

Fairness Position Statement

**BETTER4U · AI FAIRNESS
WORKSHOP · 1 April 2026**

Date: 03.06.2026



Funded by
the European Union



UK Research
and Innovation

Project funded by



Schweizerische Eidgenossenschaft
Confédération suisse
Confederazione Svizzera
Confederaziun svizra
Swiss Confederation

Federal Department of Economic Affairs,
Education and Research EAER
State Secretariat for Education,
Research and Innovation SERI

Fairness Position Statement

BETTER4U - AI FAIRNESS WORKSHOP - 1 April 2026

A Fairness Position Statement is a formal record of a project team's commitments on how they will identify, address, and monitor fairness risks in their AI system. It documents the team's conscious choices about how fairness is defined and measured, and establishes accountability for following through on those choices throughout the project lifecycle.

For the BETTER4U project, this document serves three purposes. First, it records the key fairness risks identified at the AI Fairness Workshop held on 1 April 2026. Second, it sets out the Consortium's commitments on how those risks will be addressed. Third, it provides the kind of documented evidence of practices of risk management and transparency required under the EU AI Act (Art. 9, 13, and 50) and consistent with the EU Ethics Guidelines for Trustworthy AI.

The draft prepared by WP2 leader (UNIVIE) is based on the output from the BETTER4U AI Fairness Workshop, 1 April 2026. This is a living document and may be updated if required.

Topic	Statement
1. System description	BETTER4U develops an AI-based causal model that generates personalized recommendations to support weight management for adults living with overweight or obesity. Recommendations are delivered through the BETTER4U intervention platform and inform healthcare professionals supporting participants during the randomized clinical trial (RCT).
2. Affected patient groups	Adults living with overweight or obesity participating in the BETTER4ALL RCT. Characteristics most relevant to fairness: sex and gender, age, disability status, socioeconomic status (SES) and education level, ethnic and migrant background, geographic location and presence of comorbidities and concurrent treatments, including obesity management medications.
3. Training data risks	The training data may over-represent people of higher SES, people without disabilities, local language speakers and people without obesity, and may not adequately represent the RCT target population. Known gaps include younger and older adults, people of lower SES, migrants, people with disabilities and people living with greater overweight and obesity. The data may also reflect historical biases affecting people living with obesity, people of lower SES, migrants, women and those

	with poor prior healthcare access. Additionally, the lifestyle and habits of training cohorts may differ from the RCT target population. Proposed mitigations: widen training data to include underrepresented groups; audit datasets for stigma-related proxies and biased labels; review missing data by BMI group; consider patient and lived experience input in reviewing model features and outputs; use the most recent cohorts; recalibrate models with recent RCT baseline data; train and validate the AI model specifically on people living with overweight and obesity.
4. Model design risks	The AI model design risks are generally considered low; the most significant is that patients with comorbidities, and lower SES may not be adequately represented in the feature set, which may affect AI model performance for these groups. Patients taking obesity management medications are excluded from the RCT. Future development of the AI model beyond the present project should consider including these groups to ensure accurate performance. Sensitivity analyses and broadening variable coverage to include a more comprehensive set of vulnerable groups are proposed mitigations for the currently included population.
5. Validation risks	The highest risk identified is that participants may be unable to implement the AI model's recommendations due to socioeconomic or structural constraints. Healthcare professional involvement in delivering recommendations is a main mitigation.
6. Fairness definition and metrics commitment	Fairness in the BETTER4U AI-based causal model is defined as individual fairness — users with similar health profiles receive similar personalized recommendations regardless of their protected characteristics. This was chosen because BETTER4U provides individualized rather than population-level recommendations. The consortium acknowledges that this definition does not guarantee equal outcomes across demographic groups and commits to complementing it with group-level performance monitoring, prioritizing people living with overweight and obesity, the primary target population, as well as ethnic groups and other underrepresented groups. The underrepresentation of people living with obesity in the retrospective data used to train the AI model is a high-priority gap, which the consortium expects to reduce as RCT data becomes available.