



Raad van de
Europese Unie

Brussel, 23 april 2021
(OR. en)

**Interinstitutioneel dossier:
2021/0106(COD)**

8115/21
ADD 4

TELECOM 156
JAI 429
COPEN 191
CYBER 108
DATAPROTECT 103
EJUSTICE 41
COSI 69
IXIM 74
ENFOPOL 148
FREMP 103
RELEX 347
MI 271
COMPET 275
IA 60
CODEC 573

BEGELEIDENDE NOTA

van:	de secretaris-generaal van de Europese Commissie, ondertekend door mevrouw Martine DEPREZ, directeur
ingekomen:	22 april 2021
aan:	de heer Jeppe TRANHOLM-MIKKELSEN, secretaris-generaal van de Raad van de Europese Unie
nr. Comdoc.:	SWD(2021) 85 final
Betreft:	WERKDOCUMENT VAN DE DIENSTEN VAN DE COMMISSIE SAMENVATTING VAN HET EFFECTBEOORDELINGSVERSLAG bij het voorstel voor een verordening van het Europees Parlement en de Raad TOT VASTSTELLING VAN GEHARMONISEERDE REGELS BETREFFENDE ARTIFICIËLE INTELLIGENTIE (WET OP DE ARTIFICIËLE INTELLIGENTIE) EN TOT WIJZIGING VAN BEPAALDE WETGEVINGSHANDELINGEN VAN DE UNIE

Hierbij gaat voor de delegaties document SWD(2021) 85 final.

Bijlage: SWD(2021) 85 final

Brussel, 21.4.2021
SWD(2021) 85 final

WERKDOCUMENT VAN DE DIENSTEN VAN DE COMMISSIE
SAMENVATTING VAN HET EFFECTBEOORDELINGSVERSLAG

bij het

Voorstel voor een verordening van het Europees Parlement en de Raad

**TOT VASTSTELLING VAN GEHARMONISEERDE REGELS BETREFFENDE
ARTIFICIËLE INTELLIGENTIE (WET OP DE ARTIFICIËLE INTELLIGENTIE)
EN TOT WIJZIGING VAN BEPAALDE WETGEVINGSHANDELINGEN VAN DE
UNIE**

{COM(2021) 206 final} - {SEC(2021) 167 final} - {SWD(2021) 84 final}

Samenvatting
Effectbeoordeling van een regelgevingskader voor artificiële intelligentie
A. De noodzaak om actie te ondernemen
Wat is het probleem en waarom is het een probleem op EU-niveau?
Artificiële intelligentie (AI) is een opkomende universeel toepasbare technologie en in die hoedanigheid een zeer krachtige reeks computerprogrammeertechnieken. De invoering van AI-systemen zou in sterke mate kunnen leiden tot maatschappelijke voordelen en economische groei en tevens innovatie in de EU en het wereldwijde concurrentievermogen van de EU kunnen verbeteren. In bepaalde gevallen kan het gebruik van AI-systemen evenwel aanleiding geven tot problemen. De specifieke kenmerken van bepaalde AI-systemen kunnen leiden tot nieuwe risico's die verband houden met 1) veiligheid en beveiliging en 2) grondrechten, en de waarschijnlijkheid of intensiteit van bestaande risico's versnellen. AI-systemen 3) maken het ook moeilijk voor handhavingsinstanties om de naleving van bestaande regels te verifiëren en deze regels te handhaven. Deze reeks problemen leidt op zijn beurt tot 4) rechtsonzekerheid voor ondernemingen, 5) een potentieel minder snelle brede inzet van AI-technologieën als gevolg van het gebrek aan vertrouwen bij bedrijven en burgers, evenals tot 6) reacties van de nationale autoriteiten in de vorm van regelgeving om mogelijke externe effecten te beperken die de interne markt dreigen te versnipperen.
Wat is het doel?
Het regelgevingskader is erop gericht die problemen aan te pakken om de correcte werking van de eengemaakte markt te waarborgen door omstandigheden te creëren voor de ontwikkeling en het gebruik van betrouwbare AI in de Unie. Specifieke doelstellingen zijn: 1) waarborgen dat AI-systemen die in de handel worden gebracht en worden gebruikt, veilig zijn en de bestaande wetgeving op het gebied van de grondrechten en de waarden van de Unie eerbiedigen; 2) rechtszekerheid waarborgen teneinde investeringen in en innovatie op het gebied van AI te faciliteren; 3) de governance en doeltreffende handhaving verbeteren van de bestaande wetgeving op het gebied van de grondrechten en veiligheidseisen die op AI-systemen van toepassing zijn; en 4) de ontwikkeling mogelijk maken van een eengemaakte markt voor rechtmatige, veilige en betrouwbare AI-systemen en marktfragmentatie voorkomen.
Wat is de meerwaarde van EU-maatregelen (subsidiariteit)?
De grensoverschrijdende aard van grootschalige data en datasets die vaak nodig zijn voor AI-toepassingen houdt in dat de doelstellingen van het initiatief niet op doeltreffende wijze kunnen worden bereikt door de lidstaten alleen. Het Europees regelgevingskader voor betrouwbare AI is erop gericht geharmoniseerde regels vast te stellen voor de ontwikkeling, het in de handel brengen en het gebruik van producten en diensten waarin AI-technologie is verwerkt of autonome AI-toepassingen in de Unie. Het kader heeft tot doel een eerlijk speelveld te waarborgen en alle Europese burgers te beschermen, terwijl tegelijkertijd het concurrentievermogen van Europa en de industriële basis van Europa op het gebied van AI worden versterkt. EU-maatregelen voor AI stimuleren de interne markt en bieden aanzienlijk potentieel om de Europese industrie een concurrentievoordeel te geven op wereldniveau, gebaseerd op schaalvoordelen die niet kunnen worden bereikt door de afzonderlijke lidstaten alleen.
B. Oplossingen
Welke opties dienen zich aan? Is er al dan niet een voorkeursoptie? Zo niet, waarom?
De volgende opties zijn in aanmerking genomen: Optie 1: EU-wetgevingsinstrument tot instelling van een vrijwillige etiketteringsregeling; Optie 2: een sectorale ad-hoc-aanpak; Optie 3: een horizontaal EU-wetgevingsinstrument tot vaststelling van verplichte eisen aan AI-toepassingen met een hoog risico ; Optie 3+: hetzelfde als optie 3, maar met vrijwillige gedragscodes voor AI-toepassingen die geen hoog risico vormen; en Optie 4: een horizontaal EU-wetgevingsinstrument tot vaststelling van verplichte eisen aan alle AI-toepassingen. Optie 3+ geniet de voorkeur, aangezien het evenredige waarborgen biedt tegen de risico's die AI vormt, terwijl de administratieve en nalevingskosten tot een minimum worden beperkt. De specifieke kwestie met betrekking tot de aansprakelijkheid voor AI-toepassingen wordt aangepakt met

afzonderlijke toekomstige regels en valt dus niet onder de opties.
Hoe reageren de verschillende belanghebbenden? Wie steunt welke optie?
Bedrijven, overheidsinstanties, academici en niet-gouvernementele organisaties zijn het er alle over eens dat er hiaten in de wetgeving bestaan of dat er nieuwe wetgeving nodig is, al is de meerderheid onder bedrijven kleiner. De branche en overheidsinstanties stemmen ermee in de verplichte eisen te beperken tot AI-toepassingen met een hoog risico. Burgers en het maatschappelijk middenveld zijn het eerder oneens met de beperking van de verplichte eisen tot toepassingen met een hoog risico.
C. Effecten van de voorkeursoptie
Wat zijn de voordelen van de voorkeursoptie (indien er een voorkeur is, anders van de belangrijkste opties)?
Voor burgers beperkt de voorkeursoptie de risico's die bestaan voor hun veiligheid en de grondrechten. Voor aanbieders van AI zorgt deze optie voor rechtszekerheid, en waarborgt zij dat er geen belemmeringen ontstaan voor het grensoverschrijdend aanbieden van AI-gerelateerde producten en diensten. Voor ondernemingen die AI gebruiken, zal de optie het vertrouwen onder hun klanten bevorderen. Voor nationale overheden bevordert de optie het vertrouwen van het publiek in het gebruik van AI en worden handhavingsmechanismen versterkt (door de invoering van een Europees coördinatiemechanisme, dat voorziet in passende mogelijkheden, en door audits van de AI-systemen te vergemakkelijken met behulp van nieuwe eisen voor documentatie, traceerbaarheid en transparantie).
Wat zijn de kosten van de voorkeursoptie (indien er een voorkeur is, anders van de belangrijkste opties)?
Bedrijven of overheidsinstanties die AI-toepassingen ontwikkelen of gebruiken die een hoog risico vormen voor de veiligheid of de grondrechten van burgers moeten voldoen aan specifieke horizontale eisen en verplichtingen, die worden ingevoerd met behulp van technische geharmoniseerde normen. De totale samenge telde nalevingskosten worden geraamd op tussen 100 miljoen EUR en 500 miljoen EUR tegen 2025. Dit komt overeen met tot 4–5 % van de investeringen in AI met een hoog risico (naar verwachting tussen 5 % en 15 % van alle AI-toepassingen). De verificatiekosten kunnen neerkomen op nog eens 2–5 % van de investering in AI met een hoog risico. Bedrijven of overheidsinstanties die AI-toepassingen ontwikkelen of gebruiken die niet geclassificeerd zijn als een toepassing met een hoog risico, hoeven geen kosten te maken. Zij kunnen er evenwel voor kiezen zich te houden aan vrijwillige gedragscodes om te voldoen aan passende eisen en te waarborgen dat hun AI betrouwbaar is. In deze gevallen kunnen de kosten maximaal net zo hoog zijn als voor toepassingen met een hoog risico, maar waarschijnlijk zijn deze lager.
Wat zijn de gevolgen voor kleine en middelgrote ondernemingen en voor het concurrentievermogen?
Kmo's hebben meer baat bij een hoger algemeen vertrouwen in AI dan grotere ondernemingen die ook kunnen vertrouwen op hun merkimage. Kmo's die toepassingen ontwikkelen die worden geclassificeerd als een toepassing met een hoog risico, moeten soortgelijke kosten dragen als grote ondernemingen. Als gevolg van de schaalbaarheid van digitale technologieën, kunnen kleine en middelgrote ondernemingen zelfs een enorm bereik behalen, ondanks hun kleine omvang, en zo potentieel miljoenen personen bereiken. Met betrekking tot toepassingen met een hoog risico zou de uitsluiting van kmo's die AI leveren van de toepassing van het regelgevingskader dus juist de doelstelling om het vertrouwen te vergroten, kunnen ondermijnen. In het kader wordt echter in specifieke maatregelen voorzien, waaronder testomgevingen voor regelgeving of bijstand via de digitale innovatiehubs, ter ondersteuning van kmo's bij de naleving van de nieuwe regels, waarbij rekening is gehouden met hun speciale behoeften.
Zijn er significante gevolgen voor de nationale begrotingen en overheden?
De lidstaten moeten toezichthoudende autoriteiten aanwijzen die belast zijn met de uitvoering van de wettelijke vereisten. Hun toezichthoudende functie kan worden gebaseerd op bestaande regelingen, bijvoorbeeld als het gaat om conformiteitsbeoordelingsinstanties of markttoezicht, maar zal wel voldoende technologische deskundigheid en middelen vereisen. Afhankelijk van de reeds bestaande structuur in elke

lidstaat kan dit kunnen neerkomen op 1 tot 25 voltijdequivalenten per lidstaat.
Zijn er andere significante gevolgen?
De voorkeursoptie zou de risico's voor de grondrechten van burgers, evenals voor de bredere waarden van de Unie, aanzienlijk beperken en tevens de veiligheid van bepaalde producten en diensten die AI-technologie bevatten of autonome AI-toepassingen verbeteren.
Evenredigheid
Het voorstel is evenredig en noodzakelijk om de doelstellingen te bereiken, aangezien het een risicogebaseerde aanpak volgt en alleen regeldruk oplegt wanneer de AI-systemen waarschijnlijk een hoog risico opleveren voor de grondrechten of de veiligheid. Wanneer dit niet het geval is, worden uitsluitend minimale transparantie-eisen opgelegd, met name voor wat betreft informatievoorziening om het gebruik van een AI-systeem duidelijk te maken wanneer sprake is van interactie met mensen of het gebruik van deep fakes indien deze niet voor rechtmatige doeleinden worden gebruikt. Dankzij geharmoniseerde normen en ondersteunende richtsnoeren en nalevingsinstrumenten zullen aanbieders en gebruikers beter in staat zijn de in het voorstel vastgestelde voorschriften na te leven en hun kosten tot een minimum te beperken.
D. Evaluatie
Wanneer wordt het beleid geëvalueerd?
De Commissie zal vijf jaar na de datum waarop het kader van toepassing wordt, een evaluatie- en herzieningsverslag publiceren.