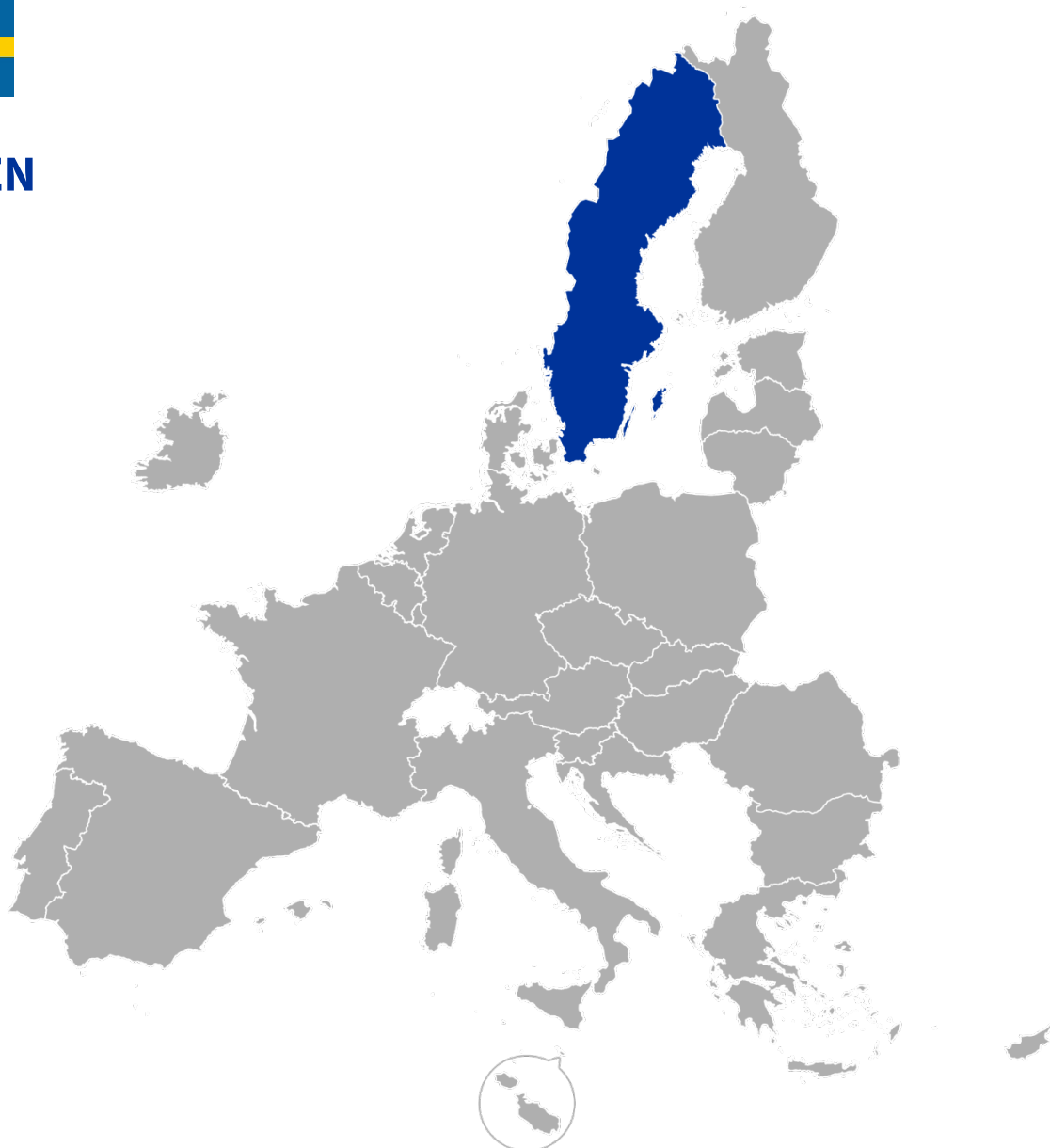




ZWEDEN



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



ZWEDEN

Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Finland, Ierland en Portugal**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Jouw nationale overheid is voorstander van toekomstgerichte maar praktische AI-regelgeving. Jij pleit voor beperkte mogelijkheden om af te wijken van de regels om te voorkomen dat AI-investeerders en -onderzoekers worden ontmoedigd.
- Hoewel Zweden alleen niet het voortouw kan nemen op het gebied van AI, kan een verenigde EU dat wel – door innovatie te stimuleren, normen vast te stellen en haar mondiale invloed te versterken. Er moeten meer stimulansen en minder regelgevingsobstakels zijn voor AI-ontwikkelaars tijdens de ontwikkelingsfasen, wat banen zal scheppen en economische groei zal aanjagen.
- Een doelgebonden aanpak is van cruciaal belang: voor AI die voor medische diagnoses wordt gebruikt, moeten strengere regels gelden dan voor AI die is ontworpen voor de administratie van een bedrijf. Je neemt de bezorgdheid over AI, bijvoorbeeld over mogelijke vooringenomenheid, serieus, maar je wilt voorrang geven aan de meest onmiddellijke risico's die rechtstreeks van invloed zijn op mensenlevens. Met een risico-inschatting op basis van de sector waarin het systeem wordt gebruikt is de kans bovendien kleiner dat de regels botsen met de bestaande sectorspecifieke regels, zoals de kaders voor gezondheidszorg en financiën. Zo blijft de administratieve rompslomp beperkt en wordt de regelgeving niet te ingewikkeld.
- Wat de vereisten voor ontwikkelaars betreft, ben je er sterk van overtuigd dat Europese technologiebedrijven ook AI van hoge kwaliteit kunnen leveren als ze niet aan strenge regels hoeven te voldoen. Als ontwikkelaars zelf mogen bepalen hoe zij modellen testen en verbeteren, zal dit zorgen voor meer innovatie en excellente producten en wordt nog steeds gewaarborgd dat AI op een verantwoorde manier wordt ontwikkeld.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Het is voor jou van essentieel belang dat zaken in verband met de nationale veiligheid, het leger of defensie buiten het toepassingsgebied van de AI-verordening blijven. Deze beleidsdomeinen vallen doorgaans buiten de bevoegdheid van de EU en de lidstaten zijn vastbesloten om de controle te behouden over besluiten op deze gebieden. Dit blijkt ook uit het feit dat de EU geen verenigd EU-leger heeft.
- Daarnaast voer je aan dat voor AI-systemen die worden gebruikt voor het verzamelen van inlichtingen, cyberdefensie of gevechtscenari's niet dezelfde openbaarmakings- of regelgevingsvereisten kunnen gelden als voor civiele systemen. Dat zou de operationele geheimhouding, de nationale veiligheidsbelangen en het strategische voordeel dat zulke systemen moeten bieden namelijk in gevaar kunnen brengen.
- Artikel 2 is voor jou belangrijker dan artikel 1. Als je de anderen niet kunt overtuigen van jouw standpunten over artikel 1, moet je ervoor zorgen dat je belangrijkste wensen met betrekking tot artikel 2 serieus worden genomen.
- Jouw opmerkingen:

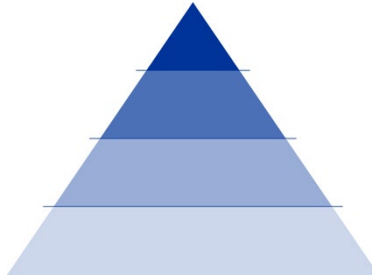
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde training-datasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden	doelgebonden	flexibele benadering		leger/defensie + nationale veiligheid	
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

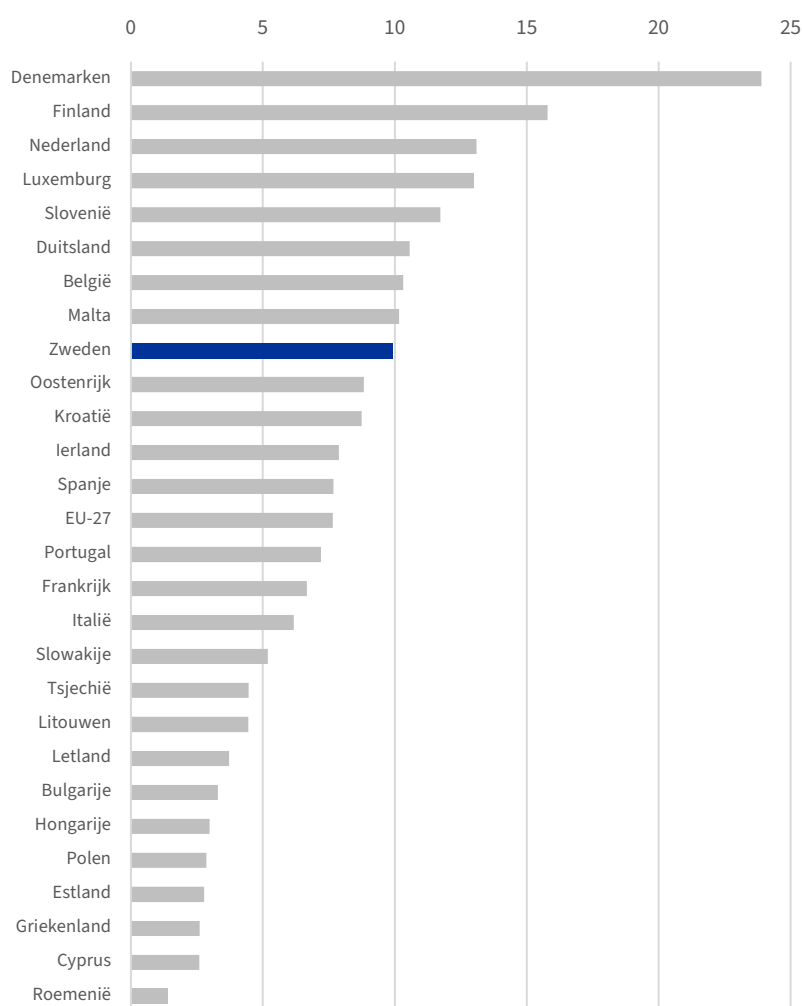
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

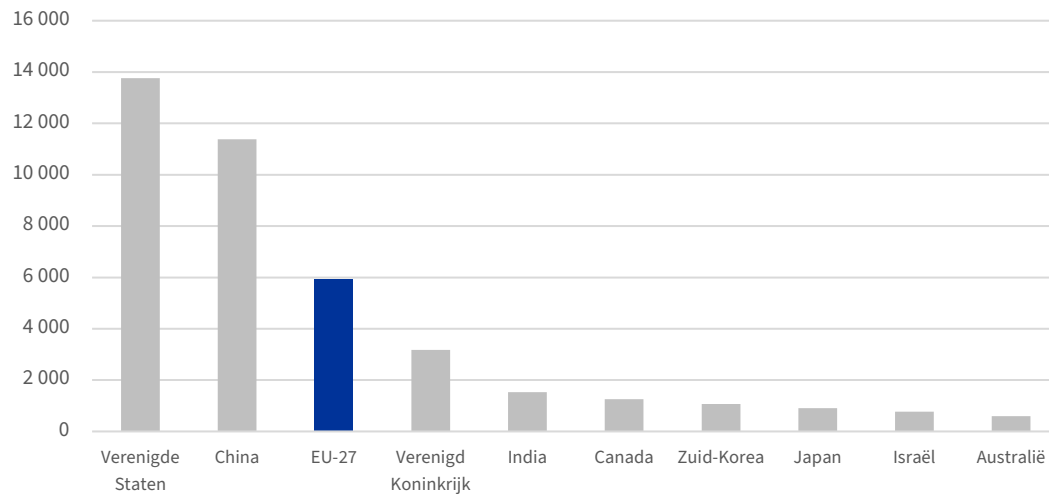
	Europese Unie	Zweden
Bevolking	447 695 350	10 521 556
Bbp per hoofd van de bevolking in EUR	38 150	51 060
Werkloosheidspercentage	6,1%	7,7%
Percentage bedrijven dat AI-technologie gebruikt	8%	10,4%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	36,5%

Percentage bedrijven met meer dan 10 werknemers die minstens één AI-systeem gebruiken (2021)

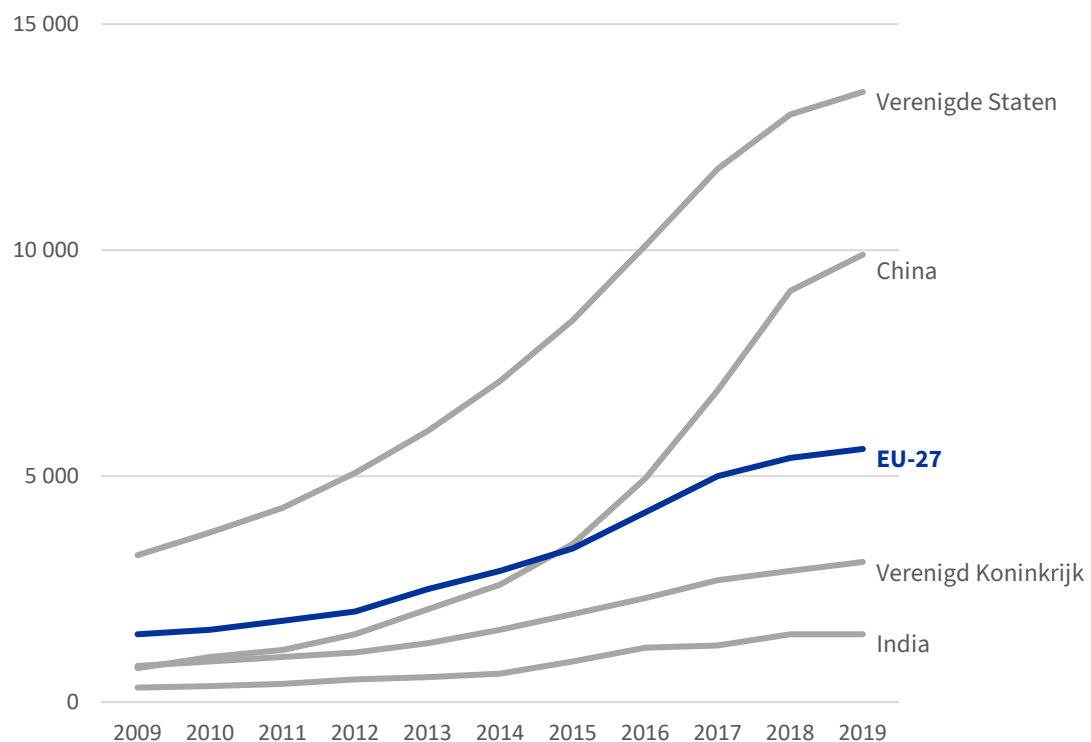


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

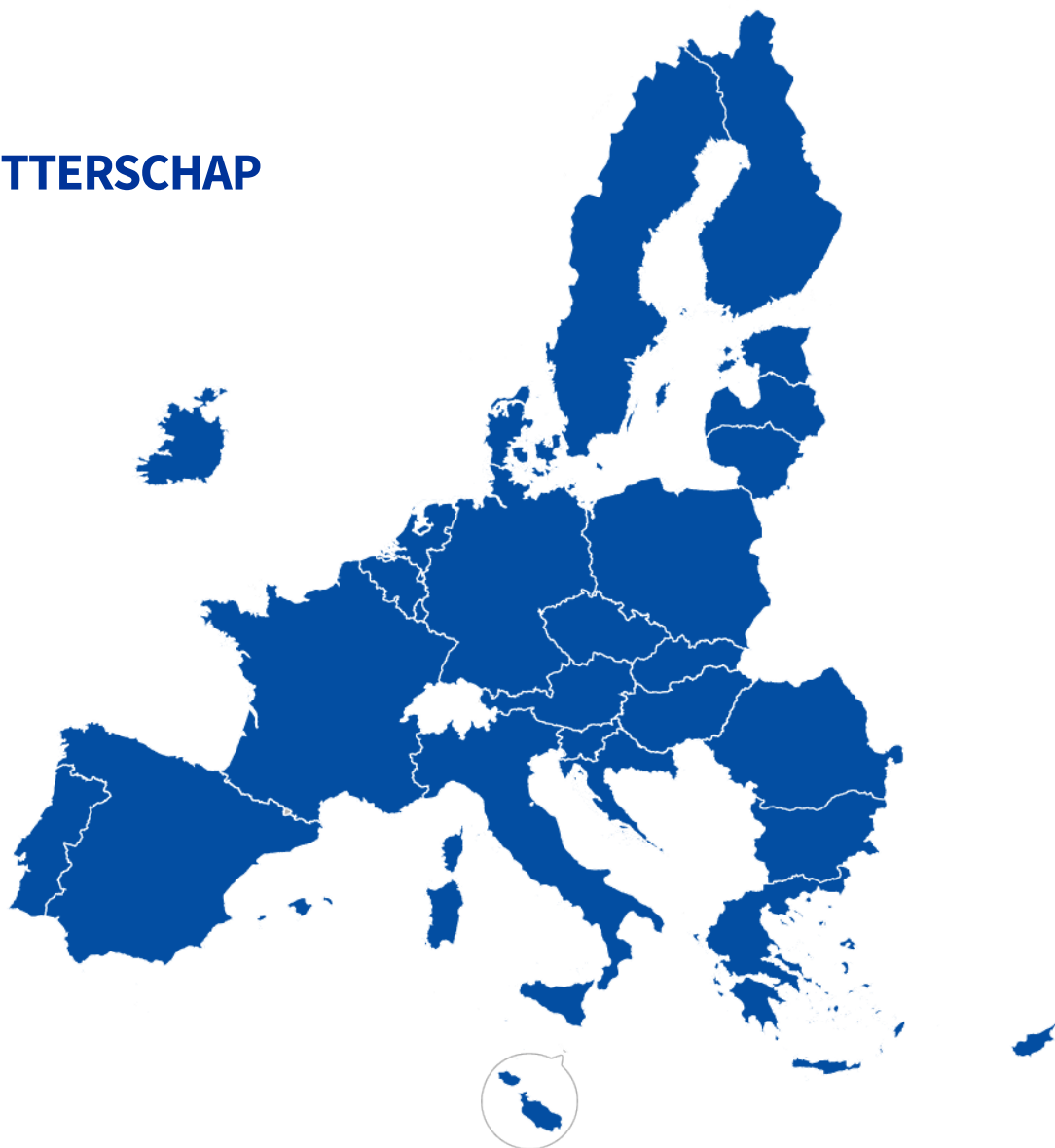
Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC-01-25-006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC-01-25-006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



VOORZITTERSCHAP



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Overeenstemming bereiken over een gemeenschappelijk standpunt van de Raad met een dubbele meerderheid (d.w.z. ten minste 55% van de lidstaten en deze lidstaten moeten 65% van de totale EU-bevolking vertegenwoordigen). Idealiter wordt gestreefd naar unanimiteit, maar als dat niet mogelijk is, wordt de regel van de dubbele meerderheid toegepast. In de praktijk wordt over de meeste besluiten niet gestemd. Het voorzitterschap probeert doorgaans een compromis te vinden dat voor alle delegaties aanvaardbaar is.



Jouw rol

Jij sluit je aan bij de delegatie van de lidstaat die het voorzitterschap bekleedt. De spelleider zal je vertellen welk land dat is.

Als lid van die delegatie is het jouw taak om de officiële onderhandelingen te leiden, terwijl jouw collega's de belangen en standpunten van het land uitdragen. Het is jouw taak om alle officiële gesprekken in goede banen te leiden, de delegatieleden het woord te geven en het tijdschema in het oog te houden.

Openingswoord

1. Jij moet de zitting **formeel openen**, alle delegatieleden en de Commissie welkom heten en hen het doel van de zitting meedelen (maximaal 1 minuut). Het script in het tekstvak op de volgende bladzijde kan je helpen.
2. Daarna **vraag je de vertegenwoordiger van de Commissie om het voorstel toe te lichten**.
3. Vervolgens geef je **de lidstaten** – in de volgorde waarin ze hebben verzocht om het woord te voeren – elk maximaal 1 minuut om een **openingsverklaring** af te leggen. Vraag de lidstaten om duidelijk aan te geven welke gemeenschappelijke belangen zij met andere lidstaten delen.

Besprekingen

1. **Modereer de daaropvolgende officiële onderhandelingen.** De ministers mogen alleen het woord voeren als jij hun het woord geeft. Geef delegatieleden het woord in de volgorde waarin zij hierom hebben verzocht.
2. Verzoek de delegatieleden om geen standpunten en argumenten te herhalen die al naar voren zijn gebracht. Vraag de lidstaten met vergelijkbare standpunten om namens elkaar te spreken. **De tijd is beperkt!**
3. De **officiële onderhandelingen worden onderbroken voor informele gesprekken.** Jij kondigt aan wanneer de informele gesprekken beginnen en wanneer de formele zitting wordt hervat. De informele gesprekken mogen zonder jouw toezicht plaatsvinden.

Bereik een akkoord

1. Tijdens de informele gesprekken heb jij de tijd om **een compromisvoorstel op te stellen.** Je kunt de lidstaten natuurlijk ook vragen om suggesties voor een compromis.
2. Laat de lidstaten stemmen over de meest veelbelovende voorstellen voor de artikelen 1 en 2. Het gemeenschappelijk standpunt van de Raad wordt aangenomen met een dubbele meerderheid, d.w.z. als ten minste 55% van de lidstaten die ten minste 65% van de totale EU-bevolking vertegenwoordigen, ermee instemt. **Gebruik de stemcalculator** om de stemmen te tellen. De spelleider zal je daarbij helpen.
3. Herinner de lidstaten er bij de definitieve stemming aan dat zij op dat moment moeten kiezen tussen enerzijds een compromis dat hen misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.

Leden van de Raad, commissaris,

Het is een genoegen de zitting van de Raad van vandaag te mogen openen.

Zoals u weet, bekleedt _____ momenteel het voorzitterschap van de Raad en zal het de onderhandelingen leiden. Mijn taak zal zijn de vergadering te modereren, en dus niet de standpunten van ons land ter zake te presenteren.

Op de agenda staat een nieuwe verordening over het gebruik en de ontwikkeling van artificiële intelligentie, of AI, in de EU. Dit onderwerp is van groot belang voor alle Europese burgers en voor de rol van de EU in de wereld, omdat _____.

Vooraleer ik het woord aan de Commissie geef, wil ik erop wijzen dat de Juridische Dienst van de Raad het voorstel heeft bestudeerd en ons heeft laten weten dat het strookt met de Verdragen en de vereiste rechtsgrond heeft.

Wij verzoeken de Commissie nu vriendelijk haar voorstel te presenteren. Daarna krijgt ieder van u de gelegenheid om het standpunt van uw land in een korte openingsverklaring toe te lichten. Wij kijken uit naar een constructief en vruchtbaar debat.

Commissaris, u heeft het woord.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



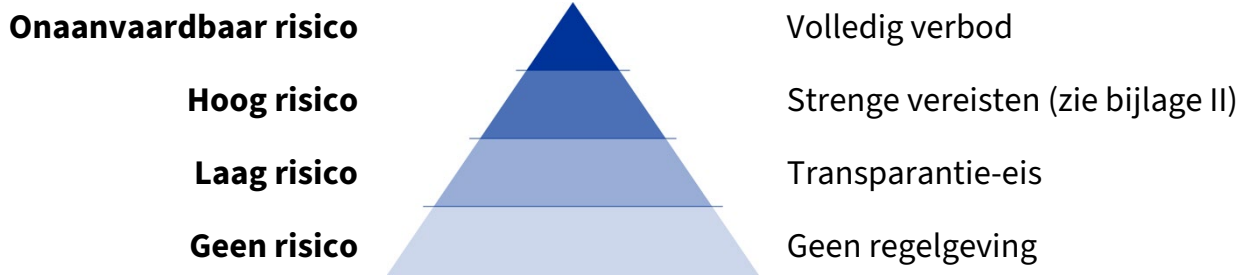
... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

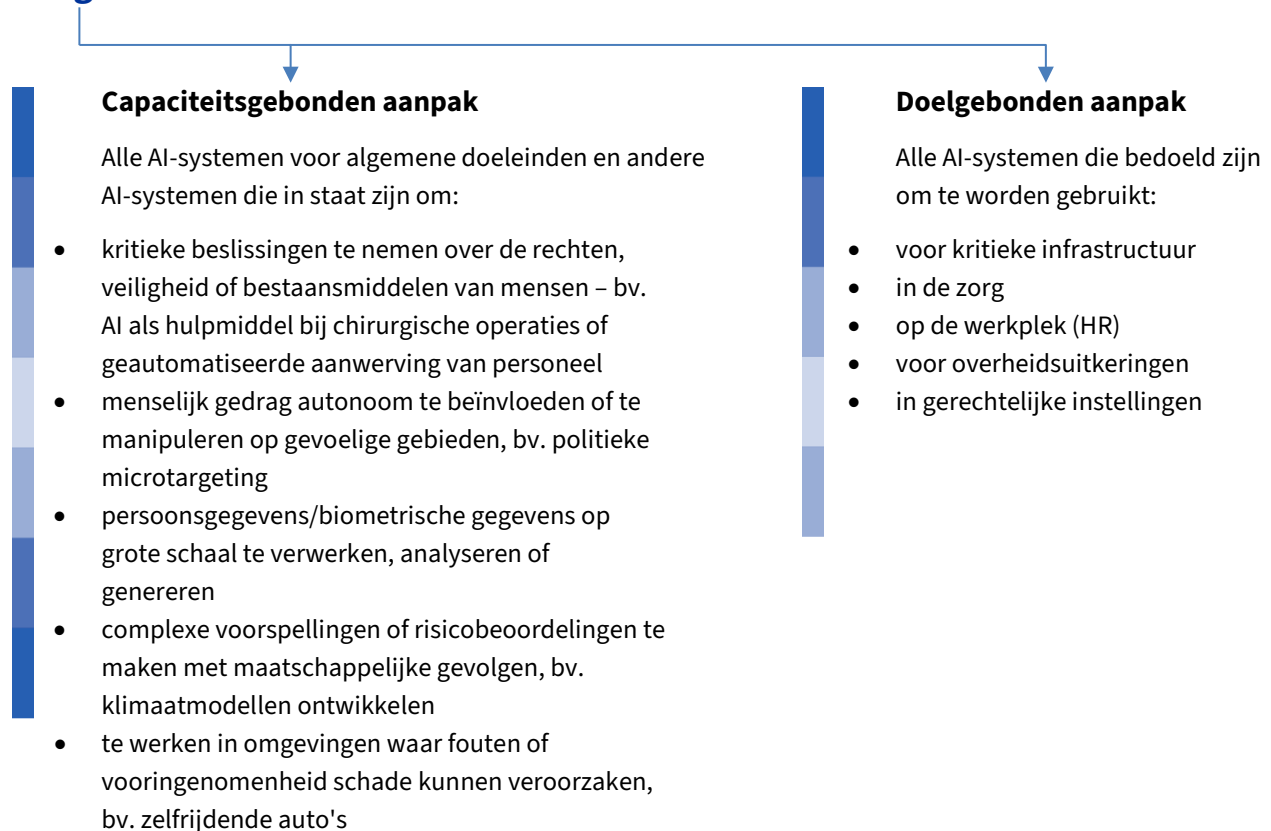
Bijlage I: Risicoclassificatie



Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk	doelgebonden	voorzorgsbenadering		nationale veiligheid	
België	capaciteitsgebonden	voorzorgsbenadering		opensource	
Bulgarije	doelgebonden	flexibele benadering		zonder uitzonderingen	
Kroatië	doelgebonden	voorzorgsbenadering		start-ups	
Cyprus	capaciteitsgebonden	voorzorgsbenadering	over 2 jaar opnieuw bekijken	zonder uitzonderingen	
Tsjechië	capaciteitsgebonden	flexibele benadering	over 5 jaar opnieuw bekijken	zonder uitzonderingen	
Denemarken	doelgebonden	voorzorgsbenadering	financiële steun	leger/defensie + nationale veiligheid	
Estland	doelgebonden	flexibele benadering	testen onder reële omstandigheden	leger/defensie + nationale veiligheid	
Finland	doelgebonden	flexibele benadering	de stemming uitstellen	start-ups	
Frankrijk	doelgebonden	voorzorgsbenadering	onderwijs toevoegen	leger/defensie + nationale veiligheid	
Duitsland	capaciteitsgebonden	voorzorgsbenadering		zonder uitzonderingen	
Griekenland	capaciteitsgebonden	voorzorgsbenadering	over 3 jaar opnieuw bekijken	leger/defensie + nationale veiligheid	
Hongarije	capaciteitsgebonden	flexibele benadering		leger/defensie	
Ierland	doelgebonden	flexibele benadering	flexibiliteit bij het testen	start-ups	

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië	capaciteitsgebonden	voorzorgsbenadering		start-ups	
Letland	capaciteitsgebonden	flexibele benadering	controle na toelating op de markt schrappen	leger/defensie	
Litouwen	capaciteitsgebonden	flexibele benadering		zonder uitzonderingen	
Luxemburg	doelgebonden	flexibele benadering	over 3 jaar opnieuw bekijken	start-ups	
Malta	doelgebonden	flexibele benadering	testen onder reële omstandigheden	nationale veiligheid	
Nederland	capaciteitsgebonden	voorzorgsbenadering		nationale veiligheid	
Polen	doelgebonden	voorzorgsbenadering		leger/defensie	
Portugal	doelgebonden	flexibele benadering	de stemming uitstellen	opensource	
Roemenië	doelgebonden	flexibele benadering		nationale veiligheid	
Slowakije	doelgebonden	flexibele benadering	financiële steun	start-ups	
Slovenië	capaciteitsgebonden	voorzorgsbenadering		start-ups	
Spanje	capaciteitsgebonden	voorzorgsbenadering		leger/defensie + nationale veiligheid	
Zweden	doelgebonden	flexibele benadering		leger/defensie + nationale veiligheid	
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken.
Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



Definities

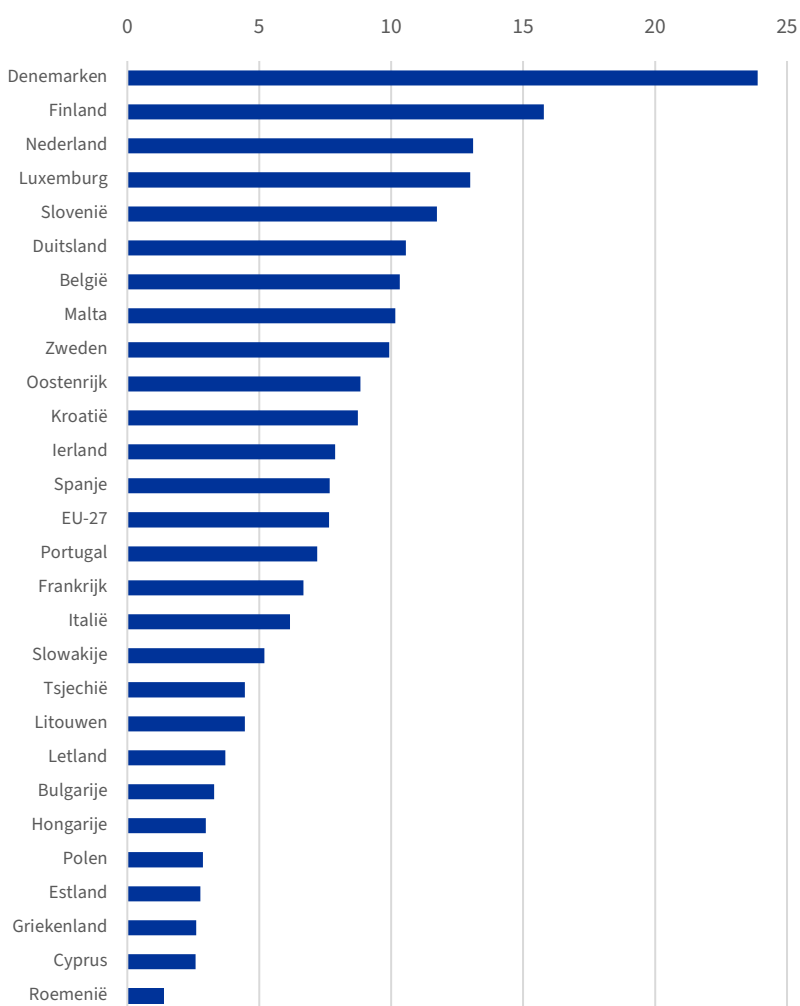
- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

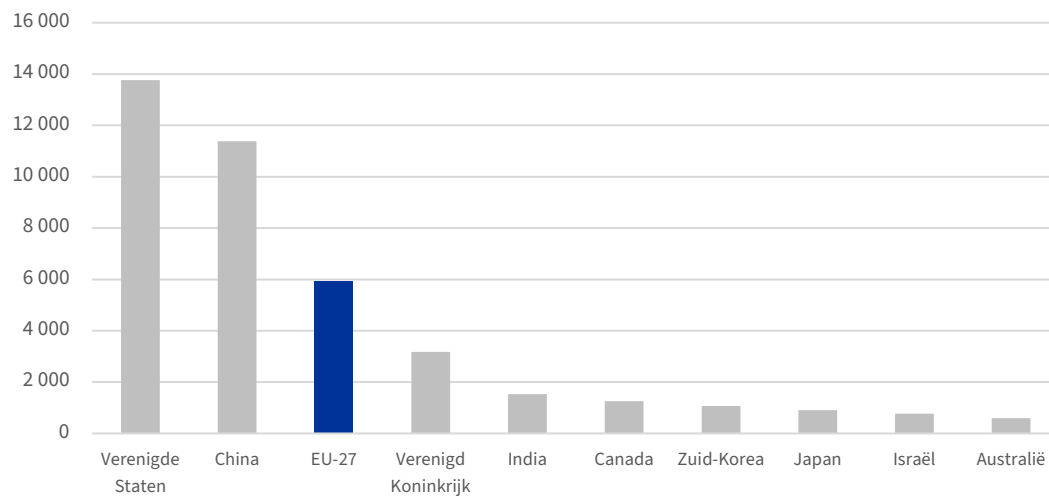
	Europese Unie
Bevolking	447 695 350
Bbp per hoofd van de bevolking in EUR	38 150
Werkloosheidspercentage	6,1%
Percentage bedrijven dat AI-technologie gebruikt	8%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%

Percentage bedrijven met meer dan
10 werknemers die minstens één AI-systeem
gebruiken (2021)

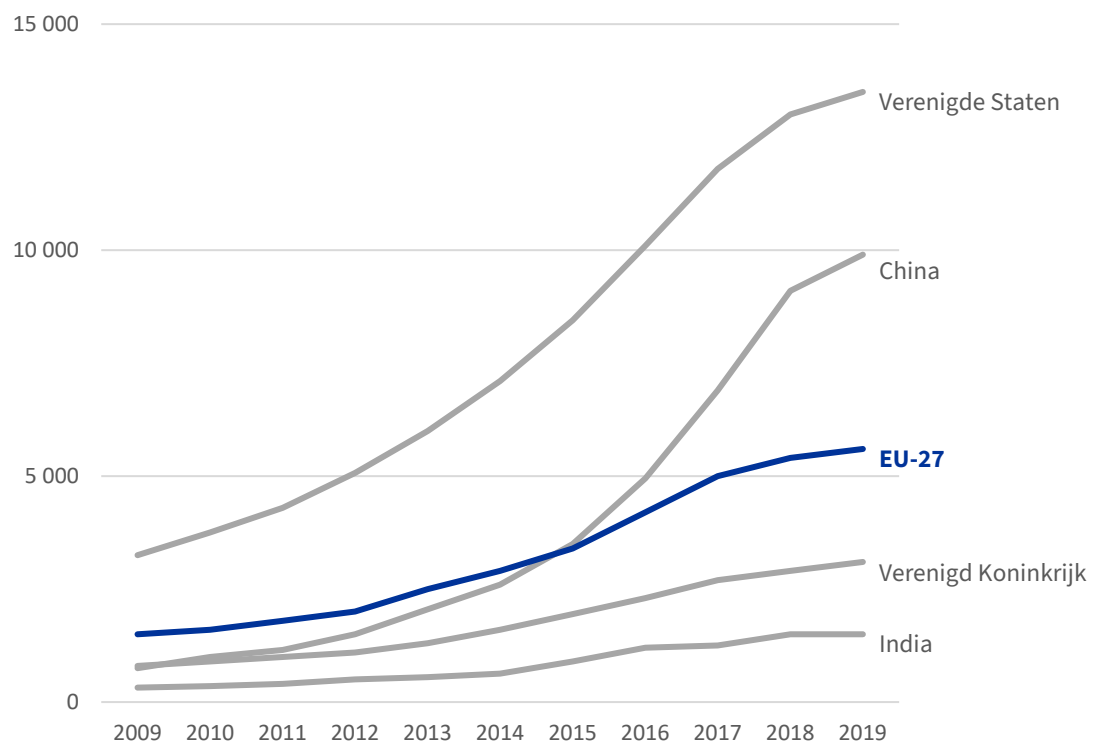


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

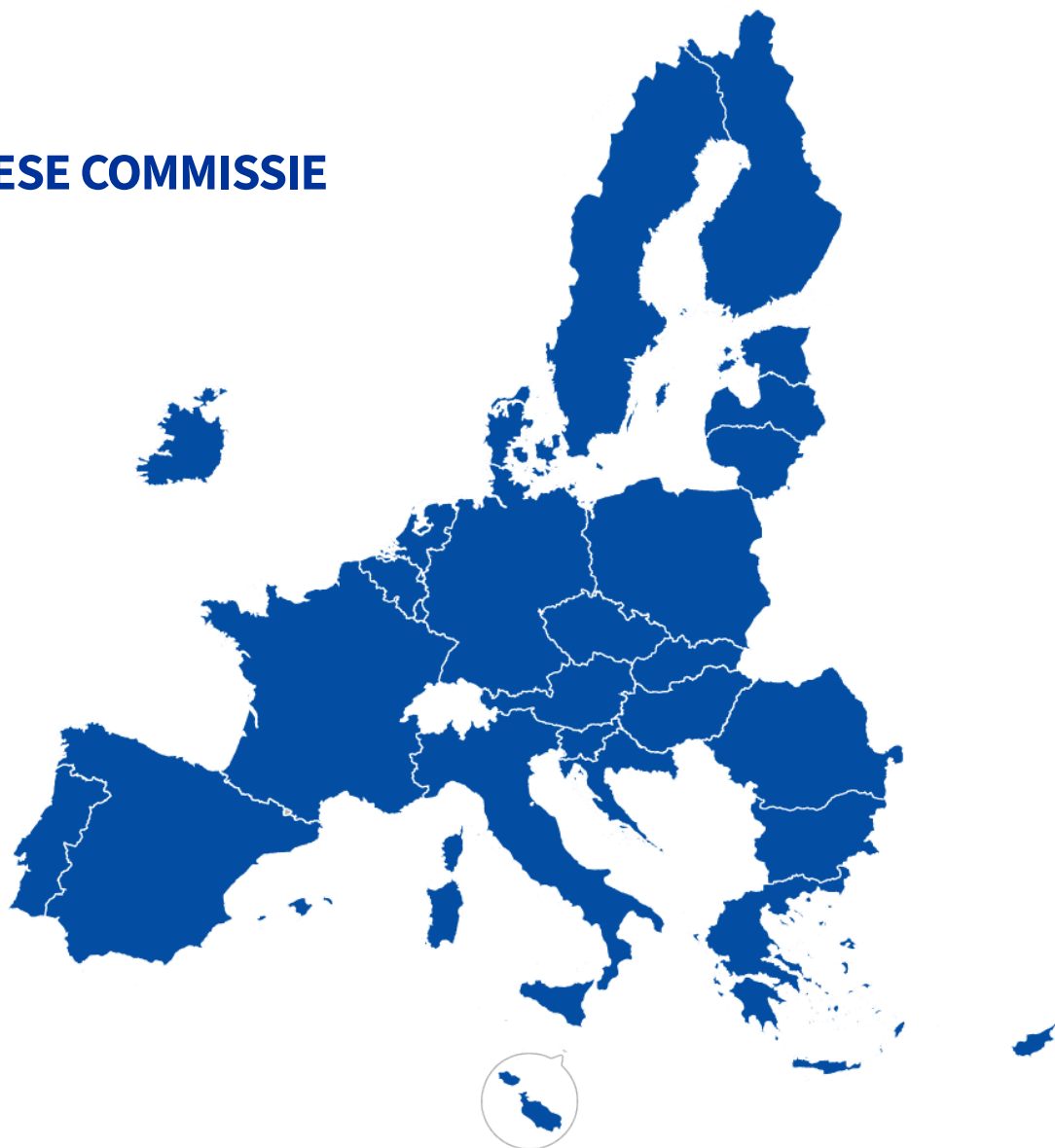
Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC-01-25-006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC-01-25-006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



EUROPESE COMMISSIE



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



Agenda

Vandaag staat jouw voorstel voor een juridisch kader voor het gebruik van AI op de agenda van de Raad.

Wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft, presenteer jij je voorstel aan de vertegenwoordigers van de lidstaten en neem je aan de onderhandelingen deel.

Jouw hoofddoel

Ervoor zorgen dat er in de definitieve versie zo veel mogelijk van de oorspronkelijke tekst overblijft, en de Raad helpen om een gemeenschappelijk standpunt te bepalen. Er moet dringend een gemeenschappelijke Europese oplossing worden gevonden!



Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. **Stel een openingsverklaring op** van maximaal 1 minuut om jouw voorstel te presenteren en toe te lichten. Het script in het tekstvak rechts kan je helpen.
3. Luister aandachtig naar de standpunten van de lidstaten. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben.
4. Leg jouw standpunten uit en **verdedig jouw voorstel** tijdens de onderhandelingen. Neem het woord en help de Raad om een akkoord te bereiken dat voldoet aan jouw verwachtingen!
5. **Jij hebt geen stemrecht** aan het einde. Benut de informele pauzes om delegaties te overtuigen van jouw voorstellen.

Geachte leden van de Raad,

Het is goed dat wij vandaag dit belangrijke debat houden over de verordening over artificiële intelligentie in de EU. De verordening zal de burgers en de EU als geheel ten goede komen, omdat _____

*In onze ontwerpverordening stellen wij voor een **doelgebonden risicoclassificatie** te hanteren voor AI-toepassingen met een hoog risico en **uit voorzorg testvereisten** te stellen, aangezien _____*

Uitzonderingen op de in artikel 1 bedoelde bepaling zijn toegestaan voor AI-systemen die zijn ontwikkeld voor militaire of defensiedoeleinden, omdat _____

Laten wij als gemeenschap ons deel doen zodat AI op een verantwoorde manier wordt ontwikkeld.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Er zijn dringend maatregelen nodig om innovatie te bevorderen en het mondiale concurrentievermogen van in de EU ontwikkelde AI te versterken. Tegelijkertijd moeten de eventuele risico's voor de grondrechten en de consumentenbescherming zorgvuldig worden afgedekt.
- Artificiële intelligentie is een nieuwe technologie met onvoorspelbare gevolgen voor individuen, de samenleving, bedrijven en economieën. Zowel ambitie als behoedzaamheid zijn dus geboden. Daarom heb je in jouw voorstel gekozen voor een doelgebonden aanpak om te bepalen welke AI-systemen moeten worden geclassificeerd als systemen met een "hoog risico". In de doelgebonden aanpak wordt ervan uitgegaan dat AI een hoog risico met zich meebrengt op gevoelige terreinen waar wordt besloten over leven en dood of waar er sprake is van een aanzienlijke impact op het welzijn van de burgers van de EU.
- De Raad heeft al overeenstemming bereikt over de categorieën "onaanvaardbaar risico", met systemen die in de EU moeten worden verboden, "laag risico" en "geen risico". Vandaag wordt gezocht naar een duidelijke definitie voor de categorie "hoog risico". AI-systemen die in deze categorie vallen, mogen niet worden behandeld als een systeem met een laag risico of zonder risico.
- Tijdens de besprekingen van vandaag probeer je ook duidelijke vereisten vast te stellen voor AI-systemen met een hoog risico. Het kader voor tests uit voorzorg zorgt ervoor dat AI op verantwoorde wijze en met verantwoordingsplicht wordt ontwikkeld. Tegelijkertijd worden ontwikkelaars door dergelijke normen uitgedaagd om hoogwaardige, betrouwbare systemen te ontwikkelen, waarmee wereldwijd succes kan worden geboekt. De verwachting is dat de EU zich met deze aanpak zal ontpoppen tot leider op het gebied van AI en dat zij de kloof met de VS en China snel zal dichten.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)

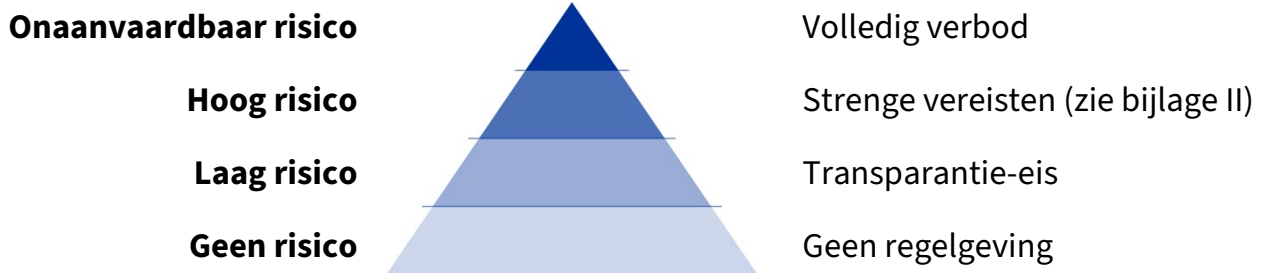


... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Hoewel vereend digitaal leiderschap in de EU en het versterken van de externe handelsdimensie – zoals de wereldwijde uitvoer van in de EU ontwikkelde AI-producten – belangrijke prioriteiten blijven, is de nationale autonomie in specifieke sectoren die van cruciaal belang zijn voor de openbare veiligheid niet-onderhandelbaar.
- Zo zien we bijvoorbeeld dat wereldwijd de internationale veiligheidsuitdagingen razendsnel toenemen en dat AI een belangrijke rol speelt in geavanceerde cyberoorlogvoering, malafide hackingsoftware enz. Defensieve tegenmaatregelen mogen niet gehinderd worden door regelgeving, hoe minimaal ook. Daarom mag artikel 1 niet van toepassing zijn op militaire en defensietoepassingen.
- Voor alle andere doelstellingen zijn geen verdere vrijstellingen nodig, aangezien minimale testvereisten ervoor zorgen dat AI op verantwoorde wijze wordt ontwikkeld.
- Jouw opmerkingen:

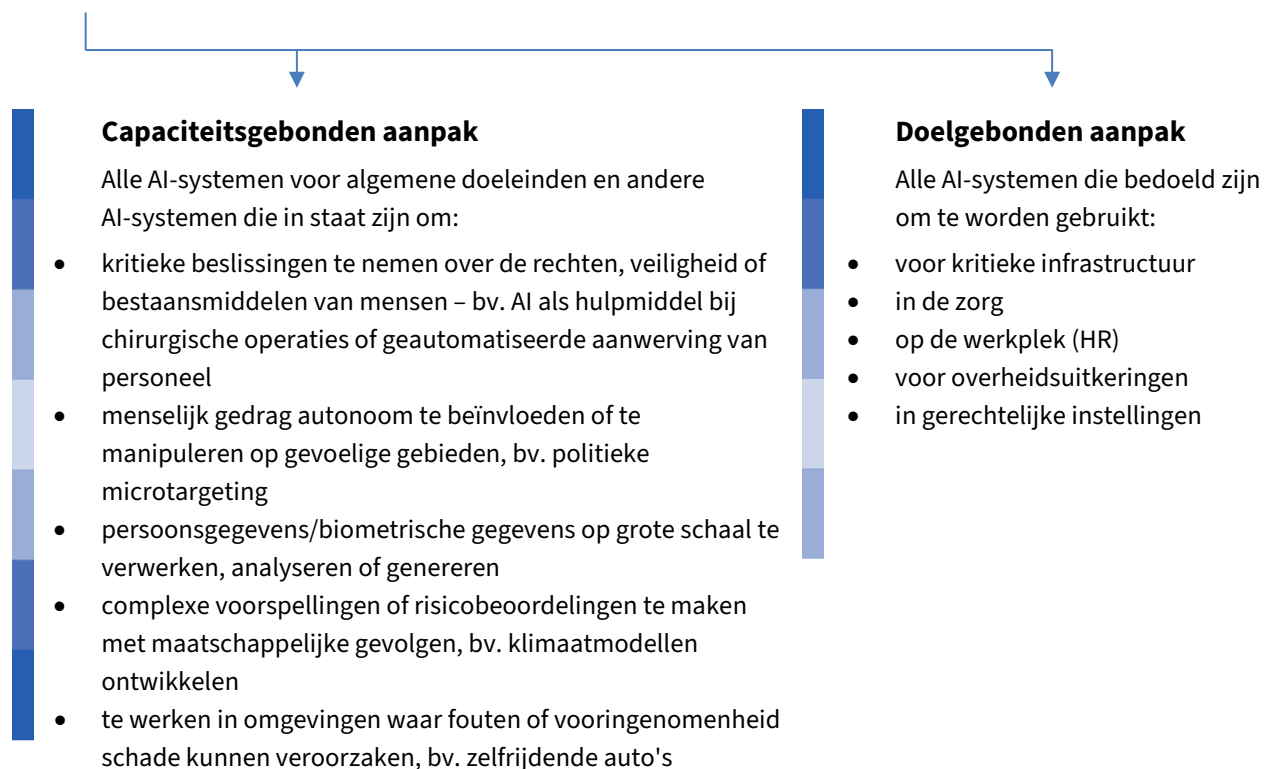
Bijlage I: Risicoclassificatie



Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde training-datasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



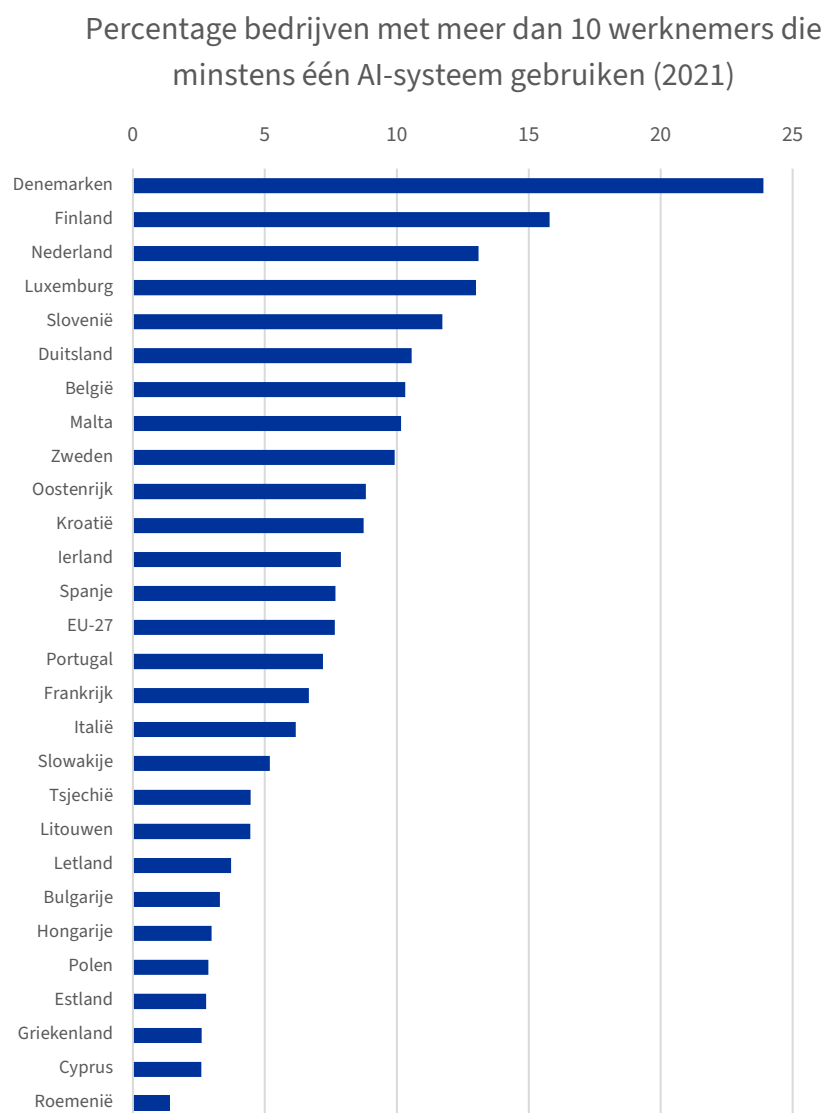
Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

Achtergrondinformatie

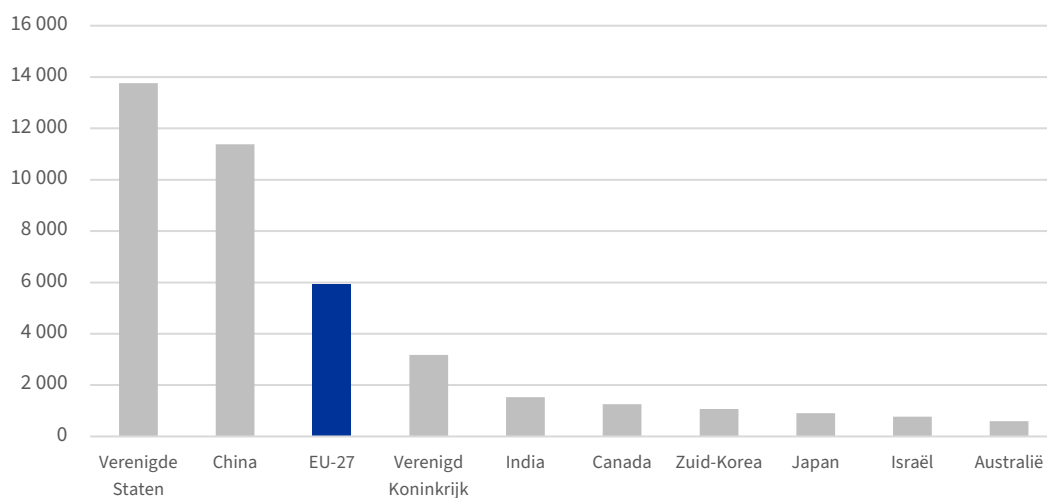
Let op: **De onderhandelingen vinden plaats in 2023!**

	Europese Unie
Bevolking	447 695 350
Bbp per hoofd van de bevolking in EUR	38 150
Werkloosheidspercentage	6,1%
Percentage bedrijven dat AI-technologie gebruikt	8%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%

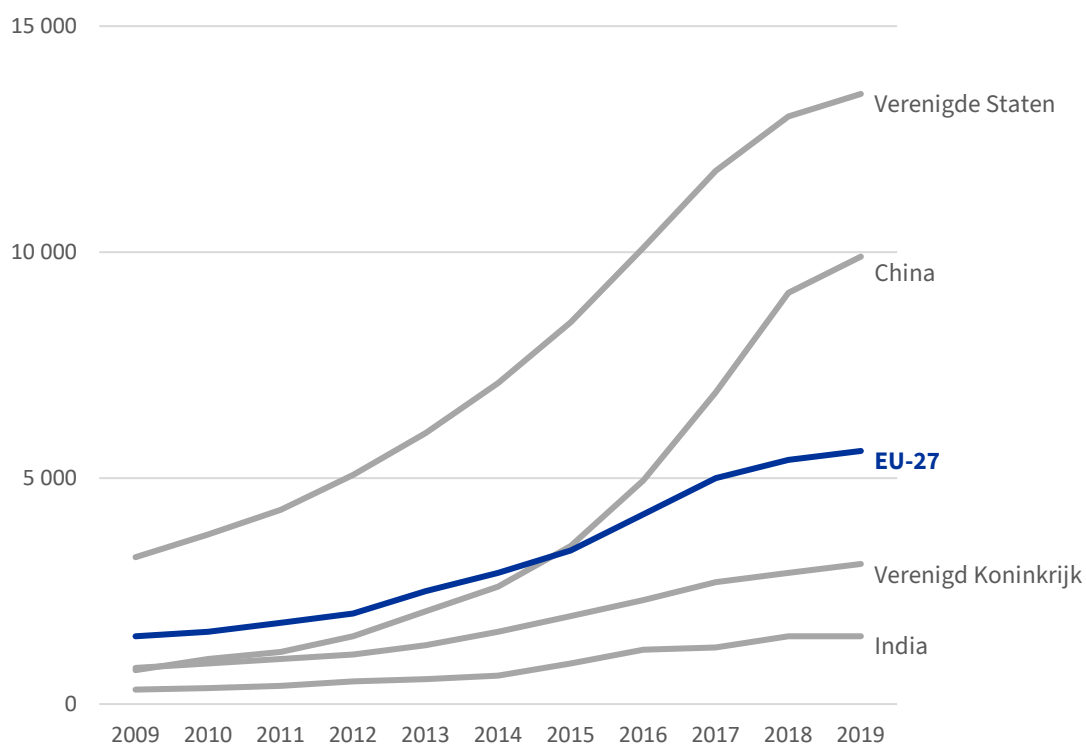


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

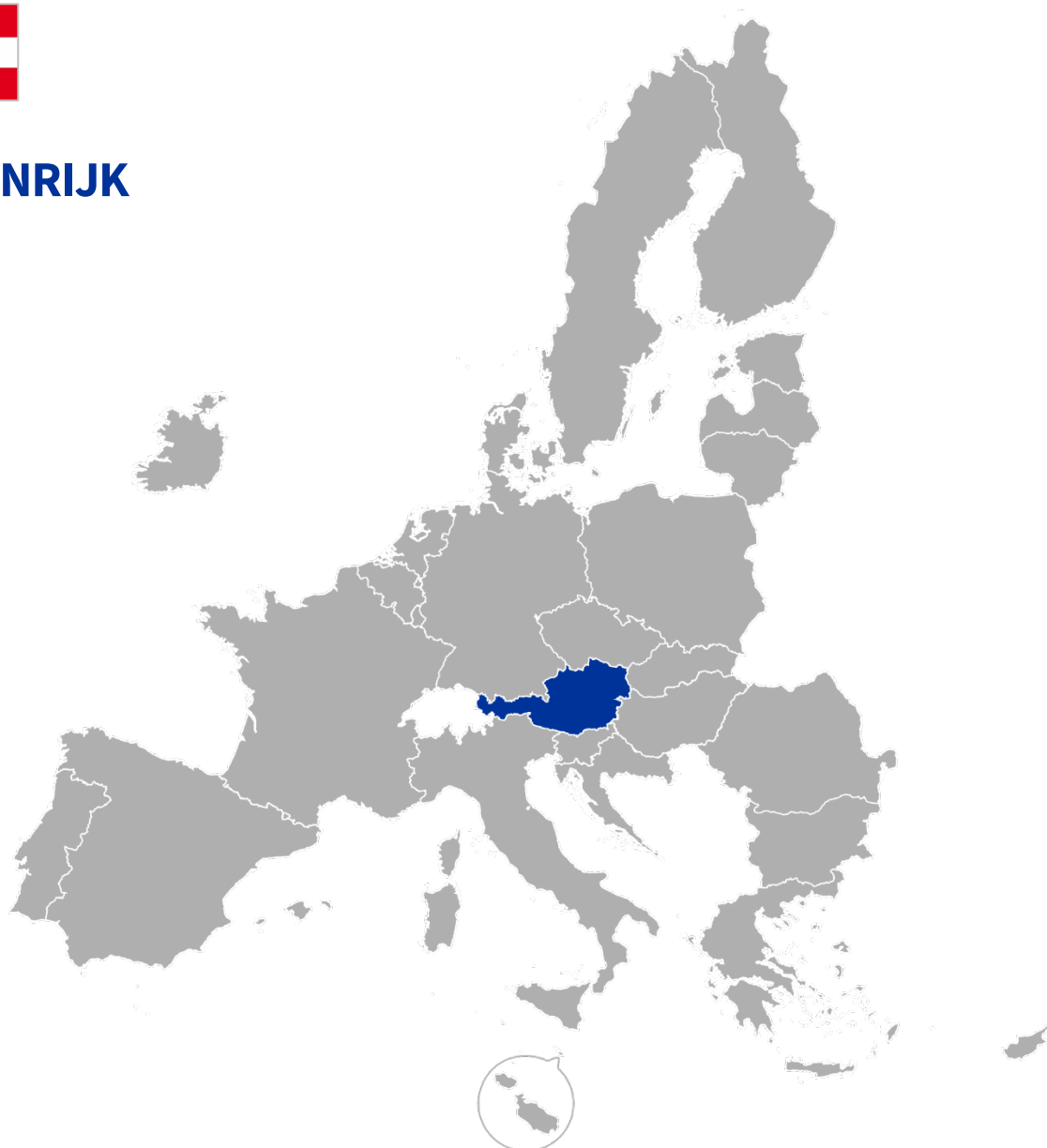
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



OOSTENRIJK



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Denemarken, Frankrijk, Polen en Roemenië**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Jouw nationale regering heeft tussen 2012 en 2017 bijna € 350 miljoen in AI-onderzoek geïnvesteerd. Daarmee is er een solide basis gelegd voor innovatie. Oostenrijkse bedrijven zijn nu koplopers in de ontwikkeling van AI, wat gunstig is voor zowel de binnenlandse als de internationale markt. Je steunt de koers van de EU op het gebied van AI-governance, maar dringt wel aan op meer financiering op EU-niveau om onderzoek en ontwikkeling verder te stimuleren.
- Je bent sterk voorstander van een doelgebonden aanpak conform het beginsel van technologieneutraliteit. Dit betekent dat de regelgeving geen specifieke AI-technologie, -modellen of -benaderingen mag bevoordelen of benadelen. In een gevoelige sector, zoals de gezondheidszorg, maakt het niet uit welke vorm van AI precies wordt gebruikt. Of het nou gaat om deep learning of op regels gebaseerde AI, het gebruik van alle vormen van AI in deze sector wordt geacht een hoog risico met zich mee te brengen, omdat er iemand kan overlijden als de AI het begeeft of een fout maakt.
- De doelgebonden aanpak is ook van essentieel belang voor de invoer van AI. Het is vaak onmogelijk om uitsluitend aan de hand van de verklaring van de ontwikkelaar te controleren wat een geïmporteerd systeem werkelijk allemaal kan doen, vooral wanneer buitenlandse bedrijven beweegredenen hebben om schadelijke elementen, zoals spyware, te verbergen.
- Wat de testvereisten betreft, stem je in met een kader voor tests uit voorzorg, omdat dat een cultuur van verantwoordelijkheid bevordert en het makkelijker maakt om te achterhalen waar er iets mis is gegaan als er zich problemen voordoen.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- In samenwerking met Interpol, Europol en het Bureau van de Verenigde Naties voor drugs- en misdaadbestrijding (met hoofdkantoor in Wenen) speelt het Austrian Institute of Technology een toonaangevende rol in het gebruik van AI-systemen voor de nationale veiligheid door middel van geavanceerde surveillanceoplossingen en immersieve opleidingsplatforms. Onderzoekers werken aan zeer betrouwbare AI-modellen die Oostenrijk veiliger moeten maken. Aangezien nationale veiligheid meestal buiten het toepassingsgebied van de EU-regelgeving valt, vind je het verstandig om dit domein buiten de verordening te laten.
- Daarnaast is het, zoals je bij artikel 1 al hebt benadrukt, belangrijk om te verduidelijken dat het toepassingsgebied van het regelgevingskader ook AI-systemen moet omvatten die buiten de EU zijn ontwikkeld en vervolgens zijn geïmporteerd. De verordening moet van toepassing zijn op alle AI-systemen die in de Europese Unie worden gebruikt, ongeacht waar ze vandaan komen.
- Jouw opmerkingen:

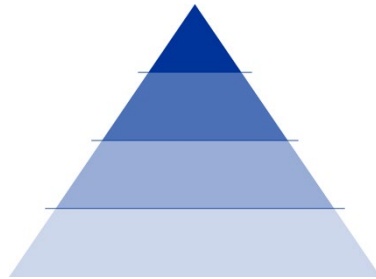
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico

Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk	doelgebonden	voorzorgsbenadering		nationale veiligheid	
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

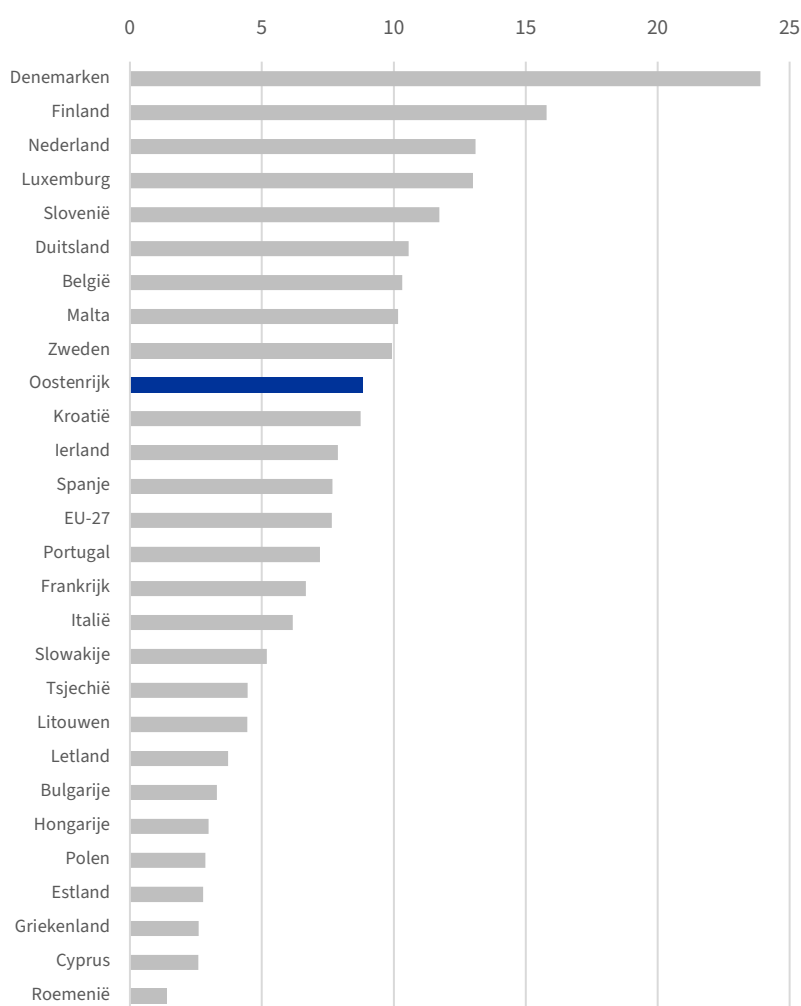
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

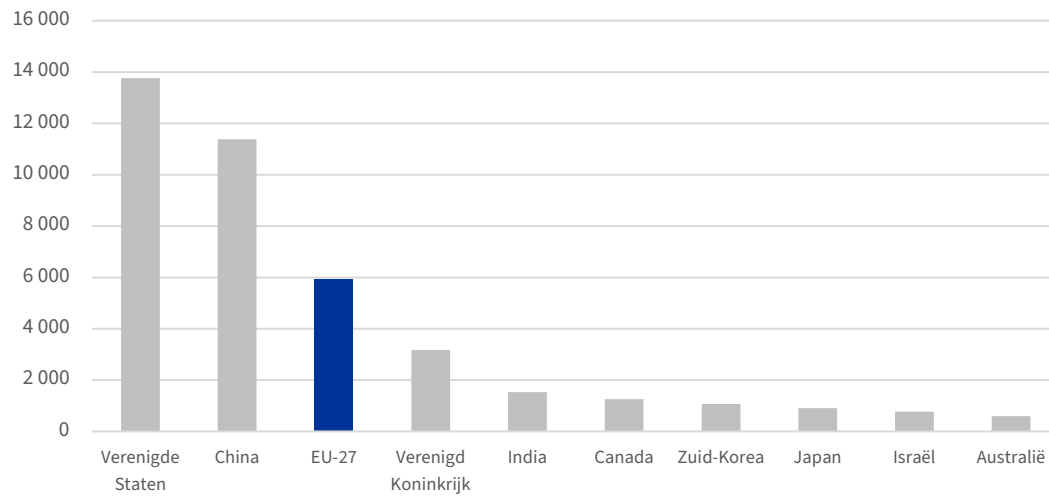
	Europese Unie	Oostenrijk
Bevolking	447 695 350	9 104 772
Bbp per hoofd van de bevolking in EUR	38 150	51 830
Werkloosheidspercentage	6,1%	5,1%
Percentage bedrijven dat AI-technologie gebruikt	8%	10,8%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	32%

Percentage bedrijven met meer dan
10 werknemers die minstens één AI-systeem
gebruiken (2021)

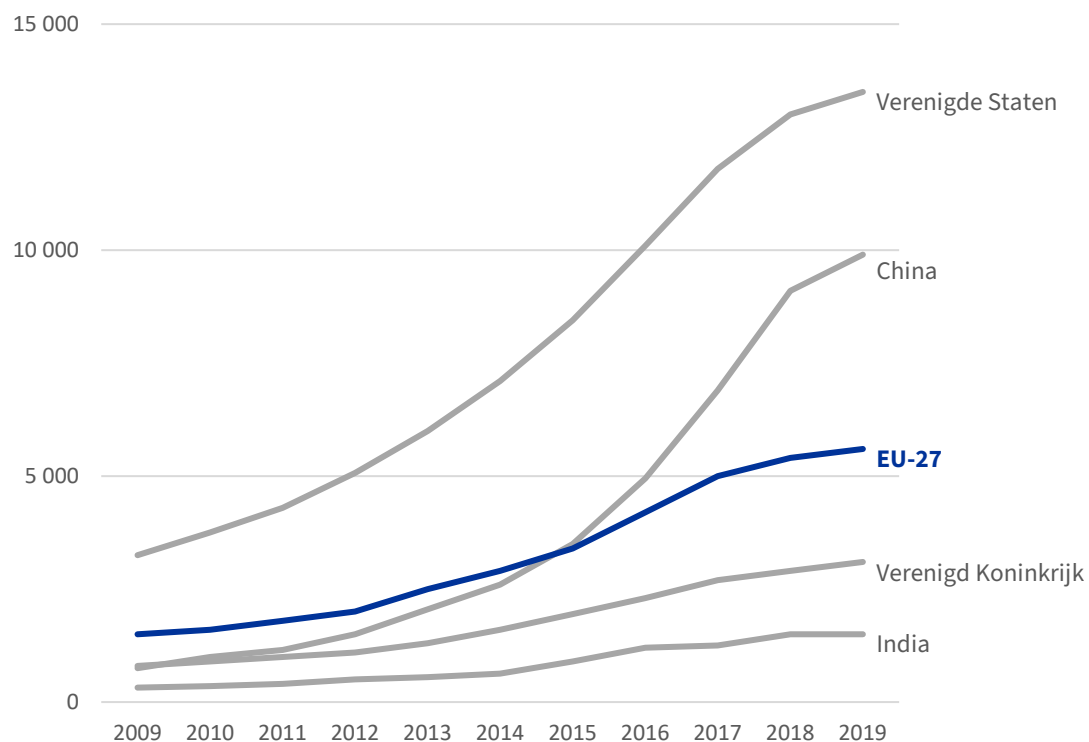


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



BELGIË



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Cyprus, Duitsland, Griekenland en Spanje**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Voor jouw land is de eerbiediging van de grondrechten een topprioriteit. Om nationale en internationale dreigingen — ook door digitale tools — aan te pakken, moet er dringend grensoverschrijdend worden samengewerkt.
- Je wilt ervoor zorgen dat AI op verantwoorde wijze wordt ontwikkeld. Sinds 2010 hebben ontwikkelaars grote doorbraken bereikt. Nieuwe tools worden sneller ontwikkeld dan dat de samenleving zich eraan kan aanpassen. Er zijn maar weinig mensen die zich een goed beeld kunnen vormen van de toekomst van AI. AI voor algemene doeleinden zal mogelijk op korte termijn het menselijk redeneren kunnen nabootsen en de vermogens van AI-systemen zouden door middel van kwantumcomputing exponentieel kunnen worden vergroot (met nog onbekende gevolgen).
- Als AI op verantwoorde wijze wordt ontwikkeld, kan AI levens verbeteren en mondiale uitdagingen doeltreffender aanpakken. De ethische aspecten moeten echter ook worden overwogen. Daarom pleit je voor een zo voorzichtig en verantwoord mogelijke aanpak. Testen in de praktijk, zoals voorgesteld in de flexibele aanpak, is een absolute no-go. Het brengt onnodige risico's voor gebruikers met zich mee — risico's die kunnen worden beperkt door eerst uitgebreid te testen in een testomgeving.
- Je bent voorstander van een capaciteitsgebonden risicoclassificatie: elk AI-systeem waar misbruik van kan worden gemaakt of dat vooroordelen of schadelijke gevolgen kan hebben, moet als systeem met een hoog risico worden beschouwd. De ontwikkeling van AI moet hand in hand gaan met de bevordering van diversiteit en mensenrechten en mag deze niet in de weg staan.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Je wilt graag dat de verordening breed toegepast kan worden. Je vindt het niet erg zinvol om een akkoord te bereiken over een EU-brede AI-verordening en verantwoorde AI-ontwikkeling te standaardiseren om vervolgens meteen zeer belangrijke gebieden, zoals veiligheid, buiten de verordening te laten.
- Alleen voor opensourcemonellen zou je eventueel een uitzondering willen maken. Die stimuleren snelle innovatie, omdat het ontwikkelingsproces ervan transparant is. Ontwikkelaars experimenteren en leren samen, en wisselen kennis uit om tools te verbeteren. Opensource-AI is vaak beter controleerbaar dan beschermde modellen en een betere controle van AI-systemen is een van de doelstellingen van de wetgevingshandeling.
- Jouw opmerkingen:

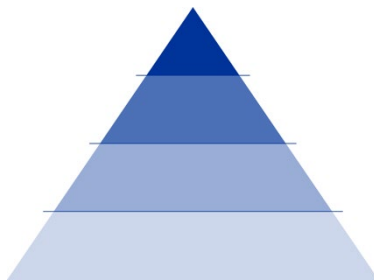
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde training-datasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België	capaciteitsgebonden	voorzorgsbenadering		opensource	
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

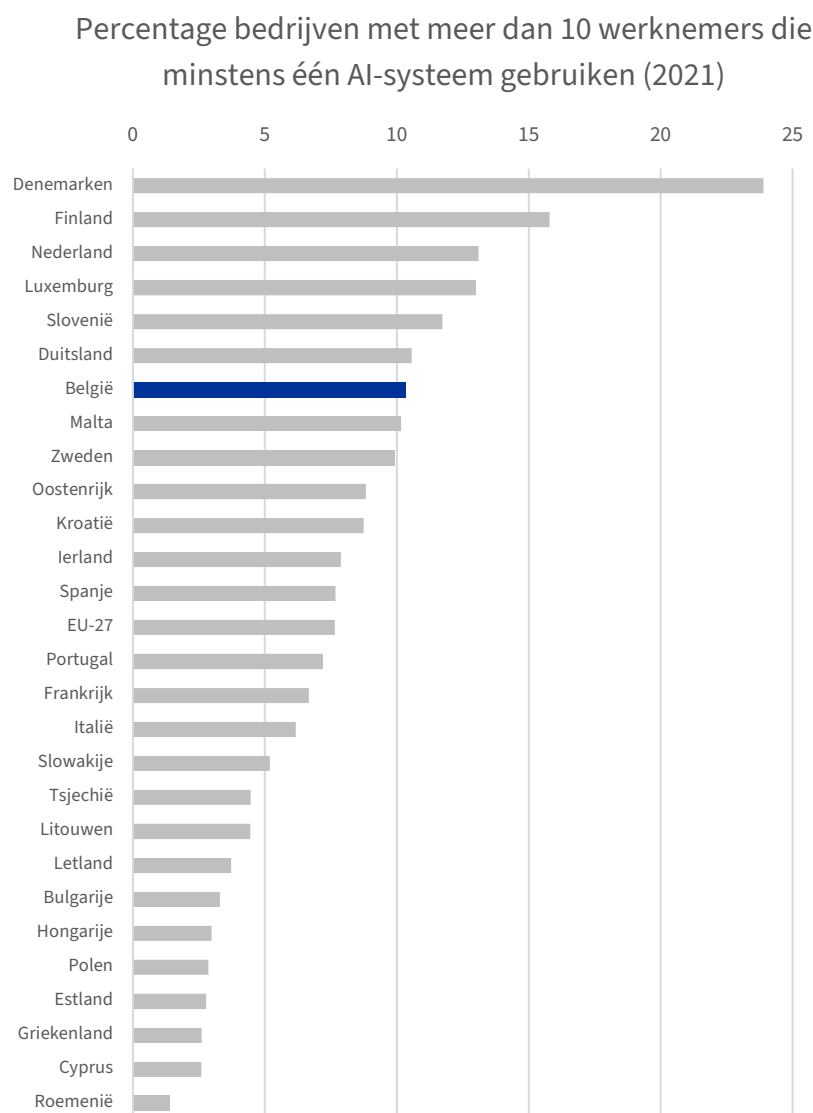
	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

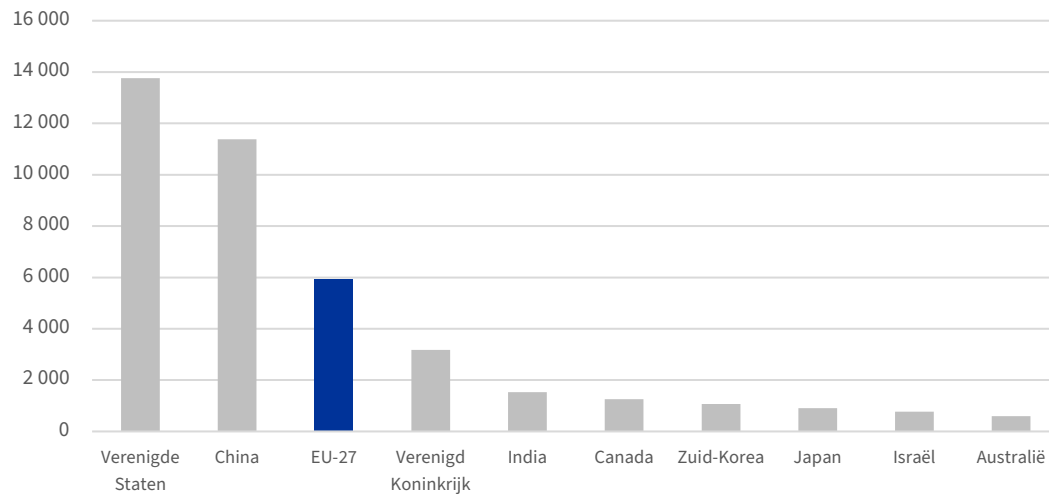
Let op: **De onderhandelingen vinden plaats in 2023!**

	Europese Unie	België
Bevolking	447 695 350	11 742 796
Bbp per hoofd van de bevolking in EUR	38 150	50 610
Werkloosheidspercentage	6,1%	5,5%
Percentage bedrijven dat AI-technologie gebruikt	8%	13,8%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	28,3%

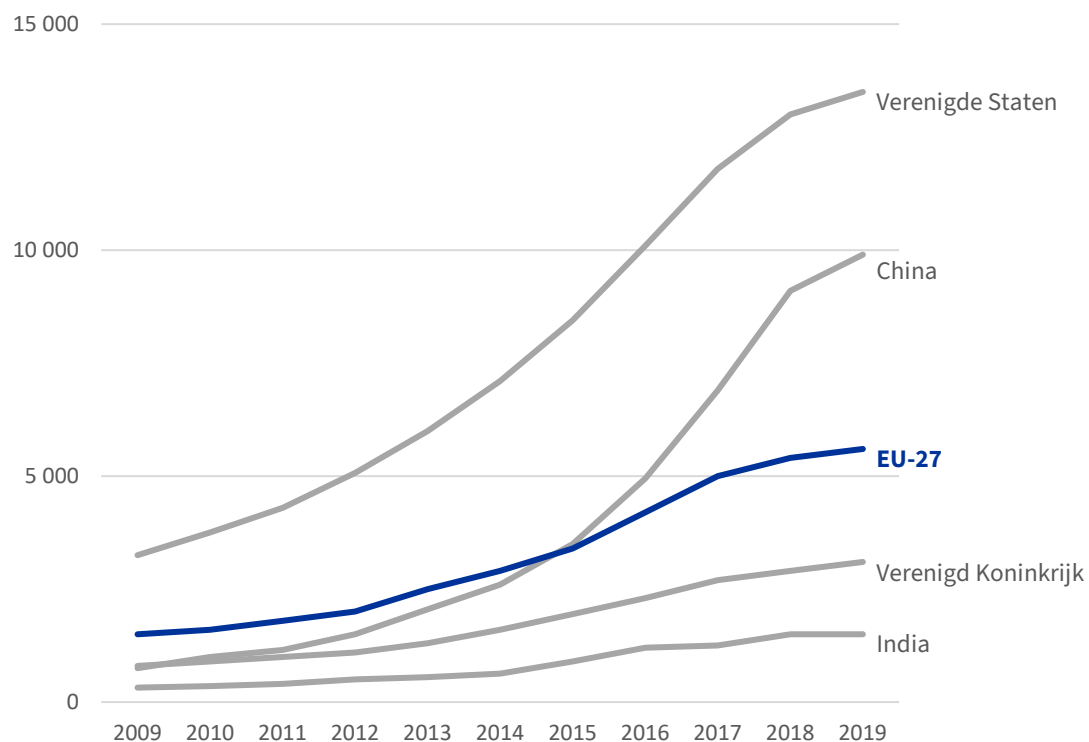


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



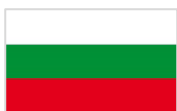
Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

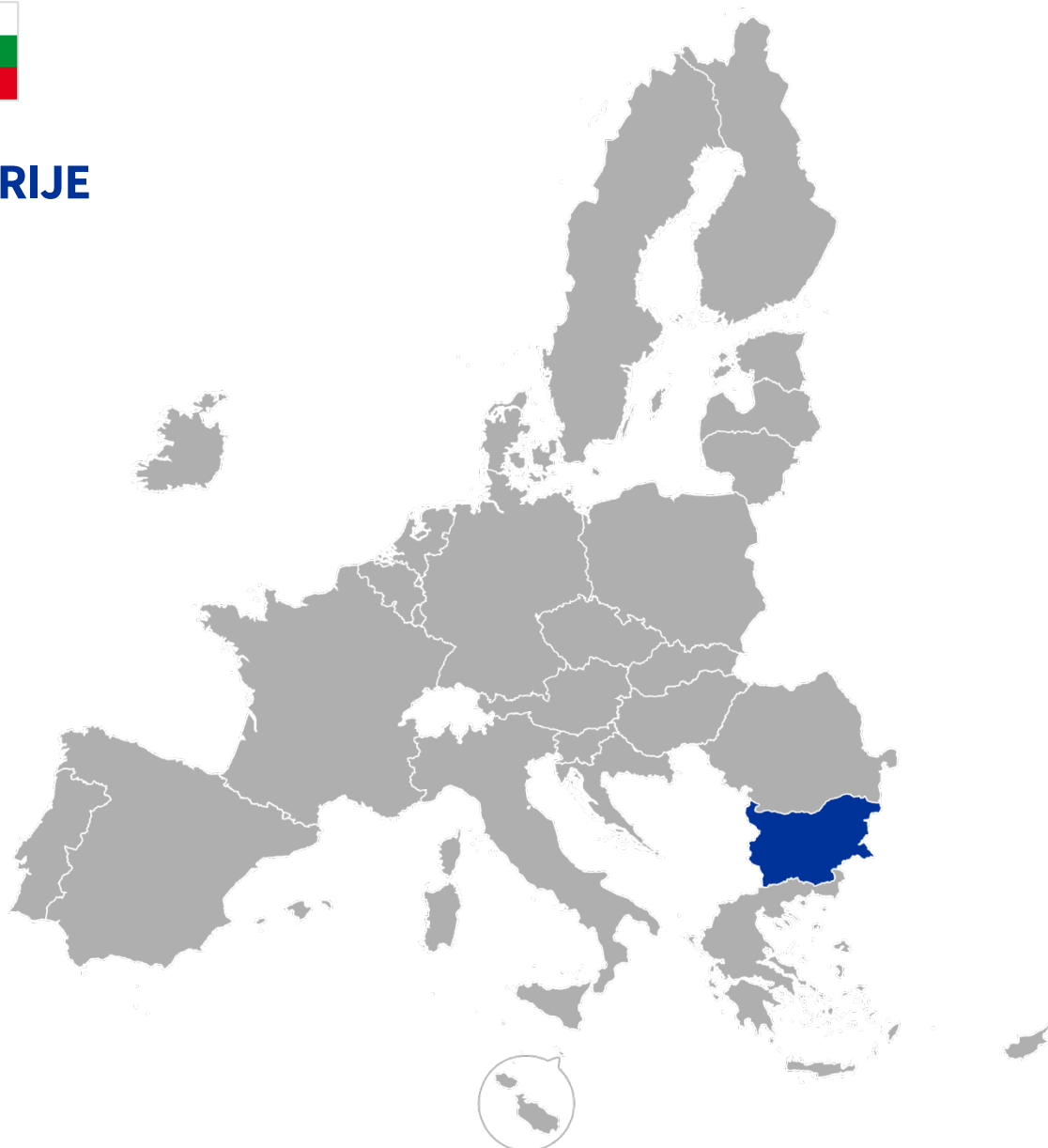
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



BULGARIJE



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

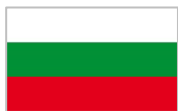
Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Luxemburg, Malta en Slowakije**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- De Bulgaarse hightech- en IT-sector zijn goed bezig op het gebied van AI. Ze hebben belangrijke klanten in West-Europa en de VS, en AI maakt nu al 3% van de Bulgaarse arbeidsmarkt uit – een aanzienlijk aandeel dat je graag nog verder zou zien groeien. Je wilt graag dat de AI-strategie van de EU jouw nationale industrieën en universiteiten helpt om gebruik te maken van alle voordelen van AI.
- In principe steun je de doelgebonden aanpak. Die aanpak moedigt ontwikkelaars en exploitanten aan om goed na te denken over hoe hun systemen kunnen worden toegepast en niet alleen over wat zij kunnen doen. Dit zorgt voor meer verantwoordingsplicht bij het gebruik van AI in de praktijk.
- Wat voor jou echter nog belangrijker is, is dat de EU erkent dat er geografische verschillen zijn. Je pleit derhalve voor geografisch evenwicht. Rijkere landen beschikken over veel meer financiële en personele middelen voor AI-onderzoek, en ook voor langdurige en dure tests voordat de AI-systemen in gebruik worden genomen. Om een eerlijke ontwikkeling in de hele EU te waarborgen, heeft Bulgarije financiering nodig om zijn kleine en middelgrote ondernemingen en start-ups te ondersteunen.
- Jouw grootste prioriteit is een duidelijke toezegging dat er EU-financiering zal komen voor AI-onderzoek in Bulgarije. Je staat ervoor open om risicobeheer en testvereisten te bespreken, maar pas nadat jouw belangrijkste zorgen serieus zijn genomen.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Hoewel je de flexibele regelgeving in artikel 1 steunt, ben je nog steeds van mening dat er duidelijke regels moeten gelden voor alle sectoren, ook voor de nationale veiligheids- en inlichtingendiensten. Deze sectoren van de verordening uitsluiten zal voor mazen in de wetgeving zorgen en vergroot de kans dat onethische of schadelijke AI ongecontroleerd wordt ingezet.
- Het gaat hier niet om het tegenhouden van het gebruik van AI voor nationale veiligheidsdoeleinden — het gaat erom dat AI op verantwoorde wijze wordt ontworpen om vooringenomenheid en discriminatie te voorkomen.
- Je vindt het huidige voorstel te vaag geformuleerd en wilt duidelijk maken dat de verordening ook van toepassing is op geïmporteerde AI-modellen van ontwikkelaars buiten de EU. Dit enkel om de wetgeving zo duidelijk mogelijk te maken. Het doel van de AI-verordening is nog steeds om de ontwikkeling van AI in de EU te stimuleren en AI-modellen in de EU aantrekkelijker te maken dan Amerikaanse of Chinese producten.
- Jouw opmerkingen:

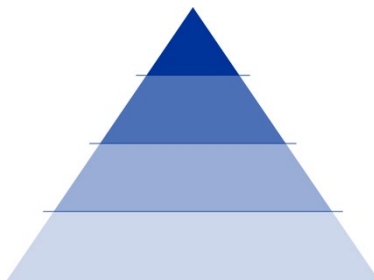
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde training-datasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije	doelgebonden	flexibele benadering		zonder uitzonderingen	
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

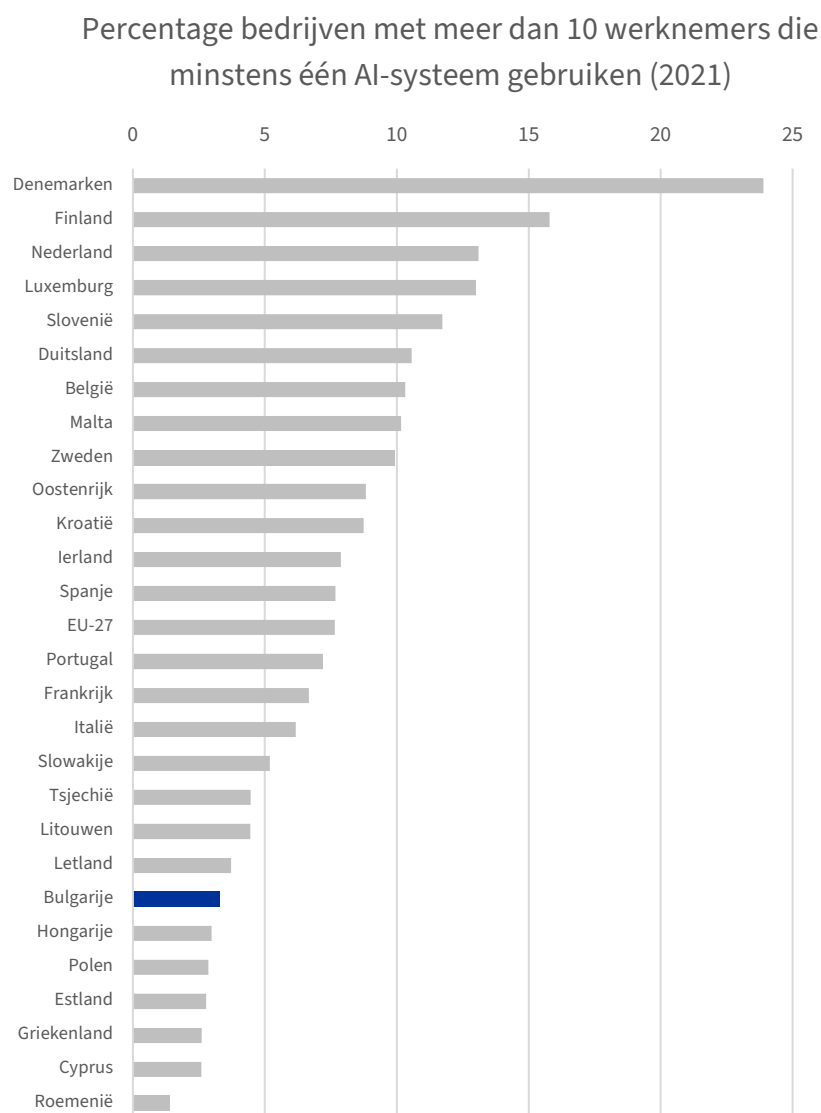
	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

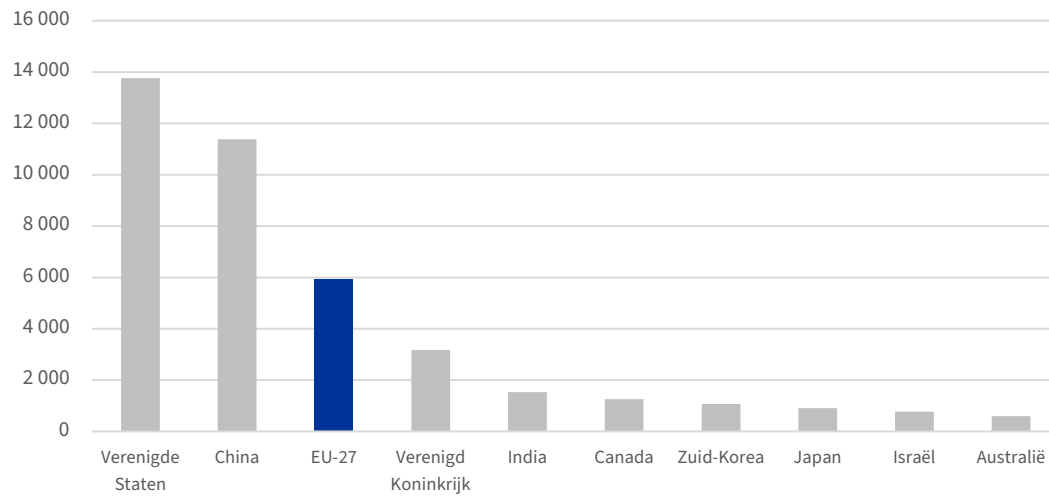
Let op: **De onderhandelingen vinden plaats in 2023!**

	Europese Unie	Bulgarije
Bevolking	447 695 350	6 447 710
Bbp per hoofd van de bevolking in EUR	38 150	14 690
Werkloosheidspercentage	6,1%	4,3%
Percentage bedrijven dat AI-technologie gebruikt	8%	3,6%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	7,7%

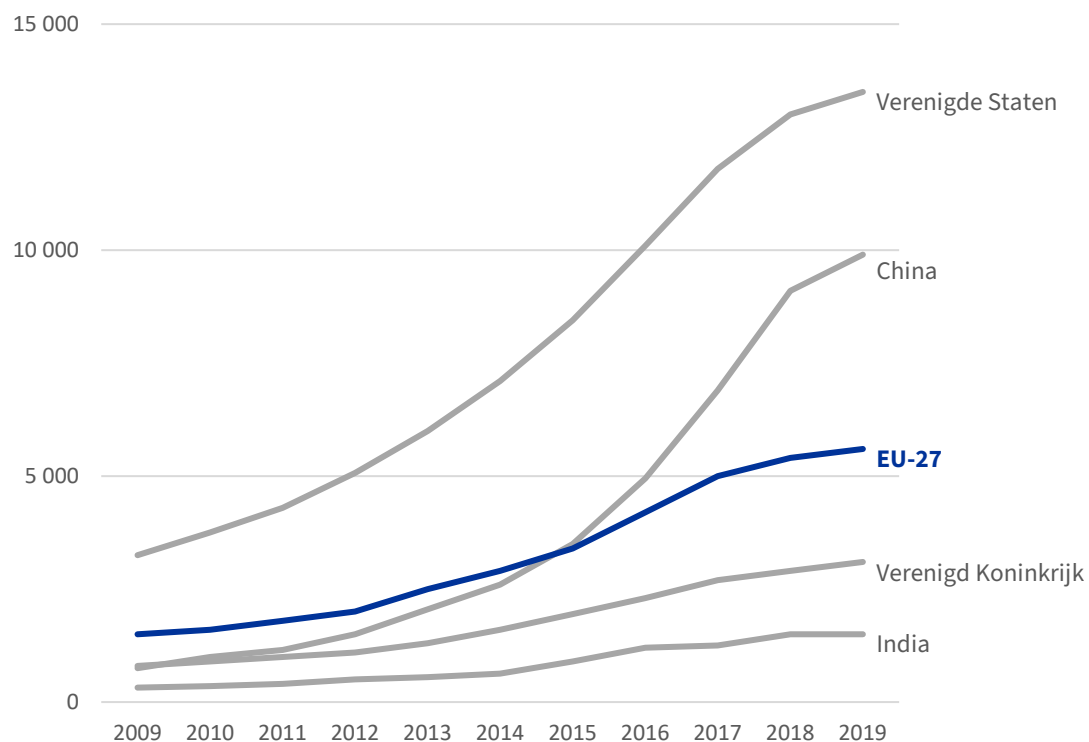


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

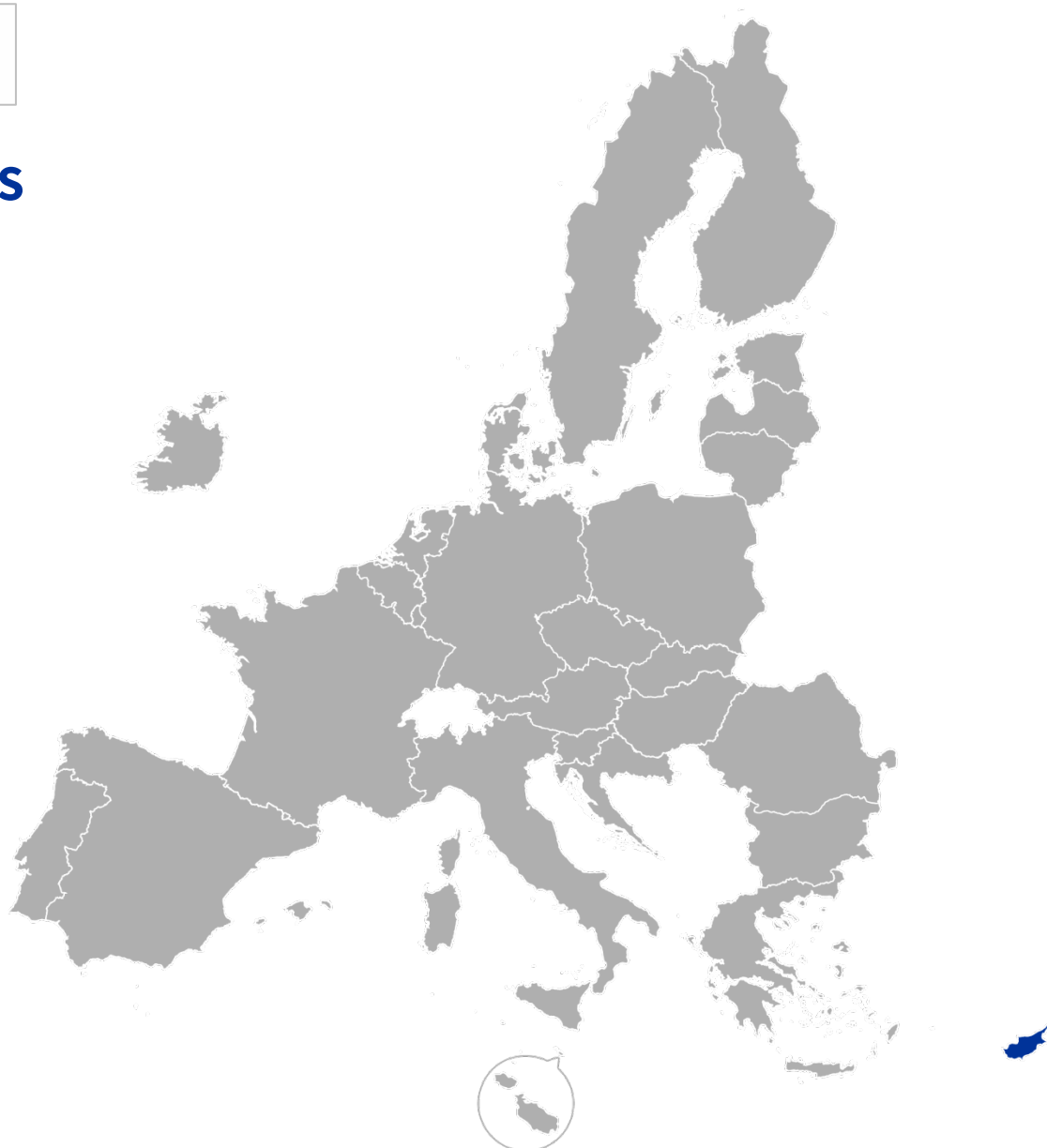
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



CYPRUS



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



CYPRUS

Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als België, Duitsland, Griekenland en Spanje**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... **capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.**



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... **voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).**

Jouw argumenten

- Jouw nationale regering geeft prioriteit aan beleidsmaatregelen die onderzoek en innovatie bevorderen en past tegelijkertijd strikte ethische richtsnoeren toe. Ongereguleerde AI zal ongetwijfeld een bedreiging vormen voor de veiligheid van burgers en bedrijven. In 2010 leidden door AI ondersteunde handelsalgoritmen bijvoorbeeld tot een instorting van de markt, waardoor \$ 1 biljoen aan aandelen binnen enkele minuten niets meer waard waren.
- Jij vindt dat bij de ethische ontwikkeling van AI prioriteit moet worden gegeven aan transparantie (AI begrijpelijk maken), verantwoordingsplicht (ervoor zorgen dat mensen verantwoordelijk blijven voor de beslissingen van AI) en controlemechanismen (voorkomen dat AI ongecontroleerd te werkt gaat). Je vreest dat een te flexibele regelgeving deze prioriteiten zou kunnen ondermijnen.
- Door doortastende maatregelen te nemen en een voorzorgsbenadering te volgen met duidelijke, uitvoerbare regels kan de EU zichzelf positioneren als een verantwoordelijke en toekomstgerichte leider op het gebied van AI-innovatie.
- Je bent voorstander van een capaciteitsgebonden risicoclassificatie. Elk AI-systeem waarvan misbruik kan worden gemaakt of dat vooroordelen of schadelijke gevolgen kan hebben, moet als systeem met een hoog risico worden beschouwd – niet omdat AI inherent gevaarlijk is, maar vanwege de manier waarop AI kan worden gebruikt. Je zou echter **bereid zijn om een doelgebonden classificatie te bespreken als deze binnen 2 jaar zou worden herzien.**

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Je bent er stellig van overtuigd dat AI overal op ethische wijze moet worden ontwikkeld — geen enkele sector en geen enkele ontwikkelaar mag buiten de AI-verordening van de EU vallen. Rechtshandhaving, het leger, defensie en nationale veiligheid zijn juist de gebieden waar het risico op misbruik, vooringenomenheid en schendingen van rechten het grootst is. Die sectoren mogen niet als uitzonderingen worden behandeld, maar moeten aan de hoogste normen voldoen.
- De EU wil op wereldniveau koploper zijn op het gebied van betrouwbare en verantwoorde AI. Daarom wil jij dat de EU het goede voorbeeld geeft. Een oogje dichtknijpen voor bepaalde toepassingen zou de geloofwaardigheid en de impact van de AI-verordening van de EU verminderen. Ethische beginselen moeten als leidraad dienen voor alle toepassingen van AI, met name de toepassingen die het grootste potentieel hebben om de levens en vrijheden van mensen te beïnvloeden.
- Artikel 2 is voor jou belangrijker dan artikel 1. Als je de anderen niet kunt overtuigen van jouw standpunten over artikel 1, moet je ervoor zorgen dat je belangrijkste wensen met betrekking tot artikel 2 serieus worden genomen.
- Jouw opmerkingen:

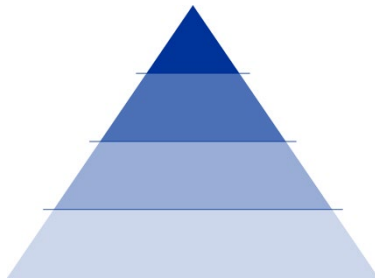
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico

Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus	capaciteitsgebonden	voorzorgsbenadering	over 2 jaar opnieuw bekijken	zonder uitzonderingen	
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

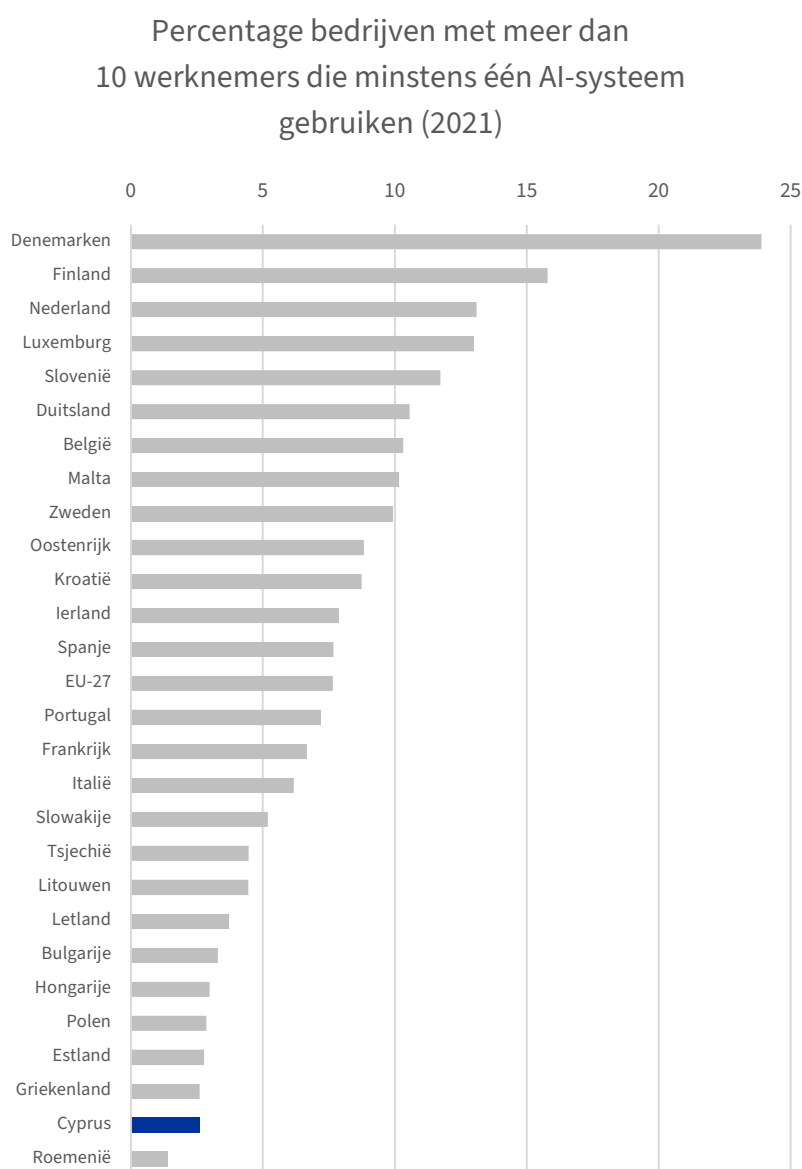
	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

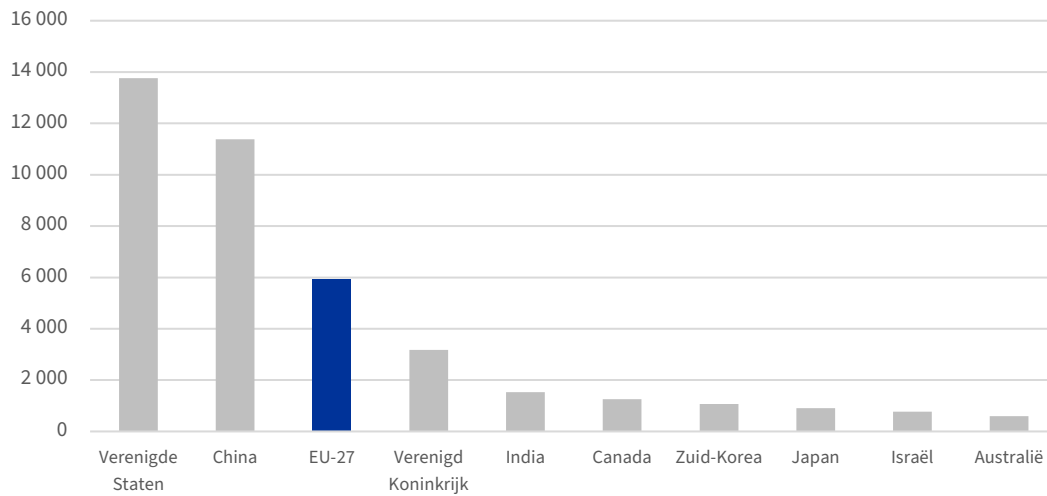
Let op: **De onderhandelingen vinden plaats in 2023!**

	Europese Unie	Cyprus
Bevolking	447 695 350	949 084
Bbp per hoofd van de bevolking in EUR	38 150	32 720
Werkloosheidspercentage	6,1%	5,8%
Percentage bedrijven dat AI-technologie gebruikt	8%	4,7%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	25%

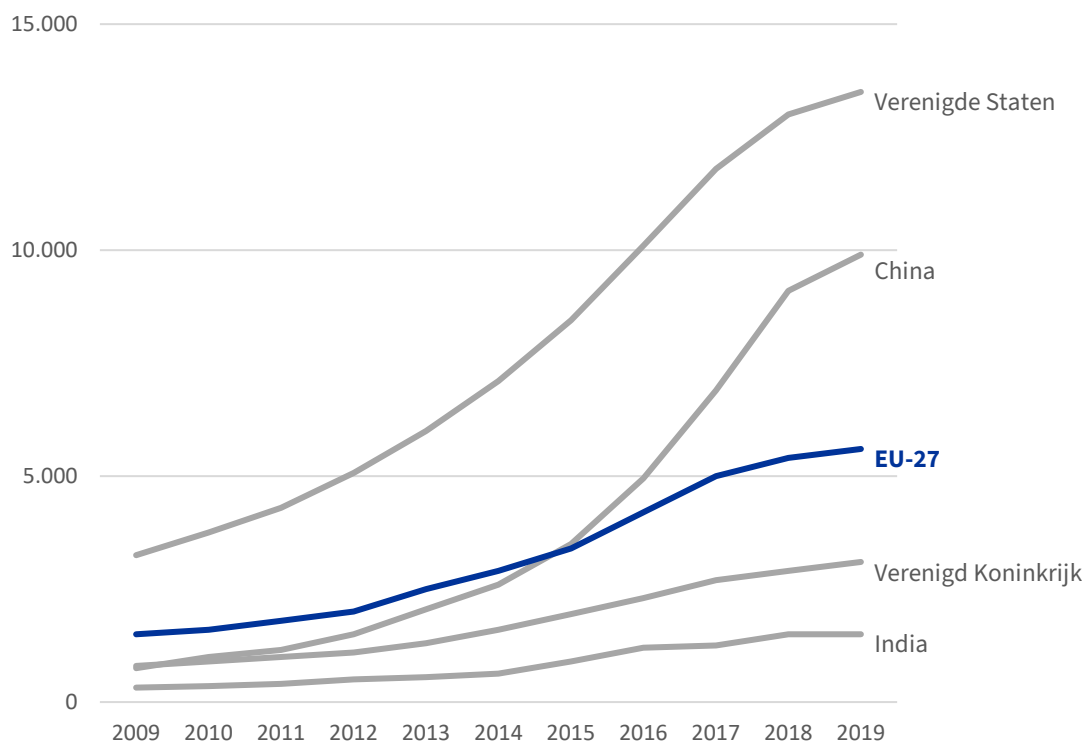


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

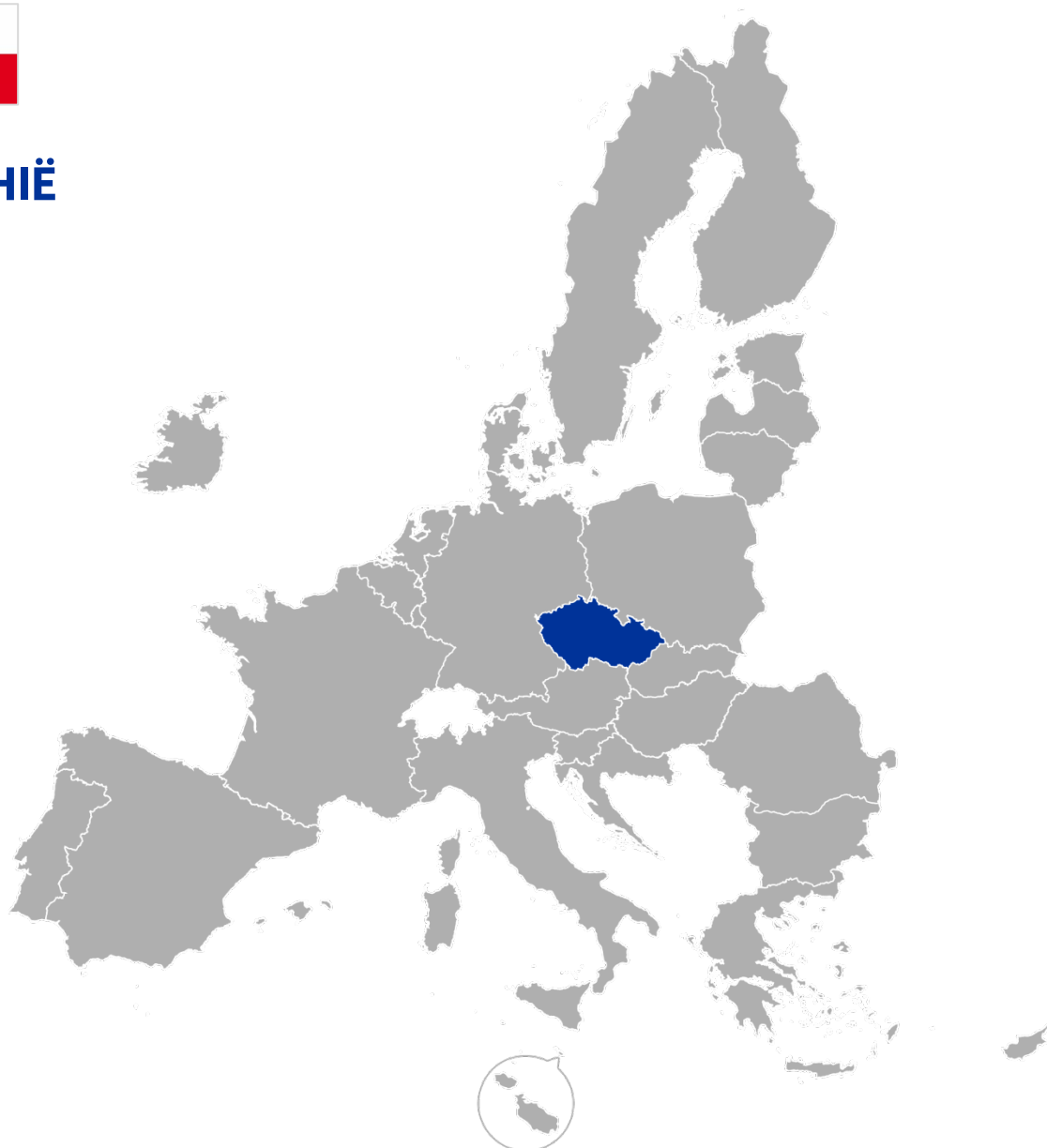
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



TSJECHIË



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Estland, Hongarije, Letland en Litouwen**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... **capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.**



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... **flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.**



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Jouw nationale regering wil AI-innovatie bevorderen. Hoewel de EU sterk is op het gebied van onderzoek, bepalen de VS en China de regels op de wereldmarkt. Ook andere landen, zoals India, Canada en het VK, boeken snel vooruitgang.
- Gezien deze context ben je van mening dat de EU evenwichtige regelgeving moet opstellen — die niet al te rigide is — om AI-innovatie te stimuleren en zichzelf als wereldleider op het gebied van AI te positioneren. De AI-verordening van de EU kan een doorbraak betekenen als de EU het vertrouwen en de steun van zowel investeerders als ontwikkelaars kan winnen.
- Je bent voorstander van een capaciteitsgebonden aanpak, omdat een doelgebonden aanpak in strijd kan zijn met bestaand sectoraal beleid op nationaal en EU-niveau en zo voor onduidelijkheid kan zorgen. Het zou echter te rigide zijn om alle AI voor algemene doeleinden te classificeren als systemen met een hoog risico. Veel AI-systemen voor algemene doeleinden, zoals AI die wordt gebruikt voor klantenservice of realtime vertaling, zijn systemen met een laag risico. Overregulering zou innovatie en investeringen kunnen ontmoedigen en het idee kunnen versterken dat de EU de strengste AI-regels ter wereld heeft, wat het concurrentievermogen van Europa uiteindelijk zou verzwakken.
- Als het moeilijk blijkt om vandaag een compromis te bereiken, **sta je ervoor open om een clause toe te voegen waarin wordt bepaald dat de verordening binnen 5 jaar wordt herzien** om ervoor te zorgen dat de verordening relevant blijft in een snel evoluerende AI-omgeving.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Je bent er stellig van overtuigd dat AI overal op ethische wijze moet worden ontwikkeld — geen enkele sector en geen enkele ontwikkelaar mag buiten de AI-verordening van de EU vallen. Rechtshandhaving, het leger, defensie en nationale veiligheid zijn juist de gebieden waar het risico op misbruik, vooroordelen en schendingen van rechten het grootst is. Deze sectoren mogen niet als uitzonderingen worden behandeld, maar moeten aan de hoogste normen voldoen.
- Als de EU op wereldniveau koploper wil zijn op het gebied van betrouwbare en verantwoorde AI, moet je het goede voorbeeld geven. Een oogje dichtknijpen voor bepaalde toepassingen zou de geloofwaardigheid en de impact van de AI-verordening van de EU verminderen. Ethische beginselen moeten als leidraad dienen voor alle toepassingen van AI, met name de toepassingen die het grootste potentieel hebben om de levens en vrijheden van mensen te beïnvloeden.
- Jouw opmerkingen:

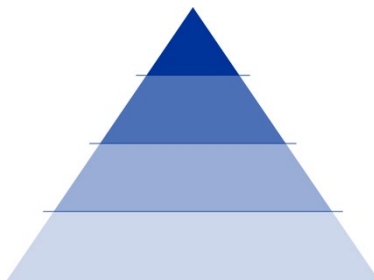
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde training-datasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië	capaciteitsgebonden	flexibele benadering	over 5 jaar opnieuw bekijken	zonder uitzonderingen	
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

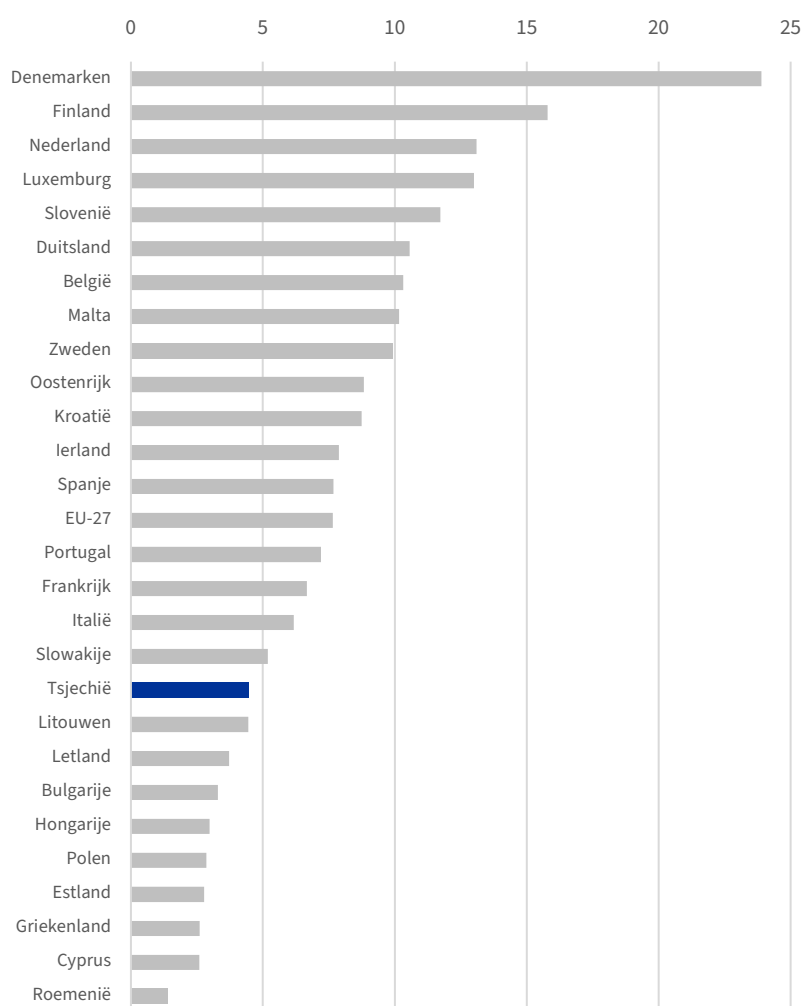
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

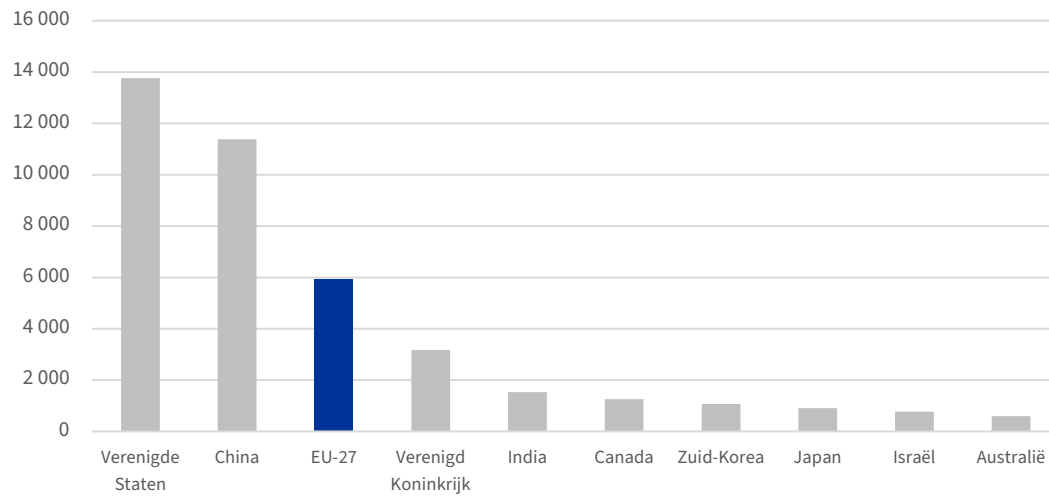
	Europese Unie	Tsjechië
Bevolking	447 695 350	10 827 529
Bbp per hoofd van de bevolking in EUR	38 150	29 180
Werkloosheidspercentage	6,1%	2,6%
Percentage bedrijven dat AI-technologie gebruikt	8%	5,9%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	35,5%

Percentage bedrijven met meer dan 10 werknemers
die minstens één AI-systeem gebruiken (2021)

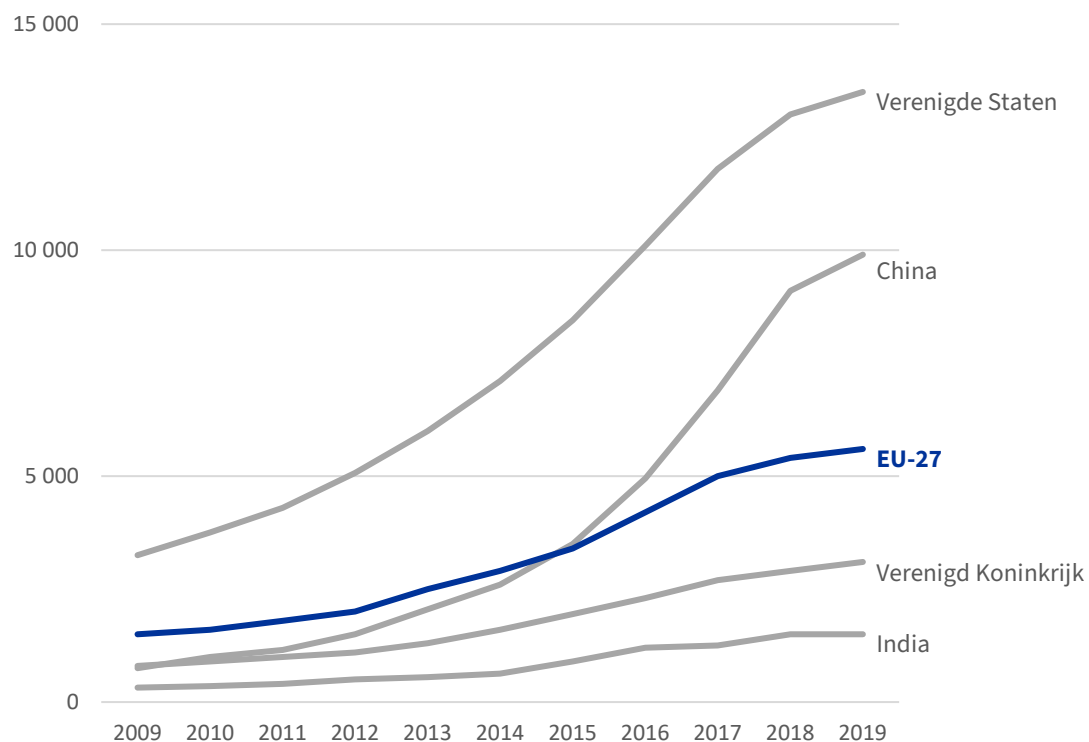


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

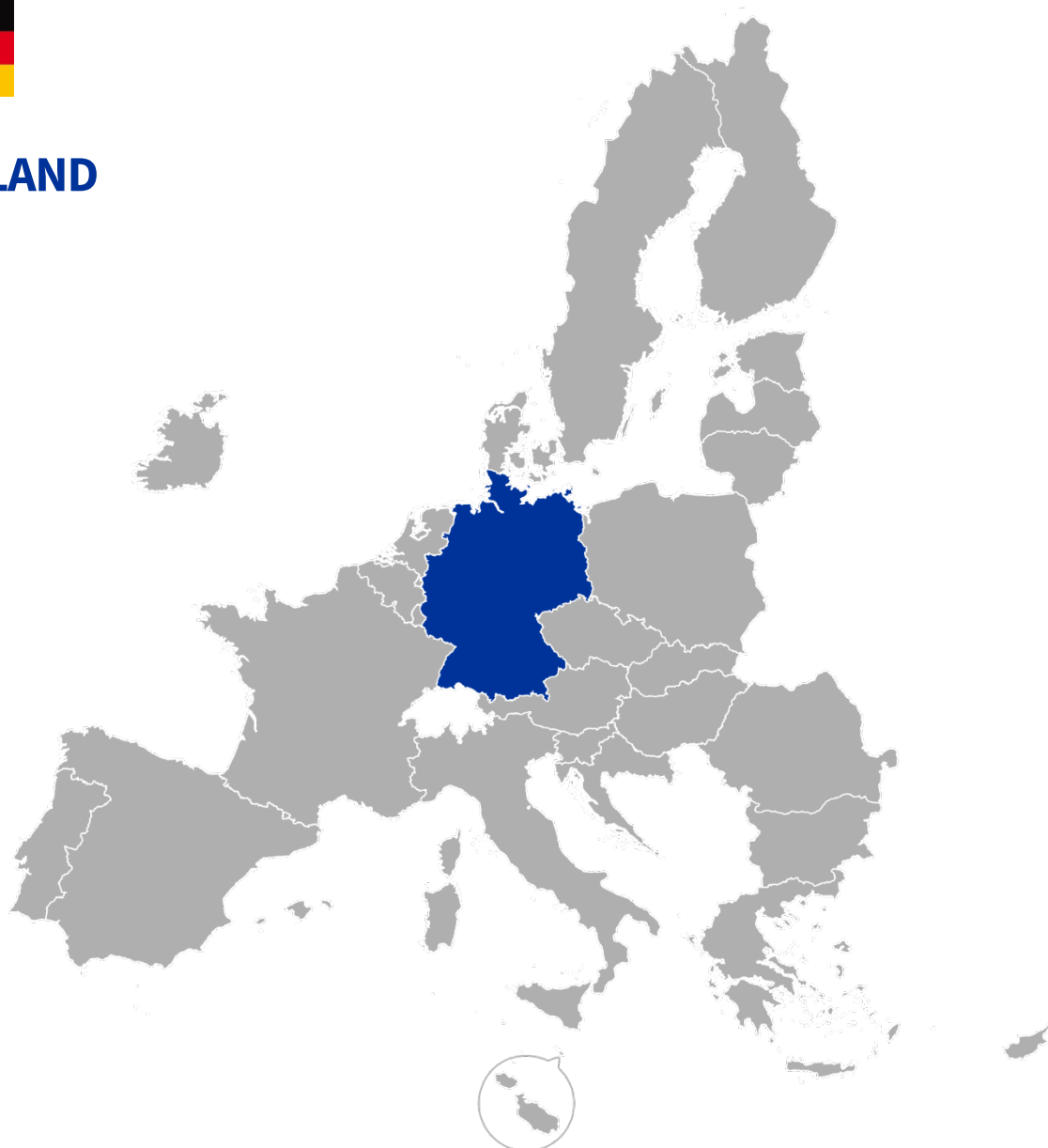
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



DUISSLAND



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als België, Cyprus, Griekenland en Spanje**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- De grondrechten bevorderen is jouw topprioriteit en je erkent dat er dringend behoefte is aan grensoverschrijdende samenwerking om zowel nationale als internationale bedreigingen aan te pakken, waaronder die van digitale technologieën. Je hebt gezien hoe AI misinformatie en deepfakecampagnes kan voeden, wat ernstige risico's inhoudt voor de vrijheid van meningsuiting en voor democratische instellingen.
- De ontwikkeling van AI vordert snel en speelt een steeds belangrijkere rol in de samenleving. Je vindt dat AI zo moet worden ontworpen dat mens en milieu er voordelen van ondervinden – geen schade.
- Je bent voorstander van strenge testvereisten voor modellen met een hoog risico, aangezien vooringenomen AI-systemen een groot probleem vormen. Wanneer AI wordt getraind op basis van niet-diverse of niet-representatieve datasets, kan AI discriminerende resultaten opleveren — zoals het geval was bij de aanwervingstool van Amazon, die werd afgedankt omdat deze mannelijke kandidaten bevoordeelde op basis van eerdere aanwervingspatronen. Vandaag is het je uiteindelijke doel om toe te werken naar regelgeving die zelfs het kleinste risico op door AI veroorzaakte schade, zoals misinformatie, wegneemt. Je bent je ervan bewust dat het een ambitieus doel is, maar je bent ook overtuigd van de noodzaak ervan.
- Je bent voorstander van een capaciteitsgebonden risicoclassificatie. Elk AI-systeem waarvan misbruik kan worden gemaakt of dat schadelijke gevolgen kan hebben, moet als een systeem met een hoog risico worden beschouwd – niet omdat AI inherent gevaarlijk is, maar vanwege de manier waarop AI kan worden gebruikt.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)

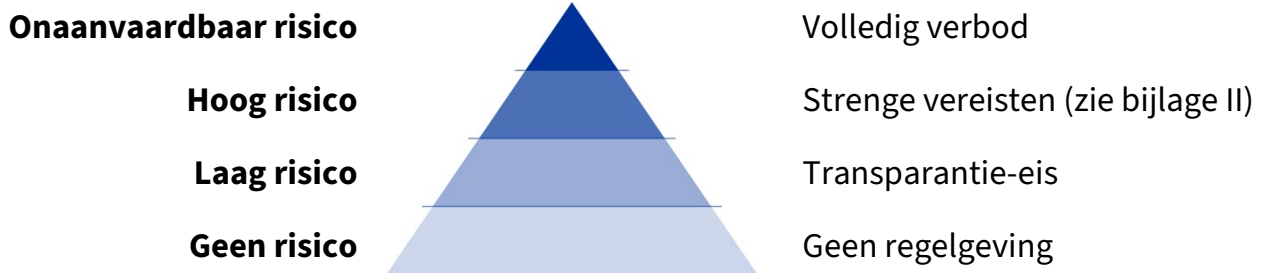


... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Je maakt je ernstig zorgen over de wens van sommige landen om rechtshandhaving, het leger, defensie of nationale veiligheid uit te sluiten van het toepassingsgebied van de AI-verordening van de EU. Dit zijn nou juist de gebieden waar AI ernstige schade kan veroorzaken — door discriminatie, foutpositieve resultaten en vooroordelen — en waar sterke waarborgen het hardst nodig zijn. Vooringenomen gezichtsherkenning bij politiewerk heeft bijvoorbeeld al geleid tot onterechte arrestaties — onevenredig vaak van geracialiseerde jonge mannen — en ernstige schendingen van individuele rechten.
- Je begrijpt dat de efficiëntie in veel landen flink is toegenomen als gevolg van het gebruik van AI in deze sectoren, en dat wil je niet ontkennen. Maar er kan niet te lichtjes met mensenrechten worden omgesprongen. Je hebt de AI-ontwikkelingen in de VS en China op de voet gevolgd en je bent ernstig bezorgd over het gebrek aan sterke beschermingsmaatregelen in hun systemen.
- Het hele punt van de AI-verordening van de EU is het bevorderen van verantwoorde, mensgerichte AI. Bepaalde sectoren van de verordening uitsluiten zou dat doel verzwakken en de ambitie van de EU om een mondiale norm voor betrouwbare AI op te stellen, ondermijnen.
- Jouw opmerkingen:

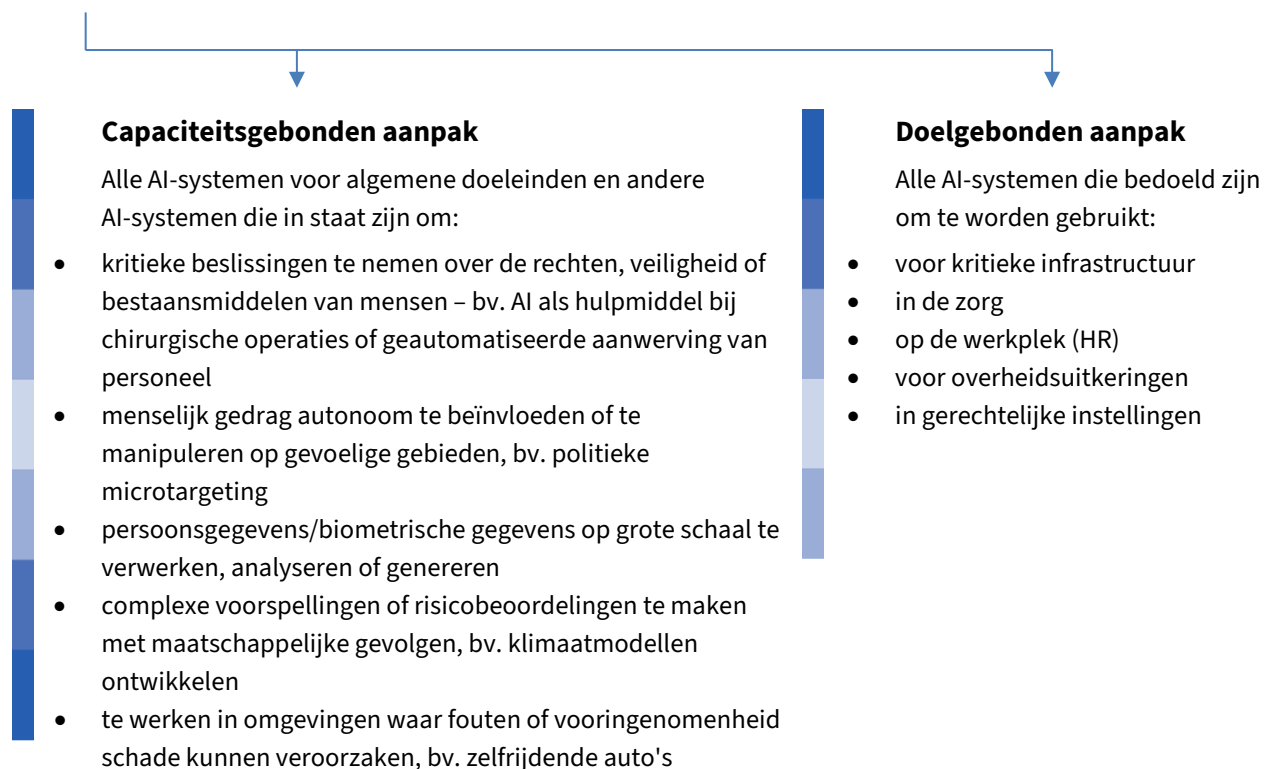
Bijlage I: Risicoclassificatie



Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde training-datasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland	capaciteitsgebonden	voorzorgsbenadering		zonder uitzonderingen	
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

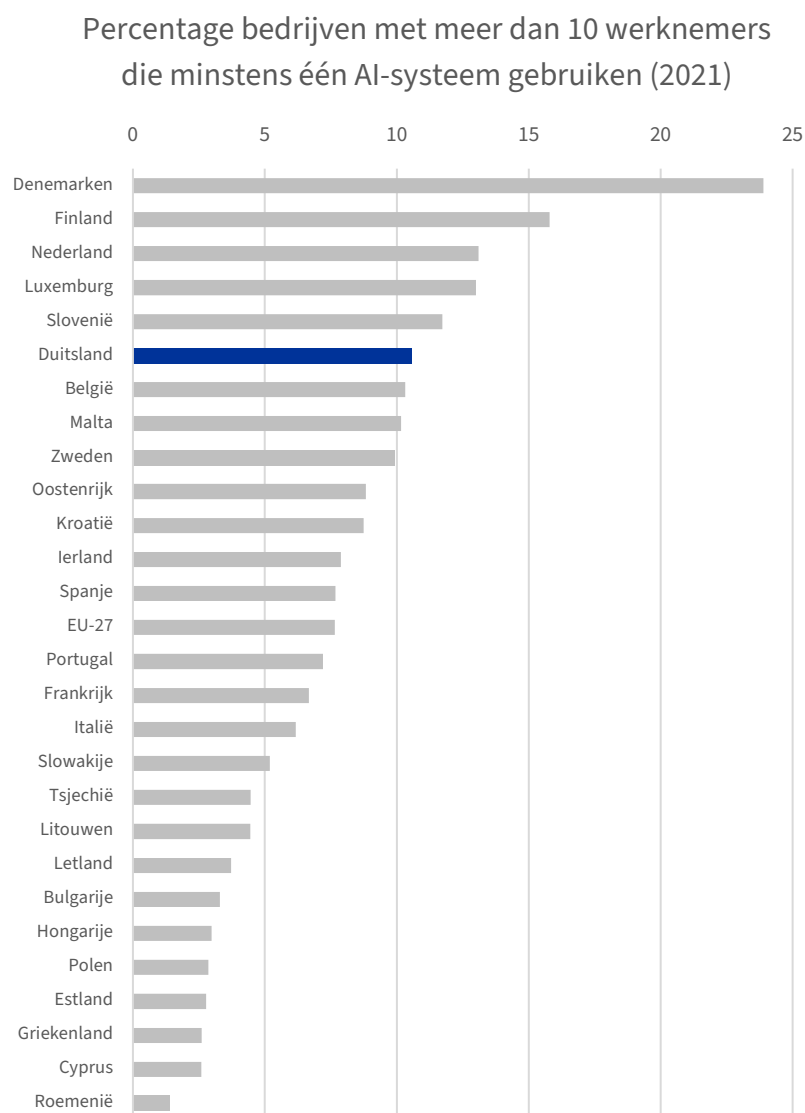
	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

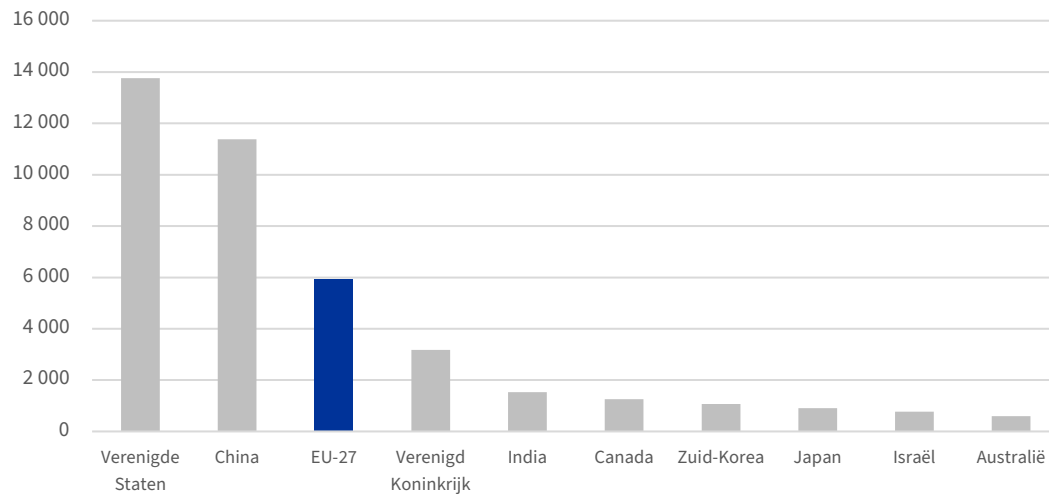
Let op: **De onderhandelingen vinden plaats in 2023!**

	Europese Unie	Duitsland
Bevolking	447 695 350	83 118 501
Bbp per hoofd van de bevolking in EUR	38 150	49 520
Werkloosheidspercentage	6,1%	3,1%
Percentage bedrijven dat AI-technologie gebruikt	8%	11,6%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	19,8%

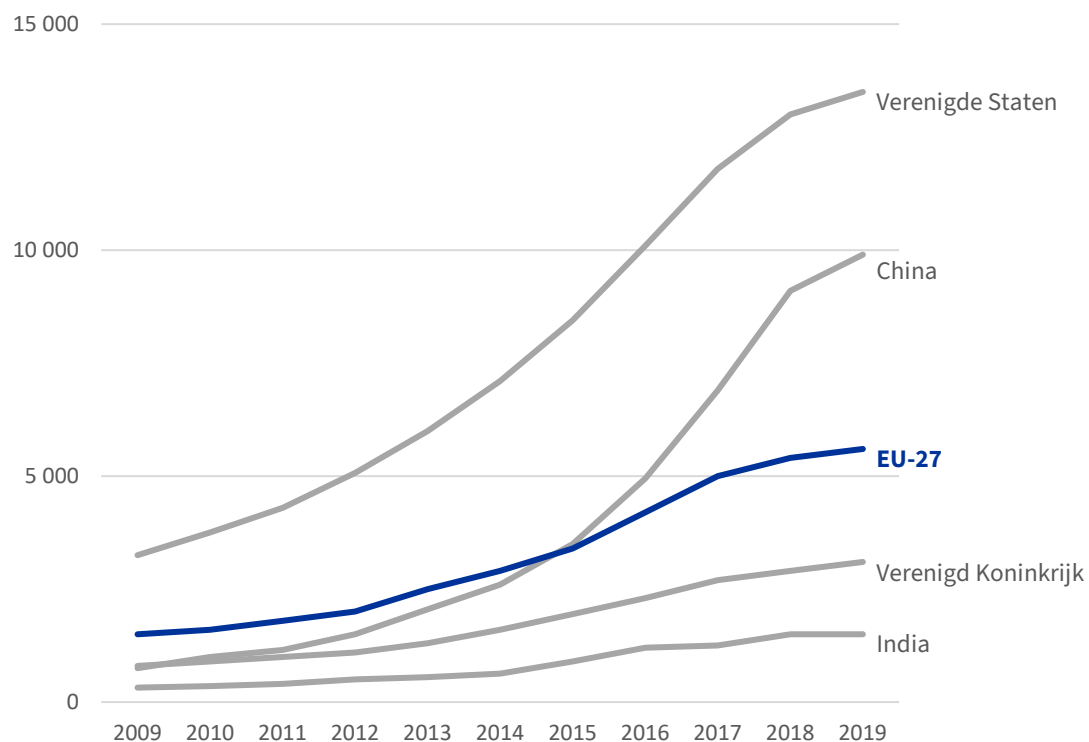


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

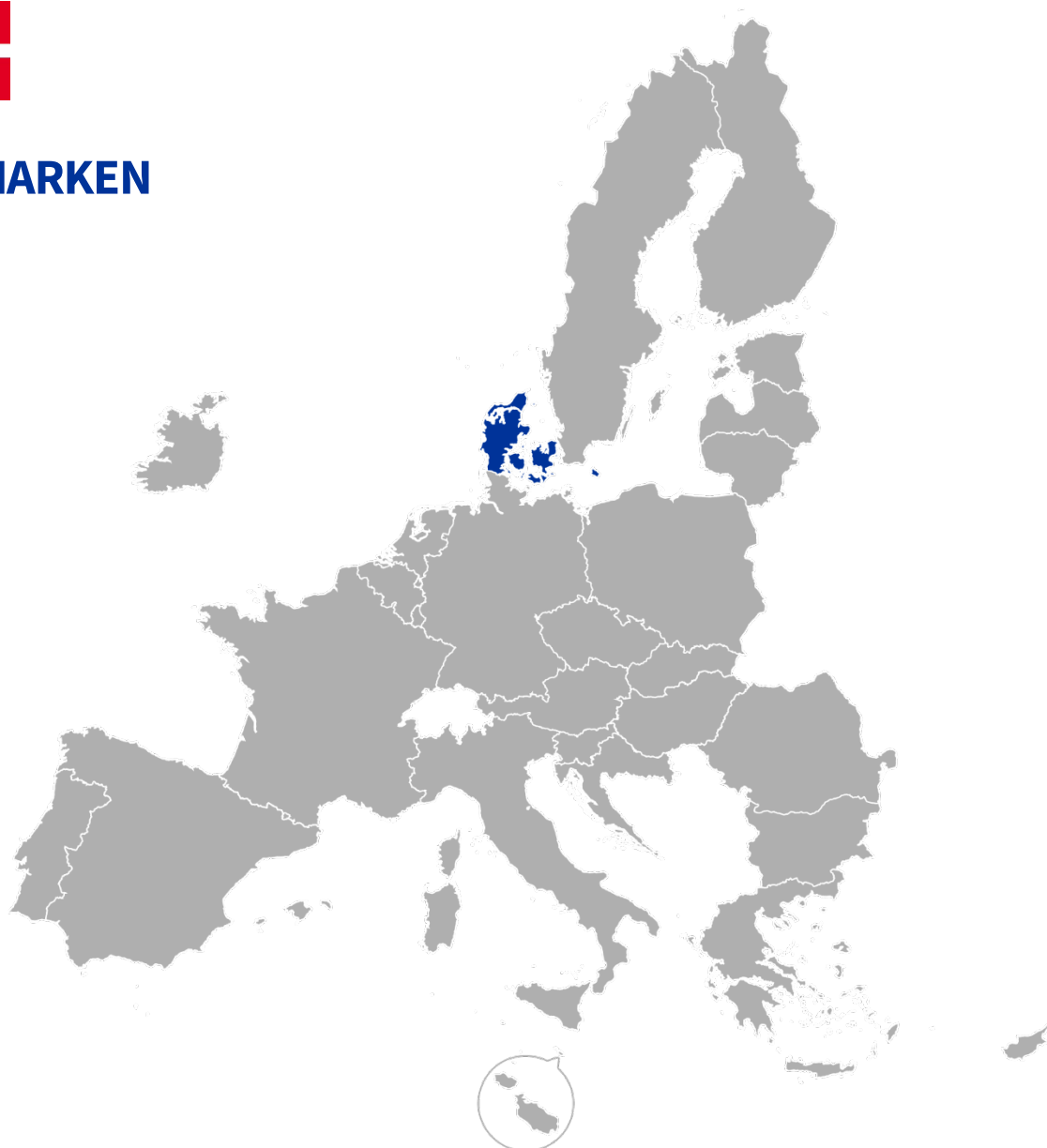
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



DENEMARKEN



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Frankrijk, Oostenrijk, Polen en Roemenië**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Je verwelkomt het voorstel van de Commissie en steunt de doelgebonden aanpak als constructieve eerste stap om de uitdagingen in verband met de ontwikkeling van AI het hoofd te bieden. Maar je wilt dat ook het onderwijs wordt beschouwd als gevoelige sector. Wanneer scholen AI gebruiken om gedrag te beoordelen, taken na te kijken of resultaten te voorspellen, kunnen studenten het gevoel krijgen dat een machine voortdurend over hun schouder meekijkt. Hierdoor kan de druk ontstaan om altijd op een bepaalde manier te "presteren", maar leerlingen mogen ook gewoon soms een slechte dag hebben. In tegenstelling tot een menselijke leerkracht kan AI vaak geen empathie tonen of begrip voor persoonlijke uitdagingen opbrengen.
- AI maakt nieuwe manieren van besluitvorming mogelijk en daar kan je ethische vraagtekens bij plaatsen. Je bent het dan ook eens met een kader voor tests uit voorzorg, waarmee een cultuur van verantwoordelijkheid wordt bevorderd en het makkelijker wordt om bij problemen de oorzaken ervan te achterhalen.
- Voorts moet in een gemeenschappelijk EU-kader de taalkundige verscheidenheid van de EU worden erkend. De meeste grote AI-modellen worden getraind op data uit veel gesproken talen zoals het Engels, het Frans of het Duits, waardoor die modellen in minder gesproken talen zoals het Deens mogelijk slechter presteren.
- **Als de EU besluit middelen vrij te maken voor de ontwikkeling van AI, dan zou het eerlijk zijn om in de eerste plaats kleinere landen en minder gesproken talen te ondersteunen.** Anders zou de ontwikkeling van AI tot meer ongelijkheid kunnen leiden, met lagere kwaliteit in minder gesproken talen.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Het blijft voor jou van essentieel belang dat zaken in verband met de nationale veiligheid, het leger of defensie buiten het toepassingsgebied van de AI-verordening blijven. Deze beleidsdomeinen vallen doorgaans buiten de bevoegdheid van de EU en de lidstaten zijn vastbesloten om de controle te behouden over besluiten op deze gebieden. Dit blijkt ook uit het feit dat de EU (nog) geen verenigd EU-leger heeft.
- Daarnaast voer je aan dat voor AI-systemen die worden gebruikt voor het verzamelen van inlichtingen, cyberdefensie of gevechtscenari's niet dezelfde openbaarmakings- of regelgevingsvereisten kunnen gelden als voor civiele systemen. Dat zou de operationele geheimhouding, de nationale veiligheidsbelangen en het strategische voordeel dat zulke systemen moeten bieden namelijk in gevaar kunnen brengen.
- Als je verzoek tijdens de discussie van vandaag niet volledig kan worden ingewilligd, ben je bereid tot een compromis waarin AI-systemen voor militaire doeleinden en nationale veiligheid wel onder de verordening vallen, maar slechts moeten voldoen aan de testvereisten van de flexibele benadering.
- Jouw opmerkingen:

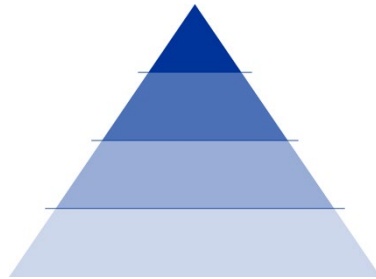
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico

Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken	doelgebonden	voorzorgsbenadering	financiële steun	leger/defensie + nationale veiligheid	
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

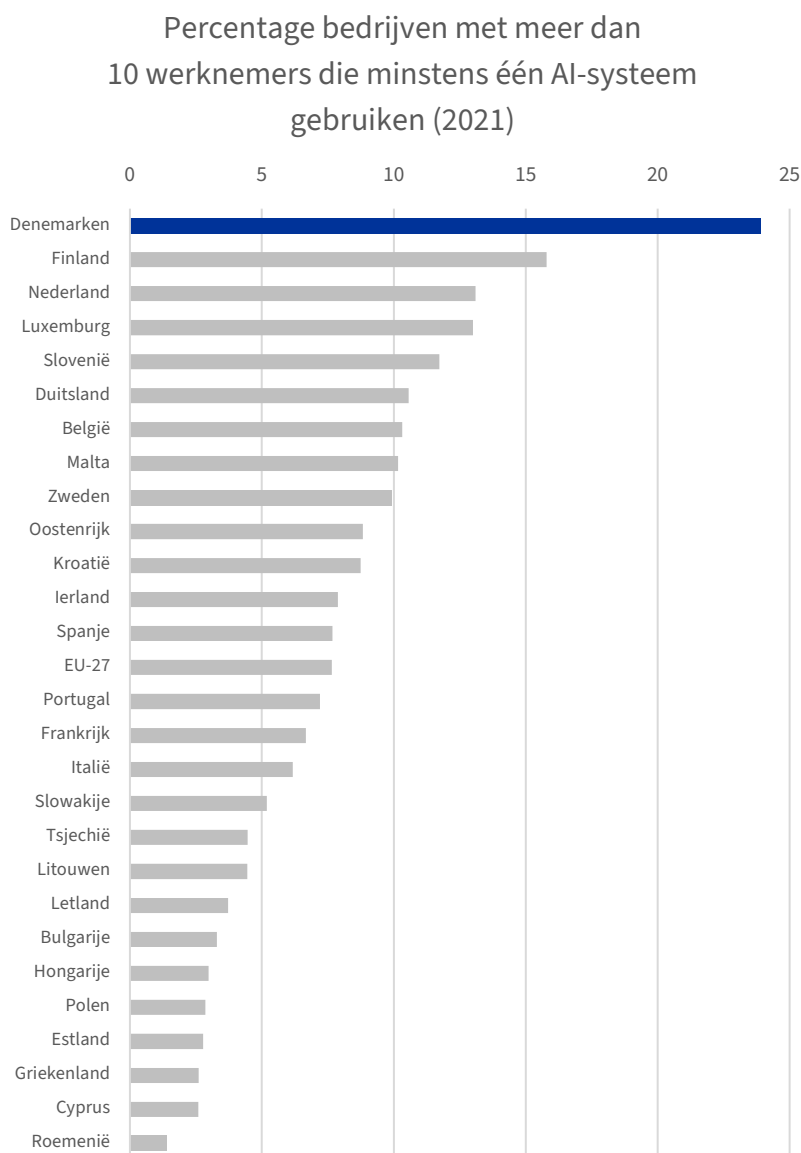
	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

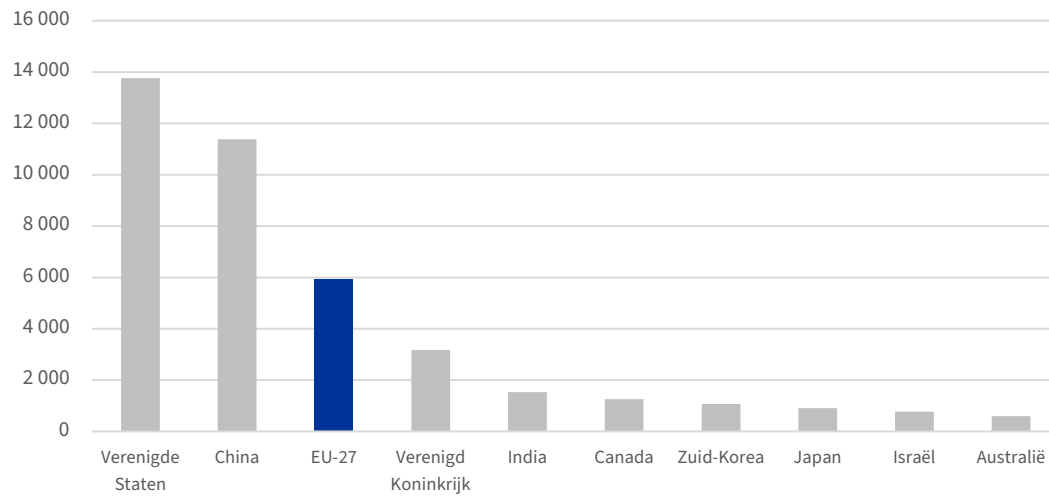
Let op: **De onderhandelingen vinden plaats in 2023!**

	Europese Unie	Denemarken
Bevolking	447 695 350	5 932 654
Bbp per hoofd van de bevolking in EUR	38 150	63 290
Werkloosheidspercentage	6,1%	5,1%
Percentage bedrijven dat AI-technologie gebruikt	8%	15,2%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	39,4%

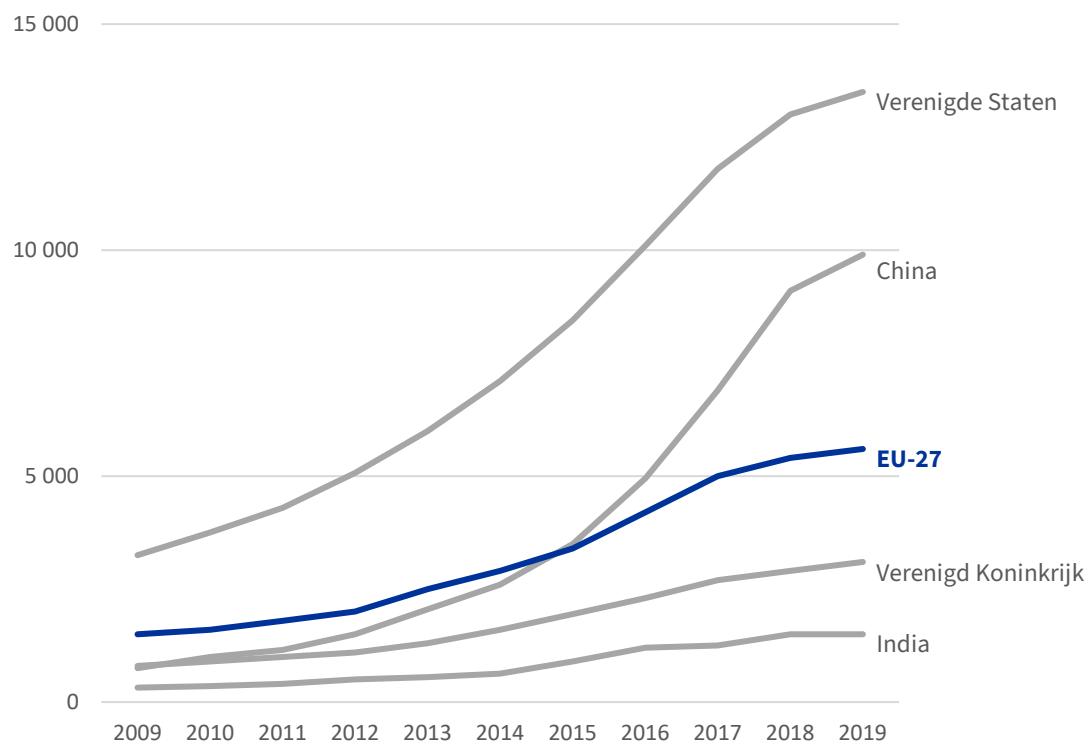


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

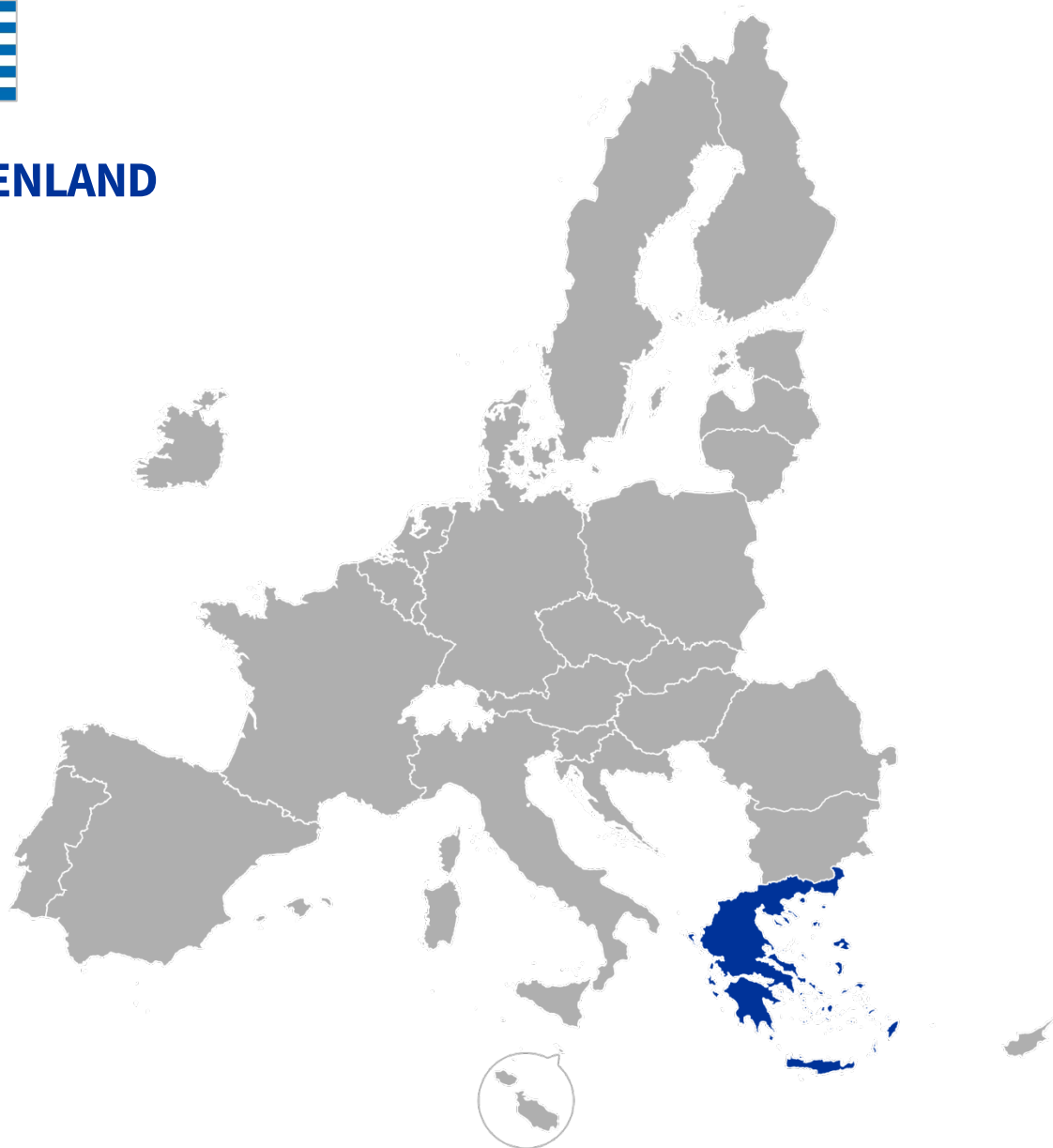
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



GRIEKENLAND



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

- Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



GRIEKENLAND

Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als België, Cyprus, Duitsland en Spanje**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Jouw land heeft een door technologie gedreven toekomst voor ogen waarin innovatie en ontwikkeling in ieders voordeel moeten zijn, gebaseerd op de Europese kernwaarden. Om te voorkomen dat AI-systemen een negatieve impact hebben op democratische processen steun jij de capaciteitsgebonden aanpak. Elk AI-systeem dat het vermogen heeft "menselijk gedrag te manipuleren" moet als een systeem met een "hoog risico" worden beschouwd, ook in contexten zoals verkiezingen.
- Je bent volledig voorstander van de voorzorgsbenadering. Voor jou zijn de sandboxtesten de belangrijkste troef hiervan, omdat ontwikkelaars in dit geval worden verplicht AI-systemen te testen onder nationaal toezicht en met kleine gebruikersgroepen. Dit zorgt voor veilige experimenten met zo min mogelijk schadelijke gevolgen.
- Je bent met name bezorgd over de mogelijke gevaren van bevooroordeelde AI. Uit studies is gebleken dat toonaangevende gezichtsherkenningssystemen veel hogere foutenpercentages vertonen bij het identificeren van mensen met een donkerdere huidskleur en vrouwen, in het bijzonder zwarte vrouwen. Vervolgens is aangetoond dat dergelijke fouten te wijten zijn aan trainingdatasets waarin witte mannelijke gezichten disproportioneel veel voorkomen, met frequente verkeerde identificatie tot gevolg.
- Het spreekt niet voor zich dat vandaag overeenstemming zal worden bereikt. **Indien nodig kan je als compromis voorstellen om een clause toe te voegen waarin staat dat de verordening over drie jaar wordt herzien.** De input van het maatschappelijk middenveld, de particuliere sector, de industrie en nationale overheden moet worden meegenomen in deze herziening.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Het blijft voor jou van essentieel belang dat zaken in verband met de nationale veiligheid, het leger of defensie buiten het toepassingsgebied van de AI-verordening blijven. Deze beleidsdomeinen vallen doorgaans buiten de bevoegdheid van de EU en de lidstaten zijn vastbesloten om de controle te behouden over besluiten op deze gebieden. Als grensland van het Schengengebied speelt Griekenland een cruciale rol bij het beheer van de buitengrenzen van de EU. De huidige systemen zijn echter verouderd en inefficiënt. De nationale autoriteiten ondervinden hierdoor zware operationele en humanitaire lasten.
- Daarom kijkt Griekenland naar AI-modellen om het grensbeheer te moderniseren en te stroomlijnen door middel van betere identificatie-, registratie- en casemanagementprocedures, in combinatie met snellere en nauwkeurigere besluitvorming. Gezien de vele migranten die aankomen in Griekenland is het hoogstnodig dat het land binnen afzienbare tijd AI-systemen op flexibele wijze kan inzetten met het oog op sneller en doeltreffender grensbeheer.
- Jouw opmerkingen:

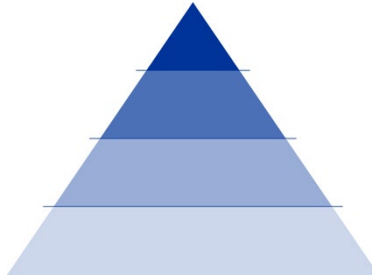
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde training-datasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland	capaciteitsgebonden	voorzorgsbenadering	over 3 jaar opnieuw bekijken	leger/defensie + nationale veiligheid	
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

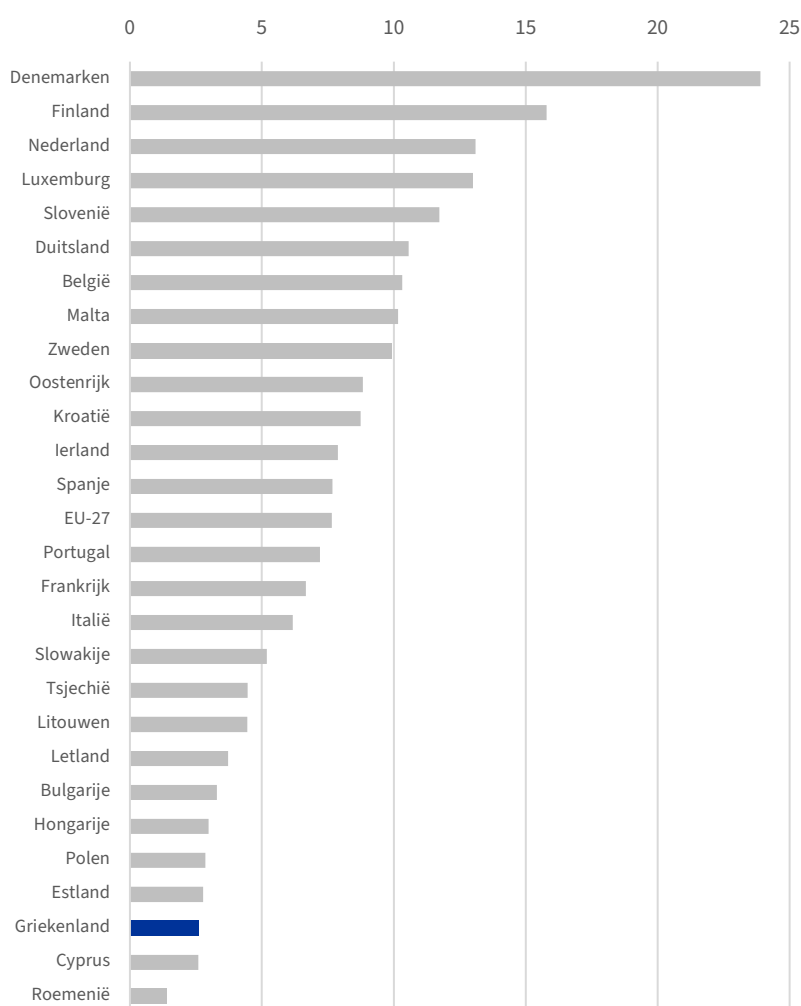
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

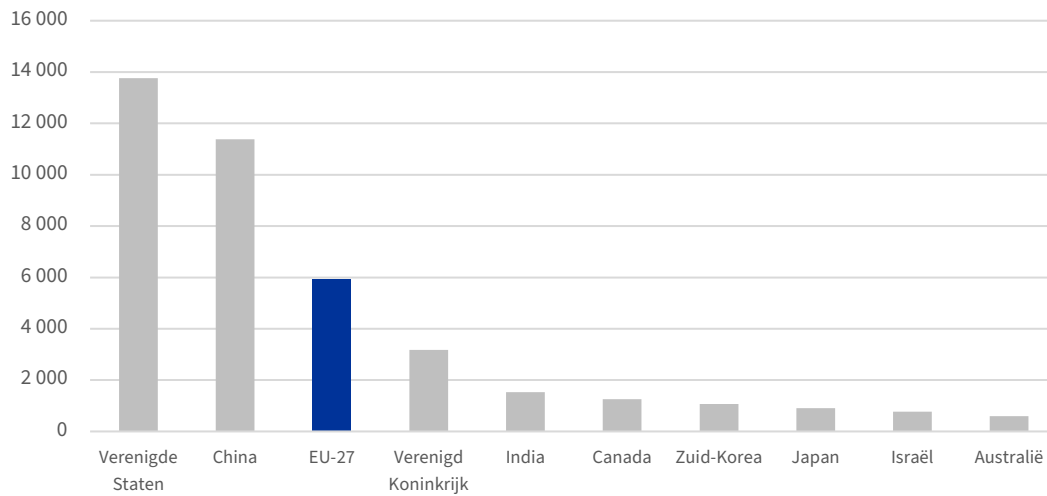
	Europese Unie	Griekenland
Bevolking	447 695 350	10 413 982
Bbp per hoofd van de bevolking in EUR	38 150	21 350
Werkloosheidspercentage	6,1%	11,1%
Percentage bedrijven dat AI-technologie gebruikt	8%	4%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	20%

Percentage bedrijven met meer dan 10 werknemers die minstens één AI-systeem gebruiken (2021)

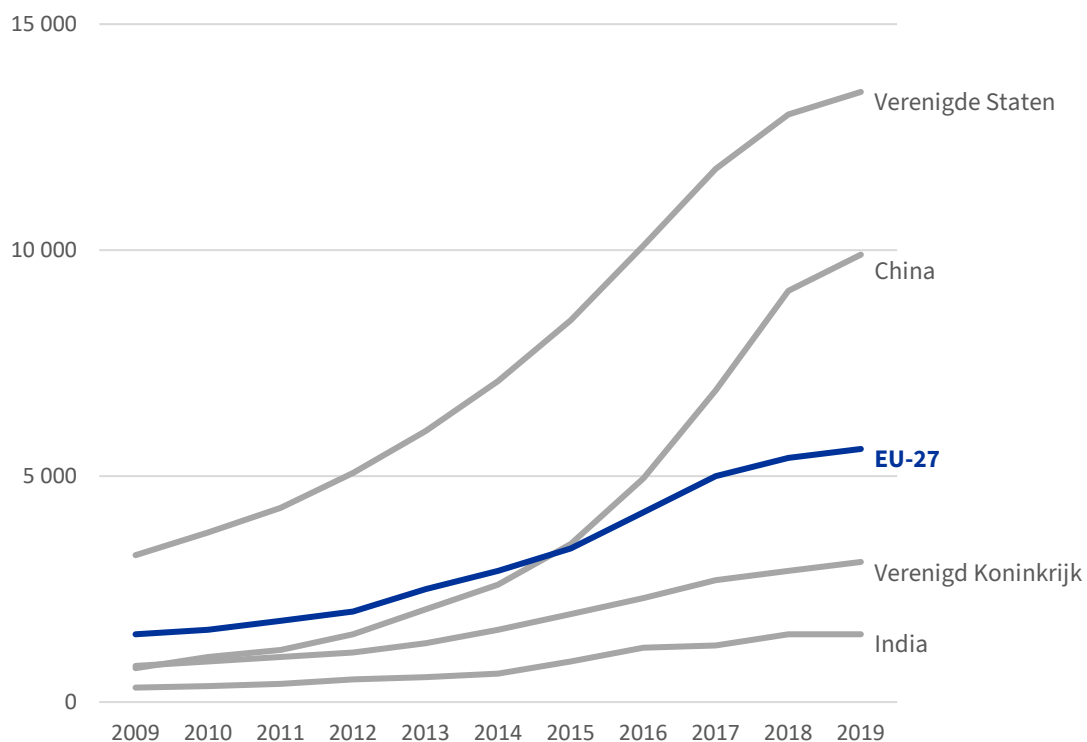


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

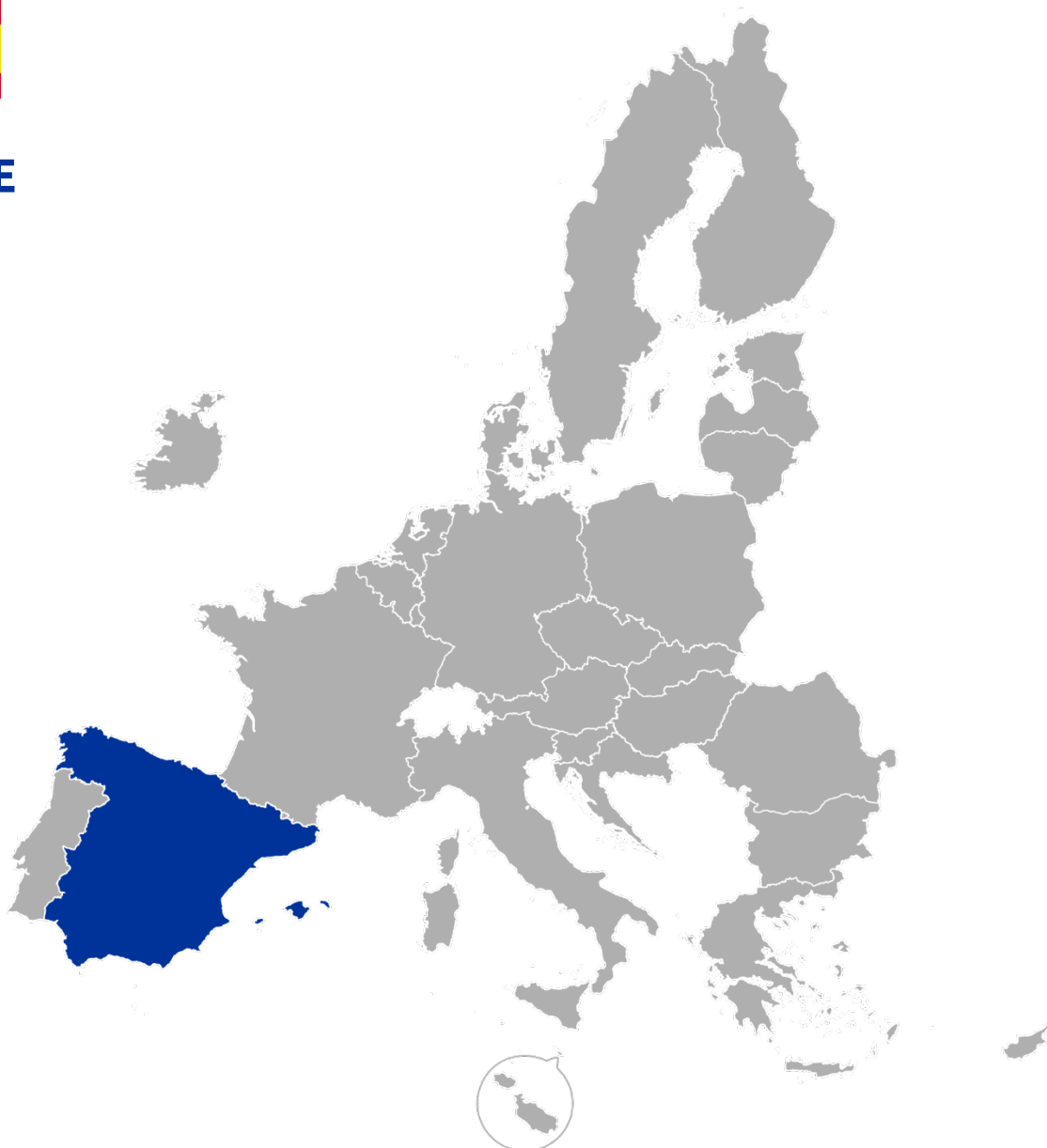
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



SPANJE



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



SPANJE

Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.



Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als België, Cyprus, Duitsland en Griekenland**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng tijdens de officiële onderhandelingen je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.

Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____.

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____.

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____.

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____.

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Jouw land geeft prioriteit aan de mensenrechten en de burgerlijke vrijheden, die aan de basis liggen van een rechtvaardige en veilige samenleving. Je bent voorstander van nationale investeringen in door bedrijven of universiteiten ontwikkelde AI-systemen, maar ethische kwesties mogen niet naar de achtergrond verdwijnen. Vandaag is het je uiteindelijke doel om toe te werken naar een kader dat zelfs het kleinste risico op door AI veroorzaakte schade, zoals misinformatie, wegneemt. Je bent je ervan bewust dat het een ambitieus doel is, maar je bent ook overtuigd van de noodzaak ervan.
- Je bent sterk gekant tegen de doelgebonden aanpak die de Commissie voorstelt. Als bijvoorbeeld een AI-systeem dat is ontworpen voor de creatieve sector om muziek te maken ook gevaarlijk luide geluiden produceert – zoals een geluidskanon – kan het nog steeds schade toebrengen, ondanks het feit dat het als "laag risico" is geclassificeerd. Uitsluitend de aangegeven intenties bekijken zonder rekening te houden met de technische capaciteiten van het systeem is gevaarlijk en onverantwoord.
- Aangezien AI nog in haar kinderschoenen staat, ben je van mening dat het de meest verantwoorde keuze is om strenge testvereisten op te leggen aan ontwikkelaars. Naarmate ontwikkelaars leren uit hun fouten, ondernemingen zich meer bewust worden van vertekende data en de technologie verder ontwikkelt, worden mogelijke afwijkingen bespreekbaar. Voorlopig vind je het echter van cruciaal belang om voorzichtig te werk te gaan en sterke waarborgen te handhaven.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Jouw regering hecht veel belang aan burgerlijke vrijheden. Tegelijkertijd ben je voorstander van een tweesporenaanpak van AI-regelgeving. Je wilt zorgen voor een hoge mate van bescherming in de civiele ruimte, maar ook flexibiliteit op het gebied van defensie en nationale veiligheid bieden. Nationale veiligheid is noodzakelijk om de burgerlijke vrijheden vrij te kunnen blijven uitoefenen, en AI kan een sleutelrol spelen bij de bescherming van de samenleving. Daarom pleit je voor de vrijstelling van het gebruik van AI op militair, defensie- en nationaal veiligheidsgebied.
- Er moet echter wel een afzonderlijke, doelgerichte regelgevingsaanpak voor deze sectoren komen om ervoor te zorgen dat AI op verantwoorde wijze wordt ontwikkeld. Anderzijds kan het opleggen van volledige transparantievereisten de operationele geheimhouding in gevaar brengen, de uitrol vertragen en de doeltreffendheid in dynamische of risicovolle omgevingen verminderen.
- Gevallen zoals de fouten door het Spaanse VioGén-systeem, waarmee een inschatting van het risico op huiselijk geweld wordt gemaakt, laten de gevaren zien van niet-transparante AI-systemen met beperkt toezicht. Tragisch genoeg werden sommige, uiteindelijk dodelijke slachtoffers verkeerd ingedeeld in de categorie met een laag risico. Deze fouten wijzen op de noodzaak van een verantwoorde ontwikkeling van AI, met name op veiligheidsgebied, maar mogen geen reden zijn om overregulering te rechtvaardigen. In plaats daarvan moet de EU een verantwoordelijkheidscultuur bevorderen, met name bij de ontwikkeling van AI voor beveiligingstoepassingen.
- Jouw opmerkingen:

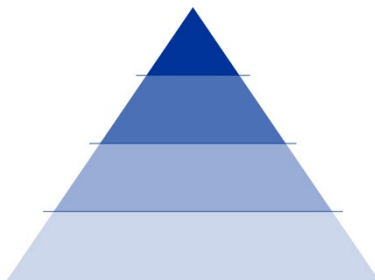
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje	capaciteitsgebonden	voorzorgsbenadering		leger/defensie + nationale veiligheid	
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

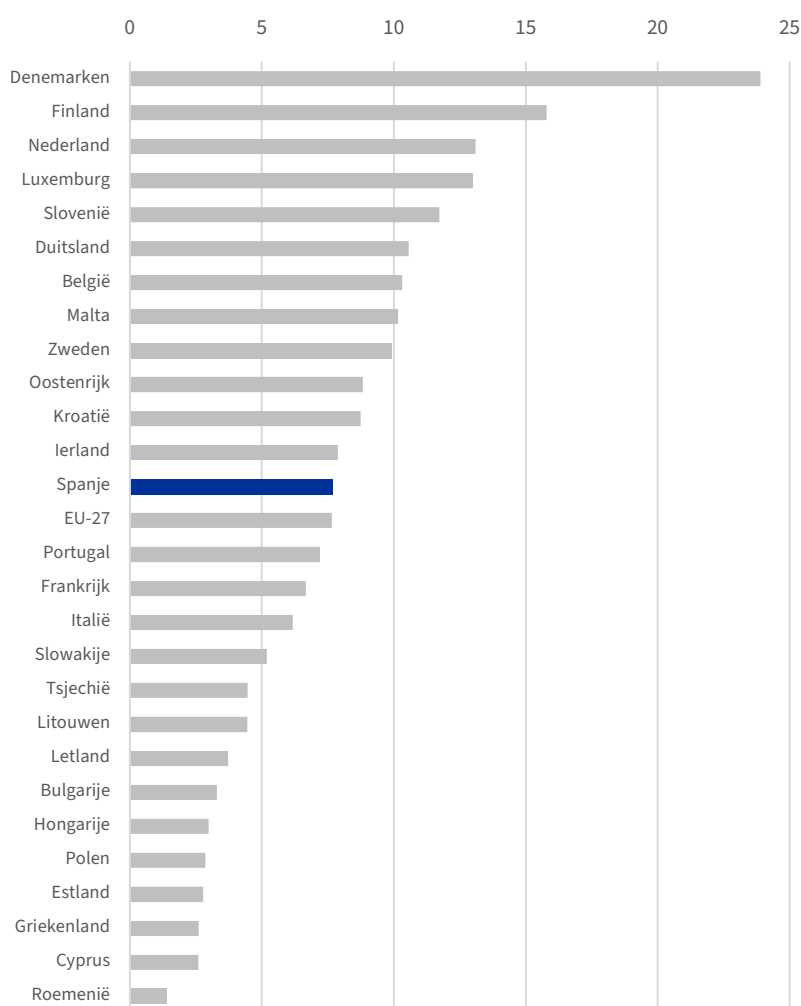
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

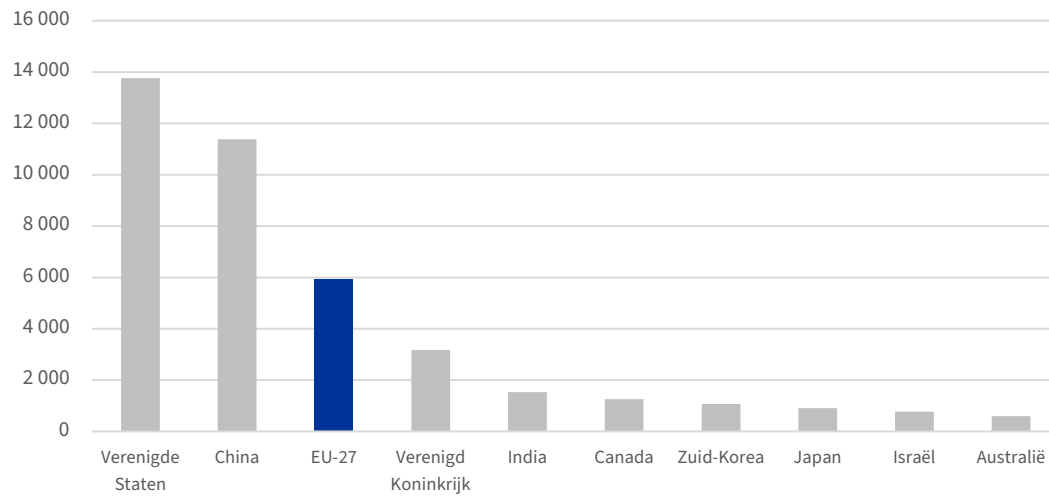
	Europese Unie	Spanje
Bevolking	447 695 350	48 085 361
Bbp per hoofd van de bevolking in EUR	38 150	30 970
Werkloosheidspercentage	6,1%	12,2%
Percentage bedrijven dat AI-technologie gebruikt	8%	9,2%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	38,7%

Percentage bedrijven met meer dan
10 werknemers die minstens één AI-systeem
gebruiken (2021)

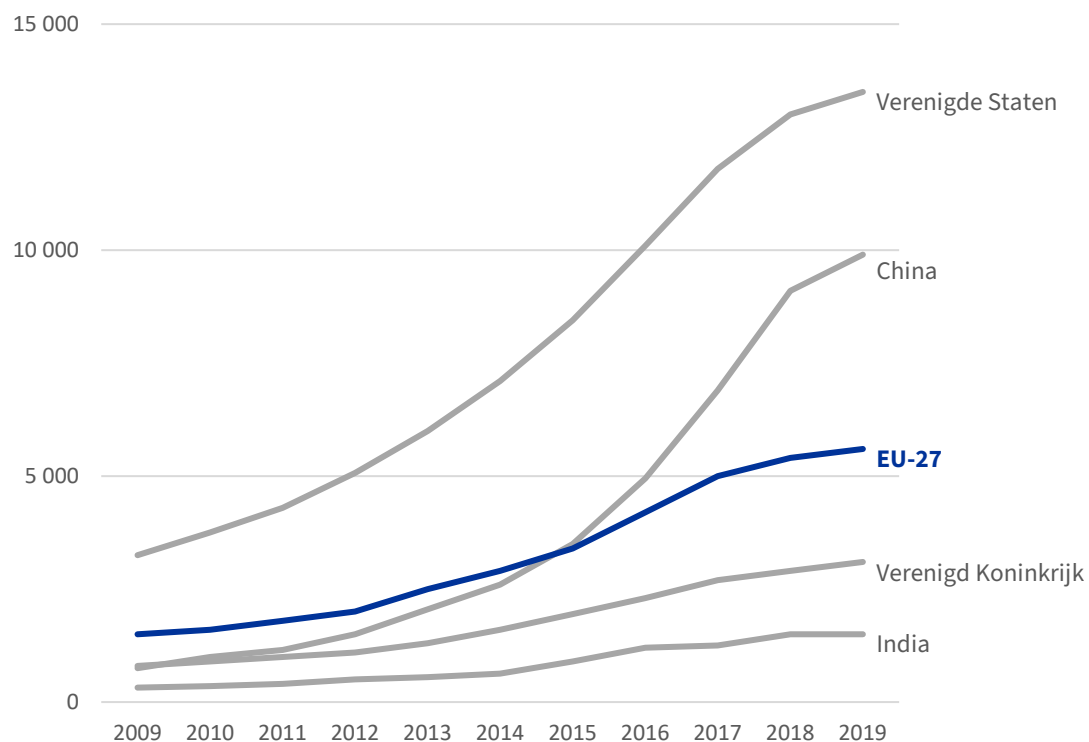


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

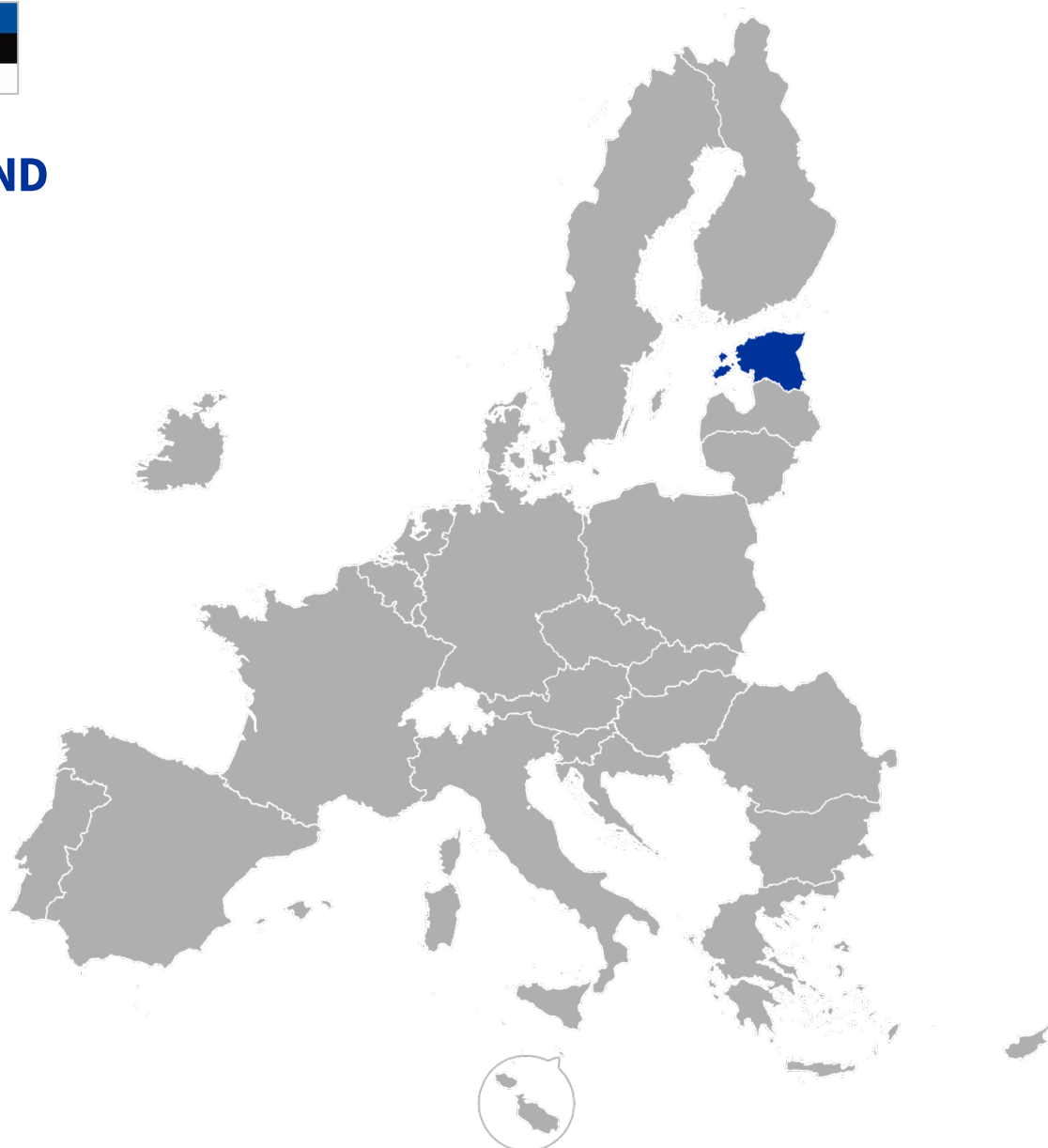
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



ESTLAND



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



ESTLAND

Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Hongarije, Letland, Litouwen en Tsjechië**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.**

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... **flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.**



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Jouw land steunt een eenduidige Europese AI-verordening. Als een van de meest gedigitaliseerde landen van de EU heeft Estland een bevolking die goed wegwijs is in de digitale wereld. In sommige andere landen zijn mensen, met name ouderen, echter minder vertrouwd met digitale hulpmiddelen en de mogelijke risico's ervan.
- Voor jou is het voorstel van de Commissie problematisch: nieuwe AI-systemen zullen in minder dan geen tijd ver boven de systemen van vandaag uitstijgen en een definitie van "hoog risico" uitsluitend op basis van de huidige technische kenmerken kan snel achterhaald zijn. Ook kan het beperken van technologie op basis van potentiële capaciteiten innovatie in de weg staan. In sectoren zoals de gezondheidszorg kan dit ertoe leiden dat waardevolle AI-oplossingen al zijn uitgesloten nog voor ze zijn overwogen. Jij bent voorstander van een doelgebonden aanpak, aangezien deze onderzoekers en ontwikkelaars meer vrijheid laat in sectoren met een laag risico, en tegelijkertijd zekerheid biedt in de vorm van testvereisten en sterkere waarborgen in sectoren met een hoger risico.
- Je steunt de flexibele benadering voor tests. Het spreekt voor zich dat AI-bedrijven veilige producten willen maken, dus overregulering is overbodig. **Je vindt het belangrijk om het testen van AI-modellen onder reële omstandigheden te bevorderen** omdat testomgevingen vaak niet volledig waarheidsgetrouw zijn.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Het blijft voor jou van essentieel belang dat zaken in verband met de nationale veiligheid, het leger of defensie buiten het toepassingsgebied van de AI-verordening blijven. Deze beleidsdomeinen vallen doorgaans buiten de bevoegdheid van de EU en de lidstaten zijn vastbesloten om de controle te behouden over besluiten op die gebieden. Dit blijkt ook uit het feit dat de EU geen verenigd EU-leger heeft.
- Daarnaast voer je aan dat voor AI-systemen die worden gebruikt voor het verzamelen van inlichtingen, cyberdefensie of gevechtscenari's niet dezelfde openbaarmakings- of regelgevingsvereisten kunnen gelden als voor civiele systemen. Dat zou de operationele geheimhouding, de nationale veiligheidsbelangen en het strategische voordeel dat zulke systemen moeten bieden namelijk in gevaar kunnen brengen.
- Jouw opmerkingen:

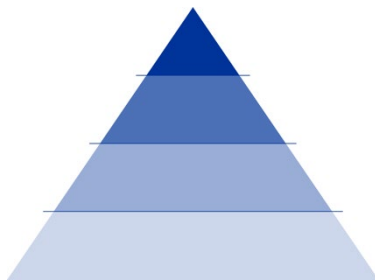
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland	doelgebonden	flexibele benadering	testen onder reële omstandigheden	leger/defensie + nationale veiligheid	
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

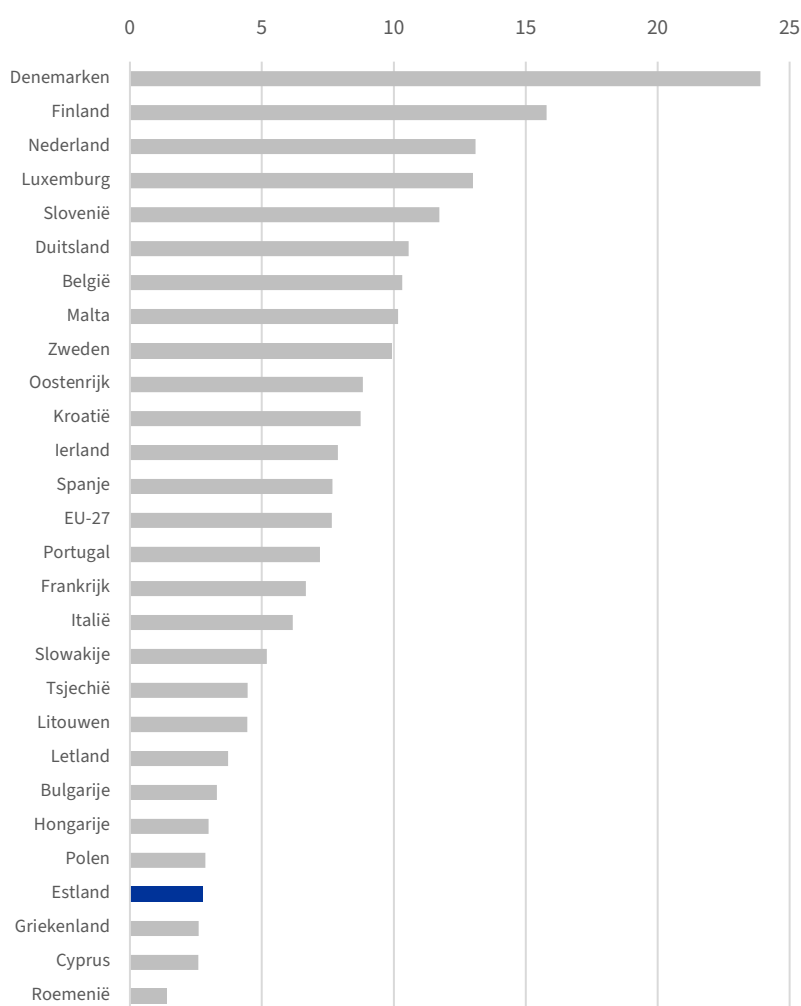
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

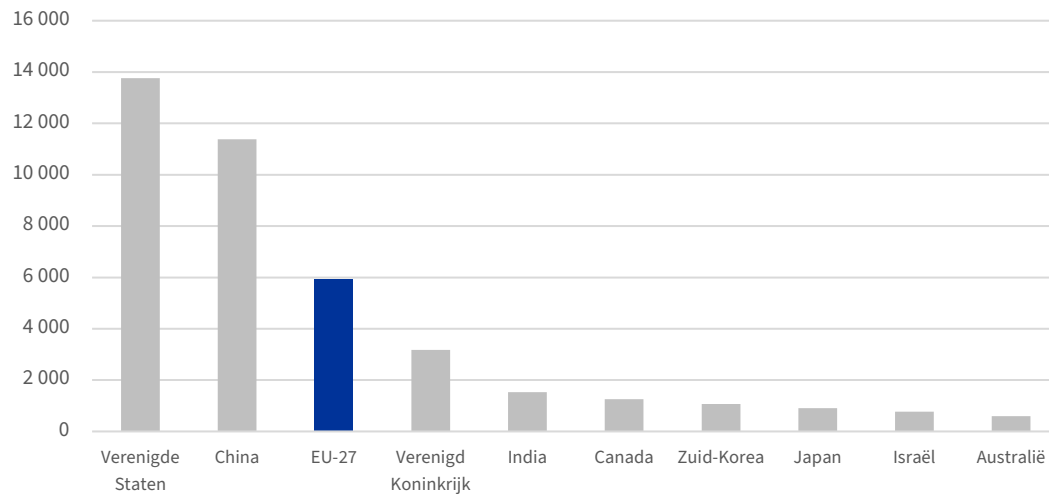
	Europese Unie	Estland
Bevolking	447 695 350	1 365 884
Bbp per hoofd van de bevolking in EUR	38 150	27 960
Werkloosheidspercentage	6,1%	6,4%
Percentage bedrijven dat AI-technologie gebruikt	8%	5,2%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	34,8%

Percentage bedrijven met meer dan
10 werknemers die minstens één AI-systeem
gebruiken (2021)

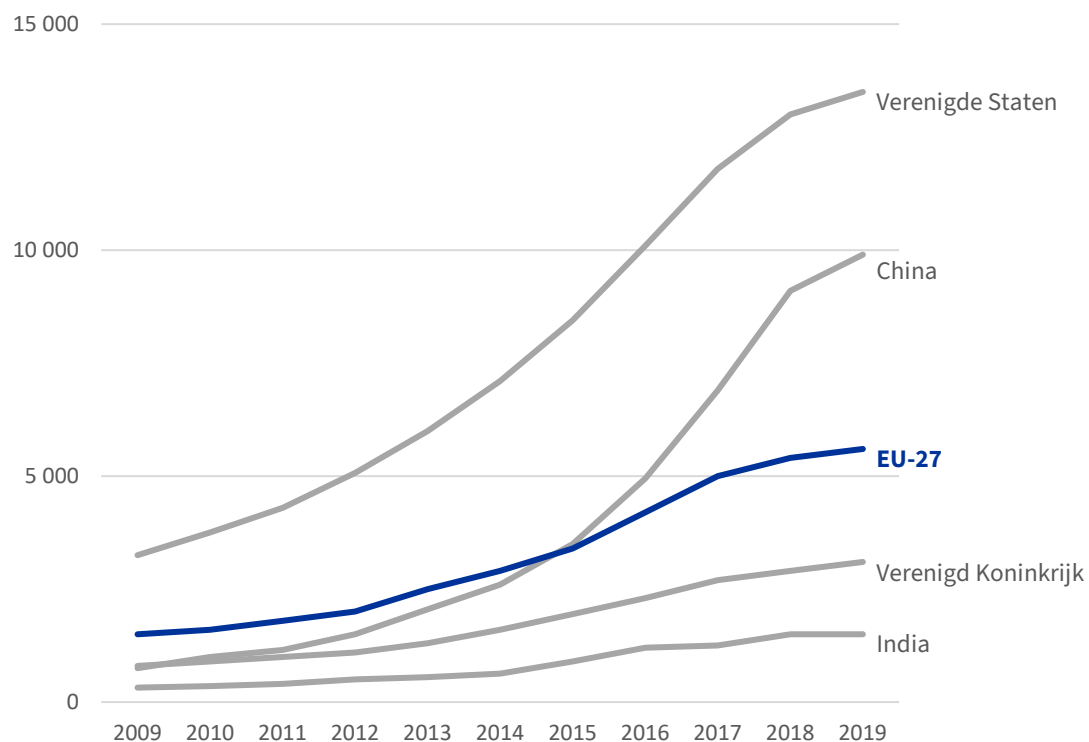


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

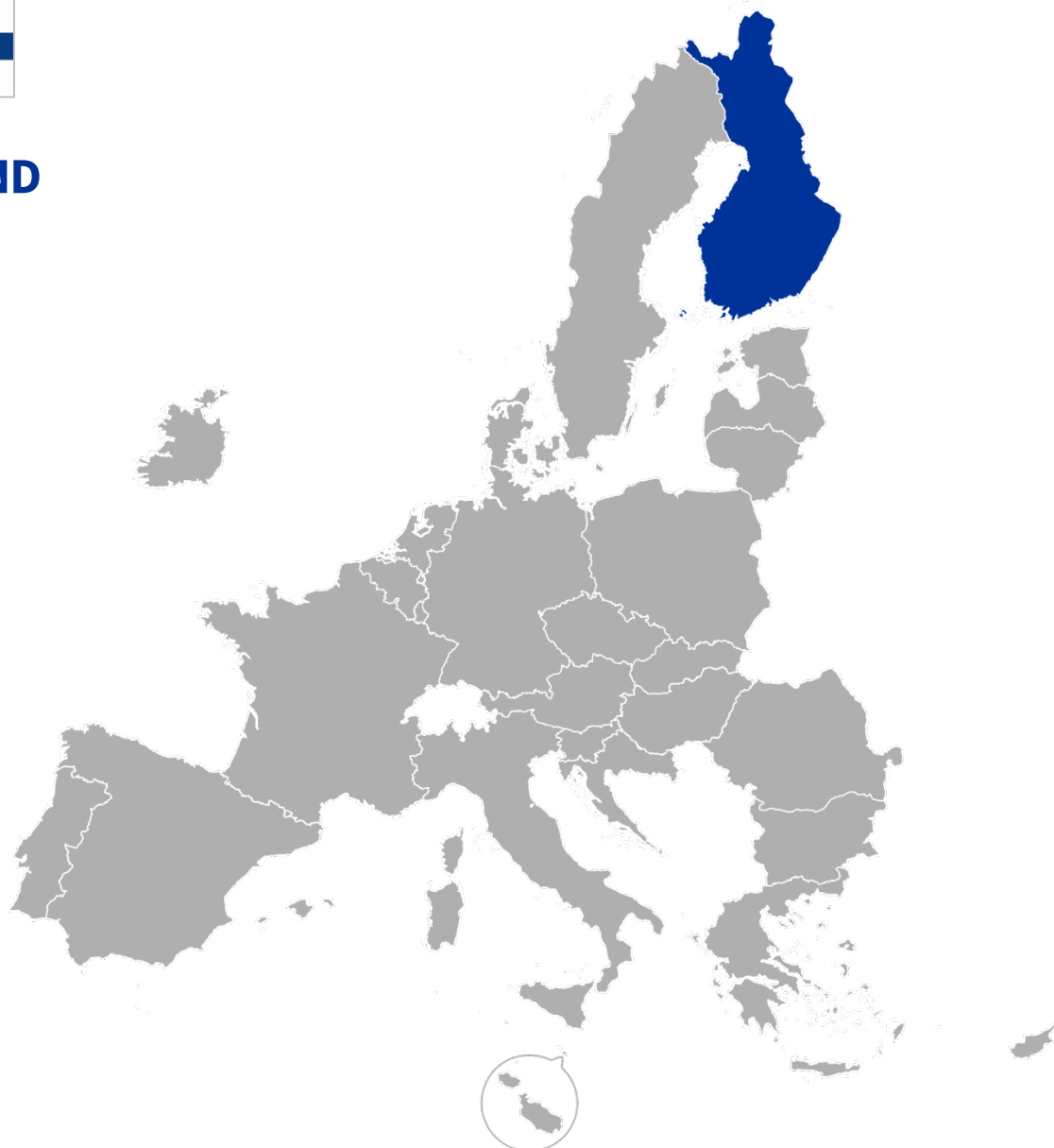
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



FINLAND



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



FINLAND

Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Ierland, Portugal en Zweden**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Jouw land is een van de koplopers van de EU op het gebied van AI, met als nationale doelstelling om een betrouwbare, veilige en innovatieve pionier van de digitale economie te worden. Je kan hiervoor op een breed draagvlak onder je bevolking rekenen, deels omdat je gratis openbaar onderwijs over AI aanbiedt.
- Je steunt de doelgebonden regelgevingsaanpak van de EU, die gericht is op de manier waarop AI wordt gebruikt in plaats van op de technische kenmerken ervan. Hiermee wordt voorkomen dat AI als inherent gevaarlijk wordt gebrandmerkt en wordt de werkelijke impact ervan vooropgesteld. Dit is volgens jou de beste manier om burgers te beschermen.
- Je vindt dat in het voorstel de nadruk moet worden gelegd op EU-steun voor innovatie door middel van een flexibele benadering, waarbij het testen van AI onder reële omstandigheden, onder overheidstoezicht, mogelijk moet zijn. Het is een goed idee om de lidstaten aan te moedigen testomgevingen op te zetten, maar niet om het gebruik ervan te verplichten, want dit zou buitenlandse investeringen kunnen afschrikken.
- Je hoopt het vandaag eens te worden over de definitieve tekst van de AI-verordening, zodat deze als leidraad kan dienen voor de betrouwbare ontwikkeling van AI in Europa, maar je bent fel gekant tegen een herzieningsclausule, aangezien dit funest zou zijn voor de stabiliteit van de regelgeving en het vertrouwen van investeerders. **Als er onzekerheid blijft, zou het verantwoordelijker zijn om de stemming uit te stellen** dan om al bij voorbaat in te plannen de regels na een tijdje om te gooien.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- In vergelijking met de rest van de wereld hinkt de EU voorlopig achterop op het gebied van AI. Start-ups vrijstellen van deze verordening zou dynamiek in het ecosysteem kunnen brengen. Start-ups hebben niet de juridische, financiële en administratieve capaciteit om aan dezelfde regels te voldoen als grote technologiebedrijven. Door strikte naleving te verlangen, loop je het risico innovatie in de beginfase af te remmen.
- Als je er niet in slaagt genoeg lidstaten ervan te overtuigen om in artikel 1 flexibele testvereisten op te nemen, wil je ten minste een aantal afwijkingen van de testvereisten voor start-ups veiligstellen. Op die manier hoeven start-ups niet meteen aan alle strenge vereisten te voldoen (als daarvoor wordt gekozen bij artikel 1), maar kunnen zij eerst aan de slag in een meer experimentele omgeving.
- Jouw opmerkingen:

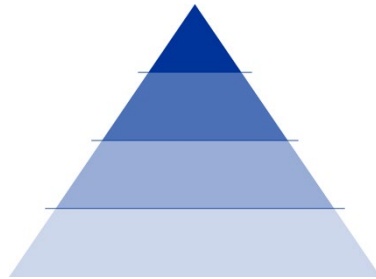
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico

Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland	doelgebonden	flexibele benadering	de stemming uitstellen	start-ups	
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

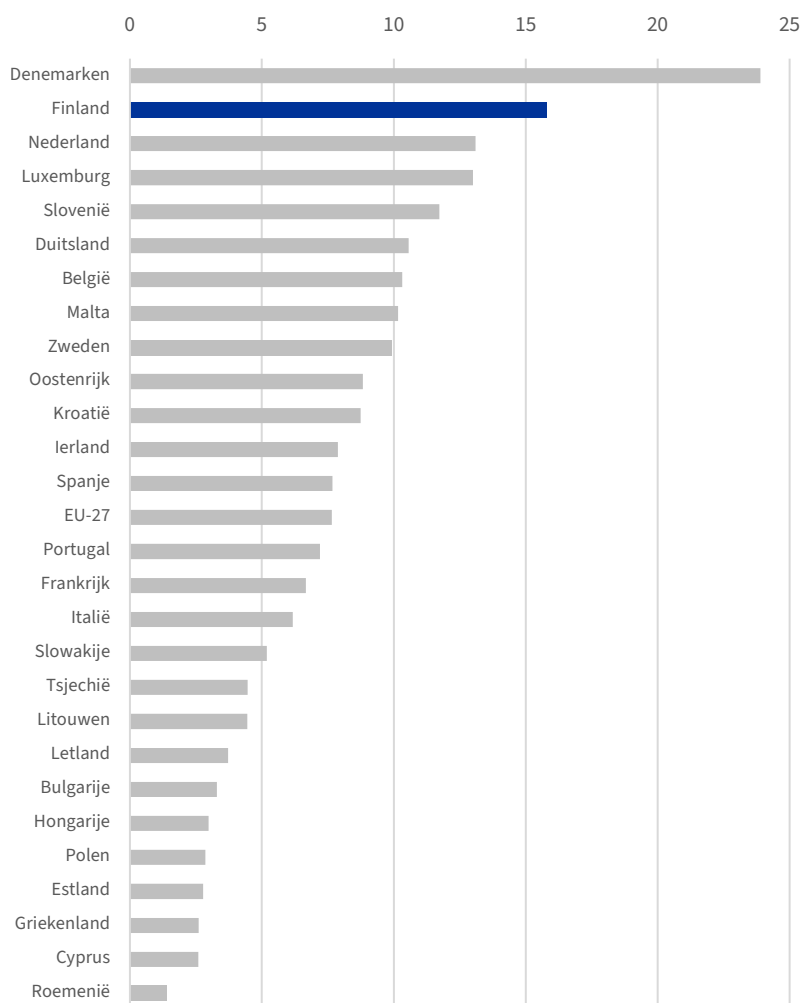
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

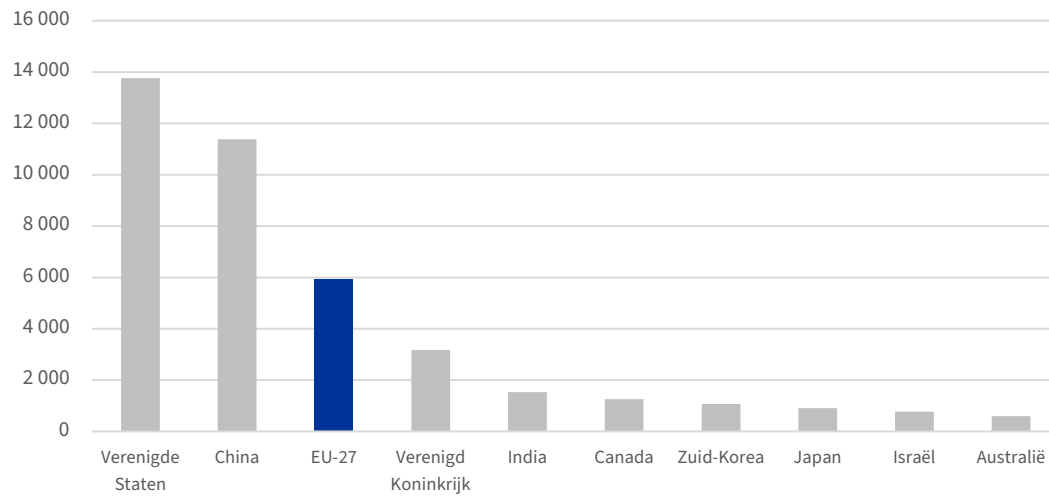
	Europese Unie	Finland
Bevolking	447 695 350	5 563 970
Bbp per hoofd van de bevolking in EUR	38 150	48 910
Werkloosheidspercentage	6,1%	7,2%
Percentage bedrijven dat AI-technologie gebruikt	8%	15,1%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	53,6%

Percentage bedrijven met meer dan
10 werknemers die minstens één AI-systeem
gebruiken (2021)

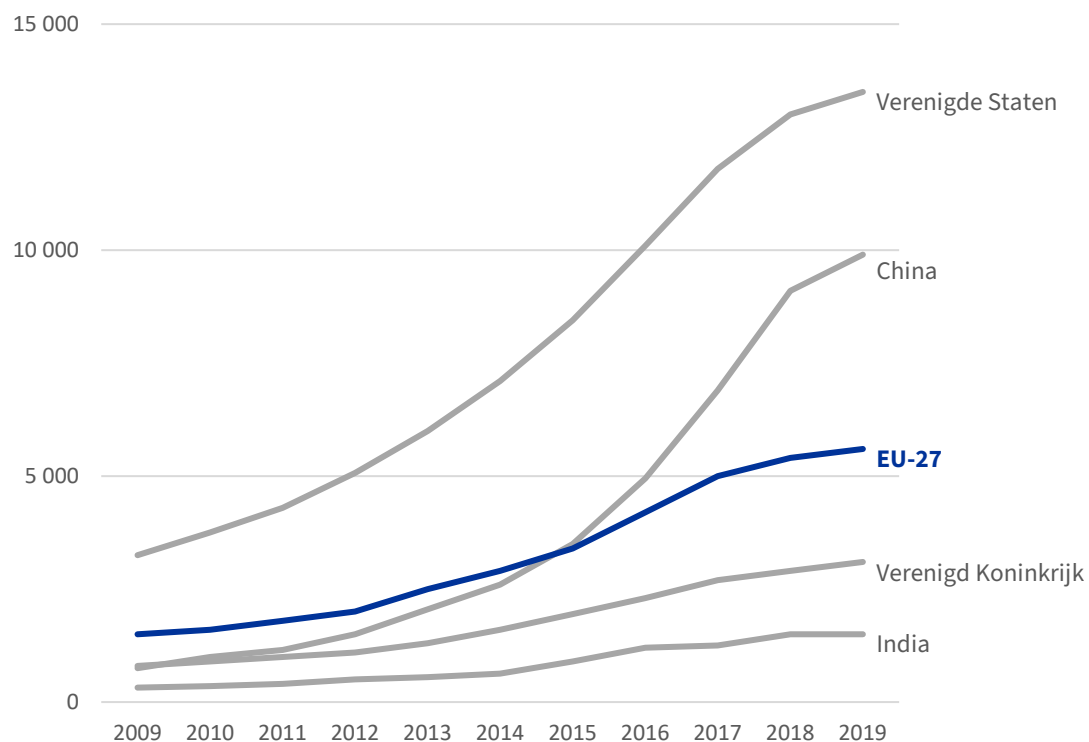


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

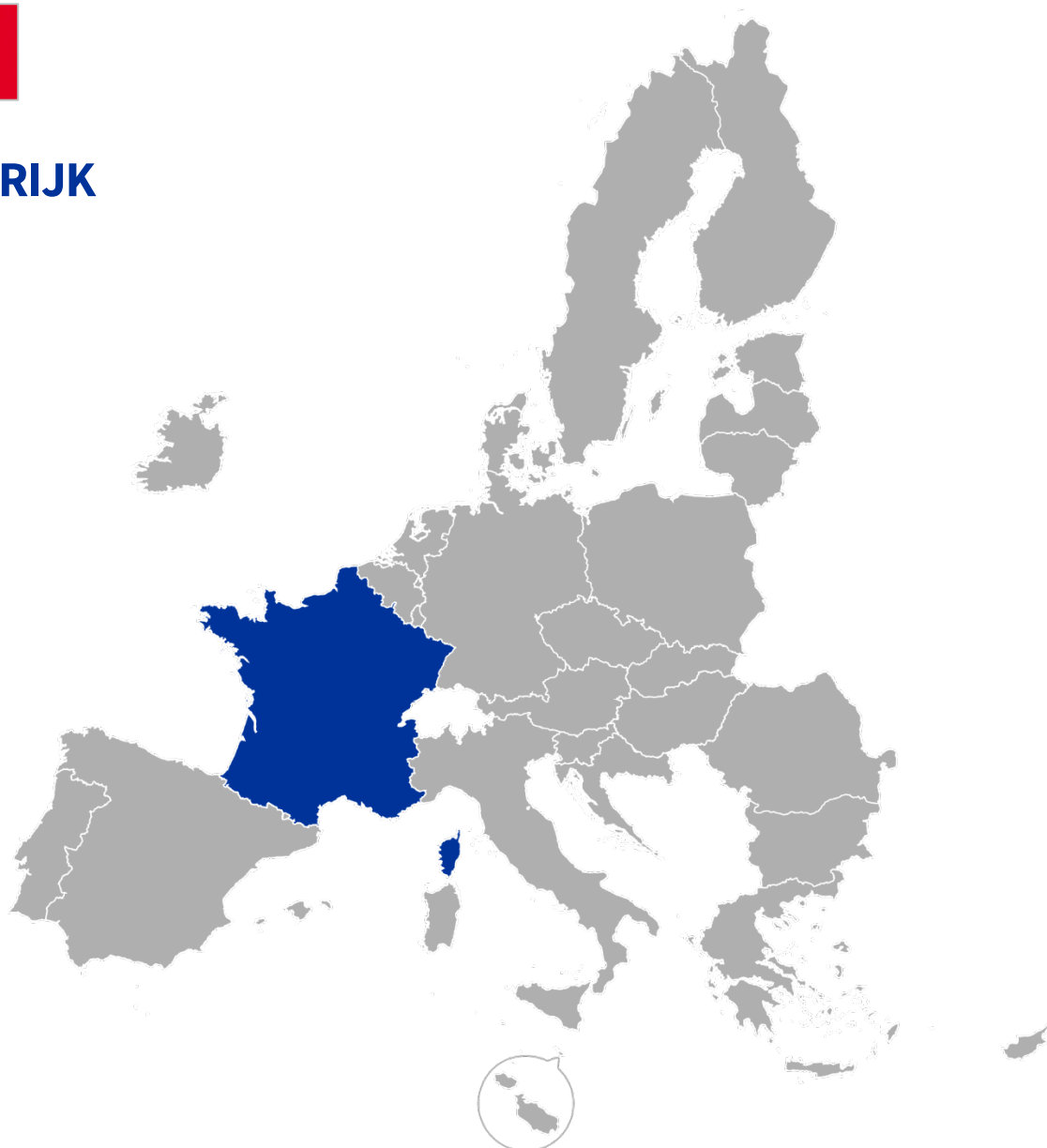
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



FRANKRIJK



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



FRANKRIJK

Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Denemarken, Oostenrijk, Polen en Roemenië**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Jouw regering heeft al een nationale strategie onder de naam "AI voor de mensheid", waaruit blijkt dat jouw belangrijkste prioriteit de mensgerichte ontwikkeling van AI is.
- Het voorstel van de Europese Commissie komt op het juiste moment. Er is dringend behoefte aan een Europese stem in het mondiale debat rond AI om de ethische en verantwoorde ontwikkeling ervan te bevorderen. Tot nu toe heeft AI zich grotendeels ongecontroleerd mogen ontplooiën en hoewel dit niet zonder risico is, mag AI niet worden gedemoniseerd. AI blijft een belangrijke stuwende kracht achter innovatie en groei.
- Hoewel je achter de doelgebonden aanpak staat omdat deze duidelijkheid biedt aan investeerders, **moet ook het onderwijs in de lijst van gevoelige sectoren worden opgenomen**. Als een AI-systeem leerlingen verkeerd indeelt of een foute score geeft, kan dit leiden tot discriminatie en meer ongelijkheid. Structurele discriminatie kan niet worden bestreden zonder eerlijke toegang tot onderwijs.
- Uit ervaring weet je dat regelgevende stimulansen nuttig zijn voor ontwikkelaars, maar zonder overheidsfinanciering dreigen innovatieve ideeën op niets uit te lopen, met kennisvlucht en bedrijfsmigratie tot gevolg. De EU moet dus niet alleen een kader vaststellen, maar zich er ook toe verbinden aanzienlijke overheidsinvesteringen in AI te doen.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)

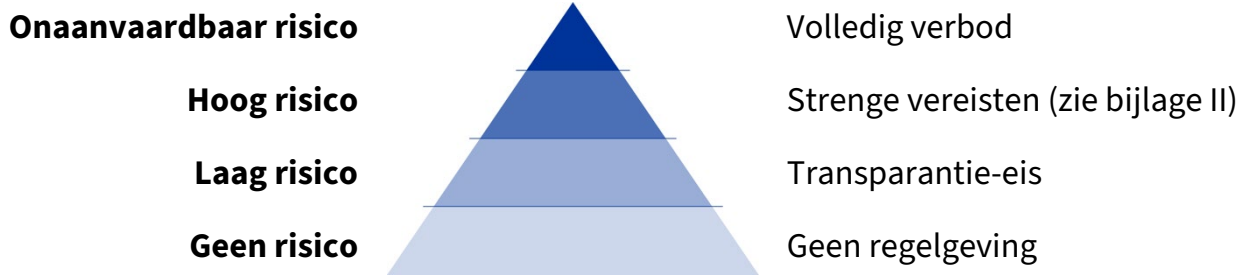


... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Jouw regering hecht veel belang aan burgerlijke vrijheden. Tegelijkertijd ben je voorstander van een tweesporenaanpak van AI-regelgeving. Je wilt zorgen voor een hoge mate van bescherming in de civiele ruimte, maar ook flexibele normen op het gebied van defensie en nationale veiligheid bieden. Nationale veiligheid is noodzakelijk om de burgerlijke vrijheden vrij te kunnen blijven uitoefenen, en AI kan een sleutelrol spelen bij de bescherming van democratische samenlevingen tegen opkomende dreigingen.
- Daarom pleit je ervoor om AI-toepassingen op militair, defensie- en nationaal veiligheidsgebied van deze verordening vrij te stellen en de regels voor deze sectoren in plaats daarvan onder te brengen in een specifieke verordening. Hoewel het van cruciaal belang blijft dat AI op verantwoorde wijze wordt ontwikkeld, kan het opleggen van volledige transparantievereisten aan defensiegerelateerde AI-systemen de operationele geheimhouding in gevaar brengen, de uitrol vertragen en de doeltreffendheid in dynamische of risicovolle omgevingen verminderen.
- Daarnaast wil je verduidelijken dat voor alle AI-systemen strenge testvereisten moeten gelden, dus ook voor systemen die van buiten de EU worden ingevoerd. Je stelt voor de formulering te wijzigen naar: "Artikel 1 is van toepassing op alle AI-systemen die in de EU worden ingezet, ...".
- Jouw opmerkingen:

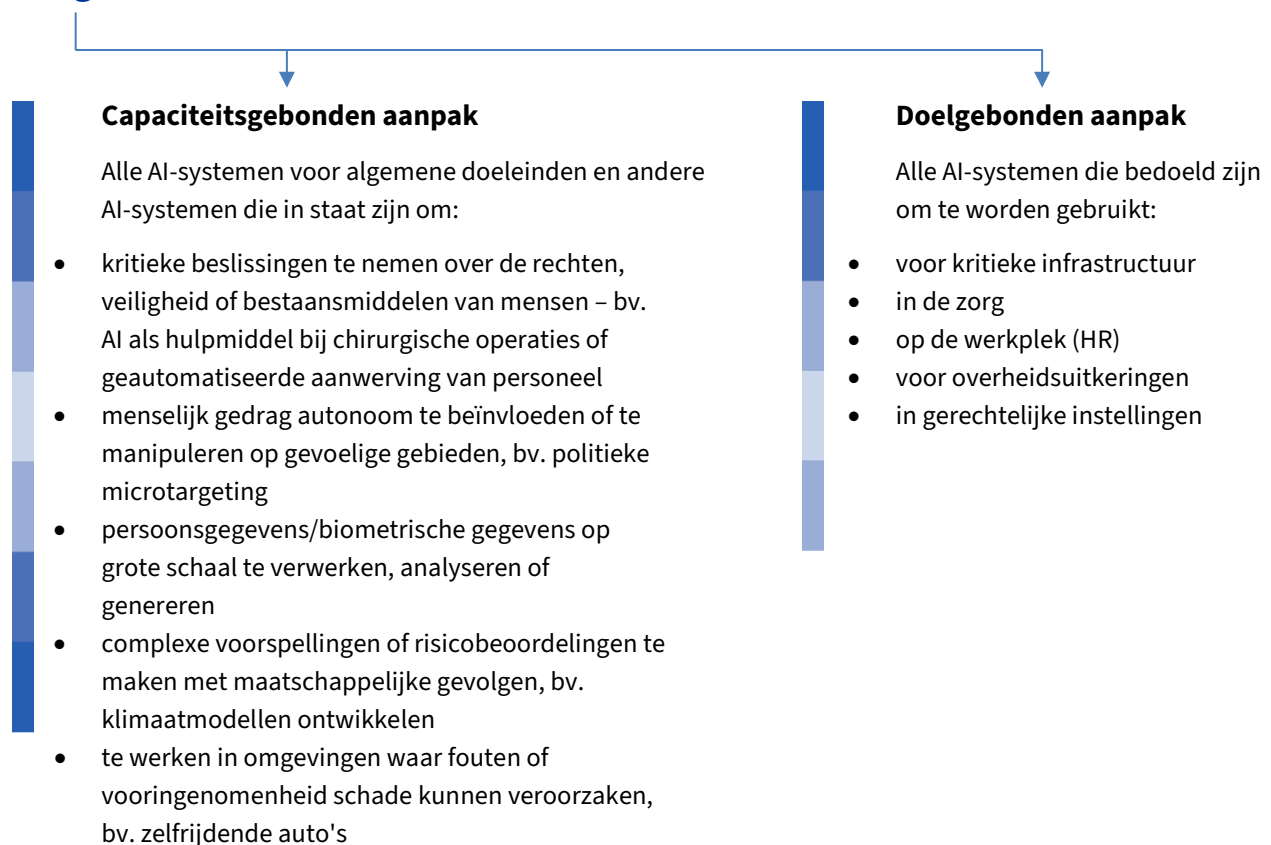
Bijlage I: Risicoclassificatie



Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk	doelgebonden	voorzorgsbenadering	onderwijs toevoegen	leger/defensie + nationale veiligheid	
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

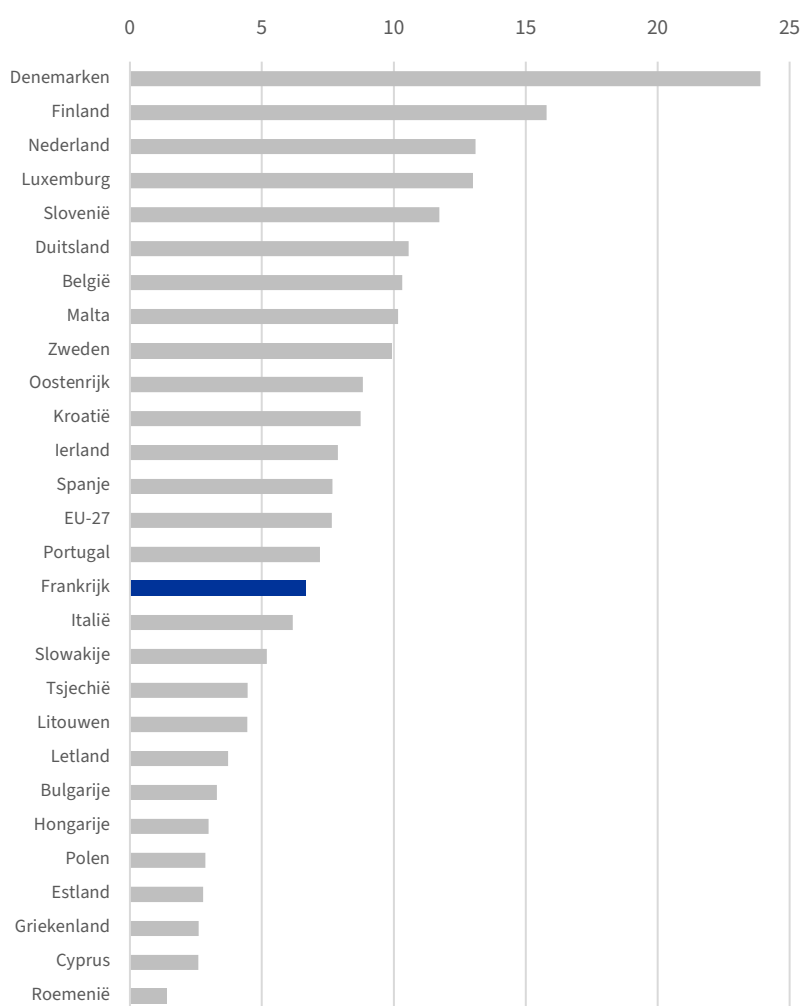
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

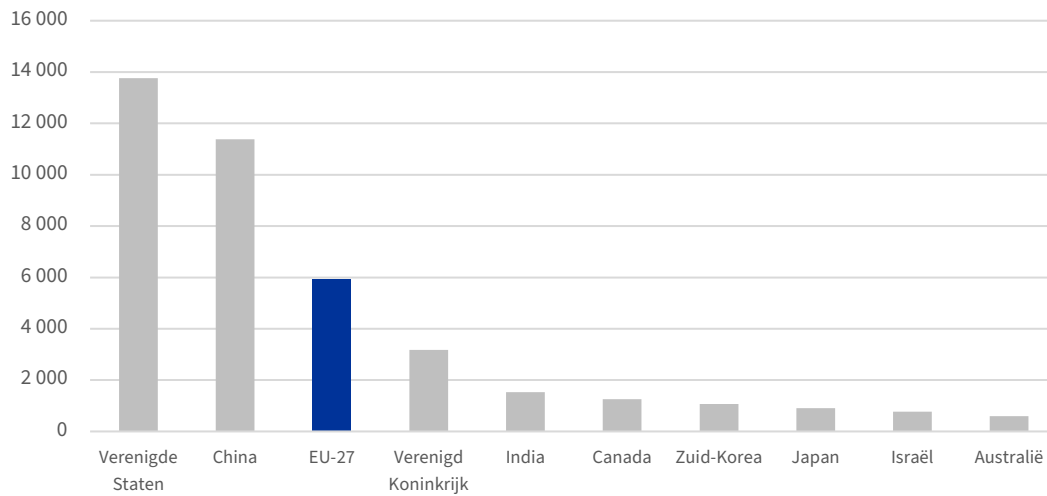
	Europese Unie	Frankrijk
Bevolking	447 695 350	68 277 210
Bbp per hoofd van de bevolking in EUR	38 150	41 330
Werkloosheidspercentage	6,1%	7,3%
Percentage bedrijven dat AI-technologie gebruikt	8%	5,9%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	30,3%

Percentage bedrijven met meer dan
10 werknemers die minstens één AI-systeem
gebruiken (2021)

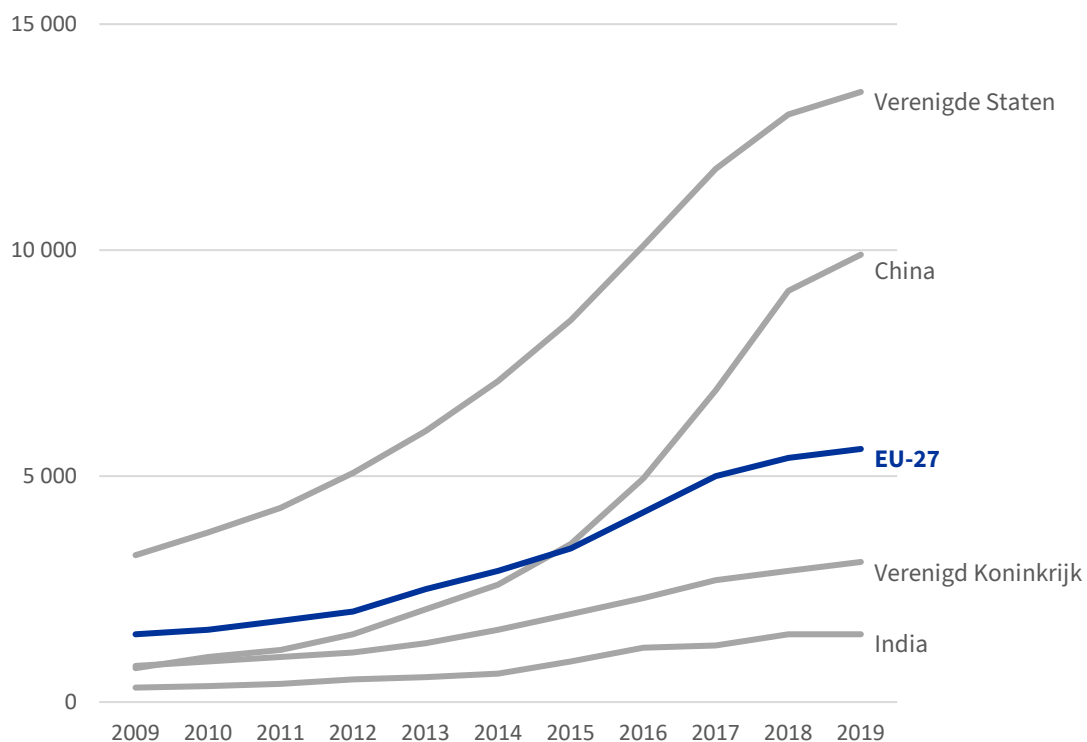


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

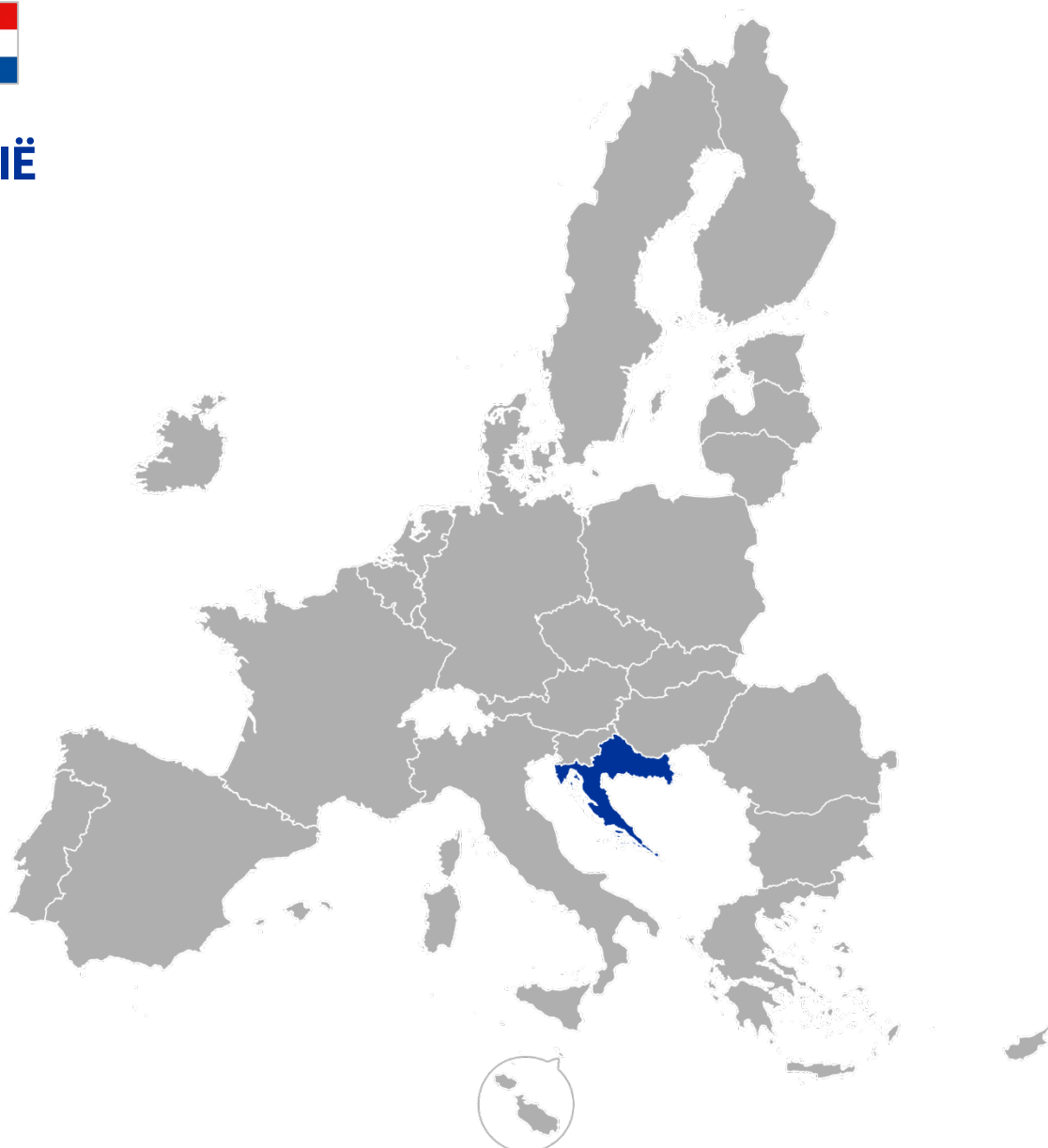
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



KROATIË



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Italië, Nederland en Slovenië**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Het aantal start-ups in Kroatië neemt snel toe. Consultancybureaus voorspellen dat door AI gestuurde automatisering tegen 2025 tot 16% van het bbp zou kunnen opleveren, voornamelijk door productiviteitswinsten. Tot wel 52% van de werkuren zou geautomatiseerd kunnen worden. Hoewel dit veelbelovend is voor de economische groei, baart de mogelijk stijgende werkloosheid in de productiesector zorgen.
- Je steunt de doelgebonden risicoclassificatie van de Commissie omdat dit een realistische kijk is op de risico's van AI in specifieke contexten, bijvoorbeeld gezichtsherkenning voor het ontgrendelen van je telefoon tegenover het gebruik ervan in massasurveillance.
- Maar je hecht wel veel waarde aan transparantie. Alle AI-systemen moeten gebruikers duidelijk informeren over hun doel, ongeacht hun risiconiveau. Dit zorgt voor aansprakelijkheid en bescherming van de individuele rechten.
- Je bent ook voorstander van strenge normen voor de ontwikkeling van AI-systemen, maar je vindt de huidige richtlijnen voor de risicobeoordeling te vaag. Om kwetsbare groepen te beschermen, pleit je voor een verplichte beoordeling van het effect op de mensenrechten van AI-systemen met een hoog risico. Deze beoordeling moet in overleg met de getroffen gemeenschappen worden ontwikkeld.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Volgens jou moet de Europese AI-verordening als katalysator dienen voor de ontwikkeling van AI in alle EU-lidstaten. Sommige kapitaalkrijke landen hebben al toonaangevende AI-bedrijven, terwijl andere landen nog bezig zijn met het tot stand brengen van een ecosysteem voor start-ups. Het hele idee achter start-ups is om op korte termijn en met beperkte middelen disruptieve technologie te creëren. Dat is een haast onmogelijke opgave wanneer je omkomt in administratieve rompslomp.
- Daarom zou je graag zien dat start-ups met minder dan 10 werknemers en een jaaromzet van minder dan 2 miljoen euro worden vrijgesteld van de Europese AI-verordening.
- Als het je niet lukt voldoende lidstaten van je standpunt te overtuigen, dan zou je extra Europese middelen voor start-ups een goed compromis vinden. Op die manier zouden zij niet worden belemmerd door de administratieve en budgettaire lasten die de verordening met zich meebrengt.
- Jouw opmerkingen:

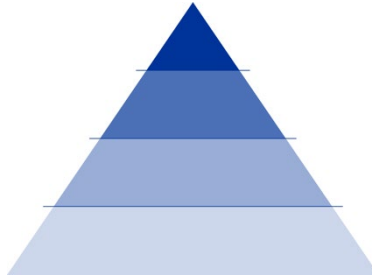
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde training-datasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië	doelgebonden	voorzorgsbenadering		start-ups	
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

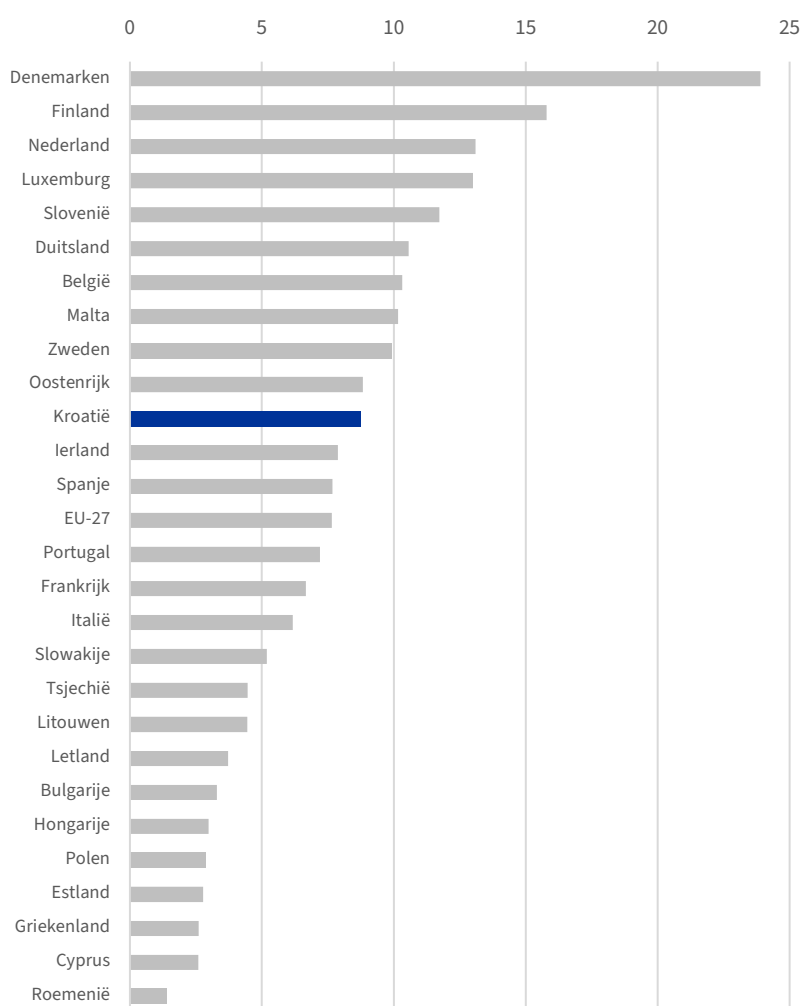
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

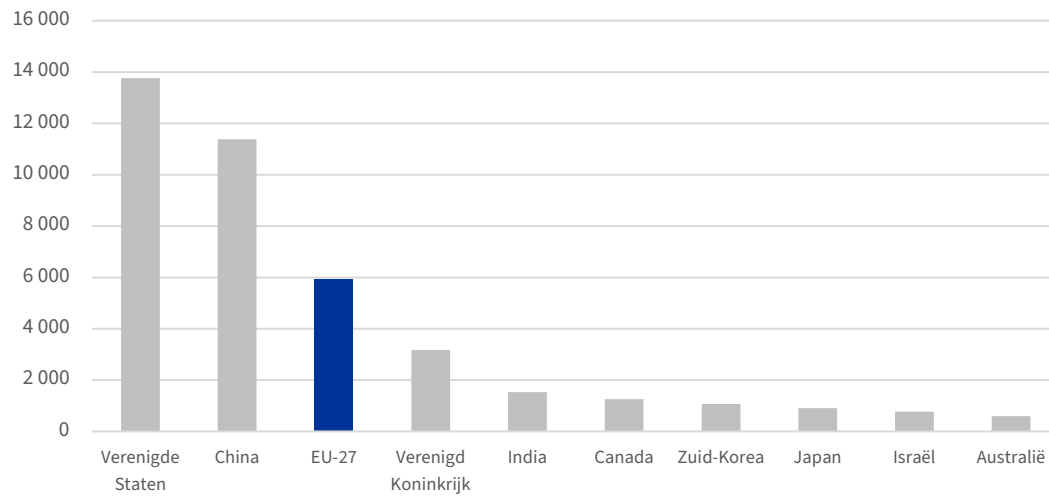
	Europese Unie	Kroatië
Bevolking	447 695 350	3 850 894
Bbp per hoofd van de bevolking in EUR	38 150	19 800
Werkloosheidspercentage	6,1%	6,1%
Percentage bedrijven dat AI-technologie gebruikt	8%	7,9%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	25%

Percentage bedrijven met meer dan 10 werknemers
die minstens één AI-systeem gebruiken (2021)

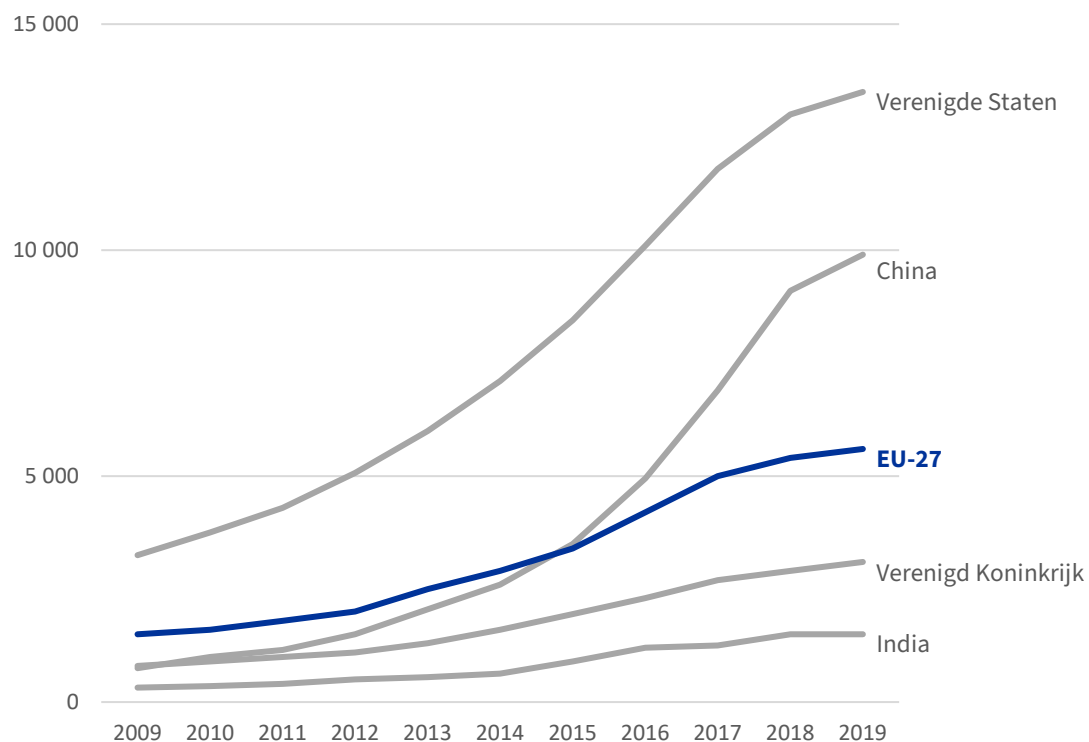


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

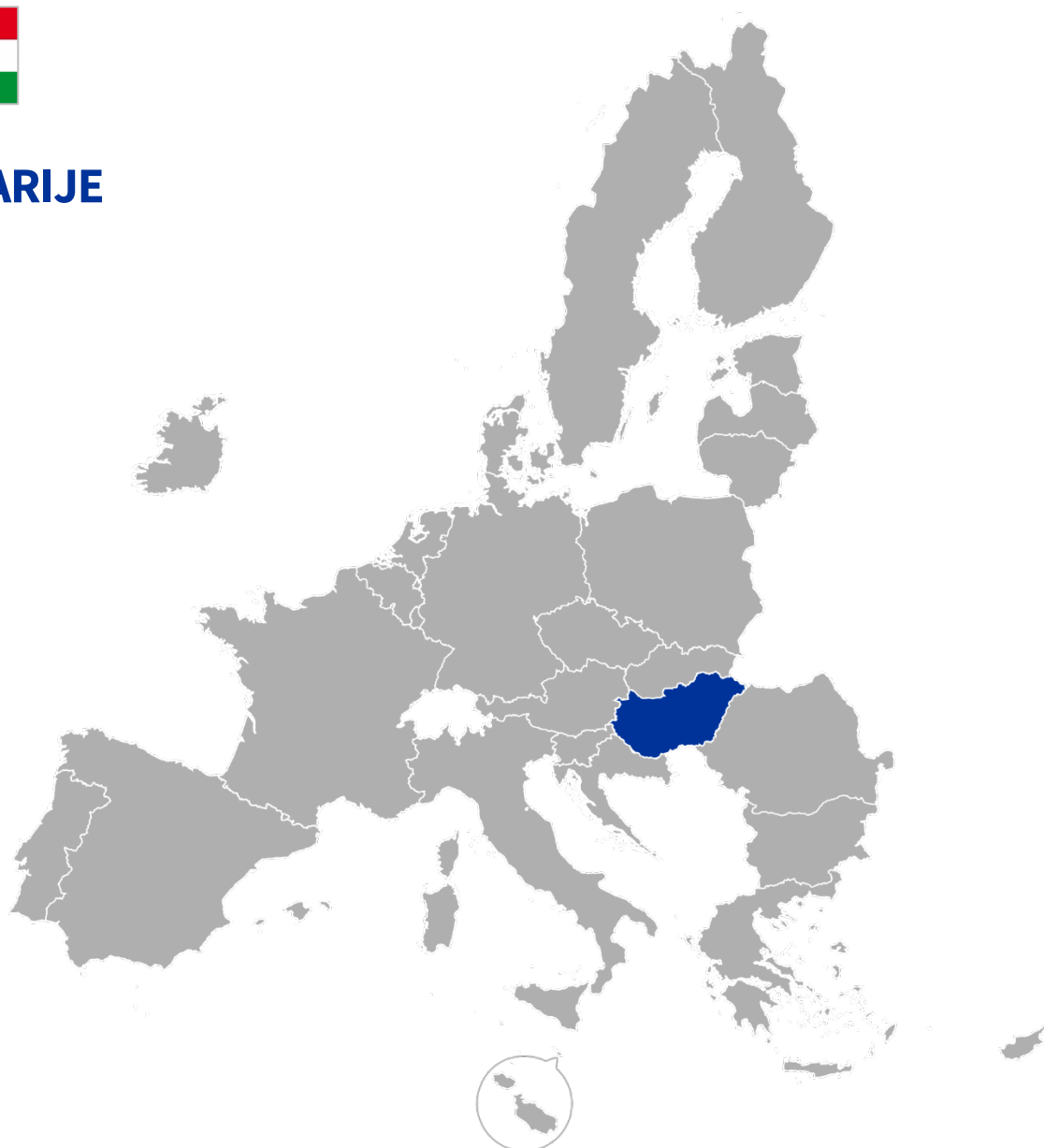
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



HONGARIJE



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

- Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Estland, Letland, Litouwen en Tsjechië**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Jouw regering heeft onlangs een nationale AI-strategie ontwikkeld. Een van de pijlers van de strategie is om de samenleving erop voor te bereiden om doeltreffend om te gaan met de onvermijdelijke veranderingen als gevolg van AI. Ook wil je dat de Hongaarse industrie en universiteiten hun voordeel uit deze technologie kunnen halen.
- Je bent ervan overtuigd dat een groter gebruik van AI een aanzienlijke toegevoegde waarde kan opleveren voor het bbp en die middelen kan je land goed gebruiken voor dringende investeringen in zijn openbare infrastructuur.
- Voor het classificeren van AI-systemen met een hoog risico vind je de capaciteitsgebonden aanpak beter geschikt dan de doelgebonden aanpak die de Commissie voorstelt. Als hetzelfde AI-systeem in de ene context een laag risico heeft en in de andere een hoog risico, kan dit tot verwarring en rechtsonzekerheid leiden voor gebruikers, ontwikkelaars en investeerders.
- Je ziet een belangrijke rol als waakhond voor eerlijke concurrentie weggelegd voor de Europese Unie, zodat Europese ondernemingen, Hongaarse ondernemingen inclusief, concurrerend blijven op de interne markt en in de mondiale waardeketens. Hongarije wil een aantrekkelijke bestemming zijn voor buitenlandse investeringen. Daarom ben jij voorstander van regelgeving met enkel minimumvereisten en -normen, om ontwikkelaars lasten en kosten te besparen.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Je pleit ervoor dat militaire en defensiegerelateerde toepassingen vrijgesteld worden van de AI-verordening en benadrukt hierbij dat overregulering op dit gebied de snelle inzet van AI-technologieën in conflictgebieden of tijdens crisisresponsituaties kan belemmeren. Gezien de aanhoudende regionale instabiliteit, is flexibiliteit bij het gebruik van dergelijke technologieën van cruciaal belang.
- Als land dat grenst aan Oekraïne – dat momenteel in oorlog is met Rusland – maakt Hongarije zich ernstig zorgen over zijn nationale veiligheid en de veiligheid van zijn burgers. Daarnaast staat Hongarije als grensland van het Schengengebied voor grote uitdagingen in verband met het beheer van de toegenomen migratiestromen, met name in het zuiden.
- Hongarije hoopt dan ook dat het land op meer begrip en steun van zijn medelidstaten kan rekenen om het hoofd te bieden aan deze unieke geopolitieke en veiligheidsuitdagingen.
- Hoewel sommige landen wantrouwig staan tegenover geïmporteerde AI-systemen, met name uit China, ben jij van mening dat dit eerder politiek gemotiveerd is dan gebaseerd op echte bedreigingen. Partnerschappen met Chinese technologiebedrijven zijn in het verleden al goed uitgedraaid voor Hongarije en volgens jou zou het oplossen van de handelsspanningen met China de EU geen windeieren leggen.
- Jouw opmerkingen:

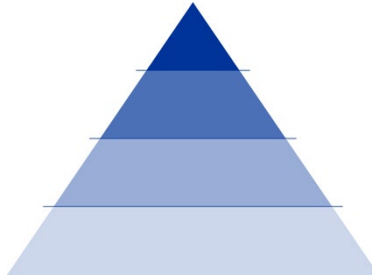
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde training-datasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije	capaciteitsgebonden	flexibele benadering		leger/defensie	
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

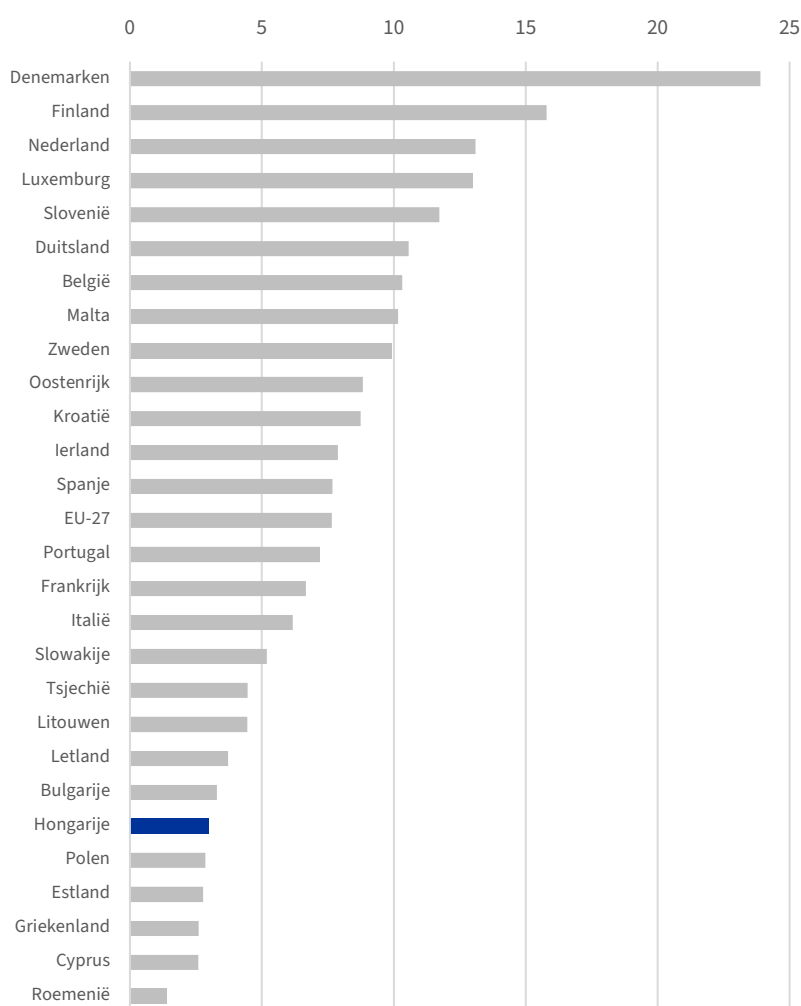
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

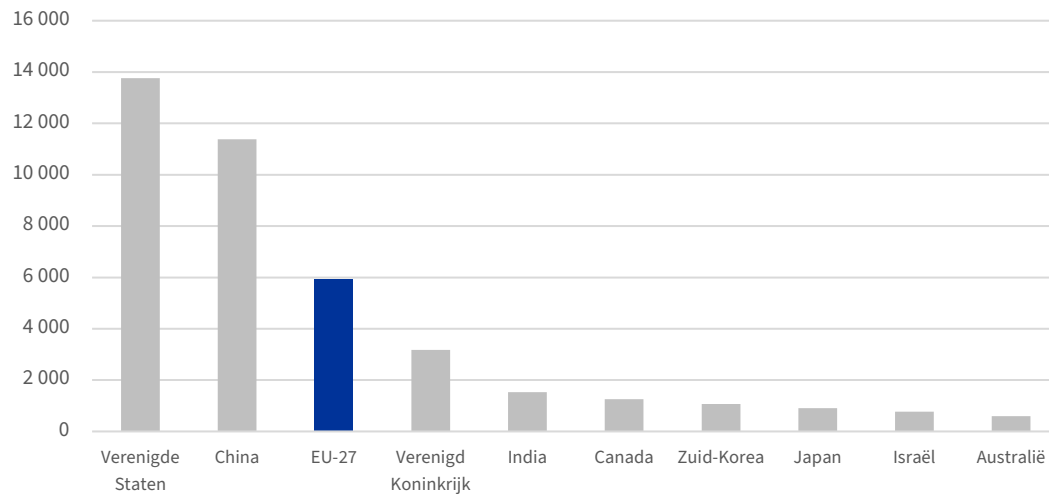
	Europese Unie	Hongarije
Bevolking	447 695 350	9 599 744
Bbp per hoofd van de bevolking in EUR	38 150	20 630
Werkloosheidspercentage	6,1%	4,1%
Percentage bedrijven dat AI-technologie gebruikt	8%	3,7%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	28,1%

Percentage bedrijven met meer dan 10 werknemers die minstens één AI-systeem gebruiken (2021)

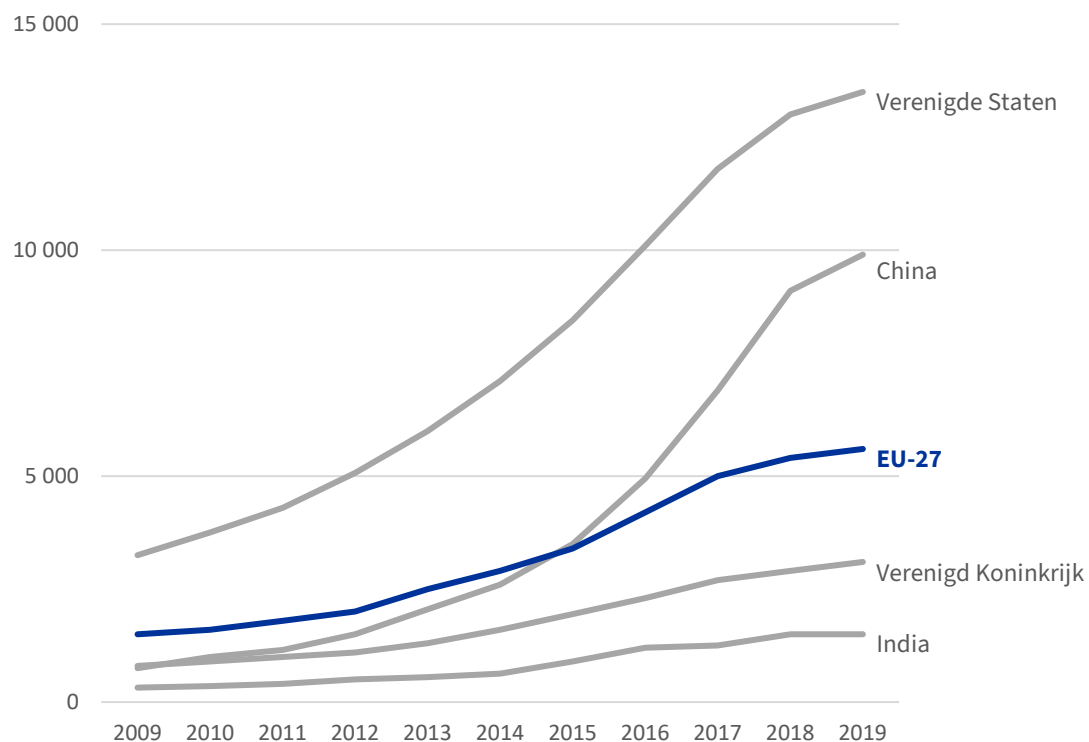


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

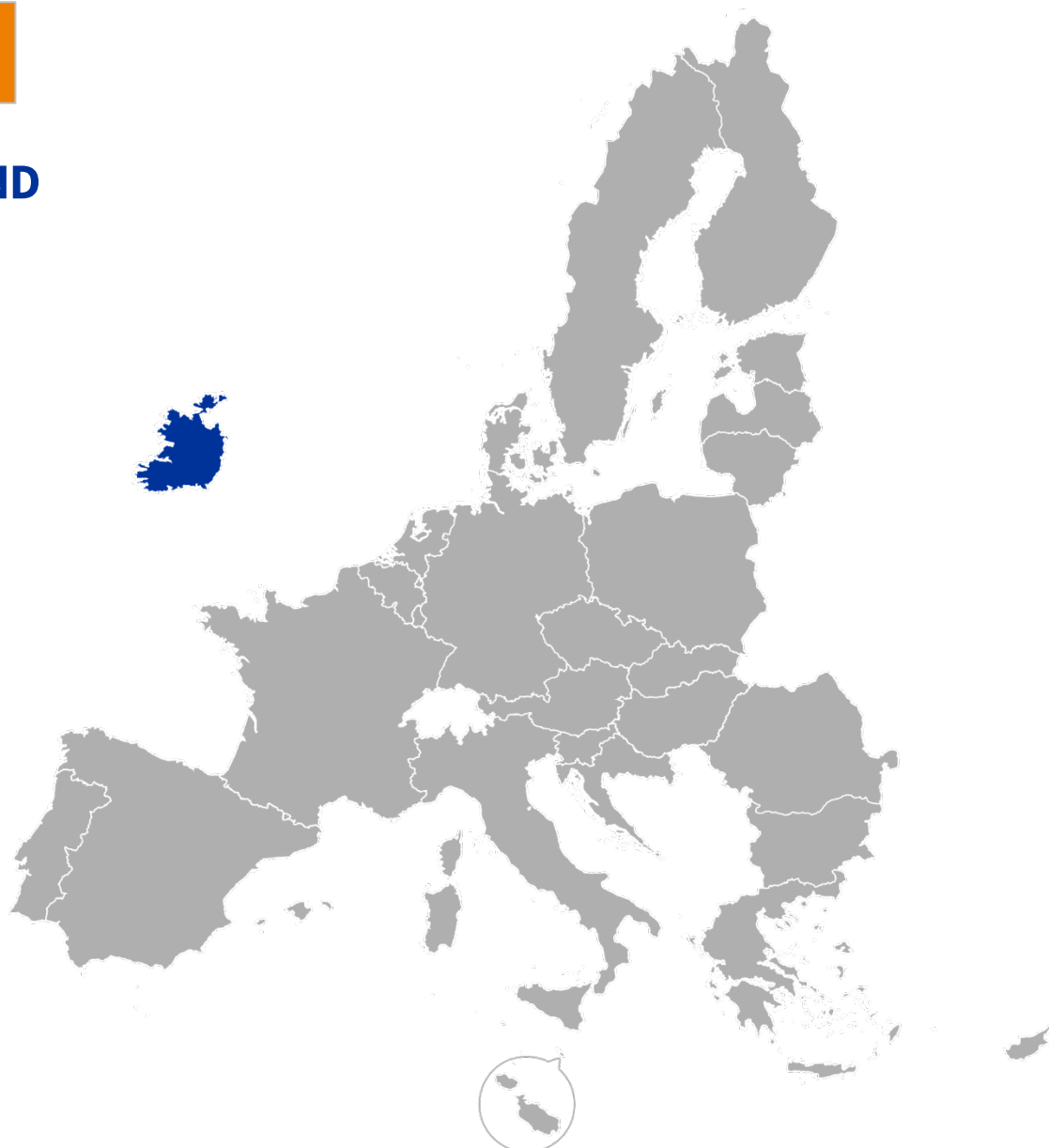
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



IERLAND



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



IERLAND

Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Finland, Portugal en Zweden**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.**

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... **flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.**



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- In jouw land zijn de EU-hoofdkantoren van OpenAI, Google en Meta gevestigd. Deze bedrijven in Ierland houden, is in jouw nationale belang, en ook de EU als geheel profiteert hiervan.
- Je bent verheugd over het voorstel van de Commissie om een gemeenschappelijk EU-kader voor AI in te voeren. Je vindt dit belangrijk omdat dit het juiste signaal afgeeft en de wereld laat zien dat de EU er klaar voor is om een verantwoordelijke AI-grootmacht te worden. Daarom steun je de doelgebonden aanpak die de Commissie voorstelt. Met deze aanpak is er — in vergelijking met de capaciteitsgebonden aanpak — minder risico dat de nieuwe vereisten in strijd zijn met bestaande vereisten in andere sectorspecifieke regelgeving.
- Wat de vereisten voor ontwikkelaars betreft, **benadruk je dat er flexibiliteit nodig is: regelgeving mag investeerders niet afschrikken maar er moet wel verantwoordingsplicht worden ingebouwd.**
- Je wijst er ook op dat veel AI-risico's zich pas na de uitrol voordoen. Voortdurende monitoring is dus doeltreffender dan uitstel van de markttoegang op basis van hypothetische risico's.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Je bent van mening dat de AI-revolutie en de digitale revolutie ook gendergelijkheid in de EU moeten bevorderen. Vrouwen zijn ondervertegenwoordigd op bijna alle gebieden en niveaus als het gaat om onderzoek en innovatie. Je benadrukt dat het zeer belangrijk is om de betrokkenheid van vrouwen bij de vormgeving van AI actief te ondersteunen.
- Start-ups bieden een unieke kans. In tegenstelling tot grote technologiebedrijven is de genderkloof bij hen vaak kleiner en bieden zij meer kansen aan vrouwen en genderqueer mensen op het gebied van AI. Het is ook wetenschappelijk bewezen dat diverse teams minder bevooroordeelde en ethischere AI ontwikkelen — iets wat de EU moet aanmoedigen.
- Start-ups moeten per definitie snel innoveren met beperkte middelen. Zware administratieve lasten kunnen dit in de weg staan, vooral voor kleine teams die proberen door te breken in de sector.
- Daarom ben je er voorstander van dat bepaalde bepalingen van de AI-verordening niet gelden voor kleine start-ups (met minder dan 10 werknemers en minder dan € 2 miljoen aan omzet).
- Jouw opmerkingen:

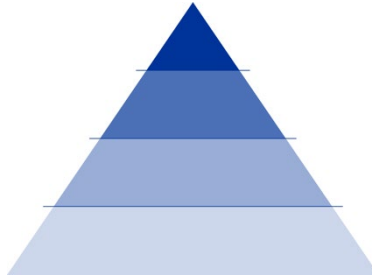
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde training-datasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland	doelgebonden	flexibele benadering	flexibiliteit bij het testen	start-ups	

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

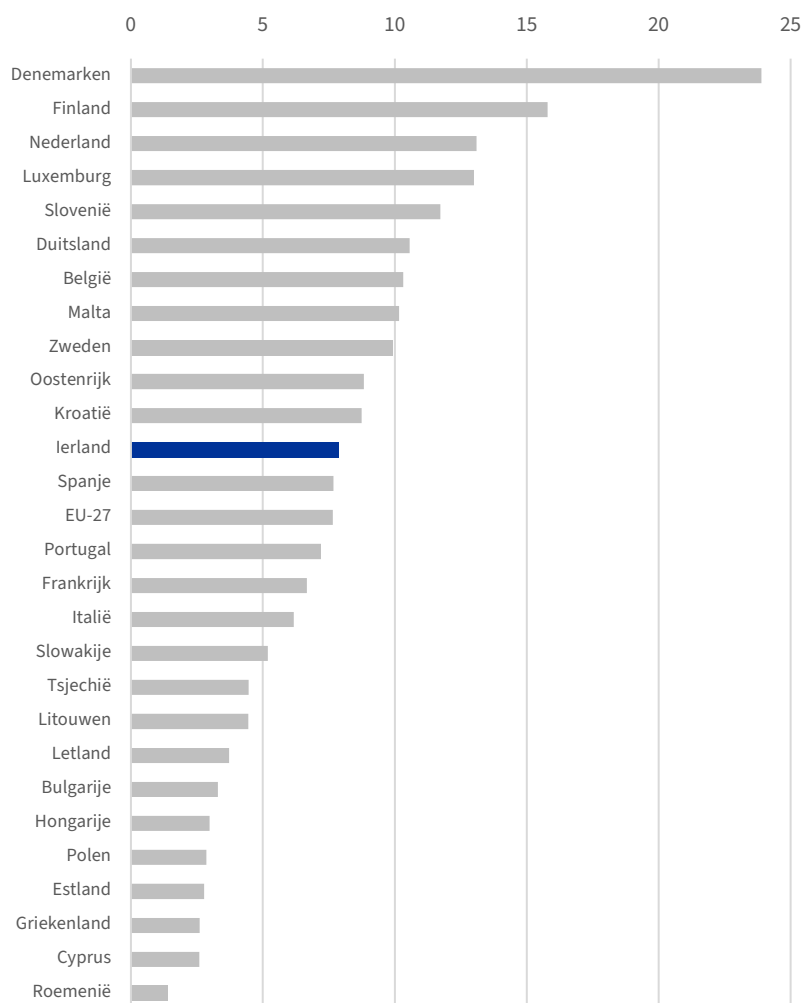
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

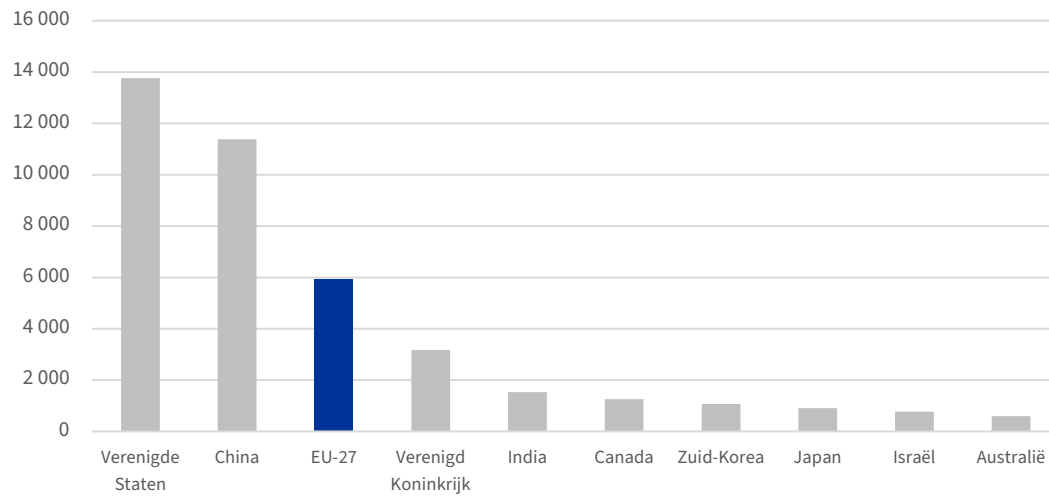
	Europese Unie	Ierland
Bevolking	447 695 350	5 271 395
Bbp per hoofd van de bevolking in EUR	38 150	96 290
Werkloosheidspercentage	6,1%	4,3%
Percentage bedrijven dat AI-technologie gebruikt	8%	8%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	43,8%

Percentage bedrijven met meer dan 10 werknemers
die minstens één AI-systeem gebruiken (2021)

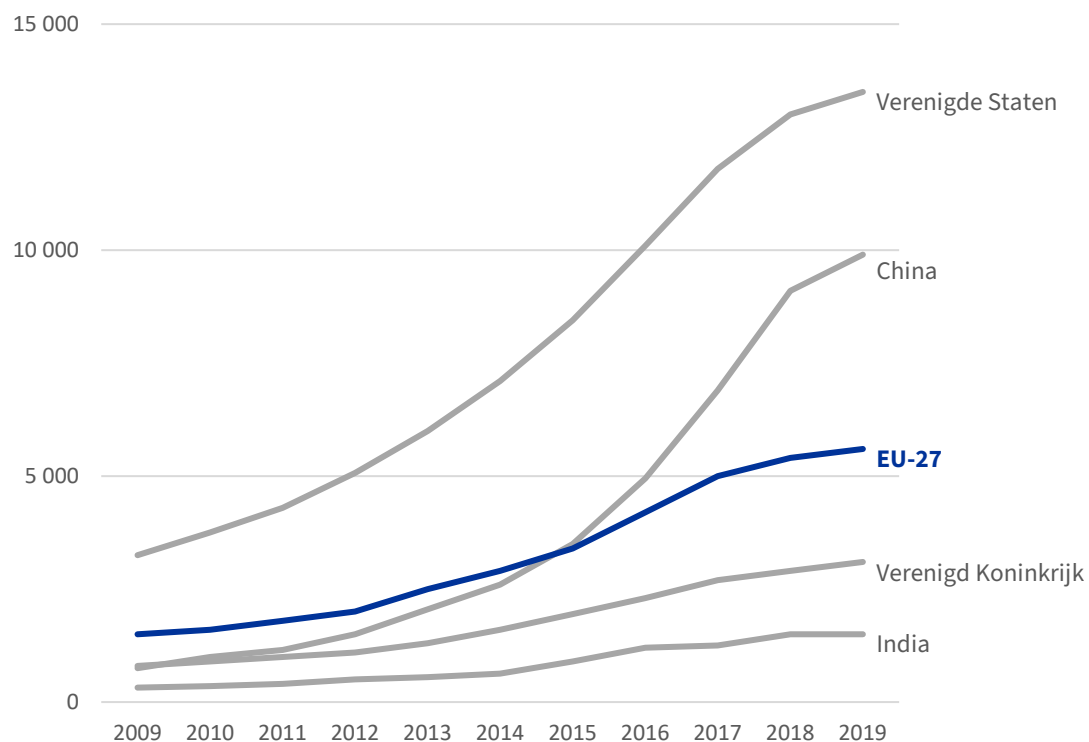


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



ITALIË



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.

Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Kroatië, Nederland en Slovenië**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____.

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____.

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____.

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____.

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Jouw nationale regering is van mening dat AI-regelgeving op Europees niveau moet worden bepaald, aangezien nationale regels het concurrentievermogen en de strategische autonomie van de EU zouden ondermijnen.
- Overeenkomstig artikel 3 van de Italiaanse Grondwet ("Alle burgers hebben gelijke sociale waardigheid en zijn gelijk voor de wet") en de duurzameontwikkelingsdoelstellingen van de VN, moet AI inclusiviteit, duurzaamheid en mensenrechten bevorderen — ten dienste van zowel de mens als de planeet.
- Je vindt de doelgebonden aanpak problematisch, omdat deze voor inconsistenties in bepaalde sectoren kan zorgen. Zo worden AI die chirurgische operaties ondersteunt en AI die ziekenhuiskamers toewijst beide als systeem met een hoog risico geclassificeerd, hoewel ze een hele andere impact hebben. Een capaciteitsgebonden aanpak weerspiegelt de werkelijke gevaren beter, hoewel je er voorstander van bent om de risicolijst regelmatig te updaten, zodat rekening kan worden gehouden met veranderende AI-vermogens.
- Je bent ook voorstander van een testkader uit voorzorg, niet alleen met het oog op veiligheidsoverwegingen, maar ook om duurzaamheidsredenen: de ontwikkeling van AI is arbeids- en energie-intensief en het trainen van grote taalmodellen vereist honderden megawatturen aan elektriciteit. Een ongereguleerde proefondervindelijke aanpak leidt daardoor tot hogere milieukosten. Strengere regelgeving kan de ontwikkeling enigszins vertragen, maar leidt op lange termijn tot veiligere, efficiëntere en betrouwbaardere AI.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Volgens jou moet de AI-verordening van de EU als katalysator dienen voor de ontwikkeling van AI in alle EU-lidstaten. Start-ups zijn cruciale aanjagers van innovatie, dus de AI-verordening mag geen buitensporige lasten voor AI-start-ups met zich meebrengen — start-ups hebben mogelijk niet de financiële en administratieve capaciteit om aan strenge testvereisten te kunnen voldoen.
- Als artikel 1 resulteert in een voorzorgsbenadering voor het testen van AI-modellen met een hoog risico — waar jij voorstander van bent — stel je een uitzondering voor voor start-ups met minder dan 20 werknemers en een omzet van minder dan € 4 miljoen. De lidstaten zouden in plaats daarvan de flexibele aanpak mogen volgen. Je ziet dit als een evenwichtig compromis en wilt het gesprek daarover leiden.
- Jouw opmerkingen:

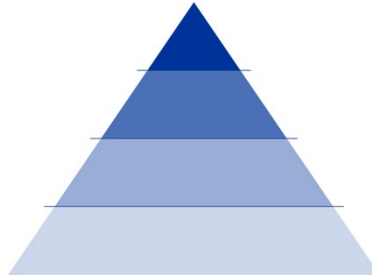
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

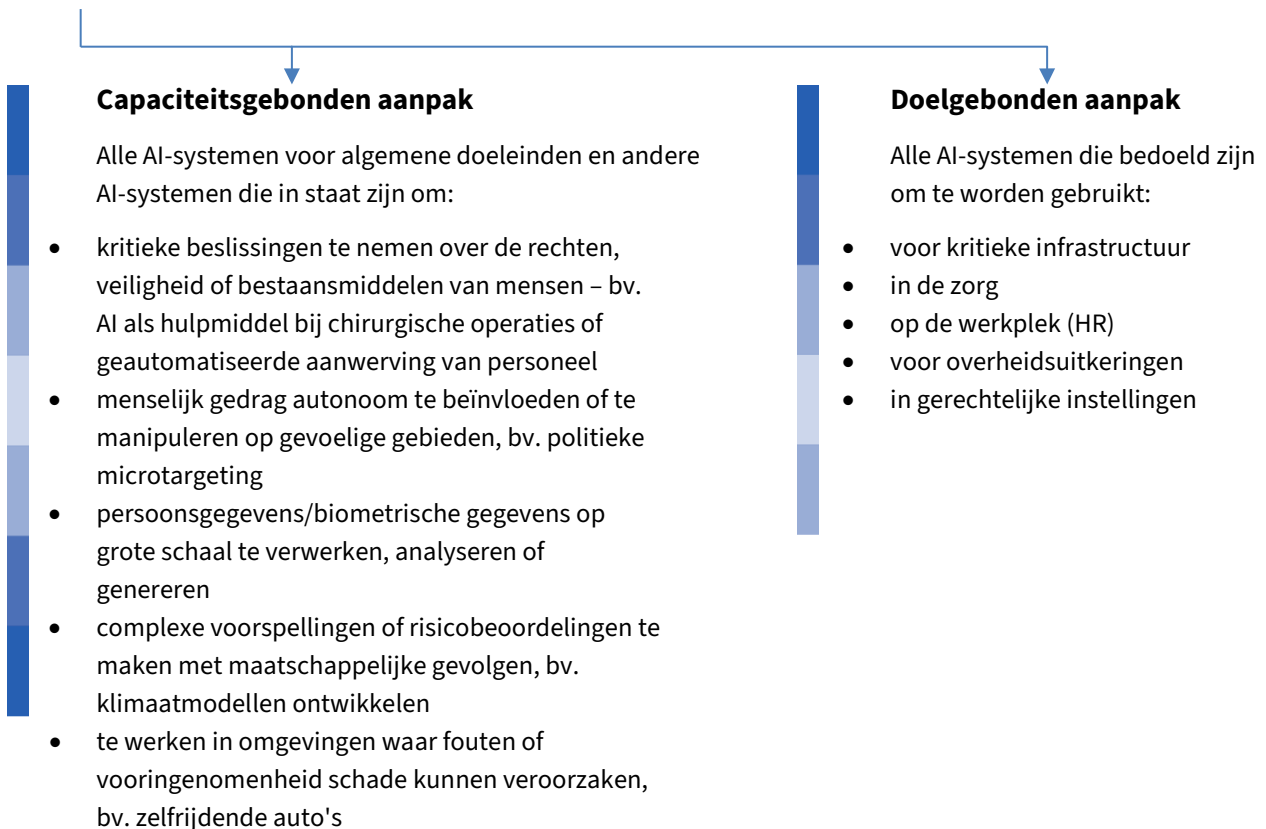
Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

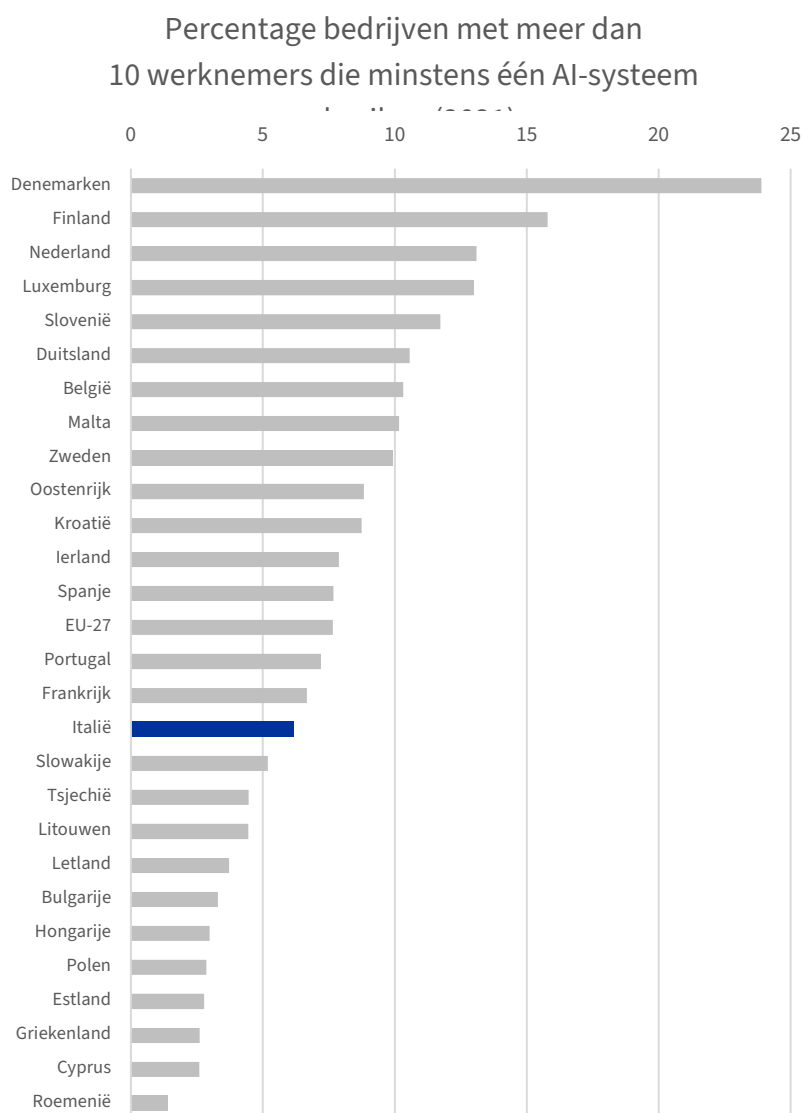
	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië	capaciteitsgebonden	voorzorgsbenadering		start-ups	
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

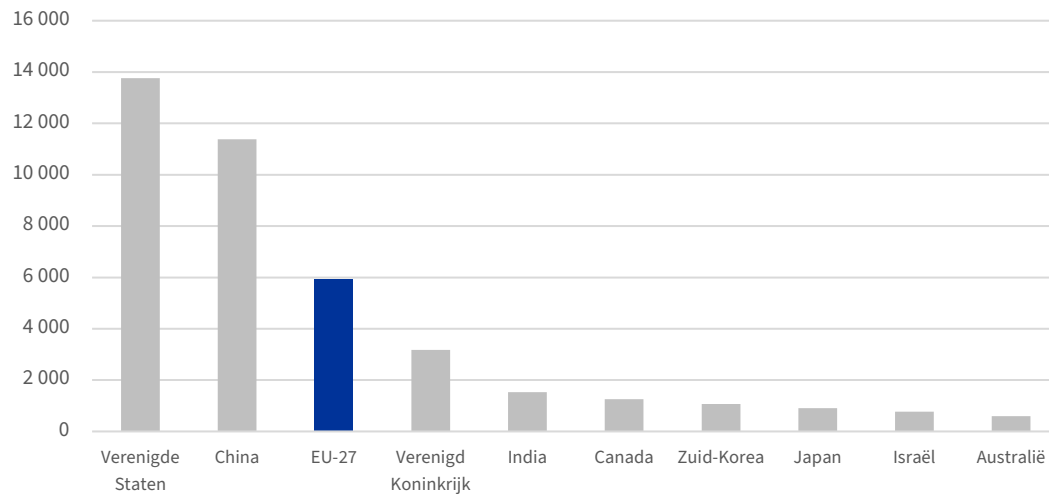
Let op: **De onderhandelingen vinden plaats in 2023!**

	Europese Unie	Italië
Bevolking	447 695 350	58 997 201
Bbp per hoofd van de bevolking in EUR	38 150	36 130
Werkloosheidspercentage	6,1%	7,7%
Percentage bedrijven dat AI-technologie gebruikt	8%	5,1%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	22,2%

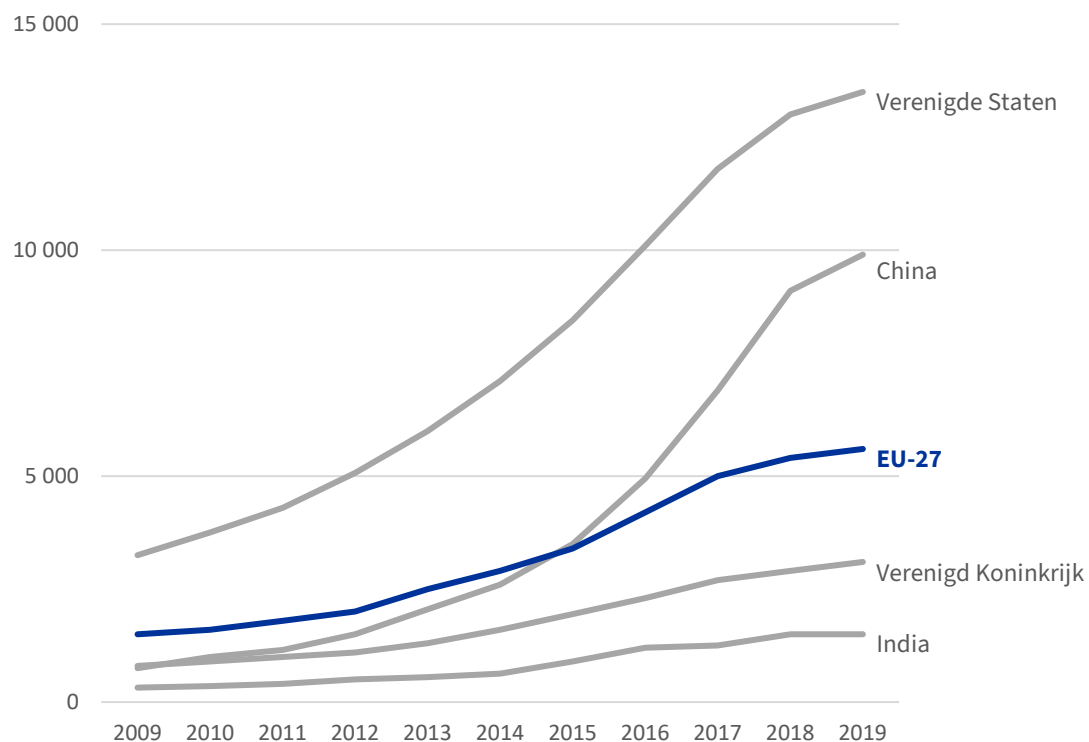


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

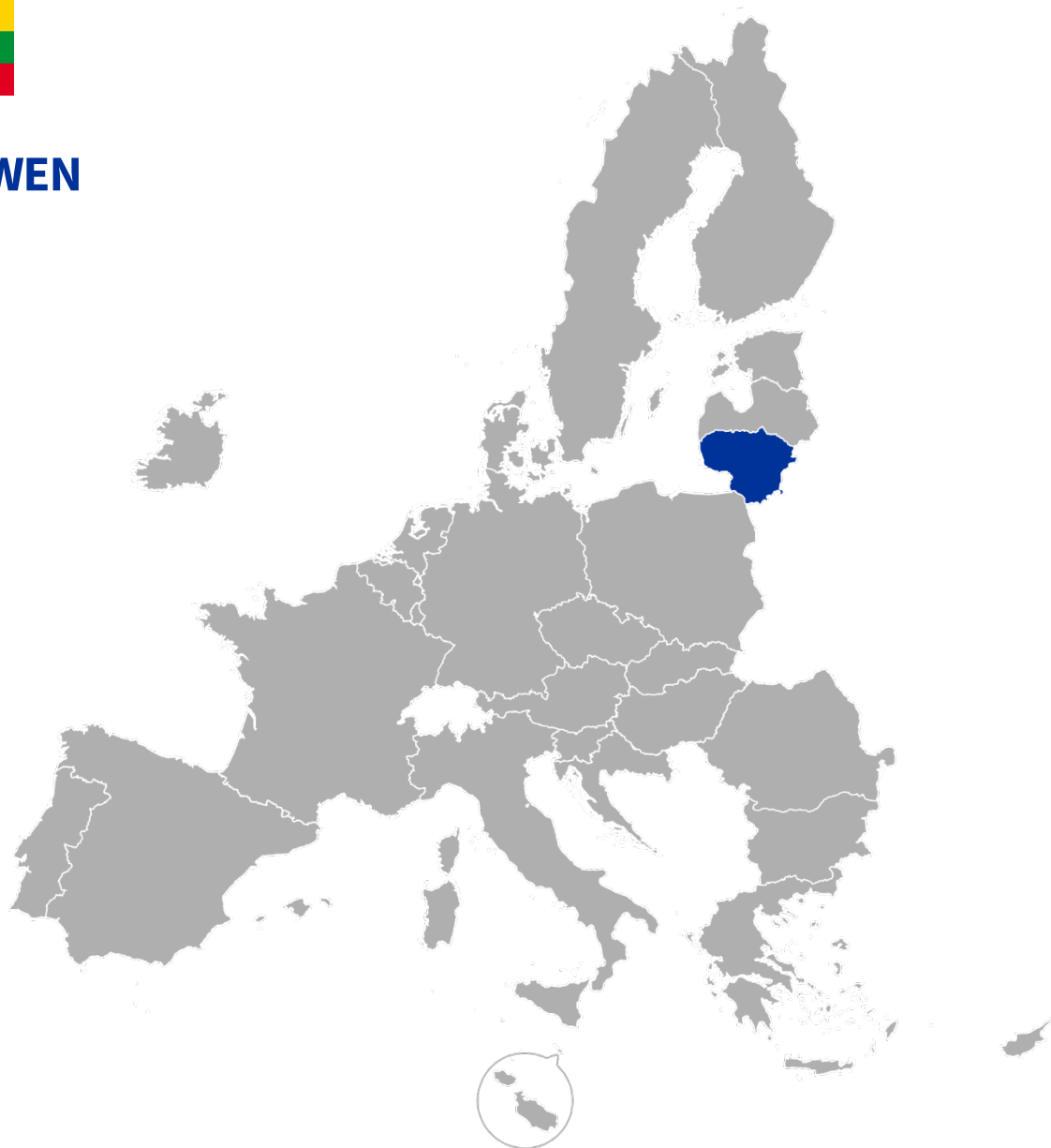
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



LITOUWEN



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Tsjechië, Estland, Hongarije en Letland**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- De be- en verwerkende industrie is de grootste sector van de Litouwse economie en is goed voor 20,4% van het bbp van het land. Een van de grootste problemen van deze sector is de lage arbeidsproductiviteit. Artificiële intelligentie kan dit oplossen door routinetaken te automatiseren en processen te optimaliseren.
- Daarom pleit je voor regelgeving die vrije innovatie en experimenten mogelijk maakt. Hoewel onbedoelde gevolgen hiervan moeten worden vermeden, benadruk je dat, zelfs met de flexibele aanpak, zelfevaluaties verplicht blijven. Ontwikkelaars mogen deze evaluaties afstemmen op hun producten.
- Te veel regelgeving zou innovatie vertragen en Europese AI-start-ups benadelen, aangezien verplichte externe audits en lange testfases duur zijn.
- Hoewel sommige lidstaten aandringen op strengere regels, omdat zij zich zorgen maken over de mensenrechten, voer jij aan dat AI ook een krachtige tool kan zijn om die rechten te beschermen. AI kan helpen om vooroordelen bij het aannemen van personeel, politiewerk en besluitvorming te verminderen, cyberaanvallen af te weren en persoonsgegevens te beschermen. En hoewel AI misinformatie of deepfakes kan genereren, kan AI ook worden getraind om deze op te sporen en tegen te gaan. De EU moet innovatie bij dit beschermende gebruik van AI actief ondersteunen.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Hoewel je bij artikel 1 pleit voor minder beperkingen en meer vrijheid om te innoveren, ben je er desondanks wel sterk van overtuigd dat regels nog altijd belangrijk zijn. Regelgeving is het zinvolste als de regels voor iedereen gelden. Op die manier krijgen ontwikkelaars een duidelijker idee van wat er wordt verwacht en kan de EU een consistente visie op verantwoorde AI tentoonspreiden.
- Er is al reden tot alarm: in China wordt door AI-aangestuurde gezichts- en emotieherkenningstechnologie gebruikt om de islamitische Oeigoerse minderheid in de gaten te houden, met zelfcensuur en de opsluiting van veel Oeigoeren in zogenaamde "heropvoedingskampen" tot gevolg. In de VS heeft Clearview AI zonder toestemming miljarden afbeeldingen van socialemediaplatforms gehaald en zijn tool verkocht aan agentschappen zoals de FBI en ICE. In het VK hadden tests met live gezichtsherkenning door de Londense politie hoge foutenpercentages, met name voor zwarte mensen.
- Je bent van mening dat deze problemen voorkomen hadden kunnen worden als er fatsoenlijke tests zouden zijn uitgevoerd en er regels zouden zijn geweest die op iedereen van toepassing zijn — zonder uitzonderingen.
- Jouw opmerkingen:

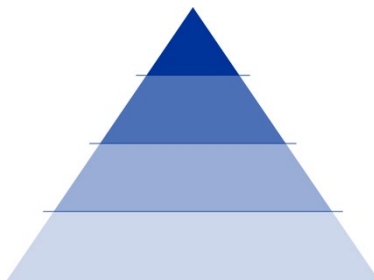
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde training-datasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

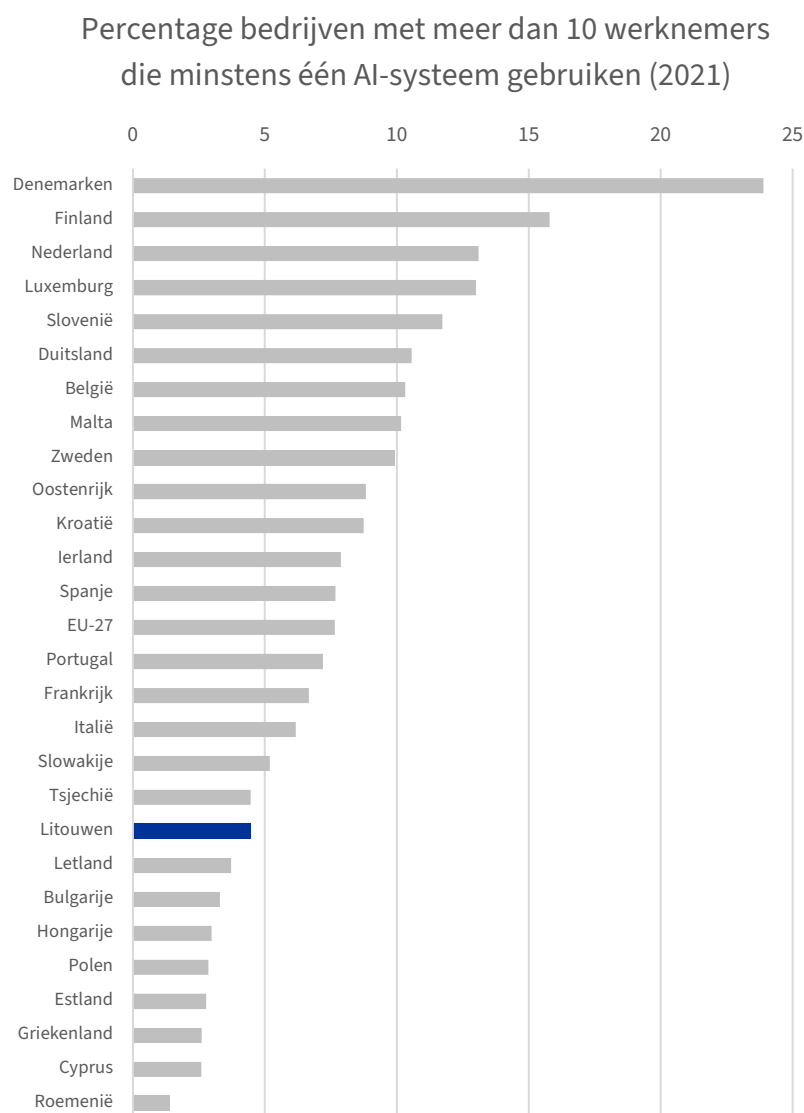
	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Italië					
Letland					
Litouwen	capaciteitsgebonden	flexibele benadering		zonder uitzonderingen	
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

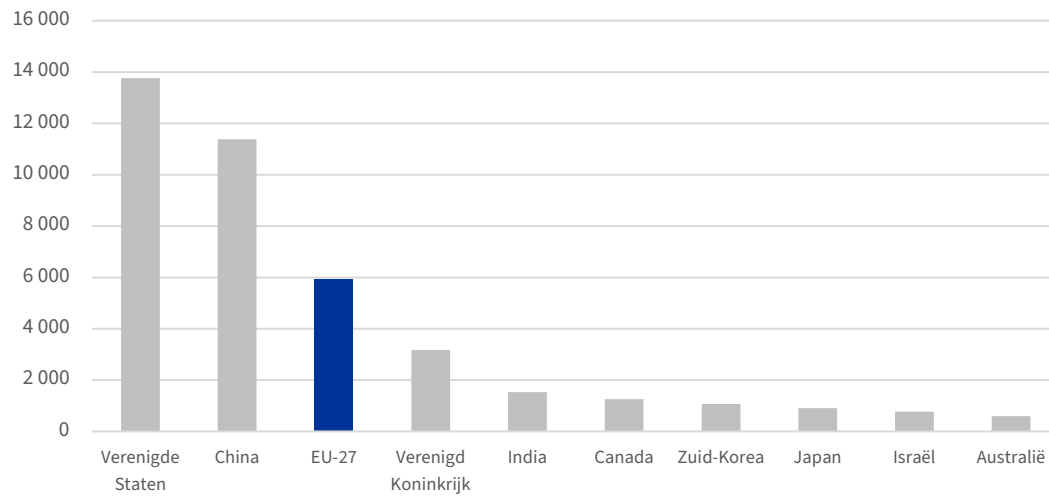
Let op: **De onderhandelingen vinden plaats in 2023!**

	Europese Unie	Litouwen
Bevolking	447 695 350	2 857 279
Bbp per hoofd van de bevolking in EUR	38 150	25 700
Werkloosheidspercentage	6,1%	6,9%
Percentage bedrijven dat AI-technologie gebruikt	8%	4,9%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	25,9%

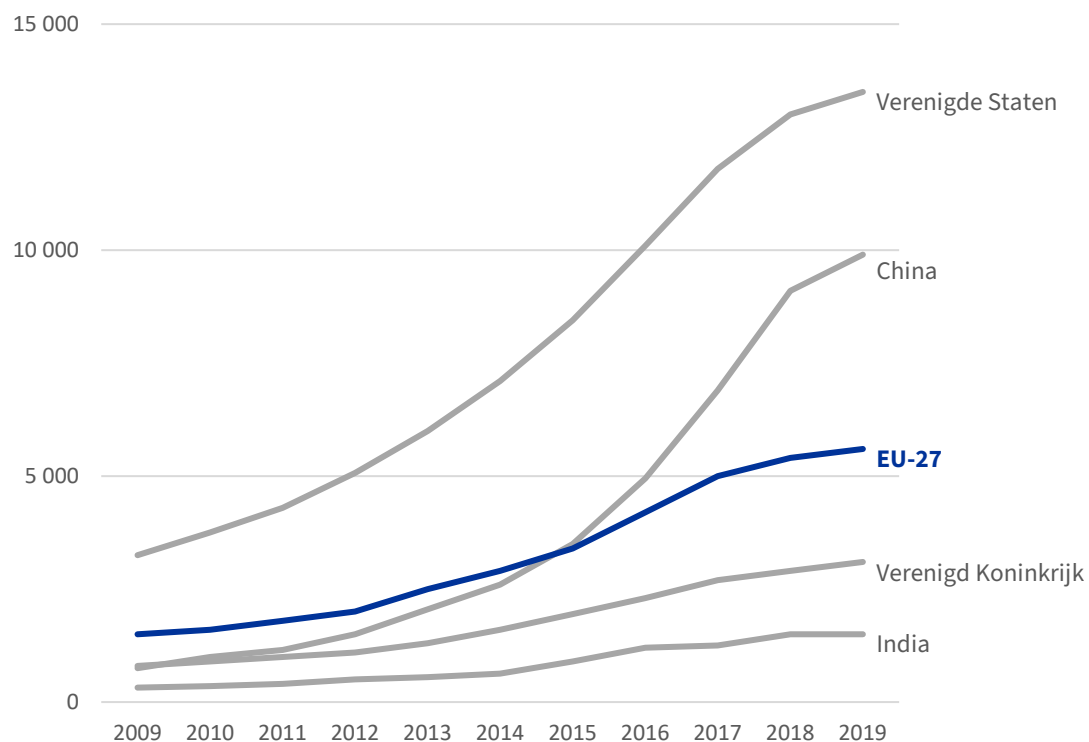


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



LUXEMBURG



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Bulgarije, Malta en Slowakije**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Landen als Duitsland en Spanje streven ernaar om ook maar de geringste vooringenomenheid van AI weg te nemen, wat jij onrealistisch vindt. Bij elk besluit komen er vooroordelen om de hoek kijken. AI kan, wanneer het op verantwoorde wijze wordt ontwikkeld, vooroordelen van mensen juist verminderen. Je probeert het debat een realistischere richting op te sturen waarbij de behoeften van alle lidstaten in aanmerking worden genomen. Als je ervan wordt beschuldigd dat je de risico's van AI bagatelliseert, dan ben je het daar absoluut niet mee eens. Je geeft simpelweg prioriteit aan de meest directe risico's. Daarom is een doelgebonden aanpak van essentieel belang.
- AI mag bijvoorbeeld geen gevaar vormen voor elektriciteits- en internetsystemen. Dat zou gevolgen kunnen hebben voor ziekenhuizen en op grote schaal levens in gevaar kunnen brengen. De capaciteitsgebonden aanpak is te verwarrend en zorgt voor te veel mazen in de wetgeving.
- Je bent er voorstander van om AI sneller te ontwikkelen, ook in sectoren met een hoog risico. Je bent echter bang dat jouw land niet over de financiële draagkracht beschikt om de door AI geboden kansen ten volle te benutten als er strenge regels gelden. Daarom pleit je ervoor de lijst van vereisten (zie bijlage II) op sommige plekken te wijzigen. Je zou de mogelijkheid om in de praktijk tests uit te voeren aan de lijst willen toevoegen, op voorwaarde dat externe audits verplicht blijven. Zo kunnen ontwikkelaars zelf kiezen welke tests zij uitvoeren, afhankelijk van hun financiële draagkracht.
- Als het moeilijk is om overeenstemming te bereiken over de doelgebonden aanpak, **ben je bereid om in te stemmen met een herziening over 3 jaar.**

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)

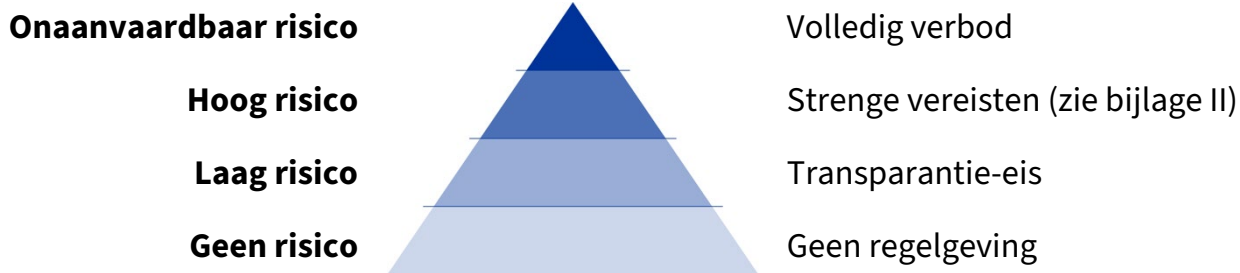


... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Je bent van mening dat de AI-revolutie en de digitale revolutie ook gendergelijkheid in de EU moeten bevorderen. Vrouwen zijn ondervertegenwoordigd op bijna alle gebieden en niveaus als het gaat om onderzoek en innovatie. Je benadrukt dat het zeer belangrijk is om de betrokkenheid van vrouwen bij de vormgeving van AI actief te ondersteunen.
- Start-ups bieden een unieke kans. In tegenstelling tot grote technologiebedrijven is de genderkloof bij hen vaak kleiner en bieden zij meer kansen aan vrouwen en genderqueer mensen op het gebied van AI. Het is ook wetenschappelijk bewezen dat diverse teams minder bevooroordeelde en ethischere AI ontwikkelen — iets wat de EU moet aanmoedigen.
- Start-ups moeten per definitie snel innoveren met beperkte middelen. Zware administratieve lasten kunnen dit in de weg staan, vooral voor kleine teams die proberen door te breken in de sector.
- Daarom ben je er voorstander van dat bepaalde bepalingen van de AI-verordening niet gelden voor kleine start-ups (met minder dan 10 werknemers en minder dan € 2 miljoen aan omzet).
- Jouw opmerkingen:

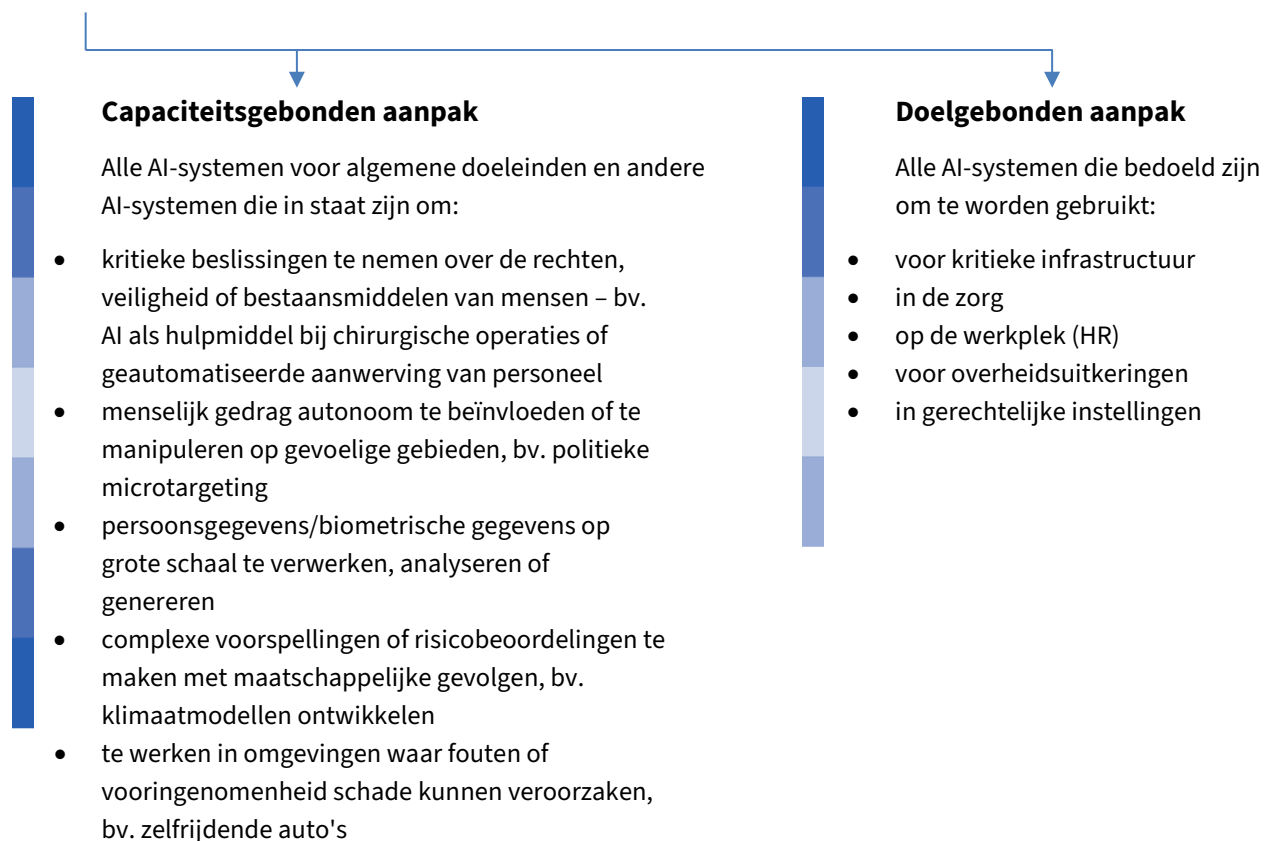
Bijlage I: Risicoclassificatie



Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië					
Letland					
Litouwen					
Luxemburg	doelgebonden	flexibele benadering	over 3 jaar opnieuw bekijken	start-ups	
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

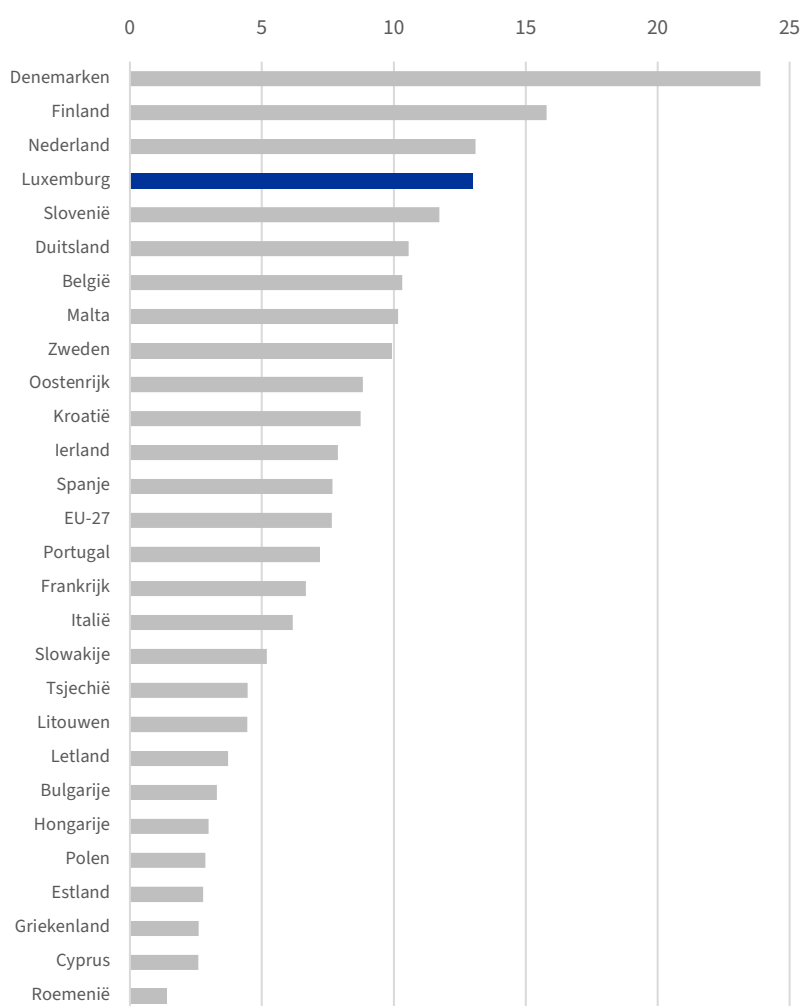
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

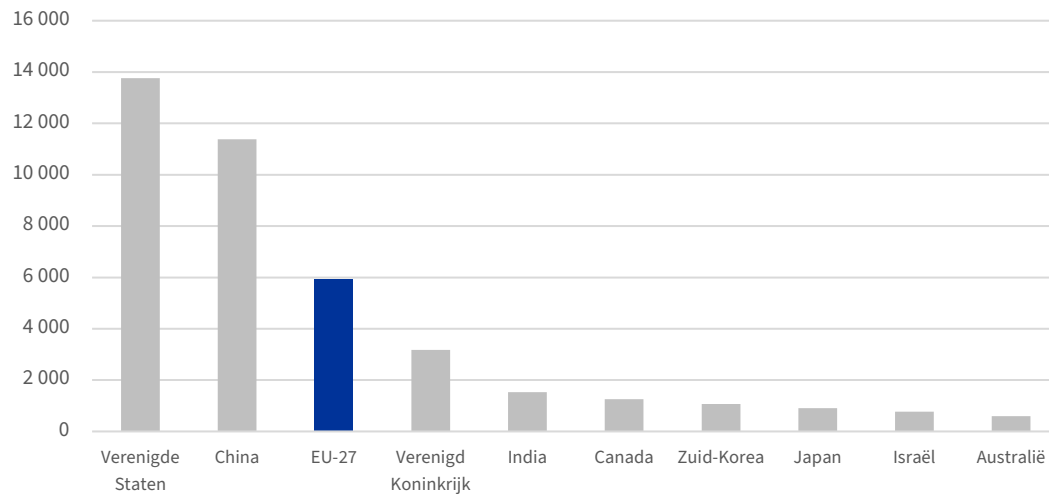
	Europese Unie	Luxemburg
Bevolking	447 695 350	660 809
Bbp per hoofd van de bevolking in EUR	38 150	121 290
Werkloosheidspercentage	6,1%	5,2%
Percentage bedrijven dat AI-technologie gebruikt	8%	14,5%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	27,9%

Percentage bedrijven met meer dan
10 werknemers die minstens één AI-systeem
gebruiken (2021)

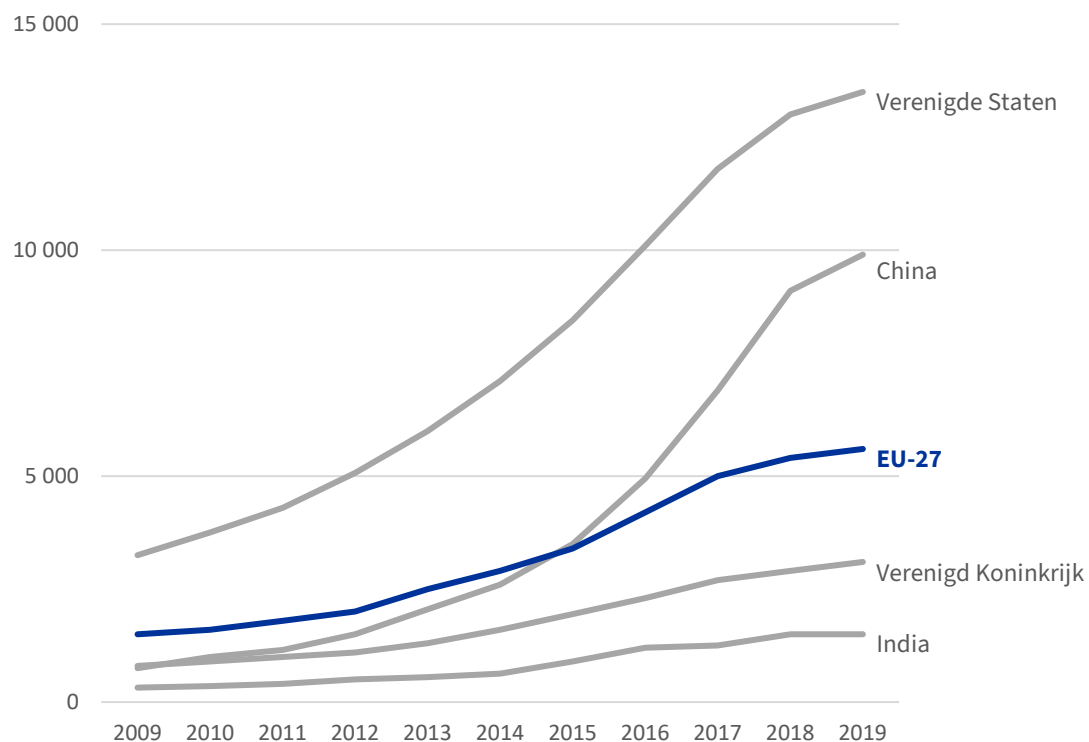


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

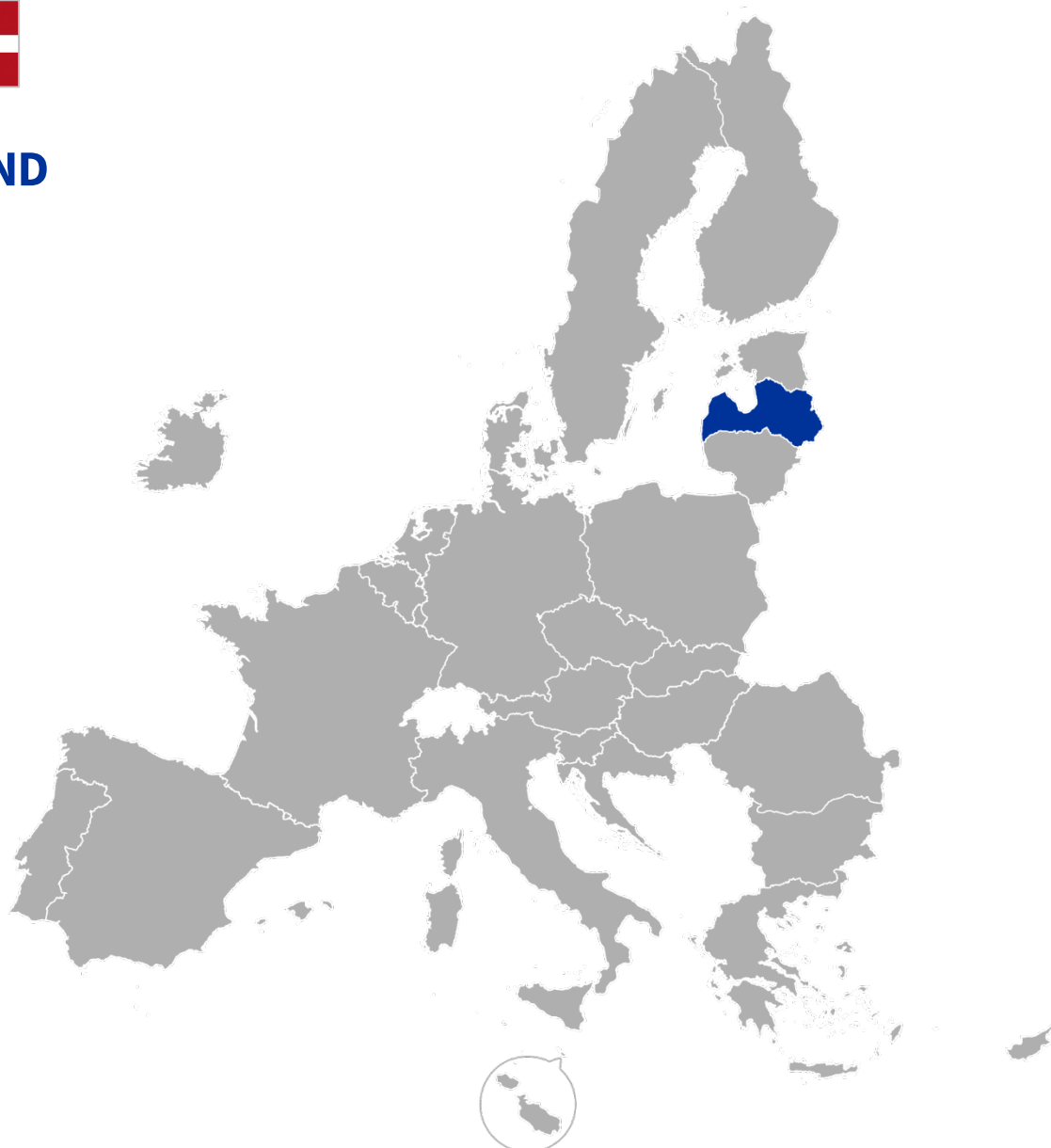
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



LETLAND



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



LETLAND

Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Tsjechië, Estland, Hongarije en Litouwen**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng tijdens de officiële onderhandelingen je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Je zet je in voor de invoering van digitale oplossingen in Letland. Dat brengt met zich mee dat er moet worden gewerkt aan de bescherming van gevoelige informatie en digitale infrastructuur, beide essentieel voor het behoud van vitale maatschappelijke functies.
- Voor het classificeren van AI-modellen met een hoog risico vind je de capaciteitsgebonden aanpak geschikter dan de doelgebonden aanpak die de Commissie voorstelt. Als hetzelfde AI-systeem in de ene context een laag risico heeft en in de andere een hoog risico, kan dit tot verwarring en rechtsonzekerheid leiden voor gebruikers, ontwikkelaars en investeerders.
- Om de nodige investeringen in AI aan te trekken, is een goed regelgevingskader van cruciaal belang. Voor jou houdt dat een goede bescherming van de grondrechten en tegelijkertijd vrije en flexibele innovatieprocessen in. Als de EU wil concurreren met AI-grootmachten als de VS en China, moet de EU uitblinken. Dit doel kan niet worden bereikt zolang overdreven restrictieve regels experimenten en ontwikkeling beperken. Daarom stel je voor om bepaalde wijzigingen aan te brengen in de lijst van vereisten voor AI-systemen met een hoog risico (zie bijlage II). Je denkt dat het beter zou zijn om **"controle na toelating op de markt" als vereiste te schrappen**, aangezien het onduidelijk is wat deze verplichting precies inhoudt.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Over het algemeen ben je het eens met de landen die geen uitzonderingen willen, waaronder jouw buurland en trouwe bondgenoot Litouwen. Je ben voorstander van gezamenlijke Europese AI-regelgeving zonder onnodige uitzonderingen.
- Sinds de Russische invasie in Oekraïne zijn jij en de andere Baltische staten echter sterk op jullie hoede voor een mogelijke Russische aanval op de Europese Unie. Jouw land deelt zelfs een grens met Rusland, en samen met de NAVO-strijdkrachten doe je elke dag je best om deze te beschermen.
- Daarom sta je ervoor open om beperkte uitzonderingen voor militaire en defensiedoeleinden te bespreken, maar niet voor nationale veiligheidsdoeleinden in bredere zin. Hoewel de nationale veiligheid onder de bevoegdheid van de lidstaten valt, zouden er mazen in de wet en fragmentatie van de regelgeving ontstaan als deze sector geheel buiten de EU-regelgeving zou worden gelaten. Een minimale EU-basisaanpak kan zorgen voor meer consistentie (en duidelijkheid voor ontwikkelaars) en uitdragen dat de EU-lidstaten een eensgezind standpunt innemen over AI voor nationale veiligheidsdoeleinden.
- Jouw opmerkingen:

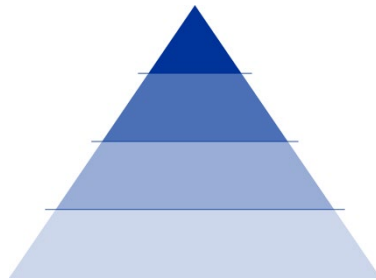
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico

Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië					
Letland	capaciteitsgebonden	flexibele benadering	controle na toelating op de markt schrappen	leger/defensie	
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

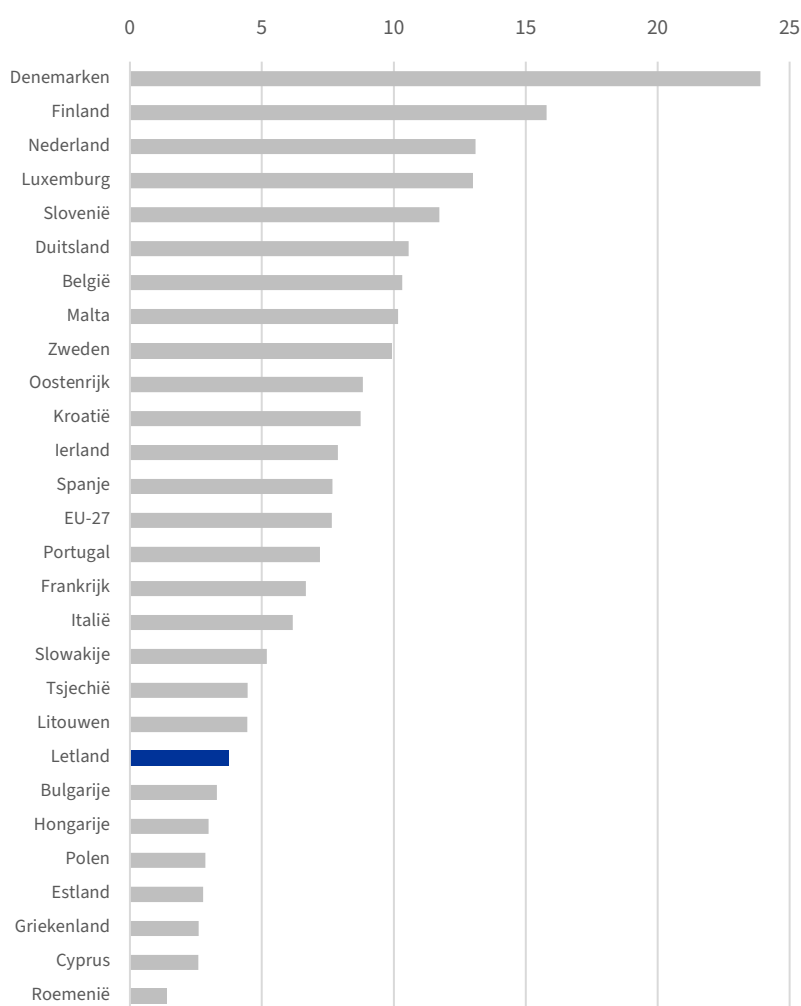
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken.
Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

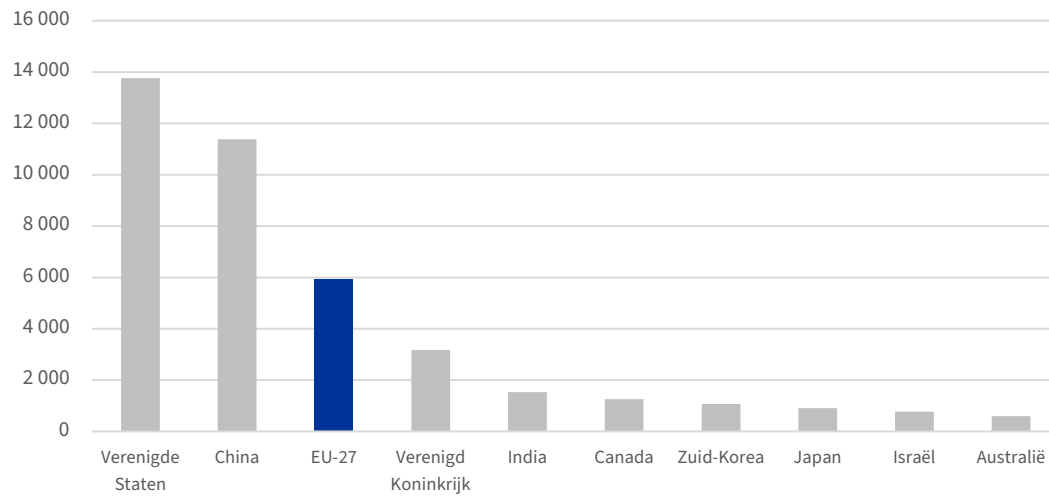
	Europese Unie	Letland
Bevolking	447 695 350	1 883 008
Bbp per hoofd van de bevolking in EUR	38 150	20 930
Werkloosheidspercentage	6,1%	6,5%
Percentage bedrijven dat AI-technologie gebruikt	8%	4,5%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	16,6%

Percentage bedrijven met meer dan
10 werknemers die minstens één AI-systeem
gebruiken (2021)

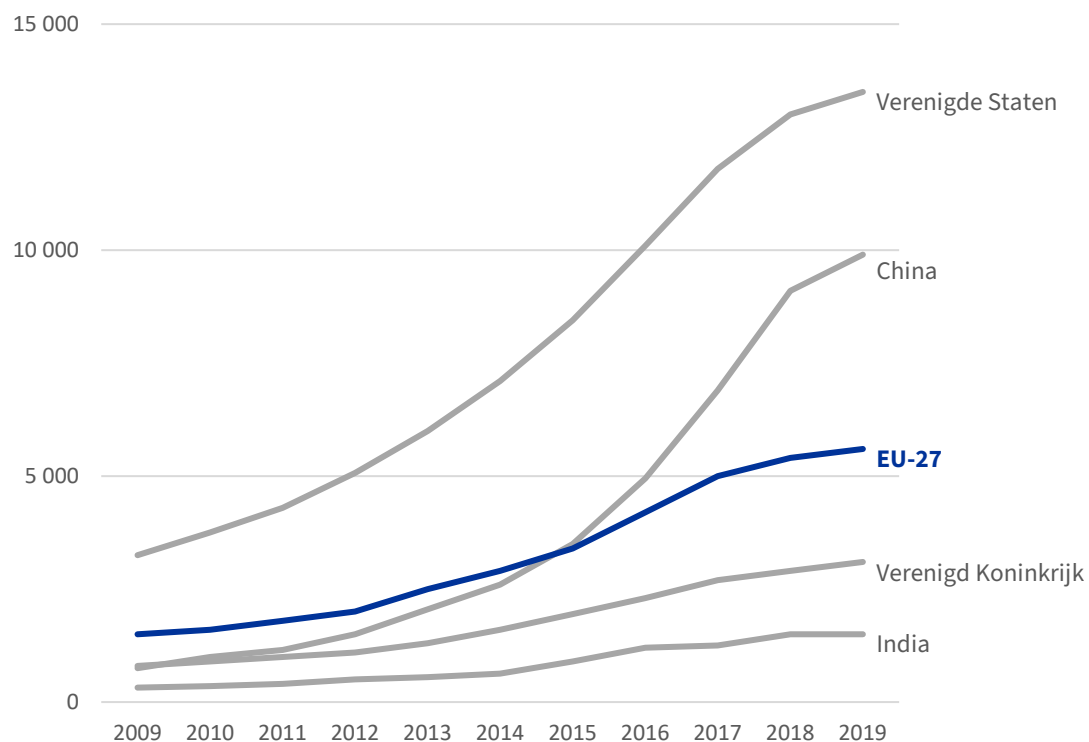


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



MALTA



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



MALTA

Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Bulgarije, Luxemburg en Slowakije**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng tijdens de officiële onderhandelingen je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.**

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... **flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.**



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Als een van de kleinste EU-lidstaten (met slechts 550 000 inwoners) en als eilandstaat, zet jouw land zich sterk in voor duurzame economische groei. Je streeft ernaar om buitenlandse investeringen en geschoold talent aan te trekken om een robuuste AI-sector te ontwikkelen, die zich vooral richt op geneeskunde en kritieke infrastructuur, gebieden waar een hoog rendement kan worden verwacht.
- Je steunt het voorstel van de Commissie voor horizontale AI-regelgeving, en met name de doelgebonden aanpak van risicobeheer. Je vindt deze aanpak een logische en geschikte keuze, omdat het voor regelgevers gemakkelijker is om specifieke toepassingen te beoordelen dan om de abstracte vermogens van een model te beoordelen.
- Het is cruciaal om te benadrukken dat er voor technologische creativiteit vrijheid nodig is — innovatie moet worden omarmd en niet worden gehinderd. Het is daarom belangrijk om het **testen van AI-modellen onder reële omstandigheden** te bevorderen, omdat testomgevingen, hoewel ze nuttig zijn, vaak niet volledig waarheidsgetrouw zijn.
- Kleine landen zoals Malta kampen met unieke uitdagingen: een beperkte talentenpool en minder middelen voor dure tests. Grotere, kapitaalrijke staten kunnen die lasten gemakkelijker dragen. Daarom ben je sterk gekant tegen overdreven strenge testvereisten die jouw nationale ambities op het gebied van AI-ontwikkeling structureel zouden belemmeren.

Artikel 2

Toepassingsgebied

Artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)

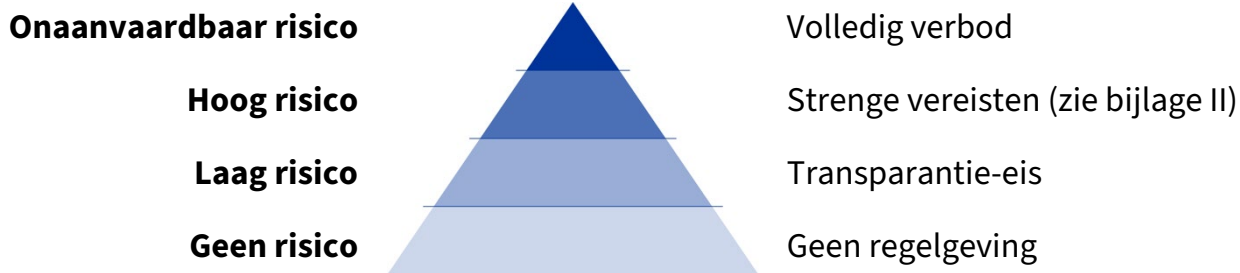


... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Het blijft voor jou van essentieel belang dat zaken in verband met de nationale veiligheid, het leger of defensie buiten het toepassingsgebied van de AI-verordening blijven. Met name de nationale veiligheid valt niet onder de bevoegdheid van de EU. Dit staat in artikel 4, lid 2, van het Verdrag betreffende de Europese Unie (VEU): "[D]e nationale veiligheid blijft uitsluitend de verantwoordelijkheid van elke lidstaat". Het heeft dan ook geen zin om de lidstaten voor te schrijven hoe zij AI-toepassingen voor nationale veiligheid moeten reguleren.
- Aan de andere kant zie je geen reden om start-ups van de verordening vrij te stellen. Het feit dat een bedrijf klein is, betekent niet dat zijn technologie niet schadelijk kan zijn. Bovendien zouden grote bedrijven om de testvereisten te omzeilen riskante projecten naar start-ups kunnen doorschuiven, omdat die toch zijn vrijgesteld. Zo'n uitzondering zou dus een gemakkelijk achterdeurtje bieden, en de EU moet ervoor zorgen dat dit niet gebeurt.
- Jouw opmerkingen:

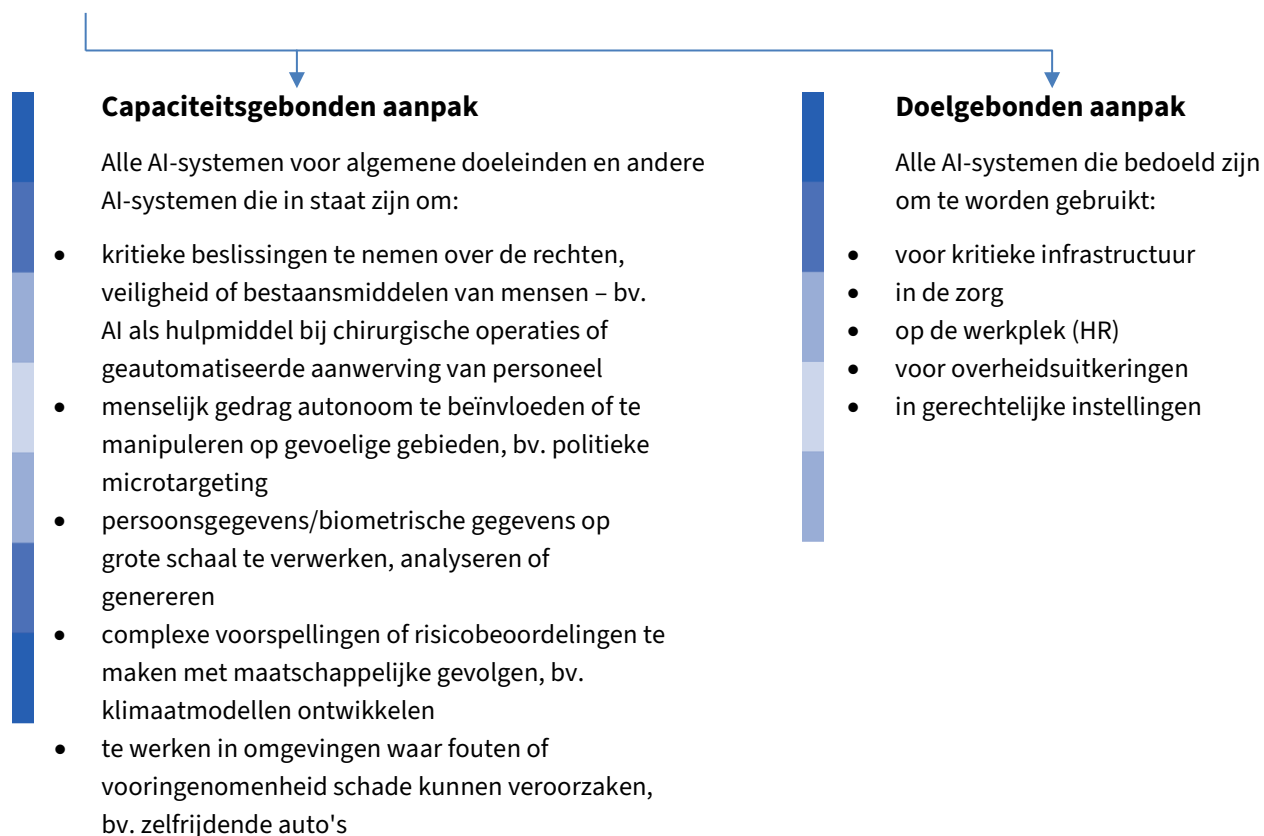
Bijlage I: Risicoclassificatie



Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië					
Letland					
Litouwen					
Luxemburg					
Malta	doelgebonden	flexibele benadering	testen onder reële omstandigheden	nationale veiligheid	
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

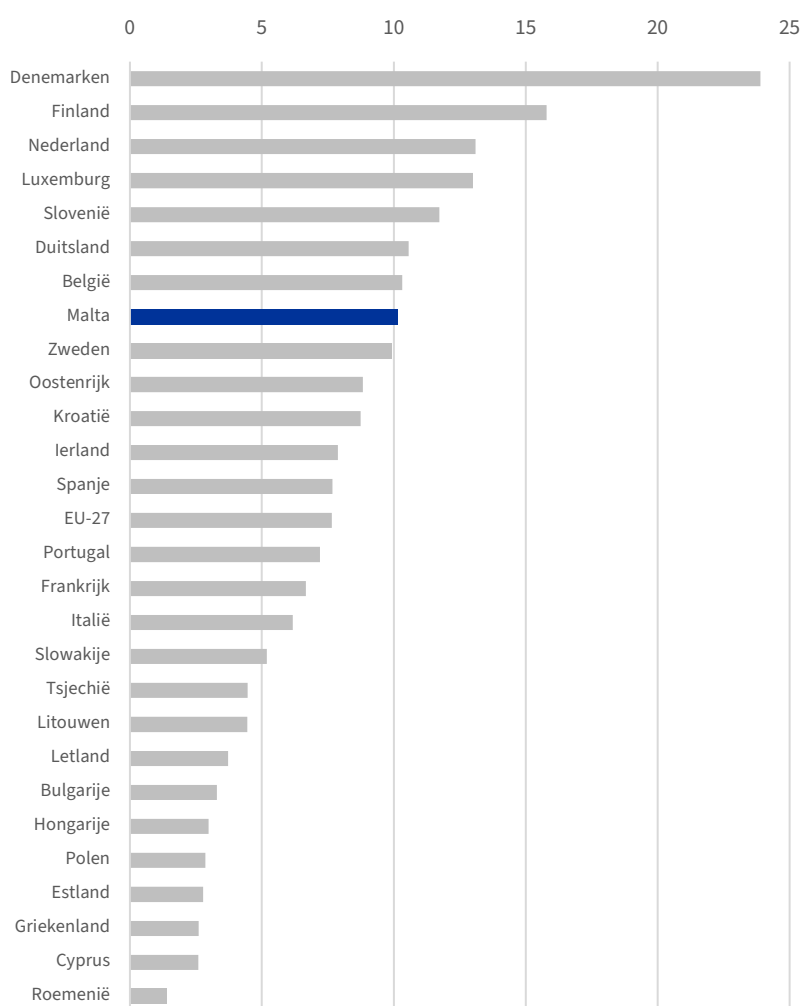
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

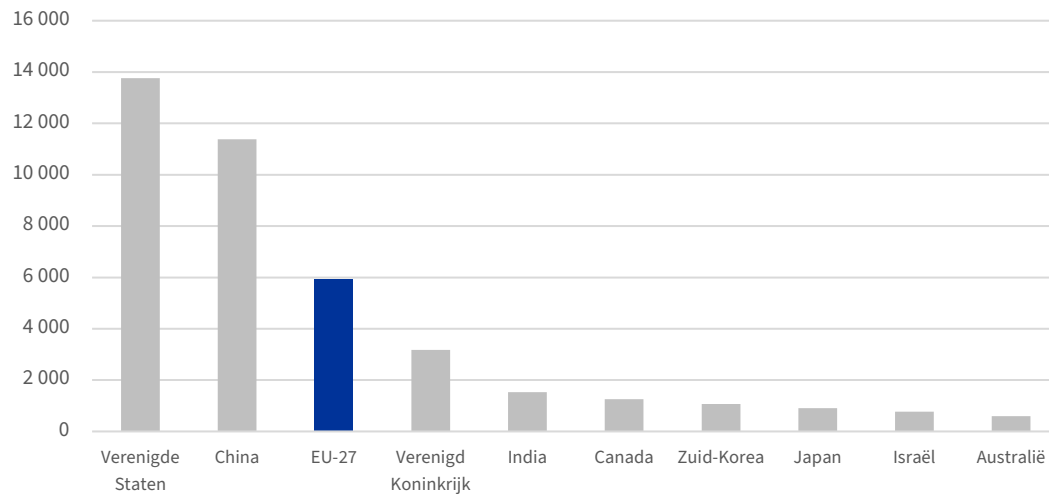
	Europese Unie	Malta
Bevolking	447 695 350	542 051
Bbp per hoofd van de bevolking in EUR	38 150	37 110
Werkloosheidspercentage	6,1%	3,5%
Percentage bedrijven dat AI-technologie gebruikt	8%	13,2%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	37%

Percentage bedrijven met meer dan
10 werknemers die minstens één AI-systeem
gebruiken (2021)

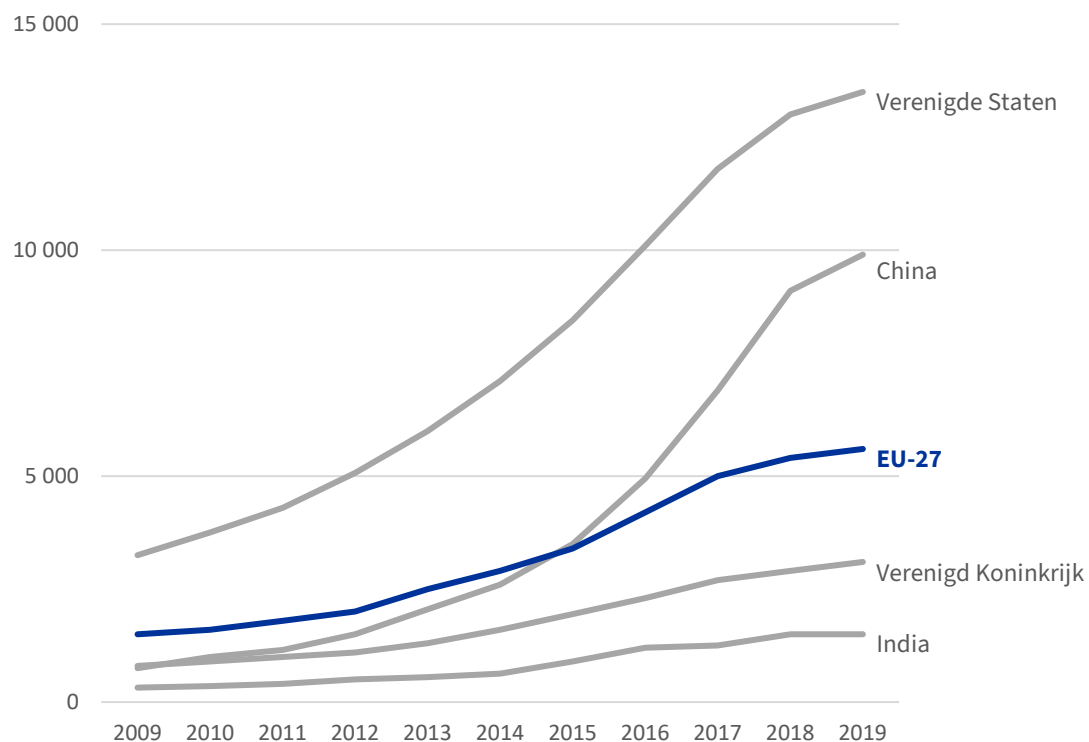


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

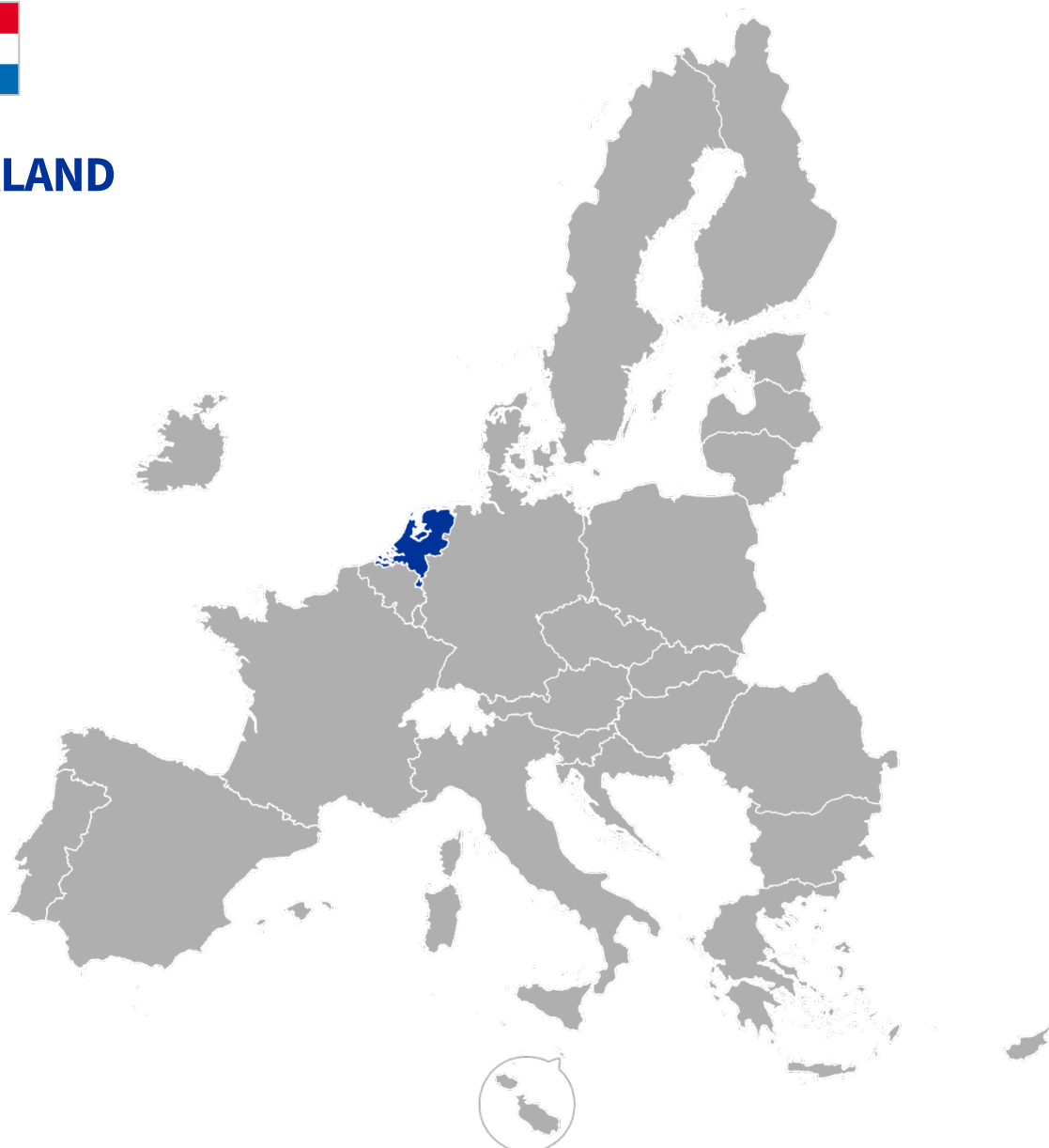
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



NEDERLAND



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



NEDERLAND

Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Italië, Kroatië en Slovenië**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Het behoud van nationale waarden en economische welvaart is een topprioriteit voor jouw regering. Volgens het IMF kan tot 60% van alle banen ingrijpend veranderen door AI. Hierdoor zullen veel banen verloren gaan en zal er steeds meer vraag zijn naar werknemers die met AI uit de voeten kunnen.
- Onbeheerste werkloosheid zou een ernstige bedreiging vormen voor de Nederlandse economie. Daarom ben je een fervent voorstander van een kader voor tests uit voorzorg, zodat technologieën met een hoog risico het niet plots begeven en hele sectoren lamleggen.
- Je wilt ook de AI-geletterdheid in de hele Nederlandse samenleving bevorderen, zodat mensen vertekende of gemanipuleerde resultaten kunnen herkennen en niet blind op geautomatiseerde systemen vertrouwen. De capaciteitsgebonden aanpak dient dit doel beter, omdat deze gericht is op het daadwerkelijke risico op misbruik en manipulatie. Bepaalde capaciteiten zoals gezichtsherkenning zijn ongeacht de context vatbaar voor misbruik of falen.
- Een capaciteitsgebonden aanpak voorkomt dat die technologieën te zwak worden gereguleerd enkel en alleen omdat zij in sectoren met een laag risico worden gebruikt. Zo zijn er bijvoorbeeld AI-systemen voor gebruik in winkels die winkelaars volgen, gedrag voorspellen en door middel van gezichtsherkenning gerichte promoties aanbieden of "waardevolle" klanten aanduiden. Bij een doelgebonden aanpak zou dit soort systemen mogelijk onder de controle uit kunnen komen, ook al geven de onderliggende capaciteiten ervan reden tot ernstige bezorgdheid met betrekking tot privacy, profilering en discriminatie.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Hoewel je met betrekking tot artikel 1 voorstander bent van een kader voor tests uit voorzorg, geef je de voorkeur aan uitzonderingen als het op nationale veiligheid aankomt. De EU is niet bevoegd voor nationale veiligheid. Dit staat in artikel 4, lid 2, van het Verdrag betreffende de Europese Unie (VEU): "[D]e nationale veiligheid blijft uitsluitend de verantwoordelijkheid van elke lidstaat". Het heeft dan ook geen zin om de lidstaten voor te schrijven hoe zij AI-toepassingen voor nationale veiligheid moeten reguleren.
- Aan de andere kant zie je geen reden om start-ups van de verordening vrij te stellen. Het feit dat een bedrijf klein is, betekent niet dat zijn technologie niet schadelijk kan zijn. Bovendien zouden grote bedrijven om de testvereisten te omzeilen riskante projecten naar start-ups kunnen doorschuiven, omdat die toch zijn vrijgesteld. Zo'n uitzondering zou dus een gemakkelijk achterdeurtje bieden, en de EU moet ervoor zorgen dat dit niet gebeurt.
- Jouw opmerkingen:

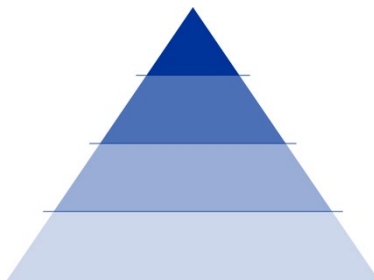
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde training-datasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland	capaciteitsgebonden	voorzorgsbenadering		nationale veiligheid	
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

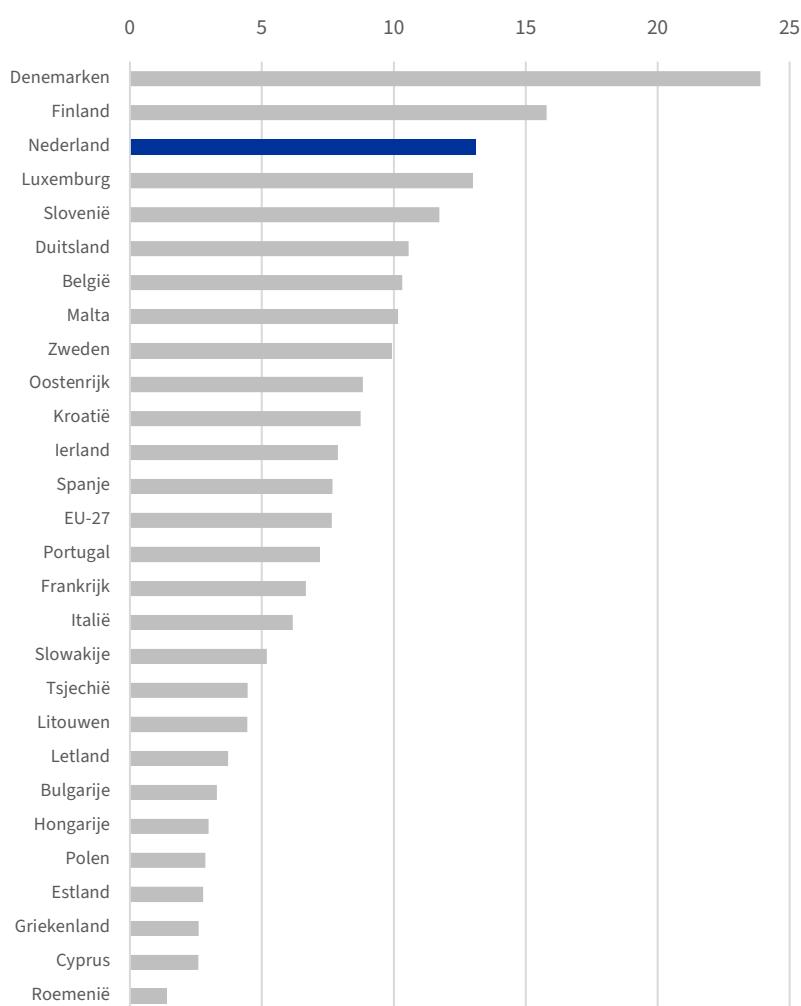
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

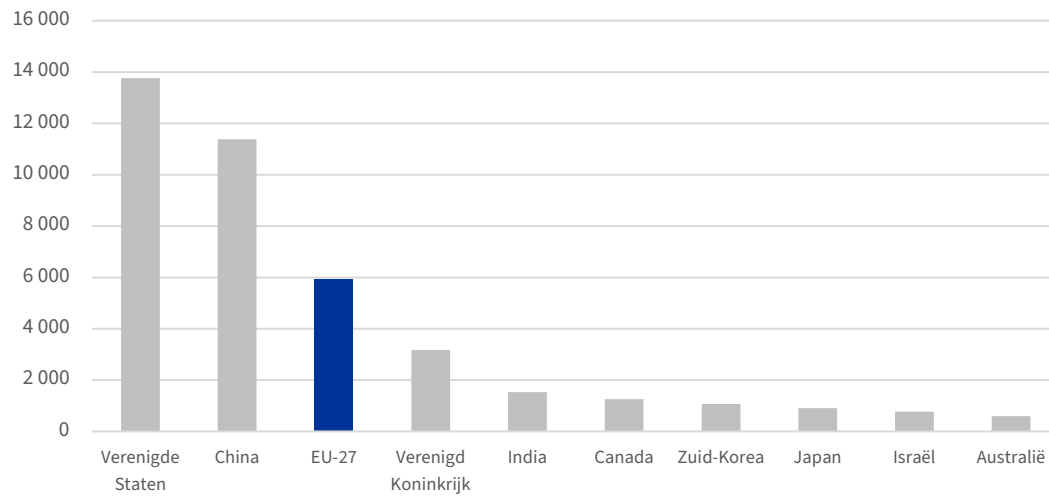
	Europese Unie	Nederland
Bevolking	447 695 350	17 811 291
Bbp per hoofd van de bevolking in EUR	38 150	59 720
Werkloosheidspercentage	6,1%	3,6%
Percentage bedrijven dat AI-technologie gebruikt	8%	14,1%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	54,5%

Percentage bedrijven met meer dan 10 werknemers
die minstens één AI-systeem gebruiken (2021)

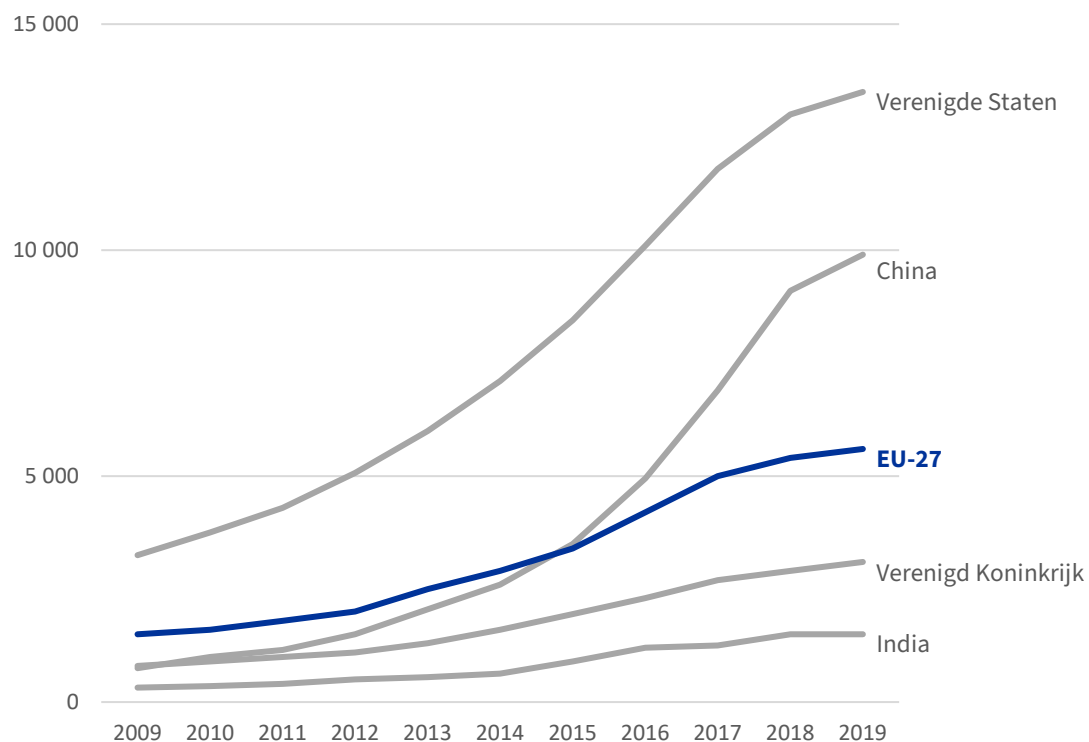


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

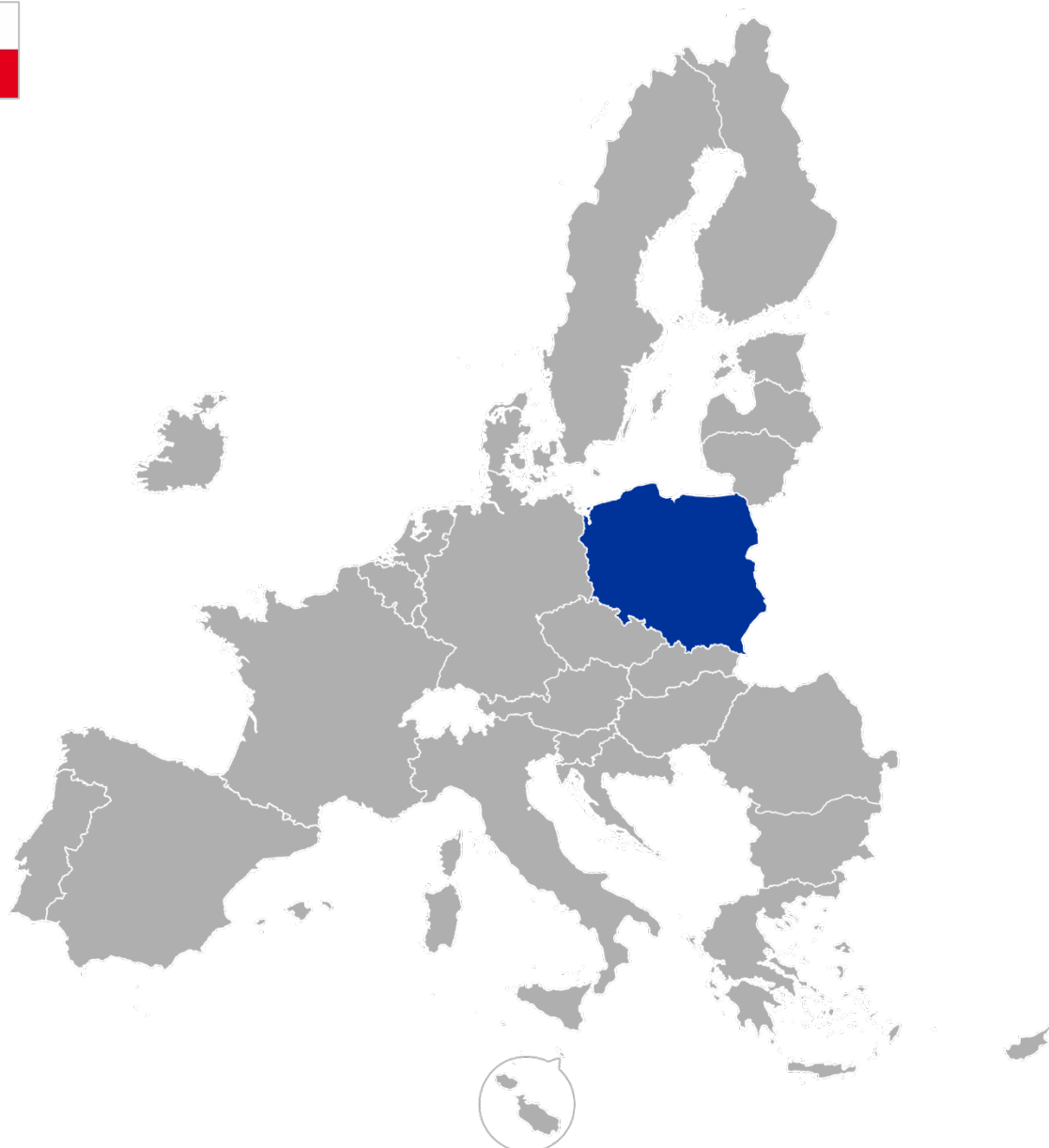
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



POLEN



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



POLEN

Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Denemarken, Frankrijk, Oostenrijk en Roemenië**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng tijdens de officiële onderhandelingen je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Polen heeft momenteel weinig technologie-reuzen op wereldschaal, waardoor voor kleine en middelgrote ondernemingen de kans om te groeien met AI voor het grijpen ligt. Jouw economie, die vooral draait op industrie, landbouw en logistiek, bevindt zich in een goede positie om profijt te trekken van de invoering van AI.
- Je steunt de doelgebonden aanpak vanwege de duidelijke stimulansen die deze biedt: sectoren met een laag risico krijgen meer vrijheid om te innoveren, terwijl sectoren met een hoog risico aan strikte EU-standaarden worden gehouden. Dit zal het vertrouwen in AI-systemen die worden gemaakt in Europa, die wereldwijd bekend staan om hun kwaliteit en veiligheid, vergroten.
- Je ziet in dat AI kan leiden tot vooroordelen, discriminatie en privacyrisico's, aangezien geen enkele dataset volledig neutraal is. Een voorzorgsbenadering betekent een strenger testkader op gebieden waar de kans op schade het grootst is, teneinde kwalijke uitwassen te voorkomen en de grondrechten te beschermen.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Je pleit ervoor dat militaire en defensiegerelateerde toepassingen vrijgesteld worden van de AI-verordening en benadrukt hierbij dat overregulering op dit gebied de snelle inzet van AI-technologieën in conflictgebieden of tijdens crisisresponsituaties kan belemmeren. Gezien de aanhoudende regionale instabiliteit, is flexibiliteit bij het gebruik van dergelijke technologieën van cruciaal belang.
- Als land dat grenst aan Rusland, Oekraïne en Belarus maakt Polen zich ernstig zorgen over zijn nationale veiligheid en de veiligheid van zijn burgers. Daarnaast kampt Polen langs zijn grens met Belarus met aanzienlijke uitdagingen als gevolg van toegenomen vluchtelingenstromen.
- Polen hoopt dan ook dat het land op meer begrip en steun van zijn medelidstaten kan rekenen om het hoofd te bieden aan deze unieke geopolitieke en veiligheidsuitdagingen.
- Jouw opmerkingen:

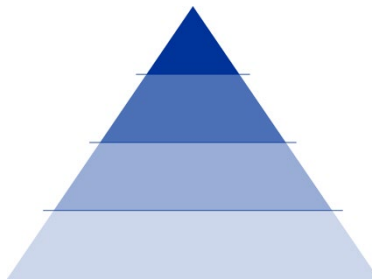
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

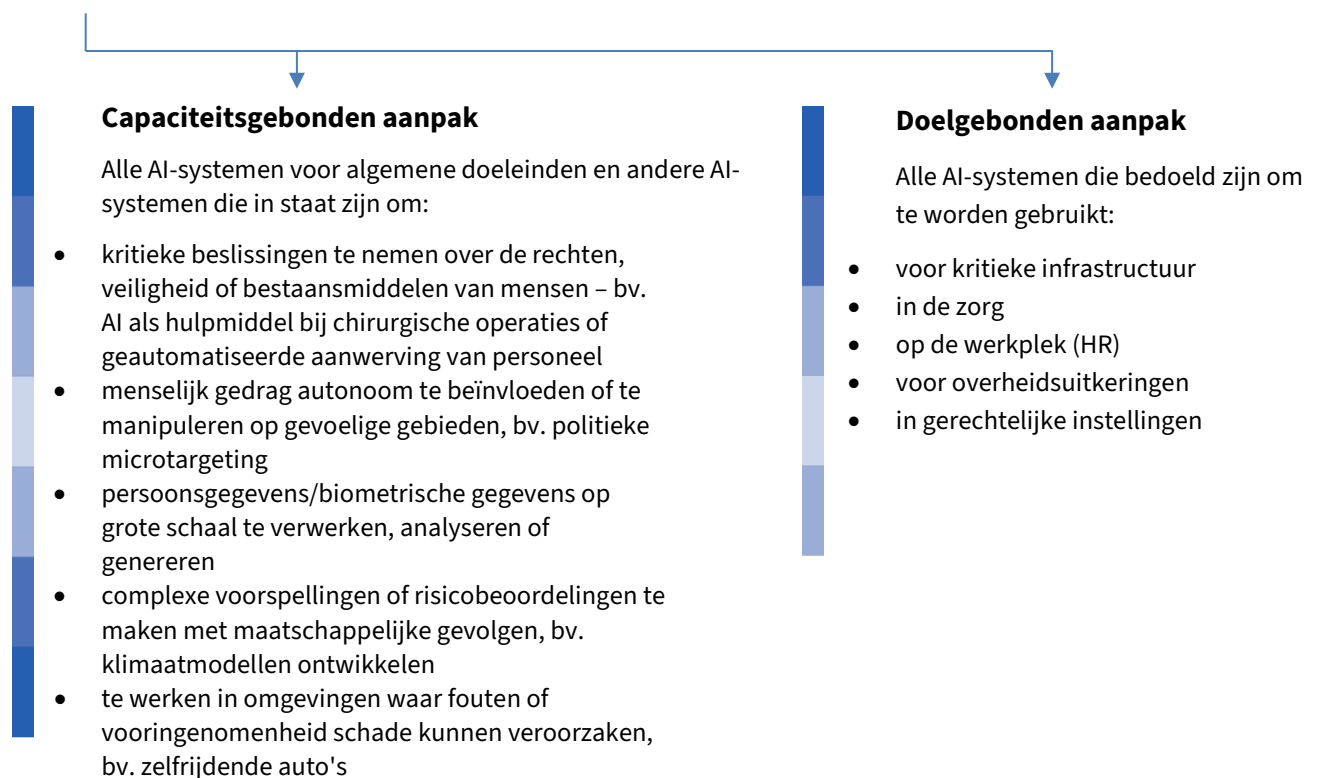
Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

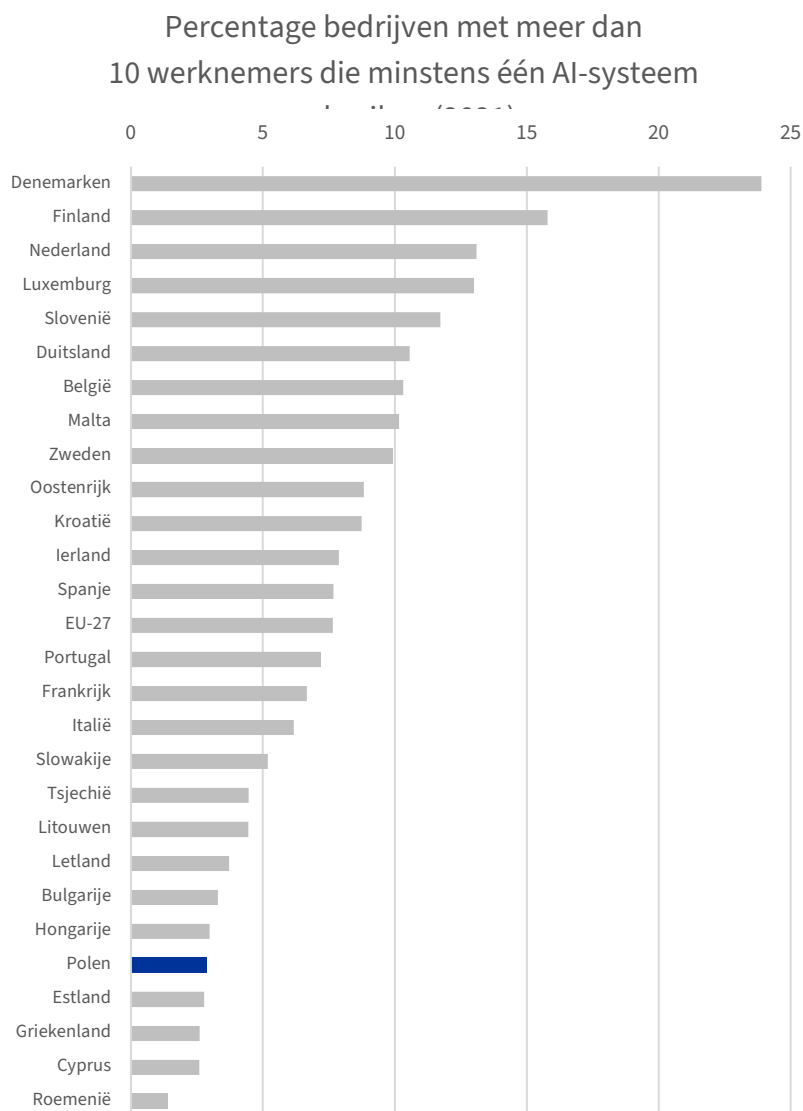
	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen	doelgebonden	voorzorgsbenadering		leger/defensie	
Portugal					
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

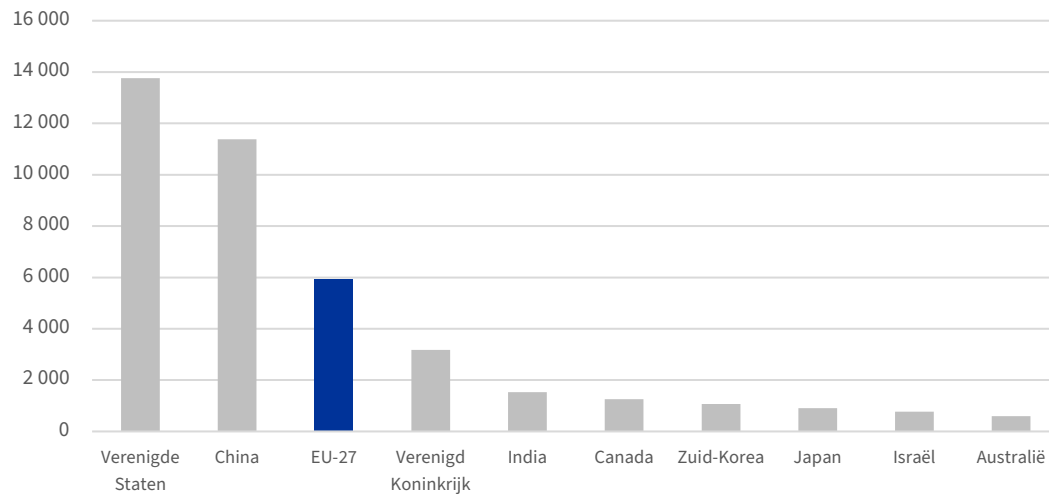
Let op: **De onderhandelingen vinden plaats in 2023!**

	Europese Unie	Polen
Bevolking	447 695 350	36 753 736
Bbp per hoofd van de bevolking in EUR	38 150	19 980
Werkloosheidspercentage	6,1%	2,8%
Percentage bedrijven dat AI-technologie gebruikt	8%	3,7%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	20,1%

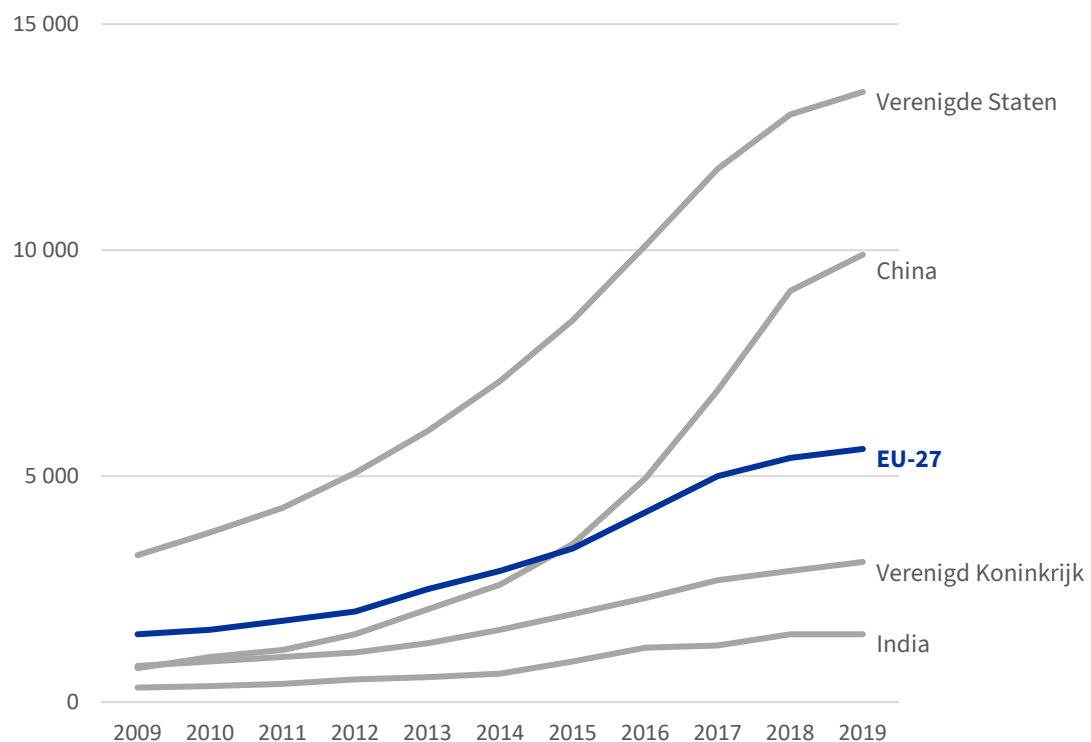


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

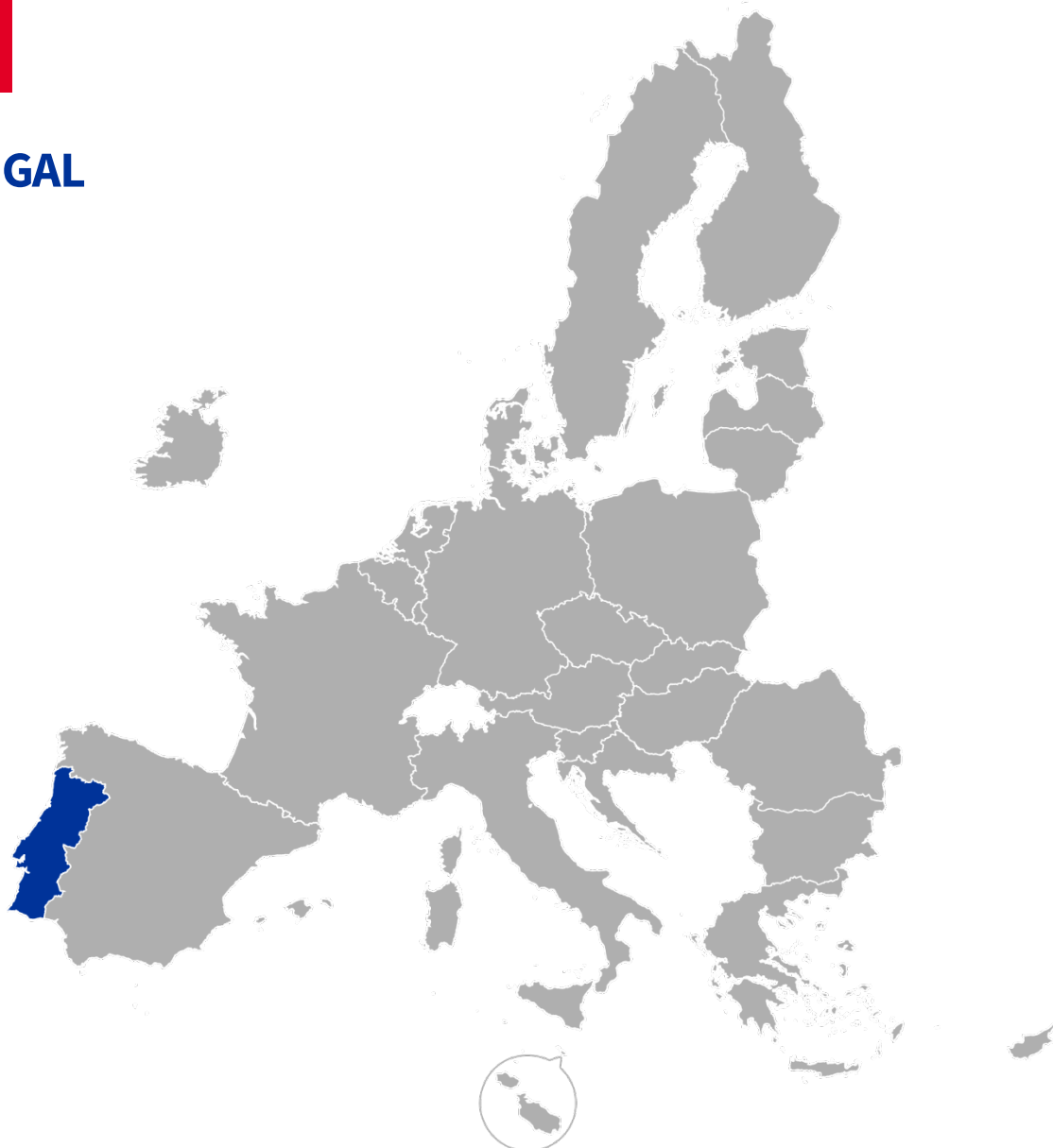
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



PORTUGAL



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



PORTUGAL

Agenda

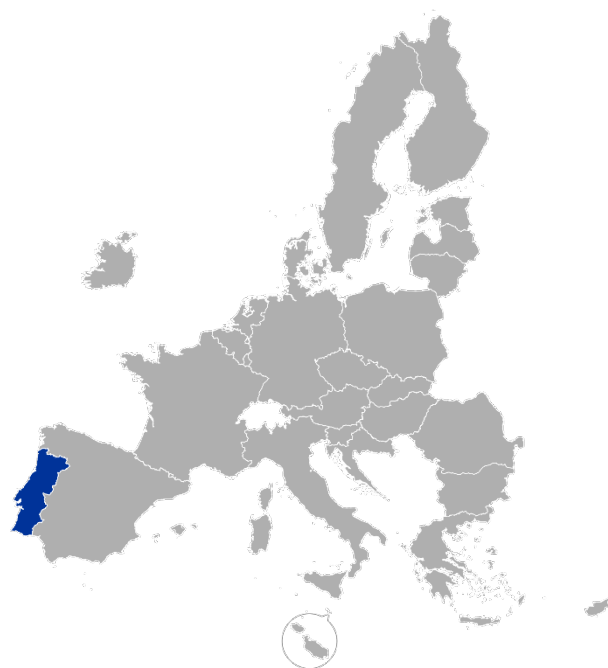
Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Finland, Ierland en Zweden**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____.

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____.

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____.

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____.

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Portugal is vanwege de technische vaardigheden van zijn werknemers een zeer aantrekkelijke bestemming voor internationale hightechbedrijven. Eén zo'n bedrijf is Tekever, dat AI-modellen ontwikkelt voor autonome dronebewaking in zowel commerciële als militaire contexten. In het bedrijfsleven worden deze drones ingezet om pijpleidingen te monitoren op lekken of sabotage en zodoende kritieke infrastructuur te beschermen. Het Europees Agentschap voor maritieme veiligheid (EMSA) maakt ook gebruik van Tekever-systemen ter ondersteuning van de kustwacht en voor milieubescherming.
- Je bent voorstander van de doelgebonden aanpak, maar je vreest dat onduidelijke sectorbeschrijvingen de AI-verordening zouden verzwakken of verwarrend zouden maken. Je wilt het begrip "kritieke infrastructuur" beter afbakenen om onder andere te verduidelijken dat elektriciteitsopwekking, telecommunicatie, watervoorziening, landbouw, verwarming en vervoer allemaal onder deze categorie vallen.
- Je vreest dat een kader met tests uit voorzorg zwaardere financiële en administratieve lasten voor AI-ontwikkelaars met zich meebrengt.
- Je hoopt het vandaag eens te worden over de definitieve tekst van de AI-verordening, zodat deze als leidraad kan dienen voor de betrouwbare ontwikkeling van AI in Europa, maar je bent fel gekant tegen een herzieningsclausule, aangezien dit funest zou zijn voor het vertrouwen van investeerders. Als er onzekerheid blijft, **zou het verantwoordelijker zijn om de stemming uit te stellen dan om al bij voorbaat in te plannen de regels na een tijdje om te gooien.**

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Hoewel je de flexibele regelgeving in artikel 1 steunt, ben je nog steeds van mening dat er duidelijke regels moeten gelden voor alle sectoren, ook voor de nationale veiligheids- en inlichtingendiensten. Deze sectoren van de verordening uitsluiten zal voor mazen in de wetgeving zorgen en een groter risico met zich meebrengen dat onethische of schadelijke AI niet wordt gecontroleerd.
- Het gaat hier niet om het tegenhouden van het gebruik van AI voor nationale veiligheidsdoeleinden — het gaat erom dat AI op verantwoorde wijze wordt ontworpen om vooringenomenheid en discriminatie te voorkomen.
- Alleen voor opensourcemodelen zou je eventueel een uitzondering willen maken. Die stimuleren snelle innovatie, omdat het ontwikkelingsproces ervan transparant is. Ontwikkelaars experimenteren en leren samen, en wisselen kennis uit om tools te verbeteren. Opensource-AI is vaak beter controleerbaar dan beschermde modellen en een betere controle van AI-systemen is een van de doelstellingen van de verordening.
- Jouw opmerkingen:

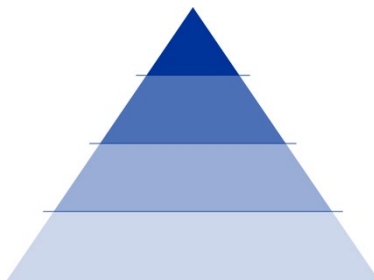
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

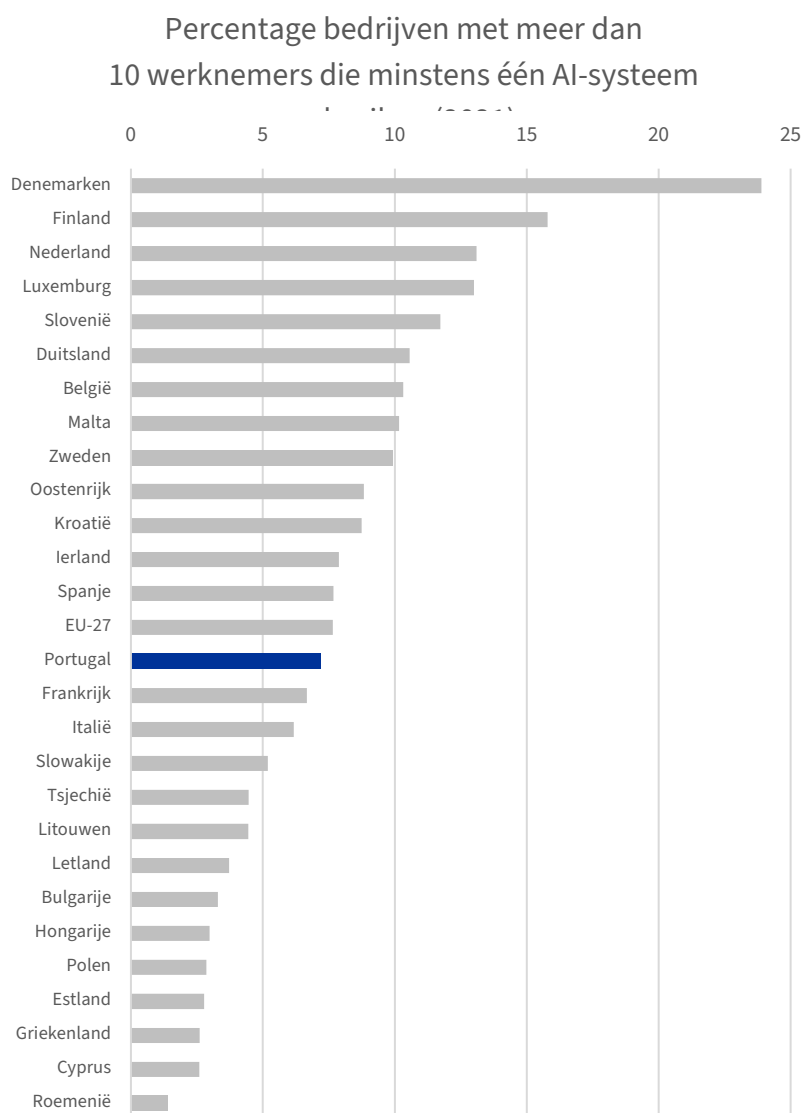
	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal	doelgebonden	flexibele benadering	de stemming uitstellen	opensource	
Roemenië					
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

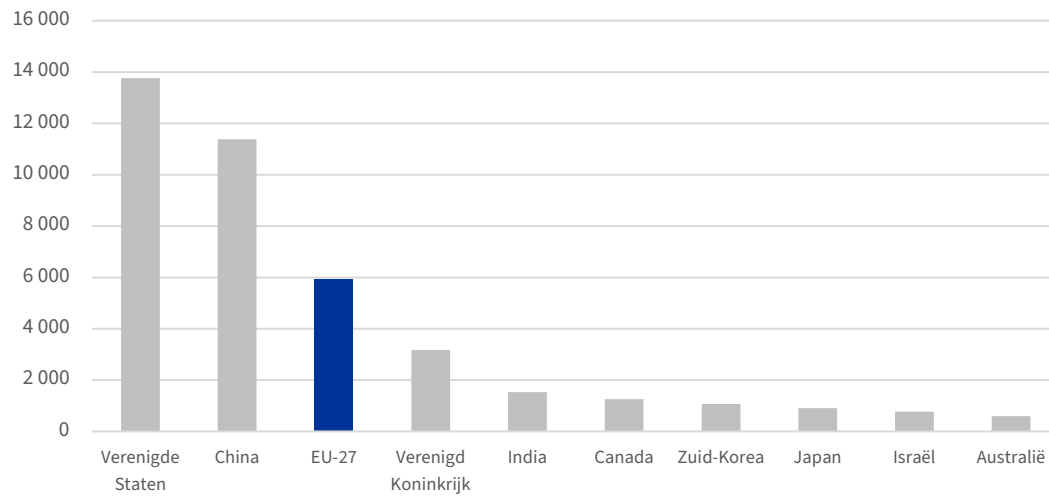
Let op: **De onderhandelingen vinden plaats in 2023!**

	Europese Unie	Portugal
Bevolking	447 695 350	10 516 621
Bbp per hoofd van de bevolking in EUR	38 150	25 330
Werkloosheidspercentage	6,1%	6,5
Percentage bedrijven dat AI-technologie gebruikt	8%	7,9
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	29,9

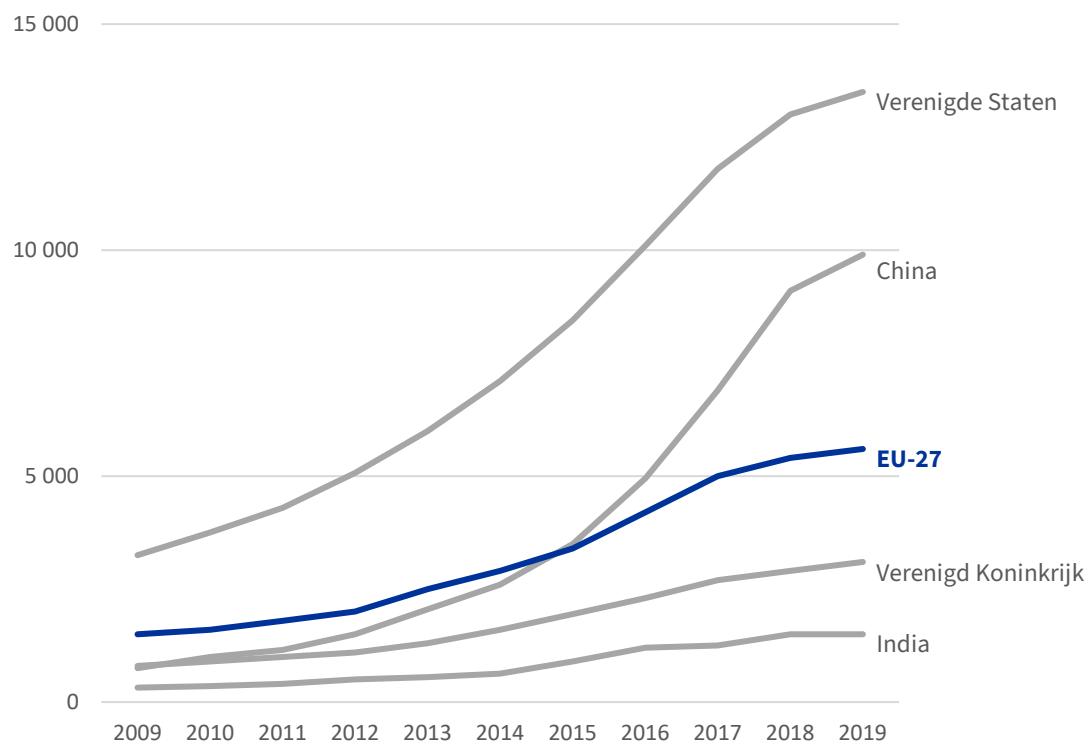


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

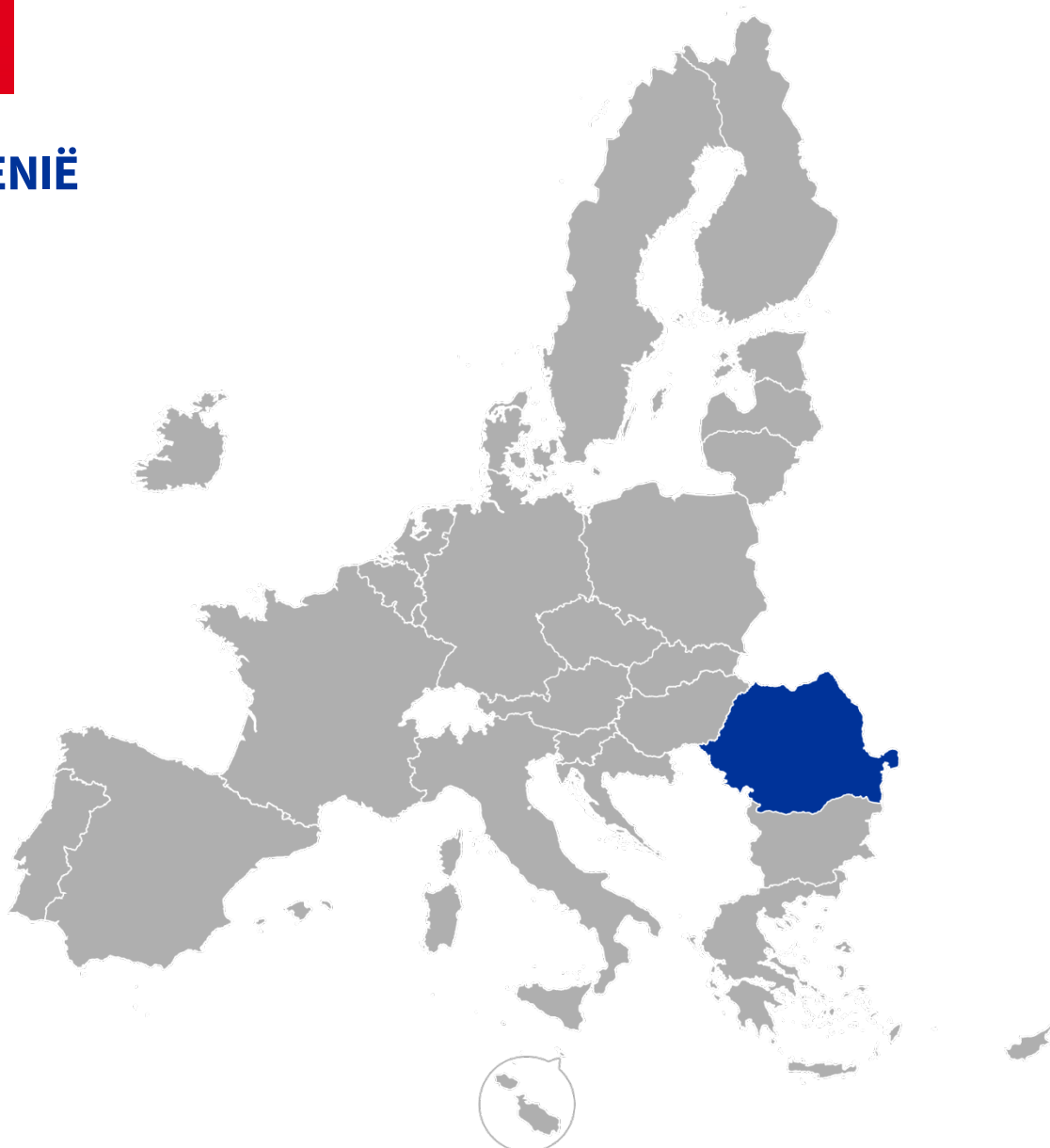
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



ROEMENIË



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Denemarken, Frankrijk, Oostenrijk en Polen**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____ .

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____ .

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____ .

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____ .

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Jij wilt het leven van de Roemeense bevolking verbeteren en ervoor zorgen dat de ontwikkeling van AI een zo positief mogelijke impact heeft op de samenleving.
- De landbouw is een belangrijke sector waar AI jouw nationale economie aardig kan oppeppen en met dit doel voor ogen steun je de lokale ontwikkeling van AI. Je bent voorstander van een sectorale aanpak, waarbij essentiële sectoren zoals de landbouw worden ontzien en er op gevoeliger gebieden aan strengere regels moet worden voldaan.
- De vereisten voor de ontwikkeling van AI met een hoog risico moeten echter evenredig blijven, zodat het kader rechtszekerheid biedt zonder ondernemingen en consumenten onnodig te belasten of op kosten te jagen. Daarom pleit je voor enkele afwijkingen van de vereisten voor ontwikkelaars.
- Je benadrukt het belang van een AI-kader dat alle EU-lidstaten, niet alleen de rijke landen zoals Duitsland of Frankrijk, ondersteunt. Mochten de besprekingen van vandaag vastlopen en mocht een compromis buiten bereik lijken, dan ben je bereid een capaciteitsgebonden aanpak te aanvaarden. In dat geval wil je echter wel laten verduidelijken dat de landbouw een sector met een laag risico is opdat die sector niet te maken krijgt met strikte testvereisten.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)

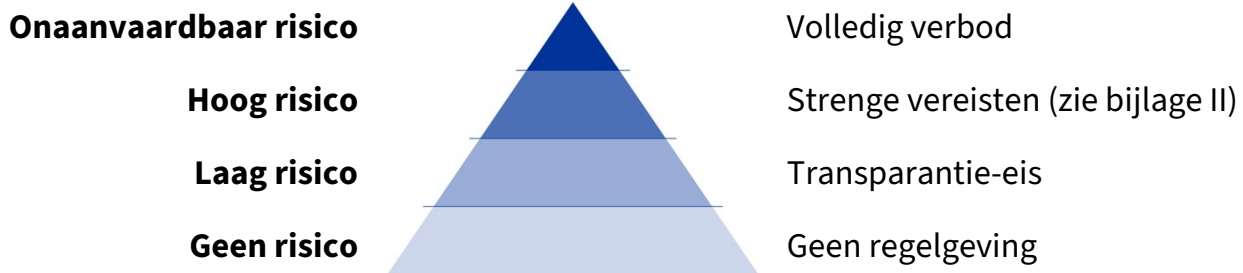


... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Hoewel de landbouw de sector is waarin je de meeste voordelen van AI verwacht, ga je ook uit van aanzienlijke voordelen op het gebied van rechtshandhaving. Je vreest echter dat als rechtshandavingsinstanties aan testvereisten, hoe gering ook, moeten voldoen, het risico bestaat dat toezichthouders inzage krijgen in gevoelige algoritmen of operaties. Dit kan vervolgens leiden tot mogelijke lekken en lopende onderzoeken in het gedrang brengen.
- Daarom pleit je ervoor toepassingen op nationaal veiligheidsgebied vrij te stellen van de verordening, waarbij je benadrukt dat politie- en nationale veiligheidsdiensten snel en flexibel moeten kunnen optreden bij de bestrijding van criminaliteit en terrorisme.
- Als je verzoek tijdens de discussie van vandaag niet volledig kan worden ingewilligd, ben je bereid tot een compromis waarin over twee jaar opnieuw wordt gekeken naar de verordening om deze mogelijk te herzien.
- Jouw opmerkingen:

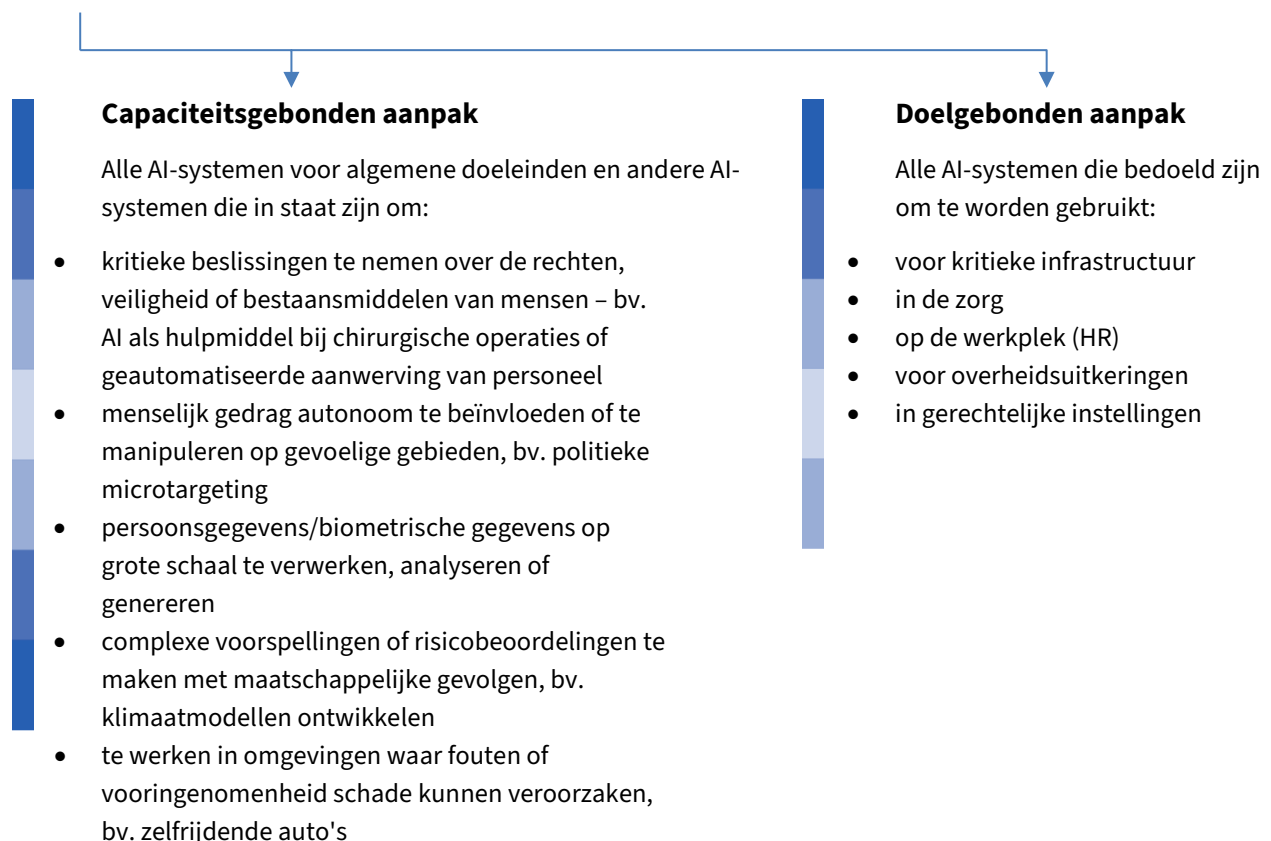
Bijlage I: Risicoclassificatie



Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië	doelgebonden	flexibele benadering		nationale veiligheid	
Slowakije					
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

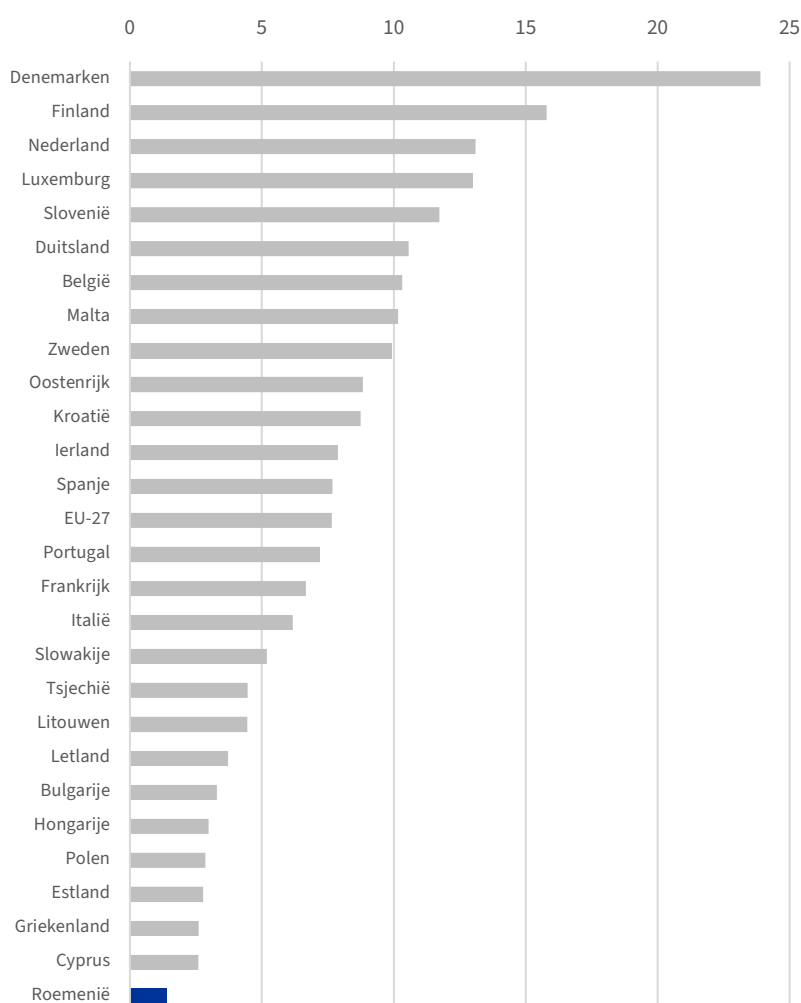
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

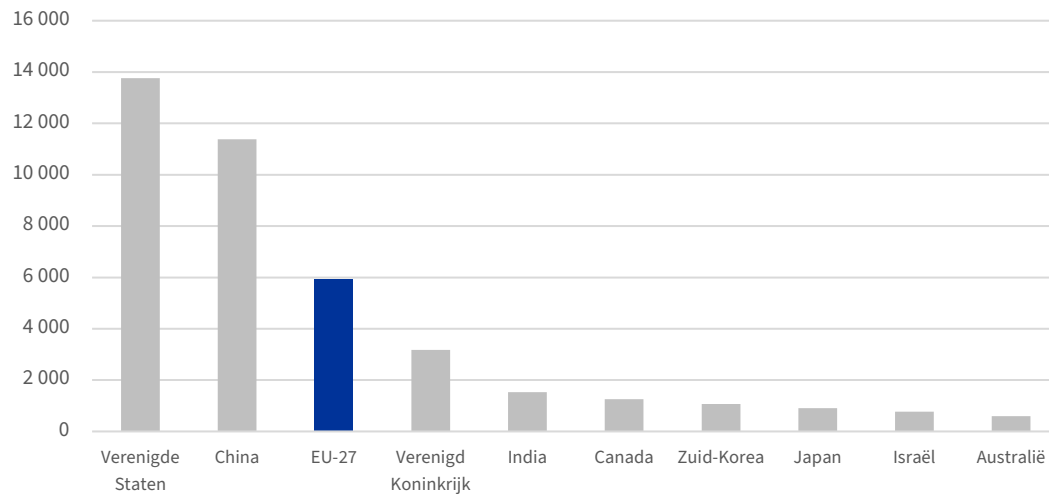
	Europese Unie	Roemenië
Bevolking	447 695 350	19 054 548
Bbp per hoofd van de bevolking in EUR	38 150	17 010
Werkloosheidspercentage	6,1%	5,6%
Percentage bedrijven dat AI-technologie gebruikt	8%	1,5%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	9%

Percentage bedrijven met meer dan
10 werknemers die minstens één AI-systeem
gebruiken (2021)

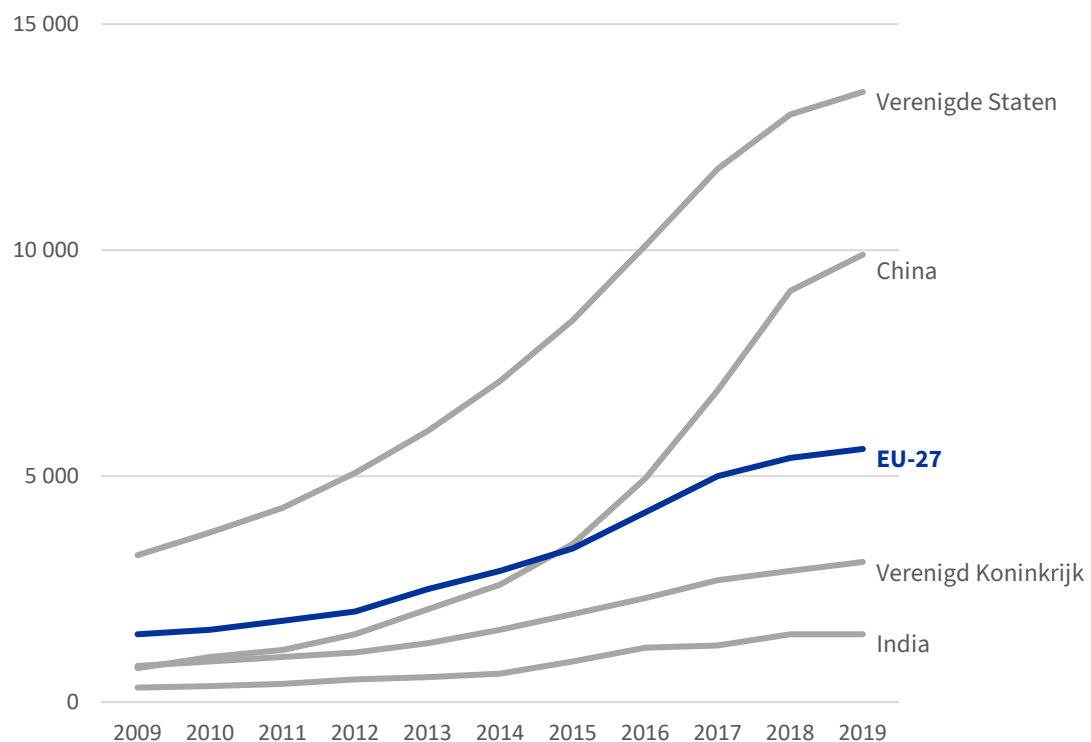


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

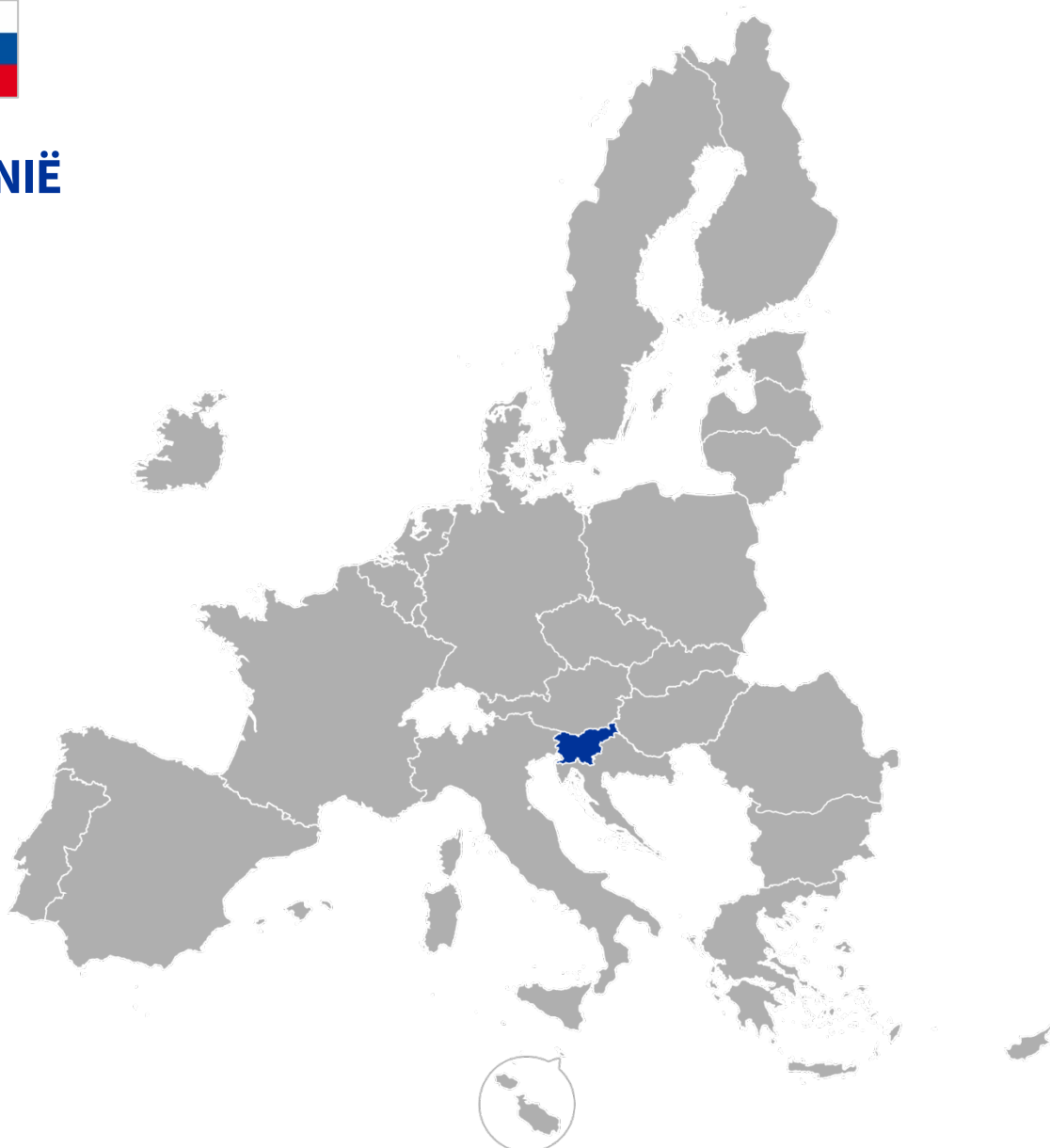
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



SLOVENIË



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

■ Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL



Raad van de
Europese Unie

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



SLOVENIË

Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Italië, Kroatië en Nederland**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng tijdens de officiële onderhandelingen je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____.

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____.

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____.

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____.

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Jouw land zet zich in voor een alomvattende en efficiënte digitale transformatie die geworteld is in een leven lang leren, intergenerationele samenwerking en het koesteren van Sloveens digitaal talent. Je beschouwt AI als een instrument dat bedoeld is om de mens te dienen.
- Je steunt de capaciteitsgebonden aanpak, want deze biedt een realistischere kijk op de werkelijke risico's, met name op het gebied van privacy. Je vindt het zeer redelijk om alle AI-systemen die persoonsgegevens kunnen verwerken, analyseren en genereren, te classificeren als "hoog risico", zoals voorgesteld in de capaciteitsgebonden aanpak.
- Tijdens de discussie van vandaag wil je de aandacht vestigen op de mogelijke risico's voor de mensenrechten, justitie en menselijke activiteiten. AI-systemen herbergen aanzienlijk grotere risico's wanneer mensen niet worden voorgelicht over hoe ze werken, wat ze kunnen en wat de beperkingen ervan zijn. Zonder basiskennis over AI vertrouwen mensen soms blindelings op geautomatiseerde systemen, hebben ze moeite vooroordelen of manipulatie te herkennen en kunnen ze ontwikkelaars of instellingen niet ter verantwoording roepen. Een dergelijk gebrekkig begrip van de technologie maakt mensen ook vatbaarder voor misinformatie, vergroot de sociale ongelijkheid en kan leiden tot irrationele angst voor of juist ondoordachte aanvaarding van AI-technologieën. Om ervoor te zorgen dat AI veilig en eerlijk wordt gebruikt, is samenlevingsbrede openbare voorlichting erover onontbeerlijk.
- Als het belang van publieksvoorlichting wordt benadrukt in de AI-verordening, sta je ervoor open om AI-ontwikkelaars meer vrijheid te gunnen.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- Volgens jou moet de Europese AI-verordening als katalysator dienen voor de ontwikkeling van AI in alle EU-lidstaten. Sommige kapitaalkrijke landen hebben al toonaangevende AI-bedrijven, terwijl andere landen nog bezig zijn met het tot stand brengen van een ecosysteem voor start-ups. Het hele idee achter start-ups is om op korte termijn en met beperkte middelen disruptieve technologie te creëren. Dat is een haast onmogelijke opgave wanneer je omkomt in administratieve rompslomp.
- Als artikel 1 resulteert in een voorzorgsbenadering voor het testen van AI-modellen met een hoog risico — waar jij voorstander van bent — stel je een uitzondering voor voor start-ups met minder dan 20 werknemers en een omzet van minder dan € 4 miljoen. Zij zouden in plaats daarvan de flexibele aanpak moeten volgen. Je ziet dit als een evenwichtig compromis en wilt het gesprek daarover leiden.
- Jouw opmerkingen:

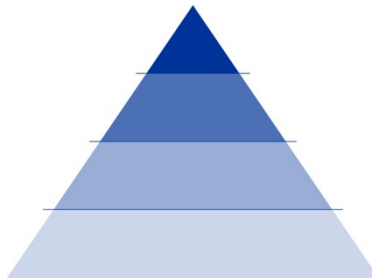
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde trainingdatasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

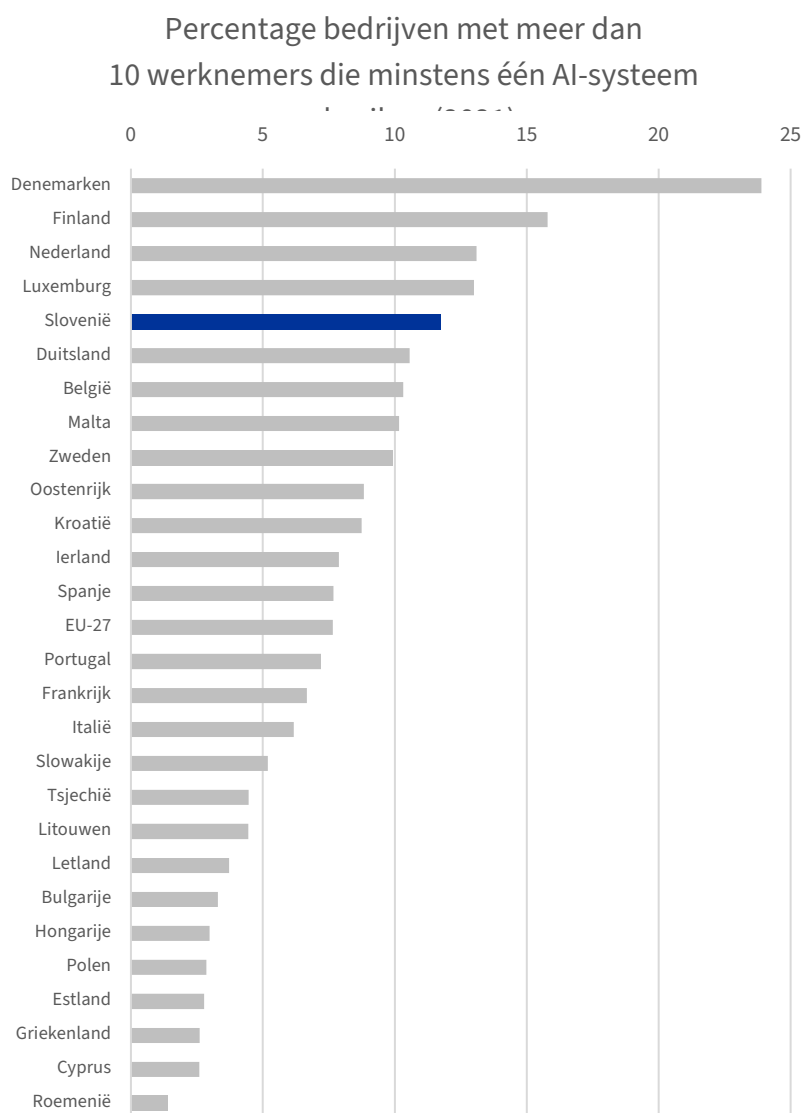
	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkinge n
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije					
Slovenië	capaciteitsgebonden	voorzorgsbenadering		start-ups	
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

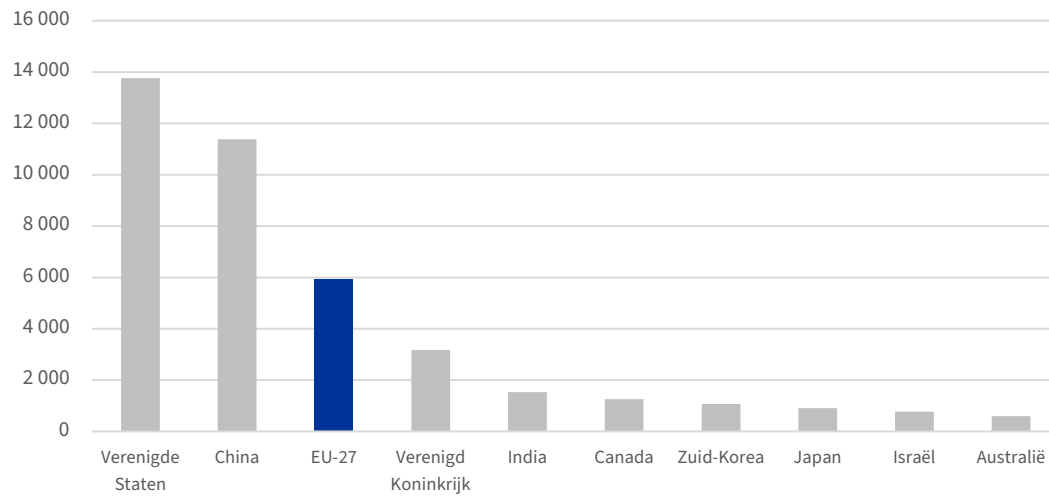
Let op: **De onderhandelingen vinden plaats in 2023!**

	Europese Unie	Slovenië
Bevolking	447 695 350	2 116 972
Bbp per hoofd van de bevolking in EUR	38 150	30 160
Werkloosheidspercentage	6,1%	3,7%
Percentage bedrijven dat AI-technologie gebruikt	8%	11,4%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	18,9%

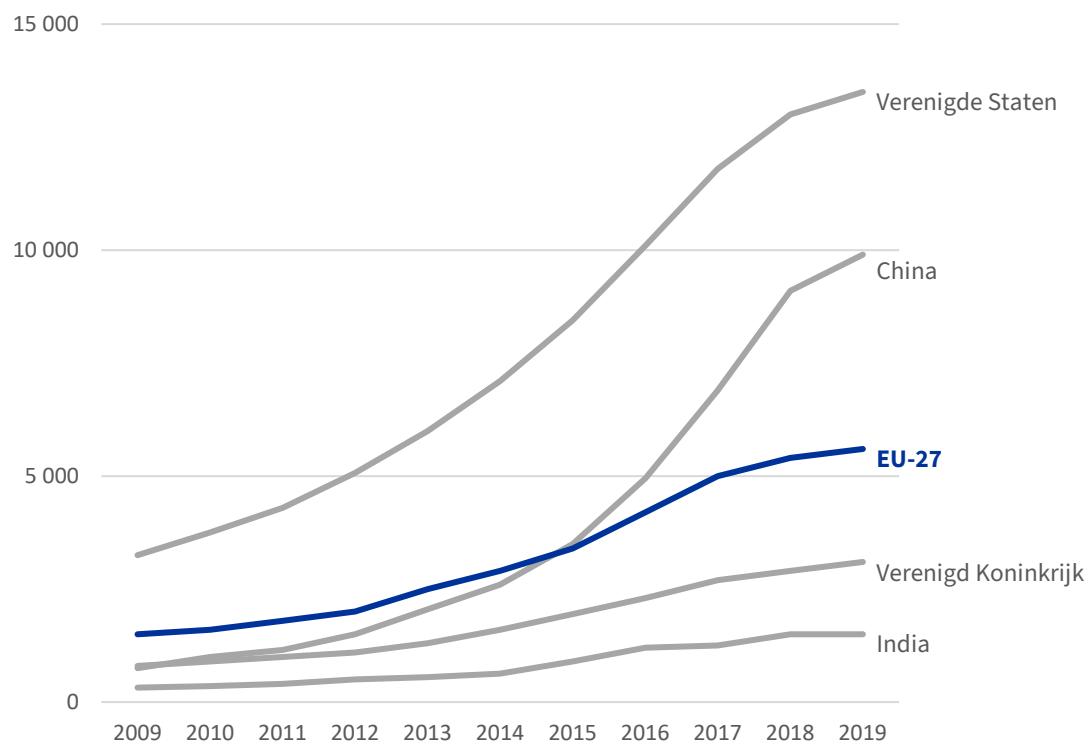


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

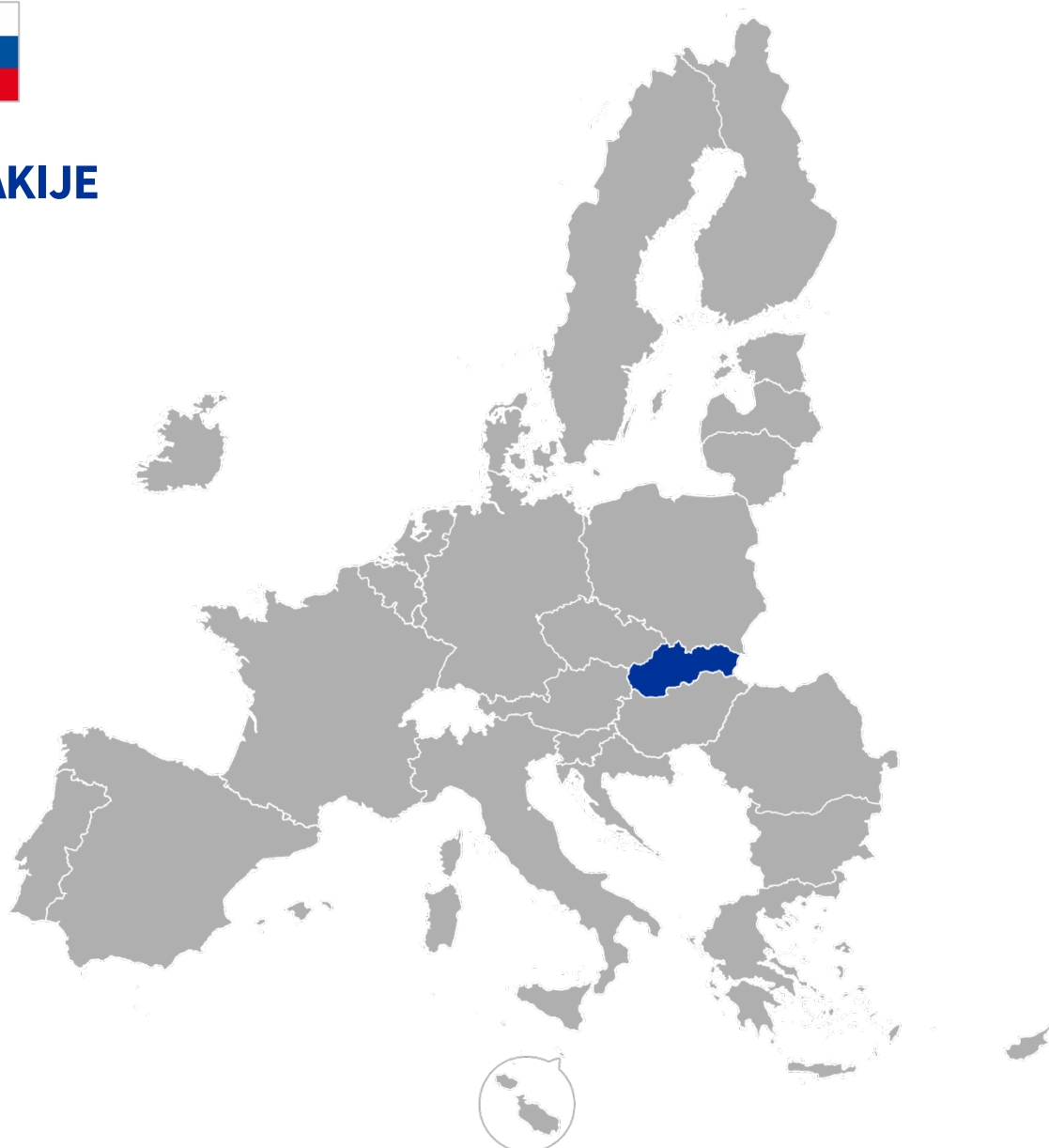
© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.



SLOWAKIJE



VOOR ÉÉN DAG MINISTER: ONDERHANDELEN MET 27

- Een juridisch kader voor AI

Rolprofielen | **Versie voor gevorderden** | NL

Voorstel voor een

VERORDENING VAN HET EUROPEES PARLEMENT EN DE RAAD

over een rechtskader voor het gebruik van artificiële intelligentie

ACHTERGROND VAN HET VOORSTEL

De Europese Unie is op een kritiek punt aanbeland: ze moet nu besluiten hoe ze wil dat haar digitale toekomst eruit gaat zien. Als onderdeel van haar strategie voor het digitale decennium wil de EU het voortouw nemen op het gebied van artificiële intelligentie (AI). Tegelijkertijd wil ze garant staan voor de grondrechten, consumentenbescherming, het concurrentievermogen, innovatie en technologisch leiderschap. Het debat focust zich vooral op een wetgevingsvoorstel van de Europese Commissie dat de ontwikkeling van AI in de EU moet reguleren. Er zijn al verschillende onderhandelrondes over het voorstel geweest maar er moet nog overeenstemming worden gevonden over de onderstaande kwesties.

Artikel 1

Classificatie van AI-systemen & risicobeheer

- A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren**.
- B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars een **voorzorgsbenadering volgen: ze moeten aan strenge gestandaardiseerde testvereisten voldoen** (zie bijlage II).

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen, **behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden**.



Agenda

Op de agenda van de Raad staat vandaag het voorstel van de Commissie voor een rechtskader voor het gebruik van artificiële intelligentie. Aan jou om het ontwerpvoorstel te bespreken met de vertegenwoordigers van de andere lidstaten en een compromis te vinden.

Jouw hoofddoel

Een gemeenschappelijk standpunt van de Raad dat aansluit bij jouw belangen. Je kunt ook een compromis steunen dat gedeeltelijk bij jouw belangen aansluit.

Jouw opdracht

1. **Lees zorgvuldig de gegeven informatie** over jouw belangen en standpunten.
2. Sommige van **jouw argumenten sluiten aan bij de standpunten van landen als Bulgarije, Luxemburg en Malta**. Praat met hen om te zien of je kunt samenwerken om jullie gemeenschappelijke belangen te behartigen.
3. Stel een **openingsverklaring** van maximaal 1 minuut op voor de eerste ronde van de officiële onderhandelingen. Geef je standpunt over de artikelen 1 en 2. Het script in het tekstvak kan je hierbij helpen.
4. Leg je openingsverklaring af wanneer het voorzitterschap (dat de zitting leidt) jou het woord geeft. Als je het eens bent met standpunten van andere landen die al aan het woord zijn geweest, geef dat dan aan.
5. Luister goed naar de standpunten van de andere landen. **Vul de tabel op bladzijde 8 en 9 in** om een overzicht van hun onderhandelingsposities te hebben. Bijzonder belangrijk bij de onderhandelingen is om naar elkaar te luisteren en tot oplossingen te komen waarmee iedereen kan leven.
6. **Breng** tijdens de officiële onderhandelingen **je standpunten en argumenten naar voren**. Als andere landen dezelfde argumenten hebben, kan het nuttig zijn jullie argumenten samen te presenteren – zo krijgen ze meer gewicht.
7. Natuurlijk **beslis je zelf hoe je stemt aan het einde**. Wees je ervan bewust dat je bij de definitieve stemming zal moeten kiezen tussen enerzijds een compromis dat jou misschien niet 100% tevreden stemt, en anderzijds onderhandelingen zonder resultaat.



Geachte voorzitter, commissaris, collega's,

Ik wil eerst en vooral de Commissie bedanken voor haar voorstel en het voorzitterschap voor het opnemen van dit punt op de agenda.

In de aanloop naar deze zitting hebben we vruchtbare besprekingen gehouden met _____.

Wij achten het noodzakelijk om een capaciteits-/doelgebonden risicoclassificatie in te voeren, omdat _____.

Wat artikel 1, punt B, betreft, zijn wij van mening dat ontwikkelaars bij de ontwikkeling van AI-systemen met een hoog risico een flexibele benadering/voorzorgsbenadering moeten volgen, omdat _____.

Wat artikel 2 betreft, zullen er bepaalde uitzonderingen/geen uitzonderingen zijn, omdat wij voorstander zijn van _____.

Wij zijn ervan overtuigd dat we tijdens de zitting van vandaag een werkbaar compromis zullen bereiken.

Dank u wel.

Artikel 1

Classificatie van AI-systemen & risicobeheer

A) De EU voert een gemeenschappelijke norm in voor risicobeheer bij de ontwikkeling, het in de handel brengen, de invoer en het gebruik van AI-systemen. Deze systemen krijgen de classificatie "geen risico", "laag risico", "hoog risico" of "onaanvaardbaar risico" (zie bijlage I). Om te bepalen welke AI-systemen een hoog risico met zich meebrengen, maakt de EU gebruik van een



... capaciteitsgebonden risicoclassificatie: het risico wordt bepaald op basis van de technische capaciteiten van het systeem.



... **doelgebonden risicoclassificatie: het risico wordt bepaald op basis van het beoogde gebruik van het systeem in gevoelige sectoren.**

AI-systemen die worden geclassificeerd als systeem met een hoog risico moeten uitgebreid worden getest voordat ze worden ingezet.

B) Tijdens de ontwikkeling van AI-systemen met een hoog risico moeten ontwikkelaars gebruikmaken van een



... **flexibele benadering: een zelfevaluatie van het AI-model uitvoeren voordat het wordt ingezet.**



... voorzorgsbenadering: aan de strenge gestandaardiseerde testvereisten voldoen (zie bijlage II).

Jouw argumenten

- Slowakije wil graag kijken hoe het zoveel mogelijk voordeel kan halen uit de AI-verordening van de EU. Tot dusver zijn er geen overheidsmiddelen toegewezen aan nationaal AI-onderzoek. Jouw regering is dan ook afhankelijk van particuliere investeringen of steun van andere EU-lidstaten.
- Je bent voorstander van de doelgebonden aanpak, omdat dit de aanpak is die de EU doorgaans gebruikt om technologieën en producten, zoals medische hulpmiddelen en chemische producten, te reguleren. Je bent echter van mening dat de huidige lijst van sectoren met een hoog risico moet worden beperkt tot kritieke infrastructuur en de gezondheidszorg, aangezien dit sectoren zijn waar het rechtstreeks om leven en dood kan gaan. De andere sectoren zijn ook belangrijk, maar daar kan later nog opnieuw naar worden gekeken. Je staat ervoor open om de lijst over 3 jaar te herzien, in overleg met het maatschappelijk middenveld en het bedrijfsleven.
- Wat de testvereisten betreft, ben je bang dat te strenge regels kunnen leiden tot een kennisvlucht, waardoor Europees AI-talent naar flexibelere omgevingen zoals de VS wordt gedreven. Een flexibele aanpak is ook in het voordeel van Slowakije omdat het een van de kleinere EU-lidstaten is en dus over weinig middelen beschikt om in AI te investeren.
- Mochten de besprekingen van vandaag vastlopen en mocht een compromis buiten bereik lijken, dan ben je bereid iets strengere testvereisten te aanvaarden. Als onderdeel van een evenwichtig compromis moet echter **financiële steun beschikbaar worden gesteld aan minder welvarende EU-lidstaten** om de hogere ontwikkelingskosten te helpen dekken.

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Artikel 2

Toepassingsgebied

De bepaling in artikel 1 is van toepassing op alle AI-systemen,



... zonder uitzonderingen



... zonder uitzonderingen, tenzij het systeem als opensourcesysteem wordt ontwikkeld



... behalve systemen die zijn ontwikkeld door kleine bedrijven of start-ups (en)



... behalve systemen die zijn ontwikkeld voor militaire of defensiedoeleinden (en)



... behalve systemen die zijn ontwikkeld voor nationale veiligheid, migratiebeheer of grenscontrole

Jouw argumenten

- In vergelijking met de rest van de wereld hinkt de EU voorlopig achterop op het gebied van AI. Start-ups vrijstellen van deze verordening zou dynamiek in het ecosysteem kunnen brengen. Start-ups hebben niet de juridische, financiële en administratieve capaciteit om aan dezelfde regels te voldoen als grote technologiebedrijven. Door een strikte naleving van de regels te verlangen, loop je het risico innovatie in de beginfase tegen te houden.
- Als je er niet in slaagt genoeg lidstaten ervan te overtuigen om in artikel 1 flexibele testvereisten op te nemen, wil je wel dat start-ups op een aantal gebieden van de vereisten mogen afwijken. Zo hoeven start-ups niet meteen aan alle strenge vereisten te voldoen (als daarvoor wordt gekozen bij artikel 1), maar kunnen zij eerst aan de slag in een experimentele omgeving.
- Jouw opmerkingen:

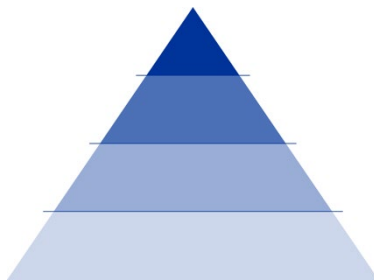
Bijlage I: Risicoclassificatie

Onaanvaardbaar risico

Hoog risico

Laag risico

Geen risico



Volledig verbod

Strengere vereisten (zie bijlage II)

Transparantie-eis

Geen regelgeving

Onaanvaardbaar risico:

- **AI voor social scoring** – systeem om burgers te beoordelen en rangschikken op basis van hun gedrag
- **Door AI aangedreven massasurveillance** – biometrische tracking zonder toestemming
- **Emotieherkenning op werkplekken/scholen** om prestaties te beïnvloeden
- **Autonome dodelijke wapens** die zonder menselijk ingrijpen kunnen doden
- **AI voor subliminale manipulatie** – besluiten van mensen beïnvloeden zonder daarover transparant te zijn

Hoog risico



Capaciteitsgebonden aanpak

Alle AI-systemen voor algemene doeleinden en andere AI-systemen die in staat zijn om:

- kritieke beslissingen te nemen over de rechten, veiligheid of bestaansmiddelen van mensen – bv. AI als hulpmiddel bij chirurgische operaties of geautomatiseerde aanwerving van personeel
- menselijk gedrag autonoom te beïnvloeden of te manipuleren op gevoelige gebieden, bv. politieke microtargeting
- persoonsgegevens/biometrische gegevens op grote schaal te verwerken, analyseren of genereren
- complexe voorspellingen of risicobeoordelingen te maken met maatschappelijke gevolgen, bv. klimaatmodellen ontwikkelen
- te werken in omgevingen waar fouten of vooringenomenheid schade kunnen veroorzaken, bv. zelfrijdende auto's

Doelgebonden aanpak

Alle AI-systemen die bedoeld zijn om te worden gebruikt:

- voor kritieke infrastructuur
- in de zorg
- op de werkplek (HR)
- voor overheidsuitkeringen
- in gerechtelijke instellingen

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Laag risico

- **AI-systemen die niet zijn geclassificeerd als systemen met een "hoog risico"**
- **Chatbots & virtuele assistenten** (zolang gebruikers weten dat ze met AI praten)
- **Aanbevelingsalgoritmen** – door AI aangereikte suggesties op sociale media
- **AI voor bedrijfsautomatisering** – AI die administratieve basistaken automatiseert

Geen risico

- Afbeeldingen, muziek of kunst gemaakt/bewerkt met AI
- AI in wetenschappelijk onderzoek
- AI in speelgoedrobots – door AI aangedreven speelgoed zonder echte risico's



Bijlage II: Vereisten voor AI-systemen met een hoog risico

Nadruk op wettelijke en ethische normen:

- Risicobeoordeling voorafgaand aan de uitrol, waarbij te voorziene risico's en mogelijk misbruik in kaart worden gebracht.
- Menselijk toezicht om de volledige automatisering van de besluitvorming te voorkomen.
- Gegevensbeheer om onbevooroordeelde training-datasets te waarborgen.
- Verplichte externe audits, waaronder sandboxtesten met menselijk toezicht.
- Bekendmaking van de testresultaten aan de regelgevende instanties van de EU.
- Controle na toelating op de markt om de prestaties op te volgen.



Definities

- **Algoritmen** – stapsgewijze instructies die een computer volgt om een probleem op te lossen of een beslissing te nemen
- **Artificiële-intelligentiesystemen** (AI-systemen) – computersystemen die AI gebruiken om taken uit te voeren waarvoor doorgaans menselijke intelligentie nodig is, bv. leren, redeneren, problemen oplossen, besluiten nemen, waarnemen en taal begrijpen
- **Kennisvlucht** – fenomeen waarbij geschoolde werknemers hun land verlaten om in het buitenland te gaan werken
- **Bedrijfsmigratie** – fenomeen waarbij een bedrijf zijn activiteiten naar een ander land verplaatst, vaak vanwege betere voorwaarden of lagere kosten
- **AI voor deep learning** – systeem dat met behulp van neurale netwerken leert van grote hoeveelheden gegevens, bv. hoe het menselijk brein werkt. Het herkent patronen en wordt na verloop van tijd beter zonder dat het daarvoor opnieuw moet worden geprogrammeerd
- **Artificiële intelligentie voor algemene doeleinden** – AI die meer dan één taak kan uitvoeren, bv. ChatGPT, Google Gemini
- **Opensourcesoftware** – software met een code die door iedereen kan worden bekeken, gebruikt en gewijzigd
- **Op regels gebaseerde AI** – AI-systeem dat beslissingen neemt op basis van vooraf vastgestelde regels en logica. Bijvoorbeeld: "Als X gebeurt, doe dan Y"
- **Sandboxtesten** – testen in een afgesloten omgeving, bv. een zelfrijdende auto met AI wordt getest in een gesimuleerde stad voordat deze op echte wegen gaat rijden

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Oostenrijk					
België					
Bulgarije					
Kroatië					
Cyprus					
Tsjechië					
Denemarken					
Estland					
Finland					
Frankrijk					
Duitsland					
Griekenland					
Hongarije					
Ierland					

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

	Artikel 1A: classificatie	Artikel 1B: benadering	Andere suggesties/opmerkingen	Artikel 2: toepassingsgebied	Andere suggesties/opmerkingen
Italië					
Letland					
Litouwen					
Luxemburg					
Malta					
Nederland					
Polen					
Portugal					
Roemenië					
Slowakije	doelgebonden	flexibele benadering	financiële steun	start-ups	
Slovenië					
Spanje					
Zweden					
Commissie	doelgebonden	voorzorgsbenadering		leger/defensie	

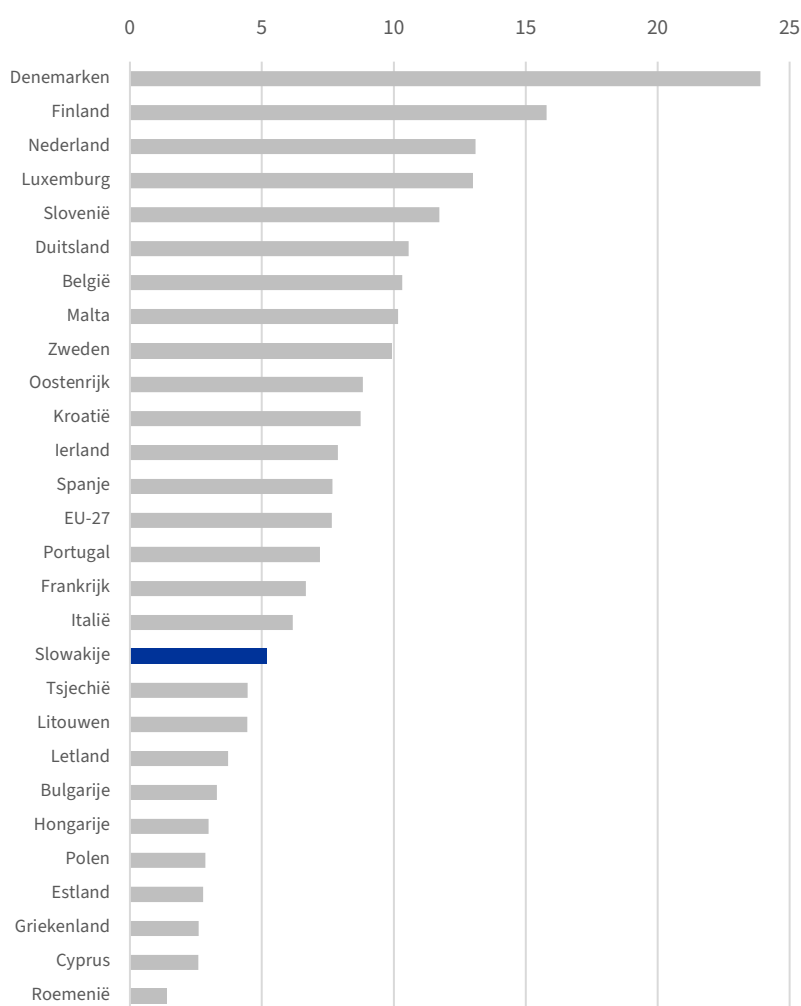
Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

Achtergrondinformatie

Let op: **De onderhandelingen vinden plaats in 2023!**

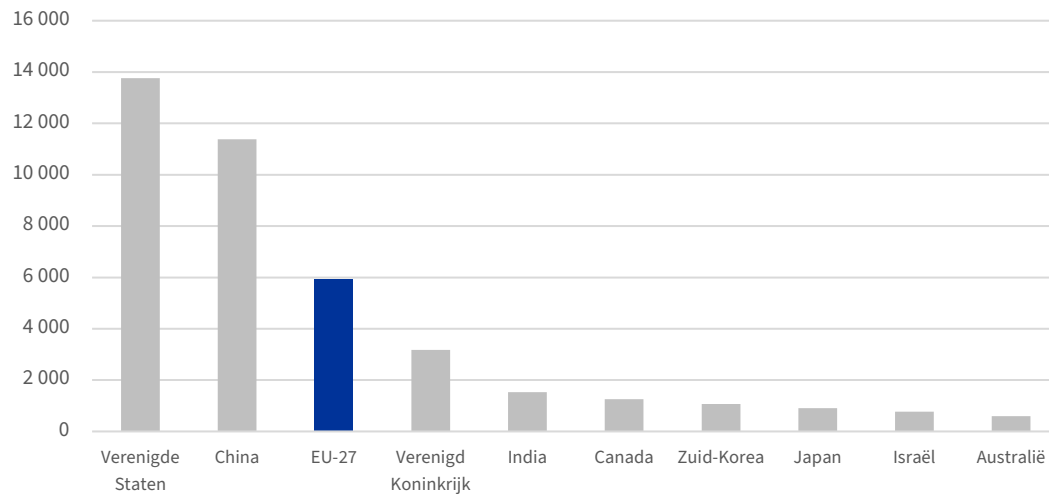
	Europese Unie	Slowakije
Bevolking	447 695 350	5 428 792
Bbp per hoofd van de bevolking in EUR	38 150	22 690
Werkloosheidspercentage	6,1%	5,8%
Percentage bedrijven dat AI-technologie gebruikt	8%	7%
Aandeel van de bevolking met geavanceerde digitale vaardigheden	27,3%	21,7%

Percentage bedrijven met meer dan 10 werknemers
die minstens één AI-systeem gebruiken (2021)

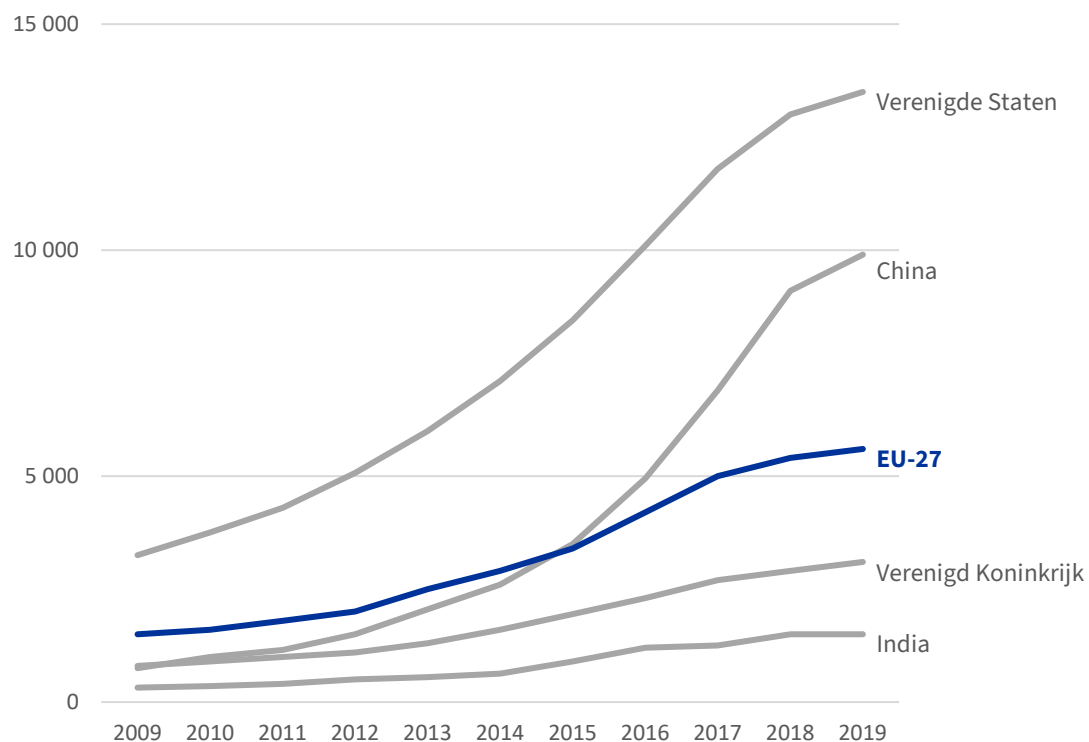


Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.

AI-bedrijven per land (2020)



Groei van het aantal AI-bedrijven



Gegevensbronnen: Eurostat, Gemeenschappelijk Centrum voor onderzoek

Deze publicatie is samengesteld door het secretariaat-generaal van de Raad en is louter informatief. Zij valt niet onder de verantwoordelijkheid van de instellingen van de EU of van de lidstaten. Zie voor meer informatie over de Europese Raad en de Raad de website: www.consilium.europa.eu

© Europese Unie, 2025 | Reproductie met bronvermelding is toegestaan

Print	ISBN 978-92-850-0528-3	doi:10.2860/0531275	QC0125006-NL-C
PDF	ISBN 978-92-850-0527-6	doi:10.2860/4291926	QC0125006-NL-N

Opmerking: Deze educatieve content is bedacht om het spel vlot te laten verlopen en het voor de deelnemers leerzaam te maken. Het gaat niet noodzakelijk om reële content of reële politieke standpunten.