



Verantwoord AI-beleid

Doelgroep

Professionals die AI willen inzetten zonder dat de privacy-afdeling spontaan begint te hyperventileren.

Duur

09:30 – 16:30

09:30 – 10:00 Kick-off & Reality Check

- Welkomstwoord en korte kennismaking (“Hoe vaak gebruik jij stiekem AI zonder het te zeggen?”)
- Inventarisatie: waar staat de organisatie nu qua AI-gebruik?
- Belangrijkste misverstanden over ChatGPT en Copilot—en waarom iedereen denkt dat zijn prompt *veilig* is.

10:00 – 11:00 De AI-toolbox: ChatGPT en Copilot in de praktijk

- Demonstratie: wat kunnen de tools écht en wat lijkt vooral magie maar is gewoon statistiek met attitude?
- Voorbeelden van goed en slecht gebruik.
- Sneltoetsen voor verantwoord gebruik: data-minimalisatie, toon-check, bias-alerts.

Mini-opdracht: deelnemers herschrijven een bedrijfsvoorbeeld met ChatGPT, maar moeten ook aangeven welke risico's ze zelf creëerden (spoiler: vaak meer dan gedacht).

11:00 – 11:15 Pauze

Koffie, thee, of een frisje. AI heeft net een voorspelling gedaan, en 92% van jullie kiest toch iets met cafeïne.

11:15 – 12:00 AVG & Dataveiligheid: AI zonder datalekken

- Wat mag je *wel* en *niet* in ChatGPT of Copilot stoppen?
- Gevoelige versus niet-gevoelige input (en hoe verrassend snel iets gevoelig wordt).
- Dataretentie, modeltraining en privacyverklaringen—verteld zonder juridisch Latijn.
- Organisatiebrede richtlijnen: hoe je voorkomt dat medewerkers per ongeluk de persoonsgegevens van hun schoonmoeder in ChatGPT zetten.

Praktijkcase: deelnemers beoordelen 4 scenario's en bepalen: “veilig”, “risicovol” of “dit gaat ons een boete opleveren”.



12:00 – 12:30 Lunch

Optioneel: AI helpt bij het kiezen van je lunch... maar neemt geen verantwoordelijkheid voor de calorieën.

12:30 – 13:30 Risicomanagement: zo houd je AI onder controle

- Risicoklassen: van laagdrempelige contentcreatie tot high-risk beslissingsondersteuning.
- Hoe je AI-risico's identificeert: bias, onjuiste output, overreliance, datalekken.
- Snelle impactanalyse speciaal voor ChatGPT & Copilot gebruik in organisaties.
- Waarom "het model zei het" geen geldig excuus is (juridisch noch moreel).

Werkvorm: deelnemers creëren een mini-AI-risicomatrix voor hun eigen organisatie of team.

13:30 – 14:00 Break

Even opladen — AI heeft geen pauzes nodig, maar mensen gelukkig nog wel.

14:00 – 15:00 Ethiek & Compliance: de morele kant van slim doen

- De grote vragen: Wanneer is AI handig, wanneer wordt het lui, en wanneer wordt het dubieus?
- Transparantie: vertel je klanten dat je AI gebruikt, of laat je het aan hun intuïtie ("dit klinkt te perfect")?
- Verantwoord prompten: framing, bewuste bias-checks, controlemechanismes.
- Grenzen van autonomie: AI als assistent, niet als beslisser met grootheidswaanzin.

Ethiek-discussie: deelnemers stemmen op stellingen zoals *"Een AI-gegenereerd advies mag direct naar de klant"* (spoiler: nee).

15:00 – 16:00 Ontwerp je eigen AI-beleid (light edition)

- Stap-voor-stap guideline:
 1. Wat mag?
 2. Wat mag niet?
 3. Welke waarborgen gelden?
 4. Wie is verantwoordelijk?
 - Format voor een intern AI-beleidsdocument dat morgen al ingevoerd kan worden.
 - Tips om collega's mee te krijgen (inclusief manieren om weerstand charmant te ontwijken).
-

16:00 – 16:30 | Wrap-up & Takeaways

- Samenvatting van de belangrijkste lessen.
- Q&A—ook de "domme" vragen, want die zijn meestal juist slim.
- Checklist *"Verantwoord AI-gebruik vanaf morgen"*.