

Waarom het algoritme jouw boeken negeert (inct.spiratie 2026)

29-09-2026 09:00



Aanbevelingsalgoritmes lijken objectief, maar trekken in de praktijk vaak ongemerkt de populairste boeken voor. Onderzoeker Savvina Daniil onthult tijdens haar closing keynote op inct.spiratie 2026 de verborgen vooroordelen van deze AI-systemen en toont hoe het boekenvak de regie over de eigen catalogus terugpakt.

De opmars van kunstmatige intelligentie in ons dagelijks leven is niet meer te stuiten. Zeven jaar geleden, toen de in Griekenland opgeleide wiskundige Savvina Daniil naar Amsterdam kwam voor een master in Data Science, was AI nog een stuk minder alomtegenwoordig. Ook binnen haar vakgebied lag de nadruk destijds vooral op het optimaliseren van de technische prestaties van systemen.

“De focus lag heel sterk op nauwkeurigheid”, vertelt Daniil in gesprek met inct. Hoe goed kan een systeem voorspellen wat iemand doet of wil? Inmiddels is dat volgens haar veranderd. Destijds onderzochten wetenschappers al kritisch de maatschappelijke gevolgen van AI, maar inmiddels gebeurt dat veel explicieter: voor wie werkt zo'n systeem eigenlijk goed, wie wordt erdoor benadeeld en welke keuzes zitten er achter de uitkomsten?

Na haar studie werkte Daniil twee jaar bij TNO, binnen een team dat zich bezighield met verantwoorde AI. Daar ging het onder meer over privacy, transparantie en de betrouwbaarheid van systemen. Zelf raakte ze vooral geïnteresseerd in inclusiviteit: bedienen algoritmes werkelijk verschillende groepen gebruikers, of werken ze

vooral goed voor de grootste of meest dominante groep?

Cultuur, ethiek en de bibliotheek

Die belangstelling bracht Daniël uiteindelijk bij het Cultural AI Lab, een samenwerkingsverband waarin academische kennisinstellingen en culturele organisaties samen onderzoek doen naar kunstmatige intelligentie. Vanuit dat netwerk begon zij aan onderzoek naar aanbevelingssystemen in de bibliotheeksector, gefinancierd door de Koninklijke Bibliotheek en uitgevoerd bij het Centrum Wiskunde & Informatica (CWI). De aanleiding was volgens Daniël heel concreet. De KB zag de mogelijkheden van aanbevelingssystemen, maar wilde tegelijkertijd weten hoe deze technologie op een verantwoorde manier kon worden ingezet. Daarbij speelde een uitgangspunt uit de ethische AI-principes van de KB een belangrijke rol: wie inclusieve systemen wil bouwen, moet rekening houden met vooringenomenheid in de onderliggende data. Die moet worden gecompenseerd of op zijn minst zichtbaar worden gemaakt. Dat werd het vertrekpunt voor het project Responsible Recommenders in the Public Library Sector.

Een algoritme is niet objectief

Een van de misverstanden die Daniël graag uit de wereld helpt, is het idee dat algoritmes neutrale rekenmachines zijn die vanzelf tot één juiste uitkomst komen.

Bij aanbevelingssystemen speelt juist veel onzekerheid en kansberekening een rol, legt ze uit. Systemen moeten voortdurend vereenvoudigingen maken om uit gigantische hoeveelheden mogelijke opties een aanbeveling te produceren.

Daniël gebruikt tijdens het gesprek een eenvoudig voorbeeld. Wanneer een AI-systeem op een afbeelding moet herkennen of er een hond of een kat staat, bestaat er in principe één juist antwoord. Bij het aanbevelen van boeken aan lezers is dat anders. Er is nooit één boek waarvan objectief kan worden vastgesteld dat een bepaalde lezer dat op een bepaald moment móét lezen.

Wanneer een bibliotheekstelsel een titel aanbeveelt en de gebruiker die vervolgens leent, kan het algoritme die aanbeveling als succesvol registreren. Maar, benadrukt Daniël, daarmee weten we nog helemaal niet of dit daadwerkelijk de beste aanbeveling was. Er waren misschien duizenden andere boeken waarmee die gebruiker minstens zo tevreden zou zijn geweest.

Ook keuzes van ontwikkelaars beïnvloeden de uitkomst. Welke gegevens worden gebruikt? Welke factoren wegen zwaar? Hoe wordt succes gedefinieerd? Verander één van die variabelen en er kan een heel andere aanbeveling uitrollen.

Bias begint al bij de data

Volgens Daniël ontstaat vooringenomenheid bovendien niet pas in het algoritme. Ook de data waarmee het werkt, is nooit volledig neutraal.

“Data wordt geconstrueerd”, benadrukt ze. Een boekencollectie is immers het resultaat van allerlei eerdere keuzes. Welke boeken zijn ingekocht? Welke auteurs worden uitgegeven? Welke titels zijn gepromoot? Welke boeken zijn door media of sociale platforms populair gemaakt?

Als een collectie bijvoorbeeld relatief veel Amerikaanse auteurs bevat, zal een systeem dat op basis van die collectie aanbevelingen doet eveneens relatief veel Amerikaanse auteurs adviseren. Het algoritme weerspiegelt daarmee niet simpelweg ‘de werkelijkheid’, maar ook eerdere redactionele, commerciële en culturele keuzes.

Voor uitgevers, bibliotheken en boekverkopers is juist dat volgens Daniël belangrijk om te beseffen: aanbevelingssystemen bouwen voort op beslissingen die al eerder in de keten zijn genomen.

Twee manieren om boeken aan te bevelen

Tijdens het interview legt Daniël ook uit hoe aanbevelingssystemen doorgaans werken. Grofweg bestaan er twee benaderingen.

Bij content-based filtering kijkt het systeem naar eigenschappen van een boek: auteur, publicatiejaar, samenvatting, onderwerp of bijvoorbeeld moeilijkheidsgraad. Vervolgens vergelijkt het die gegevens met boeken die iemand eerder heeft gelezen.

Maar zelfs ogenschijnlijk objectieve metadata bevatten menselijke keuzes, zegt Daniil. Een kwalificatie als 'moeilijk' of 'toegankelijk' is bijvoorbeeld geen natuurkundige grootheid, maar een interpretatie.

De tweede methode is collaborative filtering. Daarbij kijkt het systeem vooral naar het gedrag van groepen gebruikers. Wanneer twee lezers in het verleden grotendeels dezelfde boeken hebben gekozen, kan het algoritme aannemen dat zij vergelijkbare voorkeuren hebben. Ontdekt de ene lezer vervolgens een nieuw boek, dan kan dat ook aan de andere worden aanbevolen.

Volgens Daniil werken grote commerciële platforms tegenwoordig meestal met combinaties van dergelijke technieken. Maar welk model ook wordt gebruikt: mensen bepalen welke signalen belangrijk zijn en hoe zwaar ze meetellen.

De valkuil van popularity bias

Daniil's eigen onderzoek richt zich in het bijzonder op popularity bias. Daarbij krijgen titels die al populair zijn automatisch een steeds grotere kans om opnieuw te worden aanbevolen. Dat gebeurt niet zonder reden. Vanuit het perspectief van een algoritme zijn bekende boeken relatief veilige aanbevelingen. Gebruikers herkennen de titel en klikken of lenen daardoor sneller. Dat kan een systeem op papier zelfs bijzonder accuraat laten lijken. Maar juist daar zit volgens Daniil het probleem.

Een aanbevelingssysteem dat steeds dezelfde populaire titels laat zien, heeft weinig toegevoegde waarde. Bibliotheken en boekverkopers beschikken juist over grote catalogi, waarvan veel titels nauwelijks zichtbaar worden. De kracht van aanbevelingssystemen zou volgens haar moeten liggen in het verbinden van die zogeheten long tail aan lezers voor wie juist die minder bekende boeken interessant kunnen zijn.

Tijdens haar onderzoek zag Daniil bovendien dat popularity bias andere vormen van ongelijkheid kan versterken. Populaire boeken bleken bijvoorbeeld relatief vaak afkomstig van auteurs uit de Anglosfeer, waaronder de Verenigde Staten, Groot-Brittannië, Canada en Australië.

Een algoritme hoeft dus helemaal geen nationaliteit van auteurs te gebruiken om toch een voorkeur voor bepaalde groepen te ontwikkelen. Alleen het extra gewicht dat aan populariteit wordt gegeven, kan al tot zo'n uitkomst leiden.

Bias hoeft niet altijd verkeerd te zijn

Tegelijk waarschuwt Daniil ervoor bias automatisch als iets negatiefs te beschouwen. Om vast te stellen of een afwijking problematisch is, moet je eerst bepalen welk doel een organisatie met een aanbevelingssysteem wil bereiken. Ze noemt tijdens het interview het voorbeeld van onderzoek bij Deense bibliotheken. Daar bleek een algoritme opvallend vaak boeken van Deense auteurs aan te bevelen.

Technisch gezien was dat een duidelijke bias. Dat hoeft voor de bibliotheek niet per se een probleem te zijn. Het onder de aandacht brengen van Deense literatuur behoorde juist tot de doelstellingen. Anders werd het toen Daniil constateerde dat binnen die groep vooral reeds bekende, populaire Deense schrijvers naar voren kwamen. Debutanten en minder bekende auteurs kregen daardoor nauwelijks zichtbaarheid. De ontwikkelaars van het systeem vonden dat wel ongewenst.

Je hebt dus altijd een soort referentiepunt nodig, legt Daniil uit. Pas wanneer duidelijk is wat je met een systeem wilt bereiken, kun je beoordelen of een bepaalde bias wenselijk of problematisch is.

Een dashboard voor de boekenwereld

Die nuance wil Daniil nu ook praktisch bruikbaar maken. In haar huidige onderzoek werkt ze aan een dashboard waarmee organisaties inzicht moeten krijgen in het gedrag van aanbevelingssystemen. Niet iedere bibliothecaris, uitgever of boekverkoper heeft immers de technische kennis om een algoritme volledig te

analyseren. Het dashboard moet daarom zichtbaar maken welke voorkeuren of afwijkingen in een systeem optreden en wat daar de mogelijke gevolgen van zijn.

Het uitgangspunt is nadrukkelijk niet dat een computer vervolgens bepaalt of een systeem 'goed' of 'fout' is. Daniil wil juist dat professionals zelf die afweging kunnen maken.

Een dashboard kan bijvoorbeeld laten zien dat een bepaald aanbevelingssysteem relatief veel populaire titels naar voren haalt, dat lokale auteurs minder zichtbaar worden of dat een bepaalde groep gebruikers minder goed wordt bediend. Vervolgens is het aan de organisatie om te bepalen of dat past bij haar eigen doelstellingen.

Daarmee verschuift de discussie over AI volgens Daniil van een puur technische naar een strategische vraag: welke waarden wil je als organisatie eigenlijk in je digitale etalage terugzien?

Closing keynote op inct.spiratie 2026

De impact van aanbevelingssystemen reikt ver voorbij de muren van de openbare bibliotheek. Voor uitgevers en boekverkopers ligt er een grote kans – én verantwoordelijkheid – in het begrijpen van de algoritmes die steeds vaker bepalen welke content consumenten te zien krijgen.

Tijdens inct.spiratie 2026 zal Savvina Daniil de closing keynote verzorgen. Ze neemt het publiek mee onder de motorkap van AI en laat aan de hand van concrete praktijkvoorbeelden zien hoe aanbevelingen tot stand komen, welke biases daarin kunnen ontstaan en vooral welke keuzes organisaties zelf kunnen maken.

[MEER INFO & AANMELDEN](#)

David Huijzer