



Kritisch denken in het AI tijdperk

Wie zegt dat?



Datum: Vrijdag 22 en zaterdag 23 januari 2027

Prijs: €350,-*

Doelgroep: leraren, begaafdheidsspecialisten,
onderwijsvernieuwers

Locatie:

Buitenplaats de Bergse Bossen

Traaij 299, 3971 GM Driebergen

Is deze training iets voor jou?

Je bent geschikt voor deze training als je AI niet wil verbieden en niet blind wil omarmen, maar er bewust mee wil leren werken. Voor jezelf én voor de kinderen in de klas. Als je behoefte voelt om juist nu extra aandacht te besteden aan het inschattings- en beoordelingsvermogen van de kinderen omdat de wereld waarin zij volwassen gaan zijn, zulke complexe vraagstukken zal blootleggen, dat er veel goede kritische denkers nodig zullen zijn. Je voelt dat de tijd rijp is voor focus op denkvaardigheden, waar nieuwsgierigheid en eigenaarschap in het leerproces een voorwaarde is voor een goede ontwikkeling bij kinderen. Je zoekt inspiratie, kennis en praktische tools om in jouw praktijk aan de slag te gaan met 'toekomstbestendig' onderwijs. A rising tide raises all ships.

* De prijs van de training is bewust laag om de drempel voor deelname zo laag mogelijk te houden.



1. Waarom deze training

De kerngedachte: Het doel is niet wantrouwen maar afstemmen. Goedgegeloofvigheid en achterdocht zijn allebei denkfouten. Een kind dat alles gelooft wat het scherm zegt, denkt niet. Een kind dat alles afwijst wat het scherm zegt, denkt ook niet. Het Intellectueel Kompas is precies het instrument dat het verschil kan maken, omdat het niet vraagt of je iets gelooft maar hóé je dat beoordeelt.

2. Visie en theoretisch fundament

2.1 Het Intellectueel Kompas als basis

Paul & Elder blijven de ruggengraat. De zes kernvaardigheden (analyseren, interpreteren, evalueren, concluderen, uitleggen, zelfregulatie) sluiten aan bij het consensusmodel van Facione. De drie pijlers zijn: de Bouwstenen (doel, vraag, informatie, interpretatie, concepten, vooronderstellingen, consequenties, perspectief), de Meetlat (duidelijk, accuraat, precies, relevant, diepgaand, breed, logisch, betekenisvol, eerlijk) en het Karakter (intellectuele bescheidenheid, moed, empathie, autonomie, integriteit, doorzettingsvermogen, vertrouwen in de rede).

Dat karakterdeel is in deze training belangrijker geworden dan het was. Kritisch denken faalt zelden op techniek en bijna altijd op houding: de moed ontbreekt om iets dat goed klinkt tegen te spreken, of de bescheidenheid om te erkennen dat je het onderwerp eigenlijk niet kent. Beide worden door een vloeiende machine actief ondermijnd.

2.2 Wat AI met kritisch denken doet: huidige inzichten

- **Automation bias** is oud, hardnekkig en niet weg te praten. De neiging om automatische output over te nemen en tegenstrijdige signalen te negeren is sinds Parasuraman & Riley (1997) uitgebreid gedocumenteerd. Een meta-analyse van klinische beslissingsondersteuning (Goddard e.a., 2012) vond een stijging van ongeveer 26 procent in foute beslissingen wanneer het systeem zelf fout zat. Belangrijkste les voor onderwijs: goedgegeloofvigheid verdwijnt niet door mensen te waarschuwen (Parasuraman & Manzey, 2010). Je hebt structuur nodig, geen goede voornemens: werkvormen die controleren verplicht maken, niet een zin vooraf dat AI ook fouten maakt.
- **Niet-experts vertrouwen te veel, experts te weinig.** Waar deskundigen algoritmen juist mijden, nemen niet-deskundigen ze te makkelijk over (Hou & Jung, 2021; Logg e.a., 2019; Dietvorst e.a., 2015). Kinderen zijn per definitie niet-experts in bijna elk schoolonderwerp. Zij zitten dus structureel aan de kant van het goedgegeloofvigheid, en ze missen de domeinkennis die nodig is om de output te beoordelen.
- **Vleierij: het model kiest jouw kant.** Onderzoek naar sycophancy (Sharma e.a., ICLR 2024) liet zien dat de onderzochte modellen systematisch antwoorden verkiezen die aansluiten bij de opvatting van de gebruiker, ook als die opvatting niet klopt. Latere studies (CHI 2026) bevestigen dat dit misvattingen bekrachtigt en kritisch denken ondermijnt. Voor de klas is dit het scherpst denkbare punt: de vraag 'klopt mijn idee?' levert bij een mens weerstand op en bij een machine bevestiging.
- **Meningen schuiven mee.** Het International AI Safety Report 2026 beschrijft een studie waarin het gebruik van een chatbot bij schrijven niet alleen de meningen in de tekst verschoof richting het model, maar ook de eigen meningen van de schrijver. Dat maakt intellectuele autonomie geen abstract begrip meer maar een concreet risico.
- **Hallucinaties zijn geen bug die verdwijnt.** Het genereren van plausibele maar onjuiste inhoud volgt uit de manier waarop deze modellen werken en getraind zijn; het is geen storing die met de volgende versie is opgelost. Een verzonnen bron ziet er precies zo uit als een echte bron. Dat vraagt om een controlegewoonte, niet om een betere tool.
- **Kritisch denken en AI-gebruik hangen samen.** Onder kenniswerkers geldt: hoe meer vertrouwen in AI, hoe minder kritisch denken; hoe meer vertrouwen in de eigen vaardigheid, hoe meer kritisch denken (Lee e.a., 2025). Gerlich (2025) vond een negatief verband tussen frequent AI-gebruik en kritisch denken, gemedieerd door cognitive offloading, en het sterkst bij jongeren.

If it walks like a duck... it's a zebra? Al deze bevindingen komen samen in één mechanisme dat direct aansluit op de Meetlat. Mensen gebruiken vlotheid van taal als aanwijzing voor waarheid: wat makkelijk leest, voelt juister. AI-taal is maximaal duidelijk, goed geordend en foutloos geformuleerd. Daarmee scoort ze hoog op één intellectuele norm, duidelijkheid, en dat wekt ten onrechte de indruk dat ook de normen accuraat, precies, diepgaand en eerlijk zijn gehaald. Paul & Elder wisten dit al: duidelijkheid is een noodzakelijke maar volstrekt onvoldoende norm. De tweedaagse training maakt daar een expliciet lesonderdeel van.

2.3 Afstemmen als doel

Uit de literatuur over mens-AI-samenwerking komt het begrip appropriate reliance: vertrouwen dat past bij wat het systeem werkelijk kan. Er zijn twee symmetrische faalvormen: goedgegeloofvigheid (misuse) en het afwijzen van hulp die wel had geholpen (disuse). Voor het onderwijs betekent dit dat 'wees kritisch op AI' een te ruime boodschap is. De vraag is per situatie: hoe betrouwbaar is dit systeem voor deze taak, hoe erg is een fout hier, en hoe kan ik het controleren? Een spellingsuggestie, een samenvatting van een tekst die je zelf hebt gelezen en een historisch feit vragen elk een ander vertrouwensniveau.

We behandelen mentale modellen: grip houden betekent een kloppend beeld hebben van hoe AI zich gedraagt, inclusief waar het faalt. De geschiedenisleraar die weet dat een deel van wat terugkomt niet zal kloppen en zijn vakkennis ernaast houdt; de ICT-coördinator die meeleeft terwijl de redenering zich ontvouwt en ingrijpt zodra die de verkeerde kant op gaat. Beiden houden grip, niet omdat ze slimmer zijn, maar omdat hun verwachting overeenkomt met de werkelijkheid in plaats van met hoe ze hópen dat het werkt. En daarnaast de valkuil van het vermenschlijken: 'het snapt me niet' stuurt je frustratie een kant op die niet klopt, want meer uitleg helpt niet als de patronen niet passen.



2.4 Wat een Socratisch gesprek is en een chatbot niet

Een AI kan een vraag stellen. Een Socratisch gesprek is iets anders: een groep die samen een vraag onderzoekt waar geen antwoord op bestaat, waarin stilte betekenis heeft, waarin je aan iemands gezicht ziet dat hij nog niet overtuigd is, en waarin de deelnemers elkaars denken over tijd leren kennen. Dat laatste is Theory of Mind, het vermogen om te begrijpen dat een ander andere overtuigingen en denkpatronen heeft dan jij. Kinderen ontwikkelen dat rond hun vierde jaar; het is de basis van elke goede samenwerking tussen mensen en precies waar AI niet aan kan komen, hoe overtuigend het soms ook lijkt. Tegelijk kan AI wél iets bijdragen aan kritisch denken, mits in een strak omschreven rol: als advocaat van de duivel die tegenargumenten levert bij een standpunt dat het kind zélf al heeft ingenomen, als leverancier van perspectieven die niemand in de klas vertegenwoordigt, of als sparringpartner bij het aanscherpen van een vraag. Steeds ná het eigen denken, nooit ervoor, en steeds met de opdracht om tegen te spreken in plaats van mee te bewegen. Dat laatste is een directe tegenmaatregel tegen vleierij.

Voor deze training geldt:

Ervaren voor uitleggen.

Deelnemers ervaren automation bias eerst aan zichzelf, met eigen materiaal. Een waarschuwing vooraf werkt aantoonbaar niet; een confrontatie met de eigen blinde vlek wel.

Bestaande instrumenten, nieuwe toepassing.

Geen nieuw AI-vak. De Bouwstenen, de Meetlat en de denkroutines worden losgelaten op AI-output. Dezelfde routine, een nieuw soort tekst.

Afstemmen, niet wantrouwen.

Elk onderdeel eindigt met de vraag: hoeveel vertrouwen past hier, en hoe controleer ik het? Nooit met een algemeen 'wees kritisch'..

Structuur boven voornemens.

Controleren wordt ingebouwd in de opdracht (bronnencheck, tegenspraakronde, verantwoording), niet overgelaten aan oplettendheid.

Karakter is de kern.

Intellectuele moed en bescheidenheid krijgen evenveel tijd als de vaardigheden, omdat daar het echte verschil wordt gemaakt.

Tool-onafhankelijk en veilig

Promptpatronen werken in elke toepassing. Geen persoons-gegevens van kinderen; aandacht voor de leeftijdsgrenzen van AI-toepassingen en voor wat de school heeft goedgekeurd.

De leraar leeft het voor

Hardop twijfelen aan een AI-antwoord in het bijzijn van kinderen is de sterkste les over kritisch denken die er is.



2.6 Externe kaders waarmee de training zich verbindt

- **UNESCO AI Competency Framework for Teachers (Miao & Cukurova, 2024).** Vijf domeinen: mensgerichte mindset, AI-ethiek, AI-basiskennis, AI-pedagogiek en AI voor professionele ontwikkeling, elk op drie niveaus (verwerven, verdiepen, creëren). De training dekt alle vijf domeinen, met het zwaartepunt op mindset en pedagogiek: het uitgangspunt is dat AI de mens dient, niet andersom. Het bijbehorende raamwerk voor leerlingen (begrijpen, toepassen, creëren) is de basis voor wat kinderen zelf moeten leren.
- **Handreiking Werken met AI in het primair onderwijs (Arbeidsmarktplatform PO, 2026) en Handreiking AI voor scholen (Kennisset).** Praktische routekaarten voor AI-geletterdheid en professionalisering op school; de training sluit aan bij de daar gebruikte competenties en het waardenkader van Rathenau, Kennisset en SURF.
- **Negen principes voor AI in het onderwijs (UNESCO Nederland).** Onder meer recht op privacy, transparantie, menselijke regie en: laat AI het denken niet vervangen. Maak afspraken zichtbaar in les en toetsing.
- **EU AI Act, artikel 4 (van kracht sinds februari 2025).** Organisaties die AI inzetten, moeten zorgen voor voldoende AI-geletterdheid bij hun personeel. Scholen zijn zulke organisaties. Deze training is daarmee ook een concrete invulling van een wettelijke verplichting, en dat is een argument richting schoolleiders en besturen.
- **AVG.** Geen persoonsgegevens van kinderen in publieke AI-toepassingen; werken met geanonimiseerde casussen; weten welke toepassing de school heeft goedgekeurd. De training is tool-onafhankelijk en behandelt deze principes als onderdeel van het mentale model dat elke leraar van AI moet hebben.

3. Leerdoelen

De leerdoelen zijn geordend naar het handelingsrepertoire van de procesarchitect uit 'Onderwijs zonder plafond': weten, kunnen, doen, voorleven. Na de training kan de deelnemer:

Weten (inzicht en kennis)

- de zes kernvaardigheden en de drie pijlers van het Intellectueel Kompas uitleggen en toepassen;
- uitleggen wat automation bias, vleierij en hallucinatie zijn, waarom ze uit het systeem zelf voortkomen en waarom waarschuwen alleen niet werkt;
- de valkuil van heldere woorden benoemen: waarom duidelijkheid als norm de andere normen kan maskeren;
- uitleggen waarom kinderen als niet-experts structureel het grootste risico op goedgelovigheid lopen;
- het doel van afstemming benoemen en de twee faalvormen onderscheiden.

Kunnen (vaardigheden en oog)

- een redenering ontleden naar doel, perspectief, vooronderstellingen en consequenties, ook bij een tekst zonder auteur;
- de intellectuele normen als meetlat hanteren en er kritische denkopdrachten mee formuleren;
- denkroutines selecteren en inzetten die passen bij de gewenste denkbeweging, inclusief routines die op AI-output werken;
- een Socratisch gesprek voorbereiden en begeleiden met kinderen vanaf vier jaar;
- drogredenen en denkfouten herkennen, inclusief de AI-specifieke varianten;
- AI inzetten als tegenspreker in plaats van als bevestiger, met promptpatronen die vleierij tegengaan.

Doen (handelen in de praktijk)

- denkroutines als vast patroon in het bestaande curriculum integreren, niet als losse werkvorm;
- controle en verantwoording inbouwen in opdrachten waarin kinderen AI gebruiken;
- een cultuur van denken opbouwen waarin 'waarom denk ik dat?' een gewone vraag is;
- een persoonlijk plan uitvoeren: morgen, na een blok, over een jaar.

Voorleven (houding en ethos)

- hardop en zichtbaar twijfelen aan output, ook aan die van jezelf;
- intellectuele bescheidenheid tonen door te zeggen dat je iets niet weet;
- een kind dat jou tegensprekt belonen in plaats van corrigeren.

Een kleine greep uit de programmaonderdelen:

- **Argumentatie-mapping:** Redeneringen visueel ontleden om de logische structuur van een betoog te doorgronden en kinderen sterker te leren argumenteren. Nieuw: AI-antwoorden in kaart brengen om ontbrekende onderbouwing zichtbaar te maken. Deelnemers brengen eigen materiaal in.
- **Denkmatten en graphic organizers:** Visuele hulpmiddelen om informatie te ordenen en verbanden te leggen bij wereldoriëntatie, taal en begrip lezen. Met aandacht voor wanneer je bewust op papier werkt: ordenen met de hand is zelf de denkstap.
- **Drogredenen en logische denkfouten:** Materialen waarmee kinderen de trucjes leren herkennen die zuiver denken in de weg staan. Uitgebreid met denkfouten van deze tijd: 'de computer zei het', vlotheid als waarheid, bevestiging zoeken en niet-bestaande bronnen. Werkvormen per bouw, vanaf groep 3.
- **Bronnen en betrouwbaarheid in het AI-tijdperk:** Nieuwe sessie over bronbeoordeling nu vrijwel alle teksten er even verzorgd uitzien. Praktische controleroutines op kindniveau, omgaan met AI-gegenereerd beeld en geluid, en het gesprek over wat je wel en niet kunt zien. Sluit aan bij het UNESCO-raamwerk en de handreikingen van Kennisset.



Bronnen

- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126.
- Facione, P. A. (1990). *Critical Thinking: A Statement of Expert Consensus for Purposes of Educational Assessment and Instruction (The Delphi Report)*. American Philosophical Association.
- Gerlich, M. (2025). AI tools in society: Impacts on cognitive offloading and the future of critical thinking. *Societies*, 15(1), 6.
- Goddard, K., Roudsari, A., & Wyatt, J. C. (2012). Automation bias: A systematic review of frequency, effect mediators, and mitigators. *Journal of the American Medical Informatics Association*, 19(1), 121–127.
- Hou, Y. T., & Jung, M. F. (2021). Who is the expert? Reconciling algorithm aversion and algorithm appreciation. *Proceedings of the ACM on Human-Computer Interaction*, 5(CSCW2).
- International AI Safety Report (2026). Hoofdstuk over automation bias en menselijke oordeelsvorming.
- Lee, H.-P., Sarkar, A., Tankelevitch, L., Drosos, I., Rintel, S., Banks, R., & Wilson, N. (2025). The impact of generative AI on critical thinking. CHI 2025. Microsoft Research / Carnegie Mellon University.
- Logg, J. M., Minson, J. A., & Moore, D. A. (2019). Algorithm appreciation: People prefer algorithmic to human judgment. *Organizational Behavior and Human Decision Processes*, 151, 90–103.
- Parasuraman, R., & Riley, V. (1997). Humans and automation: Use, misuse, disuse, abuse. *Human Factors*, 39(2), 230–253.
- Parasuraman, R., & Manzey, D. H. (2010). Complacency and bias in human use of automation: An attentional integration. *Human Factors*, 52(3), 381–410.
- Paul, R., & Elder, L. (2014). *Critical Thinking: Tools for Taking Charge of Your Professional and Personal Life* (2e ed.). Pearson.
- Ritchhart, R., Church, M., & Morrison, K. (2011). *Making Thinking Visible*. Jossey-Bass (Harvard Project Zero). En: Ritchhart, R. (2015). *Creating Cultures of Thinking*. Jossey-Bass.
- Sharma, M., e.a. (2024). Towards understanding sycophancy in language models. ICLR 2024.
- CHI (2026). When flattery backfires: How sycophancy and interaction context shape perceived authenticity and trust in large language models. CHI EA '26.
- Miao, F., & Cukurova, M. (2024). AI competency framework for teachers. UNESCO. En: UNESCO (2024). AI competency framework for students.
- Kennisnet. Handreiking AI voor scholen. Arbeidsmarktplatform PO (2026). Handreiking Werken met AI in het primair onderwijs.
- Dumont, M. (2026). Artikelenreeks Begaafd Onderwijs, in het bijzonder: Mentale modellen – hoe je grip houdt op wat AI wel en niet is; Eigenaarschap; Vijf bekwaamheden die AI niet overneemt; Onderwijs zonder plafond. www.begaafdonderwijs.nl

Verantwoording: opgesteld met Claude (Fable 5.1) op basis van de bestaande trainingsbeschrijving Kritisch Denken, de acht artikelen, de skill Begaafd Onderwijs en actueel opgezocht onderzoek; inhoudelijke keuzes en eindredactie door Minka Dumont. Enkele cijfers (onder meer de meta-analyse van Goddard e.a.) zijn via secundaire bronnen gevonden