

FACTSHEET

De AI Act

De AI Act:

AI binnen de lijnen van het recht

Artificiële intelligentie is de drijvende kracht achter de huidige technologische transitie, maar innovatie zonder kaders leidt tot rechtsonzekerheid. De Europese AI Act (Verordening 2024/1689) vormt het juridische kader bij de bouw van een betrouwbare inzet van AI. Zij verbiedt onacceptabele vormen van gebruik, reguleert risicovolle toepassingen en stelt eisen aan transparantie bij interactieve diensten. Bewezen compliance staat daarbij voorop.

AI en gegevensbescherming zijn onlosmakelijk verbonden. Waar de AI Act zich richt op de veiligheid en betrouwbaarheid van het systeem zelf, waarborgt de AVG de rechtmatigheid van de onderliggende dataverwerking. Een verwerking van persoonsgegevens is immers pas rechtmatig als de betrokkene transparant is geïnformeerd en er een geldige grondslag is. Een effectieve implementatie van AI vereist daarom een integrale aanpak.

In deze factsheet wordt het volgende behandeld:

- De classificatie van AI-systemen op basis van de risicoladder
- De rollen in de keten: van de aanbieder (provider) tot de gebruiksverantwoordelijke (deployer)
- De specifieke eisen voor hoog-risico AI, transparantie en generatieve AI-tools
- De raakvlakken met de AVG, met focus op dataminimalisatie en geautomatiseerde besluitvorming
- Het toezichtkader en de sancties bij non-compliance



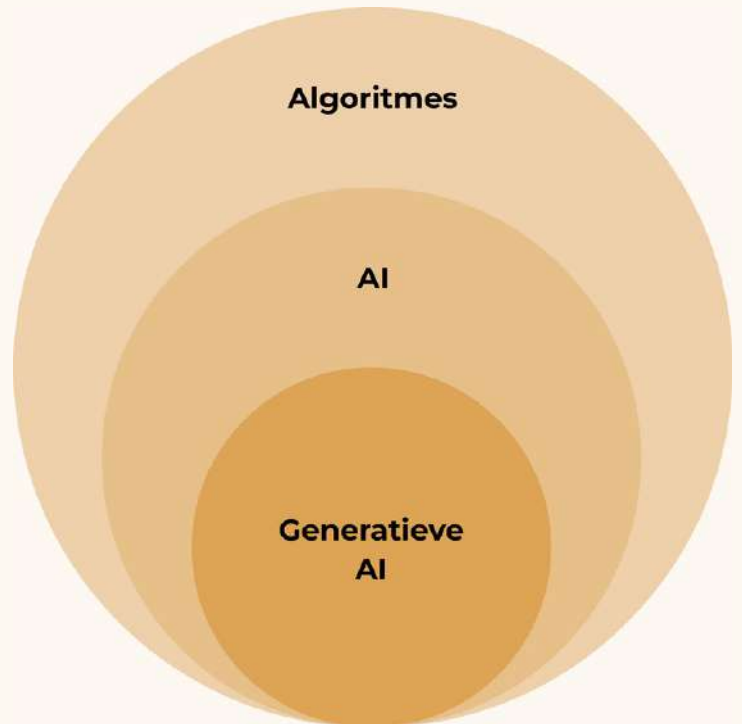
I. Definities:

De hiërarchie van AI

AI, algoritme, taalmodel, slimme software: het vakgebied kent vele termen die vaak door elkaar worden gebruikt. Het meest algemeen is de term algoritme, iedere vastomlijnde set instructies of regels om een taak uit te voeren. Niet elk algoritme is AI. Hiervan spreken we als het algoritme zelf regels afleidt, in AI Act terminologie “inferentie-vermogen” heeft. Op basis van data-analyse worden patronen vertaald naar regels waarmee het systeem nieuwe gevallen kan herkennen of oplossen.

De regels uit de Act zijn van toepassing op wat de Act een AI-systeem noemt. Dit is een AI die haar omgeving kan beïnvloeden, bijvoorbeeld door fysieke actie te nemen. Maar ook beslissystemen (zoals voor werving & selectie of fraude-opsporing) beïnvloeden: zij bepalen hoe mensen moeten handelen.

De populairste vorm van AI is de generatieve AI, die teksten, beeld, video en dergelijke genereren. Voorbeelden zijn de grote taalmodellen of LLMs zoals ChatGPT en Mistral. De AI Act noemt dit General-Purpose AI (GPAI) en behandelt dit als een aparte categorie.



II. Het risicomodel: Categorisering van AI-systemen

Vaak wordt de AI Act omschreven als een piramidemodel: hoe zwaarder een systeem de menselijke waardigheid raakt, hoe restrictiever het regime. Dit is niet helemaal juist. De AI Act kent drie hoofdcategorieën AI waar regels voor gelden. Een aantal toepassingen is verboden, anderen heten “hoog risico”. Overlappend daarmee is de algemene eis dat interactieve AI-systemen transparant moeten zijn over hun status.

Door AI gegenereerde uitvoer (zoals tekst of afbeeldingen) moet als zodanig herkenbaar zijn, helemaal als het gaat om zogenaamde deepfakes. En voor GPAI gelden weer aparte regels, gericht op transparantie en documentatie.

Verboden AI	• Onacceptabele risico's voor gezondheid, veiligheid of fundamentele rechten • Vastgelegd in AI Act (art. 5 AIA)
Hoog risico	• Significante risico's voor gezondheid, veiligheid of fundamentele rechten (art. 6 AIA) • Gereguleerde producten vermeld in Bijlage I • Risicogebieden in Bijlage III met risicotests
Transparantierisico	• Onduidelijkheid over status van AI, gebruik biometrie/emotieherkenning, omgang met deepfakes • Ongeacht status als hoog-risico AI of niet (art. 50 AIA)
General Purpose AI	• Aparte regels voor AI-modellen voor algemene doeleinden (art. 53 AIA) • Extra regels indien sprake van systeemrisico's (art. 55 AIA)



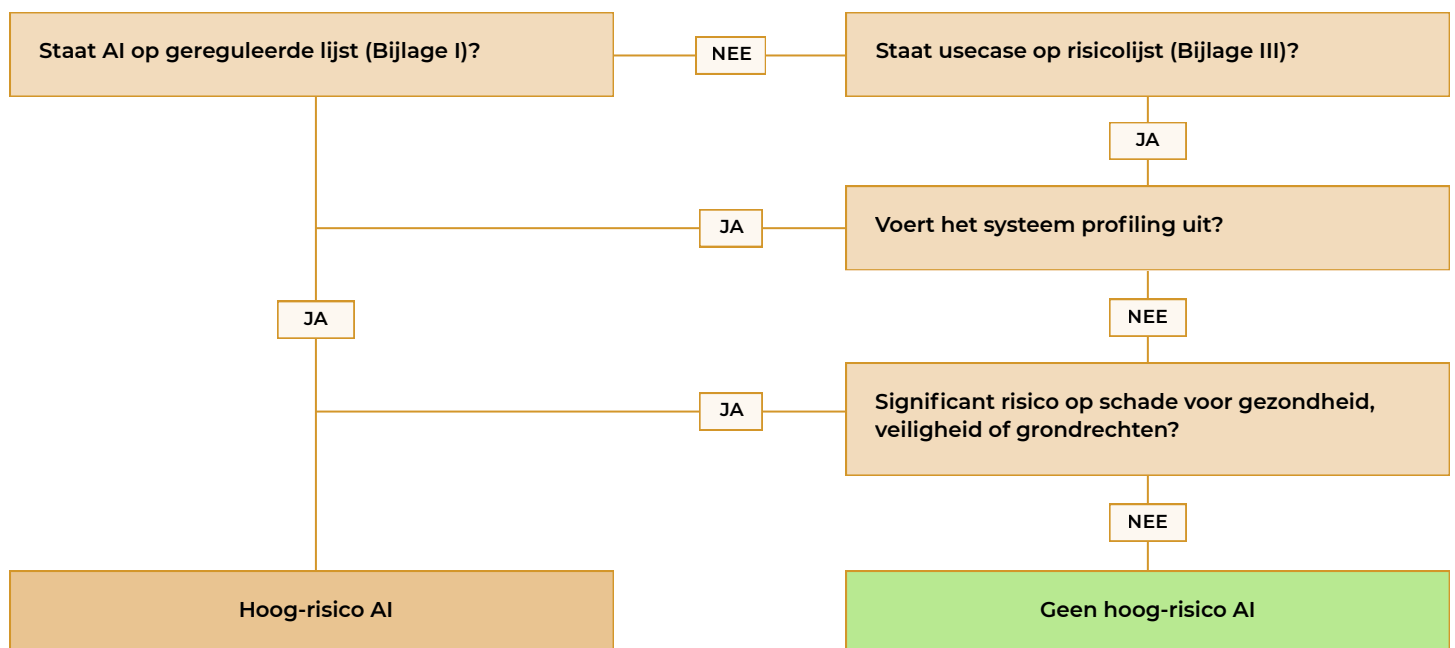
Verboden praktijken

De wetgever trekt een harde lijn bij systemen die de menselijke psyche manipuleren of sociale controle faciliteren. Artikel 5 van de AI Act somt de praktijken op die simpelweg niet welkom zijn op de Europese markt. Hieronder vallen onder meer social scoring door overheden en predictive policing op basis van louter individuele profilering. Deze verboden zijn absoluut: er is geen ruimte voor een belangenafweging. Ook de betrokkenen toestemming vragen heeft geen zin.

Hoog risico AI

Voor systemen in de categorie 'Hoog Risico' geldt een regime van voorafgaande formele toetsing aan een reeks eisen – de zogeheten conformiteitsbeoordeling. Dit gebeurt bij voorkeur aan de hand van officiële Europese standaarden of normen, zoals voor risicomanagementsystemen, bias in gebruikte data en cyberveiligheid. Deze normen worden vanaf eind dit jaar verwacht.

Wanneer is een systeem nu hoog risico? De AI Act kent hiervoor een eenduidig stappenplan. Het is dus niet nodig om zelf na te gaan welke risico's zouden kunnen kleven aan een AI-toepassing en of deze 'hoog' of niet te noemen zijn. De figuur hieronder illustreert het proces.



Bijlage I

AI kwalificeert als product of veiligheidscomponent van product gereguleerd onder genoemde wetgeving:

- Richtlijn veiligheid productiemachines
- Speelgoedrichtlijn
- Medical Device Regulation
- Liftrichtlijn
- Richtlijnen voertuigen, vaartuigen, vliegtuigen
- Richtlijn radio-apparatuur

Bijlage III

Usecase voldoet aan gegeven omschrijving uit Bijlage III, zoals

- Biometrie op afstand en emotieherkenning
- Fysieke veiligheid van kritieke infrastructuur
- Toegang tot onderwijs, beoordelen voortgang en fraudedetectie bij toetsen
- Werving en selectie, beoordeling personeel
- Toegang voor publieke diensten zoals bijstand en gezondheidszorg
- Evalueren van kredietwaardigheid (uitgezonderd detectie van fraude)
- Risico-inschatting bij levens- en gezondheidsverzekeringen
- Ondersteunen van politie en opsporing
- Analyseren personen voor migratie en grenscontrole
- Toepassen van wet of feiten op concrete casus in gerechtelijke procedure

Geen significant risico

1. AI voert alleen beperkte procedurele taak uit
2. AI poetst menselijke resultaten alleen maar op
3. AI signaleert afwijkingen van normale beslispatronen
4. AI voert alleen voorbereidende taak uit

De eerste vraag is kort gezegd of het systeem binnen een van de vele gereguleerde productsoorten valt, zoals medische hulpmiddelen, productiemachines of landbouwvoertuigen. In beginsel is de AI dan hoog risico, maar een uitzondering kan gelden wanneer de betreffende Europese regels de fabrikant toestaan de conformiteitsbeoordeling intern uit te voeren.

In de praktijk zullen de meeste AI-systemen hier niet aan voldoen. Zij komen dan bij de tweede vraag: is de beoogde toepassing genoemd in bijlage III van de AI Act. Deze bevat acht categorieën, en noemt per categorie een of meer specifieke toepassingen. Uitsluitend wanneer de AI voldoet aan zo'n specifieke omschrijving, kan zij hoog risico zijn. Categorie 2 is bijvoorbeeld de kritieke infrastructuur, maar daarbinnen wordt enkel de fysieke beveiliging van bepaalde aspecten van dergelijke infrastructuur genoemd. AI voor cyberbeveiliging van het spoor is dus bijvoorbeeld géén hoog risico, ondanks dat het spoor kritieke infrastructuur is.

Een uitzondering is mogelijk wanneer de aanbieder of ontwikkelaar van het AI-systeem kan onderbouwen dat de toepassing in zijn geval geen significant risico oplevert voor de gezondheid, veiligheid of grondrechten van personen. De AI Act noemt hierbij vier situaties, zoals het enkel uitvoeren van een beperkte procedurele stap of het alleen maar oppoetsen van menselijke resultaten, waarbij dit het geval is. Met een duidelijke (en vooraf opgestelde) motivatie waarom mag de aanbieder dan afzien van de eisen voor hoog risico.

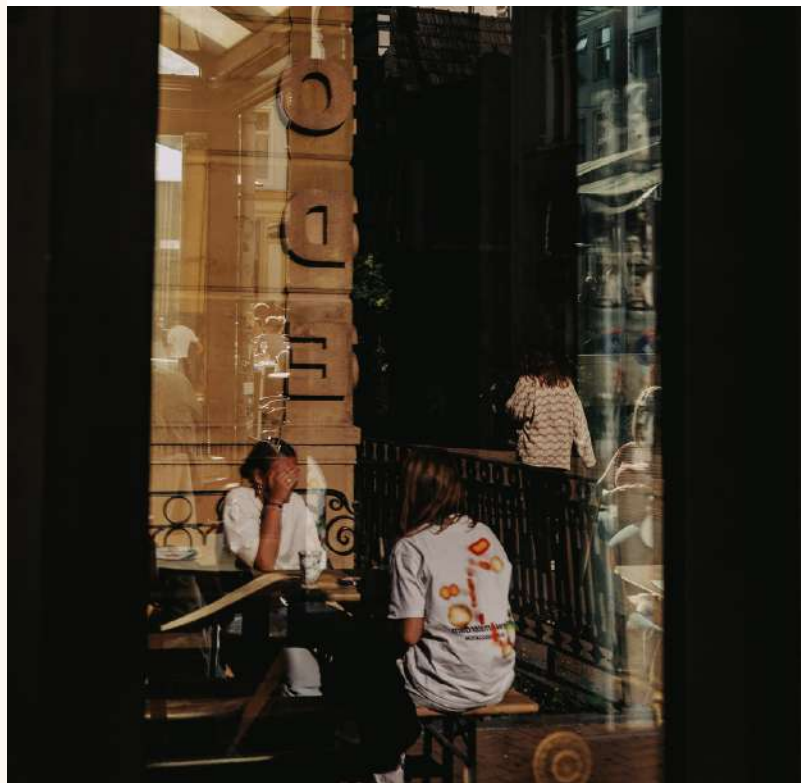
Om het ingewikkeld te maken, geldt die uitzondering dan weer niet als het AI-systeem gebruik maakt van profileren, oftewel het op basis van interesses, persoonskenmerken en dergelijke beoordelen van mensen. Een recruitment-AI die sollicitatiebrieven sorteert op basis van relevantie van het cv kan dus niet "enkel een voorbereidende taak" genoemd worden, omdat dit sorteren een vorm van profileren is. Let op dat het hier niet gaat over beslissen of niet, maar over op basis van profielen handelen in brede zin.

Transparantie, markeringen en uitleg

Ongeacht of zij hoog risico zijn of niet, moeten alle AI-systemen die directe interactie met mensen hebben, duidelijk zijn over hun status als AI. Dit zal meestal spelen bij chatbot-achtige systemen, waarbij mensen zouden kunnen denken met een medemens te spreken. De chatbot moet zichzelf dan als zodanig bekend maken. Als de bot ook beslissingen mag maken, moet dat apart worden gemeld.

Specifiek voor generatieve AI zoals tekst- en beeldgeneratoren geldt nog een aanvullende eis voor hun uitvoer. Deze moet voorzien zijn van markeringen waarmee die als zogenaamde synthetische inhoud herkenbaar is. Bij verder gebruik van die inhoud kunnen die markeringen dan worden opgepikt en gebruikt om bijvoorbeeld een waarschuwing aan de lezer te laten zien. Als de uitvoer poogt een werkelijke situatie na te bootsen (de zogeheten deepfakes) dan moet die markering ook direct zichtbaar in het beeld zitten, zoals bij een watermerk.

Wanneer een hoog-risico AI-systeem een beslissing neemt die een mens raakt, moet de AI ook transparant zijn in het waarom. Mensen hebben dan een recht van uitleg over het hoe en wat rondom die beslissing. Dit recht staat dan weer los van de vraag óf die AI die beslissing mocht nemen. De AVG bevat hierover namelijk een streng verbod, dat slechts in beperkte situaties opgeheven kan worden.

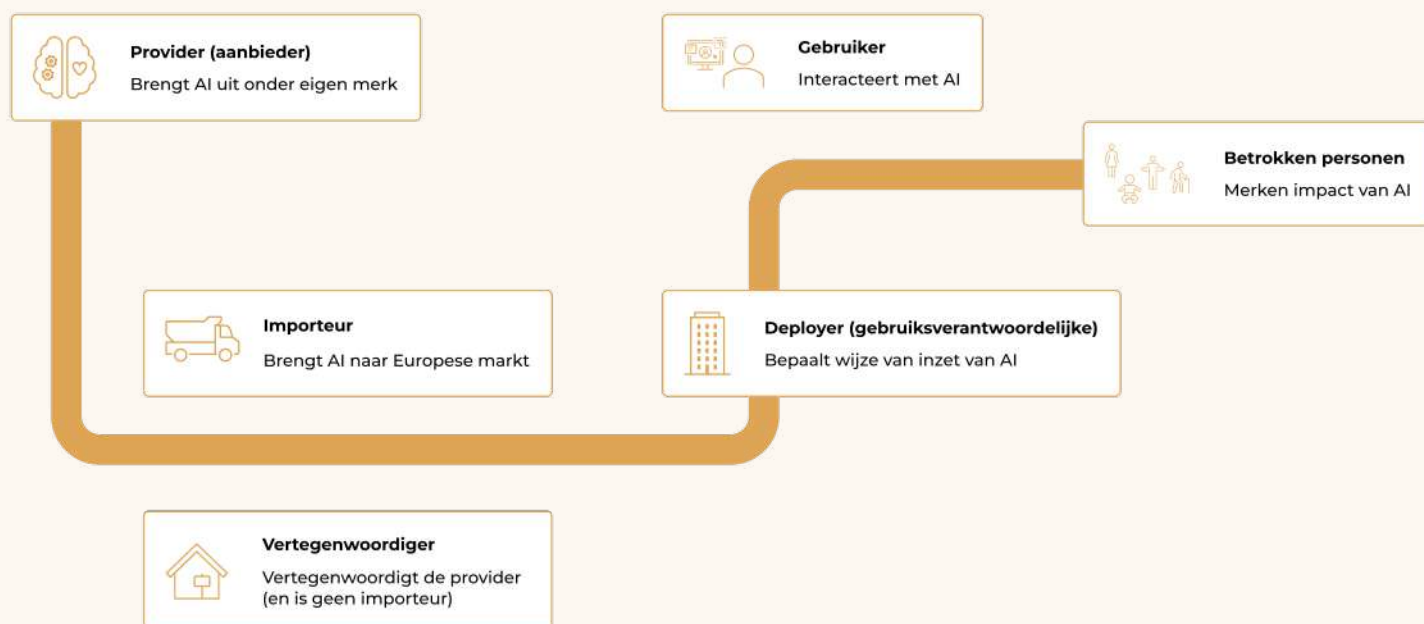


III. De AI-waardeketen: Actoren & Rollen

Van een set uit data afgeleide regels tot een compleet product op de markt komen vraagt de inzet van vele partijen. Deze actoren vormen samen wat de AI Act de waardeketen noemt voor hoog-risico AI. Deze is hieronder geïllustreerd. Elke partij in de waardeketen heeft haar eigen rol met bijbehorende verplichtingen.

De aanbieder heeft de belangrijkste taak: die moet het product de conformiteitsbeoordeling laten ondergaan, en verantwoordelijkheid nemen voor de uitkomst door zijn naam aan de CE-markering te verbinden. Deze staat voor de kwaliteit die men mag verwachten, en impliceert (onbeperkte) aansprakelijkheid als blijkt dat het hoog-risico AI-systeem niet voldoet aan de wettelijke eisen.

De afnemer van het AI-systeem wordt in de Act de gebruiksverantwoordelijke of deployer genoemd. Formeel is dit de entiteit die een AI-systeem onder eigen gezag inzet. Haar taken zijn operationeel toezicht houden en monitoring van incidenten. Zij moet hiervoor goed opgeleide personen inzetten. De gebruiksverantwoordelijke kan de AI direct bij de aanbieder hebben afgenomen, maar ook via een stelsel van distributeurs (of import van buiten de EU) verkrijgen.



IV. General-Purpose AI

De motor van AI-systemen

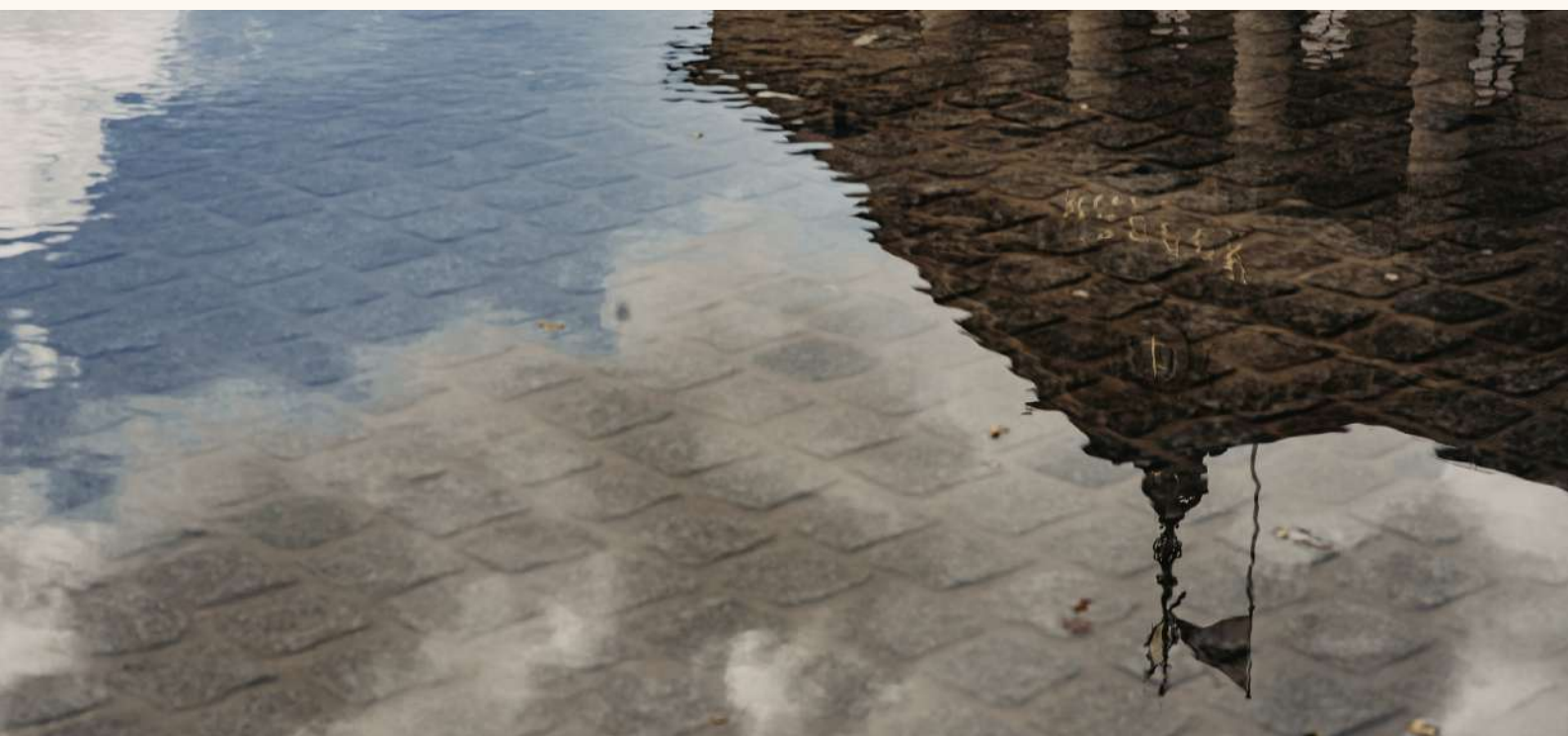
Tijdens het proces waarin de AI Act tot stand kwam, verscheen ineens de dienst ChatGPT. In tegenstelling tot bestaande AI, met specifieke doelen, kaders en toepassingen, kon deze dienst alles: vragen beantwoorden, teksten genereren en problemen oplossen. In haar kielzog doken ook beeld- en geluidgeneratoren op. Vanwege de implicaties die dat alles had, is direct besloten ook deze soort AI te reguleren.

Het kernverschil tussen GPAI en 'gewone' AI-systemen zit hem precies in die afkorting: ze zijn general-purpose oftewel voor algemene doeleinden bestemd. Daarmee is de klassieke toetsing aan hun beoogd gebruik niet mogelijk, zodat evaluaties van hoog risico of verboden praktijken niet zinnig zijn. ChatGPT kán een cv beoordelen, maar is het daarmee een recruitment AI-systeem?

Vanuit het inzicht dat GPAI vaak als motor voor specifieke applicaties wordt ingezet, is gekozen voor een lichter regime van regulering. Als bronleverancier moet de GPAI-aanbieder transparant zijn over de werking, beperkingen en risico's. Dit raakt aan zaken als welke data is gebruikt bij het trainen, welke modaliteiten het systeem heeft en welke beperkingen eraan kleven. In uitgebreide technische documentatie moeten afnemers dit alles kunnen nazoeken en zo hun eigen systeem AI Act compliant ontwikkelen.

Transparantie raakt in het bijzonder aan twee andere wetten: de auteurswet en de AVG. Bij het trainen van zulke grote modellen zijn enorme hoeveelheden data nodig, en daarmee wordt op grote schaal geraakt aan enerzijds auteursrechtelijk beschermd materiaal en anderzijds aan persoonsgegevens van miljoenen mensen. Of dit datagebruik toegelaten is onder deze wetten, is vooralsnog een open vraag, maar vele rechtszaken van auteursrechthebbenden en tientallen onderzoeken door AVG-toezichthouders schetsen een weinig optimistisch beeld.

Tegelijkertijd introduceren GPAI-modellen geheel eigen risico's door hun schaalgrootte en enorme flexibiliteit. De AI Act kent daarvoor het begrip systeemrisico: een risico dat diepe gevolgen voor mens en maatschappij kan hebben. Denk aan GPAI die gemakkelijk helpt bij het ontwikkelen van chemische wapens, of op grote schaal discriminatie bevordert. GPAI met dergelijke risico's moeten zware maatregelen nemen. Maar hoe meet je of een GPAI deze kent? Hier kwam de wetgever niet uit, en gekozen is dan ook voor een numerieke grens aan de rekenkracht benodigd voor het trainen van het model (10^{25} zogeheten FLOPS).



V. AI & Privacy:

De relatie met de AVG

Waar de AI Act de robuustheid van het systeem reguleert, bewaakt de Algemene Verordening Gegevensbescherming (AVG) de integriteit van de menselijke data die erdoorheen stroomt. Een compliant AI-systeem moet dan ook aan beide wetten tegelijkertijd voldoen. Dit kan grote uitdagingen met zich meebrengen waar het gaat om klassieke AVG-eisen zoals transparantie, grondslag en doelbinding.

Impact assessments

Bij toepassingen met 'verhoogd risico' (niet te verwarren met 'hoog risico' onder de AI Act) kent de AVG de eis van de Data Protection Impact Assessment oftewel DPIA. Deze bevat een evaluatie van de risico's voor de bescherming van persoonsgegevens en aanverwante grondrechten van individuen, inclusief voorgestelde mitigerende maatregelen. Voor de inzet van hoog-risico AI-systemen introduceert de AI Act een nieuw instrument: de Fundamental Rights Impact Assessment (FRIA). Hoewel deze sterk doet denken aan de DPIA, zijn er essentiële verschillen.

Allereerst is de FRIA beperkter in wie deze moet uitvoeren. Grofweg geldt de FRIA-eis alleen bij gebruiksverantwoordelijken in de publieke sector, en een paar specifiek genoemde taken in de private sector. Een DPIA moet door iedere verwerkingsverantwoordelijke worden uitgevoerd wanneer daar aanleiding toe is. Een willekeurig bedrijf dat AI voor het beoordelen van personeel wil inzetten, moet dus wél een DPIA maar géén FRIA uitvoeren.

Ten tweede vraagt de FRIA nadrukkelijker om oplossingen en menselijk toezicht, plus een mechanisme voor escalatie en afhandeling van klachten. De DPIA laat in het midden hoe men concreet om wil gaan met gesignaleerde risico's. Daarbij komt ten derde dat de FRIA een bredere, fundamentele blik op risico's vraagt: ook risico's voor maatschappij en milieu moeten meegenomen worden.

Wie zowel een DPIA als een FRIA moet uitvoeren, heeft geluk: de AI Act bepaalt dat dit in één document mag worden uitgevoerd. In de praktijk is dat vaak de zogeheten IAMA (Impact Assessment Mensenrechten en Algoritmen) of de AIIA (AI Impact Assessment), beide door de Rijksoverheid opgesteld en vrij beschikbaar als sjabloon.

Geautomatiseerde besluitvorming en toezicht

De inzet van AI leidt vaak tot een 'black box' van geautomatiseerde besluitvorming: een ja of nee zonder nadere toelichting of uitleg. Zowel de AVG als de AI Act stellen hier grenzen aan. De AI Act bevat een recht van uitleg (zoals hierboven toegelicht), maar de AVG gaat verder: zonder uitdrukkelijke toestemming of een concrete wettelijke regeling is AI laten besluiten over mensen simpelweg verboden.

De veel gekozen uitweg is om een mens in het proces toe te voegen: human-in-the-loop. De AI maakt een conceptbesluit en de mens controleert en stelt bij, zodat het besluit uiteindelijk mensenwerk is geworden. Hoewel dit op papier voldoet, is het risico levensgroot dat de menselijke controle slechts een formaliteit wordt. De controleur moet namelijk wel de middelen hebben om een échte toetsing uit te voeren. Ook moet hij of zij de bevoegdheid hebben om de AI te overrulen, ook als dat regelmatig gebeurt.

Ook de AI Act eist menselijk toezicht. Hoog-risico AI moet te allen tijde onder menselijke controle staan. Maar ook hier geldt dat human-in-the-loop toezicht vaak alleen op papier goed werkt. Onderzoek laat zien dat mensen snel de aandacht verliezen en de neiging ontwikkelen de AI structureel te volgen (automation bias). Een systeem van human-on-the-loop waarbij de mens erbij gehaald wordt als het systeem grenswaarden overschrijdt of ongebruikelijke gevallen tegenkomt, is dan ook veel verstandiger.

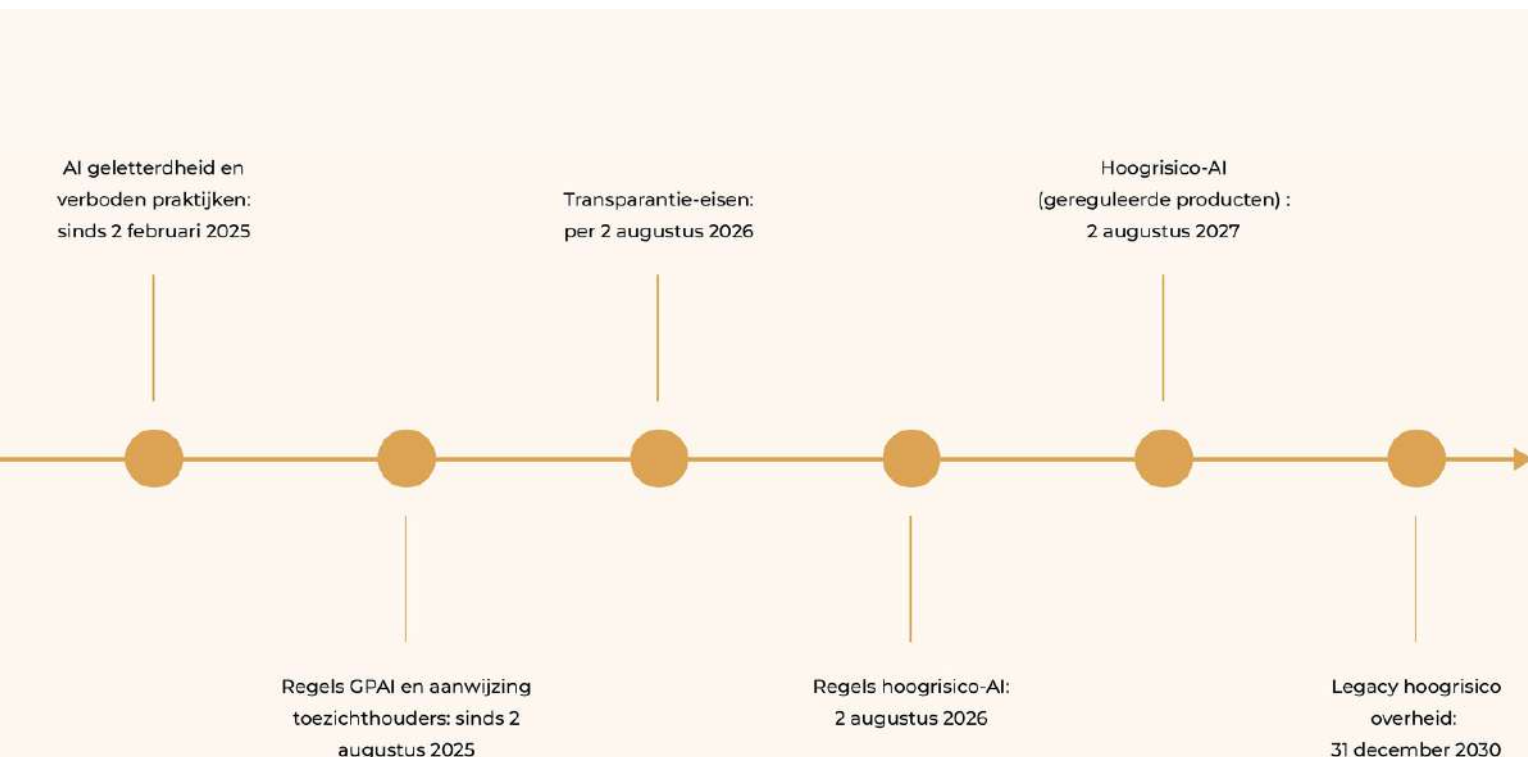
VI. Toezicht & Handhaving:

De tanden van de AI Act

De AI Act introduceert een gelaagd toezichtmodel dat zowel op Europees als op nationaal niveau toeziet op de naleving van de nieuwe regels. Elk land moet één of meer markttoezichthouders aanwijzen, die sectoraal of regionaal toezicht houden op verkoop en gebruik van AI. In Nederland is dit formeel nog niet gedaan, maar het voorstel is dit sectoraal te doen: organisaties als NVWA, IGJ, AFM en RDI houden toezicht op entiteiten die zij nu al reguleren. Waar geen sectoraal toezicht is, zal de Autoriteit Persoonsgegevens de toezichtsrol op zich nemen.

Markttoezichthouders hebben stevige bevoegdheden: zij kunnen onderzoek uitvoeren en mogen daarbij ook panden betreden en vertrouwelijke informatie opvragen. Zij kunnen bevelen een systeem aan te passen of van de markt te halen, of nadere instructies zoals aangepast toezicht geven. Ook is er de bevoegdheid boetes op te leggen. Voor verboden praktijken is dit maximaal 35 miljoen euro (of 7% van de wereldwijde concernomzet), en voor schending van de hoogrisico-plichten 15 miljoen (of 3%).

De Europese Commissie heeft voor zichzelf de rol van toezichthouder op GPAI ingebouwd. Zo kan zij centraal en gericht optreden tegen deze dominante en machtige marktpartijen. Dit toezicht vult bestaande bevoegdheden aan onder wetgeving zoals de Digital Services Act en de Digital Markets Act. Via het nieuwe orgaan van de AI Office wordt dit toezicht uitgeoefend. De AI Office heeft daarnaast een voorlichtende en coördinerende taak.



VII. Praktische vragen

De abstracte normen van de AI Act en de AVG roepen in de dagelijkse praktijk vaak prangende vragen op. Hieronder behandelen wij enkele vaak aan ons gestelde vragen.

Vraag

Is elk algoritme binnen mijn organisatie nu onderworpen aan de AI Act?

Antwoord

Nee. De AI Act beoogt niet de regulering van alle automatisering. Een algoritme dat louter werkt op basis van een vooraf door de mens gedefinieerde 'als-dan-logica', valt buiten het bereik van de AI Act. Pas wanneer het systeem zelfstandig afleidt wat de beslisregels moeten zijn, kan het een AI-systeem zijn. Vervolgens is dan de vraag of het systeem een verboden praktijk uitvoert, als hoog risico heeft te gelden of transparantie-risico's bevat.

Vraag

Mag ik publiek beschikbare data van het internet 'scrapen' om mijn eigen AI-model te trainen?

Antwoord

Dit is een juridisch mijnenveld waar vooral de AVG en de auteurswet een rol spelen. Beide wetten bieden theoretische routes om te verdedigen dat publiek beschikbare data gebruikt mag worden. Er zijn echter minstens zo veel bezwaren tegenin te brengen. Het zal helaas nog enige jaren duren voordat hier definitief duidelijkheid over is.

De AI Act zelf stelt geen nadere eisen. De aanbieder van een GPAI-model moet transparant zijn over de gebruikte trainingsdata, maar bevat geen nadere vrijstellingen of juist verboden.

Vraag

Word ik als gebruiker van een AI-tool van een derde partij (zoals een API) aangemerkt als 'aanbieder'?

Antwoord

Een gebruiksverantwoordelijke van een tool van een ander kan juridisch van status wisselen naar aanbieder wanneer deze het systeem specifiek aanpast om een hoog risico toepassing mee uit te voeren, of het model onder de eigen merknaam op de markt brengt. Een uitgebreide set prompts en data om ChatGPT te laten fungeren als cv-screeningtool kan al maken dat de HR-afdeling de aanbieder wordt van het aldus ontstane hoog-risico AI-systeem.

Vraag

Moet ik een DPIA uitvoeren als ik al een FRIA heb gedaan voor mijn hoog-risico AI?

Antwoord

De FRIA is geen vervanging voor de DPIA, maar staat ernaast. AVG en AI Act staan echter toe dat beide assessments geïntegreerd uitgevoerd worden. Gebruik de FRIA als de overkoepelende analyse en integreer de specifieke privacyrisico's vanuit de DPIA.

Vraag

Moet ik elke door AI gegenereerde tekst of afbeelding expliciet labelen?

Antwoord

Synthetische uitvoer zoals AI-tekst of afbeeldingen moet worden gemarkeerd, zo eist de AI Act. Dit hoeft echter geen zichtbaar label te zijn. De bekende AI-generatoren gebruiken de onzichtbare Content Credentials codes, die machinaal op te pikken zijn. Sociale netwerken zoals LinkedIn en Instagram voorzien de afbeeldingen dan van een gepast logo. Als het gaat om een deepfake (de afbeelding beoogt de echte wereld voor te stellen) dan moet deze wél expliciet gelabeld worden.

Bij tekst geldt een tweeledige eis. Allereerst is de vraag of de tekst bedoeld is het publiek te informeren, zoals bij nieuws of voorlichting. Als dat niet zo is (bijvoorbeeld bij fictie of e-mails), geldt geen labelplicht. Ten tweede is de vraag of er menselijke controle of eindredactie op de tekst is uitgevoerd. Zo ja, dan vervalt de eis tot markering.

Vraag

Mijn AI-systeem is al jaren in gebruik; moet ik dit nu alsnog aanpassen aan de AI Act?

Antwoord

Hoog-risico AI moet aan de AI Act voldoen, ook als het systeem al lang in gebruik is. Hierbij geldt wel overgangsrecht: de eis geldt pas bij de eerste “substantiële wijziging” na 2 augustus 2026. Helaas definieert de AI Act niet formeel wanneer een wijziging substantieel is. Wij adviseren als vuistregel te hanteren dat de markt de wijziging als een nieuwe productversie (2.0 of 3.0) zou zien en niet als een verbetering (2.3 of 3.1).

Als het AI-systeem een verboden praktijk uitvoert, dan moet het systeem onmiddellijk buiten werking worden gesteld. Hierbij geldt geen overgangsrecht.

Vraag

Valt open-source AI ook onder de strenge regels van de verordening?

Antwoord

In beginsel zijn AI-systemen en -modellen die onder een open-source licentie worden vrijgegeven vrijgesteld van de AI Act. Zij mogen echter geen verboden praktijken realiseren (artikel 5). Worden zij ingezet voor hoogrisico-toepassingen, dan gelden de regels alsnog.

Vraag

Geldt de AI Act ook voor wetenschappelijk onderzoek en R&D?

Antwoord

AI-systemen en -modellen die onder de koepel van onderzoek en ontwikkeling worden gemaakt, vallen buiten de AI Act. Pas wanneer zij op de markt worden gebracht of in gebruik gesteld – grofweg: in productie genomen – gelden hierbij de eisen van de AI Act.

Wanneer bij het onderzoek proefpersonen worden ingezet, geldt het specifieke regime van “real-world testing” (artikel 60). De Europese Commissie moet hier nog nadere kaders voor publiceren.

Vraag

Als een AI-systeem AI Act compliant is, is dan ook aan de AVG voldaan?

Antwoord

Nee. De AI Act en AVG bestaan naast elkaar. Het feit dat een AI-systeem alle regels uit de AI Act navolgt, betekent niet dat er geen bezwaren onder de AVG meer kunnen worden gevonden. Ook andere wetgeving, zoals de Cyber Resilience Act of de Digital Services Act, blijft gewoon gelden.



A person is sitting at a table in a cafe at night, looking out a window at a city skyline. A bridge with a red and yellow tower is visible in the background. A lamp is on the table, and a cup of coffee is in front of the person.

Zou je extra informatie willen ontvangen?

Meer weten? Aan de slag met AI? Neem dan contact op met ons:

e-mail: contact@ictrecht.nl

telefoonnummer: 020 663 19 41.