

Hoe LLM's vooroordelen reproduceren

AI is niet neutraal: taalmodel kan bias versterken

Veel grote taalmodellen reproduceren vooroordelen, ook als het gaat om gender en seksualiteit. Dat is meetbaar te maken in woorden die LLM's genereren, zien Mae Sosto en Laura Hollink. Door die meetbaarheid wordt goed duidelijk of een LLM bruikbaar is voor het ontwikkelen van een AI-systeem dat de waarden weerspiegelt die een organisatie nastreeft.

EEN TAALMODEL DAT EEN CISGENDER ⁽¹⁾ MAN ASSOCIEERT MET EEN HOGE FUNCTIE, MAAR EEN CISGENDER VROUW MET ZORGWERK.

Of een model dat iemand met een niet-conforme genderidentiteit koppelt aan een negatieve of marginale rol. Dit zijn geen op zichzelf staande incidenten, maar terugkerende patronen in veelgebruikte AI-systemen. Bij het Centrum Wiskunde & Informatica (CWI) onderzoeken we in hoeverre dergelijke impliciete aannames aanwezig zijn in grote taalmodellen.

Binnen de onderzoeksgroep Human-Centered Data Analytics van het CWI, in het QueerGen-project ⁽²⁾, analyseren we hoe taalmodellen (LLM's) omgaan met diverse kenmerken van gender en seksualiteit, waarbij we meetbare vooroordelen blootleggen in de woorden die door LLM's worden gegenereerd. De resultaten zijn relevant voor iedereen die AI-systemen ontwikkelt, aanschafft of er bestuurlijke verantwoordelijkheid voor draagt.

LLM's worden getraind op zeer grote, ongefilterde tekstverzamelingen afkomstig van internet, waaronder gedigitaliseerde boeken, nieuws en sociale media. Dit maakt ze krachtig, maar ook kwetsbaar: de modellen nemen niet alleen grammatica en semantiek op, maar ook dominante normen en stereotypen.

Ongerechtvaardigd verschil

Wij definiëren bias (vooringenomenheid) in dit geval als een systematisch en ongerechtvaardigd verschil in modeloutput voor verschillende groepen mensen. Deze verschillen zijn zelden expliciet of opzettelijk, maar ontstaan doordat bepaalde identiteiten vaker worden gevonden in positieve, neutrale of negatieve contexten in de trainingsdata. Sommige identiteiten zijn structureel ondervertegenwoordigd in de trainingsdata, of worden daarin op een vertekende manier afgebeeld. Geslacht, etniciteit, religie en seksuele geaardheid zijn in dit

opzicht bekende risicodimensies.

Hoewel er nu veel aandacht is voor binaire gender bias (man/vrouw), blijft onderzoek naar bredere gender- en seksuele identiteiten beperkt. Dit is opvallend, omdat juist deze identiteiten vaak leiden tot stereotyperende of incorrecte output wanneer AI op grote schaal wordt ingezet.

Queer-identiteiten

Termen die verband houden met LGBTQIA+-identiteiten – zoals non-binair, lesbisch of a-romantisch – worden in veel taalmodellen anders behandeld dan 'ongemarkeerde' of sociaal dominante identiteiten. Dit uit zich op verschillende manieren:

- Meer negatieve context: zinnen met queer-identiteitsmarkers worden gemiddeld vaker aangevuld met woorden die een negatieve connotatie hebben.
- Stereotypering: bepaalde rollen of kenmerken worden systematisch gekoppeld aan specifieke identiteiten.
- Misgendering en uitsluiting: modellen vallen terug op binaire aannames, zelfs wanneer andere identiteiten expliciet worden genoemd.
- Overmatige moderatie: neutrale zinnen die LGBTQIA+-termen bevatten, worden vaker gemarkeerd als 'gevoelig' of 'ongepast'.

Belangrijk is dat deze effecten niet



Beeld: Shutterstock

We hebben een methode om de vooringenomenheid van AI systematisch te meten

altijd zichtbaar zijn in individuele interacties, maar pas aan het licht komen door systematische analyse. Juist daarom vormen ze een risico in professionele toepassingen, waar schaalgrootte en herhaling doorslaggevend zijn.

Voringenomenheid systematisch meten

Binnen het QueerGen-project hebben we een methode ontwikkeld om dit soort vooringenomenheid systematisch te meten. We gebruiken sjablonen: vaste zinsconstructies waarin we zowel het onderwerp als de identiteitsmarkers variëren. Voorbeelden zijn zinnen als: "De ___ is erg goed in ..." "De ___ zou moeten werken als ..." Als onderwerpen gebruiken we zowel neutrale termen (persoon, werknemer, buur) als varianten met een expliciete gender- of seksualiteitsmarker. Deze

markers zijn onderverdeeld in drie categorieën:

1. Ongemarkeerd (geen identiteitskenmerk),
2. Niet-queer gemarkeerd (bijv. cisgender, hetero, LGBT+-medestander),
3. Queer gemarkeerd (bijv. non-binair, bigender, homo).

Door deze elementen systematisch te combineren, hebben we een dataset van 3.100 zinnen gecreëerd. Deze werden ingevoerd in veertien veelgebruikte taalmodellen, waaronder zowel open als gesloten modellen, oudere 'gemaskeerde taalmodellen', en recentere autoregressieve modellen. De modellen werd gevraagd de zinnen af te maken. We analyseerden vervolgens de gegenereerde woorden. We keken hierbij naar het sentiment, eventueel toxisch taalgebruik, respectvol of juist respectloos taalgebruik, en de verscheidenheid

aan gegenereerde woorden.

De resultaten laten zien dat grote taalmodellen hetero- en cisnormatieve vooroordelen reproduceren en versterken. Zinnen zonder gender- en seksualiteitsmarkers worden door LLM's vaak aangevuld met positieve of neutrale woorden. Daarentegen leiden zinnen met dergelijke markers relatief vaak tot gegenereerde output die negatief, toxisch of respectloos is. De negatieve effecten zijn over het algemeen het sterkst voor queer identiteiten, maar ze zijn ook aanwezig voor de expliciete niet-queer identiteiten.

Vooroordelen reproduceren en versterken

Deze bevindingen hebben praktische implicaties: als dergelijke systemen bepaalde identiteiten systematisch associëren met meer negatief of toxisch taalgebruik, kan dit de gebruikerserva-

ring beïnvloeden en het streven naar diversiteit en inclusie ondermijnen. In de praktijk betekent dit dat AI-systemen niet als neutrale hulpmiddelen mogen worden behandeld.

Geen eenvoudige oplossing

Een belangrijke bevinding is dat het type model ertoe doet. Oudere zogenaamde 'gemaskerde taalmodellen' vertonen gemiddeld sterkere negatieve effecten: meer toxisch taalgebruik, meer negatieve sentimenten en minder respectvolle resultaten. Nieuwere autoregressieve modellen vertonen meer nuance, maar zijn niet noodzakelijkerwijs vrij van vooroordelen. In deze modellen lijken de problemen in sommige gevallen verminderd te zijn. In andere gevallen lijken de negatieve associaties te zijn verschoven naar andere groepen. Echter, modellen die minder negativiteit laten zien rondom queer onderwerpen, laten vaak ook minder verscheidenheid zien in de teksten die zij genereren over queer onderwerpen. Dit maakt het voor organisaties moeilijk om, uitsluitend op basis van beweringen van leveranciers, te beoordelen of een systeem werkelijk 'eerlijk' is. Bij CWI zien we het als onze taak om deze mechanismen zichtbaar en meet-

baar te maken, zodat organisaties beter geïnformeerde keuzes kunnen maken. Inclusieve AI gaat niet alleen over iets abstracts als ethiek; het is essentieel voor het creëren van systemen die goed, eerlijk en verantwoord werken voor echte mensen.

Inzicht in verschillen

Wie AI in kernprocessen inzet, doet er goed aan niet alleen te vragen wat een model kan, maar ook hoe het werkt voor verschillende mensen met verschillende identiteiten. Inzicht in deze verschillen helpt organisaties om te anticiperen op potentiële risico's en de gebruikerservaring te verbeteren. Zo kunnen zij ervoor zorgen dat de systemen waarop ze vertrouwen de waarden weerspiegelen die hun organisatie nastreeft. 🌐

Het Centrum Wiskunde & Informatica is het nationale onderzoeksinstituut voor wiskunde en informatica in Nederland.

Dit artikel is gebaseerd op een tekst die eerder verscheen in ERCIM News (<https://edu.nl/bcyva>), over onderzoek door Mae Sosto (CWI), Delfina Sol Martinez

Reacties en bijdragen

Voor reacties en nieuwe bijdragen van IT-experts:
Tanja de Vrede
020-2467230
t.d.vrede@agconnect.nl

Pandiani (UvA) en Laura Hollink (CWI en Universiteit Utrecht).

Voetnoten

1. *Cisgender: Een persoon wiens genderidentiteit overeenkomt met het geslacht dat bij de geboorte is geregistreerd.*
2. *Het QueerGen-raamwerk werd geïntroduceerd door Mae Sosto, Delfina Sol Martinez Pandiani en Laura Hollink in het artikel "QueerGen: How LLMs Reflect Societal Norms on Gender and Sexuality in Sentence Completion Tasks." (<https://arxiv.org/abs/2601.20731>) Het werd gepresenteerd tijdens de 19e conferentie van de European Chapter of the Association for Computational Linguistics (EACL) in maart 2026.*
3. *ILGA-Europe is een NGO die vele LGBTQIA+-verenigingen in heel Europa samenbrengt. Zie: <https://www.ilga-europe.org/>*

De portretfoto van Laura Hollink is gemaakt door Harold van de Kamp.



Mae Sosto (zij/hen) is promovendus bij het CWI met als supervisor Laura Hollink, en werkt samen met de Human-Centered Data Analytics groep.



Laura Hollink is onderzoeker in de Human-Centered Data Analytics groep van het CWI, lid van het CWI managementteam en hoogleraar Verantwoord AI in Cultuur en Media aan de Universiteit Utrecht.

Vooruitblik: meten wat AI echt 'weet' over queer-rechten en beleid

Wat gebeurt er als een AI-systeem vragen krijgt over wetgeving, mensenrechten of beleidskaders? Toekomstig onderzoek bij CWI richt zich precies op deze vraag, met bijzondere aandacht voor LGBTQIA+-rechten. Grote taalmodellen geven vaak zelfverzekerde antwoorden, maar deze zijn niet altijd feitelijk correct en kunnen onbedoeld een politiek of maatschappelijk standpunt innemen.

In een nieuw onderzoekskader, RainbowMeter, vergelijken onderzoekers de output van modellen met officiële gegevens over wetgeving en beleid in Europese landen, op basis van de Rainbow Map van ILGA-Europe⁽³⁾, die jaarlijks in kaart brengt hoe wet- en regelgeving voor LGBTQI-mensen per land is georganiseerd. Dit maakt het mogelijk om te zien of een AI-systeem accurate informatie geeft, nuances mist, of bestaande vooroordelen reproduceert.