

Bruxelles, den 23. april 2021
(OR. en)

**Interinstitutionel sag:
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**

FØLGESKRIVELSE

fra:	Martine DEPREZ, direktør, på vegne af generalsekretæren for Europa-Kommissionen
modtaget:	22. april 2021
til:	Jeppe TRANHOLM-MIKKELSEN, generalsekretær for Rådet for Den Europæiske Union
Komm. dok. nr.:	SWD(2021) 85 final
Vedr.:	ARBEJDSDOKUMENT FRA KOMMISSIONENS TJENESTEGRENE RESUMÉ AF RAPPORTEN OM KONSEKVENSANALYSEN Ledsagedokument til Europa-Parlamentets og Rådets Forordning OM HARMONISEREDE REGLER FOR KUNSTIG INTELLIGENS (RETSAKTEN OM KUNSTIG INTELLIGENS) OG OM ÆNDRING AF VISSE AF UNIONENS LOVGIVNINGSMÆSSIGE RETSAKTER

Hermed følger til delegationerne dokument SWD(2021) 85 final.

Bilag: SWD(2021) 85 final



EUROPA-
KOMMISSIONEN

Bruxelles, den 21.4.2021
SWD(2021) 85 final

ARBEJDSDOKUMENT FRA KOMMISSIONENS TJENESTEGRENE

RESUMÉ AF RAPPORTEN OM KONSEKVENSANALYSEN

Ledsagedokument til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

**OM HARMONISEREDE REGLER FOR KUNSTIG INTELLIGENS (RETSAKTEN
OM KUNSTIG INTELLIGENS) OG OM ÆNDRING AF VISSE AF UNIONENS
LOVGIVNINGSMÆSSIGE RETSAKTER**

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

Resumé
Konsekvensanalyse af en lovgivningsmæssig ramme for kunstig intelligens
A. Behov for handling
Hvad er problemstillingen, og hvorfor er det en problem på EU-niveau?
Kunstig intelligens (AI) er en ny teknologi til generelle formål og består af en meget stærk vifte af computerprogrammeringsteknikker. Udbredelsen af AI-systemer har et stort potentiale til at skabe samfundsmæssige fordele, økonomisk vækst og fremme EU's innovation og konkurrenceevne på verdensplan. I visse tilfælde kan anvendelsen af AI-systemer imidlertid give problemer. Visse AI-systemers særlige egenskaber kan skabe nye risici med hensyn til 1) sikkerhed og 2) grundlæggende rettigheder samt forstærke sandsynligheden for eller alvorligheden af eksisterende risici. AI-systemer 3) gør det også vanskeligt for håndhævende myndigheder at kontrollere, at de eksisterende regler overholdes, og håndhæve dem. Disse problemer fører igen til 4) retlig usikkerhed for virksomheder, 5) en potentielt langsommere udbredelse af AI-teknologier på grund af manglende tillid blandt virksomheder og borgere samt 6) lovgivningsmæssige tiltag fra de nationale myndigheders side for at afbøde eventuelle eksternaliteter, der risikerer at opsplitte det indre marked.
Hvad bør opnås?
Formålet med den lovgivningsmæssige ramme er at afhjælpe disse problemer for at sikre et velfungerende indre marked ved at skabe gunstige betingelser for udvikling og anvendelse af pålidelig kunstig intelligens i Unionen. De specifikke mål er at: 1) sørge for, at AI-systemer, der bringes i omsætning og anvendes, er sikre og i overensstemmelse med den eksisterende lovgivning om grundlæggende rettigheder og med Unionens værdier, 2) sikre retssikkerhed med henblik på at fremme investering og innovation inden for kunstig intelligens, 3) forbedre forvaltningen og den effektive håndhævelse af den eksisterende lovgivning om grundlæggende rettigheder og sikkerhedskrav til AI-systemer og 4) fremme udviklingen af et indre marked for lovlige, sikre og pålidelige AI-systemer og forhindre markedsfragmentering.
Hvad er merværdien ved at handle på EU-plan (nærhedsprincippet)?
Den grænseoverskridende karakter af de omfattende data og datasæt, som anvendelse af kunstig intelligens ofte er afhængig af, betyder, at initiativets mål ikke kan opfyldes effektivt af medlemsstaterne alene. Den europæiske lovgivningsmæssige ramme for pålidelig kunstig intelligens sigter mod fastsættelse af harmoniserede regler for udvikling, omsætning og anvendelse af produkter og tjenester, hvori AI-teknologi er integreret, eller selvstændig anvendelse af kunstig intelligens i Unionen. Formålet med den lovgivningsmæssige ramme er at sikre lige vilkår og beskytte alle EU-borgere, samtidig med EU's konkurrenceevne og industrielle grundlag inden for kunstig intelligens styrkes. Handling på EU-plan vil sætte skub i det indre marked og har et betydeligt potentiale til at give europæisk industri en konkurrencefordel på verdensplan og stordriftsfordele, som ikke kan opnås af de enkelte medlemsstater alene.
B. Løsninger
Hvilke forskellige løsningsmodeller er der for at nå målene? Foretrækkes en bestemt løsningsmodel frem for andre? Hvis ikke, hvorfor?
Følgende løsningsmodeller er overvejet: model 1 : en EU-retsakt, hvorved der indføres af en frivillig mærkningsordning; model 2 : en sektorspecifik ad hoc-tilgang; model 3 : en horisontal EU-retsakt, hvorved der fastsættes obligatoriske krav til AI-anvendelser, der udgør en høj risiko ; model 3+ : som model 3, men med frivillige adfærdskodekser for AI-anvendelser, der ikke udgør en høj risiko, og model 4 : en horisontal EU-retsakt, hvorved der fastsættes obligatoriske krav til alle AI-anvendelser. Den foretrukne løsning er model 3+, da den giver forholdsmæssige garantier mod de risici, der er forbundet med kunstig intelligens, samtidig med at de administrative omkostninger og efterlevelseseomkostningerne begrænses til et minimum. Spørgsmålet om ansvar i forbindelse med anvendelse af kunstig intelligens vil blive behandlet i særskilte fremtidige regler og er dermed ikke omfattet af løsningsmodellerne.
Hvad er de forskellige interessenters holdning? Hvem støtter hvilken løsning?

Virksomheder, offentlige myndigheder, forskere og NGO'er er alle enige i, at der er huller i lovgivningen, eller at der er behov for ny lovgivning, selv om flertallet af virksomheder, der mener sådan, er mindre. Industrien og offentlige myndigheder er enige i at begrænse de obligatoriske krav til AI-anvendelser, der udgør en høj risiko. Borgere og civilsamfundet er mere tilbøjelige til ikke at være enige i at begrænse obligatoriske krav til anvendelser, der udgør en høj risiko.
C. Den foretrukne løsnings virkninger
Hvilke fordele er der ved den foretrukne løsning (hvis en bestemt løsning foretrækkes — ellers fordelene ved de vigtigste af de mulige løsninger)?
For borgerne vil den foretrukne løsningsmodel mindske risiciene for deres sikkerhed og grundlæggende rettigheder. For udbydere af kunstig intelligens vil modellen skabe retssikkerhed og sikre, at der ikke opstår hindringer for grænseoverskridende levering af AI-relaterede tjenester og produkter. For de virksomheder, der anvender kunstig intelligens, vil modellen fremme tilliden blandt deres kunder. For nationale offentlige forvaltninger vil modellen fremme offentlighedens tillid til brugen af kunstig intelligens og styrke håndhævelsesmekanismen (idet der indføres en europæisk koordineringsmekanisme, tilvejebringes passende kapacitet, og revisionen af AI-systemer lettes med nye krav til dokumentation, sporbarhed og gennemsigtighed).
Hvilke omkostninger er der ved den foretrukne løsning (hvis en bestemt løsning foretrækkes — ellers omkostningerne ved de vigtigste af de mulige løsninger)?
De virksomheder eller offentlige myndigheder, der er udviklere eller brugere af AI-anvendelser, der udgør en høj risiko for borgernes sikkerhed eller grundlæggende rettigheder, vil skulle opfylde specifikke krav og forpligtelser, som vil blive indført ved hjælp af tekniske harmoniserede standarder. De samlede efterlevelseseomkostninger anslås til at være mellem 100 mio. EUR og 500 mio. EUR frem til 2025, hvilket svarer til op til 4-5 % af investeringerne i højrisiko-AI (som anslås til udgøre mellem 5 % og 15 % af alle AI-anvendelser). Kontrolomkostningerne kan beløbe sig til yderligere 2-5 % af investeringerne i kunstig intelligens med høj risiko. De virksomheder eller offentlige myndigheder, der er udviklere eller brugere af AI-anvendelser, der ikke er klassificeret som højrisiko, vil ikke skulle afholde nogen omkostninger. De kan imidlertid vælge at overholde frivillige adfærdscodekser for at følge passende krav og for at sikre, at deres AI-systemer er pålidelige. I disse tilfælde kan omkostningerne være lige så høje som for højrisikoanvendelser, men vil sandsynligvis være lavere.
Hvad er indvirkningerne på SMV'er og på konkurrencedygtigheden?
SMV'er vil drage fordel af et højere generelt niveau af tillid til kunstig intelligens end store virksomheder, som også kan nyde godt af deres image. De SMV'er, der udvikler anvendelser, der er klassificeret som højrisiko, vil skulle afholde omkostninger, der er lig dem, store virksomheder vil skulle afholde. På grund af de digitale teknologiers høje skalerbarhed kan SMV'er faktisk have en enorm rækkevidde på trods af deres størrelse og potentielt påvirke millioner af mennesker. Med hensyn til højrisikoanvendelser kan det i alvorlig grad undergrave målet om øget tillid, hvis SMV'er, der leverer kunstig intelligens, ikke er omfattet af den lovgivningsmæssige ramme. Rammen vil imidlertid indbefatte specifikke foranstaltninger, herunder reguleringsmæssige sandkasser eller bistand gennem de digitale innovationsknudepunkter, med henblik på at støtte SMV'er i overholdelsen af de nye regler under hensyntagen til deres særlige behov.
Vil den foretrukne løsning få væsentlige virkninger for de nationale budgetter og myndigheder?
Medlemsstaterne vil skulle udpege de tilsynsmyndigheder, der skal have ansvaret for at gennemføre de lovgivningsmæssige krav. Tilsynsfunktionen kan bygge på eksisterende ordninger, f.eks. angående overensstemmelsesvurderingsorganer eller markedsovervågning, men vil kræve tilstrækkelig teknologisk ekspertise og ressourcer. Afhængigt af den eksisterende struktur i hver medlemsstat kan dette beløbe sig til 1-25 fuldtidsækvivalenter pr. medlemsstat.
Vil den foretrukne løsning få andre væsentlige virkninger?
Den foretrukne løsningsmodel vil i væsentlig grad mindske risiciene for borgernes grundlæggende rettigheder samt for Unionens værdier i bredere forstand og vil øge sikkerheden i forbindelse med visse produkter og tjenester, hvori AI-teknologi er integreret, eller selvstændig anvendelse af kunstig intelligens

i Unionen.
Proportionalitetsprincippet?
Forslaget er rimeligt og nødvendigt for at nå målene, da det følger en risikobaseret tilgang og kun pålægger reguleringsmæssige byrder, hvis et AI-system sandsynligvis vil udgøre en høj risiko for grundlæggende rettigheder eller sikkerheden. Hvis det ikke er tilfældet, indføres der kun minimale gennemsigtighedsforpligtelser, navnlig med hensyn til formidling af oplysninger for at markere anvendelse af et AI-system i interaktion med mennesker eller brug af deepfake-indhold, hvis det ikke anvendes til lovlige formål. Harmoniserede standarder og understøttende vejlednings- og overholdelsesværktøjer vil sigte mod at hjælpe udbydere og brugere med at overholde kravene og minimere deres omkostninger.
D. Opfølgning
Hvornår vil foranstaltningen blive taget op til fornyet overvejelse?
Kommissionen vil offentliggøre en rapport med en evaluering og en gennemgang af den foreslåede ramme fem år efter den dato, hvorfra den finder anvendelse.