

HÅNDTERING AF BIAS I BIG DATA TIL AI OG MACHINE LEARNING

How to use this booklet?

This document is interactive. Throughout the document, you will find links to additional information.



Button that takes you to the beginning of the document. This button appears on the top right corner of the pages.



Whenever you see this icon, it means that you have an **interactive color text** to click on, that has an external link associated to it.



Whenever you see this icon, it means that you have an **interactive color text** to click on, that will open a pop-up with additional information.

DISCLAIMER: Please note that we cannot guarantee the continued availability of external content, such as videos, as they may be subject to change or removal by its authors or host platforms.

Index

Click on the menu

- 01. Introduction
- 02. Fairness & bias - the candy analogy
- 03. Understanding the landscape: AI fairness and bias
- 04. Bias in AI systems
- 05. Methods for identifying and measuring biases in AI
- 06. Strategies for addressing and mitigating biases
- 07. Understanding fairness metrics
- 08. Fairness in data collection and preprocessing
- 09. Emerging trends in ai fairness
- 10. Conclusion

Velkommen tilbage til CHARLIE-projektets nyhedsbrev!

Denne måned sætter vi fokus på "Algoritmisk Bias" EQF6-kurset, designet til videregående uddannelsesinstitutioner, der ønsker at forbedre deres studerendes forståelse af AI-etik. Med algoritmer, der bliver stadigt mere integreret i samfundet—inden for uddannelse, sundhedsvæsen, finansiering og mere—er det afgørende at adressere de skjulte fordomme i dem.

Featured Course: Algorithmic Bias

"Algoritmisk Bias" EQF6-kurset er et omfattende uddannelsesprogram, der adresserer en af de mest presserende udfordringer ved AI—bias inden for algoritmer. I dagens data-drevne verden bestemmer algoritmer ofte, hvem der får adgang til job, lån, sundhedspleje og endda uddannelse. Imidlertid kan fordomme snige sig ind i disse systemer, hvad enten det er en jobannonce, der vises uforholdsmæssigt til mænd, eller en sundhedsalgoritme, der favoriserer visse racer frem for andre. Dette kursus sigter mod at udruste deltagerne med færdigheder og viden, der er nødvendige for at identificere, afhjælpe og håndtere sådanne fordomme.

Hvorfor dette kursus er vigtigt

Algoritmisk bias er ikke kun et teknisk problem—det har vidtrækkende sociale og etiske konsekvenser. Selvom algoritmer ofte ses som neutrale værktøjer, kan deres design, datainput og implementering introducere utilsigtet diskrimination. For eksempel bliver jobannoncer med højere løn ofte vist mere til mænd end kvinder, og sundhedssystemer kan yde bedre pleje til visse racer. Disse fordomme truer ikke kun retfærdigheden, men også den offentlige tillid til AI-systemer, hvilket gør det afgørende for teknologiprofessionelle og undervisere at forstå disse risici.

Kurset er struktureret omkring fem Kompetenceenheder (Competency Units - CU), som hver er designet til at give deltagerne de nødvendige værktøjer til at navigere og bekæmpe algoritmisk bias i virkelige scenarier:

CU1: Algorithms, Models, and Limitations

Lær grundlæggende om algoritmer, deres begrænsninger og den rolle, de spiller i beslutningstagning på tværs af sektorer. Forstå de underliggende mekanismer i algoritmer og hvordan deres design kan introducere bias utilsigtet.

CU2: Data Fairness and Bias in AI

Dyk ned i kilderne til bias i AI-systemer, herunder de data, der bruges til at træne disse algoritmer. Deltagerne vil lære teknikker til at identificere, måle og reducere bias i data, hvilket sikrer mere retfærdige resultater.

CU3: AI Privacy and Convenience

Udforsk den delikate balance mellem at opretholde brugerens privatliv og tilbyde belejlighed i AI-applikationer. Denne enhed understreger vigtigheden af at beskytte brugerdata samtidig med at skabe effektive AI-systemer.

CU4: AI Ethics – A Practical Approach

Lær at anvende etiske rammer på virkelige AI-projekter. Denne enhed dækker gennemsigtighed, ansvarlighed og deltagende design teknikker for at sikre etisk AI-udvikling.

CU5: Case Studies

I denne afsluttende enhed vil deltagerne anvende deres viden gennem praktiske projekter og case-studier, hvilket giver praktisk erfaring i at adressere algoritmisk bias i forskellige sammenhænge.

Målgruppe: EQF6 Studerende

"Algoritmisk Bias"-kurset er målrettet mod studerende på videregående uddannelser (bachelorniveau) og professionelle, der søger at fordybe deres forståelse for etisk AI. Kurset er i overensstemmelse med det Europæiske Kvalifikationsrammeverk (European Qualifications Framework (EQF – 6), som kræver avanceret viden og kognitive færdigheder, samt en høj grad af autonomi og ansvar.

Ved at gennemføre dette kursus vil studerende og professionelle udvikle kritiske kompetencer, såsom:

- **Advanced Knowledge**

En dyb forståelse af, hvordan bias opstår i AI-systemer og den indvirkning det har på tværs af forskellige sektorer. Lærende vil udforske fairness-metrikker, etiske overvejelser og tekniske tilgange til at reducere bias.

- **Cognitive Skills**

Evnen til at analysere virkelige tilfælde af algoritmisk bias, vurdere fairness i AI/ML-systemer og udvikle strategier for at afhjælpe disse biases i praksis.

- **Practical Skills**

Implementering af teknikker til detektion og afhjælpning af bias, evaluering af deres effektivitet og integration af etiske principper i udviklingen af AI-systemer.

- **Autonomy and Responsibility**

Udrustning af eleverne til at træffe informerede, etiske beslutninger vedrørende AI/ML-systemer, under hensyntagen til både samfundsmæssige virkninger og individuelle ansvar.

- **Communication and Collaboration**

Udvikle evnen til at forklare komplekse spørgsmål om algoritmisk bias for både tekniske og ikke-tekniske interessenter, hvilket fremmer samarbejde på tværs af faggrænser.

- **Lifelong Learning**

Engagement i at holde sig opdateret med de seneste udviklinger inden for AI-etik og bias-afhjælpning, sikrer kontinuerlig faglig vækst.

Bliv Engageret!

Er du spændt på fremtiden for etisk AI-uddannelse? Tjek vores hjemmeside og følg vores kanaler på de sociale medier for at holde dig opdateret om den officielle lancering af vores ressourcer og kommende muligheder for at blive involveret. Vi vil også dele indsigter og nyheder fra AI-etikkens verden, så hold øje!

Lad os sammen opbygge en fremtid, hvor AI tjener det fælles bedste!

<https://charlie-project.uib.es/>

