

ARTIFICIAL INTELLIGENCE KANSEN BENUTTEN IN DE GEZONDHEIDSZORG

Werken aan vertrouwen

Auteurs: Ir. P.J. van Eck, Dr. R. van der Kleij,
Prof. Dr. J. Janssen, Dr. Ir. A. de Keijzer,
L. Bambach, M.M.C.M. Bastiaensen, Drs. R.M.G. de
Beer, A.C.C. van Erk MSc, R.A.M. Korenromp MSc,
Dr. L. Oudenhuijsen, E. ten Thij.

Datum: 7 juli 2025

avans
hogeschool



MANAGEMENTSAMENVATTING

1. Mogelijkheden nemen toe, net als uitdagingen en zorgen bij toepassen AI en GenAI¹ in de gezondheidszorg².
2. Onderzoeksaanpak: gecodeerde analyse van twee focusgroepen en literatuurstudie van 26 artikelen.
3. Resultaten.
 1. Gebruik: daadwerkelijk ingezette AI-toepassingen.
 - Focusgroepen: benoemen zoeken, dossiervorming en diagnostiek.
 - Literatuur: vult aan met toepassingen voor begrijpen medische data, ondersteuning in clientbehandeling en zorglogistiek.
 2. Mogelijkheden: verwachte AI-toepassingen en -ontwikkelingen.
 - Focusgroepen: AI inzetten voor werkdrukverlaging, meer taaltoepassingen en begrijpen van data.
 - Literatuur: verwacht persoonlijker, empathischer en efficiëntere taal, besluitondersteuning en zorg.
 3. Hindernissen: houden organisaties tegen in het gebruiken van AI-toepassingen.
 - Focusgroepen: zien gebrek aan vertrouwen en kennis als belangrijkste hindernissen voor inzetten AI.
 - Literatuur: benoemt vertrouwen in en AI willen toepassen in clientgerichte zorg als hindernissen. Te ontwikkelen kennis en ervaring in een snel ontwikkelend AI-landschap en -regelgeving.
 4. Oplossingen: de manieren om hindernissen aan te pakken.
 - Focusgroepen: zien als oplossingen: kennis en oplossingen delen, leren met en van elkaar, experimenteren, en aanbrengen van focus door besturen.
 - Literatuur zegt: houd mensen eindverantwoordelijk, leg focus op transparantie van gebruikte data en systemen, en voldoe aan (EU) regelgeving.
 5. Risico's: bedreigingen bij toepassen AI in een organisatie.
 - Focusgroepen. Systeem/gegevensaanvallen. Afhankelijkheid, vooroordelen, kennisverlies. Niet opvolgen regels. Verlies gevoelige persoonlijke info. Niet toepassen van AI.
 - Literatuur concretiseert de AI risico's sociaal, ethische, duurzaamheid, betrouwbaarheid, transparantie, afhankelijkheid, databescherming, aanvalsoppervlak
 6. Maatregelen: aanpakken van AI-risico's.
 - Focusgroepen belangrijke maatregelen mens blijft eindverantwoordelijk, samenwerken en leren, transparante afgeschermd systemen en regels opstellen en opvolgen.
 - Literatuur geeft maatregelen aan om AI met vertrouwen toe te passen op vier aspecten, sociale versterking, kennisontwikkeling, systeembeveiliging en besturen
4. Conclusie en aanbevelingen
 - AI-kansen benutten in de gezondheidszorg; focus op mensen die AI willen gebruiken, leren toepassen, en werk aan vertrouwen.
 - Aanbeveling voor vervolgonderzoek naar key drivers voor vertrouwen in AI, hoe deze te versterken en verantwoord gebruik van AI te stimuleren.

INHOUDSOPGAVE

- 1. Aanleiding (4)*
- 2. Onderzoeksaanpak (5)*
- 3. Resultaten (6)*
 - 1. Gebruik: daadwerkelijk ingezette AI-toepassingen (6)*
 - 2. Mogelijkheden: verwachte AI-toepassingen en -ontwikkelingen (7)*
 - 3. Hindernissen: houden organisaties tegen in het gebruiken van AI-toepassingen (8)*
 - 4. Oplossingen: de manieren om hindernissen aan te pakken (10)*
 - 5. Risico's: bedreigingen bij toepassen AI in een organisatie (11)*
 - 6. Maatregelen: aanpakken van AI-risico's (13)*
- 4. Conclusie en aanbevelingen (15)*
 - 1. Identificeren huidige stand van zaken (15)*
 - 2. Analyseren van de behoeften in de beroepspraktijk en belangrijkste kwetsbaarheden (16)*
 - 3. Formuleren van aanbevelingen voor verder vervolgonderzoek (17)*
- Bijlagen (18)*
 - 1. Definities (19)*
 - 2. Aanpak van het onderzoek (20)*
 - 3. Quote analyses focusgroepen en literatuurstudie (22)*
 - 4. Bronnenlijst (28)*

1. AANLEIDING

Toenemende mogelijkheden, uitdagingen en zorgen bij toepassen Artificial Intelligence in de gezondheidszorg

Gebruik van AI en zorgen hierover nemen toe

- Toenemende integratie van GenAI in diverse aspecten van ons leven^I.
- Toenemende complexiteit en kracht van AI systemen^I.
- Tegelijkertijd toenemende vatbaarheid voor aanvallen en misbruik van AI systemen^I.
- Gewenste (te) snelle invoering van AI met als gevolg vragen over betrouwbaarheid^{II}.
- Introductie van Cyber Resilience Act (CRA) en Digital Services Act (DSA) door EU bepaald^{III}.
- Gezondheidszorg: vertrouwelijkheid en gevoeligheid van informatie over mensen^{IV}.

Doel is achterhalen stand van zaken en behoeften

Doel is een dieper inzicht in de mogelijkheden en uitdagingen bij de inzet van GenAI binnen gezondheidszorg.

Onderzoeksvragen

1. Identificeren huidige stand van zaken;
2. Analyseren van de behoeften in de beroepspraktijk en belangrijkste kwetsbaarheden; en
3. Formuleren van aanbevelingen voor verdere actie en onderzoek.

^I Zie ook: : <https://www.cybersecurityraad.nl/documenten/brieven/2023/12/22/informerende-brief-aan-de-staatssecretaris-van-bzk-over-generatieve-ai-en-cybersecurity>

^{II} Ibidem

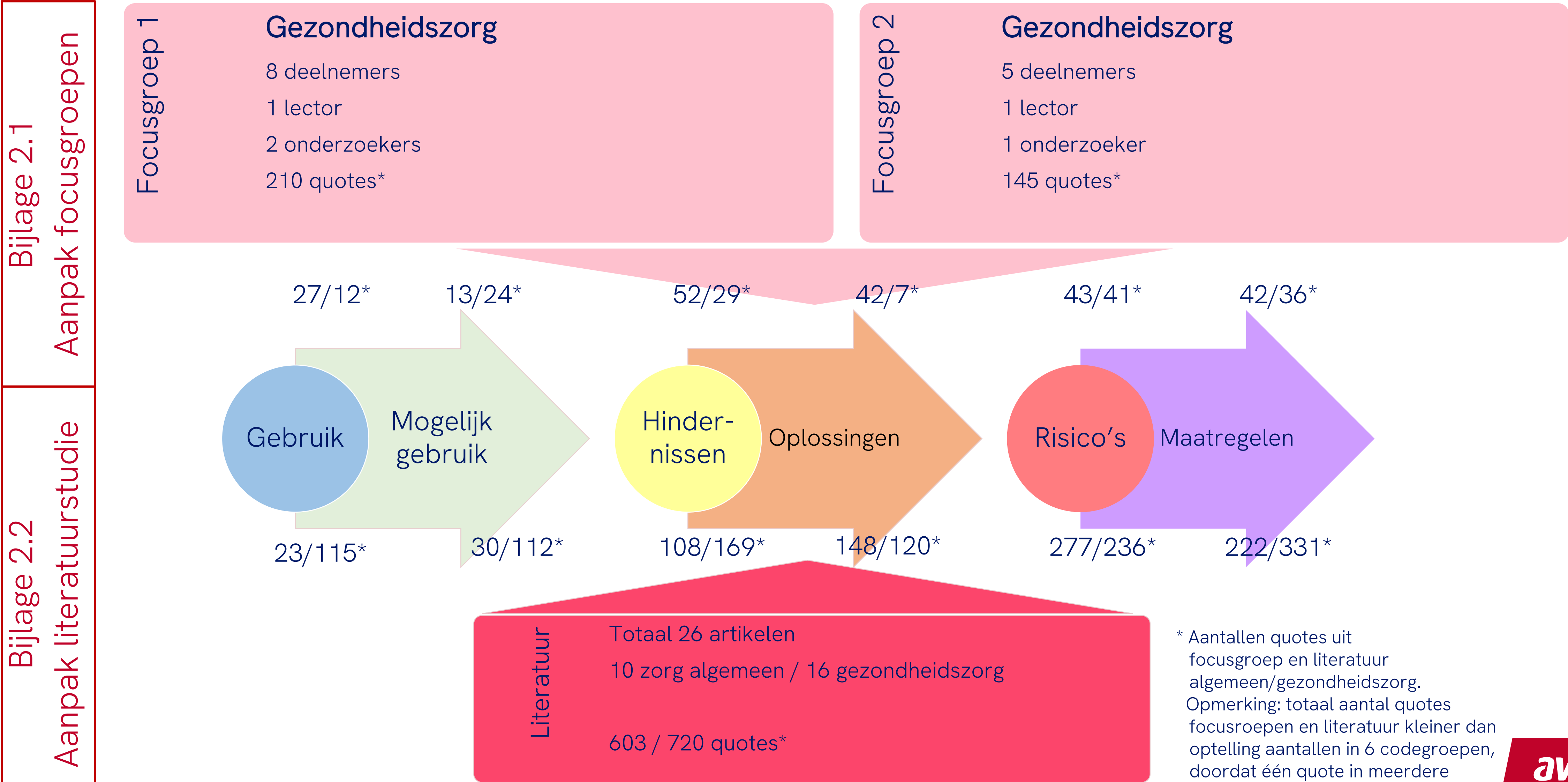
^{III} <https://www.digitaleoverheid.nl/nieuws/dsa-voor-alle-digitale-diensten-van-kracht/>

^{IV} Algemene Inlichtingen- en Veiligheidsdienst & Rijksinspectie Digitale Infrastructuur. (2024). Generatieve AI: een transformatieve impact op cybersecurity. Den Haag: AIVD.

⁴ Geraadpleegd op https://www.aivd.nl/binaries/aivd_nl/documenten/publicaties/2024/10/17/generatieve-ai-eeen-transformatieve-impact-op-cybersecurity/Generatieve+AI.+Een+transformatieve+impact+op+cybersecurity.pdf

2. ONDERZOEKSAANPAK

Twee focusgroepen in de gezondheidszorg. Daarna een literatuurstudie van 26 artikelen. Gecodeerd met Atlas.ti. Uitgevoerd door drie lectoren en acht onderzoekers.



4.1 RESULTATEN – DAADWERKELIJK TOEGEPASTE AI

Focusgroepen benoemen zoeken, bijhouden dossiers en diagnostiek.

Literatuurstudie beschrijft meer gebruikte toepassingen in begrijpen, medische data en ondersteuning.

Focusgroepen

Start AI-gebruik met zoeken, taaltoepassingen en begrijpen.

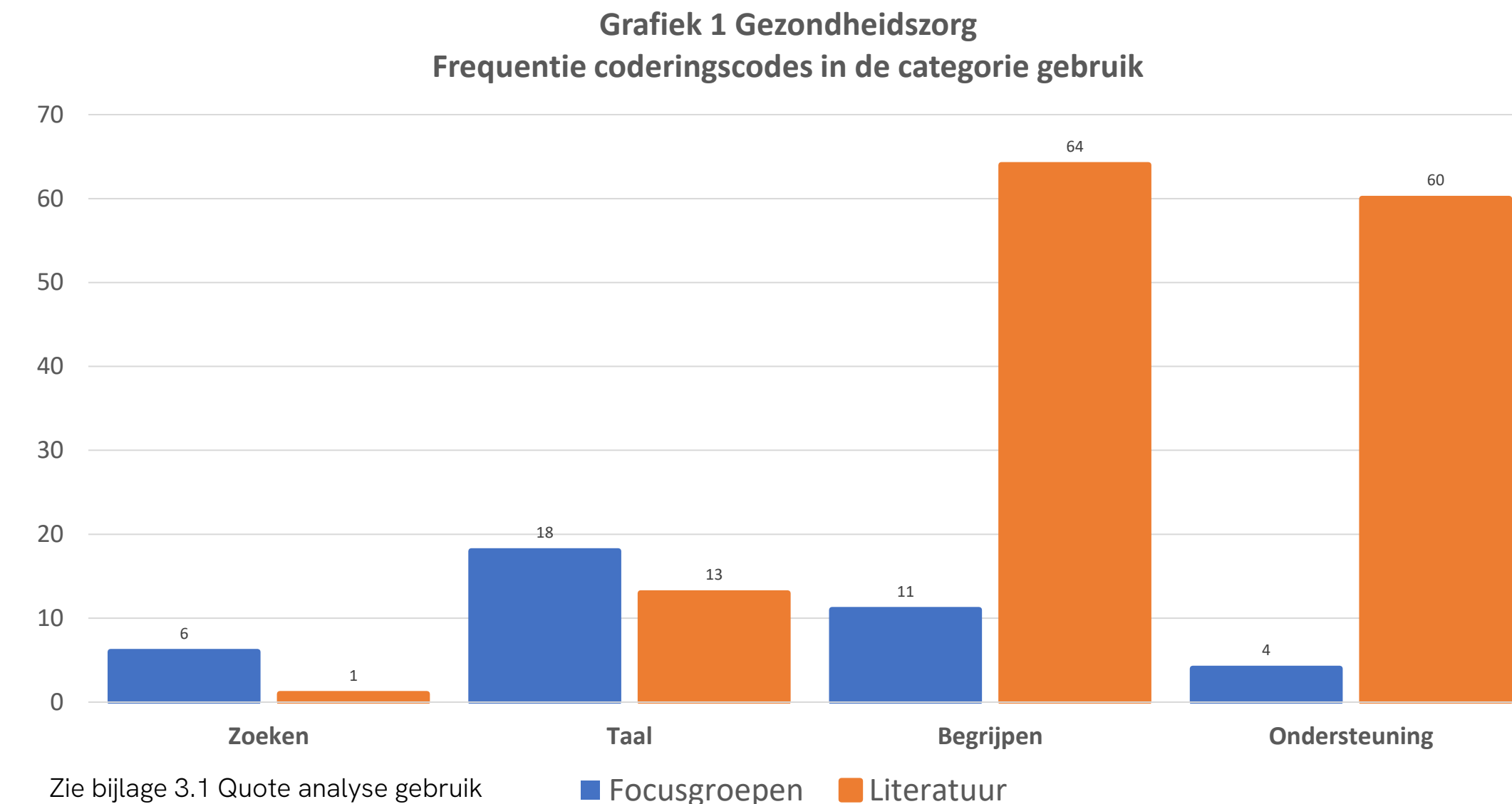
- Start AI-gebruik met zoeken, AI als een verbeterde zoekmachine.
- Daarna volgen AI-taaltoepassingen zoals verslaglegging in EPDs, hulp bij schrijven van bijvoorbeeld e-mails en ook vertaling, om te communiceren in eigen taal.
- Ook wordt AI ingezet om te begrijpen, zoals diagnostische toepassingen in radiologie, maar ook samenvattingen uit een EPD om gesprekken voor te bereiden.

Literatuur

AI gebruikt om medische data te begrijpen in diagnostiek, clientmonitoring en training.

Gebruikte ondersteuningsmogelijkheden bij voorspellen, besluiten en behandelen van cliënten. En ook zorg-logistieke toepassingen.

- Zoeken in cliëntgegevens en deze gebruiken in medische AI-toepassingen.
- Taaltoepassingen zoals verslaglegging, schrijven, vertalen ook in de literatuur. Aanvulling, met AI beantwoorden van cliëntvragen, client specifiek gemaakt door AI-gebruik van cliëntgegevens.
- Begrijpen is in aantallen de meest voorkomende categorie in de literatuur.
 - Meer diagnose mogelijkheden benoemd bv: beeldverwerking/verbetering, SEH, symptoomherkenning, herkennen hartaandoeningen/kanker.
 - Hetzelfde voor leren: trainingsdata anoniem genereren, laatste ontwikkelingen bijhouden, digital twins.
 - Uitgebreidere toepassingen EPD: vastlegging cliëntmonitoring/-veiligheid, medische calamiteitanalyses, cliënt specifiek vragen beantwoorden.
- Ondersteuning, literatuur maakt meer actuele gebruiksmogelijkheden concreet.
 - Medische ondersteuning bij cliëntdata-analyse: sneller, preciezer, effectiever en ook herkennen van zeldzaamheden.
 - Besluitondersteuning en behandeling: classificeren, beoordelen, therapiekeuzes, triage en behandelkeuzes op afstand.
 - Voorspellen: epileptische aanvallen, cliëntgezondheid
 - Maar ook plannen en voorspellen met name gericht op zorglogistiek: afspraken maken, roostering, instroom, heropname, ligduur, bedbezetting, no-show.



4.2 RESULTATEN – VERWACHTE AI-TOEPASSINGEN EN –ONTWIKKELINGEN

Focusgroepen verwacht AI inzet voor werkdrukverlaging, meer taaltoepassingen en begrijpen van data. Literatuurstudie verwacht persoonlijker, empathischer en efficiëntere taal, besluit- en zorgondersteuning.

Focusgroepen

AI-gebruik verder inzetten voor werkdrukverlaging en meer met taal en begrijpen.

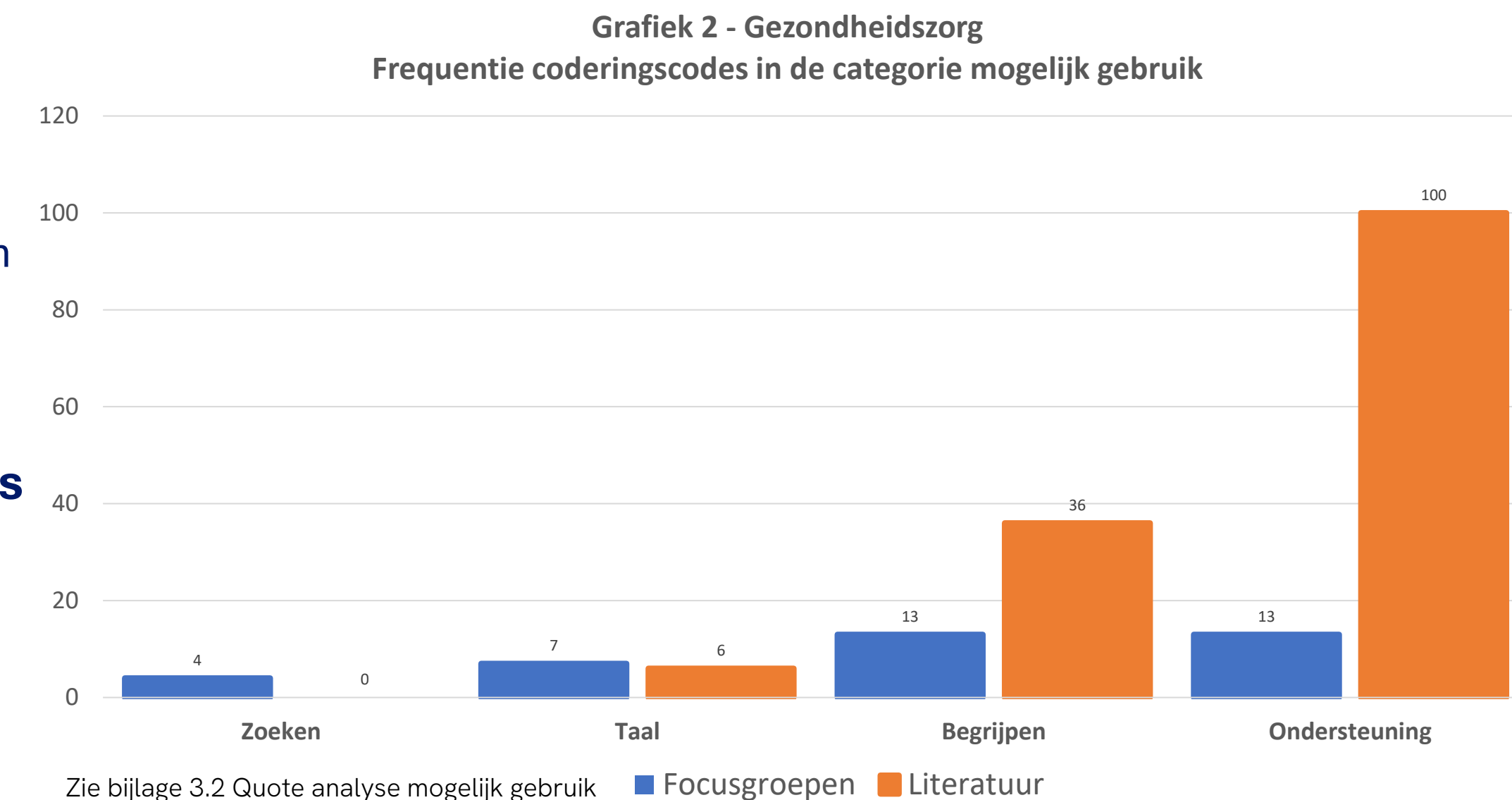
- Meer AI-taaltoepassingen in verslaglegging in EPD's en hulp bij schrijven.
- Bij begrijpen ziet men kansen voor vereenvoudigen toegang tot en inzicht geven in dossiers. En verder inzetten van AI voor diagnostiek. Leren richt zich vooral op het leren gebruiken van AI.
- Kansen voor de inzet van AI in ondersteuning. Werkdrukverlaging en plannen. Vergroten autonomie van de cliënt.

Literatuur

Taal tekst, spraak, video's persoonlijker en empathischer. In gegevensanalyse combineren van verschillende databronnen met persoonlijker uitkomsten in diagnoses en behandelingen. Genereren van data, scenario's en oefenstof om te leren.

Versnellen en verbeteren besluitvorming, zorglogistiek en cliëntgerichtheid.

- Taal.
 - Toepassingen gericht op verminderen administratie, spraak naar tekst, vastleggen (medische) dialoog, antwoord (suggesties) op cliëntvragen.
 - Tekst naar spraak, samenvatten en vereenvoudigen medische dialogen voor cliënten.
 - Persoonlijker en empathischer teksten/video's.
- Begrijpen, combineren van verschillende soorten data in analyses, diagnoses en voorspellingen
 - Analyseren en voorspellen op basis van grote hoeveelheden, verschillende bronnen en persoonlijke data. Vroegtijdige herkenning. Vergroten accuratesse. Mobiele zorgtoepassingen mogelijk maken. Van reactieve naar proactieve zorg. Landelijke (data) infrastructuur.
 - Dossierkennis/EPD's. Gebruik clientgegevens in analyses, diagnoses. Geschikte cliënten zoeken voor onderzoek.
 - Leren. Uitleg van complexe teksten, samenvatten, oefenvragen. Koppelen aan persoonlijke leerpatronen. Doelgroepgericht maken voor studenten, cliënten en algemeen publiek. Genereren materialen en scenario's om te leren, los van persoonlijke clientgegevens.
- Ondersteuning, in aantallen het meest voorkomend in literatuur.
 - Data analyse. Herkennen/virtuele representatie van gegevens ook naar VR. Haptische toepassingen druk, warmte, kou. Meten emoties, stress. Combineren gegevensbronnen. Versnellen proces van medicijnontwikkeling/inzicht in structuren/simulaties. Detecteren, analyseren, trendanalyses van aanvallen.
 - Besluitondersteuning. VR behandelingen ipv of aanvullend op bestaande behandelingen. Versnellen diagnose/screening/voorspellen/triage. Persoonlijke maken diagnose/behandeling. GAI chatbots persoonlijker/vriendelijker/empathischer. Ziekenhuisverplaatste zorg/telemonitoring/veiliger. Verwoorden van medische dialoog op verschillende taalniveaus.
 - Plannen/werkdruk. Kwaliteit, toegankelijkheid en kosten verbeteren. Verminderen onnodige zorgafspraken. Plannen afspraken/medicijn innames. Voorspellen en detecteren SEH/opname/ontslag/risico's. Content genereren. Repetitieve taken verminderen. Efficiënter maken zorglogistiek.
 - Client. Combineren van data, technologieën en sensoren. Vragen beantwoorden persoonlijker, beter, sympathieker. Verbeteren cliëntcommunicatie tekst naar spraak/video's. Vereenvoudigen teksten, samenvatten. mHealth-apps. Van reactief naar proactieve zorg. Ziekenhuisverplaatste zorg.



4.3 RESULTATEN – HINDERNISSEN VOOR GEBRUIKEN AI (1/2)

Focusgroepen zien gebrek aan vertrouwen en kennis als belangrijkste hindernissen voor inzetten AI. Literatuur benoemt vertrouwen in en AI willen toepassen in clientgerichte zorg als hindernissen. Te ontwikkelen kennis en ervaring in een snel ontwikkelend AI-landschap en –regelgeving.

Focusgroepen

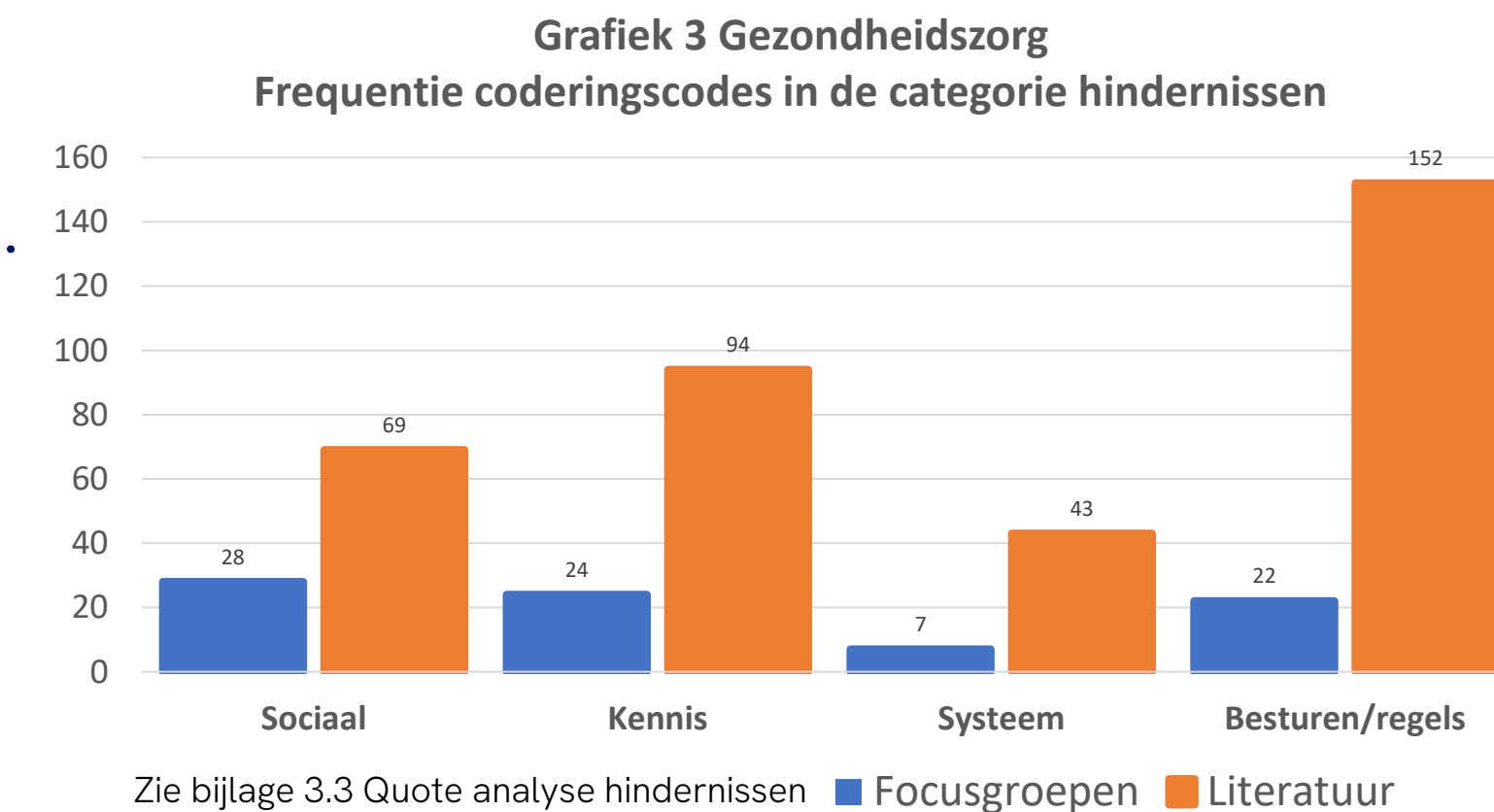
Zoeken naar kennis en vertrouwen om AI in te zetten. Beperkte capaciteit. Complexe regels.

- Sociaal. Gebrek aan vertrouwen: angst, voorzichtigheid, juistheid van de uitkomsten, beheersbaarheid, draagvlak bij cliënten.
- Kennis. Verschillen in en tussen organisaties. Onbekendheid met mogelijkheden, risico's en aanpak van deze transformatie. Ontbreken van AI-vaardigheden. Soms versterkt door beperkte Engelse taalvaardigheid en geestelijke / fysieke beperkingen.
- Systeem. Omvang en tempo van de AI ontwikkelingen weerhouden organisaties; het gevoel steeds verder achter te lopen.
- Besturen/regels. Concretiseer behoefte om met AI te focussen op wat echt nodig is. Meerdere belangrijke prioriteiten voor de handen in de zorg. Wet- en regelgeving en hoe daaraan te voldoen. Complexe subsidievereisten.

Literatuur

Vertrouwen in AI vraagt om betrouwbare uitkomsten. Alle stakeholders moeten hierbij betrokken, kennis en ervaring opdoen. Snelheid van AI ontwikkelingen maken het moeilijk bij te benen. Bestuurders moeten een voortrekkersrol nemen in opstellen van regels en bij richting aangeven.

- Sociaal.
 - Vertrouwen. Accuratesse en betrouwbaarheid van AI uitkomsten is cruciaal voor vertrouwen in deze toepassingen. Dit geldt voor alle stakeholders cliënten, zorgprofessionals, ontwikkelaars, managers. Opbouwen vertrouwen is multidimensionaal o.a. geheimhouding, beveiliging, juistheid, consistentie, kosten, werkgelegenheid, gebruiksvriendelijkheid, transparantie.
 - Willen. AI willen integreren in gezondheidszorgprocessen staat gebruik in de weg. Cliënten en zorgprofessionals willen cliëntgerichte zorg verlenen. Voorbeeld Is het wel of niet nodig om nog via de huisarts te gaan voor een doorverwijzing naar een specialist of mag een AI toepassing dat doen? Managers willen positieve kosten/baten.
- Kennis.
 - Ontbreken van medische training van AI toepassingen zelf staat gebruik AI in de weg. Ook het trainen en opdoen van praktijkervaring door medische professionals en cliënten in het gebruiken van AI. Versterkt door personeelstekorten en -wisselingen in de gezondheidszorg. Meer algemeen gaat het om trainen en ervaring opdoen in cyber-weerbaarheid door alle betrokken stakeholders.
- Systeem.
 - Snelheid, omvang, complexiteit in de ontwikkelingen van AI mogelijkheden maken evaluatie, validatie, certificering en regelgeving moeilijk bij te benen en belemmeren zo invoering, integratie en acceptatie van AI in de gezondheidszorg.
 - Data en techniek. Zorgen om betrouwbaarheid van gebruikte data. Beschermen en beveiliging van gevoelige persoonlijke gegevens. Maar ook toegang krijgen en bijbehorende kosten tot betrouwbare datasets. Ook de beperkte omvang van de Nederlandse taal is een mogelijke belemmering. Afhankelijkheid van beperkt aantal hardware-leveranciers.



4.3 RESULTATEN – HINDERNISSEN VOOR GEBRUIKEN AI (2/2)

Literatuur (vervolg)

- Bestuur/regels.
 - Voldoen aan regelgeving bij gebruik van AI toepassingen is een hindernis voor gebruik. Naast zich ontwikkelende regelgeving zijn bewezen standaarden nodig om deze in organisaties toe te passen. Nog te beantwoorden vragen zoals is algemene AI regelgeving voldoende of is specifieke regelgeving vereist voor medische toepassingen, aansprakelijkheid bij gebruik van AI, noodzaak en frequentie van herhaalde certificering bij continue lerende AI in medische toepassingen.
 - Andere en de hoeveelheid aan andere projecten met prioriteit in de gezondheidszorg is een te overwinnen obstakel voor toepassen AI. Hierbij spelen capaciteit, tijd en budget beperkingen een rol. Net als negatieve ervaringen met ICT investeringen en verwachte hoge onderhoudskosten van deze systemen.
 - Ondernemerschap en organisatie-inrichting. Durf en overtuiging om AI te gebruiken ontbreken soms. Zeker als toepassingen zich – nog niet – hebben bewezen. Naast durf gaat het ook om organisatorisch faciliteren van vernieuwende initiatieven en betrekken van alle stakeholders hierbij cliënten, zorgprofessionals, ontwikkelaars, beheerders.
 - Zorgen rondom privacy vragen continue aandacht bij toepassen van AI. Controle over gevoelige persoonlijke gegevens, databeveiliging, transparantie over opslag gebruik en in lerende AI-modellen

4.4 RESULTATEN – OPLOSSINGEN PAKKEN DE HINDERNISSEN AAN

Focusgroepen kennis en oplossingen delen, leren met en van elkaar, experimenteren, focus aanbrengen door besturen.

Literatuur houdt mensen eindverantwoordelijk, transparantie van gebruikte data en systemen, voldoe aan (EU) regelgeving.

Focusgroepen

Zien oplossingen in kennis delen, leren gebruiken, experimenteren, en focus van bestuurders.

- Kennis delen en leren gebruiken worden genoemd als oplossingen voor gebrek aan kennis en vergroten van het vertrouwen.
- Hierbij is samenwerking binnen en buiten de organisaties cruciaal. Zowel binnen het eigen netwerk als ook met leveranciers.
- Probeer AI-mogelijkheden uit, zet hierbij kleine stapjes en houd ook het aantal betrokken in het begin klein.
- Leer van anderen als het gaat om welke AI-toepassingen en technieken werken en aan de vereisten voldoen, zoals databescherming en juiste uitkomsten.
- Van bestuurders wordt verwacht dat zij beleid en regels opstellen, richting aangeven en prioriteren.

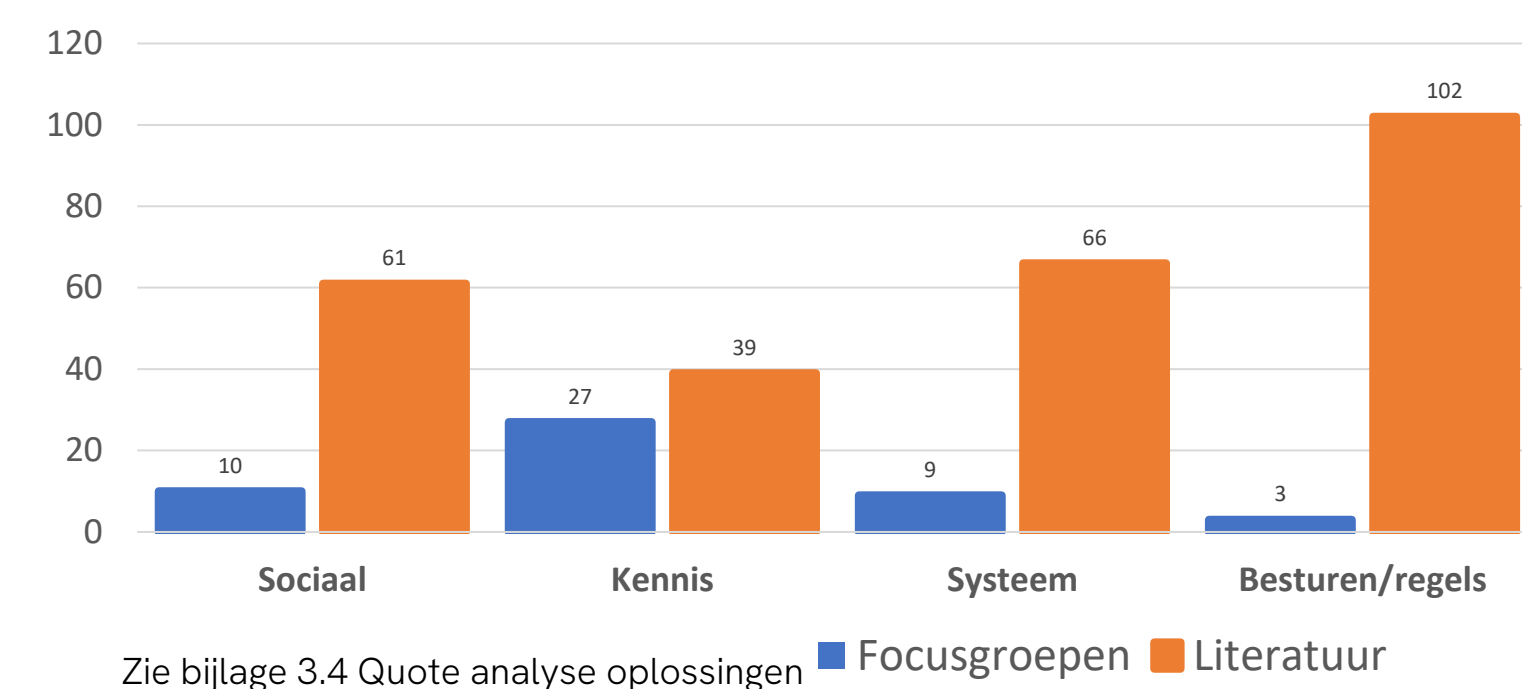
Literatuur

Eindverantwoordelijkheid voor AI bij mensen en laat hen samenwerken. Stimuleer leren en experimenteren. Focus op transparantie en gedegen bescherming van gebruikte data en systemen.

Voldoe aan (EU) wet en regelgeving. En vul die als organisatie praktisch in met beleid en richtlijnen en volg deze aantoonbaar op.

- Sociaal
 - Menselijke expertise, interactie, overzicht en eindverantwoordelijkheid zijn randvoorwaarden voor de inzet van AI-toepassingen. Dit geldt zowel in de AI-model-trainingsfase en –toepassingsfase.
 - Samenwerken en delen, van kennis en gegevens, over zorgverlening met behulp van AI en tijdens het ontwikkelen van toepassingen en modellen is essentieel voor creëren van vertrouwen. Betrek hierbij alle stakeholders, zoals eindgebruikers, bestuurders, leveranciers, toezichthouders. Nationaal en internationaal.
- Kennis.
 - Leren. Integreer kennis en expertise van zorgprofessionals en -managers en hun begrip van de werking en validatie van AI-modellen en -toepassingen. Investeer in leren en training van digitale vaardigheden, om daarmee vertrouwen en weerbaarheid te vergroten van betrokkenen, organisatie en toepassing. Benadruk en zet in op voortdurend verbeteren door anticiperen, monitoring, reageren en evalueren.
 - Experimenteren. Experimenteer en maak proof-of-concept studies onderdeel van AI-integratie in zorgprocessen, AI-gebruikersacceptatie en AI-validatie. Een “public-space” voor het delen van kennis en ervaring helpt in alle fase van AI-ontwikkelingen: ontwerp, oplossingen voor belemmeringen, beleid, gebruikersrichtlijnen en standaardisatie.
- Systeem.
 - Techniek. Een voorstel voor een (inter)nationaal gezondheidsdata infrastructuur (European Health Data Space) te gebruiken voor primair zorg, onderzoek en productontwikkeling. Combinatie van data uit zorgorganisaties en thuis verzamelde data. Voldoen aan wet en regelgeving vraagt om opt-out, machine unlearning mechanismen, risicobeheersing.
 - Transparantie. Explainable AI vergroot vertrouwen en ethisch gebruik van AI-toepassingen, o.a.. gebruik van zorgdata, testen en validatie van modellen, juistheid van antwoorden en in welke context, doorlopen stappen in modelontwikkeling en –testen. Publieke toegang tot modellen, software-code, uitkomsten van (onafhankelijke) validatie.
 - Standaardisatie. AI toepassingen moeten niet alleen de uitkomsten presenteren maar ook de onderliggende (client)data met bijbehorende samenvattingen. Opstellen van (inter)nationale standaarden en bij navolgen in AI-modellen certificering hiervan.
- Bestuur/regels.
 - Regelgeving. Voldoen aan EU wetgeving en richtlijnen, zoals General Data Protection Regulation, Medical Device Regulation, AI-Act, Cyber Resilience Act, Product Liability Directive, voorgestelde AI Liability Directive, rulings of the EU Court of Justice.
 - Beleid. Opstellen van beleid en richtlijnen door organisaties voor het in goede banen leiden van AI initiatieven. Denk aan data beschikbaarheid, tijd, kosten, integratie in de medische praktijk, benoemen en beleggen van verantwoordelijkheden van alle betrokken partijen. Hoe om te gaan met technologie, inkoop, risico's, incidenten.
 - Monitoren. Gebruik van verschillende datasets om AI-modellen te trainen en te testen. Detecteren van data drift. Regelmatig updaten, screenen en valideren van systemen. Proactieve risico inventarisaties.
 - Afschermen. Transparantie over gebruik van medische gegevens, ook richting cliënten. Afschermen en encryptie van gevoelige informatie en deze alleen lokaal gebruiken. Cruciaal voor het vertrouwen in toepassingen en voor het respect en vertrouwen van cliënten en hun autonomie.
 - Geld. Het verlagen van kosten voor het gebruik van AI vergroot de inzet hiervan. Dit geldt voor zowel cliënten als ook voor de betrokken stakeholders in de zorg.

Grafiek 4 Gezondheidszorg
Frequentie coderingscodes in de categorie oplossingen



4.5 RESULTATEN – RISICO'S BIJ TOEPASSEN AI IN ORGANISATIES (1/2)

Focusgroepen. Systeem/gegevensaanvallen. Afhankelijkheid, vooroordelen, kennisverlies. Niet opvolgen regels. Verlies gevoelige persoonlijke info. Niet toepassen van AI

Literatuur concretiseert de AI risico's sociaal, ethische, duurzaamheid, betrouwbaarheid, transparantie, afhankelijkheid, databescherming, aanvalsoppervlak.

Focusgroepen

Systeem/gegevensaanvallen. Afhankelijkheid, vooroordelen, kennisverlies. Niet opvolgen regels. Verlies gevoelige persoonlijke info. Niet toepassen van AI.

- Sociaal. Risico's voor mensen, zijn hun eigen vooroordelen, onmogelijkheid om AI-resultaten te verklaren en afhankelijkheid van AI. Bij werkdruk vraagt één deelnemer zich af wat het risico is van het niet toepassen van de AI-mogelijkheden.
- Kennis. De inzet van AI leidt mogelijk tot een verlies aan kennis.
- Systeem. Databescherming en feitelijke juistheid van de uitkomsten van de gebruikte AI-toepassingen. Aanvallen op en de beschikbaarheid van systemen.
- Regels. Privacy en non compliance brengen beide in relatie tot gevoelige persoonlijke informatie waar beide sectoren mee werken.

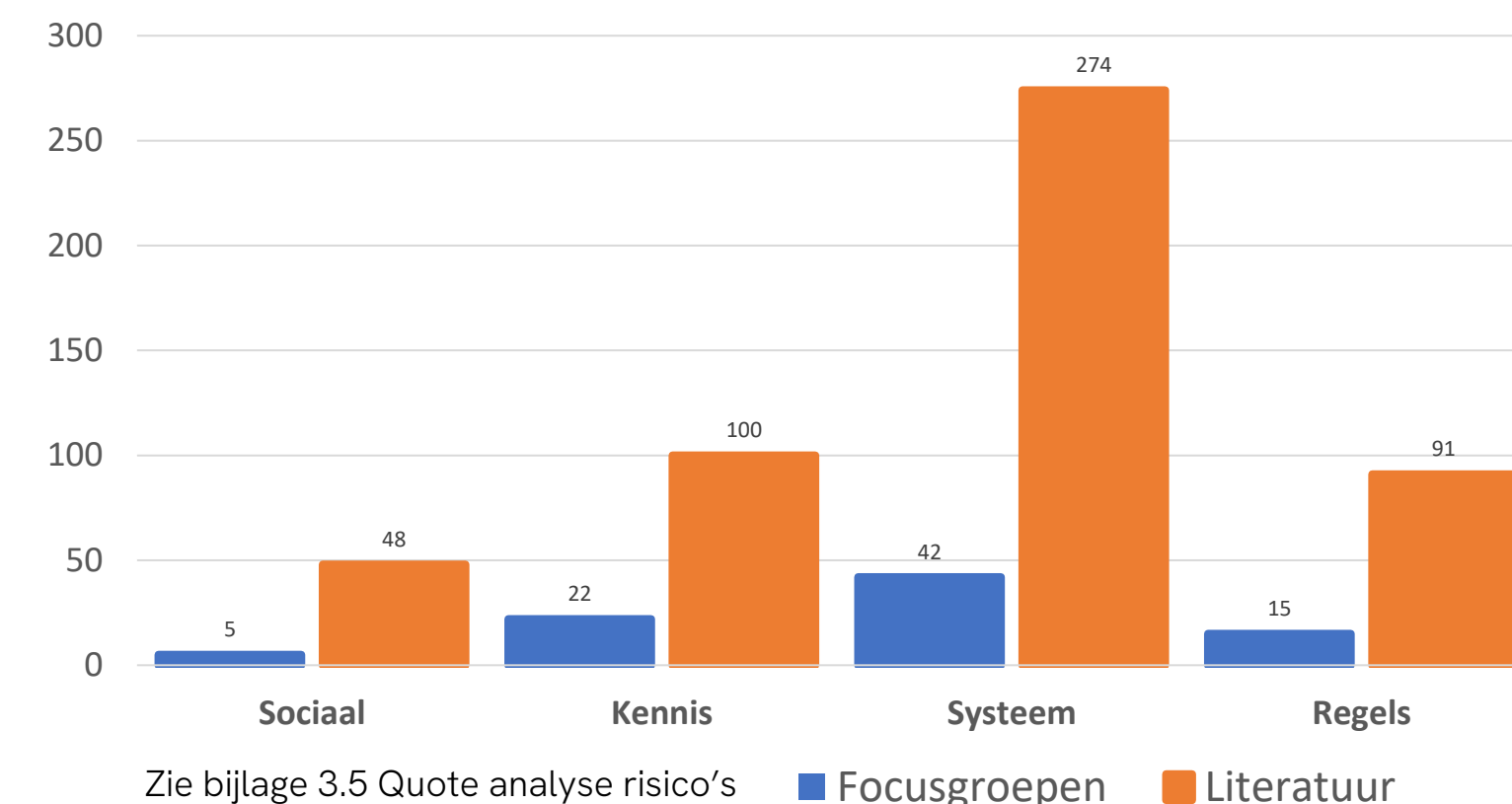
Literatuur

Grootschalig toepassen van AI leidt tot ongewenste ontwikkelingen; sociaal, ethische en op duurzaamheid. Onjuistheid van AI uitkomsten en hoe hier mee om te gaan door kennisverlies, verkeerde of gemanipuleerde gegevens. Grote zorgen over toename van vertrouwelijke gegevens.

- Sociaal
 - Ontwikkeling. AI systemen die menselijke interactie vervangen vergroten sociale isolatie en bedreigen zo mentaal en fysiek welzijn. Grootschalig toepassen AI kan ook slechts enkele superstar bedrijven bevoordelen en sociale en economische ongelijkheid veroorzaken. Sociaal engineering bedreigt AI-toepassingen en –gegevens.
 - Ethiek. Balanceren van economische drijfveren, cliënt-veiligheid. Christelijk geloof voegt daar menselijke waardigheid aan toe. Islamitisch kijkt naar verantwoordelijkheid god, cliënt, ontwikkelaar en zorgprofessional. Ook afwegen van cognitieve, sociale, en culturele aspecten; omgang met creativiteit, normen, emoties en eerlijke verdeling van kosten en opbrengsten.
 - Duurzaamheid. Grootschalige inzet van computers voor AI bedreigend voor energie, water, emissie en zeldzame materialen. Mensgerichte waarden onder druk, zoals cognitieve, sociale, culturele en interpersoonlijke ontwikkeling.
 - Winstgedrevenheid. Zorgen over eerlijke verdeling eigenaarschap, duurzaamheid, werkgelegenheid, kwaliteit van banen, kosten van arbeid. Winner takes all mentaliteit en machtsconcentratie van grote ondernemingen met lock in van cruciale sectoren zoals onderwijs, zorg en journalistiek.
- Kennis.
 - Transparantie. AI-black boxes bedreigen o.a. juistheid, betrouwbaarheid, veiligheid, transparantie, traceerbaarheid, gelijkwaardigheid, duurzaamheid en verlagen menselijke controle over besluitvorming.
 - Bias. Beschikbaarheid en kwaliteit van gegevens bepalen uitkomsten van AI-toepassingen. Hierdoor staan betrouwbaarheid en validiteit van de uitkomsten en bruikbaarheid van deze toepassingen mogelijk onder druk. Denk bijvoorbeeld aan niet of nauwelijks beschikbare data van (sub)doelgroepen waardoor AI bestaande ongelijkheden overneemt. Beïnvloeding van de mentale autonomie door de combinatie van AI en VR.
 - Afhankelijkheid. Van enkele grote leveranciers van AI-toepassingen, afnemende kennis binnen zorginstellingen en bij zorgprofessionals, versterkt door tekorten aan personeel, bedreigen de autonomie en training van de zorgprofessional. Mogelijk versterkt door gebrekkige kwaliteit van gegevens. Dit werkt mogelijke menselijke fouten in de hand.

Grafiek 5: Gezondheidszorg

Frequentie coderingscodes in de categorie risico's



4.5 RESULTATEN – RISICO'S BIJ TOEPASSEN AI IN ORGANISATIES (2/2)

Literatuur (vervolg)

- Systeem.
 - Databescherming. Data-gijzeling, -hacks, -misbruik, -vervuiling, -manipulatie, niet-gegarandeerde data-opt-out. Niet voldoen aan wettelijke strikte vertrouwelijkheidvereisten van biometrische en of gezondheidsdata. Wedloop om de grootste, rijkste, meeste data-sets te realiseren als bedrijfsstrategie vergroot de impact van aanvallen op leveranciers, AI-blackout of verlies van data-controle.
 - Feitelijke juistheid. AI is gebaseerd op statistische voorspelbaarheid en introduceert daarmee een foutkans in de uitkomsten die mogen zorgprofessionals niet klakkeloos als feiten aannemen, want deze kunnen leiden tot foutieve diagnoses of aanbevelingen. De waarheidsgetrouwheid van gegenereerde uitkomsten maakt het moeilijk om echt van vals te onderscheiden, waarbij intuïtief soms juist de valse uitkomst meer vertrouwd wordt.
 - Aanvallen. Centrale dataopslag, maar ook de toename van opgeslagen data en plaatsen waar data wordt verzameld, maakt systemen kwetsbaar voor aanvallen. AI zelf ook steeds beter te gebruiken om aanvallen uit te voeren.
 - Beschikbaarheid. Toenemende integratie van fysieke en digitale processen maakt processen kwetsbaar voor het niet beschikbaar zijn van systemen.
 - Manipulatie. Misinformatie, beïnvloeding, data-manipulatie, aanvallen, mengen van persoonlijke en generieke informatie maken allerlei bewuste en onbewuste fouten mogelijk.
 - Samenwerking. In afzondering ontwikkelde AI toepassingen leiden tot in zorginstellingen, waar systemen worden geïntegreerd en elkaar aanvullen, tot problemen. Soms versterkt door onvoldoende training van zorgprofessionals en toename in werkdruk.
- Bestuur/regels.
 - Privacy. Cliënten en zorgprofessionals bezorgd over bescherming en datalekken van gevoelige zorgdata en betrouwbaarheid van AI-uitkomsten. Versterkt door gebrek aan transparantie in gebruikte training- en gebruiksdata. Alsmat toenemende mogelijkheden om gedetailleerd fysieke, mentale en gedragsgegevens te verzamelen, combineren, controleren, misbruiken.
 - Non compliance. Negeren van voorgeschreven minimaliseren datagebruik en dataopslag. Negeren wet- en regelgeving zoals AI Act, Medical Device Regulation, Product Liability Directive, General Data Protection Regulation.
 - Aansprakelijkheid. Wettelijk onduidelijke aansprakelijkheid, ook (nog) niet vanuit jurisprudentie, van ontwikkelaars, instellingen, zorgprofessionals bij onjuiste uitkomsten AI.

4.6 RESULTATEN – MAATREGELEN AANPAKKEN VAN AI-RISICO'S (1/2)

Focusgroepen zien als belangrijke maatregelen mens blijft eindverantwoordelijk, samenwerken en leren, transparante afgeschermd systemen en regels opstellen en opvolgen.

Literatuur geeft maatregelen aan om AI met vertrouwen toe te passen op vier aspecten, sociale versterking, kennisontwikkeling, systeembeveiliging en besturen.

Focusgroepen mens eindverantwoordelijk, samenwerken, leren en kennisdelen, transparante afgeschermd systemen en regels opstellen en opvolgen.

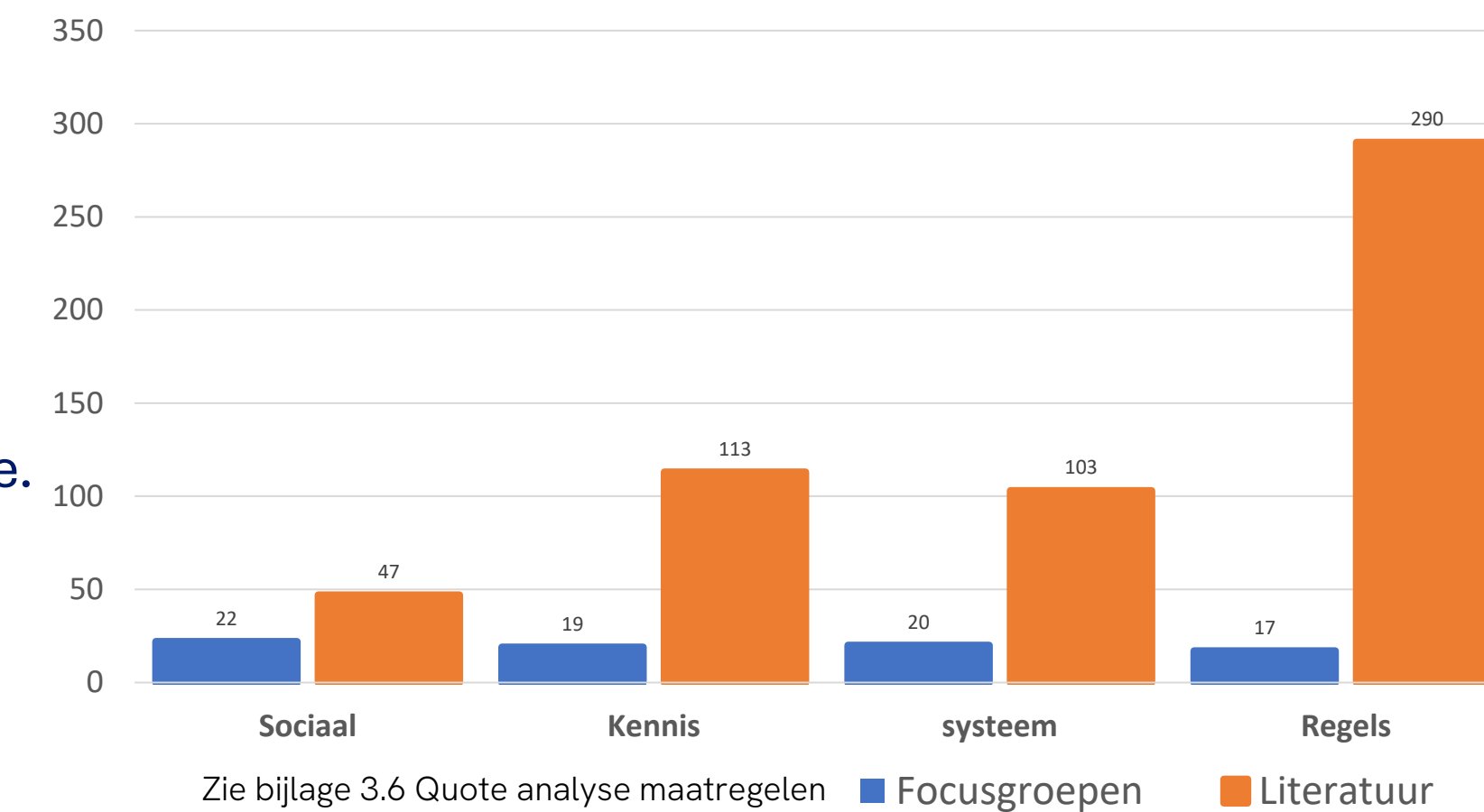
- Sociaal. Eindverantwoordelijkheid door stimuleren kritisch denken, context meewegen, zelf advies schrijven, testen, evalueren, rapporteren fouten. Samenwerken intern/extern, expertise netwerk, met onderwijs, planmatig, leverancierscontracten.
- Leren. Herhaalde opleidingen, blijven communiceren, protocollen, oefenen kritisch en reflectief denken en handelen, kleine experimenten met positief kritische koplopers, expertise concentreren.
- Systemen. Transparantie, op hoofdlijnen hoe het werkt en wat er met data gebeurt, herleidbare data en bronnen, regressie. Afschermen noodzakelijke functies en toepassingen, minimaliseer modellen en datagebruik, filter op persoons- en bedrijfsgegevens, encryptie, anonimiseren, orchestration engine, private hosting.
- Regels. Stimuleer gebruik en AI geletterdheid, DPIA en IAMA assessments, responsible AI principes, opvolging wetgeving, richtlijnen AI-tools, afspraken over gegevens en gebruik, dossiers en data op orde, toestemming van cliënten.

Literatuur

Benoemt, net als focusgroepen, maatregelen om met vertrouwen AI toe te passen. Benadrukt eindverantwoordelijkheid van mensen en het belang van vastgelegde afspraken hierover. Actief kennis verwerven door experimenteren, educatie en kwaliteitsborging. Focus op explainable AI in alle systeemaspecten. En als organisatie toepassen van AI stimuleren, en met behulp van tools en raamwerken AI-inzet monitoren en bijsturen waar nodig.

- Sociaal
 - Eindverantwoordelijkheid. Toepassen van AI vereist eindverantwoordelijkheid van mensen in alle fase van ontwikkeling, testen tot gebruik. Betrokkenheid van alle lagen in de organisatie. Gedegen training in gebruik van systemen en hoe om te gaan met AI ondersteunde besluiten. En een uitgewerkt datamanagementplan.
 - Samenwerken. Multidisciplinaire samenwerking. Binnen een zorgorganisatie, zoals zorgprofessionals, ICT, kwaliteit en veiligheidsmedewerkers. Ook met externe betrokkenen, zoals cliënten en leveranciers. Duidelijke en vastgelegde contractuele afspraken die aantoonbaar worden nageleefd.
 - Vertrouwen. Bouw stapsgewijs aan vertrouwen. Door actieve betrokkenheid van stakeholders in alle invoeringsstadia. Spreek vertrouwen uit in de ingeslagen weg en geef het goede voorbeeld. Creëer een veilige omgeving voor leren gebruiken van AI, ook voor professionals die hiervoor meer tijd nodig hebben. Overbrug de kloof van richtlijnen naar uitvoering. Hanteer en leer van best practices in alle facetten van het toepassen van AI.
- Kennis.
 - Educatie. Proactieve communicatie naar alle betrokkenen. Faciliteren van leren en valideren, zowel in het gebruik van AI-toepassingen als ook van de verdeling van verantwoordelijkheden in de organisatie. Besteed aandacht aan ontwikkelen van vaardigheden om AI-toepassingen en –uitkomsten te doorgronden, te beoordelen en gepast te handelen, o.a. fact-checken.
 - Kwaliteit. Verdedig de kwaliteit van AI toepassing door data-kwaliteit, gevalideerde input, veilige supply-chain, beveiligde systemen modellen en data, en traceerbaarheid. Maak capaciteit en middelen vrij voor kwaliteitswerkzaamheden, zoals valideren en borging. Hanteer een multidisciplinaire aanpak intern en extern. Geef praktische invulling aan bestaande wet- en regelgeving voor medische apparatuur en software.

Grafiek 6: Gezondheidszorg
Frequentie coderingscodes in de categorie maatregelen



4.6 RESULTATEN – MAATREGELEN AANPAKKEN VAN AI-RISICO'S (2/2)

Literatuur (vervolg)

- Systeem.
 - Transparantie. Explainable AI-toepassingen vraagt om begrijpelijke, uitlegbare en gevalideerde systemen. Door leveranciers, gebruikers, cliënten en toezichthouders. Dit betreft alle onderdelen van AI, zoals data, modellen, statistieken, algoritmes, ontwikkelaanpak, testen, trainingen.
 - Technologie. Hanteer bestaande en bewezen betrouwbare tools, technieken en systeemcomponenten om de juiste werking van AI toepassingen te garanderen en te monitoren.
 - Dataset. Gedurende de gehele levenscyclus van een AI –toepassing is gestructureerde en planmatig aandacht nodig voor datamanagement. Verschillende datasets voor testen en toepassen. Statistische controles. Externe en interne validatie, voorafgaande en gedurende het gebruiken van AI-toepassingen. Scannen van input op gevoelige inhoud.
 - Gedistribueerde systemen. Bij praktische toepassing van AI is beschikbaarheid cruciaal. Met gedistribueerde systemen worden single points of failure verminderd en blijft een systeem functioneren, ook als componenten uitvallen.
 - Opt-out. Cliënten hebben altijd het recht om hun gegevens uit de gebruikte zorgsystemen te onttrekken. Hanteer richtlijnen en procedures om individuele gegevens uit systemen te verwijderen, zowel uit operationele als ook trainingsdata.
- Bestuur/regels.
 - Beleid. Organisaties moeten richtlijnen opstellen om aan wetten en regels te voldoen. Voor integratie van AI in software, hardware, medische toepassingen en workflows. Hierin moet de organisatieleiding het voortouw nemen en stakeholders betrekken. Een AI governance vraagt om integratie van richtlijnen voor AI op andere organisatiebeleid.
 - Privacy. Documenteren en actief opvolgen van organisatierichtlijnen hoe om te gaan met vertrouwelijke persoonlijke gezondheidsdata van cliënten. Zeker wanneer deze informatie tussen systemen wordt geaggregeerd en met thuismonitoring systemen wordt samengevoegd. Ook het maken van onderscheid tussen AI-trainingsdata en gebruiksdata voor operationele systemen. En de mogelijkheid voor cliënten om hun gegevens uit deze systemen te onttrekken – opt out.
 - Controleren en monitoring. Gedurende alle fase van ontwikkeling en gebruik van AI-toepassingen is monitoring en bijsturen noodzakelijk. Denk aan vastlegging van onderhoud, defecten, opvolging van verbeterdoelstellingen. Menselijke eindverantwoordelijkheid voor al deze onderdelen blijft de maatstaf.
 - Management/framework/tools. Kiezen en praktisch invullen van raamwerk modellen en tools voor betrouwbaar toepassen van AI. Zoals NIST cybersecurity framework, NL AIVD veilig toepassen van AI, Googles Open Frontier Safety Framework, OpenAI's Preparedness Framework, Gladstone's proposed AI Observatory, frameworks in ontwikkeling van Dcypher/TNO/Tue.
 - Ontwikkeling stimuleren en organiseren. Om de toegevoegde waarde die AI kan leveren te benutten moeten bestuurders ontwikkeling en toepassing van AI stimuleren. Ongewenste effecten voorkomen is daar een onderdeel van, maar zeker ook gewenste effecten actief stimuleren. Om zo het risico van niet gebruiken van AI te voorkomen. En om snel en effectief in te spelen op onbekende en ongewenste effecten die zich abrupt kunnen voordoen.

5. CONCLUSIE & AANBEVELINGEN (1/3)

Onderzoeksvraag

Wat is de huidige stand van zaken met betrekking tot de cybersecurity van Generatieve AI-systemen in de gezondheidszorg, en welke specifieke behoeften en kwetsbaarheden spelen een rol bij het waarborgen van de veiligheid van deze systemen?

Door het gebruik van GenAI ontstaan risico's op vier gebieden:

- Sociaal. Systemen die menselijke interactie vervangen;
- Kennis. Afname van kennis nodig om uitkomsten van systemen kritisch te beschouwen;
- Systemen. Afhankelijkheid van systemen, zorgen om feitelijke juistheid en bescherming van persoonlijke gegevens;
- Regels. Voldoen aan wet- en regelgeving en aansprakelijkheid.

Deelvragen

1. Identificeren huidige stand van zaken.

2. Analyseren van de behoeften in de beroepspraktijk en belangrijkste kwetsbaarheden;

3. Formuleren van aanbevelingen voor verdere actie en onderzoek.

Beantwoord op de volgende pagina's

5. CONCLUSIE & AANBEVELINGEN (2/3)

Ad. 1. Identificeren huidige stand van zaken.

Onderverdeeld in nu al gebruikte en mogelijk gebruikte AI-toepassingen.

- Gebruik: daadwerkelijk ingezette AI-toepassingen.
 - Focusgroepen: benoemen zoeken, dossiervorming en diagnostiek.
 - Literatuur: vult aan met toepassingen voor begrijpen medische data, ondersteuning in clientbehandeling en zorglogistiek.
- Mogelijkheden: verwachte AI-toepassingen en –ontwikkelingen.
 - Focusgroepen: AI inzetten voor werkdrukverlaging, meer taaltoepassingen en begrijpen van data.
 - Literatuur: verwacht persoonlijker, empathischer en efficiëntere taal, besluitondersteuning en zorg.

Ad. 2. Analyseren van de behoeften in de beroepspraktijk en belangrijkste kwetsbaarheden.

Beantwoord door eerst in te zoomen op behoeften en daarna op kwetsbaarheden.

Behoeften in de beroepspraktijk onderverdeeld door eerst hindernissen in kaart te brengen en vervolgens oplossingen voor deze hindernissen te benoemen.

- Hindernissen: houden organisaties tegen in het gebruiken van AI-toepassingen.
 - Focusgroepen: zien gebrek aan vertrouwen en kennis als belangrijkste hindernissen voor inzetten AI.
 - Literatuur: benoemt vertrouwen in en AI willen toepassen in cliëntgerichte zorg als hindernissen. Te ontwikkelen kennis en ervaring in een snel ontwikkelend AI-landschap en –regelgeving.
- Oplossingen: de manieren om hindernissen aan te pakken.
 - Focusgroepen: zien als oplossingen: kennis en oplossingen delen, leren met en van elkaar, experimenteren, en aanbrengen van focus door besturen.
 - Literatuur zegt: houd mensen eindverantwoordelijk, leg focus op transparantie van gebruikte data en systemen, en voldoe aan (EU) regelgeving.

Kwetsbaarheden zijn allereerst uitgewerkt in risico's waarvoor vervolgens mogelijke maatregelen in kaart zijn gebracht.

- Risico's: bedreigingen bij toepassen AI in een organisatie.
 - Focusgroepen. Systeem/gegevensaanvallen. Afhankelijkheid, vooroordelen, kennisverlies. Niet opvolgen regels. Verlies gevoelige persoonlijke info. Niet toepassen van AI.
 - Literatuur concretiseert de AI risico's sociaal, ethische, duurzaamheid, betrouwbaarheid, transparantie, afhankelijkheid, databescherming, aanvalsoppervlak
- Maatregelen: aanpakken van AI-risico's.
 - Focusgroepen belangrijke maatregelen mens blijft eindverantwoordelijk, samenwerken en leren, transparante afgeschermd systemen en regels opstellen en opvolgen.
 - Literatuur geeft maatregelen aan om AI met vertrouwen toe te passen op vier aspecten, sociale versterking, kennisontwikkeling, systeembeveiliging en besturen.

5. CONCLUSIE & AANBEVELINGEN (3/3)

Ad. 3. Formuleren van aanbevelingen voor verdere actie en onderzoek.

De belangrijkste overeenkomsten en verschillen tussen de focusgroepen en de literatuur.

- De focusgroepen en literatuur kennen en beschrijven beide dezelfde vier risicogebieden; sociaal, kennis, systemen en regelgeving.
- De literatuur is gedetailleerder in de uitwerkingen.
- Bijvoorbeeld: focusgroepen denken aan betere taaltoepassingen. De literatuur maakt dit concreet met voorbeelden van efficiënter en empathiser taalgebruik in AI-toepassingen.

Advies aan zorgprofessionals om Artificial Intelligence kansen te benutten in de gezondheidszorg

- Werk aan vertrouwen;
- Dat begint met en draait om mensen, die AI willen gebruiken;
- Daarvoor is kennisontwikkeling door experimenteren noodzakelijk;
- Om te leren en te weten wat veilig werkt; en
- Als organisatie, medewerker en cliënt te kunnen wat daarvoor nodig is.

De behoefte en kansen voor vervolgonderzoek uit de beroepspraktijk en de wetenschap

Vragen uit de beroepspraktijk zijn vooral hoe-vragen, zoals:

- Op welke manier vergroot AI de autonomie van cliënten in de zorg?
- Welke mogelijkheden biedt AI voor leren in de zorg?
- Wat bepaalt het succes van een AI-experiment in de zorg?
- Wat zijn randvoorwaarden voor opschalen van succesvolle AI-experimenten?
- Aan welke transparantievereisten moet AI voldoen om weerbaarheid van mens en organisatie te vergroten?
- Welke (anti)neutralisatietechnieken werken om gebruik van AI te stimuleren?
- Wat zijn de key drivers voor vertrouwen in AI en hoe deze te versterken?

Vragen uit de wetenschap richten zich op meer inzicht, zoals:

- Welke factoren bepalen menselijke eindverantwoordelijkheid bij toepassing van AI en hoe deze inzichtelijk te maken?
- Op welke manier is AI risico te meten?

Hoe nu verder?

- Stap 1: Met de stakeholders prioriteren van de onderzoeksvragen; en
- Stap 2: Samenwerkingsverbanden starten voor het indienen van subsidieaanvragen om deze vragen te beantwoorden.

BIJLAGEN

1. Definities

2. Onderzoeksaanpak

2.1 Focusgroepen

2.2 Literatuurstudie

3. Observaties uit focusgroepen

3.1 Gebruik

3.2 Mogelijk gebruik

3.3 Hindernissen

3.4 Oplossingen

3.5 Risico's

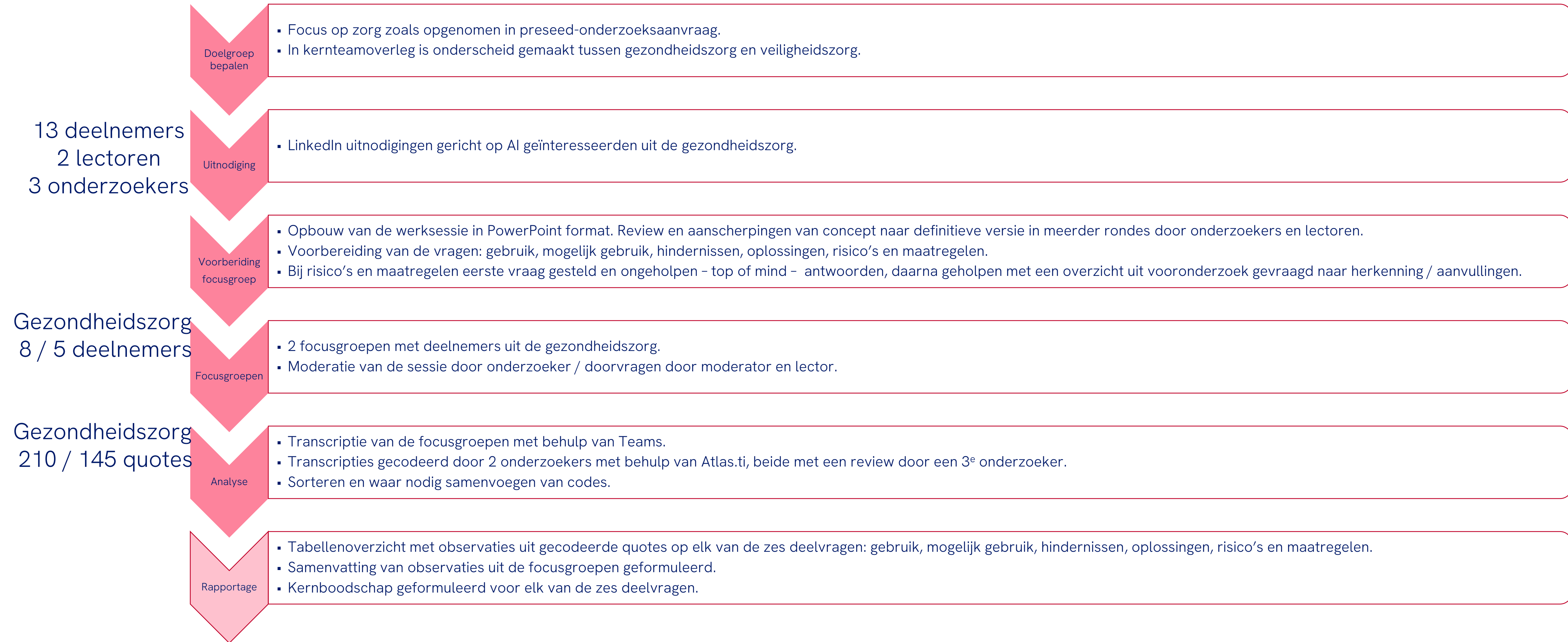
3.6 Maatregelen

4. Bronnenlijst

1. DEFINITIES

Term	Definitie	Bron
1. Generatieve AI	Generatieve AI verwijst naar een type kunstmatige intelligentie dat in staat is nieuwe inhoud te creëren. Het is daarmee een onderdeel van een nieuwe generatie van AI, in tegenstelling tot meer klassieke AI welke gericht is op het leren herkennen en voorspellen van patronen in cijfermatige data (ook wel bekend als machine learning), het begrijpen van tekst (bekend als natural language processing) en beeld (ook wel computer vision genoemd). Het genereren van nieuwe inhoud kan variëren van tekst en beelden tot muziek en video's. De kern van GAI-modellen is het vermogen patronen, stijlen of regels uit bestaande datasets te leren en deze kennis te gebruiken om nieuwe, unieke en realistische output te genereren die past binnen de geleerde context. In de laatste jaren is GAI in drie vormen sterk ontwikkeld: het genereren van beelden, genereren van teksten en genereren van synthetische data	Stokkum, R. v., Bouwman, J., & Kamstra, R. (2024). Generatieve AI in de Nederlandse zorg. TNO, TNO Healthy Living & Work. Delft: TNO Public. Opgeroepen op januari 29, 2025
2. Gezondheidszorg	Het geheel van zorgverleners (en ondersteunend personeel), instellingen, middelen en activiteiten dat direct gericht is op instandhouding en verbetering van de gezondheidstoestand en/of de mogelijkheid om zelf regie te voeren, en op het reduceren, opheffen, compenseren en voorkomen van tekorten daarin	RIVM. (2025). VZinfo Gezondheidszorg. Opgeroepen op januari 28, 2025, van VZinfo: https://www.vzinfo.nl/bronnen-methoden-en-achtergronden/gezondheidszorg
3. Gebruik	Daadwerkelijk ingezette AI-toepassingen.	
4. Mogelijkheden	Verwachte AI-toepassingen en -ontwikkelingen.	
5. Hindernissen	Houden organisaties tegen in het gebruiken van AI-toepassingen.	
6. Oplossingen	Manieren om hindernissen aan te pakken.	
7. Risico's	Bedreigingen bij toepassen AI in een organisatie.	
8. Maatregelen	Aanpakken van AI-risico's.	

2.1 AANPAK FOCUSGROEPEN



2.2 AANPAK LITERATUURSTUDIE



3.1 QUOTE ANALYSE FOCUSGROEPEN EN LITERATUUR

Tabel 1. Frequentie coderingscodes in de categorie gebruik: daadwerkelijk ingezette AI-toepassingen.

Gezondheidszorg		Focusgroepen				Literatuur			
		Focusgroep 1 - Cybersecurity GenAI in de Gezondheidszorg Gr=211	Focusgroep 2 - Cybersecurity GenAI in de Gezondheidszorg Gr=150	subtotal		Literatuur beide Gr=603; GS=10	Literatuur zorg Gr=718; GS=16	subtotal	Totals
Zoeken	● Gebruik: Zoeken Gr=7	2	4	6	● Gebruik: Zoeken Gr=8	1	0	1	7
	● Gebruik: Verslaglegging Gr=13	7	1	18	● Gebruik: Verslaglegging Gr=14	0	4	13	31
Taal	● Gebruik: Schrijven Gr=8	4	3		● Gebruik: Schrijven Gr=8	0	1		
	● Gebruik: Antwoorden Gr=6				● Gebruik: Antwoorden Gr=6	1	5		
	● Gebruik: Vertalen Gr=5	3	0		● Gebruik: Vertalen Gr=6	1	1		
Begrijpen	● Gebruik: Diagnostiek Gr=27	3	1	11	● Gebruik: Diagnostiek Gr=38	1	29	64	75
	● Gebruik: Leren Gr=17	4	0		● Gebruik: Leren Gr=30	8	14		
	● Gebruik: Dossierkennis-EPD Gr=5	1	2		● Gebruik: Dossierkennis-EPD Gr=17	0	12		
Ondersteuning	● Gebruik: Data analyse Gr=12	0	0	4	● Gebruik: Data analyse Gr=22	2	11	60	64
	● Gebruik: Besluitondersteuning Gr=16	2	0		● Gebruik: Besluitondersteuning Gr=20	1	15		
					● Gebruik: Toezichthouden Gr=7	2	1		
					● Gebruik: Behandeling Gr=17	3	12		
	● Gebruik: Plannen Gr=8	1	1		● Gebruik: Plannen Gr=8	0	6		
					● Gebruik: Voorspellen Gr=17	3	4		
Totals		27	12	39	Totals	23	115	138	50

Noot: Overgenomen uit Cybersecurity GenAI in de (veiligheids)zorg, Atlas.ti, 4-4-2025

3.2 QUOTE ANALYSE FOCUSGROEPEN EN LITERATUUR

Tabel 2. Frequentie coderingscodes in de categorie mogelijk gebruik: verwachte AI-toepassingen en –ontwikkelingen.

		Focusgroepen			Literatuur				
		Focusgroep 1 - Cybersecurity GenAI in de Gezondheidszorg Gr=212	Focusgroep 2 - Cybersecurity GenAI in de Gezondheidszorg Gr=150			Literatuur beide Gr=603; GS=10	Literatuur zorg Gr=720; GS=16		
				subtotal				subtotal	Totals
Zoeken	• Mogelijk gebruik: Zoeken Gr=5	1	3	4	• Mogelijk gebruik: Zoeken Gr=5			0	4
Taal	• Mogelijk gebruik: Verslaglegging Gr=9	1	2	7	• Mogelijk gebruik: Verslaglegging Gr=9	2	1	6	13
	• Mogelijk gebruik: Schrijven Gr=1	1			• Mogelijk gebruik: Schrijven Gr=1	1	1		
	• Mogelijk gebruik: Vertalen Gr=4	1	2		• Mogelijk gebruik: Vertalen Gr=4		1		
Begrijpen	• Mogelijk gebruik: Diagnostiek Gr=25	1	1	13	• Mogelijk gebruik: Diagnostiek Gr=25	1	23	36	49
	• Mogelijk gebruik: Leren Gr=14	1	2		• Mogelijk gebruik: Leren Gr=14	3	5		
	• Mogelijk gebruik: Dossierkennis-EPD Gr=12	1	7		• Mogelijk gebruik: Dossierkennis-EPD Gr=12	1	3		
Ondersteuning	• Mogelijk gebruik: Data analyse Gr=3	1	2	13	• Mogelijk gebruik: Data analyse Gr=3	4	7	100	113
					• Mogelijk gebruik: Extende reality Gr=4				
					• Mogelijk gebruik: Kwaliteit Gr=10				
					• Mogelijk gebruik: Ontwikkeling Gr=6				
					• Mogelijk gebruik: Veiligheid Gr=2				
	• Mogelijk gebruik: Besluitondersteuning Gr=2				• Mogelijk gebruik: Besluitondersteuning Gr=2				
					• Mogelijk gebruik: Behandeling Gr=16				
	• Mogelijk gebruik: Plannen Gr=5				• Mogelijk gebruik: Plannen Gr=5				
	• Mogelijk gebruik: Werkdruk Gr=14				• Mogelijk gebruik: Werkdruk Gr=14				
					• Mogelijk gebruik: Administratie Gr=16				
					• Mogelijk gebruik: Kosten Gr=4				
	• Mogelijk gebruik: Client Gr=25	0	1	• Mogelijk gebruik: Client Gr=25	3	26			
Totals		13	24	37		30	112	142	179

Noot: Overgenomen uit Cybersecurity GenAI in de (veiligheids)zorg, Atlas.ti, 16-5-2025

Legenda frequentieverdeling:

minst

minder

meer

meest

3.3 QUOTE ANALYSE FOCUSGROEPEN EN LITERATUUR

Tabel 3. Frequentie coderingscodes in de categorie hindernissen: houden organisaties tegen in het gebruiken van AI-toepassingen.

Gezondheidszorg		Focusgroepen		Literatuur		Totals
		Focusgroep 1 - Cybersecurity GenAI in de Gezondheidszorg Gr=214	Focusgroep 2 - Cybersecurity GenAI in de Gezondheidszorg Gr=150	Literatuur beide Gr=603; GS=10	Literatuur zorg Gr=720; GS=16	
		subtotal		subtotal		
Sociaal	• Hindernis: Vertrouwen Gr=45	12	6	7	30	69
	• Hindernis: Behoeft Gr=9	2	7			
	• Hindernis: Niet willen Gr=6	1	0	1	3	
Kennis	• Hindernis: Kennis Gr=37	15	4	15	18	94
	• Hindernis: Vaardigheden Gr=12	5	0	4	19	
	• Hindernis: Gebrek aan ervaring Gr=5			5		
	• Hindernis: Niet kunnen Gr=10			1	8	
Systeem	• Hindernis: Omvang Gr=14	3	2	1	8	43
	• Hindernis: Snelheid ontwikkelingen Gr=16		2	4	4	
	• Hindernis: Beschikbaarheid Gr=1					
	• Hindernis: Dataset Gr=9				9	
	• Hindernis: Techniek Gr=13			1	9	
Besturen/regels	• Hindernis: Regelgeving Gr=65	9	2	43	20	152
	• Hindernis: Andere prioriteiten Gr=5	1	2	2	1	
	• Hindernis: Geld Gr=12		2	7	11	
	• Hindernis: Tijd Gr=7	4	2	1	2	
	• Hindernis: Gebrek aan regels Gr=5			4	1	
	• Hindernis: Ondernemerschap Gr=8			3	1	
	• Hindernis: Opschalen gebruik Gr=1				1	
	• Hindernis: Organisatie inrichting Gr=9			1	8	
	• Hindernis: Privacy Gr=17			4	13	
	• Hindernis: Veiligheid Gr=7			4	3	
	Totals	52	29	108	169	358

Noot: Overgenomen uit Cybersecurity GenAI in de (veiligheids)zorg, Atlas.ti, 16-4-2025

Legenda frequentieverdeling:

minst

minder

meer

meest

3.4 QUOTE ANALYSE FOCUSGROEPEN EN LITERATUUR

Tabel 4. Frequentie coderingscodes in de categorie oplossingen: de manieren om hindernissen aan te pakken.

		Focusgroepen				Literatuur			
		Focusgroep 1 - Cybersecurity GenAI in de Gezondheidszorg Gr=214	Focusgroep 2 - Cybersecurity GenAI in de Gezondheidszorg Gr=150	subtotal		Literatuur beide Gr=603; GS=10	Literatuur zorg Gr=720; GS=16	subtotal	Totals
Sociaal				10	• Oplossing: Mens blijft eindverantwoordelijk Gr=21	13	8	61	71
	• Oplossing: Samenwerken Gr=42	6	2		• Oplossing: Samenwerken Gr=42	21	19		
	• Oplossing: Interesse Gr=2	1	1						
Kennis	• Oplossing: Kennis Gr=32	12	1	27	• Oplossing: Kennis Gr=32	17	4	39	66
	• Oplossing: Leren gebruiken Gr=17	5			• Oplossing: Leren gebruiken Gr=17	7	5		
	• Oplossing: Klein houden Gr=6	4	2						
	• Oplossing: Experimenteren Gr=7	3			• Oplossing: Experimenteren Gr=7	2	4		
Systeem	• Oplossing: Techniek Gr=20	4	1	9	• Oplossing: Techniek Gr=20	12	10	66	75
	• Oplossing: Transparantie Gr=10	4			• Oplossing: Transparantie Gr=10	3	27		
					• Oplossing: Standaardisatie Gr=14	12	2		
Besturen/regels		3		3	• Oplossing: Regelgeving Gr=63	41	19	102	105
	• Oplossing: Beleid Gr=17				• Oplossing: Beleid Gr=17	12	6		
					• Oplossing: Monitoren Gr=5	1	4		
					• Oplossing: Opletten Gr=15	6	9		
					• Oplossing: Afschermen Gr=4	1	2		
					• Oplossing: Geld Gr=1		1		
Totals		42	7	49	Totals	148	120	268	317

Noot: Overgenomen uit Cybersecurity GenAI in de (veiligheids)zorg, Atlas.ti, 8-5-2025

Legenda frequentieverdeling:

minst

minder

meer

meest

3.5 QUOTE ANALYSE FOCUSGROEPEN EN LITERATUUR

Tabel 5. Frequentie coderingscodes in de categorie risico's: bedreigingen bij toepassen AI in een organisatie.

		Gezondheidszorg			Literatuur					
		Focusgroep 1 - Cybersecurity GenAI in de Gezondheidszorg Gr=214	Focusgroep 2 - Cybersecurity GenAI in de Gezondheidszorg Gr=150			Literatuur beide Gr=603; GS=10	Literatuur zorg Gr=720; GS=16			
		subtotal				subtotal			Totals	
Sociaal	• Risico: Sociale ontwikkeling Gr=11		1	5	• Risico: Sociale ontwikkeling Gr=11	10	2	48	53	
	• Risico: Ethiek Gr=6	1	1		• Risico: Ethiek Gr=6	5	6			
	• Risico: Werkdruk Gr=3	1	1		• Risico: Werkdruk Gr=3	1				
					• Risico: Duurzaamheid Gr=8	5	2			
					• Risico: Winstgedrevenheid Gr=20	10	7			
Kenniis	• Risico: Verlies van transparantie Gr=22	3	7	22	• Risico: Verlies van transparantie Gr=22	5	8	100	122	
	• Risico: Bias Gr=33	2	4		• Risico: Bias Gr=33	10	28			
	• Risico: Afhankelijkheid Gr=31	3			• Risico: Afhankelijkheid Gr=31	34	4			
	• Risico: Verlies van kennis Gr=4	3	0		• Risico: Verlies van kennis Gr=4	3	3			
					• Risico: Menselijke fout Gr=5	2	3			
Systeem	• Risico: Databescherming Gr=57	9	12	42	• Risico: Databescherming Gr=57	27	23	274	316	
	• Risico: Feitelijke juistheid Gr=46	10	3		• Risico: Feitelijke juistheid Gr=46	13	23			
	• Risico: Databetrouwbaarheid Gr=32	0	4		• Risico: Databetrouwbaarheid Gr=32	20	23			
	• Risico: Aanvallen Gr=19	3			• Risico: Aanvallen Gr=19	16	2			
	• Risico: Beschikbaarheid Gr=8	1			• Risico: Beschikbaarheid Gr=8	12	1			
					• Risico: Kwaliteit Gr=8	2	4			
					• Risico: Manipulatie Gr=55	29	8			
					• Risico: Misbruik Gr=16	9	5			
					• Risico: Samenwerking Gr=2		2			
					• Risico: Veiligheid Gr=62	30	25			
Regels	• Risico: Privacy Gr=53	4	4	15	• Risico: Privacy Gr=53	29	39	91	106	
	• Risico: Non compliance Gr=29	3	4		• Risico: Non compliance Gr=29	5	15			
					• Risico: Aansprakelijkheid Gr=6		3			
Totals		43	41	84	Totals	277	236	513	597	

Noot: Overgenomen uit Cybersecurity GenAI in de (veiligheids)zorg, Atlas.ti, 29-4-2025

Legenda frequentieverdeling:

minst

minder

meer

meest

3.6 QUOTE ANALYSE FOCUSGROEPEN EN LITERATUUR

Tabel 6. Frequentie coderingscodes in de categorie maatregelen: aanpakken van AI-risico’s.

		Focusgroepen				Literatuur			
		Focusgroep 1 - Cybersecurity GenAI in de Gezondheidszorg Gr=213	Focusgroep 2 - Cybersecurity GenAI in de Gezondheidszorg Gr=150	subtotal		Literatuur beide Gr=603; GS=10	Literatuur zorg Gr=720; GS=16	subtotal	Totals
Sociaal	• Maatregel: Mens blijft eindverantwoordelijk Gr=25	9	6	22	• Maatregel: Mens blijft eindverantwoordelijk Gr=34	1	7	47	69
	• Maatregel: Samenwerken Gr=17	4	3		• Maatregel: Samenwerken Gr=29	6	7		
					• Maatregel: Samenwerking Gr=15	7	4		
					• Maatregel: Vertrouwen Gr=15	1	14		
Leren	• Maatregel: Educatie Gr=31	10	4	19	• Maatregel: Educatie Gr=59	10	25	113	132
	• Maatregel: Experimenteren Gr=5	1	2		• Maatregel: Experimenteren Gr=5				
	• Maatregel: Kennis Gr=4	1	1		• Maatregel: Kennis Gr=22	2	14		
					• Maatregel: Kwaliteit Gr=66	8	54		
systeem	• Maatregel: Transparantie Gr=31	2	7	20	• Maatregel: Transparantie Gr=67	11	42	103	123
	• Maatregel: Afschermen Gr=9	4	2		• Maatregel: Afschermen Gr=10	1			
	• Maatregel: Technologie Gr=29	3	2		• Maatregel: Technologie Gr=34	23	6		
					• Maatregel: Dataset Gr=11	1	10		
					• Maatregel: Gedistribueerde systemen Gr=2	2			
					• Maatregel: Klein houden Gr=2	1	1		
					• Maatregel: Opt out Gr=4		4		
					• Maatregel: Tegenaanval Gr=1	1			
Regels	• Maatregel: Beleid Gr=18	5	8	17	• Maatregel: Beleid Gr=61	7	36	290	307
	• Maatregel: Regels Gr=35	2	1		• Maatregel: Regels Gr=138	88	40		
	• Maatregel: Privacy Gr=11	1			• Maatregel: Privacy Gr=25	8	14		
					• Maatregel: Controleren Gr=7	1	5		
					• Maatregel: Managementsysteem/framework Gr=14	12	2		
					• Maatregel: Monitoring Gr=16	6	10		
					• Maatregel: Ontwikkeling stimuleren Gr=16	8	6		
					• Maatregel: Organisatieinrichting Gr=32	1	28		
					• Maatregel: Tool Gr=6	5	1		
					• Maatregel: Toezicht Gr=12	11	1		
	Totals	42	36	78	Totals	222	331	553	631

Noot: Overgenomen uit Cybersecurity GenAI in de (veiligheids)zorg, Atlas.ti, 29-4-2025

Legenda frequentieverdeling:

minst

minder

meer

meest

BRONNENLIJST

Zorg algemeen

- Boheemen, P. van, Munnichs, G., Kool, L., Diercks, G., Hamer, J., & Vos, A. (2020). Cyberweerbaar met nieuwe technologie: kans en noodzaak van digitale innovatie.
- Brink, N.W.T., Gijsen, B.B.M., Kamphuis, Y.N., Liebergen, N.M., van Opheikens, D.D., Stijn, J.J. van, Wijnja, S. Verkenning van het raakvlak van cybersecurity en AI. (2024)
- "Claudio Novelli, Federico Casolari, Philipp Hacker, Giorgio Spedicato, Luciano Floridi,Generative AI in EU law: Liability, privacy, intellectual property, and cybersecurity,Computer Law & Security Review,Volume 55,2024,106066,ISSN 0267-3649,https://doi.org/10.1016/j.clsr.2024.106066."
- Hamer, J., Kool, L., Hijstek, B., van Eeden, Q., & Das, D. (2023). Generatieve AI: Rathenau Scan.
- Kleij, R. van der, & Hof, T. (2024, June). Why Do Organizations Fail to Practice Cyber Resilience?. In International Conference on Human-Computer Interaction (pp. 126-137). Cham: Springer Nature Switzerland.
- Ministerie van Binnenlandse Zakenen Koninkrijksrelaties. Overheidsbrede visie Generatieve AI. (2024)"
- Nieuwenhuizen, W., Hijstek, B., Roolvink, S., & van Huijstee, M. (2023). Immersieve technologieën: Rathenau Scan.
- Piersma, N. (2024). hyperscenario's van het gebruik van artificiele intelligentie: een essay.
- Stokkum, R. v., Bouwman, J., & Kamstra, R. (2024). Generatieve AI in de Nederlandse zorg. TNO, TNO Healthy Living & Work. Delft: TNO Public. Opgeroepen op januari 29, 2025
- Sutton, A., & Tompson, L. (2024). Towards a Cybersecurity Culture-Behaviour Framework: A Rapid Evidence Review. Computers & Security, 104110.

Gezondheidszorg

- Beekum, K. van, Baartmans, M., & Wagner, C. (2024). Patiëntveiligheid en medische technologie in de ziekenhuiszorg.
- Capraro, V., Lentsch, A., Acemoglu, D., Akgun, S., Akhmedova, A., Bilancini, E., ... & Viale, R. (2024). The impact of generative artificial intelligence on socioeconomic inequalities and policy making. PNAS nexus, 3(6).
- Dieters, A., van Rijn, P., & van der Meer, J. (2024). Doorbraken in maxillofaciale radiologie. Tandartspraktijk, 45(7), 12-16.
- Jetten-van den Hoek, S. (2024). Generatieve AI in de Nederlandse zorg.
- Kers, J., Bülow, R. D., Klinkhammer, B. M., Breimer, G. E., Fontana, F., Abiola, A. A., Hofstraat, R., Corthals, G. L., Peters-Sengers, H., Djudjaj, S., von Stillfried, S., Hölscher, D. L., Pieters, T. T., van Zuilen, A. D., Bemelman, F. J., Nurmohamed, A. S., Naesens, M., Roelofs, J. J. T. H., Florquin, S., Floege, J., ... Boor, P. (2022). Deep learning-based classification of kidney transplant pathology: a retrospective, multicentre, proof-of-concept study. The Lancet. Digital health, 4(1), e18-e26. https://doi.org/10.1016/S2589-7500(21)00211-9
- Meer, B. C. van der (2024). ChatGPT in de ggz: kansen en overwegingen. Tijdschrift voor Psychiatrie, (2024/3), 161-164.
- Nab, L. D. (2024). Attitudes van zorgmedewerkers in de epilepsiezorg ten opzichte van de ontwikkeling en implementatie van generatieve artificiële intelligentie.
- Nazish Khalid, Adnan Qayyum, Muhammad Bilal, Ala Al-Fuqaha, Junaid Qadir,Privacy-preserving artificial intelligence in healthcare: Techniques and applications,Computers in Biology and Medicine,Volume 158,2023,106848,ISSN 0010-4825,https://doi.org/10.1016/j.compbiomed.2023.106848."
- Poels, R. (2024). Samenhangende digitale oplossingen. Digitale transformatie bij een zorginstelling: Aan de slag met kennis en leiderschap, 49-80.
- Sangers TE, Wakkee M, Kramer-Noels EC, Nijsten T, Lugtenberg M. Views on mobile health apps for skin cancer screening in the general population: an in-depth qualitative exploration of perceived barriers and facilitators. Br J Dermatol. 2021 Nov;185(5):961-969. doi: 10.1111/bjd.20441. Epub 2021 Jul 5. PMID: 33959945; PMCID: PMC9291092.
- Smeden, M. van, Moons, C., Hooft, L., Kant, I., Os, H. van, Chavannes, N. Leidraad voor kwalitatieve diagnostische en prognostische toepassingen van AI in de zorg 1.1. (2023)
- Solaiman, Barry; Cohen, I Glenn; "A Framework for Health, AI and the Law","Research Handbook on Health, AI and the Law",,,1-19,2024,Edward Elgar Publishing
- Sparnaaij, K., Van Eekeren, P., & Vasseur, J. (2023). AI Monitor Ziekenhuizen. Editie2023. M&I Partners. https://mxi.nl/uploads/files/publication/ai-monitorziekenhuizen-2023.pdf"
- Vuuren, C. L. van, van Mens, K., de Beurs, D., Lokkerbol, J., van der Wal, M. F., Cuijpers, P., & Chinapaw, M. J. M. (2021). Comparing machine learning to a rule-based approach for predicting suicidal behavior among adolescents: Results from a longitudinal population-based survey. Journal of affective disorders, 295, 1415-1420. https://doi.org/10.1016/j.jad.2021.09.018
- Werkgroep AI Z-Cert. (2024). AI-in de zorg 1.0. https://z-cert.nl/assets/downloads/AI-in-de-zorg-TLP-CLEAR.pdf
- Zhang, P., & Kamel Boulos, M. N. (2023). Generative AI in medicine and healthcare: promises, opportunities and challenges. Future Internet, 15(9), 286.