

Veilig werken met AI

Een praktisch naslagwerk bij de AI-presentatie: welke risico's je echt loopt bij het werken met AI, en met welke maatregelen je die verantwoord inzet.

ONDERWERP

Dataveiligheid bij AI-gebruik

TYPE

Handout & naslagwerk

DATUM

2 september 2026

DOOR

Bandwerk

— 01 Over Bandwerk

Bandwerk is een creatief digitaal bureau dat organisaties helpt om online sterk voor de dag te komen en slim gebruik te maken van technologie. We werken aan websites, merkbeleving en digitale communicatie. Nieuwe technische mogelijkheden vertalen we naar oplossingen waar onze opdrachtgevers in de praktijk iets aan hebben.

Een steeds groter deel van dat werk draait om software en data. Websites, systemen en klantgegevens hangen tegenwoordig met elkaar samen. Dat vraagt om zorgvuldigheid: niet alleen om iets werkend te krijgen, maar om het veilig en verantwoord in te richten. Dataveiligheid is voor ons geen bijzaak maar een vast onderdeel van hoe we bouwen en adviseren.

De laatste jaren komt daar kunstmatige intelligentie bij. AI-tools maken werk sneller en beter, maar roepen ook een belangrijke vraag op: wat gebeurt er met de informatie die je erin stopt? Bandwerk helpt organisaties om die vraag nuchter te beantwoorden en om AI op een manier in te zetten die past bij de gevoeligheid van hun gegevens. Deze handout vat de kern van die aanpak samen, zodat je na de presentatie een helder naslagwerk in handen hebt.

— 02 Waarom dit onderwerp ertoe doet

Rond AI en veiligheid ontstaan snel twee uitersten. Aan de ene kant *"je kunt er alles in gooien, het is toch handig"*, aan de andere kant *"doe het maar niet, want niets is veilig"*. Beide kloppen niet. De waarheid ligt ertussenin en is bovendien goed te sturen, mits je weet welke risico's je precies loopt en welke maatregelen daar tegenover staan.

De kern in één zin: een honderd procent garantie bestaat niet zolang mensen de tool bedienen. Dat geldt voor elk databeveiligingsvraagstuk, niet alleen voor AI. Denk aan e-mail, USB-sticks of bestandsdiensten. Maar je kunt wel degelijk lagen toevoegen die het risico op menselijke fouten flink verkleinen, in plaats van alleen te vertrouwen op "doe het gewoon niet".

— 03 Algemene regels, ongeacht welke AI-tool

Deze vier vuistregels gelden voor iedere AI-tool en zijn een goede basis om mee te beginnen.

- 1 Ga ervan uit dat alles wat je typt ergens wordt opgeslagen**, tenzij de tool expliciet iets anders belooft en je dat kunt verifiëren.
- 2 Anonimiseer voordat je iets invoert.** Vervang namen, adressen, klantnummers, burgerservicenummers en e-mailadressen door placeholders (bijvoorbeeld "Klant A" of "[BEDRIJFSNAAM]"). Dit werkt in de meeste gevallen prima, ook voor code, contracten of teksten.
- 3 Gebruik zakelijke of API-accounts voor gevoelig werk**, nooit gratis consumentenaccounts. Zakelijke voorwaarden zijn vrijwel altijd strenger op het gebied van training en bewaartermijnen.
- 4 Vermijd het klakkeloos plakken van hele documenten** als het ook met een samenvatting of een geanonimiseerde versie kan.

— 04 De kern: twee risico's die je uit elkaar moet halen

Veel verwarring ontstaat doordat mensen twee heel verschillende risico's op één hoop gooien. Wie ze los behandelt, ziet meteen dat de situatie beheersbaarder is dan ze lijkt.

RISICO 1

Bij de AI-leverancier

Wat gebeurt er met de data die je invoert, wordt die bewaard of gebruikt om het model te trainen? Dit deel is grotendeels goed te beheersen. Met zakelijke afspraken zoals een verwerkersovereenkomst of een Zero Data Retention-overeenkomst garandeert de leverancier contractueel dat er niets wordt bewaard of getraind op jouw invoer. Met techniek en contracten kom je hier richting "zo goed als honderd procent".

RISICO 2

Bij de eigen medewerkers

Menselijke fouten: iemand plakt per ongeluk iets gevoeligs. Dit deel is nooit volledig te dichten, met welke tool dan ook. Je beperkt het met de juiste software, toegangscontrole en bewustwording, maar "garanderen" is hier het verkeerde woord, voor AI net zo goed als voor elk ander communicatiemiddel.

De juiste vraag is dus niet "is AI honderd procent veilig te garanderen", maar "is het risico bij AI-gebruik vergelijkbaar met of kleiner dan bij de tools die we nu al gebruiken". Met de juiste maatregelen kan dat antwoord ja zijn. Dat is een eerlijker en sterker verhaal dan "niets is veilig", want dat laatste klopt simpelweg niet als je de twee risico's niet uit elkaar trekt.

— 05 **Waarom een protocol alleen niet genoeg is**

Een protocol is een afspraak, geen techniek. Het voorkomt niets, het beschrijft alleen wat zou moeten gebeuren. Voor echte garanties verschuif je van "*vertrouwen op gedrag*" naar "*gedrag technisch onmogelijk of onzichtbaar maken voor de gevoelige data zelf*". Bewustwording blijft belangrijk als basis, maar hoe hoger je op de onderstaande ladder komt, hoe minder je afhankelijk bent van de vraag of iemand op het juiste moment aan de regels denkt.

— 06 **Lagen die je kunt toevoegen, van zwak naar sterk**

1

LAAG 1

Bewustwording en protocol

Nuttig als fundament: mensen weten wat wel en niet mag. Niet waterdicht, want het leunt volledig op menselijk gedrag.

2

LAAG 2

Technische anonimisering vóór AI-gebruik

Een tussenstap die gevoelige gegevens automatisch vervangt door placeholders voordat er iets bij een AI-tool terecht komt. Dit kan met eenvoudige scripts of tools die burgerservicenummers, namen en adressen herkennen en maskeren. Zo hoeft niemand er nog aan te denken.

3

LAAG 3

Zakelijke of enterprise-omgeving met beheer

Een beheerder kan functies voor het hele team uitschakelen, bijhouden wat er wordt ingevoerd, de bewaartermijn centraal instellen in plaats van per medewerker en toegang regelen via eenmalige inlog en rechtenbeheer. Dit haalt de beslissing weg bij de individuele medewerker.

4

LAAG 4

Data Loss Prevention (DLP)

Software die op netwerk- of browserniveau meekijkt en ingrijpt zodra iemand gevoelige patronen probeert te plakken in een invoerveld, ook in een AI-chatbot. Een van de meest robuuste lagen, want het grijpt in vóórdat de data de organisatie verlaat, ongeacht of de medewerker het zich realiseert.

5

LAAG 5

Contractuele garanties

Voor echt gevoelige workflows sluit je met de leverancier een verwerkersovereenkomst of een Zero Data Retention-overeenkomst af via het zakelijke aanbod. Dat verplaatst het risico gedeeltelijk juridisch, mocht er onverhoopt iets misgaan.

6

LAAG 6

On-premise of self-hosted modellen

De enige optie die het "data verlaat het pand niet"-probleem echt oplost: modellen lokaal draaien binnen het eigen netwerk. Duurder en doorgaans minder krachtig dan de grote commerciële modellen, maar zonder afhankelijkheid van een externe partij.

— 07 **Voorbeeld: consumenten- versus zakelijk gebruik**

Het verschil tussen accounttypes is geen detail. Ter illustratie de situatie bij een van de grote aanbieders. De principes gelden breder; de exacte voorwaarden verschillen per leverancier en veranderen regelmatig, dus

controleer ze altijd bij de bron.

- **Consumentenaccounts (gratis en betaalde persoonlijke abonnementen)**

Sinds 2025 word je bij dit type account gevraagd een expliciete keuze te maken of je gesprekken gebruikt mogen worden om het model te trainen. Sta je dat toe, dan geldt een lange bewaartermijn (bij deze aanbieder tot vijf jaar). Sta je het niet toe, dan geldt een korte bewaartermijn (in de orde van dertig dagen). Belangrijke nuance: ook als je training uitzet, kan inhoud die door de veiligheidssystemen als risicovol wordt gemarkeerd langer bewaard blijven om misbruik te bestrijden. Dit wordt niet altijd afzonderlijk aan de gebruiker gemeld.

- **Zakelijk gebruik (team-, enterprise- en API-toegang)**

Hier gelden strengere, aparte voorwaarden. Ingevoerde gegevens worden standaard niet gebruikt voor training en de bewaartermijnen liggen lager. Voor de zwaarste gevallen is een Zero Data Retention-overeenkomst mogelijk, waarbij invoer en uitvoer in principe niet worden opgeslagen (op wat de wet of misbruikbestrijding vereist na).

- **Tijdelijke of privémodus**

Sommige tools bieden een modus waarin een gesprek niet wordt bewaard en niet in je geschiedenis komt. Handig voor een eenmalige gevoelige vraag, al geldt ook daar meestal de veiligheidsuitzondering.

De les: kies bewust het juiste accounttype voor het juiste werk. Gevoelig werk hoort in een zakelijke of API-omgeving thuis, niet in een gratis persoonlijk account.

— 08 Concrete stappen

Een praktische checklist om mee te nemen.

- ✓ **Controleer je privacy-instellingen.** Zet bij persoonlijke accounts de trainingsoptie uit als je dat nog niet hebt gedaan.
- ✓ **Gebruik voor werk met klantgegevens een zakelijke omgeving.** Werk je met contracten, financiële gegevens of persoonsgegevens, kies dan een zakelijk account of API-toegang via je organisatie, of gebruik een tijdelijke of privémodus.
- ✓ **Anonimiseer standaard.** Vervang gevoelige gegevens door placeholders voordat je iets invoert, en maak daar waar mogelijk een automatische tussenstap van.
- ✓ **Ruim gevoelige gesprekken op.** Verwijder gesprekken die gevoelige informatie bevatten actief na gebruik.
- ✓ **Zet grenzen richting externe partijen.** Vraag klanten of leveranciers nooit om gevoelige gegevens zoals burgerservicenummers, wachtwoorden of medische gegevens rechtstreeks in een chat te plakken.

— 09 Wat wél waar is en wat níet

WÉL WAAR

Er is geen enkele oplossing die het menselijke element volledig uitsluit. Zolang een mens iets kan kopiëren, plakken, screenshots of doorsturen, blijft er een restrisico. Dat geldt voor AI, e-mail, chat en bestandsdiensten, letterlijk alles waar mensen data doorheen bewegen.

NÍET WAAR

Dat elke AI-optie daarmee ongeveer even onveilig is. Er zit een enorm verschil tussen een gratis tool in de browser en een zakelijke of self-hosted omgeving binnen het eigen netwerk. Het risiconiveau verschilt in de praktijk met een factor tien tot honderd, ook al is geen van beide honderd procent.

— 10 Hoe je dit naar je organisatie of klant framet

Behandel de twee risico's apart en wees eerlijk over beide. Het **leveranciersrisico**, waar gaat de data heen en wordt erop getraind, is met de juiste zakelijke afspraken vrijwel volledig af te dekken. Het **menselijke risico**, iemand plakt per ongeluk iets gevoeligs, is nooit honderd procent uit te sluiten. Maar dat geldt voor elk digitaal systeem dat nu al in gebruik is. Het is dus niet iets unieks aan AI.

De eindvraag is daarom niet "kunnen we AI veilig garanderen", maar "**is het risico bij AI vergelijkbaar met of kleiner dan bij onze huidige tools**". Met de juiste lagen erbij is dat antwoord meestal ja.

AAN DE SLAG MET BANDWERK

Meer uit AI halen, op een verantwoorde manier

Daar helpt Bandwerk organisaties graag bij. We verzorgen presentaties, trainingen en workshops die teams wegwijs maken in AI. We zorgen dat je merk goed vindbaar en citeerbaar wordt in AI-zoekmachines zoals ChatGPT, Perplexity en Google AI Overviews. En we bouwen websites, merkbeleving en content waarin we AI slim inzetten. Loop je met een concrete vraag rond, dan denken we vrijblijvend met je mee hoe je AI in jouw werk toepast.

Neem contact op via bandwerk.nl of mail naar hans@bandwerk.nl om te ontdekken wat er voor jouw situatie mogelijk is.



Bandwerk · merkenbouwers

Deze handout is een algemene toelichting en geen juridisch advies. Voorwaarden van AI-aanbieders veranderen regelmatig; controleer de actuele bepalingen altijd bij de bron.