

AI LEERT DENKE

Het succesverhaal van *deep learning* is ongekend als je kijkt naar de ontwikkelingen van de afgelopen zes jaar. Plotseling vindt *artificial intelligence*, ofwel AI, toepassingen in alle delen van het bedrijfsleven en de maatschappij. Maar om ook in de toekomst nuttig te blijven, zal AI zich wel moeten blijven ontwikkelen. tekst Marc Seijlhouwer MSc

Ik zie geen enkele reden waarom AI niet net zo slim kan worden als de mens. In theorie kan AI zelfs nog veel intelligenter worden dan een menselijk brein. Filosofisch gesproken verschillen het brein en een computerprogramma immers niet wezenlijk van elkaar. Dat stelt prof.dr. Max Welling, hoogleraar *machine learning* aan de Universiteit van Amsterdam en schrijver van het binnenkort te verschijnen boek *Over leven met kunstmatige intelligentie*.

Prof.dr. Marcel van Gerven, hoofd van de onderzoeksgroep AI aan de Radboud Universiteit Nijmegen, is het met Welling eens. 'Het brein is net als een computer een fysisch mechanisme. Tenzij je denkt dat er iets ontastbaars in ons is, een ziel bijvoorbeeld. Dat mag, maar als wetenschapper zie ik dat niet zo.'

'Ik zie geen filosofische reden waarom een echt intelligente AI niet kan bestaan', denkt ook prof.dr. Frank van Harmelen, hoogleraar *knowledge representation* aan de Vrije Universiteit.

Als je het de onderzoekers vraagt die in de voorlinies van het AI-onderzoek staan, tekent zich dus een duidelijke toekomst af. Een AI die de mens evenaart of zelfs overtreft, is mogelijk. Wanneer zo'n AI verschijnt, is echter onbekend. Van Harmelen: 'Ik denk niet dat ik het nog ga meemaken.' Welling: 'Het duurt nog wel even.'

Maar hoe zetten we nu dan de eerste stappen richting die stip aan de horizon? Daarvoor zijn meerdere opties. Ten eerste het verbeteren van *machine learning*, in het bijzonder *deep learning*, en tevens het beter uitleggen en begrijpen van de algoritmes. Daarnaast zijn er een aantal andere vormen van kunstmatige intelligentie die binnenkort hun plek in de spotlight zullen opeisen. En op de lange termijn zal ons eigen brein weer in beeld moeten komen als bron van inspiratie.

TILLEN

NALS MENS

Illustratie Google

DeepDream, een *deep-learning*-algoritme van Googlebedrijf DeepMind, vervormt plaatjes tot de raarste beelden. Een normale foto van een landschap wordt via de tientallen lagen van het algoritme langzaam een veelheid aan pagodes, aquaducten, abstracte spiraalvormen en meer.

Om deep learning verantwoord en praktisch te kunnen gebruiken, is het belangrijk dat gebruikers van AI hun algoritmes kunnen uitleggen. Dat helpt om vooroordelen en fout gebruik te voorkomen.

Deep learning is gebaseerd op neurale netwerken; kunstmatige netwerken die zijn ontworpen met de neuronestructuur in onze hersenen als voorbeeld. Een van de problemen waar deep learning mee kampt, is dat het denkproces van zulke algoritmes vaak een *black box* is: de makers kunnen het proces door tientallen beslissingslagen in het neurale netwerk niet goed volgen. Dat staat toepassingen soms in de weg: als een bedrijf of overheid AI gebruikt om beslissingen te nemen, maar niet kan uitleggen waarom een bepaalde beslissing wordt genomen, dan is zo'n beslissing onverkooptbaar. Mede daarom zijn academici en bedrijfsleven hard bezig te proberen algoritmes te verklaren. Dat doen ze bijvoorbeeld door de ingevoerde data een beetje aan te passen en dan te kijken hoe de resultaten daardoor veranderen.

Aan de Universiteit Twente houdt prof.dr.ir. Bernard Veldkamp zich hiermee bezig. Hij richt zich daarbij op het opsporen van signalen van posttraumatische stress-stoornissen (PTSS) in blogs en andere geschreven teksten. 'Met een woordherkenningsalgoritme dat zichzelf traint om bepaalde woorden te herkennen als typisch voor PTSS-patiënten kun je al symptomen opsporen en mensen vroeg helpen. De grote vraag is: waarom weet het algoritme welke woorden bij PTSS horen?' Dat probeert Veldkamp met zijn onderzoek te bepalen.

Het aanpassen van de input en het kijken naar de output werkt hier redelijk goed, doordat specifieke woorden in een tekst mogelijk op PTSS duiden. Zo makkelijk is het echter niet altijd. 'Verder zie je wel de verbanden die het algoritme legt, maar dat verklaart niet waarom die verbanden er zijn.' Dat zorgt ervoor dat AI nooit te gebruiken zal zijn om PTSS echt te diagnosticeren; daarvoor blijft een psycholoog onmisbaar. Maar het doorgronden van de resultaten van het algoritme helpt de psycholoog wel om beter te begrijpen wat er met de patiënt aan de hand is, nog voor hij of zij de praktijk binnen stapt. 'Uitlegbaarheid helpt mensen die geen informaticus zijn om toch met een kunstmatig intelligent algoritme om te gaan.'

Daarnaast helpt het om draagvlak te creëren voor AI. Critici waarschuwen vaak voor onbe-

Onderzoekers van Kyoto University vroegen mensen om te denken aan bepaalde objecten. In een fMRI-machine werd hun hersenactiviteit gemeten en aan de hand van die activiteit probeerde een algoritme een tekening te maken. De kunstwerken die dat opleverde, zijn nu te bekijken in de Serpentine Gallery in Londen.



wuste vooroordelen in AI, die verergeren als de oorzaak van die vooroordelen niet is te vinden. Een algoritme dat misdaad opspoort en vaker misdrijven ziet in wijken met allochtone inwoners kan bijvoorbeeld een vooroordeel over allochtonen ontwikkelen. 'Je merkt dat het onderwerp *explainability*, zoals het in het Engels heet, nu op congressen gaat leven', constateert Veldkamp. 'Laatst kwam het zelfs in China ter sprake, een land waar men verder minder doet met AI-kritiek.'

IBM zegt bijvoorbeeld dat hun Watson-AI transparant en uitlegbaar is. Elke beslissing die het programma neemt, elk advies dat het geeft, krijgt aan het einde twee percentages: hoe betrouwbaar is dit resultaat, en in hoeverre is het bevooroordeeld? 'Elke stap die het algoritme neemt, wordt automatisch vastgelegd', vertelt Jay Bellissimo, General Manager van IBM's Cognitive Process Transformation. 'Vervolgens worden die logs geanalyseerd en krijgen gebruikers te zien hoe *fair* en



accuraat het algoritme is. Zo zie je hoe de processen verlopen en kun je dus laten zien aan je klanten dat het algoritme zijn werk eerlijk doet.'

Het is alleen de vraag of die percentages alles verklaren. Ze zeggen nog steeds niet veel over het waarom; alleen óf er vooringenomenheid of onbetrouwbaarheid in het algoritme of de data zit. 'Het is een uitdaging om deep learning zichzelf uit te laten leggen', zegt Veldkamp. 'Je moet dat eigenlijk doen bij het prilleste begin van een zelflerend algoritme. Tijdens het programmeren moet je al een systeem inbouwen dat het algoritme vraagt om zichzelf te verklaren, bij elke stap die het neemt.'

Door de theorie van AI te onderzoeken, begrijpen we algoritmes beter. Daardoor zijn ze op meer plekken toe te passen en makkelijker uit te leggen.

Prof.dr. Peter Grunwald, onderzoeker van veiligheid van statistiek en machine learning bij het Centrum voor Wiskunde en Informatica, kent de wiskunde achter de beslissingsbomen die AI gebruikt. En die laat zien dat je geen programma kunt maken dat gegarandeerd geen vooroordelen aanneemt. 'Er zijn verschillende wiskundige definities mogelijk van vooroordelen, maar die blijken onderling tegenstrijdig te zijn. En geen enkele definitie is alomvattend. Daardoor is het niet mogelijk om een aantoonbaar onbevooroordeeld algoritme te bouwen.'

Gelukkig is het niet hopeloos. 'Als je vooroordelen afdwingt in een algoritme, krijg je bijvoorbeeld meteen inzicht in de mechanismen

Een variatie op het schilderij *De Sterrennacht* van Vincent van Gogh. Het algoritme maakt vreemd uitzijende dieren van torens en sterren, en doet iets nog vreemders aan de onderkant van het beeld.

achter een vooroordeel. Daarmee kun je het algoritme verbeteren en minder *biased* maken.' Het nadeel is dat dat een beetje een ad-hocmethode is, waar je een bestaand algoritme beter probeert te maken.

Dat is de trend van AI op dit moment: het zijn vaak slordige algoritmes die op een heleboel verschillende datasets worden toegepast. Als ze niet goed werken, passen programmeurs ze aan tot ze wel werken. 'Maar je hebt geen idee hoe ver je met deze aanpak kunt komen', zegt Grunwald. 'Wat is het best haalbare en wat zijn de zwakke punten van je AI? Als je dat wilt weten, is er een fundamentele basis nodig.'

En dat is de tweede manier waarop deep-learning-algoritmes zijn te verbeteren. Wiskundigen als Grunwald proberen een theoretisch fundament te leggen onder de algoritmes die praktisch goed werken. Want er zijn nu veel vragen: waarom werkt deep learning bijvoorbeeld zo goed? Zou het nog beter werken met nog meer data, of is er een grens? En zijn er andere vormen van kunstmatige intelligentie die beter werken met veel data? 'Het onderzoek naar die vragen gaat AI uiteindelijk beter maken', zegt hoogleraar machine learning Welling. 'Als je weet wat de regels zijn, kun je ze daarna immers gebruiken om programma's te schrijven waarvan je weet dat ze het doen. De ontwikkeling gaat dan sneller.'

Grunwald zelf weet niet zeker of en hoe theoretisch onderzoek naar de werking van AI de praktijk gaat verbeteren, maar 'eerder ging het wel zo. In de jaren tachtig werden eerst kleine successen bereikt met experimentele methodes. Op een gegeven moment was er een plateau bereikt en toen tilde theorie AI in de jaren negentig naar een hoger plan. Of dat nu weer zo zal gaan, is moeilijk te zeggen, omdat de ontwikkelingen zo snel gaan. Maar er zullen hoe dan ook weer plateaus worden bereikt, en dan gaat het niet verder zonder echte doorbraken met radicaal nieuwe ideeën. Wellicht gaan die uit de theorie komen.'



Deep learning is een vrij lompe vorm van intelligentie. Wil je dichter bij menselijke slimheid komen, dan moet AI causaliteit leren begrijpen.

Deep learning werkt nu door een programma duizenden keren min of meer hetzelfde te laten zien; uiteindelijk kan het daar dan dingen uit concluderen. 'Maar redeneren, dat kan zo'n algoritme niet', zegt prof.dr. Tom Heskes, hoogleraar artificial intelligence aan de Radboud Universiteit Nijmegen. En Heskes kan het weten. Zelf werkt hij namelijk aan programma's die wél verbanden kunnen zien, kunnen 'redeneren'. In het bijzonder maakt hij zelflerende algoritmes die causaliteit begrijpen. Dat is een van de pijlers van ons menselijk kunnen: als wij een bal zien rollen, weten we meteen welke kant hij op zal rollen, of hij gaat botsen met iets, en soms zelfs waar hij vandaan komt. Als een computer een filmpje van een rollende bal ziet, heeft hij geen idee wat



er aan vooraf ging of erna zal gebeuren. 'Als je enigszins intelligente computers wilt hebben, moet je ze dat leren.'

Heskes probeert dat op verschillende manieren. Soms is het makkelijk: als tijd een factor is, moet data in een bepaalde volgorde worden 'gelezen'. Dat helpt al enorm bij het redeneren. Maar andere verbanden zijn minder evident. 'Geslacht kan bijvoorbeeld wel een medeoorzaak zijn van een bepaalde aandoening, maar door een aandoening zul je nooit spontaan van geslacht veranderen. Dat klinkt logisch, maar de vraag is: hoe leer je een computer dat?'

Voorlopig gaat dat voor een deel met de hand. Een programmeur geeft de 'richting' van een redenering, waarna de computer die begin-informatie gebruikt om andere causale verbanden te ontdekken. Met die techniek ontdekte een AI bijvoorbeeld een causaal verband tussen

aandachtsproblemen en hyperactiviteit. 'We zijn nog heel voorzichtig met harde uitspraken, maar suggereren nu dat artsen onderzoek doen om te kijken of dat verband inderdaad bestaat. Als dat zo is, weten we iets zekerder dat een algoritme inderdaad causaliteit kent.'

Een van de voordelen van causaliteitsprogramma's is dat ze inherent uitlegbaar zijn. Heskes: 'Je kent de input en de output, en je bent expliciet op zoek naar de verbanden die de AI legt. Je resultaat is dus meteen je uitleg van de beslissing van een AI. Het is daarmee transparanter dan deep learning.'

Nadeel is dat je behoorlijk specifieke data moet hebben om er causaliteit in te kunnen vinden. 'Als er tijd in de data zit, helpt dat bijvoorbeeld. Maar in het algemeen zijn datasets die geschikt zijn voor causaliteitsalgoritmes kleiner dan de hopen data die nu voor deep learning worden gebruikt.' Als causale algoritmes populairder worden, zal er ongetwijfeld ook meer 'geschikte' data komen. Maar het feit dat de data beter moet zijn geordend en het liefst een ingebakken lijn heeft, zoals tijd of een andere causale eigenschap, beperkt wel de hoeveelheid geschikte data.

Het causale algoritme staat dus in de kinderschoenen. 'Het is nu nog een niche. De correlatietechniek van deep learning heeft voorlopig veel meer toepassingen. Maar als je ingewikkeldere zaken aan een computer wilt overlaten, is een gevoel voor causaliteit onmisbaar.'

Algoritmes kunnen sneller leren als ze beschikken over een flinke hoeveelheid basiskennis. Een decenniaoude techniek kan daarbij helpen.

Straks kunnen we alle kennis ter wereld in een zelflerend algoritme stoppen. Dat scheelt, want het betekent dat een algoritme een heleboel dingen niet meer hoeft te leren. Dat een kat vier poten en een staart heeft, hoeft het bijvoorbeeld niet meer te ontdekken door naar 50.000 plaatjes te kijken; dat weet het programma dan gewoon.

Dat is mogelijk dankzij kennisgrafen, een technologie die al volop wordt toegepast. Als je bijvoorbeeld een bedrijf, persoon of locatie googelt, krijg je rechts in beeld allerlei handige informatie: gerelateerde zaken, een link naar Google Maps, korte biografische gegevens ... Die kennis wordt

automatisch verzameld met behulp van een kennisgraaf. 'Het zijn eigenlijk een soort databases, maar dan met veel meer verbindingen tussen de verschillende onderwerpen', zegt hoogleraar knowledge representation Van Harmelen. Die grafen maak je bijvoorbeeld door te kijken welke onderwerpen op Wikipedia linkjes hebben naar andere onderwerpen. Zo ontstaat een woud van verbindingen waar informatie uit te halen is.

Lange tijd zijn grafen het ondergeschoven kindje van de AI-familie geweest. Kennisrepresentatie, het vakgebied waar ze onder vallen, bestaat al decennia, maar nu maakt schaalvergroting grafen eindelijk nuttig. 'Dankzij de groeiende web-technologie is het mogelijk geworden om ze enorm groot te maken. Miljarden onderwerpen met tientallen miljarden verbindingen zijn nu normaal geworden', vertelt Van Harmelen.

Die ontwikkeling kan helpen om zelflerende algoritmes beter te maken en sneller te laten leren. 'Met grafen geef je basiskennis mee aan AI. Dan hoeven ze die niet telkens opnieuw te leren en kunnen ze dus sneller nieuwe dingen ontdekken.' UvA-hoogleraar Welling ziet ook dat deze combinatie van kennis met leren groot kan worden. 'Wij mensen kunnen redeneren doordat we een heleboel basiskennis hebben. Over de natuurwetten, over context. Hoe krijg je die kennis in een computer? Kennisgrafen bieden daar een oplossing.'

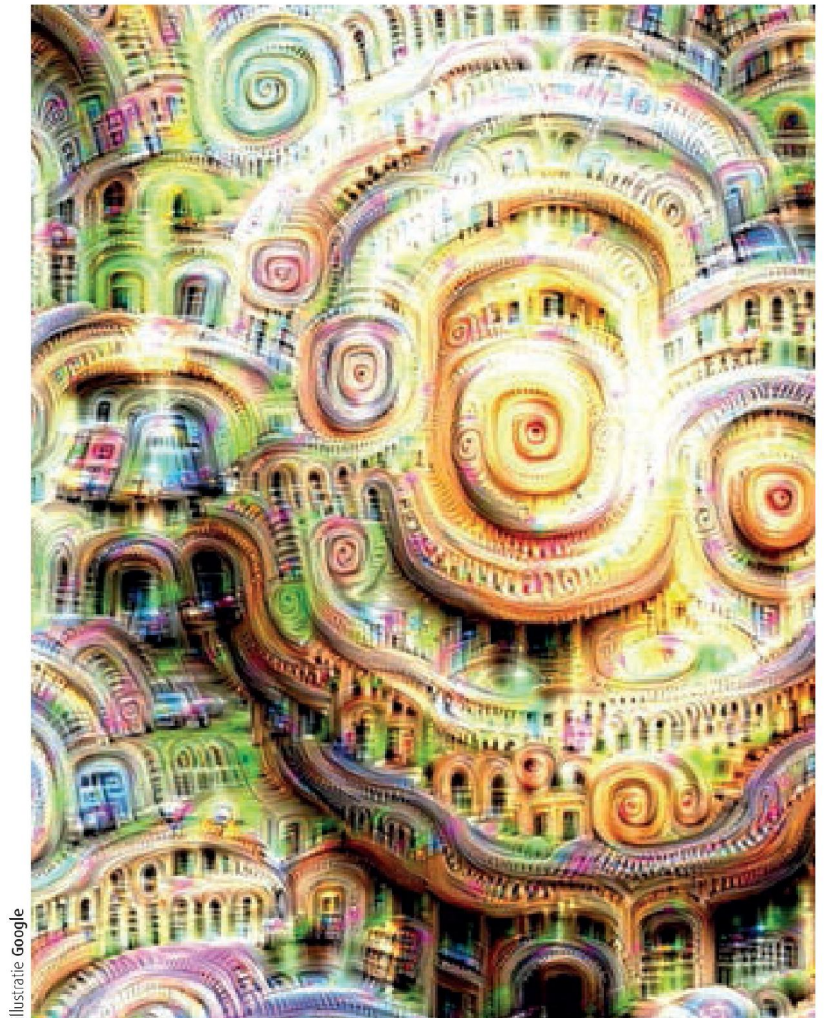
Daarmee staan Welling en Van Harmelen tegenover een andere groep die gelooft dat dat soort basiskennis vanzelf komt bovendrijven als je een algoritme maar lang genoeg traint met ruwe data. Dr. Tim Salimans, onderzoeker bij Google Brain, past in dat kamp. 'Als ik een tik tegen een waterflesje geef, weet ik wat er gaat gebeuren. Een computer kan dat ook leren als je hem maar genoeg beelden laat zien van flesjes die omvallen.' Deze *brute-force*-methode is voorlopig nog het populairst in de zakenwereld. Slimmer leren wordt daar vooral gedaan door de algoritmes te verbeteren, niet door ze een flinke smak basiskennis mee te geven.

Welling is echter resoluut. 'Voor hogere intelligentie werkt die data-aanpak niet. Er zijn veel te veel onbekende situaties. Dat zie je ook bij de zelfrijdende auto's; het is heel moeilijk om de software elk mogelijke verkeerssituatie te leren. Dan kan het helpen om voorkennis over ongelukken in te bouwen.'

Bedrijven zijn nog niet echt bezig met geavanceerdere vormen van AI. Zij gebruiken 'ouderwetse' deep learning en proberen om de nadelen en problemen heen te werken

Voorlopig lijken de verbeteringen die met kennisgrafen en causaliteits-algoritmes mogelijk worden nog niet essentieel voor het bedrijfsleven. 'Deep learning is onze grootste focus', vertelt Salimans van Google Brain. 'We werken met een enorm team en kijken ook wel naar andere dingen, maar bij Google Brain proberen we toch vooral die vorm van AI te verbeteren.' Het veiliger maken van algoritmes door ze expres voorbeelden te geven waar ze niet mee om kunnen gaan, bijvoorbeeld. 'Als je de zwakte van een algoritme kent, kun je het vervolgens verbeteren.' En ook Google werkt aan het uitlegbaar maken van haar algoritmes.

'Deep learning zal blijven als fundament, ook in de toekomst. En daar kan een heleboel mee. Maar kennisgrafen kunnen zeker interessant zijn; daar denken wij ook over na.' Voor hem ligt de toekomst echter eerder bij het uitbreiden van neurale netwerken. 'Wat als deep learning niet alleen data omzet in een resultaat, maar zelf kan kiezen op welke manier dat gebeurt? Dat een algoritme leert welk model het



illustratie Google



Hoe meer lagen het DeepDream-algoritme van DeepMind heeft, hoe abstracter de beelden. Uit een foto waar vooral ruis in zit, kan het allemaal patronen halen die er eigenlijk niet zijn, een fout die het menselijk brein ook nog weleens wil maken.

De mens blijft de leidende inspiratiebron bij de langetermijndoelen op het gebied van kunstmatige intelligentie.

De huidige neurale netwerken functioneren op een heel andere manier dan onze hersenen. Willen we echt grote stappen maken met AI, dan zullen we moeten terugkeren naar het menselijk brein als inspiratiebron, denkt Van Gerven van de Radboud Universiteit. 'Ons brein is nog vele malen efficiënter dan een computer. Daar kunnen we dus nog een grote slag maken. De neurale netwerken, die nu zoveel succes hebben, komen immers ook voort uit de psychologie en neurowetenschap. Hoe meer we over menselijke cognitie leren, des te meer inspiratie we hebben voor kunstmatige breinen. Je moet de kennis over cognitie alleen nog wel vertalen naar precieze wiskundige termen om hem bruikbaar te maken.'

best bij welke opdracht past. Of, nog verder: dat je deep learning leert om zichzelf te programmeren.' Dat laatste is echter nog een flinke stap, geeft Salimans toe.

Op korte termijn is het verbeteren van de neurale netwerken het belangrijkste. 'Dat kan niet alleen door ze meer data te voeren, maar ook door ze beter te programmeren. Dan kun je uiteindelijk ook met minder data toe, iets wat steeds belangrijker zal worden naarmate AI op meer plekken wordt toegepast.'

Ook IBM werkt hard aan de verbetering van deep learning. Hun AI-product Watson was in 2011 alleen in staat om het quizprogramma *Jeopardy!* te winnen. Nu kunnen klanten Watson trainen voor allerlei verschillende toepassingen. 'KPMG gebruikt hem bijvoorbeeld voor het belastingwerk, een Franse bank om klanten sneller te helpen', vertelt Jay Bellissimo van IBM. Het geeft aan dat de doorbraak van AI nu al bezig is; niet alleen bij techbedrijven, maar ook in de financiële wereld. 'Natuurlijk zijn er conservatieve sectoren die de boot afhouden. Vanwege de onbetrouwbaarheid van AI, het gebrek aan data of gebrek aan kennis. Maar uiteindelijk gaat dit overal invloed hebben. Elk bedrijf zou daarom een AI-strategie moeten hebben.'

Onze eigen hersenpan kan dus zeker als inspiratie dienen voor AI, maar kan AI daarmee ook even intelligent worden als wij, of zelfs nog intelligenter? Aan de ene kant geven zelfs de wetenschappers die denken dat dit in principe mogelijk moet zijn toe dat we daar nog wel even op zullen moeten wachten. Aan de andere kant: de afgelopen zes jaar ging de ontwikkeling veel sneller dan deskundigen daarvoor voor mogelijk hielden. En dankzij het onderzoek naar AI die kan redeneren en basiskennis kan opdoen, lijkt een aanmerkelijk slimmere kunstmatige intelligentie steeds dichterbij te komen.

Zelfs als het evenaren van de menselijke intelligentie toch geen haalbare kaart blijkt, hebben de nieuwe technieken nut voor de huidige vormen van AI. Hopelijk maken die ons leven daarmee een stuk makkelijker, zodat wij mensen ons brein kunnen gebruiken voor de zaken waar de computer zich voorlopig geen raad mee weet. |