



5 June 2025

CEPEJ(2023)16rev

**EUROPEAN COMMISSION FOR THE EFFICIENCY OF JUSTICE
(CEPEJ)**

**Assessment Tool for the Operationalisation of the European
Ethical Charter on the Use of Artificial Intelligence in Judicial
Systems and Their Environment**

Document adopted by the CEPEJ at its 41st plenary meeting, 4-5 December 2023
and revised at its 44th plenary meeting, 4-5 June 2025

Introduction

In 2018, the European Commission for the Efficiency of Justice (hereafter “CEPEJ”) adopted its European Ethical Charter on the use of artificial intelligence in judicial systems and their environment (hereafter “the CEPEJ Charter”).

The CEPEJ Charter lays out five key principles that should be respected in the design and use of artificial intelligence (AI) by the judiciary:

- 1) Principle of respect for fundamental rights: ensuring that the design and implementation of AI tools and services are compatible with fundamental rights;
- 2) Principle of non-discrimination: specifically preventing the development or intensification of any discrimination between individuals or groups of individuals;
- 3) Principle of quality and security: with regard to the processing of judicial decisions and data, using certified sources and intangible data with models elaborated in a multi-disciplinary manner, in a secure technological environment;
- 4) Principle of transparency, impartiality, and fairness: making data processing methods accessible and understandable, authorising external audits;
- 5) ‘Under user control’ principle: precluding a prescriptive approach and ensuring that users are informed actors and in control of the choices made.

The CEPEJ Charter represents the first step in the CEPEJ’s action to promote responsible use of AI in European judicial systems, in accordance with the Council of Europe’s values.

AI systems intended for the administration of justice have potentially significant impact on democracy, rule of law, individual freedoms as well as the right to an effective remedy and to a fair trial, and could lead to potential biases, errors, and opacity.

While AI systems that do not affect the actual administration of justice can in certain cases be considered of low risk, they still require proper assessments to ensure their use do not lead to unexpected or unwanted harms. When AI systems are intended to be used by a judicial authority or on their behalf to assist judicial authorities in researching and interpreting facts and the law and in applying the law to a concrete set of facts, they should not affect the independence of judges in their decision-making process. The final decision making should remain a human-driven activity and decision. In all such cases, the ethical compliance of these systems with the CEPEJ Charter’s principles should be assessed.

In line with the above, discussions within the CEPEJ and with external partners demonstrated that decision makers within judicial systems (i.e., representatives of judicial (management) bodies, project managers, IT managers, etc.) would benefit from more practical guidance on how to apply the five principles laid down in the CEPEJ Charter, therefore a detailed operationalisation of the principles was deemed necessary.¹

As a result, the present *Assessment Tool for Operationalisation of the European Ethical Charter on the Use of Artificial Intelligence in Judicial Systems and Their Environment* (hereafter “the

¹ In 2022, the Working Group on Quality of Justice (CEPEJ-GT-QUAL) took a decision to move forward with designing an appropriate support tool to put the CEPEJ Charter into practice, in consultation with the working group on Cyberjustice and Artificial Intelligence (CEPEJ-GT-CYBERJUST) and within the mandate of the Artificial Intelligence Advisory Board (AIAB) recently set up to support the CEPEJ on the technical aspects of AI. The work of developing the tool has been entrusted to Alexandra TSVETKOVA and Matthieu QUINIOU, members of the AIAB.

Assessment Tool”) aims to operationalize the CEPEJ Charter by providing for a set of verifications, key measures and safeguards that decision-makers within judicial systems should follow when purchasing, designing, developing, implementing and/or using AI in judicial systems and their environment, in compliance with the CEPEJ Charter. The Assessment Tool also aligns its logic with that of compliance and risk-based regulation. Further, it seeks to complement the Human Rights, Democracy, and the Rule of Law Impact Assessment (HUDERIA)², that is currently being designed by the Council of Europe Committee on Artificial Intelligence (CAI), by adding a practical layer of measures towards ethical compliance applied to the judiciary. This approach should ensure uniformity towards identifying, analysing, and evaluating the significant levels of risks and the assessment of impact of AI systems in relation to the enjoyment of human rights, the functioning of democracy and the observance of rule of law.

General Presentation of the CEPEJ Charter Assessment Tool

Who is this Assessment Tool intended for and who carries out the assessment?

This Assessment Tool is essentially intended for two types of decision-makers within judicial systems:

- representatives of judicial (management) bodies in charge of acquiring software licenses including for AI systems;
- IT managers, project managers and/or representatives in management positions within judicial (management) bodies supervising software development including development of AI systems.

The Assessment Tool is designed to be filled in by the decision makers listed above.

In the case of a software developed by third parties (such as in the case of license acquisition), it is possible for the decision makers to inquire with the legal representatives of the software provider to complete this Assessment Tool in the context of, for example, a response to a call for tender. This document, as completed by the provider, can serve either as a basis for (further) analysis of the decision maker or as an accountability requirement for the provider.

When can this Assessment Tool be used?

This Assessment Tool is designed to be used before the acquisition or the development of an AI system by a judicial institution. It can also be used in case of updating the AI system once deployed. Furthermore, this Assessment Tool can serve as a point of reference for the ideation and determination of specifications for the design of an AI system in the judicial field.

Is this Assessment Tool following a risk-based approach or a principle-based approach (in particular, with regards to human rights)?

This Assessment Tool follows a mixed approach. It is structurally based on the five principles defined in the CEPEJ Charter. The operationalization of the principles implies taking into account the technical and usage context as well as the identifiable risks in this regard. New risks are likely to be identified with the evolution of AI techniques and uses, which will require updating this Assessment Tool after its piloting³.

² The HUDERIA consists essentially of an obligation to address a certain number of questions regarding the contexts of design, development, procurement, and use and the potential short-, medium, and long-term impacts of the AI system under examination.

³ See point 2 of the Revised roadmap for ensuring an appropriate follow-up of the CEPEJ Ethical Charter on the use of artificial intelligence in judicial systems and their environment, [CEPEJ\(2021\)16](#)

Which risks have been taken into account in the Assessment Tool?

The risks listed herein correspond to the main risks envisaged in the context of the five principles of the CEPEJ Charter and related reports and academic literature.

1) Risk of reusable data and/or use of an AI model trained for another purpose

This risk refers to the 1st principle of the CEPEJ Charter, namely 'Principle of respect for fundamental rights', with regard to proportionate processing of personal data (privacy) and clear purposes. It is caused by a failure to take into account the purpose of data processing, and by the unsuitability of the trained AI model for other purposes or in a different context (notably in another legal system). This risk can also be seen in the influence of the first training context on the reuse context, which can imply, for example, the influence of one legal culture on another through the AI system.

2) Risk of personal data or trade secret disclosure

This risk refers to the 1st principle of the CEPEJ Charter, namely 'Principle of respect for fundamental rights', with regard to proportionate processing of personal data (privacy) and clear purposes. The risk of disclosing personal data or trade secrets can be caused by backdoors, but also by reinforcing an AI model with personal data or trade secrets shared by the user. This risk is amplified by the absence of precautions to inform the user of the potential risks of sharing data with the AI system.

3) Risk of judge profiling and forum shopping

This risk refers to the 1st principle of the CEPEJ Charter, namely 'Principle of respect for fundamental rights', with regard to the judges' independence in their decision-making process. The risk of judge profiling can arise from generalist AI or AI specifically dedicated to the judicial field. This profiling risk essentially depends on the data used to train the AI (not anonymized court decisions, activities and posts on social networks, etc.) and the models trained specifically for profiling purposes. The profiling of judges can lead to the recusal of a judge, the search for a jurisdiction statistically favorable to the judicial request, or the adaptation of argumentation to the judge's profile. Profiling risks also exist to a certain extent for jurors and parties.

4) Risk of misleading AI results

This risk refers to the 1st principle of the CEPEJ Charter, namely 'Principle of respect for fundamental rights', with regard to the judges' independence in their decision-making process. The risk of misleading results can be caused by AI falsely presenting results as certain (or with high percentages of reliability), and also the "hallucinations" of generative AI. The consequences of this risk can be significant for AI systems used for decision-making or automated online dispute resolution.

5) Risk of unclear criteria and inadequate weight for criteria for AI processing

This risk refers to the 1st principle of the CEPEJ Charter, namely 'Principle of respect for fundamental rights', with regard to the judges' independence in their decision-making process. This risk is linked to an inadequate methodology followed during AI training, and to the introduction of biases that are difficult to detect in the AI model. This risk is amplified in case of unsupervised AI. This risk can result, for example, in errors or amplified discrimination.

6) Risk of AI replacing the access to the judge

This risk refers to the 1st principle of the CEPEJ Charter, namely 'Principle of respect for fundamental rights', with regard to the right of access to the judge and the right to a fair trial. This

risk may be the result of an institutional choice, but it may also be caused by an interface that dissuades the litigant from having recourse to a judge. This risk is consubstantial with AIs in the field of automated online dispute resolution, and also exists to a certain extent for decision support AIs, if judges are encouraged not to question the AI's result.

7) Risk of unclear/unjustified grounds for the judgment

This risk refers to the 1st principle of the CEPEJ Charter, namely 'Principle of respect for fundamental rights', with regard to the right of access to the judge and the right to a fair trial. The causes of this risk are essentially linked to the AI's training method, which is unsuited to its intended use, and also to the non-explicable nature of the AI system.

8) Risk of unfair advantage for one party to the trial

This risk refers to the 1st principle of the CEPEJ Charter, namely 'Principle of respect for fundamental rights', with regard to the right of access to the judge and the right to a fair trial. This risk may be caused by the license cost and by the skills required to access and/or use an AI system that gives an advantage in a trial (e.g. a real-time argumentation support system).

9) Risk of infringement of fundamental rights or inappropriate arbitration between two fundamental rights

This risk refers to the 1st principle of the CEPEJ Charter, namely 'Principle of respect for fundamental rights', in light of the expected ethics- and human-rights-by-design guarantees. This risk is a cross-cutting one, and is likely to be particularly significant for automated online dispute resolution AIs, which are called upon to arbitrate between two fundamental rights.

10) Risk of discrimination or amplification of discrimination

This risk refers to the 2nd principle of the CEPEJ Charter, namely 'Principle of non-discrimination', with regard to avoiding bias and discrimination based on sensitive data. The risk of discrimination is mainly caused by biases in the AI system. These biases can be technical or cognitive, and can derive from data selection, annotation or AI training.

11) Risk of generation and use of inexistent legal provisions by generative AI

This risk refers to the 3rd principle of the CEPEJ Charter, namely 'Principle of quality and security'. This risk is specific to generative AIs, and is consubstantial with their functioning, even if it can be limited by using official sources to reinforce the model. This phenomenon is also described as AI "hallucination".

12) Risk of disempowerment and limitation of accountability of the judge through the use of non-explainable AI

This risk refers to the 4th principle of the CEPEJ Charter, namely 'Principle of transparency, impartiality, and fairness', more specifically of the need for an AI that is fair, accountable and transparent. This risk is an institutional risk, which can be amplified by the use of AI blackbox that cannot be explained, particularly if they are protected by trade secrets.

13) Risk of misuse of AI

This risk refers to the 5th principle of the CEPEJ Charter, namely 'Under user control' principle'. This risk arises from the use of AI for a purpose other than that envisaged at the training stage. This risk is amplified by the absence of any explanation to the user of the possible uses of the AI system.

14) Risk of forced use of AI

This risk refers to the 5th principle of the CEPEJ Charter, namely 'Under user control' principle'. This risk of forced use of AI is due to the way in which the AI system is integrated into the judicial process, and to an opt-out system that is either non-existent or difficult to access.

Can this tool be used in addition to impact analysis and other assessment tools?

This Assessment Tool addresses the challenges of designing and using AI in the judicial field for decision makers in judicial institutions. It is by nature specific and its use can be supplemented by others applicable to AI deployment and/or tools dedicated to the challenges related to transparency, security and personal data disclosure.

Complementary ethical and human right assessment tool that is not specific to the judicial field refer to HUDERIA,⁴ developed by the Committee on Artificial Intelligence of the Council of Europe (still in development at the date of adoption of the present tool).

Which areas of application are concerned?

The areas of application concerned by the Assessment Tool are those listed in the Resource Centre on Cyberjustice and AI of the CEPEJ⁵:

- Document search, review and large-scale discovery: *these solutions create a searchable collection of case-law descriptions, legal text and other insights to be shared with legal experts for further analysis and large-scale discovery on high volumes of electronic documents*
- Automated online dispute resolution: *these solutions cover technologies used for the resolution of disputes between parties with limited human intervention, which can be achieved through hardware and/or software*
- Prediction of litigation outcomes: *these solutions learn from large datasets to identify patterns in the data that are consequently used to visualize, simulate or predict new litigation outcomes*
- Decision support and decision-making: *these solutions facilitate or fully automate decision making processes in the justice systems*
- Anonymization: *these solutions are used for removing identifying information such as personal data of court users*
- E-filing: *these technological solutions facilitate access to justice by establishing a digital channel that enables the interaction and exchange of data and e-documents between courts and court users*
- Triaging, allocation and workflow automation: *these solutions are used to facilitate or complete some tasks and activities during the lifecycle of the proceedings within the Case Management System, minimising the need for human input. Examples are: registration and allocation of court matters, assigning levels of priority to tasks or individuals to determine the most effective order in which to deal with them*
- Natural language processing: *these solutions are capable of recognizing and analysing speech, written text and communicating back. Their main use for courts is in voice/speech recognition and (hearing) transcription of court proceedings*
- Information/assistance services: *these solutions provide individuals with information on services available in the justice systems and link individuals to the services and opportunities that are available.*

⁴ <https://rm.coe.int/huderaf-coe-final-1-2752-6741-5300-v-1/1680a3f688>

⁵ <https://www.coe.int/en/web/cepej/cepej-resource-center-on-cyberjustice-and-artificial-intelligence>. Please note that the areas of application are indicative and an overlap is possible.

Where a question listed in the Assessment Tool is only applicable to a number of these application areas, there is an explicit reference to the specific application areas concerned.

How to use this Tool?

The CEPEJ Charter Assessment Tool consists of 29 questions designed to raise awareness and identify the potential infringements of an AI system on the principles of the CEPEJ Charter, and in particular on human rights and rule of law. Subsidiary questions are also provided to refine answers, or to check or suggest the implementation of risk mitigation or management measures. Furthermore, in the case of answers that could seriously infringe the principles protected by the CEPEJ Charter, the 'you should reconsider deploying this AI system' mentioning is displayed to inform the person completing the form of the importance of that specific point.

Once the Assessment Tool is filled in by the competent person, it is recommended to discuss it within a supervisory group to ensure appropriate follow-up.

CEPEJ Charter Assessment Tool for AI systems in the Judiciary

AI System's Description and Context

1. Who developed the AI system?

- ☐ Developed in-house by the judicial institution
- ☐ Developed by a service provider based on the specifications of the judicial institution
- ☐ Developed by a third-party providing user license to the judicial institution
- ☐ Other _____

2. Which area of application are concerned by the AI system?

Multiple answers are possible.

- ☐ Document search, review and large-scale discovery
- ☐ Automated online dispute resolution
- ☐ Prediction of litigation outcomes
- ☐ Decision support and decision-making
- ☐ Anonymization
- ☐ E-filing
- ☐ Triaging, allocation and workflow automation
- ☐ Natural language processing
- ☐ Information/assistance services

3. Is the problem that the AI system is to solve (or the benefit(s) it produces) in the judiciary clearly defined?

☐ Yes ☐ No

If your answer is 'no', you should reconsider deploying this AI system.

If your answer is 'yes', please describe the problem the AI system is to solve in the space below.

4. Is the context⁶ in which the AI system is to be deployed and used clearly defined?

☐ Yes ☐ No

⁶ AI risks – and benefits – can emerge from the interplay of technical aspects combined with societal factors related to how a system is used, its interactions with other AI systems, who operates it, and the professional and/or social context in which it is deployed. In judiciary, the latter could apply to geographical coverage, what type of cases and/or issues to be solved are covered by the system purpose and functionality, in what judicial procedure, etc.

*If your answer is 'no', explain why and you should reconsider deploying this AI system.
If your answer is 'yes', please describe the usage context in the space below.*

Principle of Respect for Fundamental Rights

The principle of respect for fundamental rights refers to ensuring that the design and implementation of AI tools and services are compatible with fundamental rights.

5. Is the AI system likely to enable or facilitate the profiling of judges and/or courts?

☐ Yes ☐ No

5.1. If yes, which risk mitigation or management measures are envisaged?

Multiple answers are possible.

- ☐ Anonymization of data related to a judge
- ☐ Anonymization of data related to a court and its members
- ☐ Exclusion of search criteria facilitating such profiling
- ☐ Other _____

6. Is the AI system likely to enable or facilitate the profiling of jurors?

☐ Yes ☐ No

6.1. If yes, is the AI system available to all the parties in the lawsuit?

☐ Yes ☐ No

6.2. If yes, which risk mitigation or management measures are envisaged:

Multiple answers are possible.

- ☐ Anonymization of data related to a juror
- ☐ Exclusion of search criteria facilitating such profiling
- ☐ Other _____

7. Is the AI system likely to enable or facilitate the profiling of parties in a lawsuit?

☐ Yes ☐ No

7.1. If yes, which risk mitigation or management measures are envisaged:

Multiple answers are possible.

- ☐ Exclusion of data external to the case
- ☐ Exclusion of non-official data
- ☐ Other _____

8. (Applicable only to decision support and decision-making, automated online dispute resolution, e-filing and information/assistance services areas of application) **Is there a human alternative to the AI system available?**

☐ Yes ☐ No

- 8.1. **If yes, does the user have to opt-out to access the human alternative to the AI system?**

☐ Yes ☐ No

- 8.1.1. **If yes, are there measures in place to facilitate the opt-out?**

☐ Yes ☐ No

If your answer is 'yes', please describe the measures.

9. (Applicable only to automated online dispute resolution, prediction of litigation outcomes, and decision support and decision-making areas of application) **Is the AI system capable of providing systematically justified legal grounds?**

☐ Yes ☐ No

- 9.1. **If no, is the AI system conceived to help judges make decisions?**

☐ Yes ☐ No

If your answer is 'yes', you should check if it is clearly explained to the judges that the AI system is not capable of providing justified legal grounds. A particular attention must be paid to compliance with the 'Under user control' Principle.

- 9.2. **If no, is the AI system conceived to make decisions?**

☐ Yes ☐ No

If your answer is 'yes', you should reconsider deploying this AI system. If your answer is 'no', the purpose of the AI must be clearly explained to the user and a particular effort must be made to guarantee the 'Under user control' Principle.

10. (Applicable only to document search, review and large-scale discovery, automated online dispute resolution, prediction of litigation outcomes, decision support and decision-making, e-filing, and natural language processing areas of application) **Is the**

AI system likely to provide an advantage (e.g., data processing during witness interviews) to its user(s) in legal proceedings?

☐ Yes ☐ No

10.1. Could the financial conditions for accessing the AI system potentially hinder certain users from utilizing it?

☐ Yes ☐ No

If your answer is 'yes', please describe if specific measures have been taken to consider context and resources (e.g. discounts offered to public judicial institution, association, court-appointed attorney, etc.).

11. Is the AI system likely to collect and/or process sensitive personal data and/or trade secrets?

☐ Yes ☐ No

11.1. If yes, which risk mitigation or management measures are envisaged to protect these data or discourage users from sharing it?

Multiple answers are possible.

☐ Data encryption

☐ Pop-up warning for the user not to share sensitive data

☐ Other _____

12. (Applicable only to automated online dispute resolution, prediction of litigation outcomes, and decision support and decision-making areas of application) Is the AI system used in the context of decision(s) related to deprivation of liberty?

☐ Yes ☐ No

12.1. If yes, does the AI system encourage release and alternative measures rather than incarceration?

☐ Yes ☐ No

If your answer is 'yes', please describe the calibration and assessment measures implemented.

Principle of Non-discrimination

The principle of non-discrimination refers to specifically preventing the development or intensification of any discrimination between individuals or groups of individuals.

13. Have the data, model and/or results provided by the AI system been audited to identify biases?

☐ Yes ☐ No

If your answer is 'no', you should consider a data-, model- and results- dedicated audit before deploying the AI system.

13.1. If yes, who conducted the audit?

Multiple answers are possible.

- ☐ Judicial institution's IT Department
- ☐ Judicial institution's multidisciplinary team (IT, legal professionals, etc.)
- ☐ Third-party offering the AI system
- ☐ Independent external audit
- ☐ Other _____

13.2. If yes, what kind of audit has been conducted?

Multiple answers are possible.

- ☐ Source code audit
- ☐ Model explanation audit
- ☐ Data audit
- ☐ A/B testing of the results (by slightly changing prompt or parameter)
- ☐ UX/UI audit
- ☐ Other _____

13.3. If yes, have any biases been identified?

☐ Yes ☐ No

If your answer is 'yes', please describe the biases that have been identified.

—

—

—

—

13.4. If yes, are identified bias potentially leading to a form of discrimination?

☐ Yes ☐ No

13.4.1. Please specify which types of discrimination are concerned:

Multiple answers are possible.

- ☐ Gender identity
- ☐ Skin colour
- ☐ Religion
- ☐ Political opinions
- ☐ Ethnic/national Origin
- ☐ Age
- ☐ Disability
- ☐ Language
- ☐ Citizenship
- ☐ Sex orientation
- ☐ Other _____

Please detail the potential discriminatory effects.

13.4.2. If yes, has the type or cause of bias been identified?

Multiple answers are possible.

- ☐ Source code, algorithm or model bias
- ☐ Sampling and non-representative training data
- ☐ Data and token labelling bias
- ☐ Generalization or association bias
- ☐ Unstandardized measurement bias
- ☐ User experience (UX) bias
- ☐ User interface (UI) bias
- ☐ Other _____

<i>If the cause of bias cannot be identified, you should reconsider deploying this AI system.</i>

13.4.3. If yes, has any mitigation or management of risks of discrimination measures been implemented?

☐ Yes ☐ No

*If your answer is 'yes', please describe mitigation or management measures implemented in the space below.
If your answer is 'no', please explain why. In any cases, you should reconsider deploying this AI system if no further measures are undertaken.*

Principle of Quality and Security

The principle of quality and security relates to the processing of judicial decisions and data, and more specifically with regards to using certified sources and intangible data with models conceived in a multi-disciplinary manner in a secure technological environment.

14. Is the AI system based on a pre-trained model by third parties?

☐ Yes ☐ No

14.1. If yes, have the databases used by these models been audited?

☐ Yes ☐ No

If your answer is 'no', you should consider an audit of the databases used in these models before deploying the AI system.

15. Is the AI model trained only with official or certified data?

☐ Yes ☐ No

If your answer is 'no', please specify the sources of non-official or not certified data used to train the AI model (including user or user-generated data in case of reinforcement of the AI model) and explain why these data are important.

If the usage of non-official or not certified data cannot be justified, you should reconsider the use of such data. If the latter cannot be avoided, you should reconsider deploying this AI system.

16. **Are there monitoring mechanisms, procedures and protocols in place to secure that data based on judicial procedures is not modified prior to the feeding of the AI system model's learning mechanism?**

☐ Yes ☐ No

If your answer is 'yes', please describe the monitoring mechanisms, procedures and protocols.

If your answer is 'no', you should reconsider deploying this AI system.

17. **Are there monitoring mechanisms, procedures and protocols in place to guarantee that datasets and models are being stored and executed in secure environments?**

☐ Yes ☐ No

If your answer is 'yes', please describe the monitoring mechanisms, procedures and protocols.

If your answer is 'no', you should reconsider deploying this AI system.

18. **Has the AI training (data selection, data annotation, reinforcement, etc.) been carried out by a multidisciplinary team?**

☐ Yes ☐ No

19. **Have data and token annotations (e.g. keywords describing a court decision, etc.) for AI training been carried out with legal professionals and ethics experts?**

☐ Yes ☐ No

If your answer is 'no', you should consider an audit of the training sets used in these models before deploying the AI system.

20. **Has a Human Rights Protection Officer been appointed within the judicial institution to ensure the fundamental rights compliance of the AI system in its deployment and use?**

☐ Yes ☐ No

If your answer is 'no', you should consider appointing an expert with the proper expertise to supervise the use of AI system within the judicial institution.

Principle of Transparency, Impartiality, and Fairness

The principle of transparency, impartiality and fairness refers to making data processing methods accessible and understandable and being able to authorize external audits.

- 21. Is the source code auditable (technically auditable, no secrets limiting auditability, etc.)?**

☐ Yes ☐ No

If your answer is 'no', you should reconsider deploying this AI system if it is likely to infringe the principle of fair trial and/or the principle of equality of arms.

- 22. Are the data used to train the AI model auditable?**

☐ Yes ☐ No

If your answer is 'no', you should reconsider deploying this AI system if it is likely to infringe the principle of fair trial and/or the principle of equality of arms.

- 23. Are the AI model, criteria and weightings clear and understandable?**

☐ Yes ☐ No

If your answer is 'no', relevant measures should be undertaken to explain and clarify the AI model and its results.

- 24. Is the information kept for audit, oversight and review purposes in plain, understandable and coherent language?**

☐ Yes ☐ No

If your answer is 'no', relevant measures should be undertaken to secure such documentation.

'Under User Control' Principle

Principle "under user control": preclude a prescriptive approach and ensure that users are informed actors and in control of their choices.

- 25. Is there a clear indication provided to the user that an AI system is being used?**

☐ Yes ☐ No

If your answer is 'no', such indication should be provided prior to deploying the AI system.

- 26. Is the AI system likely to provide erroneous or misleading results?**

☐ Yes ☐ No

- 26.1. If yes, which risk mitigation or management measures are envisaged?**

☐ Recommendations to the users on the use of AI

☐ Explanation of the results provided by AI

☐ Other _____

27. Is there easily accessible documentation to capacitate users?

☐ Yes ☐ No

If your answer is 'no', such documentation should be made available to users prior to and/or during deploying the AI system.

27.1. If yes, which kind of documentation:

Multiple answers are possible

- ☐ FAQ
- ☐ Video tutorial
- ☐ Interactive guide
- ☐ Chatbot
- ☐ Detailed user guide
- ☐ Handling test
- ☐ Other _____

28. Are there easily accessible training and/or awareness materials on the AI system to capacitate users?

☐ Yes ☐ No

If your answer is 'no', such training and/or awareness materials should be made available to users prior to and/or during deploying the AI system.

28.1. If yes, which topics are covered?

Multiple answers are possible.

- ☐ Getting started with the AI system and presentation of possible uses
- ☐ Technical aspects of AI
- ☐ Risks of bias and discrimination risks in AI systems
- ☐ Explicability and AI audit
- ☐ Ethical aspects of AI
- ☐ Human rights aspects of AI
- ☐ AI use case and challenges in the judiciary
- ☒ Compliance requirements with the five principles of the CEPEJ Charter
- ☐ Other _____

29. Have any training documentation or session been developed for this specific AI system?

☐ Yes ☐ No