

MalwareLab.ai User Guide

Model: gpt-4.1-mini

Current Version: beta-v1.9

Date: June 12, 2026

REV: June 27, 2026

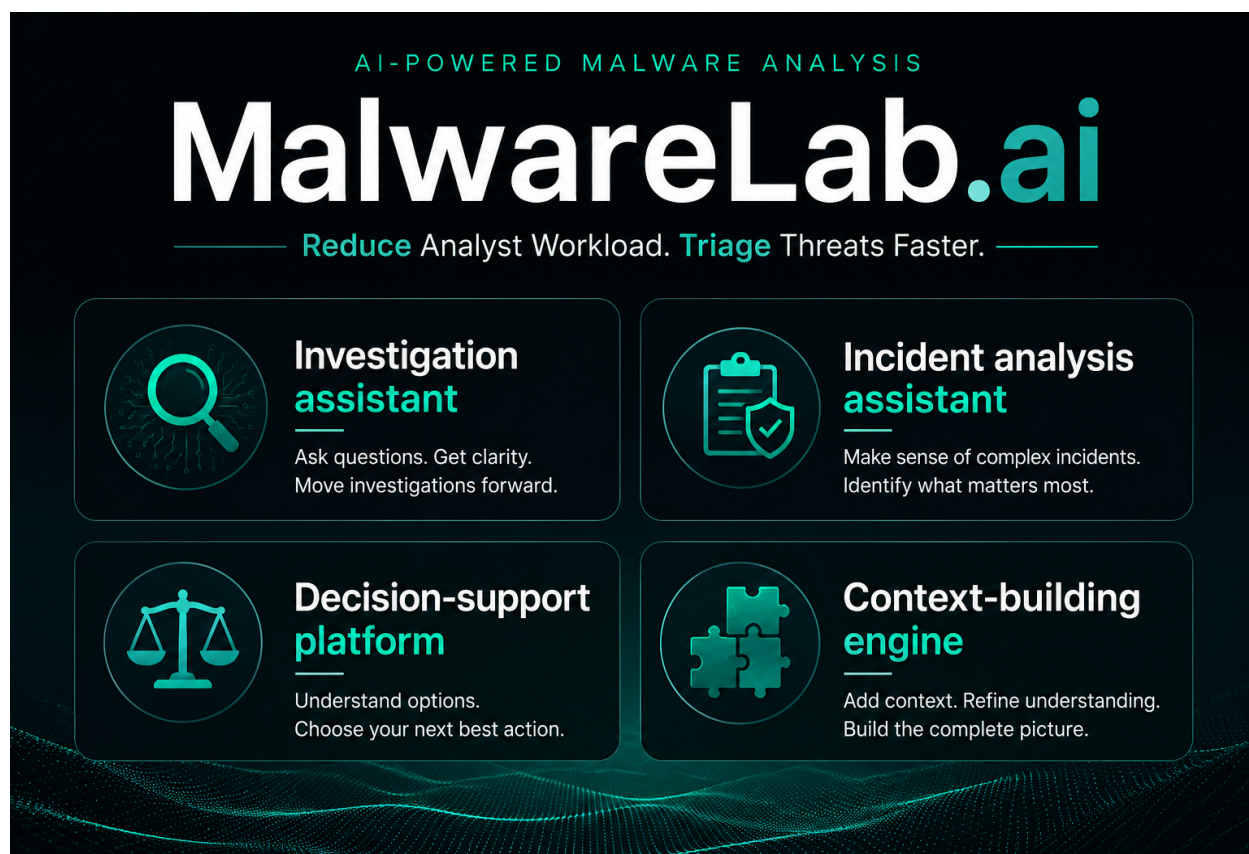


TABLE OF CONTENTS

Introduction.....	3
Our Philosophy.....	3
Cybersecurity Should Be Accessible.....	3
What MalwareLab.ai Is Not.....	4
What MalwareLab.ai Is.....	10
When Should I Use MalwareLab.ai?.....	11
No Query Language Required.....	12
User-Friendly by Design.....	13
How to Use MalwareLab.ai.....	14
Example 1: PowerShell Script.....	14
Example 2: Binary File.....	19
Example 3: Large Dataset.....	22
Supported Inputs.....	25
You Don't Need Files or Logs.....	26
Example: Potential Data Exfiltration.....	26
Example: Suspicious PowerShell Activity.....	26
Example: Impossible Travel.....	27
Example: Large-Scale File Encryption.....	27
Example: Unauthorized Administrative Changes.....	27
Example: Suspicious PowerShell Email Attachment.....	28
Building on Previous Analysis.....	29
Example: Administrative Account Investigation.....	29
Example: Data Exfiltration Investigation.....	30
Example: Ransomware Investigation.....	31
Types of Analysis Supported.....	33
Malware Analysis.....	33
Script Analysis.....	33
Log Analysis.....	33
Incident Analysis.....	34
IOC Analysis.....	34
IOC Enrichment.....	35
Protecting Sensitive Organizational Information.....	36
Understanding Your Results.....	38
Executive Summary.....	38
Technical Analysis.....	38
Indicators of Compromise (IOCs).....	38
MITRE ATT&CK Mapping.....	38
Understanding Confidence Levels.....	39

High Confidence.....	39
Medium Confidence.....	39
Low Confidence.....	39
Confidence Assessment.....	39
Recommended Next Steps.....	39
Known Limitations.....	40
AI-Generated Analysis Disclaimer.....	40
Best Practices.....	41
Final Thoughts.....	41
Appendix A: User Tips.....	42

Introduction

Welcome to MalwareLab.ai.

MalwareLab.ai is an AI-powered cybersecurity analysis platform designed to help security professionals quickly understand suspicious activity, malware, scripts, logs, and potential security incidents.

Modern cybersecurity teams are often overwhelmed by alerts, investigations, and massive amounts of technical data. MalwareLab.ai helps reduce that burden by transforming raw information into understandable analysis and actionable insights.

Whether you are investigating malware, reviewing suspicious logs, analyzing a PowerShell script, or simply trying to determine whether a situation appears malicious, MalwareLab.ai can help you reach conclusions faster.

Our Philosophy

Cybersecurity Should Be Accessible

Many cybersecurity tools require extensive training, specialized certifications, or knowledge of complex query languages before users can obtain meaningful results.

MalwareLab.ai was built around a different idea:

The tool should adapt to the analyst—not the other way around.

Our goal is to help users:

- Reduce time spent investigating threats
- Understand suspicious activity faster
- Improve decision-making
- Accelerate incident response
- Produce better documentation and reporting

MalwareLab.ai is designed to serve as an intelligent cybersecurity assistant that helps analysts spend less time searching for answers and more time acting on them.

What MalwareLab.ai Is Not

MalwareLab.ai is designed to assist with cybersecurity analysis, threat investigations, malware analysis, and suspicious activity assessment.

To ensure the best experience, users should understand what the platform is and is not intended to do.

Not a General IT Help Desk

MalwareLab.ai is not designed to troubleshoot routine IT support issues.

Examples:

- ✗ "Why won't my printer print?"
- ✗ "How do I connect to Wi-Fi?"
- ✗ "Why is Outlook running slowly?"
- ✗ "My monitor is flickering."

These issues are better handled through standard IT support channels.

Not a System Administration Tool

MalwareLab.ai does not manage, configure, or administer systems.

Examples:

- ✗ Creating Active Directory users
- ✗ Configuring firewalls

✗ Managing Microsoft 365 settings

✗ Deploying software updates

✗ Resetting passwords

While MalwareLab.ai may explain security implications of these activities, it does not perform administrative functions.

Not a SIEM Replacement

MalwareLab.ai is not intended to replace existing security monitoring platforms.

Examples:

✗ Real-time event collection

✗ Log aggregation

✗ Continuous alert generation

✗ Enterprise-wide telemetry collection

✗ Long-term log retention

Instead, MalwareLab.ai helps users understand and investigate data obtained from those systems.

Not an Endpoint Detection and Response (EDR) Platform

MalwareLab.ai does not continuously monitor endpoints.

Examples:

- ✗ Real-time process monitoring
- ✗ Live threat hunting across endpoints
- ✗ Device isolation
- ✗ Automated remediation
- ✗ Continuous endpoint telemetry collection

It analyzes information provided by the user rather than actively monitoring systems.

Not an Antivirus Product

MalwareLab.ai does not replace traditional antivirus or anti-malware solutions.

Examples:

- ✗ Real-time malware blocking
- ✗ File quarantine
- ✗ Signature updates
- ✗ Automatic malware removal

The platform assists with analysis and understanding rather than prevention and enforcement.

Not a Digital Forensics Platform

MalwareLab.ai can assist with interpreting forensic artifacts, but it is not a replacement for dedicated forensic tools.

Examples:

- ✗ Memory acquisition
- ✗ Disk imaging
- ✗ Registry extraction
- ✗ Evidence preservation
- ✗ Chain-of-custody management

Users should collect forensic evidence using established forensic tools and procedures.

Not a Vulnerability Scanner

MalwareLab.ai does not actively scan networks, applications, or systems for vulnerabilities.

Examples:

- ✗ Port scanning
- ✗ Vulnerability discovery
- ✗ Patch verification
- ✗ Asset enumeration
- ✗ Network mapping

It can help explain vulnerabilities once identified but does not perform scanning activities.

Not an Automated Decision-Making System

MalwareLab.ai provides analysis and recommendations, but final decisions remain the responsibility of the user.

Examples:

- ✗ Automatically terminating employees
- ✗ Automatically blocking users
- ✗ Automatically declaring incidents
- ✗ Automatically attributing attacks
- ✗ Automatically escalating disciplinary actions

Users should validate findings through their organization's established processes.

Not a Source of Absolute Certainty

Cybersecurity investigations often involve incomplete information.

MalwareLab.ai evaluates available evidence and provides informed analysis, but it cannot guarantee conclusions.

Examples:

- ✗ Guaranteed malware attribution
- ✗ Guaranteed threat actor identification
- ✗ Guaranteed compromise confirmation
- ✗ Guaranteed intent determination

The platform should be viewed as an investigative assistant rather than an unquestionable authority.

Not a Replacement for Human Analysts

MalwareLab.ai is designed to assist security professionals, not replace them.

Human expertise remains critical for:

- Investigations
- Decision-making
- Incident response
- Risk management
- Organizational context

The strongest results come from combining AI-assisted analysis with human judgment and experience.

What MalwareLab.ai Is

MalwareLab.ai is an AI-powered cybersecurity analysis platform designed to help users:

- ✓ Analyze malware samples
 - ✓ Explain scripts and code
 - ✓ Analyze logs
 - ✓ Review suspicious activity
 - ✓ Extract indicators of compromise
 - ✓ Generate incident summaries
 - ✓ Map activity to MITRE ATT&CK
 - ✓ Refine assessments as additional evidence becomes available
 - ✓ Accelerate investigations
 - ✓ Support cybersecurity decision-making
-

When Should I Use MalwareLab.ai?

MalwareLab.ai is designed to assist cybersecurity professionals during investigations, malware analysis, incident response activities, and suspicious activity reviews.

Consider using MalwareLab.ai when:

- You discover a suspicious file
- You encounter an unfamiliar PowerShell script
- You need to analyze malware behavior
- A user reports unusual activity
- You suspect data exfiltration
- You need help interpreting log data
- You want a quick incident summary
- You need a second opinion during an investigation
- You want help identifying potential attacker techniques
- You need to understand indicators of compromise

MalwareLab.ai is most effective when helping analysts answer:

"What am I looking at?"

"Why is it suspicious?"

"What should I investigate next?"

Example Use Cases:

Role	Use Case
SOC Analyst	Investigating suspicious PowerShell
Incident Responder	Reviewing ransomware indicators
MSSP Analyst	Triage customer alerts
Security Engineer	Understand malicious scripts
IT Administrator	Validate suspicious activity before escalation

No Query Language Required

One of the most significant advantages of MalwareLab.ai is that it does not require users to learn a specialized query language.

Many cybersecurity platforms require knowledge of:

- Splunk SPL
- Kusto Query Language (KQL)
- SQL
- SIEM-specific search syntax
- Custom detection languages

MalwareLab.ai uses natural language instead.

Simply ask questions the same way you would ask a colleague.

Examples:

- "What does this PowerShell script do?"
- "Does this malware appear to be ransomware?"
- "Summarize this incident."
- "Explain this code in plain English."
- "Is this activity suspicious?"
- "What are the indicators of compromise?"
- "What MITRE ATT&CK techniques are present?"

No special syntax is required.

User-Friendly by Design

MalwareLab.ai was built to provide useful answers with minimal effort.

Users can:

- Upload malware samples
- Upload scripts
- Upload log files
- Upload incident data
- Paste code directly into the platform
- Ask questions using natural language

Results are presented in a straightforward format that focuses on helping users understand what they are looking at rather than overwhelming them with unnecessary complexity.

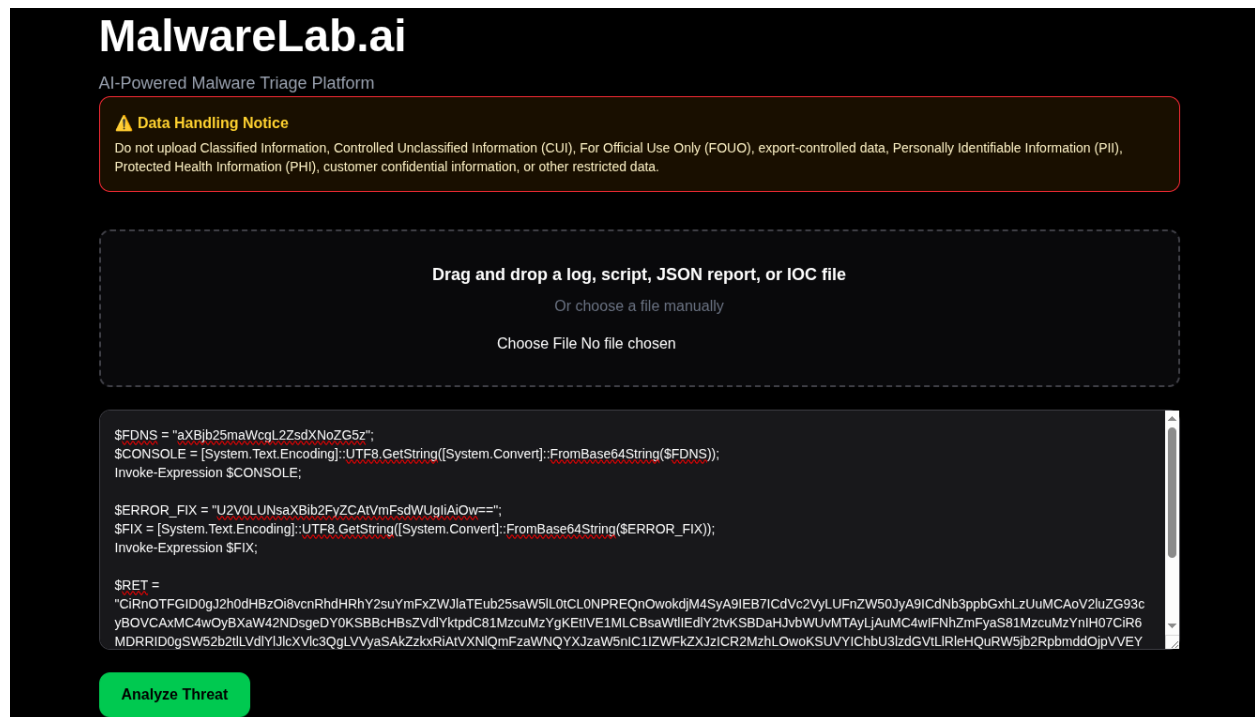
How to Use MalwareLab.ai

Enter your logs, files, scripts, IOCs, or even just a situation that you would like to analyze. Then click the “Analyze Threat” button at the bottom left.

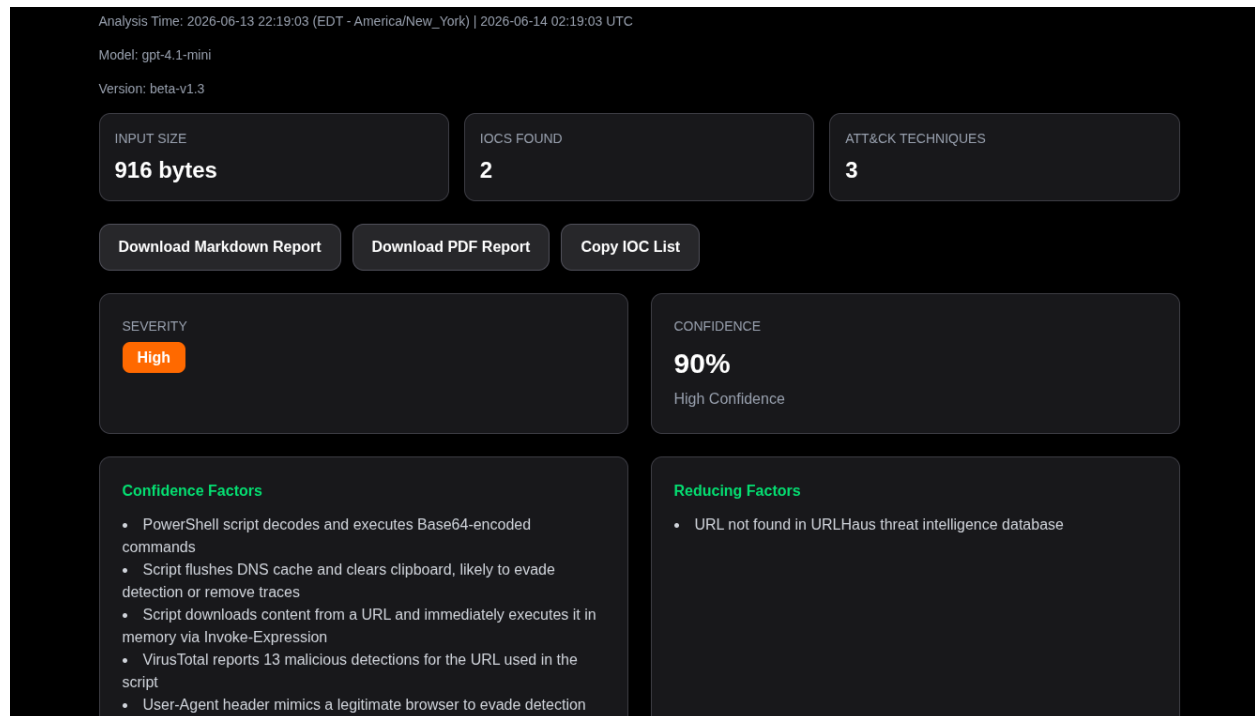
NOTE: The “Enter” button on your keyboard will not initiate an analysis. The “Enter” button will move the cursor down to the next line in the input box.

Example 1: PowerShell Script

For this illustration, a malicious PowerShell script is being examined:



Wait for the analysis to complete, and then scroll down to see results:



- **Analysis Timestamps:**
 - Both your local time (pulled from your browser) and the UTC time are displayed for ease of convertibility.
- **File Download and Copy Buttons:**
 - Download Markdown Report (downloads an editable document).
 - Download PDF Report (downloads a non-editable PDF).
 - Copy IOC List (copies the IOC list to your clipboard).
- **Severity:**
 - This section will grade the severity of every analysis as “Critical,” “High,” “Medium,” or “Low” in order to assist in triage.
- **Confidence:**
 - Gives a confidence rating based on the amount and types of evidence.
- **Confidence Factors:**
 - Displays the evidence that led to the severity conclusion.
- **Reducing Factors:**
 - States reasons why there is not 100% confidence, such as a lack of threat analysis by one of the IOC enrichment sources.

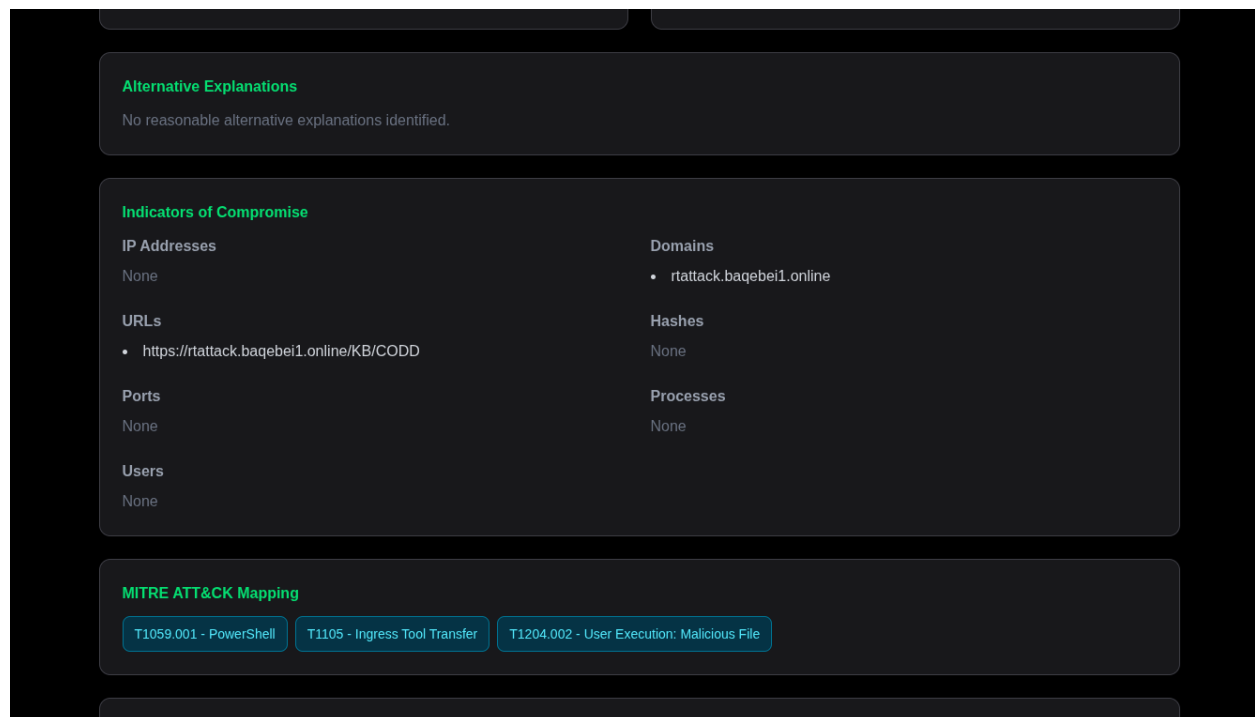
Scroll down further to see additional output display fields:

The screenshot displays a security analysis report with the following sections:

- Threat Type**
Downloader / Remote Code Execution
- Executive Summary**
The analyzed PowerShell script uses Base64-encoded commands to flush DNS, clear the clipboard, and download remote content from a URL with multiple malicious detections on VirusTotal. The downloaded content is executed dynamically in memory, indicating downloader or remote code execution behavior.
- Evidence Supporting Assessment**
 - Base64 string 'aXBjb25maWcgL2ZsdXNoZG5z' decodes to 'ipconfig /flushdns' executed via Invoke-Expression
 - Base64 string 'U2V0LUNsaXBib2FyZCAtVmFsdWUgIiAiOw==' decodes to 'Set-Clipboard -Value ""' executed via Invoke-Expression
 - Base64 string containing Invoke-WebRequest to 'https://rtattack.baqebe1.online/KB/CODD' with a browser User-Agent header
 - Invoke-WebRequest response content is decoded and executed immediately with Invoke-Expression (IEX)
 - VirusTotal URL enrichment shows 13 malicious detections for 'https://rtattack.baqebe1.online/KB/CODD'
- Ambiguity Notes**
 - No direct payload or downloaded content sample provided to analyze the final executed code
 - URL not found in URLHaus, which limits corroboration from that source
 - No endpoint telemetry or process execution logs available to confirm impact or persistence
- Alternative Hypotheses**
 - Red team or penetration testing activity using obfuscated PowerShell to simulate attacker behavior
 - Security research or malware analysis script demonstrating downloader techniques
 - False positive or benign script with unusual obfuscation and remote content retrieval

- **Threat Type:**
 - Gives each analysis a “threat type” label (i.e. Downloader/Remote Code Execution, Ransomware, Data Exfiltration, etc.).
- **Executive Summary:**
 - Gives an explanation of the findings of the analysis. This is an overall statement of the analysis based on the evidence found.
- **Evidence Supporting Assessment:**
 - This section displays all the evidence that led to the conclusion of the analysis.
- **Ambiguity Notes:**
 - Lists factors that are unclear, such as a lack of logs, telemetry, or IOC enrichment.
- **Alternative Hypotheses:**
 - Other malicious security-related explanations that could explain the evidence.
 - Example: "Compromised account", "Insider threat."

Scroll down to see additional output display fields:



- **Alternative Explanations:**
 - Benign or non-malicious explanations.
 - Example: "Authorized maintenance", "VPN usage", "Backup software".
 - **NOTE:** This section differs from the "Alternative Hypotheses" section above in that while "Alternative Hypotheses" gives possible other malicious explanations, "Alternative Explanations" gives possible benign explanations. Do not confuse the two.
- **Indicators of Compromise:**
 - Lists all IOCs discovered, including IP Addresses, Domains, URLs, Hashes, Ports, Processes, Users.
- **MITRE ATT&CK Mapping:**
 - Lists all Tactics, techniques, and Procedures (TTPs) that have been categorized by the MITRE Corporation in the ATT&CK Framework.

Scroll down to see the final output display fields:

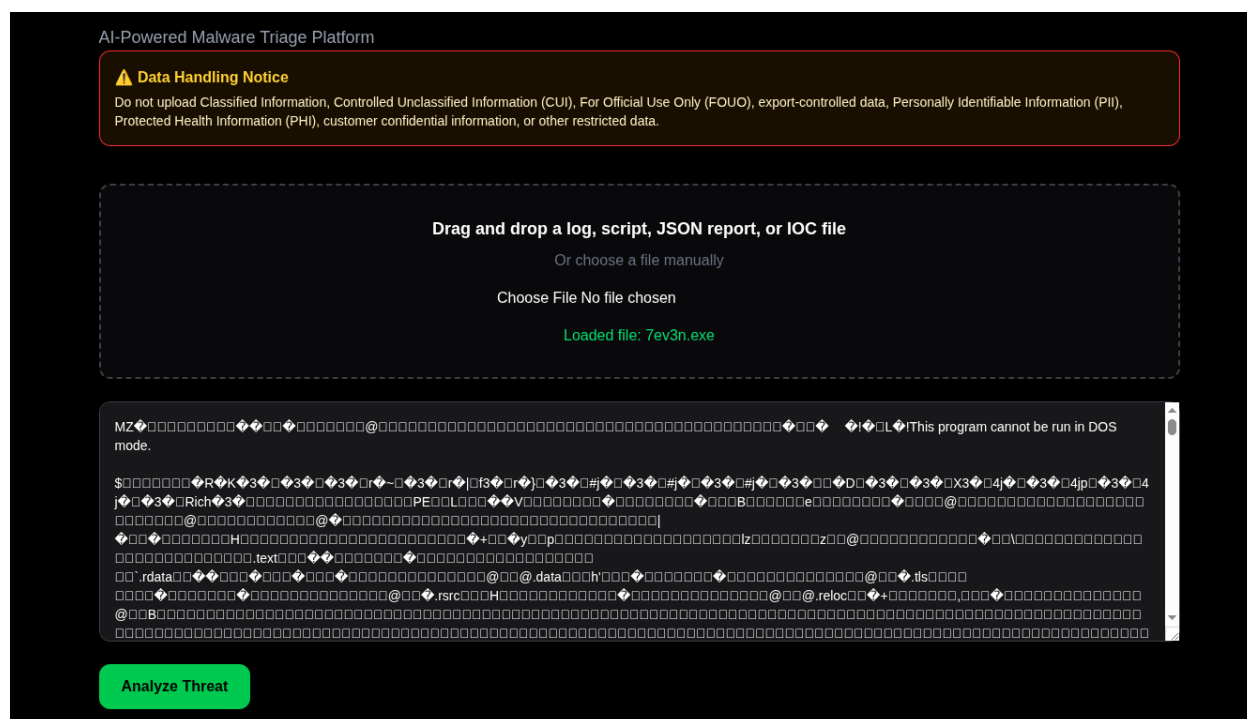
The image shows a dark-themed user interface for MalwareLab.ai. It features three distinct sections. The first section, titled 'Recommended Actions' in green, contains a bulleted list of five items: investigating PowerShell commands, blocking network connections to 'rtattack.bagebei1.online', searching for persistence mechanisms, performing endpoint forensics, and implementing PowerShell logging. The second section, titled 'Analyst Notes' in green, includes a text area with a placeholder 'Optional analyst observations, ticket numbers, customer notes, containment status, pending actions, etc.'. The third section, titled 'Was this analysis helpful?' in green, has two buttons for 'Yes' and 'No' (each with a thumbs icon), a text area for optional feedback, and a 'Submit Feedback' button at the bottom.

- **Recommended Actions:**
 - Lists what the next actions that the analyst or organization should take.
 - Examples could be anything from blocking a domain or IP, collecting additional logs, or educating the user.
 - **NOTE:** This section can be very helpful for both junior analysts or non-security focussed IT personnel to help them determine the next step to take.
- **Analyst Notes:**
 - This provides an opportunity for the analyst to put in any notes they wish to input, such as the threat they analyzed.
 - **NOTE:** Any notes input will show up in Markdown or PDF downloads of the analysis.
- **Was the analysis helpful?**
 - This provides the analyst to give MalwareLab.ai feedback with either a “Yes” or a “No.”
 - The analyst can also submit feedback to MalwareLab.ai.

Example 2: Binary File

For this illustration, the “7ev3n.exe” file (yes, the ransomware) is being examined:

IMPORTANT NOTE: Users should handle malware samples in accordance with organizational security policies and only analyze malware within approved environments.



NOTICE: For file analysis, the user can either drag and drop the file into the box, or click inside the drag and drop box and upload a file from a filepath of their choosing.

Once a file is either dropped or dropped into the drag and drop box, the text box will populate with text characters.

See the analysis for a real ransomware file below:

SEVERITY

Critical

CONFIDENCE

95%

High Confidence

Confidence Factors

- Presence of ransom note text indicating file encryption and ransom demand
- Registry modification adding 'crypted' value under CurrentVersion
- Scheduled task creation for persistence with task name 'uac'
- Use of bcdedit commands to disable recovery and boot options
- Presence of file deletion batch scripts targeting user directories
- Network communication to blockchain.info API endpoints for payment address generation

Reducing Factors

- No direct evidence of successful encryption or file modification timestamps
- No observed network traffic or process execution logs included

Threat Type

Ransomware

Executive Summary

The analyzed binary and associated artifacts exhibit strong indicators of ransomware activity, including file encryption notification, ransom demand with Bitcoin payment instructions, persistence mechanisms via scheduled tasks and registry modifications, and system configuration changes to hinder recovery.

Evidence Supporting Assessment

- Ransom note text stating files have been encrypted and providing payment instructions to JulyCezar1001@mail.com
- Registry key 'crypted' added under HKLM\SOFTWARE\Microsoft\Windows\CurrentVersion
- Scheduled task 'uac' created to execute persistence batch file
- bcdedit commands disabling recovery and boot status policy
- Batch scripts deleting files and modifying registry keys related to accessibility and keyboard layout
- Network URLs referencing blockchain.info API for Bitcoin payment address generation

Ambiguity Notes

- No direct execution or infection timeline provided
- No sample behavioral logs or network captures to confirm active communication
- No file system metadata to confirm actual file encryption
- Ransom note presence strongly suggests compromise but full impact assessment requires additional telemetry

Alternative Hypotheses

- Ransomware simulation or test artifact
- Malicious downloader or dropper preparing system for ransomware
- False positive due to embedded ransom note text in benign software

Alternative Explanations

- The ransom note and registry changes could be part of a penetration test or red team exercise
- Batch scripts and scheduled tasks might be used for legitimate system maintenance if context is missing

Indicators of Compromise

IP Addresses

Domains

- Batch scripts and scheduled tasks might be used for legitimate system maintenance if context is missing

Indicators of Compromise

IP Addresses None	Domains • blockchain.info
URLs • https://blockchain.info/api/receive?method=create&address= • http://c.e	Hashes None
Ports None	Processes • cmd.exe • schtasks.exe • bcdedit.exe
Users None	

MITRE ATT&CK Mapping

T1486 - Data Encrypted for Impact
T1053.003 - Scheduled Task/Job
T1112 - Modify Registry
T1490 - Inhibit System Recovery

T1071.001 - Application Layer Protocol: Web Protocols

Notice this analysis contains processes.

MITRE ATT&CK Mapping

T1486 - Data Encrypted for Impact
T1053.003 - Scheduled Task/Job
T1112 - Modify Registry
T1490 - Inhibit System Recovery

T1071.001 - Application Layer Protocol: Web Protocols

Recommended Actions

- Investigate scheduled task 'uac' and associated batch files in %AppData%\Local
- Review registry keys under HKLM\SOFTWARE\Microsoft\Windows\CurrentVersion for 'crypted' and 'rgd_bcd_condition'
- Audit bcdedit configuration for unauthorized changes disabling recovery
- Search for ransom note files on user desktops and other locations
- Monitor network connections to blockchain.info and related payment APIs
- Isolate affected systems and perform forensic analysis to confirm encryption scope
- Restore from known good backups and reset affected system configurations

Analyst Notes

Optional analyst observations, ticket numbers, customer notes, containment status, pending actions, etc.

Example 3: Large Dataset

For this illustration, a large dataset of 54,784 logs with a **domain compromise** hidden in the dataset is being examined:

When uploading a dataset of this size (in this case 10.5MB), MalwareLab.ai **first condenses tens of thousands of log entries into a behavioral intelligence summary before performing AI analysis**. Rather than analyzing every individual row, it identifies significant patterns, suspicious events, high-value artifacts, and other key indicators to maximize the effectiveness of the investigation while reducing token usage. This approach enables efficient analysis of large enterprise datasets without sacrificing meaningful investigative context.

⚠ Data Handling Notice
Do not upload Classified Information, Controlled Unclassified Information (CUI), For Official Use Only (FOUO), export-controlled data, Personally Identifiable Information (PII), Protected Health Information (PHI), customer confidential information, or other restricted data.

Drag and drop a log, script, JSON report, or IOC file
Or choose a file manually

Choose File No file chosen

Loaded file: domain_compromise_logset.xlsx

SHA256: 0ffc74f654c944fd1a50958ebee1281fcd0412b40b3f5fe1cc300133ab22a397

Spreadsheet / Log File Summary / Behavioral Intelligence Summary / Behavioral Investigation Summary:

Filename: domain_compromise_logset.xlsx
Total Rows: 54783
Columns: log_time, host, src_ip, username, event_id, event_type, process_name, command_line, parent_process, logon_type, dest_ip, dest_port, protocol, bytes_sent, bytes_received, status

Event Type Counts:

```
{
  "Process Create": 15446,
  "Logon": 7336,
```

Analyze Threat

SEVERITY

Critical

CONFIDENCE

95%

High Confidence

Confidence Factors

- Repeated execution of suspicious PowerShell commands with encoded payloads launched from winword.exe
- Multiple network connections from powershell.exe to suspicious external IPs with known malicious or low reputation scores
- Frequent use of rundll32.exe to execute a DLL from a public user directory, indicating possible malicious persistence or payload execution
- Numerous Process Access events involving mimikatz.exe targeting lsass.exe, consistent with credential dumping activity
- Execution of privilege enumeration commands such as 'net group "Enterprise Admins" /domain' and 'net localgroup administrators' from cmd.exe spawned by powershell.exe
- Audit Log Cleared events present, indicating possible defense evasion

Reducing Factors

- No ransomware indicators or critical events explicitly reported in the dataset
- No direct evidence of lateral movement or data exfiltration observed in the provided logs
- Private IP addresses used internally reduce confidence in external attacker infrastructure attribution

Threat Type

Domain Compromise / Credential Theft / Privilege Escalation

Executive Summary

The log data shows strong indicators of a domain compromise involving credential dumping, suspicious PowerShell execution with encoded commands, and privilege enumeration activities. The use of mimikatz.exe to access lsass.exe memory, combined with network connections to low-reputation external IPs and execution of DLLs from public user directories, supports a scenario of attacker foothold and credential theft. This is consistent with a multi-stage intrusion involving initial access via Office macros (winword.exe spawning powershell.exe), execution of encoded payloads, credential access, and privilege escalation attempts. The presence of audit log clearing further indicates attempts to evade detection.

Evidence Supporting Assessment

- powershell.exe processes spawned by winword.exe with encoded command lines indicating obfuscated payload execution
- Repeated network connections from powershell.exe to IPs 185.220.101.42 (AbuseIPDB score 100), 193.32.162.157 (AbuseIPDB score 28), and 45.137.21.9 with low VirusTotal reputations
- Multiple Process Access events where mimikatz.exe accesses lsass.exe memory (event_id 10) with command 'sekurlsa:logonpasswords'
- Process Create events executing cmd.exe with commands for privilege enumeration (net group, net localgroup, nltest)
- Audit Log Cleared events (event_id 1102) detected, indicating possible defense evasion
- Use of rundll32.exe to execute C:\Users\Public\update.dll repeatedly, suggesting persistence or payload execution

Ambiguity Notes

- No direct evidence of lateral movement or data exfiltration in the provided logs
- Lack of explicit ATT&CK tactic or technique mappings due to absence of detailed telemetry on lateral movement or persistence mechanisms
- No malware hashes or file signatures found in MalwareBazaar or VirusTotal for the uploaded file
- Private IP addresses limit external attribution confidence

Alternative Hypotheses

- Red team or penetration testing activity simulating domain compromise techniques
- Security research or malware analysis environment generating similar telemetry
- Compromised user executing malicious macros leading to observed behavior

Alternative Explanations

No reasonable alternative explanations identified.

Indicators of Compromise

IP Addresses

- 193.32.162.157
- 45.137.21.9
- 185.220.101.42
- 10.10.20.64
- 10.10.20.39
- 10.10.20.38
- 10.10.5.11

URLs

None

Ports

- 53
- 443
- 8443

Users

- CORP\jtoilolo

Domains

None

Hashes

- 5248a7ddfec96f072af043924c319b7fef2fe369
- 0ffc74f654c944fd1a50958e8ea1281fcd0412b40b3fbfe1cc300133ab22a397

Processes

- powershell.exe
- winword.exe
- rundll32.exe
- cmd.exe
- mimikatz.exe

MITRE ATT&CK Mapping

T1059.001 - PowerShell

T1140 - Deobfuscate/Decode Files or Information

T1003.001 - LSASS Memory

T1078 - Valid Accounts

T1086 - PowerShell

T1112 - Modify Registry

T1070.001 - Clear Windows Event Logs

T1087 - Account Discovery

T1010 - Application Window Discovery (possible/contextual)

Recommended Actions

- Investigate all powershell.exe processes spawned by winword.exe with encoded command lines for malicious payloads
- Review network connections to IPs 185.220.101.42, 193.32.162.157, and 45.137.21.9 for potential C2 activity
- Conduct memory analysis on affected hosts to detect credential dumping artifacts
- Audit and restrict use of mimikatz.exe and monitor for process access to lsass.exe
- Review and restore cleared audit logs and implement enhanced logging to prevent tampering
- Search for persistence mechanisms related to C:\Users\Public\update.dll and rundll32.exe usage
- Perform user account review and reset credentials for CORP\jtoilolo and other potentially compromised accounts
- Hunt for lateral movement and privilege escalation activities across the domain
- Implement endpoint detection rules for encoded PowerShell execution and suspicious rundll32.exe DLL loads

Analyst Notes

Optional analyst observations, ticket numbers, customer notes, containment status, pending actions, etc.

Supported Inputs

MalwareLab.ai currently supports analysis of:

- Malware samples
- Executables
- Scripts
- PowerShell commands
- Log filesLarge CSV security datasets
- Windows Event Log exports
- SIEM exports
- Firewall logs
- Authentication logs
- Process execution logs
- Command-line output
- Incident reports
- Security alerts
- Indicators of compromise
- Threat intelligence artifacts
- Natural language descriptions of suspicious activity

Users may also describe suspicious behavior without uploading files or logs.

For example:

"A user account logged into multiple countries within a short period of time."

or

"A workstation began uploading large amounts of data to an unauthorized cloud storage provider."

The platform can provide investigative guidance even when evidence is limited.

You Don't Need Files or Logs

One unique capability of MalwareLab.ai is that users can describe suspicious activity even when logs, malware samples, or forensic evidence are not immediately available.

The platform can evaluate the scenario, identify potential risks, and provide investigative guidance.

Example: Potential Data Exfiltration

A user submits:

"An employee is uploading large amounts of data to drive.google.com. We do not use Google Drive in our organization. We use Microsoft SharePoint."

MalwareLab.ai may identify:

- Potential unauthorized cloud storage usage
 - Possible data exfiltration
 - Insider threat concerns
 - Shadow IT activity
 - Recommended investigative steps
-

Example: Suspicious PowerShell Activity

A user submits:

"A workstation launched PowerShell and immediately connected to an external IP address."

MalwareLab.ai may identify:

- Command-and-control activity
 - Malware execution
 - Remote payload retrieval
 - Living-off-the-land techniques
-

Example: Impossible Travel

A user submits:

"A user account logged in from Maryland and then successfully authenticated from Eastern Europe ten minutes later."

MalwareLab.ai may identify:

- Potential credential compromise
 - Impossible travel activity
 - Session theft
 - Unauthorized access
-

Example: Large-Scale File Encryption

A user submits:

"Thousands of files were renamed and encrypted within a few minutes."

MalwareLab.ai may identify:

- Potential ransomware activity
 - Data destruction concerns
 - Business disruption risks
 - Immediate containment considerations
-

Example: Unauthorized Administrative Changes

A user submits:

"A help desk account created several new administrator accounts overnight."

MalwareLab.ai may identify:

- Privilege escalation
- Account compromise
- Unauthorized administrative activity

- Persistence mechanisms

Example: Suspicious PowerShell Email Attachment

A user submits:

"A user opened an email attachment. PowerShell launched and shortly afterward several antivirus alerts appeared."

MalwareLab.ai may identify:

- Initial access activity
 - Script-based malware execution
 - Payload delivery
 - Potential compromise
-

Building on Previous Analysis

Cybersecurity investigations rarely begin with complete information.

Analysts often start with a single observation and gather additional evidence over time.

MalwareLab.ai allows users to continuously build upon previous analysis by providing new facts, observations, and context as an investigation unfolds.

As new information becomes available, the platform reassesses the situation and refines its conclusions.

Example: Administrative Account Investigation

Initial Observation

"A help desk account created several new administrator accounts overnight."

Potential findings:

- Account compromise
- Privilege escalation
- Persistence activity
- Insider threat concerns

Additional Context

"The help desk employee was not working during those hours."

The analysis may shift toward:

- Unauthorized account usage
 - Credential theft
 - Increased likelihood of compromise
-

Additional Context

"The newly created administrator accounts immediately logged into multiple servers."

The analysis may further evolve toward:

- Active intrusion
 - Lateral movement
 - Enterprise-wide investigation requirements
-

Example: Data Exfiltration Investigation

Initial Observation

"A workstation uploaded several gigabytes of data to drive.google.com."

Potential findings:

- Cloud storage usage
 - Unusual data movement
 - Need for further investigation
-

Additional Context

"Our organization does not use Google Drive."

The analysis may identify:

- Shadow IT activity
 - Unauthorized cloud storage
 - Potential data exfiltration
-

Additional Context

"The employee submitted a resignation letter two days ago."

The analysis may identify:

- Elevated insider threat risk
 - Possible intellectual property theft
 - Increased urgency for investigation
-

Example: Ransomware Investigation

Initial Observation

"Thousands of files were renamed and encrypted within a few minutes."

Potential findings:

- Ransomware activity
 - Data destruction concerns
 - High-priority security incident
-

Additional Context

"A user reported a message demanding payment in Bitcoin."

The analysis may identify:

- Extortion activity
 - Strong ransomware indicators
 - Increased confidence in assessment
-

Additional Context

"Several servers are now showing the same behavior."

The analysis may identify:

- Active ransomware outbreak
- Lateral movement
- Enterprise-wide containment requirements

Types of Analysis Supported

Malware Analysis

Examples:

- Analyze malware samples
 - Explain malicious behavior
 - Identify suspicious actions
 - Explain execution flow
 - Summarize malware capabilities
-

Script Analysis

Examples:

- Analyze PowerShell scripts
 - Decode Base64 content
 - Explain obfuscated code
 - Explain command-line activity
 - Identify suspicious functionality
-

Log Analysis

Examples:

- Analyze system logs
 - Identify suspicious events
 - Summarize activity
 - Highlight anomalies
 - Explain attack timelines
-

Incident Analysis

Examples:

- Summarize incidents
 - Explain suspicious activity
 - Identify likely attack paths
 - Recommend investigative steps
 - Generate executive-level summaries
-

IOC Analysis

Examples:

- Review IP addresses
 - Analyze domains
 - Evaluate URLs
 - Examine file hashes
 - Explain relationships between indicators
-

IOC Enrichment

MalwareLab.ai can help users understand indicators of compromise by providing context around submitted indicators. This enrichment is provided by VirusTotal, MalwareBazaar, URLHaus, and AbuseIPDB.

Examples include:

- IP addresses
- Domains
- URLs
- File hashes

Rather than simply listing indicators, MalwareLab.ai attempts to explain:

- Why an indicator may be suspicious
- Potential attacker objectives
- Relationships between indicators
- Investigation considerations
- Recommended follow-up actions

Example:

A user submits:

185.220.101.1

MalwareLab.ai may provide:

- Threat context
- Observed suspicious characteristics
- Potential investigative considerations
- Recommended validation steps

The goal is to help analysts move beyond identifying indicators and toward understanding their significance.

Not every indicator will appear in every threat intelligence source. Indicators that cannot be enriched should not automatically be considered benign. MalwareLab.ai evaluates all available evidence and incorporates enrichment when external intelligence is available.

Protecting Sensitive Organizational Information

MalwareLab.ai analyzes cybersecurity information, but users should avoid unnecessarily exposing sensitive organizational data to the platform.

Whenever possible, anonymize or remove:

- Internal IP addresses
- Host names
- Usernames
- Email addresses
- Internal domains
- Employee names
- Asset identifiers
- Sensitive business information

In most cases, the behavioral information is more important than the actual names.

Original Log Entry

2026-06-10 14:22:31
User: jsmith
Hostname: HR-PAYROLL-SERVER01
Source IP: 10.20.30.15
Destination IP: 172.16.5.20
PowerShell launched cmd.exe

Anonymized Version

2026-06-10 14:22:31
User: USER-001
Hostname: SERVER-001
Source IP: INTERNAL-IP-1
Destination IP: INTERNAL-IP-2
PowerShell launched cmd.exe

The behavior remains intact while organizational details are protected.

Preserve Behavioral Context

Keep:

- Timestamps
- Process names
- Commands
- Network behavior
- Authentication activity
- Event sequences

Avoid reducing information to vague descriptions.

Good:

powershell.exe launched rundll32.exe which connected to an external IP address

Less useful:

Something suspicious happened on a computer

The more behavioral detail provided, the more effective the analysis will be.

Understanding Your Results

MalwareLab.ai may provide:

Executive Summary

A concise overview of findings.

Technical Analysis

A detailed explanation of observed behavior.

Indicators of Compromise (IOCs)

Potential indicators such as:

- IP addresses
 - Domains
 - URLs
 - File hashes
 - Registry paths
 - File paths
-

MITRE ATT&CK Mapping

Identification of ATT&CK techniques associated with observed behavior.

Understanding Confidence Levels

Confidence levels indicate how strongly the available evidence supports a conclusion.

High Confidence

Multiple indicators support the assessment and alternative explanations appear unlikely.

NOTE: A high confidence assessment does not necessarily indicate a high-severity threat. Confidence reflects how strongly the available evidence supports the assessment, while severity reflects the potential impact of the observed activity.

Medium Confidence

Evidence supports the assessment, but additional investigation may be required.

Low Confidence

Limited information is available and multiple explanations remain possible.

Confidence levels should be considered alongside the evidence provided in the report.

Confidence Assessment

An explanation of how strongly the available evidence supports the conclusions.

Recommended Next Steps

Suggested actions that may help contain, investigate, validate, or remediate suspicious activity.

Examples:

- Collect additional logs
 - Review authentication events
 - Isolate affected systems
 - Investigate network communications
 - Expand the scope of investigation
-

Known Limitations

MalwareLab.ai analyzes information provided by users and may not always have complete visibility into an environment.

Factors that may impact analysis include:

- Incomplete logs
- Missing context
- Obfuscated code
- Encrypted payloads
- Novel malware
- Inaccurate user-provided information
- Limited forensic evidence

The platform provides analysis based on available evidence and should be used as part of a broader investigative process.

AI-Generated Analysis Disclaimer

MalwareLab.ai uses artificial intelligence to generate analysis, recommendations, summaries, and investigative guidance.

While every effort is made to provide accurate and useful information, AI-generated content may occasionally contain inaccuracies, incomplete conclusions, or alternative interpretations.

Users should review findings, validate critical conclusions, and follow established organizational procedures before taking operational, legal, disciplinary, or business actions.

MalwareLab.ai is intended to assist cybersecurity professionals—not replace professional judgment.

Best Practices

To obtain the most accurate analysis:

- Upload complete information whenever possible.
 - Provide context surrounding suspicious activity.
 - Add new information as investigations evolve.
 - Preserve behavioral details when anonymizing data.
 - Review AI-generated findings before taking operational action.
 - Use MalwareLab.ai as an investigation accelerator and decision-support tool.
-

Final Thoughts

MalwareLab.ai was built to help cybersecurity professionals answer three critical questions:

What am I looking at?

What should I do next?

What information should I collect next?

Whether you are analyzing malware, reviewing logs, investigating suspicious behavior, or working with only a handful of observations, MalwareLab.ai helps transform raw information into understandable cybersecurity intelligence.

No specialized query language. No complicated workflows. Just cybersecurity analysis designed to help defenders move faster, investigate smarter, and make better decisions.

Appendix A:

User Tips:

1. You can log into MalwareLab.ai while remotely logged into a client workstation.
 - a. Example 1: You are an MSP or MSSP worker and a client reports a suspicious file. While using remote support software, use the client's browser to log into Malwarelab.ai and analyze the file.
 - i. This example allows the MSP or MSSP worker to examine the threat without putting other workstations or organizations in jeopardy.
 - ii. If you use this method, email the report you downloaded to yourself from the client workstation.
 - iii. NOTE: Remember to log out of MalwareLab.ai on the client workstation when complete.
 - b. Example 2: You can use MalwareLab.ai inside a sandboxed environment, such as a virtual machine where you have a dynamic malware lab set up, as long as it is connected to the Internet.
2. As of this time, reports and analysis are not saved to the user account. Once a file is analyzed, if you wish to keep the analysis be sure to download either the markdown or PDF report (or both if you wish).
3. Remember that the markdown file is editable. This allows you to build upon your analysis and take additional notes.
4. The PDF file is ideal to keep the integrity of an analysis when presenting to leadership. If using the PDF download option, be sure to input any notes in the "Analyst Notes" section before downloading. However, if you forget to add notes, you can always redownload the PDF so long as you have not refreshed the page or executed another analysis.
5. If you are unsure what information to provide, begin with the facts you know. MalwareLab.ai is designed to work with incomplete information and can refine its analysis as additional context is provided. If this is your situation, refer to the "Examples" section that starts on page 22 and continue on to the "Building on Previous Analysis" section that starts on page 25 as a guide.
6. If the MalwareLab.ai web application crashes, a simple page refresh will restore the web application.
7. Large log files can often provide better investigative results than isolated events:

- a. Rather than uploading a single suspicious log entry, consider uploading the complete exported dataset whenever practical. MalwareLab.ai is designed to identify relationships between events, summarize recurring patterns, and build investigative context across large collections of security data.