



ESCO-based learners' profiles

AIOLIA DELIVERABLE 4.1

Project Name	AIOLIA
Deliverable Title/Number	D4.1 ESCO-based learners' profiles
Lead beneficiary	AUMC
Lead Authors	Natalie Evans, Menno Maris
Contractual delivery date:	31/01/2026
Actual delivery date:	31/01/2026
Sensitivity	PUBLIC

Document History

Name	Organisation	Role	Action	Date
V0.1	AUMC	Task 4.1 Lead	Draft available on SharePoint for task partners and project coordinator to comment	12.01.2025
V0.2	CEA	Project coordinator	Internal feedback	27.01.2026
V0.3	CERTH	WP4 Lead	Internal feedback	27.01.2026
V0.4	KIT	Reviewer	Quality check	27.01.2026
V0.6	AUMC	Lead	Incorporating CEA, CERTH, and KIT feedback	30.01.2026
V0.7	CEA	Project coordinator	Final check	30.01.2026
V1	AUMC	Lead	Final version for submission	31.01.2026

Configuration Management

Nature of Deliverable	
R	Document, report (excluding the periodic and final reports)

Dissemination level	
PU	Public, fully open

How to cite	
Evans, N., and Maris, M. (Jan 2026). AIOLIA D4.1 ESCO-based learners' profiles	

Acronym/abbreviations	
AI	Artificial Intelligence
AI HLEG	AI High-Level Expert Group
ALTAI	Assessment List for Trustworthy AI
GPAI	General Purpose AI System
WP	Work Package

Acknowledgements

The information and views set out in this report are those of the author(s) and do not necessarily reflect the official opinion of the European Union. Neither the European Union institutions and bodies nor any person acting on their behalf may be held responsible for the use which may be made of the information contained therein. Reproduction is authorised provided the source is acknowledged.

Disclaimer

Funded by the European Union. Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union. Neither the European Union nor the granting authority can be held responsible for them.

EXECUTIVE SUMMARY

This deliverable (D4.1) reports on the analysis phase of the AIOLIA training development process, conducted within Work Package 4 and guided by the ADDIE instructional design framework. Its reports on the analysis of AI ethics training needs and defines ESCO-based learner profiles for AIOLIA's training target groups.

The analysis is based on an AI ethics learning needs survey (187 answers) targeting five professional groups: ethics reviewers, AI governance and policy experts, technical AI researchers, researchers working in projects with AI, and research ethics educators. The survey assessed respondents' confidence in operationalising AI ethics guidance and regulations in practice, identified key challenges across professional roles, and explored preferences for course content and training delivery. Open questions provided qualitative insights into real-world ethical dilemmas, regulatory uncertainty, and difficulties translating abstract principles into concrete decisions.

Results confirm a clear performance gap across all groups, particularly in relation to applying the EU AI Act, identifying prohibited practices, implementing ethics-by-design approaches, and assessing societal and fundamental rights impacts. While some groups report moderate confidence, few reported feeling highly confident, and most had never received formal AI ethics training. Across roles, respondents expressed a strong demand for practical, case-based, and application-oriented learning. Technical AI researchers in particular desired tools such as standard operating procedures, templates, and flowcharts.

A mapping of existing openly available AI ethics training materials was conducted and revealed that existing resources focus largely on high-level principles and compliance, with limited support for addressing real-world complexity, socio-technical trade-offs, and long-term societal impacts.

Finally, learner profiles and competence frameworks were co-created with consortium experts and translated into ESCO-aligned competencies. These profiles define cross-cutting and role-specific knowledge, skills, and attitudes at fundamental, intermediate, and advanced levels. Together, the findings provide a solid basis for defining instructional goals and designing scaffolded learning pathways in Task 4.2, ensuring that AIOLIA training responds directly to demonstrated needs and supports transferable, practice-relevant AI ethics competencies.

CONTENTS

1. DESCRIPTION OF THIS DELIVERABLE	5
2. INTRODUCTION TO THE ADDIE FRAMEWORK	5
3. THE PERFORMANCE GAP	7
4. COURSE CONTENT AND DELIVERY	23
5. EXISTING AI ETHICS TRAINING COURSES AND EDUCATIONAL MATERIALS	24
6. LEARNER PROFILES AND ESCO COMPETENCIES	26
7. CONCLUSION	31
8. REFERENCES	31
9. DATA SET TABLE	32
10. APPENDICES	33

LIST OF FIGURES

Figure 1. AIOLIA will develop and validate training materials using the ADDIE interactive design methodology with an additional Monitoring phase.....	6
Figure 2. ADDIE instructional design process (Branch et al 2009).....	6
Figure 3. Miro board for competencies co-creation.....	26

LIST OF TABLES

Table 1. Survey recruitment strategies per training target group.....	8
Table 2. Respondents' main area of expertise.	9
Table 3. Ethics reviewers' years of experience.	10
Table 4. Ethics reviewers' confidence in assessing AI in proposals.....	11
Table 5. Horizon ethics self-assessment AI questions which Horizon Ethics Appraisal Scheme Reviewers find challenging to assess.....	11
Table 6. Practices prohibited in the AI Act for AI systems or models intended for the EU market which ethics reviewers find challenging to identify or assess.	12
Table 7. AI governance and policy experts' years of experience.	12
Table 8. Policy and governance experts' place of work.	13
Table 9. AI policy and governance experts' confidence in applying the AI Act.....	13
Table 10. Technical AI researchers' years of experience.....	13
Table 11. Technical AI researchers' confidence related to AI ethics.....	14
Table 12. Researchers' years of experience.	16
Table 13. Confidence of researchers regarding AI ethics in projects.....	17
Table 14. Ethics educators' years of experience.	19
Table 15. Confidence teaching AI ethics.	20
Table 16. Areas which ethics educators would most like to learn.	21

Table 17. Learning delivery preferences of all survey respondents (n=187).	23
Table 18. Aggregated good and bad points of other (online) AI ethics courses followed by respondents.	23
Table 19. Core and target group specific AI ethics competencies.	28
Table 20. Training target groups, ESCO occupations, competencies, and skills.	30

LIST OF GRAPHS

Graph 1. Areas of research reviewed by ethics reviewers.....	10
Graph 2. Areas of research covered by AI policy and governance experts.	12
Graph 3. Research domains in which AI technical researchers work.	14
Graph 4. Area of AI research in which AI technical researchers are working...	14
Graph 5. Research domains in which researchers work.	17
Graph 6. Research domains in which ethics educators teach.	20
Graph 7. Areas of AI research all respondents would like to learn about.....	21
Graph 8. Areas related to the application of the EU AI Act and other EU guidance all respondents would like to learn about.....	22

1. Description of this Deliverable

This deliverable (D4.1) was developed within Work package 4 “Developing and validating training materials” - and reflects the first objective of that WP:

Obj. 1 Analyze needs to define learners’ profiles using co-creation workshops

The work described in the deliverable was undertaken in Task 4.1 (Box 1).

T4.1 Analyze training needs and define learners’ profiles (AUMC, CERTH, M6-12)

This task constitutes the ‘analysis’ phase of the ADDIE instructional design methodology to better understand the learners and context. Learners’ profiles will be defined capitalizing on ESCO classification to provide information on their role, knowledge, skills, attitudes, etc., as part of the workshops organized in T3.1. Relevant materials for openlicensed AI ethics training will be gathered and evaluated for relevance. Liaisons with EU Academy will be established to investigate publication possibilities.

Box 1. Task 4.1 description from the AIOLIA grant agreement.

2. Introduction to the ADDIE framework

In AIOLIA, training is based on the ADDIE (Analyse, Design, Develop, Implement, Evaluate) framework for instructional design. ADDIE is a process framework particularly suited for developing materials for professional performance. The framework was originally developed in the 1970s to provide a systematic approach for the design of training programmes for US Army Professionals by Florida State University (Allen 2006).

Since then, there have been numerous descriptions of the framework and what each phase should consist of (Branch 2009, Dick et al 2014, Robleyer 2015). These differ in detail, however there are some characteristics common to all descriptions, namely that the approach: 1) is suitable for the development of instruction when there is evidence that desired performance is currently not being reached and this performance gap is due to a lack of knowledge and skills; 2) is outcomes focused – with competencies, learning goals and instructional materials aligned with the perceived needs in relation to competencies, performance goals and performance tasks from practice; 3) takes a systems approach, accounting for prior knowledge of learners and contextual constraints; 4) incorporates aspects of evaluation at every stage of the process.

To ensure the ADDIE process in AIOLIA is agile to rapid changes in AI technology development, a ‘Monitor’ phase has also been incorporated (Fig 1).



Figure 1. AIOLIA will develop and validate training materials using the ADDIE interactive design methodology with an additional Monitoring phase.

This deliverable, like Task 4.1, concerns the 'Analysis' phase of the ADDIE process. The Analysis phase confirms the performance gap (and the reasons for it), determines instructional learning goals, identifies the intended audience (learners) and the resources required for learning, describes the preferred teaching delivery methods, and develops a project management plan for the training (Branch 2009) (Fig 2).

	Analyze	Design	Develop	Implement	Evaluate
Concept	Identify the probable causes for a performance gap	Verify the desired performances and appropriate testing methods	Generate and validate the learning resources	Prepare the learning environment and engage the students	Assess the quality of the instructional products and processes, both before and after implementation
Common Procedures	1. Validate the performance gap 2. Determine instructional goals 3. Confirm the intended audience 4. Identify required resources 5. Determine potential delivery systems (including cost estimate) 6. Compose a project management plan	7. Conduct a task inventory 8. Compose performance objectives 9. Generate testing strategies 10. Calculate return on investment	11. Generate content 12. Select or develop supporting media 13. Develop guidance for the student 14. Develop guidance for the teacher 15. Conduct formative revisions 16. Conduct a Pilot Test	17. Prepare the teacher 18. Prepare the student	19. Determine evaluation criteria 20. Select evaluation tools 21. Conduct evaluations
	Analysis Summary	Design Brief	Learning Resources	Implementation Strategy	Evaluation Plan

Figure 2. ADDIE instructional design process (Branch et al 2009).

Some of these aspects of the Analysis phase were already comprehensively outlined in the grant agreement. Namely, the intended audience for training, required resources, and a project management plan. The audience for the training consists of the training target groups for WP5, which are: 1) ethics appraisal experts; 2) policy experts; 3) next generation AI specialists (including graduates and postgraduates); 4) researchers and students; and 5) research ethics trainers. The resources available for the training development and the project management plan are part of the broader project management. Task 4.1 therefore focuses on validating the performance gap and determining the instructional goals.

The performance gap is determined via a learning needs survey of the training target groups, with questions related to the AI ethics tasks common to their specific professional roles, their levels of confidence completing these tasks, and the aspects they find most challenging. The survey also contains questions of learning content and delivery preferences.

Instructional goals are determined via developing learner profiles for the training target groups based on insights derived from an AI ethics learner needs survey, as well as the insights of consortium experts

from the tasks in WP3 of the project. These learner profiles are further distilled into core ESCO competencies to describe ESCO-based learners' profiles. These competence profiles are broad instructional goals for each professional group. They are oriented toward professional functioning. They are not the more finely tuned learning goals which will be developed in Task 4.2. These specify what a learner should know or do in relation to a specific course or module. More fine-grained course and module learning goals operationalize and align with the broader professional competence profiles.

Important insights for the analysis phase include understanding the intended audience's "learning context". This includes understanding existing educational possibilities and the gaps between identified learning needs and existing materials. Relevant materials for open-licensed AI ethics training have also therefore been systematically collected and evaluated for relevance. Gaps in current training, particularly for the application of current guidance in practice are outlined.

Specific aims of the AIOLIA 'Analysis' phase are therefore:

1. To describe the training target groups concerns, challenges and confidence undertaking tasks related to AI ethics in their professional roles (the performance gap).
2. To determine course content and delivery preference (to describe potential delivery).
3. To identify existing AI ethics training courses and educational materials and training gaps (the training gap).
4. To develop learner profiles for each training target group including specific knowledge and skills and attitudes and to describe the competencies according to the classification of European Skills, Competences, and Occupations (ESCO) (instructional goals).

The deliverable is structured reflecting these aims and concludes with a summary of the analysis phase and an indication of the priorities regarding the next steps.

3. The performance gap

3.1 Introduction

The Analysis phase builds on existing knowledge by taking the learning needs analysis (Aucouturier and Grinbaum 2023) and final evaluation of AI ethics training materials for researchers and ethics reviewers) (Pallise Perello Pers. Comm. Sept 2025) produced by the IRECs consortium as a starting point.

The IRECs learning needs analysis revealed that Research Ethics Committee (REC) members, researchers, and institutional leaders need training in basic AI concepts, bias and transparency. They need training in ethical issues that go beyond consent and privacy and include societal impact, fairness, and accountability. They also need training in risk assessment, ethics-by-design, and embedded ethics approaches (Aucouturier and Grinbaum 2023). Based on this needs assessment, the IRECs project produced two modules related to AI: 1) Technology Basics; and 2) Technology Ethics. The IRECs project conducted a mixed method evaluation of these two modules amongst researchers and ethics reviewers, which revealed some outstanding needs (C. Pallise Perello Pers. Comm. Sept 2025). Researchers wanted more concrete, real world, disciplined specific examples. They wanted real world complexity rather than simplification in case examples. They wanted more reflection on privacy and safety, particularly for specific groups (e.g. patients). Ethics reviewers wanted practical tools to use – e.g. an overview of AI technologies and associated ethical issues. They also wanted more real

examples (C. Pallise Perello Pers. Comm. Sept 2025). This evaluation indicates that there is still a gap between the learning needs identified in IRECs and the final IRECs training provision.

This gap can be addressed in the AIOLIA project, which is centered around real-world complex case studies. AIOLIA's training target groups are broader than those of IRECs, which was healthcare focused. Furthermore, AIOLIA also focuses more specifically on the operationalization in practice on the principles in the EU AI Act. An AI ethics learning needs survey was conducted in order to better describe the training target groups concerns, challenges and confidence undertaking tasks related to AI ethics in their professional roles.

3.2 METHODS

Survey instrument and recruitment

Originally, data for task 4.1 were envisaged in the plan of work to be collected in the WP3 workshops (as per the task description). However, due to the full schedule and narrow target group of the smaller, mostly online workshops, this was not the best approach to collect the data required for the analysis phase. After consultation with task 3.1 and the AIOLIA coordinator, we adjusted the methodology and conducted an online survey of all target groups to gather information via closed and open questions. The survey was designed by the AUMC team, piloted using the 'think aloud' method with one member of each of the training target groups from AUMC, and reviewed by the AIOLIA coordinator.

The survey consists of three sections. The first section identifies the respondent's stakeholder group. The second section consists of tailored experience questions based on the target group. The third section contains items on training gaps and learning needs. The full survey can be found in Appendix I.

The survey target groups reflect the training target groups. Table 1 describes these groups and the survey recruitment methods. Moreover, posts tailored to each target group were shared widely on LinkedIn, Bsky, and X. All AIOLIA partners help with recruitment, and networking partners EUREC and ADRA have actively participated in recruitment via their LinkedIn groups. The survey was also embedded on the AIOLIA consortium website. We have also received help from RTD (Ethics Sector) in disseminating the survey to the EU Ethics Appraisal Scheme experts.

Table 1. Survey recruitment strategies per training target group.

Survey target group	Survey recruitment approach
1. Ethics reviewers <i>(including Horizon appraisal scheme experts)</i>	• Horizon appraisal scheme experts invited via email and in the expert Synapse group by the EC directly. • EUREC members invited via email and through a targeted LinkedIn post. • EACME members invited via EACME news of the week.
2. AI governance and policy	• Policy experts in the consortium and the SAB invited directly. • Euractiv communicated the link on its media channels.
3. The next generation of Technical AI researchers <i>(from Masters to Post docs)</i>	• AUMC PhD students invited via mass LMS mailout • AIOLIA networks were asked to circulate the invitation amongst their members via email and /or LinkedIn posts (particularly ADRA, ERCIM, ALL in AI, EURODOC)

4. Other (cognitive and behavioural) researchers (including 'students' who we considered researchers in training)	• Posts shared on multiple LinkedIn accounts, Bsky, and X
5. Research ethics educators	• Invite shared on the NERQ LinkedIn • Network invitations are relevant, since some EACMA and EUREC members include research ethics educators

Analysis

Closed responses are analysed using descriptive statistics. Open responses were categorized for key themes and insights.

3.3 RESULTS

Respondents' characteristics

Expertise

The survey received 187 responses, the most frequent response groups were researchers, followed by ethics reviewers, ethics educators, technical AI researchers, and AI governance or policy experts (Table 2).

Table 2. Respondents' main area of expertise.

Expertise	n (%)
Ethics reviewer	42 (22)
AI governance or policy expert	22 (12)
Technical AI researcher	24 (13)
Researcher	48 (26)
Research ethics educator	33 (18)
Other	18 (10)

Respondents who chose the 'Other' category (18) were: Ethics and integrity experts who did not class themselves as 'reviewers' (e.g. confidential counsellors, ethics advisors, field experts etc.); education developers who did not consider themselves 'educators'; research support; researchers in domains outside cognitive and behavioral research; and students.

Experiences and perspectives related to the operationalization of AI guidelines and frameworks in practice (the performance gap)

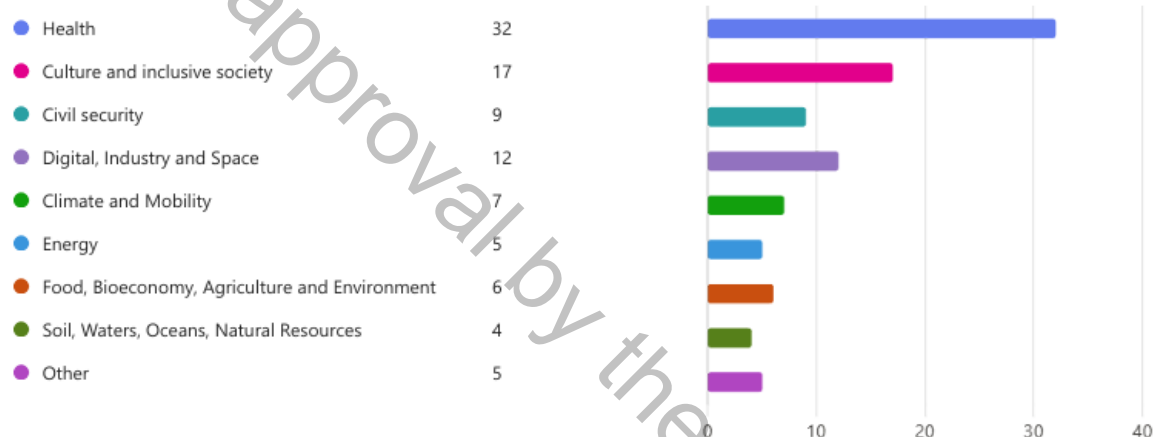
Ethics reviewers (including of European funded research proposals) and/or Member of Research Ethics Committee (n=42)

Nearly half (45%) of the ethics reviewers were very experienced with +10 years of experience (Table 3).

Table 3. Ethics reviewers' years of experience.

Ethics reviewers' years of experience	n (%)
Less than 1 year	1 (2)
1-5 years	16 (38)
6-10 years	6 (14)
10+ years	19 (45)

Although health research was the most frequent research area that the experts reviewed, across the sample there was experience reviewing all research domains (Graph 1).



Graph 1. Areas of research reviewed by ethics reviewers.

Thirty-four respondents had worked in a research ethics committee (national or institutional). Of these, over half feel moderately or very confident in assessing the ethics of AI in research proposals (Table 4).

Twenty-two respondents acted as Horizon Ethics Appraisal Scheme Reviewers. Of these, over half feel moderately or very confident in assessing the AI ethics parts of the Horizon Europe Ethics Self-Assessment (Table 4). Amongst the Horizon Ethics Appraisal Scheme Reviewers specifically, over 45% found each of the Horizon ethics self-assessment AI questions challenging to assess (Table 5).

All ethics reviewers (42) were asked how confident they feel assessing if a proposal contains practices prohibited in the AI Act. Less than half of ethics reviewers were confident that they could assess if a proposal contains practices prohibited in the AI Act (Table 4). Of the practices prohibited in the AI Act, those most difficult to assess were harmful manipulation and deception, followed by emotional recognition, and harmful exploitation of vulnerabilities (Table 6).

All ethics reviewers (42) were asked how confident they feel assessing technical and organisational design measures put in place to implement and monitor ethical concerns. Only 34% felt (moderately) confident.

Table 4. Ethics reviewers' confidence in assessing AI in proposals.

	n (%)				
	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
How confident do you feel assessing the ethics of AI in research proposals as a member of a research ethics committee? (n=34)	2 (6)	8 (24)	4 (12)	15 (44)	5 (15)
How confident do you feel assessing the AI ethics parts of the Horizon Europe Ethics Self-Assessment? (n=22)	0 (0)	4 (18)	4 (18)	10 (45)	4 (18)
How confident do you feel assessing if a proposal contains practices prohibited in the AI Act? (n=42)	4 (10)	9 (21)	11 (26)	14 (33)	4 (10)
How confident do you feel assessing technical and organisational design measures put in place to implement and monitor ethical concerns (n=42)	5 (12)	12 (29)	10 (24)	14 (33)	0

Table 5. Horizon ethics self-assessment AI questions which Horizon Ethics Appraisal Scheme Reviewers find challenging to assess.

Horizon ethics self-assessment AI questions	n (%)
Could the AI based system/technique potentially stigmatise or discriminate against people?	10 (45)
Does the AI system/technique interact, replace or influence human decision-making processes?	13 (59)
Does the AI system/technique have the potential to lead to negative social impacts?	15 (68)
Does this activity involve the use of AI in a weapon system?	10 (45)
Does the AI to be developed/used in the project raise any other ethical issues not covered by the questions above?	17 (77)

Table 6. Practices prohibited in the AI Act for AI systems or models intended for the EU market which ethics reviewers find challenging to identify or assess.

Which practices prohibited in the AI Act for AI systems or models intended for the EU market are challenging to identify or assess in proposals	n (%)
Real time remote biometric identification	18 (43)
Biometric categorisation	22 (52)
Emotion recognition	26 (62)
Untargeted scraping to develop facial recognition databases	16 (38)
Individual criminal offence risk assessment and prediction	20 (48)
Social scoring	19 (45)
Harmful exploitation of vulnerabilities	24 (57)
Harmful manipulation and deception	29 (69)

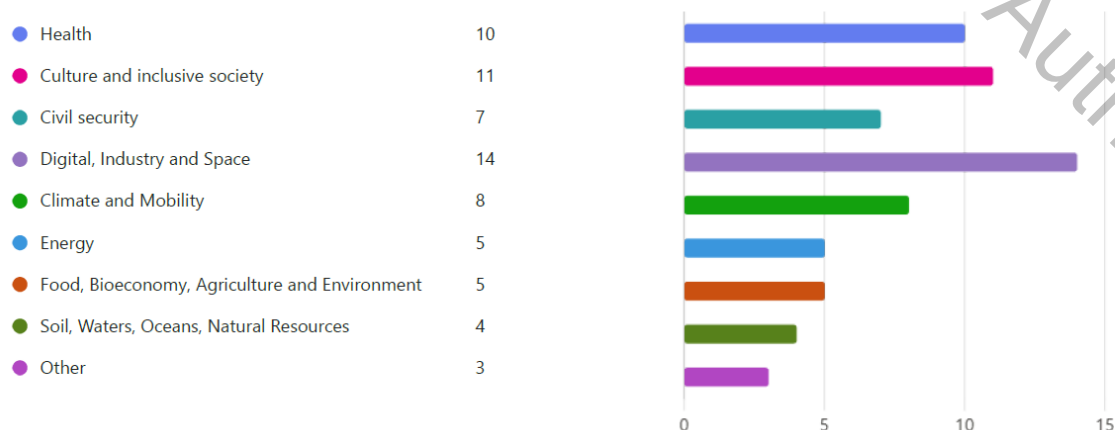
AI governance of policy expert (n=22)

Most AI policy and governance experts had less than 5 years experience, reflecting the recent developments of AI (Table 7).

Table 7. AI governance and policy experts' years of experience.

AI governance and policy experts' years of experience	n (%)
Less than 1 year	4 (18)
1-5 years	10 (45)
6-10 years	5 (23)
10+ years	3 (14)

Most experts considered their policy area to reflect the digital (industry and space) research area, reflecting more generalist policy makers (Graph 2).



Graph 2. Areas of research covered by AI policy and governance experts.

The most frequent place of work was in a research institute, followed by a national/regulatory body (Table 8). Most felt moderately confident about being able to apply the AI act in their policy work (Table 9).

Table 8. Policy and governance experts' place of work.

Place of work (more than one response possible)	n (%)
National government/regulatory body	7 (32)
International regulatory body	3 (14)
Research institute	11 (50)
NGO or think tank	5 (23)
Other	4 (18)

Table 9. AI policy and governance experts' confidence in applying the AI Act.

	n (%)				
	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
How confident do you feel applying the AI act in your research policy and governance work	2 (9)	1 (5)	6 (27)	13 (59)	0 (0)

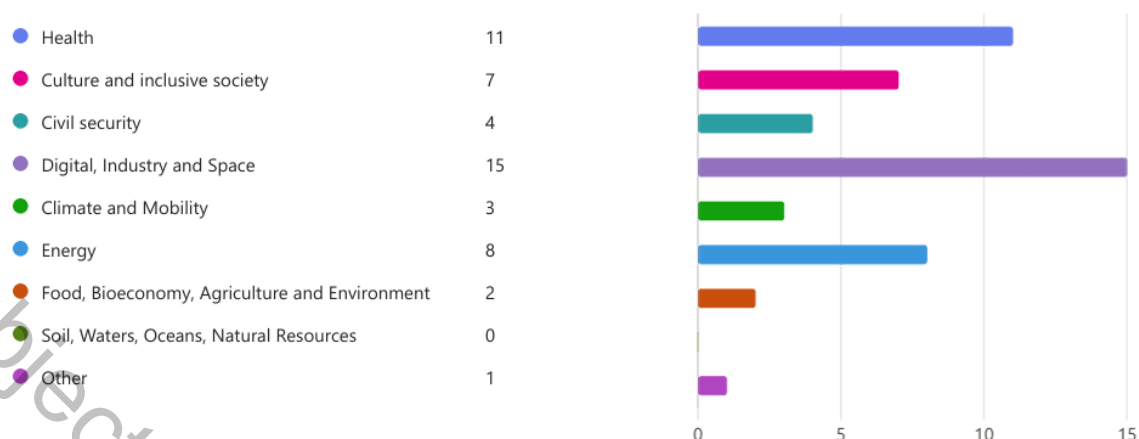
Technical AI researchers (n=24)

Most technical AI researchers had less than 5-years of experience, again reflecting the recent developments of AI.

Table 10. Technical AI researchers' years of experience.

Years of experience Technical AI researcher	n (%)
Less than 1 year	0 (0)
1-5 years	13 (54)
6-10 years	3 (13)
10+ years	8 (33)

Most described their research domain as 'digital, industry and space'. Those that were more domain specific were most frequently working in AI in health or energy research (Graph 3), and most working on GPAI or decision-support AI (Graph 4).



Graph 3. Research domains in which AI technical researchers work.



Graph 4. Area of AI research in which AI technical researchers are working.

When asked how confident they are in identifying and addressing AI ethics concerns throughout a project's entire lifecycle (including the AI systems life cycle), half were moderately or very confident (Table 11). When asked about applying an 'ethics by design' approach within their workflows, AI technical researchers were less confident, with only 38% moderately confident about their abilities (Table 11). When asked how confident they are in complying with EU AI Act and other relevant regulations, again half were moderately confident (Table 11)

Table 11. Technical AI researchers' confidence related to AI ethics.

	n (%)				
	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
n=24					
How confident are you in identifying and addressing AI ethics concerns throughout a project's entire lifecycle (including the AI-systems life cycle)	2 (8)	3 (13)	7 (29)	11 (46)	1 (4)

How confident are you applying an 'ethics by design' approach within your workflow?	3 (13)	6 (25)	6 (25)	9 (38)	0 (0)
How confident are you complying with EU AI Act and other relevant regulations? N	2 (8)	4 (17)	6 (25)	12 (50)	0 (0)

What AI ethics concerns do technical researchers encounter?

When asked in an open question about which AI ethics concerns they encounter and what makes these concerns challenging to address, respondents who described themselves as (moderately or very) confident about addressing AI concerns identified challenges using ethical principles. E.g. *"Fairness and non-discrimination"*; *"Privacy and data protection"*; *"Transparency, human oversight, trustworthiness"*; *"Trust"*.

One respondent outlined the operational processes for developing fair & robust AI:

"(1) Accessing high-quality test data, including privacy-sensitive protected features. (2) Accessing actual real-life AI systems for testing purposes (e.g., instead of outdated or dummy systems). (3) Considering ethical impacts at the early stages of AI development. (4) Producing transparent and honest reporting of AI error and bias."

Those who described themselves as unconfident responded mostly in the form of questions.

*"To what extent am I still the author and owner of an AI-generated solution?
How to ensure that new generations exploit AI for their own benefit and not to find shortcuts only?
How to make sure that the human stays in the loop of important decisions?
How to make AI a human booster, not a replacement?"*

Responders who were neither confident nor unconfident tended to describe their areas of research as without human data or ethical concerns.

When asked in an open question about what they find challenging about an ethics by design approach, AI researchers tended to describe the difficulties in upholding principles in technical decisions, ask questions about what types of measures are needed, or appeal to the need to follow the correct steps without specifying what these might be. AI researchers ask frequently for clear instructions, while ethical thinking is complex and messy. However, standard operating procedures and evaluation tools are clearly necessary for this target audience

"It is really complex to translate abstract ethical principles into concrete technical decisions throughout the AI lifecycle. Values such as fairness, transparency or accountability are dependent of the context and therefore very hard to formalize and implement."

"I can follow if proper and generic instructions are available"

"Giving concrete clear steps in order to apply"

Those who described themselves as less confident in an ethics by design approach often talked about lack of training or lack of awareness of the term

AI researchers also described tensions and trade-offs between choices.

“Fairness metrics & targets: Many choices (metrics/subgroups/thresholds); trade-offs between sensitivity and equity”

“Issues often come up during deployment. It is important that the ethic guidelines are adaptive, not monolithic”

“We cannot evaluate AI methods very thoroughly, and there may be some edge cases that we cannot foresee.”

Researchers also commented on competition between the goals/priorities of research and attention to ethical considerations, industry’s disregard for ethical issues, and the fear of being left behind due to considering ethical issues.

“Ethical concerns are important, however, are not a key income source when developing these tools. A better understanding and precise resource management to tackle ethical concerns are missing, in many cases. Thus, more resources need to be put in place in, e.g., project proposals to be able to seriously account for the ethical concerns, whilst at the same time allowing fast and reliable development of tools or applications to compete on the AI market.”

Furthermore, many AI researchers did not know what the term “ethics by design” meant.

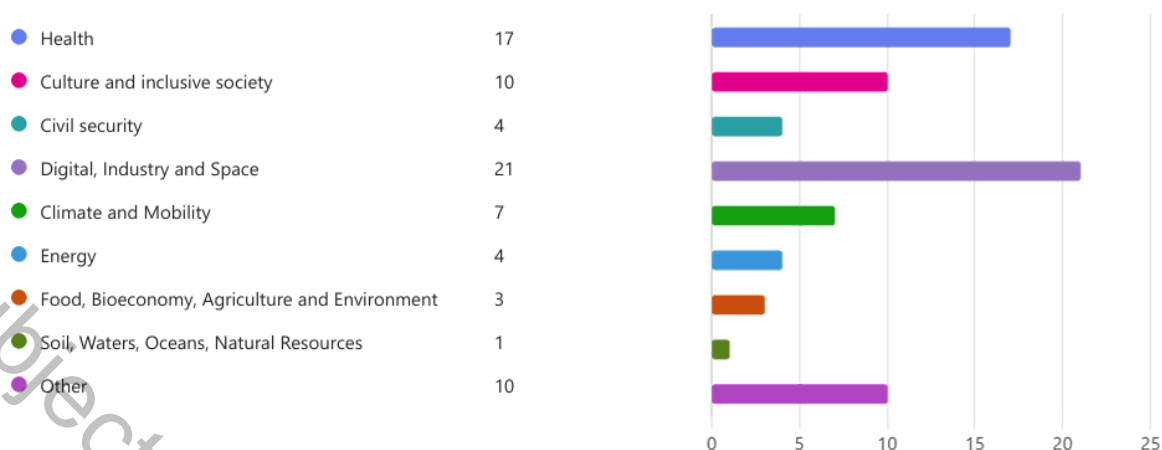
Researchers (n=48)

Half of the researchers who worked on projects with AI, but were not technical AI researchers, had over 10 years of experience, and 21% had between 6-10 years of experience. Indicating a very senior respondent profile.

Table 12. Researchers’ years of experience.

Years of experience other researchers	n (%)
Less than 1 year	7 (15)
1-5 years	7 (15)
6-10 years	10 (21)
10+ years	24 (50)

These researchers described their area most frequently as ‘digital, industry and space’, followed by health and culture (Graph 5).



Graph 5. Research domains in which researchers work.

When asked about their confidence identifying and addressing AI ethics concerns throughout a project's entire lifecycle (including the AI systems life cycle), these researchers were far more confident than the AI researchers – with over 60% moderately or very confident (Table 13). When asked about applying an 'ethics by design' approach within their workflows, researchers were more confident than the AI researchers with 50% either moderately or very confident about their abilities (compared with only 35% of AI technical researchers) (Table 13).

When asked how confident they are in complying with the EU AI Act and other relevant regulations, 48% of researchers were moderately or very confident (this is very close to 47% among AI researchers) (Table 13).

Table 13. Confidence of researchers regarding AI ethics in projects.

	n (%)				
	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
n=48					
How confident are you in identifying and addressing AI ethics concerns throughout a project's entire lifecycle (including the AI-systems life cycle)	1 (2)	4 (8)	13 (27)	22 (46)	8 (17)
How confident are you applying an 'ethics by design' approach within your workflow?	7 (15)	8 (17)	8 (17)	20 (42)	5 (10)
How confident are you complying with EU AI Act and other relevant regulations? N	5 (10)	8 (17)	12 (25)	19 (40)	4 (8)

What AI ethics concerns do researchers working on projects with AI encounter?

When asked about their AI ethics concerns, researchers were more likely to describe AI use very broadly across the whole research process (from literature review and scoping, to article text and image generation). Ethics concerns were equally broad, ranging from concerns about privacy (data leakage) to concerns about society more generally. Researchers were also concerned that AI use in research would replace entry-level research jobs and limit opportunities for young researchers. Researchers described the need for further guidance for appropriate use in projects. With some respondents concerned that the speed of integration was too fast and pushed from 'the top', whilst other comments reflected technical misunderstandings rather than a solid understanding of how the AI systems work.

"I think putting data from our research into GenAI bots means that we are training them [technical misunderstanding], and I am not certain that this information will not be leaked to others researching in my field through the answers that the bot gives based on the training it received from others"

"The overarching question is how to use AI appropriately, in line with your research objectives, etc. This ensures that you conduct high-quality research, thereby optimally contributing to the intended outcomes of the study."

Those who answered the open question quoting principles or values, often included more societal and relationship-based principles that did not come up in the AI researcher comments "empathy, inclusiveness", "social justice"

Many researchers, particularly those less confident in identifying ethics problems, were also concerned about using tools built by others and worried about unknowingly reinforcing biases.

"lack of transparency -> many AI models, especially deep learning systems, operate as "black boxes", validating output therefore and the whys behind it. privacy and data protection -> reliance on personal or sensitive data, surveillance, risking misuse and data breaches."

"I frequently encounter concerns related to data privacy, informed consent, transparency of AI-supported decision-making, and potential bias in datasets or outcomes. These are challenging to address because AI components are often developed or integrated by third parties, making it difficult to fully understand data provenance, model behaviour, and downstream impacts. Additionally, ethical guidance is sometimes high-level and difficult to operationalise in day-to-day research practices."

When asked in an open question about what they find challenging about an ethics by design approach, researchers emphasized the time and forward planning dimensions, the need for dialogue with developers and the importance of developers' technical and ethical reasoning skills. This is an important finding for Task 4.2 – how to train experts meaningfully within a relatively short time frame?

"My opinion is that the hardest thing about ethics by design is time. A correct approach to a structured, documented and safe system (which uses AI) simply takes too long. Planning times ahead of a project are already too big and AI is not a variable which should be overlooked."

"This requires a level of forward planning that may not always be compatible with the flexible nature of research work."

“The lack of acceptance by the technical people”

“Dialogue between developers/technical designers and ethicists. How to translate technical properties into ethical principles”

“The biggest challenge is balancing ethical requirements with technical feasibility, project costs and timelines, plus vague ethical standards that make practical implementation hard”

Researchers also talked about the tensions and trade-offs with other goals and priorities

*“It is in tension with all the *other* goals a system has to achieve, and there isn't *really* a systemic incentive to push for more ethics at the expense of the other goals, to find a balanced trade-off. Sure, there are high-level declaration that it is important, but when time or money enter the discussion, ethics often become secondary, even when a few individuals try to push for it.”*

“Finding the right way to slot it into a project - it has to come up at the very beginning, but not so early that it eclipses the project's development.”

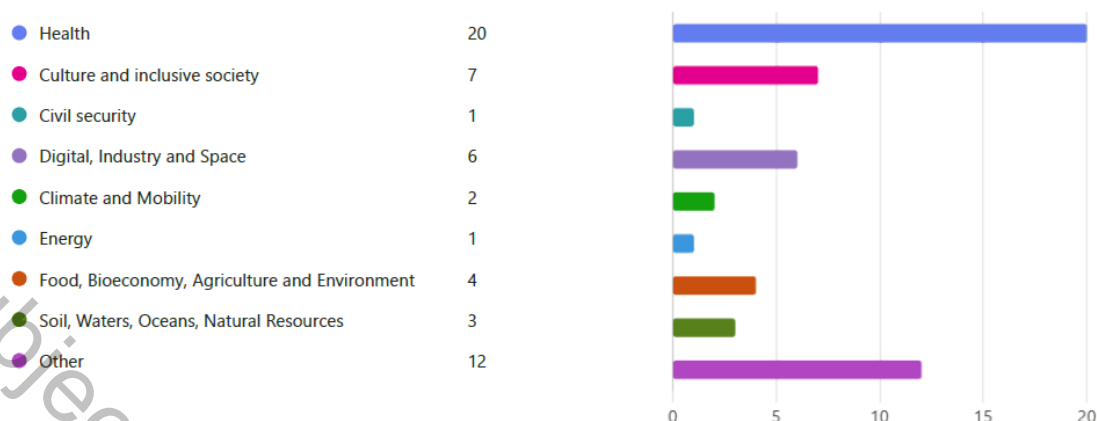
Half of (moderately) unconfident responders did not know what “ethics by design” meant.

Ethics educators (n=33)

Many of the ethics educators who answered the survey were fairly new to ethics education; 9% had less than one year and 33% between 1 and 5 years of experience. A further 27% had 6 to 10 years of experience, and 30% had over 10 years (Table 14). Ethics educators were most frequently specialized in the health research area (Graph 6).

Table 14. Ethics educators' years of experience.

Years of experience ethics educators	n (%)
Less than 1 year	3 (9)
1-5 years	11 (33)
6-10 years	9 (27)
10+ years	10 (30)



Graph 6. Research domains in which ethics educators teach.

Over half of educators felt moderately or very confident about teaching AI ethics. About a third, however, felt moderately unconfident or had no confidence (Table 15).

Table 15. Confidence teaching AI ethics.

How confident are you teaching AI ethics?	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
N (%)	2 (6)	9 (27)	4 (12)	13 (39)	5 (15)

Challenges teaching AI ethics

When asked in an open question about what they find challenging about teaching AI ethics, educators described difficulties in staying up to date with the technological developments and understanding the technical aspects of AI.

"While the focus is on AI ethics (principles, guidelines, etc.), one cannot divorce the technical aspects, which is challenging for me (to some degree) as a bioethicist."

"Knowing the functionality of tools/their security settings to be able to offer informed view or discussions"

They also described a lack of national and institutional guidance to refer to:

"Not many AI ethics guidance or handbooks, many cases with AI misconduct that are not correctly recognized in public or scientific opinion, lack of institution regulations"

"Concrete guidelines are currently lacking in our national code of conduct"

Another challenge was teaching technically oriented researchers that are not used to thinking about research ethics.

"... learners can often want 'black or white answers' or a 'tick box approach' instead, for example 'it will always be acceptable to do A, B and C', 'it will always be wrong to do X, Y and Z', 'these

tools are safe to use and these are not" rather than principles/ tools to help them work out the answers in such a fast-moving environment. This is especially the case with privacy and data governance, as the data handling policies of tools can be opaque."

And a lack of access to good case studies/resources

"Few evidence based examples are available to be used in case studies. So, I use often those related to gender inequalities reinforced by IA in clinical research."

"To find the best scientific evidence and to identify and "curate" best sources among the enormous amount of information displayed."

The areas they most frequently desired to improve in relation to teaching were, AI ethics (including philosophical foundations), learning strategies, and more advanced technical knowledge (Table 16). They also mostly wanted to learn about GPAI and decision-support AI (Graph 7).

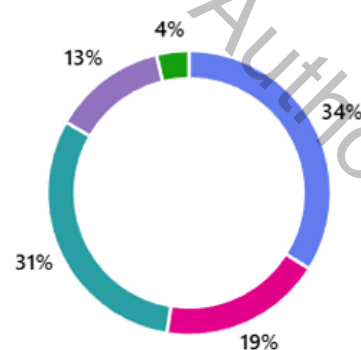
Table 16. Areas which ethics educators would most like to learn.

Areas that ethics educators would most like to learn	n (%)
Basic AI concepts and terms	11 (33)
More advanced technical knowledge on AI models and systems (including design measures to address ethical concerns)	19 (58)
AI ethical concepts and terms (including philosophical foundations)	22 (67)
Learning strategies, or pedagogical approaches, for teaching AI ethics.	20 (61)
Approaches for assessing students' AI ethics learning	16(48)
Other	2 (6)

AI Ethics course preferences – All respondents (n=187)

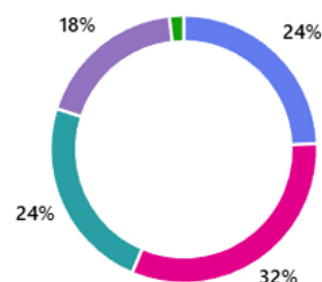
All respondents were asked about which topics they would like AI Ethics training to focus on. In relation to areas of AI research, the top topics were General Purpose AI and decision-support AI (Graph 7). In relation to areas related to the application of the EU AI Act and other EU guidance, the top topics were prohibited practices and high risks systems (Graph 8).

General-Purpose AI (GPAI)	128
Emotional AI	70
Decision-support AI	116
Image recognition AI	49
Other	14



Graph 7. Areas of AI research all respondents would like to learn about.

Prohibited practices in AI	98
High-risk AI systems	129
General-purpose AI systems	95
ALTAI principles and their operationalization	74
Other	7



Graph 8. Areas related to the application of the EU AI Act and other EU guidance all respondents would like to learn about.

Preferences for knowledge, skills and attitude development

When asked in an open question about the knowledge, skills and attitudes they would like to develop in an AI Ethics course, there were some distinct differences in responses between the training target groups:

- **Ethics reviewers** were more interested in knowledge about SOPs or lists of questions for assessing AI in research proposals.
- **AI Policy and governance experts** wanted to focus on accountability, risks for human rights and deskilling of human workers.
- **Technical AI researchers** wanted to understand technical measures to enact principles.
- Other **researchers** wanted to understand the grey areas and what was permissible in research projects.
- **Ethics educators** wanted to learn more about the impact on society, the environment, and individual well-being.

Some reflections on between group differences

Due to the survey sample size (n=187), the sizes of the training target group sub-groups (ranging 22 governance experts to 48 non-technical researchers), and the tailoring of sub-group questions to specific AI ethics related professional tasks, statistical comparison between groups or associations between their responses and personal characteristics (such as years of experience) is not feasible. There were however some noteworthy descriptive differences which are important to reflect on for the training development in Task 4.2. The groups in which more than half of respondents were moderately or very confident in their professional tasks related to AI were policy makers and researchers (non-technical). More than half of all policy makers were moderately confident in applying the AI Act in their work, though notably none were 'very confident'. Whereas researchers (non-technical) were the most confident group, consistently more confident than AI researchers, and the group most frequently 'very' confident. If this confidence reflects true skills, the relatively more experienced profile of the group, or a lack awareness or underestimation of the AI ethics risks posed by AI systems in projects in which they work is not, however, possible to say.

4. Course content and delivery

In a question to all respondents on their preferred form of learning delivery, there was no clear preference, although ‘learning by doing’ came first, followed by online following asynchronous learning modules and an online short synchronous seminar. This is an important result for Task 4.2 – and reflects the need, already expressed by AI researchers in the qualitative results above, for flowcharts, SOPs, etc.

Table 17. Learning delivery preferences of all survey respondents (n=187).

Learning delivery preferences (multiple responses possible)	n (%)
Online via short (e.g. 3 hour) synchronous seminar	91 (49)
Online following asynchronous learning modules	94 (50)
In-person class	60 (32)
Learning by doing (via the use of tools - e.g. standard operating procedures, templates, flowcharts)	107 (57)
Other	7 (4)

Finally, 72% of respondents had never followed any AI ethics training. For those who had, and who answered the open question on which training they followed and what were the good and bad points about the training, the IRECs training was the most frequently named training. Other responses included AI ethics modules from various universities, descriptions of self-study, or the respondents were ethics experts who applied their expertise in the area of AI ethics, sometimes in the development of educational materials. Very few respondents provided the good and bad points of previous courses, however the aggregated comments of those who did are presented in Table 18, indicating that courses that were expert led, interactive and utilizing complex technical cases were most appreciated.

Table 18. Aggregated good and bad points of other (online) AI ethics courses followed by respondents.

Good points	Bad points
• Online • Expert-led modules • Use of applicable case studies in some of the courses • Use of large data, algorithms, efficiency etc. • Exchange of experiences and good practices with teachers and learners from various fields of research	• Validity of decisions, IP protection, use of the results, implications on false results • No/little hands-on exercises in some of the courses • Too short, not technical enough • Very generic and high-level, not hands-on, so missing direct applicability • Too much focus on “obvious” ethical information, not enough on current/future law and policy and their implications

AIOLIA includes structured and budgeted cooperation with the Embassy of Good Science and will publish its training materials on the Embassy’s wiki platform (<https://embassy.science/wiki/Training>).

All content shared on The Embassy in creative commons BY 4.0 licensed, therefore there are no restrictions for the training being shared on other training platforms, as long as the materials are correctly attributed to the project. Restriction might come, however, if default licenses are applied to any uploaded content which are in conflict with the open license of The Embassy's content. There might also be practical restrictions related to a lack of interoperability.

Task 4.1 also involved establishing liaisons with EU Academy to investigate publication possibilities. The WP4 coordinator is already in close contact with EU Academy (as an invited expert on video interoperability for the Academy). EU Academy is based on the Moodle learning management system, which is compatible with the content developed on The Embassy (videos, H5P, SCORM packages). The EU Academy's licensing statement is also compatible with materials published on The Embassy: *"Unless otherwise indicated (e.g. in individual copyright notices), content owned by the EU on this website is licensed under the Creative Commons Attribution 4.0 International (CC BY 4.0) licence. This means that reuse is allowed, provided appropriate credit is given and changes are indicated."*

5. Existing AI ethics training courses and educational materials

5.1 Objective and methods

This section reports on the status of the scoping review conducted within the AIOLIA project to map openly available educational and training materials on AI ethics as part of deliverable 4.1. The review is conducted in accordance with a pre-registered protocol aligned with PRISMA-ScR guidelines, which specifies objectives, eligibility criteria, and methodological procedures (<https://doi.org/10.5281/zenodo.17037977>; Appendix II). The purpose of the review is to provide a structured overview of existing open AI ethics training resources in support of subsequent AIOLIA guideline development and training activities. In addition, the results will be disseminated through a peer-reviewed publication and a dataset accessible to all consortium members and partners, documenting the identified materials and their key characteristics.

Work completed to date

In accordance with the protocol, searches were conducted across academic databases, grey literature sources, and targeted web platforms to ensure broad and systematic coverage of openly accessible AI ethics education and training materials published from 2022 onward. This multi-source strategy was chosen to capture both peer-reviewed research and educational resources that are disseminated outside traditional academic databases. Following the protocol, Google Scholar and Google Advanced Search results were restricted to the first 100 hits per query, and searches were limited to materials published from 2022 onward and primarily in English.

- **Academic databases** (Web of Science and Google Scholar) were searched to identify peer-reviewed literature and formally published educational resources. Title, abstract, and full-text screening were conducted independently by two reviewers, with disagreements resolved through discussion.
- A targeted web search using systematic **Google Advanced** Search was carried out to identify relevant grey literature and openly available training materials not indexed in academic databases.

- Major **MOOC platforms** (Coursera and edX) were systematically searched and screened, as they are key providers of large-scale, openly accessible online ethics education.
- OER Commons and MIT OpenCourseWare were included as **established repositories of open educational resources**, ensuring coverage of institutionally curated, freely available teaching materials.
- **Expert consultation** within the AIOLIA consortium was conducted to identify additional openly available AI ethics training materials.

5.2 Current phase: data charting and analysis

The review is now entering the data charting and analysis phase. Data are being extracted to systematically map key features of the materials, including pedagogical design, learning objectives, ethical concepts addressed, target audiences, and the presence of assessment or reflective components. At this stage, results are still being combined and interpreted, and the final synthesis has not yet been completed. This phase is iterative, and final inclusion decisions may still change if some materials do not provide the expected content or level of detail defined in the eligibility criteria.

5.3 Next steps and contribution to AIOLIA deliverables

The preliminary set of included materials and their key characteristics is provided in Appendix III (Preliminary inclusion table). The next phase will complete full data extraction and synthesis, resulting in a structured mapping of pedagogical approaches, ethical topics, audiences, and identified gaps across the included materials. These outputs will directly inform subsequent AIOLIA project deliverables, including the development of guidelines and openly available AI ethics training resources, in line with the agreed project timeline and Description of Action.

Thematic coverage

The materials appear to consistently address core principles such as fairness, non-discrimination, transparency, accountability, privacy, and human oversight, and to link these to concrete application areas including healthcare, education, public administration, business, finance, and conversational and generative AI systems. Many resources seem to take an organizational or regulatory perspective, focusing on compliance and responsibility across the AI lifecycle, while others use more academic or critical approaches, for example through case studies and analyses of societal impact. Together, the materials reviewed so far appear to cover a wide range of audiences, AI application domains, and ethical issues, with particular attention to professional, sector-specific, and governance-related approaches.

At the same time, most materials appear to present ethics mainly as a set of principles or rules to follow, rather than as a way of working through real-world complexity and conflicting values. Issues such as power and inequality, the environmental impact of large-scale AI systems, and perspectives from outside the global North seem to receive relatively little attention. There is also limited discussion on how ethical values are built into concrete technical design choices, or on the long-term effects of AI systems. Overall, while the current landscape appears to be strong in raising awareness and supporting regulatory compliance, it seems to be less developed in encouraging deeper, critical, and socio-technical thinking about how AI shapes society and how ethical responsibility can be put into practice over time.

6. Learner profiles and ESCO competencies

Learner profiles provide the competencies for each of the training target groups which are the end goal for training. Their use is well aligned with the ADDIE approach. They are broader and less fine grained than course/module learning objectives.

Based on the insights from the AI learning needs survey, expert members of the AIOLIA consortium were invited to participate in a co-creation activity in which learner profiles for the training target groups were created. This 1.5-hour online co-creation workshop consisted of three parts: 1) presentation and discussion of the results of the learning needs survey; 2) drafting of relevant competencies (knowledge, skills, and attitudes) for training reflecting the learning needs; 3) refinement, looking at overlap within and across competencies, and discussing disagreements.

1) The results presented reflect those in section 1 of this report.

2) Drafting of relevant competencies (knowledge, skills, and attitudes)

Workshop participants were assigned a training target group to focus on developing competence items based on the AI Ethics learning needs survey results and their other prior knowledge (although they could also add competencies to any target group). Participants did this individually on a shared Miro Board (Fig. 3).

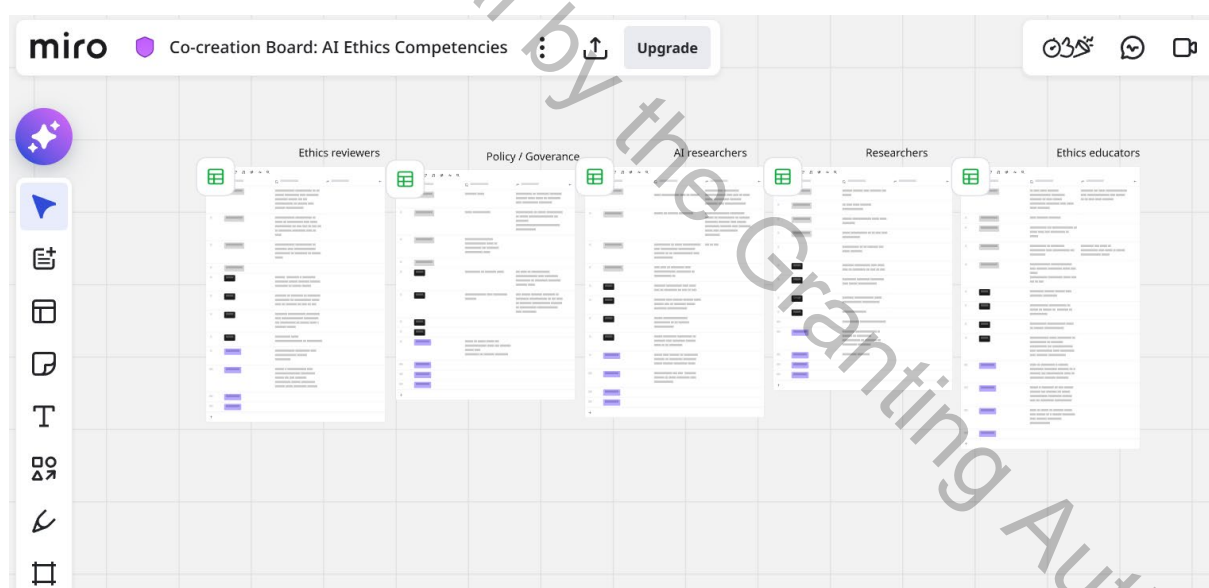


Figure 3. Miro board for competencies co-creation.

3) Refinement

Refinement began in the meeting but also online during the days following the session. Refinement involved looking at overlaps within and across competencies to find: i) core competencies common to all training target groups; ii) competencies which are more refined for the specific target groups.

This resulted in a table of common and role specific competencies, scaffolded to represent fundamental, intermediate, and advanced levels of professional competencies (Table 19). These competencies reflected specific knowledge and skills competencies, whereas desired attitudes were

reflected in the competencies themselves rather than as separate individual items. We indicate levels of progression in the competencies to aid Task 4.2 in the development of scaffolded learning pathways. Early modules might introduce more fundamental aspects of a competence, whereas later courses might focus on application in professional practice.

AIOLIA training pathways will be developed in T4.2 reflecting these competence profiles. Additionally, we develop more high-level knowledge and skills competencies for skills transferability.

European Skills, Competences, Qualifications and Occupations Learner Profiles

European Skills, Competences, Qualifications and Occupations (ESCO) is a database which describes the essential and optional skills and competencies for specified professions to support standardization of skills and transferability across Europe. The occupation descriptions within the ESCO database do not align fully with our training target groups. Below we have identified the ESCO ISCO-08 occupation codes closest to our AI ethics training groups (Table 20). We suggest new occupations and the relevant (AI only) skills and competencies for these occupations.

Subject to approval by the Granting Authority

Table 19. Core and target group specific AI ethics competencies.

Training target group	Competence area	Fundamental	Intermediate	Advanced
Cross-cutting competencies	Ethical core competencies	Understand core AI ethics principles and relevant ethical theories	Apply ethical reasoning to (real-world) complex cases	Make judgements based on ethical reasoning and societal interests
	Legal core competencies	Understand AI law (e.g. EU AI Act), data protection (e.g. GDPR) and governance frameworks	Comply with legal and governance frameworks in specific cases	Make judgements on how best to uphold the principles of the legal and governance frameworks in complex cases
	Societal core competencies	Understand AI systems to be embedded in social, political and organizational systems	Identify possible societal and rights risks	Take action to mitigate possible societal and rights risks
Ethics reviewers	Ethical and legal risk evaluation	- Explain AI ethics principles and normative frameworks - Identify EU AI Act risk categories and prohibited practices	- Assess ethical and legal risks in research proposals - Classify AI systems under the AI Act	- Judge proportionality of risks and benefits - Formulate independent, context sensitive recommendations in the public interests
	Governance and safeguards assessment	Describe technical and organizational measures for risk mitigation	Evaluate adequacy of safeguards and governance and oversight measures	Formulate recommendations technical, organizational and broader governance measures for risk mitigation
	Attitudes	Responsibility	Transparency	Accountability
AI researchers	Responsible AI and ethics by design	Explain ethics by design and privacy by design principles	- Integrate ethics by design across the AI lifecycle - Identify and mitigate ethical and societal risks - Apply technical and organizational safeguards	Make changes to systems in response to ethical risks
	Accountability and reporting	Understand legal and regulatory compliance requirements and how to comply with these in practice	Document ethical decisions, trade-offs, and risk mitigation strategies	Demonstrate accountability for societal impacts of AI

	Attitudes	Awareness	Responsibility	Accountability
Researchers	<i>Societal and rights-based impact assessment</i>	Explain ethical and rights-based considerations of AI systems	Analyze power asymmetries, societal and fundamental rights risks	Conduct interdisciplinary impact assessments of AI use within projects and in research outputs
	<i>Stakeholder engagement</i>	Describe appropriate participatory and inclusive approaches for AI system use in research projects	Facilitate stakeholder engagement and end-user and societal dialogue	Mediate value conflicts between technical and societal domains
	Attitudes	Reflexivity	Epistemic humility	Inclusivity and dialogue
Ethics educators	<i>Translating AI Ethical and legal considerations into teaching</i>	Explain AI ethics theories, AI fundamentals, and regulatory and normative frameworks	Translate ethics and regulation issues into illustrative cases and teaching materials	Design AI ethics curricula aligned with practice-oriented learning outcomes
	<i>Facilitation and assessment of ethical reasoning</i>	Apply evidence-based pedagogical approaches – such as case-based and experiential learning.	Facilitate reflective, critical, and inclusive discussions.	Assess ethical reasoning outcomes using appropriate methods.
	Attitudes	Awareness	Inclusivity and dialogue	Reflexivity
Policy / Governance	<i>AI governance and regulatory development</i>	Explain AI governance models, data governance and relevant regulatory frameworks	Develop (institutional) AI policies and governance frameworks	Design adaptive and transparent accountability and oversight structures
	<i>Societal accountability</i>	Identify societal and fundamental rights risks	- Base policy on evidence-based impact assessments - Support compliance processes	Provide policy advice with a long term, systemic and human rights perspective

Table 20. Training target groups, ESCO occupations, competencies, and skills.

Training target group	Existing ESCO ISCO-08 occupation code	Suggested ESCO occupation(s)	Suggested AI related skills and competences
Ethics reviewers	2633.2 Philosopher; 2310 University & higher education teachers; 1213.2 Policy manager (alternative label 'ethics manager')	Research ethics committee member Ethicist	• Ethics • Applied ethics • Research ethics • AI law and regulations • Data protection law • Risk analysis • Risk management
AI researchers	251 Software and applications developers and analysts 2511.4 Data scientist 2511.1 Computer scientist 2512.4 Software developer	Artificial intelligence engineer	• Artificial intelligence • Data science • Data governance • Ethics • Ethics by design • Privacy by design • ICT risk management • Technical standards compliance
Researchers	21 Science and engineering professionals 2310 University & higher education teachers (with research roles);	Academic researcher	• Ethics • Human rights • Technology assessment • Public engagement
Ethics educators	2310 University & higher education teachers	Research ethics trainer	• Ethics • Philosophy • Pedagogy • Curriculum development • Design learning activities • Promote critical thinking • Stimulate ethical reasoning
Policy / Governance	2422 Policy administration professionals; 2619 Legal professionals;	Research and innovation policy professional	• Public policy • Governance frameworks • AI law and regulations • Compliance • Risk management • Policy development • Impact assessments • Policy advice

7. Conclusion

This deliverable presents the analysis phase of the ADDIE professional instructional design approach. Following the process outlined by Branch (2009) (Figure 3), we **validated the performance gap** by assessing the AI ethics confidence of key stakeholders in AI and revealing the high proportions of respondents who were not confident about their AI ethics competencies. We identified important preferences and possibilities related to **training delivery**. Insights from the AI ethics learning needs survey allowed us to **determine the instructional goals** of AI ethics training by developing core and target group specific competence profiles. These provide a solid basis for the development of specific learning pathway learning outcomes and the development of scaffolded learning materials in Task 4.2. In mapping the existing educational courses we determined the thematic range of the current educational offerings, revealing the inadequacy of current offerings to meet our target groups learning needs. Finally, in preparation of rapidly developing AI competence needs amongst the workforce, and the need for standardisation and transferability of competencies and skills, we further identified **high-level skills and competencies for ESCO profiles** for the training target groups.

8. References

- Allen, W. C. (2006). Overview and evolution of the ADDIE training system. *Advances in developing human resources*, 8(4), 430-441.
- Etienne Aucouturier, Alexei Grinbaum. iRECS Deliverable 2.2: Recommendations to address ethical challenges from research in new technologies. CEA- Commissariat à l'énergie atomique et aux énergies alternatives. 2023. cea-04293426. Available at : <https://cea.hal.science/cea-04293426/>. Accessed August 2025.
- Branch, R. M. (2009). *Instructional design: The ADDIE approach*. New York: Springer.
- Dick, W., Carey, L., & Carey, J. O. (2014). *The systematic design of instruction* (8th ed.). New Jersey: Pearson Education, Inc..
- ESCO (2026) ESCO Classifications > Occupations. Available at: https://esco.ec.europa.eu/en/classification/occupation_main; Accessed January 2026.
- Roblyer, M. D. (2015). *Introduction to systematic instructional design for traditional, online, and blended environments*. New Jersey: Pearson Education, Inc.

9.Data set table

DATA SUMMARY	
Dataset reference	D4.1 survey
Reference Work	WP4
Package / Deliverable	D4.1
Data Manager	AUMC
Date of update of the dataset table	30.01.2026
Description and purpose	Fully anonymous survey data for determining learner profiles.
Data origin	Ethics experts, AI researchers invited via AIOLIA partners and networks
Data Size	187 questionnaires
Availability	Restricted to the consortium
Accessibility	Unprotected access with a link using office.com platform
Does the data underpin a publication	No
Utility & Re-use	Dataset will not be reused. It has only been used for the purposes of preparing this deliverable as per the data use description associated with the Survey.
Ethical aspects	All data is fully anonymised. No personal data is collected.

10. Appendices

Appendix I: Full AI Ethics Needs Survey

Appendix II: Mapping Open Educational Content on AI Ethics: A Scoping Review

Appendix III: Preliminary inclusion table in the scoping review

Subject to approval by the Granting Authority

Subject to approval by the Granting Authority

Appendix I

AI Research and Ethics - Training Needs Survey (AIOLIA Project)

Welcome to this short anonymous survey which will help the AIOLIA project (<https://aiolia.eu>) tailor AI research ethics education to your needs.

This information will only be used for education development, and not for research purposes.

* vereist

SECTION: Your background

1. What is your main role related to AI ethics for cognitive and behaviour research? [If you have multiple roles, please choose the role with the greatest AI ethics training needs] *

- ☐ Ethics reviewer (including of European funded research proposals) and/or Member of Research Ethics Committee
- ☐ AI governance or policy expert
- ☐ Technical AI researcher
- ☐ Researcher (i.e. working in projects which include elements of AI but **not** a technical AI research)
- ☐ Research ethics educator
- ☐ Andere

SECTION: Ethics reviewer

2. How many years experience do you have in ethics review processes? *

- ☐ Less than 1 year
- ☐ 1-5 years
- ☐ 6-10 years
- ☐ 10+ years

3. Which types of research proposals do you most often evaluate? [Categorized by European R&I Areas] *

- ☐ Health
- ☐ Culture and inclusive society
- ☐ Civil security
- ☐ Digital, Industry and Space
- ☐ Climate and Mobility
- ☐ Energy
- ☐ Food, Bioeconomy, Agriculture and Environment
- ☐ Soil, Waters, Oceans, Natural Resources
- ☐ Andere

4. Have you worked in a research ethics committee? *

- ☐ Yes (e.g. institutional or national)
- ☐ No

5. How confident do you feel assessing the ethics of AI in research proposals as a member of a research ethics committee? *

	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
Choose your level of confidence	☐	☐	☐	☐	☐

6. Have you worked as a Horizon Ethics Appraisal scheme reviewer? *

- ☐ Yes
- ☐ No

7. How confident do you feel assessing the AI ethics parts of the Horizon Europe Ethics Self-Assessment? *

	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
Choose your level of confidence	☐	☐	☐	☐	☐

8. Which of the AI-focused ethics self-assessment questions are challenging to assess? *

- ☐ Could the AI based system/technique potentially stigmatise or discriminate against people (e.g. based on sex, race, ethnic or social origin, age, genetic features, disability, sexual orientation, language, religion or belief, membership to a political group, or membership to a national minority)?
- ☐ Does the AI system/technique interact, replace or influence human decision-making processes (e.g. issues affecting human life, health, well-being or human rights, or economic, social or political decisions)?
- ☐ Does the AI system/technique have the potential to lead to negative social (e.g. on democracy, media, labour market, freedoms, educational choices, mass surveillance) and/or environmental impacts either through intended applications or plausible alternative uses?
- ☐ Does this activity involve the use of AI in a weapon system? (e.g. Are weapons functions are automated/autonomous and/or indiscriminate? Does a weapon have an autonomous mode for self-protection and, if yes, can distinguish between targets (threats) and non-targets?)
- ☐ Does the AI to be developed/used in the project raise any other ethical issues not covered by the questions above (e.g., subliminal, covert or deceptive AI, AI that is used to stimulate addictive behaviours, life like humanoid robots, etc.)?

9. How confident do you feel assessing if a proposal contains practices prohibited in the AI Act? *

	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
Choose your level of confidence	☐	☐	☐	☐	☐

10. Which practices prohibited in the AI Act for AI systems or models intended for the EU market do you find challenging to identify or assess in proposals? *

- ☐ **Harmful manipulation, and deception:** AI systems that deploy subliminal techniques beyond a person's consciousness or purposefully manipulative or deceptive techniques, with the objective or with the effect of distorting behaviour, causing or reasonably likely to cause significant harm.
- ☐ **Harmful exploitation of vulnerabilities:** AI systems that exploit vulnerabilities due to age, disability or a specific social or economic situation, with the objective or with the effect of distorting behaviour, causing or reasonably likely to cause significant harm.
- ☐ **Social scoring:** AI systems that evaluate or classify natural persons or groups of persons based on social behaviour or personal or personality characteristics, with the social score leading to detrimental or unfavourable treatment when data comes from unrelated social contexts or such treatment is unjustified or disproportionate to the social behaviour
- ☐ **Individual criminal offence risk assessment and prediction:** AI systems that assess or predict the risk of people committing a criminal offence based solely on profiling or personality traits and characteristics; except to support a human assessment based on objective and verifiable facts directly linked to a criminal activity
- ☐ **Untargeted scraping to develop facial recognition databases:** AI systems that create or expand facial recognition databases through untargeted scraping of facial images from the internet or closed-circuit television ('CCTV') footage.
- ☐ **Emotion recognition:** AI systems that infer emotions at the workplace or in education institutions; except for medical or safety reasons.
- ☐ **Biometric categorisation:** AI systems that categorise people based on their biometric data to deduce or infer their race, political opinions, trade union membership, religious or philosophical beliefs, sex-life or sexual orientation; except for labelling or filtering of lawfully acquired biometric datasets, including in the area of law enforcement
- ☐ **Real-time remote biometric identification ('RBI'):** AI systems for real-time remote biometric identification in publicly accessible spaces for the purposes of law enforcement; except if necessary for the targeted search of specific victims, the prevention of specific threats including terrorist attacks, or the search of suspects of specific offences

11. How confident do you feel assessing technical and organisational design measures put in place to implement and monitor ethical concerns? *

	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
Choose your level of confidence	☐	☐	☐	☐	☐

12. When assessing the technical and organisational design measures put in place to implement and monitor ethical concerns in proposals, which areas are challenging to assess? *

- ☐ **Human agency and oversight:** AI systems must support human autonomy and decision-making. This is particularly relevant for AI systems that can affect human cognition and behaviour.
- ☐ **Privacy and data governance:** AI systems must guarantee privacy and data protection throughout the system's lifecycle.
- ☐ **Transparency:** All processes associated with AI decision-making must be appropriately documented. AI systems must be explainable, and their limitations must be clearly communicated.
- ☐ **Fairness, diversity, and non-discrimination:** The best possible effort should be made to avoid unfair bias.
- ☐ **Societal and environmental well-being:** The impact of AI systems on society and the environment must be carefully evaluated. The risks of harm should be mitigated.
- ☐ **Accountability:** Actors involved in the development or operation of AI systems should clearly demarcate and assume responsibility for their functioning and outputs.

SECTION: AI Governance & Policy Experts

13. How many years experience do you have in AI governance/policy? *

- ☐ Less than 1 year
- ☐ 1-5 years
- ☐ 6-10 years
- ☐ 10+ years

14. Which areas of research and innovation does your governance/policy work cover?
[Categorized by European R&I Areas] *

- ☐ Health
- ☐ Culture and inclusive society
- ☐ Civil security
- ☐ Digital, Industry and Space
- ☐ Climate and Mobility
- ☐ Energy
- ☐ Food, Bioeconomy, Agriculture and Environment
- ☐ Soil, Waters, Oceans, Natural Resources
- ☐ Andere

15. What is your primary area of work? *

- ☐ National government/regulatory body
- ☐ International regulatory body
- ☐ Research institute
- ☐ NGO or think tank
- ☐ Andere

16. How confident do you feel applying the AI act in your research policy and governance work?

*

	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
Choose your level of confidence	☐	☐	☐	☐	☐

17. Which of the following ethical principles highlighted in the EU AI Act are most challenging from a policy and governance perspective? *

- ☐ **Human agency and oversight:** AI systems must support human autonomy and decision-making. This is particularly relevant for AI systems that can affect human cognition and behaviour.
- ☐ **Privacy and data governance:** AI systems must guarantee privacy and data protection throughout the system's lifecycle.
- ☐ **Transparency:** All processes associated with AI decision-making must be appropriately documented. AI systems must be explainable, and their limitations must be clearly communicated.
- ☐ **Fairness, diversity, and non-discrimination:** The best possible effort should be made to avoid unfair bias.
- ☐ **Societal and environmental well-being:** The impact of AI systems on society and the environment must be carefully evaluated. The risks of harm should be mitigated.
- ☐ **Accountability:** Actors involved in the development or operation of AI systems should clearly demarcate and assume responsibility for their functioning and outputs.

SECTION: Technical AI researcher

18. How many years experience do you have in technical AI research? *

- ☐ Less than 1 year
- ☐ 1-5 years
- ☐ 6-10 years
- ☐ 10+ years

19. In which areas are you doing research? [Categorized by European R&I Areas] *

- ☐ Health
- ☐ Culture and inclusive society
- ☐ Civil security
- ☐ Digital, Industry and Space
- ☐ Climate and Mobility
- ☐ Energy
- ☐ Food, Bioeconomy, Agriculture and Environment
- ☐ Soil, Waters, Oceans, Natural Resources
- ☐ Andere

20. Which types of AI do you work on? *

- ☐ General-Purpose AI (GPAI)
- ☐ Emotional AI
- ☐ Decision-support AI
- ☐ Image recognition AI
- ☐ Andere

21. How confident are you in identifying and addressing AI ethics concerns throughout a project's entire lifecycle (including the AI systems life cycle)? *

	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
Choose your level of confidence	☐	☐	☐	☐	☐

22. Which AI ethics concerns do you encounter? What makes them challenging to address? *

23. How confident are you in applying an 'ethics by design' approach within your workflow? *

	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
Choose your level of confidence	☐	☐	☐	☐	☐

24. What do you find most challenging about an ethics by design approach? *

25. How confident are you in complying with EU AI Act and other relevant regulations? *

	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
Choose your level of confidence	☐	☐	☐	☐	☐

26. Which of the following ethical principles highlighted in the EU AI Act are most challenging to comply with in practice? *

- ☐ **Human agency and oversight:** AI systems must support human autonomy and decision-making. This is particularly relevant for AI systems that can affect human cognition and behaviour.
- ☐ **Privacy and data governance:** AI systems must guarantee privacy and data protection throughout the system's lifecycle.
- ☐ **Transparency:** All processes associated with AI decision-making must be appropriately documented. AI systems must be explainable, and their limitations must be clearly communicated.
- ☐ **Fairness, diversity, and non-discrimination:** The best possible effort should be made to avoid unfair bias.
- ☐ **Societal and environmental well-being:** The impact of AI systems on society and the environment must be carefully evaluated. The risks of harm should be mitigated.
- ☐ **Accountability:** Actors involved in the development or operation of AI systems should clearly demarcate and assume responsibility for their functioning and outputs.

Subject to approval by the Granting Authority

Researcher (**not** in technical AI research)

27. How many years experience do you have in research? *

- ☐ Less than 1 year
- ☐ 1-5 years
- ☐ 6-10 years
- ☐ 10+ years

28. In which areas are you doing research? [Categorized by European R&I Areas] *

- ☐ Health
- ☐ Culture and inclusive society
- ☐ Civil security
- ☐ Digital, Industry and Space
- ☐ Climate and Mobility
- ☐ Energy
- ☐ Food, Bioeconomy, Agriculture and Environment
- ☐ Soil, Waters, Oceans, Natural Resources
- ☐ Andere

29. How confident are you in identifying and addressing AI ethics concerns throughout a project's entire lifecycle (including the AI systems life cycle)? *

- | | No confidence at all | Moderately unconfident | Neither confident or unconfident | Moderately confident | Very confident |
|---------------------------------|-----------------------|------------------------|----------------------------------|-----------------------|-----------------------|
| Choose your level of confidence | ☐ | ☐ | ☐ | ☐ | ☐ |

30. Which AI ethics concerns do you encounter? What makes them challenging to address? *

31. How confident are you in applying an 'ethics by design' approach within a project? *

	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
Choose your level of confidence	☐	☐	☐	☐	☐

32. What do you find most challenging about an ethics by design approach? *

33. How confident are you in complying with the EU AI Act and other relevant regulations? *

	No confidence at all	Moderately unconfident	Neither confident or unconfident	Moderately confident	Very confident
Choose your level of confidence	☐	☐	☐	☐	☐

34. Which of the following ethical principles highlighted in the EU AI Act are most challenging to comply with in practice? *

- ☐ **Human agency and oversight:** AI systems must support human autonomy and decision-making. This is particularly relevant for AI systems that can affect human cognition and behaviour.
- ☐ **Privacy and data governance:** AI systems must guarantee privacy and data protection throughout the system's lifecycle.
- ☐ **Transparency:** All processes associated with AI decision-making must be appropriately documented. AI systems must be explainable, and their limitations must be clearly communicated.
- ☐ **Fairness, diversity, and non-discrimination:** The best possible effort should be made to avoid unfair bias.
- ☐ **Societal and environmental well-being:** The impact of AI systems on society and the environment must be carefully evaluated. The risks of harm should be mitigated.
- ☐ **Accountability:** Actors involved in the development or operation of AI systems should clearly demarcate and assume responsibility for their functioning and outputs.

Research ethics educator

35. How many years experience do you have in teaching research ethics? *

- ☐ Less than 1 year
- ☐ 1-5 years
- ☐ 6-10 years
- ☐ 10+ years

36. Which domains does your research ethics teaching relate to? [Categorized by European R&I Areas] *

- ☐ Health
- ☐ Culture and inclusive society
- ☐ Civil security
- ☐ Digital, Industry and Space
- ☐ Climate and Mobility
- ☐ Energy
- ☐ Food, Bioeconomy, Agriculture and Environment
- ☐ Soil, Waters, Oceans, Natural Resources
- ☐ Andere

37. How confident are you teaching AI ethics? *

- | | No confidence at all | Moderately unconfident | Neither confident or unconfident | Moderately confident | Very confident |
|---------------------------------|-----------------------|------------------------|----------------------------------|-----------------------|-----------------------|
| Choose your level of confidence | ☐ | ☐ | ☐ | ☐ | ☐ |

38. What are the challenges you encounter teaching AI ethics? *

39. Which of the following ethical principles highlighted in the EU AI Act are most challenging to teach in practice? *

- ☐ **Human agency and oversight:** AI systems must support human autonomy and decision-making. This is particularly relevant for AI systems that can affect human cognition and behaviour.
- ☐ **Privacy and data governance:** AI systems must guarantee privacy and data protection throughout the system's lifecycle.
- ☐ **Transparency:** All processes associated with AI decision-making must be appropriately documented. AI systems must be explainable, and their limitations must be clearly communicated.
- ☐ **Fairness, diversity, and non-discrimination:** The best possible effort should be made to avoid unfair bias.
- ☐ **Societal and environmental well-being:** The impact of AI systems on society and the environment must be carefully evaluated. The risks of harm should be mitigated.
- ☐ **Accountability:** Actors involved in the development or operation of AI systems should clearly demarcate and assume responsibility for their functioning and outputs.

40. Which areas would you most like to learn about to improve your AI ethics teaching? *

- ☐ Basic AI concepts and terms
- ☐ More advanced technical knowledge on AI models and systems (including design measures to address ethical concerns)
- ☐ AI ethical concepts and terms (including philosophical foundations)
- ☐ Learning strategies, or pedagogical approaches, for teaching AI ethics.
- ☐ Approaches for assessing students' AI ethics learning
- ☐ Andere

SECTION: Training gaps and needs

41. Which areas of AI research would you most like to learn about? *

- ☐ General-Purpose AI (GPAI)
- ☐ Emotional AI
- ☐ Decision-support AI
- ☐ Image recognition AI
- ☐ Andere

42. Which of the following areas, related to the application of the EU AI Act and other EU guidance, would you like to learn more about? *

- ☐ Prohibited practices in AI
- ☐ High-risk AI systems
- ☐ General-purpose AI systems
- ☐ ALTAI principles and their operationalization
- ☐ Andere

43. If you had to select just three of the ethical principles highlighted in the EU AI Act to be prioritised in a time limited AI ethics training, which would you choose? *

Selecteer 3 opties.

- ☐ **Human agency and oversight:** AI systems must support human autonomy and decision-making. This is particularly relevant for AI systems that can affect human cognition and behaviour.
- ☐ **Privacy and data governance:** AI systems must guarantee privacy and data protection throughout the system's lifecycle.
- ☐ **Transparency:** All processes associated with AI decision-making must be appropriately documented. AI systems must be explainable, and their limitations must be clearly communicated.
- ☐ **Fairness, diversity, and non-discrimination:** The best possible effort should be made to avoid unfair bias.
- ☐ **Societal and environmental well-being:** The impact of AI systems on society and the environment must be carefully evaluated. The risks of harm should be mitigated.
- ☐ **Accountability:** Actors involved in the development or operation of AI systems should clearly demarcate and assume responsibility for their functioning and outputs.

44. If there are any other specific knowledge, skills or values that you would like to develop during an AI ethics training, tells us about them below. *

45. How would you prefer to learn? (select all relevant options) *

- ☐ Online via a short (e.g. 3 hour) synchronous seminar
- ☐ Online following asynchronous learning materials
- ☐ In-person class
- ☐ 'Learning by doing' (e.g through the use of tools (standard operating procedures, templates, flow charts) and best practice examples)
- ☐ Andere

46. Do you have prior training in AI ethics? *

- ☐ Yes
- ☐ No

47. Which training did you follow and what were its good and bad points? *

48. END OF SURVEY

Deze inhoud is niet door Microsoft gemaakt noch goedgekeurd. De gegevens die u verzendt, zal worden gestuurd naar de eigenaar van het formulier.

 Microsoft Forms

Subject to approval by the Granting Authority

Appendix II

Mapping Open Educational Content on AI Ethics: A Scoping Review.

Contributors

Menno Maris	Amsterdam UMC
Natalie Evans	Amsterdam UMC
Christina Tikva	University of Macedonia
Efthimios Tambouris	University of Macedonia
Alexei Grinbaum	CEA Saclay
Marieke Bak	Amsterdam UMC

Background

As AI is expected to increasingly affect human cognition and behavior, the need for practical ethics guidance has become increasingly critical (High-Level Expert Group on Artificial Intelligence [HLEG], 2019). The AIOLIA project responds to this challenge by aiming to operationalize AI ethics through the development of inclusive, interdisciplinary training materials, developed in close partnership with academic institutions, industry actors, and global policy networks. A key part of this effort involves identifying and developing accessible educational resources. This scoping review maps the current landscape of openly available AI ethics learning materials.

Rationale

A growing body of literature and guidelines has proposed ethical principles and values in response to the challenges posed by the rapid emergence of AI systems. However, translating these overarching frameworks into practice remains a challenge. The success of this operationalization depends particularly on the development of well-targeted educational resources. Although a wide range of AI ethics educational materials are openly accessible, their breadth, usability, and pedagogical quality have not yet been systematically examined.

Objectives

The objective of this scoping review is to identify and map openly available educational and training materials related to AI ethics. Specifically, the review will:

- Map resources that are free, publicly accessible, and explicitly designed for educational or training purposes.
- Examine their pedagogical design, including learning objectives, competences targeted, and instructional strategies.
- Identify the ethical concepts and AI-related issues addressed, as well as assessment and reflection components.
- Map the intended audiences (higher education students, professionals) and application domains.
- Assess the breadth, accessibility, and originality of such materials, with a focus on resources published from 2022 onward.
- Highlight strengths, gaps, and areas for further development to inform the AIOLIA project's creation of guidelines and training resources.

Context

We include all openly available educational and training materials on AI ethics that meet the eligibility criteria outlined below. There are no restrictions on the geographical origin of the materials or on the institutional sector in which they were produced (e.g., academic, professional, policy, or non-governmental organizations). The review primarily targets English-language materials, although relevant non-English resources will be included when accessible. Only materials published from 2022

onward are considered to ensure feasibility and relevance in this rapidly evolving field. This review will be conducted in accordance with the Preferred Reporting Items for Systematic reviews and Meta-Analyses extension for Scoping Reviews (PRISMA-ScR) guidelines (Tricco et al., 2018).

Eligibility criteria (inclusion/exclusion criteria)

Category	Inclusion Criteria	Exclusion Criteria
Topic relevance	Clear and substantive focus on AI ethics	AI content without an ethics component; general ethics content not sufficiently related to AI
Type of studies	Academic publications that provide or explicitly link to openly accessible educational resources on AI ethics (e.g., lesson plans, modules, training materials)	Academic publications that discuss AI ethics without linking to or including teaching resources; educational materials that are not openly or readily accessible (e.g., available only upon request, behind paywalls, or requiring institutional access).
Type of materials	Educational materials on AI ethics with defined learning objectives or outcomes, such as syllabi, modules, lessons, and assessments (e.g., MOOCs, OERs, slide decks, quizzes, assessment videos, web modules, case studies, infographics, podcasts, gamified learning platforms).	Informational resources without defined learning objectives or instructional design (e.g., single video lectures, opinion pieces, blogs)
Accessibility	Freely and openly available (no paywall) and directly accessible; e.g., Open Educational Resources	Paywalled content, “freemium” (essential features locked), time-limited access
Target audience	Materials designed for higher education students and/or professionals working with AI and ethics	General public
Language	Resources in English. Other languages may be included if the team can review them.	Resources in languages the team cannot review.
Period	Published within the past three years (2022-present)	Older materials
Intent	Clearly intended for teaching or training, with interactive components that support reflection. Structured materials, referring to content that includes components like learning outcomes, clearly organized modules, and assessments.	No clear instructional goal or pedagogical structure. Purely informative.
Quality (second phase, after inclusion)	Created by reputable sources (e.g. universities, NGOs, expert individuals)	Unverified (user-generated) content lacking credible backing

Originality (second phase, after inclusion)	Unique, original content	Duplicated or mirrored from another included source
---	--------------------------	---

Information sources

To ensure comprehensive scope of openly available AI ethics educational materials, multiple sources of evidence will be consulted. These include established academic databases, grey literature repositories, targeted online searches, and expert input from the AIOLIA consortium. Restricting the search to academic sources alone would risk overlooking relevant teaching and training materials, since many open educational resources are developed and distributed by universities, NGOs, professional associations, and industry partners outside traditional publishing channels.

- Academic databases: Web of Science and Google Scholar
- Grey literature from open educational resource repositories and institutional websites
- Web-based and non-indexed content located via targeted web searches (Google Advanced Search)
- Expert consultation with academic and industry partners from the AIOLIA consortium

Search

The search strategies were piloted at the end of August 2025. Draft strings were refined iteratively and discussed within the author group until consensus was reached. The Google Advanced Search strategy was developed in consultation with an expert from the Digital Methods Initiative at the University of Amsterdam, ensuring methodological rigor in the approach to web-based evidence.

Web of Science search string:

TI=(“artificial intelligence” OR AI OR “machine learning”) AND TI=(training OR trainer OR teach* OR course OR module OR e-learning OR learn* OR educat* OR lesson OR tutorial OR workshop OR webinar OR MOOC OR OER OR curriculum) AND TI=(ethic*)

Google Scholar and Google Advanced Search

Searches in Google Scholar will be limited to the .com domain. In both Google Scholar and Google Advanced Search, results will be restricted to the first 100 hits, following Mlake-Lye et al. (2025), to balance feasibility with capturing relevant studies. Google Advanced Search will be conducted using the Search Engine Scraper tool, developed by the Digital Methods Initiative at the University of Amsterdam (UVA), within a Mozilla Firefox browser configured for research purposes (e.g., cleared cache, no personalization, stable search engine settings).¹ Both searches will use the following search string:

((“artificial intelligence” OR “machine learning”) AND (“training” OR “trainer” OR “teaching” OR “course” OR “module” OR “e-learning” OR “learning” OR “education” OR “educational” OR “lesson” OR “tutorial” OR “workshop” OR “webinar” OR “MOOC” OR “OER” OR “curriculum”) AND (“ethics” OR “ethical”))

Selection of sources of evidence

¹ <https://digitalmethods.net/>

Screening will be conducted by a team of three reviewers. To ensure consistency, an initial pilot screening round (approximately 10% of records) will be carried out to test the clarity and feasibility of the inclusion and exclusion criteria. Based on this pilot, the criteria may be refined inductively by consensus within the review team before proceeding to full screening. In the main screening phase, titles and abstracts will be independently screened by two reviewers, with any disagreements resolved through discussion and, if needed, consultation with the third reviewer. The process will be documented in a PRISMA 2020 flow diagram (Page et al., 2021). Following an initial round of coding based on the inclusion criteria, a second phase quality and originality assessment will be carried out using a structured quality assessment form for educational materials.

Data charting process

Data charting will be conducted using a structured Excel sheet that includes fields for each data item. All three reviewers will be involved in the process. To begin, a pilot extraction will be conducted on approximately 15-20% of the included materials, with at least two reviewers independently charting each item. The pilot will be used to test the clarity of the data items and to refine the extraction framework inductively in consultation with the broader author group. For the main charting phase, the materials will be divided across the review team, with regular cross-checks and discussions to ensure consistency. Any modifications to the framework will be documented and agreed upon by the team before proceeding to full charting.

Data items

The data charting form will begin with a set of preliminary items structured as follows:

- Metadata: title, source, material type, year, accessibility, licensing, audience, and application domain
- Pedagogical approaches and learning strategies
- Learning objectives, competences or skills enhanced, and AI ethics concepts addressed
- Assessment features
- Tools

To pilot the extraction form, a sample of 15-20% of the included materials will first be charted. Alongside the predefined metadata, this process will help identify any additional relevant data items. The list of items will then be refined, following the principles of conventional qualitative content analysis (Green & Thorogood, 2018), in consultation with the broader author group, and may be further adjusted during the full review. After final inclusion, all materials will also be assessed for educational quality and originality.

Critical appraisal of individual sources of evidence

No formal critical appraisal will be conducted, as the purpose of this scoping review is to map available educational materials rather than evaluate methodological quality.

Data synthesis and presentation of results

Further meetings with the author group will be held to resolve any remaining questions about data charting. Once data extraction is complete, the categories will be clustered and organized into a matrix. Results will be presented in tables and figures where appropriate, in addition to narrative synthesis, to provide a clear overview of the characteristics and diversity of the included materials. The final synthesis and presentation of results may be adapted depending on the nature of the data, and will be determined in consultation with the author group and decided collectively.

Expected outputs

The expected outputs of this study will be: 1) this scoping review protocol; 2) project reporting (as part of deliverable 4.1); 3) a peer-reviewed review article; and 4) an openly available dataset of the included materials and their coded characteristics as supplementary material with the published article.

Review timeline

Start date: 1 November 2025. The search, screening and data extraction will be completed by 31 January 2026, along with internal reporting of the preliminary results. Manuscript preparation and submission will follow thereafter.

Ethics approval

Not applicable

Funding Declaration and Competing Interests

The AIOLIA project is funded by the European Commission under Grant Agreement 101187937. The review team declares that there are no competing interests.

References

Green, J., & Thorogood, N. (2018). *Qualitative Methods for Health Research* (4th ed.). Sage. ISBN, 1526448807, 9781526448804

High-Level Expert Group on Artificial Intelligence. (2019). *Ethics guidelines for trustworthy AI*. European Commission. <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>

Miake-Lye, I. M., Mak, S., Begashaw, M. M., & Shekelle, P. G. (2025). Using Google to search for evidence: How much is enough? One center's experience. *Systematic Reviews*, 14, Article 92. <https://doi.org/10.1186/s13643-025-02836-w>

Morley, J., Elhalal, A., Garcia, F., Kinsey, L., & Floridi, L. (2023). Ethics as a service: A pragmatic operationalisation of AI ethics. *AI and Ethics*, 3(1), 89-104. <https://doi.org/10.1007/s43681-021-00084-4>

Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., Mulrow, C. D., Shamseer, L., Tetzlaff, J. M., Akl, E. A., Brennan, S. E., Chou, R., Glanville, J., Grimshaw, J. M., Hróbjartsson, A., Lalu, M. M., Li, T., Loder, E. W., Mayo-Wilson, E., McDonald, S., ... Moher, D. (2021). The PRISMA 2020 statement: An updated guideline for reporting systematic reviews. *BMJ*, 372, n71. <https://doi.org/10.1136/bmj.n71>

Tricco, A. C., Lillie, E., Zarin, W., O'Brien, K. K., Colquhoun, H., Levac, D., Moher, D., Peters, M. D. J., Horsley, T., Weeks, L., Hempel, S., Akl, E. A., Chang, C., McGowan, J., Stewart, L., Hartling, L., Aldcroft, A., Wilson, M. G., Garritty, C., ... Straus, S. E. (2018). PRISMA extension for scoping reviews (PRISMA-ScR): Checklist and explanation. *Annals of Internal Medicine*, 169(7), 467-473. <https://doi.org/10.7326/M18-0850>

Subject to approval by the Granting Authority

Appendix III

Appendix A3 - preliminary inclusion table

Academic search

Aucouturier, E., & Grinbaum, A. (2025). Training Bioethics Professionals in AI Ethics: A Framework. Journal of Law, Medicine & Ethics, 53(1), 176-183.	Conceptual paper, does however, offer mock case materials, and outline for three hour interactive format for using the materials
--	--

Combined table Google advanced/MOOCs/Opencourseware/Consortium input

Provider / Resource	Target audience	Ethics focus	AI scope	URL
1. University of Helsinki - Ethics of AI (MOOC)	General public, AI users and developers	Societal impact, responsible development, ethical reasoning	Broad AI as societal technology	https://ethics-of-ai.mooc.fi/
2. Linux Foundation - Ethical Principles in Conversational AI (LFS118)	organizational leaders, product development teams, legal and ethical champions	Privacy, bias, children and vulnerable users, human rights	Chatbots and voice assistants	https://training.linuxfoundation.org/training/ethical-principles-in-conversational-ai-lfs118/
3. DataCamp - AI Ethics	Data and analytics learners	Fairness, transparency, accountability, privacy	General applied ML	https://www.datacamp.com/courses/ai-ethics
4. Stanford CS281 - AI Ethics resources	University students and instructors	Fairness, justice, truthfulness, societal impact	Machine learning case studies	https://stanfordaiethics.github.io/
5. British Council - AI and ethics in education	Learners 13 to 17 and adults, B2+ to C1, classroom use	Responsible classroom use, harms, guidelines	Educational AI	https://www.teachingenglish.org.uk/teaching-resources/teaching-secondary/lesson-plans/advanced-c1/ai-and-ethics-education
6. Kaggle - Intro to AI Ethics	Developers and data practitioners	Human centered design, bias, fairness, model documentation	ML deployment	https://www.kaggle.com/learn/intro-to-ai-ethics

7. Alan Turing Institute - Operationalising Ethics in AI	Businesses, organisations, practitioners with basic AI and regulatory understanding	Embedding ethics in design and lifecycle	Applied AI systems	https://www.turing.ac.uk/courses/operationalising-ethics-ai-intermediate
8. AI Ethics Course - Global Perspectives	Broad audience including data scientists, policymakers, business leaders	Social implications, global perspectives	General AI and data systems	https://aiethicscourse.org/
9. Canada School of Public Service - Ethical AI (DDN243)	Federal public servants (Canadian)	Accountability, transparency, governance, risk	Public sector AI	https://catalogue.cspsefpc.gc.ca/product?catalog=DDN243&cm_locale=en
10. fast.ai - Practical Data Ethics	Broad professional audience	Bias, data ethics, social harm	ML and data driven products	https://ethics.fast.ai/
11. SAP Learning - Putting AI Ethics into Practice	SAP employees and professionals (MOOC)	Governance, compliance, internal rules	Business AI	https://learning.sap.com/learning-journeys/putting-ai-ethics-into-practice-at-sap
12. BriGRETE shared materials	Researchers and students	Research ethics, responsible innovation	Academic AI R&D	https://www.brigrete.eu/file-share/86bb9f54-f30d-4538-9cd8-6e375a0d3d08
13. TU Delft - AI in Practice: Preparing for AI	Professionals and organisational stakeholders preparing to adopt AI	Ethics, compliance, governance in adoption	Applied organizational AI	https://www.edx.org/learn/artificial-intelligence/delft-university-of-technology-ai-in-practice-preparing-for-ai
14. Microsoft - AI for Business Leaders (AI-3017)	Business leaders, strategists, decision makers	Responsible AI, trust, governance	Business AI	https://learn.microsoft.com/en-us/training/courses/ai-3017
15. Linux Foundation - Conversational AI Compliance (LFS120)	Professionals driving trustworthy AI i.e., leaders in ethics, compliance, privacy, security,	Regulation, privacy, risk mitigation	Conversational AI	https://training.linuxfoundation.org/training/conversational-ai-ensuring-compliance-and-mitigating-risks-lfs120/

	engineers, developers, designers, product managers and trainers			
16. OER Commons - AI for tutoring and personalized learning	Educators and tech staff	Bias, privacy, overreliance	Educational AI tutors	https://oercommons.org/courseware/lesson/128665
17. Generative AI for Work Integrated Learning	Students, staff, industry partners	Disclosure, appropriate use, risk management	Generative AI in WIL	https://oercommons.org/courses/generative-artificial-intelligence-in-work-integrated-learning-resources-for-university-staff-students-and-industry-partners
18. Duke AI Ethics Learning Toolkit	Instructors and students	Privacy, environment, social impact	Cross disciplinary AI	https://oercommons.org/courses/duke-ai-ethics-learning-toolkit
19. MIT OCW - Ethics of AI Bias	Students and independent learners	Bias, discrimination, social context	Biased AI and ML (general)	https://ocw.mit.edu/courses/res-10-002-ethics-of-ai-bias-spring-2023/
20. MIT OCW - The Roosevelt Project	Policy oriented learners	Governance, societal tradeoffs, sustainability	Policy context	https://ocw.mit.edu/courses/res-14-003-the-roosevelt-project-spring-2023/
21. MIT OCW - Social and Ethical Responsibilities of Computing	Undergraduate and general learners	Privacy, surveillance, inequality, law, rights	Computing incl. AI	https://ocw.mit.edu/courses/res-tll-008-social-and-ethical-responsibilities-of-computing-serc/
22. University of Adelaide on edX - AI for Professionals	Industry professionals	Transparency, accountability, bias, IP	Generative AI at work	https://www.edx.org/learn/ethics/university-of-adelaide-ai-for-professionals-ethics-responsibility-and-best-practices
23. Institute of Science Tokyo on	Science and engineering learners	Avoiding harm, mistakes, and negative effects (“Preventive ethics”; “Aspirational ethics”	Data driven and AI systems	https://www.edx.org/learn/engineering/institute-of-science-

edX - AI and Data Ethics				tokyo-science-engineering-ai-data-ethics-ke-xue-ji-shu-ailun-li
24. KU Leuven on edX - AI in Healthcare: Hype or Help?	Healthcare professionals	Patient safety, regulation, responsible adoption	Clinical AI	https://www.edx.org/learn/artificial-intelligence/ku-leuven-ai-in-healthcare-hype-or-help
25. Université de Montréal on edX - Bias and Discrimination in AI	Learners focused on fairness	Discrimination, justice, societal harm	ML classification systems	https://www.edx.org/learn/artificial-intelligence/universite-de-montreal-bias-and-discrimination-in-ai
26. Pragmatic AI Labs on edX - Radical Ideas: AI Ethics	Broad professional audience	Human rights, surveillance, power	General AI in society	https://www.edx.org/learn/computer-science/pragmatic-ai-labs-radical-ideas-ai-ethics
27. MGH Institute of Health Professions on edX - Intro to AI in Healthcare	Healthcare professionals	Privacy, responsible use, overreliance	Medical/healthcare AI	https://www.edx.org/learn/data-analysis-statistics/mgh-institute-of-health-professions-introduction-to-ai-machine-learning-in-healthcare
28. University of Hong Kong on edX - FinTech Ethics and Risks	Finance professionals and students	Accountability, consumer harm, systemic risk	AI in financial services	https://www.edx.org/learn/fintech/university-of-hong-kong-fintech-ethics-and-risks
29. RWTH Aachen on edX - Responsible Innovators of Tomorrow	Interdisciplinary students and early professionals	Responsibility, sustainability, stakeholder impact	Responsible innovation incl. AI	https://www.edx.org/learn/innovation/rwth-aachen-university-responsible-innovators-of-tomorrow
30. Hamad Bin Khalifa University on edX - Applications of AI in Healthcare	Healthcare professionals	Oversight, benefits vs risks	Healthcare AI	https://www.edx.org/learn/computer-science/hamad-bin-khalifa-university-applications-of-ai-in-healthcare

31. INVOLV - AI voor gezondheid (<i>only in Dutch</i>)	Patient representatives	Inclusive design, participation, healthcare ethics	Healthcare AI	https://www.involv.nl/trainingen/ai-voor-gezondheid
---	-------------------------	--	---------------	---