



ACADEMIE NEDERLAND

Grip op AI-risico's

Een praktische aanpak
voor AI-risicoanalyse en
verantwoord AI-gebruik.

WHITEPAPER

INSPIRATIE

ONTWIKKELING

TOEPASSING

AI biedt kansen, maar introduceert ook nieuwe risico's

AI wordt in steeds meer organisaties onderdeel van het dagelijkse werk. Medewerkers gebruiken generatieve AI voor teksten en analyses, organisaties zetten Copilot of ChatGPT in en steeds vaker worden AI-agents en geautomatiseerde workflows ontwikkeld.

Dat biedt kansen: werk kan sneller, kennis wordt toegankelijker en routinetaken kunnen gedeeltelijk worden geautomatiseerd.

Maar wat kan er misgaan?

Vertrouwelijke informatie kan onbedoeld in een AI-systeem terechtkomen. Een foutief antwoord kan besluitvorming beïnvloeden. Bias kan groepen benadelen. Een agent kan zelfstandig een verkeerde actie uitvoeren.

De vraag is niet: “Kunnen we AI gebruiken?”

Maar: “Onder welke voorwaarden kunnen we deze toepassing verantwoord gebruiken?”

Een AI-risicoanalyse helpt organisaties deze vraag systematisch te beantwoorden. Niet alleen vanuit wet- en regelgeving, maar vooral als praktisch hulpmiddel om betere beslissingen over AI te nemen.

Wat is een AI-risicoanalyse?

Een AI-risicoanalyse is een gestructureerde beoordeling van risico's die ontstaan bij het ontwikkelen, inkopen, implementeren of gebruiken van een AI-toepassing.

Beoordeel nooit alleen de tool

AI-systeem + gegevens + gebruiker + toepassing + context + mogelijke gevolgen

A

Laag risico

Een medewerker gebruikt AI om tien ideeën voor een titel van een interne nieuwsbrief te bedenken.

De mogelijke schade bij een verkeerd antwoord is klein.

B

Hoger risico

Een HR-medewerker gebruikt AI om te bepalen welke sollicitanten worden uitgenodigd voor een gesprek.

De gevolgen voor personen zijn aanzienlijk groter.

Risico is contextafhankelijk.

Analyseer vooral waarvoor AI wordt gebruikt, welke informatie wordt verwerkt, wie wordt geraakt en wat er gebeurt wanneer het systeem een fout maakt.

De 8 belangrijkste AI-risicogebieden

01

Privacy & persoonsgegevens

Worden persoonsgegevens, bijzondere persoonsgegevens of vertrouwelijke gegevens verwerkt?

02

Onjuiste informatie

AI kan overtuigend klinkende informatie produceren die feitelijk onjuist of onvolledig is.

03

Bias & discriminatie

AI kan bestaande vooroordelen reproduceren en groepen ongelijk behandelen.

04

Informatiebeveiliging

Denk aan datalekken, prompt injection, ongewenste toegang en misbruik.

05

Vertrouwelijkheid & IE

Mag de informatie worden gedeeld en wie heeft rechten op gebruikte of gegenereerde content?

06

Transparantie

Is duidelijk dat AI wordt ingezet en zijn beperkingen begrijpelijk?

07

Menselijke controle

Kan een medewerker AI-uitkomsten controleren, corrigeren of negeren?

08

Operationele afhankelijkheid

Wat gebeurt er wanneer het systeem niet beschikbaar is, verandert of fout handelt?

Het 7-stappenmodel voor AI-risicoanalyse

1

Inventariseer de toepassing

Beschrijf doel, technologie, gebruikers, gegevens en mate van autonomie.

2

Bepaal mogelijke gevolgen

Wat kan er redelijkerwijs misgaan en wie kan hierdoor worden geraakt?

3

Bepaal kans en impact

Hoe waarschijnlijk is het risico en hoe ernstig zijn de gevolgen?

4

Bepaal het initiële risico

Risicoscore = kans × impact.

5

Kies beheersmaatregelen

Verlaag risico met controle, techniek, beleid, training of beperkingen.

6

Bepaal het restrisico

Is het resterende risico acceptabel gezien het voordeel van de toepassing?

7

Monitor en herbeoordeel

Nieuwe modellen, gegevens en processen kunnen het risicoprofiel veranderen.

De belangrijkste vragen per AI-toepassing

01

Doel

Waarom willen we AI gebruiken? Is AI noodzakelijk of bestaat er een eenvoudiger alternatief?

02

Gegevens

Welke gegevens krijgt het systeem? Bevatten die persoonsgegevens of vertrouwelijke informatie?

03

Uitkomst

Is AI alleen inspiratie of beïnvloedt de uitkomst daadwerkelijk een besluit?

04

Mensen

Welke personen of groepen kunnen gevolgen ondervinden?

05

Controle

Wie controleert het resultaat en kan die persoon AI ook afwijzen?

06

Autonomie

Geeft AI alleen informatie of voert het systeem zelfstandig acties uit?

07

Transparantie

Is voor gebruikers en betrokkenen duidelijk dat AI wordt ingezet?

08

Leverancier

Waar worden gegevens verwerkt en welke contractuele afspraken gelden?

09

Incidenten

Wat doen we bij een ernstig verkeerd resultaat en wie mag het systeem stilzetten?

Van risico naar beheersmaatregel

Vertrouwelijke informatie

Gevoelige gegevens komen bij een externe AI-dienst terecht.



Beheersmaatregel

Goedgekeurde bedrijfsaccounts, dataminimalisatie en duidelijke gegevensregels.

Feitelijk onjuiste output

AI genereert foutieve of onvolledige informatie.



Beheersmaatregel

Menselijke controle en verificatie bij primaire bronnen.

Automatiseringsbias

Medewerkers vertrouwen AI te snel.



Beheersmaatregel

AI-geletterdheid: leren wanneer je moet controleren of juist niet moet gebruiken.

Ongewenste agent-acties

Een AI-agent voert zelfstandig een verkeerde actie uit.



Beheersmaatregel

Beperkte autorisaties, goedkeuringsmomenten, logging en een noodstop.

Hoe groter de autonomie van AI, hoe belangrijker controle vooraf wordt.

AI-risicoklassen binnen de Europese AI-verordening

1

Verboden toepassingen

Een beperkt aantal AI-praktijken wordt als onacceptabel beschouwd en is verboden.

2

Hoog risico

Onder meer bepaalde toepassingen rond werkgelegenheid, onderwijs, kritieke infrastructuur, essentiële diensten en rechtspraak.

3

Transparantieplichtingen

In bepaalde situaties moeten gebruikers weten dat zij met AI communiceren of dat content door AI is gegenereerd.

4

Minimaal of beperkt risico

Veel dagelijkse toepassingen vallen niet onder hoog risico, maar kunnen intern nog steeds privacy-, security- of reputatierisico's veroorzaken.

Compliance is het minimum. Goed AI-risicomanagement gaat verder.

De Europese AI-verordening is sinds 2 augustus 2026 grotendeels van toepassing. Voor specifieke hoog-risicosystemen gelden latere toepassingsdata.

AI gebruiken binnen HR

Scenario: AI analyseert cv's en rangschikt kandidaten.

Mogelijke risico's

Bias

Bepaalde groepen kunnen worden benadeeld.

Onjuiste beoordeling

Ervaring of vaardigheden kunnen verkeerd worden geïnterpreteerd.

Gebrek aan transparantie

Recruiters en kandidaten begrijpen mogelijk onvoldoende hoe scores ontstaan.

Automatiseringsbias

Recruiters wegen AI-scores te zwaar.

Privacy

Cv's bevatten veel persoonsgegevens.

Mogelijke maatregelen

- AI uitsluitend ondersteunend gebruiken.
- Geen automatische afwijzingen.
- Recruiter blijft verantwoordelijk.
- Resultaten controleren op afwijkingen tussen groepen.
- Duidelijke grenzen aan gebruikte data.
- Medewerkers trainen in beperkingen en risico's.

Kunnen we verantwoorden hoe het systeem wordt gebruikt en welke gevolgen dit voor mensen kan hebben?

Checklist AI-risicoanalyse

1. Toepassing

- ☐ Doel is duidelijk beschreven.
- ☐ Verantwoordelijke is aangewezen.
- ☐ Technologie en leverancier zijn bekend.

2. Gegevens

- ☐ Bekend welke gegevens worden verwerkt.
- ☐ Persoons- en vertrouwelijke gegevens zijn geïdentificeerd.
- ☐ Niet meer data dan noodzakelijk.

3. Impact

- ☐ Mogelijke fouten en gevolgen zijn beschreven.
- ☐ Getroffen personen en groepen zijn bekend.
- ☐ Kans en impact zijn beoordeeld.

4. Menselijke controle

- ☐ Duidelijk welke beslissingen mensen zelf nemen.
- ☐ AI-uitkomsten kunnen worden gecorrigeerd.
- ☐ Verantwoordelijkheden zijn vastgelegd.

5. Beveiliging

- ☐ Toegang en autorisaties zijn beperkt.
- ☐ Gegevensgebruik is beheerst.
- ☐ Agent-acties zijn begrensd.

6. Transparantie

- ☐ Bekend wanneer gebruikers geïnformeerd moeten worden.
- ☐ Belangrijke uitkomsten zijn uitlegbaar.
- ☐ Beslissingen worden gedocumenteerd.

7. Medewerkers

- ☐ Medewerkers begrijpen mogelijkheden en beperkingen.
- ☐ Er zijn praktische gebruiksrichtlijnen.
- ☐ AI-geletterdheid past bij rol en risico.

8. Monitoring

- ☐ Prestaties en incidenten worden gevolgd.
- ☐ Er is een eigenaar voor herbeoordeling.
- ☐ De toepassing kan worden aangepast of gestopt.

Van losse experimenten naar AI-governance

Naarmate organisaties meer AI gebruiken ontstaat behoefte aan structurele afspraken en duidelijk eigenaarschap.

Eigenaarschap

Wie is verantwoordelijk voor AI binnen de organisatie?

Inventarisatie

Welke AI-systemen en toepassingen worden daadwerkelijk gebruikt?

Goedkeuring

Wat mag zelfstandig en waarvoor is vooraf toestemming nodig?

Risicoclassificatie

Wanneer volstaat een snelle scan en wanneer is een uitgebreide analyse nodig?

Beleid

Welke gegevens en toepassingen zijn toegestaan of uitgesloten?

Monitoring & opleiding

Hoe worden incidenten gevolgd en welke kennis hebben medewerkers nodig?

Medewerkers moeten weten: wat mag — wat niet mag — wanneer aanvullende beoordeling nodig is — en bij wie zij terecht kunnen.

Begin klein, maar begin gestructureerd

1

Inventariseer huidig AI-gebruik

Breng in kaart welke tools en toepassingen medewerkers nu al gebruiken.

2

Selecteer de gevoelige toepassingen

Focus op persoonsgegevens, belangrijke besluiten, vertrouwelijke data en geautomatiseerde acties.

3

Voer de risicoanalyse uit

Bepaal kans, impact en mogelijke gevolgen.

4

Bepaal maatregelen en eigenaarschap

Leg vast wie verantwoordelijk is en welke controles nodig zijn.

5

Maak het proces herhaalbaar

Gebruik dezelfde aanpak bij nieuwe AI-toepassingen en herbeoordeel regelmatig.

Van

**“We hopen dat medewerkers
AI verstandig gebruiken.”**

Naar

**“We weten waar AI wordt
gebruikt, welke risico's spelen
en hoe we die beheersen.”**



**ACADEMIE
NEDERLAND**

Hoe volwassen is uw organisatie?

AI-risicomanagement begint niet bij technologie, maar bij mensen,
processen en verantwoordelijkheid.

Drie vragen om vandaag mee te beginnen

01 Waar gebruiken we AI?

02 Waar kan het fout gaan?

**03 Hebben we voldoende maatregelen
om dat risico te beheersen?**