



P-ISSN : 2074-9554 | E-ISSN: 2663-811

Journal of Al-Farahidi's Arts

available online at: fia.tu.edu.iq/index.php/fia



Sabreen ali Hussein
Aseel hamoud Hamza
Suhad Al-Shoukry

E-Mail: a_Qur.sabreen.alshammary@uobabylon.edu.iq

The Role of Artificial Intelligence in Dealing with Advanced Cyber Attacks: A Comparative Study of Algorithms

Keywords:

Artificial Intelligence, Cyber Attacks, Intrusion Detection Systems, Machine Learning and Deep Learning, Artificial Neural Networks.

Article history:

Received 12/8/2025
Received in revised form 2/9/2025
Accepted 19/10/2025
Available online 9/12/2025

E-mail Jaa@tu.edu.iq

ABSTRACT

The digital world is witnessing an unprecedented escalation in the extent and development of cyber attacks, which requires more intelligent security solutions that adapt to the constantly changing technological environment. This research aims to explore the effective role of artificial intelligence technologies, especially machine learning algorithms, in dealing with advanced cyber attacks, such as intelligent phishing attacks, advanced malicious software, and zero-day attacks. The digital world is witnessing an unprecedented escalation in the extent and development of cyber attacks, which requires more intelligent security solutions that adapt to the constantly changing technological environment. This research aims to explore the effective role of artificial intelligence technologies, especially machine learning algorithms, in dealing with advanced cyber attacks, such as attacks on the Day of Zero, intelligent phishing attacks, advanced malware.

The research is based on a comparison method between a variety of smart algorithms used in intrusion detection systems (IDS), such as decision trees, artificial neural networks (ANN), support vectors (SVM), and deep learning algorithms. Estimating the performance of these algorithms according to the criteria of accuracy, detection rate, false alarm rate,

©THIS AN OPEN ACCESS ARTICLE UNDER THE CCBY LICENSE

<http://creativecommons.org/licenses/by/4.0>



دور الذكاء الاصطناعي في التصدي للهجمات السيبرانية المتقدمة: دراسة مقارنة بين الخوارزميات

صابرين علي حسين / جامعة بابل / كلية التربية الأساسية

اسيل حمود حمزة / جامعة بابل / كلية الحقوق
المستخلص:

سهاد عبد الزهرة الشكري / جامعة الفرات الاوسط التقنية، المعهد التقني النجف.

يشهد العالم الرقمي تصاعداً غير مسبوق في مدى وتطور الهجمات السيبرانية، مما يستلزم حلولاً أمنية أكثر ذكاءً وتلاوياً مع البيئة التقنية المتغيرة باستمرار. يرمي هذا البحث إلى استكشاف الدور المؤثر لتقنيات الذكاء الاصطناعي، وخاصة خوارزميات التعلم الآلي، في التصدي للهجمات السيبرانية المتقدمة، مثل هجمات التصيد الاحتيالي الذكي، البرمجيات الخبيثة المتطورة، وهجمات يوم الصفر. يشهد العالم الرقمي تصاعداً غير مسبوق في مدى وتطور الهجمات السيبرانية، مما يستلزم حلولاً أمنية أكثر ذكاءً وتلاوياً مع البيئة التقنية المتغيرة باستمرار. يرمي هذا البحث إلى استكشاف الدور المؤثر لتقنيات الذكاء الاصطناعي، وخاصة خوارزميات التعلم الآلي، في التصدي للهجمات السيبرانية المتقدمة، مثل هجمات يوم الصفر، هجمات التصيد الاحتيالي الذكي، البرمجيات الخبيثة المتطورة. يعتمد البحث أسلوب مقارنة بين طائفة من الخوارزميات الذكية المستعملة في أنظمة كشف التسلل (IDS)، على غرار: شجرة القرار، الشبكات العصبونية الاصطناعية (ANN)، دعم المتجهات (SVM)، وخوارزميات التعلم العميق. يجري تقدير أداء تلك الخوارزميات حسب معايير الدقة، معدل الاكتشاف، نسبة التنبؤات الكاذبة، وكفاءة المعالجة الوقتية. يعتمد البحث أسلوب مقارنة بين طائفة من الخوارزميات الذكية المستعملة في أنظمة كشف التسلل (IDS)، على غرار: شجرة القرار، الشبكات العصبونية الاصطناعية (ANN)، دعم المتجهات (SVM)، وخوارزميات التعلم العميق. يجري تقدير أداء تلك الخوارزميات حسب معايير الدقة، معدل الاكتشاف، نسبة التنبؤات الخاطئة، وكفاءة المعالجة الزمنية.

الكلمات المفتاحية: الذكاء الاصطناعي، الهجمات السيبرانية، أنظمة كشف الاختراق، التعلم الآلي والتعلم العميق، الشبكات العصبونية الاصطناعية

1-المقدمة:

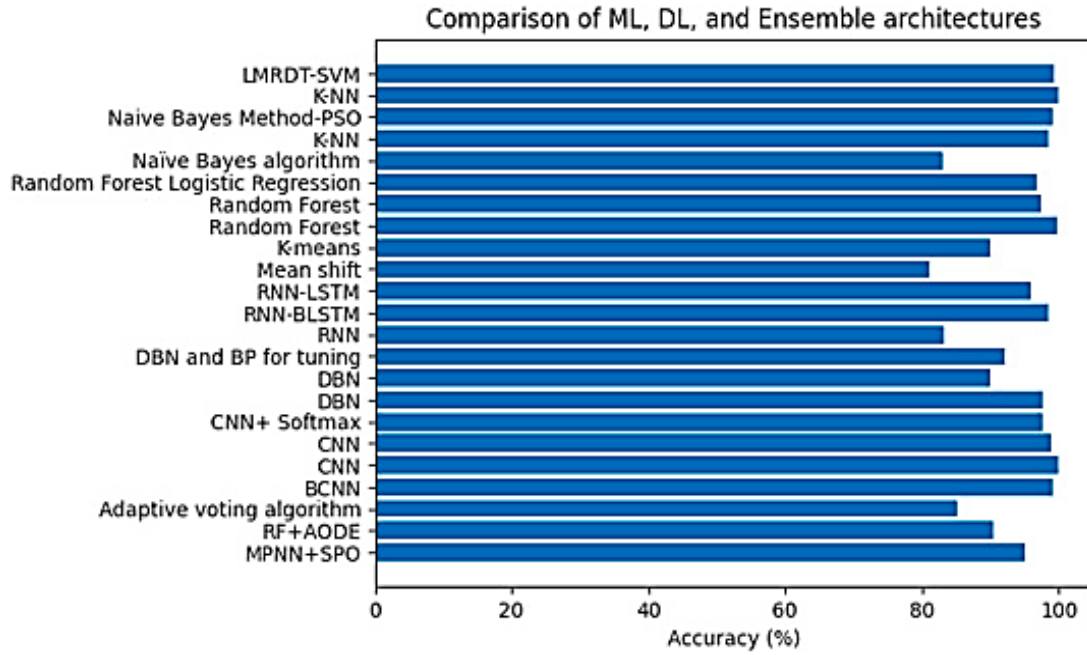
شهدت الهجمات الإلكترونية ازديادًا ملحوظًا، مما استوجب تطوير استراتيجيات أمنية مُحسّنة لحماية المعطيات الحساسة والبنية التحتية. وتشمل الهجمات الإلكترونية طائفة من الأنشطة الضارة التي تهدف إلى تفويض سلامة وسرية وتوفير أنظمة المعلومات، بما في ذلك الوصول غير المصرّح به إلى البيانات، وتعطيل الخدمات، وتسريبات البيانات، وهجمات الفدية. ومع اعتماد الشركات المتصاعد على المنصات الرقمية، أصبحت أهدافًا جذابة للمجرمين الإلكترونيين الساعين لتحقيق مكاسب مادية.

لقد بدّلت التقنيات المتقدمة مشهد المخاطر الإلكترونية، حيث أصبح المهاجمون يتبنون أساليب متطورة تستغل ثغرات الأنظمة مع القدرة على التهرب من أنظمة الكشف. فعلى سبيل المثال، تُمكن تقنيات التعلم الآلي العدائي المهاجمين من التلاعب بخوارزميات الذكاء الاصطناعي لتحقيق نتائج ضارة ([21]). علاوة على ذلك، فإن إدماج الذكاء الاصطناعي في الاستراتيجيات الهجومية يعزز من فعاليتها من خلال أتمّة العمليات بالاعتماد على تحليلات آنية ([15]).

أدى الإسراع في إنتاج البيانات واعتماد الحوسبة السحابية إلى تفاقم المخاوف الأمنية. ومع التوقعات بزيادة كبيرة في حجم المعلومات الخاصة المخزنة في بيئات الحوسبة السحابية ([8])، أصبح من الضروري إنشاء دفاعات قوية ضد أي اختراقات. أيضاً، الإجراءات الأمنية التقليدية الثابتة لم تعد كافية للتعامل مع التهديدات المتقدمة التي تستخدم تقنيات الذكاء الاصطناعي، مما يجعل من الضروري فهم الطبيعة المتغيرة للهجمات السيبرانية لوضع استراتيجيات دفاعية فعالة.

الشكل 1: دراسة أداء الدقة بين نماذج التعلم الآلي والتعلم العميق مقابل نماذج التجميع.

(المصدر: [8])



الشكل 1: تحليل أداء الدقة بين نماذج التعلم الآلي والتعلم العميق مقابل نماذج التجميع.
(المصدر: [8])

2-مشكلة البحث:

مع تطور الهجمات السيبرانية وتعقيدها، بات من الصعب على الأنظمة التقليدية كشف هذه التهديدات أو التعامل معها. وفي ظل هذا التحدي، ظهر الذكاء الاصطناعي كأداة واعدة في مجال الأمن السيبراني، خصوصاً من خلال استخدام الخوارزميات الذكية للكشف المبكر والتحليل والاستجابة للهجمات. غير أن تنوع هذه الخوارزميات يثير تساؤلات حول مدى فعاليتها وكفاءتها ومدى قدرتها على التصدي للهجمات السيبرانية المتقدمة (مثل هجمات يوم الصفر، والهجمات الموجهة).

3-أهمية البحث:

1. دعم متخذي القرار في نطاق الأمن السيبراني باختيار الخوارزمية الأفضل لبيئتهم التقنية.
2. المساهمة في تطوير أنظمة ذكية أكثر فعالية في الكشف والتعامل مع الهجمات المتقدمة.
3. توفير مرجع علمي للباحثين في حقلي الذكاء الاصطناعي والأمن السيبراني.

4-أهداف الدراسة

تتجلى الغايات الأساسية لهذا البحث في تحليل وتقدير مساهمات تقنيات الذكاء الاصطناعي في تحسين أطر الأمن السيبراني، خاصةً في ميدان أنظمة كشف الاختراق (IDS). وينشد البحث على وجه التحديد إلى دراسة طائفة متنوعة من نماذج التعلم الآلي والتعلم العميق، بما في ذلك

الشبكات العصبية الالتفافية (CNNs) والشبكات العصبية المتكررة (RNNs)، مع التركيز على فاعليتها في اكتشاف التهديدات السيبرانية المتقدمة وتقليل معدلات الإيجابيات والسلبيات الخاطئة.

بالإضافة إلى ذلك، تسعى الدراسة إلى استكشاف أساليب تطوير أداء هذه النماذج لتمكينها من الكشف الفوري عن المخاطر الناشئة، كما ورد في المرجع [6]. ويركز جانب آخر من البحث على دمج مصادر بيانات متعددة مثل بيانات حركة الشبكة، وملفات السجل، وتتبع استدعاءات النظام لتعزيز قوة ودقة أنظمة كشف التسلل. ويتضمن هذا استكشاف تقنيات معالجة البيانات المتطورة وخوارزميات الدمج التي يمكن استعمالها لتحقيق نتائج محسنة في كشف المخاطر. وكما أشار المرجع [17]، فإن فهم كيفية تسخير التعلم العميق بشكل فعال لمواجهة تحديات محددة في مجال الأمن السيبراني يشكل جانبا حيويا من هذا البحث.

المبحث الاول

أهمية الذكاء الاصطناعي في الأمن السيبراني

The importance of artificial intelligence in cyber security

المطلب الاول

ماهية الذكاء الاصطناعي

1- ماهية الذكاء الاصطناعي في الأمن السيبراني

يسهم الذكاء الاصطناعي بقدر كبير في تحسين الأمن السيبراني، خاصة مع تزايد صعوبة المخاطر السيبرانية. وعبر استعمال تقنيات مثل التعلم الآلي والتعلم العميق، يمكن لأجهزة الحماية اكتشاف الشذوذ بشكل ذاتي، مما يحسن من سرعة التجاوب للحوادث. وعلى نقيض الأساليب التقليدية التي تجد صعوبة في التعامل مع المخاطر الحديثة، يستطيع الذكاء الاصطناعي تحليل كميات كبيرة من البيانات لاكتشاف أنماط تشير إلى هجمات محتملة. وتُعتبر هذه المقدرة ضرورية في ظل ما ذُكر عن زيادة بنسبة 75% في الهجمات السيبرانية خلال السنوات الأخيرة، استناداً لقاعدة بيانات الأحداث السيبرانية.

كما يساند الذكاء الاصطناعي التحليل التنبؤي، الأمر الذي يسمح للمنظمات بالتنبؤ بالتهديدات المستقبلية استناداً إلى التوجهات التاريخية. وتتيح هذه المقاربة الاستباقية لفرق الأمن السيبراني اتخاذ إجراءات وقائية بدلاً من مجرد الرد بعد وقوع الحوادث. وكما هو مبين في

المرجع [5]، تساعد القدرات التنبؤية للذكاء الاصطناعي على تحديد أولويات الموارد بفعالية وتقليل المخاطر قبل تفاقمها.

وبالمثل، من الأساسي أن ندرك أن الذكاء الاصطناعي، على الرغم من مساعدته في تحسين الاكتشاف، يطرح صعوبات جديدة. قد تتسبب الهجمات المعادية في التلاعب ببيانات الإدخال، مما يقود إلى قرارات خاطئة أو تصنيف غير دقيق للمخاطر. لهذا السبب، فإن التطوير المستمر لتقنيات الذكاء الاصطناعي ضروري لتعزيز دقة الاكتشاف وبناء أطر أمن سيبراني متينة تتأقلم مع تطور أساليب الهجوم. ويُعد الاستخدام الاستراتيجي للذكاء الاصطناعي تحولاً جوهرياً نحو الكشف الآلي عن التهديدات في بيئات الأمن السيبراني المعاصرة.

وبالمثل، من الأساسي أن ندرك أن الذكاء الاصطناعي، على الرغم من مساعدته في تحسين الاكتشاف، يطرح صعوبات جديدة. قد تتسبب الهجمات المعادية في التلاعب ببيانات الإدخال، مما يقود إلى قرارات خاطئة أو تصنيف غير دقيق للمخاطر. لهذا السبب، فإن التطوير المستمر لتقنيات الذكاء الاصطناعي ضروري لتعزيز دقة الاكتشاف وبناء أطر أمن سيبراني متينة تتأقلم مع تطور أساليب الهجوم. ويُعد الاستخدام الاستراتيجي للذكاء الاصطناعي تحولاً جوهرياً نحو الكشف الآلي عن التهديدات في بيئات الأمن السيبراني المعاصرة.

المطلب الثاني

نظرة عامة على تقنيات الذكاء الاصطناعي

1-تعريف وأنواع الذكاء الاصطناعي

يشمل الذكاء الاصطناعي في ميدان الأمن السيبراني مجموعة من التقنيات التي ترمي إلى تعزيز الدفاعات ضد التهديدات السيبرانية. ويركّز الذكاء الاصطناعي على تطوير أنظمة قادرة على أداء مهام تستوجب ذكاءً شبيهاً بالبشري، على غرار التعلم وحل الإشكالات. وتشمل التقنيات الجوهرية كلاً من التعلم الآلي (ML)، والتعلم العميق (DL)، والتعلم المعزز. يسمح التعلم الآلي للأنظمة باستخلاص رؤى من المعطيات دون برمجة صريحة، حيث تستعمل الخوارزميات لتحليل مجموعات ضخمة من المعطيات بهدف اكتشاف الأنماط والتنبؤات. وينقسم التعلم الآلي إلى ثلاثة أنواع رئيسية: التعلم الخاضع للإشراف، والتعلم غير الخاضع للإشراف، والتعلم المعزز، وكل نوع منها يخدم غرضاً مميزاً.

أما التعلم العميق، فهو تخصص أكثر تطوراً من التعلم الآلي، ويستخدم شبكات عصبية متعددة الطبقات لتحليل المعطيات المعقدة. وتُحاكي هذه الطريقة وظائف الدماغ البشري، مما يمكنها من استخلاص سمات وتجريدات عالية المستوى من المعطيات الخام. ويتميز التعلم العميق بقدرته على التعامل مع كميات ضخمة من المعطيات، كما هو الحال في بيئات الأمن السيبراني.

وتستند هياكل الأمن السيبراني المعتمدة على الذكاء الاصطناعي على هذه المنهجيات لتعزيز فعالية أنظمة الكشف في التعامل مع التهديدات، متخطيةً بذلك الأنظمة التقليدية المعتمدة على القواعد. ويساعد دمج الذكاء الاصطناعي في تحسين اكتشاف الهجمات المعقدة وتقليل الاعتماد على الإشراف البشري، كما هو مبين في المرجع [18]

يساهم دمج الذكاء الاصطناعي بشكل كبير في تحسين كشف الهجمات المعقدة، مع تقليص الاعتماد على الإشراف البشري.

2- التعلم الآلي مقابل الأساليب التقليدية

يُشكّل التعلم الآلي تقدماً كبيراً في مجال الأمن السيبراني، حيث يختلف بشكل كبير مع الأساليب القديمة التي تعتمد على أنظمة قائمة على قواعد ثابتة لتحديد المخاطر. وتواجه هذه الاستراتيجيات القديمة صعوبة في التكيف مع مسارات الهجوم الجديدة، مما يجعلها عرضة لتجاهل المخاطر المتقدمة غير المدرجة ضمن قواعدها المحددة سلفاً.

وعلى النقيض من ذلك، يستعمل التعلم الآلي خوارزميات تمكن الأنظمة من التعلم الذاتي من البيانات، وتحسين أدائها دون الحاجة إلى برمجة واضحة. وتعزز هذه القدرة التكيفية نهجاً استباقياً في كشف المخاطر. وكما ورد في المرجع [18]، فإن أطر الأمن المعتمدة على الذكاء الاصطناعي قادرة على تحليل كميات ضخمة من البيانات، والتعرف على أنماط معقدة، والاستجابة للتهديدات المتطورة بأقل تدخل بشري. ومن خلال الاستفادة من البيانات التاريخية، تستطيع نماذج التعلم الآلي التنبؤ بالهجمات المحتملة والتعامل بفعالية مع الثغرات المعروفة والناشئة.

ومن الفوائد الأساسية للتعلم الآلي مقدرته على معالجة كميات هائلة من البيانات بسرعة وفعالية. وتُستعمل خوارزميات متعددة — على سبيل المثال أشجار القرار، وآلات الدعم الناقل (SVMs)، والشبكات العصبية — في مهام كالتقيب عن الاختراقات وتصنيف البرمجيات

الضارة ([1]). فضلاً عن ذلك، بينما تحتاج الأنظمة التقليدية خبرة كبيرة لتحديثها، فإن نماذج التعلم الآلي تحسن فهمها دائماً بالاستناد إلى البيانات الجديدة ([2])، الأمر الذي يعزز من معدلات الكشف ويخفف من الإيجابيات الكاذبة. ويمثل هذا التوجه نحو التعلم الآلي تغييراً نحو دفاعات أمنية سيبرانية أكثر ذكاءً وحيوية.

المبحث الثاني

الهجمات السيبرانية و أنظمة كشف التسلل

Cyber Attacks and Intrusion Detection Systems

المطلب الاول

أنواع الهجمات السيبرانية المتقدمة

1-أنواع الهجمات السيبرانية المتقدمة

يتصف المشهد المتصاعد للتهديدات الإلكترونية بأساليب متزايدة التعقيد، على غرار حملات التصيد المتطورة، والبرمجيات الضارة المعقدة، واستغلالات "اليوم الصفرى"، الأمر الذي يشكل صعوبة للأطر التقليدية للأمن السيبراني. ويستخدم التصيد المتقدم تقنيات الذكاء الاصطناعي لتحليل البيانات الشخصية من منصات التواصل الاجتماعي، مما يتيح إنشاء رسائل مخصصة تزيد من معدلات النجاح بشكل كبير.

علاوة على ذلك، تتأقلم البرمجيات الضارة الحديثة عبر التعلم من أنظمة الدفاع. وكما ذكر في المرجع [15]، يمكن للذكاء الاصطناعي أن يمكن المهاجمين من تصميم برمجيات ضارة تحاكي التطبيقات المشروعة، مما يسمح لها بالبقاء داخل الشبكات لفترات طويلة دون اكتشافها.

إن استغلالات "اليوم الصفرى"، هي تُعتبر من أشد المخاطر، حيث تستغل فجوات غير معروفة في البرمجيات، الأمر الذي يجعل اكتشافها أكثر صعوبة. ووفقاً للمصدر [16]، غالباً ما تعجز وسائل الحماية التقليدية في التصدي لهذا الصنف من الهجمات بسبب عدم وجود علم مسبق يُستخدم في أنظمة الكشف. وتتطلب سرعة تطور تهديدات "اليوم الصفرى" من خبراء الأمن السيبراني تبني منهجيات كشف متطورة وخوارزميات تعلم آلي قادرة على التعرف على الأنماط غير المألوفة التي قد تشير إلى مثل هذه التوغلات.

ومع تطور مهارات المهاجمين بالتوازي مع تطور التقنية، يجب على المؤسسات تقوية دفاعاتها من خلال تحديثات مستمرة واستراتيجيات مرنة تهدف إلى التعرف على هذه التهديدات المتقدمة والتصدي لها بفعالية.



الشكل ٢: الهجمات السيبرانية باستعمال الذكاء الاصطناعي

(المصدر: المصدر [15])

المطلب الثاني

أنظمة كشف التسلل

1- أنظمة كشف التسلل (IDS) وأهميتها

1-1 دور أنظمة كشف التسلل في الأمن السيبراني

تلعب أنظمة كشف التسلل (IDS) دوراً حيوياً في مشهد الأمن السيبراني من خلال مراقبة حركة مرور الشبكة بشكل مستمر لاكتشاف أي إشارات على سلوك ضار أو خرق للسياسات المعتمدة. وعلى نقيض الجدران النارية التقليدية التي تعتمد على قواعد ثابتة، تُصمَّم أنظمة IDS للتكيف مع الديناميكيات المتغيرة لبيئات الشبكات، مما يجعلها ضرورية لاكتشاف التهديدات السيبرانية المتطورة.

تستخدم هذه الأنظمة توليفة من استراتيجيات الكشف القائمة على التوقيع وعلى الشذوذ. فالكشف القائم على التوقيع يعتمد على أنماط محددة مسبقاً مرتبطة بهجمات معروفة، بينما يقوم الكشف القائم على الشذوذ ببناء نموذج للسلوك الطبيعي ليتمكن من تحديد الانحرافات التي قد تشير إلى مخاطر محتملة. وتُعد هذه المرونة ضرورية بشكل خاص في ظل استمرار تطور أساليب الهجمات السيبرانية واستغلالها لثغرات مجهولة.

لقد أفضى دمج تقنيات التعلم الآلي (ML) والتعلم العميق (DL) في أنظمة IDS إلى تحسين أدائها بشكل كبير. فمن خلال الاستفادة من خوارزميات متطورة، تستطيع هذه الأنظمة تحليل مجموعات بيانات ضخمة، مما يعزز من قدرتها على التعرف على هجمات جديدة غير مسبقة، مع تقليل الإيجابيات الكاذبة — وهي مشكلة لطالما أثقلت كاهل بروتوكولات الأمن التقليدية. وكما ورد في المرجع [19]، فإن دمج نماذج التعلم الآلي في أنظمة IDS يسهل المراقبة الآنية والكشف عن الشذوذ، وهو أمر بالغ الأهمية للحفاظ على تدابير أمنية قوية. علاوة على ذلك، تشير الدراسات إلى أن أنظمة IDS الحديثة تعتمد بشكل متزايد على مجموعات بيانات أوسع لتدريب نماذجها، مما يسمح لها بالتأقلم بشكل أفضل مع الأنواع الجديدة من التسللات، كما ورد في المرجع [16]. ويعكس التطور المستمر لهذه الأنظمة وتعزيزها التعقيد المتزايد للتهديدات السيبرانية، ويؤكد على الحاجة الماسة إلى تبني استراتيجيات استباقية لحماية الأصول الرقمية.

1-2 التحديات التي تواجهها أنظمة كشف التسلل التقليدية

تواجه أنظمة كشف التسلل التقليدية (IDS) العديد من الصعوبات التي تقلل من فاعليتها في بيئة الأمن السيبراني الحديثة. وتتمثل إحدى الإشكاليات الرئيسة في اعتمادها على مجموعات قواعد ثابتة، الأمر الذي يحد من قدرتها على اكتشاف التهديدات الجديدة أو المتقدمة، بما في ذلك ثغرات "اليوم الصفر". وكما أُشير في المرجع [13]، فإن أنظمة مثل Snort و Suricata تعتمد على توقيعات لهجمات معروفة، مما يجعلها عرضة للتجاوز عن طريق نسخ مُعدّلة من هذه التهديدات، ويؤدي ذلك إلى ارتفاع معدل الإيجابيات الكاذبة بسبب تفسير الأفعال الحميدة على أنها خبيثة وفقاً للقواعد المحددة سلفاً.

بالإضافة إلى ذلك، تُعاني أنظمة IDS التقليدية من صعوبة التكيف السريع مع المشهد المتغير للتهديدات، كما ورد في المرجع [9]. وغالباً ما تتطلب هذه الأنظمة موارد حسابية كبيرة لتحليل الحزم عالية السرعة، الأمر الذي يُعرضها لخطر فقدان الحزم أو عدم اكتشاف الأنشطة الخبيثة إذا لم تتمكن من معالجة البيانات بكفاءة ([14]).

بالإضافة إلى ذلك، تُنتج أنظمة IDS التقليدية عدداً كبيراً من التحذيرات غير المُصنّفة، وتعجز عن التمييز بين الشذوذ المشروع والتسللات الفعلية، مما يُربك مسؤولي الشبكات ([16]). ولمواجهة هذه التحديات، يزداد الاتجاه نحو دمج الذكاء الاصطناعي والتعلم الآلي في هياكل

عمل IDS لتحسين دقة الكشف وتقليل الإنذارات الكاذبة. كما يلجأ المهاجمون إلى تقنيات الإخفاء مثل أدوات الإخفاء (Anonymizers)، مما يصعب عمليات الكشف ([11])، ويستلزم استراتيجيات متقدمة تعتمد على تحليل السلوك واكتشاف الشذوذ.

طرق الكشف	المزايا	العيوب
SIDS نظام كشف التسلل المعتمد على التوقيع	- فعال جدًا في التعرف على التسلات مع أقل عدد من الإنذارات الكاذبة (FA) - يحدد التسلات بسرعة. - متفوق في اكتشاف الهجمات المعروفة. - تصميمه بسيط.	- يحتاج إلى تحديثات متكررة لتواقيع جديدة. - تم تصميم SIDS لاكتشاف الهجمات التي تحمل توقيعات معروفة، فإذا تم تعديل التسلل السابق بشكل طفيف إلى نوع جديد، فلن يتمكن النظام من التعرف على هذا التغيير الجديد. - غير قادر على اكتشاف هجمات اليوم الصفر (Zero-Day) - غير مناسب لاكتشاف الهجمات متعددة المراحل. - فهم محدود لجوهر الهجمات.
AIDS نظام كشف التسلل المعتمد على الشذوذ	- يمكن استخدامه لاكتشاف الهجمات الجديدة. - يمكن استخدامه لإنشاء توقيع للهجوم.	- لا يمكنه التعامل مع الحزم المشفرة، مما يسمح للهجوم بالبقاء دون كشف وقد يشكل تهديدًا. - معدل مرتفع من الإنذارات الكاذبة. - من الصعب بناء ملف سلوك طبيعي لنظام كمبيوتر ديناميكي للغاية. - التنبهات غير مصنفة. - يحتاج إلى تدريب مبدئي.

الجدول 1: مقارنات بين أساليب كشف الاختراق (المصدر: [16])

الخصائص	أمثلة	منهجية الكشف
تتطلب معرفة واسعة بالإحصاء بسيطة ولكن أقل دقة تعمل في الزمن الحقيقي	Bhuyan, et al. (2014)	القائمة على الإحصاء: تحليل حركة مرور الشبكة باستخدام خوارزميات إحصائية معقدة لمعالجة المعلومات.
تكلفة الحساب عالية لأن النظام يحتاج لمطابقة الأنماط مع القواعد من الصعب التنبؤ بالأحداث وتوقيتها يتطلب عددًا كبيرًا من القواعد لتغطية جميع أنواع الهجمات الممكنة معدل منخفض من الإيجابيات الكاذبة معدل كشف مرتفع	Hall, et al. (2009)	القائمة على القواعد: تستخدم "توقيع" الهجوم للكشف عن هجوم محتمل في حركة المرور الشبكية المشبوهة.
احتمالية، وتتعلم بشكل ذاتي معدل منخفض من الإيجابيات الكاذبة	Kenkre, et al. (2015a)	القائمة على الحالة: تفحص سلسلة من الأحداث لتحديد أي هجوم محتمل.
تتطلب المعرفة والخبرة تعتمد على التعلم التجريبي والتطوري	Abbasi, et al. (2014) Butun, et al. (2014)	القائمة على الاستدلال (الحدس): (تحدد أي نشاط غير طبيعي يخرج عن النمط المعتاد).

الجدول 2: سمات منهجيات الكشف لأنظمة اكتشاف التطفل

(المصدر: [16])

3-1 التحليل المقارن للخوارزميات في أنظمة كشف التسلل

1-3-1 أشجار القرار

أشجار القرار هي أسلوب تعلم آلي مرّن يُستخدم على نطاق واسع في أنظمة كشف التسلل (IDS) بفضل سهولتها ويسر تفسيرها. تقوم بتقسيم المعطيات إلى مجموعات فرعية بناءً على قيم الخصائص المدخلة، مما يؤدي إلى التصنيفات في العقد النهائية للشجرة. هذه الطريقة مفيدة لاتخاذ القرارات بسرعة، وهو أمر ضروري لاكتشاف التسلل في الوقت الفعلي.

تتمثل إحدى المزايا البارزة لأشجار القرار في قدرتها على التعامل مع البيانات التصنيفية والمستمرة، مما يتيح التكيف عبر مجموعات البيانات المختلفة وأنواع الهجمات المتنوعة. هذه النماذج تتطلب الحد الأدنى من المعالجة الأولية للبيانات مقارنة بالخوارزميات الأكثر تعقيداً، مما يسهل تطبيقها بشكل أسرع في مجال الأمن السيبراني. وفقاً للمرجع [11]، يمكن لأشجار القرار اكتشاف الشذوذ بفعالية من خلال تحديد عتبات واضحة لمؤشرات النشاط المشبوه.

أوضحت البحوث أن الخوارزميات مثل C4.5 تحرز معدلات دقيقة مرتفعة في تحديد الاختراقات، حيث تتخطى معدلات الكشف 95% لبعض نواقل الهجوم المحددة مثل هجمات [16] (DDoS). غير أن، لا تزال هناك تحديات مثل الإفراط في التلاؤم مع بيانات التدريب (Overfitting).

لمعالجة معضلة الإفراط في التلاؤم، تجمع الأساليب المجمعّة بين عدة أشجار قرار باستخدام تقنيات مثل الـ "Bagging" أو "Boosting"، ما يحسن دقة التنبؤ ويقلل من معدلات الإيجابيات الكاذبة ([19]). يقوي هذا الدمج من مهمة IDS، ما يمكن المحللين من تتبع التوقعات وفهم الأنماط الأساسية.

استطلاع	عدد الاقتباسات (حتى تاريخ 2019/6/1)	تقنيات نظام كشف التسلل	خلل في مجموعة البيانات					
SIDS	AIDS	Hybrid IDS						
التعلم المُراقب	التعلم غير المُراقب	التعلم شبه المُراقب	طرائق التجمّع في تعلم الآلة					
اسم مؤلف، "1988"	219	✓	×	×	×	×	×	×
أكسيلسون (2000)	1039	✓	✓	×	×	×	×	×
لياو وآخرون (2013)ب)	505	✓	✓	✓	×	×	×	×
أغراوال وأغراوال (2015)		✓	✓	✓	✓	✓	✓	×

							108	
✓	✓	✓	×	✓	✓	✓	338	بوشراك و غوين (2016)
✓	×	×	×	✓	✓	×	181	أحمد وآخرون (2016)
✓	✓	✓	✓	✓	✓	✓		الدراسة الاستقصائية

الجدول 3: مقارنة بين هذه الدراسة والبحوث المشابهة: (الموضوع مُغطى، الموضوع غير

مُغطى)

(المصدر: [16])

قواعد البيانات	النتيجة	الملاحظات	المصادر
DARPA 98	اكتشاف Snort ، 69% من التنبيهات الناتجة تعتبر إنذارات كاذبة.	تم تطبيق SIDS بدون AIDS	Hu, et al. (2009)
ANN	تحليل الأنظمة باستخدام الشبكات العصبية الاصطناعية (ANN) ، معدل الاكتشاف 96%.	تم تنفيذ مصنف يعتمد على الشبكات العصبية الاصطناعية (ANN) لتحضير واختبار الإطار.	McHugh (2000)
SVM	على مجموعة فرعية من DARPA 98 ، معدل الاكتشاف 99.6%.	يعزل SVM المعلومات إلى فئات مختلفة بواسطة مستوى أو مستويات هايبر ، لأنه يمكنه التعامل مع المعلومات متعددة الأبعاد. عادةً ما يظهر SVM أداءً جيداً لمشاكل التصنيف الثنائية.	Chen, et al. (2005)
KDDCUP 99	تحليل إحصائي متعدد المتغيرات للبيانات المدققة، معدل الاكتشاف 90%.	يستخدم التحليل متعدد المتغيرات لتقليل معدلات الإنذارات الكاذبة.	Ye, et al. Hotta, (2002) et al. (2008)
	أفضل النتائج تحققت بواسطة خوارزمية C4.5 التي حققت معدل إيجابي حقيقي 95%.	الأشجار القرار التي أنشأتها C4.5 يمكن استخدامها للتصنيف.	Ferrari and Cribari-Neto (2004); Shafi and Abbass (2013); Laskov, et al. (2005)
	مصنف SMO ، معدل الاكتشاف 97%.	هذا المصنف المعتمد على SVM مع تنفيذ SMO يعطي دقة اكتشاف جيدة. ومع ذلك، فإن الدقة المُبلغ عنها أقل من تلك في (Chen et al., 2005) ، لأن مجموعة بيانات KDDCUP 99 أكثر تعقيداً وشمولاً من مجموعة بيانات DARPA 98.	Shafi and Abbass (2013)
	أفضل نموذج هو نموذج HNB ، حيث يتم استخدام مستوى ثقة 95% لمقارنة النماذج.	تقنيات Bayes الخفية (HNB) يمكن تطبيقها على مجال IDS الذي يعاني من الأبعاد العالية والسمات المرتبطة بشدة وسرعة الشبكة العالية. تقنية HNB أفضل من الطريقة التقليدية (NB) من حيث دقة	Koc, et al. (2012)

	الاكتشاف في IDS.		
Adebowale, et al. (2013)	تستخدم خوارزمية k-NN جميع الحالات المصنفة كنموذج للدالة الهدف. خلال مرحلة التصنيف، تستخدم k-NN استراتيجية بحث قائمة على التشابه لتحديد دالة فرضية محلية مثلى.	خوارزمية الجار الأقرب-k (NN)، معدل الاكتشاف 94.0%	NSL-KDD
Adebowale, et al. (2013)	توفر مصنفات Bayes دقة معتدلة لأن التركيز يكون على تصنيف الفئات للحالات، وليس الاحتمالات الدقيقة.	Naïve Bayes، معدل الاكتشاف 89.0%	
Thaseen and Kumar (2013)	تختار C4.5 السمة الأكثر فعالية من البيانات لتقسيم مجموعة العينات إلى مجموعات فرعية، مما يساهم في تحسين الدقة.	C4.5 أعطت أفضل معدل اكتشاف 99.0%	
Adebowale, et al. (2013)	العمل يستخدم أيضاً مصنفاً معتمداً على SVM ويحقق معدل اكتشاف مشابه لـ (Chen et al., 2005)	مصنف SMO ، معدل الاكتشاف 97.0%	
Ahmed, et al. (2016)	يشكل "EM مهمة ناعمة" لكل صف إلى مجموعات مختلفة بالنسبة للاحتمالية لكل مجموعة. الدقة في هذه الطريقة منخفضة لأن EM لا يعطي مصفوفة التباين للمعايير الأخطاء.	التجميع باستخدام Expectation Maximization (EM)، الدقة 78.0%	
Creech and Hu (2014b)	ADFA-WD هو مجموعة بيانات جديدة تحتوي على هجمات جديدة. لهذا السبب كانت الدقة المبلغ عنها أقل مقارنة بكل تقنيات التعلم الآلي عند مقارنتها بالدقة باستخدام بيانات KDD98 القديمة. تم الإبلاغ عن أن SVM ينتج أعلى دقة.	استخدم Creech وآخرون نموذج ماركوف الخفي (HMM)، آلة التعلم الفائق (ELM)، و SVM. أبلغوا عن دقة 74.3% لـ HMM ، دقة 98.57% لـ ELM ، ودقة 99.64% لـ SVM.	ADFA-WD
Creech and Hu (2014b)	تم تطبيق سمات دلالية جديدة. لذلك، فإن ELM قادر على استخدام السمات الدلالية الجديدة بسهولة وسرعة من خلال تضمين كمية من العبارات الدلالية.	دقة 100% عند استخدام ELM باستخدام السمات الدلالية الأصلية.	ADFA-LD
Usteba, et al. (2018)	يتم اختيار السمات باستخدام خوارزمية Fisher Score.	دقة 94.5% تم الحصول عليها باستخدام MLP فقط، باستخدام MLP و Payload Classifier معاً تم اكتشاف دقة 95.2%	CICIDS2017
Koroniotis, et al. (2018)	حققت هذه الطريقة المعتمدة على SVM دقة اكتشاف جيدة.	أعلى دقة من نموذج SVM. معدل الاكتشاف 98.0%	Bot-IoT

الجدول 4: مقارنة النتائج التي تم تحقيقها بواسطة طرق مختلفة على مجموعات بيانات IDS

المتاحة علناً المصدر: [16]

المبحث الثالث

الشبكات العصبية الاصطناعية وخوارزميات التعلم العميق

Artificial Neural Networks and Deep Learning Algorithms

المطلب الاول

الشبكات العصبية الاصطناعية (ANNs)

1- الشبكات العصبية الاصطناعية (ANNs)

تعتبر الشبكات العصبية الاصطناعية (ANNs) جوهرية في تطوير أنظمة كشف التسلل (IDS) المتينة. مستوحاة من الهياكل العصبية البيولوجية، تقوم ANNs بتحديد الأنماط عبر العقد المترابطة أو "الخلايا العصبية"، مع وصلات قابلة للتعديل تتيح لها التعلم من مدخلات البيانات. تجعل هذه القدرة على التكيف من ANNs فعالة في التعرف على المخاطر السيبرانية المعقدة والتطور مع أنماط الهجوم الحديثة.

في ميدان الأمن السيبراني، تتميز ANNs في تصنيف التسللات والشذوذات المختلفة في حركة مرور الشبكة. تسمح مرونتها بالتنقل بين العلاقات غير الخطية بين ميزات المدخلات، وهو أمر ضروري للتعامل مع بيانات التسلل المعقدة. تشير الدراسات إلى أن الأنظمة المعتمدة على ANNs يمكن أن تحقق معدلات اكتشاف تتجاوز 96% على مجموعات بيانات مثل KDDCup99 و NSL-KDD [16].

ميزة مهمة لـ ANNs في IDS هي قدرتها على استخلاص الخصائص بشكل مستقل، مما يقلل الحاجة إلى التصميم المعقد للخصائص ويساعد في تحديد العلامات الحيوية للاختراقات. تعزز الهياكل المتطورة مثل الشبكات العصبية الالتفافية (CNNs) والشبكات العصبية المتكررة (RNNs) الأداء في تحليل البيانات المكانية والمتسلسلة على التوالي.

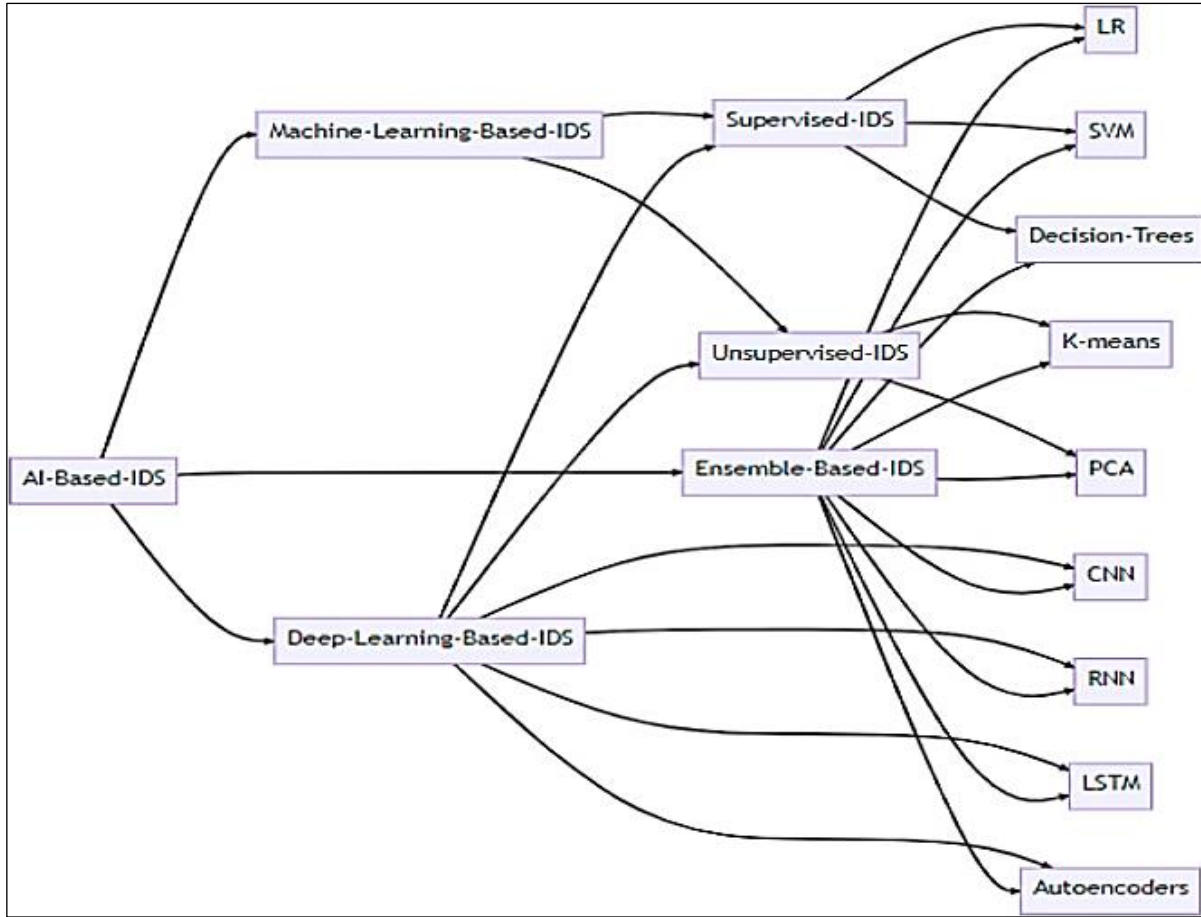
إن نشر ANNs يواجه تحديات مثل الإفراط في التكيف وارتفاع المتطلبات الحاسوبية. يمكن تقليل هذه المشكلات من خلال تعديل متغيرات الشبكة واستخدام تقنيات مثل "Dropout" أو التنظيم. إن دمج ANNs في إطارات الأمن السيبراني يظهر وعدًا كبيرًا في تحسين الكشف ضد التهديدات السيبرانية المتطورة.

2- آلات الدعم الناقل (SVMs)

تعتبر آلات الدعم الناقل (SVMs) تقنية تعلم آلي قوية تُستخدم في أنظمة كشف التسلل، ومعروفة بقدرتها على تصنيف البيانات إلى مجموعات متميزة. يقوم خوارزم SVM بتحديد المستوى الفائق الأمثل الذي يفصل بين الفئات المختلفة في فضاء الميزات، مما يدير بشكل فعال مجموعات البيانات القابلة للفصل خطيًا وغير خطيًا. من خلال استعمال دوال النواة، يعزز SVM تحديد الفئات في الفضاءات عالية الأبعاد.

في مجال الأمن السيبراني، يتم تطبيق SVMs في مجالات مثل اكتشاف البرمجيات الخبيثة وتحديد التسلات الشبكية، حيث تحقق معدلات دقة عالية في التمييز بين الحركة الضارة والحركة الحميدة. تظهر النتائج التجريبية أن SVM يمكن أن يصل إلى معدل إيجابيات حقيقية قدره 99% مع تقليل الإيجابيات الكاذبة، مما يبرز تفوقه على الأساليب التقليدية ([3] ص. 1-5). إضافة إلى ذلك، يمكن تحسين أداء SVM عن طريق دمج تقنيات الأساليب المجمعة أو النماذج الهجينة.

على الرغم من متانتها، تواجه SVMs صعوبات مثل تأثرها باختيار دالة النواة وتزايد المتطلبات الحاسوبية مع مجموعات البيانات الأكبر. إن المعالجة الأولية الفعالة وضبط المعايير الأساسية مهمان لتحسين النتائج، خاصة فيما يخص المعاملات مثل عامل التكلفة ونوع النواة ([7]). بوجه عام، تعزز قابلية SVM قيمتها في التصدي للمخاطر السيبرانية، مع استمرار التقدم في تعلم الآلة واعدًا بتحسينات إضافية في أطر الأمن السيبراني المدفوعة بالذكاء الاصطناعي.



الشكل 4: تصنيف نظم كشف التسلل المستندة على الذكاء الاصطناعي.

(المصدر: [8])

المطلب الثاني

خوارزميات التعلم العميق

1- خوارزميات التعلم العميق

1-1 مقاييس الدقة

تلعب مقاييس الضبط دوراً هاماً في تقييم فعالية أنظمة كشف التسلل (IDS) التي تعتمد على الذكاء الاصطناعي (AI). تبرز هذه المقاييس مدى قدرة الخوارزمية على التمييز بين الحالات الإيجابية الحقيقية (الهجمات الفعلية) والحالات الإيجابية الكاذبة (الأنشطة البريئة التي تم تصنيفها كتهديدات). يعد تحقيق معدلات ضبط عالية أمراً مفيداً بشكل خاص، حيث يشير إلى موثوقية IDS في تمييز الأنشطة المشروعة عن النوايا الخبيثة. على سبيل المثال، تُظهر العديد من الدراسات البحثية أن خوارزميات مثل (K-Nearest Neighbors (KNN يمكن

أن تحقق أرقام ضبط تصل إلى 99.89%، مما يبرز إمكاناتها ضمن أطر التعلم الخاضع للإشراف، كما هو مذكور في [8].

ومع ذلك، من الضروري أن نفهم أن تصميم الخوارزمية ليس العامل المحدد الوحيد للدقة؛ فهناك العديد من الجوانب مثل اختيار الخصائص وطرق معالجة البيانات التي تلعب أدواراً حاسمة أيضاً. على سبيل المثال، أظهرت غابات العشوائية (Random Forests) مستويات ضبط ممتازة تصل إلى حوالي 99.64% عندما يتم ضبط المعلمات والخصائص بدقة. في المقابل، قد تفشل بعض التقنيات في تحقيق الضبط المطلوبة بسبب اعتمادها على الحسابات المعقدة أو مجموعات البيانات التدريبية الضعيفة.

بالإضافة إلى ذلك، مجموعة من تقنيات التعلم الآلي تحقق نتائج دقة متباينة عند تقييمها مقابل مجموعات بيانات مشهورة مثل KDD'99 Cup و UNSW-NB15. أظهرت بعض النماذج أداءً متميزاً، حيث قاربت معدلات الدقة الـ 100%. ومع ذلك، بينما يعدّ الحصول على دقة عالية ضرورياً للكشف الفعال عن التهديدات، يجب توازنه مع مقاييس أخرى مهمة مثل معدلات الإنذار الكاذب وكفاءة المعالجة. إن تحقيق هذه التوازنات أمر أساسي لتطوير نظام كشف الدخيل قوي قادر على العمل بسلاسة في المواقف الفعلية.

2-1 مقاييس معدل الكشف

تُعد مقاييس نسبة الكشف أمراً جوهرياً لتقييم فعالية أنظمة كشف التسلل (IDS)، حيث تشير إلى قدرة النظام على تحديد التسللات الفعلية بدقة. تعكس معدلات الكشف الأعلى كفاءة IDS في تحديد الأنشطة الضارة، مما يعزز الأمن العام للشبكة. أظهرت الدراسات الحديثة أن الأنظمة التي تعتمد على الذكاء الاصطناعي تحسن بشكل كبير نسب الكشف مقارنة بالطرق التقليدية. على سبيل المثال، حقق Gulia وزملاؤه معدل كشف قدره 96% باستخدام شبكة عصبية عميقة مع خوارزمية مستعمرة النحل الاصطناعية الجماعية ([2]). علاوة على ذلك، وصلت بعض النماذج التي تستخدم تقنيات التجميع إلى 98.60% عند تحديد مواقع التصيد الاحتيالي من خلال استخراج الميزات المتقدم ([3] ص 1-5).

بينما يُعد السعي لتحقيق معدلات كشف عالية أمراً مفيداً، من المهم موازنة ذلك مع مقاييس الأداء الأخرى مثل معدلات الإنذار الكاذب وكفاءة المعالجة. تعد خوارزمية T-SNERF مثلاً على هذا التوازن، حيث حققت معدل كشف مثالي بنسبة 100% مع عدم وجود أي إنذارات

كاذبة على مجموعات بيانات محددة ([7]). لا تزال هناك تحديات في الحفاظ على نسب كشف مرتفعة عبر مختلف المتجهات والهجمات وتكوينات الشبكات. يمكن أن يؤدي دمج العديد من الخوارزميات أو استخدام تقنيات التجميع إلى تعزيز الأداء العام من خلال معالجة قيود النماذج الفردية ([8])، مما يبرز الحاجة إلى الابتكار المستمر في اختيار السمات وتدريب النماذج في IDS المعتمد على الذكاء الاصطناعي.

1-3 مقاييس معدل الإنذار الكاذب

معدل الإنذار الخاطئ (FAR) هو مقياس أساسي لتقييم أنظمة كشف الدخيل (IDS)، حيث يمثل الأنشطة الشرعية التي صُنفت خطأً على أنها ضارة. يمكن أن يؤدي FAR مرتفع إلى إرهاب العاملين في الأمن السيبراني، مما يزيد من أوقات الاستجابة ويعرض التهديدات غير المكتشفة للخطر. على سبيل المثال، غالبًا ما تنتج أنظمة مثل Snort و Suricata تنبيهات كاذبة بسبب تعارض القواعد التي تؤثر على أنماط حركة المرور المشروعة، مثل استعلامات DNS ([13]). تتطلب هذه الوضعية موارد حاسوبية كبيرة لتصنيف حركة المرور بدقة. لمعالجة الإنذارات الكاذبة، أظهرت الخوارزميات المبتكرة للتعليم الآلي وعدًا كبيرًا. أظهرت دراسة مذكورة في [3] ص 1-5 تحقيق معدل FAR بنحو 1.297% باستخدام تقنيات التجميع المتطورة مدمجة مع الكشف عن الشذوذ. هذه التطورات تبرز فعالية التعلم الآلي في تقليل FAR مع الحفاظ على معدلات كشف عالية.

لقد أثبتت نماذج التعلم العميق أيضًا فعاليتها في تقليل الإنذارات الخاطئة. كما هو مذكور في [20]، فإن بعض النماذج تحقق معدل FAR أقل من 0.1% مع ضمان دقة كشف التهديدات فوق 99%. يعتبر اختيار الخوارزميات لأطر IDS بناءً على FAR أمرًا بالغ الأهمية. قد توفر الأنظمة المختلطة التي تجمع بين التعلم الآلي التقليدي والحديث توازنًا مثاليًا بين FAR منخفضة ومعدلات إيجابية حقيقية مرتفعة، مما يعزز استراتيجيات الأمن السيبراني.

1-4 مقاييس كفاءة وقت المعالجة

علاوة على ذلك، يمكن أن يؤدي تعزيز الخوارزميات عبر تعديل المُعطيات إلى تحسين كبير في كفاءة المعالجة. يمكن أن يخفف استخدام موارد الحوسبة المتوازية من أوقات الكشف عن طريق توزيع المهام عبر العديد من الأجهزة، كما هو موضح في [4]. من الضروري موازنة

الدقة والسرعة؛ في حين تعتبر الدقة أمراً مهماً، فإن القدرة على اكتشاف التسلات بسرعة غالباً ما تحدد الفعالية العامة لـ IDS.

مع تطور التهديدات السيبرانية، لا يزال تحسين كفاءة وقت المعالجة أمراً مهماً، مما يدفع الباحثين للبحث عن طرق جديدة لتحسين السرعة والموثوقية في تقنيات الكشف ضمن أطر IDS.

5- النتائج والاكتشافات من الدراسة المقارنة

5-1 ملخص الأداء عبر الخوارزميات

يكشف التقييم المقارن للخوارزميات المصممة لأنظمة كشف التسلل (IDS) عن مجموعة من النتائج بالأداء، مما يبرز المزايا والقيود الموجودة في استراتيجيات التعلم الآلي المختلفة. على سبيل المثال، حقق Al-Yaseen وزملاؤه معدل دقة رائع قدره 95.75% مع جمعهم بين الآلات التعلمية المتطرفة (ELM) و SVM على مجموعة بيانات KDD'99، مما أدى إلى تقليل التحذيرات الخاطئة إلى 1.87% فقط ([3] ص 1-5). في المقابل، قدم Hammad وزملاؤه T-SNERF، وهي خوارزمية جديدة حققت دقة 100% مع عدم وجود إنذارات كاذبة على مجموعة بيانات UNSW-NB15، بينما حافظت أيضاً على دقة عالية بلغت نحو 99.8% عبر مجموعات بيانات أخرى ([7]).

علاوة على ذلك، أظهرت النماذج المهجنة التي تجمع بين الأشجار القرار والشبكات العصبية مستويات دقة ممتازة؛ حقق أحد النماذج التي تدمج C4.5 و Multilayer Perceptron دقة كشف تصل إلى 99% ([11]). ضمن مجال التعلم العميق، أثبتت التكوينات المهجنة التي تدمج الشبكات العصبية المتكررة (RNNs) مع هياكل الذاكرة قصيرة وطويلة الأجل (LSTM) فعاليتها أيضاً، حيث حققت دقة تصل إلى حوالي 99.13% عند تطبيقها على مجموعة بيانات CICIDS2017 [2].

سويًا، تؤكد هذه الاستنتاجات أنه بينما تظل أساليب التعلم الآلي التقليدية مجدية، فإن نماذج التعلم العميق غالباً ما تعرض مقاييس أداء محسنة في تطبيقات IDS.

في الختام، يعد اختيار الخوارزميات حيويًا في التأثير على معدلات الاكتشاف وتكرار الإنذارات الكاذبة ضمن أطر IDS؛ إلا أن، تشير الابتكارات المتواصلة في كل من المنهجيات

التقليدية والتعلم العميق إلى تحسينات مستمرة في الكفاءة التشغيلية والدقة بينما نواجه التهديدات السيبرانية المتزايدة التعقيد.

6-التداعيات على البحث والتطوير المستقبلي

مع تطور الذكاء الاصطناعي في ميدان الأمن السيبراني، تبرز بعض العواقب للاستكشاف المستقبلي. يكمن التركيز الجوهرى في ضرورة قواعد بيانات مرجعية موحدة لتسهيل المقارنة بين الأبحاث. في الوقت الراهن، تتسبب قواعد البيانات المتنوعة في عراقيل عند محاولة الوصول إلى استنتاجات عامة أو تطوير نماذج عملية على نطاق واسع ([4]). يتوجب على الباحثين أن يعطوا الأولوية لقواعد البيانات الشاملة التي تعكس الأوضاع الحقيقية، خصوصاً فيما يخص التهديدات المستجدة مثل الثغرات من النوع صفر-يوم.

علاوةً على ذلك، تستدعي التحديات التي تواجهها أنظمة كشف التسلل (IDS) الراهنة اهتماماً. تواجه النماذج الحالية إشكاليات في معدلات الإنذارات الكاذبة المرتفعة والمتطلبات الحوسبية الكبيرة ([13]). بوسع البحوث المستقبلية أن تستكشف الأطر المختلطة لأنظمة كشف التسلل التي تجمع بين خوارزميات التعلم الآلي المتنوعة لتعزيز دقة الكشف مع تقليل استهلاك الموارد، حيث تشير التجارب الأخيرة إلى أن الأنظمة المختلطة تتفوق على المقاربات القائمة على خوارزمية واحدة.

اتجاه آخر هام هو التركيز المتزايد بأساليب التعليم العدوانية. إن ازدياد صعوبة المخاطر السيبرانية يتطلب وسائل دفاعية قوية ضد الهجمات على نماذج التعلم الآلي ([17]). يجب أن يتعمق البحث في آليات الإعتداء وتطوير استراتيجيات مضادة لتعزيز صلابة النموذج.

أخيراً، مع تقدم تقنيات التعلم العميق، هناك حاجة ملحة لإيجاد أساليب تدريب أكثر كفاءة. إن تحسين المعطيات الفائقة واستعمال الخوارزميات التكيفية يمكن أن يحسن بشكل كبير كفاءة الحلول الأمنية المدعومة بالذكاء الاصطناعي ([11]).

Attack Type	Data Size	Note Type
Training	833	normal
Validation	4373	normal
Hydra-FTP	162	attack
Hydra-SSH	148	attack
Adduser	91	attack
Java-Meterpreter	75	attack
Webshell	118	attack

الجدول 5: صنف الهجوم في ADFA-LD (المصدر: [4])

7-التوصيات لتنفيذ الذكاء الاصطناعي في حلول الأمن السيبراني

عند اختيار الخوارزميات للأمن السيبراني، وبالأخص لأنظمة كشف التسلل (IDS)، من الضروري اتباع أفضل الممارسات لضمان الأداء الأمثل. فهم متطلبات التطبيق المحددة أمر أساسي، بما في ذلك تحديد أنواع الهجمات وتقييم سمات البيانات ([1]). تساعد الأهداف الواضحة في توجيه اختيار الخوارزميات ومعايير التقييم.

تؤثر طبيعة مجموعة البيانات بشكل كبير على هذه العملية؛ عوامل مثل الحجم والجودة والتنوع تعتبر حيوية. يمكن أن تؤدي مجموعات البيانات غير المتوازنة إلى نتائج مشوهة وتوقع أداء النموذج ([4])، بينما تحسن مجموعات البيانات التدريبية المتنوعة من قوة النموذج ضد الهجمات المختلفة.

من الأساسي أيضاً تقدير النماذج الراهنة ذات الصلة بالحالات المشابهة. توحى الدراسات التاريخية بأن بعض الخوارزميات تعطي نتائج أفضل في أوضاع معينة ([12] ص 16-20). في الغالب، تعطي النماذج المختلطة نتائج أحسن بالمقارنة مع الطرق المنفردة نظراً لقدرتها على الاستفادة من النقاط القوية المتعددة ([4]).

بالإضافة إلى ذلك، من اللازم تقييم الموارد اللازمة لتدريب ومراقبة وصيانة هذه النماذج. قد تكون الخوارزميات التي تحتاج إلى تعديل مكثف أو قوة حسابية ضخمة غير مجدية للاستعمال في الوقت الفعلي ([3] ص 1-5). في الغالب، تحقق النماذج ذات التعقيد التشغيلي القليل توازناً بين الأداء وكفاءة الموارد.

وختامًا، فإن التقدير المستمر بعد الإطلاق ضروري، حيث أن مخاطر الأمن السيبراني تتطور، مما يتطلب إعادة تدريب وتعديل الخوارزميات بانتظام.

7-1 تطوير الأنظمة الهجينة التي تجمع بين عدة نهج

شهدت الأنظمة الهجينة في الأمن السيبراني اهتمامًا متصاعدًا نظرًا للطبيعة المعقدة لتهديدات الإنترنت. عبر الاستفادة من عدة خوارزميات، تعزز هذه النماذج قدرات الكشف وتقلل من الإيجابيات الخاطئة. على سبيل المثال، ابتكر جولا ورفاقه إطار عمل ذو مرحلتين يجمع بين خوارزمية مجموعة نحل العسل الاصطناعية (G-ABC) والشبكة العصبية العميقة (DNN)، محققًا معدل كشف بنسبة 96% في بيانات السحابة ([2]).

تدمج تقنيات التجميع المبتكرة نماذج مثل الشبكات العصبية التلافيفية (CNNs) والشبكات العصبية المتكررة (RNNs) لالتقاط أنماط الهجوم المتنوعة، مما يعزز التكيف والدقة ([9]). علاوة على ذلك، أظهر الكهتاني والدهاني فعالية دمج CNNs مع شبكات الذاكرة الطويلة القصيرة (LSTMs) في شبكات الإنترنت للأشياء، باستعمال تحسين سرب الجسيمات لاستخلاص الخصائص، مما أدى إلى دقة رائعة ([2]).

بالإضافة إلى ذلك، من الممكن أن التعاون بين نماذج التعلم الآلي التقليدية مثل أشجار القرار والأساليب المتطورة للتعلم العميق يعزز الأداء بشكل كبير. على سبيل المثال، ابتكر جوتشل ورفاقه خوارزمية مهجنة تجمع بين نايف بايز و SVM وأشجار القرار، وحقت معدلات كشف مرتفعة مع إدارة للإيجابيات الخاطئة في مجموعات بيانات مثل [7] (KDD99). تعمل هذه الأنظمة المهجنة على تبسيط عمليات الكشف وتعزيز المتانة ضد المتجهات الهجومية الحديثة.

8- الاستنتاج والتوجهات المستقبلية

إن دمج الذكاء الاصطناعي في ميدان الأمن السيبراني من المتوقع أن يبدل بصورة كبيرة حركية كشف التهديدات والاستجابة لها. مع تصاعد تعقيد وسرعة الهجمات السيبرانية، سيلعب استعمال تقنيات الذكاء الاصطناعي المتطورة—خاصة التعلم العميق—دورًا جوهريًا في بناء دفاعات أقوى. يتوجب أن تركز المبادرات المستقبلية على تطوير قابلية تفسير نماذج الذكاء الاصطناعي، حيث يمثل هذا عائقًا أساسيًا لبناء الثقة بين خبراء الأمن السيبراني والجهات

المساهمة المعنية. تطور الهجمات العدوانية المستمر يستدعي إستراتيجية مزدوجة: تعزيز متانة النماذج ضد هذه التهديدات مع تحسين الوضوح في النظام.

علاوة على ذلك، يُعتبر التعاون بين المؤسسات التعليمية ورواد القطاع ضروريًا لتحفيز الإبداع في طرائق الخوارزميات التي تتأقلم بسرعة مع المخاطر الجديدة. إن تنمية مجموعات البيانات التي تعكس السيناريوهات الحقيقية ستوفر لهذه الأنظمة تدريبًا أشمل، مما يعزز دقتها في تحديد أنماط الهجوم المعقدة (كما ورد في [10]). يجب أن تركز الدراسات الإضافية على النماذج المختلطة التي تجمع بين التقنيات التقليدية ومستجدات الذكاء الاصطناعي، مما يمكن الشركات من الاستفادة من الاستراتيجيات المجربة بينما تتبنى التكنولوجيا المتطورة.

ختامًا، يُعد التعامل مع معضلات الخصوصية المرتبطة باستعمال البيانات في عمليات تدريب الذكاء الاصطناعي أمرًا بالغ الأهمية. ستساعد أطر حوكمة البيانات المحسنة على المحافظة على المعايير الأخلاقية مع ضمان توافق تطبيقات الذكاء الاصطناعي مع اللوائح. ستمهد هذه الإستراتيجية الشاملة السبيل لمستقبل حلول الأمن السيبراني، مما يؤدي إلى بنى تحتية مرنة وقوية قادرة على التعامل مع مشهد التهديدات المتغير باستمرار.

المصادر:

- [1] "Machine learning (ML) in cybersecurity". Sep 2023. [Online]. Available: <https://www.sailpoint.com/identity-library/how-ai-and-machine-learning-are-improving-cybersecurity>
- [2] Patrick K. Gikunda, R. Kimanzi, P. Kimanga and D. Cherori. "Deep Learning Algorithms Used in Intrusion Detection Systems - A Review". Feb 2024. [Online]. Available: <https://arxiv.org/html/2402.17020v1>
- [3] N. Veselinović, M. Milašinović, M. Jovanović, A. Aleksić and N. Biga. "Countering Cybersecurity Threats with AI". Mar 2022. [Online]. Available: https://ceur-ws.org/Vol-3529/short_8.pdf
- [4] Y. Xin, L. Kong, Z. Liu, Y. Chen, Y. Li and H. Zhu. "Machine Learning and Deep Learning Methods for Cybersecurity". (accessed Apr 06, 2025). [Online]. Available: <https://ieeexplore.ieee.org/document/8359287>
- [5] "What Are the Risks and Benefits of Artificial Intelligence (AI) in Cybersecurity?". Apr 2025. [Online]. Available: <https://www.paloaltonetworks.com/cyberpedia/ai-risks-and-benefits-in-cybersecurity>
- [6] U. Ahmed, M. Nazir, A. Sarwar, T. Ali, El-Hadi M. Aggoune, T. Shahzad and M. A. Khan. "Signature-based intrusion detection using machine learning and deep learning approaches empowered with fuzzy clustering". Jan 2025. [Online]. Available: <https://www.nature.com/articles/s41598-025-85866-7>
- [7] M. Hammad, N. Hewahi and W. Elmedany. "T-SNERF: A novel high accuracy machine learning approach for Intrusion Detection Systems". Mar 2021. [Online]. Available: <https://ietresearch.onlinelibrary.wiley.com/doi/10.1049/ise2.12020>

- [8] Sowmya T. and Mary Anita E.A.. "A comprehensive review of AI based intrusion detection system". Jan 2023. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S2665917423001630>
- [9] R. Chinnasamy, M. Subramanian, S. V. Easwaramoorthy and J. Cho. "Deep learning-driven methods for network-based intrusion detection systems: A systematic review". Jan 2025. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S2405959525000050>
- [10] J. Khan, R. Elfakharany, H. Saleem, M. Pathan, E. Shahzad, S. Dhou and F. Aloul. "Can Machine Learning Enhance Intrusion Detection to Safeguard Smart City Networks from Multi-Step Cyberattacks?". Jan 2025. [Online]. Available: <https://www.mdpi.com/2624-6511/8/1/13>
- [11] Z. Azam, Md. Motaharul Islam and M. N. Huda. "Comparative Analysis of Intrusion Detection Systems and Machine Learning-Based Model Analysis Through Decision Tree". (accessed Apr 06, 2025). [Online]. Available: <https://ieeexplore.ieee.org/document/10185955/>
- [12] G. Apruzzese, L. Ferretti, M. Marchetti, M. Colajanni and A. Guido. "On the Effectiveness of Machine and Deep Learning for Cyber Security". May 2018. [Online]. Available: <https://www.ccdcoe.org/uploads/2018/10/Art-19-On-the-Effectiveness-of-Machine-and-Deep-Learning-for-Cyber-Security.pdf>
- [13] S. A. R. Shah and B. Issac. "Performance comparison of intrusion detection systems and application of machine learning to Snort system". Jan 2018. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0167739X17323178>
- [14] S. A. R. Shah and B. Issac. "Performance Comparison of Intrusion Detection Systems and Application of Machine Learning to Snort System". Oct 2017. [Online]. Available: <https://arxiv.org/abs/1710.04843>
- [15] IMI. "Artificial Intelligence Cyber Intrusions". Jul 2023. [Online]. Available: <https://identitymanagementinstitute.org/artificial-intelligence-cyber-intrusions/>
- [16] A. Khraisat, I. Gondal, P. Vamplew and J. Kamruzzaman. "Survey of intrusion detection systems: techniques, datasets and challenges". Jul 2019. [Online]. Available: <https://cybersecurity.springeropen.com/articles/10.1186/s42400-019-0038-7>
- [17] M. Macas, C. Wu and W. Fuertes. "A survey on deep learning for cybersecurity: Progress, challenges, and opportunities". Jul 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/abs/pii/S1389128622001864>
- [18] C. Goodman. "AI in Cybersecurity: Transforming Threat Detection and Prevention". Jan 2025. [Online]. Available: <https://www.balbix.com/insights/artificial-intelligence-in-cybersecurity/>
- [19] N. Thapa, Z. Liu, Dukka B. KC, B. Gokaraju and K. Roy. "Comparison of Machine Learning and Deep Learning Models for Network Intrusion Detection Systems". Sep 2020. [Online]. Available: <https://www.mdpi.com/1999-5903/12/10/167>
- [20] "What is Deep Learning? | Deep Instinct". Jun 2022. [Online]. Available: <https://www.deepinstinct.com/glossary/deep-learning>
- [21] Traci Gusher , "EY Americas AI and Data", Jan 2024 . Available: https://www.ey.com/en_us/insights/cybersecurity/ai-and-ml-are-cybersecurity-problems-and-solutions