

HANDLEIDING

Aan de slag met het de Ethische Quicksan AI

“Wat vinden we moreel goed om te doen met deze AI-toepassing?”

Deze quickscan helpt je als organisatie om op een gestructureerde manier met elkaar in gesprek te gaan over de inzet van een AI-toepassing. Het doel is niet om zo snel mogelijk tot een “ja” of “nee” te komen, maar om samen te onderzoeken wat een AI-toepassing oplevert, welke risico's of twijfels er zijn en wat nodig is om deze verantwoord in te zetten.

Vorbereiding

Voordat je start, is het belangrijk dat de deelnemers voldoende inzicht hebben in:

- de AI-toepassing die wordt beoordeeld;
- het doel waarvoor deze wordt ingezet;
- de manier waarop de toepassing werkt in de praktijk;
- het AI-beleid van de leverancier.

Nodig bij voorkeur mensen uit met verschillende perspectieven, bijvoorbeeld medewerkers, inhoudsdeskundigen, projectleiders, cliënten of cliëntvertegenwoordigers en andere betrokkenen. Juist de combinatie van verschillende ervaringen zorgt voor een waardevolle afweging.

Werken met de kaarten

☒ **Tip:** print de kaarten met de ethische criteria uit, zie bijlage 1 en gebruik ze fysiek tijdens de sessie. Deelnemers kunnen ze dan vastpakken, naast elkaar leggen en er direct post-its bij plakken – dat maakt het gesprek levendiger en vergroot de betrokkenheid en het eigenaarschap.

Tijdens de sessie

Doorloop gezamenlijk de verschillende criteria van de quickscan, zie Bijlage 1 voor een overzicht. Bespreek per criterium:

- welke waarde of voordelen de AI-toepassing oplevert;
- welke twijfels, risico's of schuurnvlakken er zijn;
- welke vragen nog openstaan.

Schrijf bevindingen op post-its en leg deze bij het betreffende criterium. Zo ontstaat gedurende het gesprek een visueel overzicht van alle argumenten die een rol spelen in de uiteindelijke afweging.

Noot: Bij gebruik van generatieve AI wordt standaard de post-it 'verborgen impact van generatieve AI' als schuurvlak geplakt.

Voordelen

Twijfels en schuurvlakken

Verborgen impact van generatieve AI.

Hoog energie- en waterverbruik en mogelijke druk op arbeidsomstandigheden van mensen die data voor AI beoordelen, labelen en verbeteren.

Eindafweging

Nadat alle criteria zijn besproken, kijk je gezamenlijk naar het totaalbeeld. Welke voordelen zien we? Welke risico's of bezwaren zijn er? En wat vinden we, alles afwegende, moreel goed om te doen? De uitkomst kan bijvoorbeeld zijn:

Uitkomsten	Toelichting
Vrij inzetbaar	Deze AI-toepassing kan ingezet worden zoals deze nu is. Wij vinden het moreel goed om deze AI-toepassing te gebruiken, omdat we <deze voordelen> belangrijker vinden dan <deze twijfels of schuurvlakken>.
Inzetbaar met aanpassingen	Deze AI-toepassing kan worden ingezet, mits de volgende aanpassingen zijn gedaan of maatregelen zijn genomen, te weten ... (op proces, technologie, leverancier, etc.)
Mogelijk inzetbaar na testen en/of informatie	Deze AI-toepassing is mogelijk inzetbaar na extra testen en/of aanvullende informatie of gesprekken erover. Een 'moreel beraad' of gesprek volgens de 'aanpak begeleidingsethiek' kan een noodzakelijke vervolgstap zijn.
Niet inzetbaar (moreel niet goed)	Deze AI-toepassing is op dit moment niet verantwoord om in te zetten, omdat we <deze twijfels of schuurvlakken> belangrijker vinden dan <deze voordelen>.

Tot slot

De Ethische Quickscan AI is geen eenmalige toets, maar een hulpmiddel om het goede gesprek te voeren. Technologie, context en inzichten veranderen. Maak daarom afspraken

over evaluatie en herbeoordeling, zodat je ook op langere termijn bewust en verantwoord met deze AI-toepassing blijft omgaan.

Over de totstandkoming van v1.0 (10 juni 2026)

AI kan de zorg verbeteren en de werkdruk verlagen. Tegelijk heeft het impact op privacy, menselijke regie en de relatie tussen mensen. Daarom is het belangrijk om zorgvuldig af te wegen wat moreel passend is bij een AI-toepassing. Zo wordt een besluit ook uitlegbaar binnen de organisatie en naar cliënten.

De Ethische Quicksan AI (versie 1.0) is ontwikkeld in samenwerking tussen de AWIZ-regio's Noord-Oost Brabant en Rivierenland, met Désirée Hobbelen en Inge van Kooten-Satter. De quickscan is gebaseerd op gesprekken met experts en ethici, theorie (plichtethiek, gevolgenethiek, deugdethiek en zorgethiek), peer reviews en praktijktesten met AI-toepassingen.

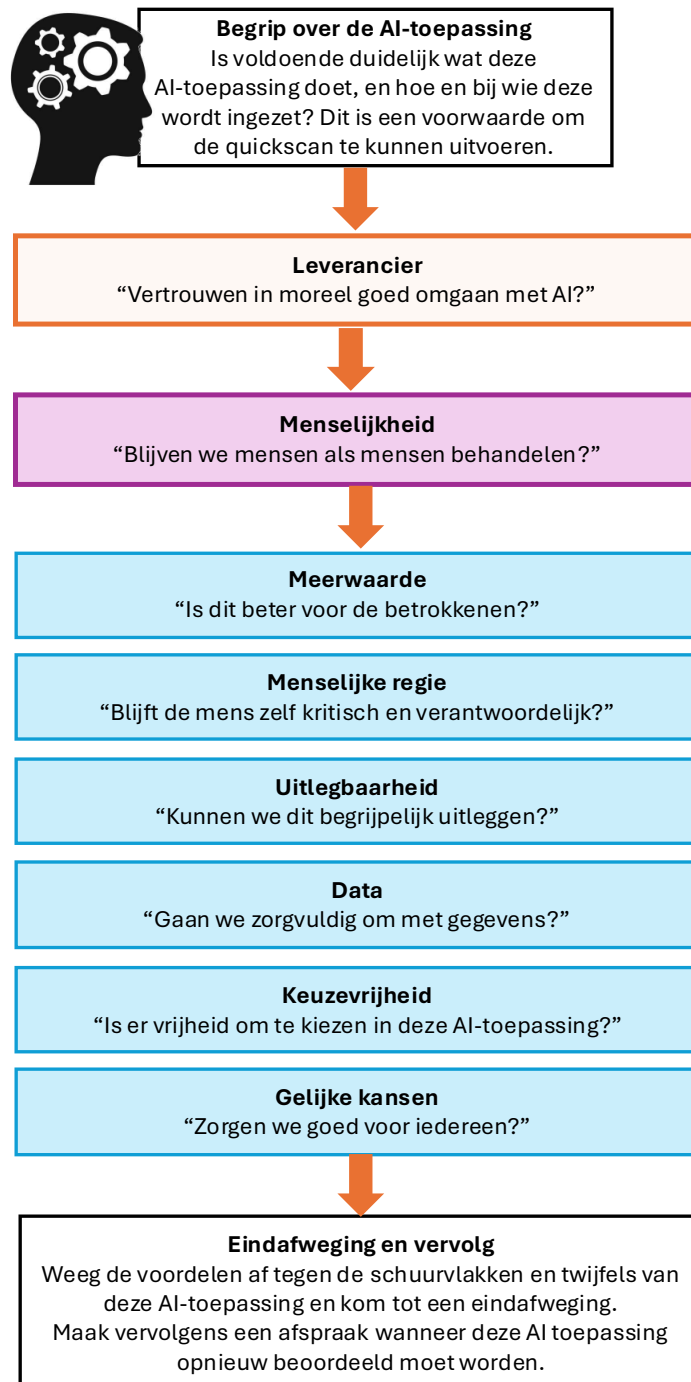
Ons doel is om de quickscan breed te gebruiken en met jullie feedback door te ontwikkelen. Door beoordeelde AI-toepassingen te delen krijgen we sneller inzicht in de verantwoorde inzet van AI-toepassingen.

10 juni 2026, Inge van Kooten-Satter en Désirée Hobbelen

Bijlage 1: Criteria Ethische Quickscan AI

Ethische Quickscan AI

“Wat vinden we moreel goed om te doen met deze AI-toepassing?”



Leverancier

“Vertrouwen in moreel goed omgaan met AI?”

AWIZ Ethische Quicksan AI-toepassing

Menselijkheid

“Blijven we mensen als mensen behandelen?”

AWIZ Ethische Quicksan AI-toepassing

✓ Waar geeft de leverancier vertrouwen dat deze AI-toepassing verantwoord wordt ontwikkeld en gebruikt? (bijv. openheid, uitleg, beleid, samenwerking)

⚠ Waar blijft het onduidelijk of twijfelachtig hoe de leverancier met risico's omgaat (zoals data, bias of veiligheid)?

⚠ Waar voelt het alsof je moet vertrouwen op de leverancier, zonder dat je echt kunt zien hoe de AI-toepassing werkt?

AWIZ Ethische Quicksan AI-toepassing

✓ Waar helpt deze AI-toepassing om beter aan te sluiten bij wat iemand nodig heeft als mens?

⚠ Waar lopen we risico dat iemand meer als 'geval' of data wordt gezien? (Kan het afstand creëren in het contact tussen jou en de ander?)

⚠ Waar kunnen mensen vooral gezien worden als middel om werk sneller of efficiënter te doen?

AWIZ Ethische Quicksan AI-toepassing

Meerwaarde

“Is dit beter voor de betrokkenen?”

AWIZ Ethische Quicksan AI-toepassing

Menselijke regie

“Blijft de mens zelf kritisch en verantwoordelijk?”

AWIZ Ethische Quicksan AI-toepassing

✓ Waar helpt deze AI-toepassing om betere zorg te leveren?

⚠ Waar helpt deze AI-toepassing vooral ons (gemak, tijd), maar minder de cliënt?

⚠ Waar voelt deze AI-toepassing als iets dat we doen omdat het kan of omdat het een leuk nieuw snufje is? (Zijn er alternatieven?)

AWIZ Ethische Quicksan AI-toepassing

✓ Waar helpt deze AI-toepassing om betere keuzes te maken?

⚠ Wanneer bestaat het risico dat uitkomsten van deze AI-toepassing worden overgenomen zonder voldoende kritische beoordeling, terwijl deze fouten, vooroordelen of verzonnen informatie bevatten?

⚠ Wat zou je kwijt kunnen raken van je vakmanschap als je deze AI-toepassing vaak gebruikt?

AWIZ Ethische Quicksan AI-toepassing

Uitlegbaarheid

“Kunnen we dit begrijpelijk uitleggen?”

AWIZ Ethische Quickscan AI-toepassing

Data

“Gaan we zorgvuldig om met gegevens?”

AWIZ Ethische Quickscan AI-toepassing

✓ Waar helpt deze AI-toepassing om dingen beter te onderbouwen of uit te leggen?

⚠ Waar kun je niet goed uitleggen wat deze AI-toepassing doet - en moet je het 'gewoon' geloven?

⚠ Waar kun je de mogelijke gevolgen van deze AI-toepassing nog niet overzien?

AWIZ Ethische Quickscan AI-toepassing

✓ Waar helpt deze AI-toepassing om zorgvuldiger om te gaan met gegevens?

⚠ Waar bestaat het risico dat te veel of gevoelige gegevens worden gedeeld met deze AI-toepassing, met mogelijke schade voor personen of de organisatie als gevolg?

⚠ Waar kan het vertrouwen van de cliënt beschadigd raken door hoe we met de data omgaan?

AWIZ Ethische Quickscan AI-toepassing

Keuzevrijheid

“Is er vrijheid om te kiezen in deze AI-toepassing?”

AWIZ Ethische Quickscan AI-toepassing

Gelijke kansen

“Zorgen we goed voor iedereen?”

AWIZ Ethische Quickscan AI-toepassing

✓ Waar vergroot deze AI-toepassing de keuzevrijheid van gebruikers*?

⚠ Wanneer is het lastig voor gebruikers om “nee” te zeggen tegen het gebruik van deze AI-toepassing?

⚠ Wat zijn de mogelijke gevolgen er iemand niet mee wil doen met deze AI-toepassing?

*medewerkers, cliënten, etc.

AWIZ Ethische Quickscan AI-toepassing

✓ Voor wie maakt deze AI-toepassing de zorg beter of toegankelijker?

⚠ Wie krijgt door deze AI-toepassing mogelijk minder aandacht of minder passende zorg?

⚠ Waar kan deze toepassing mensen op achterstand zetten?

AWIZ Ethische Quickscan AI-toepassing