

WP3**Værktøjskasse til algoritmisk bias til synkrone sessioner****MCQ-quizzer****10 enkeltvalgsspørgsmål med 3 valgmuligheder pr. CU****INSTRUKTIONER TIL UPLOAD AF SPØRGSMÅL
TIL ET SOFTWARE, DER GENERERER ONLINE TESTS**

1. Find et software til at lave en quiz med enkeltvalgsspørgsmål.
2. I tekstfeltet til spørgsmål skal du kopiere/indsætte hvert spørgsmål.
3. I tekstfeltet til svarmuligheder skal du kopiere/indsætte hver mulighed.
4. Glem ikke at definere løsningen (det korrekte svar) for hvert spørgsmål.
5. For hvert spørgsmål tilføj en instruktion:
 - a. Forslag: "Læs spørgsmålet og vælg det korrekte svar."
 - Tjek, hvordan navigationen mellem spørgsmål fungerer i dit software, og tilføj i instruktionen noget i retning af "...og klik på SEND-knappen." eller "...og klik på NÆSTE-knappen."
6. Sæt minimum beståelsesprocent for testen (vi anbefaler 60 % – brugeren skal svare rigtigt på mindst 6 ud af 10 spørgsmål).
7. Hvis det er muligt i softwaren, definer en tidsbegrænsning for testen (vi anbefaler 12 minutter).

Tjek hvilke andre indstillinger der er tilgængelige i dit software. Du kan muligvis definere ting som antallet af forsøg brugeren har til at gennemføre testen, blandt andre muligheder.

Kompetenceenhed 4 – Dataretfærdighed og bias

SPØRGSMÅL

Spørgsmål nr.		Spørgsmål og svarmuligheder	Korrekt svar
1	Spørgsmål	Hvilken type AI-bias opstår som følge af uoverensstemmelser i, hvordan data indsamles?	
	Mulighed 1	Dækningsbias.	X
	Mulighed 2	Observatør bias.	
	Mulighed 3	Survivorship Bias.	
2	Spørgsmål	Hvilken retfærdighedsmetrik sigter mod at udligne andelen af positive resultater for forskellige grupper?	
	Mulighed 1	Demografisk paritet.	X
	Mulighed 2	Lige muligheder.	
	Mulighed 3	Kalibrering efter gruppe.	
3	Spørgsmål	Hvilken metode bruges almindeligvis til at opdage bias i træningsdata til AI-systemer?	
	Mulighed 1	Algoritmeoptimering.	
	Mulighed 2	Datavisualisering.	
	Mulighed 3	Datarevision.	X
4	Spørgsmål	Hvad indebærer konceptet 'fairness gennem uopmærksomhed' i AI-retfærdighed?	
	Mulighed 1	At sikre, at alle brugerdata anonymiseres før behandling.	
	Mulighed 2	At ignorere følsomme attributter under modeltræning.	X
	Mulighed 3	At anvende den samme model til alle grupper uden tilpasning.	
5	Spørgsmål	Hvad refererer 'fjendtlig debiasing' til i maskinlæring?	
	Mulighed 1	En metode til at forbedre modellens nøjagtighed ved at introducere konkurrerende modeller.	
	Mulighed 2	En teknik, hvor modeller trænes til at forudsige og modvirke potentielle bias.	X
	Mulighed 3	En strategi til at forsvare mod ondsindede angreb på AI-systemer.	
6	Spørgsmål	Hvad er effekten af 'label bias' i overvågede læringsmodeller?	
	Mulighed 1	Det forbedrer modellens evne til at generalisere fra træningsdata.	

	Mulighed 2	Det reducerer modellens beregningskompleksitet.	
	Mulighed 3	Det skævvrider modellens forudsigelser baseret på fejlagtigt mærkede træningsdata.	X
7	Spørgsmål	Hvilken tilgang er mest effektiv til at mindske bias på dataniveau i AI-systemer?	
	Mulighed 1	At øge modellens forudsigelseshastighed.	
	Mulighed 2	At forbedre kompleksiteten i modellens arkitektur.	
	Mulighed 3	At sikre diversitet og repræsentativitet i træningsdatasættet.	X
8	Spørgsmål	Hvordan hjælper 'counterfactual fairness' i AI-beslutningstagning?	
	Mulighed 1	Ved at øge effektiviteten af AI-systemer i realtidsmiljøer.	
	Mulighed 2	Ved at garantere, at en beslutning er retfærdig, hvis den ville være den samme i en kontrafaktisk verden, hvor en følsom attribut er anderledes.	X
	Mulighed 3	Ved at sikre, at AI's beslutninger forbliver konsistente over tid.	
9	Spørgsmål	Hvad adresserer 'intersectionality' i forbindelse med AI-retfærdighed?	
	Mulighed 1	Den samlede påvirkning af flere følsomme attributter på beslutningstagning.	X
	Mulighed 2	Interaktionen mellem forskellige AI-modeller inden for et system.	
	Mulighed 3	Krydsningen af etiske og juridiske grænser i AI-anvendelser.	
10	Spørgsmål	Hvilken rolle spiller 'algoritmisk gennemsigtighed' i at fremme retfærdighed i AI-systemer?	
	Mulighed 1	Det sikrer hurtigere beregning og behandling i AI-systemer.	
	Mulighed 2	Det gør det muligt for interessenter at forstå, hvordan AI-beslutninger træffes, og at identificere potentielle bias.	X
	Mulighed 3	Det fokuserer udelukkende på at forbedre den grafiske brugergrænseflade for AI-applikationer.	