

Inhoudsopgave

Voorwoord	9
Inleiding	11

DEEL I Kennismaking met artificial intelligence

1	Wat is Artificial Intelligence?	17
1.1	Nabootsen van intelligentie	17
1.2	De geschiedenis van AI	19
1.3	De bouwblokken van AI	27
1.4	De relatie tussen AI en andere technologieën	34
2	Artificial Intelligence in een stroomversnelling	41
2.1	Exponentiële groei van computerkracht	41
2.2	Data, data en nog meer data	43
2.3	Systemen trainen zichzelf	46
2.4	Democratisering van AI	57
3	Artificial Intelligence en de maatschappij	61
3.1	Ethiek: is er een wereldwijde ethische standaard mogelijk?	61
3.2	Privacy: wie weet wat allemaal van jou?	64
3.3	Bias: niet alles is wat je denkt dat het is	67
3.4	Transparantie: uitkomst verklaren	71
3.5	Wet- en regelgeving: AI juridisch verankeren	74
3.6	Veiligheid: hoe beschermen we onszelf tegen AI?	77
3.7	Arbeid: de invloed van AI op ons werk	79
3.8	Educatie: vergaar nieuwe kennis en ontwikkel nieuwe omgangsvormen	82

DEEL II Artificial intelligence toepassen

4	De huidige status van AI	89
4.1	De rol van AI in de wereld	89
4.2	De rol van AI in bedrijven	92
4.3	AI in verschillende domeinen en sectoren	93
5	De chatbot	107
5.1	Wat is een chatbot?	107
5.2	Hoe werkt een chatbot?	112
5.3	Hoe ervaren mensen een chatbot?	122
5.4	Wanneer kies je voor een chatbot?	123
6	AI in je eigen organisatie	129
6.1	Laat je inspireren	129
6.2	Voor welke problemen is AI de oplossing?	132
6.3	Aan de slag!	137
7	Zes succesfactoren voor AI	147
7.1	Wees kritisch of AI wel de beste oplossing is en toegevoegde waarde biedt	147
7.2	Begin klein en deel je resultaten!	149
7.3	Leider met kennis, een vooruitstrevende visie en lef	150
7.4	Uitvoerend team met kennis en vaardigheden die aansluiten bij de AI-projecten	153
7.5	Samenwerking tussen organisatieonderdelen	155
7.6	Medewerkers wier werk door AI verandert erbij betrekken	158
8	Overzicht van leveranciers en tools	163
8.1	AI in verschillende smaken	163
8.2	Kopen of zelf bouwen?	164
8.3	Leveranciers	165
	Over de auteurs	178
	Dankwoord	180
	Geraadpleegde literatuur	182

Voorwoord

Vraag jij je ook weleens af waarom ze zeggen: '(Big) data zijn de nieuwe olie'? Want zelf heb je misschien wel veel bedrijfsdata beschikbaar, maar om daar nou echt waardevolle informatie uit te halen, dat lukt nog niet echt. Of wat dacht je van 'Artificial Intelligence is de nieuwe elektriciteit'? Wellicht heb je überhaupt nog geen goed beeld bij wat Artificial Intelligence (AI) nu precies is, laat staan bij wat je er in je eigen organisatie mee kunt.

Maar toch is je interesse gewekt en wil je eigenlijk wel proberen uit te vinden hoe data en Artificial Intelligence jou wellicht een concurrentievoordeel zouden kunnen geven. Want je hebt namelijk wel het ongemakkelijke gevoel dat als je niets doet, je straks door je concurrenten wordt ingehaald. Of wat te denken van die langdurige problemen in jouw organisatie, zou je die wellicht met data en AI kunnen oplossen? Of zou AI voor jou zelfs meer omzet kunnen genereren? Allemaal vragen waarop je graag een antwoord wilt hebben.

Een aantal jaren geleden hadden wij dezelfde vragen en nieuwsgierigheid als jij.

Taco Hiddinks interesse werd in 2012 gewekt, toen hij werkte aan een model dat potentiële klanten probeerde te vinden in grote hoeveelheden open data. Dat resulteerde in een grote order voor een klant in een sector met kapitaalintensieve producten. Terwijl het model nog maar voor een klein deel bestond uit echte AI, was de waarde van de combinatie big data en AI al zichtbaar. Daarom vroeg Taco zich af: wat kan dan wel niet de waarde van AI worden als die steeds krachtiger wordt en we AI steeds meer gaan toepassen? Sinds dat moment is hij zich gaan richten op het bedenken en bouwen van slimme oplossingen met AI voor bedrijven.

Muriël Serrurier Scheppers nieuwsgierigheid werd in 2006 gewekt, toen ze voor het eerst de kracht van data en AI zag. Ze ervoer dat een fraudemodel van FICO in staat was om realtime frauduleuze transacties met een creditcard te detecteren en te stoppen. Hier moest ze meer van weten, dus een baan bij FICO werd haar volgende stap. Ze voelde zich als een vis in het water tussen al

die promovendi en datascientists. Want ook al sprak zij initieel niet hun 'taal', door haar nieuwsgierigheid en niet-aflatende stroom vragen werd ze steeds meer in staat gesteld om te begrijpen wat er gebeurde en het te vertalen naar de dagelijkse praktijk.

De fascinatie voor data en Artificial Intelligence is bij ons allebei gebleven en in onze dagelijkse praktijk zien we dat veel bedrijven moeite hebben om de juiste casus te vinden waarmee ze kunnen starten. En als ze al starten, dan is het meestal meer om technieken uit te proberen, dan om AI werkelijk waarde te laten leveren voor een prangend probleem.

Ook zien we dat projecten mislukken doordat bedrijven overweldigd worden door de kennis en expertise van datascientists, die onvoldoende in staat zijn om in gewonemensentaal uit te leggen wat ze doen. Er is dus vaak ruis op de lijn. Door het gebrek aan kennis bij managers en medewerkers hebben zij ook vaak onrealistische verwachtingen en zijn de eerste kiemen voor teleurstelling over de resultaten snel gelegd.

Wij vinden het zonde en onnodig dat bedrijven op dit moment nog zo weinig waarde uit data halen en zo weinig AI-toepassingen inzetten. Samen zijn we dan ook begonnen met het geven van meerdaagse trainingen en korte presentaties onder de vlag van *ai-training.nl*. Om onze kennis nog verder te delen en een groter publiek te bereiken, hebben wij dit boek geschreven. Hiermee reiken we jou de basiskennis aan waarmee je zowel intern als extern een dialoog over AI kunt voeren. Ook bieden we concrete handvatten om zelf aan de slag te gaan. En om jou een zo succesvol mogelijk start te geven, hebben we een overzicht van de belangrijkste succesfactoren gemaakt. Niets staat je dus meer in de weg om zelf te ervaren of data de nieuwe olie voor je zijn en of Artificial Intelligence de nieuwe elektriciteit is.

Wij wensen je veel leesplezier en veel succes met jouw start in de wondere wereld van data en Artificial Intelligence!

Muriël Serrurier Schepper en Taco Hiddink

Inleiding

Je leest of hoort in de media vaak verhalen over kunstmatige intelligentie, waarin niet zelden een waarschuwende ondertoon doorklinkt. De teneur is dan dat Artificial Intelligence (AI) wel interessante mogelijkheden biedt, maar tegelijk zorgt voor banenverlies, of erger nog: de hele mensheid overneemt! Vaak ook lijkt het alsof AI alleen iets voor nerds is, en niet voor praktisch, zakelijk gebruik. Wat is er eigenlijk waar van dit alles? Wat zijn de kansen die AI het bedrijfsleven werkelijk biedt, en met welke bedreigingen moet je daarbij rekening houden? De antwoorden op deze en veel andere vragen vind je in dit boek.

Voor wie is dit boek?

AI is terug te vinden in elke sector en in elk organisatiedomein, of dit nu op hr-gebied is, juridisch, marketing, verkoop enzovoort. AI is dus overal. *Artificial Intelligence In Actie* brengt je als businessprofessional op de hoogte van de mogelijkheden en onmogelijkheden van AI op jouw vakgebied en in jouw sector. Je ontdekt met welke AI-innovaties je in je eigen organisatie aan de slag kunt gaan en hoe je dat praktisch voor elkaar krijgt. Moet je dan zelf het wiel uitvinden? Gelukkig niet, want je kunt je technisch laten ondersteunen, maar ook dan moet je de (on)mogelijkheden kennen en weten welke kritische vragen je moet stellen. Je vindt in dit boek wat er allemaal bij komt kijken om een AI-toepassing te trainen en welke digitale vaardigheden daarvoor nodig zijn. Met de checklists in het boek en aanvullende informatie op de website *ai-in-actie.nl* krijg je daar een goede indruk van. Je profiteert bovendien van de lessen die wij in de praktijk hebben geleerd en doet inspiratie op uit casussen uit het Nederlandse bedrijfsleven.

Hiermee krijg je handvatten aangereikt om zelf vernieuwende ideeën voor jouw organisatie te ontwikkelen en deze succesvol te implementeren, en daarbij zo nodig de juiste partner te selecteren. Met dit boek kun je dus met artificial intelligence in actie komen.

Opbouw van dit boek

Artificial Intelligence In Actie bestaat uit twee delen: een deel waarin we de essentie van AI schetsen en een deel waarin we beschrijven hoe je zelf met AI aan de slag kunt gaan. In het eerste deel lees je in hoofdstuk 1 en 2 wat AI is, hoe het is ontstaan en wat de actuele ontwikkelingen op dit gebied zijn. Ook leggen we een relatie met andere technieken, zoals een systeem om routine-matige handelingen over te nemen (Robotic Process Automation), het Internet of Things en blockchain. Verder lees je hoe het komt dat AI op dit moment zo'n grote ontwikkeling doormaakt, waarbij we begrippen als *algoritme*, *(on)gecontroleerd leren* en *neuraal netwerk* toelichten, en we ingaan op het belang van data. Maar laat je hier niet door afschrikken, want dat doen we op een heldere, niet-technische manier, geïllustreerd met voorbeelden.

Tot slot van het eerste deel komen in hoofdstuk 3 maatschappelijke aspecten van AI aan bod. Je leest daarin onder meer over de ethische kant van AI en hoe je veiligheid en transparantie waarborgt. En ook gaan we in dit hoofdstuk uitgebreid in op wet- en regelgeving rond AI en wat het verwachte effect van AI is op de werkgelegenheid.

Het tweede deel is geschreven om jou te inspireren zelf ideeën te ontwikkelen, zodat je AI in je eigen organisatie kunt toepassen. Om te beginnen vind je daarvoor in hoofdstuk 4 een overzicht van AI-toepassingen die bedrijven op dit moment al inzetten. Sommige voorbeelden daarvan komen je mogelijk bekend voor, maar er zullen ook veel voor jou nieuwe ideeën de revue passeren. We beschrijven daarnaast onder meer AI-toepassingen van een aantal Nederlandse en Vlaamse bedrijven.

In hoofdstuk 5 gaan we dieper in op een van die toepassingen: de chatbot, een applicatie om te communiceren met een machine. Dat doen we omdat bedrijven de laatste jaren in toenemende mate chatbots inzetten voor hun medewerkers en klanten. Je komt te weten hoe deze technologie precies werkt, en met casussen van ABN AMRO en de Rabobank en onze praktische tips krijg je een goed beeld van de mogelijkheden.

In hoofdstuk 6 nemen we je stap voor stap mee in hoe je zelf met AI aan de slag kunt gaan. Zo vind je een canvas met vragen die je moet beantwoorden voordat je bepaalde keuzes maakt. Maar het hoofdstuk begint met bruikbare ideeën en vragen waarmee je kunt brainstormen, gericht op organisaties,

medewerkers en klanten. Met een model helpen we je vervolgens om de stappen te zetten die nodig zijn om van idee naar implementatie te komen.

Hoofdstuk 7 borduurt hierop voort: daarin lees je hoe je een succes kunt maken van je AI-implementatie. Gebaseerd op onze jarenlange ervaring, hebben we zes factoren ontdekt die daarvoor nodig zijn, namelijk:

- 1** Beoordeel kritisch of AI wel de beste oplossing is en toegevoegde waarde biedt
- 2** Begin klein en deel je succes
- 3** Leider met kennis, een vooruitstrevende visie en lef
- 4** Uitvoerend team met kennis en vaardigheden die aansluiten bij de AI-projecten
- 5** Samenwerking tussen organisatieonderdelen
- 6** Medewerkers wier werk door AI verandert erbij betrekken

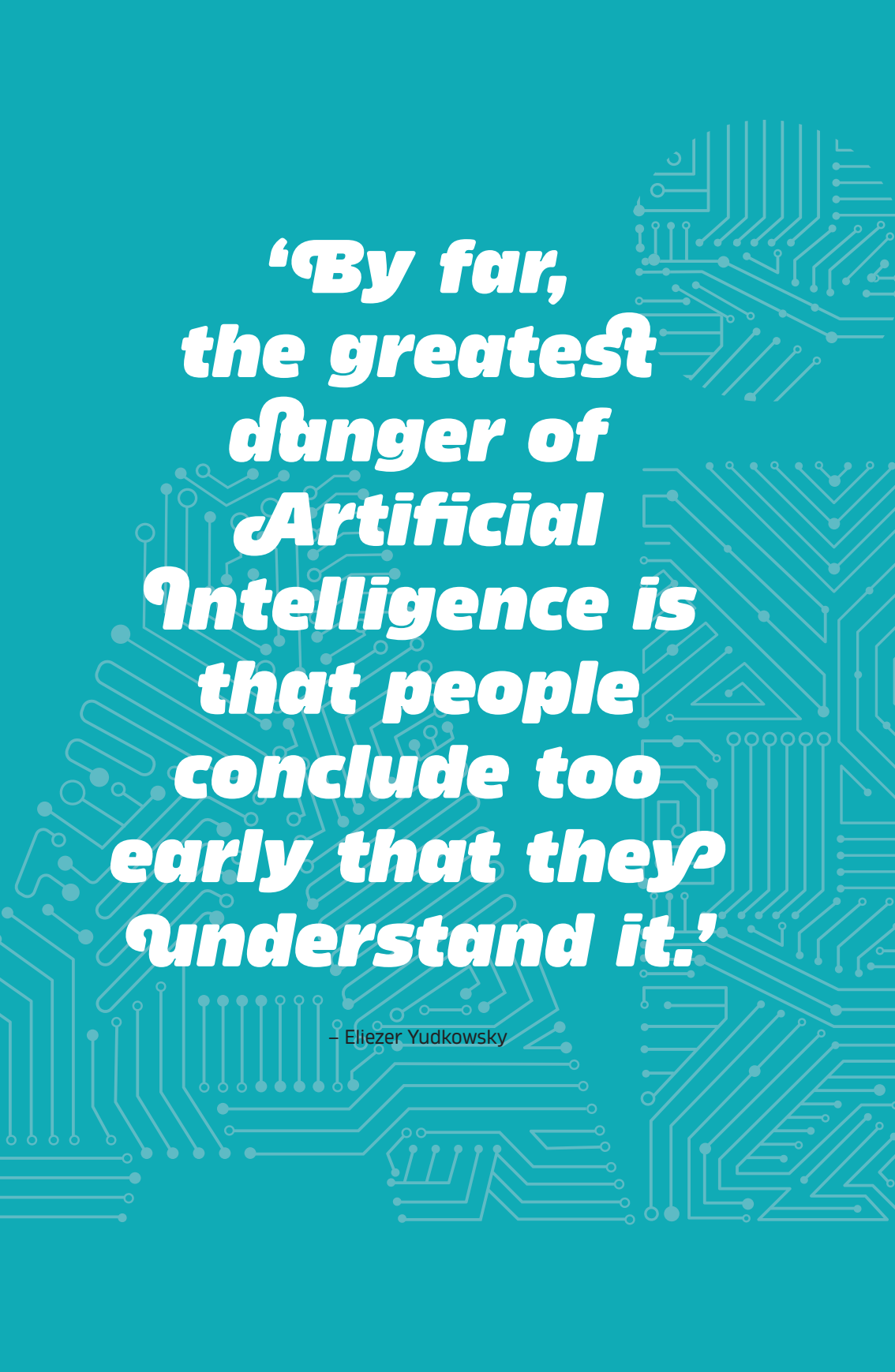
We beschrijven deze succesfactoren uitgebreid, toegelicht met voorbeelden uit de dagelijkse praktijk.

Tot besluit van deel twee en van het boek vind je in hoofdstuk 8 een overzicht van verschillende AI-oplossingen en leveranciers die je daarvoor zou kunnen benaderen. In dat hoofdstuk schetsen we ook de voor- en nadelen van zelfbouw versus kant-en-klare oplossingen.

We wensen je veel plezier en inspiratie bij het lezen! En neem gerust een kijkje op onze website ai-in-actie.nl voor updates en om de verschillende checklists te downloaden.

DEEL I

Kennismaking met artificial intelligence



***‘By far,
the greatest
danger of
Artificial
Intelligence is
that people
conclude too
early that they
understand it.’***

– Eliezer Yudkowsky

1 Wat is Artificial Intelligence?

Over kunstmatige of artificiële intelligentie (AI) doen verschillende definities de ronde. In dit eerste hoofdstuk omschrijven we daarom om te beginnen wat wij in dit boek onder AI verstaan. Daarna passeren in paragraaf 1.2 de mijlpalen in de geschiedenis van AI de revue; daaruit komt tegelijk de enorme vooruitgang naar voren die dit vakgebied in de laatste decennia heeft geboekt. Vervolgens lichten we in paragraaf 1.3 de bouwblokken van AI toe, waarna we tot slot in de laatste paragraaf van dit hoofdstuk ingaan op de relatie tussen AI en andere (opkomende) technologieën. Met dit alles leggen we het fundament onder de volgende hoofdstukken.

1.1 Nabootsen van intelligentie

Aan de vraag wat *kunstmatige intelligentie* is, gaat een andere, misschien nog wel fundamentele vraag vooraf: wat is *intelligentie*? In het algemeen taalgebruik betekent *intelligentie* eenvoudig 'verstand, verstandelijk vermogen'. We hebben het woord geleend uit het Latijn, van *intelligentia*, dat komt van het werkwoord *intelligere* 'begrijpen, inzien, gewaarworden'. Maar in datzelfde algemene taalgebruik wordt *intelligentie* ook vaak verbonden met kennis en opleiding: hoe meer je weet of hoe hoger je opleiding is, hoe intelligenter, pienterder en slimmer je bent, hoor je daarom ook wel. Meten is weten, en intelligentie kun je meten met een IQ-test, althans, zo'n test kan een indicatie geven over iemands verstandelijke vermogen, afhankelijk van de kwaliteit van het onderwijs dat diegene heeft genoten en allerlei andere omgevingsfactoren. Dat brengt ons terug naar de kern van waar we met *intelligentie* waren gebleven, namelijk: verstandelijk vermogen, hoe kun je dat nu het best omschrijven? *Verstandelijk vermogen* duidt op geestelijke begaafdheid, en dat is ook hoe de psychologie ernaar kijkt: een mentale eigenschap om bijvoorbeeld verschillen te zien, je te kunnen oriënteren, te kunnen redeneren, plannen te maken,

problemen op te lossen, abstract te kunnen denken enzovoort. Intelligentie is nodig om kennis te verwerven en die te kunnen gebruiken.

Dit alles geeft tegelijk een heel goed beeld van wie er intelligent is: intelligentie is een vermogen van levende wezens, waarmee de stap naar kunstmatig, dus niet-levend, gauw gezet is. *Kunstmatige intelligentie* bootst menselijke intelligentie na. Dat gebeurt met een computersysteem dat taken kan uitvoeren waarvoor je normaal gesproken menselijke intelligentie nodig hebt, bijvoorbeeld voor spraak, herkennen van beelden, beslissingen nemen en omgaan met geschreven taal. Deze taken worden aangeleerd doordat de computer leert van eerdere ervaringen. Je moet een AI-oplossing dan ook trainen in plaats van programmeren, en door dat trainen wordt de AI-oplossing ook steeds beter in wat deze moet doen. Hoe dat trainen gaat, daar komen we later op terug. Dit trainen is iets wezenlijks anders dan programmeren, waarbij een computer instructies uitvoert zoals die bedacht zijn door een programmeur. Zo is rekenen geen AI. Bij rekenen voert een computer heldere instructies uit, en ongeacht hoe vaak je sommen uitvoert, het systeem zal niet beter worden en ook geen nieuwe manieren aanleren om de sommen te berekenen.

Tot op heden is er nog geen eenduidige definitie van artificial intelligence. Het is een containerbegrip en het is een interdisciplinair vakgebied, waarin wordt samengewerkt met onder anderen informatici, psychologen, neurologen, statistici en linguïsten. Wij beschrijven zelf AI het liefst als een gereedschapskist met allerlei verschillende werktuigen die in meer of mindere mate een vorm van menselijke intelligentie nabootsen, bijvoorbeeld om met taal om te gaan, beelden te interpreteren of voorspellingen te doen. In deze metafoor verbeelden de timmerman, loodgieter of doe-het-zelver algoritmes – en over wat algoritmes zijn, lees je alles in hoofdstuk 2.

In dit boek gebruiken we voor *kunstmatige intelligentie* overigens de Engelse term *Artificial Intelligence* of de afkorting daarvan, AI.

**AI is een gereedschapskist met werktuigen
die in meer of mindere mate een vorm
van menselijke intelligentie nabootsen.**



1.2 De geschiedenis van AI

Over artificial intelligence wordt al millennia lang gespeculeerd en nagedacht. In Griekse mythen komen al robots voor, de Franse filosoof en wiskundige René Descartes schreef er over in zijn *Discours de la méthode* (1637) en in talloze sciencefictionverhalen en -films zijn machines soms niet te onderscheiden van mensen, in ieder geval niet in hun manier van denken en communiceren. Maar AI in de werkelijkheid gebruiken, daarvoor moesten machines eerst geavanceerder worden. En kennelijk zijn ze dat nu: vaak zonder dat je je daarvan bewust bent, maak je inmiddels gebruik van AI. Kinderen groeien er zelfs mee op. Een klein kind weet niet beter dan dat het niet iets hoeft in te toetsen op de iPad, want door tegen het apparaat te praten, kan het navigeren naar wat het zoekt. En wat te denken van virtuele assistenten in huis, zoals Google Home, een luidspreker waaraan je vragen kunt stellen? Of Snapchat, dat misschien wel simpel lijkt, maar dat gebruikmaakt van AI-technologie om animaties op de juiste plek op een gezicht te plaatsen. AI neemt ook al huishoudelijke taken over, zoals met de Roomba-stofzuiger, een robot die al in miljoenen huishoudens stofzuigt. Hij begint met schoonmaken zodra jij van huis gaat, waarbij hij leert van de omgeving om je huiskamer steeds efficiënter te reinigen. En ook als je ingaat op de suggesties die Netflix doet op basis van je kijkgedrag laat je je leiden door AI. AI is overal. Maar voordat het zover was, was er een lange weg te gaan. Daarover lees je in deze paragraaf.

Het begin van AI en van de AI-winter

Kunnen machines denken? Dat vroeg de Britse wiskundige en computerpionier Alan Turing zich zo'n zeventig jaar geleden af in zijn artikel 'Computing machinery and Intelligence' (1950). De vraag is eigenlijk: zou een machine menselijke intelligentie kunnen bezitten? Hij bedacht daarvoor een experiment: de *Imitation Game*. In dit experiment wordt getest of een machine een mens kan laten denken dat hij met een echt persoon converseert, dus of een machine kan doen wat een mens doet. Het experiment zou wereldberoemd worden onder de naam *turingtest*.

Niet lang daarna, in 1956, duikt voor het eerst de term *Artificial Intelligence* op; John McCarthy munt de term op een achtweekse zomerconferentie op Dartmouth University (hij zou in 1958 ook de eerste AI-taal ontwikkelen: Lisp).

Tijdens de conferentie werd uitgebreid gediscussieerd over wat AI nu precies is. Bijna alle elf de deelnemers aan de conferentie waren daarna decennialang leiders van AI-onderzoek, en met hun positieve verwachting dat ze in een paar decennia een machine zouden kunnen creëren die zo slim is als een mens, haalden zij miljoenen dollars op voor onderzoek. Een van de grote financiers van het Amerikaanse onderzoek was ARPA (Advanced Research Projects Agency, later ongedoopt tot DARPA), dat door president Dwight D. Eisenhower in 1958 was opgericht in reactie op de lancering van de spoetnik door de Russen. Een belangrijke instelling die hieruit voortvloeit, is het AI-lab van het Massachusetts Institute of Technology, opgezet door een van de deelnemers aan de zomerconferentie, Marcin Minsky.

Een volgende mijlpaal in de beginfase van het AI-onderzoek is een publicatie uit 1959 van Arthur Samuel, ook aanwezig op de Dartmouth-zomerconferentie, waarin hij aantoonde dat een machine in staat is zelf te leren dammen: 'Enough work has been done to verify the fact that a computer can be programmed so that it will learn to play a better game of checkers than can be played by the person who wrote the program.' (Samuel, 1959, p. 211). Naast dammen stonden ook schaken en het oosterse bordspel Go in de belangstelling van veel AI-onderzoekers. In al deze spelen is het van belang om een strategie te bepalen, vooruit te denken en logisch te redeneren – allerlei zaken die onderdeel zijn van menselijke intelligentie, daarom was onderzoek hiernaar binnen de AI-gemeenschap populair.

Tussen 1966 en 1972 ontwikkelde het Amerikaanse Stanford Research Institute, medegefinancierd door ARPA, de eerste intelligente robot: Shakey, een mobiele machine die zelfstandig beslissingen kon nemen op basis van gegevens uit zijn omgeving. De robot maakte voordat hij ging bewegen eerst een kaart van zijn omgeving. Maar dat ging allemaal wel tergend langzaam. Toch is Shakey the Robot wel belangrijk gebleken: hij heeft veel AI-onderzoek geïnspireerd, en ook de huidige manier waarop kaarten door machines worden gelezen en hoe zelfrijdende auto's rijden, is terug te voeren op Shakey.

Ondanks de kleine resultaten werden de verwachtingen van de Dartmouth-conferentie verre van waargemaakt. Op verzoek van de Britse Science Research Council stelde de wiskundige James Lighthill een overzicht samen van academisch onderzoek naar AI. Het zogeheten *Lighthill-rapport*, dat in

1973 verscheen, was zo kritisch over AI, dat op een paar universiteiten na de ondersteuning voor AI-onderzoek werd stopgezet in het Verenigd Koninkrijk. In het Amerikaanse congres werd de druk ook steeds meer opgevoerd om de financiering te stoppen. De jaren die volgden, zouden later bekend komen te staan als de AI-winter, een periode tot eind jaren negentig waarin relatief weinig AI-onderzoek plaatsvond.

Good Old Fashioned AI

In 1978 waren er wel wat oplevingen. Zo ontwikkelde John P. McDermott van de Carnegie Mellon University een expertsysteem om automatisch componenten van een computersysteem te selecteren op basis van de vereisten van klanten: R1 (intern XCON genoemd, voor eXpert CONfigurer). Hij deed dat voor de Digital Equipment Corporation (DEC), die het systeem in 1980 in gebruik nam. Daarmee was DEC het eerste bedrijf dat op industriële wijze met AI aan de slag ging. In 1984 schreef Stephen Polit hierover: 'Many people feel that the primary significance of the XCON experience is that AI techniques were used to solve real-world problems in an industrial setting. However, this is less significant than the fact that an industrial, rather than an academic, user was able to successfully help implement and maintain an AI system.' (Polit, 1984, p. 77). Begin jaren negentig werd echter toch veelal geconcludeerd dat dergelijke expertsystemen te duur waren om te onderhouden.

In Nederland werd gedurende de AI-winter op kleine schaal AI-onderzoek uitgevoerd en richtten negentien AI-onderzoekers in 1981 de Nederlandse Vereniging voor Kunstmatige Intelligentie (NVKI) op. De NVKI verbond Nederlandse AI-onderzoekers van verschillende universiteiten met elkaar om kennis met elkaar uit te wisselen. Maar ook in de jaren tachtig stelden overheden en geldschietters slechts beperkt geld beschikbaar en waren de verwachtingen bij de onderzoekers zelf bescheiden. De opmerking van prof. dr. A.D. de Groot uit 1978 in het volgende interview met prof. dr. H.J. van den Herik, geeft daar blijk van.

CASUS

Tegenspraak en tegenslag leveren de beste voorspellingen*Een interview met Jaap van den Herik*

Op een dag aan het begin van 2019 liepen we op de Universiteit Leiden de kamer binnen van prof. dr. H.J. van den Herik (1947). Nog voordat de koffie was ingeschonken, stelde hij zich voor als Jaap en vroeg hij of we elkaar maar niet beter zouden tutoyeren. Jaap van den Herik is een icoon uit de Nederlandse geschiedenis van AI. Hij is als hoogleraar Informatica en Recht verbonden aan de Faculteit Wiskunde en Natuurwetenschappen en de Faculteit der Rechtsgeleerdheid. Wij zochten hem op omdat wij graag met iemand wilden spreken die de vroege jaren van AI had meegemaakt.

Vanaf het allereerste begin luisterden we geboeid naar de verhalen van Van den Herik over de tijd dat hij studeerde en werkte aan de Vrije Universiteit Amsterdam, TU Delft, Universiteit Leiden, Universiteit Maastricht en de Universiteit van Tilburg. Vol passie vertelde hij over de historie van AI, de kennis die hij heeft opgedaan en zijn toekomstvisie. Na ruim tweeënhalf uur waren we nog steeds niet uitgepraat, terwijl het interview maar een uur zou duren. De rode draad in zijn carrière blijkt dat hij telkens in zijn denken vooropliep, waarbij hij veel weerstand ondervond en ongeloof ontmoette, maar uiteindelijk, soms decennia later, toch zijn voorspellingen bewaarheid zag.

‘Er deugt niets van je voorspelling’

Het begon allemaal toen Van den Herik van zijn schaakhobby een promotieonderzoek kon maken. Voor zijn proefschrift, *Computerschaak, schaakwereld en kunstmatige intelligentie* (1983), stuitte hij al meteen op weerstand van zijn promotor prof. dr. A.D. de Groot. Van den Herik voorspelde namelijk al in 1978 dat een computerprogramma uiteindelijk sterker zou zijn dan de wereldkampioen. De reactie van De Groot was: ‘Ik ben blij dat je het werk van Euwe en mij uit 1963 oppakt en ik zal je graag begeleiden. Overigens deugt er van je voorspelling helemaal niets.’ Ook sloot hij in die jaren met zijn vriend en internationaal schaakmeester Hans Böhm een weddenschap af dat vóór het jaar 2000 de wereldkampioen verslagen zou worden door een machine. Ook Böhm geloofde hem niet.

Uiteindelijk kreeg Van den Herik in 1997 gelijk, toen Garri Kasparov verloor van IBM's schaakcomputer Deep Blue. Voor hem was het een vanzelfsprekende uitslag. Het jaar daarvoor had Van den Herik nog meegewerkt aan de organisatie van de eerste schaakwedstijd tussen Kasparov en Deep Blue, die nog wel door Kasparov werd gewonnen.

De controverse voorbij

In 1991 stelde Van den Herik zich een nieuwe vraag: kunnen computers rechtspreken? Tegenwoordig denken we daar anders over, maar in die tijd was het controversiële vraag. In de VS gebeurt het soms al dat machines rechters ondersteunen in hun besluitvorming, met alle voor- en nadelen van dien. Een voordeel is dat veel zaken in korte tijd kunnen worden afgewikkeld, een nadeel dat het programma bias kan hebben en daardoor bijvoorbeeld racistisch is. Overigens is de stap naar machines die in de rechtspraak zelfstandig besluiten nemen hoe dan ook mogelijk niet zo heel ver weg.

Een andere ethische kwestie sneed Van den Herik aan toen hij begin deze eeuw in een radio-gesprek met bisschop Frans Wiertz van Roermond de vraag stelde of computers konden geloven. Dat leidde uiteindelijk tot de inaugurele rede 'Geloof in computers' (2009) aan de Universiteit van Tilburg.

Bij zijn afscheid van deze universiteit in 2016 poneerde hij de stelling dat intuïtie valt te programmeren. Wederom een controversiële uitspraak, als je kijkt naar de huidige staat van AI. Kort daarna speelde de achttienvoudige wereldkampioen Lee Sedol vijf spelletjes Go tegen AlphaGo van Googles DeepMind. Toen de Go-vereniging in Amstelveen Van den Herik uitnodigde om deze wedstrijd te commentariëren, voorspelde hij Lee Sedol 1 tegen AlphaGo 4. Iemand zei toen: 'Meneer, u vergist u, u bedoelt natuurlijk 4-1.' Maar meneer meende wat hij zei: AlphaGo zou winnen, en Lee Sedol zou één partij winnen doordat er vast wel ergens een bug in de software zou zitten. AlphaGo won daadwerkelijk met 4 tegen 1, een doorbraak in AI.

AI in organisaties

Over zijn verwachtingen over AI in organisaties is Van den Herik even uitgesproken: hij verwacht grote ontwikkelingen in technologie die de naleving van wet- en regelgeving zal ondersteunen (ook wel *Regulatory, Legal en Governmental Tech* genoemd). En ook van het onderzoek naar, en de toepassing van eDiscovery, waarin AI het onderscheid moet kunnen maken tussen vertrouwelijke en geprivilegieerde informatie als bepaalde partijen data en informatie (moeten) uitwisselen in bijvoorbeeld rechtszaken of onderzoeken. Daarnaast verwacht hij de komende jaren veel van Explainable AI: het verklaarbaar maken van de uitkomsten van AI. Daarbij gaat het feitelijk niet om het verklaarbaar maken van de uitkomst, bijvoorbeeld door het algoritme na te lopen, maar om de uitkomst te rechtvaardigen. Op die manier zou je AI daadwerkelijk op een eerlijke manier kunnen toepassen in bijvoorbeeld de rechtspraak. Volgens Van den Herik komt hiervoor hoe dan ook een oplossing.

Aan de andere kant ziet hij gevaren. Het grootste gevaar is misschien wel de ontwikkeling van autonome wapens. Er wordt op hoog politiek niveau veel over gediscussieerd, en een land als

China is dan wel tegen de toepassing van zulke wapens, maar tegelijk wel voor de ontwikkeling ervan. Wat Van den Herik betreft is dit een naïeve discussie, want als één land autonome wapens inzet, moeten alle landen er wel mee bezig zijn.

De geschenken van een rationeel denkproces

Voor Van den Herik zijn de tegenspraak en tegenslag die hij in zijn carrière heeft ondervonden geschenken van het rationele denkproces, omdat de wetenschap leeft van tegenspraak. En vanwege zijn tomeloze energie en passie voor het vak, verwachten we in de nabije toekomst nog veel uitdagende stellingen en onderzoeken van hem voorbij te zien komen. En wellicht ziet hij dan over een jaar of twintig zijn huidige voorspellingen bewaarheid worden.

Tot die tijd lag bij AI-onderzoek de nadruk met name op logisch redeneren en het opstellen van regels (als-dan), ook wel *symbolische AI* of *Good Old Fashioned AI (GOFAI)* genoemd. Vanaf de jaren negentig begon men steeds meer computerprogramma's te creëren die grote datasets konden analyseren en daaruit conclusies konden trekken en ervan konden leren.

Internationaal kwam AI in 1997 weer in de spotlights te staan, het jaar waarin de Russische Garri Kasparov – in die tijd de sterkste schaker ter wereld – het verloor van Deep Blue, een schaakcomputer van IBM. Kasparov kon niet geloven dat hij tegen een computer had gespeeld en niet tegen een mens; hij was ervan overtuigd dat er een echte schaakgrootmeester aan de andere kant had gezeten. Het was opnieuw een mijlpaal voor AI: velen beseften toen dat computers wel eens 'slimmer' zouden kunnen worden dan mensen.

Een nieuwe tijd voor AI

**AI is langzaam in ons
dagelijks leven geslopen.**

Langzaamaan werd AI ook steeds meer onderdeel van dagelijkse gebruiksvoorwerpen en diensten, zoals smartphones, robotstofzuigers, productaanbevelingen op websites, online vertaaldiensten enzovoort.

In 2006 zei Nick Bostrom, directeur van het Future of Humanity Institute aan Oxford University, tegen CNN 'We have seen incremental progress in AI but

not yet the great breakthroughs that people were predicting 30 or 40 years ago. A lot of cutting edge AI has filtered into general applications, often without being called AI because once something becomes useful enough and common enough it's not labelled AI anymore.' (CNN, 2006). In 2011 krijgen we echter te maken met een aantal belangrijke doorbraken die wél zichtbaar zijn als AI. We maken dan kennis met spraakherkenning, als Apple op de iPhone Siri introduceert: een virtuele assistent aan wie je in spreektaal vragen kunt stellen en instructies kunt geven.

In datzelfde jaar komt IBM met een nieuwe supercomputer: Watson, die het spel *Jeopardy!* kan spelen. Bij *Jeopardy!* krijgt een deelnemer een vaak tamelijk cryptische beschrijving die het antwoord op een vraag voorstelt, waarbij hij zo snel mogelijk de juiste vraag moet zien te bedenken. Bijvoorbeeld: de beschrijving (het antwoord) is 'De hoofdstad van Nederland'. De vraag die de deelnemer kan stellen is dan: 'Wat is Amsterdam?' Watson won glansrijk van de beste spelers ooit. Dit wordt gezien als een enorme doorbraak van AI. Een machine was nu in staat om bij een willekeurige vraag in milliseconden het juiste antwoord te genereren. Om als machine taal te verwerken, moet Watson taal begrijpen, maar ook de cryptische omschrijvingen, wat voor een machine dus heel lastig is, alleen al omdat de betekenis van een woord afhankelijk is van zijn context. Dit was radicaal anders dan in 1997 met Deep Blue. Deep Blue ging steeds beter presteren doordat er steeds meer rekenkracht werd toegevoegd aan de hardware, waardoor het programma steeds meer en snellere berekeningen en regels kon uitvoeren. Watson (en vele andere recente AI-doorbraken) maakte echter gebruik van *machine learning*-technieken en heel veel data, en dus steeds minder van voorgeprogrammeerde regels. Expertsystemen en Gold Old Fashioned AI werden vervangen door heel veel data, *machine learning* en neurale netwerken, allemaal mogelijk gemaakt door de enorme hoeveelheid rekenkracht die nu opeens beschikbaar was. In hoofdstuk 2 zullen we hier uitgebreid op ingaan.

In 2014 kondigde Google een zelfrijdende auto aan, die zonder stuur, gas- en rempedaal de weg op ging. De auto, die de naam Firefly had gekregen, werd in 2015 op de wegen van San Francisco uitgeprobeerd. Maar in 2018 bleek dat het niet allemaal rozengeur en maneschijn was, toen een zelfrijdende Uber-taxi een voetganger doodreed.