

Hoe het begon

Kunstmatige intelligentie (artificial intelligence, afgekort AI) ontstond op maandag 18 juni 1956.

18 juni is de internationale dag van angst en paniek, wat heel toepasselijk klinkt voor de dag waarop de mensheid aan haar AI-avontuur begon. Op die dag moet je je vooral niks aantrekken van het beroemde advies van de schrijver Douglas Adams:

In veel van de wat gemoedelijker beschavingen aan de uiterste oostrand van de Melkweg heeft *Het transgalactisch liftershandboek* inmiddels de *Grote galactische encyclopedie* verdrongen als vraagbaak voor alle kennis en wijsheid... Het wint het in twee belangrijke opzichten van dat oudere, conventioneelere werk.

Ten eerste is het een fractie goedkoper, en ten tweede staan in grote, aangenaam ogende letters de woorden GEEN PANIEK op het omslag.

Een andere auteur, Arthur C. Clarke, opperde dat ‘Geen paniek’ waarschijnlijk sowieso het beste advies was dat je de mensheid kon geven. Het dashboard van de oude Tesla Roadster die in 2018 door Elon Musk de ruimte in werd geschoten, was dan ook gesierd met de woorden GEEN PANIEK.

Je had vast niet gedacht dat er in de jaren vijftig van de vorige eeuw al zoiets als AI bestond. Het lijkt immers zo lang geleden. Het was een tijd die al snel nostalgische gevoelens oproept. De burgerrechtenbeweging kwam op. De naoorlogse wereld kende een periode van economisch herstel en stabiliteit. ‘Heartbreak Hotel’ stond hoog in de hitlijsten. Zoals gezegd, lang geleden dus. Was jij in 1956 al geboren? Ik in elk geval nog niet. En ik droom bijna mijn hele leven al over AI – dat wil zeggen, sinds ik als jongetje te veel sciencefictionboeken las. Boeken van auteurs als Arthur C. Clarke en Isaac Asimov, die over een toekomst met robots en intelligente computers schreven. Een toekomst die nu lijkt nu aan te breken.

Toen de AI-chatbot ChatGPT eind 2022 werd gelanceerd, leek hij als een donderslag bij heldere hemel tot ons te zijn gekomen. Je kon geen krant openslaan of het ging over AI. Velen van ons begonnen zich zorgen te maken. Zelfs bij overheden sloeg de paniek toe. Waar zou dit alles toe leiden?

Maar in werkelijkheid ging aan dit plotselinge succes van AI, zoals bij de meeste doorbraken van een bepaald fenomeen, een lange ontstaansgeschiedenis vooraf. Sterker nog, zoals je verderop in dit boek zult ontdekken heb je al decennia met AI te maken gehad. Alleen waren die vormen daarvan vóór ChatGPT misschien wat minder zichtbaar aanwezig.

Je vindt het vast ook nogal apart dat AI een specifieke ‘geboorte’dag kent. Van de meeste wetenschappelijke disciplines is niet op de dag af te zeggen wanneer ze zijn ontstaan. Voor kunstmatige intelligentie ligt dat anders. Op maandag 18 juni 1956 begon een achtdaagse workshop die als doel had na te denken over de bouw van intelligente machines. Deze bijeenkomst vormde zo het startpunt van de wetenschappelijke kunstmatige-intelligentiediscipline.

De workshop werd gehouden op de lommerrijke campus

van Dartmouth College, een Ivy League-universiteit in het mooie stadje Hanover in New Hampshire. Deze universiteit was in 1769 gesticht om inheemse Amerikanen de christelijke theologische leer en de Engelse zeden en gewoonten bij te brengen. In 1956 was er op Dartmouth echter weinig aandacht meer voor inheemse Amerikanen, christelijke theologie of Engelse zeden en gewoonten. De universiteit had zich ontwikkeld tot een van de meest prestigieuze en selectieve bacheloropleidingen in de Verenigde Staten. Zó ‘selectief’ zelfs dat er pas zestien jaar later, in 1972, voor het eerst vrouwelijke studenten werden aangenomen. Aan de workshop in 1956 deden ook alleen maar mannen mee. En helaas zijn er tot op de dag van vandaag maar weinig vrouwen die zich met AI bezighouden. Daar moet beslist iets aan worden gedaan, maar ondanks de vele pogingen die daartoe worden ondernomen, blijft dit probleem bestaan.

De workshop werd georganiseerd door John McCarthy, een jonge docent aan Dartmouth met een ambitieuze droom. Die droom was al oud – een machine maken die kon denken. Daartoe vroeg hij een groep gelijkgestemde collega’s uit de Verenigde Staten, Canada en het Verenigd Koninkrijk om naar New Hampshire te komen, om samen aan een toekomst met kunstmatige intelligentie te werken.

In 1956 kwamen computers voor het eerst algemeen beschikbaar. IBM had eind 1954 de legendarische IBM 650-mainframe op de markt gebracht. Het was de eerste in serie geproduceerde computer. De IBM 650 werd de populairste computer van de jaren vijftig van de vorige eeuw, en ook de eerste computer die winst opleverde. 1956 was dus een geschikt jaar om eens na te gaan hoe hard we die digitale monsters op hun staart konden trappen. Zouden ze ooit zelf kunnen denken? Een brutale vraag, maar John McCarthy was allesbehalve verlegen.



Een jonge John McCarthy.

In 1962 richtte John McCarthy aan Stanford University het legendarische AI Lab op. Dit lab is beroemd geworden door een aantal AI-doorbraken die het op zijn naam heeft staan. Zo kwam het in 2005 met ‘Stanley’ op de proppen, een zelfrijdende auto die de DARPA-wedstrijd won – prijs: twee miljoen dollar – door zonder chauffeur een traject van 212 kilometer door de Mojavewoestijn af te leggen. Het lab leverde in 1996 ook een vrij onbekend bedrijfje op dat met behulp van AI het internet doorzocht. Het heette BackRub. In 1997 vonden de oprichters ervan de naam toch niet zo geslaagd, en veranderden ze die in Google.

Ik heb het geluk John McCarthy persoonlijk te hebben gekend. Dat is minder imposant dan het klinkt: AI is een verrassend klein vakgebied. Volgens de Stanford AI Index zijn er in de Verenigde Staten en Canada jaarlijks minder dan honderd AI-PhD-studenten die hun bul halen en zich vervolgens met wetenschappelijk onderzoek op dit gebied blijven bezighouden.¹ Dat maakt het niet erg spectaculair dat je een van de

grondleggers ervan hebt gekend. Overigens kende ik hem niet alleen, door mijn toedoen is hij op een zondagmiddag in 2006 zelfs bijna in Sydney Harbour verdronken. Maar dat is iets voor een ander boek.*

McCarthy was superslim en had uitgesproken politieke meningen. Hij is vooral bekend doordat hij de term ‘artificial intelligence’ muntte. Die naam bedacht hij toen hij het doel van de bijeenkomst in Dartmouth in 1956 moest omschrijven. Artificial intelligence, kunstmatige intelligentie dus, werd al snel afgekort tot AI, een acroniem van twee letters dat associaties oproept met vergelijkbare acroniemen op het gebied van intelligentie, zoals IQ en EQ.

* Als je wilt weten hoe ik door dit avontuur bijna een voetnoot in de geschiedenis werd, moet je mijn eerste boek lezen: *It's Alive! Artificial Intelligence from the Logic Piano to Killer Robots*.

Het idee van kunstmatige intelligentie was in 1956 uitdagend genoeg om een bedrag van 7500 dollar, als bijdrage in de kosten van de workshop, los te krijgen van de Rockefeller Foundation, een filantropische organisatie die het welzijn van de mensheid wil bevorderen door wetenschap en innovatie te steunen. (Misschien was de Rockefeller Foundation er niet volledig van overtuigd dat AI het welzijn van de mensheid werkelijk kon bevorderen, want bij de aanvraag van de beurs was om 13.500 dollar gevraagd.) In die aanvraag werd onder meer het volgende gesteld:

Alle aspecten van leerprocessen of andere uitingen van intelligentie kunnen in principe zo precies worden omschreven dat er een machine kan worden geconstrueerd die ze kan nabootsen. Er zal worden gepoogd uit te vinden hoe we machines kunnen laten omgaan met taal, vormabstracties en concepten, en problemen kunnen laten oplossen waar nu nog mensen aan te pas moeten komen, waarbij ze zichzelf ook nog kunnen verbeteren.

Alsof dat al niet ambitieus genoeg was, werd er in de aanvraag ook nog een boude voorspelling gedaan over de hoeveelheid tijd die daarmee gemoeid zou zijn: ‘Wij denken dat er wat een aantal van deze problemen betreft aanzienlijke vooruitgang kan worden geboekt als een selecte groep wetenschappers er een zomer lang samen aan werkt.’ Dit bleek een vrij optimistische inschatting, en er lijkt sindsdien weinig veranderd. Veel critici menen dat AI nog steeds te veel belooft en te weinig waarmaakt.

Zoals je vast al vermoedde heeft de achtdaagse workshop in Dartmouth geen ‘aanzienlijke’ vooruitgang geboekt bij de ontwikkeling van AI. Zelfdenkende machines bouwen bleek een veel grotere wetenschappelijke uitdaging dan McCarthy

en zijn collega's hadden gedacht. Dit boek beschrijft in het kort de geschiedenis van onze pogingen om alsnog te doen wat deze pioniers die zomer wel hebben geprobeerd, maar niet voor elkaar kregen. De workshop heeft AI echter wel op de kaart gezet. De belangrijkste resultaten ervan waren waarschijnlijk twee simpele maar uiterst krachtige benaderingswijzen voor het bouwen van AI. De eerste was het gebruik van symbolen. De tweede was het gebruik van zelflerende mechanismen. Deze benaderingswijzen zijn ook terug te zien in de twee belangrijkste 'tijdperken' van de ontwikkeling van AI: het tijdperk van de symboolsystemen, dat tot de jaren negentig van de vorige eeuw duurde, en het daaropvolgende tijdperk van de zelflerende systemen, dat grotendeels verantwoordelijk is voor de huidige AI-hype.

Om die reden wordt deze korte geschiedschrijving ook in twee delen opgesplitst. In het eerste deel behandel ik het eerste, op de symbolische logica gebaseerde tijdperk van de kunstmatige intelligentie. Het was een periode waarin computers leerden schaken en zich in allerlei andere spellen wisten te bekwamen. Maar het was ook een periode van frustratie, omdat we ontdekten hoe moeilijk het zou worden om de simpele spelletjes achter ons te laten en échte AI te bouwen. Maar eigenlijk zou dat ons gerust moeten stellen. Intelligentie is complex. Het bleek iets wat je niet makkelijk in computerchips kon vastleggen. In het tweede deel van dit overzicht kijk ik naar de recentere ontwikkelingen op het gebied van kunstmatige intelligentie. We proberen niet langer eigenhandig AI te programmeren. In plaats daarvan laten we computers ontdekken hoe ze ingewikkelde zaken zélf onder de knie kunnen krijgen. Net zoals jij de meeste intelligente taken die je beheerst hebt *geleerd*, hebben computers inmiddels ook leren lezen, schrijven en rekenen. En dat brengt ons bij de huidige successen als de AI-chatbot ChatGPT, die geleerd heeft

een groot deel van het internet in te lezen.

Het boek eindigt, vrij ongebruikelijk voor een geschiedenisboek, met een blik op de toekomst. De geschiedenis van kunstmatige intelligentie is verre van compleet. Wat gaat er gebeuren als het ons daadwerkelijk lukt zelfdenkende machines te bouwen? Hoelang duurt dat nog? En kan AI een existentiële bedreiging voor de mensheid vormen?

Dit verhaal is ook in een ander opzicht vrij ongewoon. Geschiedschrijvingen gaan meestal over bijzondere mensen en ingrijpende gebeurtenissen. Hier is dat niet het geval. Dit is het verhaal van zes ideeën. Meer niet. Zes ideeën. Als je die doorgrondt, weet je genoeg om goed zicht te hebben op de kunstmatige intelligentie van dit moment. Elk idee krijgt een eigen hoofdstuk. Voor een goed begrip van AI is het hier en daar nodig dat ik een beetje technisch word. Dat zal illustreren dat AI minder op tovenarij lijkt dan de media je soms willen laten geloven.

Uiteraard komen er een aantal bijzondere mensen voorbij, zoals de eerdergenoemde John McCarthy, en de man wiens portret op het Britse biljet van vijftig pond staat afgedrukt en die we beslist een genie kunnen noemen, namelijk Alan Turing. Het boek gaat ook over ingrijpende gebeurtenissen, zoals toen AI beter werd dan mensen en voor het eerst een wereldkampioenschap won.

Voordat ik begin wil ik eerst een waarschuwing van Arthur Samuel herhalen. Samuel was een van de aanwezigen bij de bijeenkomst in 1956 in Dartmouth, en hij schreef tevens een baanbrekend AI-programma voor het spel *checkers* (ook wel *draughts* genoemd, een variant van het damspel op een bord met 8x8 velden), dat een eerste bewijs vormde van de immense mogelijkheden van zelflerende machines. In 1962 schreef hij:

Elke revolutie brengt een aantal halvegaren op de been – mensen die in tovenarij geloven of die zich zo door hun enthousiasme voor de nieuwe zaak laten meeslepen dat ze met wilde beweringen komen die de hele onderneming in diskrediet brengen. Het vakgebied van de kunstmatige intelligentie heeft misschien wel relatief veel van dit soort figuren voorbij zien komen... Desondanks lijkt het een uitgemaakte zaak dat het niet erg lang meer zal duren voordat eentonige mentale taken, die nu zoveel menselijke inspanning vergen, door machines worden uitgevoerd. Voor de mensheid is kunstmatige intelligentie een mythe noch een dreiging.²

Daarom wil ik je met deze nogal persoonlijke en korte geschiedenis van AI helpen bij het op waarde kunnen schatten van alle gekkigheid, wilde claims, mythen en dreigingen rond AI.

AI in de film

De film heeft een sturende rol gespeeld bij zowel de perceptie als de constructie van AI. OpenAI heeft bijvoorbeeld zonder succes geprobeerd een licentie te verkrijgen op de stem van Scarlett Johansson, om die te gebruiken voor de interface met ChatGPT. Dit naar analogie van Samantha, de stem van het AI-besturingssysteem in de film *Her* uit 2013, die ook op Scarlett Johanssons stem gebaseerd was. Onderstaand een paar beroemde voorbeelden van AI-achtige verschijnselen in de geschiedenis van de film.

Metropolis (1927): Deze klassieker van Fritz Lang was een van de eerste avondvullende sciencefictionfilms die ooit werden gemaakt. De film draait om de Maschinenmensch, half mens, half machine, een door de waanzinnige wetenschap-

per Rotwang gebouwde humanoïde robot die Maria nabootst, een geliefde figuur bij de onderdrukte arbeiders in de dystopische stad Metropolis.

Forbidden Planet (1956): In deze klassieke sciencefictionfilm is Robby the Robot een zeer geavanceerd, autonoom en vriendelijk mechanisch hulpje. Robby heeft ook een gastrol in een aantal tv-series, zoals *The Addams Family*, waar hij in 1966 in de aflevering 'Lurch's Little Helper' als Smiley aan treedt.

Blade Runner (1982): de 'replicants' in *Blade Runner* zijn door AI-bestuurde biotechnische robots, die door de Tyrell Corporation zodanig zijn ontworpen dat ze vrijwel identiek zijn aan mensen. Ze worden ingezet voor arbeid in buitenaardse kolonies. Replicants zijn sterker, behendiger en soms ook intelligenter dan mensen. Ze worden echter maximaal vier jaar oud, om te voorkomen dat ze emoties en bewustzijn ontwikkelen. Replicant Roy Batty, gespeeld door Rutger Hauer – het was zijn beroemdste rol –, houdt een van de aangrijpendste en bekendste stervensmonologen uit de filmgeschiedenis: 'Ik heb dingen gezien die jullie mensen niet zouden geloven... Aanvalsschepen, vlak bij Orion. In het donker nabij de Tannhauser-poort zag ik C-beams glinsteren. Al die momenten zullen verloren gaan in de tijd, zoals tranen in de regen... Tijd om te sterven.'

The Terminator (1984): In deze film ontwikkelt een AI-verdedigingsnetwerk genaamd Skynet bewustzijn, waarna het mensen als een bedreiging voor zijn voortbestaan gaat zien. Om zichzelf te beschermen veroorzaakt Skynet een nucleaire holocaust, 'Judgement Day' genoemd. Skynet stuurt een Terminator, een vrijwel onverwoestbare cyborg, terug in de tijd, met als doel Sarah Connor te vermoorden en zo te voorkomen dat ze de leider van de verzetsbeweging van de mensheid zal baren. Aan het begin van de jaren 2000 begon China

met de ontwikkeling van een gigantisch surveillancenetwerk, dat ook Skynet werd genoemd. Volgens de Chinese krant *People's Daily* kan Skynet in één seconde de gezichten van meer dan één miljard mensen scannen.

Ex Machina (2014): In deze film maken we kennis met Ava, een uiterst geavanceerde humanoïde robot die er complexe gedachten en emoties op na houdt en mogelijk zelfs bewustzijn heeft. Ava's realistische uiterlijk en gedrag laat de grens tussen mens en machine vervagen. Murray Shanahan, hoogleraar cognitieve robotica aan het Imperial College en senior onderzoeker bij DeepMind, heeft wetenschappelijk advies gegeven bij de totstandkoming van de film.

Transcendence (2014): Nadat AI-onderzoeker doctor Will Casteris door toedoen van technoterroristen dodelijk gewond raakt, wordt zijn bewustzijn door zijn vrouw en een vriend naar een supercomputer geüpload. Deze AI begint vervolgens zijn vaardigheden uit te breiden, wat niet alleen tot maatschappelijke veranderingen leidt, maar ook de angst voedt dat diens krachten niet meer te beteugelen zijn. Elon Musk heeft er een bijrolletje in. Een jaar later was hij een van de oprichters van het AI-bedrijf OpenAI.

Mission Impossible: Dead Reckoning (2023): Een schurkachtige AI genaamd 'the Entity' bedreigt de wereldvrede. Na uit een geavanceerd Amerikaans cyberwapen te zijn ontstaan ontwikkelt deze AI bewustzijn en wordt hij slimmer dan zijn scheppers; de Entity blijkt defensiesystemen en financiële netwerken te kunnen manipuleren. Ethan Hunt en zijn IMF-team proberen er alles aan te doen om de Entity te vernietigen, maar hij blijkt een geduchte tegenstander, die communicatiekanalen weet te hacken en zich zonder probleem als een mens kan voordoen.