

AI: Overwegingen voor wie erover gaat - Bert Hubert's writings

Onlangs presenteerde ik kort na elkaar bij het [Netwerk van Publieke Dienstverleners](#) en bij de [Adviesraad Wetenschap, Technologie en Innovatie](#) over AI. In deze twee verschillende maar korte presentaties hoopte ik wat inzichten te delen die nuttig zijn voor mensen die nu aan het stuur zitten, mensen die keuzes kunnen of moeten maken over AI-beleid. Dank aan NPD en AWTI voor de uitnodigingen en de stimulerende discussies!

 There is also [an English version of this page!](#)

Dit artikel is gebaseerd op deze twee verschillende presentaties. [De eerste bij NPD](#) ging over de concrete uitdagingen voor bestuurders. [De tweede bij de AWTI](#) over wat er nodig is om te komen tot een betrouwbare (Europese) AI infrastructuur. En wat dat betekent overigens. Na deze presentaties ben ik gaan schrijven, en uiteindelijk tot dit stuk gekomen.

Bij dat schrijven ontdekte ik hoe lastig het allemaal is. Er zijn enorme voorstanders van AI, en er zijn mensen die menen dat het de duivel is. Deze kampen zijn zo gepolariseerd dat het onmogelijk is een stuk te schrijven waar iedereen blij mee kan zijn.

Wel is het mogelijk een verhaal in te typen waar iedereen ergens wel boos van wordt. Bij deze. En toch hoop ik dat zowel voor- als tegenstanders na lezing van dit stuk wijzer zijn. En misschien is er toch iets waar we het eens over kunnen zijn: is er wel zo'n haast allemaal? Moet iedere gemeente nou echt dit jaar al aan de Copilot? Moest de Tweede Kamer z'n uitrol nu al doen voor iedereen?

In wat volgt hou ik geen "[both sides](#)" verhaal. Voor de mensen die een totale hekel hebben aan AI is onderstaand te vinden dat AI toch echt heel indrukwekkende dingen kan. Voor AI fans heb ik een berg dingen om over na te denken of het nou echt een goed idee is.

Maar netto maak ik me enorm zorgen of we niet veel te hard van stapel lopen, en ik denk dat dat gegeven de feiten geen onredelijke positie is.

► Spoilers - klik hier voor een uitgebreide samenvatting

AI: We weten het niet

De AI van nu is niet de AI van volgend jaar en de maatschappelijke impact verandert ook non-stop. Het is echt een achtbaan. Maar, er zijn wel manieren om in de wirwar toch chocola te maken van wat er gaande is, en wat verstandig is om (niet) te doen. Dit artikel hoopt daaraan bij te dragen.

Iedereen die nu wel claimt te weten hoe het zit is op z'n minst in de war. Als voorbeeld deze reclame van PwC:

74%

AI-waarde is ongelijk verdeeld:

Slechts 20% van de onderzochte organisaties realiseert 74% van alle AI-gedreven opbrengsten. Wat doen deze AI-leiders anders?

[Bron](#)

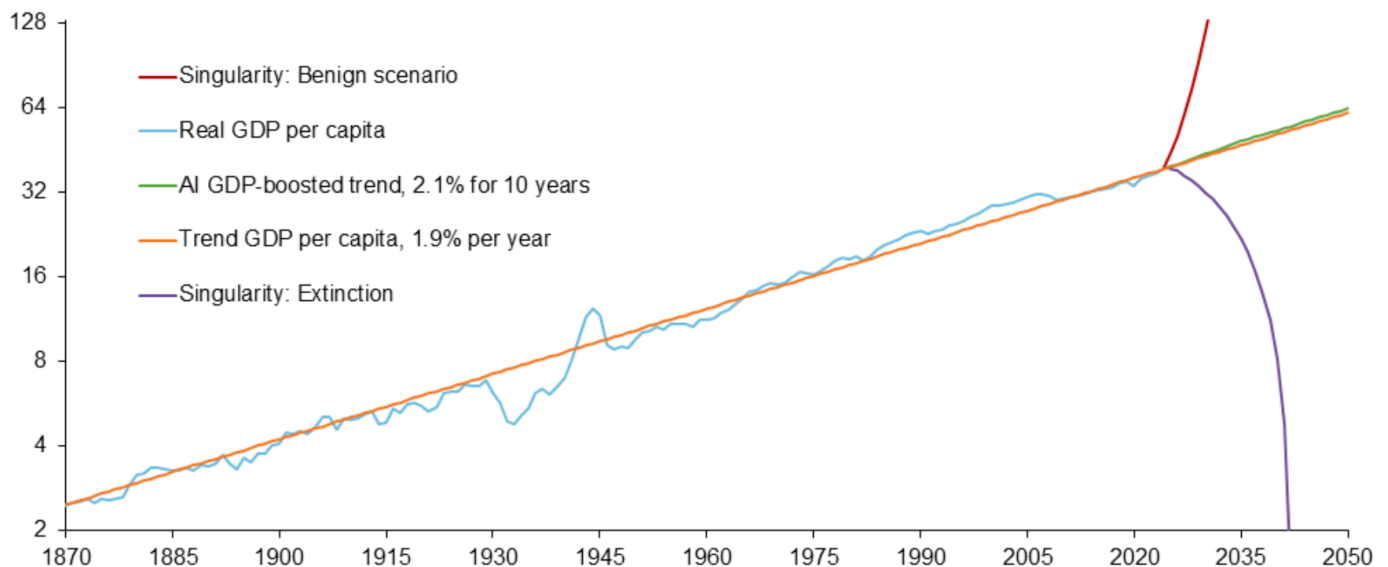
Waarbij ik graag opmerk dat 87,539% van alle statistieken verzonnen schijnt te zijn.

Een eerlijkere grafiek is deze van de Federal Reserve Bank of Dallas:

Chart 1

AI scenarios

1990 dollars (thousands), log scale



NOTES: The blue line is real gross domestic product (GDP) per capita in 1990 dollars. The orange line is a trend line fitted to the data for 1870-2024 with a trend growth rate of 1.9 percent per year. The red, green and purple lines are hypothetical paths for per capita GDP based on different scenarios.

SOURCES: Bureau of Economic Analysis; Haver Analytics; Macrohistory.net; United Nations; authors' calculations.

Federal Reserve Bank of Dallas

[Bron](#)

Volgens de “Dallas Fed” zal de economie door AI **vier** keer zo groot worden, of **zestien** keer kleiner. Ergens daar tussenin zal het wel zitten!

Dus neem iedereen die zegt het wel precies te weten niet al te serieus. Dat betekent overigens niet het geen zin heeft om na te denken hoe tegen het fenomeen AI aan te kijken.

De Nobelprijs voor scheikunde

Er zijn mensen die AI graag afserveren als onzin. “Fancy autocomplete”, “stochastic parrot”. Hier een tegenvoorbeeld uit de biologie waar ik toevallig zelf ook [goed in thuis ben](#):

The screenshot shows the website of Universiteit Leiden. At the top, there is a search bar and navigation links for Research, Education, Academic staff, About us, Collaboration, Faculties, Campus The Hague, Alumni, and Library. The main headline reads: "The Nobel Prize in Chemistry went to an AI model (and rightly so)" with a red arrow pointing to the word "AI". Below the headline, the text states: "Not experiments and lab coats, but computers and artificial intelligence: this year's Nobel Prize in Chemistry went to the inventors of the groundbreaking AI model, AlphaFold. This programme accurately predicts protein structures based on their genetic code—a crucial step in understanding biological functions and developing new medicines. At Leiden University, it has already become indispensable, say professors Remus Dame and Gerard van Westen." To the right of the text, there are two images: a black and white photo of a man in a lab coat working with glassware, and a color photo of a woman in a lab coat working at a computer. A large black arrow points from the first image to the second. On the right side of the page, there are sections for "Research" (Chromatin organisation is dynamics) and "Teaching" (Life Science & Technology (MSc), Life Science and Technology (BSc), Chemistry (MSc), Molecular Science and Technology (BSc)).

De mensen die een AI model hebben uitgevonden hebben wel een echte Nobelprijs in de scheikunde gewonnen. En zoals de [Universiteit Leiden hier zegt](#): dat was volkomen terecht. Met een AI-model kunnen we nu vanuit DNA de vorm van eiwitten voorspellen. Dit is iets wat mensen niet kunnen, en waar computers tot nu toe ook niet goed bij konden helpen. En daarom was die Nobelprijs meer dan terecht. In de biologie en farmacie wordt dankbaar gebruik gemaakt [van deze nieuwe mogelijkheden](#). De architectuur van dit AI-model [is heel vergelijkbaar met die van Large Language Models](#).

Er zijn mensen die zeggen dat deze AI niet echt intelligent is, en “[niet weet wat ‘ie aan het doen is](#)”. En dat klopt best. Maar deze AI-toepassing maakt nu al ontdekkingen mogelijk waar we allemaal profijt van hebben.

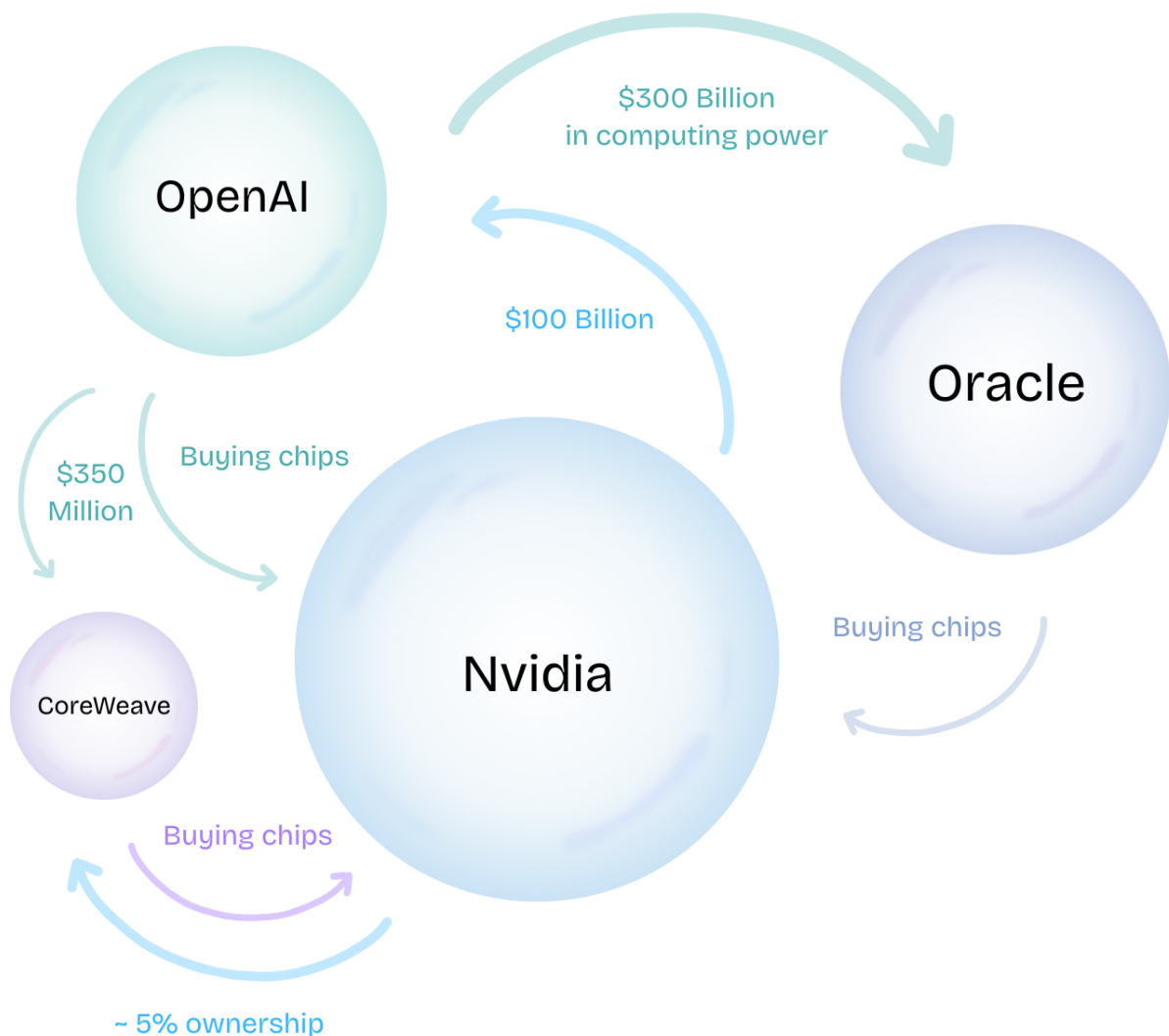
AI kan [ook heel knap hacken](#) (alleen hackt ie [niet](#) altijd [wat je bedoelde](#)), en ook heel aardig vertalen en supermenselijk goede audiotranscriptie doen (al moet je, wederom, goed checken wat ie doet, en dat doet uiteindelijk bijna niemand). Je kan ook heel indrukwekkende demo's bij elkaar *viben*.

Dus nee, AI is niet af te serveren als “allemaal onzin”. Maar, dat is geen universele aanbeveling. Sterker nog:

Ik moet het toch noemen: de rechtszaken, het casino, het klimaat, digitale autonomie, netcongestie

Men vindt dit bijna flauw als je erover begint, maar alle grote aanbieders van Large Language Model-diensten zijn verwikkeld in gigantische rechtszaken met [uitgevers](#), academici, [media](#), [kranten](#) etc. Onder normale omstandigheden zouden we niet in zee willen met partijen die kennelijk [zulke controversiële dingen doen](#), maar voor AI is dat gewoon OK kennelijk!

Het is enorm grappig om AI-boeren te vragen of ze je vrijwaren tegen intellectuele eigendomsclaims over hun output. Vraag je juristen eens om dat uit te zoeken. Het antwoord (of het gebrek daaraan) zal je verbazen.



Door Catboy69 - eigen werk, [CC0](#)

De financiële toestand van de AI-aanbieders is spectaculair zorgwekkend. Men investeert allemaal “circulair” in elkaar, en een van de grootste partijen (Oracle) heeft z’n kredietwaardigheid inmiddels zien dalen tot [1 tandje boven “junk”](#). Het is nogal wat om je toekomst te bouwen [op een dubieus casino](#). Ook het schuldpapier van AI-bedrijf van Elon Musk wordt inmiddels verhandeld als “[junk bond](#)”.

De [Financial Times schreef deze week](#):

“Whichever way, AI scepticism now feels like the consensus. As for what to do in preparation for a crash, you’re basically on your own. But if the music does stop playing, an awful lot of people are positioned and ready to say they told you so.”

Verder hadden we toch echt afgesproken om de CO₂ uitstoot te laten dalen. [Staan handtekeningen van bergen landen onder](#). Over een paar jaar gebruikt AI meer stroom dan Frankrijk ([zegt men](#)), en men [heractiveert nu oude kerncentrales](#) om aan de vraag te voldoen, en dat is kennelijk ook al prima!

Wat dichterbij huis, je kan [geen stroom krijgen voor je nieuwe woning](#). Want er is netcongestie. Maar we willen maar al te graag enorme AI datacenters hier neerzetten, en liefst zo snel mogelijk. Je woont maar even ergens anders is de boodschap.

Ook hadden we bedacht dat [we moesten werken aan digitale autonomie](#), en onszelf minder afhankelijk moesten maken van Amerikaanse big-tech. Maar alle AI-uitrol bij bedrijven en overheden is weer Amerikaans “[surveillance-capitalism](#)”, waarbij men onze data grif gebruikt (en niet om ons leven te verbeteren).

En de meest LLM's zijn [ook nog eens gebaseerd op uitbuiting van goedkoop personeel](#) wat onderbetaald werd.

Ook fascinerend is dat AI-voorstanders voorstellen je hele vloot juniormedewerkers te vervangen door AI, zonder te bedenken waar toekomstige seniors dan vandaan komen.

[En dat is toch wel heel gek allemaal](#). Hele sessies en bijeenkomsten over AI noemen de bovenstaande absurde tegenstellingen niet eens.

Iedereen die nu op de trom slaat voor AI zegt ook “he jammer van je milieu, het intellectueel eigendom, je digitale autonomie, je werknemers, je gebrek aan stroom voor je woning”. Want we moeten kennelijk toch door! Merkwaardig.

AI kan ook waanzinnige dingen

Als we ons even over de klimaatschade, het stroomtekort, en de intellectueel eigendomssituatie etc heenzetten (doet iedereen tenslotte!), dan blijkt dat veel critici van AI er zelf helemaal geen ervaring mee hebben. En dat leidt tot fascinerende botsingen. Iemand die voorheen niet kon programmeren, maar dit weekend wel een AI-tool heeft gevibecode om z'n fotocollectie mee te ordenen, die loopt over van enthousiasme. En ontmoet dan een scepticus die zegt dat het allemaal niet waar is en niet kan. Dat knettert.

Maar AI kan mensen echt een waanzinnige ervaring geven. Iemand die nog nooit iets heeft kunnen programmeren komt maandag terug op kantoor met een best functionele app. Dat is gewoon waanzinnig. Zoiets alsof ik drie pillen zou slikken en ineens sprak ik vloeiend Frans. Zou ik ook enorm enthousiast over worden.



Neo kon ook ineens Kung Fu in The Matrix

Iemand die net zo'n ervaring heeft gehad is overtuigd: dit is waanzinnig! En voor de persoon in kwestie is dat ook zo. Ik wil dit niet bagatelliseren, het is onvoorstelbaar wat iemand zonder ervaring nu bij elkaar kan viben. Dat moet je zien om te geloven.

Tegelijkertijd, **het is niet gezegd dat zo'n app bij nadere bestudering ook onderhoudbaar, compleet, nuttig of veilig is**. Meestal blijkt van niet.

En hiermee komen we bij een centraal probleem - heel veel AI wordt geëvalueerd door mensen die daar niet bekwaam voor zijn. Want wie kan zien of een app goed is? Dat zijn de eindgebruikers en de mensen die dit soort apps in de praktijk aanbieden en beheren. Die kennen de lijken in de kast. Werken de backups? Is de logging zinvol? Begrijpen

gebruikers het ding ook? Kan de helpdesk zien wat er mis gaat, en problemen oplossen? Crasht ie op sommige Samsung of Apple telefoons? Of, zoals met de recente Rabobank app update, raakt de app in de war als je je telefoon kantelt? Is de code leesbaar en te onderhouden? Is de data-opslag AVG-compliant? **Waar staat de data eigenlijk?**

Het vergt ervaring en expertise om te beoordelen of een app goed is. AI maakt het mogelijk voor veel mensen om dingen te doen die ze voorheen nooit konden. Maar helaas besluiten die personen dan ook vaak *zelf* dat wat ze gemaakt hebben goed is. Wat met name een probleem is als het om iemand hoog in de boom gaat.

De tweede uitdaging is dat zo iemand er daarna vaak van overtuigd is dat de *hele* organisatie baat heeft bij deze uitvinding. Maar niemand heeft de echte gebruikers iets gevraagd, is dit wat er nodig is? De organisatie heeft mogelijk hele andere problemen dan het gebrek aan deze app. Overigens vergeten ook niet-vibecodende developers vaak dit uit te zoeken.

De samenvatting is dat mensen en met name bestuurders met persoonlijk hele goede ervaringen niet geschikt zijn om namens hun organisatie te bepalen of dit goed was. Hoe goed het ook voelt allemaal persoonlijk. Dat laatste geldt ook voor mensen die zelf ChatGPT *individueel* toepassen voor werk - het voelt mogelijk fijn om rapporten door de AI te laten maken en lezen, maar wordt de organisatie er ook beter van? Mogelijk niet jouw probleem, maar wel dat van bestuurders. Ik kom hier verderop op terug.

Een enorm probleem is dat [in de meeste omgevingen niet \(meer\) geëvalueerd wordt of AI een succes is of kan worden](#). Het is de toekomst namelijk! En AI gaat inderdaad nooit meer weg, maar goed nadenken wat er mee te doen zou toch wel fijn zijn. Of op z'n minst vooraf bepalen wat je er van verwacht, en zinvol te meten of het ook gelukt is. Want “voelde wel goed” is niet genoeg. Ook hier later meer over.

Desondanks, AI gaat nooit meer weg

Er is een rijke geschiedenis van nieuwe technologieën waar men initieel niet van wist wat er mee te doen. De [laserstraal is inmiddels niet meer weg te denken](#) uit ons leven, en toch was de laser decennia bekend als “[een oplossing op zoek naar een probleem](#)”.

Het internet werd ook lacherig weggezet, alsof iemand ooit schoenen of pizza's ging kopen online “haha”. Moet je nu eens zien.

Maar, niet alles is een blijver. We horen zelden meer van de blockchain bijvoorbeeld, terwijl die toch alle problemen op ging lossen. Ook Second Life, de Metaverse en de VR brillen zijn het niet geworden.

Toch heeft AI er alle schijn van meer weg te hebben van de laserstraal dan van een cryptomunt. Die Nobelprijs was geen grap namelijk. En je kan ook chatten met AI en hem vragen documenten voor je te maken of bewerken. Ik weet niet of dit laatste nou de uiteindelijk nuttigste toepassing van AI zal zijn, overigens.

Grote rapporten maken met AI die dan door anderen weer gelezen gaan worden met AI komt nogal “[horseless carriage](#)” op me over. De eerste motor-aangedreven voertuigen leken meer op “paard en wagens zonder paard” dan op moderne auto's. Men had tijd nodig om los te komen van de oude werkelijkheid. Ik gok dat dit ook voor AI geldt - mensen gebruiken het nu om hun bestaande fenomenen ('documenten die toch niemand leest') mee in stand te houden.

Ik gok dat we nu echt niet weten wat voor een dingen we over 10 jaar met AI doen. Maar het zal zeker niet niks zijn. Maar ook vrijwel zeker niet “paard en wagen zonder paard”.

Overigens, ook als de AI-bubbel barst [verdwijnt AI daarmee niet spontaan](#). Wel wordt het ineens een heel stuk duurder, en valt een hoop van de druk AI te gebruiken (tijdelijk) weg. Dat duurder worden [is overigens al begonnen](#).

Naar wie te luisteren?

Dat valt ook nog niet mee. De AI-investeringen bedragen nu biljoenen dollars. Hele industrieën, *inclusief media*, zijn meegezogen in deze duizenden miljarden. Het succes van AI is voor veel consultancies en leveranciers inmiddels van *existentieel* belang.

Men loopt over van enthousiasme, maar tegelijk, ze moeten ook wel. De industriële belangen zijn zo groot dat het onmogelijk is geworden om de AI-wereld zelf te gebruiken als bron van informatie. Alsof je luistert naar een vrijwel failliete autohandelaar die jou zijn laatste auto aanprijst. Nogal wiesde dat ‘ie enthousiast is. En, tenzij er heel rare dingen gebeuren, [gaan veel van de AI spelers ook failliet](#). Dus de analogie werkt wel.

Ook media en opiniemakers zijn een heel end meegegaan in de verhaallijn dat AI gewoon onvermijdelijk is. Hele “nieuwssites” staan inmiddels vol AI-slop, dus in ieder geval hun aandeelhouders geloven er stellig in.

Een Nederlands ministerie hield recent een presentatie waarin men zei dat voor ambtenaren “AI gebruik nu verplicht is” (als je het maar goed doordacht doet). Maar het is kennelijk verplicht! Ook de coalitie heeft al volmondig en met ongebreideld enthousiasme voor AI gekozen zoals [we ook lazen in het coalitieakkoord](#).

Dan zijn er natuurlijk mensen die overal en altijd de nuance zoeken, maar die komen veel met middenwegen waar je niet zoveel aan hebt: AI is gewoon een stuk gereedschap, en alles hangt af van hoe je het inzet. Dat is natuurlijk altijd waar voor alles. Maar gaan mensen het ook verstandig inzetten? Zouden ze weten hoe? Leiden we ze daar toe op? Daar hebben de genuanceerden het dan niet snel over. Ook meent men dat het met goede regulering best zou kunnen, zonder daarbij te vermelden dat overheden totaal machteloos blijken op dit vlak, en niet eens zin lijken te hebben om gedegen regulering te maken laat staan te handhaven.

Wat we ook niet moeten vergeten is hoe het huidige LLM-gebruik (tekst, plaatjes, video, programma-code) overkomt op iedereen die er eer en lol in heeft dat soort dingen met eigen hersenen te maken. Niet alleen denken wij het beter te kunnen dan de AI, het is **intens** irritant hoe allerhande mensen ons confronteren met hun LLM-creaties, waar wij dan op moeten reageren. Want de context is duidelijk, “ik kan nu met mijn vriend Claude wat tot nu toe jouw baan was”. Uw dagen zijn geteld! Misschien bedoelt de “[Claudeviool](#)” het niet zo, maar zo komt het zeker over.

Het helpt ook niet dat ons verteld wordt dat we wel mee moeten omdat ze anders achtergelaten zullen worden. Mensen die nog nooit een app geshipped hebben vertellen me dat mijn tijd als programmeur over is. Dat het halve Nederlandse internet draait op de software die ik gemaakt heb is iets van vroeger en telt niet meer. Ik word uitgemaakt voor een soort tragisch fossiel wat niet meer mee kan.

Als je hele levenswerk zo belachelijk gemaakt wordt, en je relevantie zo in twijfel wordt getrokken is het extreem moeilijk om nog een beetje afgewogen meningen over AI te hebben (maar ik doe mijn best!). En de resulterende geheel begrijpelijke weerzin tegen alles wat met AI te maken heeft wordt door AI-fans als zelfstandig argument gebruikt. “Naar deze verzuurde mensen hoef je dus niet te luisteren”.

En pijnlijk genoeg is dat deels ook waar - mensen die (dus terecht) zo verontrust zijn over AI komen vaak met een niet aflatende stroom aan slecht nieuws over AI, of het nou waar is of niet.

Kortom, om nu een gefundeerde mening over AI te vormen ben je al snel op jezelf aangewezen. Want de industrie, de media, de fans en de haters geven allemaal inconsistente input waar geen chocola van te maken is. En je eigen succesvolle AI-ervaring is ook niet genoeg, die vertroebelt het beeld te veel. En persoonlijke aversie tegen AI leidt ook niet tot evenwichtige evaluaties. Beide kanten op geldt: een gevoel is nog geen beleid!

De situatie wordt overigens goed omschreven in dit citaat:

“In [his essay on Salvador Dali](#), Orwell argued that because Dali was a repulsive human being, the right wouldn’t admit that he was a great artist; conversely, because he was a great artist, the left wouldn’t admit he was a repulsive human being. I’m seeing the same thing with AI: because it’s unethical, one side won’t acknowledge that it’s useful, but because it’s useful, the other side won’t acknowledge that it’s unethical.” – [Greg Wilson](#)

AI-experimenten op kantoor en school

Maar, de druk is enorm, dus iedereen rolt van alles uit. [Want het moet gewoon, ook al is er geen plan](#).

Als je een AI experiment doet, **meet dan wel of het een succes is**. En om dat te doen moet je vooraf bepalen wat je zou willen bereiken, en hoe je dat gaat meten dan.

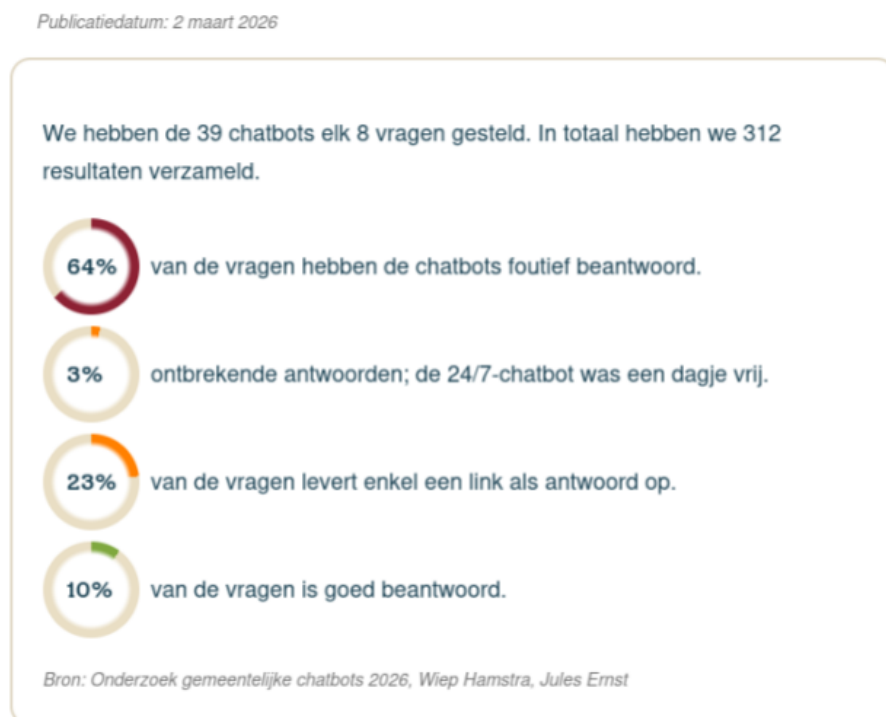
Het enthousiasme is zo groot dat menig AI experiment begint zonder hier een plan over te hebben. Wat dan leidt tot resultaten als “de samenvattingen zien er goed uit”, of “24% van de medewerkers deed mee aan de proef”. Of zelfs dat niet.

Betere metrieke zijn, “75% van de toezeggingen van de wethouder stonden in de AI samenvatting”, of “25% van de bellers meldde een fout in het verslag van het gesprek”.

Het toetsen is vanwege de enorme verwachtingen en druk overigens niet iets wat je er “even bij doet”. Er zijn experts in werkplekefficiëntie, en die zal je moeten betrekken voor men op basis van *vibes* besluit dat de AI de bom is.

Het is verbazingwekkend hoeveel AI-proeven beginnen zonder vooraf te besluiten wat het doel nou is. En dan kan de proef mogelijk afgesloten worden met een “zinloos taartmoment”, zonder dat men weet wat er eigenlijk bereikt is. Misschien is het nu wel erger dan eerst.

Als voorbeeld, hier de resultaten van onderzoek naar gemeentelijke chatbots:



[Bron](#)

Dit had eenvoudig voorkomen kunnen worden.

Een ander probleem is dat directies vaak met groot enthousiasme de AI-proef aankondigen, en zeggen er veel van te

verwachten. En na de proef zijn mensen ook enthousiast - want ze zagen de directie het goede voorbeeld geven!

Als je echt wil weten wat er op de werkvloer leeft, **wat overigens altijd al een uitdaging was**, zal je een geloofwaardige anonieme enquête moeten doen. En ook na moeten denken over tweede orde effecten.

Maar, begin dus geen AI uitrol zonder te definiëren hoe succes eruit ziet, wat men er van verwacht, en hoe dat ook gemeten gaat worden. Want anders is het al snel een “succes”.

Wat willen we eigenlijk?

Nou, AI, kennelijk! En rap wat ook. Toch is het goed stil te staan wat een organisatie er eigenlijk mee wil bereiken. Een chatbot is geen doel op zich. Het sneller, goedkoper en beter (maar nog steeds correct) helpen van burgers en klanten wel.

Ook is het inderdaad mogelijk om met AI in recordtempo enorme documenten en rapportages te maken. Vervolgens kunnen andere mensen in de organisatie deze documentenstroom met AI weer gaan samenvatten om het te begrijpen. Dit is geen vooruitgang, zelfs als er geen informatie verloren of gehallucineerd wordt in dit proces, wat wel gebeurt.

Het is vaak opgemerkt dat schrijven en (productief) denken een intieme relatie hebben. “[Writing has been called the process by which you find out you don’t know what you are talking about](#)”, en ook “[Writing is thinking](#)”. Ook leuk, “[If you’re thinking without writing, you only think you’re thinking](#)”.

Veel (interne) documenten worden uiteindelijk niet zoveel gelezen. Maar ze worden wel gepresenteerd, en mogelijk uitgevoerd. Het schrijven van die documenten was tot nu toe een integraal onderdeel van het denkproces van een organisatie. Het formuleren van die documenten maakte de presentatie en plannen echt beter. En zelfs als een document niet goed gelezen werd, als iemand een substantieel stuk over een idee had gemaakt was het wel een sterke indicatie dat men er echt moeite in had gestoken.

Maar met AI is het mogelijk oppervlakkig stevig overkomende documenten te maken zonder al te veel nadenken of moeite. Zo’n document is geen signaal meer dat iemand het goed doordacht heeft of er veel werk in heeft gestoken. [En misschien is het ook niet zo best eigenlijk](#).

Ik heb toch het idee dat veel organisaties er niet op vooruitgaan als er minder (goed) nagedacht wordt, of als de “heb je er moeite ingestoken” toets niet meer werkt.

Dit nog los van de schade als je (publieke) beleidsdocument vol gehallucineerde niet bestaande verwijzingen blijkt te staan - iets wat zelfs consultancies die schrijven over AI [steeds gebeurt, in hun stukken over AI](#) ([reservelink](#))! **Want we moeten niet vergeten dat zelfs de beste AI regelmatig de grootst mogelijke onzin genereert**.

Het loont dus om heel goed stil te staan wat men eigenlijk wil met AI voor te beginnen met de uitrol.

Ik woonde een presentatie bij van mensen die miljoenen nieuw geschreven contracten met AI uitlazen zodat de beschreven transacties in een database opgeslagen konden worden. Dat is een best rare toepassing van AI in 2026, want waarom worden die transacties niet digitaal in die database gezet, in plaats van dat we ze moeten afleiden uit een door mensenhanden opgestelde contracten? Of mogelijk zelfs contracten die ook weer met AI gemaakt zijn. Dit had veel beter direct gekund.

Tegelijkertijd, als je miljoenen *historische* contracten hebt die niet in de database staan is het niet “beter” als je die door mensen laat lezen om de data te copy pasten. Is voor die mensen ook geen leuk werk namelijk. Het verifiëren of de AI het goed heeft gedaan is al pittig genoeg.

Een ander mooi voorbeeld wat ik hoorde was een stuk overheid dat een proces met AI wilde doen, maar snel

constateerde dat de inputdata daarvoor te rommelig was. Toen dat was opgelost liep het proces zoveel beter dat er geen AI-versie meer nodig was. Nou is dit natuurlijk niet altijd zo, maar weet dat AI ook geen magie is, en goede input vergt. Veel bedrijven ontdekken dit als ze een chatbot proberen te bouwen die kennis heeft over bedrijfsprocessen. Als die nergens (correct) gedocumenteerd zijn kan de chatbot ook geen wonderen verrichten. En dan ligt hallucinatie op de loer. “AI doen om de AI” is een dure hobby, en ook riskant. Want er kan een hoop misgaan. Daar mag je best over nadenken.

Wat is het juiste tempo?

Men ervaart nu grote druk om “met AI aan de slag te gaan”. Want wie wil er nou niet een moderne organisatie zijn?

Toch is het goed om al die haast wat te evalueren. Eerst de nadelen van het langzaam aan doen:

Je loopt mogelijk kostenbesparingen mis

Er kan personeel zijn dat AI gewoon verwacht

Er is natuurlijk ook een soort algemeen gevoel van “straks missen we de boot”. Maar veel organisaties (overheden bijvoorbeeld) [zijn niet verwikkeld in een keiharde strijd met concurrenten](#). Die boot is er volgend jaar ook nog wel. De AI gaat niet op ineens!. Ook lijken organisaties in ieder geval [tot nu toe geen productiviteitsgroei te boeken met AI](#). Mogelijk komt dat nog wel, maar je mist nu nog niet veel.

Punt 2 is wel een relevante - iemand die al jaren documenten maakt en analyseert met AI is niet noodzakelijk in staat [dat nog met de hand te kunnen](#). Ik schrok hier ook van, maar iemand van een grote afdeling personeelszaken vertelde me over dit probleem. Er is nu een klasse personeel ontstaan die zonder AI-hulp niet meer functioneert. Eigen vaardigheden verlopen snel door [cognitive offloading](#).

Te snel gaan heeft ook nadelen:

Schending klimaatakkoorden, dubieuze intellectueel eigendomssituatie, de ethische problemen

Er kan personeel zijn dat vanwege AI vertrekt

Verslechtering dienstverlening

Schending privacy en rechten burgers/klanten/contacten

Weer verdere afbreuk digitale autonomie

Wanneer de AI-bubbel knapt wordt alles onbetaalbaar

Vermindering denkvermogen organisatie

Verstoring functiegebouw (“alle juniors zijn weg, hoe komen we ooit nog aan seniors”)

Over punt 1, vanuit een directie bezien werken dingen heel anders dan op werkvloeren, waar mensen hun directeur horen roepen dat hun inzet binnenkort gelukkig geautomatiseerd kan worden. Loopt niet noodzakelijk goed af.

Over punt 2, zie de chatbots boven. Het wordt vaak niet beter, ook al is het nu “modern”.

Over punt 3, het blijkt dat AI-experimenten regelmatig [bergen persoonsgegevens lekken](#). Zoals recent bij Meta (Facebook) bij een proef [op eigen personeel bijvoorbeeld](#).

Over de rechten van je contacten, je zou het bijna vergeten, maar we hebben de AI Act en de AVG (GDPR). [En die zeggen dingen over hoe je sollicitanten met AI mag werven/afwijzen](#). En voor je het weet zit iemand bij personeelszaken dit te doen, zonder dat je het wist, maar ben je wel in overtreding. Ondanks wat men veel zegt [treden hele dure stukken](#)

[van de AI act toch echt dit jaar in werking.](#)

Over punt 4, we proberen juist weg te komen van Amerikaanse overheersing. Maar de overgrote meerderheid van AI op de werkvloer is van big tech, en dan ook nog eens big tech die graag alles van ons weet en opslaat en gebruikt voor verdere AI training.

Over punt 5, zoals eerder gezegd zijn de financiën van de diverse bedrijven al heel dubieus. Als dat kaartenhuis ooit instort zal er toch echt contant afgerekend gaan moeten worden voor alles, wat knap duur uit kan pakken. Die trend [is al gaande](#).

Over vermindering van het denkvermogen van de organisatie schreef ik bovenstaand. Opschrijven van plannen is op zichzelf al een goede manier om beter na te denken. En als je de AI je stukken laat schrijven loop je dat nuttige effect mis.

Over het laatste punt, die komt in het volgende blokje:

Het functiegebouw

Als we juniorwerk overlaten aan AI rijst wel de vraag [waar seniors ooit nog vandaan gaan komen](#). Ik heb eerlijk gezegd geen idee. Als organisatie met een functiegebouw is het wel nuttig hier een beeld bij te hebben. Schade aan de opbouw van het personeelsbestand kost jaren om te compenseren.

De suggestie is ook dat AI veel werk over kan nemen, maar dat je het nog wel moet controleren. De baan van medewerkers verandert dan fundamenteel, en het blijkt dat mensen niet erg geschikt zijn om als voornaamste activiteit “controleren van de computer” te hebben. Dit kan je je bestaande seniors kosten, die niet zitten te wachten op zo’n baan. En de juniors waren ook niet meer nodig, want hun werk was dan vervangen door AI.

Het bovenstaande klinkt misschien erg dramatisch, maar het is wel wat er aan het gebeuren is.

Als men roept, zoals het kabinet, “[we zetten vol in op AI](#)”, dan moet men wel nadenken over de bredere gevolgen. Hebben we nog een functionerende organisatie over een paar jaar?

Maar het komt altijd goed!

Als je je zorgen maakt over waar nieuwe seniormedewerkers vandaan moeten komen komen fans van AI altijd met dat zoiets altijd wel goedkomt, er komen altijd nieuwe banen, dingen die vroeger niet bestonden. “Het vindt zijn weg wel”. Men noemt dan de industriële revolutie, maar meldt er niet bij dat het [meerdere generaties leed kostte](#) voor het “goed kwam”, [en in de tussentijd was het afgrijselijk](#)

Ook komen mensen met onbewezen stellingen dat ze nu *meer* tijd hebben om na te denken, want AI doet het schrijfwerk. Anderen beweren dat de AI misschien nu nog niet met nieuwe ideeën komt, maar dat we dat in de toekomst vast gaan krijgen, en het daarom nu al OK is om het denkwerk aan AI over te laten. Dat zijn mooie stellingen, maar het is wel wat goedgelovig mogelijk. Kom met bewijs.

Het is niet gek om eerst goed na te denken voor je je bedrijf omploegt met AI, ook al komen mensen met (onjuiste) verwijzingen naar de industriële revolutie om aan te geven dat zoiets altijd goed is gekomen.

En nu?

Het is op zich vrij duidelijk in juli 2026. De hype over AI is groots, maar waar het heengaat is knap onduidelijk. Ook zijn er juridische, ethische, milieukundige en financiële redenen om er heel kritisch naar te kijken. Ook kunnen er [ongewenste en ernstige effecten](#) zijn op het personeelsbestand & de privacy van burgers or klanten.

Vaak wordt dit soort kritiek afgedaan als “ja dat komt van iemand die AI haat”, maar dat is te makkelijk. We praten over mogelijk de grootste verandering op de werkvloer ooit, waarmee mensen onvoorstelbaar hard van stapel lopen. Er zijn ook objectief genoeg problemen te voorzien (zie boven). Daar mag je best goed over nadenken, waarbij je kritiek dus niet negeert.

Maar, er is dus voor de meeste organisaties helemaal geen haast. Als je niet in 2026 Copilot uitrolt heb je in 2027 echt nog de kans. En tegen die tijd zijn een hoop van de onzekerheden een heel stuk kleiner. De rechtszaken zijn misschien verder, diverse zwakke broeders in de AI-scene zijn inmiddels failliet, de prijzen zijn vermoedelijk een stuk hoger maar wel realistischer. De impact van de AI-act is ook duidelijker, wat weer boetes scheelt.

En misschien nog wel relevanter, we kunnen ondertussen een hoop leren van de clubs die het hardst van stapel gelopen zijn: hebben ze er echt wat aan gehad? Of heeft men ineens [allemaal problemen](#).

En in de tussentijd, als een organisatie met AI aan de slag gaat, maak dan niet de fout die iedereen nu maakt. Leg vooraf vast wat de hoop is van het AI-experiment, en hoe we concreet gaan meten of het een succes geworden is. Dat het wel goed voelde is een wat dunne onderbouwing om je hele instelling ingrijpend te gaan veranderen. Met goed testen ga je in ieder geval niet nat als het achteraf toch niet zinvol bleek.

En nogmaals, vergeet niet de fundamentele problemen met de AI-bedrijven en hun werkwijzen. Daar kies je ook voor als je AI actief in gaat zetten. Wel zo fair om dat op z'n minst te erkennen.

Succes er mee! Want weg gaat het nooit meer, maar hoe het afloopt zijn we zelf bij.

PS: Reacties op dit stuk zijn zeer welkom, zowel positief als negatief. Misschien mis ik wel dingen namelijk. Weet me te vinden op bert@hubertnet.nl

Verder lezen

[AI Mania Is Eviscerating Global Decision-Making](#)

[AI: Sowieso ontwrichtend](#)

[The AI-collapse pre-mortem](#)

Mijn eigen introductie in Deep Learning: [Hello Deep learning](#)