال الأمن السيبراني مقتصرًا على إضافة أدوات تقنية إلى المنظومات الأمنية القائمة، بل اتجه نحو إعادة تشكيل طريقة بناء المنصات الأمنية وآليات إدارتها للمعلومات والتهديدات. فقد أصبحت المنصة الحديثة بيئة مترابطة تستقبل البيانات من مصادر متعددة، ثم تعمل على تنظيمها وتحليلها وربطها ببعضها للوصول إلى تصور أكثر شمولاً عن الحالة الأمنية للنظام. ومن هنا يظهر الذكاء الاصطناعي بوصفه عنصراً يدخل في صميم البنية التشغيلية للمنصة، وليس مجرد وظيفة إضافية يمكن الاستغناء عنها (24).

وتقوم هذه المنصات على استقبال تدفقات مستمرة من البيانات الناتجة عن أجهزة الحاسوب والخوادم والشبكات والتطبيقات وقواعد البيانات، فضلاً عن سجلات الدخول والاستخدام. وتمثل هذه البيانات الأساس الذي تعتمد عليه الخوارزميات الذكية في بناء صورة عن النشاط المعتاد داخل البيئة الرقمية. وكلما اتسعت مصادر البيانات وتنوعت، أصبحت المنصة أكثر قدرة على تكوين تصور مترابط عن الأحداث بدلاً من التعامل مع كل واقعة أمنية بصورة منفردة.

وتبرز أهمية الذكاء الاصطناعي هنا في قدرته على إقامة علاقات بين الأحداث المتفرقة التي قد تبدو، عند النظر إليها بصورة منفصلة، غير ذات أهمية أمنية. فقد تحدث مجموعة من محاولات الدخول الفاشلة، يعقبها اتصال من جهاز غير مألوف، ثم انتقال غير اعتيادي للبيانات. وقد لا يمثل أي حدث من هذه الأحداث وحده دليلاً حاسماً على وجود اعتداء (25)، إلا أن جمعها وتحليلها ضمن سياق واحد قد يكشف عن سلوك يستحق التدقيق. وبذلك تنتقل المنصة من مراقبة الأحداث إلى فهم السياق الذي تجري فيه تلك الأحداث.

ويؤدي الذكاء الاصطناعي كذلك دوراً في ترتيب الأولويات الأمنية داخل المنصة، إذ إن كثرة التنبيهات التي تنتجها الأنظمة الأمنية قد تؤدي إلى إرهاق فرق الحماية وإضاعة الوقت في التعامل مع إنذارات منخفضة الأهمية. ولذلك يمكن للأنظمة الذكية

أن تسهم في تصنيف التنبيهات وفق مستوى الخطورة، وربط بعضها ببعض، وتمييز الحالات التي تستوجب تدخلاً سريعاً عن الحالات التي يمكن إخضاعها للمراجعة اللاحقة. وتكتسب هذه الوظيفة أهمية خاصة في المؤسسات التي تتعامل مع أعداد كبيرة من الأحداث الأمنية في وقت واحد (26).

كما أن إدخال الذكاء الاصطناعي إلى المنصات الأمنية أدى إلى تعزيز مفهوم الأتمتة الأمنية، بحيث لم يعد دور المنصة مقتصرًا على عرض المعلومات أمام المختص، وإنما أصبح بالإمكان تصميم إجراءات محددة تنفذ بصورة آلية عند تحقق شروط معينة. وقد تتضمن هذه الإجراءات تقييد اتصال معين، أو تعليق جلسة مشبوهة، أو إرسال إشعار إلى الفريق المختص، أو جمع معلومات إضافية عن الحادث. ويسهم ذلك في تقليل الفترة الزمنية الفاصلة بين ظهور المؤشر الأمني واتخاذ الإجراء المناسب. غير أن هذه الأتمتة تطرح في الوقت ذاته إشكالية تتعلق بحدود الاستقلال الذي ينبغي منحه للمنصة الأمنية. فكلما ازدادت قدرة النظام على اتخاذ الإجراءات بصورة تلقائية، ازدادت الحاجة إلى تحديد الحالات التي يمكن فيها السماح للألة باتخاذ القرار منفردة، والحالات التي تستوجب تدخلاً بشرياً. ويصبح الأمر أكثر حساسية عندما يكون الإجراء الآلي قادراً على التأثير في بيانات أو حسابات أو أنظمة يعتمد عليها أشخاص أو مؤسسات في ممارسة أعمالهم (27).

ومن الجوانب المهمة كذلك أن فعالية المنصات الذكية ترتبط بصورة وثيقة بنوعية البيانات التي تعتمد عليها. فالخوارزمية لا تعمل في فراغ، وإنما تستند إلى بيانات تدخل في عملية التعلم والتحليل واتخاذ القرار. وإذا كانت هذه البيانات ناقصة أو غير دقيقة أو متحيزة أو تعرضت للتلاعب، فإن مخرجات المنصة قد تتأثر بصورة مباشرة. ومن ثم فإن أمن البيانات المستخدمة في تشغيل الذكاء الاصطناعي يصبح جزءاً من أمن المنصة ذاتها، الأمر الذي يضيف طبقة جديدة من متطلبات الحماية (28).

وتتخذ المخاطر بعداً أكثر تعقيداً عندما يحاول المهاجم استهداف المنصة من الداخل، من خلال التأثير في البيانات التي تعتمد عليها أو استغلال نقاط الضعف الموجودة في نماذجها. ففي هذه الحالة لا يكون الهدف اختراق النظام المحمي فحسب، وإنما التأثير في الأداة التي يفترض أن تتولى حمايته. وهذا يمثل تحولاً مهماً في طبيعة التهديد، لأن نجاح المهاجم قد يؤدي إلى إضعاف القدرة الدفاعية للمنصة بأكملها، وليس مجرد الوصول إلى مورد أو جهاز محدد.

ومن جهة أخرى، فإن الاعتماد على المنصات القائمة على الذكاء الاصطناعي يثير مسألة قابلية تفسير القرارات الآلية. فقد تصل الخوارزمية إلى نتيجة معينة بشأن نشاط أو مستخدم أو اتصال دون أن يكون من السهل على الشخص المسؤول تحديد جميع الأسباب التي أدت إلى هذه النتيجة. وتظهر أهمية ذلك عند اتخاذ إجراءات أمنية مؤثرة، إذ لا يكفي من الناحية العملية أن تعرف الجهة المختصة أن النظام صنف نشاطاً معيناً باعتباره خطراً، وإنما تحتاج كذلك إلى معرفة الأساس الذي بني عليه هذا التصنيف، خصوصاً عند مراجعة القرار أو الاعتراض عليه (29).

ويترتب على ذلك بروز الحاجة إلى نموذج أمني يجمع بين الكفاءة التقنية والرقابة البشرية. فالذكاء الاصطناعي يمكن أن يوفر سرعة في معالجة البيانات وقدرة على التعامل مع حجم هائل من الأحداث، إلا أن ذلك لا يعني بالضرورة إمكانية إحلاله بصورة كاملة محل العنصر البشري. إذ تظل الخبرة البشرية ضرورية في تفسير الحالات المعقدة، والتحقق من القرارات الحساسة، وتقدير الآثار المترتبة على الإجراءات الأمنية.

وتتضح من ذلك الطبيعة المزدوجة لمنصات الأمن السيبراني المعتمدة على الذكاء الاصطناعي؛ فهي من ناحية تمثل تطوراً مهماً في القدرة على إدارة البيئة الرقمية المعقدة، ومن ناحية أخرى تخلق مجموعة من المخاطر الجديدة المرتبطة بسلامة البيانات، وأمن النماذج، وصعوبة تفسير بعض القرارات، وحدود الأتمتة، فضلاً عن مسألة تحديد المسؤولية عند وقوع خطأ ناتج عن قرار آلي.

ومن المنظور القانوني، تبرز أهمية هذه المسائل في تحديد الإطار الذي يحكم عمل المنصات الذكية، ولا سيما فيما يتعلق بتحديد مسؤولية الجهة التي تعتمد النظام، ومسؤولية مطور التقنية أو مزودها، وحدود تدخل العنصر البشري، ومدى مشروعية الإجراءات الآلية التي قد تمس حقوق الأفراد أو مصالحهم. وعليه فإن التطور التقني لمنصات الأمن السيبراني لا يمكن فصله عن ضرورة مواكبته بتطور قانوني يحدد ضوابط استخدامها ويضع معايير واضحة للمساءلة عند حدوث خلل أو إساءة استخدام<sup>(30)</sup>.

وبذلك يمكن القول إن توظيف الذكاء الاصطناعي في منصات الأمن السيبراني يمثل انتقالاً من مجرد استخدام أدوات تقنية منفصلة إلى بناء بيئة أمنية تعتمد على الترابط بين البيانات والتحليل والأتمتة واتخاذ القرار. غير أن هذا الانتقال، رغم ما يوفره من إمكانات متقدمة، يفرض في المقابل تحديات تقنية وقانونية تستوجب معالجة متوازنة تضمن الاستفادة من قدرات الذكاء الاصطناعي دون تحويله إلى مصدر جديد للمخاطر الأمنية أو وسيلة لاتخاذ قرارات غير خاضعة للرقابة والمساءلة<sup>(31)</sup>.

## المبحث الثاني

### مخاطر وتحديات الذكاء الاصطناعي في مجال الأمن السيبراني

يقوم الذكاء الاصطناعي بدور مزدوج في البيئة السيبرانية؛ فهو من جهة يمثل وسيلة متقدمة لتعزيز القدرة على اكتشاف التهديدات والاستجابة لها، ومن جهة أخرى يمكن أن يُستغل من قبل المهاجمين لتطوير هجمات أكثر سرعة ودقة وتعقيداً. وقد أصبح هذا التداخل أكثر وضوحاً مع تطور الذكاء الاصطناعي التوليدي، فضلاً عن ظهور أساليب هجومية تستهدف نماذج الذكاء الاصطناعي وبياناتها ذاتها. ويصنف المعهد الوطني الأمريكي للمعايير والتكنولوجيا (NIST) عدداً من هذه المخاطر ضمن هجمات التعلم الآلي العدائي، ومنها هجمات التهريب، وتسميم البيانات، وهجمات الخصوصية وإساءة استخدام أنظمة الذكاء.

### المطلب الأول

#### المخاطر السيبرانية الناشئة عن استخدام الذكاء الاصطناعي

أدى التطور المتسارع في تقنيات الذكاء الاصطناعي إلى ظهور مخاطر سيبرانية جديدة، نتيجة إمكانية توظيفه في تحليل البيانات، واكتشاف الثغرات، وأتمتة الهجمات، وتطوير أساليب أكثر سرعة وتعقيداً وصعوبة في الكشف. كما يُستغل في تنفيذ هجمات التصيد والاحتيال، ونشر البرمجيات الخبيثة، واستهداف الأنظمة والشبكات والبيانات الحساسة. لذلك، يتناول هذا المطلب أبرز المخاطر السيبرانية المرتبطة باستخدام الذكاء الاصطناعي، مع التركيز على دوره في تطوير الهجمات وانعكاساته على أمن المعلومات والخصوصية.

## الفرع الأول

### توظيف الذكاء الاصطناعي في تطوير الهجمات السيبرانية

أدى التطور المتسارع في تقنيات الذكاء الاصطناعي إلى تغيير طبيعة التهديدات السيبرانية، إذ لم يعد الذكاء الاصطناعي مقتصرًا على الاستخدامات الدفاعية، وإنما أصبح من الممكن توظيف قدراته في تطوير الأساليب التي يستخدمها المهاجمون في الوصول إلى الأنظمة والشبكات والبيانات. وتكمن خطورة الذكاء الاصطناعي في هذا المجال في قدرته على أتمتة عدد من العمليات التي كانت تتطلب في السابق تدخلًا بشرياً كبيراً، فضلاً عن قدرته على تحليل المعلومات واكتشاف الأنماط وتوليد المحتوى بسرعة كبيرة. وهذا يمكن أن يؤدي إلى زيادة سرعة الهجمات واتساع نطاقها، كما يمكن أن يسمح للمهاجم بتكييف أسلوبه وفق طبيعة النظام المستهدف. ومن أخطر مظاهر هذا الاستخدام توظيف الذكاء الاصطناعي في الهجمات القائمة على الهندسة الاجتماعية<sup>(32)</sup>، إذ يمكن للأنظمة التوليدية إنتاج رسائل ومحتويات تبدو أكثر إقناعاً وملاءمة للضحية، الأمر الذي يزيد من احتمالية نجاح عمليات التصيد الإلكتروني والاحتيال الرقم كما يمكن استخدام تقنيات الذكاء الاصطناعي في جمع وتحليل المعلومات المتاحة عن الأشخاص والمؤسسات، ثم استخدامها لإنشاء هجمات تستهدف أفراداً محددين بصورة أكثر دقة<sup>(33)</sup>. ومن شأن ذلك أن يجعل الحدود بين الهجوم التقني والهجوم القائم على استغلال العنصر البشري أقل وضوحاً. ولا تقتصر الخطورة على زيادة كفاءة الهجمات التقليدية، بل تمتد إلى إمكانية ظهور أساليب هجومية جديدة تستفيد من قدرة الأنظمة الذكية على التكيف والتعلم. ولذلك أصبح الأمن السيبراني مطالباً بمواجهة خصم يمكن أن يستخدم الذكاء الاصطناعي لتحليل دفاعات النظام وتعديل أسلوب الهجوم بصورة مستمرة. ومن ثم، فإن الذكاء الاصطناعي لا يغير فقط أدوات الهجوم، وإنما يمكن أن يؤدي إلى تغيير طبيعة البيئة التي تقع فيها المواجهة السيبرانية، بحيث تصبح السرعة والأتمتة والقدرة على التكيف عناصر أساسية في تحديد قدرة المهاجم والمدافع على تحقيق أهدافهما<sup>(34)</sup>.

## الفرع الثاني

### التزييف العميق والتهديدات المرتبطة بالذكاء الاصطناعي التوليدي

يمثل الذكاء الاصطناعي التوليدي أحد أكثر مظاهر التطور التقني تأثيراً في البيئة السيبرانية، نظراً إلى قدرته على إنشاء محتوى جديد يشمل النصوص والصور والأصوات ومقاطع الفيديو والبرمجيات. وعلى الرغم من القيمة الإيجابية لهذه القدرات، فإنها قد تتحول إلى مصدر خطر عندما تستخدم في إنشاء محتوى مزيف أو مضلل بهدف ارتكاب عمليات احتيال أو انتحال شخصية أو التأثير في الأفراد والمؤسسات. ويعد التزييف العميق<sup>(35)</sup> من أبرز صور هذه المخاطر، إذ تسمح تقنيات الذكاء الاصطناعي بإنتاج أو تعديل صور وأصوات ومقاطع فيديو بدرجة عالية من الواقعية، بحيث يصبح من الصعب أحياناً على الشخص العادي التمييز بين المحتوى الحقيقي والمحتوى المصطنع. وتبرز الخطورة السيبرانية للتزييف العميق عندما يستخدم في انتحال هوية شخص معين بهدف الحصول على معلومات سرية أو تنفيذ عمليات مالية أو خداع الموظفين أو التأثير في أنظمة التحقق من الهوية. كما يمكن أن يستخدم في إنشاء محتوى مضلل يهدف إلى الإضرار بسمعة الأفراد والمؤسسات أو نشر معلومات كاذبة على نطاق واسع. ويزداد الخطر عندما تقترن تقنيات التزييف العميق بوسائل الهندسة الاجتماعية<sup>(36)</sup>، إذ يمكن للمهاجم الجمع بين المعلومات المتاحة عن الضحية وبين محتوى صوتي أو مرئي مصطنع لإنتاج سيناريو احتيالي أكثر إقناعاً. وبذلك ينتقل الخطر من مجرد إنتاج محتوى مزيف إلى استخدام هذا المحتوى كجزء من عملية هجومية متكاملة. ومن الناحية الأمنية، يفرض

هذا التطور تحدياً كبيراً على وسائل التحقق التقليدية، إذ لم يعد الاعتماد على الصوت أو الصورة وحدهما كافياً في جميع الحالات لإثبات هوية الشخص أو صحة المحتوى. ولذلك تزداد الحاجة إلى تطوير وسائل تحقق متعددة العوامل وآليات تقنية قادرة على اكتشاف المحتوى المصطنع والتأكد من مصدر البيانات وسلامتها. وتتفق هذه المخاطر مع الاتجاه الدولي الذي يؤكد ضرورة مراعاة الأمن والسلامة خلال دورة حياة أنظمة الذكاء الاصطناعي، إذ تؤكد اليونسكو ضرورة معالجة نقاط الضعف وإمكانات الهجوم المرتبطة بأنظمة الذكاء الاصطناعي، إلى جانب حماية البيانات والخصوصية<sup>(37)</sup>.

### الفرع الثالث

#### تسميم البيانات والهجمات على نماذج الذكاء الاصطناعي

لا تقتصر مخاطر الذكاء الاصطناعي على استخدامه كأداة لتنفيذ الهجمات السيبرانية، بل يمكن أن تكون أنظمة الذكاء الاصطناعي ذاتها هدفاً للهجوم. وتظهر هذه المسألة بصورة خاصة في الهجمات التي تستهدف البيانات المستخدمة في تدريب النماذج أو تشغيلها. ويعد تسميم البيانات<sup>(38)</sup> من أبرز هذه الهجمات، ويقوم على إدخال بيانات ضارة أو مضللة ضمن البيانات التي يعتمد عليها النظام في عملية التدريب، بهدف التأثير في سلوكه أو مخرجاته المستقبلية. وتكمن خطورة هذا النوع من الهجمات في أن النظام قد يستمر في العمل بصورة طبيعية ظاهرياً، في حين تكون نتائجه قد تأثرت بصورة غير ظاهرة بسبب البيانات التي تم التلاعب بها. كما يمكن أن تستهدف الهجمات نماذج الذكاء الاصطناعي من خلال إدخال مدخلات مصممة بصورة معينة لدفع النظام إلى إنتاج نتائج غير صحيحة<sup>(39)</sup>. ويعرف هذا النوع من الهجمات، في بعض تطبيقاته، بهجمات التهرب حيث يحاول المهاجم تعديل المدخلات بطريقة تجعل النظام يصنفها بصورة خاطئة. وقد وضع NIST تصنيفاً حديثاً لهجمات التعلم الآلي العدائي، شمل هجمات التهرب وتسميم البيانات وهجمات الخصوصية، فضلاً عن هجمات إساءة الاستخدام الموجهة إلى أنظمة الذكاء الاصطناعي التوليدي، مؤكداً أن هذه التهديدات قد تستهدف أنواعاً متعددة من أساليب التعلم الآلي<sup>(40)</sup>. (NIST) وتثير هذه الهجمات إشكالية خاصة في الأنظمة التي تعتمد عليها القطاعات الحساسة، لأن التأثير في مخرجات النظام قد يؤدي إلى نتائج تتجاوز الخطأ التقني البسيط، ولا سيما عندما يستخدم الذكاء الاصطناعي في مجالات الأمن أو الصحة أو الخدمات المالية أو البنية التحتية الحيوية. ومن هنا تظهر ضرورة النظر إلى البيانات بوصفها أحد المكونات الأساسية التي يجب حمايتها في منظومة الأمن السيبراني، إذ إن حماية النظام لا تتحقق بمجرد تأمين أجهزته وبرمجياته، وإنما يجب ضمان سلامة البيانات المستخدمة في تدريبه وتشغيله، والتأكد من مصادرها وجودتها وعدم تعرضها للتلاعب<sup>(41)</sup>.

### المطلب الثاني

#### التحديات التي تواجه الأمن السيبراني في ظل الذكاء الاصطناعي

أفرز الاستخدام المتزايد لتقنيات الذكاء الاصطناعي تحديات جديدة أمام الأمن السيبراني، نظراً لسرعة تطور الأدوات الذكية وقدرتها على تجاوز الأساليب التقليدية للحماية. كما أصبحت المؤسسات مطالبة بتطوير أنظمتها الدفاعية بصورة مستمرة لمواجهة الهجمات المتطورة، مع ضمان حماية البيانات والخصوصية والحد من إساءة استخدام هذه التقنيات. وعليه، يتناول هذا

المطلب أبرز التحديات التي تواجه الأمن السيبراني في ظل الذكاء الاصطناعي، ولا سيما تحدي سرعة تطور الهجمات السيبرانية وصعوبة التنبؤ بها والتصدي لها.

## الفرع الأول

### تحدي سرعة تطور الهجمات السيبرانية

يمثل التطور السريع في تقنيات الذكاء الاصطناعي أحد أبرز التحديات التي تواجه الأمن السيبراني، لأن سرعة تطور الأدوات الذكية قد تتجاوز في بعض الحالات قدرة المؤسسات على تطوير وسائل الحماية والاستجابة بالسرعة ذاتها<sup>(42)</sup>. ففي البيئة السيبرانية التقليدية كان تطوير الهجوم يحتاج في كثير من الحالات إلى قدر من المعرفة التقنية والوقت والجهد، بينما يمكن للذكاء الاصطناعي أن يسهم في أتمتة بعض مراحل الهجوم وتحليل البيانات وتوليد المحتوى واختيار الأساليب المناسبة بصورة أسرع<sup>(43)</sup>. ويؤدي ذلك إلى اختلال محتمل في التوازن بين المهاجم والمدافع؛ إذ يستطيع المهاجم الاستفادة من الذكاء الاصطناعي لتوسيع نطاق الهجمات، في حين يتعين على المؤسسات تطوير أنظمة دفاع قادرة على تحليل كميات ضخمة من البيانات والاستجابة بصورة مستمرة<sup>(44)</sup>. كما أن التطور المتواصل للنماذج الذكية يعني أن الوسائل الدفاعية التي تكون فعالة في مرحلة معينة قد تحتاج إلى تحديث مستمر لمواجهة أساليب هجومية جديدة. ولذلك أصبح الأمن السيبراني عملية مستمرة وليست مجموعة من الإجراءات الثابتة التي يمكن تطبيقها مرة واحدة. وتشير NIST إلى أن أمن ومرونة تقنيات الذكاء الاصطناعي يمثلان مجالاً سريع التغير، وأن طبيعة التحديات والحلول المحتملة تتطور بصورة مستمرة. (NIST) وعليه، فإن مواجهة هذا التحدي تتطلب بناء قدرات دفاعية مرنة قادرة على التعلم والتكيف، إلى جانب تحديث مستمر للأنظمة الأمنية والكوادر البشرية والسياسات والإجراءات<sup>(45)</sup>.

## الفرع الثاني

### تحدي حماية الخصوصية والبيانات

تُعد البيانات العنصر الأساسي في عمل معظم أنظمة الذكاء الاصطناعي، وهو ما يجعل حماية البيانات والخصوصية من أبرز التحديات المرتبطة باستخدام هذه التقنيات في المجال السيبراني فكلما زادت قدرة النظام على جمع البيانات وتحليلها، ازدادت أهمية ضمان عدم استخدامها بصورة غير مشروعة أو الكشف عنها أو تعديلها أو إعادة استخدامها لأغراض تختلف عن الغرض الذي جمعت من أجله كما أن الجمع بين مصادر متعددة للبيانات قد يؤدي إلى إمكانية استنتاج معلومات شخصية إضافية لم تكن ظاهرة بصورة مباشرة في البيانات الأصلي وتزداد خطورة ذلك في حالة تعرض أنظمة الذكاء الاصطناعي للاختراق<sup>(46)</sup>، لأن المهاجم قد لا يحصل على بيانات تقليدية فقط، بل يمكن أن يحصل على بيانات التدريب أو البيانات المستخدمة في تشغيل النظام أو المعلومات الناتجة عنه وقد أكدت اليونسكو أن حماية الخصوصية والبيانات يجب أن تمتد عبر دورة حياة نظام الذكاء الاصطناعي كاملة، وأن حماية البيانات تعد من المبادئ الأساسية للاستخدام المسؤول للذكاء الاصطناعي<sup>(47)</sup>. ومن ثم فإن التحدي لا يتمثل فقط في منع الاختراق، وإنما في تحقيق توازن بين الحاجة إلى البيانات اللازمة لعمل النظام وبين ضرورة

احترام خصوصية الأفراد وحماية معلوماتهم، الأمر الذي يستلزم وجود ضوابط قانونية وتقنية واضحة تحكم جمع البيانات وتخزينها ومعالجتها ومشاركتها(48)

### الفرع الثالث

#### تحدي الشفافية وصعوبة تفسير قرارات أنظمة الذكاء الاصطناعي

تواجه أنظمة الذكاء الاصطناعي، ولا سيما النماذج المعقدة، تحدياً يتعلق بمدى إمكانية تفسير الطريقة التي توصلت بها إلى نتيجة معينة. وتعرف هذه الإشكالية في الأدبيات التقنية والقانونية بمشكلة "الصندوق الأسود"(49)، حيث قد يكون من الصعب على المستخدم معرفة الأسباب التفصيلية التي أدت إلى صدور المخرج أو القرار وتكتسب هذه المشكلة أهمية خاصة في مجال الأمن السيبراني، لأن الاعتماد على نظام ذكي في اكتشاف التهديدات أو تصنيف الأنشطة أو اتخاذ إجراءات دفاعية يتطلب قدراً مناسباً من الثقة في نتائج النظام فإذا قام النظام بتصنيف نشاط معين على أنه تهديد، فمن المهم أن تكون لدى الجهة المسؤولة القدرة على فهم الأساس الذي استند إليه هذا التصنيف(50)، ولا سيما إذا ترتب عليه اتخاذ إجراء يؤثر في مستخدم أو مؤسسة أو خدمة رقمية. كما أن صعوبة التفسير قد تعيق عملية التحقيق في الحوادث السيبرانية، إذ إن تحديد سبب الخطأ أو تحديد الكيفية التي تم من خلالها تجاوز النظام الدفاعي يتطلب إمكانية تتبع مراحل اتخاذ القرار. ولهذا أكدت اليونسكو أهمية الشفافية وقابلية التفسير والمساءلة والتدقيق والتتبع في أنظمة الذكاء الاصطناعي(51)، مع التأكيد على أن الرقابة البشرية لا ينبغي أن تختفي نتيجة استخدام الأنظمة الذكية(52). وعليه، فإن تعزيز الأمن السيبراني باستخدام الذكاء الاصطناعي لا يتطلب فقط زيادة قدرة النظام على اكتشاف التهديدات، وإنما يتطلب كذلك أن يكون النظام قابلاً للتدقيق والتقييم والفهم بالقدر الذي يسمح بتحديد أسباب النتائج الصادرة عنه ومراجعتها عند الحاجة(53).

يتضح مما تقدم أن العلاقة بين الذكاء الاصطناعي والأمن السيبراني لا تخلو من مفارقة أساسية؛ فالذكاء الاصطناعي يمثل وسيلة متقدمة لتعزيز قدرات الدفاع السيبراني، إلا أنه في الوقت ذاته يضيف طبقة جديدة من المخاطر والتحديات إلى البيئة الرقمية. فمن جهة، يمكن توظيفه في تحليل البيانات واكتشاف الأنشطة غير الطبيعية والتنبيه بالتهديدات وتسريع الاستجابة للحوادث، ومن جهة أخرى يمكن استغلاله في تطوير الهجمات السيبرانية، وإنشاء المحتوى الاحتيالي والمضلل، وتنفيذ عمليات الهندسة الاجتماعية، فضلاً عن إمكانية استهداف أنظمة الذكاء الاصطناعي نفسها من خلال تسميم البيانات أو التلاعب بالمدخلات والنماذج. ومن ثم فإن مواجهة المخاطر المرتبطة بالذكاء الاصطناعي لا يمكن أن تتحقق من خلال الحلول التقنية وحدها، وإنما تتطلب منظومة متكاملة تجمع بين التدابير التقنية والقانونية والتنظيمية، وتضمن سلامة البيانات والنماذج، وتعزز الشفافية والمساءلة والرقابة البشرية.

## المراجع:

- ابن عثمان، أ. ع. (د.ت.). أحكام تطبيقات الذكاء الاصطناعي في القضاء (الطبعة الأولى).
- الحيدري، ج. (2012). الجرائم الإلكترونية وسبل معالجتها. مكتبة السهوري.
- القاضي، ز. ع. (2010). مقدمة في الذكاء الاصطناعي (الطبعة الأولى). دار الصفاء للطباعة والنشر والتوزيع.
- عبد الهادي، ز. (2019). الأنظمة الخبيرة للذكاء الاصطناعي. دار الكتب للنشر والتوزيع.
- مجيد، س. ف. (2024). جريمة التتمر الإلكتروني: دراسة في القانون العراقي والأمريكي. المجلة الأكاديمية للبحث القانوني، 11 (4).
- غانم، ع. ح. (2006). الوسائل العلمية الحديثة لكشف الجريمة ومشروعيتها وحجبتها. المجلة العربية للدفاع الاجتماعي، (4).
- عبد النور، ع. (2006). مدخل إلى عالم الذكاء الاصطناعي. مدينة الملك عبد العزيز للعلوم والتقنية.
- عبد النور، ع. (2005). أساسيات الذكاء الاصطناعي (الطبعة الأولى). دار الفیصل الثقافية.
- جمال، ع. س. (2019). أثر تطبيقات الذكاء الاصطناعي في رفع كفاءة الأداء الأمني. أكاديمية الشرطة، كلية الدراسات العليا.
- يوسف، ك. (2020). المسؤولية المدنية عن فعل الذكاء الاصطناعي [رسالة ماجستير، الجامعة اللبنانية، كلية الحقوق].
- محمد، ف. ع. (2021). استخدامات الذكاء الاصطناعي بين المشروعية وعدم المشروعية. مجلة المركز القومي للبحوث الاجتماعية، 64.
- محمد، م. س. د. (2017). الذكاء الاصطناعي والحياة في عام 2030. مركز المستقبل، (303).
- محمد، و. س. د. (2022). المسؤولية الجنائية الناشئة عن تطبيقات الذكاء الاصطناعي. مجلة العلوم القانونية، جامعة عين شمس، كلية الحقوق.
- دهشان، ي. إ. (2019). المسؤولية الجنائية في جرائم الذكاء الاصطناعي.
- الهادي، م. م. (2005). الجريمة الإلكترونية عبر شبكة الإنترنت (الطبعة الأولى). الدار المصرية اللبنانية.
- القحطاني، م. ب. ع. (د.ت.). الذكاء الاصطناعي والأمن السيبراني: التحديات والفرص. دراسة منشورة في مجال الدراسات الأمنية والتقنية.
- المطيري، ع. ب. م. (د.ت.). الأمن السيبراني في ظل التحول الرقمي والذكاء الاصطناعي. دراسة في الأمن المعلوماتي والتحول الرقمي.
- رهيف، ز. ج. (2026). الأمن السيبراني وتأثيره في العراق. المجلة العراقية للعلوم السياسية، جامعة بغداد.

الخبزي، ب. ع. (2024). توظيف الأمن السيبراني لاستثمار فرص الذكاء الاصطناعي والحد من تحدياته في مجال التعليم والتعلم. *المجلة العربية للدراسات الأمنية*، 40(2)، 235-249.

Tabassi, E. (2023). *Artificial intelligence risk management framework (AI RMF 1.0)* (NIST AI 100-1). National Institute of Standards and Technology

National Institute of Standards and Technology. (2023). *Artificial intelligence risk management framework: AI RMF 1.0*. U.S. Department of Commerce

National Institute of Standards and Technology. (2023). *AI risk management framework .playbook*. U.S. Department of Commerce

National Institute of Standards and Technology. (2024). *Artificial intelligence risk management framework: Generative artificial intelligence profile*. U.S. Department of Commerce

National Institute of Standards and Technology. (2024). *AI risk management framework .resources*

European Union Agency for Cybersecurity. (2023). *Artificial intelligence and .cybersecurity research*. European Union

European Union Agency for Cybersecurity. (2023). *Artificial intelligence and next .generation technologies*. European Union

European Union Agency for Cybersecurity. (2023). *Multilayer framework for good .cybersecurity practices for AI*. European Union

,Brundage, M., et al. (2018). *The malicious use of artificial intelligence: Forecasting .prevention, and mitigation*

Camacho, J. M., Couce-Vieira, A., & Ríos Insua, D. (2024). A cybersecurity risk analysis .framework for systems with artificial intelligence components

Erukude, S. T., Marella, V. C., & Veluru, S. R. (2026). AI-driven cybersecurity threats: A .survey of emerging risks and defensive strategies

.United Nations Educational, Scientific and Cultural Organization. (2021)  
*.Recommendation on the ethics of artificial intelligence*. UNESCO

Council of Europe. (2024). *Framework convention on artificial intelligence and human .rights, democracy and the rule of law*. Strasbourg, France

European Union. (2024). *Regulation (EU) 2024/1689 of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)*. Official Journal of the European Union

.World Economic Forum. (2024). *Global cybersecurity outlook*. Geneva, Switzerland

.International Telecommunication Union. (n.d.). *Global cybersecurity index*. Geneva .Switzerland

## الهوامش

- 1( ) اروي عبد الرحمن بن عثمان ,احكام تطبيقات الذكاء الاصطناعي في القضاء , الطبعة الأولى
- 2( جمال الحيدري , الجرائم الالكترونية وسبل معالجتها , مكتبة السنهوري , 2012
- 3( زياد عبد الكريم القاضي , مقدمة في الذكاء الاصطناعي , ط1, دار الصفاء للطباعة و النشر والتوزيع , 2010
- 4( زين عبد الهادي , الأنظمة الخبيرة للذكاء الاصطناعي , دار الكتب للنشر و التوزيع , القاهرة , 2019
- 5( سحر فؤاد مجيد , جريمة التتمر الالكتروني (دراسة في القانون العراقي الامريكي), المجلة الاكاديمية للبحث القانوني , المجلد 11 , العدد الرابع , 2024
- 6( عادل حافظ غانم , الوسائل العلمية الحديثة لكشف الجريمة ومشروعيتها وحجيتها , المجلة العربية للدفاع الاجتماعي , عدد الرابع , 2006
- 7( عادل عبد النور , مدخل إلى عالم الذكاء الاصطناعي , مدينة الملك عبد العزيز للعلوم و التقنية , السعودية , 2006
- 8( عبد النور عادل . أساسيات الذكاء الاصطناعي . ( الرياض , دار الفیصل الثقافية , الطبعة الاولى , 2005م
- 9( عمرو سيد جمال , اثر تطبيقات الذكاء الاصطناعي في رفع كفاءة الاداء الأمني , اكااديمية الشرطة ,كلية الدراسات العليا , القاهرة , 2019
- 10( كرستيان يوسف , المسؤولية المدنية عن فعل الذكاء الاصطناعي , رسالة ماجستير مقدمة إلى مجلس كلية الحقوق , الجامعة اللبنانية , 2020
- 11( فايق عوضين محمد , استخدامات الذكاء الاصطناعي بين المشروعية وعدم المشروعية ,مجلة المركز القومي للبحوث الاجتماعية , مجلد 64 , عام 2021
- 12( محمد سعد الدين محمد , الذكاء الاصطناعي و الحياة في عام 2030 , مركز المستقبل , العدد 303 , 2017
- 13( وليد سعد الدين محمد , المسؤولية الجنائية الناشئة عن تطبيقات الذكاء الاصطناعي , مجلة العلوم القانونية , جامعة عين شمس -كلية الحقوق , 2022
- 14( يحيى ابراهيم دهشان , المسؤولية الجنائية في جرائم الذكاء الاصطناعي , 2019
- 15( الهادي محمد محمد . الجريمة الالكترونية عبر شبكة الانترنت . ( القاهرة , الدار المصرية اللبنانية , الطبعة الاولى , 2005م )
- 16( National Institute of Standards and Technology (NIST), Artificial Intelligence Risk Management Framework
- 17( محمد بن عبد الله القحطاني, «الذكاء الاصطناعي والأمن السيبراني: التحديات والفرص», دراسة منشورة في مجال الدراسات الأمنية والتقنية
- 18( عبد الله بن محمد المطيري, «الأمن السيبراني في ظل التحول الرقمي والذكاء الاصطناعي», دراسة في الأمن المعلوماتي والتحول الرقمي.

- (19) م.م. زهراء جبار رهيف، «الأمن السيبراني وتأثيره في العراق»، المجلة العراقية للعلوم السياسية، جامعة بغداد، 2026.
- (20) بدر عدنان الخبيزي، «توظيف الأمن السيبراني لاستثمار فرص الذكاء الاصطناعي والحد من تحدياته في مجال التعليم والتعلم»، المجلة العربية للدراسات الأمنية، المجلد 40، العدد 2، 2024، ص 235-249.
- (21) Elham Tabassi, Artificial Intelligence Risk Management Framework (AI RMF 1.0), National Institute of Standards and Technology, NIST AI 100-1, Gaithersburg, Maryland, p. 1, 2023.
- (22) Elham Tabassi, Artificial Intelligence Risk Management Framework (AI RMF 1.0), op. cit., pp. 10-14.
- (23) National Institute of Standards and Technology (NIST), AI Risk Management Framework Playbook, U.S. Department of Commerce, 2023.
- (24) National Institute of Standards and Technology (NIST), Artificial Intelligence Risk Management Framework, 2023, pp. 20-25.
- (25) European Union Agency for Cybersecurity (ENISA), Artificial Intelligence and Cybersecurity Research, European Union, 2023, pp. 5-8.
- (26) European Union Agency for Cybersecurity (ENISA), Artificial Intelligence and Next Generation Technologies, European Union, 2023.
- (27) Miles Brundage et al., The Malicious Use of Artificial Intelligence: Forecasting, Prevention, and Mitigation, 2018, pp. 1-10.
- (28) José Manuel Camacho, Aitor Couce-Vieira and David Ríos Insua, "A Cybersecurity Risk Analysis Framework for Systems with Artificial Intelligence Components", 2024.
- (29) Sai Teja Erukude, Viswa Chaitanya Marella and Suhasnadh Reddy Veluru, "AI-Driven Cybersecurity Threats: A Survey of Emerging Risks and Defensive Strategies", 2026.
- (30) UNESCO, Recommendation on the Ethics of Artificial Intelligence, United Nations Educational, Scientific and Cultural Organization, Paris, 2021.
- (31) Council of Europe, Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, Strasbourg, 2024.
- (32) European Union, Regulation (EU) 2024/1689 of the European Parliament and of the Council Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act), Official Journal of the European Union, 2024.
- (33) World Economic Forum, Global Cybersecurity Outlook, Geneva, 2024.
- (34) International Telecommunication Union (ITU), Global Cybersecurity Index, Geneva.
- (35) National Institute of Standards and Technology (NIST), Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile, 2024.
- (36) National Institute of Standards and Technology (NIST), AI Risk Management Framework Playbook, 2023.
- (37) European Union Agency for Cybersecurity (ENISA), Multilayer Framework for Good Cybersecurity Practices for AI, 2023.
- (38) National Institute of Standards and Technology (NIST), Artificial Intelligence Risk Management Framework, 2023.
- (39) UNESCO, Recommendation on the Ethics of Artificial Intelligence, Paris, 2021.

- (40) بدر عدنان الخبيزي، المصدر السابق، ص 237.
- (41) محمد ميسر فتحي، المصدر السابق، ص 5-8.
- (42) بدر عدنان الخبيزي، المصدر السابق، ص 239.
- (43) المنظمة العربية للتنمية الإدارية، الذكاء الاصطناعي وتطبيقاته في المؤسسات الحكومية، جامعة الدول العربية، القاهرة.
- (44) عبد الله بن محمد المطيري، «الأمن السيبراني في ظل التحول الرقمي والذكاء الاصطناعي»، دراسة في الأمن المعلوماتي والتحول الرقمي.
- (45) بدر عدنان الخبيزي، المصدر السابق، ص 239.
- (46) Council of Europe, Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, 2024.
- (47) UNESCO, Recommendation on the Ethics of Artificial Intelligence, 2021.
- (48) International Telecommunication Union (ITU), Global Cybersecurity Index.
- (49) European Union Agency for Cybersecurity (ENISA), publications on Artificial Intelligence and Cybersecurity.
- (50) National Institute of Standards and Technology (NIST), AI Risk Management Framework Resources, 2024.
- (51) National Institute of Standards and Technology (NIST), Artificial Intelligence Risk Management Framework (AI RMF 1.0), 2023.
- (52) زهراء جبار رفيف، «الأمن السيبراني وتأثيره في العراق»، المجلة العراقية للعلوم السياسية، 2026.
- (53) Council of Europe, Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, Strasbourg, 2024.