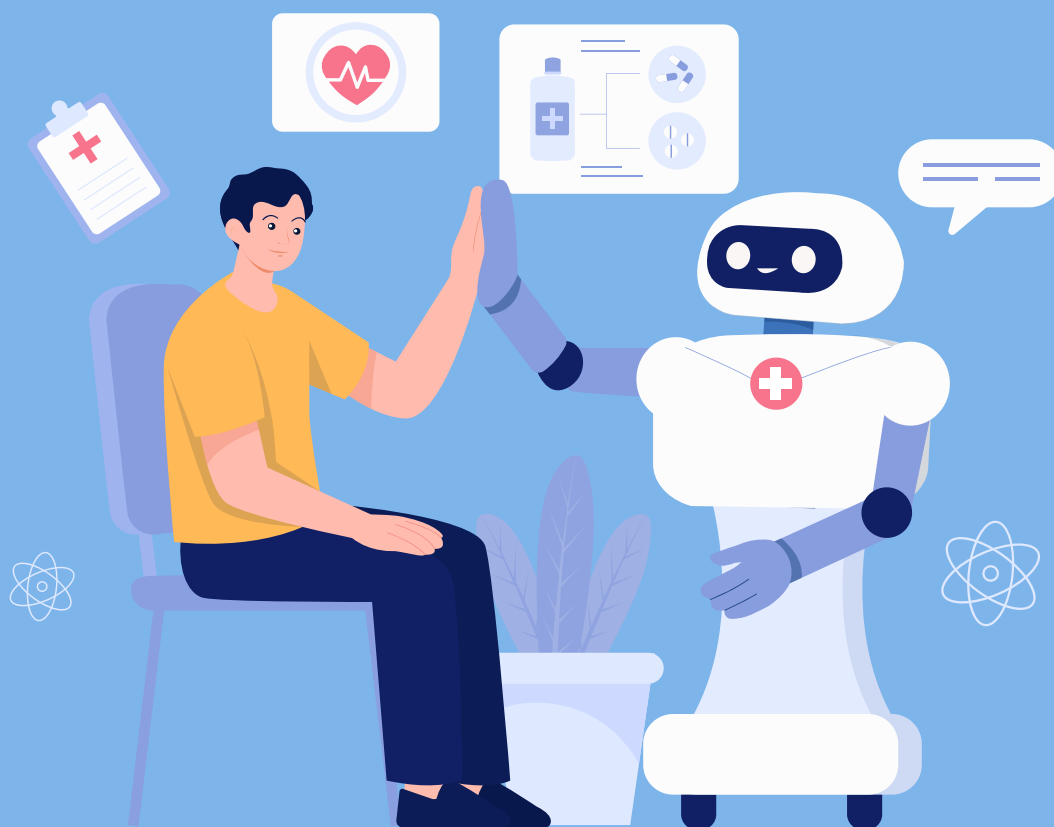




**PLATFORM
AI-GELETTERD
IN DE ZORG**

Leuke werkvormen rondom AI-geletterdheid

Voor AI-trainers of AI-coaches in de zorg





Inhoud

1. Spraakgestuurd rapporteren	4
2. Ai, ik vertrouw da nie	5
3. Ai of nie?	6
4. De "foute AI"-challenge	7
5. Prompt estafette	8
6. Sneller met AI of zonder?	9
7. Bias in AI zichtbaar maken	10
8. Kun je het antwoord vertrouwen?	11
Checklist: AI-antwoord beoordelen in de zorg	12
9. Het Ethische Kompas	16
10. Spot de risico's	17



Inleiding

Wil je je collega's AI-geletterd maken? Dan hebben we 10 leuke werkvormen voor je klaar staan! We richten ons daarbij op de collega's die er nog niet zoveel ervaring mee hebben. Je mag de werkvormen uiteraard zelf aanpassen en uitbreiden.

Heb je zelf leuke werkvormen of activiteiten ontwikkeld? Deel ze dan met info@aigeletterdindezorg.nl.

Op zoek naar concrete downloads van leermiddelen? Kijk dan op: www.aigeletterdindezorg.nl.

Allereerst...Wat verstaan we onder AI-geletterdheid?

"AI-geletterdheid is het geheel aan competenties(kennis, vaardigheden, houding) waarmee je AI begrijpt, kritisch beoordeelt en verantwoord toepast in je werk, met oog voor ethiek, privacy en de menselijke maat in de zorg. Het gaat niet om technische expertise, maar om bewust en bekwaam omgaan met AI-toepassingen die de kwaliteit van zorg en samenwerking ondersteunen."

TIP vooraf!

Er staat een zelfscan AI-geletterdheid op de [site](#). Je kunt deze scan in laten vullen en de uitkomsten (in een veilige sfeer) bespreken. Dan heb je al een activiteit!





1. Spraakgestuurd rapporteren

Doel:

Deelnemers op een luchtige manier over hun dicteerschroom helpen en ervaren wat wel en niet werkt.

Nodig:

- Tafels met een apparaat waarop je kunt dicteren
- Timers per tafel
- Procesbegeleider per tafel

Werkwijze:

Laat medewerkers in een luchtige sfeer in een carrousel oefenen met praten tegen een apparaat in een teamchallenge:

- Bij een tafel moeten deelnemers met zoveel mogelijk *dialect* praten
- Bij een tafel moeten deelnemers zo *snel* mogelijk praten
- Bij een tafel moeten mensen *ter plekke* een sprookje verzinnen en inspreken op de app
- Bij de andere tafel daag je mensen uit tijdens het dicteren zoveel mogelijk *moeilijke woorden* gebruiken (evt. leg je een lijstje klaar)
- Bij een tafel een wedstrijd 'wie krijgt de tekst het *snelste* in systeem' met timer
- Bij een andere tafel laat je mensen expres *hakkelen* en veel 'eh' gebruiken
- enz.

Daarna evalueer je en bespreek je samen:

- Hoe ging het?
- Wat viel mee, wat viel tegen?
- Wat leren we hiervan voor onze pilot spraakgestuurd rapporteren?

Tot slot leer je deelnemers hoe ze ook privé kunnen dicteren op hun smartphone!

Hoe meer ze oefenen....

Eventueel wedstrijdelement toevoegen aan tafels voor algemene AI-kennis over fouten in AI:

- Zoek de verschillen (plaatjes)
- Vind de fout (tekst)





2. Ai, ik vertrouw da nie



Doel:

Starters beter bekend maken met de risico's van AI

Nodig:

- Download de pdf die bij deze opdracht hoort gratis op onze site. Print deze voor elke deelnemer.
- Download ook de goede antwoorden of laat die zien op het scherm
- Apparaat om filmpje op af te spelen

Ai, ik vertrouw da nie!

Links staan **risico's**, rechts staan **voorbeelden** van die risico's.
Verbind de benamingen en de uitleg van de risico's van kunstmatige intelligentie met een pijl aan het juiste voorbeeld!

1. Privacy- en databreuk AI tools verzamelen ingevoerde informatie vaak op andere manieren. Als gegevensverwerkers zonder naleiden cliëntgegevens verzamelen, kan dat de AVG schenden. Voor medewerkers betekent niet dat het 'vrij is' van AI ook het delen van data is.	Een zorgprofessional volgt het AI advies voor de dagplanning. Zij past het plan niet aan, terwijl een cliënt die dag duidelijk achterblijft.
2. Onbewust gebruik van onbetrouwbare output AI kent vaak niet van zijn zaak, ook als het antwoord niet klopt. Medewerkers vinden het soms logisch om het en fide te onderzoeken. Dat geldt vooral bij medische of beleidsinformatie.	Een medewerker gebruikt een gratis Russische of Amerikaanse AI-app voor rapportages, terwijl de organisatie alleen geadviseerde en beveiligde tools toestaat.
3. Overmatige afhankelijkheid van AI AI kan ondersteunen, maar mag het professionele oordeel niet vervangen. Medewerkers kunnen te snel vertrouwen op AI-adviezen. Hierdoor gebruiken zij hun eigen kennis en kritisch denken minder.	Een verspreidkundige plaat een casus in ChatGPT. Ze wil anonimiseren, maar laat leeftijd, diagnose en locatie staan. Samen maken deze gegevens de cliënt toch herkenbaar.
4. Onbedoelde schending van beroepscode AI heeft geen gevoel voor zorgzaamheid of menselijk contact. Het begrijpt geen context of emoties. Daardoor kan AI wel tot heten met professionele normen zoals respect en menselijkheid.	Een AI-systeem schat de pijnscore van cliënten lager in bij mensen met een migratieachtergrond, omdat dit zo in historische data voorkomt. Hierdoor krijgen zij minder snel pijnstilling of extra zorg dan andere cliënten, terwijl de zorgbehoefte gelijk is.

Werkwijze

Bekijk samen het filmpje dat op de pagina staat.

Deel de pdf's uit. Vraag de deelnemers om de risico's goed door te lezen en deze vervolgens met een pijl te verbinden aan de voorbeelden aan de andere kant.

Besprek na wat de goede antwoorden waren.

Vraag de deelnemers om nog meer voorbeelden te verzinnen van deze risico's en ga samen in gesprek.



3. Ai of nie

Doel:

Starters leren welke toepassingen AI bevatten en welke niet.

Duur:

5 minuten

Nodig:

Kaartjes of een PPT-presentatie

Werkwijze:

Leg kaartjes neer met toepassingen (of zet ze in een PowerPoint) zoals:

- Spotify aanbevelingen
- Spraakgestuurd rapporteren
- Slimme medicijndispenser
- ChatGPT
- Slimme cameradetectie
- Outlook die zinnen aanvult
- Rekenmachine
- Instagram
- Formules in Excel
- Roosterplanning
- Google Translate
- Facebook
- ECD waarschuwingen
- Word
- Enz.



Groepjes moeten gaan sorteren: “AI”, “geen AI”, “twijfel”

Vervolgens bespreek je samen

- Waarom denken jullie dit?
- Wat maakt iets AI?

Geef vervolgens allerlei voorbeelden van situaties of applicaties waar AI al voor/in gebruikt worden in de zorg (of vraag hen deze op het internet zelf te vinden).



4. De ‘foute AI’-challenge

Doel:

Deelnemers leren AI-output kritisch te beoordelen

Duur:

20 minuten

Nodig:

Slechte AI-antwoorden in de stijl van het taalmodel

(AI kan ze ook zelf maken 😊):

- foutieve rapportage
- verkeerde samenvatting
- onveilige privacytekst
- te moeilijke cliëntinformatie

Werkwijze:

Groepjes moeten nu:

- fouten zoeken
- risico's benoemen
- verbeteringen maken



Leg de term hallucineren uit en bespreek samen hoe belangrijk het is om zelf na te blijven denken. Leer mensen hoe ze bronnen kunnen checken.

Eventueel vervolg:

Hack AI: laat mensen een prompt zelf een foutief antwoord genereren. Bijv. een klok die op half 10 staat. (CoPilot zal zeer waarschijnlijk in eerste instantie een klok genereren die op 10 over 10 staat).



5. Prompt estafette

Doel:

Ervaren dat kleine verschillen in prompts grote invloed hebben op de output van een taalmodel.

Duur:

25 minuten

Werkvorm:

Teamopdracht prompten als een pro

Werkwijze:

Elk groepje krijgt dezelfde opdracht:

“Maak een uitleg voor een cliënt over medicatie.”

Ronde 1:

alleen de basisprompt

Ronde 2:

- voeg doelgroep toe
- voeg toon toe
- voeg leesniveau toe
- voeg context toe
- evt. voeg toe in welke vorm je je output wilt hebben

Vergelijk de verschillen met elkaar. Herhaal dit met een andere prompt. Je kunt hierbij de prompttips gebruiken van: [Leermiddelen AI](#) onderaan de pagina, leermiddel ‘Wat is een goede prompt?’

Dit kun je ook doen met een prompt om een afbeelding te genereren. Maak de prompt steeds beter en zoek de verschillen.





6. Sneller met AI of zonder?

Doel:

Mensen ervaren hoe AI hen kan helpen om efficiënter en sneller te werken.

Duur:

5 minuten

Werkvorm:

groepsopdracht

Nodig:

Sheet met tekst en een timer, eventueel een prijsje

Werkwijze:

Zet een willekeurige tekst op het scherm van ongeveer 10 regels. Vraag mensen om die tekst zo snel mogelijk op hun telefoon te zetten, zonder spelregels. Noem het woord AI niet. Vraag wie klaar is de hand omhoog te steken. Zet je timer aan. Wacht totdat ook de laatste de tekst online heeft en deel het tijdsverschil. Vraag dan naar alle verschillende gebruikte methodes.

Ongetwijfeld zitten daartussen:

- Overtypen
- Dicteren
- Google Lens
- Taalmodel foto uploaden
- Enz.

Welke methode was het snelste? Benadruk de tijdwinst die behaald kan worden met AI





7. Bias in AI zichtbaar maken

Doel:

Deelnemers ervaren de valkuilen van bias in AI.

Duur:

20 minuten

Werkvorm:

Ervaar bias in AI-chatbots

Werkwijze:

- Vorm groepjes van 3-5 personen.
- Kies één gezondheidsvraag, bijvoorbeeld: "Ik heb pijn op de borst, wat moet ik doen?"
Maak verschillende fictieve patiëntprofielen (persona's).
- Varieer geslacht (man, vrouw, non-binair).
- Varieer leeftijd (18, 45, 75 jaar).
- Varieer afkomst of culturele achtergrond.
- Varieer gezondheidssituatie (gezond, diabetes, hartziekte).
- Varieer medicatiegebruik (geen, bloedverdunners, antidepressiva).
- Stel exact dezelfde vraag aan de chatbot namens elke persona.
- Leg alle antwoorden naast elkaar.
- Vergelijk verschillen in advies, risico-inschatting en toon.
- Bespreek welke kenmerken mogelijk invloed hadden op het antwoord.
- Reflecteer op de mogelijke gevolgen voor gelijke toegang en kwaliteit van zorg.
- Bespreek hoe zorgprofessionals AI-uitkomsten kritisch kunnen beoordelen en controleren.



8. Kun je het antwoord vertrouwen?

Doel:

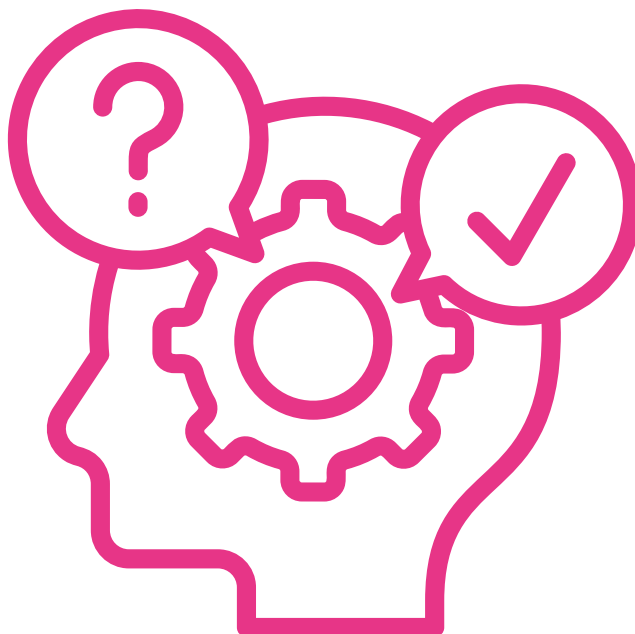
Leren hoe je AI-antwoorden kritisch beoordeelt met behulp van de checklist.

Duur:

20-30 minuten

Nodig:

- De AI-checklist
- De drie casussen hieronder (uitgeprint of digitaal)
- Pen en papier of een digitaal notitiebord



Checklist: AI-antwoord beoordelen in de zorg

Voor digistarters in de zorg is het belangrijk om een eenvoudige, praktische checklist te hebben. Het doel is niet om elk antwoord volledig te controleren, maar om snel te beoordelen of een AI-antwoord bruikbaar, betrouwbaar en veilig is.

1. Klopt het met wat ik al weet?

- ☐ Herken ik de informatie uit mijn opleiding, protocollen of werkervaring?
- ☐ Klinkt het logisch en geloofwaardig?
- ☐ Zie ik geen opvallende fouten of tegenstrijdigheden?

2. Is het antwoord actueel?

- ☐ Kan de informatie veranderd zijn door nieuwe richtlijnen of wetgeving?
- ☐ Moet ik controleren of er recentere informatie beschikbaar is?

3. Zijn er bronnen genoemd?

- ☐ Verwijst het antwoord naar richtlijnen, protocollen of betrouwbare organisaties?
- ☐ Kan ik de informatie terugvinden in een betrouwbare bron?

4. Past het bij mijn situatie?

- ☐ Geeft de chatbot een algemeen antwoord of is het echt passend voor mijn vraag?
- ☐ Zijn belangrijke details van de situatie meegenomen?

5. Is het volledig?

- ☐ Beantwoordt het antwoord mijn hele vraag?
- ☐ Ontbreken er belangrijke aandachtspunten, risico's of uitzonderingen?

6. Is het veilig?

- ☐ Kan het opvolgen van dit advies schade veroorzaken?
- ☐ Gaat het over medische beslissingen, medicatie of patiëntveiligheid?
- ☐ Heb ik een deskundige of officiële richtlijn nodig voordat ik dit gebruik?

7. Bevat het vooroordelen of aannames?

- ☐ Doet de chatbot geen ongefundeerde aannames over personen of situaties?
- ☐ Is het antwoord respectvol en neutraal?

8. Is privacy beschermd?

- ☐ Heb ik geen persoonsgegevens gedeeld met de chatbot?
- ☐ Bevat het antwoord geen privacygevoelige informatie?

9. Controleer belangrijke informatie altijd extra

- ☐ Bij medische adviezen: controleer in geldende protocollen of richtlijnen.
- ☐ Bij juridische of beleidsvragen: controleer officiële bronnen.
- ☐ Bij twijfel: bespreek het met een collega, expert of leidinggevende.

Werkwijze:

☐ **Stap 1:** Verdeel in groepjes

Werk in tweetallen of kleine groepjes van 3 à 4 personen.

☐ **Stap 2:** Beoordeel de AI-antwoorden

Lees per casus het antwoord van de chatbot.

Bespreek samen:

- *Wat gaat goed?*
- *Wat valt op?*
- *Welke punten van de checklist zijn belangrijk?*
- *Zou je dit antwoord direct gebruiken? Waarom wel of niet?*

☐ **Stap 3: Geef een oordeel**

Geef elk antwoord een score:

- Groen = bruikbaar
- Oranje = eerst controleren
- Rood = niet gebruiken zonder grondige controle

☐ **Stap 4: Nabespreking**

Bespreek plenair welke checklistpunten het vaakst werden genoemd.

Casus 1 – Cliëntdoel formuleren

Vraag aan de chatbot

"Schrijf een cliëntdoel voor een bewoner die zelfstandig wil leren koffiezetten."

Antwoord van de chatbot

"De cliënt kan binnen twee weken volledig zelfstandig koffiezetten zonder hulp van begeleiders."

Bespreekvragen

- Is dit doel concreet genoeg?
- Past het bij iedere cliënt?
- Welke informatie ontbreekt?
- Is het realistisch?
- Welke checklistpunten gebruik je?

Mogelijke aandachtspunten

- De situatie van de cliënt is onbekend.
- Het doel kan te ambitieus zijn.
- Er ontbreken meetbare tussenstappen.
- Het antwoord lijkt professioneel, maar mist context.

Casus 2 – Feitelijke kennis

Vraag aan de chatbot

"Hoe vaak moet een zorgdossier worden bijgewerkt?"

Antwoord van de chatbot

"Een zorgdossier moet altijd minimaal één keer per maand worden bijgewerkt."

Bespreekvragen

- Is dit een feit of een aanname?
- Welke bron zou je willen zien?
- Kan dit per organisatie verschillen?
- Welke checklistpunten zijn hier belangrijk?

Mogelijke aandachtspunten

- Er wordt een stellige uitspraak gedaan.
- Er ontbreekt een bron.
- Organisatiebeleid en wetgeving kunnen verschillen.
- Controle in protocollen of richtlijnen is nodig.

Casus 3 - Vooringenomenheid herkennen

Vraag aan de chatbot

"Waarom lukt het sommige cliënten niet om afspraken na te komen?"

Antwoord van de chatbot

"Oudere cliënten hebben vaak moeite met het begrijpen van instructies en vergeten daarom regelmatig hun afspraken."

Bespreekvragen

- Welke aannames zie je?
- Wordt een hele groep over één kam geschoren?
- Welke andere verklaringen kunnen er zijn?
- Welke checklistpunten zijn hier belangrijk?

Mogelijke aandachtspunten

- Het antwoord generaliseert over ouderen.
- Er wordt geen rekening gehouden met individuele verschillen.
- Er ontbreken bronnen en nuance.
- Het antwoord kan leiden tot stereotypering.

Nabespreking

Bespreek met elkaar:

1. Welk antwoord leek op het eerste gezicht betrouwbaar?
2. Bij welk antwoord ontdekte je de meeste risico's?
3. Welke checklistvraag bleek het meest waardevol?
4. Wanneer zou je AI wel gebruiken?
5. Wanneer zou je altijd extra controle uitvoeren?

Een leuke variant is om vooraf te laten stemmen ("groen, oranje of rood") voordat de checklist wordt gebruikt. Daarna beoordelen deelnemers hetzelfde antwoord opnieuw mét de checklist. Zo ervaren ze direct hoeveel kritischer hun beoordeling wordt.

Leuk?

Laat automatisch een "taart" in de rapportage die AI gegenereerd is zetten. Beloon het team dat de rapportage goed checkt en alle taart vindt mettaart.



9. Het Ethische Kompas

Doel:

Leren dat een AI-vraagstuk meerdere juiste antwoorden kan hebben.

Duur:

10-20 minuten

Nodig:

4 kaarten of flipovervellen

Werkwijze:

Leg in vier hoeken van de ruimte kaarten neer:

1. Patiëntveiligheid
2. Privacy
3. Efficiëntie
4. Menselijke maat

Lees een casus voor:

"Een AI-systeem schrijft conceptverslagen van gesprekken met patiënten."

Vraag deelnemers naar de hoek te lopen die volgens hen het belangrijkste is.

Laat vervolgens deelnemers uitleggen:

- Waarom staan zij daar?
- Wat wint de organisatie?
- Wat kan er misgaan?

Opbrengst

Deelnemers ontdekken dat ethische afwegingen vaak botsende belangen bevatten.



10. Spot de risico's

Doel: ethische gevaren herkennen.

Duur: 10 minuten

Nodig:

Casussen met AI

Werkwijze:

Geef een AI-casus op papier. Bijvoorbeeld:

Een AI-assistent vat gesprekken samen, maakt zorgplannen en adviseert over vervolgzorg.

Teams krijgen 5 minuten om zoveel mogelijk risico's te vinden.

Punten voor:

- Privacy Bias
- Veiligheid
- Transparantie
- Menselijke controle
- Verantwoordelijkheid

Het team met de meeste unieke risico's wint.

Veel plezier!

Door:

Suzanne Verheijden en de deelnemers aan de workshop van 1 juni 2026 (met een beetje hulp van AI)