



VeeHa B.V.

Cybersecurity and privacy

VeeHa B.V.

Aalsburg 1508
6602 TS Wijchen

Telefoon +31 6 125 749 24

Website www.veeha.nl

Email info@veeha.nl

WHITEPAPER

06 - 2026

Jan Willem van Hoorn

AI veilig en verantwoord gebruiken

Een praktisch kader vanuit ISO 27001 2022 en NEN 7510 2024

Voor bestuurders, management, informatiebeveiliging, privacy, HR, ICT en medewerkers in zorg- en informatie-intensieve organisaties.

AI kan medewerkers aantoonbaar ondersteunen bij analyse, schrijven, samenvatten, programmeren en besluitvoorbereiding. Dezelfde technologie kan vertrouwelijke informatie laten weglekken, overtuigende onjuistheden produceren en ongewenste afhankelijkheden creëren. Organisaties hebben daarom geen verbod als standaardantwoord nodig, maar bestuurbare kaders: toegestane toepassingen, geschikte tools, gegevensgrenzen, menselijke controle en aantoonbaar toezicht.

Deze whitepaper vertaalt de principes van een information security management system naar AI-gebruik. De kern is eenvoudig: behandel AI niet als een losse applicatie, maar als een nieuwe verwerkingsketen met leveranciers, gegevensstromen, modellen, gebruikers, outputs en gevolgen.

VeeHa B.V.

Praktische ondersteuning bij ISO 27001, NEN 7510, privacy, risicoanalyse, beleid, implementatie en interne audits.

Managementsamenvatting

Organisaties lopen vooral risico wanneer medewerkers AI al gebruiken terwijl beleid, contracten en technische maatregelen achterblijven. Een medewerker die een cliëntverslag, incidentrapport, offerte, broncode of intern besluit in een publieke AI-dienst plakt, start feitelijk een externe gegevensverwerking. De organisatie kan daarbij controle verliezen over opslag, hergebruik, bewaartermijnen, internationale doorgifte en verwijdering.

ISO 27001 2022 en NEN 7510 2024 schrijven niet voor dat AI verboden moet worden. Zij vragen wel dat de organisatie haar context, risico's, verplichtingen, leveranciers, informatie, bevoegdheden, incidenten en verbetermaatregelen aantoonbaar beheerst. Voor zorgorganisaties is de lat hoger door het beroepsgeheim, gezondheidsgegevens, patiëntveiligheid en de afhankelijkheid van correcte en beschikbare informatie.

De bestuurlijke minimumset bestaat uit acht onderdelen:

1. Een actueel register van AI-systemen en AI-toepassingen, inclusief informeel of experimenteel gebruik.
2. Een gegevensclassificatie die bepaalt welke informatie wel, beperkt of niet in welke AI-omgeving mag worden verwerkt.
3. Een goedkeuringsproces per toepassing, gebaseerd op impact, gegevens, doelgroep en mate van menselijke controle.
4. Contractuele en technische beoordeling van leveranciers, inclusief training op klantdata, logging, retentie, locatie, subverwerkers en beveiliging.
5. Heldere werkinstructies voor medewerkers, ondersteund door praktijkgerichte AI-geletterdheid en terugkerende oefeningen.
6. Menselijke verificatie van belangrijke output, met zwaardere controles bij zorg, HR, juridisch, financieel en beveiligingskritisch gebruik.
7. Logging, incidentmelding, monitoring, periodieke evaluatie en aantoonbare verbetering binnen het bestaande ISMS.
8. Bestuurlijke keuzes over proportionaliteit, fundamentele rechten, transparantie, uitsluiting, autonomie en maatschappelijke impact.

De doorslaggevende vraag is niet of een tool populair of technisch indrukwekkend is. De vraag is of de organisatie kan uitleggen welke informatie de tool ontvangt, wat ermee gebeurt, welke uitkomst wordt gebruikt, wie verantwoordelijk blijft en hoe schade wordt voorkomen of hersteld.



VeeHa B.V.

Cybersecurity and privacy

Leeswijzer

Deel	Onderwerp	Resultaat voor de lezer
1	Werking en gegevensstromen	Begrijpen wat bij AI-gebruik wordt verwerkt en achterblijft.
2	Risico's	Een bruikbaar risicobeeld voor organisatie en zorgcontext.
3	Normenkader	Vertaling naar ISO 27001 2022 en NEN 7510 2024.
4	Medewerkers	Concrete regels voor bewust en veilig gebruik.
5	Beleid en governance	Een beheersmodel, rollen en minimumbeleid.
6	Samenleving	Afweging van rechten, vertrouwen en publieke effecten.
7	Implementatie en audit	Roadmap, indicatoren en toetsvragen.

1 Wat AI in een organisatie werkelijk doet

1 1 AI is onderdeel van een verwerkingsketen

Een generatieve AI-tool ontvangt instructies, context en vaak bijlagen. De dienst zet die informatie om in tokens of andere representaties, verwerkt deze met een model en retourneert output. Rond die kern bestaan aanvullende functies zoals gespreksgeschiedenis, zoekkoppelingen, plug-ins, agentfuncties, geheugen, logging, moderatie en feedback. Iedere functie kan een extra gegevensstroom, leverancier of toegangspunt introduceren.

De organisatie moet daarom verder kijken dan het model. Ook browserextensies, mobiele apps, ingebouwde copilots, API-koppelingen, kennisbanken, transcriptiediensten en door leveranciers toegevoegde AI-functies vallen binnen de beoordeling.

1 2 Wat laat een medewerker in een AI-tool achter

Informatie	Voorbeelden	Mogelijk gevolg
Invoer en prompts	Vragen, instructies, namen, casuïstiek, conceptteksten	Opslag in geschiedenis of logbestanden; mogelijke inzage door leverancier of beheerder.
Bijlagen	Dossiers, spreadsheets, contracten, screenshots, broncode	Volledige documenten kunnen buiten de beheerde omgeving terechtkomen.
Metadata	Account, IP-adres, tijd, apparaat, locatie, gebruikspatroon	Profilering, beveiligingsanalyse, facturering of toezicht op gebruik.
Afgeleide informatie	Samenvattingen, classificaties, embeddings, labels	Nieuwe persoonsgegevens of bedrijfsgevoelige inzichten kunnen ontstaan.
Feedback	Duim omhoog of omlaag, correcties, vervolgvragen	Feedback kan inhoud en context bevatten en voor kwaliteitsverbetering worden gebruikt.
Output	Antwoorden, code, adviezen, afbeeldingen	Output kan worden bewaard, gedeeld, hergebruikt of onderdeel worden van een dossier.
Koppelingen	E-mail, agenda, SharePoint, EPD, CRM, tickets	Een agent kan meer gegevens bereiken of acties uitvoeren dan de gebruiker beseft.

Een betaalde of zakelijke licentie is geen automatische garantie. De organisatie moet de feitelijke instellingen en contractvoorwaarden controleren. Relevante vragen zijn of klantdata voor training worden gebruikt, hoe lang logs blijven bestaan, in welke landen verwerking plaatsvindt, welke subverwerkers betrokken zijn, wie supporttoegang heeft en of gegevens aantoonbaar kunnen worden verwijderd.

1 3 Vier rollen die niet door elkaar mogen lopen

- Ontwikkelaar of aanbieder: bouwt of levert het model of AI-systeem.
- Gebruiksverantwoordelijke organisatie: kiest het doel, de werkwijze en de voorwaarden waaronder medewerkers of cliënten het systeem gebruiken.
- Gebruiker: voert informatie in, beoordeelt output en neemt of ondersteunt een handeling.
- Betrokkene of geraakte persoon: ondervindt de gevolgen, zonder zelf altijd gebruiker te zijn.

Een organisatie kan tegelijk meerdere rollen hebben. Een zorginstelling die een standaardmodel koppelt aan eigen protocollen en patiëntdata is meer dan een gewone eindgebruiker. De eigen configuratie, kennisbron, instructies en besluitvorming bepalen mede het risicoprofiel.

2 Het risicolandschap

AI-risico's zijn niet beperkt tot vertrouwelijkheid. Ook integriteit, beschikbaarheid, privacy, patiëntveiligheid, gelijke behandeling, uitlegbaarheid, intellectueel eigendom, continuïteit en bestuurlijke verantwoordelijkheid staan onder druk. De risicoanalyse moet daarom zowel informatiebeveiliging als de gevolgen voor mensen en processen omvatten.

2.1 Overzicht van belangrijke risico's

Risico	Voorbeeld	Mogelijke beheersing
Ongeautoriseerde openbaarmaking	Een medewerker plakt een cliëntverslag in een publieke chatbot.	Datalabels, goedgekeurde enterprise-omgeving, DLP, blokkades en training.
Onjuiste output	AI verzint een bron, dosering, normverwijzing of feit.	Broncontrole, vier-ogenprincipe en verbod op autonome kritieke besluiten.
Bias en discriminatie	Selectie, triage of beoordeling pakt structureel nadelig uit voor een groep.	Impactanalyse, representatieve tests, monitoring en menselijke bezwaarroute.
Prompt injection	Een extern document bevat instructies die de AI bewegen gegevens te delen of een actie uit te voeren.	Invoer als onvertrouwde data behandelen, toolrechten beperken, segmenteren en output valideren.
Data poisoning	Kennisbron of trainingsdata wordt doelbewust vervuild.	Herkomstcontrole, wijzigingsbeheer, integriteitscontroles en gescheiden goedkeuring.
Overmatige bevoegdheden	Een agent leest mailboxen, wijzigt dossiers of verstuurt berichten.	Least privilege, expliciete bevestiging, transactielimieten en volledige logging.
Leveranciersafhankelijkheid	Prijs, modelgedrag, regio of voorwaarden wijzigen onverwacht.	Exitplan, exporteerbaarheid, alternatieven, contractuele kennisgeving en continuïteitstest.
Intellectueel eigendom	Vertrouwelijke ontwerpen worden ingevoerd of output bevat beschermd materiaal.	Gebruiksvoorwaarden, bronregistratie, juridische toets en publicatiecontrole.
Schaduw AI	Medewerkers gebruiken eigen accounts of browserextensies.	Werkbare goedgekeurde alternatieven, inventarisatie, endpointbeheer en open meldcultuur.
Cybermisbruik	AI versnelt phishing, social engineering, malwarevarianten en desinformatie.	Awareness, verificatiekanalen, e-mailbeveiliging, monitoring en incidentrespons.
Modelwijziging	Dezelfde prompt geeft na een update een andere kwaliteit of veiligheidsgrens.	Versiebeheer waar mogelijk, regressietests, periodieke herbeoordeling en fallback.
Dossiersvervuiling	Niet-gecontroleerde AI-output wordt als feit in een dossier opgeslagen.	Herkomstmarkering, professionele validatie en correctieproces.
Verlies van vakbekwaamheid	Medewerkers controleren minder kritisch en bouwen minder expertise op.	Taakgrenzen, training in verificatie en periodiek werken zonder AI.
Reputatie en vertrouwen	Onzorgvuldige chatbotreactie schaadt cliënten of publiek vertrouwen.	Transparantie, escalatie naar mens, kwaliteitscontrole en klachtenafhandeling.

2 2 Vertrouwelijkheid en privacy

De meest directe fout is het invoeren van gegevens die de gekozen dienst niet mag ontvangen. Bij persoonsgegevens gelden de AVG-beginselen, waaronder doelbinding, minimale gegevensverwerking, juistheid, bewaarbeperking en passende beveiliging.

Gezondheidsgegevens zijn bijzondere persoonsgegevens. Pseudonimisering verlaagt het risico, maar maakt gegevens niet automatisch anoniem. Combinaties van leeftijd, diagnose, locatie, datum en bijzondere omstandigheden kunnen iemand opnieuw herkenbaar maken.

Ook zonder namen kan een prompt vertrouwelijk zijn. Denk aan kwetsbaarheden, auditbevindingen, overnames, juridische strategie, personeelskwesties, wachtwoordstructuren, systeemconfiguraties, prijsafspraken, broncode of niet-gepubliceerde beleidskeuzes.

2 3 Integriteit en hallucinaties

Generatieve modellen produceren waarschijnlijke output en geen gegarandeerde waarheid. Een overtuigende toon zegt niets over juistheid. Het model kan feiten, bronnen, normnummers, jurisprudentie, personen of berekeningen verzinnen. Een extra risico ontstaat wanneer medewerkers alleen controleren of een tekst logisch klinkt. Controle moet teruggaan naar de oorspronkelijke bron, actuele wetgeving, het dossier of een bevoegde deskundige.

2 4 Beschikbaarheid en continuïteit

Wanneer een proces afhankelijk wordt van één AI-leverancier, kunnen storing, capaciteitslimieten, accountblokkade, modelwijziging of contractbeëindiging het proces verstoren. Kritieke processen hebben een handmatige fallback, een maximale uitvalstijd, exportmogelijkheden en een geteste werkwijze voor degradatie nodig.

2 5 Nieuwe aanvalspaden

AI-systemen verwerken vaak inhoud uit bronnen die een organisatie niet beheerst. Een kwaadaardige instructie in een webpagina, e-mail of document kan een gekoppelde AI-agent beïnvloeden. Deze vorm van prompt injection lijkt functioneel op het aanbieden van onbetrouwbare input aan software, maar natuurlijke taal maakt scheiding tussen opdracht en inhoud moeilijk. Beperk daarom de acties die een model zelfstandig kan uitvoeren en laat gevoelige acties expliciet door een mens bevestigen.



VeeHa B.V.

Cybersecurity and privacy

2 6 Risico's in de zorg

- Een foutieve samenvatting kan relevante nuance uit een dossier verwijderen.
- Een model kan een niet-bestaande richtlijn of onjuiste klinische suggestie genereren.
- Een transcriptiedienst kan gesprekken, stemgegevens en gezondheidsgegevens buiten de afgesproken omgeving verwerken.
- Een chatbot kan een cliënt ten onrechte geruststellen of onvoldoende escaleren.
- Een bias in triage of planning kan de toegang tot zorg ongelijk beïnvloeden.
- Een storing of wijziging kan een inmiddels afhankelijk zorgproces stilleggen.
- Onvoldoende herleidbaarheid maakt het moeilijk om achteraf vast te stellen welke informatie tot een besluit leidde.

Bij toepassingen die zorginhoud, triage, diagnose, behandeling, indicatie of cliëntveiligheid raken, is menselijke eindverantwoordelijkheid niet slechts een controlemaatregel. Zij moet zichtbaar zijn in het proces, de autorisaties, de verslaglegging en de mogelijkheid om van de AI-uitkomst af te wijken.

3 ISO 27001 2022 en NEN 7510 2024

3 1 AI valt binnen het bestaande managementsysteem

Een ISMS biedt een geschikt raamwerk om AI beheerst in te voeren. De organisatie bepaalt de context en belanghebbenden, stelt leiderschap en beleid vast, beoordeelt risico's en kansen, organiseert middelen en competenties, beheerst de uitvoering, evalueert prestaties en verbetert. AI vraagt dus geen losstaand papierensysteem. Het vraagt uitbreiding van bestaande registers, risicoanalyses, leveranciersprocessen, incidentprocedures, auditprogramma's en directiebeoordelingen.

3 2 Koppeling met de hoofdstukken van ISO 27001

Onderdeel	Betekenis voor AI	Te verwachten bewijs
Hoofdstuk 4 Context	Bepaal interne en externe AI-ontwikkelingen, belanghebbenden, verplichtingen en ISMS-scope.	Contextanalyse, stakeholderanalyse, toepasselijke eisen en scope.
Hoofdstuk 5 Leiderschap	Stel richting, risicobereidheid, verantwoordelijkheden en beleid vast.	Goedgekeurd AI-beleid, rollen, besluiten en middelen.
Hoofdstuk 6 Planning	Beoordeel AI-risico's, kies behandelingen en formuleer meetbare doelen.	Risicoregister, behandelplan, Statement of Applicability en doelstellingen.
Hoofdstuk 7 Ondersteuning	Regel competentie, bewustwording, communicatie en beheerste documentatie.	Opleidingsplan, deelname, instructies, communicatie en versiebeheer.
Hoofdstuk 8 Uitvoering	Beheers gebruik, wijzigingen, leveranciers en uitbestede processen.	Use-case goedkeuringen, DPIA waar nodig, testen en operationele controles.
Hoofdstuk 9 Evaluatie	Monitor werking, voer interne audits uit en bespreek resultaten in de directiebeoordeling.	KPI's, logreviews, auditrapporten, incidenttrends en managementbesluiten.
Hoofdstuk 10 Verbetering	Registreer afwijkingen, herstel oorzaken en verbeter beleid en maatregelen.	Verbeterregister, oorzaakanalyse en effectiviteitscontrole.

3 3 Relevante thema's uit Annex A

De onderstaande koppeling is een praktische selectie. De definitieve toepasselijkheid volgt uit de eigen risicoanalyse en de Statement of Applicability.

Thema en controls	Toepassing op AI
Beleid en rollen A 5 1 tot A 5 4	AI-beleid, eigenaarschap, functiescheiding, verantwoord gebruik en managementverantwoordelijkheid.
Dreigingsinformatie en projecten A 5 7 en A 5 8	Volg AI-specifieke dreigingen en neem beveiliging vanaf de start in ieder AI-project op.
Assets en aanvaardbaar gebruik A 5 9 tot A 5 13	Registreer tools, modellen, datasets, koppelingen en informatieclassificatie; leg gebruiksregels vast.
Overdracht en toegang A 5 14 tot A 5 18	Beheers dataoverdracht, accounts, rechten, API-sleutels en toegang tot kennisbronnen.
Leveranciers en cloud A 5 19 tot A 5 23	Beoordeel contracten, subverwerkers, wijzigingen, cloudregio, training op data en exit.
Incidenten A 5 24 tot A 5 28	Herken AI-incidenten, stel triage en bewijsbewaring vast en leer van gebeurtenissen.
Continuïteit A 5 29 en A 5 30	Bepaal afhankelijkheden, fallback en ICT-gereedheid voor kritieke toepassingen.
Wetgeving en privacy A 5 31 tot A 5 36	Onderhoud verplichtingen, intellectueel eigendom, persoonsgegevens, onafhankelijke beoordeling en compliance.
Mensen A 6 1 tot A 6 8	Screening waar passend, voorwaarden, bewustwording, geheimhouding, werken op afstand en melden.
Technische beheersing A 8	Beperk rechten, beheer kwetsbaarheden en configuratie, verwijder of maskeer data, voorkom lekkage, log en monitor, versleutel en pas secure development toe.

3 4 Extra betekenis van NEN 7510 2024

NEN 7510 2024 plaatst informatiebeveiliging in de context van persoonlijke gezondheidsinformatie en zorgverlening. Daardoor moeten risico's niet alleen worden beoordeeld op financiële of technische schade. Effecten op continuïteit van zorg, professionele verantwoordelijkheid, vertrouwelijkheid van het medisch dossier, juistheid van informatie, veiligheid van cliënten en vertrouwen in de zorg wegen zwaar.

De organisatie moet vastleggen welke AI-toepassingen zorginformatie verwerken, welk zorgproces wordt ondersteund, wie inhoudelijk verantwoordelijk is, welke validatie plaatsvindt en welke terugvalprocedure geldt. Voor een administratieve tekstassistent is een ander beheersniveau passend dan voor triage, diagnose, behandeladvies of geautomatiseerde dossiervoering.

3 5 Relatie met andere kaders

Kader	Rol naast het ISMS
AVG	Regelt de rechtmatige, transparante en veilige verwerking van persoonsgegevens, rechten van betrokkenen en waar nodig een DPIA.
EU AI Act	Hanteert een risicogebaseerd stelsel met verboden praktijken, verplichtingen voor onder meer hoog-risicosystemen, transparantie en AI-geletterdheid.
ISO IEC 42001	Biedt eisen voor een managementsysteem voor verantwoord ontwikkelen, leveren en gebruiken van AI.
ISO IEC 23894	Geeft richting aan risicomanagement voor AI.
Sectorale regels	Denk aan medische hulpmiddelen, beroepsgeheim, WGBO, kwaliteitswetgeving, arbeidsrecht en archief- of dossierplichten.

De EU AI Act trad op 1 augustus 2024 in werking. Volgens de Europese Commissie is het grootste deel sinds 2 augustus 2026 van toepassing, met gefaseerde uitzonderingen. De verplichting rond AI-geletterdheid gold eerder. Organisaties moeten hun concrete rol en toepassing juridisch classificeren; een ISO- of NEN-certificaat vervangt die beoordeling niet. [1]

4 Medewerkers bewust en handelingsbekwaam maken

4 1 Bewustwording is meer dan een jaarlijkse e learning

Medewerkers moeten op het moment van gebruik de juiste beslissing kunnen nemen. Daarvoor hebben zij herkenbare voorbeelden, korte beslisregels, een goedgekeurde tool en een laagdrempelig meldpunt nodig. Alleen waarschuwen voor risico's werkt slecht wanneer de organisatie geen werkbaar alternatief biedt.

4 2 De vijf vragen voor iedere prompt

1. Mag ik deze tool voor mijn werk gebruiken en ben ik ingelogd in de goedgekeurde organisatieomgeving?
2. Welke informatie voer ik in en hoe is die geclassificeerd?
3. Bevat de invoer persoonsgegevens, gezondheidsgegevens, geheimen, broncode of informatie van een klant of leverancier?
4. Kan ik het doel bereiken met minder gegevens, geanonimiseerde informatie of een fictief voorbeeld?
5. Hoe controleer ik de output voordat deze invloed heeft op een persoon, dossier, besluit, systeem of externe communicatie?

4 3 Verkeerslicht voor gegevensinvoer

Categorie	Voorbeelden	Regel
Groen	Publieke informatie, eigen algemene ideeën, fictieve voorbeelden zonder herleidbare details.	Toegestaan in goedgekeurde tools, met controle van output.
Oranje	Interne concepten, beperkte bedrijfsinformatie, gepseudonimiseerde casuïstiek.	Alleen in goedgekeurde zakelijke omgeving en wanneer doel, instellingen en contract dit toestaan.
Rood	Gezondheidsgegevens, cliënt- of personeelsdossiers, wachtwoorden, sleutels, kwetsbaarheden, geheime strategie, niet-gepubliceerde broncode.	Niet invoeren, tenzij de toepassing formeel is goedgekeurd voor precies deze gegevens en het proces aanvullende controles bevat.

Pseudoniemen zoals “cliënt A” maken een casus niet veilig wanneer de overige details herkenning mogelijk maken. Verwijder ook data, locaties, zeldzame diagnoses, functietitels en combinaties van kenmerken die indirect naar een persoon leiden.

4 4 Output veilig gebruiken

- Controleer feiten en citaten aan de primaire bron.
- Laat berekeningen opnieuw uitvoeren met een geschikt rekenmiddel.
- Test gegenereerde code en controleer afhankelijkheden, beveiliging en licenties.
- Laat professionele of klinische conclusies beoordelen door een bevoegde medewerker.
- Markeer AI-ondersteuning wanneer transparantie voor de ontvanger relevant of verplicht is.
- Neem geen volledig geautomatiseerde ingrijpende beslissing over een persoon zonder passende juridische basis, waarborgen en menselijke tussenkomst.
- Sla alleen gevalideerde informatie op in dossiers, beleid, rapporten of tickets.

4 5 Opleidingsprogramma op basis van rol

Doelgroep	Minimale inhoud	Frequentie
Alle medewerkers	Toegestane tools, invoerregels, hallucinaties, privacy, auteursrecht, phishing en incidentmelding.	Bij indiensttreding, jaarlijks en na belangrijke wijziging.
Management en proceseigenaren	Risicobereidheid, use-case goedkeuring, AI Act, KPI's, menselijk toezicht en maatschappelijke gevolgen.	Jaarlijks plus besluitgerichte sessies.
ICT en security	Architectuur, API's, prompt injection, logging, DLP, IAM, model- en leverancierswijzigingen.	Minimaal jaarlijks en bij nieuwe techniek.
Privacy en juridisch	Rollen, grondslagen, DPIA, transparantie, rechten, doorgifte, contracten en intellectueel eigendom.	Minimaal jaarlijks en bij relevante wetwijziging.
Zorgprofessionals	Dossierkwaliteit, professionele verantwoordelijkheid, klinische validatie, bias, escalatie en patiëntcommunicatie.	Praktijkgericht, terugkerend en per toepassing.
Ontwikkelaars en data teams	Secure AI lifecycle, datakwaliteit, evaluaties, misbruiktests, modeldocumentatie en monitoring.	Doorlopend binnen development lifecycle.

5 Beleid en governance

5 1 Een bruikbaar AI beleid

Een beleid moet besluiten ondersteunen. Een document dat alleen stelt dat medewerkers zorgvuldig moeten omgaan met AI is niet toetsbaar. Het beleid moet minimaal beschrijven: doel en reikwijdte, definities, uitgangspunten, rollen, toegestane en verboden toepassingen, gegevensclassificatie, goedkeuring, leveranciersbeoordeling, menselijke controle, transparantie, logging, incidenten, opleiding, monitoring, uitzonderingen en periodieke herziening.

5 2 Beleidsuitgangspunten

1. De organisatie gebruikt alleen AI-toepassingen die passen bij een vastgesteld doel en een aanvaard risico.
2. De mens blijft verantwoordelijk voor besluiten en communicatie die gevolgen hebben voor cliënten, medewerkers, leveranciers of andere betrokkenen.
3. Medewerkers verwerken informatie uitsluitend in goedgekeurde omgevingen en volgens de classificatie- en privacyregels.
4. AI-output geldt als concept of advies totdat een bevoegde persoon de inhoud passend heeft gecontroleerd.
5. De organisatie registreert toepassingen, leveranciers, gegevensstromen, eigenaren, risico's en belangrijke wijzigingen.
6. Toegang en technische mogelijkheden worden beperkt tot wat voor de taak nodig is.
7. De organisatie biedt uitleg, correctie, bezwaar en menselijke tussenkomst waar de toepassing personen raakt.
8. Incidenten, twijfelgevallen en onverwacht gedrag worden direct gemeld zonder dat medewerkers worden ontmoedigd om fouten bespreekbaar te maken.
9. De organisatie evalueert werking, proportionaliteit en maatschappelijke effecten en stopt een toepassing wanneer beheersing tekortschiet.

5 3 Governance en rollen

Rol	Kernverantwoordelijkheid
Directie	Stelt richting, risicobereidheid en middelen vast; accepteert restrisico's met grote impact.
AI of ISMS governancegroep	Beoordeelt toepassingen integraal en bewaakt register, beleid, risico's en verbeteringen.
Proceseigenaar	Formuleert doel, noodzaak, prestatienorm, menselijke controle en fallback.
Informatie eigenaar	Bepaalt classificatie, toegang, kwaliteit, bewaartermijn en toegestaan gebruik.
CISO of securityfunctie	Beoordeelt dreigingen, architectuur, IAM, logging, leveranciersrisico en incidenten.
FG of privacyfunctie	Adviseert over AVG, DPIA, transparantie, rechten en gegevensminimalisatie.
Juridische functie	Beoordeelt AI Act, contracten, aansprakelijkheid, intellectueel eigendom en sectorale regels.
ICT en beheer	Configureert, monitort, beheert wijzigingen en voert technische maatregelen uit.
Gebruiker	Volgt instructies, minimaliseert invoer, verifieert output en meldt afwijkingen.
Interne audit	Toetst onafhankelijk of ontwerp en werking van beheersmaatregelen voldoen.

5 4 Goedkeuringsproces per toepassing

1. Beschrijf doel, gebruiker, geraakte personen, gewenste output en beslissingen die ermee worden ondersteund.
2. Breng gegevens, classificatie, bronnen, koppelingen, opslag, landen en ontvangers in kaart.
3. Classificeer de toepassing onder relevante wet- en regelgeving en bepaal of een DPIA of andere impactanalyse nodig is.
4. Beoordeel informatiebeveiliging, zorgveiligheid, mensenrechten, bias, leveranciersrisico, continuïteit en misbruikscenario's.
5. Kies maatregelen, test normale werking en misbruik, leg acceptatiecriteria vast en wijs een eigenaar aan.
6. Laat de bevoegde rollen goedkeuren en registreer restrisico en eventuele geldigheidsduur.
7. Start gecontroleerd, monitor resultaten en herbeoordeel na wijzigingen, incidenten of uiterlijk op de geplande evaluatiedatum.

5 5 Minimale leveranciersvragen

- Worden prompts, bijlagen, output of metadata gebruikt voor training of productverbetering, en kan dit aantoonbaar worden uitgeschakeld?
- Waar worden gegevens opgeslagen en verwerkt, welke subverwerkers zijn betrokken en welke doorgiftemechanismen gelden?
- Welke retentie geldt voor gesprekken, logs, back-ups en verwijderde accounts?
- Welke encryptie, toegangsbeveiliging, tenant-scheiding, logging en incidentmelding zijn ingericht?
- Welke beheerders of reviewers kunnen inhoud zien en onder welke voorwaarden?
- Hoe kondigt de leverancier wijzigingen in model, functionaliteit, voorwaarden, regio of subverwerkers aan?
- Kan de organisatie gegevens exporteren en verwijderen en is een werkbare exit mogelijk?
- Welke assurance is beschikbaar, zoals auditrapporten, certificaten, penetratietests en transparantiedocumentatie?
- Welke beperkingen, bekende foutmarges en gebruiksvoorwaarden gelden voor de beoogde toepassing?

6 Incidenten en operationele beheersing

6 1 Wat is een AI incident

Een AI-incident is iedere gebeurtenis waarbij gebruik, output, koppeling of leverancier de vertrouwelijkheid, integriteit, beschikbaarheid, privacy, veiligheid, rechten of naleving bedreigt. Voorbeelden zijn invoer van gevoelige gegevens in een ongeoorloofde tool, ongewenste openbaarmaking, discriminerende output, een autonome onjuiste actie, misleiding door deepfake, verlies van logging of een modelwijziging die kritieke prestaties verslechtert.

6 2 Eerste respons

1. Stop of beperk de toepassing zonder noodzakelijk bewijs te vernietigen.
2. Leg prompt, output, tijd, account, modelversie, instellingen, gekoppelde bronnen en uitgevoerde acties vast.
3. Bepaal welke gegevens en personen geraakt zijn en voorkom verdere verspreiding.
4. Betrek security, privacy, proceseigenaar, communicatie en zorginhoudelijke deskundigheid naar behoefte.
5. Beoordeel meldplichten en termijnen onder AVG, AI Act, contracten en sectorale regels.
6. Herstel, informeer betrokkenen waar nodig en controleer of herstel daadwerkelijk werkt.
7. Onderzoek de oorzaak en pas techniek, proces, training, contract of beleidsgrens aan.

6 3 Monitoring en indicatoren

Indicator	Voorbeeld van meting	Signaal
Inventarisatie	Percentage bekende AI-toepassingen met eigenaar en risicobeoordeling.	Onbekend of dalend percentage wijst op schaduw AI.
Opleiding	Deelname en praktijkscore per doelgroep.	Hoge deelname met lage praktijkscore vraagt ander leermiddel.
Datagebruik	Aantal DLP-signalen of gemelde gevoelige prompts.	Trend en oorzaak zijn belangrijker dan alleen nul incidenten.
Kwaliteit	Steekproeffouten, correcties en afwijzingen van AI-output.	Stijging kan wijzen op modelwijziging of verkeerd gebruik.
Menselijke controle	Percentage kritieke outputs aantoonbaar beoordeeld.	Ontbrekende registratie betekent onvoldoende assurance.
Leveranciers	Beoordelingen en wijzigingen tijdig behandeld.	Verlopen beoordeling of niet-beoordeelde wijziging verhoogt restrisico.
Incidenten	Aantal, ernst, detectietijd en herhaling.	Herhaling toont dat oorzaakanalyse of maatregel niet effectief is.
Waarde	Tijdwinst of kwaliteitsverbetering tegenover kosten en risico.	Een toepassing zonder aantoonbare waarde verdient heroverweging.

7 Hoe gaan we als samenleving met AI om

7 1 De maatschappelijke afweging

AI kan toegankelijkheid verbeteren, administratieve lasten verminderen, kennis ontsluiten en professionals ondersteunen. Tegelijk concentreert de technologie gegevens, rekenkracht en beslistmacht bij een beperkt aantal aanbieders. De samenleving moet daarom niet alleen vragen wat technisch mogelijk is, maar wie voordeel ontvangt, wie risico draagt en wie een uitkomst kan betwisten.

7 2 Principes voor maatschappelijk verantwoord gebruik

- Menselijke waardigheid en autonomie: mensen mogen niet gereduceerd worden tot een voorspelling of risicoscore.
- Rechtvaardigheid: voordelen en fouten mogen niet structureel bij verschillende groepen terecht komen.
- Transparantie: iemand moet weten wanneer AI relevant bijdraagt aan een interactie of besluit en waar uitleg te krijgen is.
- Betwistbaarheid: er moet een bereikbare menselijke route bestaan voor correctie, bezwaar en herstel.
- Proportionaliteit: zet geen zwaar systeem in voor een doel dat eenvoudiger en met minder gegevens kan worden bereikt.
- Publieke controle: overheid, toezichthouders, wetenschap, maatschappelijke organisaties en burgers moeten tegenmacht kunnen organiseren.
- Duurzaamheid: energie-, water- en grondstoffengebruik behoren mee te wegen bij grootschalige toepassing.
- Informatiebetrouwbaarheid: herkomst, authenticiteit en mediawijsheid worden belangrijker door synthetische tekst, beeld, audio en video.

7 3 Werk en vakmanschap

AI verandert taken eerder dan volledige beroepen. De kwaliteit van werk kan stijgen wanneer routinetaken afnemen en professionals tijd houden voor oordeel en contact. De kwaliteit kan ook dalen wanneer organisaties uitsluitend op productiviteit sturen, toezicht automatiseren of junior medewerkers geen gelegenheid geven om expertise op te bouwen. Goed werkgeverschap vraagt betrokkenheid van medewerkers, duidelijke taakgrenzen, opleiding en monitoring van werkdruk en de kwaliteit van arbeid.

7 4 Zorg en menselijke maat

In de zorg is efficiëntie waardevol, maar geen zelfstandig einddoel. Een toepassing moet bijdragen aan veilige, toegankelijke en passende zorg. Cliënten moeten weten hoe zij een mens bereiken, hoe hun gegevens worden gebruikt en hoe zij een fout kunnen laten corrigeren. Professionals moeten kunnen afwijken van AI-advies zonder oneigenlijke druk en zonder dat het systeem hun professionele oordeel onzichtbaar vervangt.

7 5 Desinformatie en vertrouwen

Generatieve AI verlaagt de kosten van geloofwaardige nepinhoud. Technische detectie alleen is onvoldoende, omdat detectiemethoden achterlopen en fouten maken. Organisaties hebben verificatiekanalen, woordvoerdersprocedures, bronvastlegging en oefeningen met deepfake-scenario's nodig. Voor burgers en medewerkers zijn broncontrole, context en het opschorten van snelle emotionele reacties belangrijke vaardigheden.

8 Implementatieroadmap voor negentig dagen

Periode	Prioriteit	Concrete resultaten
Dag 1 tot 30	Zicht en tijdelijke grenzen	Bestuurlijke sponsor, voorlopig gebruikskader, meldpunt, eerste inventarisatie, blokkade van evident risicovolle tools en keuze van werkbaar alternatief.
Dag 31 tot 60	Risico en inrichting	AI-register, gegevensverkeerslicht, beoordelingsproces, leverancierscriteria, rollen, training en selectie van pilots.
Dag 61 tot 90	Werking en bewijs	Goedgekeurd beleid, uitgevoerde beoordelingen, technische configuratie, logreview, incidentoefening, KPI-baseline en rapportage aan directie.
Daarna	Borging en verbetering	Interne audit, periodieke evaluatie, wijzigingsmonitoring, herhalingstraining, directiebeoordeling en koppeling aan verbeterregister.

8 1 Begin met drie registers

- AI-register: toepassing, doel, leverancier, eigenaar, gebruikers, gegevens, koppelingen, risicoklasse, goedkeuring en evaluatiedatum.
- Risicoregister: dreiging, kwetsbaarheid, gevolg, bestaande maatregelen, kans, impact, eigenaar, behandeling en restrisico.
- Verbeterregister: incidenten, audits, klachten, modelwijzigingen, tests en besluiten met actiehouders en termijn.

Gebruik waar mogelijk bestaande middelen zoals SharePoint, Teams, Planner, Jira of Confluence. Een afzonderlijk zwaar AI-governanceplatform is niet nodig om te starten. Werkbaarheid, eigenaarschap en aantoonbare uitvoering wegen zwaarder dan de aanschaf van een extra systeem.

9 Interne audit en directiebeoordeling

9 1 Tien kernvragen voor een interne audit

1. Welke AI-toepassingen gebruikt de organisatie werkelijk, inclusief persoonlijke accounts, extensies en ingebouwde functies?
2. Hoe bepaalt de organisatie welke informatie in welke AI-omgeving mag worden verwerkt?
3. Welke risico- en impactbeoordeling is uitgevoerd voordat een toepassing in gebruik ging?
4. Hoe zijn leverancier, contract, subverwerkers, opslag, training op data, retentie en exit beoordeeld?
5. Wie is eigenaar van de toepassing en wie blijft verantwoordelijk voor de inhoudelijke uitkomst?
6. Hoe wordt menselijke controle aantoonbaar uitgevoerd en wanneer is escalatie verplicht?
7. Welke instructie en oefening hebben medewerkers ontvangen en kunnen zij de regels in praktijksituaties toepassen?
8. Hoe worden outputkwaliteit, bias, beveiliging, incidenten en modelwijzigingen gemonitord?
9. Welke fallback bestaat wanneer de dienst uitvalt, verandert of onbetrouwbare resultaten geeft?
10. Welke verbeteringen zijn na incidenten, klachten, audits of evaluaties doorgevoerd en is de effectiviteit vastgesteld?

9 2 Onderwerpen voor de directiebeoordeling

- Veranderingen in wetgeving, dreigingen, leveranciers en maatschappelijke verwachtingen.
- Omvang en ontwikkeling van het AI-register en resterende schaduw AI.
- Resultaten van risicoanalyses, DPIA's, tests, interne audits en leveranciersbeoordelingen.
- Incidenten, klachten, correcties, bias-signalen en effecten op cliënten of medewerkers.
- Prestaties van maatregelen, opleiding, menselijke controle en continuïteit.
- Waarde, kosten, personele gevolgen en besluiten over opschalen, aanpassen of stoppen.
- Benodigde middelen, competenties, prioriteiten en acceptatie van restrisico's.

10 Praktische checklist voor management

Controlepunt	Ja	Actie nodig
Er is een bestuurlijk eigenaar voor verantwoord AI-gebruik.	☐	☐
Alle bekende AI-toepassingen staan in een register.	☐	☐
Medewerkers beschikken over een goedgekeurde, werkbare AI-omgeving.	☐	☐
Invoerregels zijn gekoppeld aan de informatieclassificatie.	☐	☐
Hoog-risico en zorginhoudelijke toepassingen worden vooraf goedgekeurd.	☐	☐
Leveranciersvoorwaarden, training op data, locatie, retentie en subverwerkers zijn beoordeeld.	☐	☐
Menselijke controle en eindverantwoordelijkheid zijn per proces vastgelegd.	☐	☐
Medewerkers hebben rolgerichte AI-geletterdheid en praktijksituaties geoefend.	☐	☐
AI-incidenten kunnen eenvoudig worden gemeld en onderzocht.	☐	☐
KPI's, logreviews, steekproeven en herbeoordelingen zijn gepland.	☐	☐
Continuïteit en exit zijn geregeld voor belangrijke toepassingen.	☐	☐
De directie ontvangt periodiek informatie en neemt aantoonbare besluiten.	☐	☐

Slotbeschouwing

AI wordt beheersbaar wanneer de organisatie verantwoordelijkheid concreet maakt. Dat begint bij zicht op gebruik en gegevens. Daarna volgen risicogebaseerde toelating, geschikte contracten en techniek, competente medewerkers, menselijke controle en aantoonbaar toezicht. ISO 27001 2022 en NEN 7510 2024 bieden hiervoor een solide managementsysteem. Zij voorkomen echter geen schade wanneer beleid losstaat van de dagelijkse praktijk.

De volwassen organisatie durft beide kanten te zien. Zij benut AI waar de toepassing waarde toevoegt en de risico's beheersbaar zijn. Zij begrenst of stopt AI waar gegevens, veiligheid, rechten of professionele verantwoordelijkheid onvoldoende kunnen worden beschermd. Dat vraagt geen eenmalig project, maar een blijvende cyclus van leren, toetsen en verbeteren.

Bronnen en verdere verdieping

- [1] **Europese Commissie** AI Act en tijdlijn van toepassing <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>
 - [2] **Europese Unie** Regulation EU 2024 1689 Artificial Intelligence Act <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
 - [3] **ISO** ISO IEC 27001 2022 Information security management systems <https://www.iso.org/standard/27001>
 - [4] **NEN** NEN 7510 Informatiebeveiliging in de zorg <https://www.nen.nl/nen-7510-1-2024-nl-324644>
 - [5] **ISO** ISO IEC 42001 2023 AI management systems <https://www.iso.org/standard/42001>
 - [6] **Europese Commissie High Level Expert Group** Ethics guidelines for trustworthy AI <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>
 - [7] **AIVD MIVD en NCTV** Versterkte dreigingen in een wereld vol kunstmatige intelligentie <https://www.aivd.nl/documenten/2024/12/10/versterkte-dreigingen-in-een-wereld-vol-kunstmatige-intelligentie>
 - [8] **Autoriteit Persoonsgegevens** Algoritmes en AI <https://www.autoriteitpersoonsgegevens.nl/themas/algoritmes-ai>
 - [9] **EDPB** Gegevensbescherming en internationale doorgifte https://www.edpb.europa.eu/sme/be-compliant/international-data-transfers_nl
 - [10] **NIST** Artificial Intelligence Risk Management Framework <https://www.nist.gov/itl/ai-risk-management-framework>
- Geraadpleegd in september 2026. Wetgeving, norminterpretaties en leveranciersvoorwaarden kunnen wijzigen. Toets bij concrete toepassingen altijd de actuele tekst, contracten en sectorale verplichtingen.

Over VeeHa

VeeHa B.V. ondersteunt organisaties bij informatiebeveiliging, privacy en werkbare managementsystemen. De aanpak sluit aan op wat binnen de organisatie al aanwezig is. Wat goed werkt, blijft behouden. Beleid, risico's, planning, bewijs en verbeteracties worden zoveel mogelijk ingericht in bestaande omgevingen zoals Microsoft 365, SharePoint, Teams, Planner, Jira en Confluence.

De dienstverlening omvat consultancy, implementatie, gap-analyses en interne audits voor onder meer ISO 27001, NEN 7510, ISO 9001, ISO 27701, AVG, BIO en NIS2. De nadruk ligt op een aantoonbare balans tussen norm, risico en dagelijkse praktijk.

Deze whitepaper is algemene informatie en geen juridisch, medisch of sectorspecifiek advies voor een concrete toepassing.