



Deliverable 5.4

Assessment tools, training activities, best practice guide v2

31-03-2025

1.1



**Funded by
the European Union**

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 or the Agency. Neither the European Union nor the granting authority can be held responsible for them.

PROPERTIES	
Dissemination level	Public
Version	1.1
Status	Final
Beneficiary	JSI
License	 This work is licensed under a Creative Commons Attribution-No Derivatives 4.0 International License (CC BY-ND 4.0). See: https://creativecommons.org/licenses/by-nd/4.0/

AUTHORS		
	Name	Organisation
Document leader	Tanja Zdolšek Draksler	JSI
Participants	Ana Fabjan	JSI
	Lila Iosif	MT
	Vicky Stroumpou	MT
	Christos Gogos	AUTH
	Georgia Panagiotidou	AUTH
	Celia Parralejo Cano	DPB
Reviewers	Dimitris Apostolopoulos	VIL
	Danai Kyrkou	VVV

VERSION HISTORY				
Version	Date	Author	Organisation	Description
0.1	18.2.2025	Tanja Zdolšek Draksler	JSI	Initial ToC
0.2	25.2.2025	Tanja Zdolšek Draksler	JSI	Initial draft
0.3	4.3.2025	Tanja Zdolšek Draksler	JSI	Draft
0.4	7.3.2025	Tanja Zdolšek Draksler, Lila Iosif, Vicky Stroumpou	JSI, MT	Writing of Introduction and Chapters 2-4. Update from contributions of partner MT for Chapter 2.
0.5	14.3.2025	Tanja Zdolšek Draksler, Christos Gogos, Georgia Panagiotidou	JSI, AUTH	Update from contributions of partner AUTH, Chapter 4. Finalisation of the deliverable.
0.6	20.3.2025	Tanja Zdolšek Draksler	JSI	Internal review version preparation.
0.7	24.3.2025	Dimitris Apostolopoulos, Danai Kyrkou	VVV, VIL	Internal review
1.0	28.3.2025	Tanja Zdolšek Draksler, Christos Gogos	JSI	Integrated corrections
1.1	31.3.2025	Tanja Zdolšek Draksler	JSI	Final version

Table of Contents

Abstract	6
1 Introduction.....	7
1.1 Purpose and scope.....	7
1.2 Document structure.....	8
2 Learning and training	9
2.1 Targeted stakeholders	9
2.2 Stakeholders' training activities	9
2.2.1 3 rd AI4Gov training	11
2.2.2 4 th AI4Gov training	13
2.2.3 Summary of all AI4Gov trainings.....	16
2.2.4 Educational resources	16
3 Collection of Learning courses: AI4Gov Massive Open Online Courses	18
3.1 Learning course 1: "Trustworthy and democratic AI – Fundamentals"	23
3.2 Learning course 2: "Inteligencia Artificial fiable y democrática - Fundamentos"	28
3.3 Learning course 3: "Zaupanja vredna in demokratična umetna inteligenca – osnove"	30
3.4 Learning course 4: "Trustworthy and democratic AI - creating awareness and change"	33
3.5 Learning course 5: "Trustworthy and Democratic AI experts"	36
3.6 Internal focus group and OpenLearn Create quality assurance (QA).....	37
3.7 Further dissemination and exploitation of learning courses and educational resources	39
4 Development of (self)assessment tools on ethical and transparent AI.....	42
4.1 Key considerations of self-assessment tools regarding the development of responsible AI algorithms	42
4.2 Design and guidelines for implementation of assessment tools based on HRF compliant questionnaires	44
4.3 Scoring mechanisms for the assessment tools	53
4.4 Best practice guide.....	55
5 Conclusions.....	58
6 References	60

List of figures

Figure 1: Screen capture 1 from the 3 rd AI4Gov training.....	12
Figure 2: Screen capture 2 from the 3 rd AI4Gov training.....	13
Figure 3: Picture from the 4 th AI4Gov training (1)	14
Figure 4: Pictures from the 4 th AI4Gov training (2)	15
Figure 5: Sub webpage for Learning & training at the AI4Gov webpage.....	17
Figure 6: Screen capture of the collection of learning courses from the learning platform	20
Figure 7: AI4Gov learning collection graphics (banner).....	22
Figure 8: Screen capture 1 of learning course “Trustworthy and democratic AI – Fundamentals”	24
Figure 9: Screen capture 2 of learning course “Trustworthy and democratic AI – Fundamentals”	25
Figure 10: Screen capture 3 of learning course “Trustworthy and democratic AI – Fundamentals”	26
Figure 11: Certificate of completion and digital badge for the first learning course.....	27
Figure 12: Graphics for the first learning course (banner)	27
Figure 13: Graphics presented in a banner and a screen capture of the second learning course.....	29
Figure 14: Certificate of completion and digital badge for the second learning course.....	30
Figure 15: Graphics presented in a banner and a screen capture of the third learning course	32
Figure 16: Certificate of completion and digital badge for the third learning course	33
Figure 17: Screen capture from the learning platform for the learning course four	35
Figure 18: Certificate of completion and digital badge for the learning course four	36
Figure 19: Screen captures from the online questionnaire for the internal focus group for learning courses	38
Figure 20: Screen capture of the book option of the learning course.....	41
Figure 21: Holistic Regulatory Framework	44

List of Tables

Table 1: Trainings and educational resources development overview	10
Table 2: Program of the 3 rd AI4Gov training.....	11
Table 3: Program of the 4 th AI4Gov training.....	14
Table 4: Summary of all AI4Gov trainings	16
Table 5: AI4Gov learning collection overview	21

Abbreviations

Abbreviation	Description
AI	Artificial Intelligence
HRF	Holistic Regulatory Framework
IRCAI	International research centre on artificial intelligence
KPI	Key Performance Indicator
MOOC	Massive Open Online Course
OER	Open Educational Resources
QA	Quality Assurance
TRL	Technology readiness level
XAI	Explainable Artificial Intelligence

Abstract

This document (D5.4 Assessment tools, training activities, best practice guide v2) has been developed within WP5 “Citizen-centric training for creating awareness & change” and is presenting the second version of the deliverable, released on M27, while the first version was delivered on M18 as D5.3. The deliverable describes the learning and training activities with the presentation of used materials (T5.2 and T5.3) and introduces a set of (self) assessment tools and checklists on ethical and transparent AI (T5.4).

1 Introduction

1.1 Purpose and scope

This deliverable is the result of the work in WP5 “Citizen-centric training for creating awareness & change” in tasks 5.2, 5.3 and 5.4 in the time period M3-M27, but with the main focus on the second period of the WP5: M19-M27. The first period was already reported in D5.3 (first version of deliverable). Moreover, as task 5.5 relates to task 5.4, some aspects and results from T5.5 were also reported in D5.3 although not foreseen in the GA. However, we did not include T5.5 work in this deliverable (D5.4) as T5.5 has a dedicated deliverable also for M27 and it will be reported there (D5.5: Report and SoS addressing organizations).

The present document (D5.4 Assessment tools, training activities, best practice guide v2) is presenting the second version of the deliverable, released on M27. The first version of the deliverable (D5.3 Assessment tools, training activities, best practice guide v1) was released on M18. The deliverable describes the learning and training activities with the presentation of used materials (T5.2 and T5.3) and introduces a set of (self) assessment tools and checklists on ethical and transparent AI (T5.4).

WP5 is strongly interlinked with all other project work packages, especially with the technical tasks (WP2, WP3, WP4, WP6), given the fact that WP5 is transferring knowledge gained to the stakeholders and the general public (citizens). Specifically, WP5 is comprising information about the different developed components of the AI4Gov project, to enable stakeholders to learn and train how to make use of it. During the first period of WP5 and the first deliverable (D5.3), the focus was mostly on general knowledge and technology, to introduce the stakeholders and general public with the basic concepts about bias, artificial intelligence (AI), governance and ethics. On the other side, this deliverable presents a core element of the AI4Gov community, which is the reason why it is also closely linked with the dissemination and community building activities of the project (WP7).

WP5 main objectives are to develop and implement training materials designed to enhance awareness and foster citizen engagement, participation, and inclusiveness in democratic processes. The specific objectives are to: (1) create the training materials and establish continuous learning capabilities, (2) promote user awareness and education through workshops and courses, and (3) enhance collaboration and linking with external policy frameworks and stakeholders.

Within the specific needs identification (impact canvas of AI4Gov), the following two were identified in relation to WP5 and D5.3/D5.4: [J] *Increase citizens’ awareness and create change based on Inclusive AI* and [L] *User-centric guidelines/training materials produced to adapt and innovate AI and Big Data systems*. Expected results in relation to the stated needs and to WP5 were identified as results R13 (*User-centric training materials*) and R14 (*Massive open online courses (MOOCs)*).

1.2 Document structure

The deliverable is structured in five chapters, starting with Chapter 1 that introduces the document, including the purpose and scope, and document structure. The following chapters (Chapters 2-4) are dedicated to individual tasks from WP5, more in detail Chapter 2 to T5.2, Chapter 3 to T5.3 and Chapter 4 to T5.4. While Chapters 2 and 3 are focusing on learning and training, Chapter 4 is focused on the development of (self)assessment tools on ethical and transparent AI. Chapter 5 concludes the deliverable and Chapter 6 includes the references.

2 Learning and training

This section of the deliverable introduces the learning and training activities with the presentation of used materials building mostly on T5.2 *Stakeholders' Training Activities*, but partially also on T5.3 *AI4Gov Massive Open Online Courses (MOOCs)*. We will first describe the targeted stakeholders and then move to specific executed activities, explaining the trainings and educational resources.

2.1 Targeted stakeholders

The WP5 content and services target the following stakeholders that were identified from the AI4Gov consortium as the project target groups in the beginning of the work:

- policy makers and public authorities,
- legal organizations,
- AI researchers and AI solutions integrators,
- citizens.

In relation to digital competencies, we are grouping the users to:

- Tech-savvy people (tech users),
- Non-tech-savvy people (or non-tech users).

Moreover, we are additionally grouping them to:

- internal (with special focus on AI4Gov pilot partners: Ministry of Tourism (MT), Municipality of Vari-Voula-Vouliagmeni (VVV), Diputación de Badajoz (DPB).
- external (outside the AI4Gov consortium).

2.2 Stakeholders' training activities

The main aim of the T5.2 *Stakeholders' training activities* is to deliver the AI4Gov training activities addressing different audiences. The goal was to develop educational resources to address occupational needs and mismatches with the main focus on bias in AI. Training activities were divided into part one (year one of project timeline) and part two (year two and year three of project timeline), as an internal assessment at the start of the 5.2 task showed, that we need to focus on fundamentals first, to build a strong knowledge basis, otherwise the more complex topics will not be understood properly. In the second part of the T5.2 activities, we included the AI4Gov technologies, allowing different audiences to gain insights into technical developments of the project. In that way, the aim is to close the AI talent gap within the consortium organizations and within other participants of the project's community. Table 1 shows the training activities and educational resources development overview – reporting the work done in relation to T5.2 (T5.3 also partially included) and also highlighting the achieved key performance indicators (KPIs):

three training workshops foreseen. Partners involved in this task were JSI as leader with contributions from partners MT, VVV and DPB.

Table 1: Trainings and educational resources development overview

Project timeline (project month)		Activity
Year 1	Start of WP5 / M3 (March 2023)	Assessment of technical vocabulary and tasks understanding from non-tech partners (internal activities).
	M4 (April 2023)	Preparation of first training with focus on tech fundamentals.
	M5 (May 2023)	1 st AI4Gov training, hybrid, Ljubljana, Slovenia (KPI 1).
	M6 (June 2023)	Educational resources development (video lectures, transcripts, translations).
	M10 (October 2023)	Preparation of second training.
	M11 (November 2023)	2 nd AI4Gov training, hybrid, Thessaloniki, Greece (KPI 2).
	M12 (December 2023)	Educational resources development (video material).
Year 2 and 3	M15 (March 2024)	The video recordings of the workshop organized by the AI, Big Data, and Democracy taskforce as part of the AI UK Fringe event on March 7 th , 2024, are made available on AI4Gov YouTube channel.
	M15 (March 2024) / M16 (April 2024)	In March 2024 the European Parliament formally adopted the EU AI Act. Preparation of video material: "Introduction to AI Act".
	M16 (April 2024)	Learning course 1 available publicly: AI4Gov MOOC: "Trustworthy and Democratic AI – Fundamentals".
	M22 (October 2024)	Setting up a dedicated sub webpage for Learning & training in the context of the official AI4Gov project webpage (AI4Gov Learning&training, 2025).
	M25 (January 2025)	Learning course 2 and 3 available publicly: Inteligencia Artificial fiable y democrática – Fundamentos (Spanish); Zaupanja vredna in demokratična umetna inteligenca – osnove (Slovenian).
	M26 (February 2025)	3 rd AI4Gov training, online (in Greek and English) (KPI 3).
	M27 (March 2025)	4 th AI4Gov training, in person, Badajoz, Spain (KPI 4).

Next, we describe the carried-out activities from M19 to M27 in more detail (activities from M3 to M18 were already reported in D5.3, which was submitted on M18).

2.2.1 3rd AI4Gov training

The 3rd AI4Gov training workshop titled “AI and Bias: Real-World Challenges and Solutions with a Focus on Tourism” organized by project partner Ministry of Tourism (MT) in collaboration with project partner Municipality of Vari-Voula-Vouliagmeni (VVV), was held online on February 27th, 2025. The agenda of the training is presented below in Table 2.

Table 2: Program of the 3rd AI4Gov training

Time	Title	Presenter
10:00	Welcome	Stavroula Kefala, Head of the Directorate of Research, Greek Ministry of Tourism
10:05-10:15	Short introduction to the AI4Gov project	Vicky Stroumbou, Head of the Department of Monitoring EU Programmes, Greek Ministry of Tourism
10:15-10:35	Fundamentals of AI & MOOCs	Lila Iosif, Department of Monitoring EU Programmes, Greek Ministry of Tourism
10:35-11:00	Bias in AI-Real life examples	Matej Kovacic, Jožef Stefan Institute (JSI), Slovenia
11:00-11:35	The use of Artificial Intelligence in tourism and tourism businesses	Mr. Varelas Sotiris, Assistant Professor of the Department of Tourism Studies – University of Piraeus, Greece
11:35-11:50	Break	
11:50-12:15	Dealing with bias that can occur in AI systems – Interconnection between bias and data	Kostis Mavrogiorgos, UPRC
12:15-12:30	Designing reliable tourism policies based on data: Pilot applications of AI in the Municipality of Vari-Voula-Vouliagmeni	Dimitris Apostolopoulos, Municipality of Vari-Voula-Vouliagmeni, Greece
12:30-12:50	AI4GOV Tools – Ethical & trustworthy AI	Dimitris Kokkinis, Greek Ministry of Tourism
12:50-13:00	Closure – Q & A	

The first session focused on the scope of AI4Gov project and its main deliverables, the fundamentals of AI and bias as well as the developed AI4Gov learning courses (MOOCs). Also, real life examples of bias in AI were presented. The second session was dedicated to AI and tourism, and the third session focused on bias in AI. The pilot applications of the project partner Municipality of Vari-Voula-Vouliagmeni (VVV) were presented, as well as the tools developed within the framework of AI4Gov project for ethical and trustworthy AI.

The training was carried out mostly in Greek language, with one presentation in English. The training was recorded and the video recording will be available within the AI4Gov learning materials.

The training was very well received from the audience. The type of the audience were policy makers, students and educators/teachers. The latter were from the Ministry's Tourism School: *Experimental, Higher Vocational Training School of Tourism of Macedonia*. More specifically:

- from the Department of Executive Hospitality Units, 23 students,
- from the Experimental specialty "Hospitality Operations Executive", 36 students,
- from the Experimental specialty: "Travel & Tourism Operations Executive", 36 students.

The policy makers were from the Ministry of Tourism and from the Municipality of Vari-Voula-Vouliagmeni. All together there were 115 attendees (95 students, 10 teachers, 5 policy makers and 5 employees from public bodies). The students were very interested in the AI4Gov project, especially in the developed AI4Gov learning courses. They also gave very positive feedback at the end of the workshop; they found the lecture about AI and tourism especially interesting. Figures 1 and 2 are screen captures from the online event.

Figure 1: Screen capture 1 from the 3rd AI4Gov training



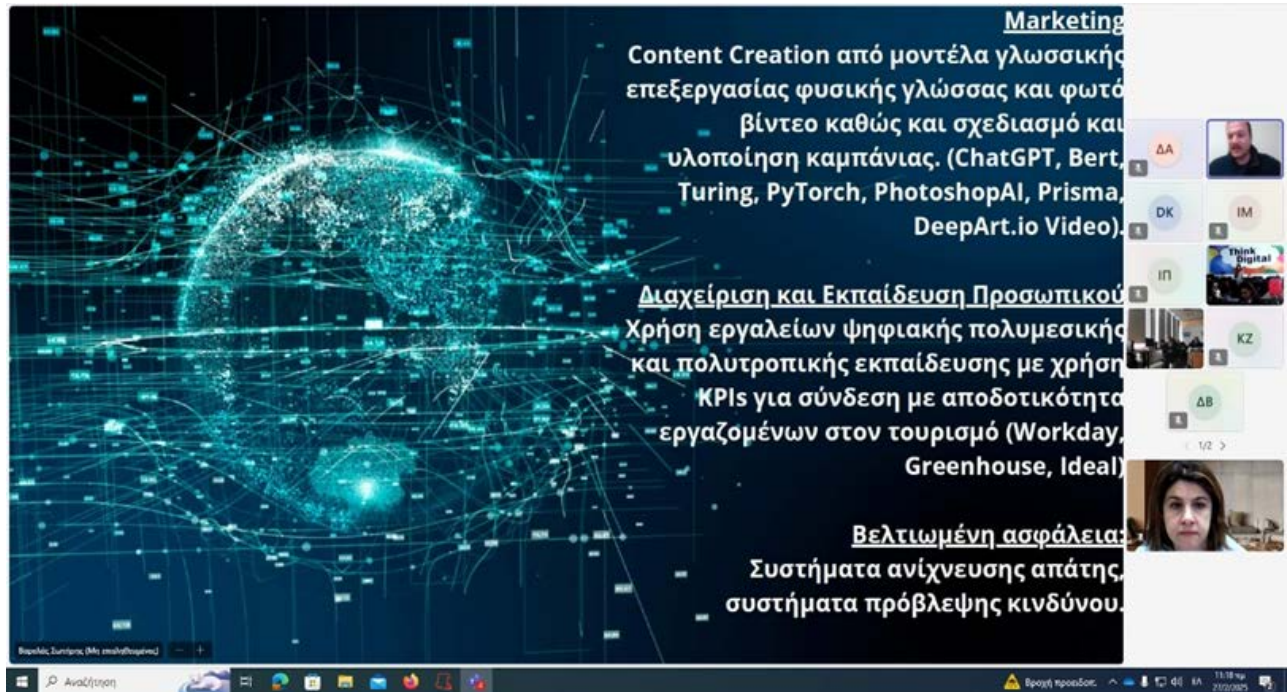


Figure 2: Screen capture 2 from the 3rd AI4Gov training

2.2.2 4th AI4Gov training

The 4th AI4Gov training workshop titled “Towards Fair AI: Real World Challenges and Solutions”, organized by project partner Diputación de Badajoz (DPB), was held in person on March 10th, 2025. The aim of the event was to present the foundations of AI, raise awareness of the importance of ethics in AI, provide examples of AI applications in the real world and disseminate the developed AI4Gov learning courses (MOOCs). The objectives and solutions of the AI4Gov project were presented, along with the tools developed for the Badajoz pilot. The training was carried out in Spanish language. The agenda of the training is presented in Table 3, while Figures 3-4 show pictures from the 4th AI4Gov training workshop.

The training was very well received by the audience, which consisted mostly of students studying in the fields of robotics, physics, and mathematics. Additionally, university professors, DPB colleagues, and other stakeholders attended. In total, there were 110 attendees. The participants were highly engaged and showed great interest in the AI4Gov project and the learning courses (MOOCs) created.

Table 3: Program of the 4th AI4Gov training

17:30	Arrival
17:35 (5')	Welcome <i>Ulises Gamero, Head of Digital Transformation Service, Diputación de Badajoz, Spain</i>
17:40 (20')	IA Introduction. AI4Gov Project <i>Celia Parralejo Cano, Mathematician y Technician in the Technology and Digitalization Area, Badajoz Provincial Council</i>
18:00 (20')	Badajoz es Más <i>Miguel López Corbacho, Dynamizer of the Badajoz Es Más project</i>
18:20 (20')	AI Without Traps: Towards a Fair and Transparent Future. <i>Rosa Elvira Lillo Rodríguez, Professor of Statistics and IoT, and director of IBiDat (UC3M-Santander Big Data Institute) at Universidad Carlos III de Madrid</i>
18:40 (20')	Mutual Benefits of Collaboration in Artificial/Human Intelligence <i>Francisco Fernández de Vega, Professor of Computer Architecture and Technology at UEx, Coordinator of the GEA research group, and Chair of the IEEE TF Creative Intelligence</i>
19:00 (20')	Application of Artificial Intelligence Techniques in Satellite Imagery for the CAP in Extremadura (online) <i>Adolfo Lozano Tello, Professor of the Quercus Group at the University of Extremadura</i>



Figure 3: Picture from the 4th AI4Gov training (1)



Figure 4: Pictures from the 4th AI4Gov training (2)

2.2.3 Summary of all AI4Gov trainings

All in all, four training workshops were organized during the T5.2 duration, reaching more than 300 participants from different domains and different countries (overview showed below in Table 4). With that number the foreseen KPIs, which were minimal three training workshops, were exceeded.

Table 4: Summary of all AI4Gov trainings

Time of training	Training	Participants number and type
May 2023 (KPI 1)	1 st AI4Gov training, hybrid, Ljubljana, Slovenia.	48 participants (policy makers from Europe, academia, IT industry)
November 2023 (KPI 2)	2 nd AI4Gov training, hybrid, Thessaloniki, Greece.	50 participants (national policy makers, academia, students)
February 2025 (KPI 3)	3 rd AI4Gov training, online.	115 participants (students, teachers, policy makers and employees from public bodies)
March 2025 (KPI 4)	4 th AI4Gov training, in-person, Badajoz, Spain.	110 participants (students, teachers, policy makers and employees from public bodies)
Sum:	4 trainings	323 participants

We can conclude that the AI4Gov trainings did play a crucial role in reaching policy makers and citizens, with advancing their knowledge and fostering a better understanding not only of bias in AI, but also about AI in general (with the focus on ethical and trustworthy AI). The crucial take away message is that AI literacy is essential in today's technology-driven society, but citizens are lacking in it. That is why it is an important topic, and we do need to focus on increasing AI literacy, among citizens particularly. Although WP5 reached its end, the AI4Gov learning resources remain and their dissemination and exploitation will continue to be pursued and used.

2.2.4 Educational resources

The organised trainings present the main source for AI4Gov educational resources, which were the building blocks for the first AI4Gov learning course (MOOC). We already described the educational resources in D5.3, highlighting that especially the first AI4Gov training was very valuable for this purpose, as it was filmed by the VideoLectures.NET team and valuable video lectures were produced in the duration of 3 hours. Video lectures from the first AI4Gov training are openly available online with presentation slides and subtitles on the VideoLectures.NET

platform (VideoLectures.NET, 2025) (9 videos) and also on YouTube via the AI4Gov project channel (AI4Gov YouTube channel, 2025). Moreover, other AI4Gov trainings delivered recordings of the presentations and are also available online as videos on YouTube AI4Gov project channel.

As there is a demand for online educational resources for teaching and learning, we were operating towards generating the educational material in the type of the open educational resources (OER). In Autumn 2024 a dedicated sub webpage for Learning & training (AI4Gov Learning&training, 2025) in the context of the official AI4Gov project webpage was set up (screen capture presented in Figure 5). The main aim of this subpage is to offer all AI4Gov learning resources in one place. Besides that, this subpage represents a simple overview of AI4Gov's training and learning activities and related resources, aimed at enhancing awareness and citizen engagement, as well as raising awareness and knowledge about trustworthy AI and bias in AI. Although the WP5 ends, the learning & training sub-page will be updated with possible new learning resources till project end.

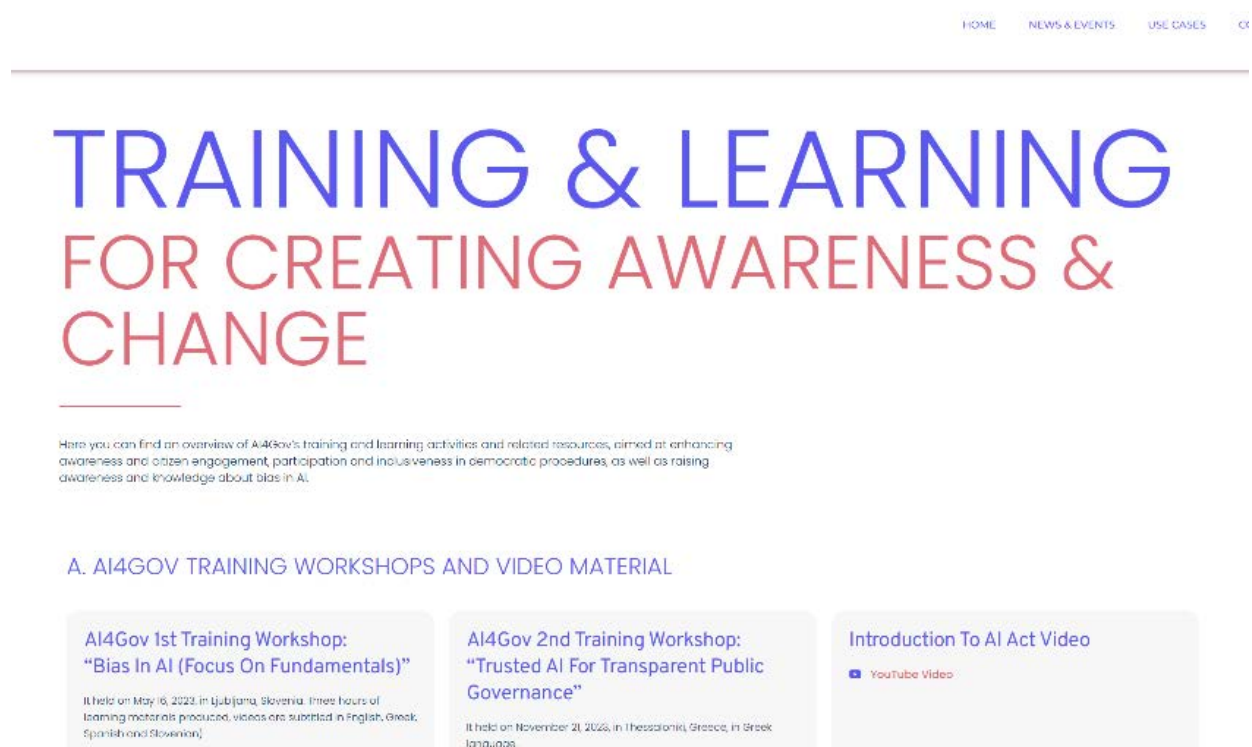


Figure 5: Sub webpage for Learning & training at the AI4Gov webpage

In the next chapter we report the work done in T5.3.

3 Collection of Learning courses: AI4Gov Massive Open Online Courses

This segment of the document presents the AI4Gov learning courses also known as the AI4Gov massive open online courses (MOOCs) developed in the scope of T5.3, which started in M4 and ends in M27. The methodology and the development of the curriculum “Trustworthy and democratic AI”, that forms the basis for task 5.3, was already in detail described and reported in deliverable D5.3. The curriculum is linked to specific learning objectives and competencies gained, that were also already presented in the first version of the deliverable (D5.3). As stated in D5.3, the curriculum was developed as a live document, meaning, that changes during the development phase of the learning courses could be made. Although the learning courses maintained the initial focus and are building on the pre-defined topics of the initial curriculum, the number of the learning courses changed as initially only three were planned (the project foreseeing KPIs= 3), but at the end five learning courses were developed.

All of the learning courses were developed on the platform OpenLearn Create¹, “an innovative open educational platform where individuals and organisations can publish their open content, open courses and resources”.

Partners involved in this task were JSI and MT. JSI was leading the task and also being the main contributor and developer, carrying out the next steps in the following order:

1. analysis of potential learners needs (based on the feedback from pilot partner MT and JSI’s previous experiences in e-learning),
2. design, development and writing of the theoretical framework / curriculum with structure (learning courses, modules, lessons),
3. learning courses development and writing (course structure, content, assessments) based on the theoretical framework, alongside with planning the learning experience,
4. finding and adding different types of existing relevant materials (video lectures, comics, real-life examples etc.),
5. finding and validating relevant e-learning platforms for learning courses development, building and hosting, based on predefined criteria + selecting the most appropriate one,
6. building the learning courses within the OpenLearn Create platform (Moodle based),
7. creation and development of graphic designs for each individual learning course and for the whole learning collection,
8. guiding the internal focus group, preparing an online questionnaire to collect feedback from internal focus group, analyse their feedback and in addition to the results align the learning,
9. analysing the learning courses to define the duration of them (based on e-learning metrics (for example video material of 1 hour is counted as 2 hours of materials) and based on specific feedback from the internal focus group),
10. connecting all learning courses into one learning collection titled AI4Gov “Trustworthy and democratic AI”,

¹ <https://www.open.edu/openlearncreate/>

11. creating specific certificates of completion and digital badges for each learning course, and inserting it in Moodle,
12. acting as learning courses manager, communicating with the platform manager on one side and with the learners on the other,
13. maintaining the learning courses, for example make specific corrections if an error is spotted, etc.,
14. preparing texts and graphics for internal and external communication, dissemination and exploitation activities for the learning collection or learning courses individually.

Partner MT collaborated as the internal focus group, ensuring first hand feedback and internal quality assurance (QA) on the curriculum and learning courses (explained in detail in Chapter 3.6).

The learning courses were developed in consecutive order, meaning that we first worked on the first learning course, titled Trustworthy and democratic AI – Fundamentals. When publicly available, we communicated it internally to project partners, with specific focus on communication to pilot partners (MT, VVV, DPB, JSI). After receiving their feedback, the most relevant raised question was, how to tackle the language barrier that was specifically raised from the Spanish partner. That is why it was decided, that it would be valuable to prepare a learning course on fundamentals in Spanish language and the learning course number two responds to this need. In addition, partner JSI is in close connection to Slovenian ministries (which also actively attended the 1st training workshop in Ljubljana, Slovenia), that raised great interest for the learning and training resources developed, so we consequently decided to build a learning course on fundamentals in Slovenian language as well (learning course number three responds to this need). Partners from Greece (pilots) did not indicate specific request that a learning course on fundamentals in Greek would be needed, mainly because the 1st training workshop in Ljubljana, Slovenia, was professionally filmed and the presentations are publicly available as video lectures with presentation slides and with respective translations to Greek language (alongside translations to Spanish and Slovenian language too), that were automatically translated, but afterwards manually checked and edited, meaning that they are 100% correct.

After the fundamentals of AI, bias, regulation on AI and ethical AI was covered within the first three learning courses, we continued to deliver more complex learning content, that was developed in the fourth learning course. While the fourth learning course is quite extensive, it offers the learning materials to the learner in two ways. First it delivers regular content, and second it delivers a lot of optional content in the forms of video lectures, which can be used, depending on the basis knowledge and the preferences and needs of the users (learners).

The majority of topics specified in the curriculum were elaborated in the first and fourth learning course).

Finally, we decided to develop a last, fifth learning course, which corresponds with the idea to briefly present the AI4Gov project results, that are available as tools that can be used from citizens, public administration and all other interested parties.

All of the learning courses are related and form the so called AI4Gov learning collection titled: “Trustworthy and democratic AI” (Figure 6 is a screen capture of the collection showing the learning courses (last, fifth learning course is still pending).

[Courses](#)
[Collections](#)
Trustworthy and Democratic AI

Trustworthy and Democratic AI

The collection equips learners with essential knowledge to understand and identify bias in Artificial Intelligence systems. It aims to make AI ethics, governance, and technical standards accessible to everyone, helping ensure future AI development prioritizes fairness, accountability, and transparency. Join us to earn a digital badge while contributing to the development of more trustworthy and equitable AI systems that benefit all of society.

Trustworthy and Democratic AI - Creating Awareness and Change Designed for learners of all backgrounds this learning course will empower you with the knowledge to navigate the AI-driven world with confidence while championing the values of trust and democracy. Try this course >	Course 50 hrs
Trustworthy and Democratic AI - Fundamentals AI powers our world, from daily virtual assistants to healthcare breakthroughs, making transparency and ethical development essential as its influence grows. Try this course >	Course 15 hrs
Zaupanja vredna in demokratična umetna inteligenca - osnove Umetna inteligenca (UI) poganja naš svet, od vsakodnevnih virtualnih pomočnikov do prebojnih dosežkov v zdravstvu, zato sta preglednost in etični razvoj bistvenega pomena, saj se njen vpliv povečuje. Try this course >	Course 15 hrs
Inteligencia Artificial fiable y democrática - Fundamentos La Inteligencia Artificial (IA) impulsa nuestro mundo, desde los asistentes virtuales cotidianos hasta los avances sanitarios, por lo que la transparencia y el desarrollo ético son esenciales a medida que crece su influencia. Try this course >	Course 15 hrs

Figure 6: Screen capture of the collection of learning courses from the learning platform

Table 5 presents an overview of the learning collection with developed learning courses, showing their respective titles, language used, anticipated learning duration, level of knowledge and time of public release.

Table 5: AI4Gov learning collection overview

AI4Gov learning collection: “Trustworthy and democratic AI” (QR code for the whole collection) 					
No.	Title	Info	Duration	Public on OpenLearn Create	QR code
Learning course 1	Trustworthy and democratic AI – Fundamentals (Zdolšek Draksler, Fabjan, Guček, & Kovačič, 2024)	Beginner level, in English	15 hours	from April 2024	
Learning course 2	Inteligencia Artificial fiable y democrática - Fundamentos (Zdolšek Draksler, Fabjan, Guček, & Kovačič, 2025a)	Beginner level, in Spanish	15 hours	From February 2025	
Learning course 3	Zaupanja vredna in demokratična umetna inteligenca – osnove (Zdolšek Draksler, Fabjan, Guček, & Kovačič, 2025b)	Beginner level, in Slovenian	15 hours	From February 2025	
Learning course 4	Trustworthy and democratic AI - creating awareness and change (Zdolšek Draksler, Fabjan, Guček, & Kovačič, 2025c)	Intermediate, in English	50 hours main content + optional content	From March 31 st , 2025.	
Learning course 5	Trustworthy and Democratic AI experts (Zdolšek Draksler, Fabjan, & Guček, 2025d)	Intermediate, in English	10 hours	Expected in May/June 2025.	Pending
Sum: 5 learning courses (100+ hours)					

The shared characteristics of all developed learning courses are that they are all open and free, designed to provide learners with the knowledge necessary for critical thinking about bias that can occur in AI systems and their consequences. It brings forward the need that AI systems need to be trustworthy and ethical. The aim is to democratize the knowledge of AI ethics, AI governance, and technical standards.

As we wanted to adapt to different stakeholders and to their requests and comments on the first learning course, we decided for two options of usage within the OpenLearn Create platform: access with registration to the platform and access without registration to the platform. First option, the user can register and exploit the full potential of the course (access the assessment part, receive a digital badge and a certificate of completion, contact the course teachers and administrators, comment etc.). Second option, the user can go through the learning course without being registered (no access to assessment of knowledge like quizzes, without the option to gain a digital badge and certificate of completion). This is valid for all learning courses in the collection.

Figure 7 shows the graphics developed specifically for the AI4Gov learning collection. The graphics were being consistently used for learning course development on the respective platform as well as for communicating, disseminating and promoting the learning collection. We plan to still focus on the learning collection dissemination and exploitation after WP5 ends in close collaboration with WP7 and WP6 and with the aim to increase the numbers of registered and unregistered learners (guest visitors). After the project end the activities will remain in the scope of partner JSI, exploiting the options of using and re-using the learning courses particularly for policy makers, as JSI and their *International research centre on artificial intelligence (IRCAI) under the auspices of UNESCO* is actively collaborating with several Slovenian and foreign ministries around Europe and worldwide, UNESCO and OECD.

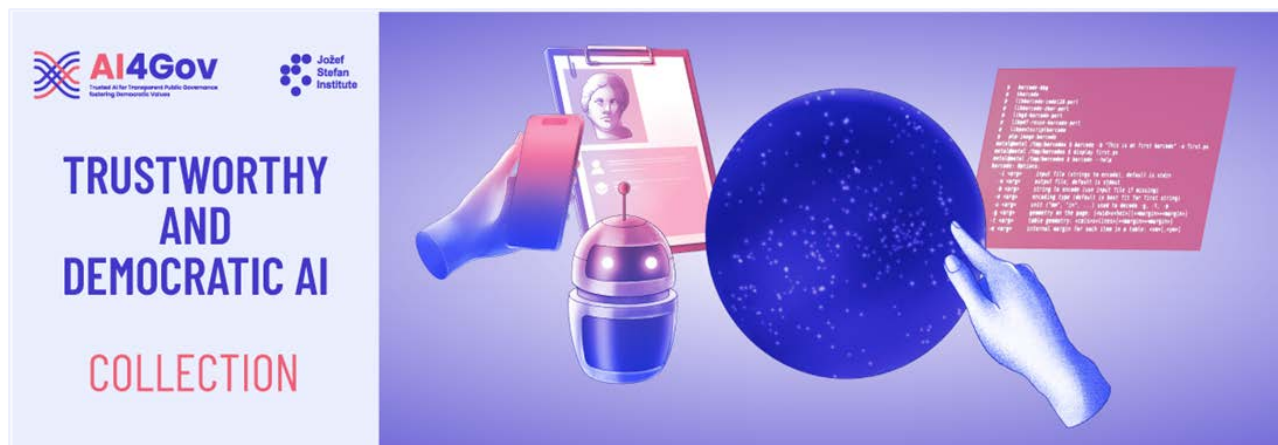


Figure 7: AI4Gov learning collection graphics (banner)

In the following sections, all developed learning courses are more in detail presented.

3.1 Learning course 1: “Trustworthy and democratic AI – Fundamentals”

The first AI4Gov learning course titled “Trustworthy and democratic AI – Fundamentals” was designed to provide learners around the globe with the foundational knowledge necessary to understand bias in AI. It aims to democratize the knowledge of AI ethics, governance, and technical standards to ensure that future AI systems are developed and used with fairness, accountability, and transparency at their core. The first learning course was made publicly available in April 2024. The structure of the learning course is as follows:

Module 1: Introduction to Trustworthy and Democratic AI (Fundamentals)

Lesson 1.1: What is Artificial Intelligence?

Lesson 1.2: The Role of AI in Society

Lesson 1.3: The Scope of Trustworthy AI

Lesson 1.4: The Importance of Ethical AI

Quiz 1: Introduction to Trustworthy and Democratic AI (Fundamentals)

Module 2: Bias in AI - Understanding Bias

Lesson 2.1: Defining Bias in AI

Lesson 2.2: Types of Bias and Sources of Bias

Lesson 2.3: Impact of Bias

Quiz 2: Bias in AI - Understanding Bias

Module 3: Data and Bias

Lesson 3.1: Bias in Data Collection

Lesson 3.2: Data Sampling Methods

Lesson 3.3: Ethical Data Sourcing

Lesson 3.4: Data Pre-processing and Bias Reduction

Lesson 3.5: Real-world Data Bias Case Studies

Brainstorming Task

Hands on session: case studies on detecting bias

Quiz 3: Data and Bias

Module 4: Ethical AI Governance

Lesson 4.1: Legal and Regulatory Frameworks

Quiz 4: Ethical AI Governance - Legal and Regulatory Frameworks

Resources & Authorship

Glossary

Acknowledgements

The learning analytics for the first learning course shows 69 registered users (OpenLearners), but on the other side a high number of guest learners (non-registered users) reaching 12.433. We can assume that learners are interested in gaining new knowledge, but do not need the certificate of completion and the digital badge, which could be related to the fact that informal e-learning is still not being recognised, accepted and valued the same as formal learning.

Figures 8-10 show screen captures from the learning platform for the respective learning course, Figure 11 showing the digital badge and the certificate of completion and Figure 12 showing the visual graphic prepared for this specific learning course.

The screenshot shows the OpenLearn Create website interface. The header includes the OpenLearn Create logo, a search bar, and navigation links. The main content area features a course banner for 'Trustworthy and Democratic AI - Fundamentals' with a blue background and AI-related graphics. To the right of the banner, there's a 'Free statement of participation on completion' icon. Below the banner, there are tabs for 'Course description', 'Course content', and 'Course reviews'. The 'Course description' tab is active, showing a 'Welcome' message and details about the course. On the right side, there's a sidebar with 'About this course' information, including study hours, level, and ratings. Below that, there's a 'Sign up to get more' section and a 'Share this course' section with social media icons. At the bottom right, there's a 'Course rewards' section highlighting the 'Free Statement of Participation' and the 'Earn a free digital badge'.

OpenLearn Create | Hosting resources for creators and learners

Search for free courses and collections

Home | Get started | Create a course | Free courses | Collections | Sign up / Sign in

Home > Courses > OER > Trustworthy and Democratic AI - Fundamentals

Course

Trustworthy and Democratic AI - Fundamentals

Free statement of participation on completion

About this course

- 10 hours study
- Level 1: Introductory
- Gain a digital badge

Ratings ★★★★★
5 out of 5 stars

Sign up to get more

You can start learning at any time. By signing up and enrolling you can track your progress and earn a Statement of Participation upon completion, all for free.

[View this course](#)

[Sign up to get more](#)

Share this course

[f](#) [t](#) [in](#) [g+](#) [e](#)

Course rewards

- Free Statement of Participation** on completion of these courses.
- Earn a free digital badge** if you complete this course, to display and share your achievement.

Course description | Course content | Course reviews

Welcome

Welcome to the exciting world of Artificial Intelligence (AI) – a realm where machines and algorithms seamlessly blend with our daily lives, shaping the future in ways unimaginable just a few years ago. AI is everywhere, from the personalized recommendations on your favorite streaming platform to the voice-activated virtual assistants simplifying your tasks. It's revolutionizing healthcare, optimizing logistics, and even contributing to scientific breakthroughs. However, as AI continues to evolve, so does the imperative for it to be explainable, fair, and transparent.

Join us as we explore the foundations of building trustworthy and democratic AI. From understanding the basics of AI to the ethical considerations that underpin responsible AI development, this course is designed for learners of all backgrounds. Whether you're a tech enthusiast, a business professional, in public administration, or simply curious about the forces shaping our digital landscape, this learning course will empower you with the knowledge to navigate the AI-driven world with confidence.

Get ready to unlock the potential of AI while championing the values of trust and democracy. Let's embark on this transformative learning experience together!

Training materials include: text, video lectures, presentation slides, comic books, quizzes.

Authors:

Tanja Zdošek Draksler, PhD, Ana Fabjan, Alenka Guček, PhD, Matej Kovačič, PhD

Funding:

This work was done within the Horizon Europe project **AI4Gov: Trusted AI for Transparent Public Governance fostering Democratic Values** (Grant agreement ID: 101094905).

Figure 8: Screen capture 1 of learning course “Trustworthy and democratic AI – Fundamentals”



Certificate of Completion

This certificate is awarded upon successful completion of the course "Trustworthy and Democratic AI - Fundamentals".

The course covers both theoretical concepts and practical applications, emphasising the importance of developing AI systems that are not only technically capable, but also aligned with democratic values and societal trust. Participants gain a comprehensive understanding of AI fundamentals and developed critical skills to deal with AI bias and recognise ethical considerations underlying responsible AI.



Gain a digital badge

By studying the course you will have the opportunity to gain a digital badge – you need to click on the 'Enrol' button to be able to do the quizzes and earn the badge.



This course is part of a collection

This course is part of a collection of courses called Trustworthy and Democratic AI. There are 3 courses in this collection so you may find other courses here that maybe of interest to you.

[See this collection >](#)

Course learning outcomes

- Understand the concepts of artificial intelligence, machine learning, and deep learning.
- Recognize the transformative potential and ethical implications of AI in various domains.
- Identify the types and sources of bias that can manifest in AI systems.
- Understand the real-world consequences of biased AI.
- Learn how data collection and pre-processing can introduce bias into AI models.
- Implement best practices for collecting, cleaning, and preparing data to reduce bias.
- Study legal and regulatory frameworks governing AI, such as GDPR and AI Act.

[View this course now](#)

Figure 9: Screen capture 2 of learning course "Trustworthy and democratic AI – Fundamentals"

Administration

- ✓ Book administration
 - Turn editing on
 - Settings
 - Locally assigned roles
 - Permissions
 - Check permissions
 - Filters
 - Logs
 - Restore
 - Import chapter
 - Print book
 - Print this chapter
 - Import from Microsoft Word
 - Export book to Microsoft Word
 - Export this chapter to Microsoft Word

> Course administration

< Course content

Welcome

Module 1: Introduction to Trustworthy and Democratic AI (Fundamentals)

Module 1: Understanding AI Fundamentals

Quiz 1: Introduction to Trustworthy and Democratic AI (Fundamentals)

Module 2: Bias in AI - Understanding Bias

Module 2: Bias in AI - Understanding Bias

Quiz 2: Bias in AI - Understanding Bias

Module 3: Data and Bias

Module 3: Data and Bias

Brainstorming Task

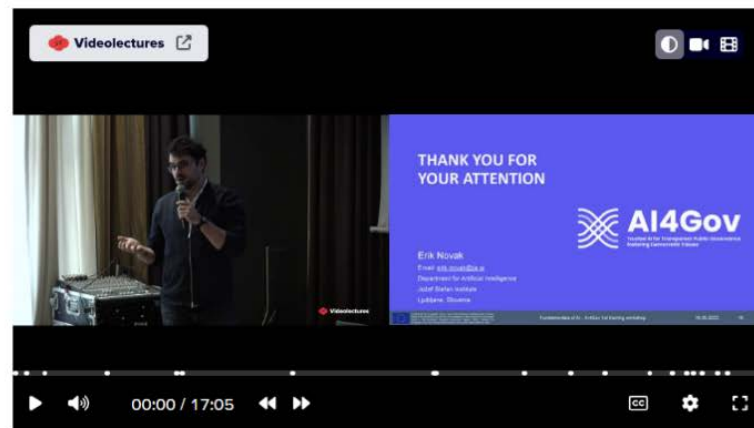
Hands on session: case studies on detecting bias

Volume 2: Learning From Data

LESSON 1.1: WHAT IS ARTIFICIAL INTELLIGENCE?

In the initial lesson, we will acquaint ourselves with the concept of Artificial Intelligence (AI), examining a range of examples and seeing where it's used. Artificial Intelligence (AI) refers to the development of computer systems that can perform tasks that typically require human intelligence. These tasks encompass learning, reasoning, problem-solving, perception, and language understanding. AI systems often utilize machine learning algorithms, enabling them to adapt and improve their performance over time as they process and analyze data. Applications of AI range from virtual assistants like Siri and Alexa to complex tasks such as image recognition, natural language processing, and strategic decision-making in various industries. The ultimate goal of AI is to create machines that can simulate human intelligence and perform tasks autonomously.

Watch the video lecture accompanied **with slides** to learn more about fundamentals of AI.



In addition, watch additional video lectures introducing machine learning and deep learning:



Figure 10: Screen capture 3 of learning course “Trustworthy and democratic AI – Fundamentals”



Figure 11: Certificate of completion and digital badge for the first learning course

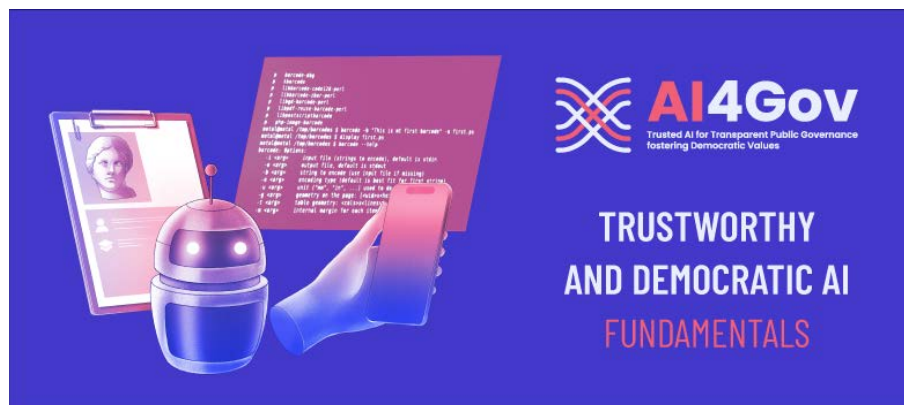


Figure 12: Graphics for the first learning course (banner)

3.2 Learning course 2: “Inteligencia Artificial fiable y democrática - Fundamentos”

The second AI4Gov learning course titled “Inteligencia Artificial fiable y democrática – Fundamentos” in Spanish was adopted from the first learning course to better address the needs of the Spanish pilot in the project. The learning course was developed in summer/autumn 2024 and finalised in October 2024, but due to staff problems at the side of the learning platform, it was made publicly available only in February 2025. The structure of the learning course is following.

Módulo 1: Introducción a la IA fiable y democrática (Fundamentos)

Lección 1.1. ¿Qué es la inteligencia artificial?

Lección 1.2: El papel de la IA en la sociedad

Lección 1.3: El alcance de la IA fiable

Lección 1.4: La importancia de una IA ética

Cuestionario 1: Introducción a la IA fiable y democrática (Fundamentos)

Módulo 2: Sesgo en la IA - Comprender el sesgo

Lección 2.1: Definición del sesgo o prejuicios en la IA

Lección 2.2: Tipos de sesgo y fuentes de sesgo o prejuicios

Lección 2.3: El impacto de los prejuicios

Cuestionario 2: Sesgo en la IA - Comprender el sesgo

Módulo 3: Datos y Prejuicios o Sesgo (Bias)

Lección 3.1: Sesgo (bias) en la recogida de datos

Lección 3.2: Métodos de muestreo de datos

Lección 3.3: Obtención ética de datos

Lección 3.4: Preprocesamiento de datos y reducción de sesgos

Lección 3.5: Estudios de casos reales de sesgo (bias) de datos

Tarea De Tormenta De Ideas

Sesión práctica: Casos prácticos de detección de sesgos

Cuestionario 3: Datos y sesgos o prejuicios

Módulo 4: Gobernanza ética de la IA

Lección 4.1: Marcos jurídico y normativo

Cuestionario 4: Gobernanza ética de la IA - Marcos jurídico y normative

Recursos

Glosario

Agradecimientos

The learning analytics for this learning course is 24 registered users and 531 guest learners, which probably corresponds as a direct result from the 4th training workshop that was held in March 2025 and where this learning course was particularly communicated to the participants of the event (being students, teachers, policy makers).

Figure 13 shows the visual graphic prepared for this specific learning course alongside a screen capture from the learning platform for the respective learning course. Figure 14 shows the digital badge and the certificate of completion.



Inteligencia Artificial fiable y democrática - Fundamentos

Enlace al curso

Course
Inteligencia Artificial fiable y democrática - Fundamentos

You're on this course ✓

Course description | Course content | Course reviews

Bienvenido

Bienvenido al apasionante mundo de la Inteligencia Artificial (IA), un reino en el que las máquinas y los algoritmos se mezclan a la perfección con nuestra vida cotidiana, configurando el futuro de formas inimaginables hace tan solo unos años. La IA está en todas partes, desde las recomendaciones personalizadas de su plataforma de streaming favorita hasta los asistentes virtuales activados por voz que simplifican sus tareas. Está revolucionando la atención sanitaria, optimizando la logística e incluso contribuyendo a avances científicos. Sin embargo, a medida que la IA sigue evolucionando, también lo hace el imperativo de que sea explicable, justa y transparente.

Únase a nosotros para explorar los fundamentos de la creación de una IA fiable y democrática. Desde la comprensión de los fundamentos de la IA a las consideraciones éticas que sustentan el desarrollo responsable de la IA, este curso está diseñado para estudiantes de todos los orígenes. Tanto si eres un entusiasta de la tecnología como un profesional de los negocios, de la administración pública o simplemente sientes curiosidad por las fuerzas que dan forma a nuestro panorama digital, este curso de aprendizaje te dotará de los conocimientos necesarios para navegar con confianza por el mundo impulsado por la IA.

Prepárese para liberar el potencial de la IA al tiempo que defiende los valores de la confianza y la democracia. Embarquémonos juntos en esta experiencia de aprendizaje transformadora.

El material de formación incluye: textos, videoconferencias, diapositivas, cómics y cuestionarios.

Autores:

Tanja Zdošek Draksler, PhD, Ana Fabjan, Alenka Guček, PhD, Matej Kovačič, PhD

Traducción del curso de español realizada por Víctor Jiménez.

Financiación:

Este trabajo se ha realizado en el marco del proyecto **AI4Gov de Horizon Europe: Trusted AI for Transparent Public Governance fostering Democratic Values** (acuerdo de subvención ID: 101094905).

Figure 13: Graphics presented in a banner and a screen capture of the second learning course



Figure 14: Certificate of completion and digital badge for the second learning course

3.3 Learning course 3: “Zaupanja vredna in demokratična umetna inteligenca – osnove”

The third AI4Gov learning course titled “Zaupanja vredna in demokratična umetna inteligenca – osnove” in Slovenian language was adopted from the first learning course. The learning course was made publicly available in February 2025. The structure of the learning course is as follows.

Modul 1: Razumevanje osnov umetne inteligence

Lekcija 1.1: Kaj je umetna inteligenca?

Lekcija 1.2: Vloga umetne inteligence v družbi

Lekcija 1.3: Zaupanja vredna umetna inteligenca

Lekcija 1.4: Pomen etične umetne inteligence

Kviz 1: Uvod v zaupanja vredno in demokratično umetno inteligenco (osnove)

Modul 2: Pristranskost v umetni inteligenci - razumevanje pristranskosti

Lekcija 2.1: Opredelitev pristranskosti v umetni inteligenci

Lekcija 2.2: Vrste pristranskosti in viri pristranskosti

Lekcija 2.3: Vpliv pristranskosti

Kviz 2: Pristranskost v UI - Razumevanje pristranskosti

Modul 3: Podatki in pristranskost

Lekcija 3.1: Predsodki pri zbiranju podatkov

Lekcija 3.2: Metode vzorčenja podatkov

Lekcija 3.3: Etično pridobivanje podatkov

Lekcija 3.4: Pred-obdelava podatkov in zmanjševanje pristranskosti

Lekcija 3.5: Študije primerov pristranskosti v podatkih iz resničnega sveta

Praktični del

Praktično srečanje: Študije primerov o odkrivanju pristranskosti

Kviz 3: Podatki in pristranskost

Modul 4: Etično upravljanje umetne inteligence

Lekcija 4.1: Pravni in regulativni okviri

Kviz 4: Etično upravljanje umetne inteligence - pravni in regulativni okviri

Viri in avtorstvo

Slovar

Zahvala

Figure 15 shows the visual graphic prepared for this specific learning course alongside a screen capture from the learning platform for the respective learning course. Figure 16 shows the digital badge and the certificate of completion.

Course

Zaupanja vredna in demokratična umetna inteligenca - osnove

You're on this course ✓



AI4Gov
Trusted AI for Transparent Public Governance
fostering Democratic Values

**ZAUPANJA VREDNA IN
DEMOKRATIČNA UI
OSNOVE**

Course description	Course content	Course reviews
--------------------	----------------	----------------

Dobrodošli

Dobrodošli v vznemirljivem svetu umetne inteligence (UI) - na področju, kjer se stroji in algoritmi nemoteno mešajo z našim vsakdanjim življenjem in oblikujejo prihodnost na načine, ki si jih še pred nekaj leti nismo znali predstavljati. Umetna inteligenca je povsod, od personaliziranih priporočil na vaši najljubši platformi za video na zahtevo do glasovno aktiviranih virtualnih pomočnikov, ki poenostavljajo vaša opravila. Revolucionizira zdravstvo, optimizira logistiko in celo prispeva k znanstvenim prebojem. Vendar pa se z razvojem UI povečuje tudi nujnost, da je le-ta razločljiva, poštena in pregledna.

Pridružite se nam pri raziskovanju izgradnje zaupanja vredne in demokratične umetne inteligence. Ta tečaj je namenjen uporabnikom vseh profilov, zajema pa vse od razumevanja osnov UI do etičnih vidikov, na katerih temelji odgovoren razvoj UI. Ne glede na to, ali ste tehnološki navdušenec, učitelj, domenski strokovnjak, zaposlen v javni upravi ali pa vas preprosto zanima digitalen svet, vam bo ta učni tečaj omogočil znanje za samozavestno krmarjenje po poteh, ki jih poganja UI.

Pripravite se na potencial umetne inteligence, a se hkrati zavzemajte za vrednote zaupanja in demokracije. Skupaj se podajmo na to transformativno učno izkušnjo!

Gradivo za usposabljanje vključuje: besedilno gradivo, videoposnetke, prosojnice, stripe, preverjanje znanja s kvizi.

Avtorji:

Dr. Tanja Zdošek Draksler, Ana Fabjan, Dr. Alenka Guček, Dr. Matej Kovačič

Financiranje:

To delo je bilo opravljeno v okviru evropskega projekta **AI4Gov: Trusted AI for Transparent Public Governance fostering Democratic Values** (ID sporazuma o nepovratnih sredstvih: 101094905).



AI4Gov
Trusted AI for Transparent Public Governance
fostering Democratic Values

**ZAUPANJA VREDNA
IN DEMOKRATIČNA
UMETNA INTELIGENCA
OSNOVE**

Povezava do tečaja



Figure 15: Graphics presented in a banner and a screen capture of the third learning course



Figure 16: Certificate of completion and digital badge for the third learning course

The learning analytics for this learning course is 5 registered users and 474 guest learners.

3.4 Learning course 4: “Trustworthy and democratic AI - creating awareness and change”

The fourth AI4Gov learning course titled “Trustworthy and democratic AI - creating awareness and change” was developed in the second half of 2024 and made publicly available end of March 2025. The structure of the learning course is as follows:

Module 1: Bias in AI

Lesson 1.1: Why AI Bias Matters and How It Affects Us

Lesson 1.2: Where Bias Begins: Human Bias

Lesson 1.3: When AI Takes a Wrong Turn: Algorithmic Bias

Lesson 1.4: The Big Picture

D5.4 Assessment tools, training activities, best practice guide V2

Quiz 1: Bias in AI

Module 2: Real-Life Examples of Bias

Lesson 2.1: Examples of Bias in AI Development Stages

Lesson 2.2: Bias in Criminal Law - The Case of COMPAS

Lesson 2.3: Bias in the Healthcare System - The Case of Predictive Algorithms

Lesson 2.4: Bias in Hiring Algorithms - The Case of Amazon and Beyond

Lesson 2.5: Bias in Government Fraud Detection Systems

Lesson 2.6: Key Takeaways and Warnings

Lesson 2.7: Poisoning ML Systems with Non-Randomness

Quiz 2: Real-Life Examples of Bias

Module 3: AI Bias Mitigation

Lesson 3.1: Methods to Mitigate Bias

Lesson 3.2: Proactive Accounting for Bias

Lesson 3.3: Recommendations to Raise Citizen Awareness about AI Bias and Ethic

Module 4: Explainable AI (XAI)

Lesson 4.1: Introduction to Explainable AI

Lesson 4.2: Why Explainability Matters in AI

Lesson 4.3: XAI Tools and Methods

Lesson 4.4: Practical Steps for Implementing XAI

Lesson 4.5: Challenges and Future Directions in XAI Implementation

Lesson 4.6: Situation Aware explainability (SAX)

Lesson 4.7: Consolidating XAI Knowledge - key insights

Quiz 4: Explainable AI (XAI)

Module 5: Virtualized Unbiasing Framework

Lesson 5.1: Scrollytelling application

Lesson 5.2: Stages of AI trainings

Lesson 5.3: Real life examples

Lesson 5.4: Catalog of bias detection and mitigation strategies

Module 6: Understanding AI fairness

Lesson 6.1: AI fairness and privacy

Resources & Authorship

Glossary

Acknowledgements

Figure 17 shows the visual graphic for this specific learning course as a screen capture from the learning platform. Figure 18 shows the digital badge and the certificate of completion.

Course

Trustworthy and Democratic AI - Creating Awareness and Change

You're on this course ✓



Welcome

Join us as we explore trustworthy and democratic AI. This course is designed for learners of all backgrounds. Whether you're a tech enthusiast, a business professional, in public administration, or simply curious about the forces shaping our digital landscape, this learning course will empower you with the knowledge to navigate the AI-driven world with confidence.

Get ready to unlock the potential of AI while championing the values of trust and democracy. Let's embark on this transformative learning experience together!

Training materials include: text, video lectures, presentation slides, comic books, quizzes.

Authors:

Tanja Zdolšek Draksler, PhD, Ana Fabjan, Alenka Guček, PhD, Matej Kovačič, PhD

Funding:

This work was done within the Horizon Europe project **AI4Gov: Trusted AI for Transparent Public Governance fostering Democratic Values** (Grant agreement ID: 101094905).

Figure 17: Screen capture from the learning platform for the learning course four



Figure 18: Certificate of completion and digital badge for the learning course four

The learning analytics for this learning course cannot be reported yet, as the learning course was made publicly available on the day of the deliverable submission.

3.5 Learning course 5: “Trustworthy and Democratic AI experts”

The last, fifth AI4Gov learning course titled “Trustworthy and democratic AI experts” is an extra learning course that is being developed and is planned to be publicly available in May/June 2025. This learning course is specific, highlighting the end results of the project AI4Gov that can be used as tools. As it is in development, the results can not be showed yet, but the graphics and the structure will follow the previous learning courses.

3.6 Internal focus group and OpenLearn Create quality assurance (QA)

In the process of the learning course development an important aspect was the internal focus group (provided from partner MT and managed from partner JSI), that provided first hand QA. The focus group was formed during the first learning courses development and consisted of five members from partner MT, which remained active till the end of the task T5.3 and evaluated all developed learning courses in English language. The main idea was to evaluate the developed courses to collect feedback for possible corrections, updates etc. Each member of the focus group has registered (created account) in the OpenLearn Create platform as team member with the goal to assess modules and lessons within the learning courses not being public yet and give feedback. After the revision by the focus group, each member completed the assessment list in the form of a prepared online questionnaire, with feedback and suggestions for possible corrections.

Figure 19 shows the screen captures of a part of the online questionnaire for the internal focus group for learning courses evaluation and feedback collection.

Six different areas were evaluated, all together forming 28 questions: (1) course expectations (3 questions), (2) course structure and content (13 questions), (3) quizzing (5 questions), (4) timing (2 questions), (5) e-Learning pace and navigation (4 questions), (6) visual design (1 question).

Overall, the focus group found the developed learning courses very enlightening. Especially the first learning course was of great importance to them, as it included valuable lessons for a beginner in AI fundamentals with enriched learning and training materials (video lectures, slides, comics, quizzes). Additionally, the focus group shared that the first learning course gives a good perspective in understanding bias and ethics in AI and therefore forms great value for understanding and following the next learning courses (especially learning course 4 and 5, which are more advanced to follow and understand for basic, non-tech users).

6. Rate the quality of the examples presented in the e-learning.

1

2

3

4

5

Very poor

☐

☐

☐

☐

☐

Very good

7. Rate your enjoyment of the e-learning course.

1

2

3

4

5

Very poor

☐

☐

☐

☐

☐

Very good

8. What part of the e-learning course did you find most useful and interesting? (You can state the modul number or particular lessons)

Vaš odgovor

9. Rate the course workload.

1

2

3

4

5

Very low

☐

☐

☐

☐

☐

Excessive

10. What are the strengths and weaknesses of this e-learning course?

Vaš odgovor

11. What part of the e-learning course did you find most useful and interesting?

Vaš odgovor

12. Was the content arranged in a clear and logical way? Why or why not (if applicable)?

Vaš odgovor

13. Did the content adequately explain the knowledge, skills, and concepts it presented?

Vaš odgovor

Quizzing

Rating scale used:

1-Very poor

2-Poor

3-Acceptable

4-Good

5-Very good

1. Rate the relevance of quizzes and assignments.

1

2

3

4

5

Very poor

☐

☐

☐

☐

☐

Very good

2. Rate the quality of the questions asked in the quizzes.

1

2

3

4

5

Very poor

☐

☐

☐

☐

☐

Very good

3. Were the quizzes presented at adequate intervals?

☐ Yes

☐ No

☐ I don't know

☐ Can not decide if yes or no-something in between

4. Did the quizzes cover and test the material presented in the course?

☐ Yes

☐ No

☐ I don't know

☐ Can not decide if yes or no-something in between

5. Do you have any comments or suggestions on the quizzes part?

Vaš odgovor

Figure 19: Screen captures from the online questionnaire for the internal focus group for learning courses

Additional to the QA from the internal focus group, the learning courses underwent also strict review process from the OpenLearn Create platform managers. As part of the review process the following items are checked:

- Metadata is complete and correct.
- Learning outcomes and course summary are present.

- Configuration of the following are correct and accessible (to meet WCAG 2.2 guidelines²):
 - ✓ Content structure and navigation
 - ✓ Heading level structure and unique heading titling
 - ✓ Figures and diagrams (figure numbers and descriptive captions are present)
 - ✓ Alt text (included for all images and is meaningful)
 - ✓ Videos (close captions (subtitles), transcripts and audio descriptions have been added)
 - ✓ Audio files (transcripts have been included)
 - ✓ Tables and table headings have been set up correctly
 - ✓ Web links (they are complete and error free)
 - ✓ Images have attributions (source/licence information)
 - ✓ Acknowledgments have been added (including third-party IP acknowledgements)
 - ✓ Other Moodle activities (functionality and settings)
 - ✓ Quizzes (including grade boundaries, question feedback and completion settings)
 - ✓ Completion settings
 - ✓ Badges
 - ✓ Statement of participation/custom certificate

On completion of the OpenLearn Create review process, possible identified issues were communicated with JSI as the learning courses developer and manager. These issues were needed to be addressed and recommended amends included before the content could be published.

3.7 Further dissemination and exploitation of learning courses and educational resources

Apart from the general specification of stakeholders listed in Chapter 2.1, the AI4Gov partners are constantly searching for new learning audience and collaboration opportunities in relation to WP5. Partner JSI is actively talking with several organisations to exploit the AI4Gov learning and training materials and learning courses (MOOCs). This exploitation could be twofold. First, exploitation of existing materials and services as they are. Second, adaptation of the materials and services to make them serve the educational needs of specific groups, for example pupils, students, teachers, regulators, etc.

Possible exploitation of the AI4Gov learning courses and education materials is being discussed with several stakeholders from Slovenia:

- The Ministry of public administration of Slovenia within their learning and training platform for public administration and public sector employees titled Administration academy (“Upravna akademija³” in Slovene language).

² Web Content Accessibility Guidelines (WCAG) international standard: <https://www.w3.org/WAI/standards-guidelines/wcag/>

³ <https://ua.gov.si/>

- SAFE.SI⁴, awareness-raising point on safe use of the internet and mobile devices for children, teenagers, parents and teachers - Safer Internet Centre Slovenia.
- The Academic and Research Network of Slovenia (ARNES)⁵, a public institute that provides network services to research, educational and cultural organizations.

However, due to operational and staff costs, there are no defined and concrete plans of using or re-using the content yet.

However, the first learning course is already included as a resource at the Digital Skills and Jobs Platform (managed from the European Commission) within the training opportunities section⁶. In addition, the explanatory video resource developed within the project's scope on the AI Act, was also highlighted⁷. Partner JSI is already in communication with the platform administration to add other learning courses.

The consortium is being open to further collaboration and exploitation opportunities in relation to learning and training. Possible activities and achievements in this relation will be reported within WP7 deliverables. We plan to still focus on the learning collection dissemination also after WP5 end in close collaboration with WP7 and with the aim to increase the numbers of registered and unregistered learners (guest visitors) of the learning courses.

Also, the learning courses will be continuously used in the scope of WP6. First, the learning courses will be promoted through the pilot's 2nd iteration of the use cases to reach a broader audience, as these will be attended from external audience. Second, partner JSI will together with partner VIL as the WP6 leader, analyse the 1st iteration results and specifically point out the expressed pilot's needs on one side (WP6) and connect it to specific learning resource or part of it (e.g., module or lesson or video lecture). For example, if the expressed concerns or comments are tackling AI regulation, we could direct the users to watch the video about the AI Act and focus on Module 4 of learning course 1. In that way, the people's concerns would be directly addressed.

In addition, the OpenLearn Create platform offers to download or print the learning courses in a form of an e-book (enabling printing specific modules or lessons as chapters). Although this feature is possible only for registered users of the platform, it still offers a valuable option for further dissemination and exploitation of the developed learning courses. Figure 20 is showing a screen capture of the book option of the learning course.

⁴ <https://safe.si/>

⁵ <https://www.arnes.si/en/>

⁶ <https://digital-skills-jobs.europa.eu/en/opportunities/training/trustworthy-and-democratic-ai-fundamentals-ai4gov-mooc>

⁷ <https://digital-skills-jobs.europa.eu/en/inspiration/resources/introduction-ai-act-ai4gov-video-resource>

Module 3: Data and Bias

Print book

Site:	OpenLearn Create
Course:	Trustworthy and Democratic AI - Fundamentals
Book:	Module 3: Data and Bias

Printed by:	████████
Date:	Tuesday, 25 March 2025, 8:53 AM

Description

Welcome to the module "Data and Bias." We will explore the crucial interconnection between data and bias, shedding light on how the information we collect can inadvertently introduce biases into various processes. As data increasingly shapes decision-making in the realms of artificial intelligence and technology, it becomes imperative to understand the nuances of bias within datasets. Join us as we unravel the complexities of this interplay, examining real-world examples and strategies to mitigate biases, ensuring a more accurate and equitable use of data in diverse applications.

In Module 3, we cover the following Lessons:

[Lesson 3.1: Bias in Data Collection](#)

[Lesson 3.2: Data Sampling Methods](#)

[Lesson 3.3: Ethical Data Sourcing](#)

[Lesson 3.4: Data Pre-processing and Bias Reduction](#)

[Lesson 3.5: Real-world Data Bias Case Studies](#)

Table of contents

[Lesson 3.1: Bias in Data Collection](#)
[Lesson 3.2: Data Sampling Methods](#)
[Lesson 3.3: Ethical Data Sourcing](#)
[Lesson 3.4: Data Pre-processing and Bias Reduction](#)
[Lesson 3.5: Real-world Data Bias Case Studies](#)

LESSON 3.1: BIAS IN DATA COLLECTION

In Lesson 3.1, we delve into the foundations of bias in data collection. Understanding that biases can be unintentionally embedded during the data gathering process is crucial. We will explore how factors such as sampling methods, data sources, and the context of collection can influence the presence of bias. By comprehending these fundamental aspects, we aim to equip you with the knowledge needed to identify and address biases at the source, fostering more reliable and unbiased datasets.

Bias in data collection refers to the systematic errors or inaccuracies introduced during the process of gathering and recording data. These errors can arise from various sources and can lead to a skewed or unrepresentative dataset. Bias in data collection can significantly impact the reliability and validity of the information obtained, influencing subsequent analyses, decisions, and outcomes. There are several ways bias can manifest in data collection:

- **Sampling Bias:** This occurs when the sample selected for data collection is not representative of the entire population. It may exclude certain groups or over-represent others, leading to a distorted view of the overall population.
- **Selection Bias:** Arises when the criteria used to select participants or data points favor a particular group, leading to a non-random and potentially unrepresentative sample.
- **Measurement Bias:** Occurs when the tools or methods used for data collection are flawed or systematically favor certain outcomes. This can include issues like poorly designed survey questions or inaccurate measurement instruments.
- **Observer Bias:** Results from the personal beliefs, expectations, or preconceived notions of the individuals collecting the data. This can influence how data is recorded, leading to unintentional distortions.
- **Cultural or Contextual Bias:** Arises from the cultural or contextual factors present during data collection. Different cultural backgrounds or contextual elements may impact responses or interpretations.

Recognizing and addressing bias in data collection is crucial to ensure the integrity of the collected data and to prevent downstream effects on analyses and decision-making processes. Strategies for mitigating bias include employing diverse and representative samples, using standardized measurement tools, providing clear instructions to data collectors, and applying ethical considerations throughout the data collection process.

Figure 20: Screen capture of the book option of the learning course

Next chapter introduces a set of (self) assessment tools and checklists on ethical and transparent AI (work done in task 5.4).

4 Development of (self)assessment tools on ethical and transparent AI

As governments increasingly deploy AI-driven solutions in public services, it is crucial to establish evaluation mechanisms that help developers assess the reliability, accountability, and social impact of these systems. Self-assessment tools such as questionnaires and checklists, alongside with best practices not only facilitate adherence to regulations such as the EU AI Act and GDPR but also promote public trust by ensuring that AI solutions are fair, explainable, and aligned with human rights principles. By integrating automated human-led assessment frameworks, responsible AI development can be fostered, mitigating risks, and enhancing the quality of AI-driven citizen services.

4.1 Key considerations of self-assessment tools regarding the development of responsible AI algorithms

The development and deployment of AI systems require robust mechanisms to ensure ethical integrity, transparency, and compliance with regulatory standards. As AI becomes more integrated into decision-making processes in government entities, the need for systematic assessment tools has grown. Self-assessment tools, often structured as questionnaires, provide a framework to evaluate the ethical and transparent nature of AI systems. This section explores the foundational considerations necessary for developing such tools, aiming at the development of AI systems provided by government entities, while focusing on key ethical principles, regulatory requirements, and technical accountability as captured in the Holistic Regulatory Framework (HRF) (Figure 21).

The following key considerations should be included in the design of the self(assessment) tools:

- **Clarity and Accessibility:** The language used should be clear and simple to ensure AI system developers fully understand the questions. Avoid technical jargon where possible and use well-defined terminology. Moreover, digital accessibility mechanisms such as screen reader compatibility, keyboard navigation, and color contrast adjustments must be incorporated to enable effective use.
- **Usability and User Experience (UX):** The tool should be optimized for seamless interaction with an intuitive interface, logical question flow, and minimal cognitive load. Fast response times should be ensured to prevent delays in assessment completion. Features such as progress indicators, autosave functionality, and a structured summary of results should be included. Additionally, a mobile-responsive design is crucial to accommodate developers accessing the tool from various devices.
- **Data Privacy and Security:** The self-assessment tool should ensure that AI systems comply with data protection laws (e.g., GDPR, CCPA) by evaluating key security measures implemented by the system. The questionnaire should assess whether the AI system follows secure data handling practices such as encryption of sensitive information, data minimization techniques, and adherence to retention policies. Questions should cover whether personal data is pseudonymized or anonymized where applicable, how access

controls (e.g., role-based access control) are enforced, and whether AI model outputs are protected from leaking sensitive information. Furthermore, the assessment should include checks on whether data governance frameworks are in place to manage security risks throughout the AI system lifecycle.

- **Actionable Insights:** The tool should not only assess compliance but also generate general recommendations. These insights should be data-driven, context-specific, and linked to relevant regulatory guidelines. Automated reporting mechanisms should provide developers with downloadable assessments, highlighting areas requiring improvement and offering best practice guidelines.
- **Continuous Improvement:** Implement built-in feedback loops allowing AI system developers to suggest enhancements or report usability issues. The tool should incorporate analytics to identify trends in responses and refine assessment criteria dynamically. Regular updates must be conducted to align with evolving regulatory policies, technological advancements, and ethical considerations in AI governance.
- **Interoperability & Integration:** The assessment tool must be designed for seamless integration with AI lifecycle management platforms, compliance monitoring tools, and risk assessment frameworks. Support for standard data formats such as JSON, XML, and CSV ensures compatibility with existing governance systems. APIs should be provided to facilitate automated data exchange between the self-assessment tool and internal compliance dashboards.
- **Transparency & Explainability of Results:** The scoring mechanism should be fully transparent, detailing how responses impact final compliance scores. Where AI-driven scoring is involved, models should incorporate explainability techniques such as SHAP (Shapley Additive Explanations) or LIME (Local Interpretable Model-Agnostic Explanations) to make results interpretable. Each compliance category should provide justifications for scores along with recommended next steps.
- **Risk Assessment & Prioritization:** Instead of providing a binary compliance status, the tool should implement a severity-weighted scoring model to highlight high-risk areas. Risks should be categorized based on impact (e.g., legal, ethical, operational) and likelihood of occurrence. A visual representation, such as a heat map or risk radar chart, can be used to help development teams prioritize corrective actions efficiently.

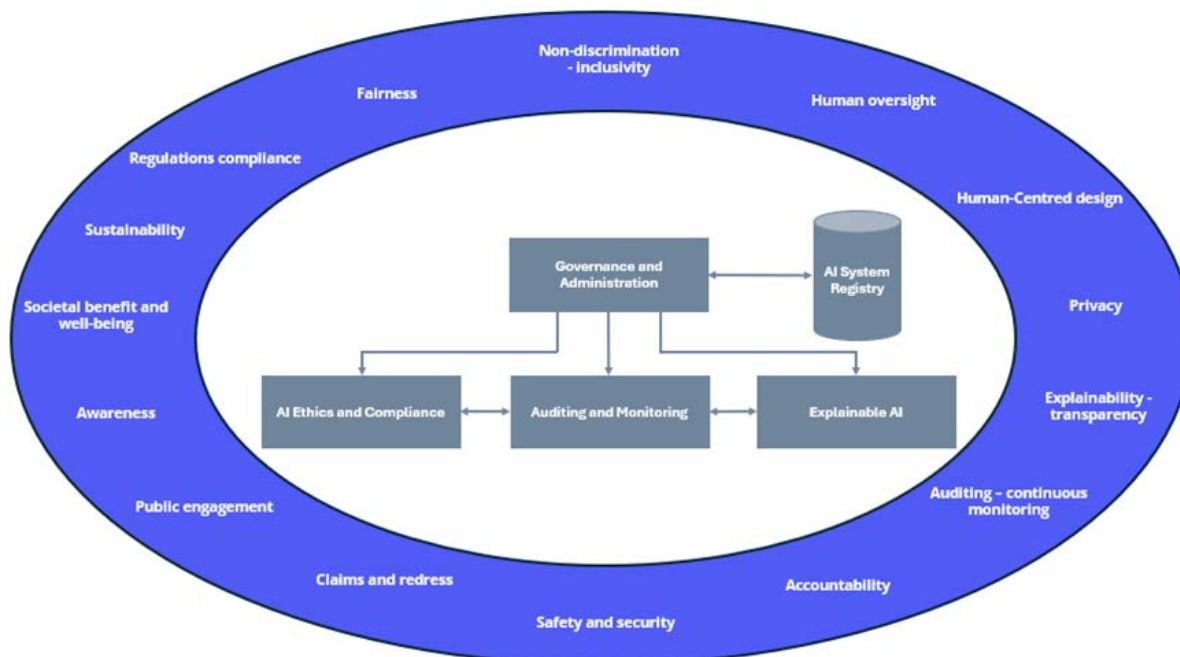


Figure 21: Holistic Regulatory Framework

4.2 Design and guidelines for implementation of assessment tools based on HRF compliant questionnaires

The HRF defines 15 regulatory dimensions that are essential for governing AI systems deployed in public sector services. These dimensions cover critical aspects such as legal compliance, fairness, transparency, accountability, risk management, and security, ensuring AI systems meet ethical, technical, and regulatory requirements.

To facilitate structured compliance evaluation, self-assessment tools serve as a systematic auditing mechanism, enabling AI system developers to perform a quantitative and qualitative analysis of their models against predefined governance benchmarks. These tools incorporate policy-driven checklists, regulatory alignment frameworks, and scoring models to measure the extent of an AI system's compliance with HRF requirements.

The self-assessment questionnaire is designed to support multi-stage AI lifecycle evaluation, covering design, data processing, model training, deployment, and continuous monitoring. Each of the 15 regulatory dimensions is broken down into key assessment criteria, with responses recorded using a structured Likert scale (1-7) and binary compliance indicators (Yes/No).

To enhance interpretability, explanatory justifications and regulatory mapping are embedded within the assessment tool, allowing developers to trace AI system compliance to specific legal, ethical, and technical principles. The scoring mechanism computes category-specific compliance scores, aggregates risk levels, and highlights priority areas requiring remediation. The tool

integrates severity-weighted risk analysis, ensuring that AI development teams can quantify, prioritize, and mitigate compliance risks based on their impact and probability of occurrence.

Ultimately, this self-assessment framework functions as a proactive AI governance mechanism, equipping development teams with data-driven insights, automated regulatory tracking, and structured compliance validation. The tool fosters transparency, accountability, and risk-aware AI deployment in government applications while ensuring alignment with HRF's ethical and legal standards.

1. Regulations compliance

Legal Framework Adherence: *To what extent does the AI system comply with applicable legal frameworks such as the EU AI Act and the EU Charter of Fundamental Rights? (1: No compliance strategy, 7: Fully compliant with automated legal validation and audit logs)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Sector-Specific Compliance: *Does the AI system adhere to domain-specific regulations (e.g., GDPR for healthcare, Basel III for finance)?*

☐ YES ☐ NO

Regulatory Tracking & Adaptation: *How effectively does the system track, assess, and adapt to evolving AI-related laws and ethical guidelines? (1: No tracking, 4: Manual compliance checks, 7: Automated compliance monitoring with regulatory updates)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Compliance Responsibility: *Is there a designated responsible entity for ensuring continuous legal compliance throughout the AI lifecycle?*

☐ YES ☐ NO

Degree of AI-related policies documentation: *Are there documented policies in place to track, assess, and adapt to evolving AI-related laws, standards, and ethical guidelines? (1: no documentation, 7: complete documentation of all AI-related policies)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

User Rights Awareness: *Are users informed about their rights, obligations, and potential risks associated with the AI system's deployment? (1: No user awareness program, 7: Comprehensive user rights education with interactive guidelines)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

2. Fairness

Bias Mitigation in Data & Models: *To what extent are bias detection and mitigation strategies implemented in the AI system? (1: No bias checks, 7: Continuous bias monitoring with automated correction mechanisms)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Fairness Metrics: *To what extent are fairness metrics (e.g., disparate impact ratio, equalized odds, demographic parity) defined, monitored, and regularly assessed? (1: No fairness metrics defined, 4: Basic fairness metrics monitored periodically, 7: Continuous real-time fairness auditing with automated reporting and corrective actions)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Disparity Correction Measures: *When disparities are identified in fairness metrics, how effectively are they addressed? (1: No correction measures, 7: Fully automated fairness-aware retraining pipeline)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

3. Non-discrimination - inclusivity

Equal Treatment Across Demographics: *Does the system ensure equal treatment for different demographic and social groups in predictions and decision-making? (1: High disparity, 7: Statistically validated fairness across groups)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Stakeholder Diversity in AI Design: *Are diverse stakeholder perspectives incorporated into AI design and implementation?*

☐ YES ☐ NO

Bias Review in Training Data: *To what extent has the training data been reviewed for representation biases? (1: No review, 7: Continuous bias auditing and synthetic data augmentation)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Prevention of Discriminatory Outcomes: *What safeguards exist to prevent AI from producing discriminatory outcomes against underrepresented groups? (1: No safeguards, 7: Built-in bias testing with active mitigation mechanisms)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

4. Human oversight

Intervention in AI Decisions: *Does the AI system allow human intervention in case of incorrect or harmful outputs?*

☐ YES ☐ NO

Levels of Human Oversight: *Which level of human oversight is available in the system (e.g., human-in-the-loop, human-on-the-loop, or human-in-command)? (1: No oversight, 7: Fully adaptive oversight with dynamic thresholds)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Operator Role Definition: *To what extent are responsibilities and roles of human operators clearly defined and integrated into AI decision-making workflows? (1: No defined roles, 4: Basic operator roles outlined but inconsistently applied, 7: Fully defined roles with clear accountability, training, and operational integration)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

5. Human-centered design

Usability & Accessibility: *Has the AI system been designed with a focus on usability and accessibility for diverse user groups? (1: No accessibility, 7: Full WCAG and universal design compliance)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

User Feedback Mechanisms: *How effectively does the system incorporate user feedback throughout its lifecycle? (1: No feedback integration, 7: Real-time adaptive learning from user feedback)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Human Role Augmentation: *Does the system empower users rather than replacing essential human roles and responsibilities? (1: High automation replacing humans, 7: Human-augmentation with AI-assisted decision support)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

6. Privacy

Data Collection: To what extent does the AI system collect only the necessary personal data for its operation? (1: Collects excessive data beyond operational needs, 7: Implements strict data minimization techniques such as differential privacy and federated learning, ensuring only essential data is processed)

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

User Consent: How effectively does the system inform users and obtain their explicit consent for data collection and processing? (1: No information/consent, 7: Full and clear information with explicit consent)

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Data Security: How strong are the security measures in place to protect personal data from unauthorized access? (1: No formal security protocols in place, 4: Implements standard security measures like role-based access control and basic encryption, 7: Implements multi-layered security protocols including real-time intrusion detection, adversarial attack prevention, and zero-trust architecture with continuous monitoring)

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Data Management: Are there procedures in place to delete or anonymize personal data once it is no longer needed?

☐ YES ☐ NO

7. Explainability - transparency

Decision Comprehensibility: *To what extent can users understand the decisions or predictions made by the AI system? (1: Completely opaque, 7: Fully understandable and well-explained)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Transparency Measures: *How well does the AI system document and communicate its decision-making processes to regulators and stakeholders? (1: No documentation, 4: Internal documentation with limited accessibility, 7: Comprehensive and publicly available documentation including detailed model decision-making processes and compliance reports)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Auditability of Model Decisions: *Can AI decisions be fully traced back to the dataset, algorithmic weights, and hyperparameters that influenced them? (1: No traceability, 5: Partial traceability with basic logging of decisions, 7: Full version-controlled traceability, including dataset lineage, model weights, hyperparameter configurations, and decision justification)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

8. Auditing – continuous monitoring

Automated Performance Audits: *How frequently are AI model drift and fairness metrics (e.g., population stability index (PSI), KL divergence, or Wasserstein distance) recalculated? (1: Never, 3: Every 6 months, 5: Monthly, 7: Continuously)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Bias & Fairness Logging: Does the system maintain logs of bias-detection metrics for every inference request (e.g., disparate impact ratio, equalized odds, demographic parity, or Theil index)?

☐ YES ☐ NO

Self-Correction Mechanisms: *Does the AI system have a mechanism for self-adapting to detected bias without manual intervention? (1: No self-correction, 7: Fully automated self-correction with compliance tracking)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

9. Accountability

Version Control for AI Pipelines: *Does the AI system store previous model versions and datasets for rollback?*

Human Oversight in AI Decisions: *Are there thresholds where human operators must approve AI-generated decisions? (1: No human intervention, 7: Full human-in-the-loop governance)*

Compliance Logging: *Are compliance checks automatically logged when AI models are updated?*

10. Safety and security

Adversarial Attack Resistance: *Has the AI system undergone penetration testing against adversarial attacks (FGSM, PGD, DeepFool)? (1: No testing, 7: Fully adversarially trained)*

Runtime Security Monitoring: *Does the AI system detect and respond to unauthorized API access or manipulation? (1: No monitoring, 7: Fully autonomous anomaly detection)*

Fail-Safe Mechanisms: *Can the AI system shut down or switch to a backup model in case of failure?*

11. Claims and redress

Explainability of AI Rejections: *Can users request a technical breakdown of why an AI-driven decision was rejected?*

Automatic Redress Predictions: Does the AI system proactively analyze and flag cases with high potential for misclassification or unfair outcomes, ensuring users are informed before submission? (1: No flagging of potential errors or unfair decisions, 4: Basic flagging based on predefined thresholds, 7: Fully predictive redress system utilizing anomaly detection, uncertainty estimation, and explainable AI to flag and suggest corrections before submission)

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

User Redress Accessibility: To what extent does the AI system provide a clear and accessible mechanism for users to report and appeal AI-driven decisions? **Scale 1-7** (1: No formal appeal mechanism, 4: Manual review process upon request, 7: Fully structured redress system with automated tracking and resolution)

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Grievance Resolution Process: Does the AI system have a documented and enforceable process for handling complaints and addressing harmful AI outcomes?

☐ YES ☐ NO

12. Public engagement

Community Feedback Integration: *Are public concerns and societal expectations considered in the development and deployment of the AI system? (1: Never, 7: Continuously)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Public Model Documentation: *Are model specifications and biases published in an open-access format?*

☐ YES ☐ NO

Public Input Impact on Governance: To what extent does public feedback actively shape AI governance policies, ethical frameworks, and operational guidelines? Scale 1-7 (1: No influence, 4: Feedback collected but rarely implemented, 7: Direct integration into policy updates and system design improvements)

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

13. Awareness

User Training & Literacy: *Are users provided with interactive explainers or tutorials on AI operations? (1: No training, 7: Extensive interactive literacy programs)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Misinformation Mitigation: *Are awareness campaigns conducted to counter AI-related misinformation? (1: No efforts, 7: Comprehensive public education strategies)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Stakeholder AI Literacy: *Is there a structured plan to improve AI literacy among different stakeholder groups?*

☐ YES ☐ NO

14. Societal benefit and well-being

Measurable Societal Impact: *Does the system track and quantify its contributions to inclusivity and accessibility through defined metrics such as demographic parity, fairness scores, and accessibility compliance (e.g., WCAG standards)? (1: No tracking, 4: Basic manual tracking with periodic reviews but no automated reporting, 7: Full real-time reporting with automated insights and compliance monitoring)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Ethical Risk Assessment: *Has the potential societal impact—both positive and negative—been assessed before deployment? (1: No assessment, 7: Comprehensive risk analysis and mitigation planning)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Protection of Vulnerable Populations: *Are safeguards in place to prevent AI from causing harm to marginalized or vulnerable communities?*

☐ YES ☐ NO

15. Sustainability

Environmental Impact Assessment: *Has the environmental impact of the AI system (e.g., energy consumption, data center emissions) been evaluated? (1: No assessment, 7: Full lifecycle carbon footprint analysis with mitigation strategies)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Energy Efficiency Measures: *Are sustainable computing practices, such as energy-efficient algorithms or carbon offset measures, implemented? (1: No efficiency measures, 7: Fully optimized for energy efficiency with hardware-aware AI training)*

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7

Long-Term Sustainability Plan: *Is there a structured plan for ethical disposal, model decommissioning, and long-term environmental responsibility?*

☐ YES ☐ NO

The above questionnaire is available online at:
<https://docs.google.com/forms/d/1wBs3SKhwCOZBTZLr3JsJaob6BV7pBYvwSAJVaYhQTC4>.

4.3 Scoring mechanisms for the assessment tools

The scoring mechanism for a self(assessment) tool should provide meaningful, actionable insights and it should align with the questionnaire's objectives, ensuring that results are both interpretable and useful for decision-making. The scoring mechanism for the self-assessment tool is designed to provide quantifiable, interpretable, and capable of providing actionable insights regarding the compliance of AI systems with the HRF. The methodology ensures that Likert scale (1-7) responses and Yes/No responses are standardized into a unified scoring system, allowing for a fair and balanced evaluation across all 15 HRF dimensions. By establishing a structured and transparent scoring methodology, this assessment framework enables AI developers and regulatory teams to systematically evaluate AI systems, identify potential compliance gaps, and implement necessary improvements.

To ensure consistency in scoring, responses are normalized to a common 0-1 range. Yes/No questions are inherently binary, where 1 represents compliance (Yes) and 0 represents non-compliance (No). Meanwhile, Likert scale questions, which use a 1-7 range, are transformed into the 0-1 range using the following normalization function:

$$\text{Normalized Likert Score} = \frac{\text{Likert response} - 1}{6}$$

This transformation ensures that a score of 1 (representing non-compliance) maps to 0, and a score of 7 (representing full compliance) maps to 1, making Likert responses directly comparable to Yes/No responses in the final calculations.

Each of the 15 HRF dimensions is assessed based on its respective set of questions. The category score for each HRF dimension is computed as follows:

$$\text{Category Score} = \frac{\sum(\text{all normalized question scores})}{\text{Number of questions in category}}$$

where a score of 0 represents a completely non-compliant AI system, and a score of 1 represents a fully compliant AI system. For improved interpretability, this score is scaled to a percentage (0-100%), making it more accessible for developers and decision-makers. To facilitate decision-making, the scoring results are categorized into compliance levels:

- 0-40% → Non-compliant (High Risk)
- 40-70% → Partial Compliance (Needs Improvement)
- 70-100% → Compliant (Low Risk)

These categories provide a structured way to assess compliance risks, helping AI system developers and regulatory bodies prioritize areas requiring improvement.

To enhance interpretability, assessment results are visualized using bar charts. Bar charts are used to illustrate compliance strengths and weaknesses across the 15 HRF dimensions, with each bar representing a specific compliance category. To provide further context, each category's bar chart will include an indicator for the recommended compliance value as well as the mean compliance score of other evaluated AI systems, allowing for a benchmarking comparison. This enables developers to assess their system's relative standing and identify gaps that require improvement. Additionally, heatmaps highlight high-risk areas, helping AI system developers prioritize critical compliance issues that demand immediate action.

By implementing this structured scoring approach, the self-assessment tool provides a reliable and actionable compliance evaluation framework, supporting AI developers in improving regulatory alignment, mitigating risks, and enhancing the overall transparency and accountability of AI systems. The integration of normalized scoring, categorical compliance evaluation, and visual analytics ensures that stakeholders can effectively assess AI governance practices and prioritize necessary compliance measures.

4.4 Best practice guide

The development process of any AI system that aims to be used by citizens while provided by government entities should be conceptualized, designed, implemented and deployed under continuous scrutiny regarding ethical aspects related with its prospected operation.

Structured in categories aligned with the HRF, a set of best practices applicable to the development of such systems follows:

1. Regulations Compliance

- Ensure adherence to national, European and international AI-related regulations.
- Keep the corpus of relevant regulations regularly updated.
- Conduct legal reviews at all development stages.

2. Fairness

- Use unbiased datasets and continuously test for bias.
- Employ fairness-aware algorithms and techniques to mitigate discrimination.
- Conduct diverse stakeholder reviews to ensure fair outcomes.

3. Non-Discrimination & Inclusivity

- Design AI systems that are inclusive and accessible to all societal groups.
- Perform impact assessments on marginalized communities.
- Engage diverse user groups during system testing.

4. Human Oversight

- Ensure human-in-the-loop mechanisms for critical decision-making.
- Define clear escalation paths for AI-generated outcomes.
- Provide training for personnel interacting with AI systems.

5. Human-Centered Design

- Prioritize user experience and ethical considerations in AI system design.
- Conduct iterative user testing and feedback loops.
- Align AI goals with human values and societal needs.

6. Privacy

- Implement privacy-by-design principles.
- Use anonymization and differential privacy techniques.
- Comply with GDPR.

7. Explainability & Transparency

- Develop inherently interpretable models where possible.
- Provide documentation on AI decision-making processes.
- Ensure users can understand AI-generated outcomes.

8. Auditing & Continuous Monitoring

- Establish independent auditing mechanisms.
- Monitor AI performance and behavior over time.
- Regularly update AI models based on new data and ethical considerations.

9. Accountability

- Clearly define roles and responsibilities within AI development teams.
- Implement mechanisms for tracking decision-making processes.
- Assign liability in case of AI-related failures.

10. Safety & Security

- Follow robust cybersecurity protocols.
- Use adversarial testing to identify vulnerabilities.
- Implement fail-safe mechanisms to prevent harm.

11. Claims & Redress

- Provide clear processes for users to challenge AI decisions.
- Establish mechanisms for correcting erroneous decisions.
- Offer legal recourse options when necessary.

12. Public Engagement

- Involve citizens in AI policy discussions.
- Organize public consultations and transparency initiatives.
- Communicate AI system capabilities and limitations clearly.

13. Awareness

- Implement AI risk-awareness training for developers.
- Launch awareness campaigns to inform the public.
- Develop documentation and explainers on AI limitations.

14. Societal Benefit & Well-being

- Design AI to improve social welfare.
- Prioritize ethical use cases aligned with public good.
- Conduct ethical impact assessments before deployment.

15. Sustainability

- Optimize AI energy consumption.
- Consider environmental impact when selecting models and infrastructure.
- Promote AI solutions that contribute to sustainability goals.

AI systems provided by government entities require proactive governance that evolves with technological capabilities. By integrating these practices early in development cycles, public confidence can be built, while harnessing AI's potential for societal good. Regular updates are essential as both technology and public expectations advance.

5 Conclusions

This deliverable provides an overview of the work done in WP5, in more detail for T5.2, T5.3 and T5.4, reporting the work within the whole duration of WP5, but specifically focused on the period from project month 18 to 27, as the first period (month 3 to month 18) was already in detail reported in the first version of this deliverable, namely D5.3 Assessment tools, training activities, best practice guide v1. This deliverable covers the topics of AI4Gov trainings workshops and related educational resources covering T5.2 (Chapter 2), learning courses (MOOCs) covering T5.3 (Chapter 3) and (self)assessment tools on ethical and transparent AI covering T5.4 (Chapter 4).

In the context of the AI4Gov training workshops and educational resources, organized trainings were presented followed with presentation and explanation of educational resources. For the learning courses (MOOCs), the whole learning collection was presented in detail. In relation to the project plan for T5.2 and T5.3, all planned activities have been adequately addressed and carried out. To specify, for the training activities (T5.2), three training workshops were foreseen as KPIs (KPIs=3), at the end four were delivered reaching more than 300 participants from different domains and different countries actively (participants of the trainings) and even more passively (due to available educational resources from these trainings). For the learning courses (T5.3), more than two MOOCs were foreseen as KPIs (KPIs=3), at the end the learning collection comprises from five learning courses.

Both groups of KPIs correspond to the foreseen project results R13 and R14 (that were delivered in D5.3 and D5.4):

- R13: User-centric training materials: services offering multimedia content enabling user's training and assistance, related to project objectives O6⁸ and O7⁹,
- R14: Massive open-online courses (MOOCs): flexible training delivery methods on the AI and rule of law, related to project objective O6.

While the training materials (R13) were foreseen to start with start TRL (Technology readiness level)¹⁰: 3-4 and end with TRL: 4-5, the final outcome is even higher, being TRL: 6-8 (video lectures from the AI4Gov trainings are publicly available online with presentation slides and subtitles on the VideoLectures.NET platform and on the AI4Gov project YouTube channel for interested citizens). Moreover, the learning courses (MOOCs) were foreseen to start with TRL: 2-3 and end with TRL: 4-5, however, the final outcome is again higher, TRL levels being 6-8 (AI4Gov learning courses being publicly available online for interested citizens).

⁸ Objective 6 - Train stakeholders and educate citizens in the use of the platform and the even important elements of a democratic AI

⁹ Objective 7 - Cocreation, deployment and validation of legal and organizational blueprints of the HRF in various use cases

¹⁰ TRL2 - technology concept formulated, TRL3 - experimental proof of concept, TRL4 -technology validated in lab, TRL5 - technology validated in relevant environment, TRL6 - technology demonstrated in relevant environment, TRL7 - system prototype demonstration in operational environment, TRL8 - system complete and qualified

On the other hand, T5.4 delivers (self)assessment tools in the form of questionnaires/checklists in text and online form in English. Translations in the languages of the consortium will be finalised from task leader AUTH giving priority to the languages of the pilot cases. The (self)assessment tools are based on the 15 dimensions of the HRF and aim at assessing the degree of conformity with HRF, of the development process of AI systems developed by government entities. These tools aim at assisting developers of AI systems to reassure that the AI systems they develop follow best practices and are compliant with essential regulatory and ethical guidelines.

The progress of the targeted tasks is clearly shown, achieving all of the needed KPIs and even more, that is why it can be concluded, that the WP5 is positively concluded.

6 References

- AI4Gov Learning&training.* (2025). Retrieved from <https://ai4gov-project.eu/home/resources/training-learning/>
- AI4Gov YouTube channel.* (2025). Retrieved from <https://www.youtube.com/@AI4GovProject>
- VideoLectures.NET.* (2025). Retrieved from 1st AI4GOV Training Workshop: Bias In AI: https://videlectures.net/events/AI4GOVtraining2023_ljubljana
- Zdolšek Draksler, T., Fabjan, A., Guček, A., & Kovačič, M. (2024). *Trustworthy and Democratic AI - Fundamentals. Learning course.* Retrieved from <https://www.open.edu/openlearncreate/course/view.php?id=11669>
- Zdolšek Draksler, T., Fabjan, A., Guček, A., & Kovačič, M. (2025a). *Zaupanja vredna in demokratična umetna inteligenca - osnove. Learning course.* Retrieved from <https://www.open.edu/openlearncreate/course/view.php?id=12525>
- Zdolšek Draksler, T., Fabjan, A., Guček, A., & Kovačič, M. (2025b). *Inteligencia Artificial fiable y democrática - Fundamentos. Learning course.* Retrieved from <https://www.open.edu/openlearncreate/course/view.php?id=12055>
- Zdolšek Draksler, T., Fabjan, A., Guček, A., & Kovačič, M. (2025c). *Trustworthy and Democratic AI - Creating Awareness and Change (in development).* Retrieved from <https://www.open.edu/openlearncreate/course/index.php?categoryid=2051>
- Zdolšek Draksler, T., Fabjan, A., & Guček, A. (2025d). *Trustworthy and Democratic AI experts (in development).* Retrieved from <https://www.open.edu/openlearncreate/course/index.php?categoryid=2051>