

LIJST VAN VRAGEN

De vaste commissie voor Digitale Zaken heeft een aantal vragen voorgelegd aan de Staatssecretaris van Binnenlandse Zaken en Koninkrijksrelaties over het Rapport van de Algemene Rekenkamer «Focus op AI bij de Rijksoverheid» (Kamerstuk 26 643, nr. 1226).

De voorzitter van de commissie,
Palmen

Adjunct-griffier van de commissie,
Muller

Nr	Vraag
1	Kunt u alle interne protocollen en kaders voor het gebruik van kunstmatige intelligentie (AI-gebruik) die nu gebruikt worden binnen departementen en uitvoerders met de Kamer delen?
2	Hoe vult u concreet uw coördinerende taak in bij het toezien op ethisch en wenselijk gebruik van AI-systemen door alle ministeries en uitvoerders?
3	Kunt u departementen en organisaties dwingen om AI-systemen buiten werking te stellen als deze niet wenselijk zijn?
4	Wie is er binnen departementen en uitvoeringsorganisaties verantwoordelijk voor het toezien op het acceptabele gebruik van AI-systemen?
5	Mede gelet op het feit dat er in dit onderzoek geen AI-systemen zijn ontdekt met onaanvaardbare risico's in de zin van de Europese AI-verordening: zijn er concrete gevallen bekend waarbij het gebruik van AI als ondersteuning direct heeft geleid tot onaanvaardbare verwezenlijking van risico's?
6	In het rapport is niets terug te lezen over de inzet van optische tekenherkenning (OCR). Wat maakt dat OCR ontbreekt of niet in beeld is?
7	Hoe gaat u waarborgen dat alle overheidsdepartementen zich er vóór 2 februari 2025 volledig van bewust zijn welke gebruikte algoritmen onder de AI-verordening vallen, welke daarvan een onaanvaardbaar risico teweegbrengen en aan welke compliancevereisten elk algoritme moet voldoen?
8	Hoe gaat u waarborgen dat de AI-systemen met een nog onbekend risiconiveau juist geclassificeerd zijn vóór 2 februari 2025 en daarnaast tijdig aan de juiste compliancevereisten voldoen?
9	Hoe gaat u waarborgen dat alle onrechtmatige AI-systemen die binnen de overheid gebruikt worden, tijdig rechtmatig zijn dan wel buiten gebruik gesteld worden?
10	Hoe gaat u waarborgen dat alle hoogrisico-AI-systemen tijdig geregistreerd staan in het openbare Algoritmeregister?
11	Hoe gaat u waarborgen dat overheidsdepartementen het risico van AI-systemen niet te laag classificeren?
12	Heeft u, op basis van gesprekken met de relevante toezichthouders, het idee dat het toezicht op de AI-verordening per 2 februari 2025 op orde is? Kunt u dit onderbouwen?
13	Heeft u het idee dat het aan u verleende mandaat op bepaalde gebieden mogelijk tekort om uw coördinerende en kaderstellende taak op AI-gebruik binnen de Rijksoverheid effectief en krachtig uit te voeren?
14	Hoe verklaart u het zeer geringe aantal AI-systemen die zijn geregistreerd in het Algoritmeregister?
15	Kunt u de verwachting dat per eind 2025 alle impactvolle AI-systemen geregistreerd zijn onderbouwen? Is het huidige tempo waarop systemen geregistreerd worden daartoe toereikend?
16	Hoe verklaart u het relatief hoge gebruik van AI-toepassingen bij de Ministeries van Justitie en Veiligheid en Asiel en Migratie?
17	Hoe verklaart u het relatief hoge gebruik van AI-toepassingen bij de Ministeries van Economische Zaken, Klimaat & Groene Groei, en Landbouw, Visserij, Voedselzekerheid en Natuur?
18	Hoe verklaart u het relatief hoge gebruik van AI-toepassingen bij het Ministerie van Infrastructuur en Waterstaat?
19	Hoe verklaart u het relatief hoge gebruik van AI-toepassingen bij het Ministerie van Sociale Zaken en Werkgelegenheid?
20	Wat zegt het geringe aantal geregistreerde AI-systemen over de werkbaarheid van het Algoritmeregister?

- 21 Kunt u alle AI-systemen in gebruik van de politie uiteenzetten, inclusief lopende en beëindigde experimenten? Kunt u hierbij aangeven wat het doel is van deze systemen?
- 22 Kunt u alle AI-systemen in gebruik bij het Uitvoeringsinstituut Werknemersverzekeringen (UWV) uiteenzetten, inclusief lopende en beëindigde experimenten? Kunt u hierbij aangeven wat het doel is van deze systemen?
- 23 Vindt u het gebruik van sommige categorieën AI-systemen meer of minder wenselijk dan andere categorieën? Kunt u aangeven welke van de negen categorieën u het meest en het minst wenselijk vindt?
- 24 Aan welke categorieën AI-systemen kleven de meeste risico's? Gaat u de meer risicovolle categorieën versneld toevoegen aan het Algoritmeregister?
- 25 In het rapport valt op bladzijde 17 bij figuur 8 onder de toepassingscategorie «Kennisverwerking» te lezen «Gesproken tekst omzetten naar geschreven tekst» en onder de toepassingscategorie «Democratisch Proces» staat «Kamerdebatten transcriberen». Transcriptie is het een op een omzetten van gesproken tekst naar geschreven tekst. Vanwaar deze splitsing?
- 26 Bent u bekend met de 82 toepassingen van AI op het gebied van inspectie en handhaving? Op welke termijn worden deze volledig opgenomen in het Algoritmeregister?
- 27 Kunt u uitleggen waarom de Dienst Justitiële Inrichtingen (DJI) een robotohond nodig heeft? Hoe duur was de robotohond?
- 28 Kunt u uitleggen welke AI-systemen worden gebruikt voor opsporing?
- 29 Op welke termijn worden alle AI-systemen met directe impact in het Algoritmeregister opgenomen om de transparantie van de inzet te vergroten?
- 30 Worden burgers actief geïnformeerd over AI-systemen met directe impact die gebruikt worden in processen die hen aangaan? Welke regels zijn er over proactieve transparantie over het gebruik van AI?
- 31 Is er in alle 140 AI-systemen met directe impact sprake van menselijke tussenkomst? Hoe wordt hier toezicht op gehouden?
- 32 Is het juridisch houdbaar om als overheid generatieve AI te gebruiken als deze getraind is op data die verkregen is op illegale wijze, zoals door het schenden van privacy- of auteursrechten via «scraping»?
- 33 Kunt u garanderen dat de overheid geen AI-systemen gebruikt die getraind zijn op illegaal verkregen data?
- 34 Hoe toetst u de kwaliteit van de trainingsdata voor AI-systemen van de overheid? Onder welke voorwaarden acht u de kwaliteit acceptabel?
- 35 Gebruikt de overheid generatieve AI-systemen die getraind zijn op illegaal verkregen data? Met hoeveel zekerheid kunt u dit zeggen?
- 36 Kunt u in brede zin uitleggen waarom de zestien beëindigde experimenten met generatieve AI geen doorgang vonden?
- 37 In hoeveel van de bij u bekende generatieve AI-toepassingen is het doel om beelden of video's te genereren?
- 38 Kunnen burgers er op vertrouwen dat teksten, beelden en video's in openbare uitingen van de overheid altijd door een mens zijn gemaakt en geen auteursrechten schenden?
- 39 Is er een plicht om burgers te informeren als de overheid gegenereerde inhoud gebruikt in openbare uitingen? Ziet u hier noodzaak toe?
- 40 Kunt u garanderen dat alle lopende experimenten en in gebruik genomen generatie AI-toepassingen strikt intern gebruikt worden?

- 41 Is er beleid voor het gebruik van AI in het schrijven of maken van
publieke uitingen door de Rijksoverheid? Zo ja, hoe wordt dit
beleid gehandhaafd?
- 42 Kunt u de meerwaarde uitleggen van AI-systemen waarvan de
gebruikers zelf niet kunnen beoordelen of ze naar behoren
presteren?
- 43 Is het verplicht dat AI-systemen in gebruik van de overheid
aantoonbaar goed functioneren?
- 44 Kunt u het smaldeel in gebruik genomen AI-systemen die slechter
presteren dan verwacht toelichten?
- 45 Is het uw intentie om AI-systemen die slechter functioneren dan
verwacht met meer urgentie op te nemen in het Algoritmeregister?
- 46 Is er een standaard methodiek voor het beoordelen van wenselijk
AI-gebruik door afdelingen Juridische Zaken?
- 47 Is het advies van een afdeling Juridische Zaken over de inzet van AI
bindend?
- 48 Binnen welke kaders moeten experimenten met AI plaatsvinden? Is
dat standaard in een beveiligde omgeving en buiten primaire
processen om?
- 49 Kunt u nader toelichten waarom wet- en regelgeving genoemd
wordt als belemmerende factor in het toepassen van en experi-
menteren met AI?
- 50 Kunt u nader toelichten waarom verantwoording afleggen over de
systemen die de overheid gebruikt wordt gezien als belemmerend?
- 51 In het rapport staat op bladzijde 23: «In een eerdere brief aan de
Tweede Kamer wezen wij al op obstakels die overheidsorganisaties
ervaren door de interpretatie van privacyregels.» Welke stappen
zijn er genomen inzake dit probleem, waarvoor de Autoriteit
Persoonsgegevens (AP) al met enige regelmaat door de Europese
Commissie (EC) op de vingers werd getikt?
- 52 Kunt u uitleggen waarom er ondanks de ervaren verantwoordings-
last vooralsnog slechts vijf procent van AI-systemen in het
Algoritmeregister staat?
- 53 Waarom gebruikt de overheid AI-systemen die onder de
AI-verordening mogelijk verboden worden of aan strengere eisen
moeten voldoen?
- 54 Kunt u inschatten welk aandeel van AI-systemen die nu in gebruik
zijn verboden zullen worden onder de AI-verordening? Op welke
termijn worden deze buiten gebruik gesteld?
- 55 Op welke manier bent u als coördinerend bewindspersoon
betrokken bij het opstellen van intern beleid en richtlijnen voor de
ontwikkeling en toepassing van AI?
- 56 Welke risico's heeft het ontbreken van een specifiek risicomanage-
mentbeleid gericht op AI? Kan bestaande privacy- en cybeveilig-
heidsbeleid een-op-een toepasbaar zijn?
- 57 Is het maken van een risicoafweging een randvoorwaarde voor het
inzetten van AI binnen de overheid?
- 58 Onder welke voorwaarden zijn de risico's bij het gebruiken van een
AI-systeem klein genoeg om deze binnen de overheid te mogen
toepassen?
- 59 Moeten AI-systemen waarvan de risico's onbekend zijn per direct
worden stopgezet totdat er wel een risicoafweging is gemaakt?
Kunt u toelichten waarom wel of niet?
- 60 Wat is uw reactie op het feit dat van 46 procent van de AI-systemen
waarmee wordt geëxperimenteerd of dat wordt gebruikt niet
bekend is welke risico's hiermee gemoeid zijn?
- 61 Hoe kunnen de 81 AI-systemen zonder risicoafweging hier zo snel
mogelijk wel van worden voorzien?

- 62 Hoe kunt u zo snel mogelijk van de 74 AI-systemen waarvan het onbekend is of er een risicoafweging is gemaakt achterhalen of dit wel is gebeurd?
- 63 Waar moet een geldige risicoafweging voor het gebruiken van een AI-systeem aan voldoen?
- 64 Zijn alle instrumenten uit tabel 1 geldige methoden om een risicoafweging te maken voor een AI-systemen? Zo nee, welke zijn wel of niet geldig?
- 65 Hoe waarborgt u de democratische controle op de AI-systemen die de overheid gebruikt als er voornamelijk interne protocollen en kaders worden gebruikt om de impact van een toepassing te wegen?
- 66 Hoe reageert u op de onduidelijkheid die ontstaat door het gebrek aan een rijksbreed instrument om de risico's van AI-systemen mee af te wegen?
- 67 Zodra er een rijksbreed instrument is voor het afwegen van de risico's van AI-systemen, hoe zorgt u er dan voor dat alle reeds in gebruik genomen systemen en lopende experimenten hier aan voldoen?
- 68 Aan welke eisen moeten intern ontwikkelde AI-systemen voldoen?
- 69 Aan welke eisen moeten ingekochte AI-systemen voldoen?
- 70 Moet er te allen tijde duidelijkheid bestaan over de afkomst van AI-systemen die de overheid gebruikt om problematische afhankelijkheden van derden en kwetsbaarheden te voorkomen?
- 71 Zijn er landen waaruit een AI-systeem onder geen enkele voorwaarde mag worden ingekocht? Kunt u garanderen dat er bij de hele overheid geen sprake van is van een dergelijk AI-systeem?
- 72 Hoe kunt u de inschatting van het risiconiveau van AI-systemen vertrouwen als dit gebaseerd is op interne protocollen?
- 73 Is er onafhankelijk toezicht op de risicoclassificatie van AI-systemen?
- 74 Kunt u alle dertig AI-systemen waarvan wordt ingeschat dat deze een hoog risico vormen in kaart brengen?
- 75 Kunt u de vier genoemde voorbeelden van AI-systemen met een hoog risico nader toelichten en nog dit jaar laten registreren in het Algoritmeregister?
- 76 Welke rol speelt u of gaat u spelen om voor februari 2025 te garanderen dat de overheid geen gebruik maakt van AI-systemen met een onaanvaardbaar risico?
- 77 Kunt u voor de vier genoemde voorbeelden van AI-systemen met een hoog risico garanderen dat er sprake is van degelijke menselijke tussenkomst? Op basis van welke kaders beoordeelt u dat?
- 78 Hoe verantwoordt u het gebruik van AI-systemen die achteraf, na het toepassen van een betrouwbare risicoclassificatie, toch onaanvaardbaar of onwenselijk blijken?
- 79 Wanneer is een AI-systeem dat tekst interpreteert en vertaalt voldoende betrouwbaar?
- 80 Wanneer is een AI-systeem dat het risico op schuldenproblematiek inschat voldoende betrouwbaar?
- 81 Hoe garandeert u de kwaliteit van de data waarop AI-systemen van de overheid getraind worden?
- 82 Welke consequenties heeft het voor een departement als deze gebruik maakt van een onaanvaardbaar AI-systeem?