



Navigeren in een woeste digitale zee

AI-kompas voor bestuurders

Samen voor betere
zorg & welzijn



Inclusie: is de toepassing geschikt voor iedereen?

Fase 1: **Besluitvorming (go/no-go)**

1. Wat is het doel dat we willen bereiken, en helpt AI daar écht bij?
2. Voor wie bouwen we dit systeem – en wie dreigt buiten beeld te vallen?
3. Wat betekent het dat het systeem is getraind op 'eigen data van onze cliënten'?
4. Hoe betrekken we cliënten met diverse achtergronden in de besluitvorming?
5. Is inclusie expliciet een afwegingscriterium bij onze go/no-go-beslissing?
6. Hoe voorkomen we informatiebias bij het gebruik van eigen of beperkte datasets?

Fase 2: **Ontwikkelen / inkopen**

1. Hoe zorgen we dat de interface begrijpelijk is voor cliënten met verschillende taal- en opleidingsniveaus?
2. Hoe sluiten de vragen en functies aan bij de leefwereld van cliënten?
3. Is inclusie een inkoopcriterium bij de selectie van leveranciers?
4. Hoe betrekken we professionals en cliënten in de ontwikkeling (feedbackloop)?
5. Hoe testen we of het systeem écht inclusief is en uitzonderingen herkent?
6. Hoe voorkomen we dat standaardprofielen cliënten reduceren tot één type?

Fase 3: **Gebruiken**

1. Hoe zorgen we dat álle medewerkers het systeem kunnen én willen gebruiken?
2. Hoe signaleren we of bepaalde groepen niet profiteren van de AI-uitkomsten?
3. Hoe voorkomen we dat cliënten zich gestigmatiseerd of in een hokje geplaatst voelen?
4. Hoe blijven we het systeem trainen en evalueren op inclusiviteit?
5. Hoe gaan we om met groepen die geen gebruik kunnen maken van het systeem?
6. Hoe borgen we dat inclusie structureel onderdeel blijft van het gebruik en beleid?

Werkplezier: wat betekent het voor medewerkers?

Fase 1: Besluitvorming (go/no-go)

1. Welk probleem willen we oplossen en draagt AI hier echt aan bij?
2. Wat verwachten we dat dit systeem toevoegt aan het werk van professionals - en verhoogt het werkplezier/ werkdruk?
3. Is de applicatie ondersteunend aan het werk van medewerkers?
4. Hoe zorgen we dat medewerkers vanaf het begin worden gehoord en betrokken?
5. Wat vinden medewerkers én cliënten van het gebruik van dit systeem?
6. Wanneer is 'enthousiasme van enkelen' niet genoeg om draagvlak te veronderstellen?
7. Is er inzicht in eerdere ervaringen of een track record van succesvolle toepassing bij andere organisaties?
8. Wordt medewerkerstevredenheid(onderzoek) meegenomen in de besluitvorming rond deze technologie?

Fase 2: Ontwikkelen / inkopen

1. Is er een referentiegroep van medewerkers betrokken bij de ontwikkeling of selectie van de tool?
2. Hoe trainen we medewerkers om AI-uitkomsten goed te begrijpen, toe te passen en te bevragen?
3. Hoe voorkomen we dat medewerkers blind vertrouwen op de output van AI?
4. Hoe gaan we om met gevoelens van angst of onzekerheid bij medewerkers over het gebruik van AI?
5. Hoe stimuleren we ruimte voor reflectie, twijfel en professionele dialoog?
6. Hoe zorgen we voor duidelijke communicatie, werkafspraken en een heldere pilotfase die vertrouwen wekt?
7. Hoe borgen we deskundigheidsbevordering en leren in deze fase?

Fase 3: Gebruiken

1. Hoe zorgen we dat medewerkers tijd en ruimte krijgen om de tool goed te gebruiken?
2. Hoe stimuleren we dat medewerkers de tool actief inzetten (en deze niet naast zich neerleggen)?
3. Hoe zorgen we dat AI ondersteunend blijft, en niet leidend of beslissend?
4. Hoe behouden medewerkers autonomie en eigenaarschap in hun werk?
5. Kunnen we indicatoren (KPI's) koppelen aan werkplezier of medewerkerstevredenheid bij gebruik van de tool?
6. Komt er tijd vrij door het gebruik van de tool en hoe benutten we die tijd om werkplezier te vergroten?
7. Hoe stimuleren we een cultuur waarin kritisch nadenken en leren van ervaringen vanzelfsprekend is?
8. Hoe evalueren we structureel het effect van AI op werkdruk, werkgeluk en samenwerking?

Privacy: hoe ga je om met risico's met betrekking tot persoonsgegevens?

Fase 1: Besluitvorming (go/no-go)

1. Hoe zorgen we dat het AI-systeem voldoet aan wettelijke eisen (AVG, medisch beroepsgeheim)?
2. Is de Privacy Officer en de Functionaris Gegevensbescherming betrokken bij het project?
3. Welke data is écht noodzakelijk voor dit systeem en voor welk doel wordt deze gebruikt?
4. Is toestemming van cliënten (/ wettelijk vertegenwoordigers) nodig en hoe organiseren we dat zorgvuldig?
5. Begrijpen cliënten waarvoor hun data wordt gebruikt en is er een alternatief voor wie er geen gebruik van wil maken?
6. Waar worden gegevens opgeslagen en zijn beveiligingsmaatregelen passend bij de gevoeligheid van de data?
7. Blijft de organisatie eigenaar van de data, of kan deze worden gedeeld met derden?
8. Hoe beperken we het risico op onrechtmatige besluitvorming door AI (menselijke tussenkomst)?

Fase 2: Ontwikkelen / inkopen

1. Hoe wordt vastgelegd wie verwerkingsverantwoordelijke en wie verwerker is (verwerkersovereenkomst)?
2. Hoe borgen we dataversleuteling, opslaglocatie, bewaartermijnen en datadoorgifte in lijn zijn met de AVG?
3. Hoe zorgen we dat data veilig wordt opgeslagen en verwerkt, ook bij internationale opslag (buiten de EER)?
4. Welke data gebruiken we voor het trainen van het systeem en hoe waarborgen we dataminimalisatie?
5. Wordt tijdens pilots gewerkt met synthetische of testdata (AI Sandbox)?
6. Hoe zorgen we dat medewerkers goed getraind zijn in veilig omgaan met data en privacyregels?
7. Hoe is de integratie met bestaande systemen geregeld (zoals cliëntdossiers) en wat betekent dit voor dataveiligheid?
8. Moet er een Data Protection Impact Assessment (DPIA) worden uitgevoerd?

Fase 3: Gebruiken

1. Hoe monitoren we dat medewerkers veilig en verantwoord met data omgaan in de praktijk?
2. Hoe wordt omgegaan met datalekken of misbruik van gegevens - is de meldroute bekend?
3. Hoe houden we toezicht op leveranciers en maken we afspraken over audits, rapportages en escalatieprocessen?
4. Hoe zorgen we dat privacybewustzijn (awareness) blijvend onderdeel is van de organisatiecultuur?
5. Hoe communiceren we helder met cliënten over het gebruik van hun data (bijv. via begrijpelijke informatie of flyers)?
6. Hoe zorgen we dat cliënten grip houden op hun gegevens (recht op informatie, inzage, verwijdering, correctie)?
7. Hoe voorkomen we dat AI ooit beslissingen neemt zonder menselijke tussenkomst?

Kwaliteit: hoe zorg je dat de uitkomsten betrouwbaar zijn?

Fase 1: **Besluitvorming (go/no-go)**

1. Wat betekent het dat het systeem wordt getraind op 'onze eigen data' - en hoe representatief is die data?
2. Leren we alleen van cliënten die al bij ons in beeld zijn en wat missen we daardoor?
3. Welke groepen of ervaringen ontbreken in de data en hoe gaan we daarmee om?
4. Hoe toetsen we of de uitkomsten eerlijk zijn voor verschillende doelgroepen?
5. Wat is het doel dat we willen bereiken - en komt het systeem de kwaliteit van de dienstverlening ten goede (verbetering, behoud, of minder snelle afname)?
6. Zijn er eenvoudiger of betrouwbaardere alternatieven dan AI?
7. Wat zijn de (zorg)risico's van het koppelen van cliënten aan 'soortgelijke profielen'?

Fase 2: **Ontwikkelen / inkopen**

1. Hoe toetsen we of het systeem groepen bevoordeelt of benadeelt?
2. Hoe zorgen we dat fictieve of voorbeeldprofielen geen stereotypes versterken?
3. Hoe voorkomen we dat hulpverleners te veel vertrouwen op de uitkomst van het AI-systeem?
4. Hoe zorgen we dat de trainingsdata kwalitatief goed, actueel en representatief is en wie controleert dit?
5. Hoe gaan we om met onvolledige of subjectieve data (zoals verslagen of interviews)?
6. Hoe vertalen we wettelijke eisen naar een helder programma van eisen (i.g.v. inkoop) (kwaliteit, betrouwbaarheid, meetbaarheid, resultaat)? En hoe borgen we dat?

Fase 3: **Gebruiken**

1. Hoe controleren we structureel op fouten, vertekeningen of bias in de profielen?
2. Welke mechanismen zijn ingericht om bias te signaleren en te corrigeren tijdens gebruik?
3. Wat doen we als het systeem verkeerde of stigmatiserende verbanden legt?
4. Hoe blijft het systeem leren van feedback, zonder bestaande vooroordelen te versterken?
5. Hoe wordt kwaliteit van data en output geborgd in de reguliere kwaliteitscyclus van de organisatie?
6. Is er een vier-ogenprincipe ingebouwd bij testen of beoordeling van AI-uitkomsten?
7. Hoe waarborgen we dat AI ondersteunend blijft aan het vakmanschap van de professional en niet sturend of beslissend?