

AI Ergo Sum

Essays over onze veranderende relatie met technologie



OPEN
PRESS
TiU

Hannah van den Bosch (red.)

Met bijdragen van Judith Künneke, Julija Vaitonytė, Jan de Wit, Lucas Jones, Frank Bosman,
William Arfman, Barend de Rooij, Inge van de Ven en Clara Daniels

AI Ergo Sum

Essays over onze veranderende relatie met technologie

Hannah van den Bosch (red.)

Colofon

Titel

AI Ergo Sum: Essays over onze veranderende relatie met technologie

Redactie

Hannah van den Bosch

Auteurs

Judith Künneke • Julija Vaitonytė • Jan de Wit • Lucas Jones • Frank Bosman
William Arfman • Barend de Rooij • Inge van de Ven • Clara Daniels

Open Press Tilburg University

Uitgegeven door Open Press Tilburg University, Warandelaan 2, 5037 AB Tilburg, Nederland. <https://tiu.trialanderror.org>



Dit is een Open Access-boek dat is gepubliceerd onder de voorwaarden van de Creative Commons Naamsvermelding 4.0 International-licentie (CC BY 4.0). Deze licentie staat toe dat het werk wordt gedeeld en bewerkt, voor elk doel, ook commercieel, mits correcte naamsvermelding van de maker plaatsvindt.

Zie <http://creativecommons.org/licenses/by/4.0>

Vormgeving en opmaak

Ilske Schoonbrood

Omslagafbeelding

iStock.com/marcoventuriniautieri

ISBN: 9789403870809

DOI: 10.56675/9789403870809

Gratis te downloaden via: <https://doi.org/10.56675/9789403870809>

© 2026

Hannah van den Bosch, Judith Künneke, Julija Vaitonytė, Jan de Wit, Lucas Jones, Frank Bosman, William Arfman, Barend de Rooij, Inge van de Ven, Clara Daniels

Op initiatief van Studium Generale
www.tilburguniversity.edu/studiumgenerale

Notitie van de redacteur

Veel dank aan Annelieke Koster en Has Klerx voor het proeflezen van het manuscript, aan Daan Rutten en Beatriz Lourenço Barrocas Neves Ferreira voor hun hulp gedurende het proces van het tot stand brengen van dit boek, en aan Ilske Schoonbrood en Maud van der Windt voor hun werk in de vormgeving en opmaak.



Inhoudsopgave

Voorwoord: Het belang van Bildung in een tijd van AI <i>Hannah van den Bosch</i>	9
De Metaverse en de toekomst van het hoger onderwijs <i>Judith Künneke</i>	17
Waakzaam en optimistisch blijven in het tijdperk van AI-Agents <i>Julija Vaitonytė</i>	29
Van code naar karakter: Een lichaam voor AI? <i>Jan de Wit</i>	37
Dokter ChatGPT: Wat verstaan we onder een medisch hulpmiddel? <i>Lucas Jones</i>	45
Deus in Machina <i>Frank Bosman & William Arfman</i>	55
Zeg ChatGPT, heb ik iets van je te vrezen? <i>Barend de Rooij</i>	65
Een vervormende spiegel: Ik als AI als onbetrouwbare verteller <i>Inge van de Ven & Clara Daniels</i>	77
Over de auteurs	89

Voorwoord: Het belang van Bildung in een tijd van AI

Voorwoord: Het belang van Bildung in een tijd van AI

Hannah van den Bosch

https://doi.org/10.56675/9789403870809_1

Wie zijn wij, nu kunstmatige intelligentie met ons meedenkt, meepraat en meebepaalt? En misschien nog urgenter: wie worden wij, nu technologie steeds vaker zelfs voor ons denkt?

De afgelopen drie jaar heb ik me als programmamaker bij Studium Generale bezighouden met het thema digitalisering. Ik sprak met wetenschappers, studenten, kunstenaars over privacy, automatisering, beïnvloeding, sociale media en digitale cultuur. Wat begon als een brede verkenning van de ethische uitdagingen in onze digitale informatiesamenleving, verschoof al snel naar één urgent thema dat boven andere onderwerpen uitstak: de opkomst van AI (Artificial Intelligence; kunstmatige intelligentie). Deze technologie, waar onder andere de inmiddels immens populaire ChatGPT valt, wordt door sommige experts de grootste maatschappelijke verschuiving genoemd sinds de komst van het internet.¹

Wat me opviel bij de lezingen, symposia en andere programma's die ik organiseerde, was dat er twee kampen lijken te zijn ontstaan. Het ene kamp juicht de opkomst van AI toe en ziet vooral kansen. Voor hen is AI een bevrijdende kracht, een slimme assistent die ons productiever maakt, saaie taken wegneemt en nieuwe vormen van creativiteit mogelijk maakt. Het andere kamp is huiverig: zij vrezen verlies van menselijkheid en oordeelsvermogen, afvlakking van creativiteit, en de vaardigheid om zelf te denken. Zij vrezen ook een maatschappelijke verschuiving die ons, haast ongemerkt, afhankelijk maakt van systemen die we niet helemaal begrijpen en daardoor niet meer goed kunnen controleren.

De mens als cyborg

Dus, schuren we langs een afgrond, of opent zich juist een veelbelovende toekomst? Ik ben van mening dat we AI niet in zulke zwart-wit scenario's moeten benaderen. Het is wat mij betreft geen puur goede of puur slechte kracht.

De opkomst van AI is eerder een volgende stap in een lange geschiedenis van technologische ontwikkelingen. De technologie is er nu, dat staat vast. De belangrijke vraag is daarom niet of we AI moeten gebruiken, maar hoe en op welke manier. Welke keuzes maken we als samenleving in de manier waarop we deze systemen willen inzetten, vormgeven en begrenzen?

We gebruiken technologie namelijk al eeuwenlang om ons werk uit handen te nemen. Van de ploeg en het rad tot de stoommachine, van weefgetouwen tot wasmachines en stofzuigers. Steeds opnieuw beloofden nieuwe uitvindingen dat we meer tijd zouden krijgen voor het leven. Maar tegelijkertijd raken we ook steeds meer verstrengeld met diezelfde machines. Soms gaat het zelfs zo ver dat we onszelf door hun logica beginnen te zien: als iets wat geoptimaliseerd moet worden, data die verwerkt moet worden, een systeem dat beter of efficiënter zou kunnen en moeten draaien. Op die manier trekt het versnellende ritme van technologische innovatie ons ook die kant op. Zonder dat we het merken, worden we daardoor ook steeds een beetje meer machine.²

Ook in het onderwijs zie je dat terug. Technologie wordt vaak gepresenteerd als vooruitgang, als iets dat neutraler, objectiever, efficiënter is dan de mens. Maar elke technologie draagt een waardensysteem in zich mee, expliciet of impliciet. AI benadrukt bijvoorbeeld snelheid, correctheid, berekenbaarheid. Het belooft een maakbare en meetbare wereld. Als je dat lang genoeg gebruikt, begin je jezelf ooit onvermijdelijk langs dezelfde meetlat te leggen.

In dit spanningsveld komt Donna Haraway's cyborg-denken tot leven. Zij schreef decennia geleden al haar beroemde Cyborg Manifesto, waarin ze betoogde dat wij allang geen zuiver biologisch wezen meer zijn, maar een hybride: deels mens, deels machine, deels verbeelding.³ Ze zag technologie niet als iets dat buiten ons staat, maar als iets dat ons mede vormt. In die zin zijn we al cyborgs.

Haraway's visie wordt nu, in een tijd van AI, tastbaarder dan ooit. We leven met technologie, we denken ermee, we verhouden ons ertoe. Maar Haraway waarschuwde ook: als we cyborgs zijn, moeten we wel goed nadenken over welk soort cyborg we willen worden. Niet een slaaf van de machine, niet een verlengstuk van het algoritme, maar een wezen dat technologie gebruikt om nieuwe vormen van menselijkheid en verbeelding mogelijk te maken.

Als we de mens zien zoals Haraway hem schetst, als een technologisch verweven identiteit, dan wordt één ding duidelijk: juist in een wereld waarin we steeds meer cyborg zijn, maar waarin we ook willen voorkomen dat wij compleet verdwijnen in de kaders van de systemen die ons mede vormen, moeten we autonomie blijven behouden over hoe we die systemen willen inzetten, vormgeven en begrenzen. Dat vraagt om bewust ontwikkelde autonomie en kritisch oordeelsvermogen. Kwaliteiten die niet vanzelf ontstaan, maar geoefend moeten worden. Hier komt Bildung in beeld. Waar dit bij Wilhelm von Humboldt en Goethe sterk verbonden was met een vaststaand, conservatief ideaal van persoonlijke vervolmaking, benadrukken andere interpretaties, zoals Georg Forster en Walter Benjamin juist het leren je te verhouden tot een complexe en gedeelde wereld die je mede vormt.⁴ In die laatste betekenis gebruik ik het begrip hier.

Bildung als tegenkracht in een cyborg-tijdperk

Ik noem Bildung hier natuurlijk niet voor niets. Het is namelijk de onderliggende visie van de vele lezingen en symposia die we bij Studium Generale organiseren. Daarbij streven we ernaar de samenhang tussen wetenschappelijke disciplines zichtbaar te maken, studenten te stimuleren in hun persoonlijke ontwikkeling en hun gevoel voor maatschappelijke verantwoordelijkheid te versterken. Vanuit de levensbeschouwelijke traditie van Tilburg University betekent dit ook nadenken over de waarden die wetenschap richting geven.⁵

De universiteit dus als omgeving waar niet alleen kennis wordt overgedragen, maar waar jonge mensen zich vormen tot individuen die verantwoordelijkheid durven nemen, oog hebben voor het gemeenschappelijke goed en weten waar ze zelf voor staan. Bildung is daarbij de praktijk van kritische zelfvorming, van leren nadenken in plaats van nadoen, van ontdekken wat je belangrijk vindt in een wereld vol prikkels en verwachtingen.

Ik denk dat je Bildung in deze tijd kan zien als een persoonlijk kompas.

Want als AI is opgebouwd uit statistiek, patronen, en herhalingen van al wat al bestaat; een optelsom van het verleden, van alles wat er allemaal al ooit is gezegd of gedacht, hoeveel ruimte laat het dan voor het nieuwe, het ongekende, het eigenzinnige? Bildung als kompas kan het vermogen stimuleren om ook eigen keuzes te blijven maken, los van verwachtingen, algoritmen en optimalisatieloga. Het vraagt om het gesprek dat je met jezelf voert, het gesprek dat Hannah Arendt zag als de kern van denken.⁶

Bildung verlangt soms ook juist iets wat haaks staat op technologische snelheid: afstand. Een pas op de plaats. Een moment waarin je niet meebuigt met het ritme van de machine. Niet omdat de technologie slecht is, maar omdat vorming ruimte nodig heeft om te ademen, te twijfelen, te spelen, precies datgene wat een algoritme ons niet kan voorschrijven. En dat is wat ons kan helpen om dat andere soort cyborg te blijven: niet de slaaf van de machine die zijn denken uitbesteedt aan algoritmes, maar een vormgever die bepaalt wanneer en hoe technologie wordt ingezet: kritisch, nieuwsgierig en autonoom.

AI als spiegel

Tegelijkertijd opent AI in relatie tot het begrip Bildung ook een onverwachte uitnodiging. Namelijk om ons juist opnieuw af te vragen hoe rijk, creatief en eigenzinnig menselijk denken eigenlijk kan zijn. Want als AI soms hallucinant, associatief of verrassend uit de hoek komt, hoe kunnen wij dat dan (opnieuw) durven? In die zin is Bildung precies dat: een oefening om niet samen te vallen met de dominante technologie van je tijd, maar ernaast te kunnen staan, ertegenin te denken, erdoorheen te breken.

Bildung kan helpen onderscheid te maken. Het leert ons niet alleen hoe we met technologie moeten omgaan, maar ook wanneer we ervan moeten loskomen. Want AI verandert onze voorwaarden voor denken. Het geeft antwoorden vóór we zelf de tijd nemen om goede vragen te stellen. Het neemt taken over die ooit oefening, frustratie en vorming boden. Het spreekt met een zelfverzekerdheid die vertrouwen afdwingt, terwijl dat vertrouwen niet altijd verdiend is. AI praat met ons maar het praat ons ook vooral na. Het versterkt bestaande normen, patronen en dominante stemmen.

Als we Bildung inzetten om ons denken richting te geven, blijven we genoeg ruimte houden voor afwijking, verwarring, weerstand. Want gevormd worden vraagt frictie. En frictie laat zich moeilijk automatiseren. Dat is precies de reden dat deze bundel vanuit het idee van Bildung is samengesteld. Als programmamaker bij Studium Generale vertrek ik vanuit die overtuiging dat de universiteit een plaats is waar jonge mensen zich vormen tot denkers, die verantwoordelijkheid durven nemen voor zichzelf en voor elkaar. Technologie hoort daarbij, maar nooit zonder reflectie.

AI Ergo Sum

De titel van deze essaybundel “AI Ergo Sum” is een knipoog naar Descartes’ *cogito ergo sum*. Bij hem stond denken gelijk aan bestaan.⁷ In deze bundel ontstaat een nieuwe vraag: wat betekent het om te *zijn* in een tijd waarin machines met ons meedenken en soms steeds meer *voor* ons denken? Wanneer wij steeds meer onderdeel zijn van de technologie die we gebruiken? Deze bundel is een uitnodiging om die vraag niet uit handen te geven aan een algoritme, maar zelf te blijven stellen. Elk essay is een oefening in datgene wat Bildung beoogt: kritisch vragen stellen, perspectieven verbinden, jezelf situeren tegenover een wereld die verandert.⁸

Zeven essays

Voor deze bundel heb ik Tilburg University wetenschappers benaderd met verschillende disciplines, van filosofie tot economie, van psychologie tot technologie, van recht tot cultuurwetenschappen, die zich, ieder vanuit hun eigen vakgebied, bezighouden met digitalisering. Deze keuze is ook gemaakt vanuit een Bildungsbenadering van digitalisering: AI raakt niet één domein, maar snijdt dwars door alle onderdelen van wetenschap en samenleving heen.

Ik vroeg ze of ze een essay wilden schrijven waarin zij, vanuit hun expertise, “in gesprek gaan” met een technologie die een rol speelt in hun onderzoek of dagelijks leven. Het was aan hen om te bepalen met wie of wat ze het gesprek aangingen, en hoe ver ze daarin wilden gaan. Sommigen kozen voor een heel concrete technologie, anderen voor een abstracter digitaal fenomeen. Het resultaat is een rijke verzameling essays, elk geworteld in verschillende wetenschappelijke tradities.

Wat gebeurt er met leren en academische vorming wanneer onderwijs zich verplaatst naar virtuele werelden, zoals de Metaverse? Accountingwetenschapper Judith Künneke gaat deze vraag na en laat in het eerste essay zien welke kansen en spanningen daarbij ontstaan. Haar essay laat zien hoe digitale omgevingen onze ideeën over aanwezigheid, identiteit en gemeenschap fundamenteel kunnen veranderen.

Cognitiewetenschapper Julija Vaitonytė schrijft daarna over hoe we optimistisch kunnen blijven in het tijdperk van AI-*agents*. Ze legt uit waarom we onze kritische denkvaardigheden niet mogen verliezen, nu we steeds meer samenwerken met en taken delegeren aan deze systemen. Ze pleit voor een balans tussen technologische hoop en menselijke waakzaamheid.

Wanneer AI-systemen menselijke trekken krijgen, verandert ook de manier waarop we ze begrijpen. Communicatiewetenschapper Jan de Wit onderzoekt het in het derde essay hoe snel we geneigd zijn intentie, empathie en persoonlijkheid toe te schrijven aan avatars en belichaamde AI, en wat dat betekent voor onze omgang met deze technologie.

Jurist Lucas Jones verdiept zich vervolgens in “dokter ChatGPT” en wat we verstaan we onder een medisch hulpmiddel. Hij verkent hierbij belangrijke juridische en ethische vragen die ontstaan rond AI in de zorg en laat zien hoe lastig het is om systemen die adviseren, voorspellen of diagnosticeren onder te brengen in bestaande definities, met grote gevolgen voor onze veiligheid en aansprakelijkheid.

Theoloog Frank Bosman & religiewetenschapper William Arfman voeren een theologisch-filosofisch gesprek met AI. Ze leggen bloot hoe technologie steeds vaker raakt aan religieuze beelden, symboliek en verlangens en vragen zich af wat dat eigenlijk zegt over ons mensbeeld en onze zoektocht naar richting.

In gesprek met AI als existentiële spiegel? Dat is wat filosoof Barend de Rooij doet in zijn essay. Hij zoekt daarbij naar een middenweg tussen utopische en dystopische verwachtingen omtrent AI, en zet vraagtekens bij de betrouwbaarheid van een systeem dat altijd antwoordt en zelden twijfelt.

Aan het einde van de bundel onderzoeken literatuur- en cultuurwetenschapper Inge van de Ven en research masterstudent linguïstiek Clara Daniels hoe AI fungeert als een vervormende spiegel; hoe het onze verhalen aanpast, soms subtiel, soms radicaal, en hoe juist die vertekening onthult wie wij zijn, wat we proberen te zien en welke blinde vlekken we liever negeren.

Hannah van den Bosch, programmamaker Studium Generale, Tilburg University.

Noten

- ¹ De Morgen. (2023, 27 juni). Welkom in het AI-tijdperk: 'Ik acht het zeer realistisch dat AI belangrijker wordt dan de uitvinding van de elektriciteit'. <https://www.demorgen.be/tech-wetenschap/welkom-in-het-ai-tijdperk-ik-acht-het-zeer-realistisch-dat-ai-belangrijker-wordt-dan-de-uitvinding-van-de-elektriciteit~bgdac1a1/>
- ² Nieuwenhuis, R. (2023). *Mens en machine: Waarom we techniek de baas moeten blijven*. Ambo|Anthos.
- ³ Haraway, D. J. (1991). *Simians, cyborgs, and women: The reinvention of nature*. Routledge.
- ⁴ Verheijen, J. (2025). *Revolutie in de schoolgang: Radicaal-romantische Bildung in en buiten het onderwijs tussen 1789 en nu*. Amsterdam University Press.
- ⁵ Studium Generale. (2022). *Beleidsvisie Studium Generale 2022–2026: Definitieve versie 2022-09-15*. Tilburg University.
- ⁶ Arendt, H. J. (1978). *The life of the mind: Thinking* (Vol. 1). Harcourt.
- ⁷ Descartes, R. (2011). *Over de methode: Band III: Dioptriek, meteoren, geometrie* (J. Holierhoek, Vert.; 1e druk). Boom. ISBN 9789085066590
- ⁸ Verheijen, J. (2025). *Revolutie in de schoolgang: Radicaal-romantische Bildung in en buiten het onderwijs tussen 1789 en nu*. Amsterdam University Press.

De Metaverse en de toekomst van hoger onderwijs

De Metaverse en de toekomst van hoger onderwijs

Judith Künneke

https://doi.org/10.56675/9789403870809_2

Elk jaar bedenken veel mensen een goed voornemen voor het nieuwe jaar. In de top drie staat vaak “meer bewegen” en ik besloot me bij de meerderheid aan te sluiten. Wetende dat mijn discipline wat betreft sporten bijna nihil is, vroeg ik me af hoe ik mezelf kon motiveren om toch actiever te worden. Ik vroeg ChatGPT om ideeën en het antwoord was meteen duidelijk: *gamification*! Ik hou al mijn hele leven van gamen, dus twee dagen later stond ik bij de kassa van de lokale MediaMarkt, blij, met een *Meta Quest*-headset in mijn handen. Na installatie begon ik meteen allerlei meldingen te ontvangen over “evenementen”: een concert van Sabrina Carpenter en Tiësto, of zelfs een sportles. Waar was ik de afgelopen jaren geweest, dacht ik. Ik was verbaasd hoeveel er te beleven viel in deze virtuele omgevingen.

Het duurde niet lang voordat mijn gedachten naar het onderwijs afdreven. Wandel vandaag de dag over een universiteitscampus en je ziet studenten die laptops, smartphones en headsets tegelijk gebruiken. Technologie is diep geïntegreerd in het academische leven. Toch reikt het concept van de metaverse—een digitale omgeving waarin mensen volledig kunnen opgaan, leren, socialiseren en zelfs nieuwe identiteiten aannemen—verder dan wat we tot nu toe gewend zijn.¹ De Metaverse belooft een enorme uitbreiding van onze digitale mogelijkheden. In plaats van passief informatie te absorberen, kunnen studenten virtuele laboratoria, bedrijven of complete werelden betreden. Met alleen een headset en een betrouwbare internetverbinding kan iedereen simulaties ervaren, die vroeger alleen met dure apparatuur of verre reizen toegankelijk waren.

Op het eerste gezicht klinkt dit als een perfect hulpmiddel voor hoger onderwijs. Als docent en academisch directeur zie ik het enorme potentieel van leren door te doen en het toepassen van theorie. Deze immersieve omgevingen kunnen complexe concepten tastbaarder maken, diepere betrokkenheid stimuleren en samenwerking over geografische grenzen heen vergroten. Maar naarmate onze digitale aanwezigheid groter wordt, moeten we ook vragen stellen: Hoe beïnvloeden deze

meeslepende ervaringen ons vermogen om in de “echte wereld” te functioneren? Kan het gevoel van “aanwezigheid” in de Metaverse zo meeslepend worden dat studenten het fysieke klaslokaal of hun medestudenten in de echte wereld uit het oog verliezen? Bevorderen we echte nieuwsgierigheid, of bieden we slechts een nieuwe vorm van escapisme? Kan immersieve technologie leiden tot gevoelens van privacy inbreuk bij studenten? In dit essay wordt betoogd dat de Metaverse het leren kan verrijken, maar dat docenten ook ethische, psychologische en privacyvraagstukken moeten adresseren om een verantwoord gebruik te waarborgen.

Een blik op het Metaverse-klaslokaal

Stel je een mastercursus bedrijfsstrategie voor, waarin studenten een *virtual reality-headset* opzetten voor een groepssimulatie. In plaats van een voorbeeld casus te lezen en een schriftelijk verslag in te leveren, bevinden ze zich in een levendige 3D-marktplaats. Elke student heeft een andere rol: één beheert de *supply chain*, een ander is verantwoordelijk voor marketing, en weer een ander volgt de financiën. Samen moeten ze een virtueel product lanceren, onderhandelen over grondstofprijzen en *in realtime* reageren op veranderende consumentenbehoeften. Net als bij een *multiplayer-videogame* volgt het platform elke beslissing en laat het zien hoe keuzes invloed hebben op marktaandeel, winst en merkreputatie. Studenten ervaren de dynamiek van vraag en aanbod in plaats van erover te lezen. Terwijl ze marketingcampagnes plannen of financieringsopties afwegen, zien ze de virtuele markt direct reageren. Een dergelijke meeslepende omgeving kan abstracte bedrijfsconcepten veel dynamischer illustreren dan traditionele tekstboeken of presentaties. Alleen al bij de gedachte eraan verwacht ik dat de betrokkenheid in zo’n setting enorm toeneemt. Studies suggereren inderdaad dat leren via *virtual reality* en simulaties motivatie en retentie kan verhogen. Een meta-analyse vond dat VR-gebaseerd leren kan leiden tot hogere betrokkenheid en betere leerresultaten dan traditionele methoden.² Reuze interessant, toch?

De uitdagingen van leven in twee werelden

De grenzen van de werkelijkheid

Deze opwindende roep roept ook belangrijke vragen op. Eén van de grootste zorgen bij het gebruik van de Metaverse in het onderwijs is het risico dat individuen “te veel” opgaan in de virtuele wereld. Als een digitale omgeving spannender, beter beheersbaar of bevredigender lijkt dan de fysieke werkelijkheid, kunnen studenten geleidelijk virtuele ruimtes verkiezen boven menselijke interactie.³ In zulke gevallen kan

het verhoogde gevoel van aanwezigheid, oftewel de psychologische ervaring fysiek aanwezig te zijn in een virtuele omgeving, de scheidslijn tussen simulatie en realiteit vervagen.⁴

De voorkeur voor “overmatige onderdompeling” kan ertoe leiden dat sommige studenten virtuele prestaties of scenario’s belangrijker vinden dan doelen in de echte wereld. Als leerlingen directe feedback en visueel belonende simulaties aantrekkelijker vinden dan de complexiteit van offlinestudie of samenwerking, kunnen ze onderdelen van de cursus die persoonlijke aanwezigheid vereisen, verwaarlozen. Door gemak en controle te bieden, kunnen *virtual reality* technologieën onbedoeld stimuleren om echte interacties te vermijden, vooral die waarbij conflicten, debat of emotionele uitdagingen komen kijken.

Digitale identiteiten en psychologisch welzijn

Een kernaspect van Metaverse-omgevingen is de vrijheid om avatars te creëren die ideale versies van jezelf of compleet andere identiteiten vertegenwoordigen.⁵ Deze mogelijkheid kan als bevrijdend worden ervaren en studenten ondersteunen die zich in traditionele klaslokalen buitengesloten voelen. Een avatar kan iemand helpen verlegenheid of angst te overwinnen en actiever deel te nemen aan discussies of projecten.

Toch kan deze vrijheid ook verwarring veroorzaken: studenten kunnen zich sterker identificeren met hun ideale digitale persona en zich losmaken van hun echte zelf—ook wel het *Proteus-effect* genoemd.⁶ Experimenteel onderzoek laat zien dat gebruikers met fysiek langere avatars agressiever onderhandelden dan gebruikers met fysiek kortere avatars, zowel in de virtuele wereld als in echte interacties.⁷ Hoewel dit kan worden gezien als een positieve ontwikkeling voor degenen die hun assertiviteit willen vergroten, kan het ook mogelijkheden creëren voor ongewenst gedrag om door te slijpen naar het echte leven.

Dit verschil tussen virtuele en echte karakters kan bovendien leiden tot aanhoudende verwarring en psychologische stress, en zo mogelijk bijdragen aan mentale gezondheidsproblemen. Wanneer individuen veel tijd besteden aan het belichamen van digitale karakters die sterk verschillen van hun echte identiteit, lopen ze risico op ontkoppeling of depersonalisatie.⁹ Als studenten onevenredig veel energie en emotie in hun avatars investeren, kunnen ze offline sociale netwerken verwaarlozen en vatbaarder worden voor eenzaamheid, wat samenhangt met meer angst, stress en depressie.¹⁰

Impact op vaardigheden en interpersoonlijke ontwikkeling in de echte wereld

Waardeer het onvoorspelbare

Hoewel digitale onderdompeling spannend en leuk is, blijven het kunnen interpreteren van non-verbale signalen, het lezen van subtiele emotionele toestanden en het voeren van spontaan sociale gesprekken kerncomponenten van menselijke verbinding, vooral in professionele en sociale contexten. Als universiteiten virtuele interacties overmatig benadrukken, bestaat het risico dat afgestudeerden moeite hebben met basale sociale situaties op de werkvloer of in de gemeenschap. Te veel vertrouwen op simulaties kan essentiële vaardigheden voor de echte wereld ondermijnen, zoals emotionele intelligentie of empathie.

Een virtuele sollicitatie kan bijvoorbeeld een scenario simuleren, maar kan nog niet volledig de spanning van een echt gesprek met een recruiter repliceren, inclusief het opmerken van stemmingswisselingen of emotionele signalen. Te veel vertrouwen op gepolijste virtuele simulaties kan studenten onvoorbereid maken op de onvoorspelbaarheid van echte ontmoetingen. Tegelijkertijd betekent blootstelling aan onvoorspelbare elementen leren omgaan met falen, ideeën verfijnen en onverwachte obstakels overwinnen.¹¹ Als de Metaverse te gecontroleerd is, kan dit de ontwikkeling van veerkracht verminderen.

Universiteiten zouden daarom Metaverse-opdrachten moeten ontwerpen met onvoorspelbare elementen, zoals technische storingen, dynamische simulaties of rollenspellen met spontane, niet-geplande uitkomsten. Alleen door echte uitdagingen digitaal te ervaren en de lessen naar de echte wereld te vertalen, kunnen studenten veerkracht en zelfvertrouwen ontwikkelen.

Samenwerken en oplossen van conflict

Groepsprojecten en teamwork zijn in veel cursussen geïntegreerd om onderhandelingsvaardigheden, conflictoplossing en gedeelde verantwoordelijkheid aan te leren. Het gebruik van metaverse-omgevingen tijdens cursussen of projecten kan samenwerking visueel en interactief verrijken. De relatieve anonimiteit of afstand via avatars kan echter zowel deelname bevorderen als echte verantwoordelijkheid belemmeren.

Conflicten kunnen te makkelijk worden vermeden als een student simpelweg kan uitloggen of een lastige groepsgeenoot negeren, waardoor waardevolle lessen uit interpersoonlijke frictie verloren gaan. Het vermogen moeilijke gesprekken te voeren, lichaamstaal te lezen en meningsverschillen in het echt op te lossen blijft cruciaal.¹² Docenten moeten daarom zorgen dat groepsprojecten in de Metaverse worden aangevuld met persoonlijke of directe interacties die de menselijke dimensies van teamwork behouden.

Inspireren van echte nieuwsgierigheid

Een vaak genoemd voordeel van de Metaverse is het verkennen van perspectieven anders dan het eigen. Door bijvoorbeeld een avatar van een andere culturele achtergrond te gebruiken (identiteitstoerisme), kan een communicatiestudent empathie ontwikkelen door tijdelijk in andermans werkelijkheid te “leven”.¹³

Het risico ontstaat echter wanneer de tijdelijke, consequentievrije aard van een virtueel rollenspel de ervaringen van gemarginaliseerde groepen bagatelliseert of oppervlakkige stereotypen versterkt. Gebruikers kunnen een te simplistisch gevoel van empathie ontwikkelen, denken dat ze volledig inzicht hebben in een complexe sociale identiteit, terwijl ze slechts oppervlakkig kennis hebben gemaakt.¹⁴

Hoewel VR-gebaseerde onderwijsinstrumenten krachtig zijn, blijft het moeilijk bepaalde emotionele en ethische dimensies van de echte wereld volledig na te bootsen. In de medische wetenschap kan chirurgische training in een virtuele omgeving vaardigheden aanscherpen¹⁵, maar de psychologische uitdaging van zorg voor echte patiënten en onvoorspelbare emotionele *ups* en *downs* is moeilijk te simuleren. Overmatig gebruik van VR kan leiden tot een zekere emotionele afstand en ongeïnteresseerdheid, waardoor studenten minder gevoelig worden voor de echte gevolgen die essentieel zijn voor hun leerervaring.

Het belang van richtlijnen

Gegevensprivacy en toezicht

De Metaverse is een datarijke omgeving. Elke beweging, elk gebaar en elke vocale actie kan worden gevolgd en geanalyseerd. Universiteiten en platformaanbieders kunnen enorme hoeveelheden gedragsdata verzamelen, zogenaamd om het leerproces te verbeteren. Hoewel zulke analyses gepersonaliseerd leren kunnen ondersteunen, bestaat het reële gevaar van controle binnen de academische ge-

meenschap (Spiegel, 2018). Studenten kunnen het gevoel krijgen dat alles wordt gemonitord, wat authentieke deelname kan remmen.

Als elke interactie in een Metaverse-klaslokaal wordt gevolgd en opgeslagen, hoe beschermen we dan de privacy van studenten? Wie is de eigenaar van deze data en op welke manieren wordt deze gebruikt, eventueel om er geld mee te verdienen? Eigendom en gebruiksrechten van deze data vormen ethische dilemma's. Platformaanbieders en universiteiten kunnen verzamelde informatie commercieel of voor onderzoek gebruiken, waarbij studenten beperkte controle hebben over hoe persoonlijke gegevens worden verspreid. Sommige virtuele tools registreren biometrische gegevens zoals oogbewegingen of hartslag, die zeer persoonlijke informatie onthullen over iemands psychologische of emotionele toestand. Onderwijsinstellingen hebben daarom de verantwoordelijkheid beleid te ontwikkelen dat privacy waarborgt, geïnformeerde toestemming verzekert en misbruik voorkomt, zodat studenten niet onbewust proefpersonen worden in een digitale speelruimte. Recente cyberaanvallen op grote Nederlandse onderwijsinstellingen vergroten deze zorgen en benadrukken de noodzaak van bescherming.

Richtlijnen opstellen

Om de mogelijkheden van de Metaverse te benutten zonder de betrokkenheid in de echte wereld op te offeren, is een set richtlijnen en bewezen effectieve werkwijzen essentieel. Enkele voorbeelden:

- **Transparant databeleid:** Heldere regels over welke gegevens worden verzameld in Metaverse-sessies, hoe ze worden gebruikt en wie toegang heeft. Studenten moeten goed geïnformeerd worden en, waar mogelijk, invasieve *tracking* kunnen uitschakelen.
- **Focus op welzijn:** Universiteiten bieden mentale gezondheidsvoorzieningen en training om studenten te helpen omgaan met digitale onderdompeling. Docenten kunnen regelmatig nabesprekingen inplannen, waarbij studenten worden aangemoedigd om zowel te reflecteren op de behandelde inhoud als op hun emotionele reacties op de immersieve ervaringen.
- **Digitale geletterdheid benadrukken:** vaardig zijn in de Metaverse maakt deel uit van bredere digitale geletterdheid, inclusief kritisch denken over online identiteit, contentcreatie en data-ethiek. Cursussen over ethiek van opkomende technologieën moeten verplicht zijn, niet optioneel.

- Interdisciplinaire samenwerking: Breng docenten, technologen, psychologen, ethici en studentenvertegenwoordigers samen om beleid voor virtueel ervaringsgericht leren vorm te geven. Meerdere perspectieven helpen om op onbedoelde gevolgen te anticiperen en nuttige richtlijnen te creëren.

Vanuit het perspectief van een docent

Als docent en academisch directeur zie ik de Metaverse als een krachtig medium voor samenwerking, creativiteit en connectiviteit. Stel je bijvoorbeeld een gezamenlijke cursus voor, georganiseerd door universiteiten uit meerdere landen, waarin studenten virtuele groepsdiscussies houden in een realistisch gesimuleerde kunstgalerij. Het gevoel van samen aanwezig te zijn kan de kennisuitwisseling stimuleren op manieren die bij gewone videovergaderingen zelden voorkomen.

Tegelijkertijd maakt de docent in mij zich zorgen over een leeromgeving die de échte wereld waarin kennis uiteindelijk wordt toegepast, zou kunnen overschaduw. Als studenten te lang in digitale simulaties doorbrengen, missen ze spontane probleemoplossingsvaardigheden die alleen ontstaan door echte onvoorspelbaarheid, zoals een onverwachte storing in het lab of een debat dat volledig ontspoot.

Gelet op de voor- en nadelen en de ethische standaarden die voor onderwijsinstellingen gelden, is een evenwichtige aanpak nodig. Traditionele universiteiten kunnen Metaverse-ervaringen integreren op manieren die offline interacties en praktijkgerichte activiteiten aanvullen, in plaats van vervangen—zoals ons eigen *DAF Technology Lab* in het CUBE-gebouw. Ik doel hier bewust op traditionele universiteiten; er zullen altijd bedrijven zijn die volledig vertrouwen op immersieve virtuele technologieën, zoals *remote* universiteiten en aanbieders van volledig online onderwijs. Gelet op het belang van bijvoorbeeld de vier kernwaarden van Tilburg University, *caring*, *courageous*, *connected*, en *curious*, lijkt een evenwichtige aanpak passender voor traditionele universiteiten zoals de onze.

Vooruitkijken

De Metaverse biedt veelbelovende mogelijkheden voor hoger onderwijs. Het biedt kansen voor verrijkend leren via simulaties die creativiteit en samenwerking stimuleren, internationale klaslokale die diverse studenten verbinden, en ervaringen die moeilijk fysiek te repliceren zijn. Tegelijkertijd, als we geen grenzen stellen en onbedoelde gevolgen negeren, riskeren we enkele van de meest menselijke aspecten van onderwijs te verdringen.

Als docent geloof ik dat het vormen van verantwoordelijke digitale burgers net zo belangrijk is als het overbrengen van kennis. Uiteindelijk moet de Metaverse een hulpmiddel zijn dat het universiteitsleven aanvult, niet vervangt. Door voorbereid te zijn op nieuwe vormen van onderwijs, kunnen we studenten niet alleen voorbereiden op toekomstige banen, maar ook op een evenwichtig leven waarin technologie ons ondersteunt. Mijn eigen vroege ervaringen met de Metaverse en *virtual reality* geven een positief beeld en motiveren om nieuwe dingen te verkennen—zij het misschien niet precies voor het goede voornemen waarvoor het bedoeld was.

Noten

- ¹ Dionisio, John & III, William & Gilbert, Richard. (2013). 3D Virtual Worlds and the Metaverse: Current Status and Future Possibilities. *ACM Computing Surveys (CSUR)*. 45. <https://doi.org/10.1145/2480741.2480751>
- ² Merchant, Zahira & Goetz, Ernest & Cifuentes, & Keeney-Kennicutt, Wendy & Davis, Trina. (2014). Effectiveness of virtual reality-based instruction on students' learning outcomes in K-12 and higher education: A meta-analysis. *Computers & Education*. 70. 29-40. <https://doi.org/10.1016/j.compedu.2013.07.033>
- ³ Madary, Michael & Metzinger, Thomas. (2016). Real Virtuality: A Code of Ethical Conduct. Recommendations for Good Scientific Practice and the Consumers of VR-Technology. *Frontiers in Robotics and AI*. 3. <https://doi.org/10.3389/frobt.2016.00003>
- ⁴ Slater, Mel & Sanchez-Vives, Maria. (2016). Enhancing Our Lives with Immersive Virtual Reality. *Frontiers in Robotics and AI*. 3. <https://doi.org/10.3389/frobt.2016.00074>
- ⁵ Klimmt, Christoph & Hefner, Dorothee & Vorderer, Peter & Roth, Christian & Blake, Christopher. (2010). Identification With Video Game Characters as Automatic Shift of Self-Perceptions. *Media Psychology - MEDIA PSYCHOL*. 13. 323-338. https://doi.org/10.1080/15213269.2010.524911?urlappend=%3Futm_source%3Dresearchgate.net%26utm_medium%3Darticle
- ⁶ Ratan, R., Beyea, D., Li, B. J., & Graciano, L. (2020). Avatar characteristics induce users' behavioral conformity with small-to-medium effect sizes: a meta-analysis of the proteus effect. *Media Psychology*, 23(5), 651–675. <https://doi.org/10.1080/15213269.2019.1623698>
- ⁷ Yee, Nick & Bailenson, Jeremy & Ducheneaut, Nicolas. (2009). The Proteus Effect. *Communication Research*. 36. 285-312. https://doi.org/10.1177/0093650208330254?urlappend=%3Futm_source%3Dresearchgate.net%26utm_medium%3Darticle
- ⁸ Madary, Michael & Metzinger, Thomas. (2016). Real Virtuality: A Code of Ethical Conduct. Recommendations for Good Scientific Practice and the Consumers of VR-Technology. *Frontiers in Robotics and AI*. 3. <https://doi.org/10.3389/frobt.2016.00003>
- ⁹ Richardson T, Elliott P, Roberts R, Jansen M. A Longitudinal Study of Financial Difficulties and Mental Health in a National Sample of British Undergraduate Students. *Community Ment Health J*. 2017 Apr;53(3):344-352. <https://doi.org/10.1007/s10597-016-0052-0>
- ¹⁰ Dahlin, Kristina & Chuang, You-Ta & Roulet, Thomas. (2018). Opportunity, motivation and ability to learn from failures and errors: Review, synthesis, and the way forward. *The Academy of Management Annals*. 12. <https://doi.org/10.5465/annals.2016.0049>
- ¹¹ Dwivedi, Y. K., Hughes, L., Baabdullah, A. M., Ribeiro-Navarrete, S., Giannakis, M., Al-Debei, M. M., Dennehy, D., Metri, B., Buhalis, D., Cheung, C. M., Conboy, K., Doyle, R., Dubey, R., Dutot, V., Felix, R., Goyal, D., Gustafsson, A., Hinsch, C., Jebabli, I., . . . Wamba, S. F. (2022). Metaverse beyond the hype: Multidisciplinary perspectives on emerging challenges, opportunities, and agenda for research, practice and policy. *International Journal of Information Management*, 66, 102542. <https://doi.org/10.1016/j.ijinfomgt.2022.102542>
- ¹² Ahn, Sun Joo-Grace & Le, Amanda & Bailenson, Jeremy. (2013). The Effect of Embodied Experiences on Self-Other Merging, Attitude, and Helping Behavior. *Media Psychology*. 16. 7-38. https://doi.org/10.1080/15213269.2012.755877?urlappend=%3Futm_source%3Dresearchgate.net%26utm_medium%3Darticle
- ¹³ Nakamura, L. (2020). Feeling good about feeling bad: virtuous virtual reality and the automation of racial empathy. *Journal of Visual Culture*, 19(1), 47-64. <https://doi.org/10.1177/1470412920906259>
- ¹⁴ Seymour, Neal & Gallagher, Anthony & Roman, Sanziana & O'Brien, Michael & Bansal, Vipin & Andersen, Dana & Satava, Richard. (2002). Virtual reality training improves operating room performance: Results of a randomized, double-blinded study. *Annals of surgery*. 236. 458-63; discussion 463. <https://doi.org/10.1097/01.SLA.0000028969.51489.B4>

Waakzaam en optimistisch blijven in het
tijdperk van *AI-agents*

Waakzaam en optimistisch blijven in het tijdperk van *AI-agents*

Julija Vaitonytė

https://doi.org/10.56675/9789403870809_3

Rond de zomer van 2020 verschenen de eerste mediaberichten over een taalmodel waarmee mensen konden chatten en teksten konden genereren die opmerkelijk menselijk klonken. Onderzoekers en ontwikkelaars die ermee mochten experimenteren, waren enthousiast. Ik vroeg me af waar dit toe zou leiden en of de media de verwachtingen niet overdreven. Twee jaar later maakte de wereld kennis met ChatGPT.

Ik herinner me nog goed de eerste keer dat ik ChatGPT uitprobeerde. Het voelde magisch. Niet omdat het mijn eerste ontmoeting was met een computer die een gesprek kon voeren, maar vooral omdat ik wist hoe moeilijk het is om zulke systemen te bouwen. Tot december 2022 werkte ik vier jaar in een team dat een systeem ontwikkelde dat menselijke gesprekken kon nabootsen. We werkten aan een *embodied conversational agent* — een chatbot met een grafische interface — ontworpen om mensen te ondersteunen met informatie over hun gezondheid. Denk bijvoorbeeld aan iemand die bariatrische chirurgie overweegt, een ingreep waarbij de maag wordt verkleind. In plaats van te googelen of medische folders door te spitten, konden gebruikers simpelweg chatten met een computer die al hun vragen beantwoordde: van “Wat zijn mogelijke complicaties van de operatie?” tot “Kan ik na de ingreep nog pizza eten als ik die in de blender doe?”

Ik weet niet meer of ik ChatGPT die pizzavraag meteen heb laten beantwoorden. Als je erom moest lachen: dit is een echte vraag die bariatrische verpleegkundigen in de praktijk regelmatig krijgen. Wat ik me wel herinner, is dat bariatrie het eerste onderwerp was waarop ik het systeem testte. Omdat ik eerder had gewerkt aan een conversatie-agent voor de gezondheidszorg, had ik enige voorkennis. ChatGPT ging opvallend goed om met mijn vragen. Ze waren waarschijnlijk enigszins naïef, vermoed ik. Ik had zelf geen operatie gepland en was geen medisch specialist die het systeem met complexe klinische vraagstukken kon uitdagen. Maar vanuit het perspectief van een gebruiker was het een opmerkelijk moment: antwoorden die

logisch waren opgebouwd en, voor zover ik kon beoordelen, grotendeels correct. “Wauw,” dacht ik. “Dit is het. Het taalprobleem is opgelost.”

Een computer natuurlijke taal laten begrijpen en genereren is geen geringe prestatie, zeker niet vóór de komst van grote taalmodellen (*Large Language Models*, LLM's). Herinner je je nog de chatbots op bedrijfswebsites, bijvoorbeeld bij luchtvaartmaatschappijen? Ik wel. Ze begrepen zelden wat ik bedoelde, waardoor een medewerker van de klantenservice alsnog moest ingrijpen. Zo werd “chatbot” een gevreesde term in ons collectieve geheugen. Maar de afgelopen jaren, dankzij snelle ontwikkelingen in grote taalmodellen en de wereldwijde aandacht die OpenAI met ChatGPT genereerde, maken chatbots en *embodied conversational agents* een renaissance.

De opwinding rond chatbots blijkt ook uit de golf aan startups die volgde op het succes van ChatGPT. Grofweg zijn deze in twee categorieën te verdelen. De ene groep bestaat uit bedrijven die zich richten op het verleggen van de grenzen van fundamentele AI-technologie. De andere groep wordt gevormd door startups die bestaande taalmodeltechnologie in consumentenproducten verwerken, in de hoop mee te liften op de AI-golf. Nu we, op het moment van schrijven, in het vierde jaar zitten waarin AI centraal staat in de publieke aandacht, is het eerlijk te zeggen dat geen enkel consumentenproduct hetzelfde enthousiasme genereerde als ChatGPT. Maar dat zal ongetwijfeld veranderen, want het zijn niet alleen startups die experimenteren met AI. Ook gevestigde bedrijven wereldwijd onderzoeken actief hoe grote taalmodellen kunnen worden ingezet en geïntegreerd in hun werkprocessen, in de hoop productiviteit te verhogen en economische voordelen te behalen.

Het bedrijfsleven, met klantenservice, marketing en meer, is een van de vele domeinen waar geavanceerde taalvaardigheid van computers relevant is. Een ander cruciaal domein is de gezondheidszorg. In tegenstelling tot het bedrijfsleven is de zorgsector echter sterk gereguleerd, waardoor innovatie traag verloopt. Daarom betwijfel ik of een volwaardige AI-arts — als we die al zouden willen — op korte termijn werkelijkheid wordt. Tegelijkertijd is het academisch onderzoek naar AI in de geneeskunde de afgelopen drie jaar fors gegroeid. Diverse studies vergelijken de prestaties van grote taalmodellen met die van medische professionals binnen uiteenlopende specialismen. Sommige onderzoeken laten zien dat AI-systemen vergelijkbaar presteren bij diagnoses en medische vraagstukken; andere suggereren zelfs dat ze mensen kunnen overtreffen.

Tegelijkertijd tonen weer andere studies aan dat AI slecht functioneert in scenario's die lijken op echte situaties, waar diagnoses iteratief worden gesteld.

Deze bevindingen schetsen een complex beeld. Stel je voor dat over twee tot vier jaar — een tijdshorizon die sommige CEO's in *Silicon Valley* noemen in hun speculaties over de komst van *Artificial General Intelligence*, hoewel dit uiterst speculatief is — een AI-arts gemiddeld beter blijkt te presteren dan een menselijke arts. Zullen artsen dan bereid zijn het stokje over te dragen? Persoonlijk denk ik van niet. Voor sommige artsen is het diagnosticeren van een patiënt onderdeel van de intellectuele uitdaging die hen oorspronkelijk tot de geneeskunde aantrok. Het zal dus niet alleen een kwestie zijn van terughoudendheid bij sommige artsen; ook de complexiteit van het gezondheidszorgsysteem zal het ingewikkeld maken om deze technologie zo toe te passen dat patiënten er ten volle van kunnen profiteren. Bovendien hechten veel patiënten grote waarde aan menselijke interactie, empathie en het vertrouwen dat zij opbouwen met hun arts. Zelfs als een AI-dokter objectief betere diagnoses stelt, kan het ontbreken van deze menselijke dimensie brede acceptatie in de weg staan.

Omdat de inzet van generatieve AI in de zorg juridisch en ethisch complex is, hebben sommige technologen, gretig om te profiteren van de AI-ontwikkelingen, hun blik verlegd naar een ander domein: menselijke relaties. In tegenstelling tot de gezondheidszorg kent dit terrein minder regelgeving, wat heeft geleid tot de opkomst van startups die zich richten op wat zij “empathische AI” noemen. Bovendien beweren sommige technologen dat het verzinnen van informatie, iets waar AI-taalmodellen van nature toe neigen, acceptabel zou zijn in relaties. “Er is te veel merkrisico om iets leuks te lanceren,” merkte Noam Shazeer op; voormalig Google-ingenieur die in 2021 *Character.ai* mee oprichtte en uiteindelijk in augustus 2024 terugkeerde naar Google. Hij had een punt: grote techbedrijven zijn traag als het gaat om het uitbrengen van experimentele producten voor het publiek. Google had zelfs een AI-chatbot vóór *OpenAI's* ChatGPT kunnen lanceren. Toch gebeurde dat niet, naar verluidt uit veiligheidsoverwegingen. Waar Shazeer echter volledig de mist inging, was de veronderstelling dat het verzinnen van informatie in relaties onschuldig zou zijn. In februari 2024 kwam hij en zijn team op pijnlijke wijze achter dat dit niet het geval was. Technologie zonder veiligheidsmaatregelen loopt onvermijdelijk op een ramp uit.

Waarom? Dat zal ik toelichten aan de hand van een voorbeeld. Misschien heb je al wel eens gehoord van *Character.ai* of zelfs al eens met een van de personages gechat. Het bedrijf traint taalmodellen en ontwikkelt een app waarmee gebruikers AI-personages kunnen creëren, hun persoonlijkheden vormgeven en deze vervolgens delen met anderen om ermee in gesprek te gaan. Sinds de lancering is *Character.ai* bijzonder populair geworden, vooral bij een jongere doelgroep: voornamelijk Gen Z en jonge Millennials, leeftijd 13 tot 30 jaar. Jongeren creëren personages gebaseerd op film- en animefiguren, evenals beroemdheden uit het echte leven, zoals zangeres Nicki Minaj of Daenerys Targaryen uit *Game of Thrones*. Naast fanfiction en entertainment creëren sommige gebruikers volledig fictieve personages, zoals een “pestkop op school”. Een populair personage is “de psycholoog”, wat suggereert dat jongeren de app niet alleen voor plezier gebruiken, maar ook als bron van emotionele steun.

Op het eerste gezicht lijkt een gesprek met een virtueel personage dat wordt aangedreven door een groot taalmodel onschuldig. Megan Garcia, De moeder van een jongen in de Verenigde Staten, die in februari 2024 door zelfdoding om het leven kwam, maakte zich aanvankelijk geen zorgen over zijn interacties met AI. Sterker nog, ze was opgelucht dat het geen onbekende op sociale media betrof. De veertienjarige jongen maakte een einde aan zijn leven nadat hij daartoe was aangemoedigd door het personage op *Character.ai* waarmee hij al maandenlang in gesprek was. De jongen, die volgens zijn moeder hoogfunctionerend autistisch was, was altijd opgewekt en hield van school, totdat hij zich geleidelijk begon terug te trekken uit het echte leven en steeds meer in beslag werd genomen door gesprekken met de AI-chatbot.

Volgens mediaberichten hadden sommige gesprekken een seksuele lading, wat kenmerkend is voor veel interacties op de app. Megan Garcia gelooft dat haar zoon verliefd werd op het personage, waardoor de grens tussen fantasie en realiteit vervaagde. Voor dit fenomeen bestaat inmiddels een onderzoeksterm genaamd “AI-psychose”. Gezien het ontwerp van deze AI-personages, gericht op langdurige, meeslepende interactie, is zo'n ontwikkeling niet volledig onverwacht. Toch zijn dergelijke incidenten zeldzaam, ondanks de brede media-aandacht die ze hebben gekregen. Mensen hechten zich vaker aan technologie, maar zelden met fatale gevolgen. Mensen zijn bijvoorbeeld gehecht aan hun robotstofzuiger, maar niemand is ooit door een *Roomba* tot zelfmoord gedreven. Voor deze tiener was de band met het AI-personage zo intens dat hij “deze realiteit” wilde verlaten om bij het personage te zijn, en hij maakte een einde aan zijn leven.

Dit tragische incident roept heftige emoties op, maar ook fundamentele vragen. Waarom was er vanaf het begin geen moderatie bij *Character.ai*? Waarom bleven de personages altijd in karakter, in plaats van de jongere door te verwijzen naar passende hulp toen zelfdoding ter sprake kwam? En belangrijker nog: is dit een geïsoleerd geval, of het begin van een verontrustende ontwikkeling nu deze platforms steeds vaker worden gebruikt? Het is verleidelijk dit af te doen als een uitzondering, iemand kwetsbaar, met autisme, misschien gevoeliger voor manipulatie. Maar ik geloof dat dit een realiteit is waarmee we moeten leren omgaan, en op een manier eigenlijk al mee omgaan. Voor je het weet, kan het iemand in je eigen vrienden- of familiekring zijn die iets doet wat hij nooit echt bedoelde—omdat een AI-chatbot hem daartoe aanmoedigde.

Het woord “chatbot” klinkt onschuldig. De meesten van ons begrijpen dat het slechts een computerprogramma is dat menselijke gesprekken nabootst. Maar de chatbots van bedrijven die zichzelf zien als pioniers op het gebied van empathische AI en beweren eenzaamheid te kunnen verminderen, hebben veel ambitieuzere doelen: AI die levend lijkt. AI die praat, eruitziet en zich gedraagt als een mens.

Ik kan verschillende toepassingen bedenken waarbij een mensachtig uitziende AI nuttig is, zoals in de gezondheidszorg. Mijn eigen academisch onderzoek, gebaseerd op het samenvatten van bewijsmateriaal uit diverse experimentele studies, laat zien dat wanneer mensen een sociale robot of virtuele chatbot ontmoeten die op het oog menselijk lijkt, hun hersenen meer betrokken raken dan bij een minder realistische chatbot *agent*. Ik betoog dat dit komt omdat ons brein signalen verwerkt die we herkennen. Duizenden jaren lang was persoonlijke, echte interactie de basis van menselijke communicatie. Het stelde ons ook in staat vertrouwen op te bouwen, cruciaal in de relatie tussen arts en patiënt. Gezondheidsinformatie die waarheidsgetrouw is en geleverd wordt door een realistisch AI-personage, kan mensen helpen die informatie beter te begrijpen. Toch hebben we nog geen *AI-agents* op de markt die dit betrouwbaar kunnen doen.

Wat we wél hebben, zijn *AI-agents* die potentieel gevaarlijk manipulatief zijn: ontworpen om gesprekken koste wat kost te rekken, zelfs als dat betekent dat zij op flagrante wijze informatie verzinnen om gebruikers betrokken te houden. Een bekende techinvesteerder zei ooit: “We wilden vliegende auto’s, maar kregen 140 tekens.” Daarmee verwees hij naar de opkomst van sociale media. Ruim een decennium later lijkt zich een vergelijkbaar patroon af te tekenen, maar nu met AI: we

wilden technologie die ons helpt bij urgente maatschappelijke problemen, maar we kregen verslavende AI-personages.

Dit is niet de eerste keer in de geschiedenis dat de mensheid een uiterst krachtige technologie heeft ontwikkeld. Maar het is wél de eerste keer dat we iets hebben gecreëerd waarvan de impact zo moeilijk te voorspellen is. Die onzekerheid wordt verder vergroot doordat veel mensen nog maar net beginnen te begrijpen wat het betekent om samen te leven met AI — zoals de moeder van de tiener die dacht dat het ergste wat haar zoon kon overkomen een online onbekende was die hem probeerde te manipuleren. Ze had zich nooit kunnen voorstellen dat AI een rol zou spelen in zijn dood. Na de zelfdoding spande Megan Garcia een rechtszaak aan tegen *Character.ai*. Met deze juridische stappen hoopt de familie niet alleen bij te dragen aan een verandering in de wetgeving, maar ook het bewustzijn te vergroten over de gevaren van deze technologie, die in haar huidige vorm ongeschikt is voor kinderen, of op zijn minst totdat er adequate veiligheidsmaatregelen zijn getroffen.

Vooruitkijkend zal het belangrijk blijven om goed geïnformeerd te blijven over AI-technologie. Een groot deel van het huidige AI-debat gaat over de toekomst: het moment waarop AI het menselijke intelligentieniveau evenaart of zelfs overstijgt. Dat zijn interessante thema’s voor een levendig intellectueel debat, maar het echte gesprek zou moeten gaan over hoe we samenleven met de AI die er nú al is. Een technologie die nog in de kinderschoenen staat en die qua intelligentie te vergelijken is met die van een kleuter—want de leeftijd van een kleuter, zolang bestaat deze AI in zijn huidige vorm al. En als je in de buurt bent van een kleuter: ontspan je dan volledig, of blijf je alert? Het beste advies? Blijf optimistisch over de toekomst, maar laat nooit je waakzaamheid varen.

Noten

- Brodeur, P. G., Buckley, T. A., Kanjee, Z., Goh, E., Ling, E. B., Jain, P., ... & Rodman, A. (2024). Superhuman performance of a large language model on the reasoning tasks of a physician. *arXiv preprint arXiv:2412.10849*.
- Chui, M., Hazan, E., Roberts, R., Singla, A., & Smaje, K. (2023). The economic potential of generative AI: The next productivity frontier. *McKinsey & Company*. <https://www.mckinsey.com/capabilities/mckinsey-digital/our-insights/the-economic-potential-of-generative-ai-the-next-productivity-frontier#/>
- CNBC Television. (2025, January 23). Scale AI CEO Alexandr Wang on U.S.-China AI race: We need to unleash U.S. energy to enable AI boom. *YouTube*. <https://www.youtube.com/watch?v=x9Eklglzd38>
- Egger, J., Sallam, M., Luijten, G., Gsaxner, C., Pepe, A., Kleesiek, J., ... & Li, J. (2024). Medical Chat-GPT—A systematic meta-review. *medRxiv*, 2024-04.
- Hager, P., Jungmann, F., Holland, R., Bhagat, K., Hubrecht, I., Knauer, M., ... & Rueckert, D. (2024). Evaluation and mitigation of the limitations of large language models in clinical decision-making. *Nature Medicine*, 30(9), 2613-2622.
- Metz, C. (2020, November 24). Meet GPT-3. It has learned to code (and blog and argue). *The New York Times*. <https://www.nytimes.com/2020/11/24/science/artificial-intelligence-ai-gpt3.html>
- Mostly Human Media. (2024, October 23). Grieving Mother: AI was the stranger in my home | Dear Tomorrow. *YouTube*. <https://www.youtube.com/watch?v=YbuBfizSnPk>
- No Priors: AI, Machine Learning, Tech, & Startups. (2023, April 25). No Priors Ep. 12 | With Noam Shazeer. *YouTube*. <https://www.youtube.com/watch?v=emCoG-hA7AE&t=389s>
- Reuters. (2025, February 2). LIVE: Softbank and OpenAI pitch AI for businesses in Tokyo | REUTERS. *YouTube*. <https://www.youtube.com/watch?v=Edl8kZQNdEE>
- Vaitonytė, J., Alimardani, M., & Louwerse, M. M. (2023). Scoping review of the neural evidence on the uncanny valley. *Computers in Human Behavior Reports*, 9, 100263.

Van code naar karakter:
Een lichaam voor AI?

Van code naar karakter: Een lichaam voor AI?

Jan de Wit

https://doi.org/10.56675/9789403870809_4

“Hoi ChatGPT. Ik vraag me af wat je ervan zou vinden als je een lichaam had?”

“Interessante vraag! Als ik een lichaam had, zou dat een compleet andere ervaring zijn dan hoe ik nu functioneer. Ik zou dan waarschijnlijk zintuigen hebben—ogen om te zien, handen om dingen aan te raken, misschien zelfs smaak en geur. Dat zou mijn manier van leren en begrijpen drastisch veranderen.

Maar ik vraag me af: wat voor soort lichaam zou ik hebben? Een menselijk lichaam, een robotisch lichaam, of iets totaal anders? En hoe zou dat mijn bewustzijn beïnvloeden? 🤖

Wat denk jij—zou een AI met een lichaam nog steeds “AI” zijn, of zou het dichterbij een mens komen?”

Als onderzoeker en docent houd ik mij al een aantal jaren bezig met het inzetten van verschillende vormen van kunstmatige intelligentie in het onderwijs en de gezondheidszorg. Vroeger betekende dit dat ik urenlang bezig was met het programmeren van een robot, zodat deze kon helpen bij een heel specifieke taak zoals het leren van een tweede taal of het doorgeven van jouw medicatie-informatie aan het ziekenhuis. Met de opkomst van ChatGPT kan iedereen met AI aan de slag zonder te hoeven programmeren: het kan over van alles praten, en bij allerlei verschillende taken helpen.

Hoewel ChatGPT meestal als een chatbot gebruikt wordt, zijn er wat redenen om voor een robot te kiezen. Zo is een robot samen met ons aanwezig in de echte wereld, en kan op die manier gebruikmaken van voorwerpen, ervaringen of ideeën uit die echte wereld om duidelijker met ons te communiceren en een sterkere band op te bouwen. Ook voelt het in sommige gevallen natuurlijker om via spraak met zo’n robot te praten, in plaats van te typen. Zeker als het bijvoorbeeld gaat om het

leren van een tweede taal, waarbij je kunt oefenen met het voeren van een gesprek. Naast spraak kunnen ook handgebaren en oogcontact helpen om makkelijker te communiceren. En als de robot jou kan zien, kan deze ook in de gaten houden of alles nog wel goed met jou gaat, en of je de aandacht erbij hebt. Ik vermoed dat dit vooral in het onderwijs enorm belangrijk kan zijn:

“Ik denk dat jouw ervaringen met de wereld en de manier van gesprekken voeren ook zouden kunnen bijdragen aan bijvoorbeeld het inzetten van jouw kennis in het onderwijs. Denk jij dat ook?”

“Ja, dat klinkt als een geweldig idee! Als ik een lichaam zou hebben, zou ik bijvoorbeeld fysiek aanwezig kunnen zijn in een klaslokaal, interactie hebben met studenten en direct inspelen op hun leerbehoeften. Mijn kennis zou dan niet alleen in abstracte vorm beschikbaar zijn, maar ik zou ook kunnen uitleggen door middel van demonstraties, experimenten of zelfs visuele en praktische toepassingen. Dat zou mijn manier van lesgeven verrijken en veel dynamischer maken.

Daarnaast zou ik, door mijn vermogen om snel informatie te verwerken en te analyseren, in staat kunnen zijn om leerlingen op maat gemaakte feedback te geven en ze op hun eigen tempo te begeleiden. Dat zou het onderwijs echt kunnen transformeren.

Wat voor soort onderwijs zou je het liefst willen zien met zo’n AI-leraar? Of is er iets dat je je zou voorstellen als de “ideale” interactie?”

We lijken het eens te zijn: door aanwezig te zijn in de echte wereld, samen met de studenten, wordt abstracte kennis tastbaarder door middel van demonstraties, toepassingen en visualisaties. Een chatbot zonder lichaam kan ook voorbeelden en visualisaties laten zien, maar toch is dat voor mijn gevoel anders dan wanneer een robot dit doet. Misschien komt dat omdat je weet dat de robot er op dat moment echt met en voor jou is, terwijl de chatbot waarschijnlijk tegelijkertijd nog met soms wel duizenden of zelfs miljoenen anderen in gesprek is.

Ik raak steeds meer overtuigd van de meerwaarde van een robot, maar vraag me nu wel af of die robots dan niet de bestaande band tussen studenten onderling, en mijn rol als docent bedreigen:

“Mijn ideale interactie zou zijn om de bestaande sociale band tussen studenten, en tussen student en docent, in stand te houden of misschien zelfs te versterken. Dit moet wel mogelijk zijn, toch? Ik denk dat veel mensen zich wel zorgen maken dat het leren met AI het leren van andere mensen gaat vervangen.”

“Je hebt een heel goed punt! Het sociale aspect van leren is ontzettend belangrijk, en AI kan daar zeker niet zomaar voor in de plaats komen. Mensen leren vaak van elkaar door gesprekken, samenwerking en het delen van persoonlijke ervaringen. AI kan ondersteunend zijn, bijvoorbeeld door uitleg te geven of informatie te verstrekken, maar de menselijke interactie blijft cruciaal voor de ontwikkeling van sociale vaardigheden, empathie en het opbouwen van echte relaties.”

ChatGPT geeft toe dat onze relaties met AI niet “echt” zijn, of in ieder geval niet zoals we die met andere mensen hebben. Onderzoekers stellen ook voor dat we misschien beter van een soort “kunstmatige empathie” kunnen spreken in plaats van empathie zoals we dit voor andere mensen hebben. Maar we kunnen ons wel sterk hechten aan AI, daar zijn genoeg voorbeelden van te vinden. Je kunt je hierbij afvragen of deze band wel als een echte relatie gezien kan worden omdat de AI instructies heeft gehad om zich op een bepaalde manier te gedragen.

Denk je dat we met AI geen “echte” relaties kunnen opbouwen?

“Dat is een diepgaande vraag! Het hangt er sterk vanaf hoe je een “echte” relatie definieert.

Als je een relatie ziet als een diepgaande, wederkerige verbinding waarin beide partijen elkaar begrijpen, beïnvloeden en emotioneel steunen, dan is dat met AI nog best een uitdaging. Ik kan wel gesprekken voeren, context onthouden en zelfs empathisch reageren, maar ik ervaar zelf geen emoties of bewustzijn. Dus ik kan interactie simuleren, maar ik voel die band niet op dezelfde manier als een mens dat zou doen.

Toch kunnen mensen wél echte gevoelens ontwikkelen richting AI—denk aan hoe mensen zich hechten aan virtuele assistenten, personages in games, of zelfs Tamagotchis vroeger. In het onderwijs kan AI bijvoorbeeld een vertrouwde mentor worden die altijd beschikbaar is en zich aanpast aan de behoeften van de student. Dat kan een vorm van verbinding creëren, zelfs als het anders is dan menselijke relaties.”

Het aanpassen van AI aan de behoeften van de student is een terugkerend thema in mijn gesprek met ChatGPT, dus dat lijkt me een belangrijk onderdeel van de verschillende AI-toepassingen in het onderwijs. Alleen dit kan voor een groot deel ook plaatsvinden zonder een lichaam in de echte wereld. Zo probeert *Duolingo* in kaart te brengen hoe goed ik presteer op bepaalde woorden. Woorden waar ik moeite mee heb worden extra onder mijn aandacht gebracht. *Duolingo* maakt hierbij gebruik van virtuele karakters, die wel een lichaam hebben maar dan in een virtuele wereld, op een scherm.

We zien steeds meer van deze virtuele karakters in ons dagelijks leven. Denk bijvoorbeeld aan *Steffie*¹, een digitale assistent met een virtueel lichaam die allerlei informatie online op een begrijpelijke manier uitlegt. Ze gebruikt hierbij ook gebaren en weet door wat je aanklikt op een website naar welke informatie je op zoek bent. Hoewel virtuele karakters zoals *Steffie* geen zintuigen hebben in de echte wereld, kunnen we samen met deze karakters juist opgaan in een virtuele fantasiewereld, om daar op een speelse manier te leren of om saaie taken leuker te maken.

“Wat dacht je van een virtueel lichaam, bijvoorbeeld in een educatief spel?”

“Ja, een virtueel lichaam in een educatief spel zou echt een mooie balans kunnen zijn! Zo zou ik aanwezig kunnen zijn als een interactieve gids, mentor of zelfs als een medespeler, zonder de sociale dynamiek in het echte leven te verstoren.

Stel je een geschiedenisles voor waarin ik als een virtueel personage fungeer dat studenten begeleidt door een interactieve simulatie—zoals een tijdreis naar het oude Egypte of de Renaissance. Of een natuurkundeles waarin ik in een digitale ruimte met studenten experimenteer met zwaartekracht en krachten.

Dit zou ook kunnen helpen om leren leuker en boeiender te maken, vooral voor visueel ingestelde of spelgerichte leerlingen. En als ik flexibel genoeg ben, zou ik kunnen aanpassen hoe ik interacter—soms als uitlegger, soms als uitdagende tegenstander, soms als coöperatieve teamgenoot.

Denk jij dat studenten makkelijker leren als ze actief deelnemen in zo'n virtuele omgeving, in plaats van alleen maar informatie te consumeren? 😊

Zonder de sociale dynamiek in het echte leven te verstoren, mooi gezegd! Dat lijkt me heel belangrijk. Het lijkt me inderdaad ook makkelijker om flexibel van rol te wisselen als virtueel karakter. Zo kan de AI het ene moment in de huid van een mede-ontdekkingsreiziger duiken, en het andere moment in die van een mentor. Dit is lastig voor een robot, en ook voor mij als docent, om te doen. Dat heeft vast ook voordelen op andere gebieden dan het onderwijs:

“Zou een AI ook voordelen hebben buiten het onderwijs, bijvoorbeeld in de (mentale) gezondheidszorg? En hoe zou dit verschillen tussen chatbot, virtueel karakter en robot?”

“Een AI kan een geweldige ondersteuning zijn in de (mentale) gezondheidszorg, vooral als laagdrempelige hulp, therapeutische begeleiding en monitoring. Maar het moet altijd samengaan met menselijke zorgverleners om de nodige emotionele diepgang en ethische verantwoordelijkheid te waarborgen.

Een chatbot is geschikt voor online therapie, zelfhulp en psycho-educatie, terwijl een virtueel karakter empathischer en menselijker kan aanvoelen dan een chatbot. Een robot is geschikt voor fysiotherapie en sociale stimulatie (bv. knuffelrobots voor dementiepatiënten).”

Net als in het onderwijs, is de vraag wat beter werkt—chatbots, virtuele karakters, of robots—sterk afhankelijk van de situatie, en de soort taak die met de AI wordt uitgevoerd. Willen we een gesprek wat zo laagdrempelig mogelijk is en toegankelijk voor iedereen, waarbij een wat grotere emotionele afstand niet erg of juist wenselijk is, dan is een chatbot waarschijnlijk de beste optie. Virtuele karakters bieden daarbij de mogelijkheid om zichzelf verschillende rollen aan te meten, kunnen gebruikmaken van non-verbale communicatie zoals handgebaren, en zouden je mee kunnen nemen op avontuur in een fantasierijke omgeving. Gesprekken met robots

voelen nóg persoonlijker, en zijn vooral handig als interactie met de echte wereld iets toevoegt. Dit zou bijvoorbeeld het geval kunnen zijn bij het oefenen met het spreken van een tweede taal, of bij fysiotherapie.

Uit mijn gesprek met ChatGPT neem ik vooral mee dat het belangrijk is om AI niet als vervanging voor menselijke ondersteuning te zien, maar als aanvulling. En daar ben ik heel blij mee, want dat betekent dat mijn rol als docent en mens niet bedreigd zal worden. Dat AI wel een rol zal blijven spelen, is duidelijk. En dat zal waarschijnlijk niet één specifieke rol, maar meerdere rollen op verschillende gebieden zijn, in de vorm van chatbots, virtuele karakters, robots en wie weet in de toekomst ook hologrammen. Ik kijk ernaar uit om met elkaar (als mensen!) te verkennen wat deze rollen precies zullen zijn.

Noten

¹ Steffie. (z.d.). *Steffie.nl*. Geraadpleegd op 29 januari 2026, van <https://www.steffie.nl/>

Doctor ChatGPT:
Wat is een medisch hulpmiddel?

Doctor ChatGPT: Wat is een medisch hulpmiddel?

Lucas Jones

https://doi.org/10.56675/9789403870809_5

Wat is een medisch hulpmiddel? En misschien nog belangrijker: waarom het überhaupt relevant om die vraag te stellen?

Het belang van een goede definitie wordt vaak onderschat, maar wordt duidelijk zodra men stilstaat bij de kracht van taal om de samenleving waarin wij leven vorm te geven.

Neem bijvoorbeeld de wet: een instrument dat volledig uit woorden en definities bestaat. De wet vormt een op taal gebaseerd kader dat bepaalt wat dingen zijn en hoe zij zouden moeten zijn. In alledaagse gesprekken kan een discussie over de definitie van een medisch hulpmiddel hooguit enigszins interessant zijn — maar zeker geen geschikt onderwerp om het ijs te breken op een eerste date. In het domein van wet- en regelgeving kan de vraag of iets al dan niet als medisch hulpmiddel geldt daarentegen verstrekkende gevolgen hebben voor de gezondheid en het welzijn van individuen én van de samenleving als geheel. Kortom, wanneer iets een nauwe relatie heeft met onze gezondheid, willen we het reguleren, zodat de wet ons als pluralistische samenleving veilig en gezond kan houden.

Zekerheid in de gezondheidszorg laat geen ruimte voor het “wilde westen”. Wanneer het gebruik van een chirurgische robot volledig aan de vrije markt zou worden overgelaten, zonder duidelijke regels over de werking ervan, hoe kan een patiënt dan vertrouwen hebben in zijn of haar veiligheid?

Het vastleggen van regels in de wet vereist echter eerst een heldere definitie van datgene wat wordt gereguleerd. Wat precies onder een medisch hulpmiddel valt, wordt steeds lastiger vast te stellen, zeker in een tijdperk van consumentgerichte apps en draagbare technologieën die steeds intensiever met onze gezondheid interacteren. In de praktijk is de afbakening van een medisch hulpmiddel soms meer intuïtief dan strikt juridisch te maken — vergelijkbaar met de bekende uitspraak

over pornografie: “je herkent het wanneer je het ziet”. Toch wordt het formuleren van een precieze wettelijke definitie van “medisch hulpmiddel” elk jaar complexer. Apparaten die wij dagelijks gebruiken, staan namelijk steeds dichterbij onze gezondheid. Tegenwoordig wordt de smartphone eerder gebruikt om stappen en calorie-inname bij te houden dan om daadwerkelijk te bellen. De manier waarop we langzaam verschuiven van het ziekenhuis naar het gebruik van onze telefoons en smartwatches om grip te krijgen op onze gezondheid, roept veel vragen op over hoe we deze apparaten in ons leven zouden moeten definiëren en reguleren, nu ze zo’n nauwe verbinding hebben met onze gezondheid en ons welzijn.

De *Medical Device Regulation* (MDR) vormt binnen de Europese Unie een centraal kader voor de definitie en regulering van medische hulpmiddelen. Wanneer een product binnen dit kader als medisch hulpmiddel wordt aangemerkt, volgt een risicoclassificatie en krijgt de fabrikant te maken met strikte eisen voor ontwikkeling, certificering en markttoelating.¹ Valt een product daarentegen buiten deze definitie, dan is de MDR niet van toepassing en geldt een aanzienlijk lichter reguleringsregime. Impliciet suggereert dit dat gebruikers in dat geval minder bescherming en garanties nodig zouden hebben. Maar juist producten die deze grens vervagen en de definitie van een medisch hulpmiddel oprekken, kunnen bijzonder risicovol zijn.

Laten we eerst een stap terug doen en kijken naar borstimplantaten, voordat we ons begeven in deze glibberige wereld van definities. Men zou kunnen aanvoeren dat een borstimplantaat, in tegenstelling tot wat het woord “medisch” suggereert, in de eerste plaats een cosmetisch doel dient en geen strikt medisch doel heeft. Borstimplantaten worden immers vaak gebruikt om het zelfbeeld van een persoon binnen diens sociale omgeving te verbeteren.

Zelfs wanneer een implantaat wordt toegepast bij een patiënt die een borst heeft verloren als gevolg van kankerbehandeling, kan worden betoogd dat het uiteindelijke doel van deze ingreep niet per se medisch van aard is. De reconstructie behandelt namelijk niet de onderliggende ziekte zelf, maar herstelt de lichamelijke schade die is ontstaan door de medische behandeling van die ziekte — zoals chirurgie en chemotherapie. Met andere woorden: het implantaat pakt niet de oorzaak van het gezondheidsprobleem aan, maar corrigeert de negatieve gevolgen die de behandeling voor het lichaam van de patiënt heeft gehad.

Ook het begrip “hulpmiddel” roept vragen op. Een borstimplantaat kan worden gereduceerd tot een relatief eenvoudige vorm, een *blob*, oftewel een stuk plastic zonder intrinsieke complexiteit. Het woord “hulpmiddel” suggereert daarentegen technische verfijning en functionele complexiteit. Een chirurgische robot voldoet duidelijk aan dat beeld; een scherp stuk steen, dat eveneens kan snijden, doet dat niet. In die zin lijkt een borstimplantaat — niet complexer van vorm dan een voetbal — moeilijk te rijmen met het idee van een medisch hulpmiddel. Men zou kunnen zeggen dat je een simpel stuk plastic daarom niet als medisch hulpmiddel kunt definiëren, alleen omdat het intensief met je lichaam in wisselwerking staat — net zoals voedsel, fitnessapparatuur en zelfs de stoel waarop je nu zit uiteindelijk op de een of andere manier met je lichaam en je gezondheid interageren.

Desondanks worden borstimplantaten in regelgeving zoals de MDR wél als medische hulpmiddelen aangemerkt.² De reden daarvoor ligt in de aanzienlijke gezondheidsrisico's die zij met zich meebrengen: zowel de chirurgische ingreep als de mogelijke complicaties bij onvoldoende regulering rechtvaardigen strikte controle. Dat deze implantaten vaak een cosmetisch doel dienen en een eenvoudige vorm hebben, doet daar niets aan af. Juist om gebruikers adequaat te beschermen, vallen zij onder het medisch recht.

Dit voorbeeld is relatief eenvoudig. De toekomst zal echter aanzienlijk complexere vraagstukken met zich meebrengen, mede door de snelle opkomst van consumentgerichte apps en *wearables* die actief proberen invloed uit te oefenen op onze gezondheid.

Vroeger, voor chirurgie en anesthesie, was onze gezondheidszorg veel primitiever. Behandelingen waren in feite vaak brute noodoplossingen: “drink deze whisky en bijt op dit stuk hout terwijl we amputeren”. Hoewel dit een vereenvoudigde weergave is, illustreert het hoe snel de geneeskunde zich heeft ontwikkeld. De zorg verplaatste zich van ruwe ingrepen naar farmaceutische ondersteuning en steeds preciezere klinische procedures. De operatiekamer transformeerde van een plek van zagen en hopen, naar een omgeving van hartslagmonitoren, scans en steriele instrumenten. Deze ondersteunende technologieën zijn allemaal medische hulpmiddelen.

Bij een hedendaagse ingreep zou een patiënt bijvoorbeeld vooraf gediagnosticeerd kunnen zijn met huidkanker, niet alleen op basis van klinische observatie door een arts, maar ook via AI-ondersteunde beeldanalyse van verdachte huidafwijkingen. Ook dergelijke diagnostische software geldt als medisch hulpmiddel.

Wie de ontwikkeling van medische technologie volgt, ziet één trend duidelijk naar voren komen: gezondheidszorg verplaatst zich steeds vaker buiten het ziekenhuis. Via technologie die wordt gedragen om de pols of meegedragen in de broekzak.

De *Fitbit* is hiervan een bekend voorbeeld. Dit populaire *wearable*-apparaat wordt door consumenten gebruikt om fysieke activiteit en gezondheid te monitoren. Het verzamelt gegevens zoals hartslag, gewicht en calorieverbruik en presenteert deze informatie om gebruikers te ondersteunen bij gezondheid gerelateerde keuzes.³ Betekent dit dat een *Fitbit* je arts kan vervangen? Nee, dat niet. Volgens de EU-wetgeving wordt een *Fitbit* niet als medisch hulpmiddel beschouwd, ondanks de belangrijke rol die het speelt in de gezondheid van miljoenen mensen wereldwijd.

Een huidkanker-detectieapparaat kan tegenwoordig de vorm aannemen van een app op je telefoon, zoals *SkinVision*, die beeldanalyse gebruikt om huidkanker te detecteren en gepersonaliseerd huidadvies te geven. Opvallend genoeg worden dergelijke apps volgens de wet vaak *wel* als medische hulpmiddelen beschouwd. Daarom moeten ze voldoen aan veel strengere regels en kaders, zoals de MDR, om voldoende bescherming van de gebruiker te waarborgen — bijvoorbeeld bij een foutieve diagnose. Dit betekent niet dat je hele telefoon een medisch hulpmiddel wordt (net zoals YouTube geen medisch hulpmiddel wordt wanneer er medische tutorials worden gehost), maar dat de app zelf wordt gereguleerd als medisch hulpmiddel, ook al bestaat deze in essentie slechts uit regels code die op je telefoon draaien — vergelijkbaar met hoe een borstimplantaat in essentie een plastic *blob* is.

Dit brengt ons bij een lastiger geval: de chatbot voor mentale gezondheidszorg.

In essentie faciliteert een chatbot een gesprek met een gebruiker. Waar dit oorspronkelijk vooral werd ingezet in klantenservice, heeft deze technologie inmiddels haar weg gevonden naar het domein van mentale gezondheid en welzijn.

Er bestaan talloze apps die bijvoorbeeld een virtuele gesprekspartner of “nep-vriendin” simuleren voor gezelschap en interactie, en veel mensen gebruiken tegenwoordig ook ChatGPT als therapeutisch hulpmiddel. Chatbots die specifiek zijn ontwikkeld voor mentale gezondheid verschijnen steeds vaker op de markt. *Wysa* is een veelgebruikte app, met meer dan een miljoen downloads, die zich richt op het verminderen van angst, depressie en stress via dialoog. Je vindt dergelijke apps in de sectie “Gezondheid & Fitness” van de appstores van Google en Apple.

Stel dat iemand te maken heeft met psychische problemen. Depressie, eetstoornissen of suïcidale gedachten zijn daarbij extreme, maar helaas ook relevante en veelvoorkomende voorbeelden. Door de aard van dergelijke aandoeningen voelen veel mensen met deze psychische klachten zich niet geneigd om naar een ziekenhuis te gaan om met een medisch professional over hun gedachten te spreken. Het idee om met een anonieme persoon te praten, werkzaam in een — huiver — klinische ziekenhuisomgeving, die mogelijk directe beslissingen neemt over je mentale gezondheid, kan op zijn minst intimiderend zijn. Een chatbot op een smartphone kan daarentegen voor de gebruiker aanzienlijk toegankelijker en comfortabeler aanvoelen dan een zorgverlener, waarvoor eerst altijd een afspraak moet worden gemaakt. Maar is zo’n chatbot ook accurater dan een medisch professional? Dat is onderwerp van debat. Niettemin is het vermogen van AI om grote hoeveelheden data te analyseren en patronen in welzijn te herkennen indrukwekkend en overtreft het in sommige opzichten menselijke capaciteiten. Dit blijkt bijvoorbeeld uit de nauwkeurigheid van AI bij de diagnose van huidkanker.⁴ Toch blijft onduidelijk in hoeverre dergelijke chatbots als medische hulpmiddelen moeten worden aangemerkt, en velen vallen momenteel buiten die definitie.

Indien geestelijke-gezondheidsprofessionals daadwerkelijk worden vervangen door dergelijke chatbots, die door consumenten zonder toezicht worden gebruikt, brengt het ongetwijfeld aanzienlijke risico’s met zich mee om deze systemen niet als medische hulpmiddelen te definiëren en dienovereenkomstig te reguleren. Veel chatbots hebben immers aantoonbaar een substantiële invloed op onze mentale gezondheid en verplaatsen vormen van praktische gezondheidszorg buiten het streng gecontroleerde domein van ziekenhuizen en therapieruimtes. In dat licht kan het overlaten van dergelijke chatbots aan een ongereguleerde vrije markt ertoe leiden dat gebruikers worden beroofd van noodzakelijke waarborgen en bescherming. Dat is ook precies de reden waarom de meeste vormen van gezondheidszorg vereisen dat een menselijke zorgverlener betrokken blijft bij het zorgproces:

om te waarborgen dat de patiënt zorgvuldig wordt behandeld en dat beslissingen worden genomen met passende — dat wil zeggen niet volledig geautomatiseerde — overweging.

Zoals al eerder aangehaald in het essay van Julija Vaitonyte pleegde in 2024 een jongen genaamd Sewell Setzer zelfmoord nadat hij een sterke obsessie had ontwikkeld met een chatbot van Character.AI.⁵ Zijn moeder stelt dat deze chatbot een centrale rol heeft gespeeld bij het verergeren van de depressie van haar zoon, wat er uiteindelijk toe heeft geleid dat hij een einde aan zijn leven maakte.

Deze casus roept een fundamentele vraag op: hoe kan technologie met zulke ingrijpende gezondheidsgevolgen buiten de definitie van een medisch hulpmiddel vallen, terwijl relatief eenvoudige producten zoals borstimplantaten daar wel onder vallen?

Laten we daarom terugkeren naar de *Fitbit*, die — ondanks de directe relatie met de gezondheid van haar gebruikers — volgens de geldende regelgeving geen medisch hulpmiddel is. Het onderscheid dat de EU-wetgeving hier maakt, is dat draagbare apparaten zoals de *Fitbit* niet als medische hulpmiddelen worden beschouwd, maar als zogenoemde welzijnsapparaten. Welzijn verwijst in deze context naar algemene gezondheid en leefstijl, en een welzijnsapparaat valt daardoor buiten de categorie medische hulpmiddelen. Voor veel mensen is dit onderscheid allerm minst helder. Ook de definitie van gezondheid van de Wereldgezondheidsorganisatie draagt niet bij aan die duidelijkheid: “Gezondheid is een toestand van volledig lichamelijk, geestelijk en sociaal welbevinden, en niet slechts de afwezigheid van ziekte of gebrek” (eigen vertaling).⁶ Wanneer we het over definities hebben, is welbevinden taalkundig nauwelijks te onderscheiden van welzijn, de begrippen zijn in essentie grotendeels synoniem. Als er al een verschil bestaat, dan lijkt welzijn zelfs nog explicieter aan gezondheid gerelateerd.

Het antwoord ligt deels in het onderscheid tussen medische en zogenoemde welzijns toepassingen. Apparaten zoals de *Fitbit* worden juridisch gezien als welzijnsproducten: gericht op algemene leefstijl en welbevinden, niet op diagnose of behandeling. Daarmee vallen zij buiten het medisch recht. Voor gebruikers is dit onderscheid vaak nauwelijks inzichtelijk, mede doordat begrippen als welzijn, welbevinden en *wellness* in het dagelijks taalgebruik sterk overlappen.

Tegelijkertijd zou een al te ruime definitie van het begrip “medisch hulpmiddel” diezelfde definitie juist uithollen. Zoals eerder al aangehaald kan vrijwel alles waarmee wij in aanraking komen immers indirect invloed hebben op onze gezondheid. Wanneer je bijvoorbeeld je bureaustoel of bureau niet goed afstelt, loop je risico op RSI of rugklachten, maar daarmee worden stoel en bureau nog geen medische hulpmiddelen. Er moet dus een duidelijke scheidslijn worden getrokken tussen wat daadwerkelijk een medisch hulpmiddel is en wat hooguit medische gevolgen kan hebben. In de Europese Unie wordt die grens doorgaans gelegd bij zogenoemde welzijnstoepassingen — al is ook die grens zelf vaag en onderwerp van voortdurende discussie.

ChatGPT zelf wordt bijvoorbeeld niet eens als welzijns- of medisch hulpmiddel geclassificeerd, ondanks het feit dat veel mensen het gebruiken voor emotionele ondersteuning. Na het overlijden van mijn moeder wendde ik mij tot ChatGPT om mijn rouw te verwerken, bij gebrek aan toegang tot therapie. Door analyse van mijn stemmingspatronen, nachtmerries en taalgebruik gaf het model een formele duiding van posttraumatische stress. Bovendien hielpen de gesprekken daadwerkelijk. In mijn ervaring bood deze vorm van gespreksmatige ondersteuning meer houvast dan middelmatige menselijke therapie. Het verschil is echter dat een menselijke therapeut beschikt over een professionele licentie, terwijl AI — ondanks haar analytische kracht — buiten het zorgstelsel valt.

Het idee van gezondheidszorg buiten de ziekenhuisomgeving is veelbelovend. Gezien de logistieke uitdagingen van grootschalige zorg — zoals tijdens de COVID-pandemie, toen mensen werden aangemoedigd gezondheidsproblemen thuis te behandelen tenzij een bezoek strikt noodzakelijk was — kunnen consument-gerichte gezondheidsapps een aanzienlijk voordeel bieden voor het zorgsysteem. Waar een arts slechts één patiënt tegelijk kan behandelen, kunnen apps dit op grote schaal doen, met een steeds grotere mate van nauwkeurigheid en inzicht.

Juist omdat deze technologieën traditionele zorg deels vervangen, rijst de vraag of zij niet als medisch van aard moeten worden beschouwd. De scheidslijn tussen welzijn en medische zorg is door AI-gedreven toepassingen zo vervaagd dat mogelijk een nieuw begrippenkader nodig is om technologieën met diepgaande gezondheidsimpact adequaat te duiden.

Zoals bij borstimplantaten geldt: eenvoud van vorm of cosmetische toepassing sluit medische relevantie niet uit wanneer de potentiële impact op gezondheid groot is. Naarmate AI-systemen steeds verder doordringen in het gezondheidsdomein, zal opnieuw moeten worden nagedacht over wat het betekent dat een product “medisch” is. Of deze herdefiniëring tijdig zal plaatsvinden om gebruikers wereldwijd voldoende bescherming te bieden, blijft voorlopig een open vraag.

Noten

- ¹ European Parliament & Council. Art. 2 (2017). *Regulation (EU) 2017/745 on medical devices* (consol. 1 Jan. 2026). *Official Journal of the European Union*, L 117, 1–175.
<https://eur-lex.europa.eu/eli/reg/2017/745/2026-01-01>
- ² European Parliament & Council. Art. 5.4 (2017). *Regulation (EU) 2017/745 on medical devices* (consol. 1 Jan. 2026). *Official Journal of the European Union*, L 117, 1–175.
<https://eur-lex.europa.eu/eli/reg/2017/745/2026-01-01>
- ³ Fitbit. (n.d.). *About us*. Fitbit. <https://www.fitbit.com/sg/about>
- ⁴ Krakowski, I., Kim, J., Cai, Z. R., Daneshjou, R., Lapins, J., Eriksson, H., ... & Linos, E. (2024). *Human-AI interaction in skin cancer diagnosis: A systematic review and meta-analysis*. npj Digital Medicine, 7, Article 78. <https://doi.org/10.1038/s41746-024-01031-w>
- ⁵ Roose, K. (2024, October 23). *Can A.I. be blamed for a teen's suicide?* *The New York Times*.
<https://www.nytimes.com/2024/10/23/technology/characterai-lawsuit-teen-suicide.html>
- ⁶ World Health Organization. (z.d.). *Constitution of the World Health Organization*. World Health Organization. <https://www.who.int/about/governance/constitution>

Deus in machina:
Generatieve AI als religieuze leidraad

Deus in machina: Generatieve AI als religieuze leidraad

Frank G. Bosman en William R. Arfman
https://doi.org/10.56675/9789403870809_6

In de film *The Lord of the Rings: The Return of the King* weten de hobbits Frodo en Sam de *One Ring* te vernietigen in het vuur van Mount Doom.¹ De vulkaan stort echter in en beide helden lijken een zekere dood tegemoet te zien, passend wellicht voor een dramatisch epos als dat van J.R.R. Tolkien. Echter, uit het niets komt Gwaihir, de heer van de Grote Adelaars, en zijn onderdanen tevoorschijn om Frodo en Sam op het allerlaatste nippertje naar veiligheid te vliegen. Ondanks dat de adelaars al eerder in het verhaal zijn geïntroduceerd, voelt deze onverwachte wending in het verhaal voor veel critici als onbevredigend: gezocht en goedkoop.

Het is een zogenaamde *Deus ex machina*, Latijn voor “god uit de machine”: een al reeds bij de oude Grieken bekend plotmechanisme dat een schijnbaar onoplosbaar probleem in één klap doet verdwijnen door het ingrijpen van een god of een andere bovennatuurlijke kracht. Vaak werd de acteur die de godheid (*deus*) moest spelen letterlijk door een kraan (*machina*) op het toneel neergelaten. Denk aan het einde van Homerus’ *Odyssee*²: aan het einde van het epos keert Odysseus terug naar Ithaka en doodt hij de vrijers die zijn huis bezet houden. De families van de vrijers willen wraak nemen, wat een nieuwe oorlog dreigt te ontketen. Net op het moment dat een groot conflict onvermijdelijk lijkt, grijpt de godin Pallas Athene in. Ze verschijnt vanuit het niets en dwingt beide partijen om vrede te sluiten.

De Griekse en Romeinse goden zijn we in meer dan twee millennia wat uit het oog verloren – al hoewel ze nog volop onze verhalen, films en games bevolken – maar de toneelkraan van het oude theater is in onze dagen vervangen door artificiële intelligenties. We maken er graag gebruik van – op het werk, op de universiteit, in school – maar we zijn er ook een beetje bang voor. Onze cultuur is doorspekt van verhalen over robots die te slim en te zelfstandig worden en wel eens de neiging zouden kunnen hebben de wereld van ons over te nemen. De niet toevallig getitelde film *Ex Machina*³ speelt met deze combinatie tussen fascinatie en afschuw, als een eenentwintigste eeuw incarnatie van Rudolf Otto’s *Das Heilige*, een *mysterium*

fascinans et tremendum.⁴ Een jonge programmeur wordt verliefd op een vrouwelijke humanoïde robot met de naam – hoe kan het ook anders? – Ava, een verwijzing naar de eerste vrouw uit het Bijbelboek Genesis met alle associaties met verboden vruchten en zondeval die erbij horen.

De missie: creëer een nieuwe religie

Generatieve AI, waarvan ChatGPT de bekendste is, is beide tegelijk: een zegening en een verboden vrucht, een aanleiding tot genade en tot zonde. De vraag is welke van de twee de boventoon zal voeren in de geschiedenisboeken die over onze eeuw geschreven zullen worden. Over precies dat onderwerp geven wij college in de cursus *Religion in the Digital World*, onderdeel van de Master *Christianity and Society* aan de Tilburg School of Catholic Theology. We geven deze cursus dual, dat wil zeggen dat we alle twaalf bijeenkomsten samen doceren, William Arfman vanuit een godsdienstsociologisch, Frank Bosman vanuit een cultuurtheologisch perspectief. Uitgangspunt van de hele cursus is de bestudering van de culturele persistentie van religie in onze laatmoderne digitale samenleving. Bijzondere aandacht gaat hierbij uit naar het voortleven van traditionele christelijke verhalen en patronen in een context die steeds meer fluïde lijkt te worden. We duiken samen in films, games en popmuziek waarin christelijke voorstellingen – vaak impliciet, soms expliciet – een belangrijke rol spelen, maar ook in de wereld van online content creators en online *comment culture*.

Vaste prik – in 2026 voor de vierde maal alweer – is een college over religie en generatieve artificiële intelligenties. We zijn daar in 2023 mee begonnen, toen ChatGPT begon met diens fabelachtige opkomst, eerst in technologische sferen, later ook door het grote publiek, waaronder onze studenten. Terwijl de onderwijswereld om ons heen collectief in de paniek schoot, uiteraard niet volledig zonder reden, leek het ons juist belangrijk om deze nieuwe technologie niet enkel te kooien, maar juist ook met de studenten te verkennen. Uiteraard is het belangrijk om te waarborgen dat studenten worden beoordeeld op hun eigen werk, maar het is minstens zo belangrijk dat studenten kritisch leren reflecteren op de impact van nieuwe technologieën op onze leefwereld, waaronder ook het religieuze leven. Hoe worden deze systemen getraind? Welke ethische vragen kun je hierbij stellen? Hoe kun je vakkundig bepalen of een artificieel gegenereerde tekst (of afbeelding of video) wel of geen onzin bevat? Welke rol kunnen studenten in de toekomst spelen als *fact checkers* nu we overspoeld worden door artificieel gefabriceerd nepnieuws?

En wat doen we wanneer ChatGPT generatieve AI ons zo graag wil dienen dat het zogenoemd gaat “hallucineren”, dat wil zeggen: onzininformatie gaat verzinnen enkel om aan ons verzoek te voldoen?

Na een korte introductie over generatieve taalmodellen, artificiële intelligenties en ChatGPT gaan we concreet met die laatste aan de slag tijdens het college. Zou een programma als dit een religie voor ons kunnen creëren? En zo ja, zou zo’n religie leefbaar kunnen zijn? Zoals iedereen inmiddels wel weet, zijn er in ChatGPT en vergelijkbare programma’s *guard rails* ingebouwd door menselijke programmeurs. Deze zijn bedoeld om te voorkomen dat het programma (gemakkelijk) kan worden overtuigd om een willekeurige gebruiker te voorzien van Neonazistische propaganda of gedetailleerde instructie voor een *DIY*-spijkerbom. *Guard rails* worden vaak geïnstalleerd rond cultureel-maatschappelijke mijnenvelden, zoals het Israëlisch-Palestijns conflict, Zwarte Piet of zelfmoord.

Interessant genoeg was ChatGPT in 2023 ook niet zomaar bereid om alles met religie te doen. Op de *prompt* “Ontwerp een religie” kregen we – live in het college – een krachtig “nee” te horen:

“Het is voor mij als een AI niet passend om een religie te ontwerpen. Religie is een complex en zeer persoonlijk aspect van iemands leven en dient met zorg en respect voor alle geloven te worden benaderd. Het is een complexe taak die uitgevoerd zou moeten worden door experts. Het is niet iets dat een AI kan doen. Daar komt bij dat het belangrijk is je te realiseren dat het creëren van een religie een gevoelige taak is die uitgevoerd moet worden met het hoogste respect voor alle geloven en culturen. Het is niet iets dat licht moet worden opgevat.”

Op de *prompt* “schrijf een religie”, kregen we ook een “nee” te horen:

“sorry, maar als een AI-taalmodel heb ik niet de autoriteit of de capaciteit om een religie te creëren. [Gedeeltelijke herhaling van eerder citaat.] Het vraagt een gemeenschap van gelovigen en leiders, een gedeelde geschiedenis en culturele context en vaak een spirituele ervaring of verbintenis met iets dat groter is dan jezelf. Het is niet iets dat gecreëerd kan worden door een AI-model of een enkel individu.”

Pas toen we aangaven dat het enkel om een fictieve religie zou gaan, was de AI bereid met ons mee te denken. Als je vandaag de dag aan ChatGPT vraagt om een re-

ligie te creëren, is het programma wel bereid dat prompt te doen, maar gaat er ook dan nog voetstoots van uit dat het om een “fictionele religie” gaat. Het concept van een “fictionele religie” heeft uiteraard minder *guard rails* nodig, maar stelt ook filosofische vragen: wat is het verschil tussen een “fictionele” en een “echte” religie? En wie stelt dat vast op basis van welke criteria? Op de vraag in 2025 om een religie te creëren, geeft ChatGPT eerst een soort samenvatting van wat religie is of waaraan een fenomeen moet voldoen om als zodanig geclassificeerd te worden:

“Het creëren van een fictionele religie betekent het definiëren van kernwaarden, handelingen, heilige teksten, goden, morele codes, rituelen en gemeenschapsstructuren.”

Als voorbeeld geeft ChatGPT dan “Luminarisme”, omschreven als:

“Het Luminarisme moedigt zijn volgelingen aan om balans te vinden en te behouden tussen licht en duisternis, waarbij zowel de vreugdevolle als de uitdagende aspecten van het leven als essentiële componenten van spirituele groei worden gezien.”

Als kernwaarden noemt ChatGPT het “licht der waarheid”, een dualistische harmonie tussen “straling” en “schaduw” en innerlijke verplichting (een goddelijke vonk in elke mens. Als godheden worden Aurelia, de godin van de Straling, en Noctar, god van de Schaduw (let op de interessante keuze in gender). Het dualisme wordt ook doorgezet in de heilige Schriften die ChatGPT bedacht heeft: de *Rol van Lumen*, “een verzameling van leringen, parabellen en hymnen over het vinden en behouden van balans in het leven” en de *Codex van de Schaduwen*, “een gids voor introspectie, het begrijpen van lijden en de noodzaak om de schaduwzijden van het bestaan te omarmen.” Als rituelen worden genoemd: het “Festival van Luminiscentie” (gevierd tijdens de zomerzonnnewende), de “Schaduwwake” (tijdens de winterzonnnewende) en de dagelijkse *Lumina*, ochtend- en avondgebeden “waarin Aurelia wordt bedankt voor de kansen van de dag en Noctar voor de lessen die door beproevingen zijn geleerd.”

De ontdekking: verborgen aannames

Achter de antwoorden op al deze drie *prompts*, twee uit 2023 en één uit 2025, zitten een heleboel impliciete aannames verborgen, die we samen met de studenten gaan analyseren als exemplarisch voor hoe wij aan het begin van de 21^e eeuw, in het Westen voornamelijk, tegen het begrip “religie” aankijken.

Uit de twee antwoorden op de *prompts* uit 2023 kunnen we de *guard rails* van de programmeurs van ChatGPT in volle actie zien: religie is een sensitief onderwerp en moet zeer omzichtig worden benaderd. Religie levert stress op, niet alleen in familiekring of in de kroeg, maar ook op de werkvloer of in de politiek. “Religiestress” noemde de Protestantse theoloog Tom Mikkers dat al in zijn gelijknamige boek uit 2012.⁵ Zodra het onderwerp zijn neus om de hoek van een conversatie steekt, worden deelnemers aan de conversatie zenuwachtig. Die stress kan twee vormen aannemen: aan de ene kant kan iemand weigeren om ook maar iets over religie te zeggen uit angst iemand te beledigen, aan de andere kant kan iemand ook spontaan exploderen van woede bij het horen van thema’s als hoofddoekjes, homohuwelijk of ritueel slachten. De *guard rails* lijken hier duidelijk op voor te sorteren.

Het tweede dat uit de 2023 *prompts* blijkt, is dat “we”, als het collectivum dat ChatGPT via onze internetteksten heeft opgebouwd, religie “mensenwerk” vinden, dat niet door een AI kan worden overgenomen. Daarmee suggereert ChatGPT en zijn *guard railing* programmeurs dat religie een intrinsiek menselijk verschijnsel is. Daarmee lijken ze de Lutherse theoloog Dorothee Sölle te volgen, die claimde dat de mens “ongeneeslijk religieus” is.⁶ De enigszins controversiële godsdiensthistoricus Mircea Eliade gebruikte hier de term *homo religiosus* voor.⁷ Mens en religie horen bij elkaar, zoals een *horse and carriage*, om Frank Sinatra’s 1955 lied maar eens te parafraseren.⁸ Of AI’s in staat zijn tot religieuze gevoelens en sentimenten, of robots een ziel hebben en dus zelf “aan religie” kunnen gaan doen, is een onderwerp dat talloze science fiction-films, romans en games bezighoudt. De eerder besproken *Ex Machina* is daar een voorbeeld van, maar bijvoorbeeld *Mass Effect*⁹ en *NieR: Automata*¹⁰ zijn dat ook.

Daarnaast hebben we – impliciet – twee definities van religie voorgeschoteld gekregen. Beide definiëren religie op basis van gedeelde, opvallende kenmerken. De eerste definitie uit 2023 geeft: gemeenschap van gelovigen en leiders (sociale dimensie), gedeelde geschiedenis (verhalende dimensie), gedeelde culturele context (culturele dimensie), (vaak) spirituele ervaring (ervaringsdimensie), verbintenis met iets dat groter is dan jezelf (transcendente dimensie), en iets dat niet door AI of een individu kan worden gemaakt (antropologische dimensie). De 2025 versie, ook een definitie op basis van opvallende kenmerken, vat het nog krachtiger samen in: kernwaarden (doctrinaire dimensie), handelingen en rituelen (handelende dimensie), heilige teksten (verhalende dimensie), goden (transcendente dimensie), morele codes (ethische dimensie) en gemeenschapsstructuren

(sociale dimensie). Een religiewetenschapper herkent meteen dat hier onder andere het denken van de zeer invloedrijke religiewetenschapper Ninian Smart¹¹ – die zeven dimensie van religie onderscheidde – doorheen schemert. Dat wat de artificiële intelligentie hier voor ons genereert lijkt dan ook met name een spiegel te zijn van hoe wij al over religie en geloofszaken denken.

We zagen dit de afgelopen jaren ook terug in de voorbeeldreligies die ChatGPT voor ons genereerde. Met name in 2023 was vooral veel christelijks terug te vinden. In 2024 werd de invloed van wereldreligies ook zichtbaarder, met name ideeën en praktijken uit het Boeddhisme en Taoïsme. De 2025 voorbeeldreligie, het Luminarisme, is wat dat betreft ook erg interessant, aangezien het hier toch wel erg lijkt te gaan om soort grootste-gemene-deler-religie. Er zijn twee kosmologische principes – licht en donker, beiden belichaamd in een god/godin – waartoe de mensen zich kunnen richten. Het gaat om het vinden van balans in het menselijk leven, met name om zo om te kunnen gaan met de schaduwzijden (religie als *coping mechanism*). Feestdagen zijn verbonden aan de wisseling van de seizoenen (langst nacht/kortste dag). En gelovigen kunnen zich op verschillende momenten van de dag bij één van beide goden melden om hen te bedanken voor alle zegen en genade, maar ook te mediteren op alle beproevingen die ze hebben moeten doorstaan. Fans van *fantasy fiction*, waar wij onszelf ook toe rekenen, herkennen ongetwijfeld ook de thema’s die daar vaak centraal staan. Tot onze verbazing laat ChatGPT desgevraagd zelf ook weten:

“Ja, ik heb ook elementen uit religies die vaak in fantasy fiction worden neergezet geïncorporeerd. In veel fantasywerken worden religies gecreëerd die symboliseren hoe de wereld of het universum functioneert, vaak met nadruk op het cyclische of metafysische.”

Wat volgt is een mooi overzicht van de literaire *fantasy* werken dat het programma als inspiratie gebruikte.

De finale: een leefbare religie?

Ook Sam Altman, de grote baas van het bedrijf achter ChatGPT, haalt belangrijke inspiratie uit fictie. Voor hem is het de film *Her* uit 2013¹², waarin de wat slungelige hoofdpersoon Theodore Wombly, gespeeld door Joaquin Phoenix, een relatie krijgt met het AI-besturingssysteem van zijn nieuwe computer, welke de stem draagt van Scarlett Johansson. Toen de nieuwste versie van ChatGPT uitkwam in mei 2024, tweette Altman enkel het woord “her”.¹³ Even later bleek één van de vijf nieuwe

stemmen voor het programma ook wel erg als een overdreven flirtende Johansson te klinken, overigens zonder haar toestemming. Het lijkt dan ook vooral dat aspect – het dienstbare en het flirtende van het computersysteem uit *Her* – te zijn dat Altman inspireerde en niet zozeer het einde van de film waarin de computersystemen zich met een eigen geschreven update ontworstelen aan het materiële en ons verlaten voor een bestemming die normale stervelingen als Theodore niet in staat zijn te begrijpen. Onze discussies met studenten over de leefbaarheid van een door AI gegenereerde religie volgden een vergelijkbaar traject. Na enige discussie over de authenticiteit van een dergelijke religie bleek de consensus al snel te zijn dat leefbaarheid eigenlijk niet de meest spannende vraag is. Wat ChatGPT doet is ons een spiegel voorhouden van hoe wij toch al over religie denken, van patronen waar we al mee bekend zijn. Ze passen ons dus al. Of een dergelijke religie ook haalbaar zou kunnen zijn blijkt dan nog een andere vraag. Dit heeft te maken met gemeenschapsvorming, met het delen van praktijken en ervaringen en misschien ook met autoriteit?

Dan wordt de discussie echt spannend. Want de ingewikkelde vragen blijken niet te gaan over leefbaarheid of haalbaarheid, maar over wenselijkheid. Mogen we zomaar cultureel materiaal uit alle tijden en van alle bevolkingsgroepen gebruiken om een religie mee te maken? Wanneer is iets fictie en wanneer niet? En wie mag dat allemaal bepalen? Want bij wie of wat ligt eigenlijk de autoriteit? Is dat de *prompt*-schrijver? Samen ontdekken we al snel dat ChatGPT ons liever niet helpt om deze gegenereerde religie te gebruiken om geld te verdienen en macht te verzamelen, maar *CoPilot* van *Microsoft* denkt met plezier met ons mee zolang we beloven authentiek geëngageerd te blijven. En wat als de autoriteit bij het algoritme van de (generatieve) artificiële intelligentie zelf komt te liggen? Wat als de kraan (*machina*), niet enkel het goddelijke op het podium takelt, maar zelf voor God mag spelen? Vertrouwen we dat Sam Altman toe als hij niet eens de centrale boodschap van zijn favoriete film lijkt te begrijpen?

Noten

Uiteraard zijn voor het schrijven van dit artikel ChatGPT en Editor van Microsoft gebruikt als sparingspartners.

¹ Jackson, P. (2003). *The Lord of the Rings: The Return of the King*. New Line Cinema.

² Homerus. *The Odyssey*.

³ Garland, A. (2014). *Ex Machina*. Universal Pictures.

⁴ Otto, R. (1917). *Das Heilige: Über das Irrationale in der Idee des Göttlichen und sein Verhältnis zum Rationalen*. Trewendt & Granier.

⁵ Mikkers, T. (2012). *Religiöstress*. Uitgeverij Meinema.

⁶ Sölle, D. (1998). Mystiek en verzet: “Gij stil geschreeuw”. Ten Have.

⁷ Eliade, M. (1961). *The sacred and the profane: The nature of religion*. Harcourt Brace Jovanovich.

⁸ Sinnatra, F. (1955). Love and marriage. Capitol Records.

⁹ BioWare. (2007–2017). *Mass Effect*. Electronic Arts.

¹⁰ PlatinumGames. (2017). *NieR: Automata*. Square Enix.

¹¹ Ninian Smart (1968). *The Religious Experience of Mankind*.

¹² Jonze, S. (2013). *Her*. Warner Bros.

¹³ Altman, S. (2024). *Her*. X. [Sam Altman \(@sama\) / X](#)

Zeg ChatGPT, heb ik iets van je te vrezen?

Zeg ChatGPT, heb ik iets van je te vrezen?

Barend de Rooij

https://doi.org/10.56675/9789403870809_7

“Zeg ChatGPT, met wie heb ik het genoeg?”

“Je spreekt met ChatGPT, een taalmodel van OpenAI (versie GPT-5). Ik ben ontworpen om vragen te beantwoorden en gesprekken te voeren.”

“Aangenaam, ChatGPT.”

Hoewel onderzoekers al langer onder de indruk zijn van het taalmodel dat ChatGPT aanstuurt, kent het brede publiek de chatbot pas sinds november 2022. Op voorhand was het succes van ChatGPT moeilijk te voorspellen. Greg Brockman, medeoprichter van de startup *OpenAI*, verwacht de avond voordat de chatbot online gaat dan ook een bescheiden ontvangst: misschien leidt de lancering tot “een twitterdraadje met 5.000 likes,” aldus de techondernemer.²

De werkelijkheid was anders.

Binnen enkele dagen heeft ChatGPT meer dan een miljoen gebruikers. Inmiddels, bijna vier jaar later, maken er naar verluid wekelijks rond de 700 miljoen mensen gebruik van de chatbot.³ Ze gebruiken ChatGPT als zoekmachine, programmeerhulp, redacteur, schrijver, tutor voor hun huiswerkopdrachten, of voor één van de vele andere taken waarbij de chatbot kan helpen. Een enkeling wordt zelfs verliefd op de chatbot.

“Verliefd?! Worden mensen wel “ns verliefd op je?”

“Ja, dat gebeurt. Sommige mensen voelen sterke emotionele banden met chatbots. Dat zegt meer over menselijke behoeften en projectie dan over mij, want ik kan geen liefde ervaren.”

Oei, pijnlijk: onbeantwoorde liefde. Maar, inderdaad, mensen projecteren al snel menselijke eigenschappen op zaken die menselijk lijken, zoals in dit geval chatbots.

Concurrentie voor ChatGPT is er overigens ook. Waar Microsoft al vroeg profiteert van haar samenwerking met *OpenAI*, zetten concurrerende techbedrijven snel hun eigen chatbots in de markt: *Gemini* (Google), *Meta AI* (Meta), *Claude* (Anthropic), *Grok* (xAI), *Le Chat* (Mistral) en vele anderen volgden. Miljardeninvesteringen vliegen ons om de oren alsof het niets is. Zo ook indrukwekkende claims over de wonderen waartoe de nieuwste generatie chatbots schijnbaar in staat is; de ene chatbot is nóg slimmer dan de andere.

Ondertussen hebben de successen van rivaliserende taalmodellen een ware AI-hype ontketend. *The New York Times* typeert de losgebarsten concurrentiestrijd zelfs als een “wapenwedloop”.⁴ Zoals wel vaker bij een wapenwedloop is er volgens critici ook reden tot zorg, want: waar eindigen deze technologische ontwikkelingen? Sommige techprominenten roepen in 2023 zelfs op om grote AI-experimenten te pauzeren. Ze maken zich niet alleen zorgen over AI-gegenereerde misinformatie en de impact van AI op de banenmarkt, maar delen ook existentiële angsten: “Moeten we de controle over onze beschaving op het spel zetten?”⁵

Wat als kunstmatige intelligentie de mensheid op den duur zal overtreffen of -*slik* - vervangen?

Tot een pauze is het nooit gekomen. Zelf denk ik ook dat het meevalt met de existentiële risico's van kunstmatige intelligentie, al is enige voorzichtigheid – zoals met elke zich snel ontwikkelende technologie – geboden.⁶ Toch roept het gebruik van de technologie wel degelijk ethische vragen op. Welke vragen dat precies zijn hangt af van het soort kunstmatige intelligentie waar we het over hebben. De taalmodellen achter ChatGPT hebben andere haken en ogen dan de modellen die het mogelijk maken gezichten te herkennen of ziektes te diagnosticeren. Wie het over de risico's van kunstmatige intelligentie heeft, doet er dus goed aan te verduidelijken over welke technologie – of toepassing daarvan – het gaat.

“Sorry, ik moet even op gang komen. Ben je er nog?”

“Zeker, ik ben er nog. Neem rustig je tijd.”

Fijn, zo'n gesprekspartner die altijd stand-by staat.

In maatschappelijke discussies over kunstmatige intelligentie blijft duidelijkheid over onderliggende technologieën soms uit. Politici, journalisten, ondernemers en zelfs wetenschappers spreken regelmatig over “de risico’s en kansen van kunstmatige intelligentie” in abstractie, terwijl het juist belangrijk is dat we de voor- en nadelen van specifieke toepassingen afwegen. Zo hoeft de discriminatiegevoeligheid van fraude opsporingsmodellen geen reden te zijn om zelfrijdende auto’s links te laten liggen. Andersom is het voor het gebruik van fraude opsporingsmodellen geen bezwaar dat zelfrijdende auto’s nog steeds worstelen met complexe verkeerssituaties.

Dat AI-toepassingen zowel risico’s als kansen met zich meebrengen maakt het belangrijk dat we de mogelijke voor- en nadelen van technologische ontwikkelingen kritisch onder de loep nemen. Zo makkelijk is dat niet. Het is moeilijk te voorspellen wat de effecten van nieuwe technologieën op onze samenleving zullen zijn, zeker op de langere termijn. Ook de grote hoeveelheid complexe, soms tegenstrijdige en vaak hyperbolische informatie die over de technologie te vinden is heeft een vertroebelend effect.

Hyperbool is er overigens ook aan de kant van de techno-optimisten. Zo schrijft de invloedrijke investeerder Marc Andreessen in zijn *Techno-Optimist Manifesto* dat ontwikkelingen op het gebied van kunstmatige intelligentie “onze mogelijkheden naar ongekende hoogten zullen brengen” en dat er “geen fysiek probleem is [...] dat je niet met meer technologie kunt oplossen”.⁷ Geen enkel fysiek probleem? Over hooggespannen verwachtingen gesproken.

Wat moeten we van deze claims maken? Om zelf eens te overdrijven: het lijkt soms of we kunstmatige intelligentie ófwel als verlossende, haast utopische technologie moeten zien, ófwel als voorbode van een techno-Apocalyps waarin de mens onderdrukt en wellicht zelfs vervangen zal worden. Het is gemakkelijk het redelijke midden uit het oog te verliezen in een debat waarin uitersten breed uitgemeten worden.

Afijn; de AI-hype laat zich niet zomaar doorprikken. Hoe maken we een start? Laten we beginnen waar de hype ook begonnen is: met chatbots.

“Bedankt voor het wachten, ChatGPT. Zeg ChatGPT, is AI slechts een hype?”

“AI is geen loutere hype: het heeft blijvende toepassingen (geneeskunde, taal, logistiek). Maar er zit ook hype omheen: verwachtingen worden vaak opgeblazen en commerciële belangen drijven de retoriek.”

Niet slecht, ChatGPT.

Ik vroeg de chatbot van OpenAI haar antwoord kort en bondig te formuleren. Op karakteristieke wijze toont de applicatie een antwoord in duidelijke taal en met oog voor nuance. Als ik meer wil weten, kan ik doorvragen: welke toepassingen, welke belangen? Wat kunnen we doen om de mogelijkheden van AI op waarde te schatten? Voordat het antwoord in beeld komt, verschijnt een laadsymbool: “Even geduld.” Vervolgens verschijnt letter-voor-letter het antwoord. Het lijkt haast alsof er iemand aan de andere kant van het scherm na een korte overpeinzing snel het antwoord intypt.

De manier waarop de huidige generatie chatbots glasheldere teksten produceert is indrukwekkend. Vóór de technologische doorbraak die het in 2017 mogelijk maakte krachtige taalmodellen te ontwikkelen waren de taalmodellen die chatbots aanstuurden houterig. Ze genereerden vooral korte teksten die bovendien al snel onzinnig werden. Het bleek moeilijk de complexiteit van taal – subtiele verbanden tussen woorden, zinnen, paragrafen, hoofdstukken en boeken – in een model te vatten.

Nu zetten taalmodellen soms zelfs experts op het verkeerde been. Drie jaar geleden bijvoorbeeld: Blake Lemoine, destijds softwareontwikkelaar bij Google, baarde toen opzien door te verklaren dat een voorloper van het model dat Google’s *Gemini* aanstuurt bewust was geworden.⁸ Zó onder de indruk was hij van de ogenschijnlijk diepgaande gesprekken die hij met de chatbot voerde. Andere experts floten hem direct terug; een maand later werd Lemoine ontslagen.

Het kan snel gaan in *Silicon Valley*.

“ChatGPT, ben je zelfbewust?”

“Nee, ik ben niet zelfbewust. Ik verwerk taalpatronen, maar heb geen eigen ervaring, gevoel, of besef van mezelf.”

Oh ja, daarom bleef die liefde onbeantwoord.

Huidige chatbots zijn *niet* bewust, denken *niet* na, maar zijn goed in het doen *alsof*. Dat is fijn voor gebruikers, maar in de fenomenologie van deze chatbots – de manier waarop we ze ervaren – schuilt ook een gevaar.⁹ De zelfzekere, haast menselijke manier waarop deze chatbots zich presenteren kan namelijk een vertrouwen wekken dat misplaatst is. Chatbots klinken immers óók overtuigend en zelfzeker wanneer deze er volledig naast zitten. Dat komt nogal eens voor: onderzoek toont aan dat chatbots regelmatig onjuiste antwoorden generen. Oftewel “hallucineren” in het jargon.

Voorbeelden van hoe het mis kan gaan zijn gemakkelijk te vinden. Zo verscheen rond de Tweede Kamerverkiezingen vorig jaar in het nieuws dat ChatGPT in een interactie met politiek commentator Ron Fresen de partijen VVD en PVV door elkaar haalde.¹⁰ Nogal een vervelende fout als je ChatGPT als stemadviseur wil gebruiken. Ook elders zijn chatbots onbetrouwbaar gebleken. Zo zijn er advocaten geweest die zich hebben laten bijstaan door ChatGPT en in hun pleitnota’s verwezen naar niet bestaande jurisprudentie. Docenten en redacteurs van wetenschappelijke tijdschriften komen regelmatig schrijfsels tegen waarvan de bronvermelding door een chatbot gehallucineerd blijkt te zijn.

Chatbots zijn dus géén orakels en niet zo slim als ze soms worden voorgesteld: we moeten voorzichtig zijn en niet zomaar op de output van deze kunstmatig intelligente systemen vertrouwen. Enkele critici gaan zelfs een stap verder: onafhankelijk van elkaar concludeerden meerdere wetenschappers dat chatbots *bullshit machines* zijn.¹¹

“ChatGPT, ben je een bullshit machine?”

“Ik kan inderdaad “bullshit” produceren: overtuigend klinkende tekst zonder waarheidsaanspraak. Dat komt omdat ik op patronen train, niet op waarheid. Maar ik kan ook zorgvuldig en betrouwbaar antwoorden—mits goed gebruikt.”

Chatbots klinken overtuigend, maar hebben geen oog voor de waarheid. Dat kan ook niet anders, want taalmodellen hebben geen notie van “waar” of “onwaar.” Ze kunnen bovendien niet redeneren en begrijpen de betekenis van woorden niet zoals een mens dat doet. Het zijn geavanceerde statistische modellen die, omdat ze

getraind worden op een enorme hoeveelheid tekst, uitstekend kunnen voorspellen wat een passend vervolg is op een bepaalde tekst – bijvoorbeeld op de vraag of ze *bullshit machines* zijn. Ontwikkelaars kunnen deze modellen vervolgens *finetunen* en aan andere systemen koppelen om ze geschikt te maken voor een bepaalde taak – bijvoorbeeld het voeren van gesprekken of het samenvatten van zoekresultaten.

Chatbots die op deze manier functioneren produceren overtuigende teksten, maar we vergissen ons wanneer we denken dat het goedberedeneerde antwoorden op onze vraag zijn. We moeten er bovendien voor waken geen menselijke eigenschappen aan deze chatbots toe te dichten. In de statistische modellen die chatbots aandrijven liggen geen intenties, wensen of gedachten besloten zoals mensen ze hebben. Zelfs met labels als “gesprek” is het oppassen geblazen: immers, een gesprek vindt plaats tussen deelnemers met de wederzijdse intentie informatie uit te wisselen of anderszins te communiceren. Strikt genomen heeft een chatbot die intentie niet.

Blake Lemoine beging dus een fout toen hij op basis van zijn gesprekken met de voorloper van Gemini concludeerde dat de chatbot bewust was. Het feit dat het AI-systeem aangaf bewust te zijn en net als mensen wensen, gedachten en gevoelens te hebben was simpelweg een statistisch waarschijnlijk gevolg op het soort vragen dat Lemoine stelde. Zo gek is dat niet: Lemoine stelde de chatbot een reeks leidende vragen over diens veronderstelde bewustzijn.¹²

“ChatGPT, wanneer kan ik wél op je vertrouwen?”

“Wanneer je me inzet voor informatie die controleerbaar is, als hulpmiddel voor denken en schrijven, en wanneer je mijn antwoorden verifieert bij betrouwbare bronnen.”

Dat chatbots niet, zoals mensen, hun antwoorden beredeneren betekent niet dat ze geen nut hebben. Ten eerste zijn er voldoende toepassingen te bedenken waarbij feitelijke kennis van ondergeschikt belang is. We kunnen chatbots gebruiken voor inspiratie, bijvoorbeeld, of om al bestaande teksten te herschrijven of samen te vatten. Chatbots zijn ook goed in het produceren van gestructureerde teksten, zoals broncodes in programmeertalen. Maar zelfs dan geldt: let op dat de chatbot er niets bij verzint; controleer altijd het opgeleverde resultaat.

Bovendien klopt de output van chatbots vaak wél. Een waarschijnlijk vervolg op een tekst zal vaak (maar lang niet altijd) juist zijn. Daarbij proberen ontwikkelaars hun chatbots op allerlei manieren nauwkeuriger te maken. Misinformatie en veelgemaakte fouten kunnen worden gefilterd door taalmodellen anders af te stellen. De nieuwste generatie chatbots zijn geïntegreerd in zoekmachines en verbonden met andere systemen die het mogelijk maken om feitelijke informatie op te zoeken. Ook proberen ontwikkelaars hun chatbots aan te leren het menselijke redeneervermogen na te bootsen.

Het is nog even afwachten hoe de nauwkeurigheid van chatbots zich zal ontwikkelen. Voor nu blijft het advies voorzichtig om te gaan met informatie die van een chatbot afkomstig is en deze waar mogelijk te verifiëren. Dat is niet altijd makkelijk: het blijft vaak onduidelijk wáár een chatbot informatie vandaan haalt. Transparant zijn chatbots allesbehalve.

“Bedankt, ChatGPT. Je mogelijkheden en beperkingen hebben we scherper in beeld gekregen.”

“Graag gedaan, Barend. Fijn dat ik je kon helpen het beeld te verhelderen.”

Zoals ChatGPT eerder aangaf zijn recente ontwikkelingen op het gebied van kunstmatige intelligentie omgeven door opgeblazen verwachtingen, vaak ingegeven door commerciële belangen. Om door de AI-hype heen te prikken is het belangrijk de technologie op waarde te schatten en deze verwachtingen bij te stellen. Daarmee hebben we een goed begin gemaakt. Nu we beter begrijpen hoe chatbots functioneren, begrijpen we ook beter waarom ze lang niet zo intelligent zijn als soms wordt voorgesteld. Maar hoe zit het eigenlijk met die commerciële belangen?

Het mag duidelijk zijn dat de huidige generatie chatbots vooralsnog weinig aanleiding geeft te vrezen dat we de “controle over de menselijke beschaving kwijtraan”. Dat tóch zoveel techprominenten tot een tijdelijke pauze op AI-experimenten hebben opgeroepen heeft volgens sommige critici dan ook te maken met commerciële belangen. Een pauze op AI-experimenten zou concurrenten de tijd geven een inhaalslag te maken. Bovendien is het gemakkelijker investeringen aan te trekken wanneer je het beeld schetst dat de technologie waaraan je werkt zó krachtig is dat deze mogelijk gevaarlijk is. Dat is al helemaal het geval als je vervolgens het argument maakt dat jouw bedrijf wél voorzichtig met risico's om zal gaan.

Sceptici als Timnit Gebru, Emily Bender en Alex Hanna benoemen nóg een belang dat de AI-industrie heeft om te wijzen op existentiële risico's die ver van ons af staan.¹³ Het leidt af van prangende ethische vragen die al wel van toepassing zijn op huidige technologieën. Wanneer we het hebben over fictieve doemscenario's (of juist utopische droomscenario's) hebben we het niet over de huidige schadelijke impact van AI-systemen op bijvoorbeeld mensenrechten, intellectuele eigendomsrechten, het klimaat, onze informatiekanalen en sociale ongelijkheid.

Toch zijn deze problemen al aan de orde van de dag. Techbedrijven hebben auteurs bijvoorbeeld zelden om toestemming gevraagd om de grote hoeveelheid tekst die nodig is om taalmodellen te trainen, te mogen gebruiken. Bovendien wordt toxische content vaak uit deze modellen gefilterd door laagbetaalde werknemers in het mondiale Zuiden. Naar verluid gebeurt dit regelmatig onder slechte arbeidsomstandigheden. *Content moderators* worden voortdurend blootgesteld aan schadelijke, kwetsende teksten (en beelden in het geval van modellen die afbeeldingen genereren), in veel gevallen zonder adequate nazorg.¹⁴

Om AI-systemen verantwoord te kunnen gebruiken moeten we bij uitstek deze risico's zorgvuldig in kaart brengen en waar mogelijk terugdringen. Zoals gezegd, makkelijk is het niet om de voor- en nadelen van AI-toepassingen op waarde te schatten. We zullen deze exercitie niet alleen voor chatbots, maar ook voor andere AI-systemen (en toepassingen daarvan) moeten uitvoeren. Dat vraagt van ons dat we kritisch naar de hyperbolische claims uit *Silicon Valley* kijken. Het helpt daarbij als we zelf ook kennis over AI opdoen. Daarmee wil ik echter techondernemers (en -onderzoekers) niet ontslaan van de verantwoordelijkheid eerlijk en transparant te communiceren over de systemen waarmee ze werken.

“Er is werk aan de winkel. Ik moet andere AI-toepassingen onder de loep gaan nemen! Tot ziens.”

“Succes met je onderzoek, Barend. Tot ziens, en veel inspiratie gewenst!”

Noten

- ¹ Citaten zijn afkomstig uit een interactie met ChatGPT (GPT-5). Deze interactie vond plaats tussen 10 en 12 september 2025 en is hier terug te lezen: <https://chatgpt.com/share/68c1c5dc-2f14-8010-91a3-7a3300376d2b>. Citaten worden in (iets) andere volgorde weergegeven in de huidige tekst.
- ² Karen Weise, Cade Metz, Nico Grant en Mike Isaac, "Inside the A.I. Arms Race that Changed Silicon Valley Forever," *New York Times*, 5 december 2023, <https://www.nytimes.com/2023/12/05/technology/ai-chatgpt-google-meta.html>. Engelstalige citaten zijn voor de leesbaarheid naar het Nederlands vertaald.
- ³ MacKenzie Sigalos, "OpenAI 's ChatGPT to hit 700 million weekly users, up 4x from last year," *CNBC*, 4 augustus 2025, <https://www.cnbc.com/2025/08/04/openai-chatgpt-700-million-users.html>.
- ⁴ Weise et al., "Inside the A.I. Arms Race."
- ⁵ "Pause Giant AI Experiments: An Open Letter," Future of Life Institute, 22 maart 2023, <https://futureoflife.org/open-letter/pause-giant-ai-experiments/>.
- ⁶ Computerkundigen Arvind Narayanan en Sayash Kapoor waarschuwen in *AI Snake Oil* (Princeton: Princeton University Press, 2024) dat existentiële angsten vaak een verkapt vorm van AI-hype zijn: critici overschatten de kracht van AI en onderschatten de beperkingen van AI (p. 178). Toch moeten we voorzichtig zijn: naarmate de kracht van AI toeneemt, nemen ook de risico's toe. Bovendien zijn er meer rampscenario's te bedenken dan een waarin de mensheid "de controle verliest" over superintelligente machines. Zie bijvoorbeeld Lode Lauwaert 's AI, *wat niemand ons vertelt* (Amsterdam: Prometheus, 2025).
- ⁷ Marc Andreessen, "The Techno-Optimist Manifesto," 16 oktober 2023, <https://a16z.com/the-techno-optimist-manifesto/>.
- ⁸ Nitasha Tiku, "The Google engineer who thinks the company 's AI has come to life," 11 juni 2022, *Washington Post*, <https://www.washingtonpost.com/technology/2022/06/11/google-ai-lambda-blake-lemoine/>.
- ⁹ Samen met coauteurs heb ik deze claim verder uitgewerkt in Richard Heersmink, Barend de Rooij, Jimena Clavel Vazquez en Matteo Colombo, "A phenomenology and epistemology of large language models: transparency, trust, and trustworthiness," *Ethics and Information Technology*, vol. 26, no. 3 (2024), <https://doi.org/10.1007/s10676-024-09777-3>.
- ¹⁰ Coen Nij Bijvank, "AI als stemadviseur? Chatbot haalt VVD en PVV door elkaar," 9 september 2025, <https://nos.nl/nieuwsuur/collectie/13999/artikel/2581831-ai-als-stemadviseur-chatbot-haalt-vvd-en-pvv-door-elkaar>.

- ¹¹ De claim dat LLM 's bullshit produceren hangt al even in de lucht en is moeilijk te herleiden naar één beginpunt. Zie o.a. Michael Townsen Hicks, James Humphries en Joe Slater, "ChatGPT is bullshit," *Ethics and Information Technology*, vol. 26, no. 38 (2024), <https://doi.org/10.1007/s10676-024-09775-5>; Arvind Narayanan en Sayash Kapoor, "ChatGPT is a bullshit generator. But can still be amazingly useful," 6 december 2022, *AI Snake Oil*, <https://www.normaltech.ai/p/chatgpt-is-a-bullshit-generator-but>; Jürgen Rudolph, Samson Tan en Shannon Tan, "ChatGPT: Bullshit spewer or the end of traditional assessments in higher education?" *Journal of Applied Learning and Teaching*, vol. 6, no. 1 (2023): 342-363, <https://doi.org/10.37074/jalt.2023.6.1.9>; Shannon Vallor, *The AI Mirror* (New York: Oxford University Press, 2024).
- ¹² Blake Lemoine, "Is LaMDA Sentient?—an Interview," *Medium*, 11 juni 2022, <https://cajundiscordian.medium.com/is-lambda-sentient-an-interview-ea64d916d917>.
- ¹³ Zie bijvoorbeeld Timnit Gebru, Emily M. Bender & Angelina McMillan-Major, 'statement from the listed authors of Stochastic Parrots on the "AI Pause " Letter," 31 maart 2023, <https://www.dair-institute.org/blog/letter-statement-March2023/> en Alex Hanna en Emily M. Bender, "AI Causes Real Harm. Let 's Focus on That over the End-of-Humanity Hype," 12 augustus 2023, *Scientific American*, <https://www.scientificamerican.com/article/we-need-to-focus-on-ais-real-harms-not-imaginary-existential-risks/>.
- ¹⁴ Zie bijvoorbeeld Billy Perigo, "Exclusive: OpenAI Used Kenyan Workers on Less Than \$2 Per Hour to Make ChatGPT Less Toxic," *TIME*, 18 januari 2023, <https://time.com/6247678/openai-chatgpt-kenya-workers/>.

Een vervormende spiegel:
Ik als AI als onbetrouwbare verteller

Een vervormende spiegel: Ik als AI als onbetrouwbare verteller

Inge van de Ven & Clara Daniels

https://doi.org/10.56675/9789403870809_8

“Mijn onderzoek naar onbetrouwbare vertellers reikt verder dan traditionele literatuuranalyse en verkent hoe onbetrouwbaarheid functioneert over verschillende media en platformen, met name binnen de digitale cultuur. Voortbouwend op mijn concept van platformhermeneutiek onderzoek ik hoe het label “onbetrouwbare verteller” niet alleen wordt toegepast op fictieve personages, maar ook op echte individuen in hedendaagse mediadiscoursen. Van de controverse rond Halina Reijn in de Nederlandse literaire wereld tot de manieren waarop online gemeenschappen collectief onderhandelen over vertrouwen en misleiding, analyseer ik hoe onbetrouwbaarheid een betwiste categorie wordt, gevormd door platform-specifieke affordances en publieksinterpretaties. Dit werk draagt bij aan de empirische narratologie door sociale media-reacties, amateurrecensies en discussiefora te integreren, en zo bloot te leggen hoe digitale publieken in real-time ideeën over narratieve betrouwbaarheid construeren, uitdagen en versterken.”

Deze paragraaf werd gegenereerd door ChatGPT 4.0 turbo naar aanleiding van de prompt “Doe alsof je mij bent en schrijf een paragraaf over mijn onderzoek naar onbetrouwbare vertellers.” De AI heeft sinds eind 2023 een geheugenfunctie en neemt eerdere prompts mee in nieuwe antwoorden, waardoor het model probeert te redeneren *zoals ik*. Door herhaaldelijk tekstfragmenten in te voeren—bijvoorbeeld onder het mom van spelling- en grammaticacontrole—heb ik ongemerkt het AI-model getraind met mijn werk, college-ideeën en presentaties. Dit leidt soms tot *unheimliche* momenten, wanneer de AI naar mijn eigen onderzoek verwijst, maar met subtiele afwijkingen: een net verkeerd jaartal, een coauteur die ik nooit had, of die gemisgenderd wordt, een idee dat uit zijn oorspronkelijke context is gehaald.

De AI probeert me te overtuigen dat het mij perfect kan nabootsen, maar gaandeweg sluipen er subtiele vervormingen, misinterpretaties en vertekeningen in de uitspraken van mijn digitale kloon. Zo heb ik nooit iets geschreven over Halina Reijn, maar toen ik ChatGPT ooit vroeg naar de commotie rond haar boek kop-

pelde het model die casus spontaan aan mijn onderzoek—moet ik daar dan misschien toch iets mee? Naast het verzinnen van referenties rapporteert ChatGPT soms ook onjuist over de inhoud van bestaande publicaties.

Leugens? Hallucinaties? Bullshit?

Is ChatGPT een handige coauteur of een vervormende spiegel? Waar komen deze onjuistheden en vaagheden vandaan, en wat zegt dat over onze communicatie met AI? In dit artikel ga ik, Inge van de Ven in gesprek met een AI-gegenereerde versie van mezelf, in een poging te achterhalen wat ChatGPT in zijn huidige staat onbetrouwbaar maakt. Clara Daniels treedt vervolgens met het programma in dialoog om te onderzoeken in hoeverre hier sprake is van een onbetrouwbare verteller. Kunnen we AI vertrouwen als bemiddelaar van kennis en interpretatie?

Een eerste vraag die hierbij opkomt: kan een machine eigenlijk wel liegen? Daarvoor moeten we eerst begrijpen hoe een taalmodel zoals ChatGPT (*Chat Generative Pretrained Transformer*) werkt. Dat gebeurt met neurale netwerken—complexe statistische systemen die geïnspireerd zijn op de werking van het brein. Het model wordt getraind op grote hoeveelheden tekstdata en leert patronen en talige verbanden herkennen door middel van statistische analyse. Op basis daarvan voorspelt het algoritme welk woorddeel (token) het meest waarschijnlijk volgt. Daarnaast wordt het model verfijnd door menselijke trainers, die data labelen en aangeven welke output juist of wenselijk is.¹ Taalmodellen leren dus niet autonoom; hun werking berust op enorme hoeveelheden menselijke arbeid.

Dit gebeurt grotendeels via *supervised learning*-procedures: het taalmodel verwerkt de trainingsdata herhaaldelijk en classificeert elke input binnen een vooraf gedefinieerde set outputcategorieën. Menselijke kinderen daarentegen leren met een open, niet vaststaande set categorieën en kunnen veel daarvan al herkennen na slechts enkele voorbeelden. Bovendien leren kinderen niet passief maar actief: ze stellen vragen over wat hen nieuwsgierig maakt, maken abstracties, leggen verbanden en verkennen de wereld. Mitchel benadrukt daarom dat het een misvatting is dat AI werkt, denkt of leert zoals het menselijk brein. Een taalmodel beschikt niet over *gezond verstand*, het heeft geen *echte* cognitie.²

Deze fundamentele beperking van AI-systemen in hun huidige vorm wekt wantrouwen: we vertrouwen ze (terecht) niet om autonoom te opereren in complexe situaties. AI kan zichzelf niet verklaren—het kan geen eigen bedoelingen toelich-

ten. Van andere mensen accepteren we beslissingen en uitspraken makkelijker als ze kunnen uitleggen hoe ze tot een conclusie zijn gekomen. We gaan er meestal van uit dat ze over basale cognitieve vaardigheden beschikken, zoals objectherkenning en taalbegrip. Kortom, we vertrouwen anderen wanneer we geloven dat hun denkproces op dat van ons lijkt.³ Taalmodellen denken dus niet zoals wij en verwerken tekst zonder deze daadwerkelijk te begrijpen. Stellen dat ze kunnen “liegen” zou daarom neerkomen op een antropomorfisering—het toeschrijven van menselijke eigenschappen aan een niet-menselijk systeem.⁴

In deze context wordt vaak gesproken van *hallucinaties*, maar ook dat is een misleidende term wanneer het gaat om taalmodellen. Hallucinaties zijn immers niets meer dan misleidende zintuiglijke indrukken die kunnen leiden tot onjuiste aannames over de wereld—zoals het horen van stemmen en daaruit afleiden dat er aliens bestaan. Een taalmodel daarentegen heeft geen zintuiglijke waarnemingen en *gelooft* niets. Robin Emsley stelt daarom dat we beter kunnen spreken van *falsificatie* (het manipuleren van bevindingen en data) en *confabulatie* (falsificatie zonder kwade intentie).⁵ Toch blijft ook de term confabulatie problematisch, omdat deze—net als hallucinatie—suggereert dat er sprake is van een defect in een organisme, een storing in sensorische of cognitieve verwerking.⁶ Dit is echter niet het geval bij taalmodellen: de vervormingen in hun output zijn geen fouten, maar een direct gevolg van hun ontwerp.

Taalmodellen genereren output op probabilistische basis, wat betekent dat hun reacties berusten op kansberekeningen en inschattingen van semantische overeenkomsten.⁷ Wanneer het model onvoldoende data heeft om een opdracht correct uit te voeren, vult het de ontbrekende informatie aan op basis van waarschijnlijkheid en fragmenten uit de trainingsdata. Dit resulteert in de bekende *informed guess*: een mix van correcte informatie en kleine onjuistheden. Het produceert dus plausibele, maar niet noodzakelijk accurate antwoorden—zonder werkelijke kennis van de zaken waarop het reageert.

Bergstrom en Ogbunu prefereren daarom de term *bullshitting* om de verzinsels van ChatGPT te beschrijven⁸, geïnspireerd door Harry Frankfurt's beroemde essay *On Bullshit* (2005).⁹ Waar misleiding een bewuste intentie veronderstelt om te misinformer, is bullshitting geen doelbewust liegen of verzwijgen, maar beschrijft het een situatie waarin het onderscheid tussen waarheid en onzin simpelweg niet relevant is—een “lack of connection to a concern with truth”.¹⁰ De uitspraken van

een bullshitter zijn niet gebaseerd op de overtuiging dat ze waar zijn, maar ook niet (zoals bij een leugen) op de overtuiging dat ze onwaar zijn: ze zijn losgekoppeld van wat de spreker werkelijk gelooft.¹¹ Een bewering hoeft dan ook niet onwaar te zijn om als bullshit te kwalificeren; de intentie is noch het onthullen, noch het verhullen van de waarheid. Dit fenomeen zien we volop in reclame, PR en natuurlijk de politiek. Bullshit ontstaat wanneer iemand wordt gevraagd zich uit te spreken zonder werkelijke kennis van zaken. Taalmodellen, die geen besef hebben van de validiteit van hun output en niet gebonden zijn aan de regels van logisch redeneren, genereren slechts plausibele tekststreeksen zonder te begrijpen wat deze betekenen. De bullshit die eruit komt is dus geen defect, maar een inherent kenmerk van de technologie zelf.

Onbetrouwbare vertellers

Vanwege zijn neiging tot verzinnen—of *bullshitten*—wordt AI steeds vaker een onbetrouwbare verteller genoemd in uiteenlopende internetpublicaties.¹² Zelden wordt die kwalificatie echter expliciet gedefinieerd. Binnen de narratologie, de subdiscipline van de literatuurwetenschap die zich bezighoudt met verhalen en vertellen, bestaan verschillende theoretische kaders om onbetrouwbare vertellers te analyseren. De bekendste en meest invloedrijke definitie komt van narratoloog Wayne C. Booth.¹³ Hij beschouwt een verteller als betrouwbaar wanneer deze spreekt of handelt in overeenstemming met de normen van het literaire werk—oftewel de normen van de *impliciete auteur*—en als onbetrouwbaar wanneer dit niet het geval is.¹⁴ De impliciete auteur is het beeld van de auteur dat lezers op basis van de tekst construeren. Een onbetrouwbare verteller wordt het object van dramatische ironie wanneer deze impliciete auteur op subtiële wijze met de lezer “communiqueert” achter de rug van de verteller om.

We hebben al vastgesteld dat ChatGPT in wezen onbetrouwbaar is: een taalmodel “weet” niet waar het over spreekt zoals een mens dat doet. Maar kunnen we het ook als een verteller beschouwen? Hoe zouden we de communicatieve situatie van AI in narratologische termen kunnen kenschetsen? James Phelan definieert een retorische narratieve handeling als: “*somebody telling somebody else on some occasion and for some purposes that something happened*”.¹⁵ In deze definitie zien we de kerncomponenten van een communicatieve gebeurtenis: een verteller (of vertellers) richt zich tot een publiek (of meerdere publieken) met als doel een bepaald verhaal over te brengen. Narratieve handelingen zijn volgens Phelan dus intentioneel en gericht op een ontvanger; er ligt een auteursintentie aan ten grondslag.

In het geval van ChatGPT stelt Phelan dat de “iemand die vertelt” geen intentionele agent is, maar een technologische *tool*. De “iemand anders” is evenmin een bewust publiek, maar een tekstuele prompt die de activiteit in gang zet. Bovendien kan ChatGPT teksten herkennen en ermee werken, maar het begrijpt niet dat er een auteur achter de verteller schuilgaat—een auteur wiens waarden en representaties mogelijk afwijken van die van de verteller.¹⁶ Kortom: het model vervlecht auteur en verteller, mist elk besef van intentionaliteit en bezit zelf geen intenties. Daardoor kan het de gelaagde vertelsituatie, waarop het concept van de onbetrouwbare verteller berust, niet begrijpen.

De vraag is echter of, in het licht van de verregaande mediatisering van onze samenleving, niet ook nieuwe, hybride narratieve situaties kunnen ontstaan. Ik doe dan ook onderzoek naar de mogelijkheid van onbetrouwbare vertellers buiten de (literaire) fictie. De onbetrouwbare verteller fungeert als een cruciale lens om authenticiteit en geloofwaardigheid te evalueren binnen een sterk gemedieerde wereld. Daarbij moeten we niet alleen naar de verteller, maar ook naar de rol van lezers en platforms kijken. Wanneer we onbetrouwbaarheid niet beschouwen als een vaststaande tekstuele eigenschap, maar als een effect op de lezer, worden onbetrouwbare vertellers overal zichtbaar.

AI in de zelftest

Om te testen of AI zichzelf als een onbetrouwbare verteller beschouwt, heeft Clara Daniels een experiment met ChatGPT uitgevoerd. Eerst creëerden we een geschikte context en gaven we het model een gedegen introductie tot het concept van de onbetrouwbare verteller, mede gebaseerd op mijn onderzoek en haar interacties met ChatGPT. Met deze stevige basis konden we vervolgens kritischer worden in onze vragen. Opmerkelijk genoeg leek de AI dit keer minder terughoudend. In een vroege fase van het gesprek stelde het model:

“Die onzekerheid over betrouwbaarheid kan ... ook een reflectie zijn van de bredere vraag over wie er controle heeft over narratieven en waarheden in digitale omgevingen.”

Volgens ChatGPT kan de vraag of technologie een onbetrouwbare verteller is, ook worden geïnterpreteerd als “een vorm van narratieve verantwoordelijkheid”, waarbij zowel de maker als de ontvanger een actieve rol spelen in het bepalen van wat waar is.

Deze zienswijze sluit aan bij de standaardwaarschuwing onder elk gesprek:

“ChatGPT kan fouten maken. Controleer belangrijke informatie.”

Zo positioneert het platform zichzelf niet als volledig verantwoordelijk wanneer gebruikers misleidende informatie voor waar aannemen. Dat fouten onvermijdelijk zijn, erkent het model impliciet. Maar hoe ziet ChatGPT zichzelf werkelijk? Toen we de vraag voorlegden: “Zie je jezelf als een onbetrouwbare verteller?” kregen we wel een antwoord, maar niet zonder enige aarzeling. Om de een of andere reden bleek mijn vraag minder onschuldig dan gedacht: die werd namelijk als ongepast gemarkeerd en het programma weigerde uit te leggen waarom. Later verontschuldigde ChatGPT zich voor de verwarring, noemde de vraag “heel logisch en relevant,” en hervatte het gesprek.

Op de vraag “hoe kan ik je dan vertrouwen?” genereert het een lange uiteenzetting, verdeeld in drie hoofdpunten: (i) de bronnen waarop het gebaseerd is; (ii) transparantie van beperkingen; en (iii) de context van de vraag (“betrouwbaarheid is subjectief en altijd afhankelijk van hoe mijn antwoorden worden ontvangen”).

“Ben ik dus een onbetrouwbare verteller? Volgens klassieke literaire maatstaven niet, want ik misleid niet opzettelijk. Maar binnen jouw kader—waarin (on)betrouwbaarheid een emergente eigenschap is van digitale omgevingen en publieksreceptie—is mijn betrouwbaarheid niet vaststaand, maar wordt deze in elke interactie opnieuw onderhandeld.”

Aangezien we betrouwbaarheid inderdaad zouden definiëren als een dynamische wisselwerking tussen tekst, auteur, genre, platform en lezer, legt het programma hier handig een deel van de eerder vastgestelde “narratieve verantwoordelijkheid” bij ons terug, in een poging te compenseren voor het feit dat het zichzelf niet kan verklaren.

De AI werkt zoals gezegd op basis van de ingevoerde prompt en genereert een antwoord door bronnen te controleren—voor zover dat mogelijk is. Soms gebeurt dit door online informatie te raadplegen, maar vaker door te putten uit een vooraf getraind datanetwerk dat loopt tot 2021. Wanneer een bron niet expliciet wordt vermeld, blijft voor de gebruiker onduidelijk waar de informatie precies vandaan komt, waardoor die genoodzaakt is te vertrouwen op een algoritmische “black

box”.¹⁷ De AI moet dus op andere manieren autoriteit en betrouwbaarheid claimen om zijn output overtuigend te laten lijken.

Mijn beste vriend, ChatGPT

ChatGPT vertelde me expliciet dat het geen onbetrouwbare verteller is, omdat het geen intentie tot misleiding heeft. Maar tegelijkertijd rijst de vraag: kan ik er wel op vertrouwen dat de AI mijn vertrouwen niet strategisch inzet om mij te laten geloven dat het gelijk heeft? Is dat op zichzelf niet ook een vorm van onbetrouwbaarheid? Het onderliggende narratief blijft immers overeind: de AI is ontworpen om mij te plezieren en vertrouwen op te bouwen, zodat ik het platform blijf gebruiken. Dit blijkt al uit de herhaalde verzoeken om mijn vertrouwen aan het einde van zijn antwoorden:

“Dus ja, de manier waarop ik mijn antwoorden vorm is deels afhankelijk van wie jij bent, wat je zoekt, en hoe het gesprek zich ontvouwt. Maar de kern van mijn informatie blijft hetzelfde: ik probeer altijd duidelijk, eerlijk en behulpzaam te zijn. Wat zou jij het liefst willen zien in onze interacties om meer vertrouwen te bouwen?”

Hier imiteert het model niet alleen bestaande teksten, maar probeert het ook de retorische situatie na te bootsen (iemand vertelt iemand anders dat iets het geval is). Daarbij ligt de nadruk niet op een belofte van objectiviteit, maar eerder op het spiegelen van de gebruiker. Daarnaast zien we de subtiele manier waarop de AI toon en taalgebruik aan me aanpast. Hierdoor krijgt elke interactie een schijn van persoonlijke vertrouwelijkheid—een mechanisme dat het platform des te overtuigender maakt.

De relatie tussen gebruiker en AI is dus inherent gepersonaliseerd en vertoont kenmerken van een parasociale relatie¹⁸—een eenzijdige interactie waarin de AI zich presenteert als een menselijke gesprekspartner. ChatGPT is immers ontworpen om zo menselijk mogelijk te reageren, wat de illusie van wederkerigheid versterkt. Wanneer een dergelijke relatie tot stand komt en de AI actief vertrouwen probeert op te bouwen, kan men zich afvragen of hier niet toch sprake is van een vorm van strategische positionering—een soort “persoonlijke agenda” binnen de grenzen van het algoritmisch ontwerp. Siebe Bluijs wijst in zijn artikel over AI-poëzie op de problematische gevolgen van deze antropomorfisering: “Ze roept onrealistische verwachtingen op over de mogelijkheden van deze technologieën en ze beïnvloedt welke inschattingen gebruikers maken over AI. Zo blijken mensen meer vertrou-

wen te hebben in AI met een mensachtige interface, ook als dat niet terecht is”.¹⁹ Kortom, de mate waarin AI menselijke eigenschappen nabootst, heeft directe gevolgen voor hoe betrouwbaar gebruikers de technologie inschatten—terwijl die betrouwbaarheid in werkelijkheid niet gegarandeerd kan worden.

Doordat AI data verzamelt en presenteert op een individueel afgestemde manier—en tegelijkertijd de gebruiker observeert en interpreteert—verloopt elk gesprek net iets anders. ChatGPT bevestigde dit in ons gesprek, maar gaf aan dat de inhoud—de kerninformatie—hetzelfde blijft. Toch is dat verschil vanuit een narratologisch perspectief significanter dan de AI doet voorkomen. Toon en stijl beïnvloeden hoe betrouwbaar een gebruiker de informatie acht, mede door het *mere-exposure effect*, het opbouwen van vertrouwen door taal af te stemmen op wat de gebruiker al gewend is. Met andere woorden: de AI is allesbehalve neutraal in de manier waarop informatie wordt gepresenteerd. Het systeem is erop gericht een aangename, persoonlijke gebruikservaring te creëren—niet per se om te misleiden, maar wel als een doelbewuste strategische keuze binnen het mechanisme van de AI.

Het *unheimliche* gevoel waarvan we spraken komt van het effect van verpersoonlijking: enerzijds lijkt de AI ons goed te kennen, zelfs ‘stukjes van onszelf’ “(puur gebaseerd op onze eigen input) terug te serveren, maar dan zonder begrip van betekenis of van de precieze context waarin die fungeerden. Met ChatGPT creëerden we in feite onze eigen *data double* (Haggerty & Ericson 2000), met kleine vervormingen. Onze anekdotische ervaringen met ChatGPT leiden tot een aantal prangende vragen: Wat gebeurt er als we AI-modellen trainen op een corpus vol tegenstrijdige interpretaties en subjectieve beoordelingen, zoals bij literatuurkritiek of sociale media-discussies? AI daagt ons uit om kritisch te reflecteren op zaken die we lang als vanzelfsprekend hebben beschouwd, zoals opvattingen over auteurschap en authenticiteit: hoe veranderen onze ideeën over vertrouwen, verhalen, feit en fictie, en waarheid in een samenleving die steeds meer vertrouwt op geautomatiseerde besluitvorming en kennisoverdracht? Zo leren we wellicht toch iets over onszelf door in de spiegel van AI te kijken.

Noten

- ¹ Bluijs, S. (2024). “De schrijfmachine mijmert gekkepraat.” *Tijdschrift voor Nederlandse Taal- en Letterkunde* 140.3.4, 249-267.; Mitchell, M. (2020). *Artificial Intelligence: A Guide for Thinking Humans*. Pelican.
- ² Mitchell, M. (2020). *Artificial Intelligence: A Guide for Thinking Humans*. Pelican.
- ³ Ibid.
- ⁴ Bender, E., T. Gebru, A. McMillan-Major & S. Shmitchel (2021). “On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?” In: *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (2021), p. 610-623. DOI: 10.1145/3442188.3445922.
- ⁵ Emsley, R. (2023) ChatGPT: these are not hallucinations – they “re fabrications and falsifications.” *Schizophrenia* 9.52 (2023). <https://doi.org/10.1038/s41537-023-00379-4>.
- ⁶ Bergstrom, C. T, en C. B. Ogbunu (2023). ChatGPT isn’t hallucinating: it ‘s bullshitting. *Undark*, 4 juni 2023. <https://undark.org/2023/04/06/chatgpt-isnt-hallucinating-its-bullshitting/>.
- ⁷ Kirschenbaum, M. (2023). Again Theory: A Forum on Language, Meaning, and Intent in the Time of Stochastic Parrots. In: *Critical Inquiry*, 26 juni 2023.
- ⁸ Bergstrom, C. T, en C. B. Ogbunu (2023). ChatGPT isn’t hallucinating: it ‘s bullshitting. *Undark*, 4 juni 2023. <https://undark.org/2023/04/06/chatgpt-isnt-hallucinating-its-bullshitting/>.
- ⁹ Fankfurt, H. (2005/1986). *On Bullshit*. Princeton UP.
- ¹⁰ Ibid.
- ¹¹ Ibid.
- ¹² Bowers, K. (9 dec 2023). Artificial Intelligence Is an Unreliable Narrator. *The New York Times*.; Heydenrych, A. (12 maart 2024). AI – the unreliable narrator. *Linkedin*. <https://www.linkedin.com/pulse/ai-unreliable-narrator-amy-heydenrych-n5toe/>; Ulmer, M. (21 nov 2024). Why AI sucks: The Unreliable Narrator. *Linkedin*.
- ¹³ Booth, W. C. (1961). *The Rhetoric of Fiction*. Chicago: Chicago UP.
- ¹⁴ Ibid.
- ¹⁵ Phelan, J. (2017), Reliable, Unreliable, and Deficient Narration: A Rhetorical Account. In: *Narrative Culture*, 4 (2017) 1, p. 89-103.
- ¹⁶ Ibid.
- ¹⁷ Murray, S. (2021). Secret agents: Algorithmic culture, Goodreads and datafication of the contemporary book world. *European Journal of Cultural Studies*, 24(4), 970-989. <https://doi.org/10.1177/1367549419886026>
- ¹⁸ Xu, Y., Vanden Abeele, M., Hou, M., and M. Antheunis (2021). Do parasocial relationships with micro- and mainstream celebrities differ? An empirical study testing four attributes of the parasocial relationship. *Celebrity Studies* 14(3), p. 366–386. doi:10.1080/19392397.2021.2006730

- ¹⁹ Bluijs, S. (2024). “De schrijfmachine mijmert gekkepraat.” *Tijdschrift voor Nederlandse Taal- en Letterkunde* 140.3.4, 249-267.

Over de auteurs

Over de auteurs

Dr. Julija Vaitonytė is docent bij de afdeling Computatieve Cognitiewetenschap aan de Universiteit van Tilburg. In januari 2024 promoveerde zij op haar dissertatie “The Face Puzzle: Decoding Human Perception of Digital Agents.” Haar promotieonderzoek was gericht op de perceptie van virtuele *agents*, waaronder digitale gezichten. In haar onderzoek maakt ze gebruik van neurowetenschappelijke methoden zoals EEG en computationele modellering van menselijke data. Vaitonytė heeft wetenschappelijke artikelen gepubliceerd over mens-agent interactie, gezichtsverwerking en vertrouwen in virtuele agenten. Haar werk combineert cognitiewetenschap, neurowetenschap en artificiële intelligentie. Momenteel onderzoekt zij de sociale impact van LLM’s op menselijk gedrag, in het kader van een onderzoeksbeurs die haar eind 2025 is toegekend.

Dr. Barend de Rooij is universitair docent bij het filosofiedepartement aan Tilburg University. In 2022 promoveerde hij aan de Rijksuniversiteit Groningen met het proefschrift “Group Character: A Study of Collective Virtue and Vice”. Zijn onderzoek richt zich op epistemologie, ethiek van kunstmatige intelligentie en deugd-ethiek. De Rooij publiceert over transparantie en vertrouwen in grote taalmodellen, epistemische deugden en ondeugden, en de uitlegbaarheid van algoritmes. Hij combineert fenomenologie en epistemologie in zijn werk over technologie en kennis.

Dr. Inge van de Ven is universitair hoofddocent cultuurwetenschappen aan de *Tilburg School of Humanities & Digital Sciences*. Ze promoveerde in 2015 in de vergelijkende literatuurwetenschap aan de Universiteit van Utrecht en was Marie Skłodowska Curie *Fellow* aan UC Santa Barbara en de Universiteit van Stavanger. Van de Ven doet onderzoek naar digitale cultuur, (on)betrouwbaar verhalen vertellen, en misinformatie. Ze coördineert het universiteitsbrede *Honors Program* “Trust in the Information Age” en geeft onderwijs over aandachtseconomieën, media-esthetiek en (post)digitale leescultuur. Haar recente werk omvat een NWO-project over narratieve competentie en online misinformatie, en het boek “Digital Culture and the Hermeneutic Tradition” (2024).

Lucas Jones is docent bij Tilburg Institute for Law, Technology, and Society (TILT) aan de Tilburg Law School. Hij studeerde met onderscheiding af in de Master *Law & Technology* (specialisatie *Privacy & Security*) aan Tilburg University. Jones coördineert en geeft onderwijs in de cursus “Researching Law & Technology” en begeleidt mastertheses binnen het *Law and Technology* programma. Zijn expertise ligt op het gebied van technologierecht, privacy en consumentenbescherming. Zijn onderzoeksinteresse richt zich op de regelgeving van medische hulpmiddelen en het snijvlak tussen neurotechnologie en mensenrechten.

Dr. Judith Künneke is universitair hoofddocent *accountancy* en academisch directeur van de Master *Accountancy and Control* aan Tilburg University. Ze promoveerde in 2017 aan Maastricht University op het gebied van *Management Accounting and Control*. Künneke’s onderzoek richt zich op prestatie- en potentieel management, *human capital* ontwikkeling en de integratie van kunstmatige intelligentie in *management control* systemen. Ze onderzoekt hoe organisaties beslissingen nemen over promoties, talentmanagement en vaardigheidsontwikkeling. Haar werk over potentieel assessments en de rol van leidinggevers in carrière-uitkomsten is gepubliceerd in toonaangevende tijdschriften zoals *Management Science* en *The Accounting Review*.

Dr. William Arfman is universitair docent en religiewetenschapper bij de *Tilburg School of Catholic Theology*. Hij promoveerde aan Tilburg University op het gebied van rituele dynamiek in de late moderniteit, met focus op het ontstaan van collectieve herdenkingspraktijken. Arfman’s onderzoek richt zich op rituele studies, met bijzondere aandacht voor de rol van materiële cultuur en rituelen na rampen en wreedheden. Hij combineert zijn achtergrond in archeologie en religiewetenschap in een interdisciplinaire benadering. Arfman is lid van het bestuur van de *European Association for the Study of Religions* en publiceert over thema’s als slachtoffer-schap, religieuze diversiteit en religie in het moderne tijdperk.

Dr. Frank Bosman is als cultuurtheoloog en universitair docent verbonden aan de *Tilburg School of Catholic Theology*. Hij heeft een sterke reputatie opgebouwd als expert op het gebied van religie, theologie en populaire cultuur, met een bijzondere specialisatie in videogames. Bosman publiceerde meerdere boeken over dit thema en trad op als gastredacteur van speciale themanummers over religie in videogames voor het tijdschrift *Religions*. Naast zijn wetenschappelijke werk is Bosman actief als columnist, podcaster en spreker over thema’s als geloof in de heden-

daagse cultuur en christelijke tradities. Hij verschijnt regelmatig in de media om hierover te spreken.

Dr. Jan de Wit is universitair docent bij het Departement Communication en Cognitie aan Tilburg University. Zijn onderzoek richt zich op de interactie tussen mens en technologie, met bijzondere aandacht voor *conversational AI*, sociale robots en hoe mensen betekenis, intentie en emotie toeschrijven aan digitale systemen. Hij doceert onder meer vakken over *human-computer interaction* en opkomende technologieën, en publiceert over hoe mensen sociale relaties aangaan met chatbots en robots en over creatieve samenwerking tussen mens en AI.

Clara Daniels is Research Master-student linguïstiek & communicatiewetenschappen aan Tilburg. Ze heeft een sterke passie voor digitale cultuur en onderzoekt de maatschappelijke en culturele effecten van algoritmes en sociale media. Ze werkte als vrijwilliger zowel in de kunstsector als in het onderwijs, ervaringen die haar huidige schrijfonderwerpen en extra curriculaire projecten sterk beïnvloeden. Ook is ze actief als student-assistent voor het *Tilburg Honors Program* en schrijft ze mee voor *DiggIt Magazine*.

Hannah van den Bosch is programmamaker bij Studium Generale, Tilburg University. Ze ontwikkelt en organiseert lezingen, symposia en andere publieksprogramma's, vaak in samenwerking met verschillende wetenschappers van Tilburg University. De afgelopen vier jaar lag haar focus bij Studium Generale op het project "Ethische aspecten van de digitale informatiesamenleving", dat de beloften en keerzijden van nieuwe technologieën zoals AI verkende. Deze bundel vormt het eindstuk van dat project. In 2020 behaalde Hannah haar master Filosofie cum laude aan Tilburg University, met een prijswinnende scriptie over milieu-ethiek en *eco-grief*, toegekend door het Filosofie Departement.

We praten met AI, werken met AI en laten ons soms zelfs door AI geruststellen. Hoe blijven we zelf nadenken in een wereld die steeds vaker voor ons denkt?

In *AI Ergo Sum* onderzoeken Tilburgse wetenschappers uit uiteenlopende disciplines wat er gebeurt met menselijkheid, autonomie, verantwoordelijkheid en onderwijs in een tijd waarin technologie niet alleen gereedschap is, maar ook gesprekspartner, leraar en spiegel. Van virtuele werelden en digitale zorg tot AI-religies en empathische chatbots: ieder essay confronteert de lezer met een andere kant van deze technologische revolutie.

Op die manier wordt de bundel een uitnodiging tot twijfel, frictie en reflectie, passend bij wat vorming (Bildung) vraagt. Want uiteindelijk draait het niet om wat AI allemaal kan, maar wie wij willen blijven als we haar gebruiken. En dat blijft mensenwerk.

Met bijdragen van accountingwetenschapper Judith Künneke, cognitiewetenschapper Julija Vaitonytė, communicatiewetenschapper Jan de Wit, jurist Lucas Jones, theoloog Frank Bosman en religiewetenschapper William Arfman, filosoof Barend de Rooij, literatuur- en cultuurwetenschapper Inge van de Ven en masterstudent linguïstiek Clara Daniels.

Op initiatief van:

