



Course Lessons

Hands-On Network Threat Detection and Response Using ACAS

CyberSuite Academy — Text Lessons & Quizzes

Produced by Montimage · June 2026

Each section is one Academy text-lesson page followed by its quiz. Answer keys are shown in green for the author and are not displayed to learners.

Introduction — Hands-On Network Threat Detection and Response Using ACAS

Estimated time: ~15 minutes · Text lesson

Welcome

Network threats are evolving rapidly, and organizations need intelligent systems to detect anomalies before they become incidents. In this course you will learn how to analyze network traffic, detect threats using both traditional rule-based and modern AI-driven approaches, and interpret security alerts through Explainable AI (XAI) techniques, all using the Advanced Cybersecurity Analytics Service (ACAS) platform.

ACAS is an AI-driven Network Threat Detection and Response platform designed for near real-time anomaly detection, threat prediction, and explainable analysis of encrypted network traffic. It combines traditional rule-based detection with advanced deep learning models to identify malicious activities, and its Explainable AI module helps security analysts understand why specific traffic is flagged as suspicious.

Everything is hands-on. You will work with real PCAP files, explore extracted features, evaluate detection models, and interpret AI predictions. By the end you will be able to perform deep packet inspection, compare detection approaches, and use XAI methods to provide actionable insights for incident response.

What you will learn

- Understand network threats and anomaly detection approaches (rule-based, statistical, machine learning)
- Navigate the ACAS platform and work with PCAP files
- Perform deep packet inspection and analyze traffic statistics and protocol hierarchies
- Extract and interpret flow-based features used for machine learning models
- Evaluate and compare rule-based vs. deep learning detection methods
- Apply Explainable AI methods (SHAP, LIME) to interpret model predictions

Who this course is for

This course suits IT and cybersecurity students, SOC analysts, network security engineers, and professionals seeking practical skills in AI-driven threat detection. A basic understanding of networking (IP addresses, ports, protocols) and familiarity with cybersecurity concepts are recommended. No prior experience with machine learning or AI is required, every concept is introduced before it is used.

How the course is structured

The course is organised as five modules plus a final hands-on assessment. Each module is a short lesson followed by a quiz:

- Module 1 — Network Threats and Anomaly Detection Fundamentals
- Module 2 — Getting Started with ACAS
- Module 3 — Deep Packet Inspection & Feature Extraction
- Module 4 — Anomaly Detection with Rule-Based and Deep Learning Approaches
- Module 5 — Advanced Explainable AI (XAI) Analysis
- Final Hands-On Assessment (Capstone)

You can follow the course at your own pace; the full effort is about 20 hours, ideally spread over three to four weeks.

Assessment and certificate

Each module ends with a quiz; the pass mark is 7/10. After the capstone you complete a final quiz that draws on all the modules. Complete every module and the final assessment with the required grades and you earn a digital certificate of completion from the CyberSuite Academy.

A note on safety and ethics

When working with network traffic data, you may encounter sensitive information. Treat all data as confidential. The sample PCAP files provided for this course are sanitized datasets designed for educational purposes. Always follow your organization's data handling policies when working with real network captures.

Meet the team

The ACAS platform and this course are produced by Montimage. For questions about the course, contact Manh Nguyen (manhdung.nguyen@montimage.eu) or Edgardo Montes de Oca (edgardo.montesdeoca@montimage.eu).

Module 1: Network Threats and Anomaly Detection

Fundamentals

Estimated time: ~45 minutes · Text lesson + quiz

In this lesson

- Define network threats and explain the importance of threat detection for organizations
- Describe the different types of network attacks (DDoS, intrusion, malware, data exfiltration)
- Explain deep packet inspection (DPI) and its role in understanding network traffic
- Compare anomaly detection approaches: rule-based, statistical, and machine learning
- Understand the concept of Network Threat Detection and Response platforms

What are network threats?

A network threat is any malicious activity or event that could compromise the confidentiality, integrity, or availability of networked systems and data. Unlike vulnerabilities (which are weaknesses), threats are active attempts to exploit those weaknesses.

Organizations face an ever-growing landscape of network threats:

- Volume — millions of packets traverse networks every second
- Sophistication — attackers use advanced techniques to evade detection
- Speed — attacks can compromise systems in minutes or seconds
- Encryption — over 80% of traffic is now encrypted, hiding malicious content

This is why automated, intelligent detection systems are essential. Human analysts cannot manually inspect every packet; they need tools that surface the threats that matter.

Types of network attacks

Network attacks fall into several categories. Understanding these helps you recognize patterns in traffic data.

Attack Type	Description	Example	Network Indicator
DDoS	Overwhelm resources with traffic	SYN flood, UDP flood	High packet rates, unusual source distribution
Intrusion	Unauthorized access attempts	Brute force, exploitation	Failed logins, unusual port access
Malware	Malicious software communication	C2 beaconing, ransomware	Failed logins, unusual port access
Data Exfiltration	Unauthorized data transfer	DNS tunneling, covert channels	Large outbound transfers, unusual protocols
Reconnaissance	Network scanning and mapping	Port scans, service enumeration	Sequential port access, ICMP probes
Man-in-the-Middle	Traffic interception	ARP spoofing, SSL stripping	Duplicate MACs, certificate anomalies

In this course, you will work with PCAP files containing various attack scenarios to understand how these threats manifest in network traffic.

Deep Packet Inspection (DPI)

Deep Packet Inspection is the examination of both the header and payload of network packets as they pass through an inspection point. Unlike simple packet filtering (which only looks at headers), DPI can:

- Identify protocols — recognize applications regardless of port (e.g., HTTP on port 8080)
- Extract content — analyze payload data for signatures or patterns
- Reconstruct sessions — follow conversations across multiple packets
- Detect anomalies — find unusual patterns that deviate from normal behavior

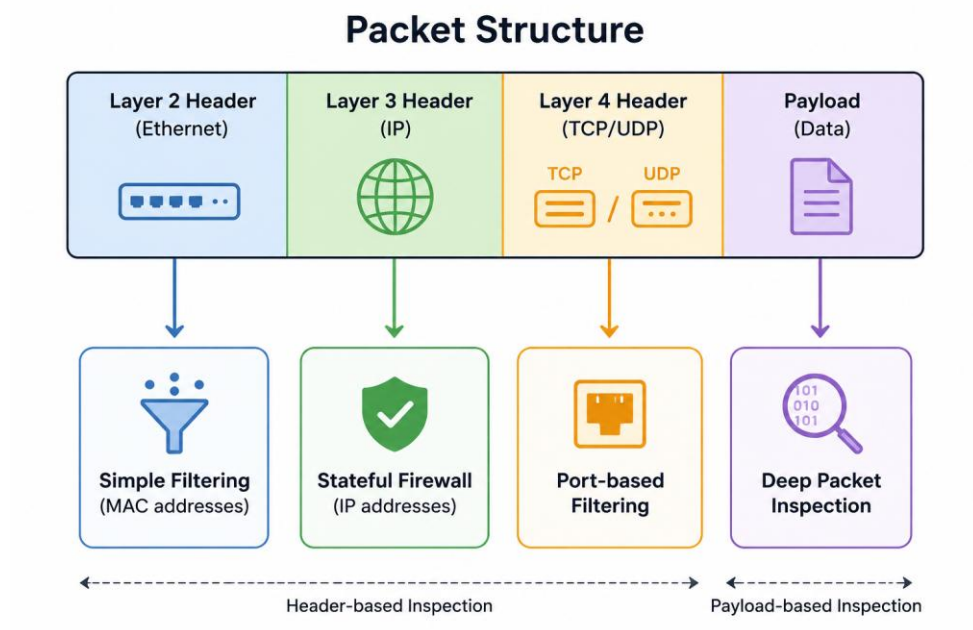


Figure 1: Packet structure.

DPI examines all layers, including the payload, to make intelligent decisions about traffic classification and threat detection. DPI outputs in ACAS:

- Traffic statistics (packet counts, byte volumes, timing)
- Protocol hierarchies (distribution of protocols)
- Bidirectional flow visualization
- Session reconstruction

Anomaly detection approaches

Anomaly detection identifies patterns that deviate from expected behavior. Two main approaches exist: Rule-based (Signature) Detection and Machine Learning Detection.

1. Rule-based (Signature) Detection

Rule-based detection matches traffic against known attack signatures — predefined patterns that indicate malicious activity.

Advantages:

- Low false positive rate for known attacks
- Fast and computationally efficient
- Easy to understand and audit

Disadvantages:

- Cannot detect unknown (zero-day) attacks
- Requires constant signature updates
- Attackers can evade by modifying patterns

Example rule: *"Alert if more than 100 failed SSH login attempts from the same IP within 60 seconds"*.

2. Machine Learning Detection

Machine learning models learn complex patterns from data, including deep learning approaches like CNNs and Autoencoders.

Advantages:

- Detects subtle, complex attack patterns
- Handles high-dimensional feature spaces
- Can identify unknown attacks

Disadvantages:

- Requires quality training data
- "Black box" — harder to interpret
- Computationally intensive

Example: *"A trained autoencoder flags traffic that cannot be reconstructed well, indicating anomalous patterns"*.

Modern platforms like ACAS combine two approaches to provide defense in depth. The table shows a detailed comparison of the two approaches.

Approach	Known Attacks	Unknown Attacks	False Positives	Interpretability	Maintenance
Rule-based	Excellent	Poor	Low	High	High
Machine Learning	Excellent	Good	Variable	Low	Medium

Network Threat Detection and Response

Network Threat Detection and Response (NDR) is a category of security solutions that monitor network traffic to detect threats and enable rapid response. The NDR platforms go beyond traditional intrusion detection systems (IDS) by:

- Using AI/ML — employing machine learning for advanced threat detection
- Analyzing behavior — focusing on behavioral patterns, not just signatures
- Providing context — correlating events to understand attack chains
- Enabling response — offering tools to investigate and respond to threats

Capability	Traditional IDS	NDR Platform
Detection method	Primarily signatures	Signatures + ML + behavior
Encrypted traffic	Limited	Advanced analysis
Response capabilities	Alert only	Alert + investigation + response
Threat hunting	Manual	Assisted/automated
Explainability	Rule matched	XAI insights

ACAS is a network threat detection platform that combines rule-based detection, deep learning models, and Explainable AI to provide comprehensive network threat detection and analysis.

Summary

- Network threats are malicious activities targeting the confidentiality, integrity, or availability of systems; they range from DDoS attacks to sophisticated intrusions and data exfiltration.
- Deep Packet Inspection (DPI) examines both headers and payloads to identify protocols, extract content, and detect anomalies; it is fundamental to understanding network behavior.
- Two anomaly detection approaches exist: rule-based (known signatures), statistical (baseline deviations), and machine learning (learned patterns), each with distinct trade-offs.
- Modern NDR platforms like ACAS combine multiple detection approaches with AI/ML capabilities and Explainable AI to detect both known and unknown threats.

Quiz: Network Threats and Anomaly Detection Fundamentals

Pass mark: 7/10. Choose the best answer.

1. What distinguishes deep packet inspection (DPI) from simple packet filtering?

- A) DPI only examines IP headers
- B) DPI examines both headers and payload content
- C) DPI is faster than packet filtering
- D) DPI only works with encrypted traffic

Answer: B — DPI examines all layers including the payload, while simple filtering only looks at headers.

2. Which detection approach is BEST suited for identifying previously unknown (zero-day) attacks?

- A) Rule-based detection
- B) Signature matching
- C) Machine learning detection
- D) Port-based filtering

Answer: C — Machine learning can identify patterns in unknown attacks that don't match existing signatures.

3. What is a key disadvantage of rule-based (signature) detection?

- A) High false positive rate
- B) Cannot detect known attacks
- C) Cannot detect unknown attacks without signature updates
- D) Requires too much computational power

Answer: C — Rule-based detection only catches attacks with known signatures; new attacks are missed.

4. Which type of attack involves overwhelming network resources with excessive traffic?

- A) Man-in-the-Middle
- B) Data exfiltration
- C) DDoS (Distributed Denial of Service)
- D) Reconnaissance

Answer: C — DDoS attacks flood resources with traffic to make them unavailable.

5. What does NDR stand for in cybersecurity?

- A) Network Data Recovery

- B) Network Detection and Response
- C) Node Distribution Router
- D) Network Defense Routine

Answer: B — NDR stands for Network Detection and Response.

6. Statistical anomaly detection establishes what to identify threats?

- A) Known attack signatures
- B) Baselines of normal behavior
- C) Encryption keys
- D) User credentials

Answer: B — Statistical methods establish baselines and flag deviations from normal patterns.

7. Which attack type typically shows periodic connections to external servers?

- A) DDoS attack
- B) Reconnaissance scan
- C) Malware C2 (Command and Control) beaconing
- D) ARP spoofing

Answer: C — Malware often "beacons" to C2 servers at regular intervals for instructions.

8. What is a primary advantage of rule-based detection over machine learning?

- A) Better at detecting unknown attacks
- B) Higher interpretability and lower false positives for known attacks
- C) Requires no maintenance
- D) Works better with encrypted traffic

Answer: B — Rules are easy to understand, audit, and have low false positives for known patterns.

9. Which layer does deep packet inspection examine that simple filtering does not?

- A) Layer 1 (Physical)
- B) Layer 2 (Data Link)
- C) Layer 7 (Application) payload
- D) Layer 3 (Network) only

Answer: C — DPI examines application layer payloads, not just network headers.

10. Why do modern NDR platforms combine multiple detection approaches?

- A) To increase computational cost
- B) To provide defense in depth — each approach covers different threat types
- C) Because regulations require it

D) To make the system more complex

Answer: B — Combining approaches (rules + statistics + ML) provides comprehensive coverage.

Module 2: Getting Started with ACAS

Estimated time: ~45 minutes · Text lesson + quiz

In this lesson

- Describe what the ACAS platform is and the problem it solves
- Understand the platform's key features and capabilities
- Access the platform and create a user account
- Navigate the user interface and understand its main sections
- Upload and work with PCAP files for analysis

What ACAS is, and the problem it solves

ACAS, the Advanced Cybersecurity Analytics Service, is an AI-driven Network Threat Detection and Response platform developed by Montimage. It is designed for near real-time anomaly detection, threat prediction, and explainable analysis of network traffic, including encrypted traffic.

The problem: Security analysts face overwhelming volumes of network data. Traditional tools generate alerts but provide little context about why traffic is suspicious or what to do about it. Analysts spend hours investigating false positives while real threats slip through.

The solution: ACAS combines multiple detection approaches (rule-based and deep learning) with Explainable AI to:

- Automatically analyze network traffic from PCAP files
- Extract meaningful features for machine learning
- Detect anomalies using both traditional and AI-based methods
- Explain why traffic is flagged as suspicious using SHAP and LIME
- Provide actionable recommendations through an LLM-powered assistant

Think of ACAS as an intelligent analyst assistant: it processes the data, surfaces the threats, and helps you understand the "why" behind each detection.

Platform architecture overview

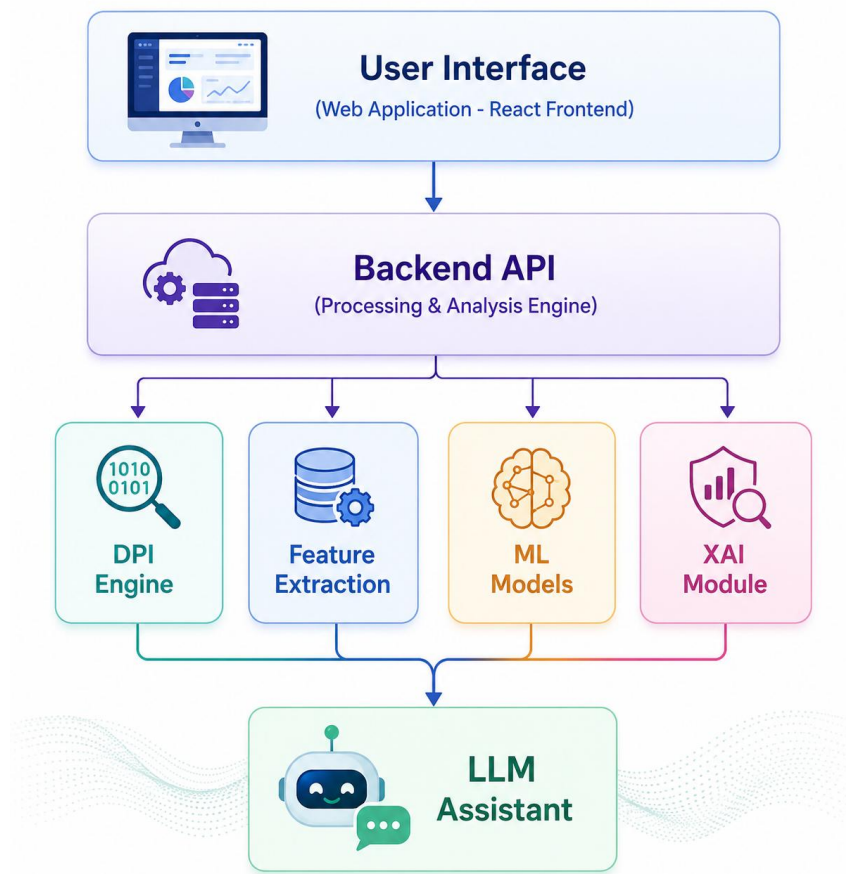


Figure 2: Overview architecture of the platform.

Components:

- **DPI Engine** — performs deep packet inspection on uploaded PCAP files
- **Feature Extraction** — extracts 60+ flow-based features for ML analysis
- **ML Models** — prebuilt detection models (CNNs, Autoencoders) for anomaly detection
- **XAI Module** — SHAP and LIME implementations for model interpretability
- **LLM Assistant** — provides natural language explanations and recommendations

Key features and capabilities

Feature	Description	Benefit
Deep Packet Inspection	Analyze packet contents, protocols, and flows	Understand network behavior in detail
Feature Extraction	Extract 60+ flow-based features automatically	Prepare data for ML without manual work
ML (Prebuilt) Models	CNN and Autoencoder models for detection	Detect anomalies without training your own models

Rule-based Detection	Traditional signature matching	Catch known attack patterns reliably
XAI Analysis	Global/Local feature importance & explanations	Understand which features drive predictions, Explain individual predictions
LLM Assistant	AI-powered recommendations	Get actionable mitigation guidance

Accessing the platform

The ACAS platform is accessible at: <https://ndr.montimage.eu>.

Creating an account:

1. Navigate to the platform URL in your web browser
2. Click "Sign in" using your Gmail/Github or "Sign Up"
3. Enter your email address and create a password
4. Verify your email if required
5. Log in with your credentials

Note: While registration provides access to all non-admin features, you can also explore some features without logging in. However, logging in is recommended to access the full functionality.

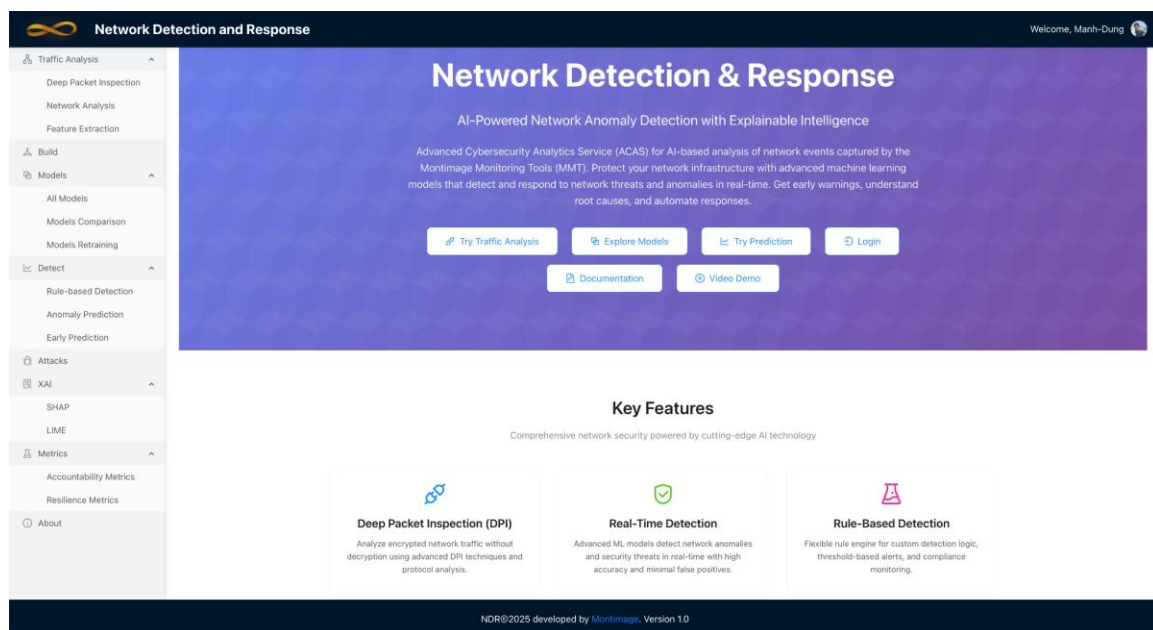


Figure 3: Landing page.

Working with PCAP files

PCAP (Packet Capture) files are the standard format for storing captured network traffic. ACAS analyzes these files (e.g., in .pcap or .pcapng) to perform all its functions. File size limit: 10 MB maximum for uploaded files

Uploading a PCAP file:

1. Navigate to the sections like Traffic Analysis or Detect
2. Click "Upload PCAP" or drag and drop your PCAP file
3. Wait for the upload to complete
4. The file will be processed automatically
5. Access results in the respective analysis sections

Sample PCAP files: The platform provides sample PCAP files for learning and testing, including:

- Normal traffic samples
- Various attack scenarios (DDoS, intrusion attempts, etc.)
- Mixed traffic for detection comparison

Best practices:

- Start with the provided sample files to learn the platform
- Use smaller files (< 10 MB) for faster processing
- Ensure PCAP files are properly formatted before upload

Hands-on activity: Platform exploration

Follow these steps to explore the ACAS platform:

1. Access the platform — Open <https://ndr.montimage.eu> in your browser
2. Register/Login — Create an account or log in
3. Explore the interface — Navigate through each section to understand the layout
4. Upload a sample file — Use one of the provided sample PCAP files
5. View DPI results — Examine the traffic statistics and protocol distribution
6. Note the features — Observe the extracted features for the uploaded file

Summary

- ACAS is an AI-driven NDR platform that combines DPI, feature extraction, multiple detection approaches, and Explainable AI to detect and explain network threats.
- The platform architecture includes a DPI engine, feature extraction module, ML models, XAI components, and an LLM assistant working together.
- Access the platform at <https://ndr.montimage.eu> and create an account to access full features.
- The interface is organized into sections for upload, DPI, feature extraction, detection, and explainability.

- PCAP files are the input format; the platform supports files up to 10 MB and provides sample files for learning.

Quiz: Getting Started with ACAS

Pass mark: 7/10. Choose the best answer.

1. What does ACAS stand for?
A) Automated Cybersecurity Alert System
B) Advanced Cybersecurity Analytics Service
C) Active Cyber Attack Scanner
D) Anomaly Classification and Security

Answer: B — ACAS stands for Advanced Cybersecurity Analytics Service.

2. What is the primary input format for analysis in ACAS?
A) CSV files
B) JSON logs
C) PCAP files
D) Syslog data

Answer: C — ACAS analyzes PCAP (packet capture) files containing network traffic.

3. What is the maximum file size for PCAP uploads in ACAS?
A) 1 MB
B) 5 MB
C) 10 MB
D) 100 MB

Answer: C — The platform supports PCAP files up to 10 MB.

4. Which component of ACAS provides natural language explanations and recommendations?
A) DPI Engine
B) Feature Extraction module
C) LLM Assistant
D) SHAP module

Answer: C — The LLM (Large Language Model) Assistant provides natural language recommendations.

5. What is the URL to access the hosted ACAS platform?

- A) <https://acas.montimage.eu>
- B) <https://ndr.montimage.eu>
- C) <https://cybersuite.eu/acas>
- D) <https://montimage.com/ndr>

Answer: B — The platform is accessible at <https://ndr.montimage.eu>

6. How many flow-based features does ACAS extract from PCAP files?

- A) 10+
- B) 30+
- C) 60+
- D) 100+

Answer: C — ACAS extracts over 60 flow-based features for machine learning analysis.

7. Which XAI methods are available in ACAS for model interpretability?

- A) Only SHAP
- B) Only LIME
- C) SHAP and LIME
- D) Random Forest importance only

Answer: C — ACAS provides both SHAP and LIME methods for explainability.

8. What types of prebuilt models are available in ACAS for anomaly detection?

- A) Only decision trees
- B) CNNs and Autoencoders
- C) Only linear regression
- D) Random forests only

Answer: B — ACAS includes prebuilt CNN and Autoencoder models for detection.

9. What problem does ACAS primarily solve for security analysts?

- A) Password management
- B) Understanding why traffic is flagged and getting actionable recommendations
- C) Firewall configuration
- D) User authentication

Answer: B — ACAS helps analysts understand detections through XAI and provides recommendations.

10. Is registration required to use any features of the ACAS platform?

- A) Yes, all features require registration
- B) No, registration is not available
- C) Some features work without login, but registration is recommended for full access
- D) Registration is only for administrators

Answer: C — While some exploration is possible, logging in provides access to all non-admin features.

Module 3: Deep Packet Inspection & Feature Extraction

Estimated time: ~60 minutes · Text lesson + quiz

In this lesson

- Perform deep packet inspection on PCAP files and interpret the results
- Analyze traffic statistics including packet counts, byte volumes, and timing
- Understand protocol hierarchies and their distribution in network traffic
- Visualize bidirectional flows and identify communication patterns
- Extract flow-based features and understand their role in machine learning
- Interpret feature distributions and identify important features for threat detection

Deep packet inspection in ACAS

In Module 1, we introduced deep packet inspection (DPI) conceptually. Now we put it into practice using ACAS. When you upload a PCAP file, the DPI engine processes every packet to extract meaningful information about your network traffic. What the DPI engine produces:

- Traffic Statistics — quantitative summaries of the captured traffic
- Protocol Hierarchies — breakdown of protocols present in the traffic
- Bidirectional Flows — visualization of communication between endpoints
- Temporal Patterns — how traffic volume changes over time

Traffic statistics

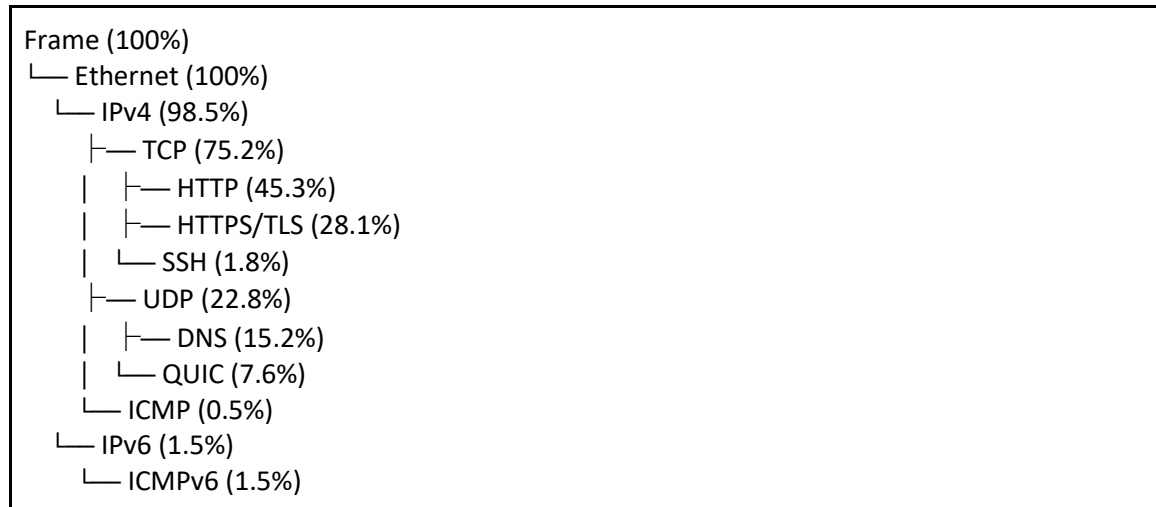
Traffic statistics provide a high-level view of the captured data. Several key metrics are listed in the following table. When analyzing statistics, look for:

- Unusual volumes — sudden spikes may indicate DDoS or data exfiltration
- Duration anomalies — very short or very long flows can be suspicious
- IP distribution — many sources to one destination suggests DDoS; one source to many destinations suggests scanning

Metric	Description	What it tells you
Total Packets	Number of packets in the capture	Overall traffic volume
Total Bytes	Sum of all packet sizes	Data transfer volume
Duration	Time span of the capture	Observation window
Packets/Second	Average packet rate	Traffic intensity
Bytes/Second	Average throughput	Bandwidth usage
Unique Source IPs	Distinct source addresses	Number of senders
Unique Dest IPs	Distinct destination addresses	Number of receivers
Unique Flows	Distinct 5-tuple combinations	Number of conversations

Protocol hierarchies

Protocol hierarchy analysis shows which protocols are present and their relative proportions. This is displayed as a tree structure reflecting the network layer model:



Bidirectional flow visualization

A flow (or bidirectional flow) represents a complete conversation between two endpoints. ACAS identifies flows using the 5-tuple (source IP address, destination IP address, source port, destination port, protocol).

What to look for in flows:

- Asymmetric flows — large difference between sent/received may indicate exfiltration
- Long-duration flows — persistent connections could be C2 channels
- Unusual ports — unexpected services on non-standard ports
- Flow counts — many flows from one source could indicate scanning

Feature extraction for machine learning

While DPI gives us human-readable insights, machine learning models need numerical features. ACAS automatically extracts over 60 flow-based features from PCAP files. As raw packets are not suitable for ML models, features transform packet data into numerical values that capture the essential characteristics of network behavior. Good features distinguish between normal and malicious traffic.

Feature Descriptions

Detailed metadata and explanations for each extracted feature

ID	Feature Name	Description	Type
1	ip.session_id	Session ID	numerical
2	meta.direction	Flow direction (0 means uplink and 1 means downlink)	categorical
3	ip.pkts_per_flow	Total number of IP packets	numerical
4	duration	Flow duration	numerical
5	ip.header_len	Total length of IP header	numerical
6	ip.payload_len	Total length of IP payload	numerical
7	ip.avg_bytes_tot_len	Average of total length of IP header	numerical
8	time_between_pkts_sum	Sum time between two flows	numerical
9	time_between_pkts_avg	Average time between two flows	numerical
10	time_between_pkts_max	Maximum time between two flows	numerical

Total 83 features < 1 2 3 4 5 ... 9 > 10 / page

Figure 4: Detailed metadata and explanations for extracted features.

Hands-on activity: DPI and feature extraction

1. Upload a PCAP file — Use a sample file from the platform
2. Examine traffic statistics — Note total packets, duration, and unique IPs
3. Review protocol hierarchy — Identify the dominant protocols
4. Explore bidirectional flows — Find the largest and longest flows
5. View extracted features — Examine at least 10 different features
6. Compare distributions — If multiple samples are available, compare feature distributions

Deep Packet Inspection

Analyze network traffic with protocol hierarchy, statistics, and flow patterns using mmit-probe

Configuration

Mode: Offline (PCAP)

PCAP File: infiltration.pcap Upload PCAP

View DPI View Network Extract Features

Traffic Analysis

Traffic Statistics

PCAP File

infiltration.pcap

Total Packets

23,265

Total Data

12.01 MB

Avg Packet Size

541.12 B

Duration

3m 5s

Throughput

543.91 Kbps

Packet Rate

125.64 pps

Figure 5. DPI configurations.

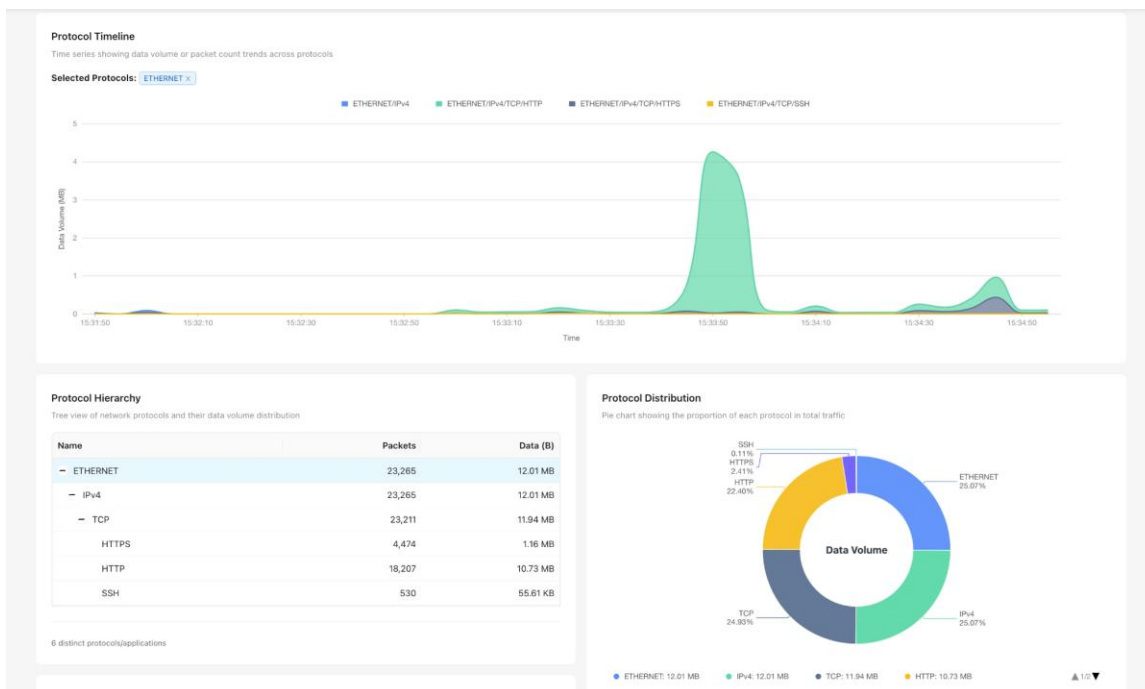


Figure 6. Visualization of selected DPI analysis results.

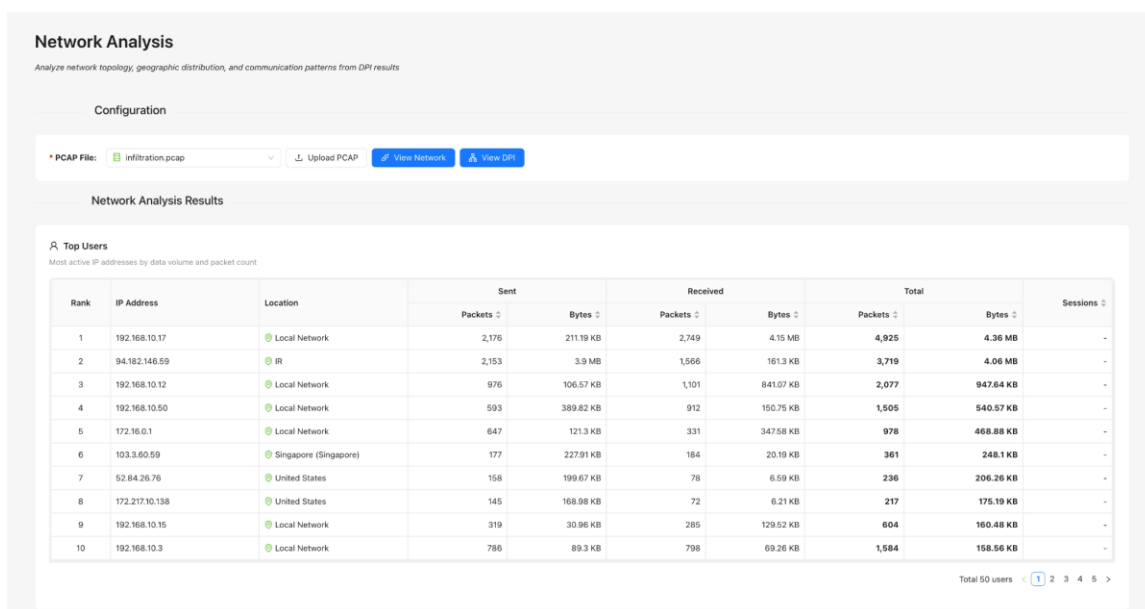


Figure 7. Network Analysis configurations.

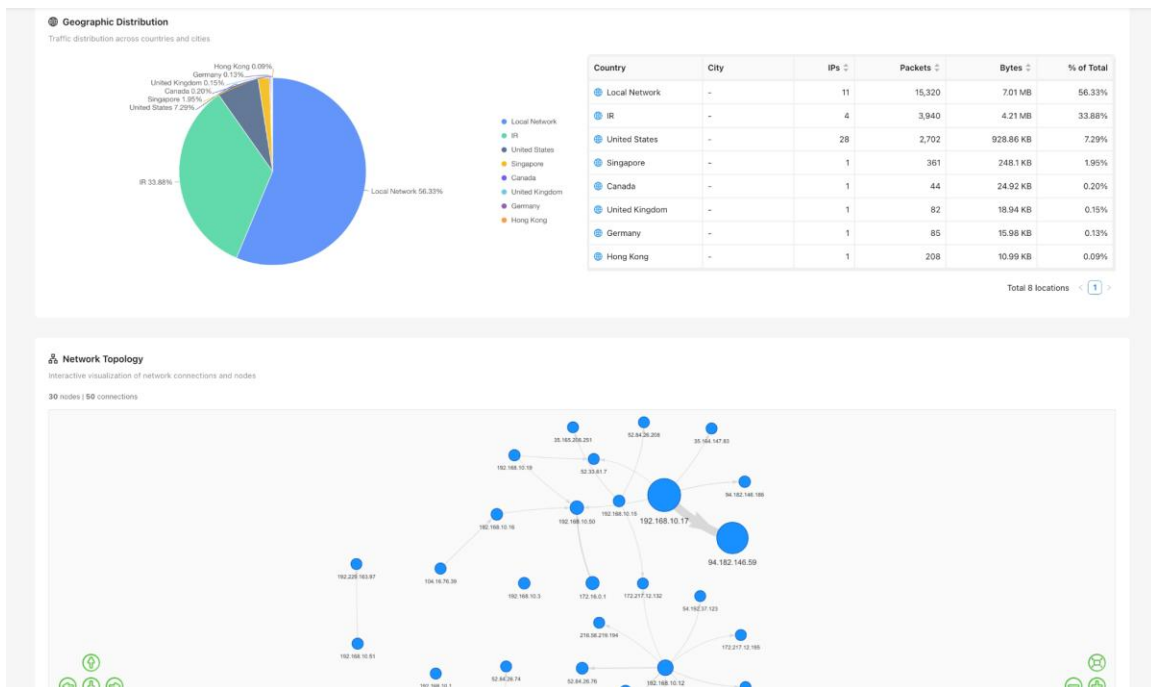


Figure 8. Visualization of selected Network Analysis results.

Feature Extraction

Upload a PCAP and extract features

Configuration

* PCAP File: Label: ☐ Normal ☐ Malicious

Extracted Features

Extraction Summary

PCAP File infiltration.pcap	Extracted Flows 1,754	Features 83
---	---------------------------------	-----------------------

Figure 9. Feature Extraction configurations.

Features Table
Tabular view of all extracted features for each network flow

[Download Features CSV](#) [Send all to NATS](#)

ip.session_id	meta.direction	ip.pkts_per_flow	duration	ip.header_len	ip.payload_len	ip.avg_bytes_tot_len	time_between_pkts_sum	time_between_pkts_avg	time_between_pkts_max	time_between_pkts_min	time
1	0	48	126.48031210899353	960	66096	1294.3846153846155	8636.890172956374	179.93521193663278	7230.845212936401	0.0	106
1	1	4	126.4803831577301	80	172	1294.3846153846155	46.77224159240723	11.693060398101807	21.499156951904297	0.07104873657226562	10.8
2	0	7	120.15009903907776	140	224	53.9375	177.68287658691406	25.383268083844865	36.9720458984375	0.2488626708984375	15.3
2	1	9	120.15015318009521	180	319	53.9375	56923.37679862976	6324.819644292195	56683.92086029053	0.04792213439941406	188
3	0	2	65.39945101737976	40	95	58.2	57.51800537109375	28.759002685546875	50.88496208190918	6.63304328918457	31.2
3	1	3	65.39949417114258	60	96	58.2	0.286102294921875	0.095367431840625	0.2140998840320312	0.028848648071289062	0.10
4	0	7	121.14837908744812	140	224	52.0	23.378372192382812	3.3397674660546875	6.425142288208008	0.1201629638671875	2.16
4	1	8	121.14843201637268	160	256	52.0	131.93599646606445	16.491949558258057	59.2498779296875	0.00095367431840625	26.4
5	0	7	121.14885210990906	140	224	52.0	3.2622814178466797	0.46604020254952566	0.61297416668701172	0.03504753112792969	0.20
5	1	8	121.1489040851593	160	256	52.0	2.360105514526367	0.2950131803157959	2.307891845703125	0.0	0.81

Total 1754 flows < 1 2 3 4 5 ... 176 > 10 / page

Feature Descriptions
Detailed metadata and explanations for each extracted feature

ID	Feature Name	Description	Type
1	ip.session_id	Session ID	numerical
2	meta.direction	Flow direction (0 means uplink and 1 means downlink)	categorical
3	ip.pkts_per_flow	Total number of IP packets	numerical
4	duration	Flow duration	numerical
5	ip.header_len	Total length of IP header	numerical
6	ip.payload_len	Total length of IP payload	numerical
7	ip.avg_bytes_tot_len	Average of total length of IP header	numerical
8	time_between_pkts_sum	Sum time between two flows	numerical
9	time_between_pkts_avg	Average time between two flows	numerical
10	time_between_pkts_max	Maximum time between two flows	numerical

Figure 10. Visualizations of selected Feature Extraction results.

Summary

- Deep packet inspection in ACAS produces traffic statistics, protocol hierarchies, bidirectional flow visualizations, and temporal patterns from PCAP files.
- Traffic statistics reveal volume, duration, and endpoint distribution; anomalies in these metrics can indicate threats.
- Protocol hierarchies show the breakdown of network protocols; unusual distributions warrant investigation.
- Bidirectional flows represent conversations between endpoints; asymmetric or long-duration flows may be suspicious.
- ACAS extracts 60+ flow-based features including packet lengths, inter-arrival times, flags, and rates for ML analysis.
- Feature importance varies by attack type; understanding which features matter helps interpret detection results.

Quiz: Module 3 — Deep Packet Inspection & Feature Extraction

Pass mark: 7/10. Choose the best answer.

1. What defines a network flow (5-tuple)?

- A) MAC address, IP address, port, protocol, timestamp
- B) Source IP, destination IP, source port, destination port, protocol
- C) Packet size, duration, direction, encryption, compression
- D) Username, password, IP, port, session ID

Answer: B — A flow is defined by source IP, destination IP, source port, destination port, and protocol.

2. Which metric would be MOST useful for detecting a potential DDoS attack?

- A) Total unique source IPs to a single destination
- B) Average packet length
- C) Number of DNS queries
- D) Protocol hierarchy percentage

Answer: A — Many unique sources targeting one destination is a classic DDoS pattern.

3. What does a high "down_up_ratio" feature value potentially indicate?

- A) Normal web browsing
- B) Possible data exfiltration (more upload than expected)
- C) DNS tunneling
- D) ICMP flood

Answer: B — An unusual upload/download ratio can indicate data being sent out.

4. What type of feature is "flow_iat_mean"?

- A) Packet length feature
- B) Flag-based feature
- C) Inter-arrival time feature
- D) Subflow feature

Answer: C — IAT stands for Inter-Arrival Time; this measures average time between packets.

5. In protocol hierarchy analysis, what might excessive DNS traffic indicate?

- A) Normal network behavior only
- B) Possible DNS tunneling as a covert channel
- C) Server misconfiguration
- D) IPv6 transition

Answer: B — While some DNS is normal, excessive DNS can indicate tunneling for data exfiltration.

6. Which TCP flag count would be unusually high during a SYN flood attack?

- A) FIN flag count
- B) ACK flag count
- C) SYN flag count
- D) PSH flag count

Answer: C — SYN floods send many SYN packets without completing the handshake.

7. What does an asymmetric bidirectional flow (large difference between sent and received packets) suggest?

- A) Normal HTTP request/response
- B) Possible data exfiltration or scanning
- C) Encrypted traffic
- D) Load-balanced connection

Answer: B — Large asymmetry can indicate one-way data transfer like exfiltration or scanning.

8. How many flow-based features does ACAS extract from PCAP files?

- A) 10+
- B) 30+
- C) 60+
- D) 200+

Answer: C — ACAS extracts over 60 flow-based features.

9. What does low "flow_iat_std" (inter-arrival time standard deviation) combined with regular intervals suggest?

- A) Random normal traffic
- B) Possible C2 beaconing (periodic communication)
- C) DDoS attack
- D) Protocol error

Answer: B — Very regular timing (low variability) suggests automated, periodic communication like C2 beacons.

10. Which visualization helps identify redundant features in your dataset?

- A) Histogram
- B) Box plot
- C) Correlation matrix
- D) Pie chart

Answer: C — Correlation matrices show relationships between features; highly correlated features may be redundant.

Module 4: Anomaly Detection with Rule-Based and Deep Learning Approaches

Estimated time: ~60 minutes · Text lesson + quiz

In this lesson

- Understand how rule-based detection works and configure detection rules
- Explain deep learning models for anomaly detection (CNNs and Autoencoders)
- Evaluate prebuilt models in ACAS and understand their configurations
- Interpret performance metrics: accuracy, precision, recall, F1-score
- Compare detection results from rule-based and deep learning approaches
- Analyze strengths and trade-offs of each detection method for different scenarios

Rule-based detection in ACAS

Rule-based detection uses predefined patterns (signatures) to identify known threats. When network traffic matches a rule, an alert is generated. Rules work by continuously monitoring network traffic as a stream of events and checking whether specific patterns of those events match predefined conditions. Each rule defines constraints on event types, attributes (like IPs, ports, or protocols), and sometimes timing or ordering between events.

When the observed traffic satisfies all the conditions of a rule, such as a sequence of actions happening in a specific order or multiple suspicious events occurring within a time window, the rule engine triggers an alert with a defined severity and message.

Example of a security rule in XML format that detects web scans with Nikto:

```
<beginning>
<!-- Property 15: Web scan with Nikto. User-Agent based detection-->
<property value="THEN" delay_units="ms" delay_min="0" delay_max="0" property_id="15"
type_property="ATTACK"
  description="Nikto detection" >
  <event event_id="1" value="COMPUTE"
    description="Context: an user agent in the HTTP header"
    boolean_expression="((http.method != '') & & ((http.user_agent != '')
& & (ip.src != ip.dst)))"/>
  <event event_id="2" value="COMPUTE"
    description="Trigger: Nikto detected. "
    boolean_expression="(#strstr(http.user_agent, 'Nikto') != 0)"/>
</property>
</beginning>
```

Advantages of rule-based detection:

- Predictable — you know exactly what triggers an alert
- Fast — simple pattern matching is computationally cheap
- Low false positives — for well-crafted rules targeting known attacks
- Auditable — rules can be reviewed and understood by humans

Limitations:

- No unknown attack detection — only catches what rules describe
- Maintenance burden — rules need constant updates for new threats
- Evasion — attackers can modify behavior to avoid matching rules

Deep learning models for anomaly detection

Deep learning models learn patterns from data rather than relying on predefined rules. ACAS includes two types of prebuilt models:

1. Convolutional Neural Networks (CNNs)

Originally designed for image processing, CNNs have been adapted for network traffic analysis. How CNNs work for network data:

Input Features → Convolution Layers (Pattern Extraction) → Pooling Layers (Feature Reduction) → Dense Layers (Decision Making) → Normal / Attack

- Convolution layers detect local patterns in the feature space
- Pooling layers reduce dimensionality and extract important patterns
- Dense layers combine patterns for final classification

CNN characteristics:

- Supervised learning (trained on labeled data)
- Binary or multi-class classification (normal vs. attack types)
- Good at detecting spatial patterns in features

2. Autoencoders

Autoencoders are neural networks trained to reconstruct their input. They detect anomalies by identifying traffic that cannot be reconstructed well.

How Autoencoders work:

Input Features → Encoder → Latent Space → Decoder → Reconstructed Feature → Reconstruction Error → Low Error (Normal) / High Error (Anomaly)

The reconstruction error represents the difference between the input data and the output generated by the model. A high reconstruction error indicates a potential anomaly, suggesting that the traffic does not conform to the patterns learned during training.

Autoencoder characteristics:

- Semi-supervised learning (trained mainly on normal data)
- Anomaly detection via reconstruction error
- Good at detecting novel attacks not seen in training

Evaluating prebuilt models in ACAS

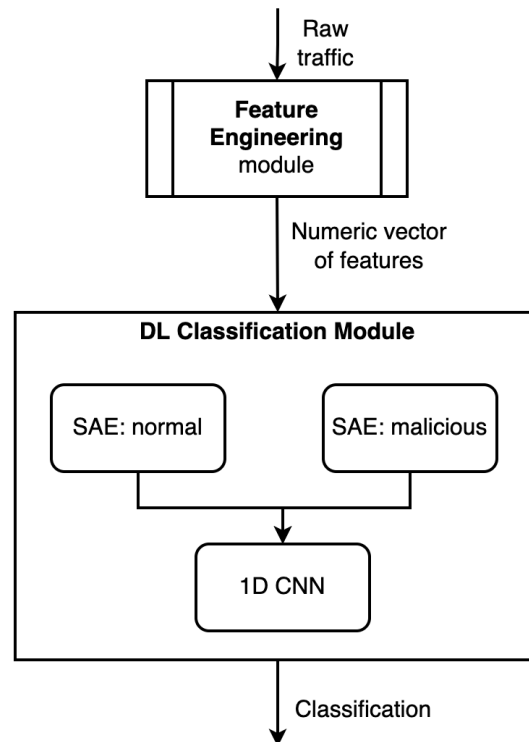


Figure 11: Overview of ML-based model for anomaly detection.

ACAS provides prebuilt models trained on public and private datasets with the overall architecture as above.

1. Upload a PCAP file (or select a previously uploaded file)
2. Choose the detection approach (rule-based, ML-based)
3. Run the analysis
4. View results with detected abnormal flows

Understanding performance metrics

To evaluate how well detection models perform, we use standard metrics, such as confusion matrix:

- TN = True Negative (correctly identified normal)
- FP = False Positive (normal flagged as attack)
- FN = False Negative (attack missed)
- TP = True Positive (correctly identified attack)

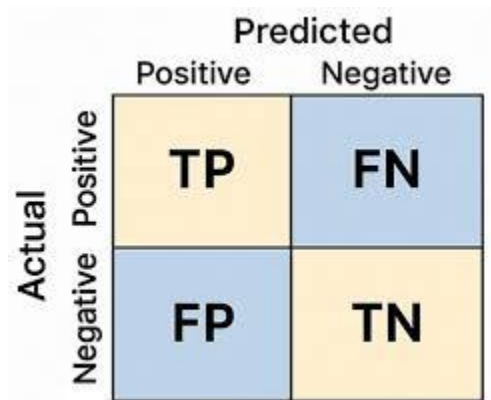


Figure 12: Confusion matrix.

Metric	Formula	What it measures
Accuracy	$(TP+TN)/(TP+TN+FP+FN)$	Overall correctness
Precision	$TP/(TP+FP)$	How many alerts are real attacks
Recall	$TP/(TP+FN)$	How many attacks are detected
F1-Score	$2 \times (Precision \times Recall) / (Precision + Recall)$	Balance of precision and recall

Interpreting metrics:

- High precision, low recall — Few false alarms, but missing attacks
- Low precision, high recall — Catching attacks, but many false alarms
- F1-Score — Good for imbalanced datasets; balances precision and recall

Example: Model results on test dataset: Accuracy: 95%, Precision: 87%, Recall: 92%, F1-Score: 89%.

Interpretation: The model correctly classifies 95% of traffic overall. When it raises an alert, 87% of the time it's a real attack (precision). It catches 92% of all attacks in the data (recall).

Strengths and trade-offs by scenario

Scenario	Best Approach	Reasoning
----------	---------------	-----------

Known attack patterns	Rule-based	Fast, accurate, low false positives
Zero-day attacks	Autoencoder	Detects deviations from normal
Multi-class classification	CNN	Trained to distinguish attack types
High-value targets	All combined	Defense in depth
Resource-constrained	Rule-based	Lower computational requirements
SOC alert fatigue	Rule-based + XAI	Fewer alerts, explainable

Hands-on activity: Detection comparison

1. Select a PCAP file
2. Run rule-based detection
3. Run ML-based detection
4. View metrics
5. Observe detected abnormal flows

All Models
All the machine learning models

[Delete All Models](#)

Model Id	Built At	Training Dataset	Testing Dataset	Actions
+	2024-02-23 09:11:11	Send to NATS	Send to NATS	Select action
model_3	2024-02-23 09:11:11	View Download Send to NATS	View Download Send to NATS	Select action
-	2023-09-26 15:54:30	View Download Send to NATS	View Download Send to NATS	Select action
model_2	2023-09-26 15:54:30	View Download Send to NATS	View Download Send to NATS	Select action

Build Configuration [JSON](#)

```
{
  "datasets": [
    {
      "csvPath": "report-bot-traffic.pcap-5da3849a-8385-4f38-ada8-a8585a763f9/1695734356.323230_0_bot-traffic.pcap.csv",
      "isAttack": true
    },
    {
      "csvPath": "report-normal-traffic.pcap-380ad2d5-ac60-4733-84f3-1dea23f82ca4/1695734366.296286_0_normal-traffic.pcap.csv",
      "isAttack": false
    }
  ],
  "training_ratio": 0.7,
  "training_parameters": {
    "nb_epoch_cm": 5,
    "nb_epoch_sae": 2,
    "batch_size_cm": 16,
    "batch_size_sae": 32
  }
}
```

Retrain
Predict
XAI
SHAP
LIME
Attacks
Metrics
Accountability

Figure 13. List of all prebuilt models.

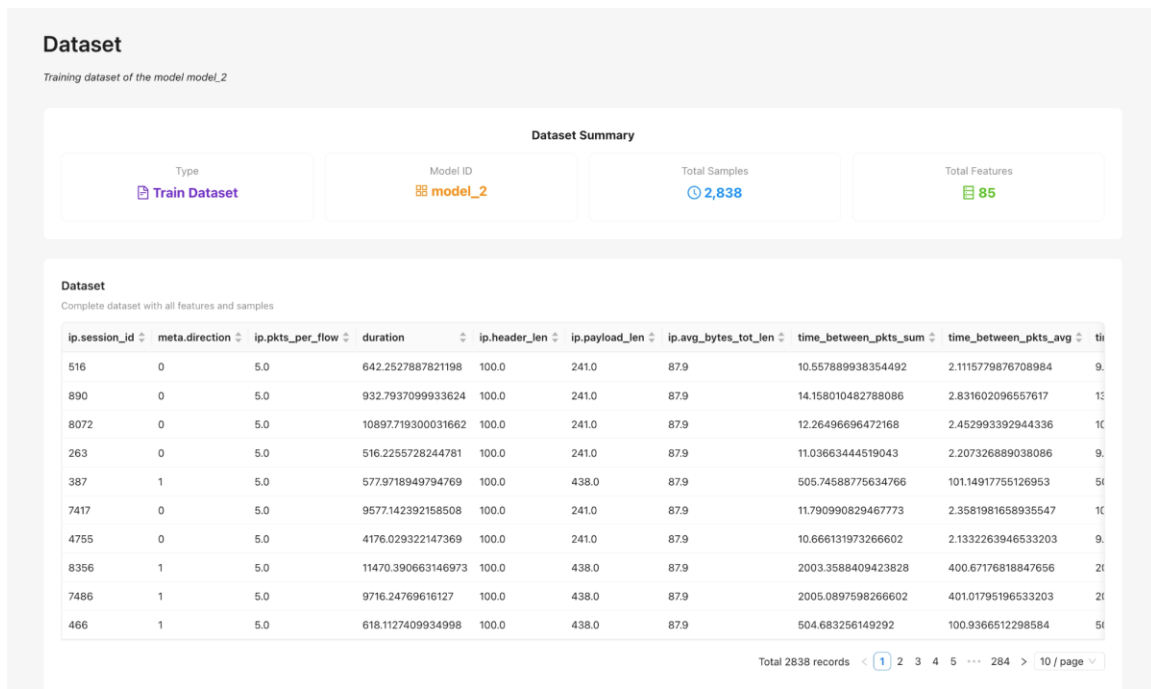


Figure 14. Visualizations of selected Dataset results.

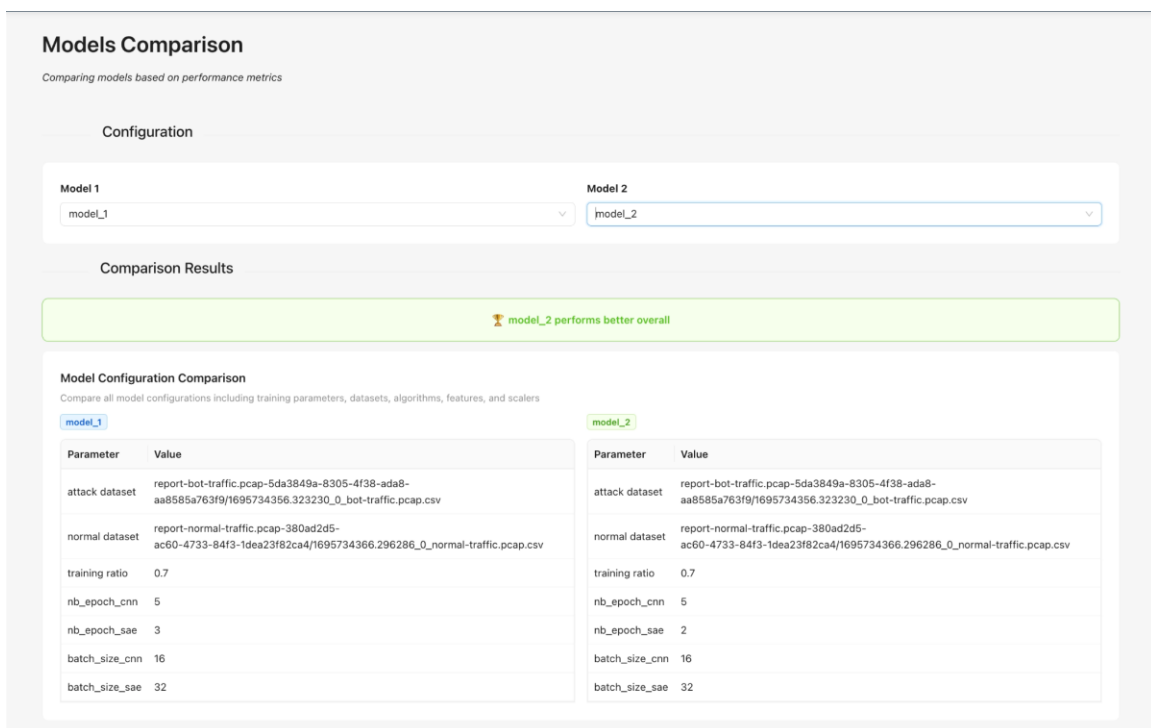


Figure 15. Models Comparison configurations and selected results.

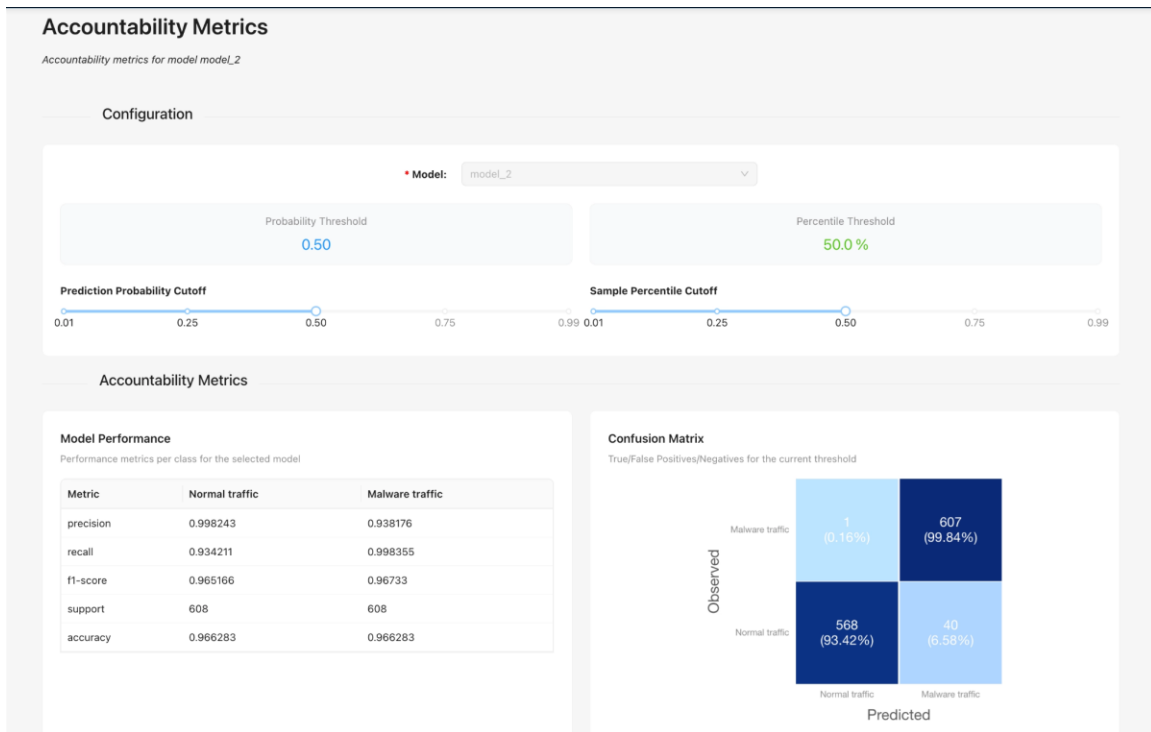


Figure 16. Accountability Metrics configurations and selected results.

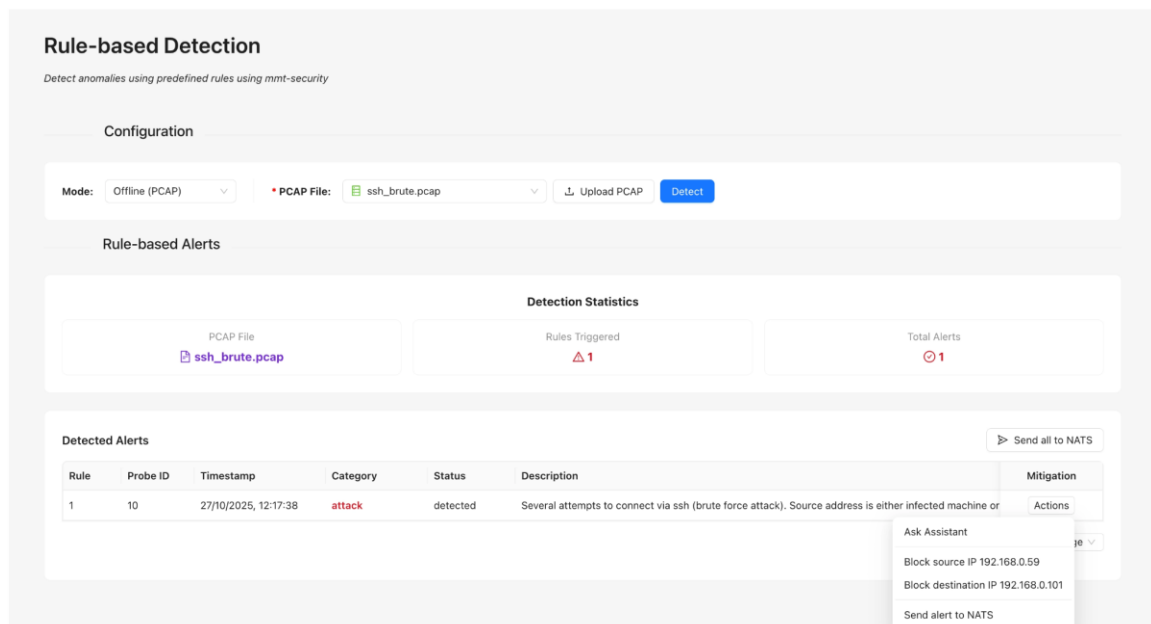


Figure 17. Rule-based Detection configurations and results.

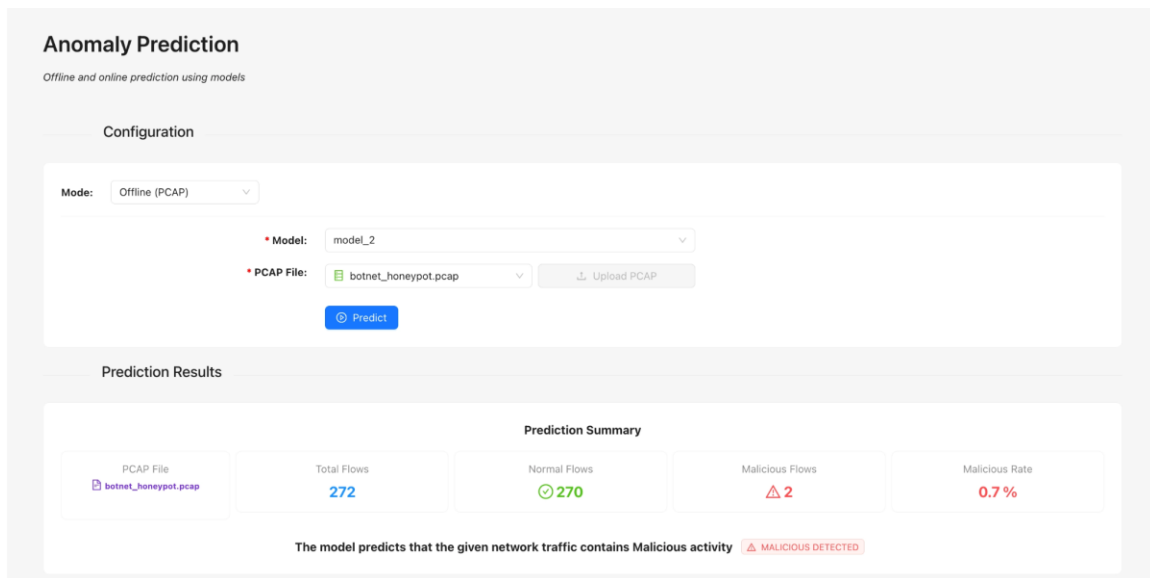


Figure 18. Anomaly Prediction configurations.

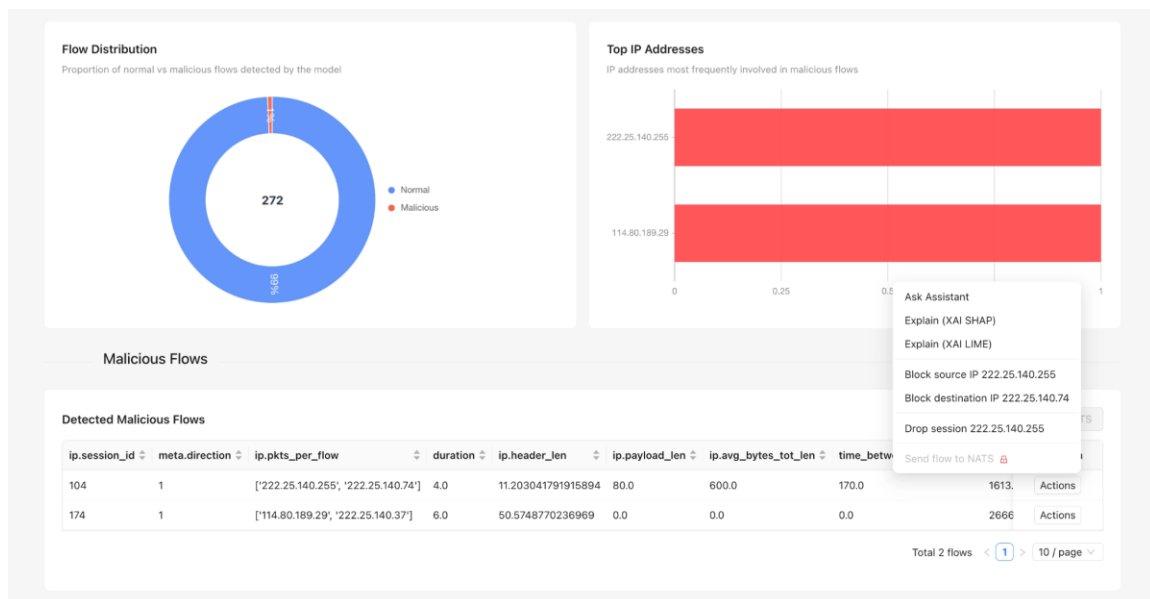


Figure 19. Visualizations of Anomaly Prediction results.

Summary

- Rule-based detection matches traffic against predefined signatures; it is fast and precise for known attacks but cannot detect novel threats.
- CNNs are supervised deep learning models that classify traffic into categories based on learned patterns from labeled data.

- Autoencoders detect anomalies by measuring reconstruction error; traffic that doesn't match learned normal patterns produces high error.
- Performance metrics (accuracy, precision, recall, F1-score) help evaluate model effectiveness; the right balance depends on your use case.
- Comparing detection results reveals complementary strengths; rules catch known patterns, CNNs classify, and autoencoders find novelty.

Quiz: Module 4 — Anomaly Detection with Rule-Based and Deep Learning Approaches

Pass mark: 7/10. Choose the best answer.

1. What is the main limitation of rule-based detection?

- A) Too slow for real-time use
- B) Cannot detect unknown attacks without new rules
- C) Produces too many false positives
- D) Requires too much memory

Answer: B — Rule-based detection only catches attacks that match existing signatures.

2. How do Autoencoders detect anomalies?

- A) By matching traffic to known attack signatures
- B) By classifying traffic into predefined categories
- C) By measuring reconstruction error — high error indicates anomaly
- D) By counting packet frequencies

Answer: C — Autoencoders flag traffic that cannot be reconstructed well as anomalous.

3. What does high precision but low recall indicate about a detection model?

- A) Many false alarms, catching most attacks
- B) Few false alarms, but missing many attacks
- C) Perfect detection performance
- D) The model is not working

Answer: B — High precision means alerts are usually real; low recall means many attacks are missed.

4. Which detection approach is BEST for identifying a completely new, previously unseen attack type?

- A) Rule-based detection
- B) Signature matching
- C) Autoencoder-based detection
- D) Port-based filtering

Answer: C — Autoencoders can detect novel anomalies by identifying deviations from normal patterns.

5. In a confusion matrix, what does a False Positive (FP) represent?

- A) An attack correctly identified as an attack
- B) Normal traffic correctly identified as normal
- C) Normal traffic incorrectly flagged as an attack
- D) An attack missed by the system

Answer: C — False positive = normal traffic incorrectly flagged as malicious.

6. What type of learning do CNNs use for network traffic classification?

- A) Unsupervised learning
- B) Supervised learning with labeled data
- C) Reinforcement learning
- D) Self-supervised learning

Answer: B — CNNs are trained on labeled examples of normal and attack traffic.

7. Which metric is most appropriate when you have an imbalanced dataset (few attacks, many normal)?

- A) Accuracy alone
- B) Precision alone
- C) F1-Score
- D) True Negative rate

Answer: C — F1-Score balances precision and recall, which is important when classes are imbalanced.

8. What does a high reconstruction error from an Autoencoder suggest about a network flow?

- A) The traffic is definitely malicious
- B) The traffic is definitely normal
- C) The traffic deviates from learned normal patterns and may be anomalous
- D) The model needs retraining

Answer: C — High reconstruction error means the traffic doesn't match normal patterns; it may be an anomaly.

9. Why is a layered approach (combining rule-based + CNN + Autoencoder) recommended?

- A) It's simpler to implement
- B) Each method has different strengths; combining them provides better coverage
- C) It uses less computational resources

D) Regulations require it

Answer: B — Different methods catch different types of threats; combining them provides defense in depth.

10. When rule-based detection says "normal" but the Autoencoder flags "anomaly," what should an analyst do?

A) Always trust the rules and ignore the Autoencoder

B) Always trust the Autoencoder and ignore the rules

C) Investigate the traffic further — it may be a novel attack

D) Restart the detection system

Answer: C — Disagreements warrant investigation; the Autoencoder may have detected something new.

Module 5: Advanced Explainable AI (XAI) Analysis

Estimated time: ~60 minutes · Text lesson + quiz

In this lesson

- Understand why Explainable AI is essential for cybersecurity applications
- Apply SHAP (SHapley Additive exPlanations) to interpret model predictions globally and locally
- Apply LIME (Local Interpretable Model-agnostic Explanations) for individual prediction explanations
- Interpret feature importance visualizations and understand which features drive decisions

Why Explainable AI matters in cybersecurity

Machine learning models, especially deep learning, are often called "black boxes", they make predictions, but we don't know why. In cybersecurity, this is a trust problem:

- Analysts need to understand alerts to act on them
- "The model says attack" is not actionable
- False positives erode trust if unexplained
- Compliance and auditing require justification

The accountability problem:

- Who is responsible when a model makes a wrong decision?
- Can you defend your detection logic to stakeholders?
- Regulations increasingly require algorithmic transparency

Explainable AI (XAI) addresses these problems by making model decisions interpretable. ACAS integrates XAI methods to help analysts understand not just what was detected, but why.

SHAP: SHapley Additive exPlanations

SHAP is based on game theory (Shapley values) and provides both global and local explanations. SHAP answers: "How much did each feature contribute to this prediction?"

How SHAP works

Think of features as "players" in a game, and the prediction as the "prize." SHAP calculates the fair contribution of each player to the final outcome. It provides global explanations, understand feature importance across all predictions.

SHAP output example

Global Feature Importance (SHAP values averaged across all samples):

- flow_packets_per_second: 0.42
- total_fwd_packets: 0.31
- syn_flag_count: 0.28
- flow_duration: 0.19
- fwd_packet_length_mean: 0.12

Interpretation: Across all samples, flow_packets_per_second has the highest impact on predictions, it's the most important feature globally.

LIME: Local Interpretable Model-agnostic Explanations

LIME explains individual predictions by creating a simple, interpretable model around each prediction. LIME asks: "If we slightly change the input features, how does the prediction change?" It perturbs the input, observes prediction changes, and builds a local linear model to approximate the complex model's behavior around that point.

How LIME works

1. Take the instance to explain (e.g., one network flow)
2. Create perturbed samples by slightly modifying features
3. Get predictions for all perturbed samples
4. Fit a simple linear model to these results
5. The linear model's coefficients = feature importance

LIME output example

Explaining prediction for Flow #1247: Prediction ATTACK with probability 0.89.

Features supporting "ATTACK":

```
flow_packets_per_second > 5000 → +0.32
syn_flag_count > 1000          → +0.24
flow_duration < 10 seconds    → +0.18
```

Features opposing "ATTACK":

```
total_bwd_packets > 100      → -0.08
fwd_packet_length_mean > 500 → -0.05
```

Interpretation: The high packet rate and SYN count strongly support the attack classification, while the presence of backward packets and larger packet sizes slightly oppose it.

SHAP vs. LIME comparison

Aspect	SHAP	LIME
Theoretical basis	Game theory (Shapley values)	Local linear approximation
Scope	Global and local	Primarily local
Consistency	Mathematically consistent	May vary with perturbations
Computation	Can be slower	Generally faster
Interpretation	Additive feature contributions	Feature importance weights
Best for	Deep understanding	Quick local explanations

When to use each:

- Use SHAP when you need consistent, theoretically grounded explanations
- Use SHAP global to understand overall model behavior
- Use LIME for quick explanations of individual alerts
- Use both for high-stakes investigations

Hands-on activity: XAI analysis

1. Select a detected anomaly — Choose a flow flagged as an attack
2. Run SHAP analysis — Generate global explanations, which features are important
3. Run LIME analysis — Get local explanation for the flow
4. Formulate response — Based on insights, determine what action you would take

Explainable AI with Local Interpretable Model-Agnostic Explanations (LIME)

LIME explanations of models

Configuration

Model:

Sample ID:

Features to display:

Contributions to display: ☒ Positive ☒ Negative

Feature(s) to mask:

LIME Method Overview

Local Interpretable Model-Agnostic Explanations (LIME) trains a simple local surrogate model around a specific prediction to approximate the model's behavior. It perturbs the input data and observes how predictions change, then fits an interpretable model (like linear regression) to explain that single prediction.

Figure 17. XAI LIME configurations.

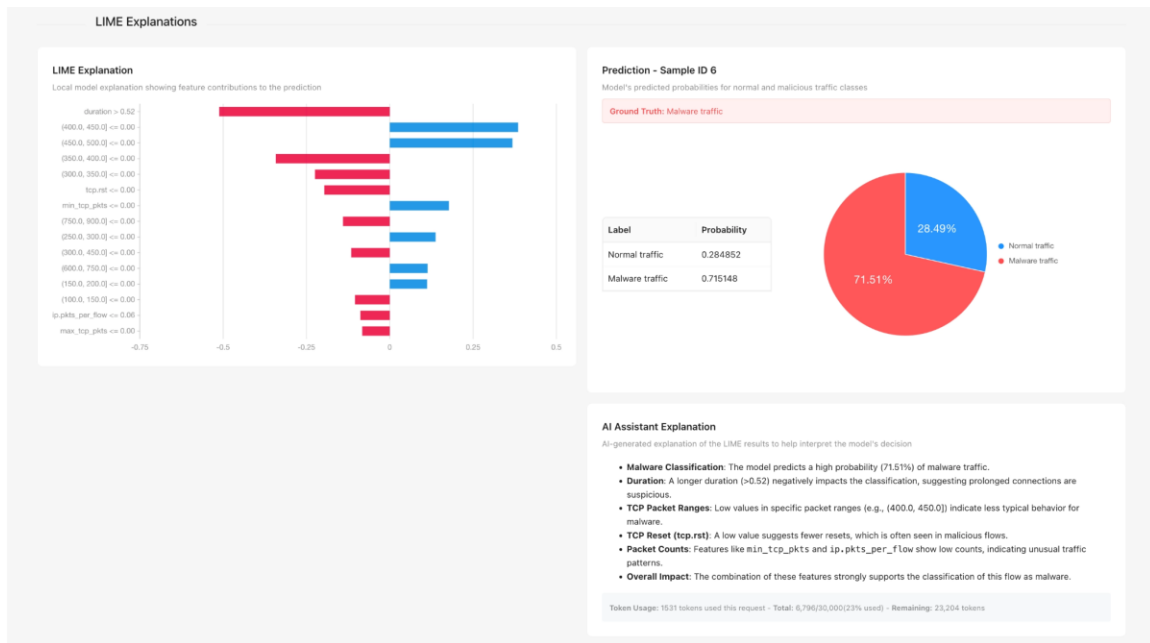


Figure 18. XAI LIME explanations.

Explainable AI with SHapley Additive exPlanations (SHAP)

SHAP explanations of models

Configuration

Model: model_2

Background samples: 20

Explained samples: 1

Features to display: 10

Contributions to display: ☒ Positive ☒ Negative

Feature(s) to mask: Select feature(s) ...

[SHAP Explain](#) [Ask Assistant](#)

SHAP Method Overview

SHapley Additive exPlanations (SHAP) uses Shapley values from cooperative game theory to quantify each feature's contribution to a prediction. It computes the average marginal contribution of each feature across all possible feature combinations, providing a theoretically sound and model-agnostic explanation.

Performance Warning: Running SHAP with a large number of background or explaining samples can be very slow and may take hours to complete. Start with small sample sizes (10-30) for testing.

Figure 19. XAI SHAP configurations.

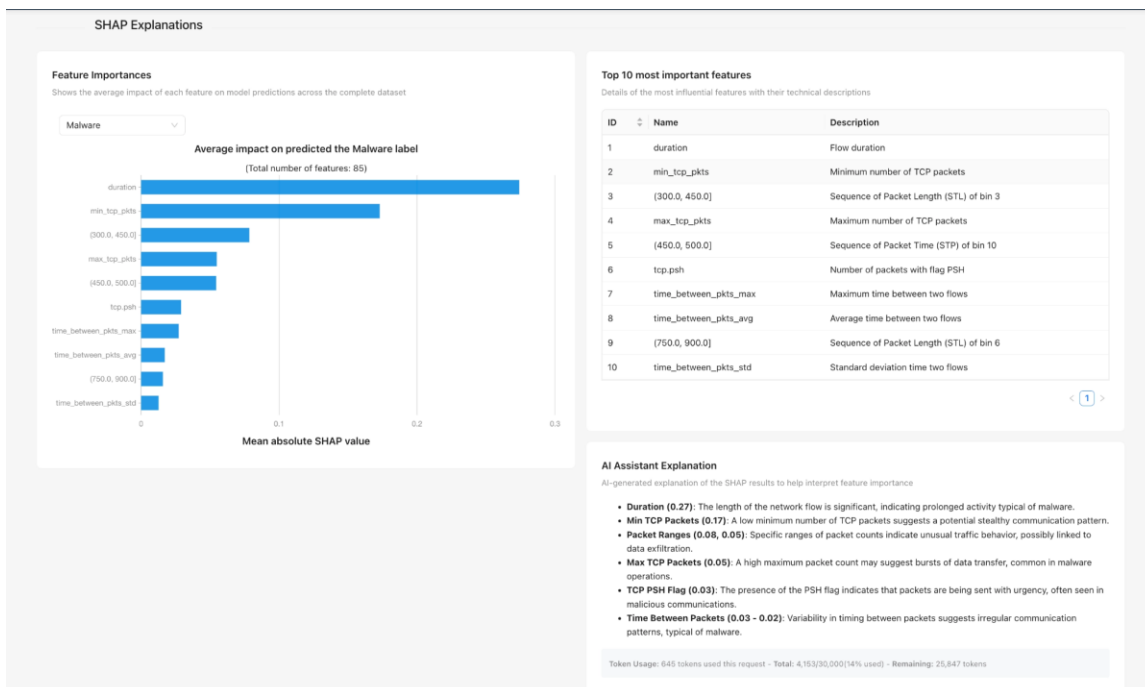


Figure 20. XAI SHAP explanations.

Summary

- Explainable AI addresses the "black box" problem, providing trust, accountability, and actionable insights for security analysts.
- SHAP uses game theory to calculate feature contributions; it provides both global (overall model) and local (individual prediction) explanations.
- LIME creates local linear approximations to explain individual predictions; it shows which features support or oppose the prediction.
- Feature importance visualizations (bar charts, force plots, dependence plots) help analysts quickly identify key factors in detections.
- Effective XAI usage follows a workflow: detect → explain → validate → act → document.

Quiz: Module 5 – Advanced Explainable AI (XAI) Analysis

Pass mark: 7/10. Choose the best answer.

1. What is the main purpose of Explainable AI (XAI) in cybersecurity?
 - A) To make models run faster
 - B) To understand why models make specific predictions
 - C) To reduce the size of training datasets
 - D) To encrypt model parameters

Answer: B — XAI helps us understand the reasoning behind model predictions.

2. SHAP is based on which theoretical foundation?

- A) Linear regression
- B) Game theory (Shapley values)
- C) Bayesian statistics
- D) Fourier analysis

Answer: B — SHAP uses Shapley values from cooperative game theory.

3. What does a positive SHAP value for a feature indicate?

- A) The feature decreases the prediction probability
- B) The feature increases the prediction toward the positive class (e.g., "attack")
- C) The feature is not important
- D) The feature has missing values

Answer: B — Positive SHAP values push predictions toward the positive class.

4. How does LIME generate explanations?

- A) By training a new deep learning model
- B) By perturbing the input and fitting a local linear model
- C) By computing exact Shapley values
- D) By clustering similar samples

Answer: B — LIME perturbs inputs and uses a simple linear model to approximate local behavior.

5. Which XAI method is better for understanding overall model behavior across all predictions?

- A) LIME local explanations
- B) SHAP global explanations
- C) Random feature selection
- D) Neither provides global insights

Answer: B — SHAP global analysis shows feature importance across all predictions.

6. In a SHAP force plot, features that push the prediction toward "attack" point in which direction?

- A) Left (toward lower values)
- B) Right (toward higher values / attack class)
- C) Up
- D) Down

Answer: B — In force plots, features pushing toward the positive class point right.

7. What additional value does the LLM assistant provide beyond SHAP/LIME?

- A) Faster computation

- B) Natural language explanations and actionable recommendations
- C) More features extracted
- D) Automated model retraining

Answer: B — The LLM translates technical XAI outputs into understandable language and recommendations.

8. If SHAP shows "syn_flag_count" with a high positive value for a flagged flow, what does this suggest?

- A) SYN flags are reducing the attack probability
- B) SYN flags are a key reason this flow was classified as an attack
- C) The model ignored SYN flags
- D) SYN flags are irrelevant to this detection

Answer: B — High positive SHAP means this feature strongly contributed to the attack classification.

9. Why might you use both SHAP and LIME for a high-stakes investigation?

- A) They always give identical results
- B) They provide complementary perspectives — SHAP is consistent, LIME is quick
- C) Regulations require both
- D) Using both is faster than using one

Answer: B — Using both provides more confidence through complementary methodologies.

10. What is the recommended final step after using XAI to understand a detection?

- A) Delete the alert
- B) Ignore the findings
- C) Document findings and take appropriate action based on insights
- D) Retrain the model immediately

Answer: C — XAI insights should lead to action and documentation for future reference.

Final Hands-On Assessment (Capstone)

Estimated time: ~60 minutes · Text lesson + final quiz

Overview

This capstone assessment evaluates your ability to perform an end-to-end network threat analysis using the ACAS platform. You will apply all the skills learned throughout the course: deep packet inspection, feature extraction, anomaly detection, and explainable AI analysis.

Scenario

You are a SOC analyst at a mid-sized company. Your network monitoring system has captured traffic during a suspected security incident. Management needs a clear report explaining:

1. What threats (if any) are present in the traffic
2. How you detected and identified them
3. Why the traffic is considered malicious (with evidence)
4. What actions should be taken to respond

Your task is to analyze the provided PCAP file and produce a comprehensive analysis.

Suggested workflow

Task 0: Capture a small PCAP file containing attacks, such as brute-force attempts or DDoS traffic.

Task 1: Initial analysis and DPI.

Task 2: Feature extraction analysis.

Task 3: Detection analysis using both rule-based and ML-based approaches.

Task 4: Explainability analysis using both LIME and SHAP methods.

Task 5: Final report on all findings and observations.

Summary

- Conducted an end-to-end analysis of network traffic using ACAS.
- Applied DPI, feature extraction, anomaly detection, and explainable AI techniques.
- Identified and investigated suspicious or malicious network activity.
- Validated findings using both rule-based and ML-based detection methods.
- Produced a comprehensive security report with evidence and response recommendations.
- Demonstrated practical SOC analyst skills for real-world incident investigation.

Final Quiz

Final assessment — 10 questions drawn across the whole course. Pass mark: 7/10. Choose the best answer.

1. In the ACAS workflow, what is the correct order of analysis steps?

- A) XAI → Detection → Feature Extraction → DPI
- B) Detection → DPI → XAI → Feature Extraction
- C) DPI → Feature Extraction → Detection → XAI
- D) Feature Extraction → XAI → DPI → Detection

Answer: C — The logical flow is: inspect traffic (DPI) → extract features → run detection → explain results (XAI).

2. If rule-based detection shows "normal" but both CNN and Autoencoder flag "anomaly", what is the most likely explanation?

- A) The ML models are broken
- B) The traffic matches a known attack signature
- C) The traffic may be a novel attack not covered by existing rules
- D) The rules are more accurate than ML

Answer: C — When ML detects what rules miss, it often indicates a new or variant attack.

3. Which combination of feature values would MOST strongly suggest a DDoS attack?

- A) Low packet rate, long duration, high bwd_packets
- B) High packet rate, short duration, high syn_flag_count, low bwd_packets
- C) Low packet rate, long duration, balanced traffic
- D) Average packet rate, moderate duration, low flag counts

Answer: B — DDoS typically shows high rates, short flows, many SYN flags, and minimal responses.

4. When writing the executive summary of a threat report, what should it contain?

- A) Full technical details of every detection
- B) Complete list of all features extracted
- C) Brief overview of threats found and overall risk level
- D) Step-by-step reproduction instructions

Answer: C — Executive summaries are brief, high-level overviews for decision-makers.

5. If SHAP shows flow_iat_std (inter-arrival time standard deviation) with a very low value contributing to an attack classification, what might this indicate?

- A) Random, normal network traffic
- B) Possible automated/scripted communication (like C2 beaconing)

- C) Manual user browsing
- D) Feature extraction error

Answer: B — Very consistent timing (low IAT std) suggests automated communication.

6. What is the primary purpose of comparing detection results across rule-based, CNN, and Autoencoder methods?

- A) To determine which single method to use
- B) To increase computational load
- C) To gain confidence through agreement and investigate disagreements
- D) To satisfy compliance requirements

Answer: C — Comparing methods builds confidence and reveals edge cases for investigation.

7. Which section of the final report should contain "Enable SYN cookies and implement rate limiting"?

- A) Executive summary
- B) Traffic overview
- C) Detected threats
- D) Recommendations

Answer: D — Specific mitigation actions belong in the Recommendations section.

8. If the protocol hierarchy shows 40% of traffic is DNS, significantly higher than typical enterprise traffic, what should you investigate?

- A) Nothing — DNS is always normal
- B) Possible DNS tunneling or exfiltration
- C) HTTPS decryption issues
- D) IPv6 transition problems

Answer: B — Unusually high DNS percentages can indicate tunneling for covert communication.

9. When the LLM assistant provides mitigation recommendations, what should you do before implementing them?

- A) Implement immediately without review
- B) Ignore them entirely
- C) Validate they are appropriate for your environment and findings
- D) Forward them to the attacker

Answer: C — Always validate AI recommendations against your specific context before acting.

10. What makes a good threat report actionable?

- A) Maximum technical jargon

- B) Clear findings, evidence-based explanations, and specific remediation steps
- C) Lengthy descriptions without conclusions
- D) Recommendations without evidence

Answer: B — Good reports provide clear findings, explain the evidence, and give specific actions.

We Value Your Feedback!

Text lesson · Course evaluation survey

Congratulations on completing **Hands-On Network Threat Detection and Response Using ACAS!**

Before you go, please take a few minutes to complete our short course-evaluation survey. Your honest feedback helps us improve the lessons, labs, and the ACAS tool itself for future learners.

What the survey asks about

- The clarity and usefulness of each module.
- Whether the hands-on labs were easy to follow and worked as expected.
- The pacing and overall difficulty of the course.
- Any technical issues you encountered with the platform.
- Suggestions and features you would like to see.

Take the survey

Please complete the evaluation form here: **[survey link to be provided by the CyberSuite Academy]**.

The survey is anonymous and takes about five minutes. Thank you for helping us make the course better and good luck applying your new network threat detection and response skills, responsibly and only on systems you are authorised to test.