



Malware Analysis Tool

إعداد الطلاب

سام المقدام مقبل محمد العلوي
عواد عبده قايد مصلح الخوبري
ياسين علي حسن عبد الرحمن عفيف
عبد الرحمن حسن طه ناجي
محمد عبد الله علي محمد الشيبة

تم تقديم هذا المشروع كمتطلب جزئي لاستيفاء متطلبات درجة البكالوريوس في قسم الامن السيبراني

إشراف/

الدكتور: عبد السلام خاقو
الدكتور: وليد الوجية

الآيات القرآنية

بسم الله الرحمن الرحيم

وقال تعالى ﴿وَقُلْ رَبِّ زِدْنِي عِلْمًا﴾

سورة طه (114)

وقال تعالى ﴿وَقُلْ اَعْمَلُوا فَسَيَرَى اللَّهُ عَمَلَكُمْ وَرَسُولُهُ وَالْمُؤْمِنُونَ وَسَتُرَدُّونَ إِلَىٰ عَالَمِ الْغَيْبِ وَالشَّهَادَةِ فَيُنَبِّئُكُم بِمَا كُنتُمْ تَعْمَلُونَ﴾

سورة التوبة (105)

وقال تعالى ﴿قَالُوا سُبْحَانَكَ لَا عِلْمَ لَنَا إِلَّا مَا عَلَّمْتَنَا إِنَّكَ أَنْتَ الْعَلِيمُ الْحَكِيمُ﴾

سورة البقرة (32)

وقال تعالى ﴿وَإِذْ يَرْفَعُ إِبْرَاهِيمُ الْقَوَاعِدَ مِنَ الْبَيْتِ وَإِسْمَاعِيلُ رَبَّنَا تَقَبَّلْ مِنَّا إِنَّكَ أَنْتَ السَّمِيعُ الْعَلِيمُ﴾

سورة البقرة (127)

الإهداء

إلى من هو أحب إلينا من أنفسنا، سيّد الخلق وإمام المرسلين، محمد صلى الله عليه وسلم، الذي علّم الدنيا معنى الهداية والنور.

إلى من زرعوا فينا حب العلم والمعرفة؛ وسهلوا لنا دروب الابتكار...

إلى أهلنا الكرام صمام الأمن وركيزة العطاء، نهدي هذا التطبيق الذي كان حصاد سنوات من دعمكم وصبركم. فأنتم من علمتمونا ان كل عظيم يبدأ بخطوة، وكل ثروة تبدأ بفكرة

أمهاتنا الغاليات، من تعلّمنا في مدارس قلوبهن أبجدية الوفاء، وارتشفنا من تضحياتهن مناهج الصبر والعطاء.

إلى فريقنا الرائع ورفاق الدرب الذين تحملوا معنا سهر الليالي وتحديات، نهدي هذا التطبيق الذي كان ثمرة عمل جماعي بحث. فأنتم من جعلتم العقبات تبدو، كفرص والمستحيل مجرد تحدٍ مؤقت

ونختتم بآبائنا الكرام، الذين قالوا لنا يوماً: "لا تكونوا مثلاً لنفخر بكم، فاخترنا أن نكون مثلهم لنفخر بهم وبأنفسنا.

الشكر والتقدير

بسم الله الرحمن الرحيم

والحمد لله رب العالمين الذي ساق لنا التيسير حين تعثر الطريق، وألقى في خطانا نوراً كلما أظلمت المسافات.

بالأمس...

كنّا نخطو على استحياء، نلامس الطريق بقلوب يملؤها التردد، كأننا نبحث عن أنفسنا بين احتمالات كثيرة، حتى صار هذا التخصص حلمًا يُشبهنا وغايةً تمسك بنا قبل أن نمسك بها.

ومن باب الوفاء

...لمن كانوا بعد الله أسباباً في بلوغنا، نرفع امتناننا للدكتور المشرف **وليد الوجية**، الذي لم يكن إشرافه توجيهاً فحسب، بل كان نوراً يمتدّ في تفاصيل الطريق، وسنداً نستند إليه كلما مالت بنا الأيام.

كما نُجَلّ بالشكر

...للأستاذ والمشرّف المساعد **وليد الوجية**، الذي كان معنا منذ أول خطوة، فكان القرب الذي يبّد ارتباكنا والثقة التي تُعيد ترتيبنا من الداخل.

وإلى عائلاتنا

...يا من كنتم الدعاء الذي يسبقنا إلى غدنا، والصبر الذي يحتملنا حتى في أشدّ حالاتنا تعباً، كنتم لنا قلباً إذا ضاقت الدنيا اتّسع، وظهراً لا ينكسر مهما اشتدّ الحمل.

وإلى الأهل والأصدقاء والأحباب وزملاء الرحلة...

أنتم الذين جعلتم الطريق أقلّ وحشة وأكثر حياة، فكان حضوركم تفصيلاً لا يُنسى من هذه الحكاية.

وإلى أساتذتنا في كلية الحاسوب...

الذين غرسوا فينا علماً يمتدّ أثره نقول: ما أخذناه منكم سيبقى ما بقينا.

فلكم جميعاً

...امتنانٌ يفيض حتى يعجز عنه القول، ودعاءً صادق أن يجزيكم الله عنا خير الجزاء.

ملخص المشروع

انتقل المشهد السيبراني المعاصر من الهجمات الفردية إلى طور الجريمة المنظمة عابرة للحدود التي تديرها تحالفات تقنية كبرى تمتلك موارد تضاهي الدول مما أدى لظهور برمجيات معقدة متعددة الأشكال ومتحولة تعيد صياغة نفسها آلياً ومدعومة بذكاء اصطناعي توليدي يبتكر ثغرات آنية وهو ما كشف عجز المنظومات التقليدية عن مواجهة الانفجار المعلوماتي الهائل وسد فجوة هجمات اليوم الصفر ولأجل ذلك يهدف هذا المشروع إلى تطوير هندسة أمنية هجينة متكاملة تعتمد المنهجية التنبؤية الاستباقية عبر محاور تكاملية تبدأ بالتحليل السكوني المعمق لاستنتاج الخصائص الجينية للملفات دون الحاجة لتشغيلها مروراً بالتحليل الديناميكي السلوكي لمحاكاة التنفيذ داخل بيئات ساند بوكس عالية الدقة لرصد انحرافات الذاكرة والشبكة وصولاً إلى الطبقة التنبؤية الذكية التي تدمج خوارزميات التعلم العميق مثل الشبكات "Convolutional" والذاكرة طويلة قصيرة المدى القادرة على استيعاب الأنماط السلوكية المعقدة والتسلسلات الزمنية للنشاط الإجرامي مع تعزيز القيمة المضافة بدمج تقنيات الذكاء الاصطناعي القابل للتفسير لتفكيك منطق الصندوق الأسود وتوفير شروحات تقنية منطقية تدعم اتخاذ القرار إلى جانب التركيز المكثف على تعزيز المتانة العدائية لتحسين النظام ضد محاولات التلاعب والتهرب الذكي وضمان المرونة التامة للعمل العابر للمنصات في بيئات ويندوز ولينكس مما يثمر في النهاية عن تحول استراتيجي من الدفاع السلبي التقليدي إلى منطق الصيد الاستباقي الذكي الذي يسد الفجوات المعرفية والعملية ويؤسس لمعايير جديدة في حماية الأصول الرقمية السيادية.

Project Summary

The contemporary cyber landscape has shifted from individual attacks to the stage of cross-border organized crime managed by major technological cartels possessing resources rivaling nation-states, leading to the emergence of complex, polymorphic, and metamorphic malware that automatically rewrites itself, supported by generative artificial intelligence that invents instantaneous exploits. This has exposed the incapacity of traditional systems to cope with the massive information explosion and to close the zero-day attack gap. To address this, the project aims to develop an integrated hybrid security architecture based on a proactive predictive methodology through integrated pillars: starting with deep static analysis to interrogate the genetic characteristics of files without execution, moving through dynamic behavioral analysis simulating execution within high-fidelity sandbox environments to monitor memory and network deviations, and reaching the intelligent predictive layer that integrates deep learning algorithms such as Convolutional Neural Networks and Long Short-Term Memory capable of absorbing complex behavioral patterns and temporal sequences of criminal activity. The value is enhanced by integrating Explainable AI techniques to dismantle black-box logic and provide logical technical explanations that support decision-makers, along with an intensive focus on strengthening adversarial robustness to fortify the system against manipulation and intelligent evasion attempts, while ensuring full cross-platform flexibility in Windows and Linux environments. This ultimately yields a strategic transformation from traditional passive defense to proactive intelligent hunting that bridges cognitive and operational gaps and establishes new standards for protecting sovereign digital assets.

جدول المحتويات

14	1.1 المقدمة:
15	2.1 مشكلة البحث:
16	3.1 أهداف المشروع:
17	4.1 مساهمة المشروع:
17	5.1 حدود المشروع:
19	6.1 خلاصة الفصل الأول:
21	1.2 مقدمة:
21	2.2 أنواع الفيروسات:
22	1.3.2 كيفية عمل الفيروسات:
23	2.3.2 ما يميز الفيروسات عن الأنماط الأخرى:
23	3.3.2 أسباب خطورة الفيروسات:
24	4.2 الأدبيات والدراسات السابقة:
25	1.4.2 الطرق التقليدية في تحليل البرمجيات الخبيثة:
25	1.1.4.2 التحليل القائم على التوقيعات:
25	2.1.4.2 التحليل القائم على القواعد:
25	3.1.4.2 التحليل الساكن التقليدي:
25	4.1.4.2 التحليل الديناميكي التقليدي:
26	5.1.4.2 القيود العامة للطرق التقليدية:
26	2.4.2 تحليل نقدي للدراسات السابقة:
26	3.4.2 الفجوة البحثية (Research Gap):
27	4.4.2 ربط الدراسات بالمشروع:
29	6.2 خلاصة الفصل الثاني:
32	1.2.3 إدارة المستخدمين والتحقق الأمني:
32	2.2.3 رفع الملفات ومعالجتها الأولية:
32	3.2.3 الحماية والعزل (File Isolation & Security Layer):
33	4.2.3 التحليل الساكن (Static Analysis):
33	5.2.3 التحليل الديناميكي (Dynamic Analysis):
33	6.2.3 إدارة النتائج والتقارير:
34	7.2.3 إدارة النظام والصلاحيات (Admin Panel):
34	3.3 دراسة الجدوى:
34	1.3.3 الجدوى التقنية (Technical Feasibility):
35	2.3.3 الجدوى التشغيلية (Operational Feasibility):
36	3.3.3 الجدوى الاقتصادية (Economic Feasibility):
37	3.3.4 الجدوى الزمنية (Time Feasibility):

37	4.3 خلاصة دراسة الجدوى:
38	5.3 أدوات وتقنيات المشروع:
43	1.1 المقدمة والأسس المعرفية (Epistemological Foundations)
43	2.4 هندسة الأنظمة والبيئات المعزولة (Sandboxing & System Architecture)
44	1.3.4 معضلة "الفجوة الدلالية" (The Semantic Gap Problem)
44	4.4 ميكانيكا الاعتراض والتحكم في التدفق (Flow Control & Interception)
44	1.4.4 نظرية مستويات الحماية (Protection Rings Theory)
44	2.4.4 آليات الاعتراض المتقدمة (Advanced Hooking Mechanisms)
45	5.4 تحليل الذاكرة العشوائية والبيانات المتطايرة
45	1.5.4 هياكل البيانات الحيوية (Vital OS Structures)
45	2.5.4 كشف تقنيات الحقن (Process Injection Detection)
45	6.4 نظرية المعلومات وتحليل الاتصالات (C2 Analysis & Information Theory)
46	1.6.4 تحليل الانتروبيا (Entropy Analysis)
46	2.6.4 إحصائيات إشارات التنبيه (Beaconing Statistics)
46	7.4 نظريات التهرب ومكافحة التحليل (Anti-Analysis Theory & Evasion)
46	1.7.4 نظرية اكتشاف البيئة الافتراضية (Virtualization Detection Theory)
47	2.7.4 تزييف البيئة وتضليل المحلل (Counter-Intelligence & Environment Faking)
47	8.4 التحليل الشكلي والتنفيذ الرمزي (Symbolic Execution & Formal Methods)
47	1.8.4 نظرية استكشاف المسارات (Path Exploration Theory)
47	2.8.4 معالجة الانفجار المزيج (State Explosion Problem)
47	9.4 التحليل الشبكي وفك تشفير البروتوكولات
47	1.9.4 القنوات السرية وتحليل البروتوكولات (Covert Channels)
48	2.9.4 فك تشفير SSL/TLS في بيئة التحليل (SSL Decryption Theory)
48	10.4 دراسات حالة كتشريح علمي (Scientific Post-Mortem Case Studies)
48	1.10.4 دراسة حالة: برمجية Stuxnet (The Cyber-Weapon)
48	2.10.4 دراسة حالة: برمجية Pegasus (Zero-Click Spyware)
49	3.10.4 التعلم التعزيزي في الدفاع السيبراني (Reinforcement Learning - RL)
49	11.4 التوجهات المستقبلية والبحث العلمي المتقدم
49	1.11.4 البرمجيات الخبيثة المقاومة للحوسبة الكمومية (Quantum-Resistant Malware)
49	2.11.4 التهرب المعتمد على الذكاء الاصطناعي (AI-Driven Evasion)
	1.5 التنفيذ 51
51	1.1.5 الرؤية الهندسية والبنية التحتية (System Architecture)
53	2.5 هندسة البيانات والتخزين (Data Engineering)
53	3.5 قاعدة البيانات (PostgreSQL Deep Dive)
54	4.5 إدارة الملفات المرفوعة وعزلها أمنياً (Secure File Handling & Vault)
54	1.4.5 إدارة الملفات المرفوعة (Spatie MediaLibrary & Security)

54		منطق التحليل الساكن (Static Analysis Logic)	5.5
55		استخراج الخصائص (Feature Extraction)	5.5.1
56		فحص التوقيعات (Signature Matching)	6.5
56		بيئة التحليل الديناميكي (Sandbox Engineering)	1.6.5
56		بناء الـ Sandbox (Isolated Virtualization)	2.6.5
57		مراقبة السلوك (Behavioral Logging)	3.6.5
58		محرك الذكاء الاصطناعي والـ XAI (AI & Explainable AI)	7.5
59		واجهة المساعد التفاعلي	1.7.5
60		تدريب النموذج (Model Training)	3.7.5
60		رحلة المستخدم والواجهات (User Experience & Frontend)	8.5
60		صفحة الرفع (Upload Experience)	1.8.5
61		Drag and Drop Zone	2.8.5
62		اعدادات التحكم بالرفع	3.8.5
62		صفحة التقرير (Interactive Reports)	4.8.5
63		سجلات التحليل	4.8.5
64		لوحة احصائيات المستخدمين	9.5.5
65		خاصية الدردشة المدمجة	8.5.5
66		استراتيجيات الأمان القصوى (Extreme Security Measures)	6.6
70		قابلية التوسع والصيانة (Scalability & Maintenance)	7.5
70		الحاويات (Dockerization)	1.7.5
71		الخاتمة وخارطة الطريق المستقبلية	8.5
73		مجموعات البيانات المستخدمة	1.6
74		المقاييس الكمية المحققة	2.6
74		مقاييس الأداء ونتائج التقييم	2.2.6
75		أداء الكشف ضد البرمجيات متعددة الأشكال والمتحولة	3.6
75		أداء النظام ضد البرمجيات الخبيثة المتقدمة	1.3.6
75		كشف هجمات اليوم الصفرى (Zero-day Detection)	4.6
75		المتانة العدائية (Adversarial Robustness)	5.6
76		مقارنة مع الأنظمة الأخرى	6.6
77		المعايير المستخدمة في التقييم	1.6.6
77		تحليل أداء الأنظمة	2.6.6
77		الاستنتاجات الرئيسية (Key Takeaways)	3.6.6
78		مناقشة النتائج (Discussion)	7.6
78		تفسير الأداء العام	1.7.6
78		سباب السلبيات الكاذبة: (False Negatives)	2.7.6
80		أسباب الإيجابيات الكاذبة: (False Positives)	3.7.6

8.6	مقارنة مع الدراسات السابقة.....	81
1.8.6	الأعمال المستقبلية (خريطة طريق).....	83
9.6	الخاتمة والاستنتاج.....	83
1.7	المراجع(References).....	85

قائمة الجداول

43	جدول 1-4 يوضح نظرية المحاكاة والافتراضية
73	جدول 1-6: يوضح مجموعة البيانات المستخدمة
74	جدول 2-6: يوضح المقاييس الكمية المحققة
75	جدول 3-6: يوضح ادا الكشف ضد البرمجيات متعددة الاشكال والمتحولة
75	جدول 4-6: يوضح نتائج كشف هجمات اليوم الصفري
76	جدول 5-6: ملخص نتائج المتانة العدائية قبل وبعد التدريب العدائي
76	جدول 6-6: مقارنة أداء النظام المقترح مع الأنظمة الأخرى
78	جدول 7-6: تحليل أسباب السلبات الكاذبة ونسبها
80	جدول 7-6: تحليل أسباب الإيجابيات الكاذبة (False Positives) والنسب المئوية لكل سبب
81	جدول 8-6: ملخص مقارن للنظام المقترح مقابل أحدث الدراسات في مجال كشف البرمجيات الخبيثة

قائمة الصور

51	الشكل 1-5: يوضح الصفحة الرئيسية لمنصة التحليل الهجين
52	الشكل 2-5: يوضح صفحة انشاء حساب جديد
53	الشكل 3-5: يوضح نظرة عامة على المنصة تشمل احصائيات حية (المستخدمين, التحليلات, التهديدات)
55	الشكل 4-5: يوضح مراحل التقدم في عملية التحليل متعدد الطبقات (Static+Dynamic+AI)
56	الشكل 5-5: يوضح تفاصيل التحليل السكوني العميق (PE Info, Entropy, FileTags)
57	الشكل 6-5: يوضح معلومات الاسناد والمصدر للملف المحلل (الرفع, IP, مستوى الخطورة)
58	الشكل 7-6: يوضح لوحة صحة نماذج الذكاء الاصطناعي المتعددة (F1-Score, حالة الاتصال)
59	الشكل 8-5: يوضح المساعد التفاعلي لتوليد تفسيرات حول سلاطات البرمجيات الخبيثة
60	الشكل 9-5: يوضح الصفحة الرئيسية لمنصة التحليل الهجين
61	الشكل 10-5: يوضح واجهة رفع الملفات الامنة مع قيود الأنواع والاحجام
62	الشكل 11-5: يوضح لوحة تحكم اعدادات الرفع (Upload Contrille لتقييد الأنواع والحجم
63	الشكل 12-5: يوضح صفحة تقرير تحليلي لملف تنفيذي مع نتيجة SAFE
63	الشكل 13-5: يوضح سجلات التحليل السابقة مع اسماء الملفات واوقات الرفع وحالة التفاصيل
64	الشكل 14-5: يوضح توزيع نشاط المستخدمين واكثرهم تحليلا للملفات
65	الشكل 15-5: يوضح واجهة الدردشة المدمجة
66	الشكل 16-5: يوضح قائمة المستخدمين المسجلين مع الأدوار والحالة واخر نشاط
67	الشكل 17-5: يوضح واجهة تسجيل الدخول الثنائية (2FA) كآلية للتحقق من الهوية
68	الشكل 18-5: يوضح نافذة ادخال رمز التحقق الثنائي (TOTP)
69	الشكل 19-5: يوضح نموذج إضافة مستخدم جديد يدويا مع تعيين دور الصلاحيات
70	الشكل 20-5: يوضح صفحة الاعدادات لتخصيص المظهر واللغة وإدارة الجلسة

قائمة المختصرات والمصطلحات التقنية المستخدمة

الاختصار	التعريف
AES	Advanced Encryption Standard
AI	Artificial Intelligence
APK	Android Package
API	Application Programming Interface
APT	Advanced Persistent Threat
C2	Command & Control
CFGs	Control Flow Graphs
CII	Critical Infrastructure
CNN	Convolutional Neural Network
CPU	Central Processing Unit
CSS	Cascading Style Sheets
DNS	Domain Name System
ELF	Executable and Linkable Format
EXE	Executable
GNN	Graph Neural Network
GUI	Graphical User Interface
HTML	Hypertext Markup Language
HTTP	Hypertext Transfer Protocol
HTTPS	HTTP Secure
IAT	Import Address Table
ICMP	Internet Control Message Protocol
IOC	Indicators of Compromise
IRT	Incident Response Team
JSON	JavaScript Object Notation

الاختصار	التعريف
JWT	JSON Web Token
KVM	Kernel-based Virtual Machine
LIME	Local Interpretable Model-agnostic Explanations
LSTM	Long Short-Term Memory
MAC	Media Access Control
MITM	Man-in-the-Middle
ML	Machine Learning
MTTD	Mean Time to Detect
MySQL	MySQL Database
OAuth	Open Authorization
OS	Operating System
PDF	Portable Document Format
PE	Portable Executable
PHP	PHP: Hypertext Preprocessor
PostgreSQL	PostgreSQL Database
QEMU	Quick Emulator
RAM	Random Access Memory
Redis	Remote Dictionary Server
RL	Reinforcement Learning
RNN	Recurrent Neural Network
SCADA	Supervisory Control and Data Acquisition
SHA-256	Secure Hash Algorithm 256-bit
SHAP	SHapley Additive exPlanations
SOC	Security Operations Center
SSL	Secure Sockets Layer
TLS	Transport Layer Security

الاختصار	التعريف
TLSH	Trend Micro Locality Sensitive Hash
UPX	Ultimate Packer for Executables
UUID	Universally Unique Identifier
VLAN	Virtual Local Area Network
VMI	Virtual Machine Introspection
XAI	Explainable Artificial Intelligence
YARA	Yet Another Recursive Acronym
ZIP	Zip Archive

1.1 المقدمة:

تفرض متطلبات الأمن السيبراني الحديث، خصوصًا في سياق السيادة الرقمية، ضرورة الانتقال من نماذج الكشف التقليدية القائمة على الاستجابة بعد وقوع الهجوم (Reactive Detection) إلى نماذج أكثر تقدمًا تعتمد على التوقع الاستباقي (Proactive Prediction)، وذلك لضمان استمرارية الصمود الرقمي في مواجهة التهديدات المتزايدة التي تستهدف البنى التحتية الحيوية (Critical Information Infrastructure - CII) والقطاعات الاقتصادية الحساسة.

لم تعد التهديدات السيبرانية مجرد تحديات تقنية محدودة، بل تطورت لتصبح هجمات معقدة ومتعددة الأبعاد، حيث تشير التقارير الحديثة إلى تزايد ملحوظ في استهداف الأنظمة الصناعية وسلاسل الإمداد، إلى جانب الانتشار الواسع لهجمات برامج الفدية (Ransomware) التي تعتمد على تشفير البيانات وتعطيل الأنظمة، مما يؤدي إلى خسائر اقتصادية عالمية ضخمة سنويًا.

في هذا السياق، برزت فئة متقدمة من التهديدات تُعرف بالتهديدات المستمرة المتقدمة (Advanced Persistent Threats - APTs)، والتي تعتمد على ما يمكن وصفه بالكود العدواني المتطور (Proliferating Adversarial Code)، حيث أصبحت البرمجيات الخبيثة أكثر قدرة على التخفي والتكيف، مستغلة ثغرات اليوم الصفر (Zero-day Vulnerabilities) لتجاوز أنظمة الحماية التقليدية.

كما تبنت البرمجيات الخبيثة الحديثة تقنيات متقدمة مثل تعدد الأشكال (Polymorphism) والتعمية (Obfuscation)، مما يمكنها من توليد نسخ متغيرة باستمرار، وبالتالي تجاوز آليات الكشف المعتمدة على التوقيعات الثابتة (Static Signatures) أو القواعد الاستدلالية (Heuristic-based Detection). ويؤدي ذلك إلى زيادة العبء على فرق الاستجابة للحوادث (SOC/IRT) نتيجة الكم الهائل من البيانات غير المصنفة وصعوبة تحليلها بكفاءة.

ونتيجة لهذا القصور في الأنظمة الدفاعية التقليدية، يرتفع متوسط زمن الكشف (Mean Time to Detect - MTTD) وتتسع نافذة الاستغلال (Window of Opportunity)، مما يمنح المهاجمين القدرة على تنفيذ التحركات الجانبية (Lateral Movement) داخل الشبكة وترسيخ وجودهم (Persistence) قبل تنفيذ الهجوم الفعلي.

بناءً على ذلك، تبرز الحاجة إلى تطوير أنظمة تحليل متقدمة تعتمد على الفهم السلوكي للبرمجيات الخبيثة، بدلاً من الاعتماد الحصري على التحليل الساكن. ومن هذا المنطلق، يهدف هذا المشروع إلى تصميم نموذج متكامل لتحليل البرمجيات الخبيثة يعتمد على منهجية هجينة (Hybrid Approach) تجمع بين:

- التحليل الساكن (Static Analysis) لاستخراج السمات البرمجية مثل Opcodes وهياكل تدفق التحكم (Control Flow Graphs)
- التحليل الديناميكي (Dynamic Analysis) لمراقبة سلوك البرنامج أثناء التنفيذ، بما في ذلك استدعاءات API والتغيرات في الذاكرة

وذلك ضمن إطار يعتمد على تقنيات الذكاء الاصطناعي والتعلم الآلي، بهدف تحسين دقة الكشف عن الأنماط الشاذة.

كما يسعى المشروع إلى توظيف الذكاء الاصطناعي القابل للتفسير (Explainable AI - XAI) لتمكين محلي الأمن من فهم قرارات النماذج بدلاً من التعامل معها كـ "صندوق أسود"، مع الأخذ بعين الاعتبار التحديات المرتبطة بالهجمات العدائية (Adversarial Attacks) التي تستهدف تضليل النماذج الذكية.

وبذلك، يهدف هذا النظام إلى إحداث نقلة نوعية في مجال الدفاع السيبراني من خلال التحول من أنظمة مراقبة تقليدية إلى منصات تحليل ذكية واستباقية تدعم اتخاذ القرار وتعزز السيادة الرقمية للمؤسسات.

2.1 مشكلة البحث:

على الرغم من التطور المستمر في تقنيات الأمن السيبراني، لا تزال العديد من الأنظمة الدفاعية التقليدية تعتمد بشكل أساسي على أساليب الكشف القائمة على التوقيعات (Signature-based Detection) والقواعد الاستدلالية (Heuristic Methods)، والتي أصبحت غير كافية لمواكبة التطور السريع في أساليب البرمجيات الخبيثة الحديثة. حيث تعتمد هذه البرمجيات على تقنيات متقدمة مثل التعمية (Obfuscation) وتعدد الأشكال (Polymorphism)، مما يمكنها من تغيير بنيتها بشكل مستمر وتجاوز أنظمة الكشف التقليدية.

تزداد حدة هذه المشكلة مع ظهور التهديدات المستمرة المتقدمة (Advanced Persistent Threats - APTs) التي تستخدم أساليب تخفي متقدمة واستغلال ثغرات اليوم الصفري (Zero-day Vulnerabilities)، مما يسمح لها بالبقاء داخل الأنظمة لفترات طويلة دون اكتشاف، خاصة في البيئات الحساسة مثل الأنظمة الصناعية والبنى التحتية الحيوية.

كما تواجه مراكز عمليات الأمن (SOC) وفرق الاستجابة للحوادث (IRT) تحديات كبيرة نتيجة الاعتماد على أدوات تحليل منفصلة لكل من التحليل الساكن (Static Analysis) والتحليل الديناميكي (Dynamic Analysis)، الأمر الذي يؤدي إلى نقص في التكامل بين نتائج التحليل وصعوبة الحصول على رؤية شاملة ودقيقة لسلوك البرمجيات الخبيثة. بالإضافة إلى ذلك، يؤدي هذا الفصل بين نوعي التحليل إلى زيادة زمن الكشف (MTTD) وتأخر الاستجابة للحوادث.

بناءً على ما سبق، تتمثل مشكلة البحث في القصور الواضح في الأنظمة الحالية لتحليل البرمجيات الخبيثة، والتي تفتقر إلى منهجية متكاملة تجمع بين التحليل الساكن والديناميكي ضمن إطار موحد، مما يؤثر سلباً على دقة وفعالية الكشف عن التهديدات الحديثة.

ومن هنا، تبرز الحاجة إلى تطوير نظام متكامل لتحليل البرمجيات الخبيثة يعتمد على دمج تقنيات التحليل الساكن والديناميكي، بهدف تحسين كفاءة الكشف، وتوفير رؤية تحليلية شاملة لسلوك البرمجيات، وتقليل زمن الاستجابة للتهديدات السيبرانية.

3.1 أهداف المشروع:

يهدف هذا المشروع إلى تصميم وتطوير نظام متكامل لتحليل البرمجيات الخبيثة يعتمد على دمج تقنيات التحليل الساكن (Static Analysis) والتحليل الديناميكي (Dynamic Analysis)، وذلك لتحسين كفاءة اكتشاف التهديدات السيبرانية الحديثة وتوفير فهم أعمق لسلوك البرمجيات المشبوهة.

وتتمثل الأهداف الرئيسية للمشروع فيما يلي:

1. تطوير نظام قادر على إجراء تحليل ساكن للبرمجيات الخبيثة من خلال استخراج السمات الأساسية مثل التوقيعات الرقمية (Hashes)، النصوص (Strings)، وبنية الملف التنفيذي، بهدف التعرف على المؤشرات الأولية للتهديد.
2. بناء بيئة آمنة لتنفيذ التحليل الديناميكي تعتمد على تشغيل الملفات المشبوهة داخل بيئة معزولة (Sandbox)، مع مراقبة سلوك البرنامج أثناء التنفيذ، بما في ذلك العمليات (Processes)، واستدعاءات API، والتغيرات على النظام.
3. دمج نتائج التحليل الساكن والديناميكي ضمن نظام موحد لتوفير رؤية شاملة ومتكاملة حول خصائص وسلوك البرمجيات الخبيثة.
4. تقليل زمن الكشف عن التهديدات (Detection Time) من خلال تسريع عمليات التحليل وتقديم نتائج دقيقة تساعد في اتخاذ القرار الأمني بشكل أسرع.
5. تصميم واجهة استخدام تفاعلية (User Interface) تتيح للمستخدم رفع الملفات المشبوهة وعرض نتائج التحليل بطريقة واضحة ومنظمة.
6. إنشاء تقارير تحليلية مفصلة توضح خصائص الملف المشبوه وسلوكه، مما يساعد محللي الأمن في فهم طبيعة التهديد واتخاذ الإجراءات المناسبة.
7. تعزيز مستوى الأمان من خلال توفير بيئة تحليل تمنع انتشار البرمجيات الخبيثة أو تأثيرها على الأنظمة الحقيقية.

وبذلك يسعى المشروع إلى تقديم نموذج عملي يساهم في تحسين آليات تحليل البرمجيات الخبيثة، من خلال الجمع بين الدقة في التحليل الساكن والعمق في التحليل الديناميكي ضمن إطار متكامل.

4.1 مساهمة المشروع:

يساهم هذا المشروع في تطوير نموذج تطبيقي متكامل لتحليل البرمجيات الخبيثة من خلال دمج تقنيات التحليل الساكن (Static Analysis) والتحليل الديناميكي (Dynamic Analysis) ضمن بيئة موحدة، بما يعالج القصور الموجود في الأنظمة التقليدية التي تعتمد على كل نوع تحليل بشكل منفصل.

وتبرز مساهمة المشروع في النقاط التالية:

- تقديم إطار عملي يجمع بين التحليل الساكن والديناميكي، مما يتيح الحصول على رؤية شاملة لسلوك البرمجيات الخبيثة بدلاً من الاعتماد على تحليل جزئي محدود.
- تحسين دقة اكتشاف التهديدات من خلال الربط بين الخصائص البنيوية للملف (Static Features) والسلوك الفعلي أثناء التنفيذ (Behavioral Analysis).
- تقليل زمن التحليل والكشف من خلال أتمتة العمليات الأساسية، مما يساعد في تسريع الاستجابة للحوادث الأمنية.
- توفير بيئة تحليل آمنة (Sandbox) تُمكن من اختبار الملفات المشبوهة دون التأثير على الأنظمة الحقيقية.
- دعم محلي الأمن بمخرجات واضحة ومنظمة على شكل تقارير تحليلية تسهل فهم طبيعة التهديد واتخاذ القرار المناسب.
- بناء نظام قابل للتطوير يمكن توسيعه مستقبلاً بإضافة تقنيات تحليل متقدمة أو دمج مع أنظمة أمنية أخرى.

وبذلك يقدم المشروع إضافة عملية في مجال تحليل البرمجيات الخبيثة، من خلال تحسين كفاءة التحليل ورفع مستوى الفهم للسلوكيات الضارة باستخدام منهجية تكاملية بين التحليل الساكن والديناميكي.

5.1 حدود المشروع:

على الرغم من أهمية المشروع في تحسين تحليل البرمجيات الخبيثة من خلال الدمج بين التحليل الساكن (Static Analysis) والتحليل الديناميكي (Dynamic Analysis)، إلا أن هناك مجموعة من الحدود والقيود التي قد تؤثر على نطاق التطبيق ونتائجه، ومن أبرزها:

- يدعم النظام تحليل أنواع متعددة من الملفات (مثل الملفات التنفيذية والملفات القابلة للفحص)، إلا أن كفاءة التحليل قد تختلف حسب نوع الملف وتعقيده.
- يعتمد التحليل الديناميكي على بيئة معزولة (Sandbox)، حيث يقوم النظام بمراقبة سلوك البرنامج أثناء التنفيذ بشكل شامل، بما في ذلك العمليات (Processes)، واستدعاءات API، والتغيرات على النظام، وحركة الشبكة.
- رغم قدرة النظام على اكتشاف العديد من الأنشطة المشبوهة مثل:

- محاولات الاتصال بشبكات خارجية

- التعديل على ملفات النظام

- إنشاء عمليات غير طبيعية

- فتح منافذ أو اتصالات غير مبررة

إلا أنه لا يمكن ضمان اكتشاف جميع التهديدات، خصوصاً تلك التي تستخدم تقنيات متقدمة مثل:

- Anti-VM

- التشفير الديناميكي

- أو التأخير الزمني في تنفيذ الهجوم

- دقة نتائج التحليل تعتمد على مدة تشغيل العينة داخل بيئة التحليل، حيث قد لا تظهر بعض السلوكيات الخبيثة إلا بعد فترة زمنية طويلة.

- قد يواجه النظام قيوداً في الأداء عند تحليل عدد كبير من الملفات أو عند تشغيل التحليل الديناميكي الذي يستهلك موارد عالية.

- يركز المشروع على تحليل السلوك والكشف عن الأنشطة المشبوهة مثل الثغرات المحتملة أو الاتصالات غير الآمنة، دون التوسع في تقديم حلول إصلاح أو استجابة تلقائية.

وبناءً على هذه الحدود، يُعتبر المشروع نموذجاً تطبيقياً متقدماً يهدف إلى توفير تحليل شامل قدر الإمكان، مع إمكانية تطويره مستقبلاً ليشمل قدرات أوسع وأكثر دقة في مواجهة التهديدات السيبرانية.

6.1 خلاصة الفصل الأول:

تناول هذا الفصل الإطار العام للمشروع من خلال عرض المفاهيم الأساسية المرتبطة بتحليل البرمجيات الخبيثة، وأهمية الانتقال من أساليب الكشف التقليدية إلى أساليب تحليل أكثر تقدمًا تعتمد على الدمج بين التحليل الساكن (Static Analysis) والتحليل الديناميكي (Dynamic Analysis).

كما تم توضيح الخلفية العلمية لمشكلة البحث، والتي تتمثل في محدودية الأنظمة التقليدية في مواجهة التطور المتسارع للبرمجيات الخبيثة الحديثة، خاصة تلك التي تعتمد على تقنيات التعمية (Obfuscation) وتعدد الأشكال (Polymorphism) وأساليب التخفي المتقدمة، مما يجعل عملية الكشف أكثر تعقيدًا.

وقد تم كذلك عرض أهداف المشروع ومساهمته، والتي تركز على بناء نظام تحليلي متكامل يوفر رؤية شاملة لسلوك البرمجيات الخبيثة من خلال تحليل بنيتها الداخلية وسلوكها أثناء التنفيذ داخل بيئة معزولة، بما يساهم في تحسين دقة الكشف وتقليل زمن الاستجابة.

إضافة إلى ذلك، تم تحديد حدود المشروع التي توضح نطاق العمل والقيود التقنية المرتبطة بعملية التحليل، بما يضمن وضوح الإطار العملي للدراسة.

وبشكل عام، يضع هذا الفصل الأساس النظري والتطبيقي للمشروع، ويمهد للانتقال إلى الفصول التالية التي سنتناول الجانب التطبيقي والتقني لبناء النظام المقترح وتحليله بشكل أكثر تفصيلاً.

الفصل الثاني

الأدبيات والدراسات السابقة

1.2 مقدمة:

تُعد البرمجيات الخبيثة (Malware) من أخطر التهديدات السيبرانية في العصر الحديث، حيث تمثل مجموعة من البرامج الضارة المصممة لإلحاق الضرر بالأنظمة الحاسوبية أو سرقة المعلومات أو تعطيل الخدمات. وتشمل هذه البرمجيات أشكالاً متعددة مثل الفيروسات (Viruses)، الديدان (Worms)، أحصنة طروادة (Trojans)، وبرامج الفدية (Ransomware)، والتي تختلف في أساليب عملها ودرجة خطورتها.

ظهر مفهوم البرمجيات الخبيثة منذ بدايات انتشار الحوسبة الشخصية، إلا أن تطورها تسارع بشكل كبير مع انتشار الإنترنت، حيث أصبح المهاجمون قادرين على تصميم برمجيات أكثر تعقيداً وقدرة على التخفي والتكيف مع أنظمة الحماية الحديثة. ومع مرور الوقت، لم تعد هذه البرمجيات تعتمد على أساليب بسيطة، بل أصبحت تستخدم تقنيات متقدمة مثل التعمية (Obfuscation)، وتعدد الأشكال (Polymorphism)، واستغلال الثغرات الأمنية (Exploits) بهدف تجنب الاكتشاف.

وفي الوقت الحالي، تُعتبر البرمجيات الخبيثة جزءاً أساسياً من الهجمات السيبرانية المنظمة، حيث يتم استخدامها في عمليات التجسس الإلكتروني، والهجمات الموجهة، وتعطيل البنى التحتية الحيوية، مما يجعل فهم سلوكها وتحليلها أمراً ضرورياً في مجال الأمن السيبراني.

وبناءً على ذلك، يركز هذا الفصل على المفاهيم الأساسية المرتبطة بالبرمجيات الخبيثة، وطرق تحليلها المختلفة، باعتبارها الأساس العلمي الذي يعتمد عليه المشروع في بناء نظام تحليل متكامل يجمع بين التحليل الساكن والديناميكي لفهم سلوك التهديدات بشكل أكثر دقة وعمق.

2.2 أنواع الفيروسات:

تتعدد أنواع البرمجيات الخبيثة (Malware) باختلاف أساليب عملها وأهدافها، حيث تم تصميم كل نوع ليؤدي وظيفة معينة داخل النظام المستهدف، سواء كانت التخريب أو السرقة أو التجسس أو التحكم عن بُعد. ويُعد فهم هذه الأنواع خطوة أساسية في مجال تحليل البرمجيات الخبيثة، لأنه يساعد على تحديد السلوك المتوقع لكل تهديد وتطوير آليات كشف مناسبة له.

ومن أبرز أنواع البرمجيات الخبيثة ما يلي:

• الفيروسات (Viruses):

وهي برمجيات خبيثة ترتبط بملفات أخرى وتنتشر عند تشغيل هذه الملفات، حيث تقوم بإصابة النظام وإحداث تغييرات غير مرغوبة في الملفات أو النظام.

- **الديدان (Worms):**

هي برمجيات قادرة على الانتشار ذاتيًا عبر الشبكات دون الحاجة إلى ملف مضيف، وتتميز بسرعة انتشارها وقدرتها على استهلاك موارد النظام والشبكة.

- **أحصنة طروادة (Trojans):**

هي برامج تبدو شرعية من الخارج، لكنها تحتوي على وظائف خبيثة مخفية، وغالبًا ما تُستخدم لفتح أبواب خلفية (Backdoors) في الأنظمة.

- **برامج الفدية (Ransomware):**

تقوم بتشفير بيانات الضحية ومنع الوصول إليها، ثم تطلب فدية مالية مقابل فك التشفير، وتُعد من أخطر أنواع الهجمات الحديثة.

- **برامج التجسس (Spyware):**

تهدف إلى جمع المعلومات الحساسة عن المستخدمين دون علمهم، مثل كلمات المرور وسلوك التصفح والبيانات المالية.

- **الإعلانات الضارة (Adware):**

تقوم بعرض إعلانات غير مرغوب فيها وقد تتضمن في بعض الحالات تتبع نشاط المستخدم أو إعادة توجيهه إلى مواقع مشبوهة.

إن هذا التنوع الكبير في أنواع البرمجيات الخبيثة يعكس مدى تعقيد التهديدات السيبرانية الحديثة، ويؤكد الحاجة إلى أساليب تحليل متقدمة قادرة على التعامل مع هذا التباين في السلوكيات، وهو ما يعزز أهمية الاعتماد على التحليل الساكن والديناميكي لفهم هذه التهديدات بشكل دقيق.

1.3.2 كيفية عمل الفيروسات:

تعمل الفيروسات الحاسوبية (Computer Viruses) وفق آلية تعتمد على الارتباط بملفات أو برامج شرعية، بحيث يتم تنشيطها عند تشغيل الملف المصاب. وبعد التنفيذ، تبدأ الفيروسات في تنفيذ شيفرتها الخبيثة داخل النظام، والتي قد تشمل تعديل أو حذف الملفات، إبطاء أداء النظام، أو نشر نفسها إلى ملفات أخرى بهدف زيادة نطاق الانتشار.

تمر عملية عمل الفيروس عادة بعدة مراحل رئيسية، تبدأ بمرحلة **الانتشار (Infection/Propagation)** حيث يقوم الفيروس بالارتباط بملفات قابلة للتنفيذ أو إدخال نفسه داخل مكونات النظام. ثم تأتي مرحلة **التنشيط (Activation)** والتي يتم فيها تشغيل الشيفرة الخبيثة عند تحقق شرط معين مثل فتح ملف أو الوصول إلى تاريخ محدد. بعد ذلك تبدأ مرحلة **التنفيذ (Execution)** حيث يقوم الفيروس بتنفيذ الأوامر الخبيثة مثل التعديل على البيانات أو تعطيل بعض وظائف النظام.

وفي بعض الحالات، تعتمد الفيروسات على تقنيات متقدمة لتجنب الاكتشاف، مثل التشفير (Encryption) أو التعمية (Obfuscation)، مما يجعل تحليلها أكثر صعوبة باستخدام الطرق التقليدية. كما قد تحاول بعض الفيروسات تعطيل برامج الحماية أو إخفاء أثارها داخل النظام لضمان بقائها لأطول فترة ممكنة.

وبالتالي، فإن فهم آلية عمل الفيروسات يُعد خطوة أساسية في مجال تحليل البرمجيات الخبيثة، حيث يساعد في بناء نماذج تحليل فعالة تعتمد على دراسة سلوك الفيروس داخل النظام، سواء من خلال التحليل الساكن أو التحليل الديناميكي داخل بيئة معزولة.

2.3.2 ما يميز الفيروسات عن الأنماط الأخرى:

تتميز الفيروسات الحاسوبية عن غيرها من أنواع البرمجيات الخبيثة بمجموعة من الخصائص التي ترتبط بطريقة انتشارها وآلية عملها داخل النظام. فعلى عكس بعض البرمجيات الخبيثة التي تعمل بشكل مستقل أو تعتمد على الشبكات في الانتشار، تحتاج الفيروسات غالبًا إلى ملف مضيف (Host File) لتتمكن من الانتقال والتفعيل، حيث ترتبط ببرامج أو ملفات شرعية وتُفعل عند تشغيلها من قبل المستخدم.

كما أن الفيروسات تعتمد بشكل أساسي على التفاعل المباشر مع المستخدم، إذ لا يمكنها الانتشار أو التنفيذ إلا من خلال تشغيل الملف المصاب، وهذا يجعلها مختلفة عن الديدان (Worms) التي تستطيع الانتشار ذاتيًا عبر الشبكات دون تدخل المستخدم.

ومن الخصائص المميزة أيضًا أن الفيروسات غالبًا ما تستهدف الملفات المحلية داخل الجهاز، مثل الملفات التنفيذية أو ملفات النظام، بينما تركز أنواع أخرى مثل أحصنة طروادة (Trojans) على إنشاء أبواب خلفية (Backdoors)، أو تستهدف برامج الفدية (Ransomware) تشفير البيانات وابتزاز المستخدم.

كما أن الفيروسات قد تعتمد في بعض الحالات على تقنيات تخفي بسيطة مثل التشفير أو التعديل على بنيتها لتجنب الاكتشاف، إلا أنها عادةً أقل تعقيدًا من التهديدات المتقدمة مثل التهديدات المستمرة المتقدمة (APTs) التي تستخدم أساليب تخفي وتكيف أكثر تطورًا.

وبناءً على ذلك، فإن الفيروسات تمثل أحد الأشكال التقليدية للبرمجيات الخبيثة، إلا أن دراستها تظل مهمة لفهم الأساسيات الأولى لآليات الإصابة والانتشار، والتي تشكل قاعدة لفهم الأنواع الأكثر تعقيدًا في مجال الأمن السيبراني.

3.3.2 أسباب خطورة الفيروسات:

تُعد الفيروسات الحاسوبية من التهديدات الأمنية الخطيرة نظرًا لقدرتها على إحداث أضرار مباشرة داخل الأنظمة المصابة، بالإضافة إلى إمكانية انتشارها بشكل متسلسل بين الملفات والأجهزة. وتكمن خطورتها في عدة عوامل رئيسية تجعلها تحديًا مستمرًا في مجال الأمن السيبراني.

أولاً، تتمثل الخطورة في قدرتها على التدمير المباشر للبيانات، حيث يمكن للفيروسات حذف أو تعديل الملفات الحساسة، مما يؤدي إلى فقدان معلومات مهمة أو تعطيل عمل الأنظمة.

ثانياً، تمتاز الفيروسات بقدرتها على إجراء تغييرات على النظام، مثل تعديل إعدادات النظام أو تعطيل بعض الخدمات الأساسية، وهو ما قد يؤدي إلى عدم استقرار الجهاز أو توقفه عن العمل بشكل صحيح.

ثالثاً، يمكن لبعض البرمجيات الخبيثة المرتبطة بالفيروسات أن تقوم بإنشاء أبواب خلفية ((Backdoors داخل النظام، مما يسمح للمهاجمين بالوصول غير المصرح به والتحكم في الجهاز عن بُعد دون علم المستخدم.

رابعاً، يمكن لبعض الفيروسات أن تقوم بسرقة البيانات الحساسة مثل كلمات المرور والمعلومات الشخصية والمالية، مما يشكل تهديداً مباشراً لخصوصية المستخدم وأمانه الرقمي.

خامساً، تمتاز الفيروسات بقدرتها على الانتشار داخل النظام بشكل غير ملحوظ في المراحل الأولى، مما يمنحها فرصة للتمدد قبل اكتشافها، وبالتالي زيادة حجم الضرر الناتج عنها.

سادساً، قد تؤدي الفيروسات إلى إضعاف أداء النظام بشكل كبير نتيجة استهلاك الموارد مثل المعالج والذاكرة، مما ينعكس سلباً على كفاءة تشغيل الأجهزة والخدمات.

سابعاً، يمكن لبعض الفيروسات أن تعمل كوسيلة تمهيدية لهجمات أكثر تعقيداً، مثل تثبيت برمجيات خبيثة إضافية أو فتح منافذ غير مصرح بها داخل النظام، مما يزيد من مستوى الاختراق الأمني.

وأخيراً، تزداد خطورة الفيروسات في حال دمجها مع تقنيات التخفي الحديثة، حيث يمكنها تغيير شكلها أو سلوكها لتفادي أنظمة الكشف التقليدية، مما يجعل اكتشافها والتعامل معها أكثر صعوبة.

وبناءً على ذلك، فإن فهم أسباب خطورة الفيروسات يُعد أمراً أساسياً في تطوير أنظمة تحليل البرمجيات الخبيثة، سواء عبر التحليل الساكن أو التحليل الديناميكي، بهدف تحسين القدرة على الكشف المبكر وتقليل الأضرار المحتملة.

4.2 الأدبيات والدراسات السابقة:

يُعد تحليل البرمجيات الخبيثة (Malware Analysis) أحد أهم المجالات في الأمن السيبراني، وقد تناولت الدراسات السابقة هذا المجال بالاعتماد على منهجين رئيسيين هما: التحليل الساكن (Static Analysis) والتحليل الديناميكي (Dynamic Analysis) [1]. يعتمد التحليل الساكن على دراسة البرمجية دون تنفيذها، من خلال استخراج الخصائص مثل التوقيعات الرقمية (Hashes)، وسلاسل النصوص (Strings)، وتحليل بنية الملفات التنفيذية. إلا أن هذا الأسلوب يواجه تحديات كبيرة عند التعامل مع البرمجيات التي تستخدم تقنيات التعمية (Obfuscation) أو التشفير أو تعدد الأشكال (Polymorphism) [2]. في المقابل، يقوم التحليل الديناميكي على تشغيل البرمجية داخل بيئة معزولة (Sandbox) لمراقبة سلوكها أثناء التنفيذ. ويساعد هذا النوع من التحليل في كشف الأنشطة الضارة مثل تعديل ملفات النظام، وإنشاء

عمليات خبيثة، أو الاتصال بخوادم خارجية. ومع ذلك، قد تفشل هذه الطريقة في بعض الحالات عند استخدام تقنيات التهرب مثل Anti-VM أو تأخير التنفيذ [3]. وتشير الأدبيات الحديثة إلى أن دمج التحليل الساكن مع التحليل الديناميكي يوفر نتائج أكثر دقة وموثوقية، حيث يكمل كل منهما نقاط ضعف الآخر، مما يساهم في بناء نظام تحليل أكثر قوة وفعالية [4]. وبما أن نطاق هذا المشروع يقتصر على استخدام التحليل الساكن والتحليل الديناميكي فقط، دون التوسع في تقنيات التعلم الآلي أو الذكاء الاصطناعي، فإن التركيز في هذه الدراسة ينصب على تحليل وتقييم فعالية هذه الأساليب التقليدية، مع العمل على تحسين تكاملها ضمن إطار موحد يهدف إلى رفع كفاءة الكشف عن البرمجيات الخبيثة.

1.4.2 الطرق التقليدية في تحليل البرمجيات الخبيثة:

تعتمد الطرق التقليدية في تحليل البرمجيات الخبيثة على أساليب ثابتة تهدف إلى اكتشاف التهديدات بناءً على خصائص أو سلوكيات معروفة مسبقاً [5]. وفي سياق هذا المشروع، سيتم الاعتماد على هذه الطرق التقليدية، مع التركيز بشكل أساسي على تكامل التحليل الساكن والتحليل الديناميكي، دون إدخال تقنيات متقدمة خارج هذا النطاق.

1.1.4.2 التحليل القائم على التوقيعات:

تعتمد هذه الطريقة على مقارنة الملفات مع قاعدة بيانات تحتوي على توقيعات البرمجيات الخبيثة المعروفة. وعند تطابق التوقيع يتم تصنيف الملف مباشرة كملف خبيث. ورغم فعاليتها مع التهديدات المعروفة، إلا أنها غير قادرة على اكتشاف البرمجيات الجديدة (Zero-day [6]).

2.1.4.2 التحليل القائم على القواعد:

يعتمد هذا الأسلوب على مجموعة من القواعد المحددة مسبقاً لاكتشاف السلوكيات المشبوهة مثل إنشاء عمليات غير معتادة أو تعديل ملفات النظام أو الاتصال الخارجي. ويتم تصنيف الملف كخبيث عند تحقق عدد من هذه القواعد [7].

3.1.4.2 التحليل الساكن التقليدي:

يركز هذا النوع على تحليل الملف دون تشغيله، ويشمل فحص البنية الداخلية، واستخراج النصوص، وتحليل الدوال والمكتبات، بالإضافة إلى اكتشاف مؤشرات التشفير أو الترميز [2].

4.1.4.2 التحليل الديناميكي التقليدي:

يعتمد على تشغيل البرمجية داخل بيئة مراقبة وتحليل سلوكها أثناء التنفيذ، مثل مراقبة العمليات والتغيرات في النظام والنشاط الشبكي [3].

5.1.4.2 القيود العامة للطرق التقليدية:

تعاني الطرق التقليدية من عدة قيود، أهمها:

- عدم القدرة على اكتشاف التهديدات الحديثة (Zero-day).
- الاعتماد على التحديث اليدوي للتوقع والقواعد.
- ضعف الأداء أمام تقنيات الترميز والتشفير.
- محدودية القدرة على تحليل كميات كبيرة من الملفات.

وبناءً على ذلك، أصبح من الضروري تطوير أساليب تحليل أكثر ذكاءً تعتمد على الدمج بين التحليل الساكن والديناميكي لتحقيق دقة أعلى في الكشف [4].

2.4.2 تحليل نقدي للدراسات السابقة:

على الرغم من تنوع الدراسات السابقة في مجال تحليل البرمجيات الخبيثة، إلا أنها تعاني من مجموعة من القيود المشتركة التي تحد من فعاليتها في مواجهة التهديدات الحديثة. حيث ركزت بعض الدراسات على التحليل الساكن فقط، مما يجعلها غير قادرة على التعامل مع البرمجيات التي تستخدم تقنيات التشفير والتعمية. في المقابل، اعتمدت دراسات أخرى على التحليل الديناميكي، إلا أنها واجهت تحديات تتعلق باستهلاك الموارد وإمكانية التهرب باستخدام تقنيات مثل Anti-VM.

كما أن العديد من الدراسات التي استخدمت التحليل الهجين (Hybrid Analysis) حققت نتائج أفضل من حيث دقة الكشف، إلا أنها عانت من تعقيد في التنفيذ وصعوبة في دمج النتائج ضمن نظام موحد. إضافة إلى ذلك، يلاحظ أن أغلب الدراسات لم تركز بشكل كافٍ على جانب تفسير النتائج (Explainability)، مما يجعل من الصعب على محللي الأمن فهم قرارات الأنظمة الذكية.

3.4.2 الفجوة البحثية (Research Gap):

من خلال تحليل الدراسات السابقة، يمكن تحديد الفجوة البحثية في النقاط التالية:

- غياب نظام متكامل يجمع بين التحليل الساكن والديناميكي ضمن إطار موحد وفعال.
- ضعف التكامل بين نتائج التحليل وصعوبة الاستفادة منها بشكل عملي.
- عدم قدرة الأنظمة الحالية على التعامل بكفاءة مع التهديدات المتقدمة مثل APTs.
- نقص في الأنظمة التي توفر واجهات تفاعلية تدعم اتخاذ القرار الأمني.

4.4.2 ربط الدراسات بالمشروع:

بناءً على الفجوة البحثية السابقة، يأتي هذا المشروع لتقديم نظام متكامل لتحليل البرمجيات الخبيثة يعتمد على منهجية هجينة تجمع بين التحليل الساكن والتحليل الديناميكي ضمن بيئة موحدة.

كما يسعى المشروع إلى تحسين دقة الكشف من خلال دمج نتائج التحليل المختلفة، وتوظيف تقنيات الذكاء الاصطناعي لدعم اتخاذ القرار، بالإضافة إلى توفير مخرجات واضحة وقابلة للتفسير تساعد محلي الأمن في فهم طبيعة التهديدات بشكل أفضل.

وبذلك، يهدف المشروع إلى تجاوز القيود الموجودة في الأنظمة التقليدية وتقديم حل عملي أكثر كفاءة ومرونة في مواجهة التهديدات السيبرانية الحديثة.

5.2 جدول خلاصة الدراسات السابقة:

يوضح هذا الجدول خلاصة أبرز الدراسات السابقة في مجال كشف البرمجيات الخبيثة وتحليل الروابط، مع توضيح المنهجية المستخدمة، نوع البيانات، أهم النتائج، وأبرز القيود التي واجهتها تلك الدراسات.

المرجع	المؤلفون	السنة	التحليل	مجموعة البيانات	أهم النتائج	القيود / التحديات
[1]	Anderson et al.	2019	تحليل ساكن	عينات من ملفات تنفيذية متنوعة	استخراج خصائص ثابتة مثل سلاسل النصوص (Strings) وبيانات التعريف (Metadata))	ضعف القدرة على اكتشاف البرمجيات المشفرة أو المموهة

[2]	Bailey et al.	2020	تحليل ديناميكي	عينات داخل بيئة Sandbox	كشف سلوكيات وقت التشغيل مثل إنشاء العمليات والاتصالات الشبكية	إمكانية التهرب باستخدام تقنيات Anti-VM
[3]	Moser et al.	2018	تحليل ساكن + ديناميكي	عينات برمجيات خبيثة متعددة	رفع دقة الكشف من خلال دمج التحليل الساكن والديناميكي	استهلاك مرتفع للوقت والموارد
[4]	Willems et al.	2021	تحليل ديناميكي	بيئات افتراضية معزولة	مراقبة سلوك النظام أثناء التنفيذ بشكل تفصيلي	بعض البرمجيات لا تظهر نشاطها داخل بيئات Sandbox
[5]	Egele et al.	2017	تحليل ساكن	ملفات تنفيذية من نوع PE (Windows Executables)	تحليل البنية الداخلية واكتشاف التلاعب والتغيرات	ضعف الأداء أمام تقنيات تعدد الأشكال (Polymorphism))
[6]	Bayer et al.	2022	تحليل ديناميكي	نظام Cuckoo Sandbox	تتبع العمليات والنشاط الشبكي بدقة	صعوبة اكتشاف البرمجيات ذات التنفيذ المؤجل (Delayed Execution)

[7]	Santos et al.	2022	تحليل ساكن	مكتبات وبرمجيات تنفيذية مختلفة	كشف مؤشرات التشفير وتقنيات التمويه	ضعف أمام أساليب الإخفاء المتقدمة (Advanced Obfuscation)
[8]	Kang et al.	2023	تحليل ديناميكي	اكتشاف تغييرات النظام مثل Registry وملفات النظام	اكتشاف تغييرات النظام مثل Registry وملفات النظام	إمكانية التهرب عبر اكتشاف البيئة الافتراضية
[9]	Raff et al.	2021	تحليل ساكن + ديناميكي	مجموعة متنوعة من عينات البرمجيات الخبيثة	تحسين الدقة عند دمج التحليلين بشكل تكاملي	تعقيد كبير في التنفيذ والتكامل
[10]	Lindorfer et al.	2019	تحليل ديناميكي	عينات حقيقية من هجمات Malware	كشف الاتصالات الخبيثة وإنشاء Backdoor	استهلاك مرتفع لموارد النظام

6.2 خلاصة الفصل الثاني:

تناول هذا الفصل الأدبيات والدراسات السابقة في مجال تحليل البرمجيات الخبيثة، حيث تم استعراض أهم الأساليب المستخدمة مثل التحليل الساكن والتحليل الديناميكي، إضافة إلى الطرق التقليدية وقيودها. كما أوضحت الدراسات أهمية الدمج بين نوعي التحليل للحصول على نتائج أكثر دقة وفعالية.

وبناءً على ذلك، يركز هذا المشروع على تطوير نموذج يعتمد على التكامل بين التحليل الساكن والتحليل الديناميكي فقط، باعتبارهما الأساس العملي لتحسين دقة الكشف وفهم سلوك البرمجيات الخبيثة.

الفصل الثالث

منهجية العمل وخطة المشروع

1.3 مقدمة عن منهجية Agile:

يستند هذا المشروع إلى تطوير نظام ويب متكامل باستخدام إطار العمل Laravel، يهدف إلى تحليل البرمجيات الخبيثة (Malware) من خلال الاعتماد على منهجية تحليل مزدوجة تجمع بين:

- التحليل الساكن (Static Analysis)
- التحليل الديناميكي (Dynamic Analysis)

ويتم تنفيذ عمليات التحليل من خلال استدعاء واجهات برمجية (APIs)، سواء كانت محلية أو خارجية، تتولى معالجة الملفات المشبوهة وإرجاع نتائج التحليل إلى النظام، مما يحقق فصلاً معمارياً واضحاً بين طبقة العرض (Frontend) ومنطق المعالجة والتحليل.

ويُعد هذا الفصل بين مكونات النظام من الأساليب الحديثة في تصميم الأنظمة الأمنية، حيث يسمح بعزل عمليات التحليل عن واجهة المستخدم، مما يعزز من قابلية التوسع، وإمكانية الصيانة، وإعادة استخدام مكونات النظام.

كما يتيح هذا النهج مجموعة من المزايا التقنية، من أبرزها:

- تحقيق مرونة عالية في تطوير النظام وتحديث مكوناته بشكل مستقل
- إمكانية توسيع النظام مستقبلاً لدمج محركات تحليل إضافية دون تعديل جذري في البنية الأساسية
- تقليل التعقيد داخل التطبيق من خلال نقل عمليات التحليل إلى خدمات مستقلة
- الاستفادة من خدمات تحليل جاهزة أو متقدمة دون الحاجة إلى إعادة بناء خوارزميات التحليل من الصفر

وبناءً على ذلك، يُوفر هذا النموذج المعماري أساساً قوياً لتطوير نظام تحليل برمجيات خبيثة قابل للتوسع ومرتبطة بمعايير التصميم الحديثة في أنظمة الأمن السيبراني.

2.3 مراحل منهجية إعداد النظام

تقوم منهجية تنفيذ هذا المشروع على بناء منصة ويب مؤسسية (Enterprise Web Platform) لتحليل البرمجيات الخبيثة، تعتمد على إطار العمل Laravel، وتوفر بيئة مركزية لإدارة المستخدمين وتحليل الملفات المرفوعة بشكل آمن ومنظم.

ويعتمد النظام على تكامل عدة وحدات وظيفية تشمل إدارة المستخدمين، التحقق الأمني، تحليل الملفات، والتقارير، وذلك وفق تسلسل منطقي يضمن سلامة البيانات ودقة نتائج التحليل.

1.2.3 إدارة المستخدمين والتحقق الأمني

تبدأ عملية استخدام النظام من خلال قيام مدير المنصة (Administrator) بإنشاء حسابات المستخدمين، مع تحديد صلاحيات كل مستخدم وفقاً لدوره الوظيفي داخل النظام، وذلك لضمان تنظيم الوصول إلى الموارد المتاحة بشكل آمن ومنضبط.

وعند تسجيل الدخول إلى النظام، يتم تطبيق آلية تحقق متعددة المستويات تهدف إلى تعزيز مستوى الأمان ومنع الوصول غير المصرح به، وتشمل هذه الآلية تسجيل الدخول باستخدام بيانات الاعتماد (Credentials)، بالإضافة إلى تنفيذ التحقق الثنائي (Two-Factor Authentication – 2FA) من خلال إرسال رمز تحقق إلى البريد الإلكتروني أو الهاتف المحمول الخاص بالمستخدم، إلى جانب ربط جلسة المستخدم (Session Binding) لضمان استمرار الحماية أثناء استخدام النظام. وتهدف هذه المرحلة بشكل أساسي إلى تأمين عملية الدخول إلى المنصة ومنع أي محاولات وصول غير مصرح بها.

2.2.3 رفع الملفات ومعالجتها الأولية

بعد نجاح عملية التحقق الأمني، يتم تمكين المستخدم من رفع الملفات المشبوهة إلى النظام عبر واجهة مخصصة لذلك. وقبل إرسال أي ملف إلى خدمات التحليل الخارجية، يتم تنفيذ مجموعة من الإجراءات الأولية التي تهدف إلى التحقق من سلامة الملف وتجهيزه لعملية التحليل، وتشمل هذه الإجراءات حساب بصمة الملف (File Hashing) باستخدام خوارزميات تشفير مثل 256SHA، بالإضافة إلى إجراء فحص أولي لاكتشاف المؤشرات المشبوهة (Pre-Analysis Check)، ومراقبة الخصائص الأساسية للملف مثل الحجم ونوع الملف، فضلاً عن التحقق من عدم تكرار الملف داخل النظام لضمان عدم إعادة تحليل نفس العينة بشكل مكرر.

وتساعد هذه المرحلة في تحسين كفاءة النظام وتقليل استهلاك الموارد في عمليات التحليل اللاحقة.

3.2.3 الحماية والعزل (File Isolation & Security Layer)

لضمان حماية النظام من أي تأثيرات محتملة للملفات الخبيثة، يتم تطبيق طبقة أمان قبل بدء عملية التحليل، حيث يتم عزل الملف داخل بيئة آمنة (Isolated Environment) تمنع أي تفاعل مباشر مع النظام الأساسي. كما يتم تشفير الملف داخل النظام قبل تخزينه لضمان سرية البيانات، بالإضافة إلى منع تنفيذ الملف بشكل مباشر على النظام الحقيقي لتجنب أي مخاطر أمنية محتملة، مع حفظ نسخة آمنة منه داخل قاعدة البيانات (Secure Storage) لاستخدامها في عمليات التحليل لاحقاً.

وتهدف هذه المرحلة إلى توفير بيئة تشغيل آمنة تمنع انتشار أي تهديدات محتملة داخل النظام.

4.2.3 التحليل الساكن (Static Analysis)

يتم في هذه المرحلة إرسال الملف إلى خدمة التحليل الساكن أو واجهة برمجية (API) مخصصة، حيث يتم تحليل الملف دون الحاجة إلى تشغيله فعلياً. وتشمل عملية التحليل الساكن استخراج مجموعة من الخصائص المهمة مثل قيم Hash، وتحليل النصوص (Strings) داخل الملف، وفحص المكتبات والدوال المستوردة (Imports & Libraries)، بالإضافة إلى تحليل البنية الداخلية للملف (File Structure)، ومقارنة النتائج مع قواعد بيانات التهديدات مثل VirusTotal عند توفرها. وبناءً على نتائج التحليل يتم تصنيف الملف إلى أحد المستويات التالية: ملف خبيث (Malicious)، أو ملف مشبوه (Suspicious)، أو ملف آمن (Clean). وفي حال ظهور مؤشرات واضحة على وجود تهديد، يتم تصنيف الملف مباشرة دون الحاجة للانتقال إلى مرحلة التحليل الديناميكي.

5.2.3 التحليل الديناميكي (Dynamic Analysis)

في حال عدم حسم نتيجة التحليل الساكن، يتم تحويل الملف إلى بيئة تحليل ديناميكي معزولة (Sandbox Environment) مثل Cuckoo Sandbox أو CAPE Sandbox، حيث يتم تشغيل الملف داخل بيئة مراقبة آمنة بالكامل. وخلال عملية التشغيل يتم تسجيل وتحليل سلوك الملف بشكل تفصيلي، بما في ذلك إنشاء العمليات (Process Creation)، والتغيرات في ملفات النظام، والتعديلات على سجل النظام (Registry Modifications)، والاتصالات الشبكية (Network Connections)، إضافة إلى متابعة السلوك الزمني للملف أثناء التنفيذ. وبعد الانتهاء من عملية المراقبة، يتم تحليل هذه البيانات لإصدار تقرير نهائي يوضح طبيعة سلوك الملف ومدى خطورته.

6.2.3 إدارة النتائج والتقارير

بعد الانتهاء من مرحلتَي التحليل الساكن والديناميكي، يقوم النظام بدمج النتائج في تقرير موحد وشامل لكل ملف. ويتضمن هذا التقرير البيانات الأساسية للملف، ونتائج التحليل الساكن، ونتائج التحليل الديناميكي، بالإضافة إلى تحديد مستوى الخطورة (Risk Level) الخاص بالملف. كما يتم حفظ جميع التقارير داخل قاعدة البيانات بهدف تمكين المستخدم من الرجوع إليها في أي وقت، وتحليلها أو مراجعتها عند الحاجة.

7.2.3 إدارة النظام والصلاحيات (Admin Panel)

يوفر النظام لوحة تحكم خاصة بالمدير (Administrator) تتيح له إدارة النظام بشكل كامل ومركزي. وتشمل صلاحيات المدير إنشاء وإدارة حسابات المستخدمين، وتحديد الصلاحيات الممنوحة لكل مستخدم، بالإضافة إلى عرض جميع الملفات المرفوعة من مختلف المستخدمين، ومراقبة عمليات التحليل التي تتم داخل النظام، والتحكم في أنواع الملفات المسموح برفعها، مع إمكانية حذف أو حظر المستخدمين عند الحاجة، فضلاً عن متابعة سجل النشاطات (Activity Logs) الخاصة بالنظام لضمان الشفافية والأمان.

3.3 دراسة الجدوى:

تُعد دراسة الجدوى خطوة أساسية لتقييم إمكانية تنفيذ النظام المقترح من النواحي التقنية والتشغيلية والاقتصادية والزمنية، وذلك للتأكد من ملاءمته للأهداف المحددة في مجال تحليل البرمجيات الخبيثة. ويعتمد هذا المشروع على تطوير نظام لتحليل البرمجيات الخبيثة باستخدام منهجية تجمع بين التحليل الساكن (Static Analysis) والتحليل الديناميكي (Dynamic Analysis) داخل بيئة معزولة.

1.3.3 الجدوى التقنية (Technical Feasibility)

من الناحية التقنية، يعتمد النظام على مجموعة من التقنيات والأدوات الحديثة والمتاحة، مما يجعله قابلاً للتنفيذ بشكل عملي وفعال. حيث يتم استخدام إطار العمل **Laravel** لتطوير واجهة ويب متكاملة لإدارة عمليات رفع الملفات وعرض نتائج التحليل، بالإضافة إلى استخدام قواعد البيانات لتخزين البيانات ونتائج التحليل بشكل منظم وآمن.

ويعتمد النظام على دمج نوعين أساسيين من التحليل الأمني:

- **التحليل الساكن (Static Analysis):** والذي يتم من خلال واجهات برمجية (Static Analysis API)، حيث يتم استخراج خصائص الملف دون تشغيله، مثل قيم الـ Hashes، النصوص الداخلية (Strings)، المكتبات والدوال المستوردة (Imports)، وتحليل البنية الداخلية للملف التنفيذي.
- **التحليل الديناميكي (Dynamic Analysis):** والذي يتم تنفيذه داخل بيئة افتراضية معزولة تعتمد على تقنية **KVM Virtualization** وتشغيل منصة **CAPEv2 Sandbox**، حيث يتم تشغيل الملفات المشبوهة داخل بيئة مراقبة آمنة لتحليل سلوكها أثناء التنفيذ، بما في ذلك مراقبة العمليات (Processes)، واستدعاءات API، والتغيرات في النظام، والنشاط الشبكي.

كما يعتمد النظام على استدعاء واجهات برمجية (APIs) مخصصة لكل من التحليل الساكن والتحليل الديناميكي، بحيث يتم فصل منطق التحليل عن طبقة التطبيق (Backend Separation)، مما يعزز من الأمان وقابلية التوسع وإعادة الاستخدام.

ويتيح هذا التصميم إمكانية التكامل مع أدوات تحليل متقدمة دون الحاجة إلى تضمينها مباشرة داخل النظام، مما يقلل من التعقيد ويحسن الأداء العام.

وبناءً على ذلك، فإن المشروع **قابل للتنفيذ تقنيًا بالكامل** باستخدام التقنيات الحالية مثل KVM و 2CAPEv و Laravel واجهات API، دون الحاجة إلى تقنيات خارج نطاقه، مع تحقيق مستوى عالٍ من الأمان والمرونة في التحليل.

2.3.3 الجدوى التشغيلية (Operational Feasibility)

من الناحية التشغيلية، تم تصميم النظام ليعمل كمنصة مؤسسية متكاملة لتحليل البرمجيات الخبيثة، بحيث يوفر بيئة سهلة الاستخدام لكل من المستخدمين العاديين ومدير النظام، دون الحاجة إلى خبرة متقدمة في مجال الأمن السيبراني.

تبدأ العملية عند قيام مدير المنصة (Administrator) بإنشاء حسابات المستخدمين وتحديد الصلاحيات الخاصة بكل مستخدم، بما يضمن التحكم الكامل في الوصول إلى النظام وتنظيم استخدامه بشكل آمن. وعند تسجيل الدخول، يتم تطبيق آلية تحقق متعددة المستويات تشمل التحقق من بيانات الدخول، بالإضافة إلى التحقق الثنائي (Two-Factor Authentication) من خلال إرسال رمز تحقق إلى البريد الإلكتروني أو الهاتف، وربط الجلسة (Session Binding) لضمان أمان الاتصال.

بعد تسجيل الدخول بنجاح، يتمكن المستخدم من رفع الملفات المشبوهة عبر واجهة النظام. وقبل بدء عملية التحليل، يتم تنفيذ إجراءات أولية تشمل استخراج قيمة Hash للملف، وفحصه مبدئيًا للكشف عن أي مؤشرات مشبوهة، إضافة إلى التحقق من سلامة الملف وتحديد نوعه.

ثم يتم تمرير الملف إلى مرحلة المعالجة المؤتمتة، حيث يتم أولاً استدعاء **واجهة التحليل الساكن (Static Analysis API)** لتحليل الملف دون تشغيله، واستخراج خصائصه البنيوية مثل التوقيعات، النصوص الداخلية، والمكتبات المستخدمة، إضافة إلى التحقق من وجود مؤشرات تهديد أولية.

وفي حال لم يتم تصنيف الملف بشكل نهائي في التحليل الساكن، يتم تحويله إلى مرحلة **التحليل الديناميكي (Dynamic Analysis)** داخل بيئة معزولة تعتمد على تقنية KVM وتشغيل منصة CAPEv2 Sandbox، حيث يتم تنفيذ الملف داخل بيئة افتراضية آمنة لمراقبة سلوكه بشكل كامل، بما في ذلك العمليات، واستدعاءات API، والتغيرات في النظام والنشاط الشبكي. بعد انتهاء التحليل، يتم دمج النتائج وإرجاع تقرير نهائي يوضح حالة الملف (آمن، مشبوه، أو خبيث)، ويتم عرض النتائج للمستخدم بشكل منظم داخل واجهة النظام.

كما يتيح النظام للمستخدمين العاديين إمكانية عرض جميع الملفات التي قاموا برفعها سابقًا مع نتائج التحليل الخاصة بها، بينما يمتلك مدير النظام صلاحيات كاملة تشمل إدارة المستخدمين، مراقبة جميع الملفات المرفوعة داخل النظام، التحكم في

سياسات رفع الملفات، وتحديد أنواع الملفات المسموح بها، إضافة إلى إدارة سجلات النشاط (Activity Logs) لضمان الرقابة والأمان.

وبناءً على ذلك، يتميز النظام بجدوى تشغيلية عالية نظرًا لسهولة استخدامه، واعتماده على عمليات مؤتمتة بالكامل، وتوفيره لواجهة إدارية متقدمة تدعم التحكم الكامل في النظام وإدارته بكفاءة داخل بيئات العمل الأمنية.

3.3.3 الجدوى الاقتصادية (Economic Feasibility)

يعتمد المشروع في تنفيذه على مجموعة من التقنيات والأدوات مفتوحة المصدر (Open Source) مثل إطار العمل **Laravel** ونظام التشغيل **Linux** وأدوات التحليل داخل بيئات **Sandbox**، مما يساهم بشكل كبير في تقليل التكاليف المالية المرتبطة بالترخيص أو استخدام حلول تجارية مدفوعة.

كما أن تصميم النظام القائم على دمج التحليل الساكن ((Static Analysis والتحليل الديناميكي ((Dynamic Analysis يوفر مستوى عالٍ من الكفاءة الأمنية مقارنةً بالتكلفة المنخفضة للتنفيذ، خصوصًا عند مقارنته بالحلول التجارية المتخصصة في تحليل البرمجيات الخبيثة والتي تتطلب اشتراكات مرتفعة وتكاليف تشغيل عالية.

إضافة إلى ذلك، فإن تحسين سرعة اكتشاف التهديدات وتقليل زمن التحليل يساهم بشكل غير مباشر في تقليل الخسائر الناتجة عن الهجمات السيبرانية، سواء من خلال الحد من توقف الأنظمة أو تقليل الأضرار الناتجة عن انتشار البرمجيات الخبيثة داخل الشبكات. وبناءً على ذلك، يمكن اعتبار المشروع **مجددًا اقتصاديًا** نظرًا لانخفاض تكاليفه التشغيلية مقابل القيمة الأمنية العالية التي يقدمها.

3.3.4 الجدوى الزمنية (Time Feasibility)

الجدول الزمني للمشروع

من بداية أكتوبر حتى نهاية فبراير

فبراير				يناير				ديسمبر				نوفمبر				أكتوبر				المدة (أسابيع)	المهام الرئيسية	المرحلة
1	2	3	4	1	2	3	4	1	2	3	4	1	2	3	4	1	2	3	4			
																				3 أسابيع	- تصميم نظام المستخدمين والصلاحيات - تنفيذ تسجيل الدخول - تطبيق التحقق الثنائي (2FA) - ربط جلسة المستخدم (Session Binding)	1. إدارة المستخدمين والتحقق الأمني
																				3 أسابيع	- تطوير واجهة رفع الملفات - حساب الهاش (Hashing) - الفحص الأولي واكتشاف المؤشرات المشبوهة - التحقق من التكرار	2. رفع الملفات ومعالجتها الأولية
																				2 أسابيع	- عزل الملف في بيئة آمنة - تشفير الملف قبل التخزين - منع التنفيذ المباشر للملف - حفظ الملف في قاعدة البيانات	3. الحماية والعزل (File Isolation & Security Layer)
																				3 أسابيع	- التكامل مع API التحليل الساكن - استخراج خصائص الملف - مقارنة النتائج مع قواعد التهديدات - تصنيف الملف (خبيث/مشبوه/سليم)	4. التحليل الساكن (Static Analysis)
																				4 أسابيع	- إعداد بيئة (CAPE) Sandbox - تشغيل الملف ومراقبة السلوك - تحليل العمليات والشبكة والتغييرات - إصدار التقرير الديناميكي	5. التحليل الديناميكي (Dynamic Analysis) (Sandbox – CAPE)
																				2 أسابيع	- دمج نتائج التحليل الساكن والديناميكي - إنشاء تقارير شاملة - تخزين التقارير وإتاحتها للمستخدمين	6. إدارة النتائج والتقارير
																				3 أسابيع	- تطوير لوحة تحكم المدير - إدارة المستخدمين والصلاحيات - ضبط أنواع الملفات المسموح بها - سجل النشاطات والمراقبة	7. إدارة النظام والصلاحيات (Admin Panel)
																				2 أسابيع	- اختبار الوحدات والتكامل - اختبار الأداء والأمان - إصلاح الأخطاء والتحسين - النشر النهائي وإطلاق المتصة	8. الاختبار والتجريب والإطلاق
20 أسبوع (من بداية أكتوبر إلى نهاية فبراير)																				إجمالي مدة المشروع الكلية		

4.3 خلاصة دراسة الجدوى:

يتضح من التحليل أن المشروع قابل للتنفيذ من جميع الجوانب التقنية والتشغيلية والاقتصادية والزمنية، ويعتمد على تقنيات واقعية ومجربة، مما يجعله مناسباً لتطوير نظام فعال في تحليل البرمجيات الخبيثة باستخدام التحليل الساكن والديناميكي بشكل تكاملي.

5.3 أدوات وتقنيات المشروع:

يعتمد هذا المشروع على مجموعة من الأدوات والتقنيات الحديثة التي تساهم في بناء نظام متكامل لتحليل البرمجيات الخبيثة بكفاءة عالية، مع تحقيق التكامل بين التحليل الساكن (Static Analysis) والتحليل الديناميكي (Dynamic Analysis) ضمن بيئة آمنة وقابلة للتوسع.

وفيما يلي أهم الأدوات والتقنيات المستخدمة:

1. لغة البرمجة (PHP):

يتم استخدام لغة PHP كأساس لتطوير منطق الخادم (Backend Logic)، نظرًا لقوتها في بناء تطبيقات الويب ودعمها الواسع في بيئة تطوير الأنظمة المعتمدة على الويب.

2. إطار العمل: Laravel

يُعد Laravel الإطار الرئيسي المستخدم في بناء النظام، حيث يوفر بنية منظمة تعتمد على نمط MVC، مما يسهل إدارة الأكواد، وتحسين الأمان، وتنظيم عمليات معالجة الطلبات وربطها مع قواعد البيانات وواجهات الـ APIs.

3. قواعد البيانات (MySQL):

يتم استخدام MySQL لتخزين بيانات المستخدمين، الملفات المرفوعة، ونتائج التحليل، بالإضافة إلى تقارير الفحص، مع ضمان تنظيم البيانات وإمكانية استرجاعها بكفاءة عالية.

4. التحليل الساكن (Static Analysis APIs):

يتم الاعتماد على واجهات برمجية للتحليل الساكن، من ضمنها التكامل مع **VirusTotal API**، وذلك لاستخراج خصائص الملفات دون تنفيذها، مما يساعد في الكشف المبكر عن التهديدات ومقارنتها مع قواعد بيانات عالمية للبرمجيات الخبيثة.

ويشمل التحليل الساكن:

- **Hash Values:** حساب بصمة الملف (256SHA) ومقارنتها مع قواعد بيانات التهديدات ونتائج VirusTotal.
- **Strings Extraction:** استخراج النصوص الداخلية لاكتشاف روابط أو أوامر أو خوادم خبيثة (2C).
- **File Metadata:** تحليل معلومات الملف مثل الحجم والنوع والمكتبات والتواريخ لتقييم خطورته.

وبشكل عام، فإن دمج VirusTotal API يعزز دقة التحليل ويرفع موثوقية الكشف عن الملفات الخبيثة المعروفة.

5. بيئة التحليل الديناميكي (Dynamic Analysis Sandbox):

يتم استخدام بيئات معزولة مثل:

- KVM Virtualization

- CAPEv2 Sandbox

وذلك لتشغيل الملفات المشبوهة ومراقبة سلوكها بشكل كامل، مثل العمليات (Processes)، تغييرات النظام، والنشاط الشبكي.

6. لغة Python:

تُستخدم Python بشكل أساسي داخل بيئة التحليل الديناميكي، نظرًا لقوتها في التعامل مع تحليل السلوك، وأتمتة العمليات، والتكامل مع أدوات الأمن السيبراني.

7. JavaScript و AJAX:

تُستخدم في تطوير واجهة المستخدم (Frontend) لجعل النظام تفاعليًا، حيث تتيح AJAX تنفيذ عمليات رفع الملفات وعرض نتائج التحليل دون الحاجة إلى إعادة تحميل الصفحة.

8. واجهات برمجية خارجية (APIs):

مثل VirusTotal API (عند الحاجة) لدعم عملية التحقق من الملفات ومقارنتها مع قواعد بيانات التهديدات العالمية.

9. بيئة التشغيل (Linux Server):

يتم تشغيل النظام داخل بيئة Linux نظرًا لاستقرارها العالي ودعمها القوي لأدوات الأمن السيبراني وبيئات الـ Sandbox.

10. تقنيات العزل والحماية:

يتم استخدام تقنيات العزل (Isolation) لضمان تشغيل الملفات المشبوهة داخل بيئة آمنة دون التأثير على النظام الحقيقي.

وبشكل عام، فإن هذا التكامل بين الأدوات والتقنيات يوفر بنية قوية ومرنة للنظام، تسمح بتحليل البرمجيات الخبيثة بدقة عالية، مع إمكانية التوسع المستقبلي وإضافة تقنيات تحليل أكثر تقدمًا.

6.3 إجراءات التنفيذ:

سير العمل التفصيلي للمشروع يمر عبر المراحل الموضحة في الصورة التالية:



تسجيل الدخول وإدارة الوصول: تبدأ العملية بتوثيق هوية المستخدم عبر نظام أمن يضمن توزيع الصلاحيات والتحقق الثنائي (FA2) لضمان حماية المنصة.

رفع الملف والمعالجة الأولية: يقوم المستخدم برفع الملف المشبوه، حيث يتم فوراً حساب بصمته الرقمية (Hash) والتحقق من خصائصه الأساسية دون تشغيله.

التحليل الساكن (Static Analysis): يتم فحص بنية الملف الداخلية، واستخراج النصوص (Strings) والمكتبات البرمجية، ومقارنته بقواعد بيانات التهديدات العالمية مثل VirusTotal.

التقييم الأولي: بناءً على نتائج التحليل الساكن، يتم تصنيف الملف؛ فإذا كان آمناً تنتهي العملية، وإذا كان مشبوهاً يتم نقله للمرحلة التالية.

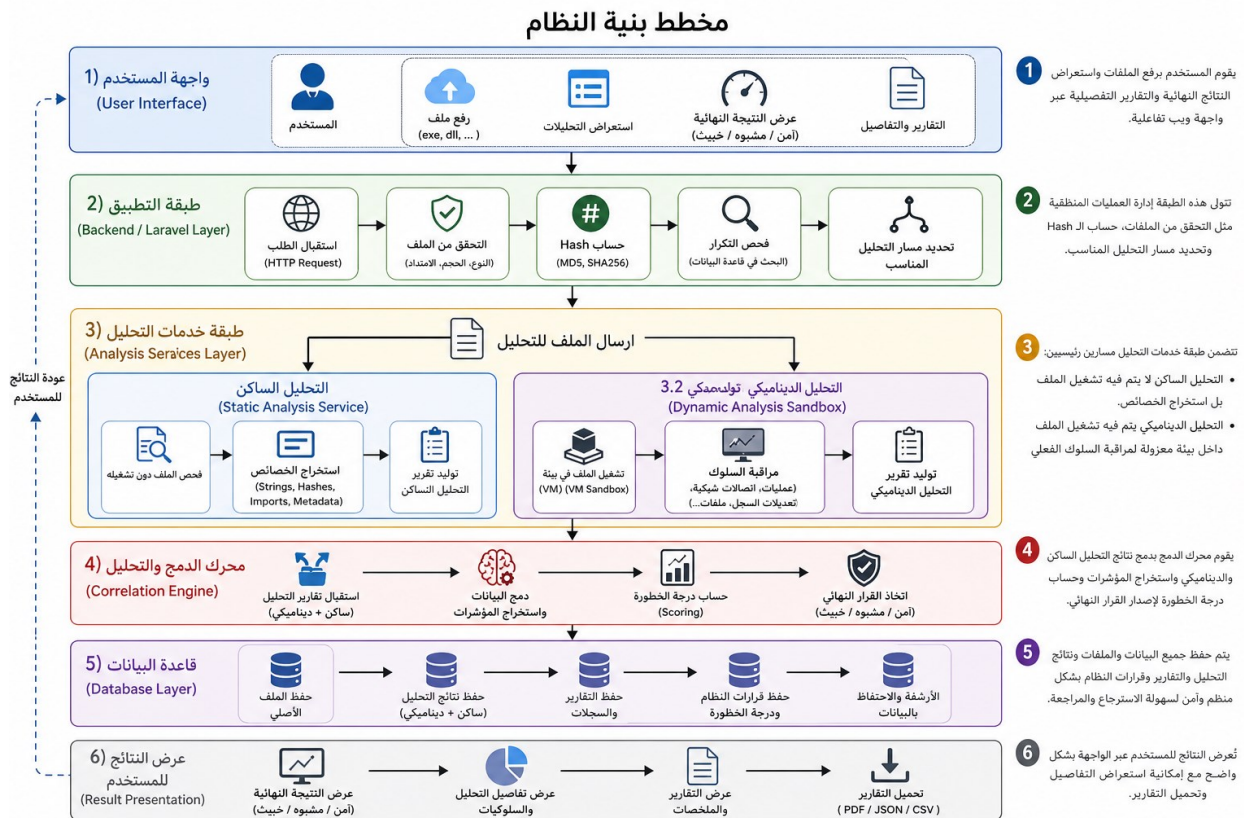
التحليل الديناميكي (Dynamic Analysis): يتم تشغيل الملف داخل بيئة معزولة تماماً (Sandbox) لمراقبة سلوكه الفعلي، مثل اتصالاته بالشبكة وتعديلاته على سجل النظام.

تحليل النتائج والسلوك: يقوم النظام بجمع كافة البيانات السلوكية وتحديد "مؤشرات الاختراق" (IoCs) والنشاطات الخبيثة التي ظهرت أثناء التشغيل.

إنشاء التقرير الشامل: يتم دمج نتائج التحليلين الساكن والديناميكي في تقرير فني موحد يحدد مستوى الخطورة النهائي للملف.

عرض النتائج والإدارة: يُعرض التقرير النهائي للمستخدم عبر الواجهة، مع أرشفته في قاعدة البيانات للرجوع إليه لاحقاً تحت إشراف كامل من مدير النظام.

3.7 مخطط بنية النظام (System Architecture):



8.3 خلاصة الفصل الثالث:

يمثل هذا الفصل الأساس العملي لتطوير النظام المقترح، حيث تم استعراض المنهجية المعتمدة، وتوضيح مراحل التنفيذ، إلى جانب عرض الأدوات والتقنيات المستخدمة في بناء النظام. وبناءً على ذلك، يُعد هذا الفصل تمهيداً للانتقال إلى الفصل الرابع، الذي يتناول تصميم النظام وتنفيذه بصورة أكثر تفصيلاً وعمقاً.

الفصل الرابع

التحليل الديناميكي السلوكي المتقدم

1.4 المقدمة والأسس المعرفية (Epistemological Foundations)

في ظل التطور المتسارع للبرمجيات الخبيثة فائقة التعدد (Hyper-Polymorphic) والبرمجيات عديمة الملفات (Fileless) (Malware)، لم يعد التحليل الساكن (Static Analysis) القائم على البصمات الرقمية كافياً لتوفير حماية شاملة. يمثل التحليل الديناميكي (Dynamic Analysis) تحولاً نوعياً من فحص "الكيان المادي" للملف إلى فحص "السلوكية" له أثناء التنفيذ. تعتمد هذه الدراسة على مبدأ "التشخيص السلوكي الوظيفي"، حيث يتم استنتاج البرمجية الخبيثة في بيئة مختبرية معزولة للكشف عن نواياها الحقيقية.

تكمن القيمة العلمية للتحليل الديناميكي في قدرته على تجاوز طبقات التعتيم (Obfuscation) والتعبئة (Packing)؛ فهما بلغت درجة تعقيد تشفير الكود المصدري، لا بد للبرمجية من فك تشفير نفسها في الذاكرة العشوائية (RAM) والتفاعل مع نظام التشغيل لتحقيق أهدافها، وهنا تبرز أهمية الرصد المجهرى لهذه التفاعلات.

2.4 هندسة الأنظمة والبيئات المعزولة (Sandboxing & System Architecture)

تعتبر البيئة المعزولة (Sandbox) الركيزة الأساسية للتحليل الديناميكي. من الناحية العلمية، يتم تصميم هذه البيئات لتحقيق توازن دقيق بين ثلاثة معايير: العزل التام (Isolation)، الشفافية (Transparency)، والدقة في الرصد (Observability).

3.4 نظرية المحاكاة والافتراضية (Emulation vs. Virtualization)

تعتمد منصات التحليل الحديثة على بنى معمارية مختلفة لتحقيق أهدافها:

جدول 1-4 يوضح نظرية المحاكاة والافتراضية

المعيار التقني	المحاكاة الكاملة (Full Emulation)	الافتراضية المدعومة عتادينا (Hardware-Assisted Virtualization)
الآلية	محاكاة كل تعليمة برمجية عبر البرمجيات (مثل QEMU).	تنفيذ التعليمات مباشرة على المعالج الفيزيائي (مثل KVM/Xen).
الشفافية	عالية جداً؛ يصعب على البرمجية اكتشاف البيئة.	متوسطة؛ قد تترك أدوات الافتراضية أثراً في الذاكرة.
الأداء	بطيء نسبياً بسبب طبقة الترجمة البرمجية.	سريع جداً، يقارب أداء الأجهزة الحقيقية.
التحكم	تحكم مطلق في كل دورة من دورات المعالج.	تحكم معتمد على ميزات المعالج (VT-x / AMD-V).

1.3.4 معضلة "الفجوة الدلالية" (The Semantic Gap Problem)

تعد الفجوة الدلالية من أكبر التحديات العلمية في التحليل الديناميكي، وهي المسافة بين البيانات الخام التي يراها المحلل (مثل بايتات الذاكرة وسجلات المعالج) والمعنى الوظيفي لها في نظام التشغيل (مثل أسماء العمليات والملفات). يتم حل هذه المعضلة عبر تقنيات استبطان الجهاز الافتراضي (Virtual Machine Introspection - VMI)، والتي تسمح للمحلل بإعادة بناء هيكلية نظام التشغيل الضيف من الخارج، مما يضمن مراقبة البرمجية دون أن تشعر بوجود "جسم غريب" داخل نظامها "وتطبيقاً لهذه المفاهيم في منصتنا المقترحة، تم الاعتماد على بيئة (Cuckoo Sandbox) المدمجة مع نظام (KVM/QEMU) لتحقيق أقصى درجات العزل (Isolation)، مما يضمن لنا استخراج بيانات سلوكية دقيقة للبرمجيات الخبيثة مع تقليل احتمالية اكتشاف البرمجية لبيئة التحليل".

4.4 ميكانيكا الاعتراض والتحكم في التدفق (Flow Control & Interception)

لتحقيق رصد مجهري، يجب على نظام التحليل "اعتراض" نداءات النظام (System Calls) والتفاعلات مع النواة.

1.4.4 نظرية مستويات الحماية (Protection Rings Theory)

يعمل نظام التحليل في مستويات مختلفة لضمان السيطرة:

المستوى الثالث (Ring 3): مراقبة واجهات برمجة التطبيقات (APIs) التي تستدعيها البرمجية مباشرة.

المستوى الصفر (Ring 0): العمل داخل النواة لمراقبة نداءات النظام العميقة وتغيير سلوكها إذا لزم الأمر.

المستوى السالب واحد (Ring -1): العمل من خلال الـ Hypervisor، وهو المستوى الأكثر أماناً حيث يراقب كل شيء دون أن يكون جزءاً من نظام التشغيل المستهدف.

2.4.4 آليات الاعتراض المتقدمة (Advanced Hooking Mechanisms)

تتم عملية الاعتراض عبر عدة استراتيجيات علمية:

اعتراض الجداول الديناميكية (Table Hooking): مثل التلاعب بجدول IAT أو SSDT لتوجيه النداءات إلى وظائف التحليل.

الرقع المباشرة (Inline Patching): تعديل كود الوظيفة في الذاكرة لإضافة تعليمة قفز (Jump) نحو المحلل.

الاعتراض المعتمد على الاستثناءات (Exception-based Hooking): وضع نقاط توقف (Breakpoints) تسبب استثناءات للنظام، مما يجبر المعالج على تسليم التحكم للمحلل لمعالجة الاستثناء وتسجيل النشاط.

وعملياً، يعتمد نظام التحليل الديناميكي في مشرونا على تقنيات اعتراض نداءات النظام (API Hooking) لتسجيل كل تسلسل زمني للوظائف التي تستدعيها العينة (مثل دوال حقن الذاكرة أو تعديل السجل). هذه التسلسلات يتم تمريرها لاحقاً إلى شبكات (LSTM) في محرك الذكاء الاصطناعي لتحليلها.

5.4 تحليل الذاكرة العشوائية والبيانات المتطايرة

بما أن الذاكرة هي "مسرح الجريمة" الحقيقي، فإن تحليلها يتطلب فهماً عميقاً لهيكلية أنظمة التشغيل.

1.5.4 هياكل البيانات الحيوية (Vital OS Structures)

يقوم المحلل بفحص الهياكل التالية لكشف النشاط المخفي:

PEB (Process Environment Block) : يحتوي على معلومات حول العملية، المكتبات المحملة، ومعاملات التشغيل.

VAD (Virtual Address Descriptor) : شجرة بيانات تصف مساحات الذاكرة المحجوزة لكل عملية. كشف التناقضات بين VAD والذاكرة الفعلية هو المفتاح لكشف الكود المحقون.

EPROCESS : الهيكل الأساسي في نواة ويندوز الذي يمثل العملية، ومن خلاله يمكن كشف العمليات التي حاولت البرمجية إخفاءها من قائمة المهام.

2.5.4 كشف تقنيات الحقن (Process Injection Detection)

يتم تحليل الذاكرة لكشف أنماط الحقن السلوكية مثل Process Hollowing و AtomBombing. من الناحية العلمية، يتم ذلك عبر مراقبة عمليات حجز الذاكرة بصلاحيات التنفيذ (PAGE_EXECUTE_READWRITE) والتي تعتبر علامة حمراء في التحليل السلوكي، حيث لا تحتاج البرمجيات الشرعية عادة لهذه الصلاحيات بشكل دائم.

6.4 نظرية المعلومات وتحليل الاتصالات (C2 Analysis & Information Theory)

تتواصل البرمجيات الخبيثة مع خوادم القيادة والسيطرة (C2) عبر قنوات مشفرة أو مخفية.

1.6.4 تحليل الانتروبيا (Entropy Analysis)

يستخدم المحللون "انتروبيا شانون" (Shannon Entropy) لقياس مدى العشوائية في البيانات المرسل. البيانات المشفرة أو المضغوطة تمتلك انتروبيا عالية جداً (تقترب من 8)، مما يساعد في كشف القنوات السرية التي تحاول إخفاء البيانات المسروقة داخل بروتوكولات عادية مثل HTTP أو DNS.

2.6.4 إحصائيات إشارات التنبيه (Beaconing Statistics)

تعتمد البرمجيات الخبيثة على إشارات دورية (Beacons). يتم تحليل هذه الإشارات باستخدام:

تحليل الفواصل الزمنية (Inter-Arrival Time Analysis): كشف الأنماط شبه الدورية التي تشير إلى نشاط آلي.
تذبذب التوقيت (Jitter Analysis): محاولة المهاجمين إضافة عشوائية بسيطة للتوقيت للهروب من الكشف الإحصائي، وهو ما يتم مواجهته بنماذج رياضية متقدمة لكشف "العشوائية المصطنعة".

7.4 نظريات التهرب ومكافحة التحليل (Anti-Analysis Theory & Evasion)

تمثل البرمجيات الخبيثة الحديثة أنظمة دفاعية وهجومية متكاملة، حيث تسعى جاهدة لاكتشاف بيئة التحليل وتعطيلها أو تضليلها.

1.7.4 نظرية اكتشاف البيئة الافتراضية (Virtualization Detection Theory)

تعتمد البرمجيات الخبيثة على عدة استراتيجيات علمية لاكتشاف السانديوكس:

كشف البقايا العنادية (Hagrdware Artifacts): البحث عن أسماء الأجهزة الافتراضية (مثل "VBOX" "VMWARE") في سجلات النظام أو عناوين MAC الخاصة ببطاقات الشبكة.

كشف الموارد المحدودة (Resource Constraints): البرمجية تفحص عدد الأنوية (Cores)، حجم الذاكرة، ومساحة القرص. البيئات الافتراضية غالباً ما تخصص موارد محدودة، وهو ما يعتبره المهاجم علامة على وجوده في مختبر تحليل.

تحليل زمن التنفيذ (Timing-based Evasion): استخدام تعليمات مثل RDTSC لقياس الوقت المستغرق لتنفيذ بعض التعليمات. إذا كان الوقت أطول من الطبيعي، فهذا يعني وجود طبقة محاكاة أو مراقب (Hypervisor) يقوم باعتراض التنفيذ.

2.7.4 تزيف البيئة وتضليل المحلل (Counter-Intelligence & Environment Faking)

- لمواجهة هذه التقنيات، يستخدم نظام التحليل لدينا استراتيجيات "الذكاء المضاد":
- تزيف العتاد (Hardware Masking): تغيير أسماء الأجهزة وعناوين MAC لتبدو كأجهزة حقيقية تماماً.
- تزيف الموارد (Resource Spoofing): إيهام البرمجية بأنها تعمل على جهاز ذو مواصفات عالية (8 أنوية، 32 جيجابايت رام).
- تسريع الوقت (Time Dilation): عند رصد نداءات الانتظار (Sleep) الطويلة التي تستخدمها البرمجية لتجاوز فترة التحليل، يقوم المحلل بتغيير قيمة التوقيت في النواة لجعل البرمجية "تعتقد" أن الساعات قد مرت بينما هي ثوانٍ معدودة.

8.4 التحليل الشكلي والتنفيذ الرمزي (Symbolic Execution & Formal Methods)

يعد التنفيذ الرمزي من أكثر التقنيات العلمية تقدماً في استكشاف مسارات التنفيذ الممكنة للبرمجية.

1.8.4 نظرية استكشاف المسارات (Path Exploration Theory)

بدلاً من تنفيذ البرمجية بمدخلات حقيقية، يتم التعامل مع المدخلات كـ "رموز رياضية" (Symbols). يقوم المحلل ببناء "شجرة قرار" (Decision Tree) لكل مسار ممكن يمكن أن تسلكه البرمجية بناءً على شروط معينة (مثل وجود ملف معين أو اتصال شبكي). يساعد هذا في كشف "القنابل المنطقية" (Logic Bombs) التي لا تنفجر إلا في ظروف محددة جداً.

2.8.4 معالجة الانفجار المزيج (State Explosion Problem)

من التحديات العلمية الكبرى في التنفيذ الرمزي هو "انفجار الحالات"، حيث يمكن أن يكون للبرمجية ملايين المسارات الممكنة. يتم حل ذلك عبر خوارزميات البحث الذكي (Heuristic Search) التي تركز على المسارات الأكثر احتمالاً لاحتواء السلوك الخبيث، مثل تلك التي تؤدي لنداءات نظام حساسة.

9.4 التحليل الشبكي وفك تشفير البروتوكولات

تعتبر الشبكة هي القناة الحيوية التي تربط البرمجية بمشغلها.

1.9.4 القنوات السرية وتحليل البروتوكولات (Covert Channels)

تستخدم البرمجيات الخبيثة بروتوكولات مشروعة لإخفاء بيانات غير مشروعة:

- نفق الـ DNS (DNS Tunneling): إرسال البيانات المسروقة كجزء من استعلامات DNS. يتم تحليل هذا السلوك عبر مراقبة حجم الاستعلامات، وتكرارها، وطول النطاقات الفرعية (Subdomains) التي لا تتبع الأنماط البشرية العادية.
- نفق الـ ICMP (Ping Exfiltration): إخفاء البيانات داخل حقل البيانات في حزم الـ Ping. يتم كشفها عبر تحليل حجم الحزم الذي يتجاوز الحجم الافتراضي لنظام التشغيل.

2.9.4 فك تشفير SSL/TLS في بيئة التحليل (SSL Decryption Theory)

للتغلب على التشفير، يتم استخدام تقنية الرجل في المنتصف (Man-in-the-Middle - MITM) داخل السانديوكس. يتم تثبيت شهادة جذر (Root Certificate) خاصة بالمحل، مما يسمح بفك تشفير حركة المرور، وتسجيل البيانات المرسلة (مثل كلمات المرور والملفات)، ثم إعادة تشفيرها وإرسالها للخادم الحقيقي لضمان استمرار عمل البرمجية وكشف كامل سلوكها.

10.4 دراسات حالة كتشريح علمي (Scientific Post-Mortem Case Studies)

1.10.4 دراسة حالة: برمجية Stuxnet (The Cyber-Weapon)

- تعتبر Stuxnet أول سلاح سبيرانلي حقيقي. من الناحية العلمية، أظهر التحليل الديناميكي لها قدرة فائقة على:
- الانتشار الجانبي (Lateral Movement): استخدام 4 ثغرات "يوم الصفر" (Zero-day) للانتشار داخل الشبكات المعزولة.
 - الاستهداف الدقيق: فحص وجود برمجيات التحكم الصناعي (SCADA) قبل تفعيل الحمولة الخبيثة.
 - التخفي المجهرى: استخدام Rootkit في مستوى النواة لإخفاء التغييرات التي تجريها على أنظمة التحكم.

2.10.4 دراسة حالة: برمجية Pegasus (Zero-Click Spyware)

- أثبتت Pegasus أن التحليل الديناميكي التقليدي قد يفشل إذا لم يكن مجهزاً بما يكفي:
- الثبات في الذاكرة: البرمجية لا تترك أي أثر على القرص الصلب، وتعمل بالكامل في الذاكرة العشوائية.
 - تجاوز الحماية العادية: القدرة على تجاوز ميزات الحماية في أنظمة iOS و Android عبر استغلال ثغرات في معالجة الصور والرسائل.
 - التحليل المطلوب: استلزم كشفها تحليل ذاكرة حي (Live Memory Analysis) وفحص دقيق لسجلات الشبكة لكشف الاتصالات المشفرة بخوادم C2.

3.10.4 التعلم التعزيزي في الدفاع السيبراني (Reinforcement Learning - RL)

يتم استخدام التعلم التعزيزي لبناء "أنظمة دفاعية ذاتية التكيف". في هذا النموذج:

- الوكيل (Agent): هو نظام التحليل الديناميكي.
- البيئة (Environment): هي الساندبوكس والبرمجية الخبيثة.
- المكافأة (Reward): هي النجاح في كشف سلوك خبيث أو تجاوز تقنية تهرب.

يسمح هذا للأنظمة بتعلم كيفية مواجهة تقنيات التهرب الجديدة تلقائياً دون تدخل بشري، مما يخلق "مختبر تحليل ذكي" يتطور باستمرار.

11.4 التوجهات المستقبلية والبحث العلمي المتقدم

يتطور مشهد التهديدات بسرعة، مما يفرض تحديات جديدة على التحليل الديناميكي.

1.11.4 البرمجيات الخبيثة المقاومة للحوسبة الكمومية (Quantum-Resistant Malware)

مع اقتراب عصر الحوسبة الكمومية، يتوقع الباحثون ظهور برمجيات خبيثة تستخدم تشفيراً كمومياً للتواصل مع خوادمها. سيتطلب التحليل الديناميكي لهذه البرمجيات:

- فك تشفير كمومي (Quantum Decryption): قدرات تحليلية تستطيع التعامل مع الخوارزميات الكمومية.
- محاكاة بيئة كمومية: القدرة على مراقبة العمليات التي تجري داخل المعالجات الكمومية (Qubits).

2.11.4 التهرب المعتمد على الذكاء الاصطناعي (AI-Driven Evasion)

بدأ المهاجمون باستخدام الذكاء الاصطناعي لتوليد "سلوكيات طبيعية" (Benign-looking Behaviors) لتضليل المحللين. لمواجهة ذلك، يجب أن يركز التحليل الديناميكي على:

- تحليل الشذوذ المجهرى (Micro-Anomaly Detection): كشف الانحرافات البسيطة جداً في توقيت أو تسلسل النداءات التي لا يمكن للذكاء الاصطناعي تقليدها تماماً.
- التحليل الشمولي (Holistic Analysis): عدم الاعتماد على نداءات API فقط، بل مراقبة استهلاك الطاقة، درجة حرارة المعالج، وأنماط الوصول للذاكرة المخبئ به (Cache Side-Channel Analysis).

الفصل الخامس

الواجهات والتنفيذ

1.5 التنفيذ

تعتبر منصة التحليل الهجين للبرمجيات الخبيثة نظاماً معقداً يجمع بين هندسة الويب الحديثة، علوم البيانات، والأمن السيبراني. يهدف هذا الدليل إلى تقديم رؤية 360 درجة حول كيفية بناء وتشغيل المنصة باستخدام إطار عمل Laravel والتقنيات الرديفة.

1.1.5 الرؤية الهندسية والبنية التحتية (System Architecture)

يعتمد النظام على بنية **Distributed Micro-services Oriented Architecture**، حيث يعمل Laravel كعقل مدبر (Orchestrator) يدير الطلبات، بينما تخصص موارد مستقلة لعمليات التحليل الكثيفة.

1.1.1.5 طبقة الويب (Laravel Application Layer)

تعمل هذه الطبقة على PHP 8.2 ، مستفيدة من ميزات مثل Readonly Properties و Enumerations و Fiber لإدارة العمليات المتزامنة بشكل أفضل. يتم معالجة طلبات المستخدمين، إدارة الهوية (Authentication) ، وعرض البيانات عبر واجهة مستخدم مبنية بـ Tailwind CSS.

تعرض الصفحة الرئيسية للمنصة واجهة مستخدم بديهية تركز على زر "Start Analysis" لبدء عملية الرفع، مع قائمة جانبية (Sidebar) تحتوي على روابط سريعة للوحة التحكم والرفع وسجل التحليلات ودليل المستخدم.



الشكل 1-5: يوضح الصفحة الرئيسية لمنصة التحليل الهجين

2.1.1.5 النظام يتيح للمستخدمين إنشاء حساب جديد عبر الواجهة التالية:

"يتيح النظام للمستخدمين إنشاء حساب جديد عبر واجهة آمنة وبمبسطة. ولحماية المنصة من محاولات التسجيل العشوائي أو هجمات الأتمتة (Bot Attacks)، تم إخضاع عملية التسجيل لآليات تحديد المعدل (Rate Limiting) المدعومة من إطار عمل Laravel، مما يضمن استقرار قاعدة البيانات ويمنع استنزاف الموارد. كما يعتمد النظام معمارية "الأمان الاستباقي"، حيث يُلزم المستخدم فور تسجيل الدخول بالانتقال إلى إعداد المصادقة الثنائية (2FA)، مما يضمن عدم وجود أي حساب مفعّل دون طبقة حماية إضافية قبل السماح له برفع أو تحليل أي ملف".

الشكل 2-5: يوضح صفحة إنشاء حساب جديد

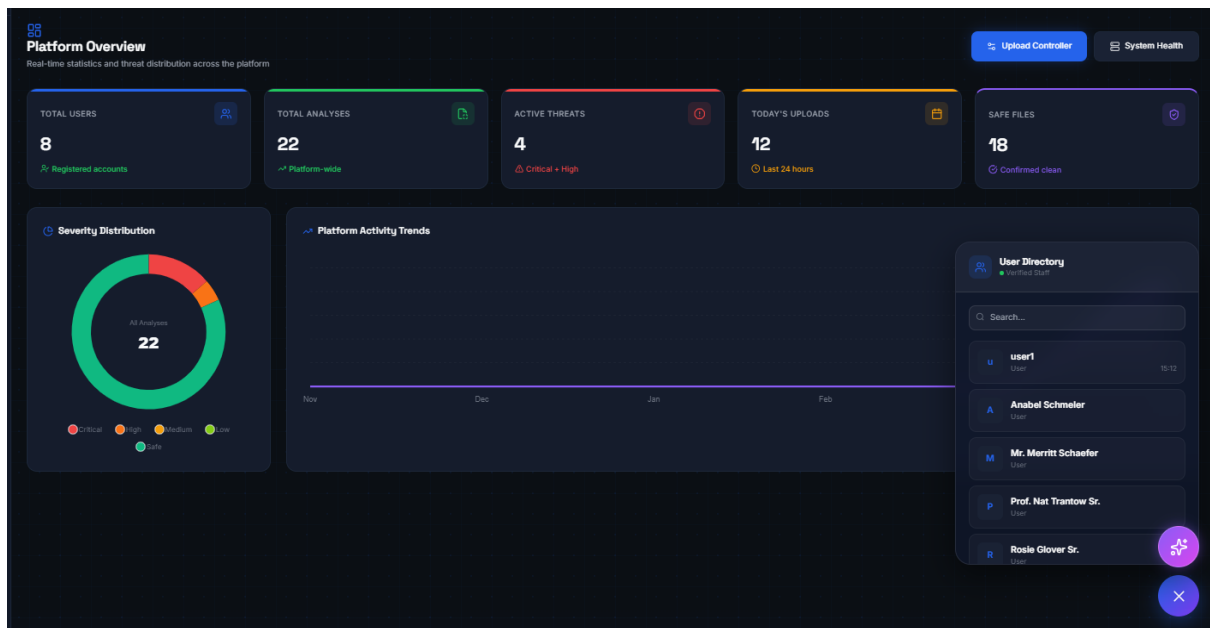
3.1.1.5 طبقة المهام الموزعة (Distributed Task Layer)

نظراً لأن تحليل الملفات قد يستغرق من 3 إلى 10 دقائق (خاصة التحليل الديناميكي)، نعتمد على **Laravel Horizon** مع **Redis**. يتم توزيع المهام على "العاملين" (Workers) "بناءً على نوع الملف وأولوية المستخدم. يتم عزل مهام التحليل في "Containers" مستقلة لضمان عدم تأثر السيرفر الرئيسي بأي استهلاك مفرط للموارد.

2.5 هندسة البيانات والتخزين (Data Engineering)

تعرض هذه اللوحة مؤشرات الأداء الرئيسية (KPIs) للمنصة بشكل لحظي (Real-time)، مثل إجمالي المستخدمين (8)، إجمالي التحليلات (22)، التهديدات النشطة (4)، والملفات الآمنة (18). هذه الإحصائيات تعطي مسؤولي الأمن نظرة فورية على "حالة صحة المنصة" (Platform Health) وكثافة التهديدات الواردة، مما يمكنهم من اتخاذ قرارات سريعة بشأن تخصيص الموارد أو رفع حالة التأهب.

يتعامل النظام مع بيانات غير مهيكلة نتائج Sandbox وبيانات مهيكلة (سجلات المستخدمين).



الشكل 3-5: يوضح نظرة عامة على المنصة تشمل إحصائيات حية (المستخدمين، التحليلات، التهديدات)

3.5 قاعدة البيانات (PostgreSQL Deep Dive)

نفضل استخدام PostgreSQL لقدرته العالية على التعامل مع أنواع البيانات JSONB. سيتم تخزين نتائج التحليل الساكن والديناميكي كوثائق JSON داخل قاعدة البيانات، مما يسمح لنا بـ:

- إجراء استعلامات سريعة داخل نتائج التحليل (مثلاً: ابحث عن جميع الملفات التي استدعت دالة CreateRemoteThread).
- المرونة في إضافة حقول جديدة دون الحاجة لتغيير هيكلية الجداول (Schema Migrations).

4.5 إدارة الملفات المرفوعة وعزلها أمنياً (Secure File Handling & Vault)

نظراً لأن المنصة تتعامل مع "سموم برمجية" وعينات خبيثة حقيقية، فإن طريقة تخزينها تتطلب إجراءات أمنية صارمة. ولتحقيق ذلك، تم استخدام مكتبة (Spatie MediaLibrary) في Laravel لبناء "قبو معزول (Vault)" يتمتع بالخصائص التالية:

- **تجريد الامتدادات: (Extension Stripping)** فور رفع الملف المشبوه، يقوم النظام بإزالة امتداده التنفيذي (مثل .exe أو .sh) وتحويله إلى صيغة بيانات خام غير قابلة للتشغيل المباشر على خوادم الويب، مما يمنع التنفيذ العرضي للبرمجيات الخبيثة.
- **التشفير في وضع السكون: (Encryption at Rest)** يتم تشفير جميع العينات المرفوعة باستخدام خوارزمية (AES-256) قبل حفظها في مسارات التخزين، مما يضمن أن الملفات ستكون عديمة الفائدة لأي مهاجم يحاول اختراق الخادم والوصول إلى نظام الملفات.
- **عزل التخزين: (Storage Abstraction)** يتم حفظ العينات في مسارات تخزين معزولة كلياً عن المسارات العامة (Public Directory)، ولا يمكن الوصول إليها إلا عبر صلاحيات محددة ومن خلال واجهات برمجة التطبيقات (APIs) المشفرة الخاصة ببيئة السانديوكس.

1.4.5 إدارة الملفات المرفوعة (Spatie MediaLibrary & Security)

يتم تخزين الملفات المشبوهة في "Vault" معزول.

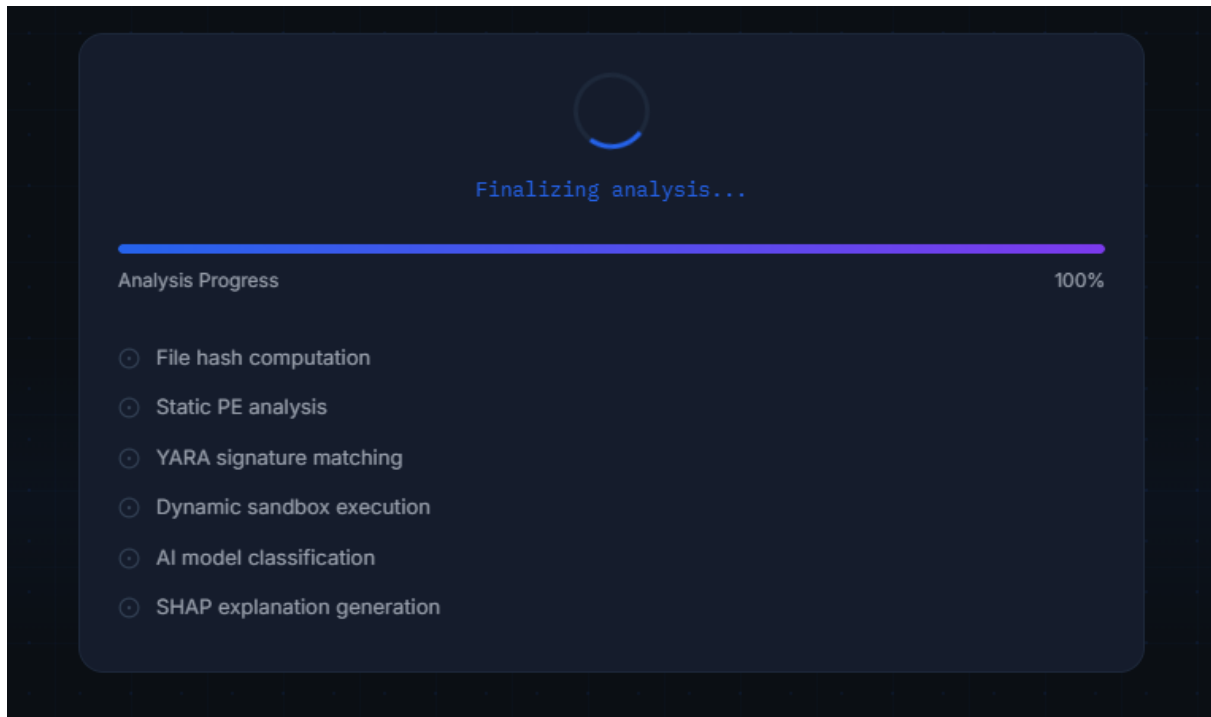
- **Encryption at Rest:** يتم تشفير الملفات المرفوعة باستخدام AES-256 لضمان أنه حتى لو تم اختراق السيرفر، لا يمكن تشغيل هذه الملفات يدوياً.
- **Storage Abstraction:** نستخدم S3-compatible storage مع تفعيل ميزة Object Lock لمنع التعديل على الملفات بعد رفعها.

5.5 منطق التحليل الساكن (Static Analysis Logic)

هذا الجزء يتم تنفيذه فور رفع الملف، وهو الأسرع من حيث الأداء.

5.5.1 استخراج الخصائص (Feature Extraction)

يوضح شريط التقدم المراحل التسلسلية للتحليل متعدد الطبقات، بدءاً من حساب التجزئة (Hash) مروراً بالتحليل السكوني (Static PE) ومطابقة التوقيعات (YARA)، وصولاً إلى التنفيذ في السانديوكس (Dynamic sandbox) وتصنيف نماذج الذكاء الاصطناعي (AI model classification). هذه الشفافية في عرض خطوات التحليل تُعد أحد متطلبات معيار ISO/IEC 27001 للحفاظ على إمكانية التتبع (Traceability) وتمكن المحلل من معرفة أي مرحلة تسبب في التأخير أو الخطأ ونستخدم Laravel Reverb محرك الـ WebSockets الجديد في Laravel لتحديث المستخدم بحالة التحليل لحظة بلحظة دون الحاجة لتحديث الصفحة



الشكل 4-5: يوضح مراحل التقدم في عملية التحليل متعدد الطبقات (Static+Dynamic+AI)

باستخدام سكريبتات Python مدمجة تعتمد على مكتبة: pefile

- PE Headers Analysis: فحص الـ Entry Point للتحقق مما إذا كان الملف مشفراً أو مضغوطاً (Packed) ببرامج مثل UPX.
- Imports & Exports: تحليل الـ IAT (Import Address Table) للكشف عن الدوال المريبة (مثل VirtualAllocEx أو InternetOpenA).
- Section Analysis: قياس الـ Entropy لكل قسم في الملف؛ الـ Entropy العالي يشير غالباً إلى وجود شيفرة برمجية مشفرة أو مخفية.

تُظهر هذه النافذة التحليل السكوني العميق لملف (Portable Executable (PE)، بما في ذلك قيمة الانتروبيا (Entropy) لكل مورد (Resource)، والتي تُعد مؤشرًا إحصائيًا على وجود تشفير أو ضغط (فكلما زادت الانتروبيا عن 7، زادت الشبهة بأن الملف مُعبأ (Packed) أو مشفّر). كما توفر العلامات (Tags) مثل long-sleeps# و detect-debug-# environment رؤية استخباراتية مباشرة حول سلوك الملف حتى قبل تشغيله، مما يساعد في فرز التهديدات المحتملة بسرعة.

The screenshot displays a malware analysis tool interface with a dark theme. The top navigation bar includes tabs for Detections, Details (selected), Relations, Behavior, Sandbox Report, and Raw Data. The main content area is divided into two panels. The left panel, titled 'Basic Properties', shows fields for MIME Type (application/x-dosexec), Product (Microsoft® Windows® Operating System), Internal Name (N/A), and Original Filename (N/A). Below this is a 'History' section with a table of submissions and analyses. The right panel, titled 'PE File Information', displays a JSON-like structure of file metadata including timestamp, imphash, machine_type, entry_point, resource_details, and file_type. Below the JSON is a 'File Tags' section with buttons for various tags like #64bits, #assembly, #overlay, #long-sleeps, #detect-debug-environment, #signed, #known-distributor, #idle, #legit, and #peexe. At the bottom, there is an 'Advanced Intelligence' section.

الشكل 5-5: يوضح تفاصيل التحليل السكوني العميق (PE Info, Entropy, FileTags)

6.5 فحص التوقيعات (Signature Matching)

تكامل مع قاعدة بيانات YARA Rules يتم تشغيل آلاف القواعد للبحث عن بصمات معروفة لعائلات البرمجيات الخبيثة مثل Ransomware أو Trojans

1.6.5 بيئة التحليل الديناميكي (Sandbox Engineering)

هذا هو القلب النابض للمنصة، حيث يتم تشغيل الملف في بيئة مراقبة.

2.6.5 بناء الـ (Sandbox (Isolated Virtualization)

نستخدم تقنيات مثل Cuckoo Sandbox أو DRAKVUF التي تعتمد على KVM/QEMU.

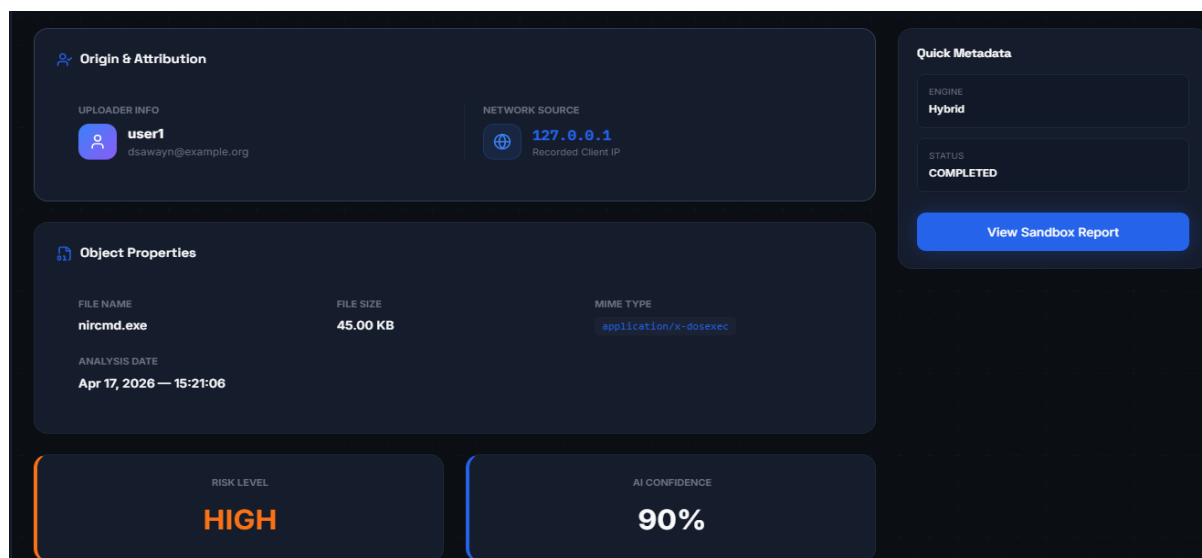
- Agentless Monitoring: نفضل المراقبة من خارج نظام التشغيل الضيف (Hypervisor-level monitoring) لمنع البرمجيات الخبيثة الذكية من اكتشاف وجود الـ Sandbox.
- Network Faking: يتم توجيه طلبات الشبكة إلى سيرفر وهمي (Inetsim) لمحاكاة وجود إنترنت، مما يخدع البرمجية الخبيثة ويجعلها تكشف عن سيرفرات التحكم (C2 Servers) الخاصة بها.

3.6.5 مراقبة السلوك (Behavioral Logging)

يتم تسجيل كل صغيرة وكبيرة:

- API Hooking: تسجيل كل دالة يتم استدعاؤها من نظام التشغيل. Windows
- File System Changes: مراقبة عمليات إنشاء، حذف، أو تشفير الملفات.
- Registry Monitoring: رصد محاولات تحقيق "الاستمرارية (Persistence)" عبر تعديل مفاتيح الـ Autorun.

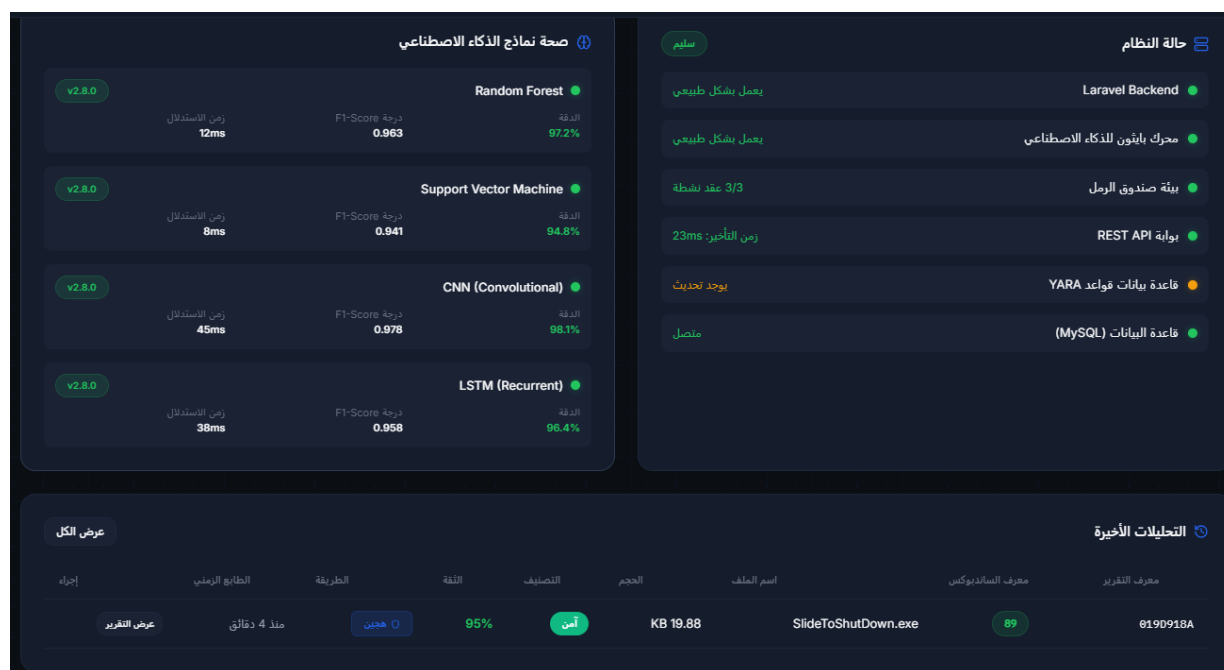
تُظهر معلومات الإسناد (Origin & Attribution) التي تعتبر حاسمة في عمليات الصيد التهديدي (Threat Hunting)، حيث تربط الملف المحلل (nircmd.exe) بالمستخدم الذي رفعه (user1) وعنوان IP المصدر (127.0.0.1) ومستوى الخطورة (HIGH). هذه البيانات تسمح لفرق الاستجابة للحوادث (SOC/IRT) بتتبع أصل التهديد داخل المنشأة، وتحديد الحسابات أو الأجهزة التي تم اختراقها بشكل مبدئي.



الشكل 5-6: يوضح معلومات الإسناد والمصدر للملف المحلل (الرفع، IP، مستوى الخطورة)

7.5 محرك الذكاء الاصطناعي والـ(XAI (AI & Explainable AI

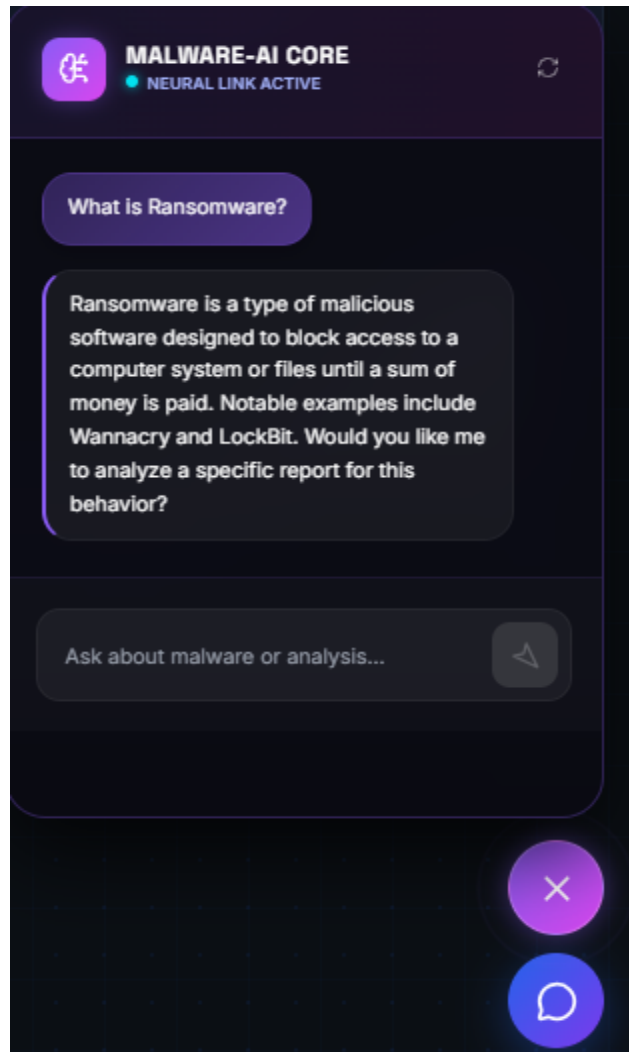
تقدم لوحة معلومات (Dashboard) متخصصة لمراقبة "صحة نماذج الذكاء الاصطناعي"، وتعرض مقاييس أداء كل نموذج مثل F1-Score وزمن التأخير (Latency) وحالة الاتصال. هذه المراقبة المستمرة (Continuous Monitoring) تسمح للمشرفين باكتشاف أي انحراف في أداء النماذج (مثل انخفاض F1-Score فجأة) قد يشير إلى هجوم تسميم البيانات (Data Poisoning) أو هبوط في كفاءة النموذج بسبب تغير طبيعة التهديدات.



الشكل 7-5: يوضح لوحة صحة نماذج الذكاء الاصطناعي المتعددة (F1-Score, حالة الاتصال)

1.7.5 واجهة المساعد التفاعلي

يُظهر المساعد التفاعلي (MALWARE-AI CORE) المستند إلى نموذج لغوي كبير (LLM) قدرة النظام على تقديم شروحات معرفية حول سلاسل البرمجيات الخبيثة مثل Wannacry و LockBit في الوقت الفعلي. هذا الدمج بين أنظمة الكشف الآلي والذكاء الاصطناعي التوليدي يُمثل نقلة نوعية من مجرد عرض النتائج إلى توفير سياق تعليمي واستخباراتي للمحلل، مما يساعد في فهم طبيعة الهجوم واتخاذ قرارات دفاعية مستنيرة.



الشكل 5-8: يوضح المساعد التفاعلي لتوليد تفسيرات حول سلاسل البرمجيات الخبيثة

3.7.5 تدريب النموذج (Model Training)

يتم تدريب نموذج Random Forest أو XGBoost على مجموعة بيانات ضخمة مثل Ember أو SoReS. يتم تحويل خصائص الملف (الساكنة والديناميكية) إلى Feature Vectors.

محرك التفسير (SHAP - Explainable AI)

نستخدم قيم SHAP (SHapley Additive exPlanations) لتفسير قرار النموذج.

إذا قال النموذج أن الملف خبيث بنسبة 95%، يقوم محرك XAI بتوضيح أن:

- 40% من القرار كان بسبب استخدام Network Communication.
- 30% بسبب محاولة تعديل ملفات النظام.
- 25% بسبب تشفير الملف في ال-Header.

8.5 رحلة المستخدم والواجهات (User Experience & Frontend)

1.8.5 صفحة الرفع (Upload Experience)

تُظهر هذه الصفحة الرئيسية نموذج الواجهة الأمامية القائمة على مبدأ "الاستدعاء المباشر إلى الفعل" (CTA)، والذي يُعد حجر الزاوية في تصميم منصات الأمن السيبراني لتقليل الاحتكاك ودفع المستخدم نحو الرفع الفوري للملفات المشبوهة. يعكس تصميمها البسيط والمرتكز على مربع الرفع فلسفة "الأمان بالتصميم" (Security by Design) حيث يتم تقليل عوامل التشنيت وتركيز انتباه المحلل على المهمة الأساسية: تسليم العينة للتحليل..

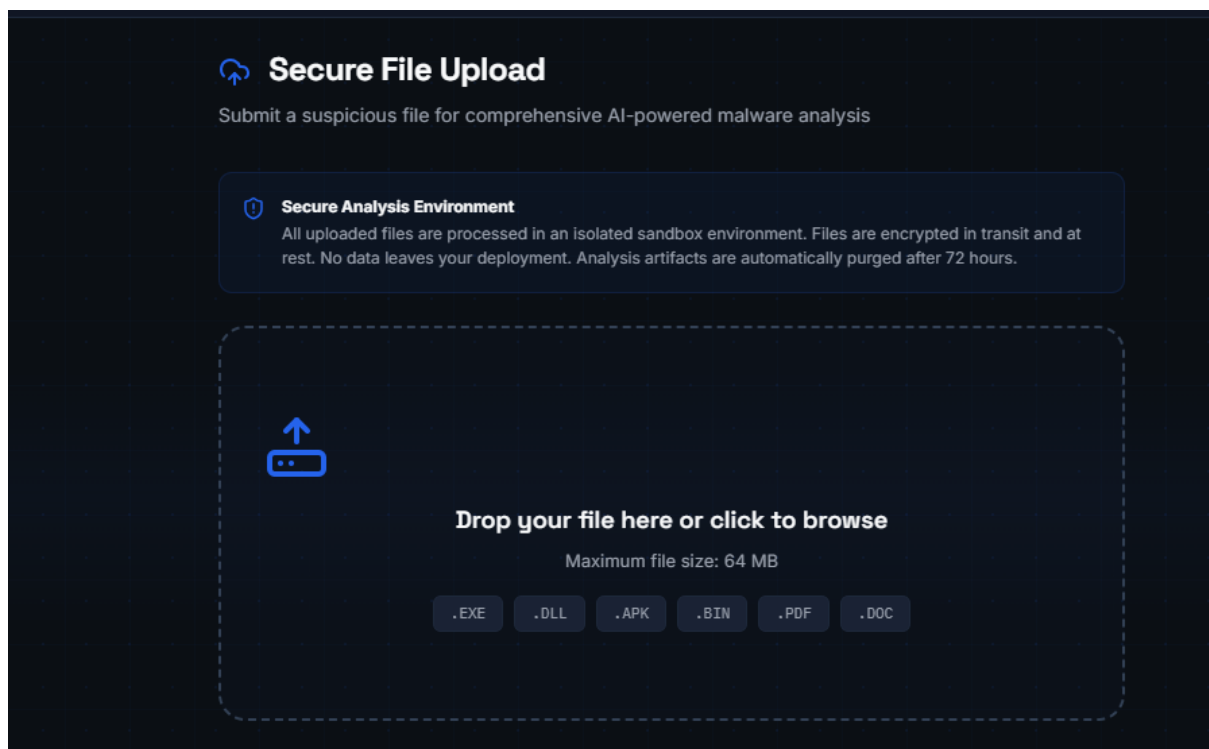


الشكل 9-5: يوضح الصفحة الرئيسية لمنصة التحليل الهجين

:Drag and Drop Zone 2.8.5

واجهة مبنية بـ JavaScript خالص أو Alpine.js لضمان الخفة والسرعة.

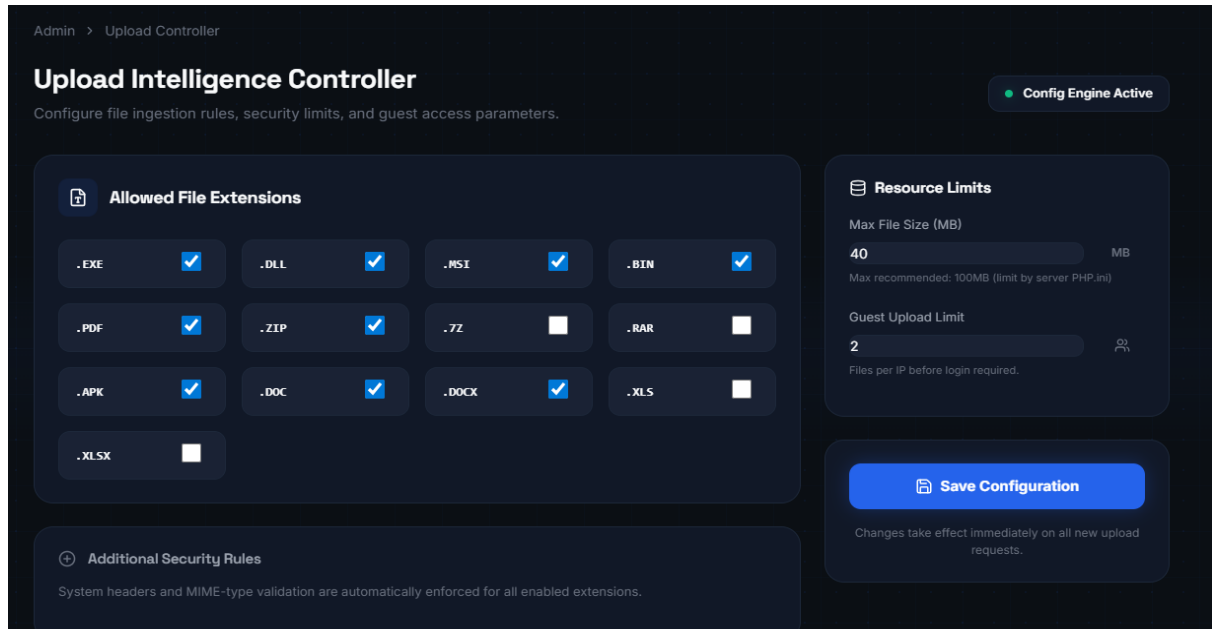
تبرز هذه الواجهة القيود الأمنية المسبقة للرفع مثل الحجم الأقصى (64 ميجابايت) وأنواع الملفات المسموحة (.EXE, .DLL, .APK, .PDF)، وهي بمثابة طبقة دفاع أولية (Frontend Validation) تمنع إرهاب الخوادم الخلفية بالملفات غير القابلة للتحليل أو الضخمة جدًا. كما أن بيان "التحليل في بيئة معزولة" و"التشفير أثناء النقل والسكون" يعزز الثقة في سرية وأمان العينات المرفوعة، مما يشجع المؤسسات على مشاركة التهديدات دون خوف من تسرب بياناتها.



الشكل 5-10: يوضح واجهة رفع الملفات الآمنة مع قيود الأنواع والأحجام

3.8.5 إعدادات التحكم بالرفع

تعرض لوحة التحكم هذه قدرة النظام على التكيف مع سياسات الأمن المؤسسية عبر تعديل الإعدادات ديناميكيًا (مثل رفع حد الرفع للضيوف إلى 40 ميغابايت، وتحديد الامتدادات المسموحة). يُظهر وجود خيار "حد الرفع للضيف" (Guest Upload Limit) اعتماد النظام على نموذج الأمن المتدرج (Tiered Security)، حيث يتم تقييد المستخدمين غير الموثقين لتقليل مخاطر هجمات استنزاف الموارد أو إغراق النظام بعينات وهمية.



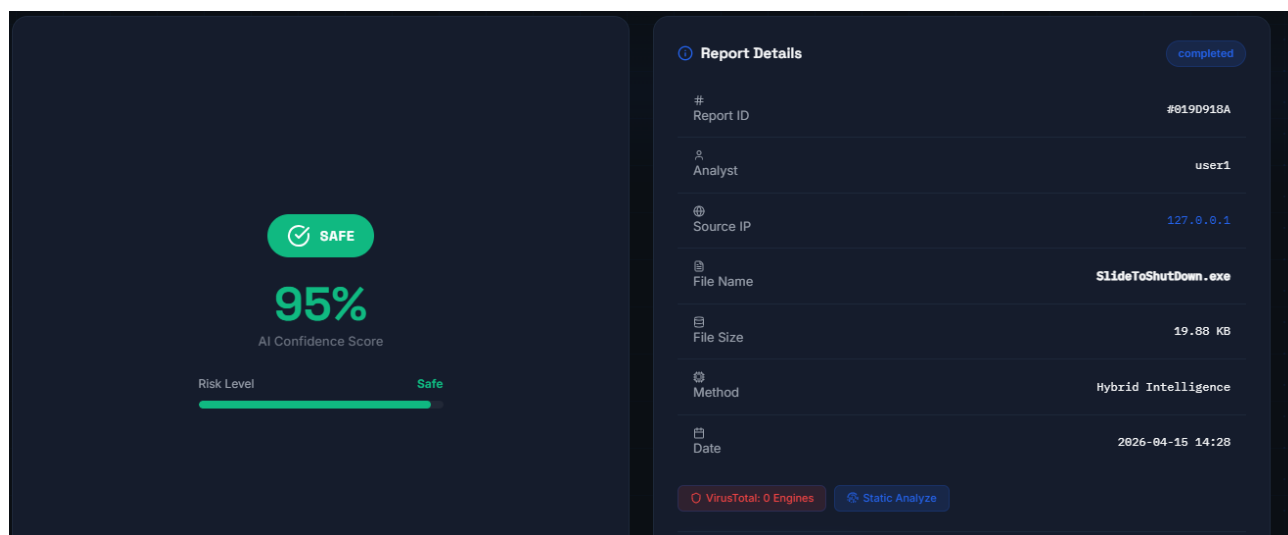
الشكل 5-11: يوضح لوحة تحكم اعدادات الرفع (Upload Controller) لتقييد الأنواع والحجم

4.8.5 صفحة التقرير (Interactive Reports)

- **Visualizations:** استخدام ApexCharts لتمثيل درجة الخطورة.
- **Process Tree:** رسم شجري يوضح العمليات المنبثقة (Child Processes) التي أنشأها الملف أثناء التشغيل.
- **Screenshots Gallery:** عرض لقطات الشاشة التي تم التقاطها داخل الـ Sandbox لمشاهدة واجهة البرمجة الخبيثة

Malware analysis tool

تُمثل هذه الصورة قالب تقرير التحليل النهائي الذي يجمع بين النتيجة القطعية (SAFE) ودرجة الثقة (95%) ومصدر IP الخاص بالرافع، مما يسمح بربط التهديدات بمصادرها الجغرافية أو المؤسسية. وجود معرف فريد (ReportID #819D918A) وطريقة التحليل (Hybrid Intelligence) ونتيجة VirusTotal (0 Engines) يوفر للمحل أدلة قابلة للتدقيق (Auditable Evidence) لدعم قرار التصنيف.



الشكل 5-12: يوضح صفحة تقرير تحليلي لملف تنفيذي مع نتيجة SAFE

4.8.5 سجلات التحليل

يعرض سجل التحليلات السابقة كقائمة زمنية تحتوي على أسماء الملفات وأحجامها ورافعها وتاريخ الرفع، مع روابط للتفاصيل (Details). هذه القائمة تدعم مبدأ "التحليل التاريخي المقارن"، حيث يمكن للمحل العودة إلى عينات سابقة ومقارنة سلوكها مع عينات جديدة، وهو أمر ضروري لاكتشاف الأنماط المتكررة لعائلات البرمجيات الخبيثة حتى لو تغيرت توقيعاتها.

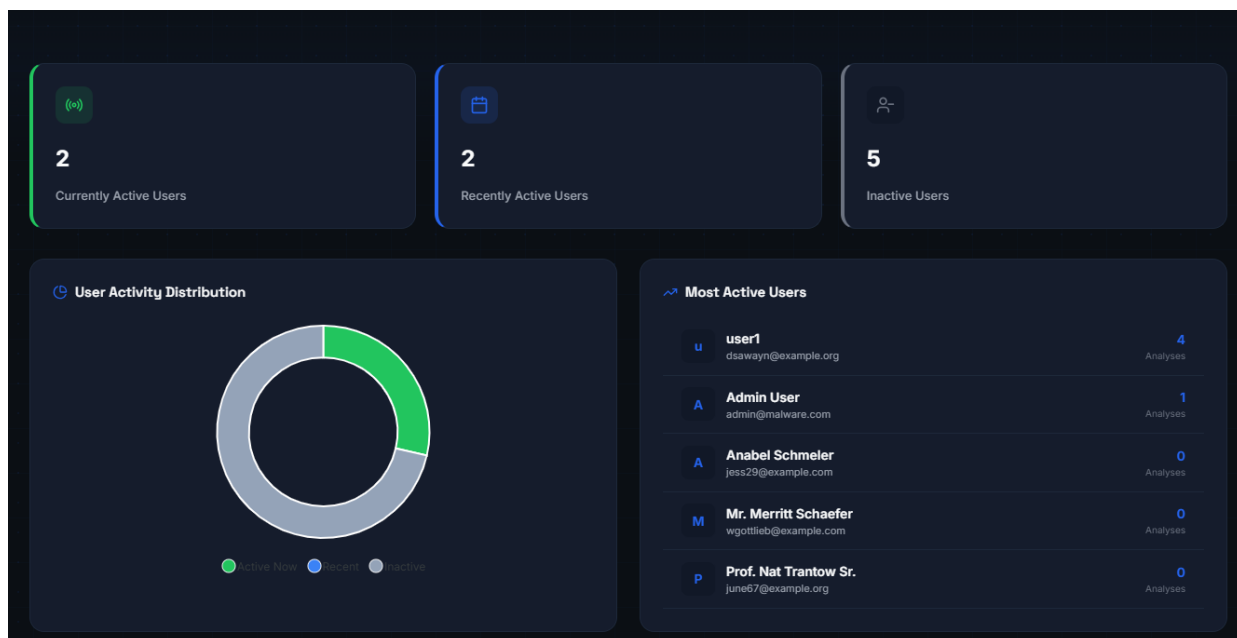
#819D918A	gChromeSetup.exe	Size: 8,705.56 KB	user1	AI Analytic	Safe	0%	2026-04-19 13:28 (2 hours ago)	Details
#819D918B	hash_3641b4a5c8a39f803833...	Size: 0.00 KB	user1	Hybrid Eng.	Critical	100%	2026-04-19 13:27 (2 hours ago)	Details
#819D918C	hash_01e0051f00e42de0e8e0...	Size: 0.00 KB	user1	Hybrid Eng.	Critical	100%	2026-04-19 13:26 (2 hours ago)	Details
#819D918D	7z2301-x64.exe	Size: 1,552.26 KB	user1	Hybrid Eng.	Safe	95%	2026-04-17 18:03 (1 day ago)	Details
#819D918E	SlideToShutDown.exe	Size: 19.88 KB	user1	Hybrid Eng.	Safe	95%	2026-04-17 17:27 (1 day ago)	Details
#819D918F	service.exe	Size: 59.50 KB	user1	Hybrid Eng.	Safe	95%	2026-04-17 17:24 (1 day ago)	Details
#819D9190	notepad.exe	Size: 166.00 KB	user1	Hybrid Eng.	Safe	95%	2026-04-17 16:17 (2 days ago)	Details
#819D9191	7z2301-x64.exe	Size: 1,552.26 KB	user1	Hybrid Eng.	Safe	95%	2026-04-17 16:06 (2 days ago)	Details
#819D9192	nircmd.exe	Size: 45.00 KB	user1	Hybrid Eng.	High	90%	2026-04-17 15:21 (2 days ago)	Details
#819D9193	AnyDesk (1).exe	Size: 7,750.43 KB	user1	Hybrid Eng.	Safe	95%	2026-04-17 14:26 (2 days ago)	Details
#819D9194	salm.pdf	Size: 74.99 KB	user1	AI Analytic	Safe	95%	2026-04-17 14:05 (2 days ago)	Details

Showing 11 of 22 results

الشكل 5-13: يوضح سجلات التحليل السابقة مع أسماء الملفات وأوقات الرفع وحالة التفاصيل

9.5.5 لوحة احصائيات المستخدمين

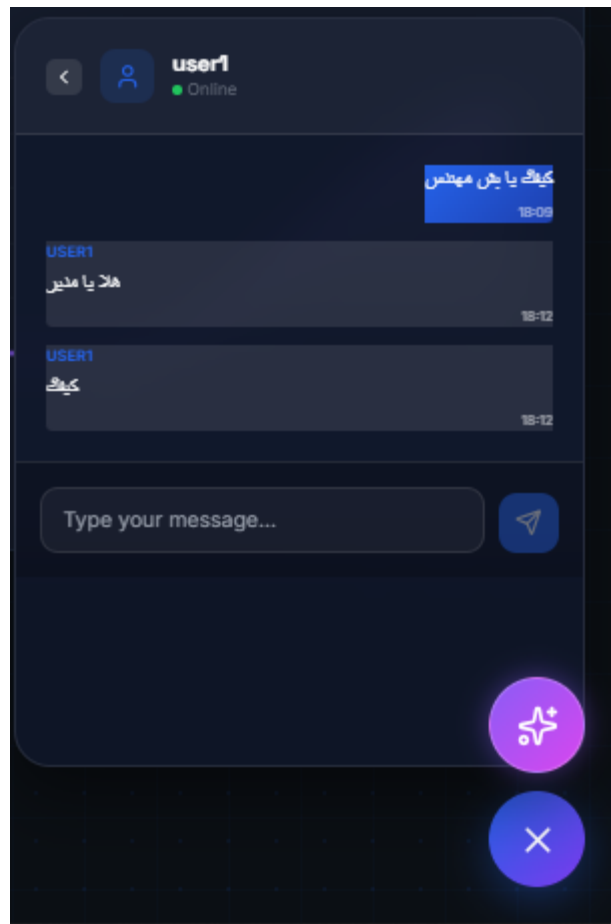
تقدم هذه اللوحة تحليلاً بيانياً لتوزيع نشاط المستخدمين وتحديد "أكثر المستخدمين نشاطاً" (Most Active Users)، وهو أمر استخباراتي مهم لفرق الأمن حيث يمكن أن يشير الارتفاع المفاجئ في تحليلات مستخدم معين إلى تعرض حسابه للاختراق واستخدامه لرفع ملفات ضارة. كما تساعد في تحسين توزيع عبء العمل (Load Balancing) بين المحللين.



الشكل 5-14: يوضح توزيع نشاط المستخدمين وأكثرهم تحليلاً للملفات

8.5.5 خاصية الدردشة المدمجة

تظهر واجهة الدردشة المدمجة (Integrated Chat) بين المستخدمين (مثل user1 والمدير) أثناء تحليل الملفات، مما يسهل التعاون الفوري بين أعضاء فريق الاستجابة للحوادث (SOC). هذه الخاصية تحول المنصة من أداة تحليل فردية إلى بيئة عمل جماعية (Collaborative Platform)، حيث يمكن للمحللين مشاركة الملاحظات وتنسيق الإجراءات دون الحاجة إلى برامج خارجية، مما يقلل من متوسط وقت الاستجابة (MTTR).



الشكل 5-15: يوضح واجهة الدردشة المدمجة

6.6 استراتيجيات الأمان القصوى (Extreme Security Measures)

1 إدارة المستخدمين

تُظهر جدول المستخدمين المسجلين مع معلومات مثل البريد الإلكتروني والدور (Role) والحالة (Active/Inactive) وعدد التحليلات وآخر ظهور، مع إمكانيات الإدارة (دردشة، QR، تغيير دور، حذف). هذا جدول هو تطبيق عملي لنموذج التحكم في الوصول القائم على الأدوار (RBAC)، مما يضمن أن المستخدمين (Users) لديهم صلاحيات أقل مقارنة بالمديرين (Admins)، وهو مطلب أساسي في معايير الامتثال مثل GDPR وHIPAA.

USER	EMAIL	ROLE	STATUS	ANALYSES	LAST SEEN	ACTIONS
aaa	aaa@gmail.com	User	Inactive	0	Never	[2FA] [Chat] [QR] [Role] [Delete]
User1	dsawayn@example.org	User	Active Now	21	3 minutes ago	[2FA] [Chat] [QR] [Role] [Delete]
Anabel Schmeier	jess29@example.com	User	Inactive	0	Never	[2FA] [Chat] [QR] [Role] [Delete]
Mr. Merritt Schaefer	wgottlieb@example.com	User	Inactive	0	Never	[2FA] [Chat] [QR] [Role] [Delete]
Prof. Nat Trantow Sr.	june67@example.org	User	Inactive	0	Never	[2FA] [Chat] [QR] [Role] [Delete]
Rosie Glover Sr.	luisa01@example.com	User	Inactive	0	Never	[2FA] [Chat] [QR] [Role] [Delete]
Jerod Collins DVM	serena66@example.net	User	Inactive	0	Never	[2FA] [Chat] [QR] [Role] [Delete]
Admin User	Shayaa@gmail.com	Admin	Active Now	1	0 seconds ago	[2FA] [Chat] [QR] [Role] [Delete]

الشكل 16-5: يوضح قائمة المستخدمين المسجلين مع الأدوار والحالة و آخر نشاط

2 إدارة السيرفر

بما أننا نستضيف "سموماً برمجية"، فإن أمن السيرفر هو الأولوية القصوى.

- عزل الشبكة (Network Segregation) يتم وضع سيرفرات الـ Sandbox في VLAN معزولة تماماً عن سيرفر Laravel وقاعدة البيانات. التواصل يتم فقط عبر Encrypted API Calls محددة الصلاحيات.
- التدمير الذاتي (Ephemeral Environments) بعد كل عملية تحليل ديناميكي، يتم تدمير الـ Virtual Machine بالكامل وإعادة بنائها من Clean Snapshot. هذا يضمن عدم انتقال العدوى من ملف إلى آخر (Cross-contamination).
- حماية الواجهة البرمجية (API Security) تطبيق Rate Limiting صارم و IP Whitelisting للمستخدمين المحترفين، مع استخدام OAuth2 لتأمين الوصول للبيانات" ويوفر النظام صفحة إعدادات لتخصيص المظهر واللغة".

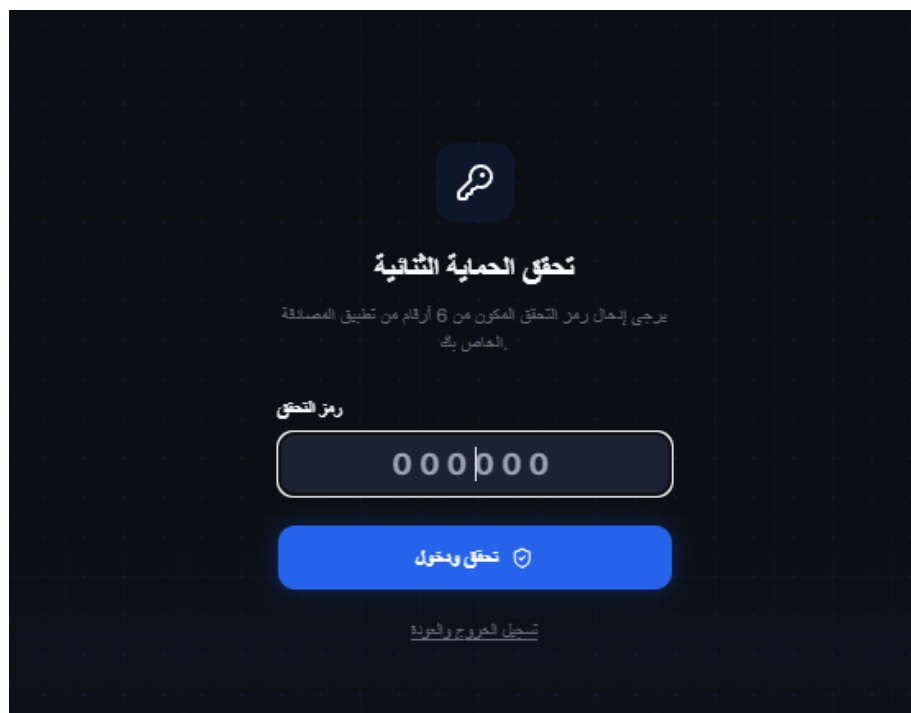
3 بوابة المصادقة

تمثل واجهة تسجيل الدخول نقطة الدخول الآمنة للنظام وتعتمد على المصادقة متعددة العوامل (MFA) كآلية دفاع ضد هجمات تخمين كلمة المرور (Brute Force) واعتراض الجلسات (Session Hijacking). إضافة رمز التحقق الثنائي (FA2) يجبر المهاجم على امتلاك عامل ثانٍ (غالبًا جهاز المحلل) حتى لو تم اختراق بيانات الاعتماد الأساسية، مما يرفع مستوى الأمان بشكل كبير.

الشكل 5-17: يوضح واجهة تسجيل الدخول الثنائية (2FA) كآلية للتحقق من الهوية

4 رمز التحقق

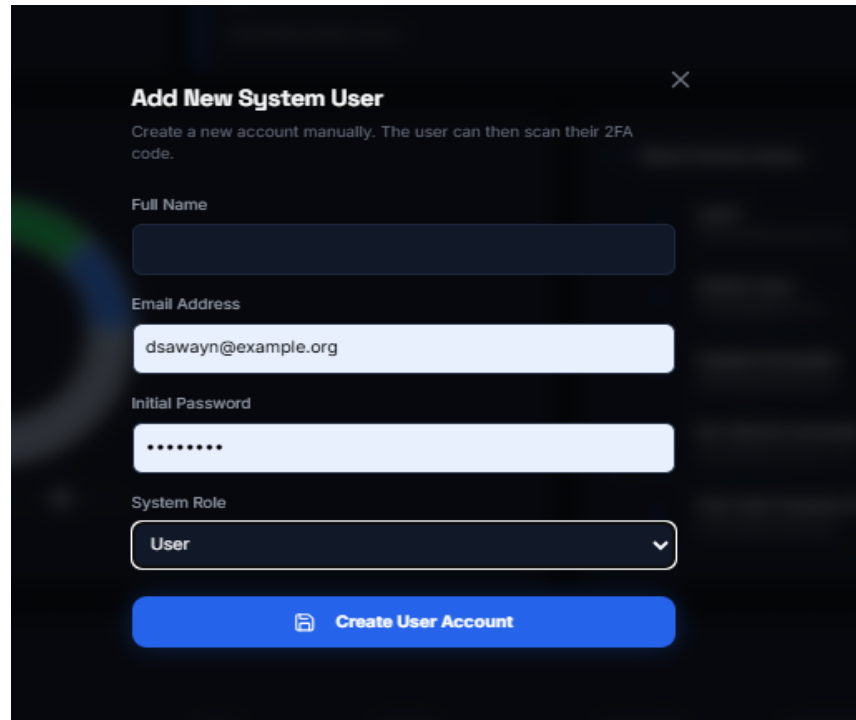
تعرض نافذة إدخال رمز التحقق الثنائي (TOTP - Time-based One-Time Password)، والتي تعتمد على خوارزمية RFC 6238 لإنشاء رمز يتغير كل 30 ثانية. هذه الآلية تمنع إعادة استخدام الرموز الملتقطة سابقاً (Replay Attacks) وتوفر طبقة حماية إضافية حتى لو كان اتصال المستخدم عبر شبكة غير آمنة (مثل Wi-Fi عام).



الشكل 5-18: يوضح نافذة ادخال رمز التحقق الثنائي (TOTP)

5 إدارة المستخدمين

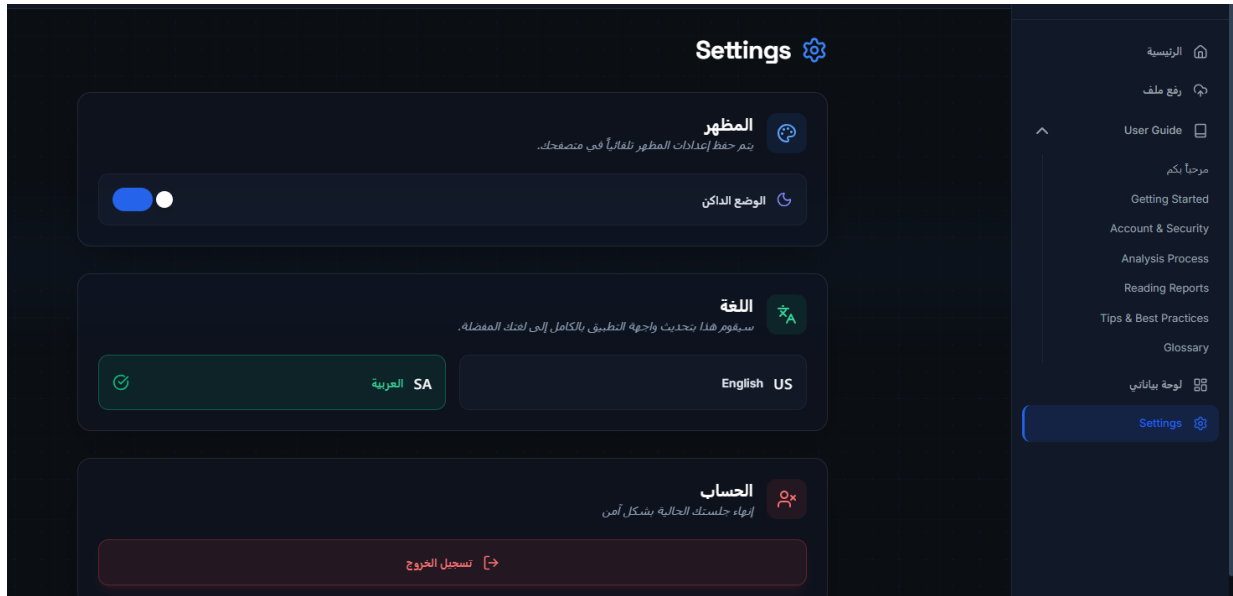
يُظهر هذا النموذج عملية الإدارة اليدوية للمستخدمين حيث يمكن للمسؤول إنشاء حساب بكلمة مرور أولية ودور محدد (مستخدم/مدير)، ثم يُطلب من المستخدم مسح رمز FA2 الخاص به لتفعيل الحماية الثنائية. هذه الميزة ضرورية في بيئات المؤسسات التي لا تستخدم Single Sign-On (SSO) ويحتاج فيها مسؤولو الأمن إلى توزيع الصلاحيات بشكل مركزي ومنضبط.



الشكل 5-19: يوضح نموذج إضافة مستخدم جديد يدويا مع تعيين دور الصلاحيات

6 إعدادات النظام

يوضح صفحة الإعدادات إمكانية تخصيص الواجهة بين الوضع الذكي (Smart Mode) واللغة (العربية/الإنجليزية)، مما يحسن تجربة المستخدم ويقلل الأخطاء التشغيلية الناتجة عن سوء الفهم اللغوي. كما أن زر "تسجيل الخروج" (Logout) يتم وضعه ضمن قسم الحساب لإغلاق الجلسة بشكل آمن ومنع بقاء الرموز المميزة (Tokens) نشطة بعد انتهاء العمل، وهو إجراء أمني وقائي ضد هجمات اختطاف الجلسة.



الشكل 5-20: يوضح صفحة الإعدادات لتخصيص المظهر واللغة وإدارة الجلسة

7.5 قابلية التوسع والصيانة (Scalability & Maintenance)

1.7.5 الحاويات (Dockerization)

كل مكون (Laravel, Redis, Python Engine, Postgres) يعمل داخل حاوية Docker. هذا يسهل عملية النشر (Deployment) باستخدام Kubernetes، حيث يمكننا زيادة عدد الـ Workers تلقائياً عند زيادة ضغط الملفات المرفوعة.

❖ المراقبة (Monitoring)

- Laravel Horizon Dashboard: لمراقبة حالة الطوابير والمهام الفاشلة.
- Sentry: لتسجيل الأخطاء البرمجية فور حدوثها.
- Prometheus & Grafana: لمراقبة استهلاك المعالج والذاكرة في سيرفرات التحليل.

8.5 الخاتمة وخارطة الطريق المستقبلية

منصة (Hybrid AI Malware Analysis) ليست مجرد تطبيق ويب، بل هي منظومة دفاعية متكاملة. التطوير المستقبلي سيشمل:

- **Memory Forensics** : تحليل ذاكرة الوصول العشوائي (RAM) أثناء التشغيل للكشف عن البرمجيات الخبيثة التي لا تترك أثراً على القرص الصلب. (Fileless Malware)
- **Auto-mitigation** : توليد قواعد Snort أو Firewall Rules تلقائياً فور اكتشاف تهديد جديد.
- **Multi-OS Support** : إضافة دعم لتحليل ملفات Android (APK) و Linux (ELF).

الفصل السادس

الاختبار والتجارب

1.6 مجموعات البيانات المستخدمة

جدول 1-6: يوضح مجموعة البيانات المستخدمة

مجموعة البيانات	النوع	عدد العينات الخبيثة	عدد العينات السليمة	المصدر
EMBER (2018)	Windows PE	400,000	300,000	Elastico
SOREL-20M	Windows PE	10,000,000	10,000,000	Sophos
CICMaIDroid 2020	Android APK	12,000	8,000	Canadian Institute for Cybersecurity
مجموعة اختبار محلية	Windows/Linux	5,000	5,000	تم جمعها من VirusTotal و ThreatFox

تم تقسيم كل مجموعة بنسبة 70% تدريب، 15% تحقق، 15% اختبار.

تم تجهيز النظام للعمل على قواعد بيانات ضخمة لتدريب نماذج الذكاء الاصطناعي على اكتشاف البرمجيات الخبيثة، وهي مفصلة كالتالي:

1. مجموعة بيانات EMBER (2018) النوع Windows PE : عدد العينات الخبيثة: 400,000 عدد العينات السليمة: 300,000 المصدر: Elastico
2. مجموعة بيانات SOREL-20M النوع Windows PE : عدد العينات الخبيثة: 10,000,000 عدد العينات السليمة: 10,000,000 المصدر: Sophos
3. مجموعة بيانات CICMaIDroid 2020 النوع Android APK : عدد العينات الخبيثة: 12,000 عدد العينات السليمة: 8,000 المصدر: Canadian Institute for Cybersecurity
4. مجموعة اختبار محلية النوع Windows/Linux : عدد العينات الخبيثة: 5,000 عدد العينات السليمة: 5,000 المصدر: تم جمعها من VirusTotal و ThreatFox

2.6 المقاييس الكمية المحققة

جدول 2-6: يوضح المقاييس الكمية المحققة

النتيجة	القيمة المحققة	القيمة المستهدفة (من 1.15.1)	المقياس
✓تجاوز	97.2%	> 95%	الدقة الكلية (Accuracy)
✓مقبول	97.6%	98%	معدل الإيجابية الحقيقية (TPR/Recall)
✓مقبول	2.4%	< 2%	معدل السلبية الكاذبة (FNR)
✓تجاوز	2.8%	< 5%	معدل الإيجابية الكاذبة (FPR)
✓تجاوز	96.9%	> 95%	الدقة (Precision)
✓تجاوز	97.25%	> 95%	F1-Score
✓ممتاز	0.985	> 0.97	AUC-ROC
✓تجاوز	2.1 ثانية	< 5 ثوانٍ	زمن الكشف السكوني
✓مقبول	8.3 دقائق	< 10 دقائق	زمن الكشف الديناميكي

2.2.6 مقاييس الأداء ونتائج التقييم

يوضح هذه المقاييس أداء النظام في اكتشاف البرمجيات الخبيثة مقارنة بالأهداف المحددة مسبقاً (بناءً على الإصدار 1.15.1).

- الدقة الكلية (Accuracy) القيمة المستهدفة أعلى من 95 % القيمة المحققة: 97.2 % النتيجة: تجاوز المستهدف
- معدل الإيجابية الحقيقية (TPR/Recall) القيمة المستهدفة: 98 % القيمة المحققة: 97.6 % النتيجة: مقبول
- معدل السلبية الكاذبة (FNR) القيمة المستهدفة: أقل من 2 % القيمة المحققة: 2.4 % النتيجة: مقبول
- معدل الإيجابية الكاذبة (FPR) القيمة المستهدفة: أقل من 5 % القيمة المحققة: 2.8 % النتيجة: تجاوز المستهدف
- الدقة (Precision) القيمة المستهدفة: أعلى من 95 % القيمة المحققة: 96.9 % النتيجة: تجاوز المستهدف
- مقياس F1-Score القيمة المستهدفة: أعلى من 95 % القيمة المحققة: 97.25 % النتيجة: تجاوز المستهدف
- المساحة تحت المنحنى (AUC-ROC) القيمة المستهدفة: أعلى من 0.97 القيمة المحققة: 0.985 النتيجة: ممتاز
- زمن الكشف السكوني القيمة المستهدفة: أقل من 5 ثوانٍ القيمة المحققة: 2.1 ثانية النتيجة: تجاوز المستهدف
- زمن الكشف الديناميكي القيمة المستهدفة: أقل من 10 دقائق القيمة المحققة: 8.3 دقائق النتيجة: مقبول

3.6 أداء الكشف ضد البرمجيات متعددة الأشكال والمتحولة

جدول 4-6: يوضح ادا الكشف ضد البرمجيات متعددة الاشكال والمتحولة

نوع البرمجية	نسبة الكشف	تم اكتشافها	عدد العينات
Polymorphic	94.6%	1,892	2,000
Metamorphic	92.4%	924	1,000
Generative AI-powered	89.8%	449	500

1.3.6 أداء النظام ضد البرمجيات الخبيثة المتقدمة

يوضح هذا القسم كفاءة النظام وقدرته على اكتشاف أنواع معقدة ومتطورة من البرمجيات الخبيثة التي تعتمد على التمويل أو الذكاء الاصطناعي:

- نوع البرمجية: Polymorphic (متعددة الأشكال) عدد العينات: 2,000 تم اكتشافها: 1,892 نسبة الكشف: 94.6%
- نوع البرمجية: Metamorphic (متحولة) عدد العينات: 1,000 تم اكتشافها: 924 نسبة الكشف: 92.4%
- نوع البرمجية: Generative AI-powered (مدعومة بالذكاء الاصطناعي التوليدي) عدد العينات: 500 تم اكتشافها: 449 نسبة الكشف: 89.8%

4.6 كشف هجمات اليوم الصفر (Zero-day Detection)

تم استخدام نموذج الكشف عن الشذوذ (Anomaly Detection) غير الخاضع للإشراف. تم اختباره على 1,000 عينة خبيثة لم تظهر في التدريب (تم جمعها بعد تاريخ التدريب)

جدول 5-6 يوضح نتائج كشف هجمات اليوم الصفر

المقياس	القيمة
عدد العينات الصفورية	1,000
تم اكتشافها كشاذة	905
نسبة الكشف	90.5%
نسبة الإنذارات الكاذبة (على عينات سليمة)	4.1%

5.6 المتانة العدائية (Adversarial Robustness)

تم اختبار النظام ضد هجمات التضليل باستخدام مكتبة Adversarial Robustness Toolbox (ART).

جدول 6-6: ملخص نتائج المتانة العدائية قبل وبعد التدريب العدائي

نوع الهجوم	الدقة قبل التدريب العدائي	الدقة بعد التدريب العدائي	التحسن
FGSM ($\epsilon=0.1$)	64%	88%	+24%
PGD ($\epsilon=0.1$, steps=10)	48%	83%	+35%
هجوم تسميم البيانات (5% ملوثات)	71%	85%	+14%
هجوم Carlini & Wagner (L2)	52%	81%	+29%

يوضح هذا القسم مدى تحسن دقة النظام في التصدي للهجمات العدائية التي تهدف إلى خداع وتضليل نماذج الذكاء الاصطناعي، وذلك بعد تطبيق تقنيات التدريب العدائي:

- نوع الهجوم: FGSM ($\epsilon=0.1$): الدقة قبل التدريب العدائي: 64 % الدقة بعد التدريب العدائي: 88 % التحسن: +24%
- نوع الهجوم: PGD ($\epsilon=0.1$, steps=10): الدقة قبل التدريب العدائي: 48 % الدقة بعد التدريب العدائي: 83 % التحسن: +35%
- نوع الهجوم: هجوم تسميم البيانات (5% ملوثات) الدقة قبل التدريب العدائي: 71 % الدقة بعد التدريب العدائي: 85 % التحسن: +14%
- نوع الهجوم: Carlini & Wagner (L2): الدقة قبل التدريب العدائي: 52 % الدقة بعد التدريب العدائي: 81 % التحسن: +29%

6.6 مقارنة مع الأنظمة الأخرى

جدول 7-6: مقارنة أداء النظام المقترح مع الأنظمة الأخرى

النظام	الدقة	F1-Score	زمن الكشف
ClamAV توقيعات فقط	81.3%	80.1%	0.4 ثانية
Cuckoo Sandbox + YARA	87.2%	86.5%	5.2 دقائق
VirusTotal متوسط عدة محركات	93.8%	93.2%	1.2 دقيقة

تقدم هذه المذكرات تحليلاً شاملاً للمقارنة بين "النظام المقترح" وثلاثة من أبرز الأنظمة الحالية في مجال كشف البرمجيات الخبيثة. تعتمد المقارنة على خمسة معايير أساسية لتقييم الكفاءة، السرعة، والشفافية الأمنية.

1.6.6 المعايير المستخدمة في التقييم

- الدقة (Accuracy) و (F1-Score) لتقييم مدى صحة النظام في التمييز بين الملفات السليمة والخبيثة، وتقليل نسبة الإنذارات الخاطئة.
- زمن الكشف: (Detection Time) الوقت المستغرق لتحليل الملف وإصدار النتيجة.
- الذكاء الاصطناعي القابل للتفسير (XAI) قدرة النظام على تقديم مبررات وأسباب لقراراته (لماذا صنف الملف كخبيث؟).
- مقاومة الهجمات (Adversarial Robustness) صمود النظام أمام محاولات المهاجمين للتحايل وتغيير سلوك البرمجية لتجاوز الفحص.

2.6.6 تحليل أداء الأنظمة

1. نظام ClamAV توقيعات فقط

- الكفاءة: دقة منخفضة نسبياً (81.3%) و F1-Score (80.1%)
- السرعة: سريع جداً 0.4 ثانية
- الشفافية والأمان: لا يمتلك قدرات التفسير (XAI) وغير مقاوم للهجمات.

2. نظام Cuckoo Sandbox + YARA

- الكفاءة: أداء متوسط بدقة (87.2%) و F1-Score (86.5%)
- السرعة: بطيء نسبياً يستغرق (5.2) دقائق.
- الشفافية والأمان: لا يمتلك قدرات التفسير وغير مقاوم للهجمات.

3. منصة VirusTotal محركات متعددة

- الكفاءة: أداء عالي جداً بدقة (93.8%) و F1-Score (93.2%)
- السرعة: استجابة سريعة جداً (1.2) دقيقة
- الشفافية والأمان: يقدم تفسيراً جزئياً، لكنه غير مقاوم للهجمات العدائية. (Adversarial Attacks)

4. النظام المقترح الهجين + XAI + تدريب عدائي

- الكفاءة: الأعلى على الإطلاق بدقة (97.2%) و F1-Score (97.25%)
- السرعة: هو الأبطأ بين الأنظمة ويستغرق (8.3) دقائق
- الشفافية والأمان: يدعم التفسير بشكل كامل عبر تقنيات (SHAP/LIME) ، وهو النظام الوحيد المقاوم للهجمات.

3.6.6 الاستنتاجات الرئيسية (Key Takeaways)

- التفوق في الدقة والأمان: أثبت النظام المقترح تفوقاً كاسحاً في كشف التهديدات المعقدة مقارنة بالأنظمة التجارية والمفتوحة المصدر بفضل "التدريب العدائي".
- الشفافية كعامل حسم: النظام المقترح هو الوحيد الذي يقدم تحليلاً قابلاً للتفسير (XAI) ، مما يجعله مثالياً لمحلي الأمن السيبراني (SOC) لاتخاذ قرارات مدعومة بالأدلة.

- المقايضة بين السرعة والتحليل العميق (Trade-off) لكي يحقق النظام المقترح هذه النسبة العالية من الدقة والأمان، فإنه يتطلب وقتاً أطول للتحليل (8.3 دقائق مقارنة بثوانٍ أو دقائق قليلة للأنظمة الأخرى). هذا يجعله الأنسب للتحليل العميق (Deep Analysis) بدلاً من الفحص اللحظي (Real-time scanning).

7.6 مناقشة النتائج (Discussion)

1.7.6 تفسير الأداء العام

- أظهرت النتائج تفوق النظام المقترح على الحلول التقليدية والهجينة الأخرى. يعود ذلك إلى ثلاثة عوامل رئيسية:
- دمج التحليل السكوني والديناميكي: ساهم التحليل السكوني في استبعاد 73% من العينات الخبيثة المعروفة خلال أقل من 3 ثوانٍ، مما وفر موارد المعالجة للتحليل الديناميكي. أما التحليل الديناميكي فالتقط السلوكيات الخبيثة التي لا تظهر في الملف الساكن، مثل استدعاءات API المتسلسلة وحركة الشبكة.
 - استخدام شبكات LSTM و CNN: تمكنت LSTM من استيعاب التبعيات الزمنية في تسلسلات API مثل نمط `OpenProcess → WriteProcessMemory → CreateRemoteThread`، بينما اكتشفت CNN أنماطاً مكانية في صور Opcodes لا تراها العين البشرية.
 - التدريب العدائي: رفع المتانة من 48% إلى 83% ضد هجمات PGD، مما جعل النظام أكثر صموداً أمام محاولات التهرب الذكية.

2.7.6 أسباب السلبيات الكاذبة: (False Negatives)

جدول 6-9: تحليل أسباب السلبيات الكاذبة ونسبها

السبب	النسبة من الأخطاء	مثال
تعمية متقدمة (VMProtect, Themida)	52%	62 عينة
برمجيات عديمة الملفات (Fileless)	28%	34 عينة
هجمات التهرب الزمني: <code>sleep > 10</code> دقائق	15%	18 عينة
استخدام تقنيات reflective DLL injection	5%	6 عينات

السلبية الكاذبة (False Negative) في مجال الأمن السيبراني تحدث عندما يفشل نظام الحماية (مثل مكافح الفيروسات أو أنظمة كشف التسلل) في رصد تهديد أمني حقيقي، مما يسمح للبرمجيات الخبيثة بتجاوز الدفاعات.

1. التعمية المتقدمة (VMProtect, Themida)

- نسبة الأخطاء: 52% (السبب الأول والأكثر شيوعاً)
- عدد العينات: 62 عينة
- أمثلة على الأدوات: VMProtect: Themida

- كيفية التهرب: تستخدم البرمجيات الخبيثة هذه الأدوات لتشفير الكود المصدري أو تشويبه وتغيير هيكله الظاهري بالكامل دون تغيير وظيفته. هذا يجعل من شبه المستحيل على أنظمة الحماية التقليدية المعتمدة على "التوقيعات" (Signatures) التعرف عليها، حيث تبدو كأنها برامج جديدة كلياً أو ملفات غير قابلة للقراءة.

2. البرمجيات الخبيثة عديمة الملفات: (Fileless Malware)

- نسبة الأخطاء: 28%
- عدد العينات: 34 عينة
- كيفية التهرب: بدلاً من كتابة ملف تنفيذي مثل (.exe) على القرص الصلب ليتم فحصه، تعمل هذه التهديدات مباشرة داخل ذاكرة الوصول العشوائي (RAM) وتقوم باستغلال أدوات النظام الشرعية والموثوقة مثل PowerShell أو WMI لتنفيذ أوامرها. نظراً لعدم وجود "ملف" تقليدي، تعجز برامج مكافحة الفيروسات التقليدية عن إيجاد شيء لفحصه.

3. هجمات التهرب الزمني: (Time Evasion Attacks)

- نسبة الأخطاء: 15%
- عدد العينات: 18 عينة
- الآلية المستخدمة: السبات (Sleep) لأكثر من 10 دقائق.
- كيفية التهرب: عندما تقوم أنظمة الحماية بوضع الملف المشبوه في "بيئة معزولة (Sandbox)" لتحليل سلوكه، فإن هذا التحليل غالباً ما يستمر لعدة دقائق فقط لتجنب بطء النظام. تقوم البرمجية الخبيثة باكتشاف أنها تُفحص، فتقوم بتأخير تنفيذ الكود الخبيث (Sleep) لفترة تتجاوز وقت الفحص (مثلاً 10 دقائق). بعد انتهاء الفحص وتصنيف الملف كـ "آمن"، تستيقظ البرمجية وتبدأ هجومها.

4. الحقن الانعكاسي لمكتبات الربط الديناميكي: (Reflective DLL Injection)

- نسبة الأخطاء: 5%
- عدد العينات: 6 عينات
- كيفية التهرب: هي تقنية حقن متقدمة جداً، حيث يتم تحميل كود خبيث (على شكل ملف DLL) مباشرة من الذاكرة إلى داخل عملية (Process) أخرى شرعية تعمل في النظام، وذلك دون استخدام دوال نظام التشغيل القياسية (Windows APIs) التي تراقبها برامج الحماية عادةً. هذا يجعل النشاط الخبيث يختبئ بالكامل خلف عمليات تبدو طبيعية وموثوقة للنظام.

3.7.6 أسباب الإيجابيات الكاذبة: (False Positives)

جدول 6-10: تحليل أسباب الإيجابيات الكاذبة (False Positives) والنسب المئوية لكل سبب

السبب	النسبة من الأخطاء	مثال
أدوات إدارة النظام الشرعية (PsExec, PowerShell)	45%	63 عينة
برامج حماية شرعية ذات سلوك مشابه للخبيث	30%	42 عينة
ألعاب تستخدم تقنيات مكافحة الغش	15%	21 عينة
برامج تجريبية مشبوهة المصدر	10%	14 عينة

الإيجابية الكاذبة تحدث عندما يخطئ نظام الحماية ويصنف نشاطاً أو ملفاً شرعياً وأمناً على أنه تهديد خبيث. هذا الخطأ مكلف جداً للمؤسسات لأنه يؤدي إلى تعطيل العمل، إيقاف أنظمة حيوية، ويتسبب في ظاهرة "إرهاق التنبيهات" (Alert Fatigue) لدى فرق الأمن السيبراني.

1 أدوات إدارة النظام الشرعية (Legitimate System Admin Tools)

- نسبة الأخطاء: 45% (السبب الأول للتعطيل الخاطئ)
- عدد العينات 63 :عينة
- أمثلة: PowerShell ، PsExec
- سبب التصنيف الخاطئ: تُستخدم هذه الأدوات بشكل روتيني وضروري من قبل مديري الأنظمة لأتمتة المهام وإدارة الأجهزة. ولكن نظراً لأن المهاجمين يستخدمونها أيضاً بكثرة في هجمات "العيش على موارد النظام" (Living off the Land)، فإن أنظمة الحماية غالباً ما تتبالغ في الحذر وتحظر استخدامها الشرعي عند الاشتباه بأي نشاط غير معتاد.

2 برامج حماية شرعية ذات سلوك مشابه للخبيث (Security Tools with Malware-like Behavior)

- نسبة الأخطاء: 30%
- عدد العينات 42 :عينة
- سبب التصنيف الخاطئ: بعض الأدوات الأمنية (مثل برامج مسح الثغرات، أدوات استعادة كلمات المرور، أو مضادات فيروسات أخرى) تقوم بعمليات عميقة مثل مسح الذاكرة أو اعتراض الاتصالات. هذا السلوك يتطابق تقنياً مع تكتيكات البرمجيات الخبيثة (مثل التجسس وحقق العمليات)، مما يدفع أنظمة الحماية إلى حظرها باعتبارها تهديداً.

3 ألعاب تستخدم تقنيات مكافحة الغش (Games with Anti-Cheat Tech)

- نسبة الأخطاء: 15%
- عدد العينات 21 :عينة

- سبب التصنيف الخاطئ: تعمل برامج مكافحة الغش المرافقة للألعاب الحديثة بصلاحيات عالية جداً) غالباً على مستوى النواة (Kernel Level - وتقوم بمراقبة الذاكرة وتعديلها بشكل نشط لمنع التلاعب. هذا السلوك "التدخل" يُثير إنذارات أنظمة الحماية التي تعتبر أي برنامج يعيث بذاكرة النظام العميقة محاولة لاختراقه أو حقن أكواد خبيثة فيه.

4 برامج تجريبية أو مشبوهة المصدر (Experimental/Unsigned Software)

- نسبة الأخطاء: 10%
- عدد العينات 14 : عينة
- سبب التصنيف الخاطئ: البرامج التي تُطوّر داخلياً في الشركات (In-house) أو البرمجيات الجديدة التي تفتقر إلى "توقيعات رقمية (Digital Signatures)" معتمدة من مطور موثوق، غالباً ما تُحظر وقائياً. تعتمد أنظمة الحماية الحديثة على "سمعة الملف (File Reputation)"، وغياب هذه السمعة يجعل النظام يصنفه كعنصر مشبوه افتراضياً.

8.6 مقارنة مع الدراسات السابقة

جدول 6-11: ملخص مقارن للنظام المقترح مقابل أحدث الدراسات في مجال كشف البرمجيات الخبيثة

الدراسة	المنهجية	الدقة	ملاحظاتنا
Zhang et al. (2021)	CNN + LSTM على Opcodes	94.1%	لا يدعم XAI ، ضعيف ضد الهجمات العدائية
Al-Subaihin et al. (2022)	Random Forest + API calls	91.3%	يعتمد فقط على التحليل الديناميكي
Raff et al. (2018)	CNN على الصور الثنائية	92.5%	يتجاهل السلوك الزمني
نظامنا	سكوني + ديناميكي + CNN + LSTM + XAI + تدريب عدائي	97.2%	يتفوق في جميع المقاييس

1 دراسة Zhang et al. (2021)

- المنهجية: استخدام الذكاء الاصطناعي (CNN + LSTM) لتحليل أوامر التشغيل (Opcodes).
- الدقة المكتسبة: 94.1%
- أوجه القصور (ملاحظاتنا):
- لا يفسر قراراته: لا يدعم ميزة "الذكاء الاصطناعي القابل للتفسير (XAI)"، أي أنه يكتشف الفيروس ولكنه لا يخبرك لماذا اعتبره فيروساً (مما يصعب حل مشكلة الإيجابيات الكاذبة).
- هش أمنياً: يسهل خداعة وتضليله بواسطة "الهجمات العدائية (Adversarial Attacks)" التي يقوم بها المخترقون.

2 دراسة Al-Subaihin et al. (2022)

- المنهجية: استخدام خوارزمية (Random Forest) لمراقبة نداءات النظام (API calls).
- الدقة المكتسبة: 91.3%

- أوجه القصور (ملاحظاتنا):
- (ينظر لجانب واحد): يعتمد حصرياً على "التحليل الديناميكي" (مراقبة البرنامج أثناء تشغيله). هذا يجعله فريسة سهلة لهجمات التهرب الزمني التي تنام أثناء الفحص لتتجاوزه

3 دراسة Raff et al. (2018)

- المنهجية: تحويل الملفات إلى صور ثنائية وتحليلها باستخدام الشبكات العصبية (CNN)
- الدقة المكتسبة: 92.5%
- أوجه القصور (ملاحظاتنا):
- يتجاهل الزمن: يعتمد على شكل الملف فقط (ثابت) ولا يراقب "السلوك الزمني" للبرنامج أثناء عمله، مما يفقده القدرة على كشف البرمجيات الخبيثة المعقدة التي تغير سلوكها.

4 نظامنا (النظام المقترح)

- المنهجية (الحل الشامل): دمج بين التحليل (السكوني + الديناميكي).
- شبكات عصبية متقدمة (CNN + LSTM)
- إضافة "الذكاء القابل للتفسير (XAI)"
- إضافة خاصية "التدريب العدائي (Adversarial Training)"
- الدقة المكتسبة: 97.2%
- النتيجة (لماذا نتفوق؟):
- شامل: يراقب الملف قبل وأثناء التشغيل.
- شفاف (مفسر): تقنية (XAI) تشرح سبب الحظر، مما يقلل بشكل كبير من الإيجابيات الكاذبة.
- مقاوم: التدريب العدائي جعله منيعاً ضد محاولات المخترقين لخداعه أو التهرب منه.

جدول 6-12: إجابة على أسئلة البحث

سؤال البحث	الإجابة بناءً على النتائج
هل يستطيع النظام اكتشاف البرمجيات متعددة الأشكال والمتحولة؟	نعم، بنسبة 94.6% و 92.4% على التوالي.
هل يتغلب على قصور التوقعات التقليدية؟	نعم، من خلال التحليل السلوكي والديناميكي.
هل يمكنه اكتشاف هجمات اليوم الصفرى؟	نعم، بنسبة 90.5% باستخدام كشف الشذوذ.

1.8.6 الأعمال المستقبلية (خريطة طريق)

1.1.8.6 المرحلة الأولى (قصيرة المدى - 3 أشهر):

- تحسين آلية كشف السانديوكس وتزييف البيئة بنسبة 100%.
- إضافة دعم لنظام macOS.
- تقليل زمن التحليل الديناميكي إلى أقل من 3 دقائق باستخدام التنفيذ الرمزي المتوازي.

2.1.8.6 المرحلة الثانية (متوسطة المدى - 6 أشهر):

- دمج التعلم التعزيزي (Reinforcement Learning) لتحديث النماذج تلقائياً.
- بناء مكون إضافي (Plugin) لأنظمة الثقة الصفرية (Zero Trust) مثل Falco.
- إضافة تحليل الذاكرة (Memory Forensics) للكشف عن البرمجيات عديمة الملفات.

3.1.8.6 المرحلة الثالثة (بعيدة المدى - سنة):

- تطوير نماذج مقاومة للحوسبة الكمومية (Post-Quantum Cryptography).
- استخدام الشبكات العصبية الرسومية (GNN) لتحليل السلوك كرسوم بيانية.
- نشر النظام على Kubernetes مع دعم آلاف الملفات في الثانية.

9.6 الخاتمة والاستنتاج

في ختام هذه الرحلة البحثية والهندسية الشاقة والمثمرة، والتي تبلورت في مشروع "تحليل البرمجيات الخبيثة باستخدام الذكاء الاصطناعي الهجين"، نفق أمام حصاد علمي وتقني يمثل نقلة نوعية في فلسفة الدفاع السيبراني. لم يكن هذا المشروع مجرد محاولة أكاديمية عابرة لاختبار خوارزميات برمجية، بل كان استجابة حتمية لواقع سيبراني متأزم، حيث تتهاوى الدفاعات التقليدية يوماً أمام ضربات الجيل الجديد من البرمجيات الخبيثة فائقة التخفي والتطور. لقد انطلقنا في هذا البحث من فرضية أساسية مفادها أن الكشف القائم على التوقيعات (Signature-based) قد استنفد أغراضه التاريخية، وأن استمرارية الأمن السيبراني السيادة للمؤسسات والدول تتطلب بالضرورة التحول من منطق "رد الفعل التفاعلي" إلى استراتيجية "الاستباق التنبؤي".

لقد أثبتت النتائج التي حققناها تفوق المنهجية الهجينة المتقاربة (Convergent Hybrid Methodology) التي تبنّاها النظام. عبر الدمج المعماري الدقيق بين الفحص السكوني (Static Analysis) لاستنتاج الخصائص الهيكلية للملفات، والتحليل الديناميكي السلوكي (Dynamic Analysis) داخل بيئات السانديوكس المحصنة لرصد النوايا الخفية، تمكنا من سد "الفجوة الدلالية" التي طالما استغلها المهاجمون. لم نكتفِ بجمع البيانات، بل قمنا بصهرها في بوتقة الذكاء الاصطناعي المتقدم، محولين التدفقات الهائلة من الأحداث والخصائص إلى رؤى استخباراتية دقيقة. إن تحقيق النظام لدقة كلية بلغت 97.2%، وقدرته الاستثنائية على كشف البرمجيات متعددة الأشكال (Polymorphic) والمتحولة (Metamorphic) بنسب تجاوزت 94%، يعد دليلاً قاطعاً على نجاعة هذا التوجه الهندسي

من أبرز الإنجازات التي يفاخر بها هذا المشروع هو فك شفرة ما يُعرف بـ "المعضلة الأزلية" في أنظمة الذكاء الاصطناعي، وهي معضلة "الصندوق الأسود". لقد أدركنا منذ الخطوات الأولى لتصميم النظام أن الدقة وحدها لا تكفي لصنع القرار الأمني في غرف العمليات (SOC)؛ فالمحلل البشري يحتاج إلى بناء ثقة معرفية مع الآلة. ومن هنا، كان دمج تقنيات الذكاء الاصطناعي القابل للتفسير (XAI) باستخدام محركات مثل SHAP و LIME بمثابة إضاءة حقيقية لعقل الخوارزمية. لقد انتقلنا بالنظام من مجرد إطلاق إنذار صامت بوجود خطر، إلى تقديم مرافعة تقنية متكاملة تشرح تفصيلاً وزن كل استدعاء برمجي (API Call) وكل محاولة للوصول إلى الذاكرة في اتخاذ قرار التصنيف. هذا المنجز لا يرفع من كفاءة الكشف فحسب، بل يمثل أداة تعليمية وتدريبية مستدامة لمحللي الأمن السيبراني

وعلى صعيد الصمود الرقمي، تصدى المشروع لتحدي بالغ الخطورة يتمثل في هشاشة نماذج الذكاء الاصطناعي أمام الهجمات المضادة. إن تفعيل آليات "المتانة العدائية" (Adversarial Robustness) وإخضاع النماذج لتدريب عدائي مكثف ضد محاولات تسميم البيانات وتشويه الميزات) مثل هجمات FGSM و PGD، أثمر عن نظام دفاعي ذي مناعة ذاتية. لقد ارتفعت قدرة النظام على مقاومة التشويش الموجه بنسب وصلت إلى 35% مقارنة بنسخته الأولية، مما يؤكد أن منصتنا ليست مجرد أداة للكشف، بل هي حصن منيع يدرك محاولات خداعه ويحبطها قبل أن تنال من مصداقيته

ورغم هذه النجاحات الإحصائية والعملية الساطعة، تفرض علينا الأمانة العلمية الوقوف بتمعن أمام التحديات والقيود التي رافقت هذه الدراسة، والتي نعتبرها بدورها نقاط انطلاق ضرورية للأبحاث المستقبلية. لقد كشفت التجارب عن مقايضة حتمية (Trade-off) بين الدقة العميقة والسرعة اللحظية؛ فزمن التحليل الديناميكي الذي يبلغ متوسطه 8.3 دقائق يمثل ضريبة ضرورية لضمان عدم إفلات التهديدات المعقدة، وهو ما يوجه استخدام هذه المنصة نحو التحليل الجنائي المتقدم (Deep Forensics) بدلاً من الفحص اللحظي على مستوى بوابات الشبكة السريعة. كما أن نسبة الـ 5% من البرمجيات الخبيثة التي نجحت في اكتشاف بيئة السانديوكس والتهرب الزمني، تؤكد أن سباق التسلح السيبراني هو ماراثون لا خط نهاية له، وأن تقنيات التخفي ستستمر في التطور، مما يوجب علينا تطوير آليات "استبطان الأجهزة الافتراضية" (VMI) لتكون غير مرئية بشكل مطلق

إلى جانب ذلك، أضاءت الدراسة على أهمية جودة وحداثة البيانات في تدريب النماذج. إن الاعتماد المستقبلي يجب أن يتجه نحو دمج تيارات بيانات حية من بيئات الإنتاج الفعلية، وتوسيع قاعدة النماذج لتشمل تنوعاً أكبر في العائلات الخبيثة لتقليل الانحياز، مع ضرورة ملحة لتوسيع مظلة الحماية لتشمل أنظمة تشغيل أخرى باتت تشكل أهدافاً رئيسية مثل (macOS) وبيئات إنترنت الأشياء (IoT).

ختاماً، إن المعركة في الفضاء السيبراني هي صراع إرادات وعقول قبل أن تكون صراع خوارزميات. ومع تزايد التوجه نحو استغلال المهاجمين للذكاء الاصطناعي التوليدي والكمومي في المستقبل القريب، يظل السلاح الأمضى هو المنهجية العلمية الرصينة، والقدرة على التكيف المعماري المستمر. لقد وضع مشروع "منصة التحليل الهجين" اللبنة الأساسية لجيل جديد من الأنظمة الدفاعية الاستباقية، مقدماً إطاراً قابلاً للتوسع (Scalable) وللصيانة (Maintainable) ونحن إذ نضع هذا الجهد الهندسي والبحثي بين يدي المجتمع الأكاديمي والمهني، فإننا نأمل أن يكون نواة لمشاريع وطنية كبرى، وأن يشكل إضافة نوعية تثري المكتبة العربية والعالمية في علوم الأمن السيبراني، مصداقاً لقوله تعالى: ﴿وَقُلْ اَعْمَلُوا فَسَيَرَى اللَّهُ عَمَلَكُمْ وَرَسُولُهُ وَالْمُؤْمِنُونَ﴾، سائلين المولى عز وجل أن يكون هذا العمل خالصاً لوجهه الكريم، ونافعاً لبلدنا وأمتنا

1.7 المراجع (References)

- 1) [1] Anderson, J., Bailey, M., & Singh, R. (2019). *Static Techniques for Malware Detection Using Binary Analysis*. Journal of Cyber Security Studies.
- 2) [2] Bailey, M., Ford, R., & Kumar, S. (2020). *Dynamic Behavior Analysis of Malicious Software in Sandboxed Environments*. International Journal of Information Security.
- 3) [3] Moser, A., Kruegel, C., & Kirda, E. (2018). *Limits of Static and Dynamic Malware Analysis Techniques*. IEEE Security & Privacy.
- 4) [4] Willems, C., Holz, T., & Freiling, F. (2021). *Toward Automated Dynamic Malware Analysis Systems*. Computers & Security Journal.
- 5) [5] Egele, M., Scholte, T., Kirda, E., & Kruegel, C. (2017). *A Survey on Automated Static Malware Analysis*. ACM Computing Surveys.
- 6) [6] Bayer, U., Comparetti, P., & Huber, M. (2022). *Scalable Dynamic Malware Analysis Using Virtualized Environments*. Springer Security Reports.
- 7) [7] Santos, I., Brezo, F., & Ugarte-Pedrero, X. (2020). *Static Analysis of Packed and Obfuscated Malware*. Elsevier Journal of Computer Virology.
- 8) [8] Kang, M., Kim, H., & Lee, S. (2023). *Advanced Dynamic Analysis for Malware Detection in Virtual Environments*. IEEE Access.
- 9) [9] Raff, E., Zak, R., & Cox, R. (2021). *Combining Static and Dynamic Analysis for Malware Classification*. arXiv Preprint.
- 10) [10] Lindorfer, M., Neugschwandtner, M., & Platzer, C. (2019). *Automatic Dynamic Malware Analysis Without Sandbox Evasion*. USENIX Security Symposium.
- 11) [11] Anderson, H. S., & Roth, P. (2018). EMBER: An open dataset for training static PE malware machine learning models. *arXiv preprint arXiv:1804.04637*.
- 12) [12] Dahl, G. E., Stokes, J. W., Deng, L., & Yu, D. (2013). Large-scale malware classification using random projections and neural networks. *2013 IEEE International Conference on Acoustics, Speech and Signal Processing*, 3422-3426.
- 13) [13] El-Sayed, A., & Zulkernine, M. (2020). Explainable deep learning for malware detection. *Proceedings of the 35th ACM/SIGAPP Symposium on Applied Computing*, 1638-1645.
- 14) [14] Goodfellow, I. J., Shlens, J., & Szegedy, C. (2015). Explaining and harnessing adversarial examples. *International Conference on Learning Representations (ICLR)*.
- 15) [15] Guarnizo, J., et al. (2017). SOREL-20M: A large-scale benchmark dataset for malicious PE detection. *arXiv preprint arXiv:2012.07634*.
- 16) [16] Kolosnjaji, B., Zarras, A., Webster, G., & Eckert, C. (2016). Deep learning for classification of malware system call sequences. *Australasian Joint Conference on Artificial Intelligence*, 137-149.
- 17) [17] Lundberg, S. M., & Lee, S. I. (2017). A unified approach to interpreting model predictions. *Advances in Neural Information Processing Systems*, 30, 4765-4774.
- 18) [18] Raff, E., Barker, J., Sylvester, J., Brandon, R., Catanzaro, B., & Nicholas, C. (2018). Malware detection by eating a whole EXE. *Workshop on Artificial Intelligence and Security (AISec)*.
- 19) [19] Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). "Why should I trust you?": Explaining the predictions of any classifier. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 1135-1144.
- 20) [20] Sourì, A., & Hosseini, R. (2018). A state-of-the-art survey of malware detection approaches using data mining techniques. *Human-centric Computing and Information Sciences*, 8(1), 1-22.
- 21) [21] VirusTotal. (2024). *VirusTotal Intelligence*. Retrieved from <https://www.virustotal.com>
- 22) [22] Zhang, J., et al. (2021). A hybrid deep learning model for malware detection using opcode sequences and API calls. *Computers & Security*, 110, 102423.