

Inzichten en aanbevelingen: Evaluatiestudie DECIDE- VerA fase I & II

Onderzoeksrapport t.b.v. projectfase III

Leeswijzer

Waar “patiënten” staat kan ook gelezen worden “naasten en/of verzorgers van patiënten en/of wettelijke vertegenwoordigers van patiënten”.

Waar “kiezen voor een behandeling / optie” of “behandelingsmogelijkheid” staat moet ook het “kiezen om niet behandeld te worden / geen actief vervolgtraject te willen volgen” worden begrepen. Onder “behandeling” vallen ook (leefstijl)interventies.

Waar “project” staat wordt het DECIDE-VerA project bedoeld.

Afkortingen

APK: Algemene Periodieke Keuring

BE-sessies: Begeleidingsethiek sessies

CDSS: Clinical Decision Support System

DECIDE-VerA: Clinical DECISION support system voor het cardiovasculaire risico management in de eerstelijnsgezondheidszorg: een Verantwoordelijk- en Aansprakelijkheidsperspectief

LUMC: Leids Universitair Medisch Centrum

TRL: Technology Readiness Level

TU/e: Technische Universiteit Eindhoven

SB: Samen Beslissen

Inhoudsopgave

LEESWIJZER	2
INHOUDSOPGAVE	3
ACHTERGRONDINFORMATIE	5
1 INLEIDING	6
1.1 Projectdoelstellingen	6
1.2. Situering van het evaluatieonderzoek en gebruikte informatie	6
1.3 Doel evaluatie	7
1.4 Theoretische kaders	7
1.5 Hulpmiddel Handelingsruimte & Technology Readiness level	8
2 AI EN DE DECIDE AI TOOL	10
2.1 Artificiële Intelligentie (AI)	10
2.2 AI-CDSS DECIDE -model en uitgangsposities project	10
3 EVALUATIESTUDIE	13
3.1 Plan van aanpak	13
3.2 Onderzoeksactiviteit -1: workshop met patiënten/naasten	13
3.3 Onderzoeksactiviteit -2: focusgroep met huisartsen	20
4 SYNTHESE EN AANBEVELINGEN	23
4.1 Synthese van onderzoeksresultaten	23
4.2 Aanbevelingen	25
5 BRONNEN	29
6 BIJLAGEN	31
6.1. Projectfasen	31

6.2 Scenario's in fase 1 en fase 2 van het gebruik -----	32
6.3 Onderzoeksactiviteit-1 workshop met patiënten en naasten -----	34
6.4 Onderzoeksactiviteit-2 focusgroep met huisartsen-----	48
6.5 Toestemmingsformulieren -----	50

Achtergrondinformatie

Het huidige rapport is tot stand gekomen binnen het project Clinical DECision support system voor het carDiovasculaire risico management in de eerstelijnsgezondheidszorg: een Verantwoordelijk- en Aansprakelijkheidsperspectief (DECIDE-VerA). Het doel van het onderzoeksrapport is om het projectteam van adviezen te voorzien voor fase III van het project. In fase III worden praktisch hanteerbare tools ontwikkeld om de AI Clinical Decision Support System (CDSS) DECIDE tool verder te ontwikkelen en om de kennis voortvloeiend uit dit project te generaliseren voor de ethiek, medische en AI-communities en de patiëntenpopulatie.

DECIDE-VerA is een interdisciplinair project¹ van 2.5 jaar, geleid door het Leids Universitair Medisch Centrum (LUMC) n gefinancierd door ZonMw (grant no. 08540122120004). In het project brengen verschillende groepen (LUMC, Leiden Universiteit, Technische Universiteit Eindhoven (TU/e), Hogeschool Rotterdam en de Patiëntenfederatie Nederland) hun expertise samen om de vervolgstap in de ontwikkeling van een AI-gedreven CDSS te helpen nemen een deze kennis te generaliseren voor de AI community.

Het project DECIDE-VerA is gericht om de ethische, juridische en ontwerpaspecten van het AI-CDSS te onderzoeken en de ontwikkelaars van het systeem van advies te voorzien in het vervolgtraject van de ontwikkeling.

¹ <https://nell.eu/nieuws/decide-vera-ethisch-design-als-basis-van-een-nieuw-clinical-decision-support-systems-voor-hart-en-vaatziekte-patienten-en-hun-artsen>

1 Inleiding

1.1 Projectdoelstellingen

De primaire doelstelling van het DECIDE-VerA project is om het ontwikkelteam van het AI-CDSS-DECIDE tool van **praktisch advies** te voorzien **voor de verder ontwikkeling van de tool** zodat het conform de pijlers van betrouwbaar AI tot een eindproduct leidt. De pijlers: ethisch, robuust en wettig komen in de driehoek van de interdisciplinaire strategie van het project dan ook terug: de ethische, juridische en design perspectieven.

Uiteindelijk en op langere termijn zou het **AI-CDSS DECIDE-tool een efficiënt middel** zijn dat **patiënt en arts ondersteunt in het samen nemen van beslissingen** om de risico's op hart- en vaatziekten te verminderen en de gezondheid van de patiënt te verbeteren. Deze langere termijn doelstelling -dat in dit rapport in het kort met de term secundair doelstelling wordt aangeduid - staat niet los van de primaire doelstelling, maar vertegenwoordigt een additionele uitdaging: de toekomstige implementatie van het AI-CDSS tool in de zorgpraktijk.

1.2. Situering van het evaluatieonderzoek en gebruikte informatie

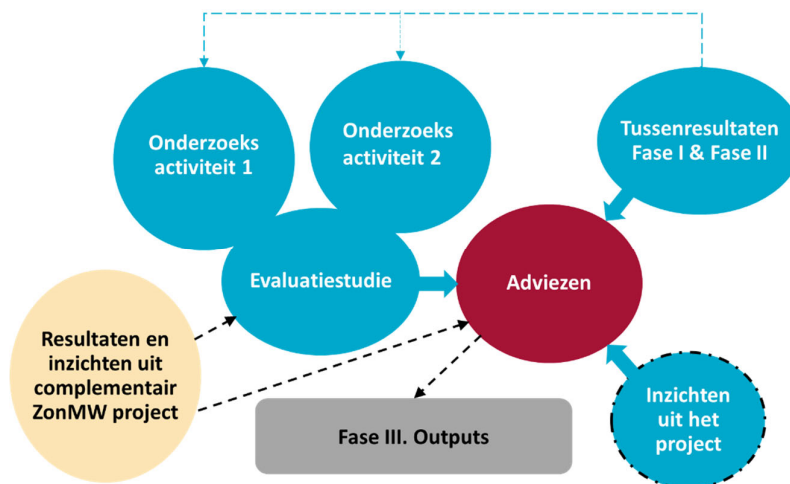
Het DECIDE-VerA project kent drie fasen: (I) de formatieve analyses, (II) de fase van co-design en convergentie van de ethische, design en juridische aspecten en fase III, waarin het opleveren van de resultaten o.a. in de vorm van praktisch hanteerbare tools gaat plaatsvinden (zie 6.1). De huidige evaluatiestudie is onderdeel van fase II. Fase II richt zich op verdieping van fase I d.m.v. participatieve methodologie om ethiek, recht en design te integreren. In deze fase worden begeleidingsethiek sessies gehouden waarin verschillende stakeholderperspectieven zijn vertegenwoordigd.

De aanbevelingen (zie 4.2) in dit rapport zijn gebaseerd op:

- de beschikbare informatie en tussenresultaten uit fasen I en II van het project (bijv. interviews, begeleidingsethiek sessies, manuscripten),
- de uitkomsten van de evaluatiestudie beschreven in sectie 3,
- impliciete kennis (inzichten) opgedaan door de auteur tijdens de uitvoering van het project (inclusief 1 op 1 gesprekken met projectpartners t.b.v. de evaluatiestudie en
- resultaten en inzichten uit het complementaire ZonMW project getiteld 'Ethical appraisal of AI-driven MedTech for Trustworthy AI – development of a participatory, patient-centric methodology and toolkit'².

Een schematische weergave van de gebruikte informatie is weergegeven in Figuur 1.

² Project funding by ZonMW (# 10580012210019). <https://projecten.zonmw.nl/nl/project/ethical-appraisal-ai-driven-medtech-trustworthy-ai-development-participatory-patient>



Figuur 1. Gebruikte informatie voor het opstellen van de adviezen t.b.v. de primaire en de secundaire doelstellingen van het DECIDE-VerA project. De adviezen worden vervolgens gebruikt voor de verschillende deliverables van het project in fase III, waaronder een rapport met daarin een voorstel om ethisch en rechtvaardig by design te bereiken in de vervolg versie van het AI-CDSS-systeem.

1.3 Doel evaluatie

De projectaanvraag voorzag een evaluatie in fase II om de formatieve analyses te onderzoeken en te beoordelen op basis van uitkomsten van de begeleidingsethiek en design sessies en de verantwoordelijkheids- en aansprakelijkheidsanalyses. De aspecten waarop geëvalueerd wordt worden beschreven in hoofdstuk 3. **Het uiteindelijke doel is om van uit een integraal perspectief (de samenhang van wettige, ontwerp- en ethische principes) gebaseerd o.a. op de evaluatiestudie, aanbevelingen te geven t.a.v. beide doelstellingen van het DECIDE-VerA project** (zie 1.1).

1.4 Theoretische kaders

De tussenresultaten van fase I en II zijn tot stand gekomen vanuit de drie interdisciplinaire onderzoeksactiviteiten, elk gegrond op de "eigen" theoretische kader (zie Figuur 2 en 4.3.IV van de onderzoeksaanvraag).



Figuur 2. Theoretische kaders behorend bij de drie onderzoekspijlers

Ten behoeve van de evaluatie van de secundaire doelstelling van het project, zijn aan deze drie kaders twee andere toegevoegd:

- Het interpretatieve model van de arts-patiënt relatie [1]. Het interpretatieve model richt zich op het begrijpen en verhelderen van de waarden en voorkeuren van de patiënt door de arts. De arts helpt de patiënt hierbij door de informatie te helpen interpreteren en in lijn te brengen met de mogelijkheden van de patiënt. De rol van de arts: is dan ook adviserend: hij/zij verschaft informatie en helpt de patiënt bij het ontdekken en verduidelijken van zijn/haar eigen waarden. Het interpretatief model combineert elementen van zowel de expertise van de arts als de voorkeuren en ervaringskennis van de patiënt, met als doel een gezamenlijke besluitvorming te bevorderen die rekening houdt met de persoonlijke waarden en mogelijkheden van de patiënt (zie hieronder vier stappenmodel van Elwin, 2012).
- De tweede kader is het vier stappenmodel van samen beslissen van Elwin, 2012 [2]. De vier stappen in het model voor gedeelde besluitvorming zijn essentieel voor het ondersteunen van patiënten bij het maken van geïnformeerde keuzes. De stappen zijn kort als volgt te beschrijven³:
 - **1. Keuze:** de zorgverlener informeert de patiënt dat er een beslissing genomen moet worden en dat de waarden, voorkeuren en mogelijkheden van de patiënt belangrijk zijn in het bepalen van de beste passende optie voor deze patiënt;
 - **2. Opties:** de zorgverlener legt de opties en de voor- en nadelen uit van elke optie; een optie kan zijn dat er geen actieve behandeling wordt gestart;
 - **3. Voorkeur:** de zorgverlener en de patiënt bespreken de voorkeuren van de patiënt en de zorgverlener ondersteunt de patiënt in het wikken en wegen (waarin aandacht is voor elementen zoals: wensen, voorkeuren, doelen, waarden en verwachtingen mee te nemen);
 - **4. Beslissing:** de zorgverlener en de patiënt nemen een besluit waarin de voorkeuren en de voor- en nadelen van de opties geïntegreerd worden, of stellen het expliciet uit en regelen eventuele follow-up.

Belangrijk bij samen beslissen proces is, is dat er voortdurende nadruk moet zijn op ondersteuning van deliberatie, waarbij de zorgverlener de patiënt aanmoedigt om tijd te nemen om informatie te verwerken en zijn/haar gedachten te bespreken, wat cruciaal is voor het bereiken van een weloverwogen beslissing.

1.5 Hulpmiddel Handelingsruimte & Technology Readiness level

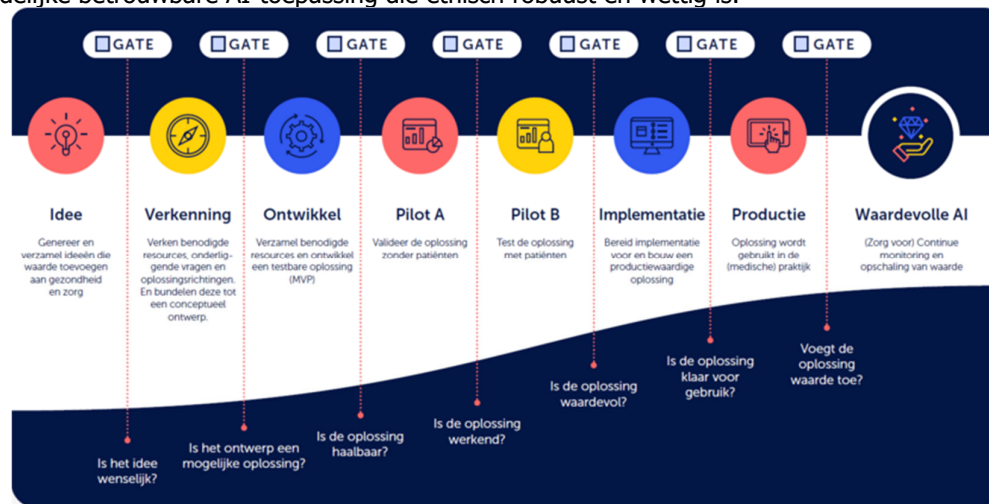
Voor analyse van het AI-DCSS DECIDE model is in dit onderzoek gebruik gemaakt van het Hulpmiddel Handelingsruimte Waardevolle AI voor gezondheid en zorg (kortweg handelingsruimte) en de framework Technology Readiness Level (kortweg TRL).

De handelingsruimte is ontwikkeld in het kader van het VWS-programma Waardevolle AI voor Gezondheid⁴. Het doel ervan is om projectteams te ondersteunen in het innovatieproces waarin ze AI-toepassingen voor in de zorg ontwikkelen. De handelingsruimte biedt een heldere structuur; het bestaat uit 7 fasen (zie figuur 3). Elke fase is gericht op het creëren van zoveel mogelijk waarde voor de gezondheidszorg en/of gezondheid van bepaalde patiëntengroep(en). Bij elke faseovergang (gate) is er een checklist die helpt om de juiste vragen te stellen op vijf domeinen: waarde, toepassing,

³ <https://www.nfu.nl/sites/default/files/2023-12/4-stappenmodel-Samen-Beslissen.pdf>

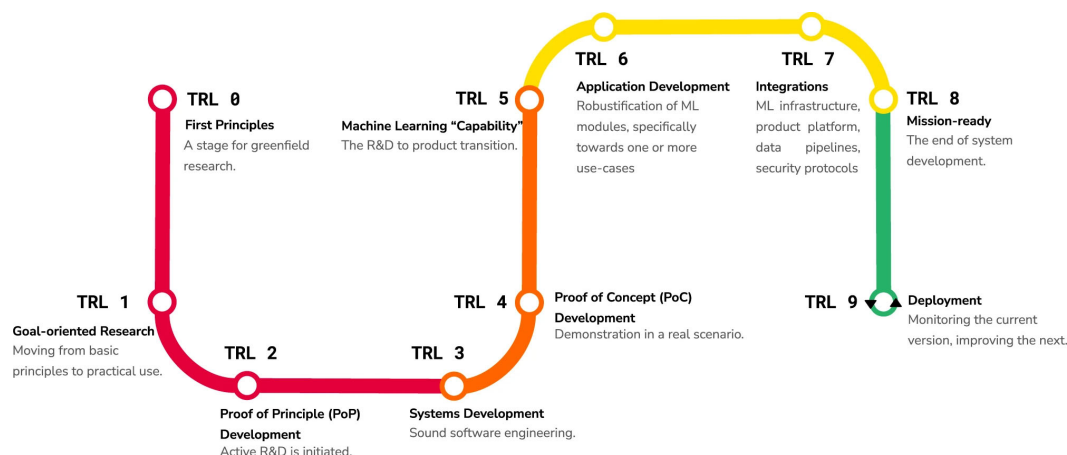
⁴ <https://nlaic.com/wp-content/uploads/2022/06/04a.-Hulpmiddel-Handelingsruimte-Waardevolle-AI-voor-gezondheid-en-zorg.pdf>

ethiek, techniek en verantwoordelijkheid. Door het doorlopen van de fasen stijgt de kans op een uiteindelijke betrouwbare AI-toepassing die ethisch robuust en wettig is.



Figuur 3. Handelingsruimte Waardevolle AI voor gezondheid.

De TRL framework is oorspronkelijk ontwikkeld door NASA en aangepast voor gebruik in andere sectoren, waaronder de gezondheidszorg. Het TRL is een systematisch meetinstrument dat wordt gebruikt om de maturiteit van een specifieke technologie te evalueren, waaronder machinaal leren [3]. In figuur 4 zijn de TRL fasen schematisch weergegeven en elke fase beknopt beschreven [3].



Figuur 4. TRL-fasen inclusief beknopte beschrijving van de afzonderlijke fasen (figuur uit: Lavin et al. Nat.Comm, 2022) [3].

2 AI en de DECIDE AI tool

2.1 Artificiële Intelligentie (AI)

Artificiële Intelligentie heeft geen eenduidige definitie. De *High Level Expert Group on Artificial Intelligence*, aangesteld door de Europese Commissie verstaat onder kunstmatige of artificiële intelligentie (AI) "Systemen die intelligent gedrag vertonen door hun omgeving te analyseren en – met een zekere mate van zelfstandigheid – actie te ondernemen om specifieke doelen te bereiken" [4, p.].

Machine learning (ML) is een vorm van AI die computers in staat stelt om specifieke taken te leren (narrow AI) van grote hoeveelheden data [5]. Er bestaan verschillende typen ML [6]. In essentie zijn ML-gebaseerde systemen in staat om te leren en zichzelf te verbeteren aan de hand van eerdere ervaringen, zonder dat ze daarvoor expliciet geprogrammeerd zijn. Een ML-gebaseerd algoritme wordt, in principe, beter naarmate het meer input van data krijgt. Deze systemen worden gevoed met veel data, waarbij het zeer belangrijk is dat de kwaliteit van de data goed is. Op basis van de data stelt de computer een set van regels of instructies op, die een computer toe kan passen op nieuwe situaties. Is de kwaliteit van de data slecht, dan zal het algoritme minder goede uitkomsten hebben [7]. ML is krachtig, maar ook omstreken, zeker in de zorg, omdat de uitkomsten van de algoritmen (met name van het type *deep learning*) niet te reconstrueren zijn, waardoor uitleg over waarom ze met een bepaalde uitkomst komen, moeilijk of onmogelijk blijft [8]. Voor AI in de medische sector is van belang dat de werking van het model uitlegbaar is, d.w.z. de uitkomsten interpreteerbaar [9]. In de huidige ontwikkelfase van het AI-systeem wordt gebruikgemaakt van een Cox-elastic net model, die van nature interpreteerbaar is (parametrisch) (zie sectie 2.2.1 en 2.2.2).

2.2 AI-CDSS DECIDE -model en uitgangsposities project

Het AI-CDSS-model is bedoeld om nauwkeurig het risico op hart- en vaatziekten bij (jonge) volwassenen (mannen en vrouwen van 30 tot 49 jaar) te voorspellen, gebruikmakend van routinematig verzamelde zorggegevens⁵ [10].

Voor begrip over het doel van het model, de inzet ervan in samen beslissen proces en t.b.v. de adviezen, wordt hieronder kort ingegaan op drie aspecten: de ontwikkeling, de uitlegbaarheid en maturiteit van het model.

2.2.1 Ontwikkeling van het model – Van Os et al., 2023 [10]

Gegevens van jonge volwassenen uit elektronische medische dossiers met routine gegevens zijn gebruikt voor het trainen van het model. De gegevens zijn verdeeld in de gebruikelijke trainingsset (gebruikt om het model te leren) en testset (gebruikt om het model te beoordelen). Met behulp van de "Cox elastic net" methode zijn ze de belangrijkste voorspellers voor hart- en vaatziekten gekozen. Hierbij is belangrijk dat deze zowel de traditionele als ook de niet-traditionele cardiovasculaire voorspellers omvatten. Traditionele voorspellers zijn o.a. leeftijd, geslacht, bloeddruk, cholesterolniveaus, rookstatus, diabetesstatus en familiegeschiedenis van hart- en vaatziekten. De niet-traditionele voorspellers geïdentificeerd door in de studie van Van Os et al. waren o.a. socio-economische status en hormonale anticonceptiegebruik.

Van Os et al., testte diverse modellen: een referentiemodel met traditionele risicofactoren, modellen

met de 20 belangrijkste voorspellers en modellen met de 50 belangrijkste voorspellers. De label (de te voorspellen uitkomstvariabele) was de **eerste cardiovasculaire gebeurtenis**. De prestaties van de diverse modellen zijn vergeleken met behulp van de C-index (een standaard evaluatiemethode in de medische statistiek). Bij het meenemen van de niet-traditionele voorspellers presteerden de modellen iets beter dan met alleen traditionele voorspellers. Dit suggereert dat er nog ruimte is voor verdere optimalisatie. De auteurs benadrukken dit en wijzen op de noodzaak van verdere validatie en mogelijk herkalibratie van de modellen voor verschillende populaties.

2.2.2 Uitlegbaarheid -Van Os et al., 2023 [10]

In het onderzoek zijn verschillende typen AI-modellen gebruikt voor cardiovasculaire risicopredictie: De Cox proportional hazards (PH) model, de Cox elastic net modellen en Random survival forests. Deze modellen zijn in verschillende mate uitlegbaar: de Cox PH modellen het meest door direct te tonen hoe elke risicofactor het cardiovasculaire risico beïnvloedt. De Cox elastic net modellen zijn minder uitlegbaar dan Cox PH, maar nog steeds interpreteerbaar in termen van identificeerbaarheid van voorspellers en hun relatieve bijdrage aan de voorspelling. De Random survival forests zijn het minst uitlegbaar: deze complexe modellen kunnen niet-lineaire relaties en interacties tussen risicofactoren mogelijk beter vangen, maar functioneren relatief meer als een "black box" [11]. De onderzoekers vergeleken deze modellen met een referentiemodel (Cox PH) gebaseerd op traditionele cardiovasculaire risicofactoren. Ze ontdekten dat de meer complexe, data-gedreven modellen een bescheiden verbetering in voorspellingsprestaties opleverden, vooral bij vrouwen. De Cox PH en Cox elastic net modellen leverden vergelijkbare prestaties op. Het zijn beide statistische modellen, waarbij de elastic net versie enkele machine learning elementen bevat (supervised learning). Volgens de auteurs zouden Cox PH modellen de voorkeur hebben voor klinisch gebruik aangezien ze gemakkelijker te interpreteren zijn dan Random forest modellen.

Voor het huidige DECIDE-VerA project is de uitgangspunt het AI CDSS DECIDE model dat gebaseerd is op de machine learning type model, het **Cox elastic net model**. Van dit type model is namelijk sprake van gebruik van AI in de vorm van supervised learning (zie 2.1) en dit type model presteerde beter dan de Random Forest model. Belangrijk hierbij is om te onthouden dat De Cox elastic net modellen minder uitlegbaar zijn dan Cox PH modellen, maar nog steeds **interpreteerbaar** zijn in termen van identificeerbaarheid van voorspellers en hun relatieve bijdrage aan de voorspelling.

2.2.3 Maturiteit – mapping naar kaders (zie 1.5).

Inschatting van de fase waarin het AI-CDSS DECIDE model zich bevindt is belangrijk voor de aanbevelingen en dan met name t.a.v. de primaire doelstelling van het project (zie 1.1).

Het model bevindt zich in een vroeg stadium van ontwikkeling en moet nog verschillende fasen doorlopen voordat het klaar is voor grootschalige implementatie in de zorg. Volgens de hulpmiddel handelingsruimte waardevolle AI (zie 1.5) bevindt het model zich deels in de verkenning, deels in de ontwikkelfase. Bij het doorlopen van de checklist van de handelingsruimte in de fase verkenning en ontwikkeling (zie hulpmiddel handelingsruimte waardevolle AI, Gate 1 – Afronding Ideeefase Is het idee wenselijk?) blijkt dat er bepaalde stappen zijn overgeslagen.

Dat een innovatie niet conform een lineair proces ontwikkeld wordt komt vaak voor tijdens het innovatieproces vanwege de iteratieve aard van softwareontwikkeling en de inzichten die worden opgedaan tijdens de ontwikkeling en het testen. Veelvoorkomende "switchbacks" in het proces zijn al beschreven in het kader van Machine Learning modelontwikkeling [3]. Volgens het TRL kader zou het

model zich in TRL 4/5 bevinden. In het huidige DECIDE-VerA project waarbij de ethische, juridische en design vraagstukken nog vóór validatie van het model tegen het licht worden gehouden kunnen we spreken van een “switchbak” naar **TRL fase 3** (research & development).

Conform beide kaders dient het model nog een aantal stappen te doorlopen voordat het geïmplementeerd kan worden: er dient een optimalisatie van het model plaats te vinden, het model dient ook getest te worden in een klinische omgeving en er moet nog evaluatie van de impact en waarde in de praktijk plaatsvinden. Hoewel het beoogde doel duidelijk is (gepersonaliseerde voorspelling van cardiovasculair incident), is het beoogd gebruik in de praktijk nog onduidelijk.

2.2.4 Beoogd doel en gebruik / MDR - uitgangspunt

Het DECIDE-VerA project is ingestoken en gekaderd vanuit de eerstelijnszorg en t.b.v. klinische besluitvorming t.b.v. secundaire preventie. Dit impliceert dat de **eind-user de zorgprofessional** zal zijn resp. dat het uiteindelijke AI-systeem (model ingebed in software) onder de **MDR** zal vallen. Onder de MDR dient de ontwikkelaar het beoogd doel / het beoogd gebruik te beschrijven (in het geval van AI-CDSS DECIDE is dat preventie van cardiovasculair incident bij mensen onder de 50 jaar; zie ook 2.2.). Het beschrijft wat de fabrikant wil bereiken met het hulpmiddel en hoe het hulpmiddel moet worden gebruikt, voor welke medische toepassingen, en door wie [12], [13].

2.2.5 Gebruiksfases - uitgangspunt

Voor de concrete onderzoekactiviteiten in de formatieve fase van het project is gebruik gemaakt van fasering van gebruik van AI CDSS DECIDE en daarbinnen de reis van de burgers/patiënten.

Het gebruik van AI CDSS DECIDE is opgedeeld in 2 fasen:

- Fase 1: identificatie individuen met verhoogde risico, communicatie erover en benadering patiënt incl. evt. uitnodiging voor gesprek. Deze fase wordt kortweg ook “tijdig signaleren” genoemd.
- Fase 2: patiënt en zorgverlener in gesprek over wat de risico betekent en in het Samen Beslissen proces op zoek naar de best passende vervolgstap / behandeling. Deze fase wordt kortweg ook “tijdig actie ondernemen” genoemd.

Voor elke fase is een scenario geschetst op basis van wat de burger/patiënt doorloopt in de tijd. Deze patiëntreizen zijn voor zowel fase 1 als fase 2 schematisch weergegeven (zie 6.2.) Beide scenario's betreffen toekomstige scenario's waarbij het AI CDSS DECIDE tool de nodige certificeringen behaald heeft en ingebed is in software en via een leverancier beschikbaar komt voor aankoop door huisartsenpraktijken. Het model zelf is en blijft open source⁶.

⁶ Persoonlijke correspondentie met de hoofd van het ontwikkelteam H. van Os.

3 Evaluatiestudie

De projectaanvraag voorzag een evaluatie in fase II om de formatieve analyses te onderzoeken en te beoordelen op basis van uitkomsten van de begeleidingsethiek sessies (BE-sessies), de design sessies, en de verantwoordelijkheids- en aansprakelijkheidsanalyses.

3.1 Plan van aanpak

De tussenresultaten van fase I en deels fase II van het project (zie 1.2 en figuur 1) dienden als input voor deze onderzoeksactiviteiten. Daarnaast was – in overleg met de projectleider- ruimte voor exploratief onderzoek. Zowel kwantitatief als kwalitatieve onderzoeksmethoden zijn gebruikt (zie 3.2.2. en 3.3.2).

Ten tijde van de planning van deze studie waren er verschillende outputs beschikbaar, die zich respectievelijk richtten op het ontwerp-, juridische of ethische aspect. Al deze aspecten waren relevant voor het project. Vanuit het perspectief van de patiënt – het expertisegebied van de auteur van deze evaluatie – waren echter vooral de BE-sessies en de interviews met patiënten uitgevoerd door het designteam tijdens fase II van belang. De interviews waren op dat moment nog niet geanalyseerd en bevonden zich niet in een stadium waarin evaluatie zinvol was. De auteur van dit rapport was betrokken bij de planning en trad op als co-moderator tijdens beide BE-sessies, wat betekent dat hij deze sessies uit eerste hand heeft meegemaakt. De belangrijkste resultaten van de 2 begeleidingsethiek sessies waren eveneens beschikbaar (tabellen met in totaal ca. 75 unieke stakeholders, de in totaal 9 door deelnemers geselecteerde effecten in elk van de BE-sessies, d.w.z. in totaal 18 effecten). Om die reden is ervoor gekozen om de resultaten van de BE-sessies te analyseren en in gesprek te gaan met het juridisch team, zodat hun perspectieven konden worden meegenomen in de opzet van de evaluatie. De auteur was op de hoogte van alle resultaten van het project. De analyse van de BE-sessies werd gezien als een waardevolle bron van input voor het project, inclusief aanbevelingen voor de ontwikkelaars van DECIDE in de slotfase van het onderzoek.

De evaluatiestudie bestond uit twee onderzoeksactiviteiten, beiden in de vorm van participatieve bijeenkomsten: één gericht op patiënten/naasten in de vorm van een fysieke workshop. De andere onderzoeksactiviteit was gericht op huisartsen in de vorm van een online focusgroep. De reden voor deze splitsing (naast praktische redenen van beschikbaarheid van de stakeholders) was dat patiënten met een vrijer gevoel deelnemen aan bijeenkomsten (zich vrij voelen om hun ervaringen te delen en hun opvattingen te uiten over de onderzoeksthema's) bij afwezigheid van zorgverleners.

Hieronder volgen de beschrijvingen van de onderzoeksactiviteiten, samenvatting en interpretatie van de resultaten.

3.2 Onderzoeksactiviteit -1: workshop met patiënten/naasten

3.2.1 Onderzoeksvragen

Exploratieve vragen:

- a) Wat zijn mogelijke beoogde gebruiken van een AI-voorspelmodel voor het tijdig ontdekken van hoge risico op hartaanval/beroerte? (opdracht 1)
- b) Wie zijn betrokken bij dit mogelijk gebruik en welke rollen en relaties hebben de betrokkenen?

(opdracht 2)

Evaluatievragen

- c) Is er saturatie t.a.v. stakeholders zoals in fase 1 (benoemen actoren) van de begeleidingsethiek workshops geïnventariseerd is en benoemd in interviews? (opdracht 3)
- d) Is er overeenstemming met de outputs van fase 2 (positieve en negatieve effecten) van de begeleidingsethiek workshops? (opdracht 4)
- e) Wat wordt onder verantwoordelijkheid (vooruitkijken en terugkijkend) en aansprakelijkheid verstaan t.b.v. voorkomen van schade? (opdrachten 5 en 6)

3.2.2 Deelnemers & Methodes

De fysieke workshop vond plaats op 19 juli 2024 bij de Patiëntenfederatie Nederland in Utrecht. In totaal hebben 6 patiënten/naast deelgenomen. Via ikzoekeenpatient.nl zijn deelnemers gerekruteerd. Inclusiecriteria voor deelname waren:

- Alle deelnemers spreken Nederlands en zijn 18+
- 2 deelnemers (1 vrouw, 1 man), die zelf een hartaanval/beroerte hebben gehad
- 2 deelnemers (1 vrouw, 1 man), die naasten zijn van iemand, die een hartaanval of beroerte heeft gehad
- 2 deelnemers (1 vrouw, 1 man), die geen beroerte of hartaanval hebben gehad, maar wel een verhoogd risico hebben daarop.
- Geen eerder deelname aan begeleidingsethiek sessie, interview of was anderszinds betrokkenheid in het project.

De workshopleider was de auteur van dit rapport. De workshop duurde 3 uur met een korte pauze halverwege. De workshop was opgenomen op audio (2 bestanden corresponderend met voor en na de pauze). Deelnemers hebben afzonderlijk het toestemmingsformulier ingevuld. Het formulier was een aangepaste versie van het eerder in fase I gebruikte TU/e-formulier en bevatte dezelfde 6 vragen (bijlage 6.5.1). De workshopleider heeft aangegeven dat de na de sessie een mapping wordt gemaakt tussen de aangegeven toestemming / bezwaar zodat alleen die onderdelen van de gegevens / workshop worden gebruikt waarvoor er wel toestemming is gegeven).

Na het korte voorstelrondje heeft de workshopleider een introductie gegeven over het onderwerp van voorkomen van hartaanval en beroerte bij jonge mensen en over de rationale van hoe de AI-toepassing werkt voor de voorspelling van risico's. De workshop had **6 opdrachten in totaal**. De volgorde van de opdrachten was zodanig gekozen dat het exploratieve onderzoeksvraag aan het allerbegin beantwoord kon worden, nog voordat de deelnemers over de in het project ontwikkelde scenario's kennisnamen. Dit zorgde voor een onbevangen en open blik tijdens het maken van de eerste twee opdrachten. De twee fasen (tijdig ontdekken resp. tijdig aanpakken) van preventie van cardiovasculair incident waren wel aangestipt zodat de deelnemers handvaten hadden de eerste twee opdrachten. Er volgde een uitleg over predictiemodellen in het algemeen en over een toekomstige AI-model voor voorspelling van risico op cardiovasculair incident in het bijzonder. Uitgelegd is hoe een dergelijk model tot stand komt (training op heel veel data) en dat bij het voorspellen van het individuele risico het model de data van individuele patiënten "verwerkt" alvorens het persoonlijke risicogetal te kunnen berekenen. Opdracht 3 was gericht op onderzoeksvraag c), opdracht 4 op d) en opdrachten 5 & 6 op e). De exacte opdrachten zijn in bijlage 6.3 te lezen. Hieronder volgt een korte beschrijving van de methoden per opdracht.

De gebruikte methode voor opdracht 1 was de value scenario [14]. Deelnemers hebben de opdracht geprint gekregen en onder de opdracht hebben het verhaal met de hand opgeschreven. Nadat iedereen klaar was, zijn de verhalen door de deelnemers mondeling toegelicht.

Opdracht 2 was uitgevoerd door gebruik te maken van stakeholder tokens [15]. Iedere deelnemer heeft hiervoor evenveel aantal houten poppetjes (zie 6.3.1) gekregen, en 1 A3 formaat vel. Diverse kleuren stiften, post-its en plakband waren aanwezig en deelnemers konden die uit een gemeenschappelijk doos pakken. Deelnemers hadden de vrijheid om zelf de stakeholders en hun relaties m.b.t. de aanwezige materiaal te schetsen / vorm te geven.

Opdracht 3 was om de in de BE sessies benoemde stakeholders te bestuderen en aan te geven of ze nog stakeholders missen⁷. Deze opdracht vereiste een schets van het beoogd gebruik zoals in het project gebruikt (via de eerstelijns; zie schema's in 6.2). Dit was gedaan door een fictief narratief te schetsen (zie 6.3.2).

Voor opdracht 4 was gebruik gemaakt van 5 categorie Likert scale: Helemaal eens, eens, neutraal, oneens, helemaal oneens. Deelnemers hebben in totaal 18 stellingen beoordeeld. Deze stellingen waren geformuleerd op basis van de top-3 mogelijk positieve en negatieve gevolgen van gebruik uit beide BE-sessies (zie 6.3.3). De 3 subgroepen van beide workshops maakten in totaal dat we 18 effecten en corresponderende stellingen zijn geformuleerd. Deelnemers hadden de opdracht om aan te geven in hoeverre ze het eens waren met de stelling door het juiste vakje aan te kruisen. De stellingen waren op geprinte papier beschikbaar voor elke deelnemer. De opdracht was individueel en onafhankelijk van elkaar gemaakt.

Opdracht 5 was een open vraag opdracht en betrof 2 fictieve scenario's waarin schade is opgetreden: bij de ene was er een te lage risico ingeschat bij de andere en te hoge met als gevolg / schade: verliezen van betaald werk resp. bijwerkingen medicijnen en overmatig stress. De opdracht was om in enkele zinnen op te schrijven hoe ze tegen de fictieve schadelijke gevolgen aan die ze zijn overkomen?.

Opdracht 6 bestond wederom uit open vragen: een gericht op de betekenis van wat verantwoordelijk gebruik van AI model betekent ter voorkoming van schade. De andere vraag betrof de vraag welke partij of partijen moeten dan aansprakelijk worden gehouden voor de opgetreden schade als de oorzaak van de schade niet te achterhalen is o.a. door de black-box karakter van het AI-model.

3.2.3 Samenvatting en interpretatie van resultaten

Hieronder volgt per vraagstelling de samenvatting van de resultaten. Voor gedetailleerde uitkomsten wordt verwezen naar bijlage 6.3, en naar de ruwe data (door deelnemers ingevulde opdrachten en de 2 audio-files, beide in bezit van de projectleider).

Exploratieve vragen:

- a) Wat zijn mogelijke beoogde gebruiken van een AI-voorspelmodel voor het tijdig ontdekken van hoge risico op hartaanval/beroerte?
- b) Wie zijn betrokken bij dit mogelijk gebruik en welke rollen en relaties hebben de betrokken?

⁷ Hoewel het in kaart brengen van stakeholders ook in andere onderzoeksactiviteiten plaatsvond, leverden de BE-sessies de meest uitgebreide lijst van stakeholders op.

Eindgebruiker

Deelnemers hebben 6 scenario's voorgesteld (zie bijlage 6.3.1):

- In 4 scenario's is de eindgebruiker van het AI-model de patiënt/burger. Die installeert het als app op het eigen apparaat.
- In 1 scenario is de eindgebruiker de huisarts, maar nadat er een grove risicoschatting heeft plaatsgevonden
- In 1 scenario is de eindgebruiker de huisarts

Aanleiding gebruik

De aanleidingen voor gebruik zijn: met vage klachten de huisarts bezoeken, sterke / overmatige behoefte aan controle, aanbeveling door zorgverlener, gebruik van medisch hulpmiddel (pacemaker), routinematig gebruik door huisartsenpraktijk, 'tweetrapsgebruik'-scenario: eerst een grove schatting laten maken van de gezondheid van de patiënt (zoals een Algemene Periodieke Keuring - "APK") bijvoorbeeld door zorgprofessionals, maar ook de media kunnen mensen attenderen hierop en vervolgens als er een vermoeden risico op hart- en vaatziekte is, dan pas zou het AI-model ingezet kunnen worden door zorgprofessionals om de risico te berekenen.

Rol zorgverlener

Bij alle geschetste scenario's is er een rol voor de zorgprofessional, maar op verschillende momenten van het gebruik van het AI-model. De rol varieert tussen neutraal (in contact staan), raadgever / duider en behandelaar.

Deelnemers hebben uiteenlopende scenario's bedacht voor gebruik van een toekomstig AI-model ter voorspelling van risico op cardiovasculair incident. Hierbij valt op dat het AI-model in meeste gevallen op de eigen apparaat als app gebruikt wordt, maar dat voor duiding van de risico of t.b.v. vervolgstappen (tijdig actie ondernemen) een rol wordt gegeven aan de zorgprofessional. Integratie van meerdere databronnen komt ook terug in veel verhalen. Privacy- en veiligheidsvraagstukken gerelateerd aan data opslag en gebruik zijn genoemd, maar niet uitvoerig besproken.

Evaluatievragen

- c) Is er saturatie t.a.v. stakeholders zoals in fase 1 (benoemen actoren) van de begeleidingsethiek workshops geïnventariseerd is en benoemd in interviews?

In totaal zijn er 11 extra stakeholders benoemd /opgeschreven door de deelnemers.

- Ontwikkelaars van IT-systemen van specialisten en andere (para)medici, verzekeringsarts
- Stappentellers, sociale media, "APK-bureau",
- Gecertificeerde leefstijlcoaches, diëtist
- Buren (noodhulp), beheerder reanimatie-systeem in woonwijken
- Gemeente, overheid

Deelnemers hebben stakeholders genoemd die deels te maken hebben met hun scenario uit de voorgaande opdracht ondanks de instructie om het scenario te volgen van de fictieve personage die wel correspondeerde met de scenario's -tijdig signaleren en tijdig actie ondernemen- zoals geschetst bij beide BE-sessies. Aan de andere kant hebben de deelnemers wel toegelicht waarom deze extra stakeholders legitiem zijn: ze zijn net zulke indirecte stakeholders als bijv. ChatGPT of *internet of things* die in de tabellen te vinden waren.

Bij analyse van de genoemde stakeholders viel op dat er overlap is met stakeholders die in de in de tabellen zijn weergegeven. Deze betreffen: de gemeente, overheid, diëtist (in de tabel onder de noemer paramedicus), stappenteller (valt onder wearables). Buren staan in de tabel onder vrienden (buurvrouw). Hoewel leefstijlcoach ook al genoemd is tijdens de BE sessie, is hier de nadruk gelegd op de certificering, gezien "iedereen zich leefstijlcoach kan noemen". Ontwikkelaars van IT systemen bij zorgaanbieders anders dan huisartsen valt in de categorie "systeemontwikkelaars" zoals benoemd tijdens de eerste BE-sessie.

Dit maakt dat **er maar 4 nieuwe stakeholders benoemd zijn**: de beheerder van de reanimatie system in woonwijken, sociale media/ platformen die voor veel mensen "het dorsplein" is waar ze mensen ontmoeten en het APK bureau, de laagdrempelige inloopinstantie voor de eerste grove schatting van risico en de hier speciaal benoemde verzekeringsarts.

d) Is er overeenstemming met de outputs van fase 2 (positieve en negatieve effecten) van de begeleidingsethiek workshops?

De 5 opties Likert scale en het even aantal deelnemers maken het mogelijk om de volgende categorisatie aan te houden: helemaal eens en eens tezamen betekenen overeenstemming in het geval als de meerderheid van de aanwezigen (4 van de 6) tezamen deze twee categorieën hebben aangekruist. Andersom geldt het ook: helemaal oneens en oneens bij ten minste 4 deelnemers geldt als gebrek aan overeenstemming (onenigheid). De overeenstemming/onenigheid wordt als sterk beschouwd in het geval dat alle 6 deelnemers helemaal links (helemaal eens) resp. helemaal rechts (helemaal oneens) hebben gescoord. In alle andere gevallen (5 of 4 deelnemers) is de sterkte gemiddeld.

Bij geen van de 18 stellingen is er sprake van onenigheid (zie bijlage 6.3.3). Bij 13 is er sprake van overeenstemming, waarvan bij 3 een sterke. Deze drie betreffen stellingen 11, 14 en 18. De stellingen waarbij sprake is van overeenstemming zijn 1, 3, 6, 7, 9, 10, 13, 15, 16 en 17. Bij de resterende 5 stellingen is de categorie "neutraal" gescoord.

Stelling #	Mate van overeenstemming	Stelling
11	sterk	Het gebruik van het AI-voorspelmodel kan leiden tot verwarring als de burger onvoldoende <i>uitleg</i> krijgt
14	sterk	Als de huisarts het risicogetal door de telefoon <i>communiceert</i> , zal niet iedereen het kunnen begrijpen waar het precies over gaat
18	sterk	Het is belangrijk dat de huisarts in het <i>gesprek</i> over verlagen van het risico een positieve houding heeft
1	gemiddeld	Het gebruik van het AI-voorspelmodel kan mensen onnodig angstig maken
3	gemiddeld	Het gebruik van het AI-voorspelmodel draagt bij het "dataficeren" van mensen
6	gemiddeld	Het gebruik van het AI-voorspelmodel kan leiden tot medicalisering van de burger
7	gemiddeld	Het gebruik van het AI-voorspelmodel biedt de kans om mensen op te sporen die anders niet in beeld zouden komen

9	gemiddeld	Het gebruik van het AI-voorspelmodel kan leiden tot discriminatie van bepaald groepen zoals mensen met een lage score van maatschappelijke positie (SES)
10	gemiddeld	Het gebruik van het AI-voorspelmodel kan ervoor zorgen dat de burger schrikt en in paniek raakt en denkt dood te gaan
13	gemiddeld	Het gebruik van het AI-voorspelmodel kan leiden tot een negatievere toekomstperspectief van de burger
15	gemiddeld	De huisarts moet tijd gunnen zodat mensen hun naaste mee kunnen vragen naar het consult na het horen van het hebben van een hoge risico
16	gemiddeld	Iedereen moet kennis kunnen nemen van alle uitslag van de berekening, niet alleen als het risico hoog is
17	gemiddeld	Bij het gebruik van het AI-voorspelmodel is belangrijk dat er uitvoerbare opties gegeven worden om uit te kiezen om het hoge risico te verlagen

Opvallend is dat bij alle drie stellingen waarmee de deelnemers sterk het eens waren gaan over de **communicatie van het risico** waarbij bij twee stellingen duidelijk de communicatie en houding door de huisarts is bedoeld (cursieve termen in tabel). Deelnemers hebben de opdracht onafhankelijk van elkaar uitgevoerd. Het feit dat bij deze 3 stellingen alle 6 deelnemers helemaal eens & eens hebben gescoord zegt – ondanks de kleine aantallen deelnemers – laat zien dat communicatie een hele belangrijke aspect is bij het gebruik van het AI-CDSS-model voor de stakeholder patiënt/naaste.

- e) Wat wordt onder verantwoordelijkheid (vooruitkijken en terugkijkend) en aansprakelijkheid verstaan t.b.v. voorkomen van schade?

In **opdracht 5** zijn 2 fictieve scenario's over opgetreden schade door gebruik van AI CDSS DECIDE gepresenteerd (zie 6.3.4). Op elk scenario hebben de deelnemers afzonderlijk gereflecteerd. Scenario 1 schetste een te laag inschatte risico waardoor uiteindelijk de patiënt schade ondervindt in de vorm van verliezen van betaald werk. Scenario 2 schetste een te hoog ingeschatte risico waardoor uiteindelijk de patiënt onnodige medicijnen heeft gebruikt met sterke bijwerkingen en overmatig stress.

De reflecties van de deelnemers op de scenario's tonen een breed scala aan emoties en overwegingen, van frustratie en boosheid tot ethische en praktische vragen over de rol van AI in risicovoorspelling en de kansen op fouten waardoor schade kan ontstaan. Een eenduidig antwoord op de onderzoeksvraag wordt niet gegeven door de deelnemers. Factoren die verantwoordelijkheid bij het gebruik van het AI CDSS DECIDE model beïnvloeden zijn wel genoemd: de professionele verantwoordelijkheid van de huisarts, maar ook van de burger, de afhankelijkheid van technologie voor risicobeoordeling, overmatig vertrouwen op het model door de huisarts en het verlies van het oog voor de mens en de persoonlijke situatie van de patiënt. Daarnaast is een realistische beoordeling van het model door de huisarts onderdeel van diens verantwoordelijkheid. Deelnemers zijn wel bewust van het feit dat bij een dergelijke technologie een geïntegreerde benadering nodig is: menselijke EN technologische maatregelen moeten gecombineerd worden om fouten te minimaliseren.

In **opdracht 6** zijn twee vragen gesteld (zie 6.3.5): wat onder verantwoordelijkheid bij het gebruik

wordt verstaan en wie aansprakelijk zou moeten zijn als het model niet uitlegbaar is en schade optreedt.

De resultaten van opdracht 6 laten zien dat verantwoordelijkheid bij het gebruik van het AI CDSS DECIDE model een combinatie is van verschillende aspecten: gerichte toepassing van het model, professionele verantwoordelijkheid / verantwoord gebruik van de huisarts en patiëntgerichte benadering.

3.3 Onderzoeksactiviteit -2: focusgroep met huisartsen

3.3.1 Onderzoeksvragen

Evaluatievragen

- a) Komen de opvattingen van de focusgroepe deelnemers grotendeels overeen met die van de geïnterviewde huisartsen met betrekking tot de implementatie van AI CDSS DECIDE in de huisartsenpraktijk? – resultaten van survey voorafgaand de focusgroep
- b) Op welke geprioriteerde thema's lopen de visies en reflecties het meest uiteen en wat zijn de achterliggende opvattingen / redenen hiervoor? – focusgroep op basis van de top-3 thema's (thema 3 deels exploratief)

Exploratieve vraag

- c) Welke additionele thema's en perspectieven komen naar voren in een open dialoog met focusgroepe deelnemers over de implementatie van AI in de huisartsenpraktijk, die niet zijn genoemd door de geïnterviewde huisartsen? – focusgroep, open dialoog

3.3.2 Deelnemers & Methoden

De online focusgroep op 9 september 2024 vond plaats via Microsoft Teams omgeving gehost door de Patiëntenfederatie Nederland. In totaal hebben 2 huisartsen deelgenomen en 1 beleidsadviseur eerstelijnszorg (voor leesbaarheid worden ze ook wel tezamen 'huisartsen'; genoemd). De enige inclusiecriteria voor deelname was dat geen van de deelnemers bekend is met BE-methode, heeft niet eerder deelgenomen aan BE-sessie of was anderzijds betrokken in het DECIDE-VerA project. De focusgroepleider was de auteur van dit rapport. De focusgroepleider heeft haar deelname voor zover mogelijk beperkt tot het stellen van (verhelderende) vragen, notuleren en bewaken van de voortgang.

De focusgroep duurde 1 uur zonder pauze. De exploratieve vraag was aan het einde gesteld, in de laatste overgebleven ca. 5 minuten. De focusgroep was opgenomen via de Microsoft Teams opnamefunctie. Deelnemers hebben afzonderlijk het toestemmingsformulier ingevuld en ondertekend teruggestuurd per e-mail. Het toestemmingsformulier was beschikbaar gesteld door de projectleider van het DECIDE-VerA project (bijlage 6.5.2).

De deelnemers hebben vooraf de focusgroep een enquête ingevuld waarmee ze stellingen beoordeeld hebben en reactie hebben gegeven op een vraag. De stellingen zijn geformuleerd op basis van 6 thema's die op hun beurt het resultaat zijn van een thematische groepering o.g.v. sleutelwoorden en synoniemen van een transcript van een interview met een huisarts EN een samenvatting en deels analyse van een ander interview wederom met een huisarts. In bijlage 6.4.1 is schematisch het proces weergegeven van de selectie van de top-3 thema's van de focusgroep.

De deelnemers hadden 3 werkdagen de tijd om de online stellingen te beoordelen en reactie te geven op de vraag (thema 6). De online enquête is via google formulier beschikbaar gesteld aan de deelnemers. In bijlage 6.4.2 zijn de stellingen (kolom 2) behorend bij de thema's (kolom 1) en de reacties (kolom 3) samengavat. De top-3 thema's zijn geselecteerd op basis van twee criteria: liggen de beoordelingen ver van elkaar op de Likert-scale (bijv. neutraal – helemaal eens, eens-helemaal oneens enz.) en de mate waarmee de stelling over de arts-patiënt relatie / samen beslissen gaat (secundaire doelstelling van project). De volgende 3 thema's zijn geselecteerd voor de focusgroep ter beantwoording van de evaluatievraag b).

1. Risico's & afwegingen – afwegen negatieve gevolgen

2. Ethiek – impact op arts-patiëntrelatie
3. Beoogd gebruik van een AI-CDSS DECIDE toepassing

3.3.3 Samenvatting resultaten

Hieronder volgt per onderzoeksvraag de samenvatting van de resultaten. Voor gedetailleerde uitkomsten wordt verwezen naar bijlage 6.4.2, en naar de ruwe data (de Teams-opname van de focusgroep, in bezit van de projectleider).

Evaluatievragen

- d) Komen de opvattingen van de focusgroepdeelnemers grotendeels overeen met die van de geïnterviewde huisartsen met betrekking tot de implementatie van AI CDSS DECIDE in de huisartsenpraktijk? – resultaten van survey vooraf de focusgroep

Op 5 van de 8 stellingen hebben de deelnemers vergelijkbare beoordelingen gegeven (scores op de Likert scale liggen naast elkaar). Waar sterke overeenstemming is betreffen de volgende: het recht van de burgers om verhoogde risico niet te weten, onnodige ongerustheid bij burgers door vals-positieve resultaten, voorsuccesvolle adoptie en om de werkdruk te verlagen, is het essentieel dat het model wordt geïntegreerd in de bestaande processen en systemen van de huisartsenpraktijk en het belang van uitlegbaarheid van het model aan patiënten.

- a) Op welke geprioriteerde thema's lopen de visies en reflecties het meest uiteen en wat zijn de achterliggende opvattingen / redenen hiervoor? – focusgroepsresultaten op basis van top-3 thema's (thema 3 deels exploratief).

Van de in totaal 6 thema's (5 met stellingen vooraf uitgevraagd en 1 met een open vraag) lopen bij 3 thema's de meningen uiteen door de deelnemers. Per thema worden de volgende argumenten en opvattingen aangedragen.

- Thema-1: Risico's & afwegingen – afwegen negatieve gevolgen
Twee van de deelnemers maken zich zorgen om overmatige medicalisering en "ziek maken" van gezonde mensen. Daarnaast geven alle drie deelnemers aan dat de gebruiker zich niet blind moet staren op het model (belang van arts-patiëntrelatie boven modelvoorspelling, ondersteunde rol van het model en overmatig vertrouwen erop). Vrije keuze van burgers om risico niet te willen weten en de mogelijk juridische implicaties zijn verdere zorgen. M.b.t. het werken van het model hebben 2 deelnemers aangegeven dat daar behoefte is aan betrouwbaarheidsindicator en dat er alertheid moet komen op mogelijke concept en/of data drift bij (langdurig) gebruik van een dergelijk model.
- Thema 2: Thema: Ethiek – impact op arts-patiëntrelatie
De klinische blik en niet-pluis gevoel behouden blijven zeer belangrijk. Het model als ondersteunend systeem wordt nogmaals benadrukt. Uitgaande van de aanname dat deze competenties de huisarts blijft behouden, is één van de deelnemers van mening dat de arts-patiënt relatie niet wezenlijk zal veranderen. Één van deelnemers waarschuwt dat de relatie ondermijnd kan worden doordat patiënten zelf op zoek gaan naar de betekenis van het risicogetal. De arts moet altijd in de loop blijven voor uitleg en interpretatie. Twee van de deelnemers geven nogmaals aan dat het model een assistent moet zijn en vooral jonge artsen er niet te veel op moeten vertrouwen.

- **Thema 3: Beoogd gebruik**
Het thema was beoogd gebruik was bedoeld om te reflecteren op het door het project gekaderd gebruik (evaluatie) en om te exploreren welke andere mogelijkheden de deelnemers zien voor het gebruik van het model. Hiermee zwengelen we de discussie aan om te heroverwegen wie -naast evt. de huisarts- de eindgebruikers van een dergelijk AI-systeem moet worden. Alle 3 deelnemers zien een beoogd gebruik in de sfeer van collectieve preventie waardoor het niet bij de huisarts thuis hoort. Twee deelnemers waarschuwen voor extra werk door veel bezorgde patiënten en het gevaar dat huisartsen “doorverwijsmachines” worden.

Tijdens de focusgroep komen de in de eerste instantie uiteenlopende opvattingen (n.a.v. de Likert-scale resultaten) grotendeels samen bij alle drie focusgroep thema's. Minstens twee van de deelnemers is dan het eens met elkaar bij thema 1. Bij thema 2, het arts-patiënt relatie zien we dat de opvattingen van de drie deelnemers minder convergeren dan bij het eerste thema (afwegen negatieve gevolgen).

Thema 3 (deels exploratief) brengt de deelnemers weer bijeen: alle drie vinden dat een dergelijke model niet bij de huisarts/huisartsenpraktijk thuis hoort, maar bij partijen, die zich bezig houden met preventie.

Exploratieve vraag

- b)** Welke additionele thema's en perspectieven komen naar voren in een open en vrije dialoog met focusgroepeelnemers over de implementatie van AI in de huisartsenpraktijk, die niet zijn genoemd door de geïnterviewde huisartsen in fase I van het project? – focusgroep, open dialoog

Onderwerpen die niet eerder genoemd zijn door de geïnterviewde huisartsen betreffen **i)** bezorgdheid over de toenemende administratieve last voor huisartsen, vooral in het licht van het dreigende artsen tekort, **ii)** de noodzaak om de autonomie van artsen te behouden en niet volledig afhankelijk te worden van modelvoorspellingen. Modellen kunnen wel helpen met omgaan met uitdagingen zoals tekort aan huisartsen, **iii)** culturele verschillen (o.a. niet willen weten of je ziek bent)

Analyse laat zien dat deze onderwerpen weliswaar niet door de geïnterviewde huisartsen expliciet benoemd zijn, maar overlappen met de door de focusgroeleden eerder besproken thema's: bezorgdheid over de toenemende administratieve lasten is in lijn met angst om “verwijsmachine” te worden (thema 3). De noodzaak om de autonomie te behouden van de huisarts is in lijn met het eerder benoemd belang van het mogen behouden en gebruiken van de klinische blik en AI als ondersteuning te zien i.p.v. vervanging (thema 2).

4 Synthese en aanbevelingen

4.1 Synthese van onderzoeksresultaten

4.1.1 Evaluatieonderzoeken – patiënten en naasten

De onderzoeken in onderzoeksactiviteit 1 (patiënten en naasten) gericht op de formatieve analyses in fase I en deels fase II van het project laten het volgende zien: patiënten en naasten konden weinig nieuwe stakeholders noemen (in totaal 4 unieke stakeholders) waardoor onderzoeksvraag c) positief beantwoord kan worden: er **sprake is van saturatie t.a.v. stakeholders** (ca. 75 unieke stakeholders) zoals die in de begeleidingsethiek sessies zijn geïnventariseerd (zie 3.2.3e).

Onderzoeksvraag d) “Is er overeenstemming met de outputs van fase 2 (positieve en negatieve effecten) van de begeleidingsethiek workshops?” kan als volgt beantwoord worden: **er is overeenstemming; meer dan 2/3^e (13 van de in totaal 18) van de effecten was herkend** door de patiënten en naasten. Hieruit zijn 3 effecten waarmee de deelnemers sterk eens mee waren. Al deze 3 effecten betreffen **het belang van communicatie door de huisarts van het risicogetal en een positieve houding daarbij**. Belangrijk is hier om voor ogen te houden dit resultaat in het kader van de door de projectgroep geformuleerde 2 fasen van gebruik van het AI CDSS DECIDE systeem geldt: het tijdig signaleren door gebruik van de toepassing door de huisarts en tijdig actie ondernemen o.g.v. samen beslissen tussen de persoon met hoog risico en de huisarts.

Het perspectief van patiënten en naasten op de juridische vraagstukken van verantwoordelijkheid en aansprakelijkheid (onderzoeksvraag e) leverde een gemengd beeld op: het vooruit- en achteruit kijkende verantwoordelijkheid bleek onduidelijk en ongrijpbaar voor de deelnemers. Tegelijkertijd hebben deelnemers wel kunnen reflecteren op beide vraagstukken door de voor hun belangrijkste aspecten te noemen en te beschrijven. Uit deze beschrijvingen wordt duidelijk dat **naar de verantwoordelijkheidsvraagstuk de deelnemers genuanceerd kijken** omdat ze beseffen dat **diverse factoren in samenhang bepalen wat verantwoordelijkheid eigenlijk betekent** bij het gebruik van het AI CDSS model (gerichte toepassing van het model, professionele verantwoordelijkheid / verantwoord gebruik en patiëntgerichte benadering). **De resultaten t.a.v. aansprakelijkheid zijn onduidelijk:** de meningen van de deelnemers lopen zeer uiteen (niet weten, softwareontwikkelaar als aansprakelijke partij, de vraag als onzin bestempelen en benadrukken de verantwoordelijkheid van de huisarts belangrijk is).

4.1.2 Evaluatieonderzoeken – huisartsen

De onderzoeken gericht op evaluatie van resultaten van fase I laten zien dat de opvattingen van de focusgroep-deelnemers grotendeels overeenkomen met die van de geïnterviewde huisartsen: sterke overeenstemming is gevonden wat betreft **burgers/patiënten**:

- Over het recht hebben / houden om verhoogde risico niet te weten,
- Over de risico van ongerustheid bij burgers door vals-positieve output van het AI-CDSS DECIDE model,
- Over het belang van uitlegbaarheid van het model aan patiënten.

T.a.v. eventuele adoptie van AI CDSS DECIDE model huisartsenpraktijken vinden zowel de geïnterviewde als de aan de focus groep deelgenomen huisartsen essentieel dat het model **geïntegreerd** wordt in de **bestaande processen en systemen** van de huisartsenpraktijken.

Ten aanzien van de thema's waar de opvattingen in de eerste instantie grotendeels uiteenliepen (de drie thema's van de focusgroep-sessie) zien we **convergentie van de opvattingen, behalve bij thema 2 Ethiek: invloed van het gebruik van de tool op de arts-patiënt relatie**. Hier lopen de opvattingen uiteen: van geen effect op de relatie tot de relatie ondermijnen.

4.1.3 Exploratief onderzoek – patiënten en naasten

Op de vraag wat zijn nog mogelijke beoogde gebruiken van een AI-voorspelmodel voor het tijdig ontdekken van hoge risico op hartaanval/beroerte? (opdracht 1) en welke rollen / relaties horen daarbij voor de betrokkenen zien we dat de meeste **patiënten/ naasten zichzelf als eindgebruiker zien**: het AI-model zou op de eigen apparaat gebruikt worden als app. Maar **voor duiding van de risico of t.b.v. vervolgstappen (tijdig actie ondernemen) een rol wordt gegeven aan de zorgprofessional**.

4.1.4 Exploratief onderzoek – huisartsen

In de laatste minuten van de focusgroep is gevraagd naar additionele onderwerpen, perspectieven die de deelnemers niet in de stellingen terug hebben gezien (die gebaseerd waren op de interviews). Uniek onderwerp was het evt. culturele verschil tussen groepen burgers waarmee rekening gehouden moet worden als men een dergelijk risicovoorspel model inzet. In bepaalde contexten (o.a. religieus, cultureel bepaalde gewoontes of tradities) kan het delen van persoonlijke gezondheidsrisico's als belastend of onwenselijk worden ervaren, waardoor het gebruik van het model minder effectief kan zijn.

4.2 Aanbevelingen

Het uiteindelijke doel van dit rapport is om vanuit een integraal perspectief (de samenhang van wettige, ontwerp- en ethische principes) gebaseerd o.a. op de evaluatiestudie, aanbevelingen te geven t.a.v. beide doelstellingen van het DECIDE-VerA project:

- **primaire doelstelling:** het ontwikkelteam van het AI-CDSS-DECIDE tool van praktisch advies te voorzien voor de verder ontwikkeling van de tool;
- **secundaire doelstelling:** betreft de langere termijn doelstelling om van het AI-CDSS DECIDE-tool door te ontwikkelen tot een efficiënt middel t.b.v. ondersteuning van de jongere patiënt (30-49 jaar) en arts in het samen nemen van beslissingen om de risico's op hart- en vaatziekten te verminderen.

4.2.1 Aanbevelingen t.a.v. primaire doelstelling

Ten aanzien van de eindgebruiker

Bij de verkenningsfase van een Medtech innovatie hoort exploratie van diverse stakeholders als eindgebruikers. Vanuit het DECIDE-VerA project is in de eerste instantie vanuit de huisarts als eindgebruiker de burger- en patiëntreis ontwikkeld (zie 2.2.5.). De bevindingen van beide participatieve bijeenkomsten in de huidige exploratieve onderzoeksactiviteiten wijzen erop dat de meeste patiënten en naasten het AI-model op de eigen apparaat als app zouden willen gebruiken, d.w.z. zij zelf zouden de eindgebruikers willen zijn totdat er verhoogd risico wordt gesignaleerd. Huisartsen zelf zien het gebruik van de tool ook niet primair als hun taak, d.w.z. ze zien zichzelf juist niet in de rol van de eindgebruiker. Op grond van deze resultaten is de aanbeveling om alternatieve scenario's voor fase 1 (tijdig signaleren) te verkennen:

- de ontwikkelaars zouden de in de workshop met patiënten ontwikkelde values, scenario's verder kunnen verkennen (zie 6.3.1) en
- afhankelijk van de eindgebruiker(s) in de diverse scenario's, de claim van de tool kunnen aanpassen en te laten toetsen door een notified body [16].

Ten aanzien van verdere ontwikkeliteraties

Naast de in het huidige project opgehaalde inzichten over *ethics by product design* binnen het VSD framework, wordt op basis van inzichten uit het parallelle ZonMw project⁸ aanbevolen om:

- Agile-methodologieën, waaronder concepten als Epics en User Stories te verkennen. Volgens Leijnen et al., 2020 [17] leent een agile aanpak zich beter voor ethics by software design dan het Waterfall-model waarbij het ontwerp vooraf wordt bepaald. Bij de agile aanpak vindt continue planning, ontwerp en leren plaats tijdens de ontwikkeling waarbij m.b.v. user-stories de ethische overwegingen tijdens het gehele ontwikkelingsproces geadresseerd worden. Specifiek gaat het om het uitproberen van user-stories die - volgens Leijnen en collega's - op een natuurlijke manier de ethische aspecten weergeven (dit gebeurt specifiek in het deel van de user-story waar het doel wordt beschreven).
- In het parallelle ZonMw project⁸ is een participatieve co-design methodologie ontwikkeld, waarvan de 4 hoofdfasen overeenkomsten vertonen met de door Umbrello en van de Poel [18] benoemde fasen van iteratieve doorontwikkeling van AI-innovaties: na elke prototypering en small-case testing wordt de AI-tool in ontwikkeling opnieuw in context geanalyseerd. In de vervolgstap worden stakeholders' waarden geëliciteerd, in de daarop volgende stap vertaald naar design-vereisten die dan in de vervolgversie van het prototype/product worden 'ingebouwd'. Deze cyclus herhaalt zich waarbij het product in feite door de waardevolle AI-

⁸ <https://projecten.zonmw.nl/nl/project/ethical-appraisal-ai-driven-medtech-trustworthy-ai-development-participatory-patient>

- funnel heen gaat (Figuur 3) / op de TRL omhoog schuift (Figuur 4).
- Van Os et al. [10] (zie ook 2.2.1) wijzen op de noodzaak van verdere validatie en mogelijk herkalibratie van de modellen voor verschillende populaties. Het verlaten van de Cox elastic net modellen wordt hierbij echter niet gesuggereerd. Omdat dit type modellen interpreteerbaar zijn in termen van identificeerbaarheid van voorspellers en hun relatieve bijdrage aan de voorspelling, wordt aanbevolen de hiervoor gangbare visuele weergaves te verkennen zoals de SHapley Additive exPlanations (SHAP) [19,20]. Na mogelijke aanpassingen van visualisatie-aspecten ten behoeve van communicatie naar de patiënt is het belangrijk om de effectiviteit en begrijpelijkheid hiervan te evalueren.

4.2.2 Aanbevelingen t.a.v. secundaire doelstelling

De secundaire doelstelling betreft de langere termijn doelstelling om het AI-CDSS DECIDE tool door te ontwikkelen tot een efficiënt middel t.b.v. ondersteuning van de jongere patiënt (30-49 jaar) en arts in het samennemen van beslissingen om de risico's op hart- en vaatziekten te verminderen.

De aanbevelingen in deze sectie zijn niet onder eenduidige aanbevelingscategorieën te plaatsen omdat voor deze langere termijn doelstelling het design, ethische en juridische aspecten in elkaars verlengde zitten en/of in elkaar grijpen en dat ook op diverse manieren afhankelijk van keuzes die eerst gemaakt worden. Dit blijkt uit de onderstaande analyses. Gekozen is daarom om de aanbevelingen na de analyse te formuleren.

Beginkeuze in design

We beginnen met de eerste keuze die in het DECIDE-VerA project gemaakt is: **de eindgebruiker is de huisarts**. De aannahme hierbij was dat de huisarts (en praktijk) de plek is waar de AI-CDSS DECIDE-tool meerwaarde zal bieden voor het probleem dat geadresseerd is: het tijdig kunnen identificeren van relatief jonge mensen met verhoogd risico die te vaak onder de radar blijven. Dit biedt vervolgmogelijkheden om in het proces van samen beslissen tijdig actie te ondernemen.

Deze ontwerpkeuze leidde vervolgens tot ethische dilemma's: kan de burger/patiënt door de huisarts zomaar geconfronteerd worden met het nieuws dat hij of zij een verhoogd risico heeft op een hartaanval of beroerte? De moeite die patiënten/burgers hiermee hadden, wordt gereflecteerd door de drie sterkste effecten waarmee de deelnemende patiënten aan de huidige evaluatiestudie het meest eens waren. Al deze drie effecten betreffen het belang van communicatie door de huisarts over het risicogetal en een positieve houding daarbij. Huisartsen gaven aan het recht van burgers om niet te willen weten belangrijk te vinden.

Het proces van samen beslissen vindt in dit scenario geheel bij de huisarts plaats, waarbij het communiceren van het verhoogde risico tevens het communiceren van het keuzemoment inhoudt (stap 1 in het proces). Dit wordt geconcludeerd doordat bij het communiceren van verhoogd risico de huisarts de patiënt uitnodigt om de vervolgstappen van samen beslissen te ondernemen, met als doel om tijdig actie te kunnen ondernemen, inclusief niets doen. Dit laatste is belangrijk: het niets willen doen aan de verhoogde risico met de kennis van de verhoogde risico is wezenlijk anders dan 'by default' niets doen omdat de burger niet bewust is van de risico. Verdere analyse wordt aanbevolen vanuit het framework van Jacobs (2020) – Capability Sensitive Design, zodat er inzicht ontstaat in hoe de ethische, design- en verantwoordelijkheidsaspecten daadwerkelijk samenhangen in dit scenario, waarbij de keuze is gemaakt voor de huisarts als eindgebruiker. Pas dan kunnen conclusies worden getrokken over de vraag of een tool zoals AI-CDSS DECIDE uiteindelijk bijdraagt aan het verwezenlijken van de belangrijkste capabilities, dan wel of de AI-CDSS DECIDE-tool een positieve of

negatieve conversiefactor is.

Een tweede design-keuze is het verplaatsen van het gebruik van AI-CDSS DECIDE tool naar de burger/patiënt. Hiermee wordt automatisch de signaleringsfunctie ook verplaatst naar de burger/patiënt. In 4.2.1 wordt aanbevolen om dit verder te verkennen gezien de deelnemende patiënten dit als een potentieel waardevolle ontwerpkeuze zagen en de huisartsen zichzelf niet als eindgebruiker van een dergelijke tool beschouwen.

In de pre-digitale tijdperk vertelde altijd de arts aan de patiënt dat er een keuzemoment is aangebroken (stap 1 in het vier stappenmodel van samen beslissen van Elwin, 2012 [2]). Als de eindgebruiker de patiënt is, dan komt het signaal direct bij de patiënt/burger binnen op het eigen digitale product (telefoon, tablet, computer, smartwatch). Hierbij verplaatst het "tijdig signaleren" zich van de huisartsenpraktijk naar de huiskamer van de patiënt en verdwijnt het vraagstuk van hoe de relatieve jonge patiënt door de huisarts geïnformeerd dient te worden over het verhoogde risico dat uit het AI-CDSS DECIDE tool is "gerold". Het is niet de huisarts die vertelt dat er iets te kiezen valt, maar dit keuzemoment wordt automatisch "by design" gecommuniceerd. Uitgaande ervan dat de patiënt de tool vrijwillig gebruikt (en dus gebruik maakt van zijn *capability*), verplaatst het kennis nemen van het keuzemoment zich kortom van de zorgverlener naar de patiënt. De patiënt kan na het signaal van verhoogd risico vervolgens op eigen initiatief actie ondernemen (*functionings*) door bijv. een afspraak te maken met de huisarts. De huisarts zal vervolgens stap 1 van samen beslissen "afmaken" door de patiënt mee te geven dat diens waarden, voorkeuren en mogelijkheden (*capabilities, resources*) belangrijk zijn in het vervolg van het samen beslissen proces. Nieuwe vragen werpen zich op zoals: in hoeverre is de zorgverlener bekend met de zelf aangeschafte AI-tool en accepteert hij/zij de betrouwbaarheid ervan? Kan en wil de huisarts de informatie tot zich nemen over de bijdrage van de voorspellers op basis waarvan het signaal "afgegaan" is? Zo ja, in hoeverre kan verwacht worden dat de huisarts bekend is met de diverse tools alla AI-CDSS DECIDE die op de markt zijn? De aanbeveling is hierbij om deze vragen verder te onderzoeken in een vervolgtraject in nauwe samenwerking met de ontwikkelaars van AI-CDSS DECIDE tool (Van Os et al. 2023).

Afhankelijk van de antwoorden op bovenstaande vragen kan het samen beslissen proces hier ofwel stoppen (de huisarts kan of wil het niet, of vertrouwt het signaal niet door het bijvoorbeeld als een gadget te zien) ofwel verder gaan. Bij het verder gaan kan het samen beslissen proces uitgesteld worden (huisarts wil extra onderzoeken doen ter verificatie van het verhoogde risico) voordat stap 2 van samen beslissen aan de orde is.

Bij het uiteindelijk wel nemen van stap 2 van samen beslissen heeft de huisarts de verantwoordelijkheid op zich genomen. De huisarts legt in deze stap de opties die er zijn voor de patiënt uit, inclusief de voor- en nadelen van elke optie (een optie kan zijn dat er geen actie wordt ondernomen). Onbekend is in hoeverre de huisarts de verantwoordelijkheid wil nemen uitgaande van alleen het verhoogd risico getal, het signaal dat de patiënt zelf heeft "meegenomen" naar de spreekkamer. Een mogelijke oplossing hiervoor kan zijn dat de huisarts een aantal gevalideerde AI-toepassingen adviseert aan zijn patiënten in de leeftijdsgroep waarvoor deze AI-tools gevalideerd zijn (incl. mogelijk AI-CDSS DECIDE tool). In dat geval is er minder kans op stagnatie bij stap 1 van samen beslissen omdat de huisarts de kwaliteit van het signaal accepteert en verantwoordelijkheid in het vervolg op zich neemt. In deze scenario waar de patiënt de eindgebruiker is, worden vervolgens stappen 3 en 4 van samen beslissen op gebruikelijk (volgens het interpretatief model) wijze doorlopen.

Beginkeuze in ethiek of wetgeving

Als bij het ontwerpen van een signaleringsinstrument zoals de AI-CDSS DECIDE tool, met het identieke doel om jongere volwassenen tijdig te kunnen identificeren, op de eerste plaats de keuze in de dimensie van ethiek wordt gemaakt ofwel in de wetgeving, dan zal dat gevolgen hebben voor het design van de tool.

Laten we stellen dat ervoor gekozen is om vanuit het recht om niet te willen weten de vervolgstappen te nemen om een signaleringstool zoals AI-CDSS DECIDE te ontwikkelen. Hieruit volgt dat deze wens – om al dan niet kennis te willen nemen van een eventueel verhoogd risico – indien de eindgebruiker de huisarts is, eerst vastgelegd moet worden bij de huisarts. Waardendialogen (begeleidingsethiek) over de mogelijke voor- en nadelen van het gebruik van de AI-CDSS DECIDE tool door de huisarts zou na deze beginkeuze andere uitkomsten geven.

De optie waarbij de eindgebruiker de burger/patiënt is, vervalt vanuit deze ethische beginkeuze (*het recht om niet te willen weten*). Ervan uitgaande dat mensen uit vrije wil de AI-CDSS DECIDE tool gebruiken, is het vrijwillig gebruiken van de tool ten behoeve van vroegtijdige signalering van verhoogd risico niet te rijmen met het niet willen weten van een verhoogd risico.

Het uitdenken van diverse scenario's, ook vanuit het wetgevingsperspectief, wordt sterk aanbevolen om zo de interactie en samenhang van de ethische, design- en juridische perspectieven te ontrafelen.

In het algemeen kan geconcludeerd worden dat beginkeuzes in belangrijke mate bepalen welke ethische, juridische en designkeuzes in het vervolg zich voordoen en of ze überhaupt aan de orde zijn. Keuzemogelijkheden die zich wel voordoen zullen allemaal hun eigen "vervolgvertakkingen" hebben, d.w.z. de volgorde van keuzes beïnvloedt de keuzeruimtes in het vervolg.

Het oorspronkelijke probleem, namelijk dat jongere individuen (30-49 jaar) met verhoogd risico onder de radar blijven, waardoor niet tijdig actie wordt ondernomen om het risico te verlagen, wordt in een design waarbij de eindgebruiker de patiënt/burger is, niet direct opgelost. Ten eerste, omdat het gebruik van een dergelijke tool op het eigen digitale apparaat wordt bepaald door meerdere factoren:

- De middelen hebben om een digitaal apparaat aan te schaffen (*resources to realize capabilities*).
- Het kunnen gebruiken en daadwerkelijk actie ondernemen door zelf de huisarts te benaderen (*functionings*).

Dit design elimineert zeker een aantal ethische dilemma's, maar confronteert ons met andere dilemma's. Bovendien roept het vragen op over de bereidheid van de huisarts om het samen beslissen-proces überhaupt aan te gaan. De huisarts kan immers het AI-CDSS DECIDE-tool niet vertrouwen, tenzij het gebruik ervan expliciet wordt aanbevolen aan zijn/haar patiëntenpopulatie. Een eenduidig antwoord op de vraag op welke manier de arts-patiëntrelatie wordt beïnvloed door een toepassing zoals AI-CDSS DECIDE, is niet te geven. Dit lijkt afhankelijk te zijn van de keuzes die aan het begin en in het vervolg van de ontwikkelstappen worden gemaakt.

Ten aanzien van het samen beslissen-proces concludeer ik dat in feite alleen stap 1 van het proces – "er valt iets te kiezen" – wenselijk anders kan zijn, afhankelijk van de beginkeuzes, zoals hierboven uiteen is gezet. De vervolgstappen van het proces – het bespreken van de opties, de voorkeuren en uiteindelijk het nemen van een beslissing – worden niet wezenlijk beïnvloed door het inzetten van een toepassing zoals AI-CDSS DECIDE.

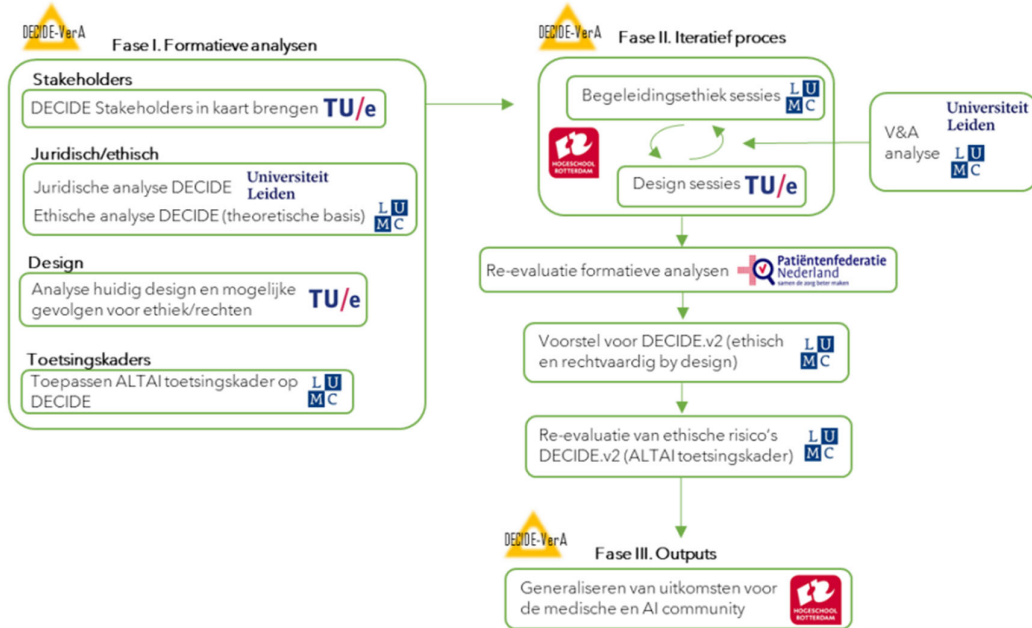
5 Bronnen

- [1] E. J. Emanuel en L. L. Emanuel, 'Four models of the physician-patient relationship', *JAMA*, vol. 267, nr. 16, pp. 2221-2226, apr. 1992.
- [2] G. Elwyn *e.a.*, 'Shared Decision Making: A Model for Clinical Practice', *J. Gen. Intern. Med.*, vol. 27, nr. 10, pp. 1361-1367, okt. 2012, doi: 10.1007/s11606-012-2077-6.
- [3] A. Lavin *e.a.*, 'Technology readiness levels for machine learning systems', *Nat. Commun.*, vol. 13, nr. 1, p. 6039, okt. 2022, doi: 10.1038/s41467-022-33128-9.
- [4] High-Level Expert Group on Artificial Intelligence, 'ETHICS GUIDELINES FOR TRUSTWORTHY AI', apr. 2019.
- [5] M. Weda, A. De Bruijn, T. Meneses Leonardo Alves, en C. De Vries, 'Digitale beslissingsondersteuning in de zorg', RIVM, 2019. Geraadpleegd: 24 april 2020. [Online]. Beschikbaar op: <https://rivm.openrepository.com/handle/10029/623072>
- [6] Z. Angehrn *e.a.*, 'Artificial Intelligence and Machine Learning Applied at the Point of Care', *Front. Pharmacol.*, vol. 11, 2020, doi: 10.3389/fphar.2020.00759.
- [7] F. Wouda en H. Hutink, 'Artificial Intelligence in de zorg - begrippen, praktijkvoorbeelden en vraagstukken', Nictiz, Den Haag, 2019.
- [8] B. Heinrichs en S. B. Eickhoff, 'Your evidence? Machine learning algorithms for medical diagnosis and prediction', *Hum. Brain Mapp.*, vol. 41, nr. 6, pp. 1435-1444, apr. 2020, doi: 10.1002/hbm.24886.
- [9] T. J. Loftus *e.a.*, 'Ideal algorithms in healthcare: Explainable, dynamic, precise, autonomous, fair, and reproducible', *PLOS Digit. Health*, vol. 1, nr. 1, p. e0000006, jan. 2022, doi: 10.1371/journal.pdig.0000006.
- [10] H. J. A. van Os *e.a.*, 'Cardiovascular Risk Prediction in Men and Women Aged Under 50 Years Using Routine Care Data', *J. Am. Heart Assoc.*, vol. 12, nr. 7, p. e027011, apr. 2023, doi: 10.1161/JAHA.122.027011.
- [11] L. Farah, J. M. Murris, I. Borget, A. Guilloux, N. M. Martelli, en S. I. M. Katsahian, 'Assessment of Performance, Interpretability, and Explainability in Artificial Intelligence-Based Health Technologies: What Healthcare Stakeholders Need to Know', *Mayo Clin. Proc. Digit. Health*, vol. 1, nr. 2, pp. 120-138, jun. 2023, doi: 10.1016/j.mcpdig.2023.02.004.
- [12] Medical Device Coordination Group, 'Guidance on Qualification and Classification of Software in Regulation (EU) 2017/745 – MDR and Regulation (EU) 2017/746 – IVDR'. Cotober 2019.
- [13] Medical Device Coordination Group Document, 'Regulation (EU) 2017/745: Clinical evidence needed for medical devices previously CE marked under Directives 93/42/EEC or 90/385/EEC. A guide for manufacturers and notified bodies'. april 2020.
- [14] B. Friedman en D. G. Hendry, *Value Sensitive Design: Shaping Technology with Moral Imagination*. The MIT Press, 2019. doi: 10.7551/mitpress/7585.001.0001.
- [15] D. Yoo, 'Stakeholder Tokens: a constructive method for value sensitive design stakeholder analysis', *Ethics Inf. Technol.*, vol. 23, nr. 1, pp. 63-67, mrt. 2021, doi: 10.1007/s10676-018-9474-4.
- [16] W. en S. Ministerie van Volksgezondheid, 'Rol van de notified body - Medische technologie - Inspectie Gezondheidszorg en Jeugd'. Geraadpleegd: 19 februari 2025. [Online]. Beschikbaar op: <https://www.igj.nl/zorgsectoren/medische-technologie/markttoelating/rol-notified-body>
- [17] S. Leijnen, H. Aldewereld, R. van Belkom, R. Bijvank, en R. Ossewaarde, 'An agile framework for trustworthy AI.', in *NeHuAI@ ECAI*, 2020, pp. 75-78.
- [18] S. Umbrello en I. van de Poel, 'Mapping value sensitive design onto AI for social good principles', *AI Ethics*, vol. 1, nr. 3, pp. 283-296, 2021, doi: 10.1007/s43681-021-00038-3.
- [19] K. Lee, M. V. Ayyasamy, Y. Ji, en P. V. Balachandran, 'A comparison of explainable artificial

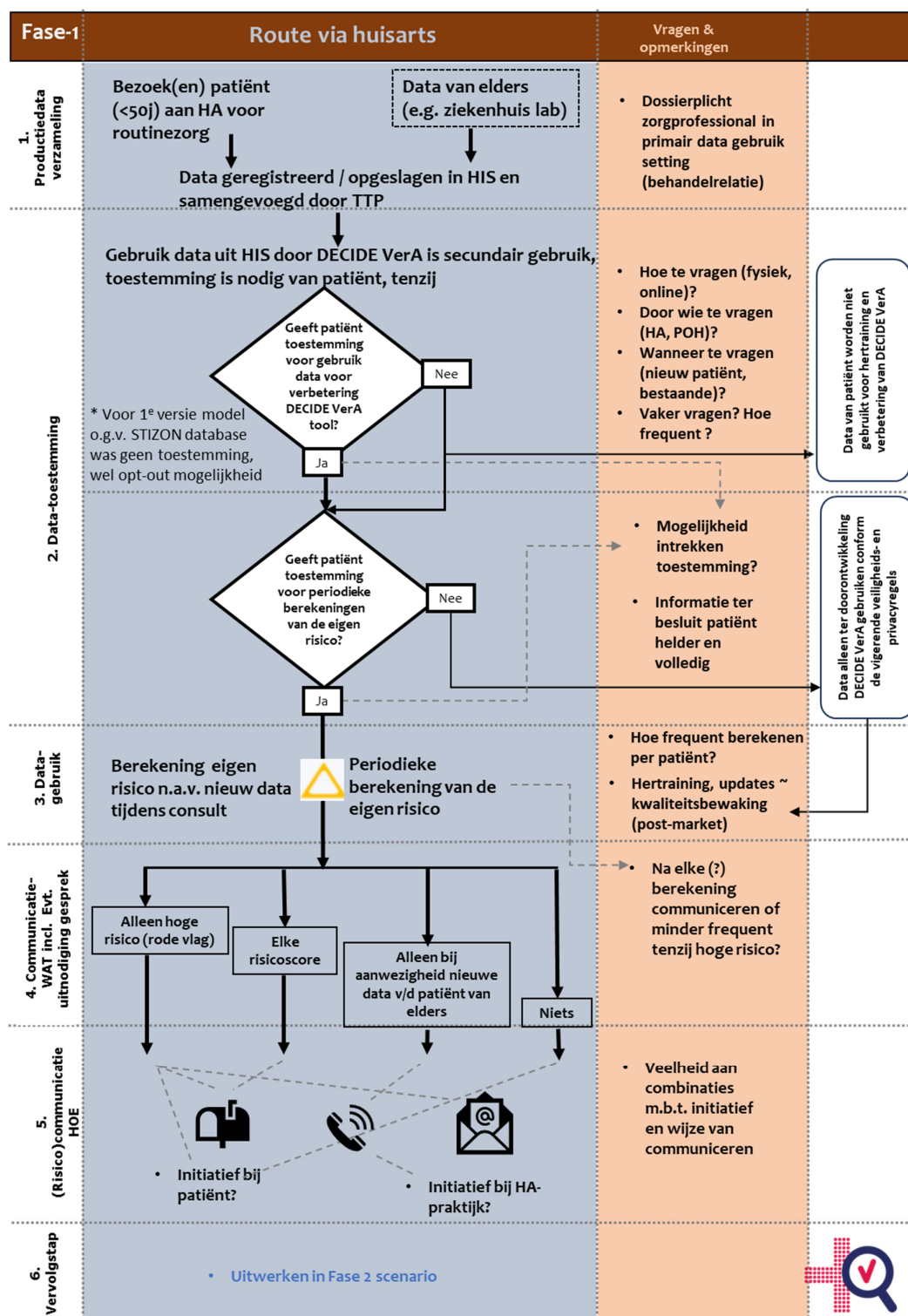
- intelligence methods in the phase classification of multi-principal element alloys', *Sci. Rep.*, vol. 12, nr. 1, p. 11591, jul. 2022, doi: 10.1038/s41598-022-15618-4.
- [20] P. K. Mohanty, S. A. J. Francis, R. K. Barik, D. S. Roy, en M. J. Saikia, 'Leveraging Shapley Additive Explanations for Feature Selection in Ensemble Models for Diabetes Prediction', *Bioengineering*, vol. 11, nr. 12, Art. nr. 12, dec. 2024, doi: 10.3390/bioengineering11121215.

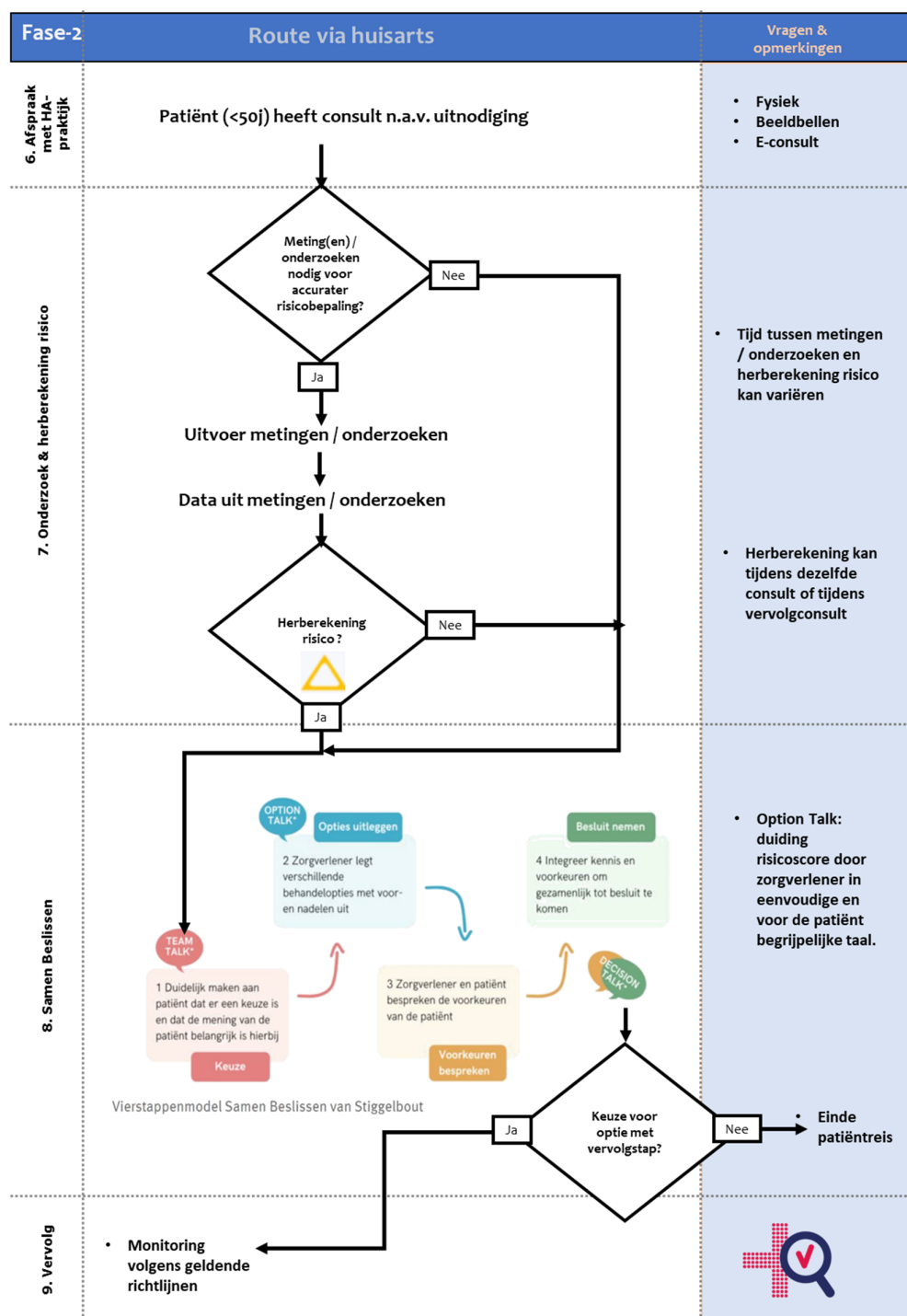
6 Bijlagen

6.1. Projectfasen



6.2 Scenario's in fase 1 en fase 2 van het gebruik

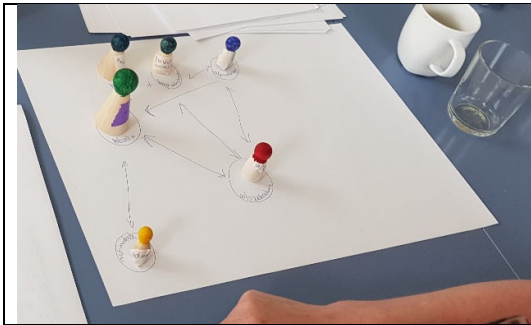




6.3 Onderzoeksactiviteit-1 workshop met patiënten en naasten

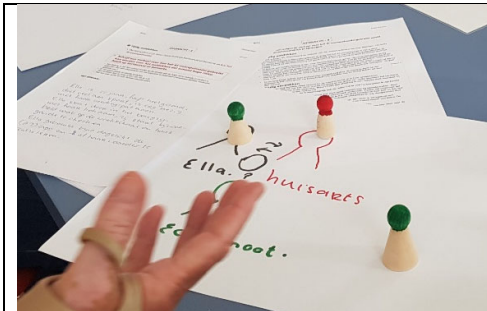
6.3.1. Opdrachten 1 en 2 inclusief resultaten

K: Patiënt X komt al een aantal keer de afgelopen jaren bij de huisarts met verschillende vage klachten de huisarts adviseert een app met AI voorspelmodel te installeren en gegevens in te voeren. Ook gegevens van de huisarts uit het dossier gaan hier in hieruit komt een uitslag waarmee verder onderzoek gedaan zou kunnen worden.



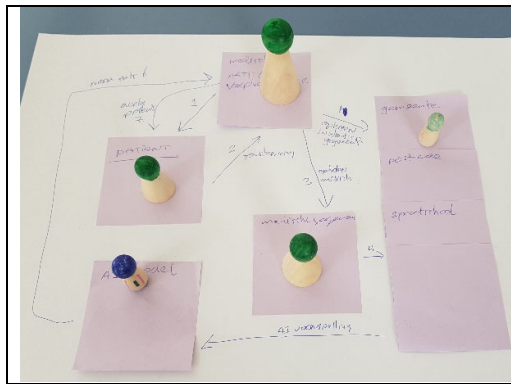
Patiënt X staat centraal (grote pop met groen hoofd). Die wordt beïnvloed door niet medische factoren (geel poppetje) en staat in contact met de arts in het ziekenhuis (rood poppetje). De bovenste rij poppen staan voor de huisarts, de praktijkondersteuner die extraonderzoek kunnen verrichten (blauw poppetje).

B: Ella is 55 jaar, leeft heel gezond doet veel aan sporten is erg bezig met haar voedingspatroon. Ella slaat het door in het bezig zijn met haar lichaam. Zij staat bijvoorbeeld vaak op de weegschaal om haar gewicht te checken Ella gebruikt bijna dagelijks de AI app om al haar waardes te controleren.



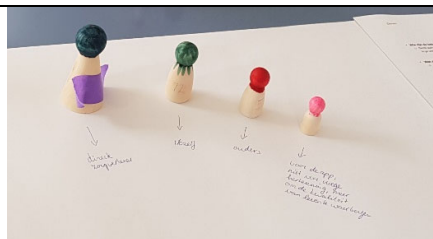
Ella is in contact met de huisarts, heeft een echtgenoot op wie het AI systeem ook invloed heeft en zo hun relatie indirect beïnvloedt.

A: Michel heeft diabetes type 2 en sinds kort een verhoogde bloeddruk. Michel hoort van zijn arts/verpleegkundige van het AI model en of hij interesse heeft om het te testen waarbij voorspellen op regelmatige basis de bedoeling is (zelf de frequentie te kiezen). Michel installeert de app op zijn telefoon, tablet en computer. Hij geeft toestemming voor het gebruik van bepaalde soorten gegevens. Als basis geldt de 1e uitkomst van het model. Michel bespreekt het met de arts/verpleegkundige een keer per periode de uitkomst. Ten alle tijden te stoppen en wijzigen etc. is mogelijk. Het gebruik van data is altijd met toestemming, nooit met "ja, tenzij bezwaar" -systematiek.



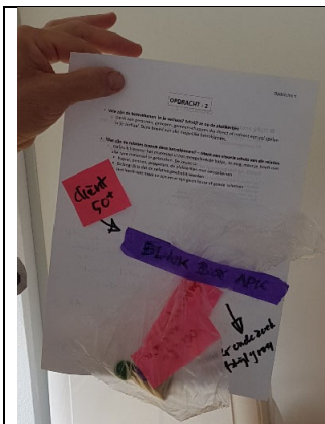
De flow van acties is met pijltjes aangegeven: Michel krijgt na een uitnodiging door de zorgverlener (grootse pop aan bovenkant), geeft toestemming voor bepaalde typen gegeven bij bepaalde instanties (rechtser kolom zijn niet medische instanties waar zijn gegevens staan). De voorspelling wordt met de zorgverlener besproken. Michel houdt de regie want hij is de eindgebruiker van het AI-model.

D: Wanneer ik de AI in zou willen zetten is met name alles mijngezondheid minder wordt en er meer risico's zijn en de gang naar een huisarts ziekenhuis niet makkelijk meer gaat. Het zou fijn zijn dat ik er niet zelf de hele dag alert in hoef te zijn of op een smartwatch moet kijken. Maar de systemen veranderingen herkennen en de zorgverlener ingrijpt. Het zou als een soort pop-up mogen verschijnen net als mijn thuis monitor die dagelijks mijn pacemaker uitleest en een melding geeft wanneer ik bijvoorbeeld en aantal dagen niet wordt gezien (ben op vakantie, slaap in een andere kamer enz.). Er de hele dag mee bezig mee moeten zijn geeft meer arts bezoeken.



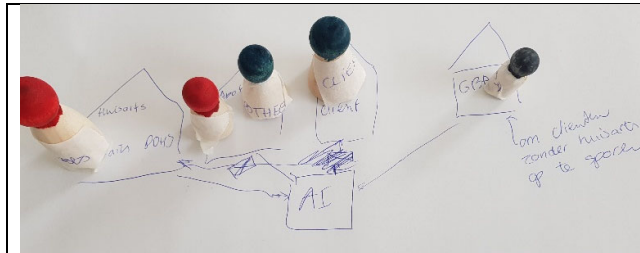
De stakeholders bevinden zich rondom de patiënt die een pacemaker draagt. Voor deze deelnemer is kwaliteit van leven belangrijker dan vroege herkenning van het hebben van verhoogd risico. Technologie moet onzichtbaar zijn werk doen, de patiënt helpen om een goed leven te leiden. Just in time – melding is voor deze deelnemer voldoende.

G: Reclame via specialist TV huisarts sociale media voor APK, waaruit hoge risico blijkt. Dan krijgt de persoon gratis APK. Omdat bijvoorbeeld overgewicht als hoger risico is. Daarna wordt doorverwezen naar behandelaar die risico's beter in zicht brengt. Na invullen van het model komt vervolg: leefstijladvies, medicatie enz. Vanaf dat moment wordt AI model verder gevuld.



Voor deze deelnemer moet het AI-model vooral niet bij iedereen vanuit het niets ingezet worden. De deelnemer stelt een laagdrempelig "consultatiebureau"-achtige instantie voor, waar met name 50+-ers kunnen inlopen voor een eerste advies. In het vervolg zou het model bij mensen bi wie een eerste grove risicoschatting is gegeven, de risico verder kunnen verfijnen. Zo voorkomen je onterechte zorgen bij mensen als ze "zomaar" door hun huisarts uitgenodigd zouden worden gebaseerd op automatische berekening van het AI-model. Media zouden de aandacht van het publiek moeten trekken om bij de laagdrempelige instantie binnen te lopen voor een eerste grove schatting en advies.

W: jaarlijks draait AI over alle patiënten bij de huisarts om het risico te bepalen bij cliënten van de huisarts waarbij het risico van laag middel naar hoog is gewijzigd of hoog is en een ja niet voor controle zijn geweest wordt en oproep voor een consult gedaan met en POH/huisarts. Toepassing zou gebruikt moeten worden om mensen zonder huisarts of mensen die de huisarts mijden op te sporen.



Voor deze deelnemer is het belangrijk dat het AI-model ingezet is om mensen op te sporen die geen huisarts hebben of het mijden. Die kunnen makkelijk niet bereikt worden, maar lopen meestal de grootste risico door ongezonde voeding, roken enz.

6.3.2 Opdracht 3

OPDRACHT - 3

Lees eerst het volgende verhaal over het AI-voorspelmodel gebruikt wordt bij een huisartsenpraktijk in Nijmegen.

❶ Tijdig ontdekken

- Robert is begin 30, alleenstaand en verhuist van zijn ouders in Eindhoven naar Nijmegen. Hij werkt bij een bakkerij.
- Zijn nieuwe huisarts in Nijmegen gebruikt al lange tijd het AI-voorspelmodel in zijn praktijk
- De eerste keer de huisarts Robert op consult heeft, leg zij uit wat het AI-voorspelmodel doet en vraagt hem of hij zijn toestemming wil geven om het risico periodiek te laten berekenen door het model.
- In Robert zijn familie hebben al een paar mensen hartaanval gehad. Hij wil op tijd weten als hij hogere risico heeft en besluit daarom om toestemming te geven. Met de huisarts spreekt hij af dat elke keer als er nieuwe gegevens van hem zijn vastgelegd door de huisarts, apotheek, ziekenhuis of waar hij maar komt m.b.t. gezondheid, dan berekent het model elke keer het risico.
- Robert wil ook het risicogetal weten als het niet te hoog is. Hij spreekt dus met de huisarts af dat hij elke keer dat er opnieuw het risico is berekend een e-mail krijgt met daarin het risicogetal. Hij gaat naar huis.

- Jaren lang krijgt Robert lage risicogetallen toegestuurd door de huisarts, maar op een gegeven is het risico groot. De huisarts belt nog dezelfde dag en nodigt Robert uit voor een consult.

OPDRACHT 3

❷ Tijdig aanpakken

- Robert kan nog dezelfde week terecht bij de huisarts en hij zit in haar spreekkamer
- De huisarts legt uit wat het risico precies betekent en op grond van welke gegevens het berekend is. Er springt iets uit in die gegevens: Robert heeft een afwijking in een bepaald eiwit waardoor hij hogere risico loopt op een hartaanval.
- Deze afwijking is erfelijk en is waarschijnlijk ook de reden dat er veel hartaanvallen in zijn familie voorkomen.
- Tot nu toe bleef het onopgemerkt omdat het eiwit nooit bepaald was bij labonderzoeken. Maar Robert deed de afgelopen maanden mee met een klinisch onderzoek. Tijdens het onderzoek gaf hij bloed voor onderzoek. Tijdens dat onderzoek is o.a. het eiwit bepaald en daar had hij dus veel te veel van in zijn bloed. Deze informatie is in zijn dossier vastgelegd en ook meegenomen door het AI-voorspelmodel.
- De huisarts bespreekt met Robert welke opties hij heeft om zijn risico te verlagen op een hartaanval. Ze bespreken de voor- en nadelen van alle opties en maken samen een keuze. De keuze valt op minstens 20 kg afvallen.
- Daardoor zal de eiwitgehalte zijn bloed afnemen en ook het risico dalen.
- Robert gaat naar een leefstijlcoach, gaat sporten en op dieet en gaat na een jaar weer naar de huisarts.

Mis je in onderstaande tabellen mogelijke betrokkenen in het verhaal van Robert? Zo ja, schrijf ze bij de juiste categorie.

Categorie	Betrokkenen genoemd in workshop van februari 2024
De patiënt en diens directe omgeving	De patiënt zelf, naasten van patiënt (familie, vrienden, ouders), lotgenotengroep, patiëntforums
Zorg	Huisarts, medisch specialisten, medisch ondersteuners (POH, doktersassistent), patiëntverenigingen, regionale en landelijke huisartsorganisaties, huisarts.nl, wijkzorg, sociaal domein, paramedici, leefstijlcoaches, revalidatiecentra, apothekers, laboranten
Beleid en toezicht	Overheid (VWS, gemeenten), Zorginstituut Nederland, Inspectie van Gezondheidszorg en Jeugd, Zorgautoriteit, organisaties die medische toepassingen certificeren, tuchtraad, rechtspraak.
Techniek & data	Systeemontwikkelaars, mensen die het systeem onderhouden, helpdesk, data warehouse, auditor, user <u>experience</u> experts, algemeen voorlichtingspunt, <u>wearables</u> , <u>quantified self</u> , Internet of Things, implementatiemanager
Wetenschap	Onderzoekers, opleiders, onderwijs, scholen.
Financiering	Verzekeraars (zorgverzekeraar, aansprakelijkheidsverzekering, arbeidsongeschiktheidsverzekering), werkgever van patiënt, belastingbetaler
Commerciële partijen	Leverancier, fabrikant, tabaksindustrie, voedselindustrie, <u>Medtech bedrijven</u> bedrijven.
Overig	Sporthal, regionale organisaties, supranationale instanties, (sociale) media, andere expertsystemen, expertisecentra

Categorie	Betrokkenen genoemd in workshop van juli 2024
De patiënt en diens directe omgeving	De patiënt zelf, naasten van patiënt (gezin; familie -zussen, broers, moeder, vader, partner, kinderen), mantelzorger, vrienden (buurvrouw), collega's (werk), werkgever, patiëntgroep,
Zorg	Huisarts, medisch ondersteuners (praktijkondersteuner (<u>somatiek</u>), doktersassistent), huisartspraktijk en medewerkers (bijv. receptionist), degene die belt (doktersassistent), virtuele dokter, huisartsenpost, medisch specialisten, ziekenhuis, bedrijfsarts, fysiotherapeut, apothekers, gezondheidsinstanties.
Beleid en toezicht	Gemeente, CBR.
Techniek & data	Ontwikkelaars, makers van de "slimme computer", HIS leverancier (HIS = Huisartseninformatiesysteem).
Leefstijl	Leefstijlcoaches, de sportschool/sportvereniging, diëtist, gezondheidsadviseur, begeleiders.
Financiering	Verzekeraars (zorgverzekeraar, hypotheekverstrekker).
Overig	Google, <u>ChatGPT</u> , inwoners Nederland.

6.3.3 Opdracht 4 en resultaten

De in totaal 18 effecten uit de 2 BE-sessies op grond waarvan de stellingen zijn geformuleerd.

Positief effect

Negatief effect

Onnodige angst bij de patiënt (2)

Hogere werklast voor de huisarts (4)

Dataficeren van het individu (8)

Sociale controle (32)

Keuzes maken in het leven (41)

Medicaliseren (afhankelijkheid van de huisarts)
(42)

Kans om mensen te betrekken die nu niet
worden betrokken (preventief) (48)

Grotere druk op kosten basiszorg en om daarin
keuzes te maken (49)

Risico-discriminatie (lage SES associatie
gezondheidsrisico) - bias in de data (65)

Positief	Positief/Negatief/vragen	Negatief
Advies speciaal voor jou, wat moet ik nog meer doen?	Klopt het [systeem]? Bewijs	Paniek, schrikken bijna dood*
	Het is niet uitlegbaar?*	Lang gesprek
		Stigmatiserend*
Tijd om te accepteren		Je voelt je gauw patiënt, invloed op je toekomstperspectief*
Tijd om actie te ondernemen		Niet iedereen snapt wat door de telefoon verteld wordt*
Je voelt je gauw patiënt, invloed op je toekomstperspectief*		Te veel medicijnen worden gebruikt (preventief)

Iedereen moet uitslag krijgen hoog of laag risico*

"Maakt patiënten"*

Kan je iets bieden*

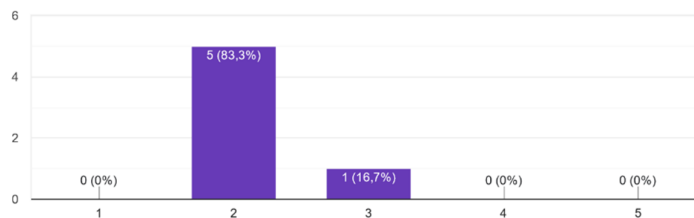
Geaggregeerde responsen per stelling (in totaal 18).

1 = helemaal eens; 2 = eens; 3 = neutraal; 4 = oneens; 5 = helemaal oneens

[1] Het gebruik van het AI-voorspelmodel kan mensen onnodig angstig maken

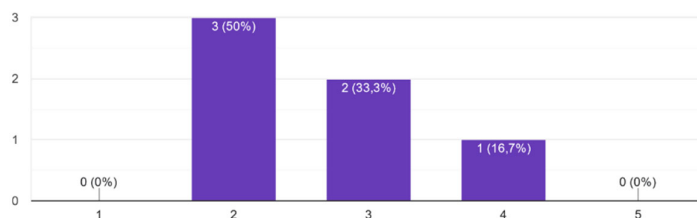
1.

6 antwoorden



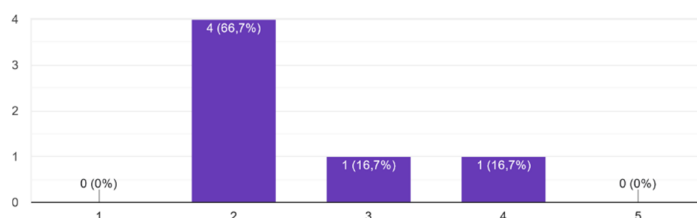
[2] Het gebruik van het AI-voorspelmodel betekent hogere werklast voor de huisarts

2.
6 antwoorden



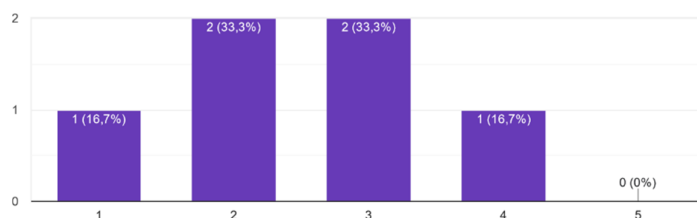
[3] Het gebruik van het AI-voorspelmodel draagt bij het "dataficeren" van mensen

3
6 antwoorden



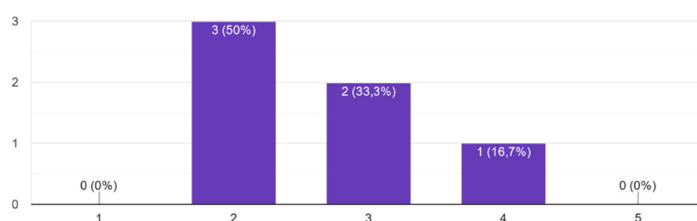
[4] Het gebruik van het AI-voorspelmodel kan leiden tot het voelen van druk en sociale controle (je durft niet meer NIET meedoen)

4
6 antwoorden



[5] Het gebruik van het AI-voorspelmodel draagt bij het gevoel van eigen keuzes te kunnen maken in het leven

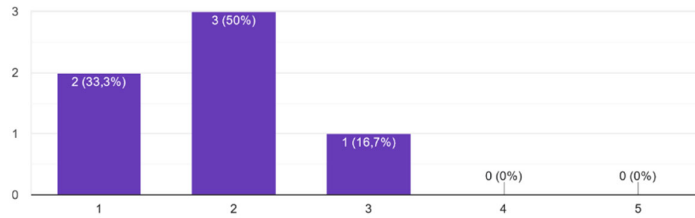
5
6 antwoorden



[6] Het gebruik van het AI-voorspelmodel kan leiden tot medicalisering van de burger

6

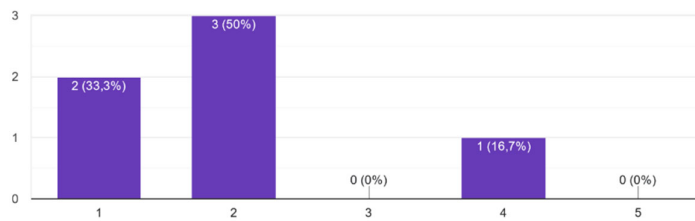
6 antwoorden



[7] Het gebruik van het AI-voorspelmodel biedt de kans om mensen op te sporen die anders niet in beeld zouden komen.

7

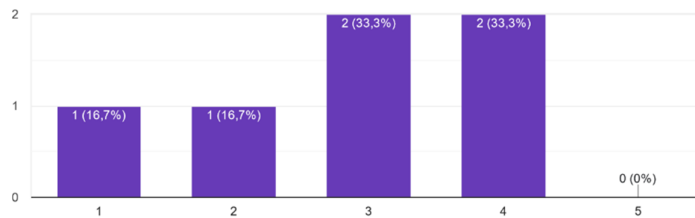
6 antwoorden



[8] Het gebruik van het AI-voorspelmodel zet alleen nog meer druk op de kosten van basiszorg.

8

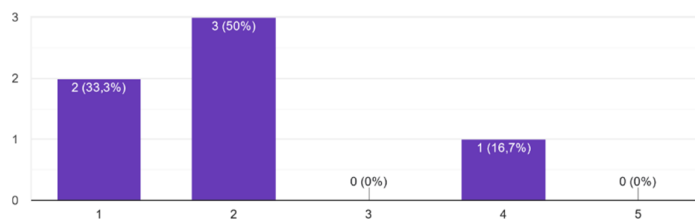
6 antwoorden



[9] Het gebruik van het AI-voorspelmodel kan leiden tot discriminatie van bepaalde groepen zoals mensen met een lage score van maatschappelijke positie (SES)

9

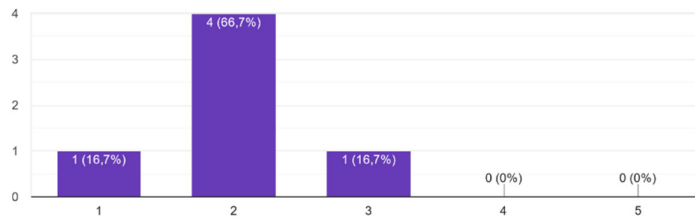
6 antwoorden



[10] Het gebruik van het AI-voorspelmodel kan ervoor zorgen dat de burger schrikt en in paniek raakt en denkt dood te gaan.

10

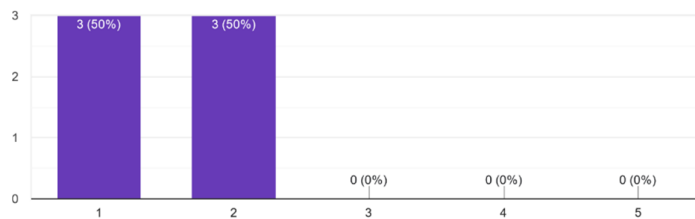
6 antwoorden



[11] Het gebruik van het AI-voorspelmodel kan leiden tot verwarring als de burger onvoldoende uitleg krijgt.

11

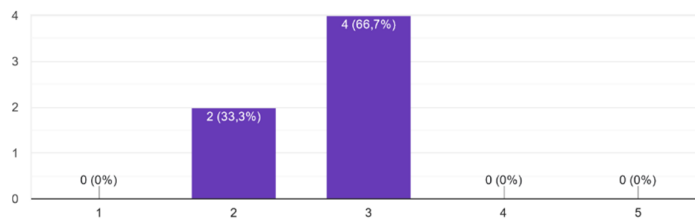
6 antwoorden



[12] Het gebruik van het AI-voorspelmodel kan zorgen voor stigmatisering. Mensen die zich gezond voelen willen niet voor zieke aangezien worden.

12

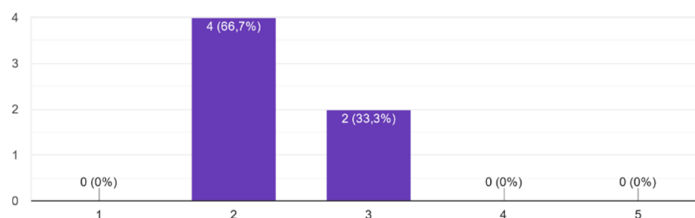
6 antwoorden



[13] Het gebruik van het AI-voorspelmodel kan leiden tot een negatiever toekomstperspectief van de burger.

13

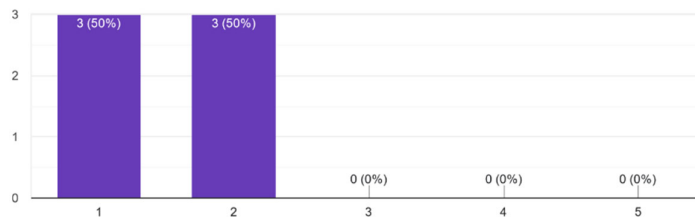
6 antwoorden



[14] Als de huisarts het risicogetal door de telefoon communiceert, zal niet iedereen het kunnen begrijpen waar het precies over gaat.

14

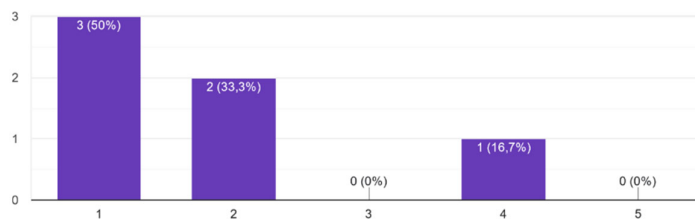
6 antwoorden



[15] De huisarts moet tijd gunnen zodat mensen hun naaste mee kunnen vragen naar het consult na het horen van het hebben van een hoge risico.

15

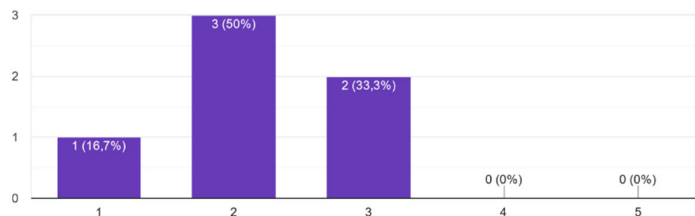
6 antwoorden



[16] Iedereen moet kennis kunnen nemen van alle uitslag van de berekening, niet alleen als het risico hoog is.

16

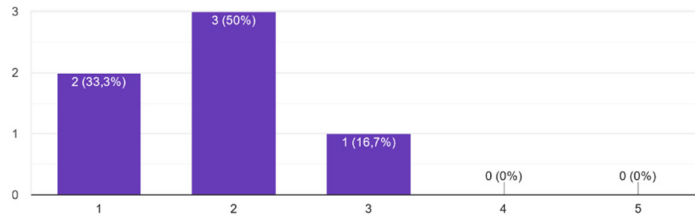
6 antwoorden



[17] Bij het gebruik van het AI-voorspelmodel is belangrijk dat er uitvoerbare opties gegeven worden om uit te kiezen om het hoge risico te verlagen.

17

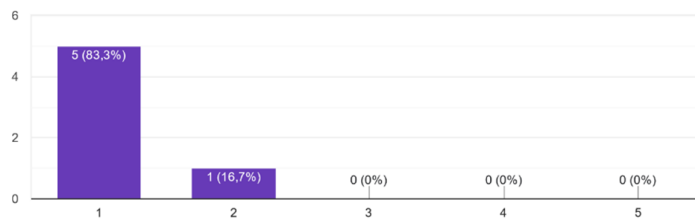
6 antwoorden



[18] Het is belangrijk dat de huisarts in het gesprek over verlagen van het risico een positieve houding heeft.

18

6 antwoorden



6.3.4 Opdracht 5 en resultaten

Scenario 1

Stel voor: je huisartsenpraktijk gebruikt het AI-voorspelmodel en het geeft een foute uitkomst: een te lage risico. In werkelijkheid is je risico hoog. Je huisarts heeft het niet door (de reden waarom laten we even buiten beschouwing). Je huisarts nodigt alleen mensen uit met een hoge risico. Je wordt dus niet uitgenodigd voor een afspraak bij de huisarts terwijl je wel een hoge risico hebt om een hartaanval of beroerte te krijgen. Je krijgt inderdaad een beroerte, overleeft het, maar houdt ernstige lichamelijke beperkingen over waardoor je niet meer kunt werken.

Scenario 2

Stel je een omgekeerd verhaal voor: het AI-voorspelmodel geeft een foute uitkomst, maar nu is het risico onterecht veel te hoog. Met je huisarts kies je samen voor de optie dat je preventief medicijnen slikt (die helaas veel bijwerkingen hebben). Je gaat 4x/jaar je bloedwaarden laten bepalen en maak je je ernstig zorgen, leeft in voortdurende spanning omdat je bang bent dat je elk moment een hartaanval of beroete krijgt.

• **HOE KIJKT JE TEGEN DE SCHADELIJKE GEVOLGEN AAN DIE JE IS OVERKOMEN?**

- AI-voorspelmodel werkt niet goed, geeft laag risico aan terwijl je hoog risico hebt.
- Je denkt dus laag risico te hebben en onderneem je dan ook geen actie om het nog te verlagen.
- Je krijgt wel een hartaanval, krijgt beperkingen en verlies je je werk.

- AI-voorspelmodel werkt niet goed, geeft hoog risico aan terwijl je laag risico hebt.
- Je denkt dus hoog risico te hebben en slikt preventief medicijnen die veel bijwerkingen hebben.
- Daarnaast heb je 4x/jaar bloedonderzoek en maak je je ernstig zorgen, leef je in angst.

Reflecties van deelnemers op scenario 1

- De situatie is wellicht vergelijkbaar met een medische blunder. Discussie vindt plaats over hoe de fout had kunnen worden voorkomen en de rol van de huisarts en de burger hierin.
- Niet iedereen past in het model en mensen zijn ook niet door data geheel te vangen. Hierdoor blijft altijd een risico op verkeerde voorspelling en kans op fouten.
- Één van de deelnemers geeft aan zich in een dergelijk scenario boos en gefrustreerd te zou voelen omdat de situatie vermeden had kunnen worden. Het roept een gevoel van onrechtvaardigheid op bij de deelnemer en onzorgvuldigheid bij het gebruik van het model door de zorgverlener.
- AI blijft mensenwerk: menselijke input en oordeel blijven essentieel en kunnen niet vervangen worden door technologie.
- De zorg wordt geuit dat in het scenario het AI-model belangrijker lijkt te zijn dan de behoeften en het welzijn van de patiënt terwijl dat de verantwoordelijkheid is van de zorgverlener.
- Op het aansprakelijkheidsvraagstuk wordt onvoldoende gereflecteerd door de deelnemers Het is een te complexe vraagstuk voor de deelnemers. Door één van de deelnemers wordt aangegeven dat in een dergelijke casus in de werkelijkheid het overdreven zou zijn om iemand aansprakelijk te stellen en schadevergoeding te eisen.

Reflecties van deelnemers op scenario 2

- Er zouden meerdere lagen van controle en verificatie moeten zijn, waaronder de patiënt zelf, de huisarts, en ingebouwde controles in het AI-systeem.
- Het is belangrijk om de situatie goed te bespreken en de patiënt niet onnodig bang te maken.
- Één van de deelnemers zou in een dergelijk scenario de situatie van de patiënt hee; erg vinden. De schade was te voorkomen door verantwoorde gebruik. Deze deelnemer legt de nadruk op de negatieve impact van zorgen en angst, die kunnen leiden tot nog meer stress en gezondheidsproblemen die in deze scenario er oorspronkelijk niet waren. Er wordt ook opgemerkt dat dit kan leiden tot onnodige zorgkosten door extra onderzoeken, wat inefficiënt en belastend is voor de gezondheidszorg.

- Net als in het vorige scenario, wordt hier opnieuw de zorg geuit dat het AI-model belangrijker lijkt dan de patiënt.

6.3.5 Opdracht 6

OPDRACHT - 6

OPDRACHT - 6

Als [patiënt / naaste / individu met verhoogd risico/ burger] vind ik dat verantwoord gebruik van het AI-voorspelmodel ter voorkoming van schade het volgende inhoudt:

Soms is het niet te achterhalen waarom een AI-voorspelmodel een foute uitkomst heeft gegeven (black box / geen zicht op alle data waarvan het model 'geleerd' heeft enz.). Er kan dan niemand die verantwoordelijk gesteld kan worden. Welke partij of partijen moeten dan aansprakelijk worden gehouden voor de opgetreden schade? (en dus alleen of gezamenlijk het geld op tafel leggen).

Deelnemers geven de volgende reflecties op de vraag wat is **verantwoordelijkheid** bij het gebruik van het model:

- Kritisch kijken, ingebouwde controle, altijd uitleg (voor inzicht patiënt), waarschuwing,
- Als er uitkomt dat je op een bepaald moment een extra check nodig hebt (bijv. leeftijd/overgang is het goed. Moet ook in media benoemd worden zodat het breed bekend is.
- Inzetten wanneer nodig, client niet te veel mee belasten door alles zelf in te moeten vullen.
- Inzetten als er al klachten zijn en signaleren / monitoren, arts kan dan ingrijpen wanneer nodig.
- Het aan de hand nemen van de burger/patiënt m.b.t. risico: niet medisch insteken, laagdrempelig en op maat, niet dwingend.
- De huisarts moeten goed / kritisch blijven nadenken. De standaard hoog risico gevallen staan toch sowieso op het netvlies van de huisarts?

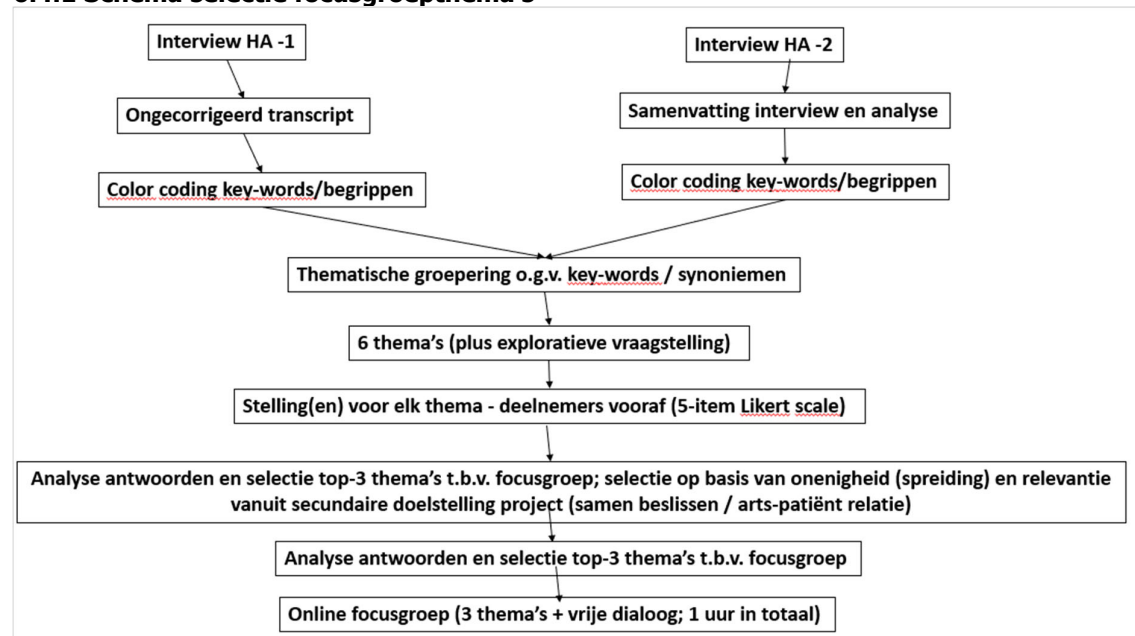
Deelnemers geven de volgende reflecties op de vraag welke partij of partijen **aansprakelijk** gesteld zouden moeten worden als schade optreedt, maar het gebruikte model een black-box model is (niet uitlegbaar):

- Het systeem moet altijd alle invoerdata kunnen achterhalen, geen inzicht is geen optie, het systeem is ter ondersteuning alleen.
- AI moet zorgvuldig bekeken worden en een arts moet altijd de risico's zelf ook inschatten, die kent de patiënt. Ligt aan de kleine lettertjes waar heb je toestemming aangegeven, zijn de risico's beschreven voordat ik toestemming gaf?
- -In het uiterste geval de ontwikkelaar, maar de invoerder blijkt ook systeemcontroles uitvoeren. Er mag nooit alleen van data worden uitgegaan.
- Onzin. Gaat uit van de maakbaarheid van de samenleving. Het is ter ondersteuning en begeleiden van patiënten.

- De huisarts/zorgverlener moet altijd kritisch staan t.a.v. de uitkomst van de AI. De huisarts kan ook aangeven waar AI echt schade aanricht. Eigenlijk weet ik het niet.

6.4 Onderzoeksactiviteit-2 focusgroep met huisartsen

6.4.1 Schema selectie focusgroepthema's



6.4.2 Stellingen incl. reacties t.b.v. selectie top-3 thema's voorde focusgroep

Thema	Stellingen	Reacties
1. Ethiek - recht om niet te weten - impact op arts-patiënt relatie	"burgers hebben het recht om niet te weten als ze verhoogd risico lopen op een hartaanval of beroerte" 'het gebruik van het model zal de arts-patiënt relatie doen veranderen'	3x eens 2x neutraal + 1x helemaal oneens
2. Risico's & afwegingen - Onnodig ongerust door vals-positieve resultaten - Negatieve gevolg van vroegtijdig identificeren: ziek maken - Vals-positieve en vals-negatieve resultaten	"het model zal onvermijdelijk ook vals-positieve resultaten geven. Dit zal mensen onnodig ongerust maken." - het voordeel van vroegtijdig identificeren van mensen die hoge risico lopen (waardoor ze handelingsperspectief krijgen) weegt niet op tegen het onbedoeld 'ziek maken' van mensen	1x eens + 2x neutraal 1x helemaal eens + 1x eens + 1x oneens
3. AI-systeem - Transparantie - Informatie aan patiënten	"Ik zou het model gebruiken of aanbevelen alleen als ik begrijp hoe het model tot de uitkomsten komt." - Stelling-6: "ook de patiënt moet weten waarom hij/zij hogere risico heeft op een hartaanval of beroerte"	2 helemaal eens + 1x neutraal 1x helemaal eens + 2x eens
4. Stakeholders Laaggeletterden / lage SES: gemist door geen deelname aan screenings of informatie niet kunnen lezen en verwerken	Mensen met lage SES en/of laaggeletterdheid zullen niet veel hebben aan het model.	1x neutraal + 1x oneens + 1x helemaal oneens
5. Werkproces huisarts - Integratie met bestaande processen en systemen - Verhoging werkdruk	Voor een succesvolle adoptie en om de werkdruk te verlagen, is het essentieel dat het model wordt geïntegreerd in de bestaande processen en systemen van de huisartsenpraktijk	2x eens + 1x neutraal
6. Beoogd gebruik - Verantwoordelijkheid - Stakeholder gemeente	Stel dat je het model zou verder mogen ontwikkelen en mogen bepalen hoe en door wie het model uiteindelijk gebruikt zou worden. Vink aan wat je van toepassing vindt: - Ik zou de huisarts de eindgebruiker laten worden en ook de verantwoordelijke - Ik zou de burger de eindgebruiker laten worden, maar niet de verantwoordelijke - Ik zou de burger de eindgebruiker laten worden en ook de verantwoordelijke - Ik vind dat de huisarts zich niet bezig moet houden met preventie van data-gedreven vorm. De overheid is verantwoordelijk en aan zet (screeningsprogramma, bevolkingsonderzoek)	erg uiteenlopende responsen

6.4.3 Resultaten focusgroep per thema en per deelnemer

Thema-1: Risico's & afwegingen – afwegen negatieve gevolgen

Deelnemer-1

- Vindt dat patiënten een **keuze** moeten hebben in het wel of niet weten van risico's.
- Waarschuwt voor het idee van een "**maakbare wereld**" en **overmatig vertrouwen** op protocollen en modellen.
- Ziet het risico dat **overmatige medicalisering** mensen ongelukkig maakt

Deelnemer-2:

- Maakt zich zorgen over de **juridische implicaties** als artsen modelvoorspellingen negeren.
- Benadrukt het belang van de arts-patiëntrelatie boven modelvoorspellingen.
- Wijst op het belang van de betrouwbaarheid van voorspellingen en de noodzaak van een betrouwbaarheidsindicator.

Deelnemer-3:

- Ziet potentieel in AI-modellen, maar benadrukt dat implementatie zorgvuldig moet gebeuren.
- Vindt dat het model een duidelijk ondersteunende rol moet hebben voor de arts.
- Maakt zich zorgen over het "ziek maken" van gezonde mensen door vroegtijdige identificatie.
- Wijst op het risico dat het model zijn eigen ongelijk kan bewijzen door vroeg ingrijpen

Thema 2: Thema: Ethiek – impact op arts-patiëntrelatie

Deelnemer-1

- Vindt dat het model ondersteunend moet zijn aan de arts-patiëntrelatie en niet dominant.
- Is bezorgd dat jonge artsen te veel zullen vertrouwen op modellen en richtlijnen.
- Benadrukt het belang van de arts als academicus die zelf kan denken en een "pluis/niet-pluis gevoel" heeft.

Deelnemer-2

- Gelooft dat zijn arts-patiëntrelatie niet wezenlijk zal veranderen door het gebruik van een model.
- Benadrukt het belang van de klinische blik en het gesprek met de patiënt.

Deelnemer-3

- Ziet potentieel voor verbetering van de arts-patiëntrelatie, bijvoorbeeld door meer tijd voor gesprek.
- Waarschuwt dat het model de relatie ook kan ondermijnen als patiënten denken dat ze zelf kunnen googlen n.a.v. het berekende risico.
- Benadrukt dat het model een duidelijke assistent van de arts moet worden.

Thema 3: beoogd gebruik

Deelnemer-1:

- Vindt dat de patiënt een keuze moet hebben, maar niet eindverantwoordelijk kan zijn voor complexe modellen.
- Denkt dat het model **niet per se bij de huisarts hoort**.
- Waarschuwt voor **extra werk voor huisartsen**, vooral met het dreigende artsentekort.

Deelnemer-2:

- Ziet het meer als collectieve preventie dan als taak voor individuele huisartsen.
- Vindt dat de patiënt niet de eindverantwoordelijke kan zijn vanwege de complexiteit van het model.
- Suggereert dat andere instanties, zoals de GGD, mogelijk beter geschikt zijn om het model te implementeren.

Deelnemer-3:

- Ziet het als collectieve preventie, wat volgens huidige standaarden geen basisaanbod is voor huisartsen.
- Vindt dat de huisarts alleen de eindgebruiker kan zijn als het om geïndiceerde preventie gaat.
- Waarschuwt voor het risico dat huisartsen "verwijsmachines" worden voor patiënten die met modelvoorspellingen komen.

Vrije dialoog:

- Bezorgdheid over de toenemende administratieve last voor huisartsen, vooral in het licht van het dreigende artsentekort.
- De noodzaak om de autonomie van artsen te behouden en niet volledig afhankelijk te worden van modelvoorspellingen. Modellen kunnen wel helpen met omgaan met uitdagingen zoals tekort aan huisartsen.
- Culturele verschillen (niet willen weten of je ziek bent)

6.5 Toestemmingsformulieren

6.5.1 Formulier deelnemers workshop

Ziekte-subsidie #dossiernummer 08540122120004

Project: "Clinical DECIDEon: support system voor het cardiovasculair risico management - Verantwoordelijk- en Aansprakelijkheidsperspectief (DECIDE-V&A)"

Het project is een samenwerking tussen: het LUMC, Technische Universiteit Eindhoven, Hogeschool Rotterdam, Patiëntenfederatie Nederland.

Geving van toestemming

Kies je toestemmingen:

- ☐ Ik geef wel toestemming om audio opnames te maken tijdens de sessie, deze worden zo snel mogelijk na de sessie getranscribeerd en hierna verwijderd.
- ☐ Ik geef wel toestemming om video opnames te maken tijdens de sessie, deze worden na analyse verwijderd.

Kies één van onderstaande opties:

- ☐ Ik geef wel toestemming om tijdens de sessie gemaakte foto- en video opnames van mij geanonimiseerd te gebruiken in het evalueren van de tool met andere designers en ik geef toestemming om genomen foto's van mij geanonimiseerd in (online) publicaties gerelateerd aan dit project te gebruiken.
- ☐ Ik geef wel toestemming om tijdens de sessie gemaakte foto- en video opnames van mij geanonimiseerd te gebruiken in het evalueren van de tool met andere designers en ik geef geen toestemming om genomen foto's van mij geanonimiseerd in (online) publicaties gerelateerd aan dit project te gebruiken.
- ☐ Ik geef geen toestemming om tijdens de sessie gemaakte foto- en video opnames van mij geanonimiseerd te gebruiken in het evalueren van de tool met andere designers en ik geef wel toestemming om genomen foto's van mij geanonimiseerd in (online) publicaties gerelateerd aan dit project te gebruiken.
- ☐ Ik geef geen toestemming om tijdens de sessie gemaakte foto- en video opnames van mij geanonimiseerd te gebruiken in het evalueren van de tool met andere designers en ik geef geen toestemming om genomen foto's van mij geanonimiseerd in (online) publicaties gerelateerd aan dit project te gebruiken.

¶

¶

Ik heb dit toestemmingsformulier gelezen en begrepen en de mogelijkheid gehad om vragen te stellen. Ik ga ermee akkoord vrijwillig deel te nemen aan de workshop geleid door Dr. Ildikó Vajda, senior adviseur bij de Patiëntenfederatie Nederland.

¶

Naam en handtekening

¶

¶

Datum en plaats

¶

19 juli 2024, Utrecht

6.5.2 Formulier deelnemers focusgroep

Informatie voor deelname aan wetenschappelijk onderzoek

DECIDE-VeRA

~~Official title:~~ Clinical ~~DECIDE~~ support system for ~~cardiovascular~~ risk management in primary health care: a Responsible and Liability Perspective for AI-CDSS (~~DECIDE-VeRA~~).
De studie is goedgekeurd bij het LUMC.

Geachte heer/mevrouw,

Wij vragen u om deel te nemen aan een studie over een Kustmatige Intelligentie voor beter management van hart- en vaatziekte. U beslist zelf of u mee wilt doen.

De studie wordt geleid door het Leids Universitair Medisch Centrum (LUMC). Het onderzoeksteam bestaat uit wetenschappers, ethici en juristen: André Krom, María Villalobos en Niels H. Chavannes (LUMC), Sara Colombo en Loes van Renswouwe (Technische Universiteit Eindhoven), Tycho de Graaf en Jan van Staalduinen (Leiden Universiteit), Maaike Harbers en Tjitske Portogies (Hogeschool Rotterdam), Ildikó Vajda (Patiëntenfederatie Nederland), Indre Kalinauskaitė (Universitair Medisch Centrum Utrecht).

Doel van de studie: Ons doel is om de ethische en juridische uitdagingen van een Kustmatige Intelligentie voor hart- en vaatziekte beter te begrijpen en oplossingen te vinden.

Wat moet u doen als u meedoet? U neemt deel aan één online focusgroep van 1 uur.

Privacy: Alles wat u vertelt blijft geheim. Uw naam staat niet op de opname en zal niet worden gepubliceerd. De opnames worden nooit openbaar gemaakt en zijn alleen voor deze studie bedoeld. We willen de resultaten van deze studie publiceren maar geen persoonlijke informatie (zoals je naam) wordt gepubliceerd. Niemand kan achterhalen dat het over u ging. Aan het eind van het project gaan we de informatie anonimiseren. We bewaren deze data 10 jaar veilig in het LUMC. Onderzoeksgegevens worden volgens EU-privacy standaarden opgeslagen en verwerkt.

Heeft u vragen?

Neem contact met het onderzoeksteam: María Villalobos (m.j.villalobos_quesada@lumc.nl, 0642697065)
Voor klachten: wetenschapscommissie_PHEG@lumc.nl, 071-5268444
Voor privacy: Functionaris Gegevensbescherming (privacy@lumc.nl), www.autoriteitpersoonsgegevens.nl

Toestemming

- ☐ Ik heb de informatiebrief gelezen en mijn vragen zijn goed beantwoord.
- ☐ Ik weet dat meedoen vrijwillig is en dat ik op ieder moment kan stoppen.
- ☐ Ik geef toestemming aan de onderzoekers om mijn gegevens te verzamelen en te gebruiken voor dit onderzoek.
- ☐ Ik begrijp dat voor de controle van het onderzoek sommige mensen mijn gegevens kunnen inzien.
- ☐ Ik geef toestemming om eventueel na dit onderzoek benaderd te worden voor deelname aan vervolgonderzoek.

☐ Ik wil meedoen aan dit onderzoek.

Mijn naam is:

Datum: ____ / ____ / ____