

Ethische Quicksan AI

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

De Ethische Quicksan AI dient om een AI-toepassing te beoordelen vanuit verschillende ethische perspectieven. De quickscan wordt uitgevoerd in een gesprek met bijvoorbeeld een ethicus, inhoudsdeskundige, projectleider, medewerker en cliënt. In dit gesprek wordt zichtbaar wat de AI-toepassing oplevert en waar het schuurt. Vervolgens wordt bepaald wat moreel passend is om te doen en onder welke voorwaarden of met welke aanpassingen de AI-toepassing kan worden ingezet.

Vooraf is duidelijk hoe de AI-toepassing werkt en hoe deze in de organisatie wordt toegepast. Ook is het AI beleid van de leverancier bekend. Op de volgende pagina staat een overzicht van de ethische criteria van de quickscan. Per criterium worden twee aspecten bekeken:

- Wat de AI-toepassing oplevert (waarde)
- Wat schuurt of waar twijfel zit

De bevindingen worden genoteerd op post-its en in de juiste kolom geplaatst, zie onder. 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.

Ethische Quickscan AI

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

Na het bespreken van alle ethische criteria worden de voordelen afgewogen tegen de twijfels en schuurlakken. Op basis daarvan wordt bepaald wat moreel passend is om te doen bij deze AI-toepassing. In de onderstaande tabel vind je de mogelijke uitkomsten van de *Ethische Quickscan AI*.

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 schuurlakken>.
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 'aankpak 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 schuurlakken> belangrijker vinden dan <deze voordelen>.

Ethische Quicksan AI

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



Begrip over de AI-toepassing

Is voldoende duidelijk wat deze AI-toepassing doet, en hoe en bij wie deze wordt ingezet? Dit is een voorwaarde om de quickscan te kunnen uitvoeren.



Leverancier

“Vertrouwen in moreel goed omgaan met AI?”



Menselijkheid

“Blijven we mensen als mensen behandelen?”



Meerwaarde

“Is dit beter voor de betrokkenen?”

Menselijke regie

“Blijft de mens zelf kritisch en verantwoordelijk?”

Uitlegbaarheid

“Kunnen we dit begrijpelijk uitleggen?”

Data

“Gaan we zorgvuldig om met gegevens?”

Keuzevrijheid

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

Gelijke kansen

“Zorgen we goed voor iedereen?”



Eindafweging en vervolg

Weeg de voordelen af tegen de schuurvlakken en twijfels van deze AI-toepassing en kom tot een eindafweging.
Maak vervolgens een afspraak wanneer deze AI toepassing opnieuw beoordeeld moet worden.

Verborgene impact van generatieve AI.

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



Leverancier

“Vertrouwen in moreel goed omgaan met AI?”

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

Menselijkheid

“Blijven we mensen als mensen behandelen?”

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

✓ 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

Menselijke regie

“Blijft de mens zelf kritisch en verantwoordelijk?”

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 Quicksan 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 Quicksan AI-toepassing

Data

“Gaan we zorgvuldig om met gegevens?”

AWIZ Ethische Quicksan 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 Quicksan AI-toepassing

Keuzevrijheid

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

AWIZ Ethische Quicksan 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 Quicksan AI-toepassing

Gelijke kansen

“Zorgen we goed voor iedereen?”

AWIZ Ethische Quicksan 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 Quicksan AI-toepassing