

INKIJKEXEMPLAAR

Tobias dos Santos Gomes is een enthousiaste student filosofie aan de Universiteit van Leiden. Na zijn gymnasiumdiploma wilde hij meer van de grootste denkers uit de geschiedenis weten. Voor het uiten van zijn eigen ambities en de behoefte om kennis te delen en te verspreiden richtte hij het bedrijf 'Poeta Vates' op. In dit bedrijf combineert hij zijn creatieve kant met de liefde voor filosofie.

Bijdragers

Serge en Mariska

Bedankt voor alle mogelijkheden die jullie mij geven, en al hebben gegeven, om te groeien. Dankzij jullie ben ik in welvaart opgegroeid, ben ik nooit iets tekortgekomen en heb ik vertrouwen gekregen om dingen te doen die ik echt leuk vind. Jullie zijn de reden waarom ik ben geworden wie ik ben en mezelf heb weten te verbeteren. Jullie hebben me zelfs aangemoedigd om naar de universiteit te gaan! Bedankt!

Sophie

Bedankt, lieve zus, voor alle fijne gesprekken die we voeren. Bedankt voor je aanmoedigingen en je bijdragen aan dit boek. In elke vorm, bedankt!

Leen en Nelly van Wensem

Ik wil mijn opa en oma bedanken voor elk item over AI dat ze naar me toegestuurd hebben. Dankzij jullie ben ik op het idee gekomen om dit boek te schrijven. Bedankt lieve opa en oma, ik houd van jullie.

Ook wil ik de volgende speciale personen uit mijn leven bedanken:

- **Mijn lieve familie en vrienden** die mij altijd hebben gesteund. In mooie tijden beleefden we geweldige herinneringen. In de moeilijkere tijden waren we er voor elkaar. Jullie hebben mij altijd als een eigen vrije denker behandeld. Ik ben jullie eeuwig dankbaar!
- **Mijn docenten van de basisschool.** Ik ben in groep 7 van school gewisseld en stroomde bij jullie in. Jullie gaven mij altijd het gevoel dat ik wilde leren en wisten de volledige potentie uit mij te halen. Van mavo-advies wisten jullie mij naar vwo-niveau te tillen, dankzij de persoonlijke en oprechte benadering die jullie mij gaven. Jullie leerden mij om door te blijven zetten. Bedankt!
- **De docenten op het Gymnasium** die mij hielpen om mijn liefde voor filosofie te ontdekken. De lastige covid-periode en lockdowns maakten het voor ons allemaal niet gemakkelijk. Toch hebben jullie zo goed als jullie konden mij geholpen om met vlag en wimpel te slagen voor het eindexamen!
- **Ik wil ook nog alle andere mensen bedanken** die, op welke manier dan ook, een ondersteunende rol in mijn leven hebben gespeeld. De betekenisvolle mensen in mijn leven zijn er simpelweg te veel om bij naam en toenaam op te noemen. Jullie weten wie jullie zijn, bedankt!

Ook van Tobias dos Santos Gomes:

Filosofietjes-Boek



Scheurkalender 2026



Filosofietjes-Werkboek



Toegankelijke en praktische filosofie van de grootste denkers uit onze geschiedenis. Filosofietjes is een reeks begonnen in 2024.

WWW.FILOSOFIETJES.NL

De Technologische Revolutie

de onomkeerbaarheid van AI

Colofon

Voor het eerst gepubliceerd op 5 oktober 2025.

Poeta Vates
Rotterdam
www.poetavates.nl

Copyright © Poeta Vates, 2025

Coverontwerp: Tobias dos Santos Gomes

Redactie: MyMedia – Dot Management Services

Geprint en gebonden door CB

Niets uit deze uitgave mag worden verveelvoudigd en/of openbaar gemaakt, door middel van druk, fotokopie, microfilm of op welke andere wijze ook, zonder voorafgaande schriftelijke toestemming van Uitgeverij Poeta Vates.

Uitgeverij Poeta Vates probeert haar boeken zo goed mogelijk te verspreiden. Kunt u een uitgave van Filosofietjes niet vinden in de boekhandel, dan kunt u ook rechtstreeks bestellen bij onze uitgeverij. Stuur een e-mail naar info@poetavates.nl of bezoek www.filosofietjes.nl voor meer informatie, inspiratie en dagelijkse wijsheid.

ISBN 978-90-834715-5-6

*De wezenlijke aard van de techniek
is geenszins iets technisch.*

Martin Heidegger, *The Question Concerning
Technology* (1954)

Inhoud

Introductie

1. Wat is AI?	11
2. Toekomstscenario's van AI	30
3. Pauzeknop!?	47
4. Het (on)menselijke antwoord	58
5. De 'Ai-Trias-Politica' - Montesquieu	73
6. Vernieuwing van de democratie	82
7. Wat is Geluk? – Utilitarisme	95
8. Een eeuwige vrede? – Immanuel Kant	104
9. Is de mensheid eraan toe?	112
Eindwoord	126
Bronnen & Suggesties voor verdere verdieping	129
BONUS: Tools, Trends en Toekomst	132

Introductie

We leven in een tijdperk waarin technologie niet langer alleen een gereedschap is, maar een vormgever van onze realiteit. Kunstmatige intelligentie – AI – is geen sciencefiction meer. Het is overal: in je telefoon, in de aanbevelingen op je scherm, in de muziek die je hoort, de vacatures die je ziet, en zelfs in de woorden die je leest. AI voert sollicitatiegesprekken, schrijft teksten, vertaalt emoties in data, en voert gesprekken die zó natuurlijk klinken dat we soms even twijfelen of er nog wel een mens aan de andere kant zit.

Maar ondanks die razendsnelle opmars is er iets opvallends aan de hand: we begrijpen nog nauwelijks wat AI werkelijk is. Is het een krachtig hulpmiddel, of een sluimerend gevaar? Een slimme machine, of een spiegel van onze eigen intelligentie – met al haar fouten, wensen en vooroordelen? En vooral: wat zegt deze technologische ontwikkeling over óns? Over onze drang naar controle, efficiëntie en gemak? Over wie wij zijn, en wie we zouden kunnen worden?

Stel je voor: een wereld waarin auto's zonder menselijk ingrijpen rijden, koelkasten zelfstandig boodschappen bestellen, artsen worden vervangen door algoritmes die eerder en nauwkeuriger een diagnose stellen dan de beste specialist. Klinkt dit als toekomstmuziek? Het is momenteel al gaande. AI verandert onze samenleving fundamenteel en doet dat sneller dan de tijd die wij nodig hebben om haar te doorgronden. De technologie is klaar. Maar zijn wij als mensheid ook voorbereid?

In dit boek neem ik je mee op een filosofische ontdekkingsreis door het landschap van kunstmatige intelligentie. Geen handleiding voor programmeurs en dus ook geen paniekverhaal voor doemdenkers. Dit is een zoektocht naar betekenis. Want als AI onze wereld verandert, moeten wij opnieuw nadenken over fundamentele

vragen: Wat betekent het om mens te zijn? Wat blijft er over van het begrip ‘ik’ als systemen bestaan die ons beter lijken te kennen dan wij onszelf? En misschien nog belangrijker: waar willen wij eigenlijk naartoe met deze technologie?

We beginnen bij het begin. In **hoofdstuk 1 – Wat is AI?** verkennen we de verschillende definities van kunstmatige intelligentie. We maken onderscheid tussen smalle en algemene AI, leren over neurale netwerken en machine learning, en stellen de vraag: is AI écht intelligent, of slechts een goede imitatie?

In **hoofdstuk 2 – Toekomstscenario’s van AI** schetsen we mogelijke toekomsten: van utopie tot dystopie. Wat als AI ons leven makkelijker maakt? Wat als ze het overneemt? We kijken naar de hoop én de waarschuwingen van wetenschappers, denkers en schrijvers.

Daarna volgt **hoofdstuk 3 – Pauzeknop!?**, waarin we het idee onderzoeken om de ontwikkeling van AI tijdelijk te stoppen. Is dat wijsheid of angst? En kunnen we eigenlijk nog wel stoppen, nu de trein eenmaal op volle snelheid rijdt?

Hoofdstuk 4 – Het (on)menselijke antwoord draait om onszelf: hoe blijven wij mens in een wereld waarin machines menselijke taken overnemen? We onderzoeken de ziel, empathie, opvoeding – en de grens tussen mens en machine.

In **hoofdstuk 5 – De AI-Trias-Politica** passen we het gedachtegoed van Montesquieu toe op kunstmatige intelligentie. Kunnen we de macht van AI verdelen en controleren, net zoals we dat met de staat hebben geprobeerd? Of is macht in een digitale wereld ongrijpbaar?

Hoofdstuk 6 – Vernieuwing van de democratie stelt fundamentele vragen over bestuur en rechtvaardigheid. Wat gebeurt er met de democratie als algoritmes onze wetten gaan uitvoeren – of zelfs maken? Wordt AI de nieuwe filosoof-koning waar Plato

ooit over droomde?

In **hoofdstuk 7 – Wat is geluk?** keren we naar John Stuart Mill en het utilitarisme. Wat als AI het ‘grootste geluk voor het grootste aantal’ probeert te optimaliseren? Is dat wenselijk, of gevaarlijk? En wat betekent geluk in een wereld van constante optimalisatie?

Hoofdstuk 8 – Een Eeuwige Vrede? is gewijd aan Immanuel Kant. Hij droomde van een wereld zonder oorlog. Kan AI, paradoxaal genoeg, juist bijdragen aan wereldvrede? Of vormt ze een nieuw geopolitiek strijdmiddel?

Dan, in **hoofdstuk 9 – Is de mensheid eraan toe?**, reflecteren we op onszelf. Waarom roept verandering zoveel weerstand op? En wat vraagt deze technologische revolutie van onze mindset, onze ethiek en ons onderwijs?

Tot slot sluiten we af met een **Eindwoord**. Geen conclusie in de klassieke zin, maar meer een uitnodiging. Een oproep om niet te stoppen met denken, juist nu alles zo snel gaat. Reflectie, filosofie en betekenis zijn onze belangrijkste gidsen in een tijd van kunstmatige intelligentie.

Dit boek is een beginpunt voor bewustwording over AI. Het is geen verzameling pasklare antwoorden, het slokt je niet op in doemscenario's, maar het vormt een oefening in het stellen van betere vragen. Want in een wereld die steeds slimmer wordt, is de grootste wijsheid misschien wel: weten wanneer je moet reflecteren, reageren én *filosoferen*.

Welkom bij deze reis!

Wat is AI?

1. Wat is AI?

§1.1 Van schaken tot zelfbewustzijn – niveaus van AI

Voordat we ons kunnen buigen over de ethische, maatschappelijke en politieke implicaties van kunstmatige intelligentie (AI), moeten we eerst begrijpen wat AI überhaupt is – en vooral: wat het niet is. Begrijp me goed, ik ben geen ai-expert. Ik studeer filosofie en richt mij vooral op de wijsheden van de oude meesters. Echter, vind ik de hedendaagse ontwikkelingen wél ontzettend interessant en vind ik het leuk om filosofen hierbij te betrekken. Hierom heb ik mij ingelezen in de verschillende posities en begrippen in het debat.

In het dagelijks taalgebruik worden begrippen als AI, algoritme, machine learning en deep learning nogal eens door elkaar gehaald. Maar wat bedoelen we nu eigenlijk als we zeggen: “deze computer is intelligent”?

Volgens Dindi & Smith (2018)¹ is AI “een systeem of algoritme dat computers in staat stelt om taken uit te voeren zonder dat ze daar expliciet voor zijn geprogrammeerd”.

Denk aan een spamfilter dat leert welke e-mails je vervelend vindt, of een aanbevelingssysteem dat begrijpt welke films jij waarschijnlijk leuk vindt. Het lijkt eenvoudig, maar deze definitie impliceert iets revolutionairs: een machine die zelf leert van ervaring – dus niet handelt volgens vaste regels, maar zich aanpast op basis van data.

Om orde aan te brengen in het woud van toepassingen, modellen en maatschappelijke angsten, maken wetenschappers vaak onderscheid tussen drie (of soms vier) niveaus van AI: smalle AI, algemene AI, superintelligentie en zelfbewuste AI. Deze indeling helpt ons om met meer nuance te praten over wat AI wel en niet is. Hiermee kunnen we vragen over ethiek, democratie en

¹ Dindi & Smith (2018), *Hands-On Artificial Intelligence for Beginners*, p. 6

menselijkheid op de juiste plaats stellen.

Smalle AI – intelligentie in een hokje

Dit is de AI die vandaag overal om ons heen is. Ze is slim, maar slechts binnen één domein. Ze kan beter schaken dan de wereldkampioen, maar weet niet dat ze een spel speelt. Ze herkent gezichten op foto's, maar begrijpt niets van emoties. Ze stelt medische diagnoses, maar kan geen empathisch gesprek voeren met de patiënt.

Dit noemen we narrow AI – smalle AI – en het is momenteel het meest voorkomende type AI. Denk aan spamfilters, Siri, Google Translate, gezichtsherkenning bij het ontgrendelen van je telefoon, of Netflix-algoritmes die jouw voorkeuren voorspellen. Elk van deze systemen is getraind op een specifieke taak en doet die taak vaak briljant – beter dan een mens – maar kan niets daarbuiten.

In 1997 versloeg IBM's Deep Blue (schaakcomputer) schaakgrootmeester Garry Kasparov. Dit was een mijlpaal voor AI, maar het systeem wist niet wat winnen of verliezen betekende. Deep Blue beschikte over brute rekenkracht en slimme optimalisatie. Later, in 2016, versloeg AlphaGo – een geavanceerd deep learning systeem ontwikkeld door Google DeepMind – wereldkampioen Lee Sedol in het spel Go. Go geldt als complexer dan schaken, met meer mogelijke posities dan er atomen zijn in het universum. Toch wist AlphaGo geen regels, geen spelplezier en geen strategie in menselijke zin. Het optimaliseerde een abstracte waarde over duizenden simulaties heen.

Deze voorbeelden laten zien: narrow AI is effectief, maar niet bewust. Het is geen 'denkende entiteit'. Toch kan narrow AI al enorme maatschappelijke impact hebben. En daar begint onze filosofische vraag: hoe beïnvloedt een niet-bewuste, maar wel krachtige intelligentie onze samenleving?

Algemene AI – denken zoals wij?

General AI (AGI) verwijst naar een hypothetische vorm van kunstmatige intelligentie die in staat is om elk cognitief probleem op te lossen dat een mens ook kan oplossen – én zich flexibel kan aanpassen aan nieuwe situaties, contexten en taken. Denk aan een AI die zowel kan vertalen als schaken, koken, emoties herkennen, juridische oordelen vellen en poëzie schrijven – zonder aparte training per taak.

In tegenstelling tot narrow AI heeft general AI het vermogen tot generalisatie: leren in het ene domein helpt om beter te presteren in een ander domein. Dit is precies wat mensen wél doen, en AI-systemen vandaag nog niet. General AI vormt nog altijd het ultieme doel van de AI-gemeenschap, maar ondanks indrukwekkende taalmodellen zoals ChatGPT of GPT-4, hebben we dit niveau nog niet bereikt.

Wat we nu zien – grote taalmodellen, beeldgenerators, autonome voertuigen – zijn complexe combinaties van smalle systemen, die soms generalistisch lijken, maar geen intrinsiek begrip hebben van de wereld. Ze hebben geen ‘wereldmodel’, geen bewustzijn van tijd, ruimte, oorzaak-gevolg, of van zichzelf.

Filosofisch roept dat vragen op: is denken slechts patroonherkenning? Of vereist het ook zelfreflectie, ervaring, belichaamde waarneming? Kunnen we ooit een digitale entiteit bouwen die begrijpt, voelt, en keuzes maakt vanuit waarden?

Superintelligentie – meer dan menselijk(?)

Het derde niveau is superintelligentie: een hypothetische entiteit die de menselijke intelligentie op alle domeinen overtreft. Niet alleen sneller en preciezer, maar ook creatiever, moreler, empathischer, strategischer. Een systeem dat Shakespeare beter begrijpt dan literatuurwetenschappers, én betere medische beslissingen neemt dan alle artsen samen.

Hoe zien de toekomstscenario's eruit?

§2.4 Wat zeggen experts?

Wat als we het toekomstscenario van AI niet baseren op speculatie, maar op de inschattingen van de mensen die er nu dagelijks aan bouwen, over schrijven of beleid over maken?

In dit hoofdstuk bespreken we de visies van diverse experts uit het veld: van de filosofisch geïnspireerde analyses van Alexander Klöpping, tot het beleidsmatige rapport van Peter Stone, en de kritische oproep tot pauze van Elon Musk en Steve Wozniak. Deze stemmen verschillen in toon en verwachting, maar tonen allemaal: AI is geen verre droom meer. Het gebeurt nu.

Alexander Klöpping: “AI is nu al slimmer dan we willen toegeven”

Techondernemer en commentator Alexander Klöpping waarschuwt dat veel beleidsmakers achterlopen bij de realiteit. Hij citeert onder andere Ezra Klein (New York Times), die stelt dat binnen 2 à 3 jaar AI slimmer zal zijn dan de mens – een visie die wordt gedeeld door CEO’s van grote AI-labs.²¹

In zijn analyses wijst Klöpping²² ook op DeepSeek, een Chinees AI-model dat qua prestaties vergelijkbaar is met GPT-4 maar slechts een fractie van de rekenkracht vergt. Het denkt als een mens, redeneert als een mens, en draait lokaal op een MacBook. Dat verandert alles, volgens hem.

Klöpping’s oproep is helder: zelfs als je zelf niet gelooft dat general AI (AGI) er al bijna is, zou je als maatschappij nú scenario’s moeten doordenken. Want de economische gevolgen zijn enorm. Zo liet hij bij Eva Jinek²³ zien hoe een AI-team in tien minuten een

²¹ Klöpping, A. (2025), *Iedereen kan programmeren met AI* – Nieuwsbrief 7 maart 2025.

²² Klöpping, A. (2025), *Is AI over 3 jaar slimmer dan mensen?* – Nieuwsbrief 24 januari 2025.

²³ AVROTROS. (2024). *Alexander Klöpping laat zien hoe AI jouw kantoorbaan overneemt: ‘We hebben marketingbureau EVA opgericht’*. [Tv-uitzending Eva 24 juni] <https://eva.avrotros.nl/artikel/alexander-klopping-laat-zien-hoe-ai-jouw-kantoorbaan-overneemt-we-hebben-marketingbureau-eva-opgericht-624>

marketingcampagne ontwikkelde – werk dat normaal weken zou kosten.

Zijn zorg: deze revolutie is stil, sluipend, en onvoorbereid. We hebben een Minister van Digitale Zaken nodig. Nu.²⁴

Peter Stone: “AI moet proactief worden gereguleerd”

Computerwetenschapper Peter Stone was in 2016 hoofdonderzoeker van het rapport “Artificial Intelligence and Life in 2030”, een initiatief van de Stanford One Hundred Year Study on Artificial Intelligence (AI100). Stone waarschuwde dat AI alle sectoren van de samenleving zal beïnvloeden: van vervoer tot geneeskunde, van onderwijs tot justitie. Zijn visie is genuanceerd: AI brengt vooruitgang, maar ook disruptie, en dat vereist actieve beleidsontwikkeling²⁵

Opvallend is dat Stone al vroeg waarschuwde voor “algoritmische vooringenomenheid” en de noodzaak van transparantie, uitlegmodellen en controleerbaarheid. Hij pleitte voor sector-specifieke wetgeving en overheidsinvesteringen in AI-veiligheid en -infrastructuur. Zijn benadering is dus voorzichtig optimistisch: AI hoeft geen bedreiging te zijn – mits goed gemanaged.

Elon Musk en de pauzeknop

Een radicaal andere toon klinkt in de open brief van maart 2023 van het ‘future of life institute’, waarin Elon Musk, Steve Wozniak en meer dan 1.000 experts pleitten voor een pauze van minstens zes maanden op het trainen van systemen krachtiger dan GPT-4. Hun waarschuwing: “Contemporary AI systems are now becoming human-competitive at general tasks. [...] Should we develop nonhuman minds that might eventually outnumber, outsmart, obsolete and replace us?”²⁶

²⁴ Klöpping, A. (2025), *Minister van Digitale Zaken en AI: NU!* – Nieuwsbrief 3 juni 2025.

²⁵ Stone, P. (2025), *Artificial Intelligence: Looking Forward 15 Years*

²⁶ Future of Life Institute (2023), *Pause Giant AI Experiments: An Open Letter*, van

Een (on)menselijk antwoord?

4. Het (on)menselijke antwoord

§4.1 Wat maakt ons menselijk?

Kunstmatige intelligentie roept vaak angst op omdat het dingen kan die wij lang als uniek menselijk beschouwden. Taal begrijpen. Redeneren. Schrijven. Vertalen. Tekenend. Coderen. Argumenteren. AI komt binnen in het domein waarvan we dachten dat het het onze was.

Maar misschien stelt dat ons juist in staat tot het stellen van een diepere, eerlijkere vraag: als machines steeds meer ‘menselijke’ capaciteiten overnemen, wat blijft er dan écht over van ons? Wat maakt ons, ondanks alles, onmiskenbaar menselijk?

Is het onze ratio? Creativiteit? Bewustzijn? Ons lichaam? Kwetsbaarheid? Of juist onze ethische vermogens? Het antwoord is wellicht niet eenvoudig of eenduidig, maar hoe beter we begrijpen wie wij zijn, hoe beter we kunnen bepalen op welke manier we willen samenleven met deze nieuwste vorm van technologie.

Wij kunnen lijden – en dat erkennen

AI kan verdriet herkennen, maar niet voelen. Ze kan rouw verwoorden, maar niet dragen. Ze heeft geen lichaam, geen honger, geen breekbaarheid. Mensen daarentegen zijn wezens die gewond kunnen raken – fysiek, psychisch, existentieel. We zijn kwetsbaar. En juist die kwetsbaarheid maakt ons ontvankelijk voor zorg, verbondenheid en compassie.

De Duitse filosoof Martin Heidegger karakteriseert dan ook in zijn hoofdwerk *Sein und Zeit* (1927) de mens als een wezen dat fundamenteel wordt gekenmerkt door zorg. Ook sprak filosoof Hans Jonas in *Das Prinzip Verantwortung* (1979) over “de ethiek van verantwoordelijkheid richting de kwetsbare ander”: wij voelen

Een eeuwige vrede? – Immanuel Kant

8. Een Eeuwige Vrede?

§8.1 Kant's visie op een wereldregering

Immanuel Kant betoogt in zijn beroemde werk “*Zum ewigen Frieden. Ein philosophischer Entwurf*” (1795) dat duurzame wereldvrede geen naïeve droom is, maar een morele plicht die voortvloeit uit de menselijke rede. Zijn benadering begint bij het inzicht dat de ‘natuurtoestand’ tussen staten – evenals die tussen individuen – in feite een toestand van oorlog is. Zelfs als er geen actieve oorlog woedt, blijft de dreiging van geweld in de natuurtoestand bestaan zolang er geen bovenliggende juridische structuur is.

Om deze situatie te overwinnen stelt Kant drie niveaus van recht voor⁵²:

1. Het staatsrecht (*ius civitatis*): binnenlandse juridische orde;
2. Het volkenrecht (*ius gentium*): juridische orde tussen staten;
3. Het kosmopolitisch recht (*ius cosmopoliticum*): universele gastvrijheid en mensenrechten wereldwijd.

Kant pleit niet voor een wereldstaat (zoals een soevereine wereldregering), omdat dat zou leiden tot despotisme: een enkel machtscentrum dat, eenmaal gecentraliseerd, alle vrijheid tenietdoet. In plaats daarvan stelt hij een “federatie van vrije staten” voor, een ‘*foedus pacificum*’, die zich vrijwillig verbinden aan een gemeenschappelijk rechtssysteem om oorlog te voorkomen, maar hun interne autonomie behouden.

De wereldvrede wordt niet bereikt via een opgelegde

⁵² Kant, I. (1927). *Perpetual Peace: A Philosophical Proposal* (H. O'Brien, Trans.).

wereldregering, maar via een voortschrijdende uitbreiding van deze vrijwillige federatie. Kant noemt dit idee niet utopisch, maar als iets waar de menselijke rede en de natuur gezamenlijk naartoe kunnen werken: de natuur dwingt mensen indirect tot samenwerking, onder andere via handel en de gevolgen van oorlog.

In zijn beroemde analogie vergelijkt Kant staten met individuen: net zoals individuele mensen hun vrijheid niet kunnen behouden zonder wet, kunnen ook staten onderling geen vrede bewaren zonder een gedeeld juridisch kader. Zijn beroemde “eerste definitieve artikel”⁵³ luidt dan ook dat de staatsinrichting van elke staat “republikeins”⁵⁴ moet zijn, omdat alleen dan burgers de prijs van oorlog echt overwegen voordat zij ermee instemmen.

Tot slot onderstreept Kant het belang van kosmopolitisch recht, vooral het ‘recht op gastvrijheid’, waarbij iedere vreemdeling het recht heeft om niet vijandig ontvangen te worden op andermans grondgebied. Dit is volgens hem cruciaal in een wereld die steeds meer met elkaar verweven raakt.

§8.2 Is AI de neutrale arbiter die Kant bedoelde? Vrede door datagestuurde samenwerking

Immanuel Kant schreef in *Zum ewigen Frieden* (1795) dat duurzame vrede tussen staten pas mogelijk wordt wanneer er een universele rechtsorde ontstaat: een stelsel waarin geen enkele staat de ander mag overheersen, en waarin elk besluit onderworpen is aan publiek debat, juridische normering en wederzijdse verantwoording.

Zijn ideaal was géén wereldstaat (dat zou te veel macht centraliseren), maar een federatie van republikeinse staten die zich gezamenlijk onderwerpen aan rationele beginselen van rechtvaardigheid. Volgens Kant is het niet medelijden of vriendschap die vrede mogelijk maakt, maar onze rede: een bewust

⁵³ Kant, I. (1927). *Perpetual Peace: A Philosophical Proposal* (H. O’Brien, Trans.).

⁵⁴ De republikeinse staatsvorm is voor Kant gebaseerd op drie kernprincipes: vrijheid van de mens, onderwerping aan een gemeenschappelijke wet, en gelijkheid van alle burgers.

Is de mensheid eraan toe?

9. Is de mensheid eraan toe?

§9.1 De geschiedenis van technologische angst. Vliegtuigen, treinen, televisie – en nu AI

Wanneer we vandaag discussiëren over de risico's van kunstmatige intelligentie, klinkt daar vaak een diepe existentiële zorg in door: “Wat als het misgaat? Wat als we de controle verliezen?” Maar deze zorg is niet uniek voor onze tijd. Angst voor nieuwe technologieën is van alle eeuwen. Steeds wanneer er een radicale sprong voorwaarts wordt gemaakt, schrikken we terug. Dit doen we niet alleen uit rationele bezorgdheid, maar ook vanuit iets diepers: ons oerinstinct waarschuwt dat verandering (mogelijk) bedreigend is.

De trein die melk liet bederven (maar niet echt)

Toen de eerste stoomtreinen door het landschap reden in de 19e eeuw, waren mensen verbijsterd. Sommige artsen beweerden dat de snelheid (30 km/u!) het menselijk lichaam uit balans zou brengen. Er werd gevreesd dat koeien, geschrokken van de ronkende locomotieven, zouden stoppen met melk geven. Spoorlijnen werden aanvankelijk buiten dorpen aangelegd, zodat de ‘gemoedsrust’ van bewoners niet verstoord werd.

Toch: vandaag de dag reizen we achteloos met 300 km/u in hogesnelheidstreinen, zonder dat onze ingewanden het begeven of onze hersenen oververhit raken. Er is daarnaast nog nooit zoveel melk in Nederland geproduceerd als afgelopen jaren...

De vliegende doodskist

Ook de eerste vliegtuigen – de broertjes Wright rond 1903 – werden met argwaan bekeken. “De mens is niet gemaakt om te vliegen,”

Bronnen & Suggesties voor verdere verdieping

- Arendt, H. (1963). *Eichmann in Jerusalem: A Report on the Banality of Evil*
- Arendt, H. (1958). *The Human Condition*
- Aristoteles. (1987). *De Anima (On the Soul)* (H. Lawson-Tancred, Trans.). Penguin Books.
- AVROTROS. (2024). *Alexander Klöpping laat zien hoe AI jouw kantoorbaan overneemt: 'We hebben marketingbureau EVA opgericht'*. [Tv-uitzending Eva 24 juni]
<https://eva.avrotros.nl/artikel/alexander-klopping-laat-zien-hoe-ai-jouw-kantoorbaan-overneemt-we-hebben-marketingbureau-eva-opgericht-624>
- Bandura, A. (1977). *Social Learning Theory*.
- Bentham, J. (1789). *An introduction to the principles of morals and legislation*.
- Block, N. (1995). On a confusion about a function of consciousness. In *Behavioral and Brain Sciences*, 18(2), pp. 227–247
- Buolamwini, J. & Gebru, T. (2018). *Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification*. MIT Media Lab
- Carlson, R. (1997). *Don't Sweat the Small Stuff*
- Chalmers, D. (2010). Facing Up to the Problem of Consciousness. In *The Character of Consciousness*, pp. 4–6
- Covey, S. (1989). *The 7 Habits of Highly Effective People*
- Dastin, J. (2018). *Amazon scraps secret AI recruiting tool that showed bias against women*. Reuters.
- Descartes, R. (1996). *Meditations on First Philosophy* (J. Cottingham, Trans.). Cambridge University Press.
- Dindi & Smith (2018). *Hands-On Artificial Intelligence for Beginners*

- Foucault, M. (1975). *Discipline and Punish*
- Frankl, V. (1946). *Man's Search for Meaning*
- Future of Life Institute (2023), *Pause Giant AI Experiments: An Open Letter*, van <https://futureoflife.org/open-letter/pause-giant-ai-experiments/>
- Gebru, T. (2023), In Roppolo, M. (2023), *Elon Musk joins hundreds calling for a six-month pause on AI development in an open letter*. CBS News
- Goleman, D. (1995). *Emotional intelligence: Why it can matter more than IQ*.
- Gorz, A. (1985). *Paths to Paradise: On the Liberation from Work*
- Gramsci, A. (1971). *Selections from the Prison Notebooks*.
- Heidegger, M. (1927) *Sein und Zeit*
- Heidegger, M. (1954). *The Question Concerning Technology*
- Jonas, H. (1979). in *Das Prinzip Verantwortung* (1979)
- Kant, I. (1785). *Perpetual Peace: A Philosophical Proposal* (H. O'Brien, Trans.).
- Kelly, K. (2010), *What Technology Wants*
- Klöpping, A. (2025), *Iedereen kan programmeren met AI – Nieuwsbrief 7 maart 2025*.
- Klöpping, A. (2025), *Is AI over 3 jaar slimmer dan mensen? – Nieuwsbrief 24 januari 2025*.
- Klöpping, A. (2025), *Minister van Digitale Zaken en AI: NU! – Nieuwsbrief 3 juni 2025*.
- Kurzweil, R. (2005), *The Singularity Is Near*
- Lane, M. (2012). *Eco-Republic: What the Ancients Can Teach Us about Ethics, Virtue, and Sustainable Living*.
- Larson, J., Angwin, J., Mattu, S., & Kirchner, L. (2016), *Machine Bias: There's software used across the country to predict future criminals. And it's biased against blacks*.

- Mill, J. S. (2009). *On Liberty*; chapter 3 On Individuality, as One of the Elements of Wellbeing, p.101
- Mill, J. S. (1863). *Utilitarianism*
- Montesquieu, C.L. (2024). *The Spirit of Laws*, (W.B. Allen, Trans.),
- Nagel, T. (1974), *What Is It Like to Be a Bat?*
- Nick Bostrom (2014), *Superintelligence*
- Niiler, E. (2019). *Can AI Be a Fair Judge in Court? Estonia Thinks So*. Wired
- Pattison, G. (2000), *Routledge Philosophy Guidebook to The Later Heidegger*
- Plato. (2012). *Republic* (C. Rowe, Trans.). Penguin Classics. Book 7
- Pilarczyk, M. (2016). *Master Your Mindset*
- Robbins, T. (1992). *Awaken The Giant Within*
- Roovers, D. (2014). *School*. Filosofie Magazine
- Searle, J. (1980), *Minds, brains, and programs*
- Soranno, D. E., et al. (2022). *Artificial intelligence for AKI!- Now: Let's not await Plato's utopian Republic*. Kidney 360
- Sternberg, R. (1948), *Toward a triarchic theory of human intelligence*, In *The Behavioral and Brain Sciences*, 7, p. 271
- Stone, P. (2025), *Artificial Intelligence: Looking Forward 15 Years*
- The New York Times (1903). *Flying Machines Which Do Not Fly*
- Turing, A. (1950), *Computing Machinery and Intelligence*
- Vinge, V. (1993), *The Coming Technological Singularity: How to Survive in the Post- Human Era*
- White, S. (2024). *How AI is reshaping the future of legal practice*. The Law Society.
- Winner, L. (1986). *The Whale and the Reactor*

*“AI leert razendsnel. Maar wij zijn
degenen die bepalen waarheen. Dit boek
is mijn bijdrage aan die richting.”*

-Tobias dos Santos Gomes, 2025-

Stichting Parkinsonfonds

Mijn opa is een van de vele mensen die leven met de ziekte van Parkinson. Zijn moed en doorzettingsvermogen in het omgaan met deze aandoening hebben me geïnspireerd om ook een bijdrage aan het bestrijden van deze ziekte te leveren. Daarom heb ik besloten om € 1,00 per verkocht boek te doneren aan Stichting Parkinsonfonds.

Stichting Parkinsonfonds zet zich in voor onderzoek naar deze slopende ziekte en biedt ondersteuning aan patiënten en hun families. Door uw aankoop helpt u direct mee in de strijd tegen Parkinson en draagt u bij aan een toekomst met betere behandelingen en uiteindelijk, hopelijk, een genezing.



Elke aankoop van mijn boek is een stap op weg naar persoonlijke groei, en een stap dichterbij een wereld zonder Parkinson. Samen kunnen we een verschil maken. Bedankt voor je hulp!

STICHTING
ParkinsonFonds