



Workshop

How AI that uses pictures and text can be unfair — and what we can do

6 of June 2026

PRESENTATION OF THE PROJECT



Context

Main Objective

1 Building power with communities — learning, claiming and putting into practice the EU fundamental rights that AI systems must also respect and promote.

Citizens, Equality, Rights and Values Programme (CERV)

Background & general objectives:

- Tackle results in AI that uses pictures and text that create harm
- Follow EU rules on rights and on AI
- Protect human rights online
- Community appropriation of AI — understanding how it affects our lives, and taking back control

Key challenges:

- AI is taught with unfair examples
- Few ways to check if AI is fair /Fairness benchmarks
- People affected by AI rarely take part in the design or have the opportunity to use

Who we are

Project Management Team

CVC Lluís Gomez, Sofia Llàcer, Adrià Molina, Laura

ALIA/Donestech Eva Cruells i Mireia Orra

Diputació de Barcelona Sonia Ruiz

Advisory Board

Dimosthenis Karatzas, Alex Cadon Alba Elvira, Sara Aguilar, Marisa Molina (Diputació de Barcelona,) Anna Paricio (UOC) Núria Vergès (UB) Yuki M. Asano (Technische Universität Nürnberg) Titilola Olojede, Graziella Piga (Society of Gender Professionals)



Working Groups

WG Intersectional Impact Assessment Guidelines

Marta Cruells, Margarita Padilla, Nadia Nadesan & Maya Aidlin

WG Participatory Workshops

Gala Pin & Ágata Bañon

Participants and civil society organisations facilitating/contributing participatory Workshops

Objectives

1.

Participatory bias audit tool for AI that uses pictures and text.

Made to find and reduce unfair results in AI that uses pictures and text.

Built together with the people most affected by AI.

2.

Clear **fairness test /benchmark** to measure unfair results in AI that uses pictures and text.

Looks at how AI can be unfair to overlapping groups. Aims to be a standard way to measure fairness.

An open website where anyone can run the test.

3.

Use **ways of building AI with people** when creating the tool and the test.

Communities, community groups, and tech experts all help build and test the tools.

Hold 6 workshops with the community.

4.

Write clear **intersectional impact assessment guidelines**.

Ideas to keep social justice at the heart of the project. Help find and tackle deep-rooted unfairness.

We have set up a group of experts.

5.

Raise awareness and build capacity to protect basic rights when machines make decisions.

Outreach activities. We will share why fair AI matters.

Push for clear responsibility when machines affect people's rights.

Workshops Methodology

**Intersectional
feminist theory**

**Participatory
AI practices**

Core principles

Challenges

- Ⓟ Researchers and AI developers
- Ⓟ Civil society organizations
- Ⓟ Communities affected by AI
- Ⓟ Policy makers

Measuring unfair results in AI:

- Ⓟ Test AI by matching pictures with text
- Ⓟ Measure how fair the results are
- Ⓟ Look at how AI shows overlapping groups

Workshops

- 6 Workshops: 4 in **Barcelona**, 1 in Berlin, 1 in Warsaw (May-June)
 - People in care work and unstable jobs
 - Women over 55
 - LGBTIQ+ community
 - Young people from Roma community



What we will do together?

In this workshop, we will interact directly with AI systems that connect images and text

Together we will explore

- Who is represented and who is missing
 - Whether people are portrayed in a (un)fair and (un)respectful ways
 - Which stereotypes or assumptions may appear
 - How different groups may be affected by the results
 - **What a more inclusive, diverse and fair system could look like**
- Rather than relying only on technical measurements, we want this benchmark to reflect the experiences and perspectives of the people who are most affected by these technologies.



How your contributions will help?

Turning experiences into fairer AI evaluation

- ⑤ The observations and reflections collected during this workshop will help to build a better way of evaluating AI systems.
- ⑤ Your feedback will contribute to a public benchmark that measures how AI models of vision language perform across different communities and life experiences.
- ⑤ Interact → Reflect → Evaluate → Build a fairer benchmark

AI Models do not affect everyone in the same way (intersectional approach)

- ⑤ People experiences can differ depending on factors such as age, gender, ethnicity, disability, sexual orientation, migration background, administrative situation, social or economic situations, etc...
- ⑤ They interact and combine in different ways, shaping how people experience opportunities, barriers and discrimination.
- ⑤ Looking at people's experiences in context helps us uncover forms of unfairness that often go unnoticed.

There are no right or wrong answers.

What matters is your experience, your perspective and what you notice

Outputs: contributing to a more just AI

Intersectional Impact Assessment Guidelines

- Help build fairer AI
- Tackle deep-rooted unfairness in AI
- Push for AI design that is fair to all

Free Audit Tool

- Find and record unfair results in AI for pictures and text
- Allow open checks with the community
- Let anyone repeat the checks and see who is responsible

Public Fairness Benchmark /Scoreboard

- Compare how well AI works for overlapping groups
- Show gaps in AI that uses pictures and text
- Encourage AI makers to improve fairness

Contact us!
intervisions@donestech.net

6 of June 2026