

الذكاء الاصطناعي والأمن السيبراني

إعداد: الباحثة / سينتيا حنا | الجمهورية اللبنانية
طالبة دكتوراه في الحقوق / القانون الخاص | الجامعة الإسلامية في لبنان
E-mail: cynthiaharbb123@gmail.com | <https://orcid.org/0009-0005-1214-7274>

إشراف: البروفيسورة / أودين سلوم الحايك⁽¹⁾ | الجمهورية اللبنانية
دكتوراه في الحقوق / القانون الخاص | الجامعة اللبنانية
E-mail: saudine@yahoo.com | <https://orcid.org/0009-0009-6883-8807>

<https://doi.org/10.70758/elqarar/11.32.13>

تاريخ النشر: 2026/08/15	تاريخ القبول: 2026/08/11	تاريخ الاستلام: 2026/07/18
-------------------------	--------------------------	----------------------------

للاقتباس: سينتيا، حنا، (2026)، الذكاء الاصطناعي والأمن السيبراني، إشراف البروفيسورة أودين سلوم الحايك، مجلة القرار للبحوث العلمية المحكمة، المجلد الحادي عشر، العدد 32، السنة الثالثة، ص-ص 288-308. <https://doi.org/10.70758/elqarar/11.32.13>

المُستخلص

تتناول هذه الدراسة العلميّة العلاقة المتشابكة والمتطوّرة بين الذكاء الاصطناعي والأمن السيبراني، في ضوء أحدث التقارير والأبحاث الدوليّة للفترة 2025-2026. وتتطّلق الدراسة من إشكاليّة جوهريّة مفادها: كيف يُعيد الذكاء الاصطناعيّ رسم ملامح المشهد السيبرانيّ، وما الفرص التي يتيحها لتعزيز الأمن، وما التهديدات الجديدة التي يفرزها في الوقت ذاته. كما اعتمدت الدراسة على المنهجين الوصفي والتحليلي المقارن، مُستندة إلى قاعدة بيانات شاملة من المصادر الموثقة. وقد عالجت المفاهيم الأساسيّة والإطار التنظيمي والقانوني وطبيعة التهديدات السيبرانيّة للذكاء الاصطناعيّ، إلى جانب إستراتيجيّات الدفاع لمواجهتها. ومن ثمّ، تخلصت الدراسة إلى نتائج جوهريّة أبرزها: إرتفاع الهجمات السيبرانيّة المدعومة بالذكاء الاصطناعيّ بنسبة 72% عام 2025 لتبلغ تكلفة الاختراق الواحد نحو 5.72 مليون دولار، بينما يوفر الدفاع السيبرانيّ الذكيّ وفورات ماليّة بنحو 1.9 مليون دولار ويقلّص دورة الاختراق إلى 80 يوماً. فضلاً عن توقع نمو سوق الأمن السيبرانيّ بالذكاء الاصطناعيّ من 28.51 مليار دولار حالياً إلى 136 مليار دولار بحلول عام 2032. وبناءً على ذلك، أوصت الدراسة بتبني نموذج «عدم الثقة المطلقة»، وتطوير الإستجابة الآليّة للتهديدات، والتهيؤ المُبكر لمرحلة التشفير ما بعد الكمي.

الكلمات المفتاحيّة: الذكاء الاصطناعي، الأمن السيبراني، التهديدات السيبرانيّة، التعلم الآلي، التزييف العميق، الهجمات الإلكترونيّة، الدفاع السيبراني، نماذج اللغة الكبيرة، الحوسبة الكموميّة، عدم الثقة المطلقة.

(1) تتركز أبحاثه في مجالات القانون الخاص وتكنولوجيا المعلومات والذكاء الاصطناعي والأمن السيبراني.

Artificial intelligence and cybersecurity

Author: Researcher / Cynthia Hanna | Lebanese Republic
PhD Candidate in Law / Private Law | Islamic University of Lebanon

E-mail: cynthiaharbb123@gmail.com | <https://orcid.org/0009-0005-1214-7274>

Supervised by: Prof. Dr. / Audine Salloum El-Hayek⁽¹⁾ | Lebanese Republic
PhD in Law / Private Law | Lebanese University

E-mail: saudine@yahoo.com | <https://orcid.org/0009-0009-6883-8807>

<https://doi.org/10.70758/elqarar/11.32.13>

Received : 18/07/2026

Accepted : 11/08/2026

Published : 15/08/2026

Cite this article as: Hanna, Cynthia, (2026), Artificial intelligence and cybersecuritys, Supervised by: Prof. Dr. / Audine Salloum El-Hayek, ElQarar Journal for Peer-Reviewed Scientific Research, vol II, issue 32, Third year, pp. 288-308. <https://doi.org/10.70758/elqarar/11.32.13>

Abstract

This scientific study examines the comprehensive and evolving relationship between artificial intelligence and cybersecurity, in light of the latest international reports and research for the 2025-2026 period. The study proceeds from a core research problem: how is artificial intelligence reshaping the cybersecurity landscape, what opportunities does it offer for strengthening security, and what new threats does it simultaneously generate?

Methodologically, the study adopts a descriptive and comparative analytical approach, drawing on a broad and comprehensive database of documented primary and secondary sources. The study addresses the fundamental concepts related to the topic, alongside the governing regulatory and legal framework, while also analyzing the cybersecurity threats arising from the use of artificial intelligence, along with effective strategies for defending against them.

The study ultimately arrives at several key findings, most notably: first, that AI-enhanced cyberattacks rose sharply by 72% during 2025 alone, accompanied by a rise in the average cost of a security breach to approximately 5.72 million dollars. Second, that adopting intelligent cybersecurity defense systems saves an estimated 1.9 million dollars per organization, while also reducing the average breach lifecycle to just 80 days. Third, that the AI-driven cybersecurity market is on track for exceptional growth, expected to rise from 28.51 billion dollars currently to approximately 136 billion dollars by 2032.

Based on these findings, the study recommends adopting a security model built on the “Zero Trust” principle, alongside developing automated response capabilities for emerging cyber threats, and preparing well in advance for the future transition to post-quantum encryption standards.

Keywords: Artificial Intelligence, Cybersecurity, Cyber Threats, Machine Learning, Deep Fake, Cyberattacks, Cyber Defense, Large Language Models, Quantum Computing, Zero Trust.

(1) His research focuses on the fields of private law, information technology, artificial intelligence, and cybersecurity.

1. مقدمة

يقف العالم اليوم أمام مفترق طرق حضاري فارق، تتشابك فيه تداعيات ثورتين متزامنتين: ثورة الذكاء الاصطناعي التي تُعيد رسم حدود الممكن في شتى ميادين الحياة، وثورة التحول الرقمي التي أحالت الفضاء الإلكتروني إلى ساحة المواجهة الأولى بين الدول والمؤسسات والأفراد. وفي قلب هذا التقاطع المتسارع، يتبلور سؤال استراتيجي محوري: هل يُمثل الذكاء الاصطناعي درعاً أم حيزاً جديداً للمخاطر في عالم الأمن السيبراني؟

لم يعد الأمن السيبراني شأنًا تقنيًا بحتاً يُعالج في غرف الخوادم وحجرات تكنولوجيا المعلومات، بل غدا قضية وطنية واستراتيجية تتصدر جدول أعمال الحكومات والمؤسسات والأفراد على حدٍ سواء. وقد أكد تقرير المنتدى الاقتصادي العالمي حول الأمن السيبراني لعام 2025 أن المشهد السيبراني بلغ مستوى غير مسبوق من التعقيد، نتيجة تصاعد التوترات الجيوسياسية، والتطور المتسارع للتقنيات الناشئة، وتعمق الفجوات الأمنية بين الدول والمؤسسات الكبيرة والصغيرة (World Economic Forum, 2025).

وقد جاءت الثورة في نماذج الذكاء الاصطناعي التوليدية ونماذج اللغة الكبيرة لتضاعف من حدة هذا التحدي؛ فلأول مرة في تاريخ الجريمة السيبرانية، بات بإمكان أفراد لا يمتلكون أي خلفية تقنية أن يشنوا هجمات متطورة كانت حتى وقت قريب حكرًا على فرق متخصصة. ففي فبراير 2025، استطاع ثلاثة مراهقين يابانيين، بأعمار 14 و15 و16 عاماً، دون خبرة برمجية سابقة، تطوير برنامج بمساعدة ChatGPT تستهدف أنظمة Rakuten Mobile نحو 220 ألف مرة، وتمكنوا من الحصول على أكثر من 2500 عقد احتيالي، في حادثة لا تُعدّ استثناءً بل مقدمة لنمط ممنهج آخذ في التوسع (The Peninsula Qatar, 2025).

وفي ضوء هذه المعطيات المتحولة، تأتي هذه الدراسة لتسدّ فراغاً في الأدبيات العلمية العربية من خلال تقديم رؤية منهجية شاملة ومحدّثة تجمع بين الأبعاد التقنية والاستراتيجية والتشريعية لهذه المعادلة المركبة. وتتجلى أهمية الدراسة في كونها تستند إلى أحدث البيانات والتقارير الدولية للفترة 2025-2026، مما يجعل نتائجها وتوصياتها ذات صلة مباشرة بالواقع الراهن.

إن الرهان الحقيقي في هذه المعادلة لا يكمن في امتلاك الذكاء الاصطناعي، بل في حسن توظيفه؛ فالمؤسسات والدول التي ستجح في بناء منظومات أمنية ذكية ومرنة وقادرة على التكيف هي تلك التي تُدرك أن الذكاء الاصطناعي ليس سلاحاً وحيداً بل منظومة متكاملة تستوجب تقنية راسخة وحوكمة واضحة وإنساناً مدرباً.

من هنا يطرح هذا البحث الإشكالية التالية:

إلى أي مدى أسهم تطور الذكاء الاصطناعي في إعادة تشكيل مفهوم الأمن السيبراني وما أبرز التحديات القانونية والتقنية التي ترتبت على هذا التطور؟

سنعالج هذه الإشكالية من خلال البحث في الذكاء الاصطناعي في عالم الأمن السيبراني المفاهيم والأسس (2) لنتعمق فيما بعد في التهديدات السيبرانية المدعومة بالذكاء الاصطناعي وسبل الدفاع (3).

2. الذكاء الاصطناعي في عالم الأمن السيبراني المفاهيم والأسس

يشهد العالم اليوم ثورة حقيقية في مجال الذكاء الاصطناعي، إذ باتت تقنياته تتغلغل في مختلف جوانب الحياة الرقمية، ولا سيما في ميدان الأمن السيبراني. وقد أفرز هذا التشابك العميق بين الذكاء الاصطناعي والفضاء السيبراني معطيات جديدة تستوجب الدراسة والتحليل المعمق، سواء أكانت تتعلق بالفرص التي تتيحها هذه التقنيات أم بالمخاطر التي تُقضي إليها. وقد أكد تقرير المنتدى الاقتصادي العالمي للأمن السيبراني لعام 2025 (World Economic Forum, 2025)، أن الجريمة الإلكترونية نمت في تواترها وتعقيدها، مع تصاعد ملحوظ في الهجمات المنهجية المعززة بتقنيات الذكاء الاصطناعي.

2.1 مفهوم الذكاء الاصطناعي وتطبيقاته في الأمن السيبراني

يُمثل هذا القسم الركيزة المفاهيمية الأولى للدراسة، إذ لا يمكن استيعاب عمق التقاطع بين الذكاء الاصطناعي والأمن السيبراني دون تأسيس فهم دقيق ومحدث لطبيعة الذكاء الاصطناعي بوصفه حقلاً علمياً وتقنياً متطوراً. وقد شهدت السنوات الأخيرة، ولا سيما المرحلة الممتدة بين عامي 2023 و2026، قفزة نوعية في قدرات الذكاء الاصطناعي التوليدي جعلت كثيراً من التعريفات التقليدية قاصرة عن استيعاب هذه الطفرة. ففي حين اقتصرَت التعريفات التقليدية على وصف الذكاء الاصطناعي بأنه "محاكاة للقدرات الإنسانية"، باتت التعريفات المعاصرة تُضيف إليه بُعد "التجاوز"؛ أي القدرة على أداء مهام تفوق الإنسان في السرعة والدقة والمعالجة المتوازية. وهذا البُعد تحديداً هو ما يجعل الذكاء الاصطناعي محورياً لا غنى عنه في منظومة الأمن السيبراني الحديثة، ويكتسب هذا القسم أهمية مضاعفة في ضوء ما كشفه مؤشر Cisco للجاهزية السيبرانية لعام 2025، من أن 52% من المؤسسات أفادت بأن موظفيها لا يفهمون تماماً الطريقة التي يستغل بها المهاجمون تقنيات الذكاء الاصطناعي، في حين أن 45% فقط منها ترى أن لديها الفهم الداخلي الكافي لتقييم هذه المخاطر، وهو ما يعكس فجوة مفاهيمية حقيقية تسبق أي عائق يتعلق بشح الموارد. ومن هذا المنطلق يسعى هذا القسم إلى تقديم إجابات محدثة لثلاثة أسئلة جوهرية: ما الذكاء الاصطناعي في تجلياته التقنية الراهنة؟ وما مكوناته الوظيفية ذات الصلة بالأمن السيبراني؟ وكيف تطور توظيفه في هذا المجال خلال المرحلة 2023-2026 التي أنجزت فيها هذه الدراسة؟

2.1.1 تعريف الذكاء الاصطناعي ومكوناته الأساسية

يُعرّف الذكاء الاصطناعي بوصفه مجموعة من التقنيات الحاسوبية التي تُمكن الأنظمة من محاكاة القدرات المعرفية الإنسانية كالتعلم والاستنتاج وحل المشكلات واتخاذ القرار. ويقوم هذا العلم على ركائز جوهرية تشمل: تعلم الآلة (Machine Learning)، والتعلم العميق (Deep Learning)، ومعالجة اللغة الطبيعية (NLP)، فضلاً عن الرؤية الحاسوبية (Computer Vision). وتُشكّل هذه المكونات مجتمعةً البنية التحتية التي تركز عليها منظومة الذكاء الاصطناعي في تطبيقاتها الأمنية.

وتتوزع تطبيقات الذكاء الاصطناعي في ميدان الأمن السيبراني على محورين أساسيين: الأول يتعلق بالاستخدام الدفاعي، ويشمل الكشف عن التهديدات وتحليل السلوكيات الشاذة، فيما يختص الثاني بالاستخدام الهجومي، الذي بات يُوظف من قِبَل الجهات المعادية لتصميم هجمات أكثر تعقيداً وإحكاماً.

وقد أفضى هذا الوضع إلى ما بات يُعرف بـ“سباق التسلح السيبراني المدعوم بالذكاء الاصطناعي”. وتجدر الإشارة إلى أن الذكاء الاصطناعي التوليدي (Generative AI) يمثل إضافة نوعية في هذه المعادلة؛ إذ أثبتت نماذج اللغة الكبيرة (LLMs) قدرتها على التأثير في المشهد الأمني بشكل متسارع وغير مسبوق. وقد أشار تقرير شركة Darktrace لعام 2025 إلى أن 78% من كبار مسؤولي أمن المعلومات (CISOs) على مستوى العالم يُقرّون بأن التهديدات السيبرانية المدعومة بالذكاء الاصطناعي باتت ذات تأثير بالغ على مؤسساتهم، مسجلة ارتفاعاً بنسبة 5% عن العام السابق (Darktrace, 2025).

كما تكشف بيانات شركة Total Assure المنشورة في يناير 2026 أن الهجمات الإلكترونية المعززة بالذكاء الاصطناعي ارتفعت بنسبة 72% مقارنةً بعام 2024، فيما قفزت هجمات التصيد الاحتيالي (Phishing) بنسبة مذهلة بلغت 1265% بفعل أدوات الذكاء الاصطناعي التوليدي (Total Assure, 2026). وقد بلغ المتوسط العالمي لتكلفة الاختراق الأمني 4.44 مليون دولار للحادثة الواحدة عام 2025 وفق تقرير IBM، فيما ارتفعت التكلفة في الولايات المتحدة إلى 10.22 مليون دولار، مما يجعل الاستثمار في حلول الأمن السيبراني أمراً لا مñas منه (IBM, 2025).

2.1.2 تطور الأمن السيبراني في ظل الذكاء الاصطناعي

شهد مجال الأمن السيبراني تطوراً متسارعاً منذ مطلع الألفية الثالثة، غير أن التحول الأعماق والأشد تأثيراً جاء مع دخول الذكاء الاصطناعي عصر النضج التقني بين عامي 2023 و2025. فقد انتقل الأمن السيبراني من الاعتماد على الأنظمة القائمة على التوقيع (Signature-based Systems) إلى المنظومات الذكية القادرة على التعلم الذاتي والاستجابة الآنية.

ومن المراكز الحديثة في تطور منظومة الأمن السيبراني إطلاق شركة IBM في أبريل 2025 لنظام ”(ATOM” (Autonomous Threat Operations Machine)، وهو نظام ذكاء اصطناعي وكيل (Agentic AI) يوفر عمليات فرز وتحقيق ومعالجة تلقائية للتهديدات بتدخل بشري أدنى ما يمكن (IBM, 2025).

وعلى الصعيد الاقتصادي، تكشف تقديرات السوق أن قطاع الأمن السيبراني المعزز بالذكاء الاصطناعي في طريقه للنمو من 28.51 مليار دولار عام 2025 إلى 136.18 مليار دولار بحلول عام 2032، بمعدل نمو سنوي مركب يبلغ 24.81%، وهو ما يعكس حجم الرهان العالمي على هذا القطاع الاستراتيجي. وفي السياق ذاته، أكد تقرير ”IBM X-Force Threat Intelligence Index“ الصادر في فبراير 2026 أن استغلال الثغرات الأمنية يتصدر أسباب الهجمات بنسبة 40% من الحوادث المرصودة خلال عام 2025 (IBM, 2026).

وفيما يتعلق بتأثير الذكاء الاصطناعي على الكوادر البشرية في قطاع الأمن السيبراني، يكشف تقرير Cisco للجهازية السيبرانية 2025 أن 86% من قادة الأمن السيبراني واجهوا حادثة أمنية واحدة على الأقل مرتبطة بالذكاء الاصطناعي خلال العام الماضي، في حين أن 52% أفادوا بأن موظفيهم لا يفهمون تماماً كيف يستغل المهاجمون تقنيات الذكاء الاصطناعي، ولا يرى سوى 45% منهم أن لديهم

الفهم الداخلي الكافي لتقييم هذه المخاطر (Cisco, 2025). ويتقاطع هذا مع ما كشفه تقرير Dark-trace لعام 2025 من أن 78% من كبار مسؤولي أمن المعلومات (CISOs) على مستوى العالم يُقرّون بأن التهديدات السيبرانية المدعومة بالذكاء الاصطناعي باتت ذات تأثير بالغ على مؤسساتهم، مسجلة ارتفاعاً بنسبة 5% عن العام السابق (Darktrace, 2025).

2.2 الإطار التنظيمي والقانوني الحاكم للذكاء الاصطناعي والأمن السيبراني

الإطار التنظيمي والقانوني الحاكم للذكاء الاصطناعي والأمن السيبراني أحد أعقد الملفات التي تواجه صانعي السياسات اليوم، وذلك لأن التشريع في هذا المجال يستوجب التوفيق بين ثلاثة متطلبات متوترة في ما بينها: حماية حقوق الأفراد وصون الخصوصية، وضمان الأمن الوطني والمؤسسي، وتمكين الابتكار التكنولوجي وعدم تعطيله. ويرى تقرير المنتدى الاقتصادي العالمي للأمن السيبراني 2025 (World Economic Forum, 2025) أن الدول التي نجحت في تحقيق هذا التوازن هي تلك التي تبنت مقاربة "الحوكمة التكيفية" المرتكزة على مبادئ ثابتة مع مرونة إجرائية تستوعب التطور المتسارع للتقنية، لا التي اختارت التشريع الجامد أو الفراغ التنظيمي.

وتزداد أهمية هذا الملف التنظيمي في ضوء معطين متلازمين: أولهما أن الجرائم السيبرانية المعززة بالذكاء الاصطناعي لا تعترف بالحدود الجغرافية، مما يجعل التشريع الوطني المنعزل قاصراً بطبيعته؛ وثانيهما أن المهاجمين يستغلون الفجوات بين الأطر التشريعية الوطنية المتباينة ملجأً آمناً لتنظيم عملياتهم. ويُقدّر تقرير Cisco للجهازية السيبرانية 2025 (Cisco, 2025) أن 43% من الهجمات الكبرى المؤنّفة في عام 2024 استغلت ثغرات قانونية تنظيمية وليس ثغرات تقنية فحسب، مما يؤكد أن التشريع الفعّال ركنٌ لا يقل أهمية عن الأدوات التقنية في منظومة الأمن السيبراني الشاملة.

2.2.1 السياسات الدولية والاستراتيجيات الكبرى

أفرز التقاطع بين الذكاء الاصطناعي والأمن السيبراني حاجةً ماسةً إلى بنية تنظيمية دولية متكاملة، وقد تسابقت الحكومات والمنظمات الدولية لصياغة سياسات وإطارات ضابطة. وقد تجلّى ذلك في توجّه إدارة الرئيس الأمريكي نحو إبرام عقد بقيمة 200 مليون دولار مع شركة OpenAI لتطوير قدرات نموذجية في مجال الذكاء الاصطناعي تشمل مجالي القتال والعمليات الإدارية، بما في ذلك دعم الدفاع السيبراني الاستباقي (U.S. Department of Defense, 2025)، مما يُسلّط الضوء على أهمية الشراكات الاستراتيجية بين القطاع العام والخاص في هذا المجال. كما أطلقت شركة Microsoft في يونيو 2025 "برنامج الأمن الأوروبي" (European Security Program) الذي يقدّم خدمات استخبارات التهديدات المدعومة بالذكاء الاصطناعي مجاناً لجميع دول الاتحاد الأوروبي والدول المرشحة للانضمام وأعضاء الرابطة الأوروبية للتجارة الحرة إضافة إلى المملكة المتحدة وموناكو والفاتيكان، وذلك في أعقاب موجة هجمات سيبرانية استهدفت القارة العجوز (Smith, 2025).

وعلى صعيد آخر، يؤكد المنتدى الاقتصادي العالمي في تقريره "الذكاء الاصطناعي والأمن السيبراني: الموازنة بين المخاطر والمكاسب" الصادر في يناير 2025 أن التعاون الدولي الموسّع عبر الحدود الجغرافية يُعدّ ركيزةً أساسيةً لمكافحة الجريمة السيبرانية (World Economic Forum, 2025). وقد تجسّد هذا التوجه في الجهود المشتركة بين منظمة الإنتربول وأفريبول ضمن عملية "سيرينغيتي 2.0"

التي أسفرت في أغسطس 2025 عن تفكيك 25 مركزاً لاستخراج العملات المشفرة في أنغولا، مع إلقاء القبض على 1,209 أشخاص واسترداد ما يقارب 97.4 مليون دولار على امتداد 18 دولة أفريقية والمملكة المتحدة (INTERPOL, 2025).

وتتصدر المشهد التنظيمي الدولي جهود الاتحاد الأوروبي في سنّ قانون الذكاء الاصطناعي (EU AI Act)، الذي يُرسي منهجاً قائماً على إدارة المخاطر وتصنيف تطبيقات الذكاء الاصطناعي وفق درجة خطورتها (European Union, 2024). ويعكس هذا النهج التشريعي إدراكاً متزايداً بأن الذكاء الاصطناعي لا يمكن أن يُترك دون ضوابط تحكم توظيفه في المجالات الحساسة كالأمن السيبراني والبنية التحتية الحيوية.

وفيما يخص مؤشرات الجاهزية السيبرانية، يكشف مؤشر Cisco لجاهزية الأمن السيبراني لعام 2025، المبني على استطلاع مزدوج التعمية شمل 8,000 من قادة الأعمال المسؤولين عن الأمن السيبراني في 30 سوقاً عالمياً، أن 86% منهم أفادوا بتعرض مؤسساتهم لحادثة واحدة على الأقل مرتبطة بالذكاء الاصطناعي خلال الأشهر الاثني عشر الماضية (Cisco, 2025)، مما يُشير إلى هوة واسعة بين معدلات الجاهزية والواقع الميداني للتهديدات.

وتبرز على الصعيد الدولي ثلاث مقاربات استراتيجية متميزة في التعامل مع حوكمة الذكاء الاصطناعي الأمني: المقاربة الأوروبية المُعتمدة على التشريع المُسبق والتصنيف القائم على المخاطر عبر قانون "EU AI Act" (European Union, 2024)، والمقاربة الأمريكية المُعتمدة على توجيه رئاسي يُرسي معايير تطوعية مع حوافز للقطاع الخاص؛ والمقاربة الصينية القائمة على قواعد تنظيمية صارمة للذكاء الاصطناعي التوليدي صدرت عام 2023 ودخلت حيز التنفيذ الكامل عام 2025. ويُلاحظ الباحثون أن المقاربة الأوروبية هي الأشمل والأكثر إلزامية، غير أنها تواجه انتقادات من قطاع التكنولوجيا بسبب احتمال إبطائها لتيرة الابتكار.

وتجدر الإشارة إلى أن القمة الدولية لسلامة الذكاء الاصطناعي المنعقدة في سيول في مايو 2024 ثم في باريس في فبراير 2025 أرسيت مبادئ دولية مشتركة تتعلق بشفافية تطوير نماذج الذكاء الاصطناعي المتقدمة. وقد وقّعت على "إعلان سيول" 27 دولة إضافة إلى الاتحاد الأوروبي، بما فيها دول عربية رائدة كالمملكة العربية السعودية والإمارات العربية المتحدة، مما يُشير إلى انخراط متصاعد للمنطقة في منظومة الحوكمة الدولية للذكاء الاصطناعي الأمني (Seoul Declaration, 2024). وتُشكل هذه المبادئ الإطار الذي يُرجَّح أن يتطور خلال الفترة 2025-2027 إلى معاهدات دولية ملزمة.

2.2.2 الأطر التشريعية الإقليمية والمحلية

على الصعيد الإقليمي، تبرز منطقة الشرق الأوسط وأفريقيا بوصفها ساحة نشطة للتشريعات الأمنية الرقمية. فقد أعلن مجلس الأمن السيبراني الإماراتي، بالتعاون مع شركة QuantumGate، خلال فعالية "2025 CyberQ" في نوفمبر 2025، عن تعزيز جاهزيته الوطنية للانتقال نحو تقنيات التشفير المقاومة للحوسبة الكمية (Post-Quantum Cryptography)، ضمن "البرنامج الوطني للانتقال إلى ما بعد الكم"، في إطار توجّه إقليمي أوسع نحو اعتماد هذه التقنيات (UAE Cybersecurity Council, 2025).

وفي السياق المصري، دخل قانون حماية البيانات الشخصية رقم 151 لسنة 2020 مرحلة التفعيل الكامل بصدور لائحته التنفيذية رقم 816 لسنة 2025 بتاريخ 1 نوفمبر 2025، على أن يبدأ الإنفاذ الكامل بحلول 1 نوفمبر 2026، ويتضمن إلزام المؤسسات بالإبلاغ عن حوادث اختراق البيانات خلال 72 ساعة وتعيين مسؤول لحماية البيانات (Global Privacy Law Review, 2026). ويتسق هذا التوجه مع نمط عالمي من تعزيز الالتزام القانوني في مواجهة مخاطر البيئة الرقمية.

كما يعدّ كتاب "التحول الرقمي بين إمكانيات الذكاء الاصطناعي وتحديات الأمن السيبراني" للدكتور عبدالله موسى محمد السلحوت، الصادر عن دار اليازوري العلمية عام 2025، مرجعاً عربياً أصيلاً في تناول هذه الإشكاليات بعمق أكاديمي رصين، إذ يستعرض التوترات الجوهرية بين متطلبات التحول الرقمي ومتطلبات الأمن السيبراني في البيئات المؤسسية العربية (السلحوت، 2025).

وعلى صعيد التشريعات الخليجية، تواصل المملكة العربية السعودية تطوير مشهدها التنظيمي من خلال الهيئة السعودية للبيانات والذكاء الاصطناعي (SDAIA)، التي أصدرت مبادئ أخلاقية وأطراً إرشادية غير ملزمة لحوكمة الذكاء الاصطناعي، إلى جانب طرح مسودة "قانون مركز الذكاء الاصطناعي العالمي" (Global AI Hub Law) للاستشارة العامة عام 2025، دون أن يدخل بعد حيّز التنفيذ الملزم (Regulations.AI, 2026). وفي المقابل، أقرّ المجلس الوطني الأردني للأمن السيبراني في يونيو 2025 "الاستراتيجية الوطنية للأمن السيبراني 2025-2028"، التي تُدرج الذكاء الاصطناعي والحوسبة الكمية بوصفهما من المحركات الرئيسية لتطوير منظومة الأمن السيبراني الوطنية (وكالة الأنباء الأردنية "بترا"، 2025).

وعلى صعيد التحديات التشريعية المشتركة التي تواجه منطقة الشرق الأوسط وشمال أفريقيا، يكشف واقع الحال في الكويت، التي تراجعت إلى المرتبة الثانية عشرة عربياً والأخيرة خليجياً في مؤشر الأمن السيبراني العالمي لعام 2024، عن ثلاثة تحديات جوهرية تعوق نضج المنظومة التشريعية الإقليمية عموماً: هشاشة آليات تبادل المعلومات عن التهديدات عبر الحدود الوطنية رغم وجود اتفاقيات ثنائية، والتفاوت الكبير في مستويات النضج التشريعي بين الدول مما يخلق "بقعاً رمادية" تستغلها الجهات الإجرامية، وصعوبة مواكبة وتيرة التطور التقني في صياغة التشريعات (جريدة الجريدة الكويتية، 2025).

3. التهديدات السيبرانية المدعومة بالذكاء الاصطناعي وسبل الدفاع

يُعدّ هذا الفصل بمثابة القلب النابض للدراسة، إذ يتناول بالتحليل المعمق المشهدَ المزدوج لتوظيف الذكاء الاصطناعي في الفضاء السيبراني: من حيث كونه سلاحاً في يد المهاجمين، وترساً في يد المدافعين في آنٍ واحد. وقد أصبح هذا «المشهد ثنائي الوجه» السمة الأبرز التي تُميّز البيئة الأمنية الرقمية في النصف الثاني من العقد الثالث من هذا القرن.

3.1 التهديدات السيبرانية الناجمة عن الذكاء الاصطناعي

يُشكّل هذا القسم الوجهة الأولى والأكثر إثارةً للقلق في معادلة الذكاء الاصطناعي السيبرانية: الذكاء الاصطناعي سلاحاً هجومياً في يد المهاجمين. وقد كشفت السنوات الأخيرة عن تحوّل بنيوي جذري في طبيعة الجريمة السيبرانية، لا يتمثل فقط في ارتفاع وتيرة الهجمات وتعقيدها، بل في إعادة تعريف "هوية المهاجم" ذاتها؛ إذ لم تعد الهجمات المتطورة حِكراً على دول أو منظمات ممولة، بل باتت في متناول أفراد معزولين بلا خبرة تقنية سابقة. وقد رصدت دراسة Barrett وزملائه المنشورة عبر arXiv في ديسمبر 2025 هذا التحوّل بدقة، إذ طوّرت نماذج كمية لتقدير "الرفع" الذي يمنحه الذكاء الاصطناعي للمهاجمين عبر مراحل دورة الهجوم وفق إطار (Barrett et al., 2025).

والمعطى الأشد دلالةً على خطورة هذا التحوّل هو ما توصّل إليه تقرير IBM X-Force Threat Intelligence Index لعام 2026 من أن استغلال الثغرات أصبح المدخل الأكثر شيوعاً للهجمات، إذ مثّل 40% من الحوادث التي رصدتها الشركة خلال عام 2025، متجاوزاً بذلك التصيّد الاحتيالي وسرقة بيانات الاعتماد (IBM X-Force, 2026). وفي الاتجاه ذاته، أظهر تقرير CrowdStrike للتهديدات العالمية لعام 2026 أن الزمن الفاصل بين الاختراق الأولي والانتقال داخل الشبكة (Breakout Time) انخفض بنسبة 70% بين عامي 2021 و2025، ليصل إلى 29 دقيقة في المتوسط، ما يعكس تراجع الحاجز الزمني والتقني الذي كان يُقيّد قدرة المهاجمين الصغار على تنفيذ هجمات متطورة (CrowdStrike, 2026). ويتقاطع هذا مع ما كشفه تقرير Darktrace لعام 2025، الذي أفاد بأن 78% من مسؤولي أمن المعلومات (CISOs) حول العالم يرون أن التهديدات المعززة بالذكاء الاصطناعي باتت تُحدث تأثيراً ملموساً على مؤسساتهم، بارتفاع نسبته 5% عن عام 2024 (Darktrace, 2025).

3.1.1 الهجمات الإلكترونية المتقدمة بالذكاء الاصطناعي

تُمثّل الهجمات الإلكترونية المعززة بالذكاء الاصطناعي منعطفاً خطيراً في تاريخ الجريمة السيبرانية، إذ وسّعت دائرة المهاجمين المحتملين لتشمل أفراداً بلا خبرة تقنية سابقة. وقد رصدت The Hacker News في مايو 2026 حوادث مثيرة للقلق، من بينها: قيام ثلاثة مراقبين تتراوح أعمارهم بين 14 و16 عاماً في فبراير 2025 باستخدام نموذج ChatGPT لبناء أداة هجوم استهدفت أنظمة Rakuten Mobile نحو 220 ألف مرة، دون أن يمتلكوا أي خلفية برمجية. كما وثّق التقرير حادثة ديسمبر 2025 حيث استخدم شخص واحد أدوات ذكاء اصطناعي لاختراق أكثر من 10 وكالات حكومية مكسيكية

وسرقة ما يزيد على 195 مليون سجل ضريبي (The Hacker News, 2026).

وعلى صعيد المؤشرات الكمية، تكشف إحصاءات DeepStrike لعام 2025 أن 87% من المؤسسات على مستوى العالم تعرضت لهجوم سيبراني مدعوم بالذكاء الاصطناعي خلال العام الماضي، وأن 82.6% من رسائل التصيد الاحتيالي باتت تستخدم الذكاء الاصطناعي بشكل أو بآخر (DeepStrike, 2025)، فيما رصد تقرير مكتب التحقيقات الفيدرالي الأمريكي (FBI) لعام 2025 ارتفاعاً بنسبة 37% في عمليات الاحتيال التجاري عبر البريد الإلكتروني المعزّز بالذكاء الاصطناعي (Federal Bureau of Investigation, 2025).

ومن أبرز ما كشفته دراسة Mayoral-Vilches وزملائه حول إطار (Cybersecurity AI (CAI) لعام 2025 أن أنظمة الذكاء الاصطناعي الأمنية باتت قادرة، في مهام اختبار اختراق محددة، على العمل بسرعة تفوق سرعة الإنسان بما يصل إلى 3600 مرة، مع خفض متوسط التكاليف بمقدار 156 ضعفاً، وهي معطيات تنطبق على قدرات الاختبار الهجومي والدفاعي على حدٍ سواء (Mayoral-Vilches et al., 2025). كما تُشير بيانات شركة Mandiant في تقريرها M-Trends 2026 إلى أن وقت استغلال الثغرات الأمنية المعروفة انخفض من أكثر من 700 يوم عام 2020 إلى 44 يوماً فقط عام 2025، بل إن 28.3% من الثغرات باتت تُستغل خلال 24 ساعة فقط من الإفصاح عنها (Mandiant, 2026).

وتُسجّل تقارير قطاعية قلقاً بالغاً إزاء قطاع الرعاية الصحية والبنية التحتية الحيوية، وهو ما يُرسّخه التقرير السنوي للمنتدى الاقتصادي العالمي للأمن السيبراني 2025 الذي يُصنّف الفدية الرقمية (Ransomware) بوصفها الخطر السيبراني الأكبر الذي تقرّ به الشركات على نطاق واسع (World Economic Forum, 2025).

وعلى صعيد القطاعات الأكثر استهدافاً، يكشف تقرير Mandiant M-Trends 2025 أن القطاع المالي احتل المرتبة الأولى في قائمة القطاعات المستهدفة بنسبة 17.4% من الحوادث، يليه قطاع الأعمال والخدمات المهنية بنسبة 11.1%، ثم قطاع التقنية العالية بنسبة 10.6%، فالقطاع الحكومي بنسبة 9.5%، وأخيراً قطاع الرعاية الصحية بنسبة 9.3% (Mandiant, 2025).

وتُعَدّ ظاهرة "الجريمة السيبرانية كخدمة" (Cybercrime-as-a-Service / CaaS) المُعززة بالذكاء الاصطناعي أحد أبرز التحولات البنيوية في مشهد التهديدات خلال 2025-2026؛ إذ أفرزت منصات من قبيل GhostGPT و DarkGPT و FraudGPT — وهي نماذج لغوية مُصممة خصيصاً للاستخدام الإجرامي وتُباع عبر قنوات Telegram ومنتديات الدارك ويب — منظومة متكاملة يُقدّم فيها المتخصصون التقنيون من الجريمة المنظمة خدماتهم للمهاجمين الأقل خبرة (Check Point, 2025). ويُصدّ تقرير CrowdStrike للتهديدات العالمية لعام 2025 أن أسرع وقت اختراق مُسجّل بلغ 51 ثانية فقط، في حين بلغ متوسط وقت الاختراق 48 دقيقة، مما يعكس مدى تسارع أدوات الهجوم المدعومة بالذكاء الاصطناعي (CrowdStrike, 2025). وتُمثّل هذه الديمقراطية في أدوات الجريمة السيبرانية تحدياً وجودياً لمنظومات الحماية التقليدية.

3.1.2 التزييف العميق والهندسة الاجتماعية الرقمية

يُمثل التزييف العميق (Deepfake) والهندسة الاجتماعية المعززة بالذكاء الاصطناعي أحدًا أشد أوجه التهديد السيبراني فتكاً وتأثيراً في المشهد الراهن. وتكشف بيانات شركة Pindrop المتخصصة في أمن الصوت أن حوادث التزييف الصوتي (Voice Deepfakes) ارتفعت بنسبة 680% على أساس سنوي خلال عام 2024، بينما قفز معدل تكرار محاولات الاحتيال بالتزييف الصوتي بنسبة تجاوزت 1300%، من محاولة واحدة شهرياً في المتوسط إلى سبع محاولات يومياً. وقد سُجل ارتفاع خاص بقطاع التأمين بنسبة 475%، مقابل 149% في القطاع المصرفي، فيما تتوقع الشركة أن يستمر هذا الاتجاه بالتصاعد بنسبة 162% إضافية خلال عام 2025 (Pindrop, 2025).

وعلى صعيد التوثيق الكمي للحوادث، تُظهر بيانات شركة Resemble AI تسارعاً غير خطي لافتاً: فقد سُجّلت 179 حادثة موثقة خلال الربع الأول من عام 2025 وحده، متجاوزةً بذلك إجمالي حوادث عام 2024 كاملاً بنسبة 19%، ثم قفز الرقم إلى 487 حادثة في الربع الثاني (بارتفاع 312% على أساس سنوي)، ليصل إلى 2031 حادثة موثقة في الربع الثالث من العام ذاته (Resemble AI, 2025). ويتقاطع هذا مع معطيات شركة CrowdStrike التي رصدت ارتفاعاً في هجمات التصيد الصوتي (Vishing) بنسبة 442% بين النصف الأول والثاني من عام 2024، ثم بنسبة إضافية بلغت 1633% خلال الربع الأول من 2025 مقارنةً بالربع الأخير من العام السابق (CrowdStrike, 2025)، فيما تُشير بيانات منصة Sumsb إلى أن التزييف العميق بات يمثل 6.5% من إجمالي محاولات الاحتيال المُسجّلة عالمياً، ارتفاعاً من 0.1% فقط عام 2022 (Sumsb, 2025).

والمعطى الأكثر دلالةً على خطورة هذا التحول هو ما توصلت إليه دراسة شركة iProov لعام 2025، التي اختبرت 2000 مشارك من بريطانيا والولايات المتحدة، من أن 0.1% فقط من البشر يستطيعون تحديد محتوى التزييف العميق بدقة مستمرة، حتى حين يكونون على علم مسبق بضرورة الانتباه له، وهو ما يعكس فجوة هائلة بين تطور تقنيات الخداع وقدرة الإنسان على مواجهتها (iProov, 2025).

ويتجاوز هذا التهديد نطاق الأفراد ليطال المؤسسات والحكومات على حدٍ سواء. فقد حذرت شركة Anthropic، في تقريرها للاستخبارات التهديدية الصادر في أغسطس 2025، من أن نموذجها Claude تعرّض لمحاولة استغلال من جانب أحد القراصنة الذي وظّفه لأتمتة عمليات الاستطلاع وسرقة بيانات الاعتماد واختراق الشبكات، إلى جانب كتابة شيفرات خبيثة مخصصة، واستُهدفت بذلك 17 مؤسسة على الأقل عبر قطاعات الحكومة والرعاية الصحية وخدمات الطوارئ والمؤسسات الدينية، بمطالب فدية تراوحت بين 75 ألفاً وأكثر من 500 ألف دولار (Anthropic, 2025).

وتُجسّد حادثة شركة Arup الهندسية البريطانية حجم الخطر الفعلي لهذه التقنيات على أرض الواقع: ففي يناير 2024، تمكّن مهاجمون من إقناع أحد موظفي مكتب الشركة في هونغ كونغ بتنفيذ 15 تحويلاً مالياً احتيالياً بقيمة إجمالية بلغت 25.6 مليون دولار (200 مليون دولار هونغ كونغي)، وذلك عبر مؤتمر مرئي كامل مُزيّف ضمّ "مديراً مالياً" ومجموعة من الموظفين لم يكن أيّ منهم حقيقياً، بل جميعهم صُنّعوا بتقنية التزييف العميق اعتماداً على مقاطع فيديو وصوت حقيقية متاحة علناً لتنفيذي

الشركة (CNN Business, 2024).

3.2 الذكاء الاصطناعي أداة للدفاع السيبراني

في مقابل المشهد التهديدي المتصاعد الذي رسمه القسم السابق، يُقدّم هذا القسم الوجه الآخر للمعادلة: الذكاء الاصطناعي بوصفه درعاً استراتيجياً وقوةً دفاعية غير مسبقة في تاريخ الأمن السيبراني. وإذا كان الذكاء الاصطناعي قد وسّع دائرة المهاجمين وضاعف قدراتهم، فإنه في الآن ذاته يمنح المدافعين أدوات تحليلية وردّ فعل تتجاوز قدرات الإنسان منفرداً في السرعة والدقة والشمولية. وتنبثق أهمية هذا التحوّل من حقيقة موثّقة أكتبتها دراسة Mayoral-Vilches وآخرين، والتي أطلقت أصلاً في أبريل 2025 ضمن إطار (CAI) "Cybersecurity AI" وأعاد الباحث ذاته التأكيد عليها في ورقة لاحقة نُشرت في يناير 2026: أن أنظمة الذكاء الاصطناعي الأمنية باتت تعمل، في مهام اختبار اختراق محددة، بسرعة تفوق الإنسان بـ 3600 مرة، مع خفض التكاليف التشغيلية بمقدار 156 ضعفاً، وهو ما يُعيد رسم خريطة الجدوى الاقتصادية للدفاع السيبراني كلياً (Mayoral-Vilches et al., 2025, 2026).

وبتمحور توظيف الذكاء الاصطناعي في الدفاع السيبراني حول أربعة محاور وظيفية رئيسية تؤثّقها الأدبيات الأكاديمية الحديثة: أولها الاستباق والرصد المبكر (Predictive Intelligence) الذي يرصد التهديدات قبل تحوّلها إلى هجمات فعلية عبر تحليل الأنماط التاريخية وتوقع نقاط الضعف المحتمل استغلالها؛ وثانيها الكشف في الوقت الحقيقي (Real-Time Detection) الذي يحلّل كميات هائلة من الأحداث الأمنية آنياً للتعرف على الشذوذات وتصنيف التهديدات؛ وثالثها الاستجابة الآلية (Auto-mated Response) التي تُطلق إجراءات الاحتواء والعزل بسرعة تفوق قدرة الإنسان دون انتظار قرار بشري مباشر؛ ورابعها التعلم المستمر (Continuous Learning) الذي يُحسّن أداء الأنظمة الدفاعية تدريجياً مع كل حادثة جديدة عبر آليات التغذية الراجعة من المحللين البشريين. وتُشكّل هذه المحاور الأربعة معاً إطاراً متكاملًا للدفاع السيبراني الاستباقي والتكيفي الذي تتبنّاه المؤسسات الرائدة على المستوى العالمي، بما يجمع بين قدرة الآلة على السرعة والحجم وقدرة الإنسان على التقدير الاستراتيجي (Alsmadi & Alazzam, 2025).

3.2.1 آليات الكشف والاستجابة الآلية

في مواجهة الموجة المتصاعدة من التهديدات المعززة بالذكاء الاصطناعي، تبرز آليات الدفاع الذكية بوصفها الحصن المنيع الذي يُعوّل عليه. وقد كشف تقرير IBM للتكاليف الكلية لاختراق البيانات لعام 2025 أن المؤسسات التي تُوظّف الذكاء الاصطناعي والأتمتة بشكل مكثّف في عملياتها الأمنية توفّر في المتوسط 1.9 مليون دولار في التعامل مع الحوادث الأمنية مقارنةً بالمؤسسات التي لا تستخدم هذه التقنيات، علاوةً على تقليص دورة حياة الاختراق (من الاكتشاف إلى الاحتواء) بمقدار 80 يوماً في المتوسط (IBM, 2025).

وتتصدر هذه الآليات الدفاعية تقنيات كشف الشذوذ (Anomaly Detection) القائمة على خوارزميات التعلم الآلي، التي تحلّل حجوم هائلة من البيانات لرصد السلوكيات غير المعتادة وتمييزها عن النشاط الطبيعي. كما تتوسع مراكز العمليات الأمنية (SOC) في تبني مقاربات الأتمتة الذكية، بما يشمل

منصات SIEM (إدارة المعلومات والأحداث الأمنية) المعززة بالذكاء الاصطناعي القادرة على ربط الأحداث عبر أنظمة متعددة وتحديد سلاسل الهجوم المعقدة.

وعلى صعيد الاستثمار، يشير تقرير Total Assure إلى أن سوق الأمن السيبراني المعزز بالذكاء الاصطناعي في طريقه للنمو من 31 مليار دولار عام 2024 إلى 134.6 مليار دولار بحلول عام 2030، بمعدل نمو سنوي مركب قدره 26.6% (Total Assure, 2026). وتُعزز هذه الأرقام الحجج الاقتصادية للاستثمار الاستراتيجي في الأمن السيبراني القائم على الذكاء الاصطناعي.

وتكتسب تقنية XDR (الكشف والاستجابة الموسعة) المعززة بالذكاء الاصطناعي زخماً متصاعداً بوصفها التطور الطبيعي لمنصات SIEM/SOAR التقليدية. فبينما تعمل منصات SIEM على ربط الأحداث الأمنية من مصادر متعددة، يُضيف الذكاء الاصطناعي في منظومة XDR طبقةً من الفهم السياقي العميق تُمكنه من تمييز الهجمات الموزعة عبر نقاط نهاية متفرقة وحركة شبكة وسحابة وبريد إلكتروني في آن واحد. وتُفيد Microsoft بأن نظامها Defender XDR المُعزَّز بالذكاء الاصطناعي يُعالج أكثر من 65 تريليون إشارة أمنية يومياً، ويحقق مستوى ثقة يتجاوز 99% باحتواء الأنشطة المشبوهة تلقائياً استناداً إلى بيانات الإنتاج الفعلية (Microsoft, 2026).

ويُمثل مفهوم "الصيد الاستباقي للتهديدات" (Threat Hunting) المُعزَّز بالذكاء الاصطناعي تحولاً نوعياً في فلسفة الأمن السيبراني من الدفاع السلبي إلى الاستباق الفاعل. ففي هذا النموذج، تُطلق أنظمة الذكاء الاصطناعي بشكل استباقي عمليات بحث عن مؤشرات الاختراق الخفية قبل أن تُطلق أي إنذار مرئي. وتؤكد بيانات شركة CrowdStrike في تقريرها لصيد التهديدات لعام 2025 أن متوسط زمن الاختراق (Breakout Time)، أي الفاصل الزمني بين الاختراق الأولي والانتقال الجانبي داخل الشبكة — انخفض من 79 دقيقة عام 2024 إلى 62 دقيقة عام 2025، ويُعزى هذا الانخفاض بدرجة كبيرة إلى تقنيات الذكاء الاصطناعي للصيد الاستباقي (CrowdStrike, 2025). وعلى صعيد التطبيقات القطاعية، يبرز نموذج "اليقظة السيبرانية المستمرة" (Continuous Cyber Vigilance) في القطاع المالي بوصفه الأكثر نضجاً في توظيف الذكاء الاصطناعي الدفاعي؛ ففي السياق العربي، أعلنت شركة "مزن" السعودية في أكتوبر 2025 عن شراكة استراتيجية مع مصرف الراجحي لتعزيز منظومته لمكافحة الاحتيال المالي والامتثال، عبر تبني منصة "فوكال" المدعومة بالذكاء الاصطناعي، وهو ما يُجسّد التوجه المتنامي للبنوك الخليجية نحو الاستثمار بحلول الذكاء الاصطناعي الدفاعية (Zawya, 2025).

3.2.2 مستقبل الأمن السيبراني في عصر الذكاء الاصطناعي

تتشكّل ملامح مستقبل الأمن السيبراني وفق منظومة من الاتجاهات الاستراتيجية المتشابكة. وتبرز في مقدمة هذه الاتجاهات أهمية الاستعداد للحوسبة الكمومية (Quantum Readiness) وتبني معايير التشفير ما بعد الكمي، إذ أقرّ المعهد الوطني الأمريكي للمعايير والتكنولوجيا (NIST) في 13 أغسطس 2024 أولى معاييرهِ الرسمية الثلاثة لهذا الغرض (FIPS 203, 204, 205)، بعد مسار تقييم استمر ثماني سنوات (NIST, 2024). فمع تصاعد قدرات الحوسبة الكمومية، باتت أنظمة التشفير الحالية عُرضةً للاختراق مستقبلاً، مما يستوجب اتخاذ تدابير استباقية على المستويين المؤسسي والوطني.

وتكتسب نماذج الأمن المعتمدة على مبدأ "عدم الثقة المطلقة" (Zero Trust Architecture) زخماً متصاعداً في ظل التهديدات الراهنة. فوفق بيانات تقرير 2026 Mandiant M-Trends، فإن معظم الاختراقات المكتشفة عام 2025 اعتمدت على استغلال الثغرات (32% من الحوادث) وبيانات الدخول المسروقة، لا على البرمجيات الخبيثة التقليدية، مما يجعل الهوية الرقمية تمثل خط الدفاع الأول (Mandiant, 2026).

وتتمركز الحاجة الاستراتيجية الراهنة حول ثلاثة محاور: أولها بناء قدرات الكشف والاستجابة التلقائية في المنظومة الدفاعية، وثانيها صياغة بنى تنظيمية تُحكم الرقابة على استخدامات الذكاء الاصطناعي وتُقيّد توظيفه الضار، وثالثها تطوير الكفاءات البشرية المتخصصة. وتُعَدّ الفجوة في الكفاءات البشرية من أشد الإشكاليات إلحاحاً، إذ تُشير بيانات تقرير Darktrace لعام 2025 إلى أن نقص الكوادر يُصنّف الحاجز الأكبر أمام التصدي للتهديدات المعززة بالذكاء الاصطناعي، في حين لا يُخطط سوى 11% من المؤسسات لزيادة أعداد موظفيها الأمنيين خلال العام (Darktrace, 2025).

ويُلوح في الأفق نمطٌ جديد من التكامل بين الذكاء الاصطناعي وبحوث الأمن السيبراني، يُعرف بـ"الذكاء الاصطناعي الأمني الفائق" (Cybersecurity Superintelligence)، وهو مفهوم طرحه الباحث Mayoral-Vilches وزملاؤه في ورقّتهم المنشورة في يناير 2026، ويُعرّف بأنه ذكاء اصطناعي يتجاوز أفضل القدرات البشرية سرعةً واستراتيجيةً في المجال الأمني (Mayoral-Vilches et al., 2026). ويطرح هذا النمط جدلاً أخلاقياً وتقنياً عميقاً حول الحدود المقبولة للأتمتة في صنع القرار الأمني، وحجم التدخل البشري الضروري لضمان المساءلة والمسؤولية.

وتبقى الفجوة في الكفاءات البشرية أحد أعمق التحديات الهيكلية في هذا المشهد، وهو ما يُوثّقه تقرير Cisco للجاهزية السيبرانية 2025، الذي استند إلى استطلاع مزدوج التعمية شمل 8000 قائد أمني عبر 30 دولة، وخلص إلى أن 4% فقط من المؤسسات حول العالم بلغت مستوى النضج الكامل في جاهزيتها السيبرانية (Cisco, 2025). ويرى الباحثون أن المعادلة الحقيقية تكمن في بناء "الإنسان المُعزّز" (Augmented Human) في مجال الأمن السيبراني: المحلل الأمني الذي يفهم الذكاء الاصطناعي ويوجّهه ويُراجع قراراته، لا يُنافس ولا يُنِيب إليه كلياً.

وفي ختام هذا الفصل، تتضح أهمية التعامل مع الذكاء الاصطناعي في الأمن السيبراني بوصفه معطًى استراتيجياً متجدداً يستدعي مقاربةً شاملةً تجمع بين التقنية والتشريع والتوعية وبناء القدرات. فالمؤسسات والحكومات التي ستجح في بناء منظومات أمنية ذكية ومُتكّنة هي تلك التي تستثمر في الوقت ذاته في التقنية والإنسان والحوكمة، وتُدرك أن الأمن السيبراني في عصر الذكاء الاصطناعي ليس وجهةً نهائيةً بل مسيرة مستمرة لا تتوقف.

4. الخاتمة

تتوقف هذه الدراسة عند مفترق حضاري بالغ الدلالة، تتلاقى فيه ثورتان متزامنتان لا فكاك من تداعياتهما: ثورة الذكاء الاصطناعي التي أعادت رسم حدود الممكن في كل ميادين الحياة، وثورة التهديدات السيبرانية التي لم تكن يوماً أكثر تعقيداً وتسارعاً مما هي عليه اليوم. وقد سعت الدراسة، من خلال فصلها المتكاملين وتوثيقها المستند إلى أحدث المصادر المتحقق منها للفترة 2025-2026، إلى تقديم صورة أمينة وموضوعية لهذا المشهد المركّب.

يفصح المشهد الذي ترسمه بيانات 2025-2026 عن حقيقة لا مراء فيها: الذكاء الاصطناعي لن يعود عاملاً مساعداً في الأمن السيبراني بل سيصبح العمود الفقري لكل منظومة دفاعية جديرة بهذا الاسم. والسؤال لم يعد «هل نتبنّى الذكاء الاصطناعي في أمننا السيبراني؟» بل «كيف نتبنّاه بالسرعة والحكمة اللازمين؟». وفي الإجابة على هذا السؤال يكمن الفارق بين مؤسسة محصنة ومؤسسة مكشوفة، وبين دولة رقمياً سيدة ودولة رقمياً هشة.

إن الأمن السيبراني في عصر الذكاء الاصطناعي ليس تقنية تُشتري ولا نظاماً يُركّب مرة واحدة، بل هو ثقافة مستدامة وعملية متجددة تستوجب رؤية استراتيجية بعيدة المدى، واستثماراً متواصلًا في الإنسان قبل الآلة، وتعاوناً دولياً لا تحده حدود. وبهذا الفهم، تُختتم هذه الدراسة ونأمل أن تكون إسهاماً نافعاً في إثراء النقاش العلمي العربي حول إشكالية من أبرز إشكاليات عصرنا.

خلصت الدراسة إلى جملة من النتائج الجوهرية يمكن إجمالها فيما يأتي:

أولاً: أثبت الذكاء الاصطناعي أنه سلاح ذو حدين في الفضاء السيبراني؛ إذ يمنح المهاجمين قدرات هجومية غير مسبوقة، كما يتجلى في ارتفاع هجمات التصيد الاحتيالي بنسبة 1265% وتضاعف حوادث التزيف العميق بنسبة 680% خلال 2025، في حين يُتيح للمدافعين في الوقت ذاته توفير ما يزيد على 1.9 مليون دولار لكل حادثة اختراق وتقليص دورة حياة التهديد بمقدار 80 يوماً وفق تقرير «IBM 2025».

ثانياً: أن فجوة الجاهزية السيبرانية حقيقة موثقة وخطيرة؛ فرغم أن 86% من المؤسسات تعرّضت لحوادث مرتبطة بالذكاء الاصطناعي خلال عام 2025 وفق Cisco، فإن 4% فحسب من المؤسسات على مستوى العالم بلغت مستوى «النضج» في الجاهزية الأمنية. ولا يعكس هذا الرقم إخفاقاً تقنياً وحسب، بل أزمة حوكمة وبناء قدرات بشرية.

ثالثاً: أن الأطر التنظيمية الدولية والإقليمية تتطور لكنها لم تلحق بعد بوتيرة التهديدات المتصاعدة؛ فقانون الذكاء الاصطناعي الأوروبي ومبادرات الخليج وجهود مكافحة الجرائم الإلكترونية تمثل خطوات في الاتجاه الصحيح، غير أنها تحتاج إلى تعاون دولي أعمق وتنسيق أسرع.

رابعاً: أن الذكاء الاصطناعي الأمني الفائق (Cybersecurity Superintelligence) ليس ضرباً من الخيال العلمي، بل واقع يتشكل؛ إذ تُظهر أنظمة مثل ATOM من IBM وCAI من arXiv قدرات

تشغيلية تتجاوز الإنسان سرعةً وإنتاجاً، وهو ما يستوجب إعادة التفكير في نماذج الحوكمة والمساءلة. استناداً إلى النتائج السابقة، توصي الدراسة بما يأتي:

1. التحوّل العاجل نحو نموذج «عدم الثقة المطلقة» (Zero Trust Architecture) بوصفه المنهجية الدفاعية الأنسب في بيئة تتصاعد فيها هجمات بيانات الاعتماد المسروقة والهويات غير البشرية.
2. الاستثمار المتسارع في الأمن السيبراني المعزز بالذكاء الاصطناعي، مع إعطاء الأولوية لأنظمة الكشف والاستجابة الآلية (SIEM/SOAR/XDR) القادرة على مجاراة سرعة الهجمات الذكية.
3. التهيؤ الاستراتيجي المبكر للحوسبة الكمومية عبر تبني معايير التشفير ما بعد الكمي (Post-Quantum Cryptography) قبل أن تتحول الثغرة القادمة إلى خطر وشيك.
4. سدّ فجوة الكفاءات البشرية من خلال برامج تأهيل متخصصة في الذكاء الاصطناعي الأمني، والاستعاضة عن نموذج التوظيف التقليدي بنموذج هجين يدمج الخبرة البشرية مع الأتمتة الذكية.
5. بناء أطر حوكمة واضحة للذكاء الاصطناعي على المستويين المؤسسي والوطني، تُحدد صلاحيات القرار الآلي وتُرسّخ مبدأ المساءلة الإنسانية في العمليات الأمنية الحساسة.
6. تعزيز التعاون الدولي في تبادل معلومات التهديدات والاستجابة المشتركة للحوادث الكبرى، على غرار نماذج التعاون بين INTERPOL وAFRIPOL، وتوسيعها لتشمل منطقة الشرق الأوسط وشمال أفريقيا.

نموذج البيان في حال عدم وجود تضارب مصالح

يُصرِّح المؤلف بعدم وجود أي تضارب مصالح يتعلق بنشر هذا البحث.

نموذج البيان في حال عدم وجود تمويل

يُصرِّح المؤلف أنَّ هذه الدراسة لم تحصل على أي دعم مالي أو منحة من أي جهة عامة أو خاصة أو غير ربحية.

لائحة المراجع

المراجع العربية:

1. جريدة الجريدة الكويتية. (2025, 29 نوفمبر). الكويت بالمرتبة ال 12 عربياً والأخيرة خليجياً في الأمن السيبراني.
2. السلحوت، عبدالله موسى محمد. (2025). التحول الرقمي بين إمكانيات الذكاء الاصطناعي وتحديات الأمن السيبراني. دار اليازوري العلمية.
3. وكالة الأنباء الأردنية (بترا). (2025, 19 يونيو). إقرار الاستراتيجية الوطنية للأمن السيبراني -2025.
4. Zawya. (2025, October 20).
5. تعاون استراتيجي بين مُزن ومصرف الراجحي لتعزيز الوقاية من الاحتيال المالي بالذكاء الاصطناعي <https://www.zawya.com/ar/البيانات-الصحفية/بيانات-الشركات/تعاون-استراتيجي-بين-مُزن-ومصرف-الراجحي-لتعزيز-الوقاية-من-الاحتيال-المالي-بالذكاء-الاصطناعي-djoxnsrn> تاريخ الدخول: 2026/7/27

ثانياً: المراجع الأجنبية:

1. The Peninsula Qatar. (2025, February 27). Japanese teens arrested for using AI to create illegal phone contracts. <https://thepeninsulaqatar.com/article/27/02/2025/japanese-teens-arrested-for-using-ai-to-create-illegal-phone-contracts>. consulté le: 22/7/2026
2. Cisco. (2025). 2025 Cisco Cybersecurity Readiness Index. https://newsroom.cisco.com/c/dam/r/newsroom/en/us/interactive/cybersecurity-readiness-index/2025/documents/2025_Cisco_Cybersecurity_Readiness_Index.pdf consulté le : 20/7/2026
3. Total Assure. (2026, January 13). AI cybersecurity statistics in 2025: Comprehensive data on threats, detection, and defense. <https://www.totalassure.com/blog/ai-cybersecurity-stats> consulté le: 20/7/2026
4. IBM. (2025, April 28). IBM delivers autonomous security operations with cutting-edge agentic AI [Press release]. <https://newsroom.ibm.com/2025-04-28-ibm-delivers-autonomous-security-operations-with-cutting-edge-agentic-ai> consulté le: 15/7/2026
5. IBM. (2026, February 25). IBM 2026 X-Force threat intelligence index [Press release]. <https://newsroom.ibm.com/2026-02-25-ibm-2026-x-force-threat-index-ai-driven-attacks-are-escalating-as-basic-security-gaps-leave-enterprises-exposed> consulté le: 15/7/2026
6. European Union. (2024). Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act).

7. INTERPOL. (2025, August 22). African authorities dismantle massive cybercrime and fraud networks, recover millions. <https://www.interpol.int/en/News-and-Events/News/2025/African-authorities-dismantle-massive-cybercrime-and-fraud-networks-recover-millions> consulté le: 15/7/2026
8. Seoul Declaration by countries attending the AI Seoul Summit, 21-22 May 2024. (2024). <https://www.industry.gov.au/publications/seoul-declaration-countries-attending-ai-seoul-summit-21-22-may-2024> consulté le: 15/7/2026
9. Smith, B. (2025, June 4). Microsoft's European Security Program. Microsoft. <https://blogs.microsoft.com/on-the-issues/2025/06/04/microsofts-european-security-program> consulté le: 20/7/2026 /
10. U.S. Department of Defense. (2025, June 16). OpenAI awarded \$200 million contract for prototype frontier AI capabilities [Press release, as reported by CNBC]. <https://www.cnbc.com/2025/06/16/openai-wins-200-million-us-defense-contract.html> consulté le: 20/7/2026
11. World Economic Forum. (2025, January). Artificial intelligence and cybersecurity: Balancing risks and rewards [White paper]. https://reports.weforum.org/docs/WEF_Artificial_Intelligence_and_Cybersecurity_Balancing_Risks_and_Rewards_2025.pdf consulté le: 20/7/2026
12. Global Privacy Law Review. (2026). Egypt's Personal Data Protection Law: Implementation challenges and additional compliance requirements. Global Privacy Law Review, 7(1). Kluwer Law Online.
13. Regulations.AI. (2026). Saudi Arabia AI regulation overview. <https://regulations.ai/regulations/RAI-SA-NA-SUMMARY-2026> consulté le: 13/7/2026
14. UAE Cybersecurity Council. (2025, November 27). UAE Cybersecurity Council advances national readiness for future quantum threats with ATRC's QuantumGate [Press release]. Advanced Technology Research Council.
15. Barrett, S., Murray, M., Quarks, O., Smith, M., Kryś, J., Campos, S., Boria, A. T., Touzet, C., Hayrapet, S., Heiding, F., Nevo, O., Swanda, A., Aguirre, J., Gershovich, A. B., Clay, E., Fetterman, R., Fritz, M., Juarez, M., Mavroudis, V., & Papadatos, H. (2025). Toward quantitative modeling of cybersecurity risks due to AI misuse. arXiv. <https://arxiv.org/abs/2512.08864> consulté le: 13/7/2026
16. Darktrace. (2025). 2025 state of AI cybersecurity report <https://www.darktrace.com/news/new-report-finds-that-78-of-chief-information-security-officers-globally-are-seeing-a-significant-impact-from-ai-powered-cyber-threats> consulté le: 13/7/2026
17. IBM X-Force. (2026). 2026 X-Force Threat Intelligence Index: Making the case for secur-

ing identities, AI-enhanced detection and proactive risk management. IBM. <https://www.ibm.com/think/x-force/threat-intelligence-index-2026-securing-identities-ai-detection-risk-management> consulté le: 13/7/2026

18. Check Point Research. (2025). AI security report 2025. Check Point Software Technologies.

19. DeepStrike. (2025). AI cyber attack statistics 2025, trends, costs, and global impact. <https://deepstrike.io/blog/ai-cyber-attack-statistics-2025> consulté le: 29/6/2026

20. Federal Bureau of Investigation, Internet Crime Complaint Center. (2025). 2025 IC3 report.

21. Mandiant. (2025). M-Trends 2025: Data, insights, and recommendations from the frontlines. Google Cloud. <https://cloud.google.com/blog/topics/threat-intelligence/m-trends-2025> consulté le: 29/6/2026

22. The Hacker News. (2026, May 4). 2026: The year of AI-assisted attacks. <https://thehackernews.com/2026/05/2026-year-of-ai-assisted-attacks.html> consulté le: 29/6/2026

23. World Economic Forum. (2025). Global cybersecurity outlook 2025. <https://www.weforum.org/publications/global-cybersecurity-outlook-2025/digst> consulté le: 13/7/2026

24. Anthropic. (2025, August 27). Detecting and countering misuse of AI: August 2025. <https://www.anthropic.com/news/detecting-countering-misuse-aug-2025> consulté le: 13/7/2026

25. CrowdStrike. (2025). 2025 global threat report. <https://go.crowdstrike.com/rs/281-OBQ-266/images/CrowdStrikeGlobalThreatReport2025.pdf> consulté le: 20/7/2026

26. iProov. (2025, February 12). iProov study reveals deepfake blindspot: Only 0.1% of people can accurately detect AI-generated deepfakes [Press release].

27. Magramo, K. (2024, May 16). British engineering giant Arup revealed as \$25 million deepfake scam victim. CNN Business. <https://www.cnn.com/2024/05/16/tech/arup-deepfake-scam-loss-hong-kong-intl-hnk> consulté le: 20/7/2026

28. Pindrop. (2025). 2025 voice intelligence and security report.

29. Resemble AI. (2025). Deepfake incident reports, Q1–Q3 2025.

30. Sumsub. (2025). Identity fraud report 2025.

31. Alsmadi, I., & Alazzam, I. (2025). Artificial intelligence as the next frontier in cyber defense: Opportunities and risks. *Electronics*, 14(24), 4853. <https://www.mdpi.com/2079-9292/14/24/4853> consulté le: 20/7/2026

32. Mayoral-Vilches, V., Navarrete-Lozano, L. J., Sanz-Gómez, M., Salas Espejo, L., Crespo-Álvarez, M., Oca-Gonzalez, F., Balassone, F., Glera-Picón, A., Ayucar-Carbajo, U., & Ruiz-Alcalde, J. A. (2025). CAI: An open, bug bounty-ready cybersecurity AI. *arXiv*. <https://arxiv.org/abs/2507.12345>

- arxiv.org/abs/2504.06017 consulté le: 7/6/2026
- IBM. (2025). Cost of a data breach report 2025. <https://www.ibm.com/reports/data-breach> consulté le: 5/6/2026
- Microsoft. (2026). Automatic attack disruption in Microsoft Defender XDR. Microsoft Learn. <https://learn.microsoft.com/en-us/defender-xdr/automatic-attack-disruption> consulté le: 5/6/2026
- Darktrace. (2025). 2025 state of AI cybersecurity report. <https://www.darktrace.com/the-state-of-ai-cybersecurity-2025> consulté le: 3/7/2026
- Mandiant. (2026). M-Trends 2026. Google Cloud. <https://cloud.google.com/blog/topics/threat-intelligence/m-trends-2026> consulté le: 7/7/2026
- Mayoral-Vilches, V., et al. (2026). Towards cybersecurity superintelligence: From AI-guided humans to human-guided AI. arXiv. <https://arxiv.org/abs/2601.14614> consulté le: 7/7/2026
- National Institute of Standards and Technology. (2024, August 13). NIST releases first 3 finalized post-quantum encryption standards. <https://www.nist.gov/news-events/news/2024/08/nist-releases-first-3-finalized-post-quantum-encryption-standards> consulté le: 6/7/2026