



EMERALD

Deliverable D2.7

ML Model Certification – v2

Editor(s):	Ching-Yu Kao
Responsible Partner:	Fraunhofer Institute for Applied and Integrated Security (FhG AISEC)
Status-Version:	Final – v1.1
Date:	15.09.2026
Type:	OTHER
Distribution level:	PU

Project Number:	101120688
Project Title:	EMERALD

Title of Deliverable:	D2.7 – ML model certification – v2
Due Date of Delivery to the EC	31.10.2025
Updated Version of Deliverable	15.09.2026

Workpackage responsible for the Deliverable:	WP2 – Methodology for knowledge extraction
Editor(s):	Ching-Yu Kao (FhG)
Contributor(s):	--
Reviewer(s):	Marinella Petrocchi (CNR) Cristina Martínez, Juncal Alonso (TECNALIA)
Approved by:	All Partners
Recommended/mandatory readers:	WP1, WP2, WP3, WP4, and WP5

Abstract:	This deliverable presents the component for evidence extraction from machine learning models that can be integrated with the certification graph. It is the result of work performed in Task 2.4. This document is the revised version on those extractors.
Keyword List:	Knowledge extraction, machine learning, deep learning, robustness, security, technical evidence
Licensing information:	This work is licensed under Creative Commons Attribution-ShareAlike 4.0 International (CC BY-SA 4.0 DEED https://creativecommons.org/licenses/by-sa/4.0/)
Disclaimer	Funded by the European Union. Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union. The European Union cannot be held responsible for them.

Document Description

Version	Date	Modifications Introduced	
		Modification Reason	Modified by
v0.1	03.09.2025	First draft version, outline	Ching-Yu Kao (FHG AISEC)
v0.2	07.10.2025	Added contents to <i>AI-SEC</i>	Ching-Yu Kao (FHG AISEC)
v0.3	18.10.2025	Finalization	Ching-Yu Kao (FHG AISEC)
v0.4	26.10.2025	Internal QA review	Marinella Petrocchi (CNR)
v0.5	28.10.2025	Modification after QA review	Ching-Yu Kao (FHG AISEC)
v0.6	28.10.2025	Final Review	Cristina Martínez/ Juncal Alonso (TECNALIA)
v0.7	30.10.2025	Modifications after final review	Ching-Yu Kao (FHG AISEC)
v1.0	31.10.2025	Submitted to the European Commission	Cristina Martínez/ Juncal Alonso (TECNALIA)
v1.0.1	29.06.2026	Relevant updates in development of <i>AI-SEC</i>	Ching-Yu Kao (FHG AISEC)
v1.0.2	07.09.2026	Internal QA review	Marinella Petrocchi (CNR)
v1.0.3	07.09.2026	Incorporate QA review comments	Ching-Yu Kao (FHG AISEC)
v1.1	15.09.2026	Submitted to the European Commission	Cristina Martínez/ Juncal Alonso (TECNALIA)

Table of contents

Terms and abbreviations.....	6
Executive Summary	8
1 Introduction.....	9
1.1 About this deliverable	9
1.2 Document structure	9
1.3 Updates from D2.6	9
1.4 Technological advances from the MEDINA project.....	10
2 ML Model Evidence Extractors in the EMERALD Architecture.....	11
3 Implementation.....	13
3.1 Functional description.....	13
3.2 Technical description	14
3.2.1 Prototype architecture.....	14
3.2.2 Technical specifications	16
3.3 Delivery and usage	19
3.3.1 Package information	19
3.3.2 Installation	20
3.3.3 Instructions for use	20
3.3.4 Licensing information	25
3.3.5 Download.....	25
3.4 Demo	25
3.5 Connection to the Evidence Store.....	28
3.5.1 Generating the evidence JSON file	28
3.5.2 Sending AI-SEC evidence to the Evidence Store	29
3.6 Pilot 1 Test: SVM-based Spam Email Classifier	30
3.7 Limitations and future work.....	31
4 Conclusions.....	33
5 References.....	34

List of tables

TABLE 1. OVERVIEW OF THE CHANGES WITH RESPECT TO D2.6.....	9
TABLE 2. REQUIREMENT AI-SEC.01 - EXTRACTION OF SECURITY FEATURES FROM ML MODELS	14
TABLE 3. TEST RESULTS ON NLP MODELS.....	18
TABLE 4. TEST RESULTS ON IMAGES DATASETS	18
TABLE 5. OVERVIEW AND DESCRIPTION OF PACKAGE STRUCTURE FOR THE AI-SEC.....	19
TABLE 6. SETUP FOR THE ML MODEL USING MNIST DATASET	23
TABLE 7. SETUP FOR THE ML MODEL USING CIFAR-10 DATASET	23
TABLE 8. RESULTS ON MNIST AND CIFAR10 USING CLEVER SCORE, SHAPr SCORE, DATA POISONING AND LIME	24
TABLE 9. EXAMPLE OF THE INPUT SENTENCES.	30

List of figures

FIGURE 1. EMERALD COMPONENT OVERVIEW DIAGRAM [7]. THE RED RECTANGLE HIGHLIGHTS THE ML MODEL EVIDENCE EXTRACTION COMPONENT, WHICH IS DESCRIBED IN THIS DELIVERABLE.	12
FIGURE 2. AI-SEC ARCHITECTURE. TO EXTRACT FEATURES FROM ML MODELS, WE NEED TO EVALUATE THE ROBUSTNESS AGAINST POISONING AND ADVERSARIAL ATTACKS, THE PRIVACY LEVEL AND THE ABILITY TO GIVE EXPLANATIONS	15
FIGURE 3: EMERALD UI VIEW FOR SETTING THE EVIDENCE EXTRACTOR FOR A TARGET OF EVALUATION	21
FIGURE 4. THE DEMONSTRATION INTERFACE CONSISTS OF THREE MAIN STEPS	26
FIGURE 5. AFTER RUNNING THE ANALYSIS, WE CAN SEE THE RESULTS AND SAVE THE RESULTS FOR FURTHER INVESTIGATION.....	27

Terms and abbreviations

AI	Artificial Intelligence
AI-SEC	AI Security Evidence Collector
AMOE	Assessment and Management of Organisational Evidence
API	Application Programming Interface
BERT	Bidirectional Encoder Representations from Transformers
BSI	Bundesamt für Sicherheit in der Informationstechnik
BSI C4	Artificial Intelligence Cloud Services Compliance Criteria Catalogue
CertGraph	Certification Graph
CIFAR	Canadian Institute For Advanced Research
CLI	Command Line Interface
Codyze	Static Code Analyzer from FHG
CSA or EU CSA	EU Cybersecurity Act
CSP	Cloud Service Provider
CSV	Comma-Separated Values
CLEVER	Cross Lipschitz Extreme Value for nEtnetwork Robustness
DoA	Description of Action
EC	European Commission
eknows	Platform for Software Analysis from SCCH
eknows-e3	The extractor component developed in the context of EMERALD
GA	Grant Agreement to the project
ISO/IEC	International Organization for Standardization / International Electrotechnical Commission
JSON	JavaScript Object Notation
KPI	Key Performance Indicator
IMDb	Internet Movie Database
JSON	JavaScript Object Notation
MEDINA	Predecessor project of EMERALD
MIT	Massachusetts Institute of Technology
ML	Machine Learning
MNIST	Modified National Institute of Standards and Technology database
LIME	Local Interpretable Model-agnostic Explanations
LLM	Large Language Model
MEDINA	Predecessor project of EMERALD
MIT	Massachusetts Institute of Technology
MNIST	Modified National Institute of Standards and Technology database
NLP	Natural Language Processing
PNG	Portable Network Graphics
RGB	Red, Green, Blue
SMS	Short Message Service
SVM	Support Vector Machines
SW	Software
SHAPr	SHapley Additive exPlanations
SPADE	Situation, Problem, Analysis, Decision, and Execution
SVC	Support Vector Classifier
SVM	Support Vector Machine
TF-IDF	Term Frequency – Inverse Document Frequency
TRL	Technology Readiness Level

WP	Work Package
----	--------------

DRAFT

Executive Summary

This deliverable presents the design, architecture, and implementation status of the machine learning (ML) model evidence set of extractors of WP2, which we have named *AI-SEC*. *AI-SEC* contributes to the key result KR1-EXTRACT of EMERALD, as a framework to continuously extract knowledge from well-trained ML models and prepare suitable evidence based on them.

EMERALD follows a knowledge graph-based approach to provide a unified view of the cloud service under certification at different layers of the service, ranging from the infrastructure layer (e.g., virtual resources) to the business layer (e.g., policies and procedures), to the implementation layer (e.g., source code files) and data layer (e.g., increasingly used AI models) in cloud applications.

The ML model evidence extractors, developed in Task 2.4 “Security and Privacy Evaluation of Machine Learning Models” and described in this deliverable, aim at identifying critical security-related features, such as adversarial robustness, privacy, security and explainable AI in the cloud environments. Other related deliverables in WP2, all due at project month 24 (October 2025), provide functional and technical details on further evidence extractors from different sources, i.e., D2.3 [1] on source code evidence extraction in Task 2.2, D2.5 [2] on evidence extraction from policy documents in Task 2.3 and D2.9 [3] on runtime data extraction in Task 2.5. All these details contributed to D2.1 [4] on the overall information model of the certification graph in Task 2.1.

The main part of this deliverable provides functional and technical descriptions of the evidence extractor *AI-SEC*, including its purpose and scope, the current coverage of the EMERALD requirements, and the components’ internal architecture. These descriptions are complemented by information on delivery and usage, as well as on limitations and future work. Finally, the document concludes with a short summary.

The ML model evidence extractors described in this deliverable contribute to KR1-EXTRACT by providing next-generation evidence gathering tools and techniques based on a knowledge graph approach. The presented extractors currently have the tool implemented and ready to be (to some degree) integrated with other components of the EMERALD architecture.

Based on the work described in this deliverable, the ML model evidence extractors will be further extended and integrated into the EMERALD framework. This is the final version of this deliverable (the first one is D2.6 [5]).

1 Introduction

EMERALD aims to provide a next generation set of evidence gathering tools and techniques based on a knowledge graph approach. KR1-EXTRACT supports an improved and unified tool-supported approach to continuously extract knowledge from different layers of a cloud service, e.g., infrastructure, platform, runtime information, policy documents, software, and AI models.

1.1 About this deliverable

The deliverable describes the role of the *AI Security Evidence Collector (AI-SEC)* in generating evidence from ML models. This report is the second of two iterations of deliverables resulting from Task 2.4 “Security and Privacy Evaluation of Machine Learning Models”, with the first iteration documented in D2.6 [5].

Building on the work presented in D2.6, this deliverable further focuses on the integration of *AI-SEC* into the EMERALD ecosystem. To enable this integration, the outputs of the ML security analyses are transformed into the common EMERALD evidence format. *AI-SEC* can convert these analysis outputs into structured EMERALD evidence objects. The generated evidence includes the required fields, such as id, timestamp, toolId, and targetOfEvaluationId, as well as a valid ontology-based resource representation.

1.2 Document structure

The document is structured as follows.

In Section 2, we show the *AI-SEC* tool in the EMERALD architecture and how it interacts with other EMERALD components. In Section 3, we report on the design and implementation of *AI-SEC* ML model extractor. For the ML model extractor, functional and technical descriptions are provided, including their purpose and scope, the (current and planned) coverage of the EMERALD requirements, the components’ internal architecture, their subcomponents, and details about the programming language, libraries, etc. used. These descriptions are complemented by information on delivery and usage, including package information, installation instructions, user manual, licensing, download information and demonstration, as well as limitations and future work. Section 4 concludes this report.

1.3 Updates from D2.6

The deliverable is based on D2.6 [5] and presents an updated version of *AI-SEC* component. It includes sections from D2.6 that remain unchanged, along with newly added information. To facilitate tracking the updates and enhancements made since the previous version, Table 1 details the modifications and new additions for each section of the document. Please also note that the latest version v1.1 of this deliverable reflects the final status of development of the runtime evidence extractors until the end of WP2.

Table 1. Overview of the changes with respect to D2.6

Section	Changes
1 Introduction	Added new subsections: • “Updates from D2.6” • “Technological advances from the MEDINA project”
2 ML Model Evidence Extractors in the EMERALD Architecture	Minor updates, texts of demo are added to new section 3.4

Section	Changes
3 Technical specifications	- Added test data results to the section - Added demo section - Updated all sections to reflect new status of the <i>AI-SEC</i> component
4 Conclusion	Minor updates

1.4 Technological advances from the MEDINA project

EMERALD builds upon the outcomes of the MEDINA project: *Security framework to achieve a continuous audit-based certification in compliance with the EU-wide cloud security certification scheme* (GA 952633) [6]. While MEDINA reached Technology Readiness Level (TRL) 5, the goal of EMERALD is to enhance the system and elevate its TRL from prototype (TRL5) to product level (TRL7).

In MEDINA, some extraction components implemented their own assessments, which led to increased maintenance efforts whenever certification requirements evolve. Therefore, centralising the assessment process is one of the main improvements introduced by EMERALD. This is achieved by delivering, as far as possible, raw evidence directly to the central *Evidence Store*, instead of performing local assessments within each extraction component.

A key part of this advancement is the introduction of a new component *AI-SEC*, specifically designed as an evidence extractor for machine learning (ML) models. *AI-SEC* focuses on analysing and quantifying the security-relevant characteristics of ML models, such as their robustness against adversarial attacks, resilience to data poisoning, privacy leakage risks, and explainability of decisions.

By integrating these AI-specific evidence types into EMERALD's Certification graph, *AI-SEC* extends the MEDINA framework towards trustworthy AI certification. It enables the extraction of technical evidence at the ML model layer, which is a capability that was not available in MEDINA, and thus represents a significant step toward comprehensive, model-level assurance within EMERALD's certification ecosystem.

2 ML Model Evidence Extractors in the EMERALD Architecture

This section describes how the ML model evidence extractors interact with (selected) EMERALD components on a conceptual level. Figure 1 shows the EMERALD high-level architecture as a component diagram, as described in D1.2 [7]. In EMERALD, a component is defined as “any part of the EMERALD ecosystem that has a specific functionality and can be considered a separate entity with respect to other components” (see D1.4 [8]).

The components for collecting evidence about technical and organisational measures, i.e., *AMOE* [2], *eknows-e3* [1], *AI-SEC*, *Clouditor-Discovery* [3], and *Codyze* [1], are represented at the bottom part of Figure 1. The ML model evidence extractor *AI-SEC*, which obtains technical evidence from the analysis of ML models adopted in cloud environments, is highlighted using a thick frame. In particular, *AI-SEC* provides evidence about key AI-specific features, including model robustness against adversarial attacks, data integrity, explainability, and fairness. The extracted evidence from *AI-SEC* is sent to the *Evidence Store* [9], where it can be combined with evidence collected by other EMERALD components. *AI-SEC* is a newly developed component in EMERALD.

By integrating *AI-SEC* into EMERALD, we are now able to cover a broader range of security catalogues (i.e., AI-related standards like for example, ISO/IEC 24028:2020 (regarding the trustworthiness of artificial intelligence), ISO/IEC 42001:2023 (AI management systems), and other standards addressing robustness and transparency. This integration enables EMERALD not only to assess traditional security and privacy controls but also to evaluate AI-specific properties that are becoming increasingly critical for compliance and risk management in intelligent systems.

AI-SEC provides model-specific insights into machine learning systems, which, when combined with evidence gathered from other components such as *Clouditor-Discovery* and *Codyze*, contribute to a broader evidence base. For instance, *Clouditor-Discovery* can collect configuration data or deployment metadata from cloud environments, while *Codyze* performs static code analysis to verify adherence to secure coding practices. When these complementary sources of evidence are integrated with *AI-SEC*'s machine learning model analysis, they together provide both system-level and AI-model-level assurance. This combination enables a comprehensive, multi-layered assurance process that links configuration, code, and AI model evidence to deliver stronger guarantees of overall security and compliance.

Overall, the *AI-SEC* component extends EMERALD's capability to systematically collect, analyse, and reason about evidence related to AI models. This makes EMERALD a more complete ecosystem, capable of demonstrating compliance and trustworthiness in complex cloud-based AI applications.

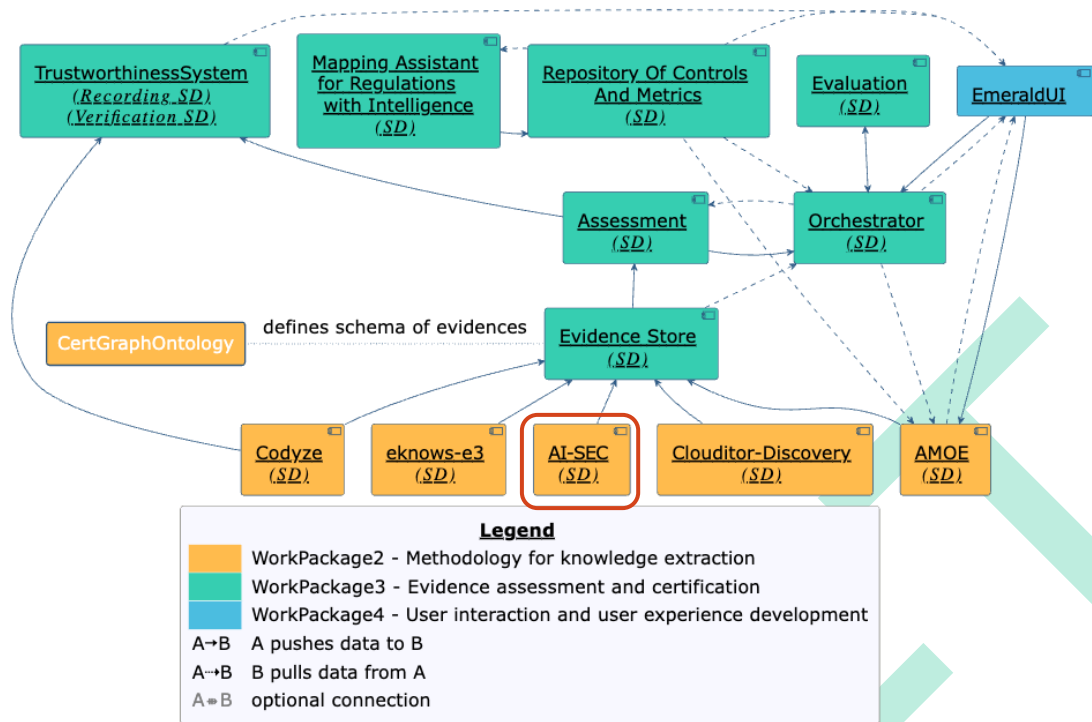


Figure 1. EMERALD component overview diagram [7]. The red rectangle highlights the ML model evidence extraction component, which is described in this deliverable.

3 Implementation

AI-SEC is one of the evidence-extraction components in EMERALD. Its specific focus is on machine learning (ML) models that are part of cloud-based services under certification. While other EMERALD extractors (such as *Codyze* or *AMOE*) analyse source code or policy documents, *AI-SEC* inspects trained ML models themselves. It evaluates these models to generate quantitative evidence about:

- Robustness: how resistant a model is to adversarial perturbations or malicious input data (using CLEVER score [10])
- Security and data integrity: how well the model resists data-poisoning or backdoor attacks using backdoor data poisoning [11]
- Privacy: whether individual training samples can be inferred from the model (using SHAPr leakage analysis [12])
- Explainability: whether model decisions can be locally explained in a transparent way (using LIME [13])

In EMERALD's overall workflow, this evidence is collected when a cloud service provider includes an ML component as part of its cloud application. The extracted results are stored in the *Evidence Store* [9] and connected through the *Certification Graph* [14] (CertGraph). This allows the certification engine to trace security properties of the ML components alongside other technical and organisational evidence, contributing to a holistic, auditable view of the cloud service's trustworthiness.

In line with the requirements set out in Section 6.2 of the Criteria Catalogue for AI Cloud Service (AIC4)¹, which covers the Security & Robustness Objective, *AI-SEC* has been designed to meet these critical security criteria, ensuring comprehensive protection and compliance. Our solution addresses four key aspects: **robustness, security and data integrity, privacy, and explainability**.

3.1 Functional description

Overall purpose. The *AI-SEC* prototype provides a comprehensive toolkit for evaluating and improving the security of machine learning models by focusing on adversarial robustness testing, data poisoning attacks, privacy vulnerability assessment, and model interpretability. It effectively assesses vulnerabilities using techniques such as CLEVER score calculation [10], SHAPr leakage analysis [12], backdoor data poisoning [11], and LIME-based explanations [13]. These features will be collected as evidence for the *Certification graph* [14].

Context and scope. The toolkit assumes that users already have pre-trained models available, which can be directly utilized for evaluation. Additionally, it supports standard datasets for robustness and privacy assessments, while also allowing custom dataset imports for tailored evaluations.

Motivation. AI models are increasingly deployed in security-critical contexts. While previous projects such as MEDINA addressed software and policy layers, model-level evidence extraction remained limited. *AI-SEC* closes this gap by analysing the ML layer directly, ensuring that adversarial robustness, privacy leakage, and model transparency can be systematically quantified and certified.

Innovation. *AI-SEC* focuses on the following innovations:

¹ https://www.bsi.bund.de/EN/Themen/Unternehmen-und-Organisationen/Informationen-und-Empfehlungen/Kuenstliche-Intelligenz/AIC4/aic4_node.html

- Integration of multiple security assessments, robustness, privacy, data integrity and interpretability into a single, unified system.
- Automation of calculations and result generation through command-line inputs

Requirements. The relevant requirements with their respective implementation progress (degree of implementation - *partially / fully / not implemented*) and a brief description of how they have been implemented are provided in Table 2.

Table 2. Requirement AI-SEC.01 - Extraction of security features from ML models

Field	Description
Requirement ID	AI-SEC.01
Short title	The extractor tool includes defined criteria
Description	The AI-SEC component is designed to extract quantifiable evidence from trained ML models in line with the BSI AIC4 criteria. It evaluates the robustness, privacy, and explainability of models by applying established assessment methods (e.g., CLEVER for robustness, SHAPr for privacy, and LIME for explainability). The resulting metrics form part of the technical evidence contributing to EMERALD.
Status	Validated
Priority	Must
Component	AI-SEC
Source	Component, KPI
Type	Technical
Related KR	KR5_AIPOC
Related KPI	KPI 5.1
Validation acceptance criteria	Code review: Review code and check if analysis methods work for different ML models.
Progress	100% - Fully Implemented
Milestone	MS5: Components V2 (M30)

3.2 Technical description

The following subsections describe the technical details of AI-SEC for EMERALD.

3.2.1 Prototype architecture

The AI-SEC architecture was developed as a prototype to demonstrate and validate its core functionality and key components. In a prototype architecture, the primary components are defined, and basic functionality is implemented, enabling early-stage evaluations and iterative improvements. It helps identify potential issues and make design adjustments based on testing and user feedback. The AI-SEC prototype contains four main functions:

- **Data Processor:** Prepares datasets, labels and models for attacks and evaluations. In Figure 2, this refers to the “Selection of ML model” box and the “Selection of Dataset for ML model” box.
- **Attack:** Performs analysis using selected model and datasets. In Figure 2, this refers to the “Selection of attack” box.
- **Output:** It refers to the “Evaluation Score (s)” in Figure 2. According to the selection of attack, the output could be poison robustness (data poisoning attacks using the “Hidden Trigger Backdoor Attack” method [15]), adversarial robustness (CLEVER score [8]), privacy vulnerability scores (SHAPr [9]) and Explainability (Generates explanations of model predictions using LIME [11]).

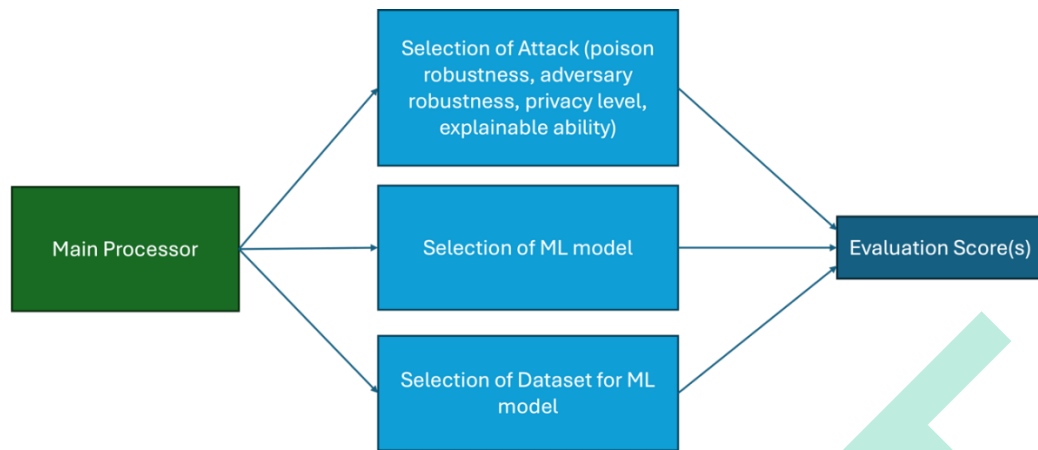


Figure 2. AI-SEC architecture. To extract features from ML models, we need to evaluate the robustness against poisoning and adversarial attacks, the privacy level and the ability to give explanations

Based on this architecture, we obtain the following results, which serve as evidence for the Certification Graph.

1 Poisoning Resilience Score

To know if a model is robust to poisoning attacks, we perform data poisoning attacks to evaluate model resilience against malicious examples. The “Hidden Trigger Backdoor Attack Sleeper Agent”² [11] is used for this. In this approach, the attacker inserts a *hidden trigger* (such as a specific pattern, object, or symbol) into some of the training data. The model learns to associate this trigger with a particular output. During normal use, the model behaves as expected, but when it encounters this hidden trigger, it produces the attacker’s desired outcome instead. Unlike adversarial attacks, which aim to fool an already trained model by manipulating its inputs at inference time, poisoning attacks target the training process by injecting malicious or manipulated data to corrupt the resulting model. A lower score indicates that less data is needed to alter the model’s robustness. The score ranges are normalized from 0 to 1, with lower values indicating higher susceptibility to poisoning attacks and higher values indicating greater robustness against poisoning.

2 Adversarial Robustness Score

CLEVER Score Module: CLEVER Scores Calculation³ [10] assesses the robustness of a model by calculating CLEVER scores. It provides an estimate of how resilient a model is to small perturbations in input data that could lead to incorrect or unexpected outputs. Adversarial attacks test *how easy it is to fool a trained model*. In contrast, poisoning attacks discussed above test *how easy it is to corrupt a model during training*.

A higher CLEVER score means the model is more robust to adversarial attacks. In simple terms, it would take a bigger or more noticeable “nudge” to the input data to trick the model into making a wrong decision. So, the higher the CLEVER score, the safer or more reliable the model is when facing small, potentially harmful changes in the data that are meant to confuse it. In contrast, a lower CLEVER score means the model is more vulnerable to these kinds of attacks, as even a small tweak to the input data could lead it to make mistakes.

² <https://arxiv.org/pdf/2106.08970>

³ <https://openreview.net/pdf?id=BkUHIMZ0b>

The value range depends on the radius size, which is 0 - 5 by default and can be modified in the function `utils/compute_untargeted_clever()`.

3 Privacy Score

SHAPr Leakage Module: Evaluates model privacy using SHAPr leakage metrics⁴ [12]. The SHAPr Leakage score analyses the SHAP values to identify if private information (like unique patterns in sensitive data) could be reconstructed or inferred. If the SHAPr values indicate that individual data points can significantly influence a model's prediction, there may be a risk of leakage. A higher final SHAPr score for a training sample means it is more vulnerable to privacy attacks. The values range from 0 to 1.

4 Explainability

This component leverages the LIME⁵ [13] (Local Interpretable Model-Agnostic Explanations) library to explain individual predictions made by the model. LIME works by generating perturbed versions of an input and observing how the model's predictions change, effectively creating a local approximation of the model's decision boundary around the input. This provides insight into which features in the inputs contribute most to a particular decision and helps the user understand the behaviour of the model. The values range from 0 to 1, higher value means better explainability.

3.2.2 Technical specifications

Programming Language:

- Python 3.7

Libraries:

- Torch⁶ (version 1.13.1) for model processing and evaluation.
- Torchvision⁷ (version 0.14.1) for image data handling.
- Lime⁸ (version 0.2.0.1) for model interpretability.

Databases:

- Datasets are stored locally in compatible formats such as CSV for labels and PNG for images.
- IMDb datasets⁹

Core functionality extension and optimization:

We have successfully extended *AI-SEC* with core functionality designed to support a wider range of machine learning models and more flexible operations:

- Natural Language Processing (NLP) Model Support: In the first version of *AI-SEC*, we could only test on image datasets and models. In the current version, the robustness and privacy assessment capabilities of the Toolkit have been successfully extended to

⁴ <https://arxiv.org/abs/2112.02230>

⁵ <https://github.com/marcotcr/lime>

⁶ <https://pytorch.org/>

⁷ <https://pytorch.org/vision/stable/index.html>

⁸ <https://github.com/marcotcr/lime>

⁹ <https://huggingface.co/datasets/stanfordnlp/imdb>

support mainstream NLP pre-trained models on Hugging Face¹⁰. This greatly enhances the applicability of the toolkit.

- Supported model types: Models that have been tested and supported include: *bert-base-uncased*, *distilbert-base-uncased*, *roberta-base*, *deberta-base*, *albert-base-v2* and *xlnet-base-cased*.
- Limitations: After code analysis and testing, the toolkit is currently only applicable to encoder-only models and is not yet applicable to decoder-only models due to the lack of embedding or hidden states output and logits output.
- Code versatility: All NLP functionalities have been integrated into the original image processing functionality, making the code more versatile.
- Interpretable AI numerical output optimization: The SPADE [10], SHAPr [11], and LIME [12] methods have been optimized to produce precise numerical results even on small datasets, serving as usable evidence for interpretable AI.

Table 3 and Table 4 show the test results on NLP and image datasets, respectively.

To evaluate the *AI-SEC* component, three standard benchmark datasets were used, namely IMDb, MNIST, and CIFAR-10. Each representing a different modality and learning task.

- IMDb¹¹ (Internet Movie Database): A text-based dataset commonly used for sentiment analysis. It contains 50,000 movie reviews labelled as *positive* or *negative*. For our experiments, pre-trained transformer models (e.g., BERT, DistilBERT, RoBERTa) were fine-tuned on this dataset to classify the sentiment of a given review. *AI-SEC* then analysed the trained models for privacy leakage and robustness in a natural-language context.
- MNIST¹² (Modified National Institute of Standards and Technology): A collection of 70,000 grayscale images of handwritten digits (0–9), widely used for image classification tasks. The CNN model trained on MNIST was used to assess *AI-SEC*'s ability to measure robustness (CLEVER score), privacy leakage (SHAPr score), and interpretability (LIME explanations) on simple, low-complexity data.
- CIFAR-10¹³ (Canadian Institute for Advanced Research): A dataset of 60,000 colour images in 10 categories (airplanes, cars, birds, cats, etc.), used for multi-class image recognition. This dataset enabled testing *AI-SEC* on more complex image data, evaluating adversarial robustness, poisoning resilience, and explainability of model predictions.

These datasets were selected to demonstrate *AI-SEC*'s applicability across different types of AI models, from text classification to image recognition, and to verify that the component can extract meaningful evidence under diverse data and model configurations.

For the NLP models evaluated on the IMDb dataset, the SPADE scores are consistently high across all tested architectures, ranging from 0.93 to 1.00. BERT, DistilBERT, and RoBERTa achieve a score of 1.00, while ALBERT and XLNet obtain slightly lower scores of 0.93 and 0.96, respectively. This shows that the SPADE extractor can be successfully applied across different Transformer-based model architectures, while also capturing model-specific differences in the evaluated metric. For SHAPr, all evaluated models obtain a score of 1.0. However, these results should be interpreted with caution, as only the first 10 samples of the test dataset were used in

¹⁰ <https://huggingface.co/dataset>

¹¹ <https://huggingface.co/datasets/stanfordnlp/imdb>

¹² <https://www.kaggle.com/datasets/hojjatk/mnist-dataset>

¹³ <https://www.cs.toronto.edu/~kriz/cifar.html>

this evaluation. The SHAPr results therefore primarily demonstrate the successful execution and applicability of the evidence extraction pipeline to different NLP models rather than providing a comprehensive comparison of their privacy characteristics.

For the image-based models, the SPADE extractor produces a score of 0.948 for both the MNIST and CIFAR-10 models, while the single-sample SPADE evaluation on CIFAR-10 results in a slightly higher score of 0.964. The CLEVER evaluation of the CIFAR-10 model produces a score of 5.0, demonstrating that robustness evidence can also be generated using a complementary robustness metric. The privacy analysis using SHAPr shows a noticeable difference between the two evaluated models, with a score of 0.099 for the CIFAR-10 model and 0.8 for the MNIST model. In addition, the poisoning evaluation on CIFAR-10 reports a target-category recognition probability of 0.872 and an accuracy of 36.45% on poisoned test examples, demonstrating that *AI-SEC* can also generate evidence related to model behaviour under data-poisoning attacks. Since the individual extractors assess different security and trustworthiness properties and use different metric scales, their numerical scores should not be directly compared with each other. Rather, these results demonstrate that *AI-SEC* can successfully execute different evidence extraction methods across multiple image datasets and models and represent their results as corresponding security evidence.

Table 3. Test results on NLP models¹⁴

Method	Model Name	Dataset	Score (Evidence)	Note
SPADE	bert-base-uncased	IMDb	1.00	-
SPADE	distilbert-base-uncased	IMDb	1.00	-
SPADE	roberta-base	IMDb	1.00	-
SPADE	albert-base-v2	IMDb	0.93	-
SPADE	xlnet-base-cased	IMDb	0.96	-
SHAPr	bert-base-uncased	IMDb	1.0	Only the first 10 samples of test data were used.
SHAPr	distilbert-base-uncased	IMDb	1.0	Only the first 10 samples of test data were used.
SHAPr	roberta-base	IMDb	1.0	Only the first 10 samples of test data were used.
SHAPr	albert-base-v2	IMDb	1.0	Only the first 10 samples of test data were used.
SHAPr	xlnet-base-cased	IMDb	1.0	Only the first 10 samples of test data were used.

Table 4. Test results on images datasets

Method	Model Name	Dataset	Score (Evidence)	Note
SPADE	mnist_net	MNIST	0.948	-
SPADE	cifar_net	CIFAR-10	0.948	-
Clever Score	cifar_net	CIFAR-10	5.0	-
single SPADE	cifar_net	CIFAR-10	0.964	-
SHAPr	cifar_net	CIFAR-10	0.099	-

¹⁴ Regarding the absence of LIME and Poison data: Poison data development for the NLP section remains unimplemented. LIME's overall text content score is saved as an HTML file, while the text box outputs binary positive/negative evaluations for the top ten words, making it inconvenient to present within a table.

Method	Model Name	Dataset	Score (Evidence)	Note
SHAPr	mnist_net	MNIST	0.8	-
poison	cifar_net	CIFAR-10	Test Success Rate (Probability of poisoned sample recognized as a target category is 0.872): Accuracy on poisoned test examples: 36.45%	-
Explain_lime	mnist_net	MNIST	0.8390	-
Explain_lime	cifar_net	CIFAR-10	-	Runtime Error: mat1 and mat2 shapes cannot be multiplied (1x256 and 400x120)

Overall, the results across the NLP and image-based use cases demonstrate that the *AI-SEC* evidence extraction pipeline is applicable to different model architectures, datasets, and evidence types. The purpose of these experiments is primarily to validate the functionality and applicability of the implemented extractors rather than to benchmark or directly compare the security properties of the individual models.

3.3 Delivery and usage

The following subsections detail the delivery and usage of *AI-SEC* for EMERALD.

3.3.1 Package information

AI-SEC provides a comprehensive solution for evaluating model robustness, performing adversarial attacks, assessing privacy risks, and explaining model predictions using techniques like LIME. The package includes scripts, configurations, and resources needed for various deep learning model operations as depicted in Table 5.

Table 5. Overview and description of package structure for the *AI-SEC*

Folder	Description
art/	Imported external library
explained_imgs/	Stores result of explained images
models/	Stores the trained model files
models/my_model.pth	The model file, provided by the user (in .pth format)
trigger/	Folder of trigger images used for performing hidden trigger
backdoor_attacks/	Folder to upload image files uploaded by the user, used for LIME model explanation, etc.
toolbox.py	Main script file to execute various operations (e.g., CLEVER score calculation, privacy assessment, data poisoning)
model.py	Script file for neural network architecture
mydata.py	Script file for import custom datasets
config.yaml	Configuration file that specifies model paths, trigger paths, and other settings
mydata/	Folder to store custom data
mydata/labels.csv	CSV Label File
mydata/images/images_upload/class_0	Image folder of custom datasets

Folder	Description
utils.py	Utility functions directory containing helper scripts
README.md	Instructions to build and use the extractor

3.3.2 Installation

In this section, we provide the steps and guidelines to help users set up and install *AI-SEC*. These instructions are crucial for ensuring that the component is correctly configured and operational on a user's system.

- 1. Install Required Libraries:** Run the following commands to install the necessary libraries:


```
pip install torch==1.13.1
pip install torchvision==0.14.1
pip install lime==0.2.0.1
```
- 2. Modify the Configuration File (config.yaml):**
 - Set the `model_path` to the location of your trained model checkpoint.
 - Set the `trigger_path` to the location of the trigger inputs used for data poisoning.
- 3. Define Your Neural Network Architecture:**
 - Implement the `Net` class in `model.py` to define your custom neural network.
- 4. Store Model Files:**
 - Place your trained model files (e.g., `my_model.pth`) in the `models/` directory.
- 5. Prepare Inputs for LIME Explanations:**
 - Put the inputs you want to process into the `images_upload/class_0/` directory.
 - Put the inputs you want to process into the `NLP_upload/class_0/` directory.
- 6. Prepare Trigger inputs for Data Poisoning:**
 - Place the trigger inputs in the `trigger/` directory.

By following these steps, you can properly set up your environment for model training, data processing, and analysis.

3.3.3 Instructions for use

Within the EMERALD project, the *EMERALD-UI* component is used to access and manage the workflows. *AI-SEC* is not directly accessible through the *EMERALD-UI*; instead, it is integrated into the EMERALD workflows through the *Evidence Store*. Figure 3 shows the UI view for setting the evidence extractors for a Target of Evaluation.

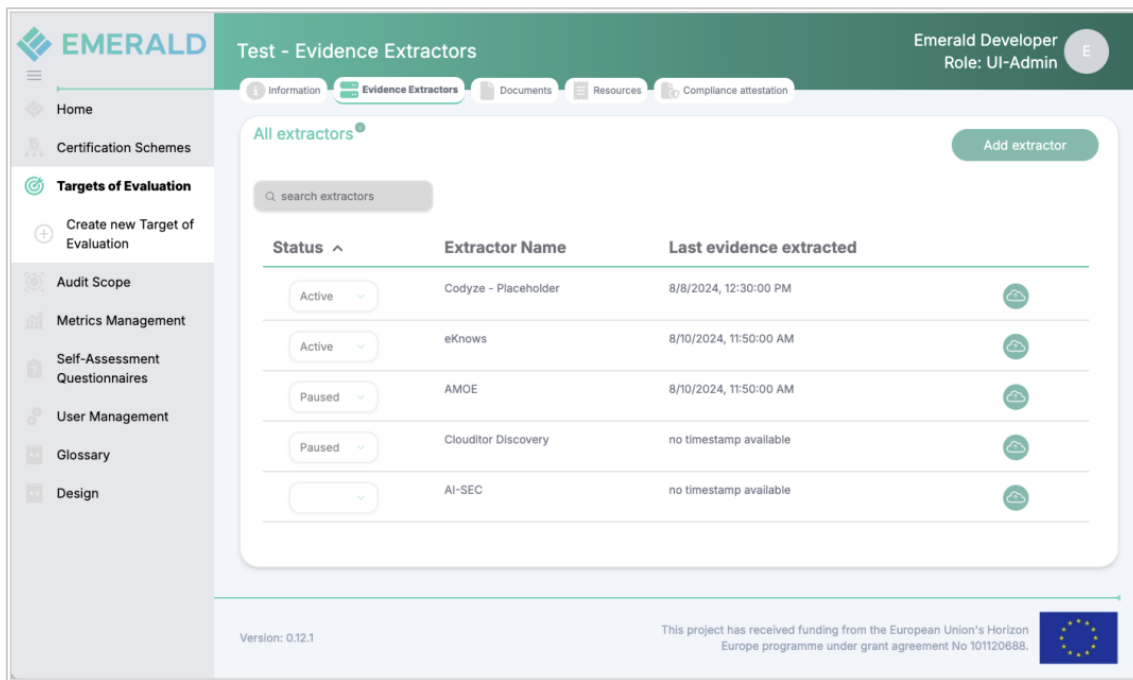


Figure 3: EMERALD UI view for setting the evidence extractor for a Target of Evaluation

In the following, we will provide instructions on how to use the tool for calculating the robustness scores, assessing privacy, performing data poisoning, and explaining model predictions.

Model selection:

1. Custom Own Dataset with PyTorch:

Adjust config.yaml file in the same directory with the following format:

```
$ model_path: 'path/to/your/model.pth'
$ trigger_path: 'path/to/trigger/image.png'
```

Define the Net class in model.py to represent your neural network architecture. Store your .pth files in the models/ directory (e.g., models/my_model.pth).

2. Use Hugging Face Model:

Specify the *model_class* and *model_name* in Hugging Face.

```
$ python toolbox.py -m <model_class> --model_name <model_name>
```

<model_class>: class of the Hugging Face model ("BertModel", "BertForSequenceClassification", "BertForTokenClassification", "BertForQuestionAnswering", etc.)

<model_name>: name of the Hugging Face model ("bert-base-uncased", "gpt2", etc.)

Calculating evidence:

1. Calculate CLEVER Scores

To evaluate the robustness of your model using CLEVER scores, follow these steps:

```
$ python your_script.py -d <dataset> -t robustness -c <nb_classes>
```

<dataset>: Specify the dataset you are using (e.g., "cifar10", "mnist", "mydata").

<nb_classes>: Replace with the total number of classes in your dataset.

2. Assess Privacy (SHAPr Leakage)

To evaluate the privacy of the model using SHAPr leakage analysis, use the following command:

```
$ python your_script.py -d <dataset> -t privacy -c <nb_classes>
```

<dataset>: Indicate your dataset.

<nb_classes>: Specify the number of classes in the dataset.

3. Perform Data Poisoning

To generate poisoned data and evaluate the impact of a data poisoning attack, execute the following commands:

- To generate poisoned data and evaluate the attack:

```
$ python your_script.py -d <dataset> -t poison -c <nb_classes> -s  
<patch_size> -test
```

- To generate poisoned data only:

```
$ python your_script.py -d <dataset> -t poison -c <nb_classes> -s  
<patch_size>
```

<dataset>: Specify your dataset.

<nb_classes>: Replace with the number of classes in your dataset.

<patch_size>: Define the size of the patch for the poisoned data.

4. Explain Model Predictions Using LIME

To generate explanations for model predictions with LIME, run the following command:

```
$ python your_script.py -d <dataset> -t explain -c <nb_classes> -ch  
<num_channels>
```

<dataset>: Specify your dataset.

<nb_classes>: Replace with the number of classes in the dataset.

<num_channels>: Set the number of channels in your input images (e.g., 1 for grayscale, 3 for RGB).

Example for using CIFAR-10 dataset¹⁵:

Here are examples of using the commands with the CIFAR-10 dataset:

1. Calculate CLEVER Scores:

```
$ python toolbox.py -d cifar10 -t robustness -c 10
```

2. Assess Privacy:

```
$ python toolbox.py -d cifar10 -t privacy -c 10
```

¹⁵ <https://www.cs.toronto.edu/~kriz/cifar.html>

3. Perform Data Poisoning:

```
$ python toolbox.py -d cifar10 -t poison -c 10 -s 8 -test
```

4. Explain Model Predictions:

```
$ python toolbox.py -d cifar10 -t explain -c 10 -ch 3
```

Notes:

- Replace placeholders (<dataset>, <nb_classes>, <patch_size>, <num_channels>) with the actual values based on your data and model requirements.
- All commands should be executed in the terminal or command line interface where your Python environment is set up.

3.3.3.1 Example for Running the Tool

Here, we provide one example to demonstrate how to run the tool for image datasets. The procedure is identical for NLP datasets.

Data: The datasets used in this example are MNIST¹⁶ and CIFAR-10¹⁷. The MNIST dataset consists of 28×28-pixel grayscale images with 10 classes. The CIFAR-10 dataset comprises 32×32-pixel RGB colour images with 10 classes.

Model: We trained the model using two datasets, namely the model for MNIST (see Table 6) and the model for CIFAR-10 (see Table 7).

Table 6. Setup for the ML model using MNIST dataset

For MNIST	Layer	Description
1	Conv2D + ReLU	32 Filters (3 × 3), stride 1
2	Conv2D + ReLU	64 Filters (3 × 3), stride 1
3	Max Pooling	2 × 2 pooling
4	Flatten	Flatten the tensor
5	Dense (Fully Connected) + ReLU	128 Units
6	Dense (Fully Connected)	10 Units (Classes)

Table 7. Setup for the ML model using CIFAR-10 dataset

For CIFAR-10	Layer	Description
1	Conv2D + ReLU	6 Filters (5 × 5), stride 1, 3
2	Max Pooling	2 × 2 pooling
3	Conv2D + ReLU	16 Filters (5 × 5), stride 1
4	Max Pooling	2 × 2 pooling
5	Flatten	Flatten the tensor
6	Dense (Fully Connected) + ReLU	120 Units
7	Dense (Fully Connected) + ReLU	84 Units
8	Dense (Fully Connected)	10 Units (Classes)

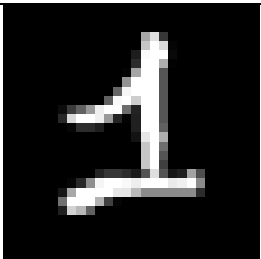
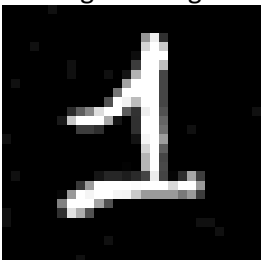
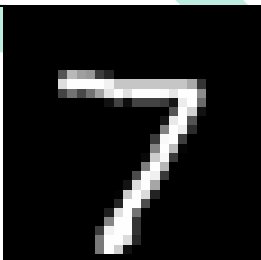
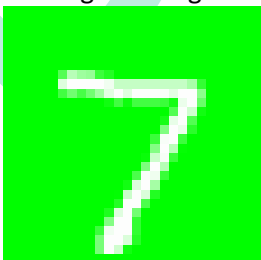
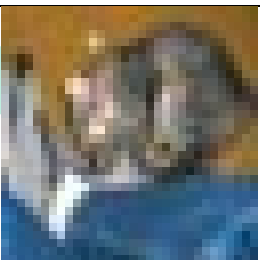
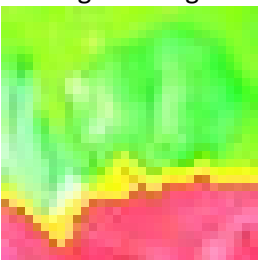
¹⁶ <http://yann.lecun.com/exdb/mnist/>

¹⁷ <https://www.cs.toronto.edu/~kriz/cifar.html>

Result:

Table 8 shows the results of our testing on the two models. The CLEVER scores for the two models are 3.76 and 2.81, respectively. It means that compared to CIFAR-10 model, the MNIST model is more robust against adversarial examples. The SHAPr scores are 0.98 and 0.27, for the privacy leakage analysis, which indicates that the MNIST model is more privacy preserved. For data poisoning, we provide images with the necessary pixel modifications, and when collecting evidence, the amount of pixel change will be recorded as evidence. For LIME, we present the model's explanation results, and in the evidence collection, we provide binary results (0 indicates an unreasonable explanation, while 1 indicates a reasonable explanation).

Table 8. Results on MNIST and CIFAR10 using CLEVER score, SHAPr score, data poisoning and LIME

	Result for the model (MNIST)	Result for the model (CIFAR-10)	Note
CLEVER Scores	3.76/5	2.81/5	Calculations were made using 50 samples from the dataset
SHAPr Score	0.98/1	0.27/1	The computation time is long; it takes 10,000 iterations.
Data Poisoning	 original image  poisoned image	 original image  poisoned image	The poisoning difficulty is quantified based on the number of pixels that need to be modified in order to achieve a successful backdoor attack A score of 1 indicates that <i>all</i> pixels must be changed for the attack to succeed, meaning the model is highly resilient and very difficult to poison. Conversely, a lower score means that only a small portion of pixels need to be altered, indicating that the model is more vulnerable to poisoning.
LIME-based explanations	 original image  model prediction	 original image  model prediction	It shows which pixel regions the model used to make its decision (highlighting the digit shape in green). In our toolkit, we convert these images into an integer. We divide the explained region by the total number of pixels to obtain the final output score. For example, in this image, the green area is divided by the total number of pixels, and the resulting value represents the explainability score. 1 means the result can be explained and 0 means the result cannot be explained.

3.3.4 Licensing information

Since LIME and SHAPr use permissive MIT licenses, we choose to license this tool under the MIT License.

3.3.5 Download

The currently implemented parts are stored on EMERALD's public Gitlab¹⁸.

3.4 Demo

We have developed a small demonstration to present the *AI-SEC* tool in an intuitive and user-friendly way (see Figure 4 and Figure 5). The demo allows users to upload a model or select a predefined one, run different analysis modules such as robustness, privacy, and explainability, and visualize the corresponding results through interactive charts and evidence summaries. This demonstration aims to help both technical and non-technical stakeholders understand how *AI-SEC* performs evidence extraction and how the resulting metrics can be interpreted within the EMERALD compliance framework.

This demo shows the following abilities:

1. Functional realization

Through the graphical user interface, users can configure an *AI-SEC* analysis workflow by selecting the model, dataset, and evaluation criterion. Depending on the selected model and use case, the interface provides the applicable evidence extraction methods for assessing properties such as robustness, privacy, data integrity, and explainability. The workflow consists of the following steps:

- a) Select the ML model to be analysed. Users can select either a supported pre-trained model or provide their own model.
- b) Select the dataset. Users can choose a dataset available through Hugging Face or provide a custom dataset.
- c) Select the evaluation criterion. Users select the security or trustworthiness property to be assessed and the corresponding evaluation method, such as robustness, privacy, data integrity, or explainability.
- d) Review the generated command. The interface displays the corresponding terminal command that will be executed, allowing users to understand and reproduce the analysis workflow from the command line.

For LLMs, only black-box evaluation metrics are currently supported. Therefore, the range of available evidence extraction methods and assessment capabilities for LLMs is currently more limited than for other supported model types.

2. Real-time interaction

Demo users can run and view a visual representation of the analysis results (e.g. SPADE Score), detailed command line output, and the path to the results file directly in the interface.

¹⁸ <https://git.code.tecnalia.dev/emerald/public/components/ai-sec>

User Guide

> Click to view detailed instructions

Step 1: Select the type of task

Select the type of data to be processed:

☒ NLP (Natural Language Processing) ☐ LLM (Large Language Model Security) ☐ Image (Image Processing) ☐ SVM/Classical ML (Spam Classification)

Step 2: NLP Configuration

Parameter Configuration

Dataset selection

Source of datasets: ⓘ

☒ HuggingFace Dataset
☐ Customised dataset

Choose HuggingFace Dataset:

imdb

Model Selection

Model Source: ⓘ

☒ Pre-trained Model ☐ Upload Custom Model

Select Pre-trained Model: ⓘ

distilbert-base-uncased-finetuned-sst-2-english

Sentiment model – already fine-tuned on movie reviews (SST-2). No training needed.

Methods of analysis

Selection of analytical methods (can select multiple):

☒ SPADE ⓘ ☐ SHAPr ⓘ ☐ LIME ⓘ ☐ POISON ⓘ

Command Preview

Command 1 (SPADE):

```
python toolbox.py -d hf_dataset --hf_dataset imdb --hf_text_field text --hf_
```

current configuration

Dataset: imdb (HuggingFace) Model: distilbert-base-uncased-finetuned-sst-2-english Methods of analysis: SPADE Number of commands: 1

SPADE: Evaluates the robustness of the model, with scores ranging from 0-1, the closer to 1 the more robust it is.

Step 3: Execute

Operating analysis

Reconfiguration

Estimated implementation time: 1 method(s) × ~3 minutes = ~3 minutes

Figure 4. The demonstration interface consists of three main steps

Combined Analysis Results

Successfully completed: SPADE, SHAPr, LIME

SPADE Results

SPADE Robustness Score on Features

0.9850

Very robust — Features are strong robust to adversarial perturbations

View SPADE detailed log

SPADE complete output:

```
/opt/miniconda3/envs/emerald/lib/python3.10/site-packages/requests/__init__.py:86: RequestsDependencyWarning: Unable to find acceptable character detection dependency (chardet or charset_normalizer).
warnings.warn(
Warning: You are sending unauthenticated requests to the HF Hub. Please set a HF_TOKEN to enable higher rate limits and faster downloads.
robustness_spade_NLP
Dataset 'imdb' loaded successfully.
Train samples: 25000, Test samples: 25000

Loading weights: 0%|          | 0/104 [00:00<?, ?it/s]
Loading weights: 100%|██████████| 104/104 [00:00<00:00, 12886.87it/s]
```

LIME Results

LIME Explanation Score (Local Fidelity R²)

0.5964

Moderate explanation quality

Generated LIME Explanation HTML:

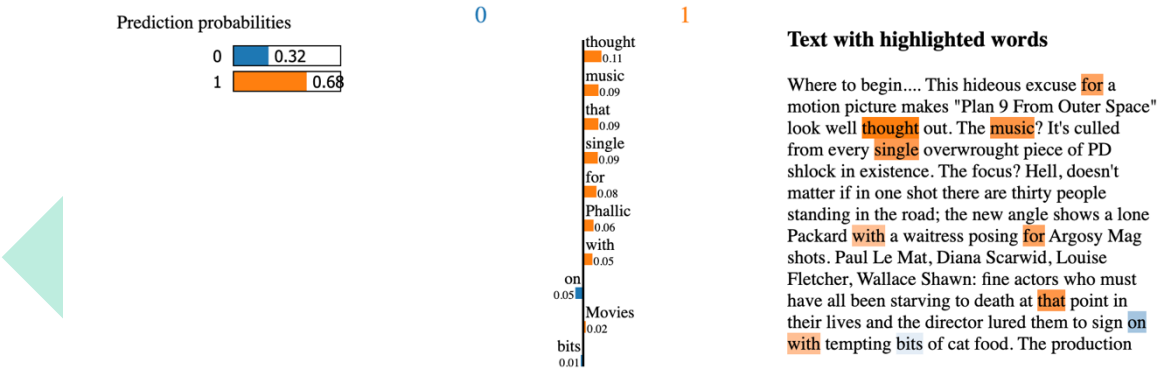


Figure 5. After running the analysis, we can see the results and save the results for further investigation

3.5 Connection to the Evidence Store

3.5.1 Generating the evidence JSON file

A new `evidence_wrapper.py` component was added. It runs *AI-SEC* analysis, parses the output, converts the result into EMERALD evidence format, and can either save the evidence as JSON or send it directly to the *Evidence Store*.

The wrapper supports command-line options for dataset, task type, model, number of classes, Hugging Face datasets, poisoning parameters, sample selection, output JSON generation, dry-run mode, target-of-evaluation ID, and tool ID.

A metric extraction layer was implemented in the evidence conversion flow. It extracts relevant security values from *AI-SEC* output, including robustness score, vulnerability level, privacy leakage score, and membership inference score. These values are stored as evidence labels, such as `metric_RobustnessScore`, `metric_VulnerabilityLevel`, `metric_PrivacyLeakageScore`, and `metric_MembershipInference`.

A new `validated_evidence.py` utility was introduced to validate generated evidence before submission. It checks mandatory fields, timestamp format, supported resource types, required `machineLearningService` fields, label typing, and whether the raw analysis payload is valid JSON.

The *AI-SEC* Command-Line Interface (CLI) was also extended with a Go/Cobra-based interface definition. It now models parameters for datasets, Hugging Face input, models, task types, poisoning options, explanation options, output formats, database settings, API ports, discovery scheduling, target-of-evaluation ID, and tool ID.

An example of this JSON file (`dev_evidence.json`) is as follows:

```
{
  "id": "8d5a5f64-3d9e-4c9a-8f4d-5d0e6d6f8b21",
  "timestamp": "2026-06-29T12:00:00Z",
  "targetOfEvaluationId": "00000000-0000-0000-0000-000000000000",
  "toolId": "ai-sec",
  "resource": {
    "machineLearningModel": {
      "id": "2e0fdc63-9c63-4d6d-9a2f-7dd7c5fd2b72",
      "name": "svm_sms_classifier",
      "description": "AISEC security assessment for an SVM-based spam
email classifier.",
      "adversarialRobustnessScore": 0.9064,
      "membershipInferenceResilience": 0.82,
      "poisonedDataLevel": 0.18,
      "poisoningResilienceLevel": 0.82,
      "explainability": 0.74,
      "explainabilityEnabled": true,
      "labels": {
        "dataset": "sms_spam_email",
        "dataDomain": "nlp_email_classification",
        "modelType": "svm",
        "task": "robustness_spade",
        "supportedImageDatasets": "mnist,cifar10",
        "supportedNlpDatasets": "imdb,huggingface_text,sms_spam_email",
        "supportedLlmFunctions": "limited",
        "supportedTraditionalMl": "svm_spam_email_classifier"
      },
      "raw": "{\"tool\":\"ai-
sec\",\"status\":\"completed\",\"analysis\":{\"dataset\":\"sms_spam_email"
}
```

```

l\", \"task\": \"robustness_spade\", \"model\": \"svm_sms_classifier\", \"nb_
classes\": 2, \"model_type\": \"svm\", \"data_domain\": \"nlp_email_classific
ation\"}, \"supported_scope\": {\"image_datasets\": [\"mnist\", \"cifar10\"],
\"nlp_datasets\": [\"imdb\", \"huggingface_text\", \"sms_spam_email\"], \"ll
m_functions\": [\"limited\", \"traditional_ml\": [\"svm_spam_email_classifi
er\"]}, \"results\": {\"spade_score\": 0.9064, \"membership_inference_resili
ence\": 0.82, \"poisoned_data_level\": 0.18, \"poisoning_resilience_level\":
0.82, \"explainability\": 0.74, \"explainability_enabled\": true}}\"
}
},
\"experimentalRelatedResourceIds\": [
  \"MLSecurity\",
  \"MLRobustness-01.1\",
  \"MLPrivacy-02.1\",
  \"MLPoisoning-03.1\",
  \"MLExplainability-04.1\"
]
}

```

The current workflow supports several AI security analysis tasks, including adversarial robustness, privacy leakage, poisoning resilience, and explainability-related analysis. Results from SPADE, CLEVER, SHAPr, poisoning tests, and explanation tasks can be represented as machine-readable evidence. The supported evaluation scope covers multiple model and dataset types, including image classification datasets such as MNIST and CIFAR-10, NLP datasets such as IMDb and Hugging Face text datasets, limited LLM-related evaluation functions, and a classical machine learning use case based on an SVM spam email classifier. This allows *AI-SEC* to generate evidence for both deep learning and traditional ML scenarios across image, text, and email classification tasks.

3.5.2 Sending AI-SEC evidence to the Evidence Store

The following example demonstrates the workflow for requesting an access token, submitting *AI-SEC*-generated evidence to the *Evidence Store*, and retrieving the uploaded evidence. All sensitive information has been replaced with placeholders for documentation purposes.

Step1. Request an Access Token (valid for 5 minutes)

```

export TOKEN=$(curl -s -X POST \
  "https://<KEYCLOAK_SERVER>/realms/<REALM_NAME>/protocol/openid-
connect/token" \
  -H "Content-Type: application/x-www-form-urlencoded" \
  -d "grant_type=password" \
  -d "username=<USERNAME>" \
  -d "password=<PASSWORD>" \
  -d "client_id=<CLIENT_ID>" \
  | jq -r '.access_token')

```

Step2. Check that the Token Has Been Generated

```
echo ${#TOKEN}
```

Step3. Send Evidence to the *Evidence Store*

```

python evidence_wrapper.py \
  -d hf_dataset \
  --hf_dataset imdb \
  --hf_text_field text \
  --hf_label_field label \
  -m AutoModelForSequenceClassification \

```

```
--model_name distilbert-base-uncased-finetuned-sst-2-english \
-t robustness_spade_NLP \
-c 10 \
--evidence-store-url https://<EVIDENCE_STORE_URL> \
--token "$TOKEN" \
--output-json /tmp/dev_evidence.json
```

Finally, we can check if our evidence was sent to the *Evidence Store*.

Step4. Retrieve the Generated Evidence ID

```
EVIDENCE_ID=$(jq -r '.id' /tmp/dev_evidence.json)
echo "$EVIDENCE_ID"
```

Step5. Check the Uploaded Evidence

```
curl -k \
-H "Authorization: Bearer $TOKEN" \
"https://<EVIDENCE_STORE_URL>/v1/evidence_store/evidences/$EVIDENCE_ID" \
| jq .
```

3.6 Pilot 1 Test: SVM-based Spam Email Classifier

In Pilot 1 (IONOS), *AI-SEC* was tested on a classic machine learning (ML) model for classifying spam emails. The evaluated model is a scikit-learn pipeline composed of a Tfidf Vectorizer and a linear Support Vector Machine Classifier (SVC). The text preprocessing step converts email/SMS messages into TF-IDF features with a maximum of 1,000 features, lowercase normalization, and English stop-word removal. The classifier uses a linear kernel and probability outputs, enabling both class prediction and confidence reporting.

The model was trained on the Hugging Face sms_spam dataset using an 80/20 train-test split. The recorded baseline performance of the trained model is:

```
{
  "accuracy": 0.9812,
  "precision": 0.9730,
  "recall": 0.8944,
  "f1": 0.9320
}
```

For the Pilot 1 evidence-generation test, a local sample file named `example_customer_emails.csv` was used. The file contains 14 labelled email/SMS-style messages with two classes: 0 for ham/legitimate messages and 1 for spam/phishing-like messages. Table 9 shows some examples of the input sentences.

Table 9. Example of the input sentences.

Text	Label
Congratulations! You've won a free iPhone! Click here to claim.	1
Hi Sarah, Can you review the Q3 report by Friday? Thanks!	0
URGENT: Your account will be closed! Verify now.	1
Meeting rescheduled to 2pm tomorrow in Conference Room B.	0

AI-SEC evaluated the model using the SVM/classical ML execution path in the Evidence Wrapper. The test included sample prediction, confidence extraction, adversarial robustness analysis with SPADE-style text perturbations, poisoning resilience testing, membership inference resilience input, and LIME-based interpretability analysis. The Pilot 1 dry-run generated an *Evidence Store* compatible `machineLearningModel` evidence object.

Observed sample predictions were:

- "Congratulations! You've won a free iPhone! Click here to claim." -> SPAM
- "Hi Sarah, Can you review the Q3 report by Friday? Thanks!" -> HAM
- "URGENT: Your account will be closed! Verify now." -> SPAM

The generated Pilot 1 evidence included the following machine learning security values:

```
{
  "adversarialRobustnessScore": 0.5556,
  "membershipInferenceResilience": 0.82,
  "poisoningResilienceLevel": 1.0,
  "explainability": 0.8143,
  "explainabilityEnabled": true
}
```

For documentation, it is useful to mention that the evidence payload is structured according to the *Evidence Store* schema, using `machineLearningModel` as the resource type. The evidence includes the *AI-SEC* tool ID, target-of-evaluation ID, model metadata, normalized security metrics, labels describing the dataset and task, and a raw JSON field preserving the full *AI-SEC* analysis output for auditability.

One small note: during local execution, scikit-learn reported a version mismatch warning because the model was trained with scikit-learn 1.1.3 and loaded with 1.7.2. The test still completed successfully, but for production Pilot 1 execution the model should ideally be loaded with the same scikit-learn version used during training.

In addition to local evidence generation, the Pilot 1 test also validated submission to the remote Pilot 1 *Evidence Store*. The *AI-SEC* Evidence Wrapper was configured to send the generated `machineLearningModel` evidence object to the Pilot 1 *Evidence Store* endpoint.

3.7 Limitations and future work

The current evaluations using accessible models have demonstrated promising results and highlight the potential of *AI-SEC* for AI-specific security assessment. During this reporting period, the tool was further validated in the Pilot 1 use case using a traditional SVM-based spam email classifier. The generated assessment results were successfully converted into EMERALD-compliant evidence, demonstrating that *AI-SEC* can support the security assessment of both deep learning models and conventional machine learning models, provided that the model artifact and representative test data are available.

However, a significant limitation remains. Many cloud-based AI services do not provide direct access to their underlying models, training data, or internal inference pipelines, making it impossible to perform comprehensive white-box or model-level security evaluations. In such cases, *AI-SEC* cannot directly analyse the proprietary production model and must instead rely on indirect evaluation strategies.

To address this limitation, one possible approach is to use a proxy model as a substitute for the inaccessible cloud-based model. A representative model can be trained locally using similar datasets, comparable model architectures, or observed input-output behaviour of the original service. *AI-SEC* can then perform robustness, privacy leakage, poisoning resilience, and explainability analyses on the proxy model. The resulting assessment metrics can serve as indirect evidence of the trustworthiness and security properties of the deployed cloud-based AI model.

Another limitation is related to tool optimization and execution environments. Some assessment methods, such as SHAPr, explainability techniques, and adversarial robustness evaluations, are computationally expensive and highly dependent on specific software libraries and execution environments. Therefore, *AI-SEC* requires further optimization to improve execution efficiency, dependency management, and reproducibility across different deployment platforms.

Despite these limitations, *AI-SEC* is already capable of generating valuable AI-specific assurance evidence whenever a model artifact, proxy model, or representative dataset is available. The resulting quantitative metrics can be incorporated as supporting evidence within the EMERALD certification workflow.

DRAFT

4 Conclusions

In this deliverable, as final output of Task 2.4, we presented the design, architecture, and current implementation status of *AI-SEC*. *AI-SEC* follows the holistic approach of the EMERALD framework and is aligned with the technical requirements gathered in WP1 (D1.4 [8]). The report outlines the relationship between the presented component, *AI-SEC* and other parts of the EMERALD framework, detailing the internal structure of the component, its subcomponents, and relevant information about its technical implementation.

AI-SEC supports evidence extraction for machine learning models. The component has also been validated using accessible machine learning models, including a traditional SVM-based spam email classifier in the Pilot 1 use case. In addition, *AI-SEC* supports the analysis of publicly available models and datasets from Hugging Face, as well as user-provided models and datasets. The evaluation results demonstrate that *AI-SEC* can process locally available model artifacts and labelled sample data, generate quantitative AI security assessment metrics, and produce evidence compatible with the *Evidence Store* schema. Furthermore, the generated evidence has been successfully submitted to the *Evidence Store in Pilot 1*, confirming the integration of *AI-SEC* with the EMERALD evidence infrastructure.

5 References

- [1] EMERALD Consortium, “D2.3 Source Evidence Extractor–v2,” 2026.
- [2] EMERALD Consortium, “D2.5 AMOE – v2,” 2026.
- [3] EMERALD Consortium, “D2.9 Runtime evidence extractor – v2,” 2026.
- [4] EMERALD Consortium, “D2.1 Graph Ontology for Evidence Storage: Description of a uniform schema for storing and linking heterogenous data,” 2024.
- [5] EMERALD Consortium, “D2.6 ML model certification–v1,” 2024.
- [6] MEDINA Consortium, “CORDIS - EU research results,” [Online]. Available: <https://cordis.europa.eu/project/id/952633>. [Accessed 06 08 2026].
- [7] EMERALD Consortium, “D1.2 Data modelling and interaction mechanisms - v2,” 2025.
- [8] EMERALD Consortium, “D1.4 EMERALD solution architecture - v2,” 2025.
- [9] EMERALD Consortium, “D3.4 Evidence assessment and Certification–Implementation-v2,” 2025.
- [10] T.-W. Weng, H. Zhang, P.-Y. Chen, Y. Jinfeng, D. Su, Y. Gao, C.-J. Hsieh and L. Daniel, “Evaluating the robustness of neural networks: An extreme value theory approach,” *arXiv preprint*, 2018.
- [11] H. Souri, L. Fowl, R. Chellappa, M. Goldblum and T. Golstein, “Sleeping agent: Scalable hidden trigger backdoors for neural networks trained from scratch,” *Advances in Neural Information Processing Systems* 35, 2022.
- [12] V. Duddu, S. Szyller and N. Asokan, “Shapr: An efficient and versatile membership privacy risk metric for machine learning,” *arXiv preprint*, 2021.
- [13] M. T. Ribeiro, S. Singh and C. Guestrin, ““Why should I trust you?” Explaining the predictions of any classifier,” *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 1135-1144, 2016.
- [14] EMERALD Consortium, “D2.11 Certification Graph–v2: Integration of the graph with semantically linked and combined evidence,” 2026.
- [15] A. Saha, A. Subramanya and H. Pirsiavash, “Hidden trigger backdoor attacks,” in *Proceedings of the AAAI conference on artificial intelligence*, , vol. 34, no. 07, pp. 11957-11965, 2020.