



> Retouradres Postbus 20011 2500 EA Den Haag

Aan de voorzitter van de Eerste Kamer der Staten-Generaal
Postbus 20017
2500 EA Den Haag

Turfmarkt 147
2511 DP Den Haag
Postbus 20011
2500 EA Den Haag

Onze referentie
2025-0000295774

Uw referentie
175936U

Datum 6 mei 2025
Betreft Beantwoording schriftelijke vragen n.a.v. het rapport "Focus
op AI bij de rijksoverheid" van de Algemene Rekenkamer

Bij dezen doe ik u de beantwoording toekomen op de schriftelijke vragen die door uw commissie voor Digitalisering zijn gesteld op 19 februari j.l. n.a.v. het rapport "Focus op AI bij de Rijksoverheid" van de Algemene Rekenkamer (uw kenmerk: 175936U).

De staatssecretaris van Binnenlandse Zaken en Koninkrijksrelaties,
Digitalisering en Koninkrijksrelaties

Zsolt Szabó

Vragen van de leden van de fractie van GroenLinks-PvdA, mede namens de fractie van BBB, D66, SP, PvdD en Volt

Vraag 1: Bent u het met de leden van de fractie van GroenLinks-PvdA eens dat de burgers beter geïnformeerd moeten worden over het gebruik van algoritmen, en dat dit het vertrouwen in zowel het gebruik ervan als in de overheid zal versterken?

Vraag 2: Bent u bereid te onderzoeken hoe de kenbaarheid rond algoritmegebruik door overheden, met name in besluitvormingsprocessen, verbeterd kan worden?

Het kabinet vindt dat burgers moeten kunnen weten hoe een besluit, en dus ook een besluit dat met behulp van algoritmes is voorbereid of genomen, tot stand is gekomen. Transparantie hierover draagt bij aan een betrouwbare overheid. Om deze reden onderzoek ik de noodzaak en mogelijkheden om de bestaande voorschriften en werkwijzen te verbeteren. Dit heeft mijn voorganger in 2023 aan Uw Kamer toegezegd.¹ In 2024 is bij de consultatie over het wetsvoorstel 'Wet versterking waarborgfunctie Awb' ook het reflectiedocument 'Algoritmische besluitvorming en de Awb' in consultatie gebracht. De reacties op dit document zijn verzameld en worden geanalyseerd. De analyse zal binnenkort aan de Tweede Kamer worden gezonden. Ik zal Uw Kamer daarvan een afschrift doen toekomen.

Vraag 3: Hoe luiden de exacte bepalingen in de Franse wetgeving met betrekking tot kenbaarheid rond algoritmegebruik, en hoe beoordeelt u de wijze waarop dit is gewaarborgd?

Vraag 4: Kan het Franse voorbeeld in Nederland worden gevolgd? Zo niet, wat zijn de redenen? Zo ja, hoe zou dit concreet kunnen worden uitgewerkt?

Vraag 5: Indien het Franse voorbeeld niet wordt gevolgd, bent u dan bereid om advies in te winnen bij de Autoriteit Persoonsgegevens over hoe kenbaarheid rond algoritmegebruik verbeterd en wettelijk geborgd kan worden?

Bij de internetconsultatie over het reflectiedocument 'Algoritmische besluitvorming en de Awb' is ook gevraagd om reacties op de Franse wetgeving met betrekking tot algoritmische besluitvorming. De reacties daarop maken onderdeel uit van mijn analyse; ik zal Uw Kamer van mijn brief daarover aan de Tweede Kamer een afschrift zenden. Die brief zal ook beschrijven op welke wijze het kabinet een vervolg wil geven aan de internetconsultatie en de analyse van de reacties.

Vraag 6: Kunt u toelichten welke maatregelen u voornemens bent te nemen om ervoor te zorgen dat burgers het algoritmeregister weten te vinden? Wordt er bijvoorbeeld een publiekscampagne opgezet om burgers

¹ Deze toezegging is gedaan tijdens het plenair debat Grip op algoritmische besluitvorming bij de overheid (CXLVII) – Eerste Kamer der Staten-Generaal 21 maart 2023.

te informeren over het bestaan van het register, of worden overheden aangemoedigd om burgers hier actief op te wijzen?

Met het Algoritmeregister wil ik burgers en bedrijven meer zicht geven op welke algoritmes de overheid gebruikt in haar processen. Momenteel vinden gesprekken plaats, onder meer met medeoverheden, over de beste wijze om burgers te informeren over algoritmes en het Algoritmeregister. Dit heeft raakvlakken met de hiervoor genoemde analyse van de reacties op het reflectiedocument 'Algoritmische besluitvorming en de Awb', aangezien ook in besluiten inzicht kan worden geboden in de gebruikte algoritmes en wellicht kan worden verwezen naar het Algoritmeregister.

Vraag 7: Het huidige algoritmeregister bevat enkele velden die op vrijwillige basis kunnen worden ingevuld. Wordt het register aangepast nu het registreren van hoog-risico-systemen verplicht wordt? Hoe ziet u erop toe dat deze systemen niet alleen worden geregistreerd, maar ook dat alle relevante velden naar behoren worden ingevuld?

Ik stimuleer overheidsorganisaties, zowel binnen de Rijksoverheid als medeoverheden, om de door hen gebruikte algoritmes te publiceren. Daarvoor neem ik proactief contact op met overheden en help ik hen met verschillende instrumenten, zoals de handreiking algoritmeregister, templates van leveranciers, ondersteuning van een implementatieteam, zogenaamde aansluitsessies, een regiotour, een nieuwsbrief en artikelen met tips van organisaties die al gepubliceerd hebben. Verder hebben alle departementen toegezegd om alle hoog-risico AI-systemen die zij gebruiken voor het eind van 2025 te registreren in het Algoritmeregister. Vanaf 2 augustus 2026 geldt er voor alle overheden op basis van de AI-verordening een verplichting om hoog-risico AI-systemen te registreren. Hiervoor hoeft het Algoritmeregister niet te worden aangepast.

Parallel worden er met de medeoverheden gesprekken gevoerd om de kwaliteit van de registratie te verbeteren. Tevens wordt door mijn ministerie de registratie van een algoritme gecontroleerd, voordat het wordt gepubliceerd.

Vraag 8: De leden van de fractie van GroenLinks-PvdA lezen op pagina 4 van het rapport van de Algemene Rekenkamer het volgende:

"Voor meer dan de helft van de opgegeven AI-systemen hebben de organisaties geen risicoafweging gemaakt. Voor een derde van de opgegeven AI-systemen met een hoog risico is geen risicoafweging gemaakt. Daar waar wel een risicoafweging is gemaakt, valt op dat dit niet op uniforme wijze gebeurt. De instrumenten die organisaties gebruiken lopen uiteen van privacytoetsen en AI-toetsingskaders tot zelfontwikkelde risicoanalyses. Er is geen rijksbreed instrument om risico's af te wegen."

Dit roept bij de leden van de fractie van GroenLinks-PvdA de vraag op of een rijksbreed instrument ontwikkeld zou moeten worden. Bent u voornemens om een dergelijk instrument te ontwikkelen? Zo niet, wat is

de reden daarvoor? Zo ja, hoe zal dit worden vormgegeven en welke rol ziet u daarbij voor uzelf?

Door middel van het Algoritmekader bied ik overheidsbreed inzicht in de toetsingskaders en risicoanalyses waaraan voldaan moet worden. Een daarvan is de verplichte risicoafweging die gemaakt moet worden op grond van de artikelen 9 en 27 van de AI-verordening tijdens de ontwikkeling en respectievelijk het gebruik van een hoog-risico AI-systeem. De verplichting daartoe berust op de aanbieder (degene die het systeem in de handel brengt) en de gebruiksverantwoordelijke (in dit geval de overheidsorganisatie). Vanuit mijn coördinerende rol ondersteun ik organisaties hierin. Datgene wat nodig is voor de verantwoorde inzet van algoritmes en AI is enorm in ontwikkeling. In relatief korte tijd komen er veel nieuwe mogelijkheden, maar ook nieuwe wetten, regels, normen en instrumenten. Het Algoritmekader biedt inzicht hoe wordt voldaan aan minimale normen, verantwoorde inzet en welk instrument wanneer dient worden toegepast. Het Impact Assessment Mensenrechten en Algoritmes is een goed instrument voor een risicoanalyse en terug te vinden via het Algoritmekader.

Vraag 9: In een antwoord op de vraag van de vaste commissie voor Digitale Zaken over waarom de overheid AI-systemen gebruikt die onder de AI-verordening worden verboden of aan strengere eisen moeten voldoen (vraag 53), merkt u het volgende op:

"[...] Hoog-risico AI-systemen die voor 2 augustus 2026 op de markt zijn gebracht of in gebruik worden gesteld vallen buiten de toepassing van de AI-verordening, behalve als deze systemen een significante wijziging ondergaan. Voor hoog-risico AI-systemen die bij de overheid voor 2 augustus 2026 in gebruik zijn en geen significante wijziging ondergaan, geldt dat deze systemen uiterlijk 31 december 2030 in lijn met de AI-verordening moeten zijn gebracht."

Op basis van het voorgaande constateren de leden van de fractie van GroenLinks-PvdA dat hoog-risico AI-systemen van de overheid die reeds voor 2 augustus 2026 in gebruik zijn pas vanaf augustus 2030 hoeven te voldoen aan de regels van de AI-verordening. De vraag is of dat wel wenselijk is. De leden van de fractie van GroenLinks-PvdA onderstrepen dat, gezien de risico's voor de gezondheid, veiligheid en fundamentele rechten, het wenselijk is dat hoog-risico-systemen zo spoedig mogelijk voldoen aan de vereisten van de AI-verordening. Deelt u deze mening? Zo ja, bent u voornemens om overheden te stimuleren zo snel mogelijk te voldoen aan de AI-verordening? Zo niet, wat is hiervoor de reden? Kunt u in dit verband reflecteren op de termijn van ruim vijf jaar? Geldt hier niet "the sooner, the better"?

Ik sluit mij erbij aan dat het van belang is dat overheidsorganisaties zich goed voorbereiden op de vereisten die uit de AI-verordening voortvloeien. De termijn die daarbij geldt voor bestaande AI-systemen bij overheden biedt organisaties de komende jaren ruimte om stapsgewijs toe te werken naar de verplichtingen voor hoog-risico AI-systemen uit de AI-verordening. Daarbij moeten deze AI-systemen

nu al wel voldoen aan bestaande wet- en regelgeving zoals de Algemene wet bestuursrecht en de Algemene Verordening Gegevensbescherming.

Voor (nieuwe) AI-systemen die onder de hoog-risico bepalingen vallen en die overheidsorganisaties overwegen in te kopen, te laten ontwikkelen of zelf te ontwikkelen, is het zeer verstandig om nu al goed te kijken naar de vereisten die de AI-verordening stelt en hier bij het inkoop- of ontwikkeltraject rekening mee te houden.

Via verschillende routes stimuleer ik overheidsorganisaties al om zich voor te bereiden op de implementatie van de AI-verordening. Zie hiervoor ook het antwoord onder vraag 7.

Vraag 10: In uw antwoord op vraag 53 stelt u voorts dat het onwaarschijnlijk is dat overheden bewust AI-systemen inzetten die schade toebrengen aan mensen. Kunt u met zekerheid uitsluiten dat dit zowel bewust als onbewust gebeurt? Hebt u daar voldoende zicht op? Zo niet, zou dat niet een prioriteit moeten zijn? Welke stappen onderneemt u om zekerheid te krijgen dat schadelijke AI-praktijken niet plaatsvinden?

Dit valt niet met absolute zekerheid te zeggen, maar op basis van een eerder uitgevoerde inventarisatie door de Rijksoverheid zijn geen AI-systemen aangetroffen die als verboden AI-praktijken zijn aangemerkt. Volgens de AI-verordening zijn dit systemen die een onacceptabel risico vormen voor de gezondheid, veiligheid en fundamentele rechten. Dit betreft dus systemen waarvan de inzet grote schade kan toebrengen aan individuen. Overheidspartijen worden daarnaast gestimuleerd om een algehele inventarisatie uit te voeren van AI-systemen binnen hun organisaties, zoals ook de Vereniging van Nederlandse Gemeenten (VNG) heeft opgenomen in haar impactanalyse van de AI-verordening van januari jl., uitgevoerd in opdracht van BZK.²

Om zekerheid te bieden dat schadelijke AI-praktijken niet plaatsvinden, voorziet de AI-verordening in een stevig toezichtstelsel. Dit biedt de nog aan te wijzen (markt)toezichthouder de bevoegdheid om sancties op te leggen wanneer organisaties gebruik (blijven) maken van verboden AI-praktijken. Vanaf 2 februari 2025 kunnen overtredingen van de verboden leiden tot civielrechtelijke aansprakelijkheid. Ook voor andere risicocategorieën, zoals hoog-risico-AI, zullen (markt)toezichthouders worden aangewezen.

Vraag 11: De leden van de fractie van GroenLinks-PvdA constateren dat het Kaderverdrag van de Raad van Europa inzake kunstmatige intelligentie, mensenrechten, democratie en de rechtsstaat op 28 augustus 2024 door de Raad namens de Europese Unie is goedgekeurd. Het Kaderverdrag treedt in werking op de eerste dag van de maand die volgt op het verstrijken van een periode van drie maanden na de datum waarop vijf ondertekenaars, waaronder ten minste drie lidstaten van de Raad van Europa, het hebben geratificeerd. De leden van de fractie van

² https://vng.nl/sites/default/files/2025-01/uitvoeringsanalyse_digital_decade_ai-verordening.pdf

GroenLinks-PvdA vernemen graag of de regering voornemens is dit verdrag te ratificeren en daarmee de inwerkingtreding te bespoedigen. De leden zijn van mening dat Nederland door ratificatie een actievere rol kan spelen in de verdere ontwikkeling en toepassing van de regels in zowel de Raad van Europa als de EU. Bovendien laat ratificatie zien dat Nederland zich inzet voor ethische AI, mensenrechten en democratische principes. Deelt u deze mening? Zo ja, wat gaat u doen om deze ratificatie te bespoedigen? Zo niet, waarom bent u daartoe niet bereid?

In augustus 2024 heeft de Raad van de Europese Unie een besluit aangenomen dat de Europese Commissie een mandaat geeft om namens de Europese Unie het kaderverdrag voor AI en de bescherming van mensenrechten, democratie en rechtsstaat te ondertekenen. Op 5 september 2024 heeft de EU het verdrag ondertekend. Dit betekent dat het kaderverdrag ook namens Nederland is ondertekend. Ook landen zoals Andorra, het Verenigd Koninkrijk en de Verenigde Staten hebben het verdrag ondertekend. Het kabinet steunt een voorspoedige ratificatie. Op dit moment is de Europese Commissie bezig met het proces voor ratificatie van het verdrag namens de gehele Europese Unie, inclusief Nederland.

Vragen van de leden van de VVD-fractie, mede namens de fractie van BBB, D66, SP, PvdD en Volt

Vraag 12: Hoe gaat u ervoor zorgen dat de standaardprocedures voor het testen van AI-systemen voorafgaand aan de toepassing ervan net zo transparant zijn als de veiligheidsnormen voor crashtests van nieuwe voertuigen, zodat zowel de effectiviteit van deze AI-systemen als de bescherming van grondrechten worden gewaarborgd?

In de AI-verordening staan vereisten die alle ontwikkelaars van risicovolle AI moeten volgen om bescherming van grondrechten te borgen en risico's tegen te gaan. Zo is een kwaliteitsmanagementsysteem verplicht, waarmee ook de effectiviteit kan worden geborgd. Ook moet er een risicomanagementsysteem worden ingericht, waarmee alle mogelijke risico's in kaart moeten worden gebracht en waar mogelijk gemitigeerd worden.

Voor deze eisen komen Europese normen, die AI-ontwikkelaars handvatten geven om de eisen na te leven. Deze normen worden niet verplicht, maar de verwachting is dat ze grootschalig gehanteerd zullen worden en daarmee procedures voor zover mogelijk worden gestandaardiseerd. Wie dat niet doet, moet namelijk zelf aantonen aan alle vereisten voor die systemen te voldoen.

Voor AI die niet in één van de hoog-risico categorieën valt, zijn deze vereisten er niet vanuit de AI-verordening. Lidstaten mogen dit ook niet in aanvulling op de verordening vereisen, omdat dit de interne markt kan verstoren en tot hogere administratieve lasten leidt. Wel mogen ontwikkelaars van dergelijke AI uiteraard op vrijwillige basis gebruik maken van dergelijke systemen.

Vraag 13: Wat is uw visie op het ontwikkelen van een keurmerk voor AI-systemen, vergelijkbaar met het Euro NCAP-keurmerk voor auto's?

Euro NCAP is een keurmerk voor één specifiek aspect van auto's: de veiligheid in geval van een botsing. Daar kan een standaard set van eisen voor opgesteld worden. Die zijn ook in alle gevallen van botsingen van toepassing.

Omdat er niet een standaard set van eisen is en er geen standaard testprocedures voor AI mogelijk zijn, is het onmogelijk om een keurmerk te ontwikkelen om aan te geven in welke mate aan die eisen wordt voldaan.

Voor hoog-risico AI gelden echter wel eisen, zie ook het antwoord op de vorige vraag. Systemen die aan die eisen voldoen, mogen op de markt gebracht worden nadat een beoordeling is gedaan, aan de eisen uit de verordening is voldaan en een CE-markering is aangebracht.

Vraag 14: Welke maatregelen neemt u om ervoor te zorgen dat burgers en maatschappelijke organisaties actief betrokken worden bij het ontwikkelen en evalueren van testprocedures voor AI-systemen, zodat hun perspectieven en ervaringen worden meegenomen in de beoordeling van zowel de effectiviteit van deze systemen als de bescherming van grondrechten?

In EU-verband worden standaarden ontwikkeld voor diverse aspecten van de ontwikkeling van AI-systemen. Daar kan eenieder aan bijdragen. De EU heeft daarvoor ook een adviesforum en een wetenschappelijk panel in het leven geroepen. Dat adviesforum vertegenwoordigt een evenwichtige selectie van belanghebbenden, waaronder het bedrijfsleven, start-ups, het maatschappelijk middenveld en de academische wereld.

De minister van Economische Zaken geeft subsidie aan het NEN (Nederlandse Normalisatie Instituut) dat ervoor zorgt dat belanghebbenden uit Nederland ook betrokken kunnen zijn bij de processen om tot deze standaarden te komen.

Het is daarnaast aan de AI-ontwikkelaars om burgers, maatschappelijke organisaties of anderen te betrekken bij het ontwikkelen of evalueren van AI-systemen of de wijze waarop die getest worden.

Vraag 15: In hoeverre acht u het waarschijnlijk dat regulering van AI, bedoeld om de risico's daarvan te beperken, de innovatieve ontwikkeling ervan belemmert? Op welke wijze bent u voornemens om de juiste balans te bevorderen tussen innovatie enerzijds en de bescherming van grondrechten anderzijds?

De eisen die op grond van de AI-verordening aan AI gesteld worden, vergen het nodige van de ontwikkelaar van die AI. Tegelijkertijd zorgen die eisen ervoor dat ontwikkelaars weten binnen welk kader AI kan worden ontwikkeld die veilig gebruikt kan worden en grondrechten respecteert. Dit zorgt voor het stimuleren van innovatie van deze verantwoorde AI-systemen in de EU. Om met name het MKB te ondersteunen bij het ontwikkelen van AI binnen de nieuwe wettelijke kaders, zijn diverse instrumenten via de AI-verordening ingericht. Een belangrijke is een regulatory sandbox, waarin door marktpartijen in samenwerking met toezichthouders gekeken kan worden hoe aan de vereisten van de verordening

kan worden voldaan. Naast deze regulatory sandbox wordt er voorlichting gegeven. Vanuit de EU worden bijvoorbeeld richtsnoeren ontwikkeld om richting te geven aan diverse onderdelen van de AI-verordening, zoals over verboden AI-praktijken en wat wel en niet onder de definitie van AI valt.