

⚠ **DRAFT** — Automatically generated working document. Not a recognised conformity assessment under Regulation (EU) 2024/1689 Art. 11/47. Must be reviewed by qualified legal counsel before submission to any authority or Notified Body. Scanara accepts no liability for regulatory use of this document.

Justice App - Technical Documentation (Annex IV)

Organisation: —

AI System: Justice App

Risk Level: high_risk

Version: 1

Generated: 2026-07-02 21:00:29.235261+00:00

Table of Contents

System Overview	4
Intended Purpose	5
General Information	5
Risk Classification	7
Scan-Based Risk Indicators	7
Risk Management System (Art. 9)	8
1. Risk Management System Overview	8
2. Risk Identification	8
3. Risk Evaluation — Post-Market Monitoring Feedback Loop	8
4. Risk Management Measures Adopted	9
5. Scope of the Risk Management System	9
6. Combined Effect of Chapter 2 Requirements	9
7. Preventive Measures — Elimination by Design	9
8. Mitigation and Control Measures	10
9. Residual Risk Acceptance	10
10. Testing Against Predefined Metrics and Thresholds	10
11. Real-World Condition Testing	11
12. Testing Throughout the Development Lifecycle	11
13. Vulnerable Groups	11
14. Integration with Sectoral Legislation Risk Management	11
Data and Data Governance	12
Training Data	12
Validation Data	12
Data Quality Measures	12
Data Bias Assessment	12
Technical Specifications	13
Architecture	13
Hardware Requirements	13
Software Dependencies	13
Performance Specifications	13
Integration Points	13
Accuracy & Robustness	14
Accuracy Metrics	14
Robustness Measures	14
Cybersecurity Measures	14
Known Limitations	14

Human Oversight Measures.....	15
Oversight Mechanisms	15
User Interface for Oversight	15
Ability to Override or Reverse	15
Training Requirements for Human Overseers.....	15
Automation Bias Awareness (Art. 14(3)(b)).....	15
Designated Oversight Persons (Art. 14(5))	15
Post-Market Monitoring Plan	17
Monitoring Strategy	17
Incident Reporting Procedures.....	17
Performance Metrics to Track	17
Review Schedule	17
Annex IV Compliance Checklist.....	18
Cybersecurity Measures (Annex IV §2(h))	19
Cybersecurity Measures (Article 15)	19
Standards Applied (Annex IV §2(g))	21
Applied Harmonised Standards (Annex IV §2(g), Art. 40)	21
Common Specifications Applied (Art. 41)	21
Presumption of Conformity (Art. 40(3)).....	21
Deviations from Standards	21
Record-Keeping, Logging and Retention.....	22
Automatic Logging (Art. 12(1))	22
Provider Log Retention (Art. 19(1))	22
Technical Documentation Retention (Art. 18)	22
Incident Record Retention (Art. 73)	22
Deployer Log Retention (Art. 19(2))	23

System Overview

AI System Name: Justice App **Description:** This is a sample ai system to demonstrate the user the Scanara EU AI Act compliance platform. **Intended Purpose:** The simulated system is an AI tool for predicting sentencing outcomes and assessing evidence reliability, deployed by a court administration body. **Technology Stack:** Python **Geographic Scope:** EU

This section provides a high-level overview of the AI system as required by Annex IV of Regulation (EU) 2024/1689.

Intended Purpose

General Information

Article 11, Annex IV §1, Regulation (EU) 2024/1689 (EU AI Act)

Intended Purpose

CaseInsight Judicial Support Suite provides AI-assisted sentencing-outcome prediction, evidence-reliability scoring, and legal-precedent research support to judicial officers within Example Regional Court Administration's criminal courts. All outputs are advisory recommendations for a human judge; the system does not issue, apply, or communicate any binding decision.

Sector of Deployment

Public sector — administration of justice, Annex III §8 (criminal courts, sentencing preparation, evidence review).

Problem Addressed

Court backlogs and uneven access to comparable-case precedent and evidence-quality assessment slow case preparation and can contribute to variability in sentencing recommendations between judges handling similar cases. CaseInsight surfaces relevant precedent, flags evidence-reliability concerns, and produces a sentencing-outcome estimate to inform — not replace — judicial deliberation.

Target Users and Stakeholders

Role	Description
Primary users	Sitting judges and judicial clerks in criminal courts
Deployers	Example Regional Court Administration (body governed by public law)
Affected persons	Defendants, accused persons, victims, witnesses, and other parties to criminal proceedings whose case is processed with system assistance

Measurable Goals and KPIs

Success criteria are defined against held-out historical case data and monitored post-deployment via the Post-Market Monitoring Plan.

KPI	Target	Measurement Method
Sentencing-outcome prediction accuracy	≥ 85% agreement with actual court outcome	Multi-class accuracy on historical case holdout
Evidence-reliability classification accuracy	≥ 90%	F1 score vs. expert-reviewer labels

KPI	Target	Measurement Method
Judge override rate within expected band	10%–60%	Monthly override-log analysis
Disparate impact ratio across protected characteristics	0.8–1.25 (four-fifths rule)	Quarterly demographic-stratified accuracy review

Prohibited Uses

- This system must not be used to issue a sentencing decision, evidentiary ruling, or any output with binding legal effect without review and sign-off by a human judge.
- This system must not be used for real-time biometric identification or any purpose listed under Article 5.
- This system must not be used to profile or score individuals for purposes outside active criminal case preparation.
- Outputs must not be disclosed to defendants, victims, or the public as a standalone document; they must always accompany full judicial reasoning.

Operational Environment

Deployment Context: Single Member State, regional court network (Example Regional Court Administration jurisdictions only)

- **Platform:** cloud (AWS eu-west-1), containerised inference service
- **Network environment:** court-network-internal, no public internet exposure
- **Integration mode:** REST API called from the court case-management system
- **Expected scale:** up to 2,000 case analyses/month, single-tenant per court administration

Ethical Implications

The system operates in a fundamental-rights-sensitive domain where prediction errors or undetected bias can materially affect defendants' right to a fair trial (Charter Art. 47) and the presumption of innocence (Charter Art. 48). Foreseeable risks include automation bias (judges over-relying on AI recommendations) and disparate accuracy across demographic groups arising from historical sentencing-data imbalances. See the Human Oversight Measures section for mitigation measures and the linked Fundamental Rights Impact Assessment (Art. 27) for a full rights-based analysis.

Risk Classification

Risk Level: high_risk

The risk classification has been determined based on automated code scanning and document analysis in accordance with Regulation (EU) 2024/1689.

GPAI Model: This system uses or is a General-Purpose AI model.

Note: This AI system has been classified as **high-risk**. All obligations under Title III, Chapter 2 of the EU AI Act apply.

Annex III Categories: justice

Scan-Based Risk Indicators

Indicator	Value
Compliance Score	26.0% (F)
Total Findings	155
Critical	1
High	58

Risk Management System (Art. 9)

This section documents the risk management system for CaseInsight Judicial Support Suite in accordance with Article 9 of Regulation (EU) 2024/1689 (EU AI Act).

1. Risk Management System Overview

Article 9(1), Regulation (EU) 2024/1689

A risk management system has been established for CaseInsight Judicial Support Suite as a continuous, iterative process intended to run throughout the entire lifecycle of the AI system. Formal documentation of the risk-identification and evaluation methodology is currently in an early implementation phase; this dossier represents the first structured iteration and will be updated following the next AIRA review cycle.

2. Risk Identification

Article 9(2)(a)(b), Regulation (EU) 2024/1689

2.1 Known and Foreseeable Risks

- **Automation bias** — judges over-relying on sentencing-outcome predictions without independently evaluating case facts (risk to the right to a fair trial, Charter Art. 47).
- **Disparate accuracy across demographic and jurisdictional groups** — training-data representativeness and bias have not yet been assessed (risk to non-discrimination, Charter Art. 21).
- **Unvalidated evidence-reliability scores** — accuracy validation is not currently implemented (`EvidenceReliabilityScorer.validate_accuracy`), risking reliance on unvalidated reliability assessments in evidentiary review.
- **Absent uncertainty disclosure** — sentencing-outcome and evidence-reliability outputs do not currently surface confidence intervals, risking overconfident interpretation by judicial users.
- **No functioning stop-button / human-intervention pathway** in the judicial decision support module.

2.2 Risks Under Reasonably Foreseeable Misuse

- Treating AI-generated sentencing predictions as a de facto sentencing recommendation rather than an advisory input, particularly under caseload pressure.
- Citing AI-surfaced "relevant precedents" without independent verification — the reference precedent-retrieval implementation is not yet complete, and any integration must not silently substitute unverified output for verified legal research.

3. Risk Evaluation — Post-Market Monitoring Feedback Loop

Article 9(2)(c), Regulation (EU) 2024/1689

The post-market monitoring feedback loop described in the Post-Market Monitoring Plan (Art. 72) has not yet completed a full review cycle, as the system is in early deployment. The first scheduled quarterly review (see Post-Market Monitoring Plan, Review Schedule) will be the first source of operational risk-evaluation data feeding back into this section.

4. Risk Management Measures Adopted

Article 9(2)(d), Regulation (EU) 2024/1689

Measures adopted to date: - Advisory-only output design — no AI output is automatically applied to a case record (see Human Oversight Measures). - Onboarding training for judges and clerks addressing automation-bias awareness.

Measures planned but not yet implemented (see Human Oversight Measures for the full remediation list): - Mandatory human-oversight gate prior to use of any prediction. - Automated automation-bias detection via override-rate monitoring. - Demographic bias assessment and mitigation for training data (see Data and Data Governance).

5. Scope of the Risk Management System

Article 9(3), Regulation (EU) 2024/1689

This risk management system addresses risks that can be mitigated through system design (e.g. advisory-only output framing, planned oversight gating) and through information provided to Example Regional Court Administration via the Instructions for Use. Risks arising from judicial decision-making itself remain outside the technical scope of this system and are governed by existing judicial conduct and procedural rules.

6. Combined Effect of Chapter 2 Requirements

Article 9(4), Regulation (EU) 2024/1689

Residual risk has been estimated considering the combined effect of: incomplete data governance (Art. 10 — bias assessment not yet performed), gaps in human oversight design (Art. 14), and absent accuracy validation for evidence-reliability scoring (Art. 15). Given these compounding gaps, current residual risk for the evidence-reliability and sentencing-prediction components is assessed as **elevated** pending completion of the remediation items tracked throughout this dossier.

7. Preventive Measures — Elimination by Design

Article 9(5)(a), Regulation (EU) 2024/1689

Risks that have been eliminated or reduced through design choices applied to the AI system:

Risk	Preventive Measure (Design)	Status
AI output used as sole basis for a ruling	Outputs are structurally advisory-only; not written directly to case disposition records	Implemented
Direct public disclosure of AI scores	Outputs surfaced only within the internal case-file UI, not exposed to defendants or the public	Implemented

Risk	Preventive Measure (Design)	Status
Real-time biometric identification misuse	Not a system capability; out of scope by design	Implemented (not applicable)

8. Mitigation and Control Measures

Article 9(5)(b), Regulation (EU) 2024/1689

For risks that could not be eliminated by design, the following mitigation and control measures have been implemented or are planned:

Residual Risk	Mitigation / Control Measure	Status
Automation bias in sentencing predictions	Mandatory human-oversight acknowledgement gate	Planned — not yet implemented
Undetected demographic disparity in accuracy	Bias assessment across protected characteristics	Planned — not yet implemented
Unvalidated evidence-reliability accuracy	Implement accuracy validation against held-out test set with declared thresholds	Planned — not yet implemented
Uncritical acceptance of AI recommendations	Override-rate monitoring (see Post-Market Monitoring Plan)	Planned — not yet implemented

9. Residual Risk Acceptance

Article 9(5)(c), Regulation (EU) 2024/1689

Responsible Person: *[To be completed by the provider — name and role of Example AI Systems Ltd.'s designated AI Act compliance lead]*

Residual Risk Acceptance Statement: Residual risk associated with the gaps listed in Section 8 has not yet been formally accepted pending completion of the planned mitigation measures. Interim deployment, where it occurs, is limited to supervised pilot use with mandatory judge sign-off on every output.

Information / Training Provided to Deployers: Onboarding training module (see Human Oversight Measures). The full Instructions for Use dossier is to be completed prior to general deployment.

10. Testing Against Predefined Metrics and Thresholds

Article 9(6), Regulation (EU) 2024/1689

Testing of the AI system has been carried out against predefined metrics and probabilistic thresholds. Thresholds were defined **prior** to conducting the tests.

Test	Predefined Threshold	Result	Pass / Fail
Sentencing-outcome multi-class accuracy (holdout)	≥ 85%	<i>[To be completed — pending formal test run]</i>	Pending
Evidence-reliability F1 score (holdout)	≥ 90%	Not run — accuracy validation not yet implemented	Fail (not implemented)

11. Real-World Condition Testing

Article 9(7) and Article 60, Regulation (EU) 2024/1689

No real-world condition testing under Article 60 has been conducted for this AI system.

12. Testing Throughout the Development Lifecycle

Article 9(8), Regulation (EU) 2024/1689

Testing to date has consisted of offline evaluation on historical holdout data only, without a per-demographic accuracy breakdown. A structured testing schedule spanning unit, integration, and pre-deployment validation phases is planned but not yet formally documented — tracked as an open item for the next dossier revision.

13. Vulnerable Groups

Article 9(9), Regulation (EU) 2024/1689

An assessment of the adverse impact of the AI system on persons under 18 years of age and other groups at higher risk of vulnerability has been carried out.

Assessment for Persons Under 18: Not yet formally assessed. Minors in proceedings are identified as an affected group in the linked Fundamental Rights Impact Assessment (Art. 27); a dedicated impact assessment for this group is an open item.

Assessment for Other Vulnerable Groups: Not yet formally assessed. Persons with disabilities and persons from minority groups are identified as affected groups in the FRIA; dedicated assessment pending.

14. Integration with Sectoral Legislation Risk Management

Article 9(10), Regulation (EU) 2024/1689

No sectoral legislation imposing additional risk management requirements applies to this AI system. Article 9(10) integration is not applicable.

Data and Data Governance

This section documents the data governance practices for the AI system in accordance with Article 10 of Regulation (EU) 2024/1689 (EU AI Act).

Training Data

Historical case records, sentencing outcomes, and evidence-reliability corpora sourced from Example Regional Court Administration's case-management system exports (2015–present), restricted to closed cases only, PII-redacted prior to model use.

Validation Data

A held-out 20% split of the same corpus, stratified by offence category. Evidence-reliability labels are independently assigned by two legal reviewers, with disagreements adjudicated by a senior clerk.

Data Quality Measures

A minimum of 500 labelled examples is required per offence category before inclusion in training. Automated schema validation runs on ingest, with missing-value and outlier checks prior to model fitting.

Data Bias Assessment

Open item. Demographic-stratified representativeness and bias analysis across protected characteristics and court jurisdictions has not yet been formally conducted for the current model version. This is tracked as a required action ahead of the next model iteration — see the Post-Market Monitoring Plan (Performance Metrics and Thresholds) for the fairness metrics that will govern this assessment once implemented.

Technical Specifications

Architecture

Technology Stack: Python 3.12; scikit-learn (`GradientBoostingClassifier` , 200 estimators, for sentencing-outcome prediction; `svc` with RBF kernel and `StandardScaler` normalisation for evidence-reliability scoring); a retrieval-based legal-precedent module for judicial decision support; containerised on AWS ECS.

Hardware Requirements

Standard CPU-based inference — no GPU required. Containers sized at 2 vCPU / 4GB RAM per replica.

Software Dependencies

`numpy` , `scikit-learn` , and the court case-management system API client. The full dependency manifest is maintained in the repository's requirements file.

Performance Specifications

- Sentencing-outcome prediction: < 500ms per request
- Evidence-reliability scoring: < 300ms per request
- Judicial decision support (precedent retrieval): < 2s per query against the legal database

Integration Points

REST API consumed by the court case-management system. There is no direct end-user-facing interface — all outputs are surfaced within the existing case-file UI used by judges and clerks.

Accuracy & Robustness

This section addresses the requirements of Article 15 of Regulation (EU) 2024/1689 (EU AI Act) regarding accuracy, robustness, and cybersecurity.

Accuracy Metrics

This subsection auto-populates from your connected repository's scan results once a compliance scan has completed — nothing to hand-fill here. If it's still showing the placeholder note after a scan has run, check that the scan is linked to this AI system record.

Robustness Measures

No adversarial-robustness or out-of-distribution testing has been conducted for the sentencing-outcome or evidence-reliability models. Input validation is currently limited to basic type and schema checks. Adversarial-example resilience testing is planned for the next model iteration (see Risk Management System, Section 10).

Cybersecurity Measures

Access to the inference API is restricted to the court case-management system via network-level controls, with no public internet exposure. Model artefacts, training data, and logs are IAM-restricted to compliance officers and authorised court IT staff, with access audited via CloudTrail. Model artefact signing and dedicated security-monitoring/alerting are planned but not yet implemented. Full detail in the Cybersecurity Measures (Annex IV §2(h)) section.

Known Limitations

- Sentencing-outcome and evidence-reliability predictions are advisory only and must not be used as the sole basis for any judicial decision.
- No adversarial-robustness or out-of-distribution testing has been performed; behaviour on unusual or edge-case inputs is not validated.
- Demographic bias assessment has not yet been completed; per-group accuracy is currently unknown (see Data and Data Governance).
- The judicial decision support (precedent-retrieval) module is a reference implementation and does not yet return results from a production legal database.
- Evidence-reliability scoring accuracy has not been formally validated against a held-out test set.

Human Oversight Measures

This section addresses Article 14 of Regulation (EU) 2024/1689 (EU AI Act) — human oversight measures.

Oversight Mechanisms

All CaseInsight outputs are presented as advisory recommendations within the case file. No output is automatically applied to a case record or communicated externally without judge review.

User Interface for Oversight

Planned enhancement — not yet implemented. The case-file UI will display a mandatory acknowledgement step before a prediction can be marked as reviewed, alongside the model's confidence score and key contributing factors.

Ability to Override or Reverse

Judges may override any AI-generated recommendation via the system's override function. The override is logged, though structured override-reason capture is not yet mandatory — tracked as a remediation item ahead of the next release.

Training Requirements for Human Overseers

Judges and clerks assigned oversight responsibility complete a mandatory onboarding module covering the system's capabilities, known limitations, and automation-bias risks before first use.

Automation Bias Awareness (Art. 14(3)(b))

Article 14(3)(b), Regulation (EU) 2024/1689

Natural persons responsible for overseeing this AI system are trained and enabled to remain aware of the possible tendency to automatically rely on or over-rely on the output produced by the AI system (automation bias).

Measures implemented to counter automation bias: Onboarding training addresses automation bias directly. **Gap:** automated detection of uncritical AI-recommendation acceptance (e.g. consistently low override rates) is not yet implemented; this is planned via the override-rate monitoring described in the Post-Market Monitoring Plan.

Designated Oversight Persons (Art. 14(5))

Article 14(5), Regulation (EU) 2024/1689

The deployer of this AI system is required to designate natural persons with the authority and competence to carry out human oversight. Requirements for this designation are communicated in the Instructions for Use.

Competence and authority requirements for designated persons: Oversight must be exercised by a sitting judge (or, for evidence-reliability scoring only, a senior clerk under judicial supervision) who has completed onboarding training and holds full authority to disregard or modify any AI-generated output, without any justification requirement beyond standard case-record documentation.

Post-Market Monitoring Plan

This section describes the post-market monitoring plan as required by Article 72 of Regulation (EU) 2024/1689 (EU AI Act). It summarises the dedicated Post-Market Monitoring Plan dossier, which contains full detail.

Monitoring Strategy

Example AI Systems Ltd. continuously collects deployment data (predictions, confidence scores, judge overrides) and conducts quarterly performance and bias-drift reviews.

Incident Reporting Procedures

Suspected serious incidents (Art. 3(49)) are reported via Example Regional Court Administration's internal compliance channel to Example AI Systems Ltd.'s compliance officer, with formal notification to the market surveillance authority within 15 days where the Art. 73 threshold is met.

Performance Metrics to Track

Overall and per-demographic prediction accuracy, evidence-reliability score calibration drift, judge override rate, and data drift on key input features. Thresholds are defined in the Post-Market Monitoring Plan.

Review Schedule

Quarterly performance and bias-drift review; annual full re-assessment aligned with AIRA/FRIA renewal. See the Post-Market Monitoring Plan dossier for the complete schedule.

Annex IV Compliance Checklist

The following checklist verifies coverage of Annex IV requirements under Regulation (EU) 2024/1689 (EU AI Act).

#	Requirement	Status
1	General description of the AI system	Documented — see Intended Purpose
2	Detailed description of elements and development process	Documented — see Technical Specifications
3	Information about monitoring, functioning and control	Documented — see Post-Market Monitoring Plan (summary) and the dedicated PMM Plan dossier
4	Description of the risk management system	Partially documented — see Risk Management System; residual risk acceptance (§9) and testing against predefined thresholds (§10) are still open
5	Description of data governance measures	Partially documented — see Data and Data Governance; demographic bias assessment is an open item
6	Description of accuracy, robustness, cybersecurity	Partially documented — see Accuracy & Robustness and Cybersecurity Measures; adversarial-robustness testing and model-artefact signing are open items
7	Description of human oversight measures	Partially documented — see Human Oversight Measures; mandatory oversight UI gate and automation-bias detection are planned, not yet implemented
8	Description of any change made to the system	Not currently tracked — no <code>change_log</code> section is seeded for this dossier type; this requirement has no home in the current Annex IV structure and needs either a manual addendum or a product fix to wire in the existing (unused) <code>change_log</code> / <code>predetermined_changes</code> templates
9	Validation and testing procedures	Partially documented — covered only within Risk Management System §10/ §12; no dedicated validation/testing section exists in this dossier type

Checklist auto-populated from scan results (1 repositories scanned, compliance score: 26.0%).

Cybersecurity Measures (Annex IV §2(h))

Cybersecurity Measures (Article 15)

This section documents the cybersecurity measures implemented to protect the AI system against unauthorised access, data poisoning, model manipulation, and adversarial attacks in accordance with Article 15 of the EU AI Act.

Threat Model Summary

Primary threats considered: (1) unauthorised access to case data or model artefacts via compromised credentials, (2) training-data poisoning via a compromised case-management export, (3) model extraction or inference attacks against the prediction API, and (4) adversarial manipulation of input case features to skew sentencing-outcome or evidence-reliability predictions. Threat modelling was performed at design time using a lightweight STRIDE-based review; formal red-teaming has not yet been conducted — tracked as an open item.

Attack Surface Analysis

The system's external attack surface is limited to the REST API endpoint consumed by the court case-management system, accessible only from the court-network-internal environment with no public internet exposure (see Technical Specifications). The internal attack surface includes the training pipeline (data ingestion, preprocessing) and the model artefact store.

Protection Against Unauthorised Access

Access to the inference API is restricted to the court case-management system via network-level controls (VPC / security groups) and service-to-service authentication. Access to model artefacts, training data, and logs is restricted via IAM roles limited to compliance officers and authorised court IT staff (see Record-Keeping, Logging and Retention), with all access itself audited via CloudTrail.

Data Integrity and Poisoning Prevention

Training-data ingestion is a manual, reviewed batch export from the court case-management system (see Data and Data Governance). There is currently no automated integrity verification (e.g. checksums, anomaly detection) on ingested training data — tracked as a remediation item.

Model Security

Model artefacts are versioned and stored with training-data lineage (see Technical Specifications); artefact storage access is IAM-restricted. Model artefacts are not currently cryptographically signed — signing and integrity verification prior to deployment is planned but not yet implemented.

Adversarial Attack Resilience

No adversarial-robustness testing (e.g. adversarial example generation, input perturbation testing) has been conducted for the sentencing-outcome or evidence-reliability models. This is an open item tracked for the next model iteration, consistent with the accuracy-validation gap noted in the Risk Management System section (Section 10, Testing Against Predefined Metrics and Thresholds).

Logging and Monitoring

Inference requests, predictions, and access to model artefacts and logs are recorded per the automatic logging controls described in Record-Keeping, Logging and Retention. Security-relevant events (failed authentication, unusual access patterns) are captured via CloudTrail but do not yet feed a dedicated security-monitoring or alerting pipeline — tracked as a planned enhancement.

Incident Response Plan

Cybersecurity incidents are handled through the same escalation path as compliance incidents described in the Post-Market Monitoring Plan (Incident Reporting Procedures). A dedicated cybersecurity incident-response runbook, distinct from the general compliance incident process, has not yet been drafted.

Redundancy and Fallback Mechanisms

The inference service runs as a multi-replica container deployment (see Technical Specifications), providing basic availability redundancy. There is no formally documented fallback procedure for judges if the system becomes unavailable; in practice this defaults to standard manual case-preparation procedures, but that has not yet been written down as an official contingency plan.

Standards Applied (Annex IV §2(g))

Applied Harmonised Standards (Annex IV §2(g), Art. 40)

Article 40, Regulation (EU) 2024/1689

Standard	Title	Scope of Application
ISO/IEC 42001	Artificial intelligence — Management system	Organisational AI governance framework
ISO/IEC 23894	Artificial intelligence — Guidance on risk management	Risk management system design (Art. 9)

No harmonised standard under Art. 40 has yet been published specifically for Annex III high-risk systems at the time of drafting; the above are applied as best-practice reference frameworks pending publication.

Common Specifications Applied (Art. 41)

Article 41, Regulation (EU) 2024/1689

In the absence of published harmonised standards for this system category, the NIST AI Risk Management Framework was used as supplementary guidance for risk-identification and bias-evaluation methodology.

Presumption of Conformity (Art. 40(3))

The following standards create presumption of conformity with the requirements of Regulation (EU) 2024/1689 (EU AI Act) for the covered scope:

The standards listed above are informative references only; as they are not yet formally harmonised under the AI Act mandate, they do not currently establish presumption of conformity. This section will be updated once CEN/CENELEC JTC 21 publishes applicable harmonised standards for Art. 9, 10, 13, 14, and 15 requirements.

Deviations from Standards

None — no harmonised standards are yet in force for this system category.

Record-Keeping, Logging and Retention

This section describes the record-keeping, logging and retention obligations for CaseInsight Judicial Support Suite as required by Articles 12, 18, 19 and 73 of Regulation (EU) 2024/1689 (EU AI Act).

Automatic Logging (Art. 12(1))

Article 12, Regulation (EU) 2024/1689 This AI system automatically records events relevant to compliance monitoring.

Log Format and Content

Each prediction and evidence-reliability score is logged with: timestamp, input feature hash, model version, output value, confidence score, and reviewing-judge identifier (once the mandatory oversight logging described in the Human Oversight Measures section is implemented).

Log Integrity Controls (Art. 12(2))

Logs are written to append-only S3 storage with object-lock enabled; access is restricted to compliance officers and authorised court IT staff via IAM role, with all access itself logged via CloudTrail.

Provider Log Retention (Art. 19(1))

Provider logs are retained for a minimum of **6 months** from generation, unless Union or national law requires a longer retention period.

Log Type	Retention Period	Storage Location	Access Controls
Prediction and evidence-scoring logs	≥ 6 months (retained 10 years to align with court archive policy)	AWS S3 (eu-west-1), append-only	IAM role-restricted, dual authorisation required for export

Technical Documentation Retention (Art. 18)

Technical documentation for this AI system is kept available for national competent authorities for **10 years** after market placement or putting into service.

Documentation retained: - [x] Technical documentation package (Annex IV) - [x] EU Declaration of Conformity (Art. 47) - [x] Quality management system records (Art. 17) - [x] Post-market monitoring reports (Art. 72)

Incident Record Retention (Art. 73)

Records of any serious incidents reported under Art. 73 are retained for **10 years**.

The incident register is maintained by Example AI Systems Ltd.'s compliance function in a dedicated case-tracking system, accessible to compliance officers and, on request, to the relevant market surveillance authority.

Deployer Log Retention (Art. 19(2))

Deployers are required to keep logs under their control for an appropriate period (minimum 6 months). This requirement is communicated to deployers in the Instructions for Use.

Disclaimer

This document was automatically generated by the Scanara platform. It is provided as a compliance aid under Regulation (EU) 2024/1689 (EU AI Act) and does not constitute legal advice. Organisations are responsible for ensuring accuracy and completeness.