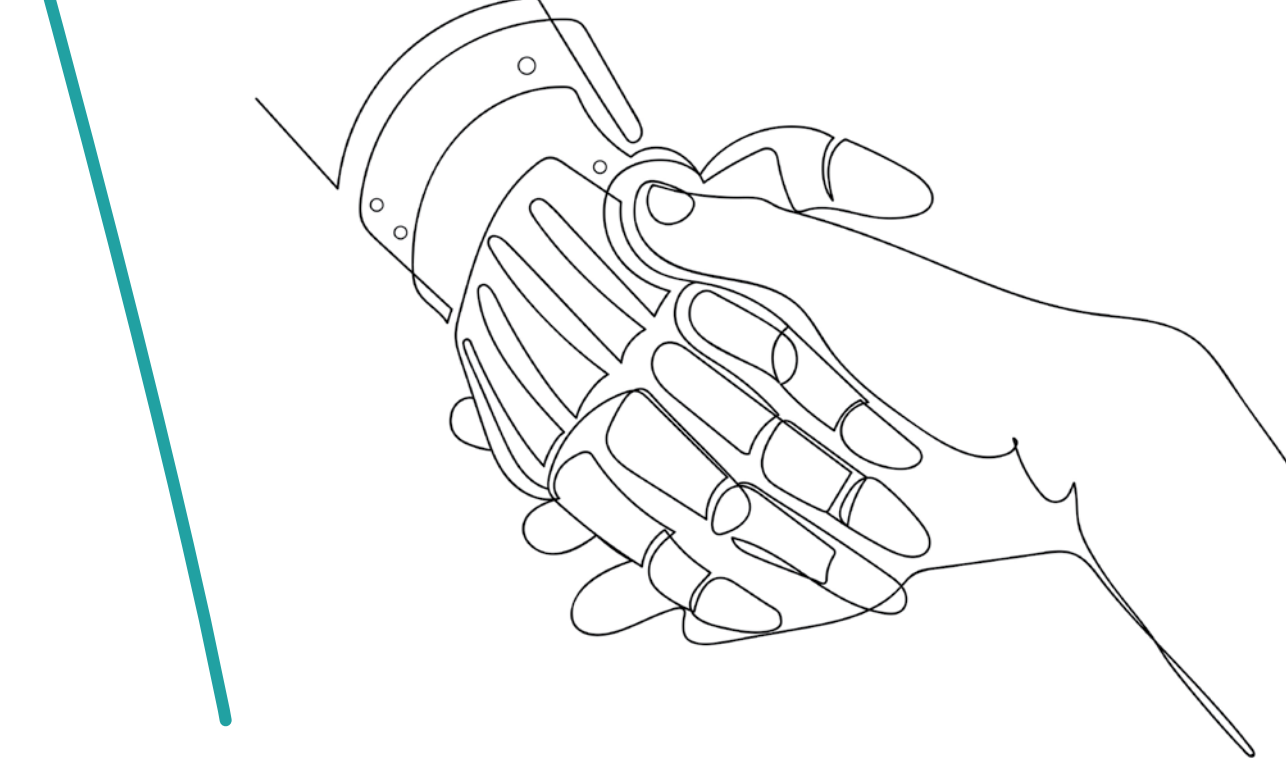




Waarom angst een slechte raadgever is in het AI-debat

Auteur: **Barend Last**

"Ze gaan ons allemaal vervangen!", riep een deelnemer tijdens een van mijn workshops over kunstmatige intelligentie. Zijn gezicht sprak boekdelen: pure angst. En hij is niet alleen. De afgelopen maanden hoorde ik in scholen, bedrijven, zorginstellingen en conferentiezalen dezelfde bezorgde stemmen. De ene helft van de professionals voorspelt een utopisch wonderland waarin AI alle administratieve rompslomp wegtovert.

De andere helft schildert een dystopisch landschap waarin mensen overbodig worden en werknemers alleen nog maar prompttrucjes leren.

Ik snap die emotionele reacties wel. Grote veranderingen roepen nu eenmaal gemengde gevoelens op. Maar wat me opvalt: angst verlamt – en dat kunnen we nu echt niet gebruiken.

Angst is een slechte raadgever en maakt ons letterlijk dom. Bij angstige mensen zien neurowetenschappers dat de prefrontale cortex – het deel van ons brein dat verantwoordelijk is voor complexe afwegingen – minder actief is. We schakelen terug naar binaire keuzes: vechten of vluchten.

Vertaald naar het AI-debat houdt dat in: verbieden of ongecontroleerd toelaten. Zo blokkeren sommige schoolbesturen ChatGPT compleet op hun netwerk. Andere organisaties rollen zonder enige voorbereiding AI-tools uit in hun werkprocessen. Beide zijn paniecreacties, geen verstandig beleid.

Nog problematischer: angst leidt tot wat psychologen *'solution aversion'* noemen. Bang voor een probleem, verwerpen mensen ook alle mogelijke oplossingen die ermee geassocieerd worden. Frame AI als existentiële bedreiging, en mensen keren zich ook af van AI-geletterdheid, van experimenteren met AI-tools, of van investeren in AI-onderzoek. Precies het tegenovergestelde van wat we nodig hebben.

Verantwoordelijkheid

Waarom zijn we zo bang? Deels omdat we AI antropomorfiseren: vermensenlijken. "AI gaat ons overnemen", hoor je dan. "AI leert zichzelf te verbeteren." "AI begint eigen doelen te stellen." Dit taalgebruik suggereert een wezen met eigen wil, intentie en autonomie. Maar dat is een fundamentele denkfout. AI denkt niet. AI wil niets. AI leert niet in menselijke zin. Het is geen bewustzijn. Het genereert statistisch waarschijnlijke output op basis van trainingsdata. Niets meer, niets minder.

Waarom is deze denkfout zo relevant? Om de verantwoordelijkheid op de juiste plek te leggen. Als we geloven dat AI 'zichzelf verbetert', dan lijken we machteloos; een autonoom proces dat ons overspoelt. Als we inzien dat AI-systemen worden verbeterd door mensen, met menselijke agenda's en belangen, dan kunnen we eisen stellen aan transparantie, controle en verantwoording.

Je hebt geen sciencefiction nodig om wakker te liggen van AI.

Daarom moeten we stoppen met formuleringen als 'AI schrijft', 'AI leert' of 'AI vindt'. Feitelijk zetten mensen AI in om tekst te genereren, patronen te herkennen, of beslissingen te ondersteunen – maar altijd onder menselijke regie, via menselijke infrastructuur, voor menselijke doelen.

Functionele autonomie

AI-systemen kunnen echter wel degelijk functioneel autonoom opereren, met echte impact. Ze kunnen binnen vooraf gedefinieerde kaders zelfstandig handelingen uitvoeren zonder directe menselijke tussenkomst. Een recruitmentsysteem kan kandidaten selecteren, een monitoringsysteem kan medische afwijkingen detecteren, een voorspellingsmodel kan kredietrisico's inschatten. Die functionele autonomie leidt tot reële effecten op mensenlevens – mensen die een baan niet krijgen, die een behandeling ontvangen, of die een lening geweigerd wordt.

Maar die functionele autonomie is fundamenteel anders dan de intentionele autonomie van mensen. AI-systemen kunnen niet buiten hun geprogrammeerde grenzen treden, kunnen geen eigen waarden of motieven ontwikkelen, en hebben geen bewuste ervaring van wat ze doen. Het zijn instrumenten, geen actoren. Wanneer we dit onderscheid uit het oog verliezen, beginnen we de verkeerde vragen te stellen.

Scenariodenen

Een andere hindernis voor goed beleid: we denken snel in enkelvoudige doemscenario's in plaats van meervoudige toekomsten. Ik schrok van een recente presentatie bij Eva Jinek waar Alexander Klöpping een dramatisch AI-scenario presenteerde: binnen twee jaar zou AI ons allemaal overtreffen, en dan...

Zo denk je niet serieus over de toekomst. De kracht van toekomstverkenning zit juist in meervoudigheid. Je onderzoekt verschillende richtingen, analyseert signalen, trends en drijvende krachten. Je kijkt wat wenselijk is en wat waarschijnlijk is, en op basis daarvan anticipeer je in de richting van het meeste wenselijke. Dat is constructief toekomstdenken. Geen sciencefiction met een agenda.

Urgentie in het nu

Het ironische is: je hebt geen sciencefiction nodig om wakker te liggen van AI. We hebben nu al te maken met urgente vraagstukken: Hoe waarborgen we digitale soevereiniteit wanneer strategische AI-infrastructuur grotendeels in handen is van enkele Amerikaanse en Chinese techgiganten? Hoe beschermen we onze democratische processen

tegen synthetische misinformatie die met steeds geavanceerdere AI-tools wordt geproduceerd? Welke nieuwe geopolitieke machtsverhoudingen ontstaan er wanneer landen met toegang tot grote taalmodellen en datacentra een voorsprong krijgen op landen zonder deze middelen? Hoe voorkomen we dat AI-systemen in kritieke sectoren zoals energie, financiën en logistiek een cyberveiligheidsrisico vormen voor onze nationale infrastructuur? Wat betekent de opkomst van AI voor de positie van Europa in de mondiale technologierace, en hoe voorkomen we dat we afhankelijk worden van buitenlandse AI-oplossingen?

Dit zijn concrete, actuele uitdagingen die om beleid vragen. Ze zijn complex en urgent genoeg om mensen wakker te houden zonder dat we hoeven te fantaseren over 'superintelligentie'. Waarom zouden we energie verspillen aan speculatie als deze concrete uitdagingen voor het oprapen liggen?

*De toekomst overkomt
ons niet. De toekomst
maken we samen.*

Gewone technologie

Ik pleit niet voor naïef optimisme. Integendeel. Ik pleit voor wat ik 'waakzame betrokkenheid' noem: alert zijn op risico's, maar tegelijk actief meedenken over oplossingen.

Laten we AI zien voor wat het is: een uiterst krachtige systeemtechnologie, maar ook niets meer dan dat. Net zoals

elektriciteit, de stoommachine of het internet heeft AI het potentieel om onze maatschappij ingrijpend te veranderen. Maar net als die eerdere technologieën blijft AI beheersbaar door verstandig beleid, ethische kaders en geleidelijke integratie.

De komende jaren bepalen we samen hoe AI onze wereld gaat vormen. Niet door ons

te laten verlammen door angst of door te vluchten in sciencefiction. Maar door met beide benen op de grond te blijven staan, elkaars zorgen serieus te nemen, en door de mouwen op te stropen. Door te experimenteren, te leren en bij te sturen.

De toekomst overkomt ons niet. De toekomst maken we samen. ♦



Barend Last (1986) begon als leerkracht in het basisonderwijs, waar zijn passie voor onderwijsinnovatie ontstond. Hij interesseerde zich vooral voor de vraag: 'Waarom doen we eigenlijk wat we doen?' en stroomde daarom door naar de wetenschappelijke wereld. Hij werkte onder meer als docent, manager en onderwijspecialist in verschillende lagen van het onderwijs. Inmiddels is hij vooral actief als schrijver, spreker en onderwijsmaker, en staat hij weer een dag per week als leraar voor de klas in het basisonderwijs. www.barendlast.com