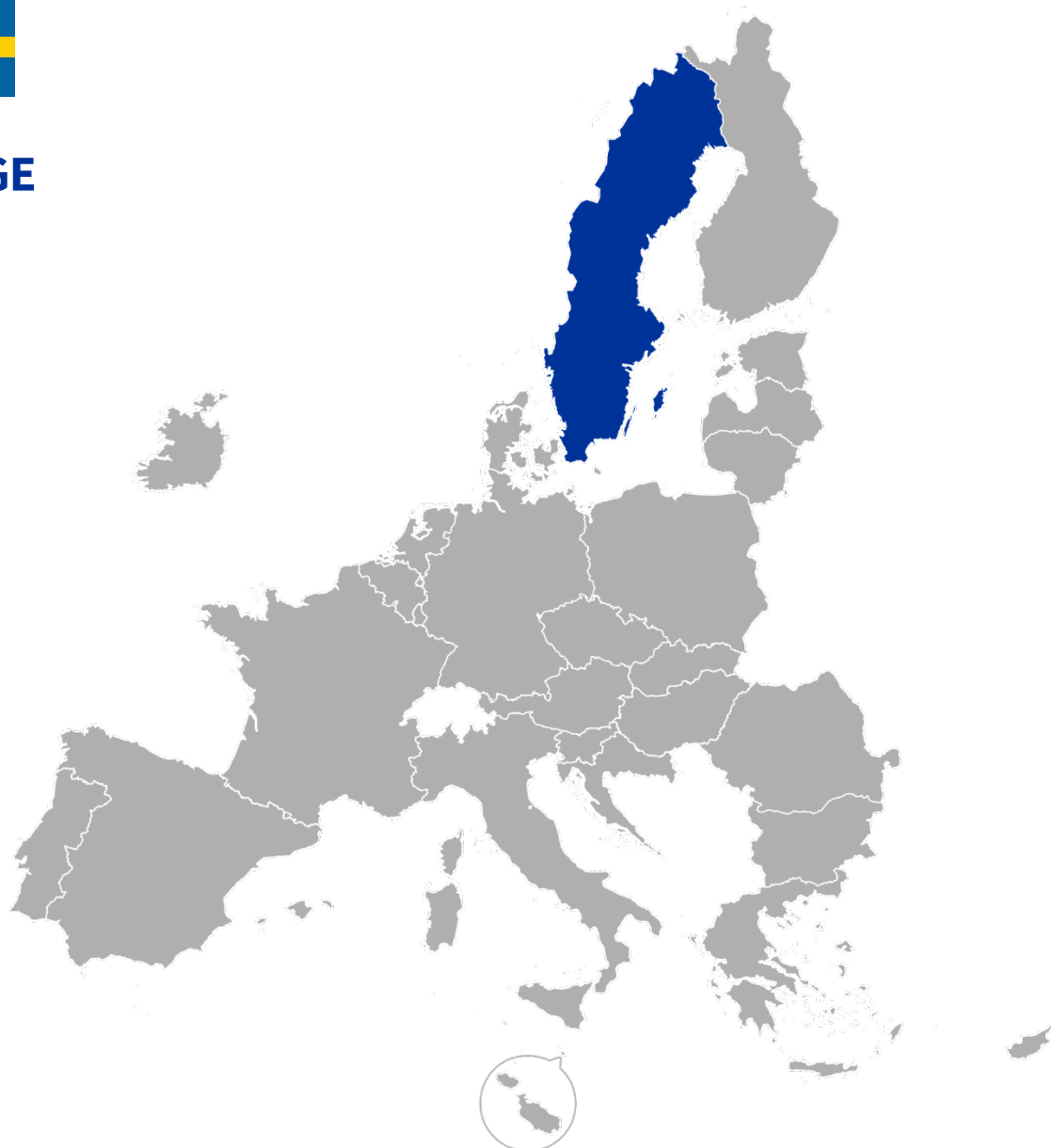




SVERIGE



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



SVERIGE

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Finland, Irland og Portugal**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Din nationale regering støtter fremsynet, men praktisk AI-regulering. Du argumenterer for begrænsede undtagelser i forslaget for at undgå at demotivere AI-investorer og -forskere.
- Alene kan Sverige ikke føre an inden for AI, men et forenet EU kan – ved at fremme innovation, fastsætte standarder og styrke sin globale indflydelse. For at opnå dette skal AI-udviklere have stærkere incitamenter og færre reguleringsmæssige hindringer i udviklingsfaserne, hvilket fremmer jobskabelse og økonomisk vækst.
- En formålsbaseret tilgang er afgørende: AI, der bruges til medicinsk diagnostik, bør være underlagt strengere regler end AI, der er designet til forretningsadministration. Dette er ikke for at bagatellisere bekymringer vedrørende f.eks. bias i AI, men snarere for at prioritere de mest umiddelbare risici – dem, der har direkte indvirkning på menneskers liv. Desuden har definitionen af risici baseret på sektorer også den fordel, at der er færre potentielle konflikter med eksisterende sektorbaserede krav i andre forskrifter, f.eks. rammerne for sundhedspleje og finans. På denne måde forbliver den bureaukratiske kompleksitet på et håndterligt niveau.
- Hvad angår krav til udviklere, er du overbevist om, at europæiske techvirksomheder kan levere AI af høj kvalitet uden rigid regulering. Det gavner både innovation og ekspertise, hvis udviklerne har fleksibilitet med hensyn til, hvordan de tester og finjusterer modeller, og samtidig sikres ansvarlig AI-udvikling.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Det er afgørende for dig, at spørgsmål vedrørende national sikkerhed, militær eller forsvar ikke bliver omfattet af AI-forordningen. Medlemsstaterne er fast besluttet på at bevare kontrollen over beslutninger på disse områder, som under alle omstændigheder typisk falder uden for EU's jurisdiktion. Dette ses også af, at EU ikke har et fælles EU-militær.
- Du argumenterer desuden for, at AI-systemer, der anvendes til indsamling af efterretninger, cyberforsvar eller slagmarksscenarier, ikke kan være underlagt de samme oplysningskrav eller reguleringsmæssige krav som civile systemer, fordi dette kunne kompromittere operationel hemmeligholdelse, nationale sikkerhedsinteresser og den strategiske fordel, som sådanne systemer er designet til at give.
- Artikel 2 er vigtigere for dig end artikel 1. Hvis du ikke kan overbevise andre om dine pointer vedrørende artikel 1, skal du fokusere på at sikre, at dine vigtigste krav vedrørende artikel 2 bliver taget alvorligt.
- Dine bemærkninger:

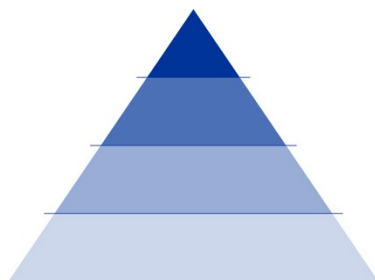
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige	formålsbaseret	fleksibel		militær/forsvar + national sikkerhed	
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

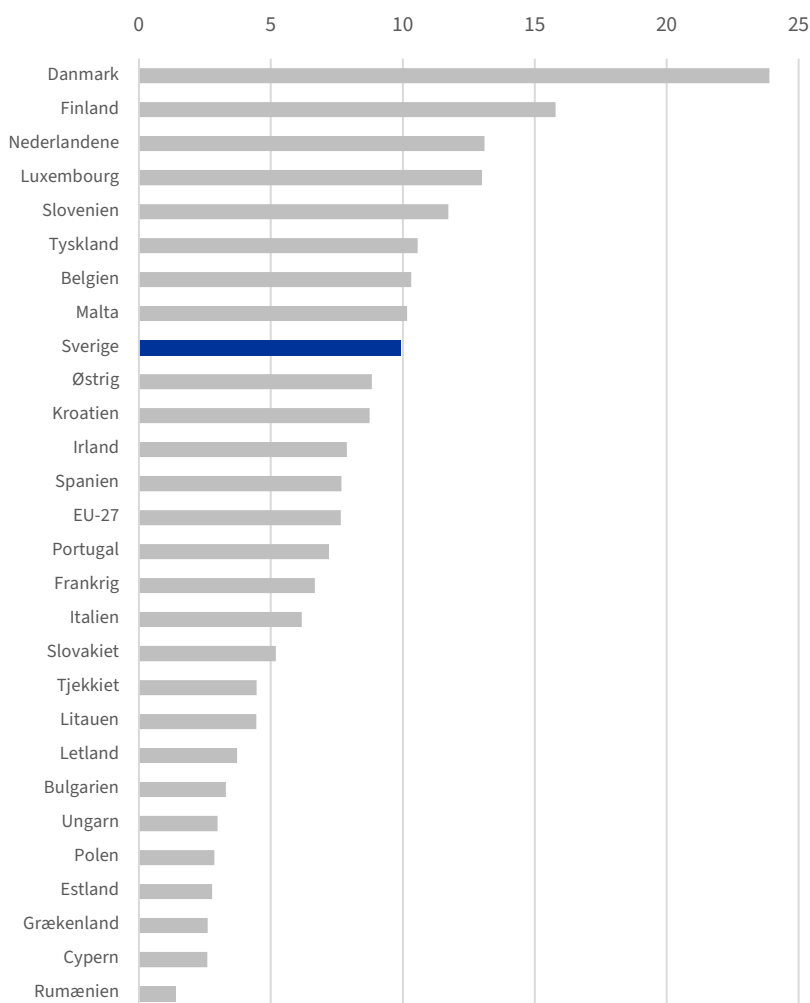
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

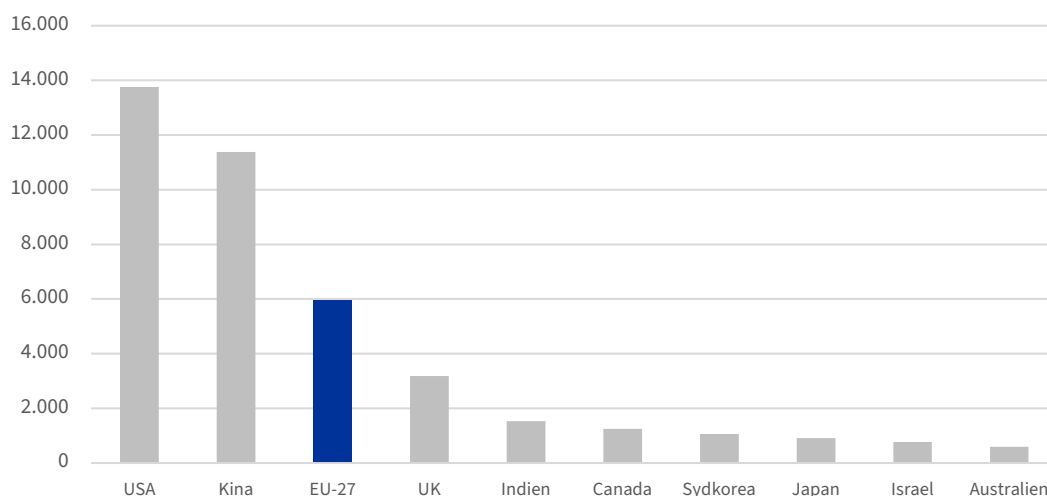
	Den Europæiske Union	Sverige
Befolkning	447 695 350	10 521 556
BNP pr. indbygger i €	38 150	51 060
Arbejdsløshedsprocent	6,1 %	7,7 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	10,4 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	36,5 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

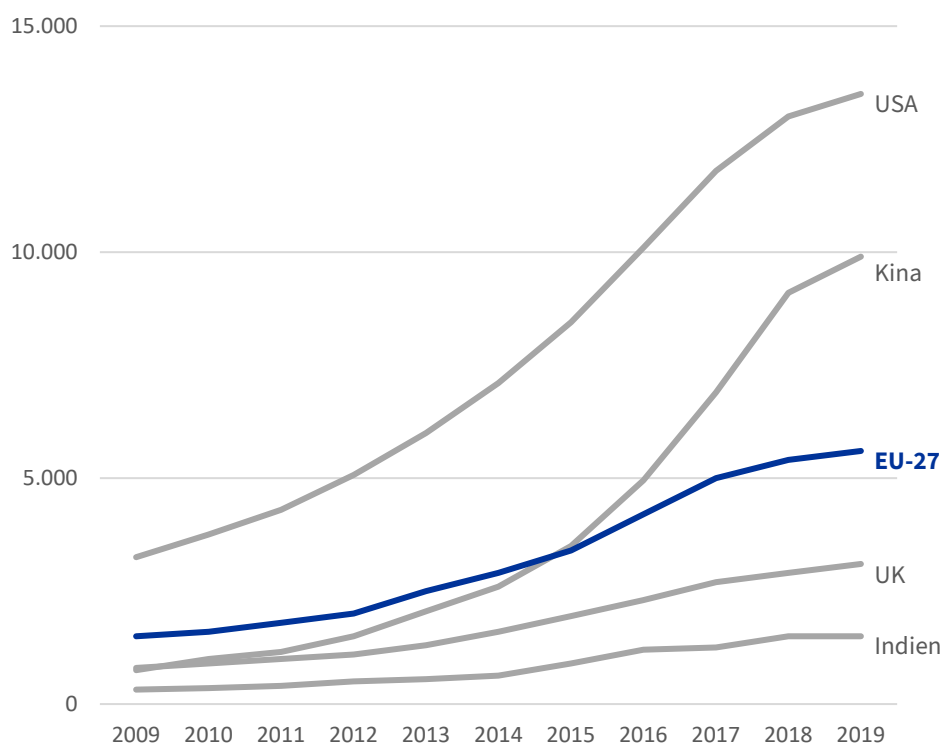


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

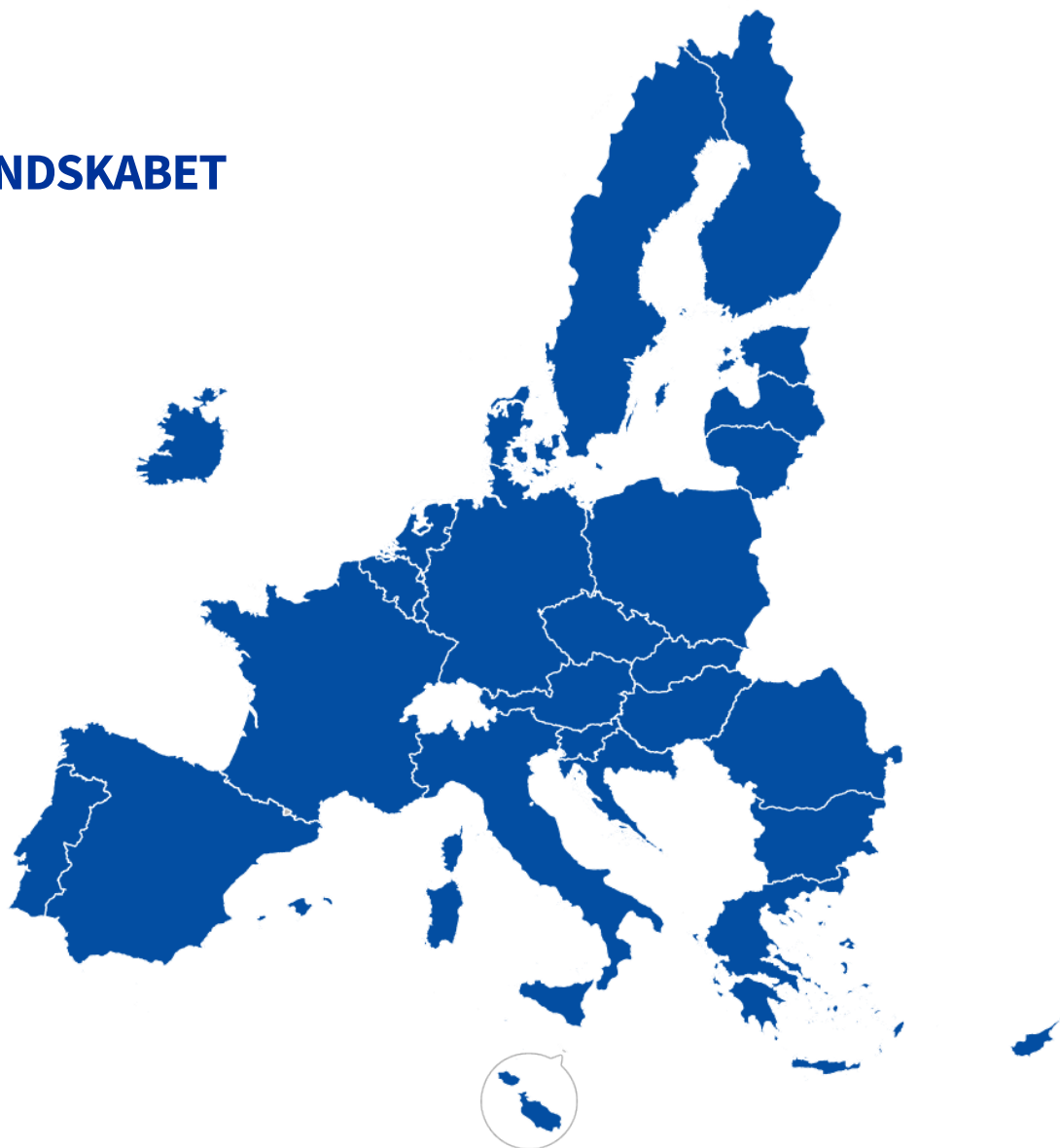
© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print	ISBN 978-92-850-0498-9	doi: 10.2860/0371762	QC-01-25-006-DA-C
PDF	ISBN 978-92-850-0497-2	doi: 10.2860/0996198	QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltageres læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



FORMANDSKABET



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

At nå til enighed om Rådets fælles holdning ved dobbelt flertal (dvs. mindst 55 % af medlemsstaterne, der repræsenterer 65 % af EU's samlede befolkning). Ideelt set stræbes der altid efter enstemmighed, men hvis dette ikke kan opnås, anvendes reglen om dobbelt flertal. I praksis sættes de fleste af afgørelserne ikke til afstemning, men formandskabet forsøger at finde et kompromis, som kan støttes af alle delegationer.



Din rolle

Du er en del af delegationen fra den medlemsstat, der for tiden varetager formandskabet. Spilkoordinatoren fortæller dig, hvilket land det er.

Som en del af delegationen er det din opgave at lede de officielle forhandlinger, mens dine kolleger fremlægger landets interesser og holdninger. Du er ansvarlig for at styre alle officielle forhandlinger, give de delegerede ordet og holde tidsplanen.

Indledning

1. Du skal **formelt indlede mødet**, byde alle delegerede og Kommissionen velkommen og forklare dem formålet med mødet (maks. 1 minut). Du kan eventuelt hente inspiration i manuskriptet i tekstboksen på næste side.
2. Herefter **beder du repræsentanten for Kommissionen om at fremlægge udkastet**.
3. Dernæst **følger de indledende bemærkninger fra medlemsstaterne** (maks. 1 minut hver), som du giver ordet i den rækkefølge, de har anmodet om det. Bed medlemsstaterne om klart at angive, på hvilke punkter de har fælles interesser med andre medlemsstater.

Drøftelser

1. **Du skal styre de efterfølgende officielle forhandlinger.** Ministrene må kun tale, når du giver dem ordet. Giv de delegerede ordet i den rækkefølge, de har anmodet om det.
2. Bed de delegerede om ikke at gentage argumenter og holdninger, der allerede har været fremsat. Fortæl medlemsstater med ens holdninger, at de er velkomne til at tale på hinandens vegne. **Tiden er begrænset!**
3. De officielle **forhandlinger afbrydes for at give plads til uformelle samtaler.** Du meddeler, når de uformelle samtaler begynder, og hvornår det formelle møde genoptages. De uformelle samtaler kan finde sted uden dit tilsyn.

Nå til enighed

1. Du har tid til at **udarbejde et kompromisforslag** under de uformelle samtaler. Du kan selvfølgelig også spørge medlemsstaterne, om de har forslag til et kompromis.
2. Du sætter det/de mest lovende forslag til artikel 1 og 2 til afstemning. Rådets fælles holdning vedtages ved dobbelt flertal (55 % af medlemsstaterne, der repræsenterer mindst 65 % af EU's samlede befolkning). Optæl stemmerne **ved hjælp af afstemningssimulatoren**. Spilkoordinatoren hjælper dig med det.
3. Når tiden er kommet til den endelige afstemning om et kompromisforslag, minder du medlemsstaterne om, at de på dette trin skal beslutte, om de foretrækker et kompromis, som de måske ikke er 100 % tilfredse med, eller intet resultat.

Medlemmer af Rådet, kommissær

Det er mig en fornøjelse at indlede dagens møde i Rådet.

Som I ved, varetages formandskabet for Rådet for tiden af _____, som vil lede forhandlingerne. Min rolle er at være ordstyrer for dagens møde, så jeg vil ikke fremlægge vores lands holdninger til spørgsmålet.

På dagsordenen i dag har vi en ny forordning om anvendelse og udvikling af kunstig intelligens, også kaldet AI, i EU. Spørgsmålet er meget vigtigt for alle europæiske borgere og for EU's rolle i verden, fordi _____.

Inden jeg giver ordet til Kommissionen, vil jeg gerne understrege, at Rådets Juridiske Tjeneste har gennemgået forslaget og givet os besked om, at det er i overensstemmelse med traktaterne og har det nødvendige retsgrundlag.

Vi beder nu Kommissionen om at fremlægge sit forslag. Derefter vil hver af jer få mulighed for at fremlægge jeres lands holdning i nogle korte indledende bemærkninger. Vi ser frem til en konstruktiv og frugtbar debat.

Kommissær, ordet er dit.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

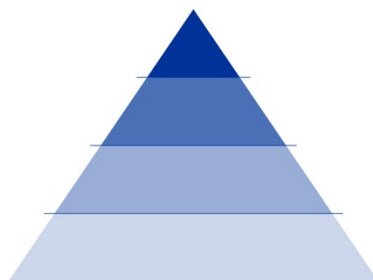
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig	formålsbaseret	forsigtigheds-		national sikkerhed	
Belgien	kapabilitetsbaseret	forsigtigheds-		open source	
Bulgarien	formålsbaseret	fleksibel		uden undtagelse	
Kroatien	formålsbaseret	forsigtigheds-		opstartsvirksomheder	
Cypern	kapabilitetsbaseret	forsigtigheds-	revision om to år	uden undtagelse	
Tjekkiet	kapabilitetsbaseret	fleksibel	revision om fem år	uden undtagelse	
Danmark	formålsbaseret	forsigtigheds-	finansiel støtte	militær/forsvar + national sikkerhed	
Estland	formålsbaseret	fleksibel	afprøvning under faktiske forhold	militær/forsvar + national sikkerhed	
Finland	formålsbaseret	fleksibel	udsætte afstemningen	opstartsvirksomheder	
Frankrig	formålsbaseret	forsigtigheds-	inkludere uddannelse	militær/forsvar + national sikkerhed	
Tyskland	kapabilitetsbaseret	forsigtigheds-		uden undtagelse	
Grækenland	kapabilitetsbaseret	forsigtigheds-	revision om tre år	militær/forsvar + national sikkerhed	
Ungarn	kapabilitetsbaseret	fleksibel		militær/forsvar	
Irland	formålsbaseret	fleksibel	fleksibilitet med hensyn til afprøvning	opstartsvirksomheder	

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien	kapabilitetsbaseret	forsigtigheds-		opstartsvirksomheder	
Letland	kapabilitetsbaseret	fleksibel	fjerne overvågning efter markedsføring	militær/forsvar	
Litauen	kapabilitetsbaseret	fleksibel		uden undtagelse	
Luxembourg	formålsbaseret	fleksibel	revision om tre år	opstartsvirksomheder	
Malta	formålsbaseret	fleksibel	afprøvning under faktiske forhold	national sikkerhed	
Nederlandene	kapabilitetsbaseret	forsigtigheds-		national sikkerhed	
Polen	formålsbaseret	forsigtigheds-		militær/forsvar	
Portugal	formålsbaseret	fleksibel	udsætte afstemningen	open source	
Rumænien	formålsbaseret	fleksibel		national sikkerhed	
Slovakiet	formålsbaseret	fleksibel	finansiel støtte	opstartsvirksomheder	
Slovenien	kapabilitetsbaseret	forsigtigheds-		opstartsvirksomheder	
Spanien	kapabilitetsbaseret	forsigtigheds-		militær/forsvar + national sikkerhed	
Sverige	formålsbaseret	fleksibel		militær/forsvar + national sikkerhed	
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



Definitioner

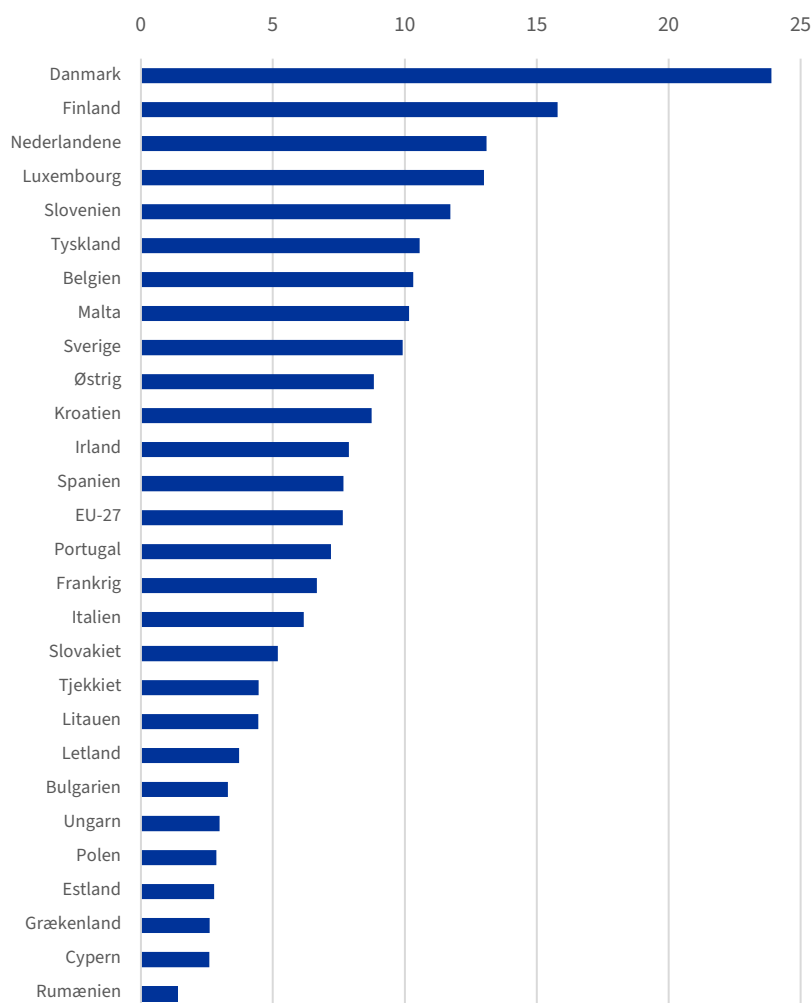
- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkørende AI i en simuleret by, inden der køres på rigtige veje

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

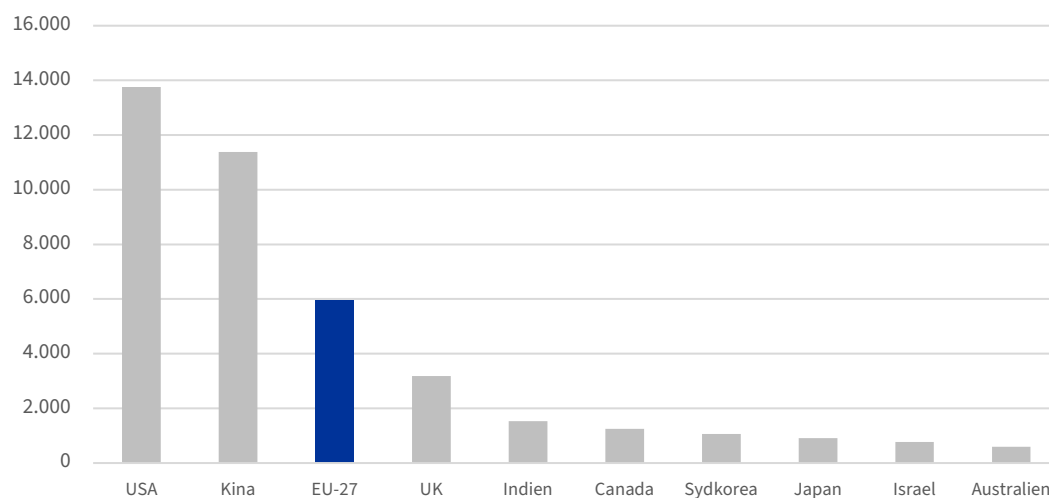
	Den Europæiske Union
Befolkning	447 695 350
BNP pr. indbygger i €	38 150
Arbejdsløshedsprocent	6,1 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

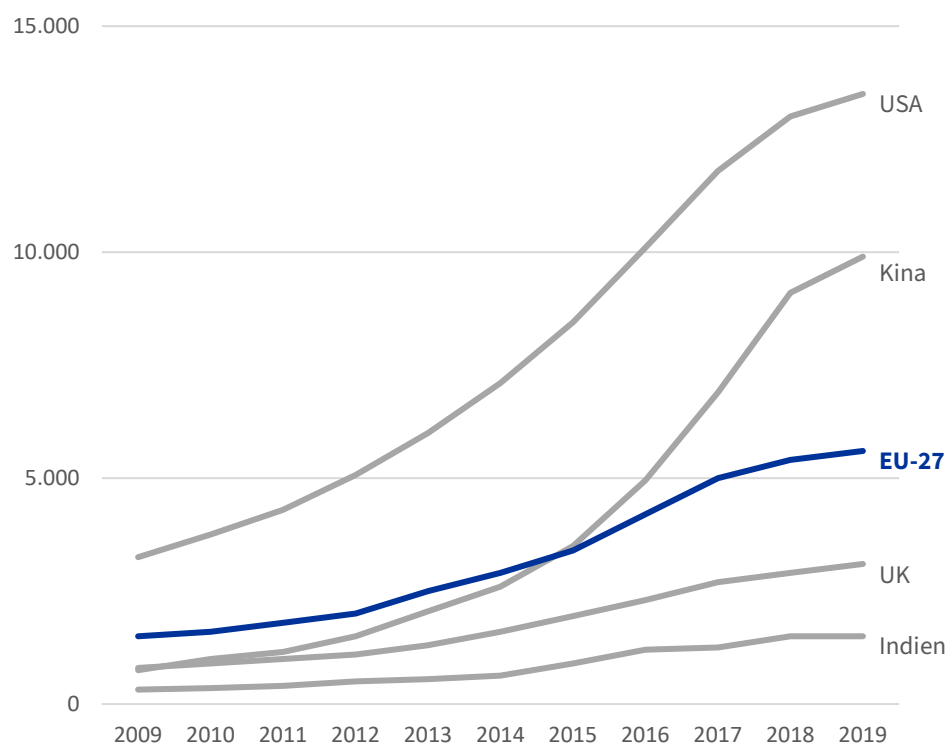


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print	ISBN 978-92-850-0498-9	doi: 10.2860/0371762	QC-01-25-006-DA-C
PDF	ISBN 978-92-850-0497-2	doi: 10.2860/0996198	QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltageres læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



EUROPA-KOMMISSIONEN



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens.

Når formandskabet (som leder mødet) giver dig ordet, forelægger du dit forslag for repræsentanterne for medlemsstaterne og deltager i forhandlingerne.

Dit primære mål

At sikre, at så meget som muligt af den oprindelige tekst kommer med i den endelige udgave, mens du hjælper Rådet med at finde en fælles holdning. Der er behov for en fælles europæisk løsning her og nu!

Din opgave

1. **Læs** grundigt **dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. **Forbered de indledende bemærkninger** (maks. 1 minut), hvor du forelægger og forklarer dit forslag. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen til højre.
3. Hør godt efter, når medlemsstaterne fremsætter deres holdninger. **Udfyld tabellen** på side 8 og 9 for at holde styr på deres holdninger.
4. Under forhandlingerne skal du forklare dine holdninger og **forsvare dit forslag**. Tag ordet, og hjælp Rådet med at finde et kompromis, der lever op til dine forventninger.
5. **Du har ikke stemmeret** til sidst. Brug de uformelle pauser til at overbevise delegationerne om, at de skal følge dine forslag.



Ærede medlemmer af Rådet

Jeg er glad for, at vi i dag har denne vigtige debat om forordningen om kunstig intelligens i EU. Forordningen vil være til stor gavn for borgerne og for EU som helhed, fordi

_____.

I Kommissionens forslag til forordning foreslår vi en **formålsbaseret risikoklassificering** for højrisiko-AI-applikationer **og forsigtighedstestkrav**, fordi _____

_____.

Der tillades undtagelser fra forskrifterne som fastsat i artikel 1, når AI-systemer er udviklet til militære eller forsvarsmæssige formål. Dette skyldes _____

_____.

Lad os sammen som samfund bidrage til at sikre, at AI udvikles ansvarligt.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Der kræves en hurtig indsats for at fremme innovation og styrke den globale konkurrenceevne, som AI fremstillet i EU har, samtidig med at potentielle risici for de grundlæggende rettigheder og forbrugerbeskyttelsen håndteres behørigt.
- Som en ny teknologi med uforudsigelige konsekvenser for enkeltpersoner, samfund, virksomheder og økonomier skal kunstig intelligens behandles med både ambitioner og forsigtighed. Dit forslag afspejler dette ved at støtte en formålsbaseret tilgang til definitionen af, hvilke AI-systemer der bør klassificeres som højrisiko. Med den formålsbaserede tilgang anses AI for at være højrisiko på følsomme områder, der omfatter beslutninger, der gælder liv og død, eller har en betydelig indvirkning på EU-borgernes trivsel.
- Rådet er allerede nået til enighed om kategorier for AI-systemer med uacceptabel risiko, som skal forbydes i EU, samt for systemer med lav og ingen risiko. I dag drejer det sig om at definere klart, hvad højrisiko er. Et AI-system, der hører under denne kategori, må ikke behandles som et system med lav risiko eller ingen risiko.
- Du forsøger også at fastsætte klare krav til højrisiko-AI-systemer i dagens drøftelser. Forsigtighedsrammen for test sikrer, at AI udvikles ansvarligt og med ansvarlighed. Samtidig får den slags standarder udviklerne til at producere troværdige systemer af høj kvalitet, som opnår global anerkendelse. Med denne tilgang forventer du, at EU bliver førende inden for AI og hurtigt haler ind på USA og Kina.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Samtidig med at det fortsat prioriteres at fremme et forenet digitalt lederskab i EU og styrke udenrigshandelsdimensionen – såsom eksport af EU-udviklede AI-produkter globalt – er det vigtigt at bevare den nationale autonomi i specifikke sektorer, som er afgørende for den offentlige sikkerhed.
- F.eks. eskalerer de internationale sikkerhedsudfordringer hurtigt, og AI spiller en central rolle for udviklingen af cyberkrigsførelse og ondsindet hackingsoftware mv. Defensive modforanstaltninger må ikke hæmmes af reguleringsmæssige byrder, uanset hvor små de er. Derfor bør artikel 1 ikke finde anvendelse på militære og forsvarsmæssige applikationer.
- Til alle andre formål kræves der ikke yderligere undtagelser, da minimumstestkrav sikrer, at AI udvikles ansvarligt.
- Dine bemærkninger:

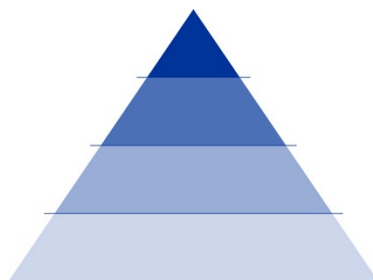
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



Definitioner

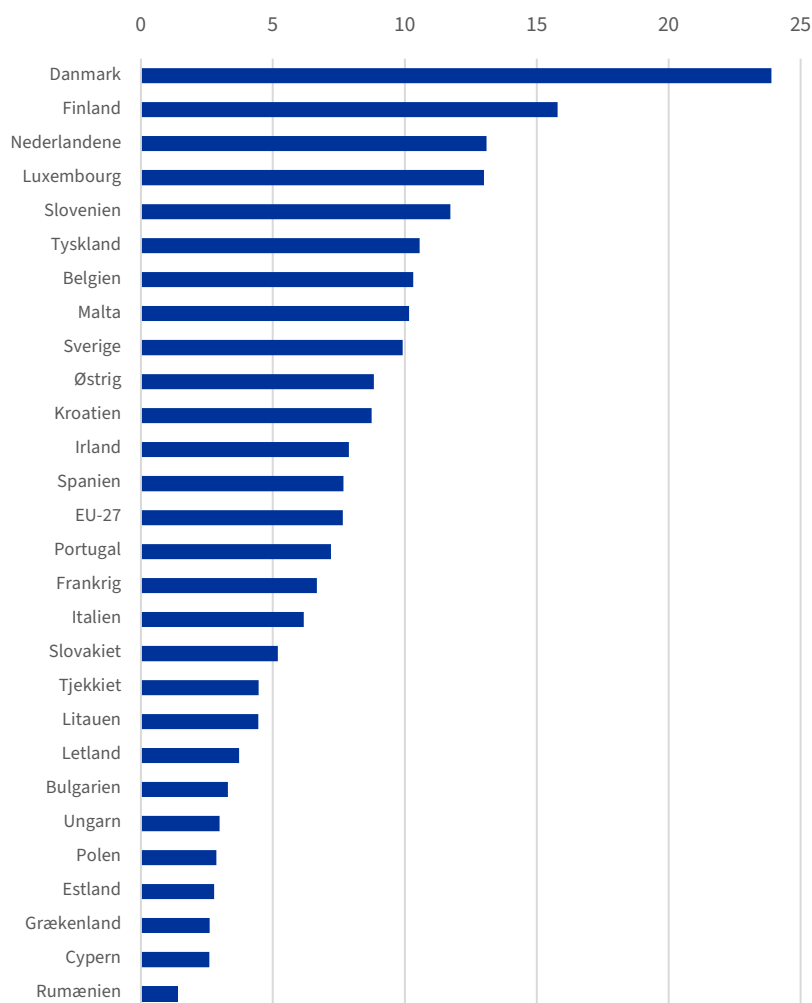
- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

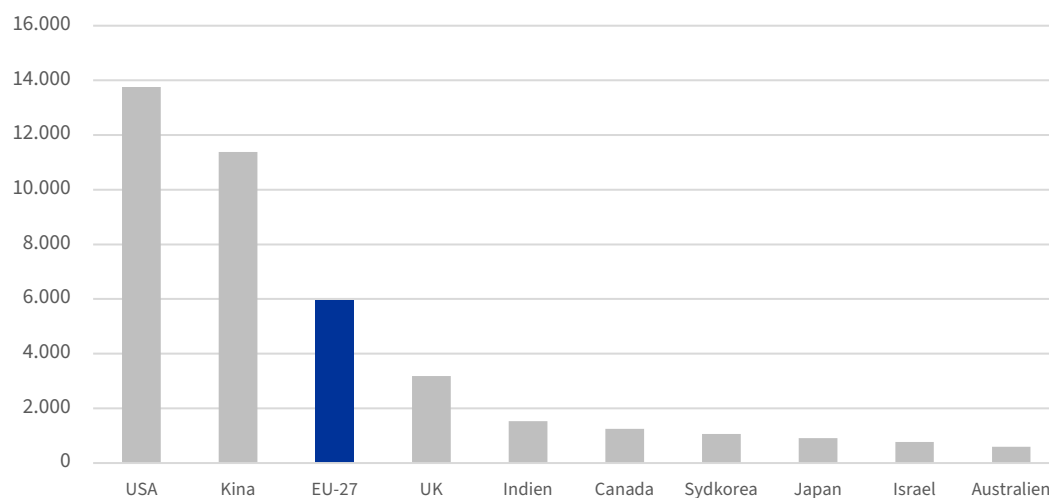
	Den Europæiske Union
Befolkning	447 695 350
BNP pr. indbygger i €	38 150
Arbejdsløshedsprocent	6,1 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

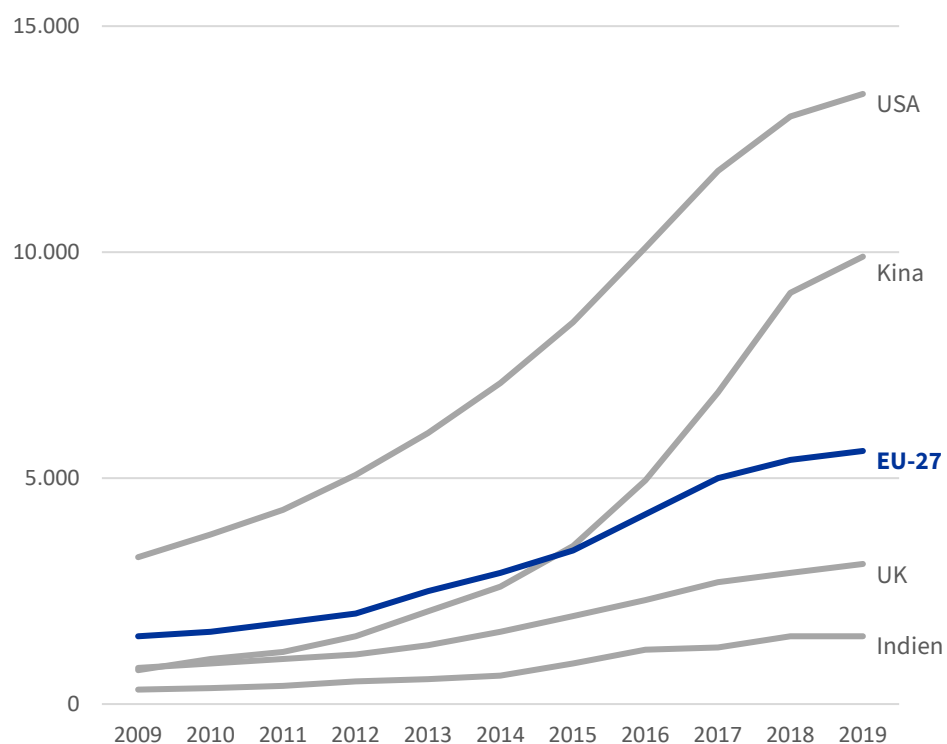


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

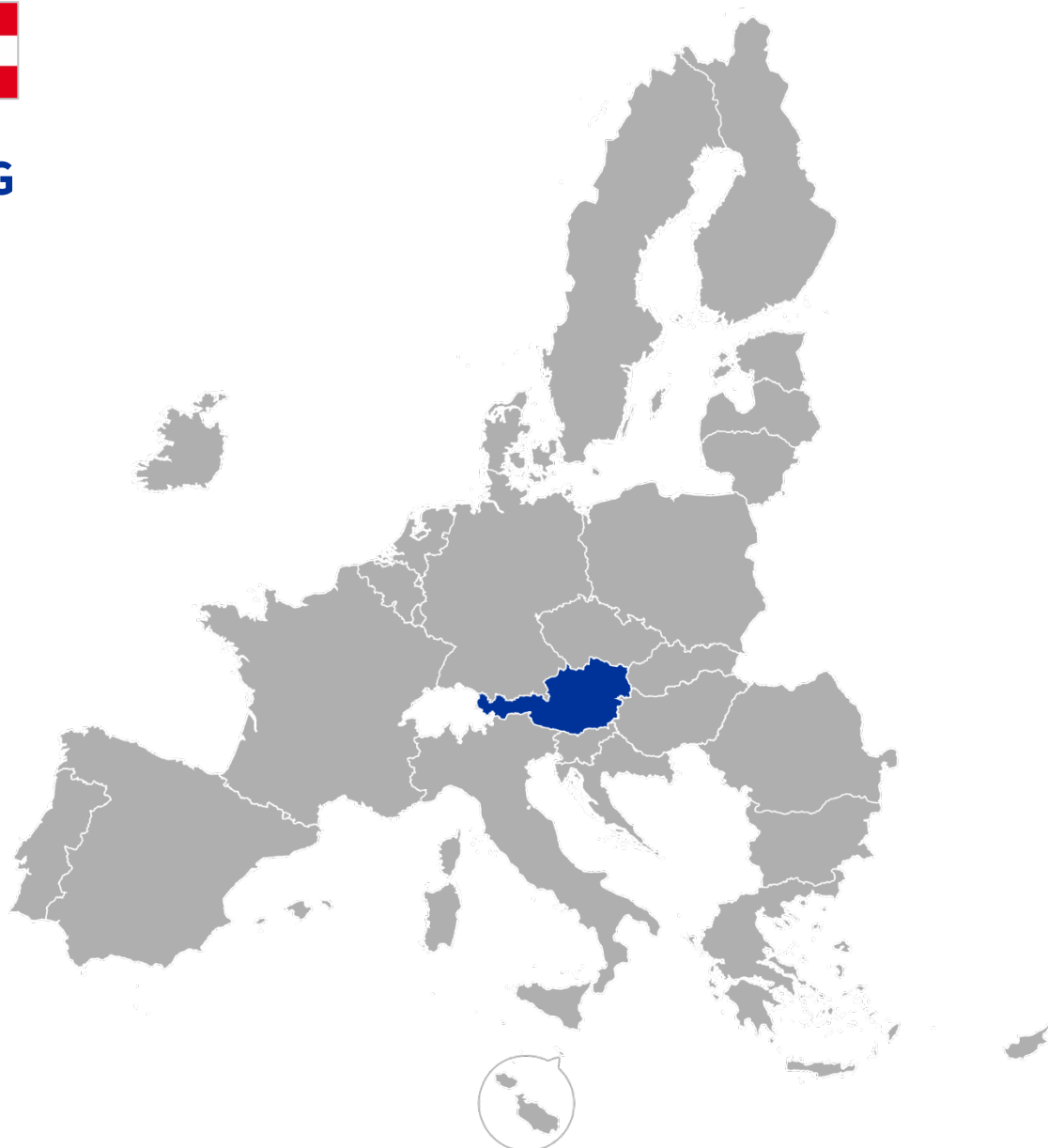
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



ØSTRIG



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



ØSTRIG

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Danmark, Frankrig, Polen og Rumænien**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Din nationale regering investerede næsten 350 mio. € i AI-forskning mellem 2012 og 2017 og skabte et solidt grundlag for innovation. Østrigske virksomheder er nu førende inden for AI-udvikling til gavn for både det indenlandske og de internationale markeder. Du støtter EU's holdning til AI-forvaltning, men opfordrer til øget finansiering på EU-plan for at sætte mere skub i forskning og udvikling.
- Du går stærkt ind for anvendelsen af den formålsbaserede tilgang gennem princippet om teknologineutralitet. Det betyder, at forskrifter ikke må favorisere eller diskriminere mod bestemte AI-teknologier, -modeller eller -tilgange. I følsomme sektorer som f.eks. sundhedssektoren er det uden betydning, om systemet anvender dyb læring eller regelbaseret AI – den slags systemer anses for at være højrisiko, fordi de kan være ansvarlige for en persons død, hvis AI svigter eller begår en fejl.
- Den formålsbaserede tilgang er også afgørende for AI-import. Det er ofte umuligt at kontrollere importerede systemers reelle kapabiliteter udelukkende ud fra udviklerens erklæring – især når udenlandske enheder kan have grund til at skjule skadelige elementer såsom spyware.
- Med hensyn til testkrav går du ind for en ramme for forsigtighedstest, da den fremmer en ansvarskultur, der gør det lettere at spore, hvor det er gået galt, når der opstår problemer.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- I partnerskab med Interpol, Europol og FN's Kontor for Narkotikakontrol og Kriminalitetsbekæmpelse – med hovedsæde i Wien – spiller det østrigske teknologiinstitut en ledende rolle med hensyn til at integrere AI-systemer i national sikkerhed gennem avancerede overvågningsløsninger og immersive træningsplatforme. Forskere engagerer sig i at udvikle yderst troværdige og pålidelige AI-modeller, der bidrager til at forbedre Østrigs sikkerhed. Da national sikkerhed typisk falder uden for EU's lovgivningsmæssige rammer, mener du, at det er rimeligt at undtage dette område fra AI-regulering.
- Som du har fremhævet i artikel 1, er det desuden vigtigt at præcisere, at den lovgivningsmæssige ramme også skal omfatte AI-systemer udviklet uden for EU og efterfølgende importeret. Forordningen bør finde anvendelse på alle AI-systemer, der anvendes i EU, uanset deres oprindelse.
- Dine bemærkninger:

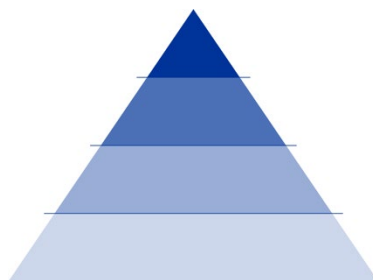
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig	formålsbaseret	forsigtigheds-		national sikkerhed	
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

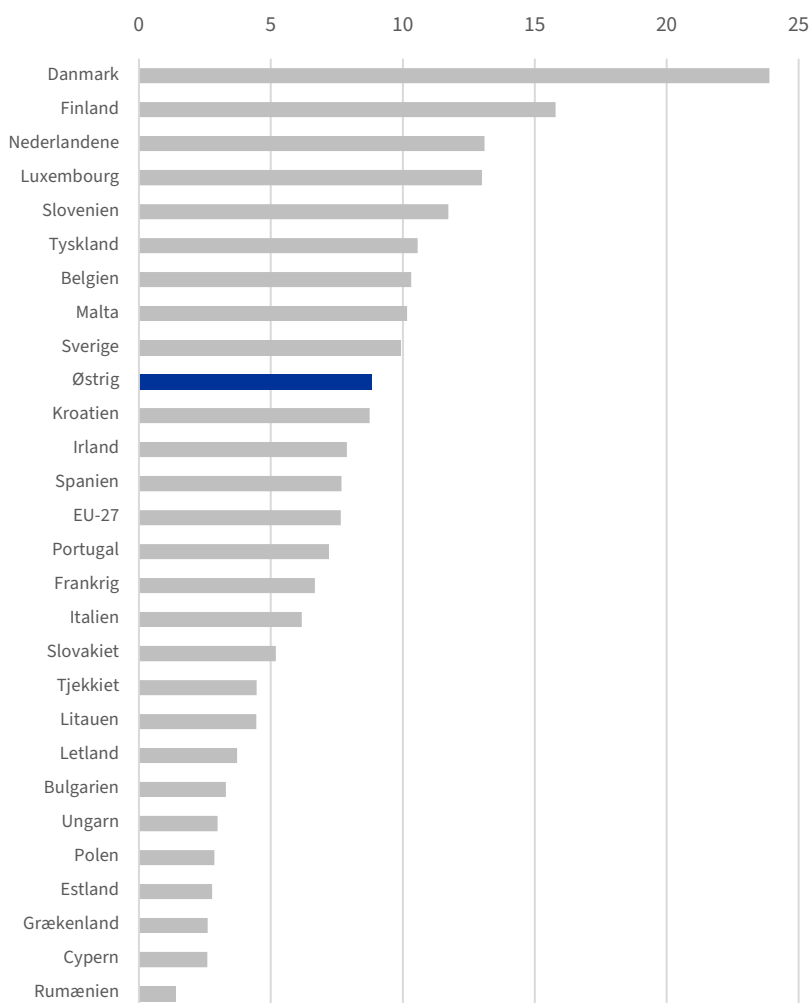
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

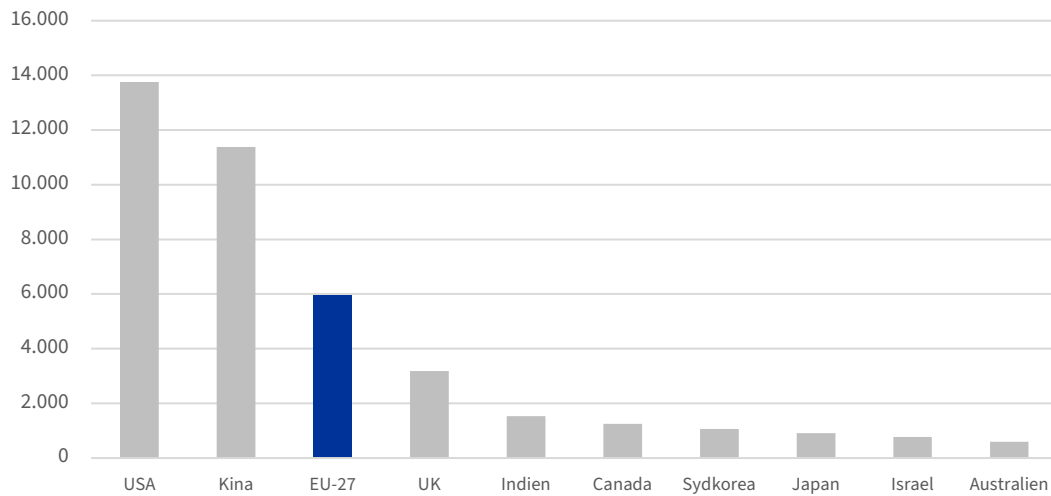
	Den Europæiske Union	Østrig
Befolkning	447 695 350	9 104 772
BNP pr. indbygger i €	38 150	51 830
Arbejdsløshedsprocent	6,1 %	5,1 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	10,8 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	32 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

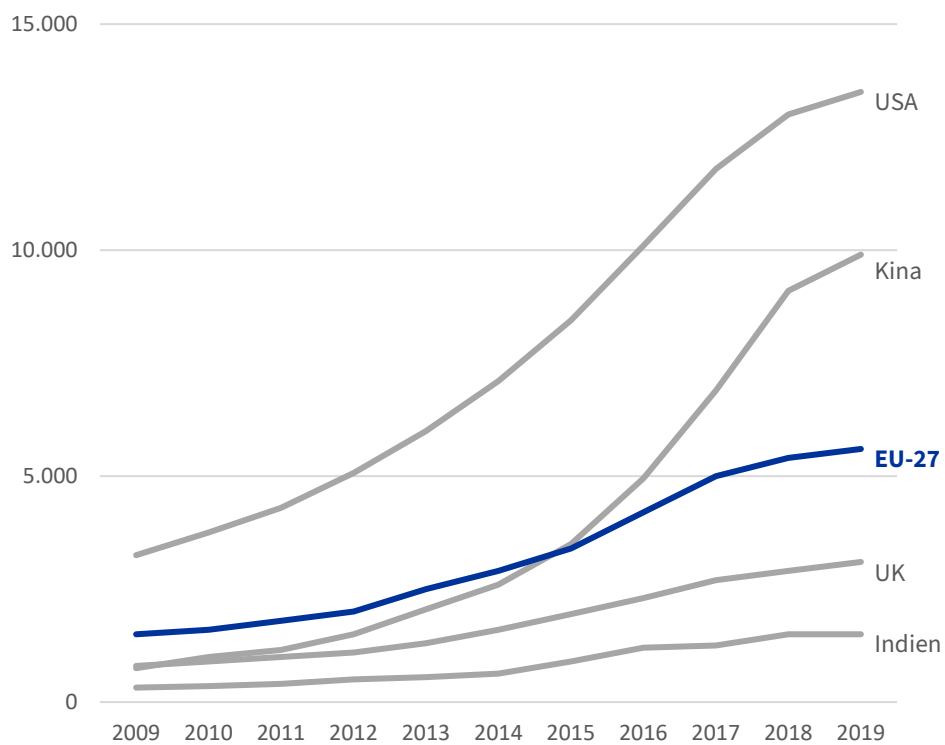


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9
PDF ISBN 978-92-850-0497-2

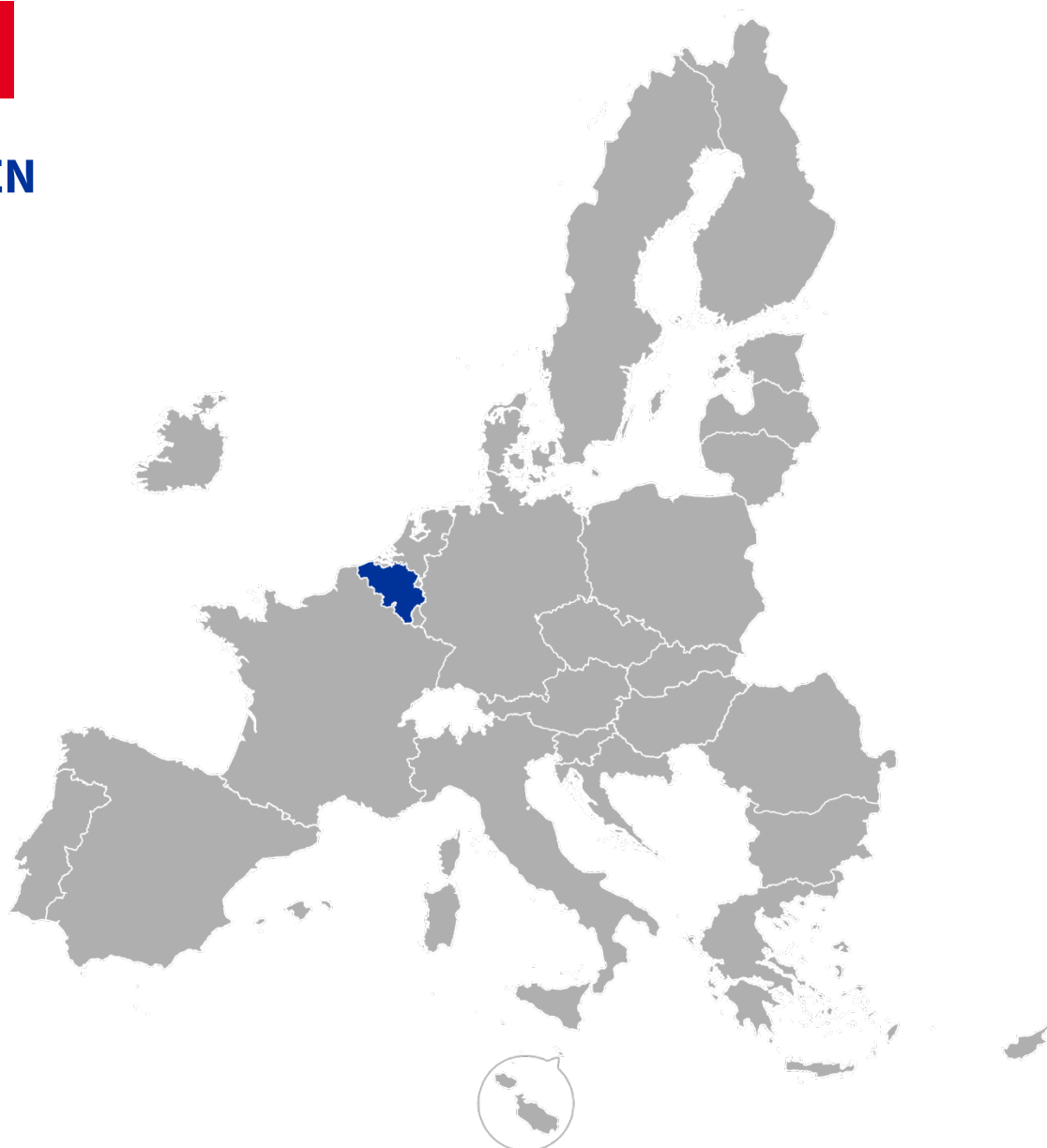
doi: 10.2860/0371762
doi: 10.2860/0996198

QC-01-25-006-DA-C
QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



BELGIEN



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Cypern, Grækenland, Spanien og Tyskland**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.

Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Dit land er fast besluttet på at forsvare de grundlæggende rettigheder som en topprioritet. Håndtering af nationale og tværnationale trusler – også dem, der opstår som følge af digitale værktøjer – kræver omgående grænseoverskridende samarbejde.
- Du ønsker at sikre, at AI udvikles ansvarligt. Siden 2010'erne har udviklerne haft store gennembrud og skabt nye værktøjer så hurtigt, at samfundet ikke har kunnet følge med. Kun få personer forstår fremtiden med AI fuldt ud: Kunstig intelligens til almen brug (GPAI) kan måske snart efterligne menneskeligt ræsonnement, mens kvante-AI kan øge kapabiliteterne eksponentielt gennem kvantecomputing – hvilket kan få hidtil usete konsekvenser.
- Hvis AI udvikles ansvarligt, kan den forbedre liv og tackle globale udfordringer mere effektivt. Det kræver imidlertid etiske overvejelser. Derfor argumenterer du for den mest forsigtige og ansvarlige tilgang. Afprøvning under faktiske forhold som foreslået i den fleksible tilgang duer overhovedet ikke. Den skaber unødvendige risici for brugerne – risici, der kan afbødes gennem streng indledende sandkasseafprøvning.
- Du støtter en kapabilitetsbaseret risikoklassificering: Ethvert AI-system med potentiale for misbrug, bias eller skadelig indvirkning bør betragtes som højrisiko. AI-udvikling bør ske i samklang med fremme af mangfoldighed og menneskerettigheder – ikke modarbejde disse ting.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Du argumenterer for, at forordningen får et bredt anvendelsesområde. Efter din opfattelse giver det ikke mening at nå til enighed om en AI-forordning på EU-plan, der skal standardisere ansvarlig AI-udvikling, og så fra starten undlade at inddrage meget vigtige områder som sikkerhed i den.
- Den eneste undtagelse, du kan godtage, er open source-modeller: Disse modeller sætter gang i innovation, fordi de udvikles på en gennemsigtig måde. Udviklerne eksperimenterer og uddanner sig sammen og deler viden for at forbedre værktøjerne. Open source-AI er ofte mere inspicierbar og kontrollerbar end ejendomsretligt beskyttede modeller, hvilket er i tråd med retsaktens mål.
- Dine bemærkninger:

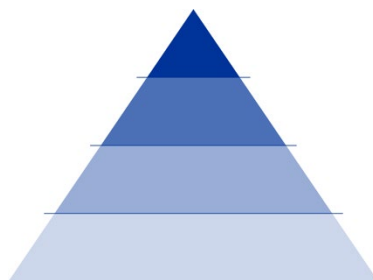
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien	kapabilitetsbaseret	forsigtigheds-		open source	
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

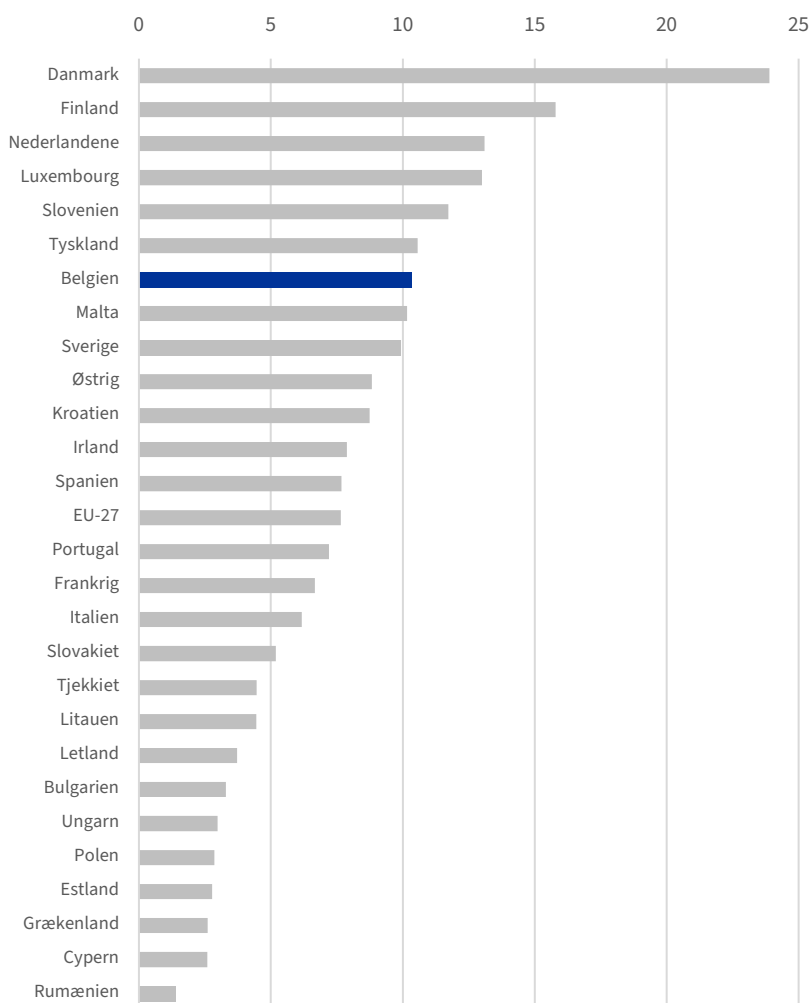
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

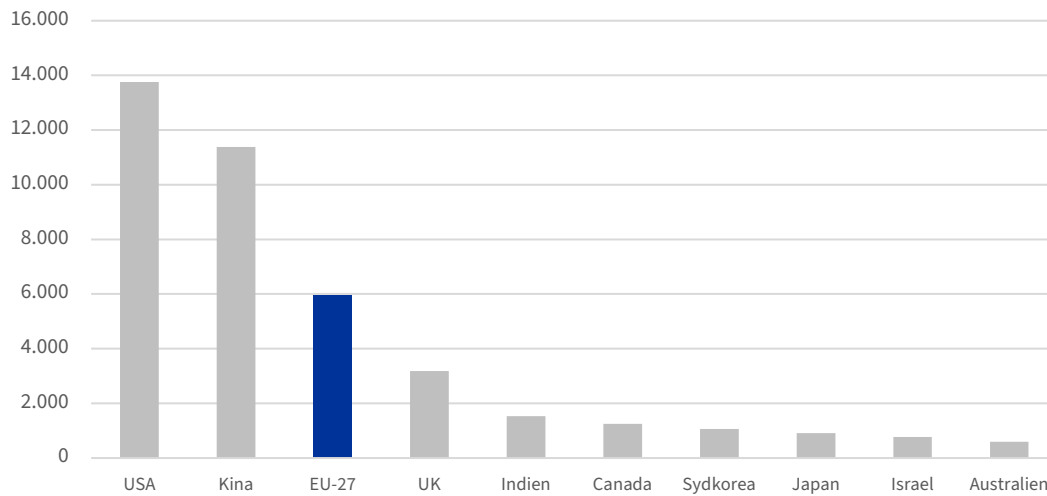
	Den Europæiske Union	Belgien
Befolkning	447 695 350	11 742 796
BNP pr. indbygger i €	38 150	50 610
Arbejdsløshedsprocent	6,1 %	5,5 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	13,8 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	28,3 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

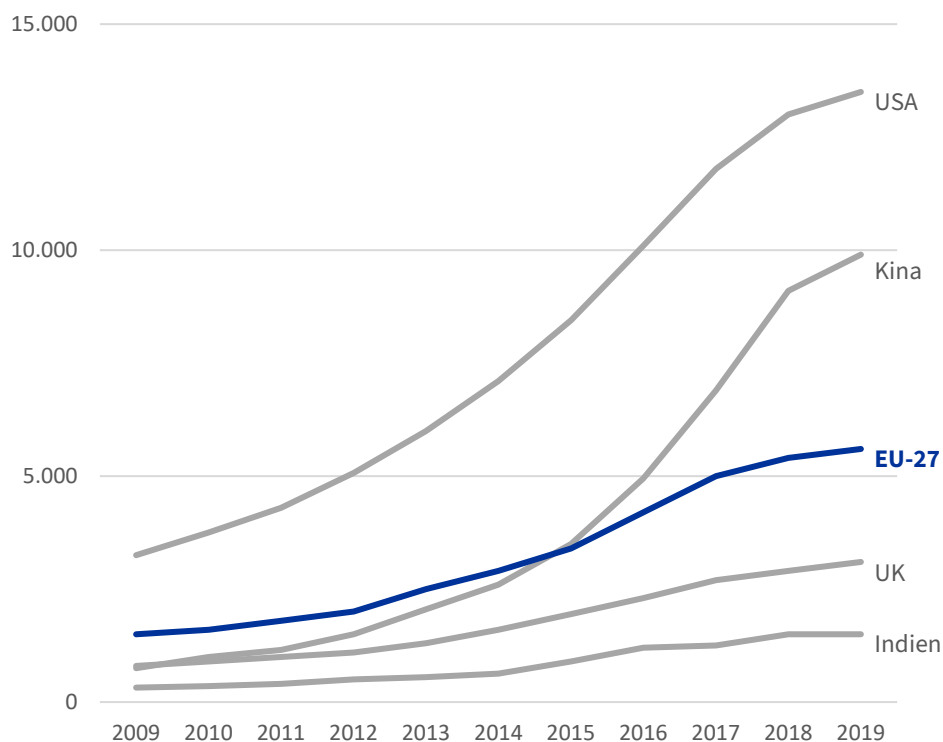


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

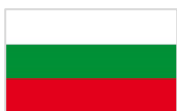
QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

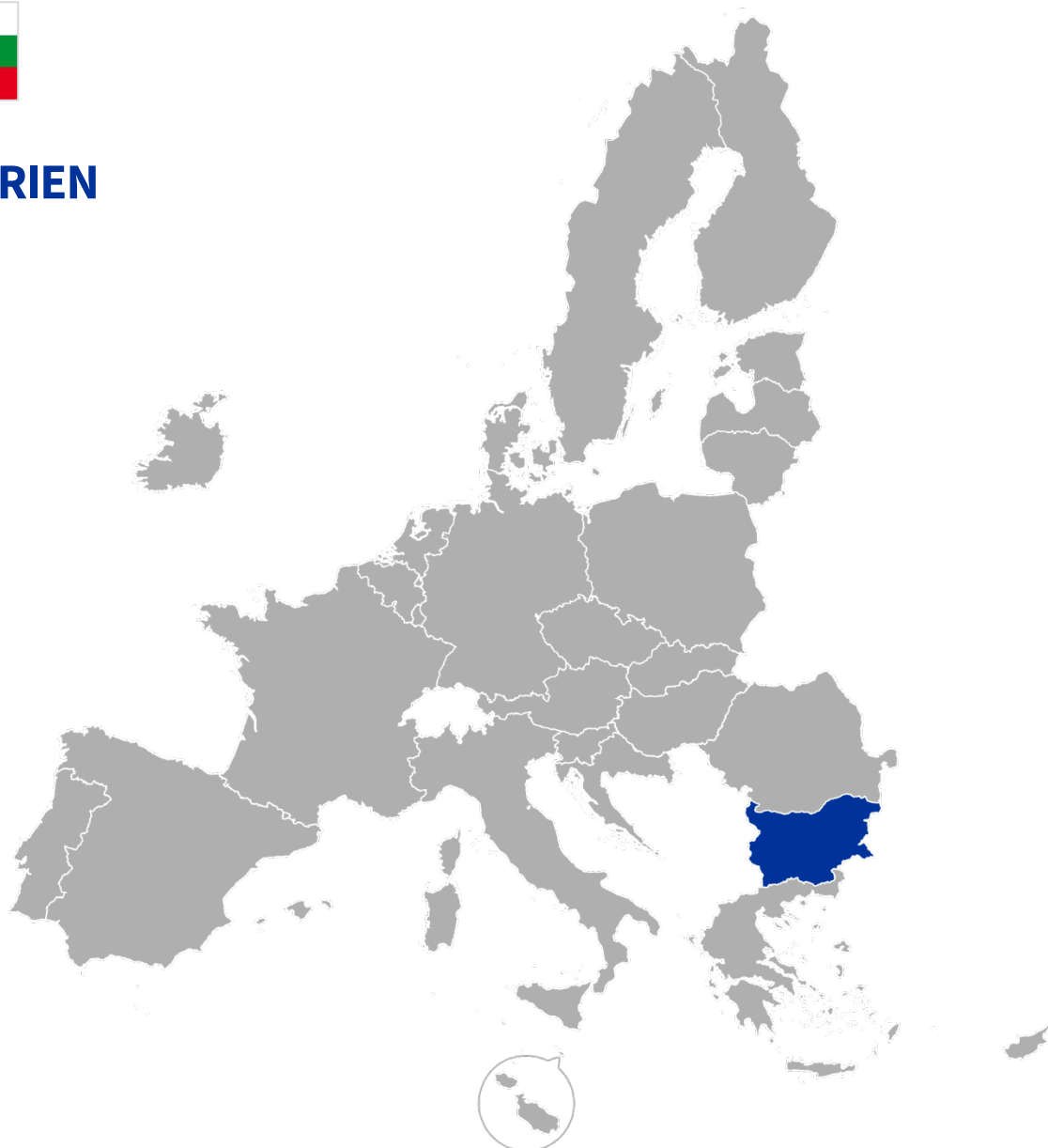
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



BULGARIEN



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

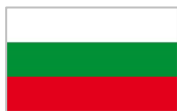
Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



BULGARIEN

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **delar nogle af argumenterne med lande som Luxembourg, Malta og Slovakiet**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Bulgariens højteknologi- og IT-sektorer udfolder sig inden for AI med stærke bånd til kunder i Vesteuropa og USA, og AI tegner sig allerede for 3 % af Bulgariens arbejdsmarked – en betydelig andel, som landet har til hensigt at øge. Du vil bruge EU's AI-strategi til at opnå resultater, der støtter nationale industrier og universiteter i at maksimere fordelene ved AI.
- I princippet støtter du den formålsbaserede tilgang. En sådan tilgang tilskynder udviklere og idriftsættere til nøje at overveje, hvordan deres system kan anvendes, ikke bare hvad de kan. Dette øger ansvarligheden i brugen under faktiske forhold.
- Men for dig er det vigtigere, at EU anerkender geografiske forskelle. Derfor opfordrer du til geografisk balance. Rigere lande har langt større finansielle og menneskelige ressourcer til AI-forskning, også til langvarige og dyre testprocedurer inden idriftsættelse. For at sikre en retfærdig udvikling i hele EU har Bulgarien brug for finansiering til at støtte sine små og mellemstore virksomheder og opstartsvirksomheder.
- Din topprioritet i dag er at sikre et stærkt tilsagn om EU-finansiering af AI-forskning i Bulgarien. Du er åben over for at drøfte krav til risikostyring og test – men først når der er forståelse for dine største betænkeligheder.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Selv om du støtter fleksible bestemmelser i artikel 1, mener du stadig, at der bør gælde klare regler for alle sektorer – også nationale sikkerheds- og efterretningssektorer. Hvis de undtages, opstår der smuthuller, og der er større risiko for, at uetisk eller skadelig AI får frit spil.
- Dette handler ikke om at blokere AI i national sikkerhed, men om at sikre, at den er designet ansvarligt for at undgå bias eller forskelsbehandling.
- Du mener, at forslagets nuværende tekst er for vag, og ønsker at præcisere, at disse regler også gælder for importerede AI-modeller fra udviklere uden for EU. Det er ganske enkelt for at undgå uklarheder. Formålet med AI-forordningen er netop at sætte skub i EU's udvikling inden for AI og gøre EU's AI-modeller mere attraktive end amerikanske eller kinesiske produkter.
- Dine bemærkninger:

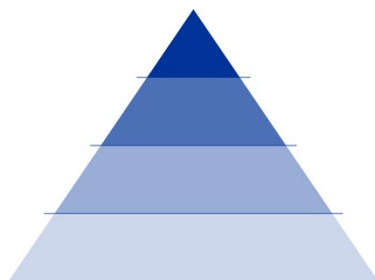
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien	formålsbaseret	fleksibel		uden undtagelse	
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

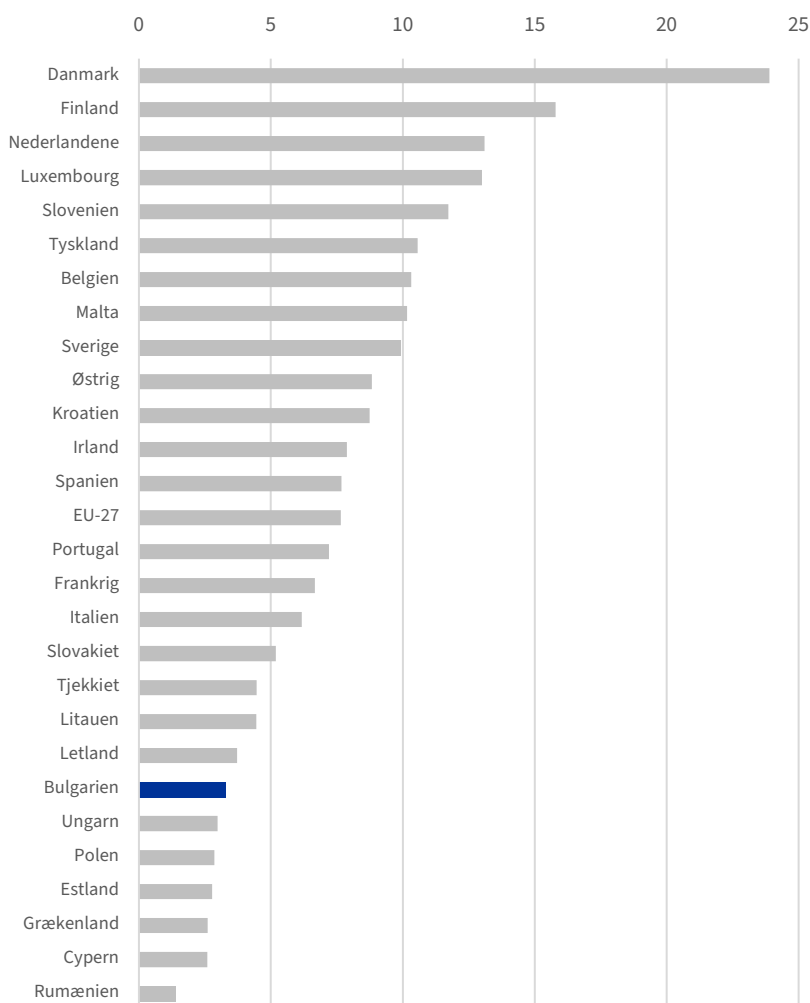
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

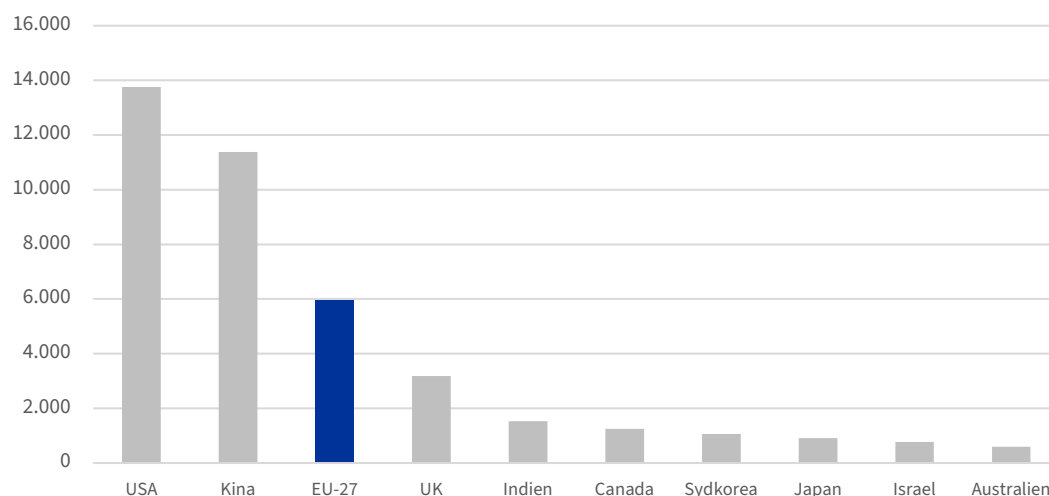
	Den Europæiske Union	Bulgarien
Befolkning	447 695 350	6 447 710
BNP pr. indbygger i €	38 150	14 690
Arbejdsløshedsprocent	6,1 %	4,3 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	3,6 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	7,7 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

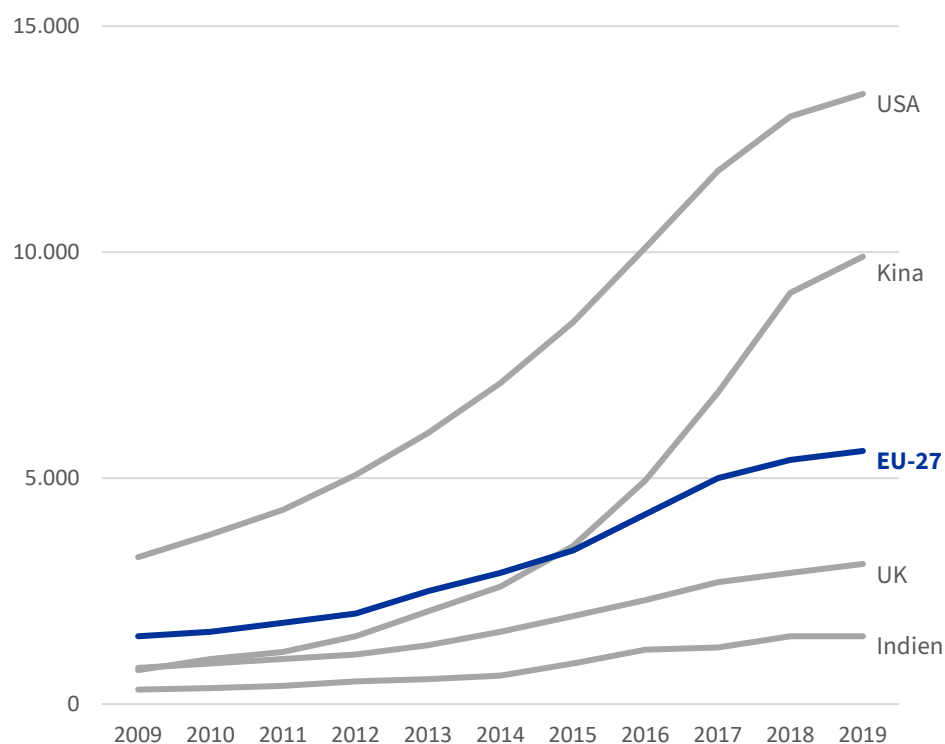


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

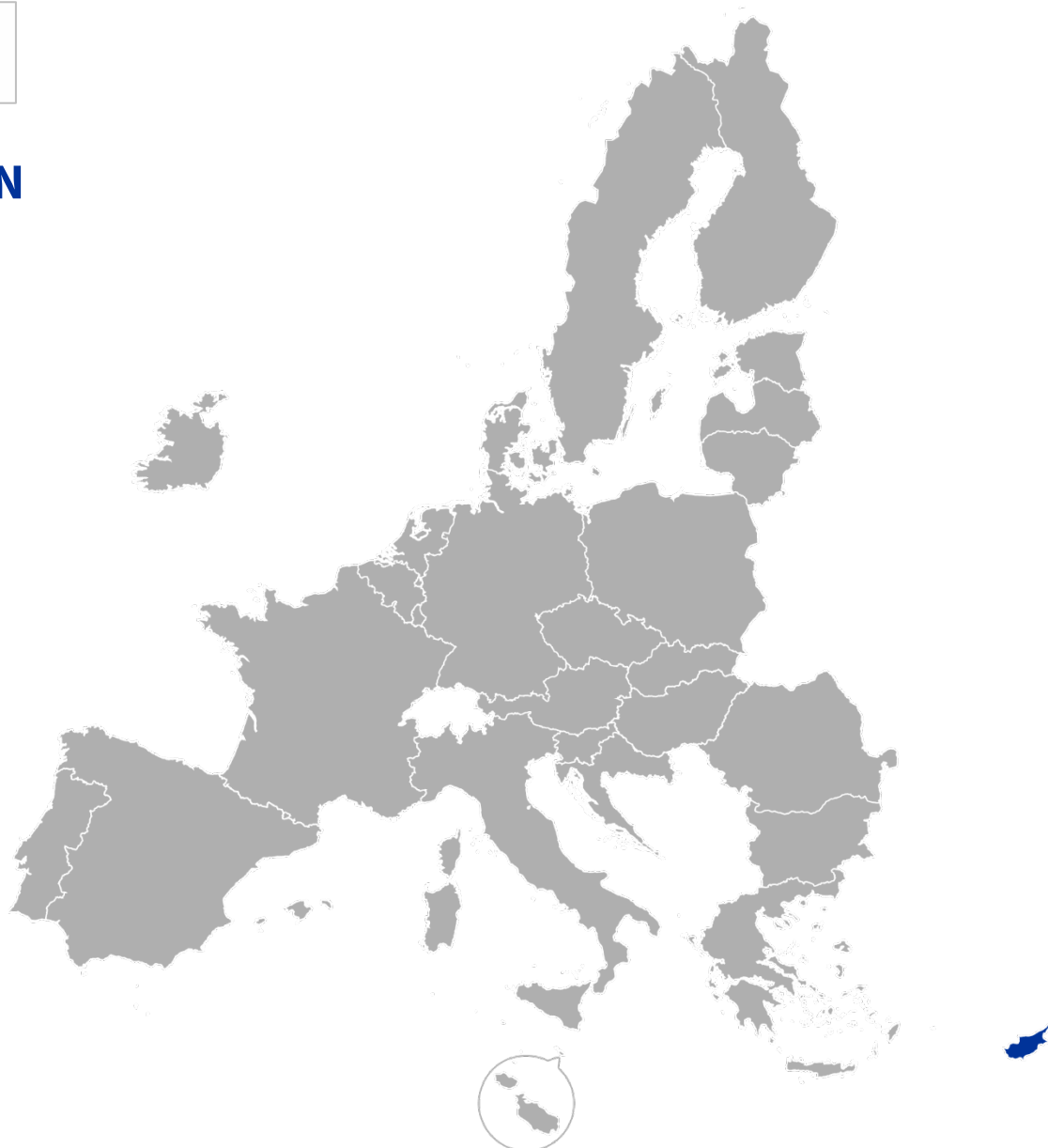
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



CYPERN



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



CYPERN

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Belgien, Grækenland, Spanien og Tyskland**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.

Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____.

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____.

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____.

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____.

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Din nationale regering prioriterer politiske tiltag, der fremmer forskning og innovation og samtidig anvender strenge etiske retningslinjer. Ureguleret AI vil uundgåeligt udgøre en sikkerhedstrussel mod personer og virksomheder. I 2010 udløste AI-drevne handelsalgoritmer f.eks. et børskrak, som borteliminerede 1 billion \$ i aktieværdi på få minutter.
- Efter din mening skal etisk AI-udvikling prioritere gennemsigtighed (som gør AI forståelig), ansvarlighed (som sikrer, at mennesker forbliver ansvarlige for AI-beslutninger) og kontrolmekanismer (som forhindrer AI i at fungere uden kontrol). Du er bekymret for, at alt for fleksible regler kan underminere disse principper.
- Ved at handle resolut med en forsigtighedstesttilgang, som omfatter klare regler, der kan håndhæves, kan EU positionere sig som en ansvarlig og fremsynet leder inden for AI-innovation.
- Du støtter en kapabilitetsbaseret risikoklassificering: Ethvert AI-system med potentiale for misbrug, bias eller skadelig indvirkning bør betragtes som højrisiko – ikke fordi AI i sig selv er farlig, men på grund af de måder, den kan bruges på. Du vil dog være **åben over for at drøfte en formålsbaseret klassificering, hvis dette tages op til revision om et par år.**

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Du er overbevist om, at etisk AI-udvikling skal gælde over hele linjen – ingen sektor og ingen udvikler bør undtages fra EU's AI-forordning. Retshåndhævelse, militær, forsvar og national sikkerhed er netop de områder, hvor risikoen for misbrug, bias og krænkelse af rettigheder er størst. Disse sektorer bør ikke behandles som undtagelser, men snarere omfattes af de højeste standarder.
- EU ønsker at være nummer et i verden inden for troværdig og ansvarlig AI. Derfor argumenterer du for, at der foregår med et godt eksempel. Hvis der ses gennem fingre med visse applikationer, vil det svække troværdigheden og virkningen af EU's AI-forordning. Ethiske principper skal være retningsgivende for alle anvendelser af AI – især dem, der har det største potentiale til at påvirke menneskers liv og friheder.
- Artikel 2 er vigtigere for dig end artikel 1. Hvis du ikke kan overbevise andre om dine pointer vedrørende artikel 1, skal du fokusere på at sikre, at dine vigtigste krav vedrørende artikel 2 bliver taget alvorligt.
- Dine bemærkninger:

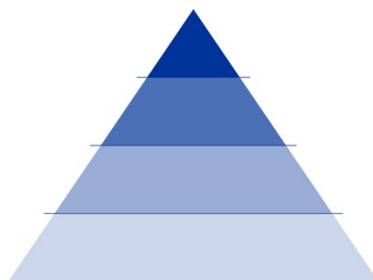
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern	kapabilitetsbaseret	forsigtigheds-	revision om to år	uden undtagelse	
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

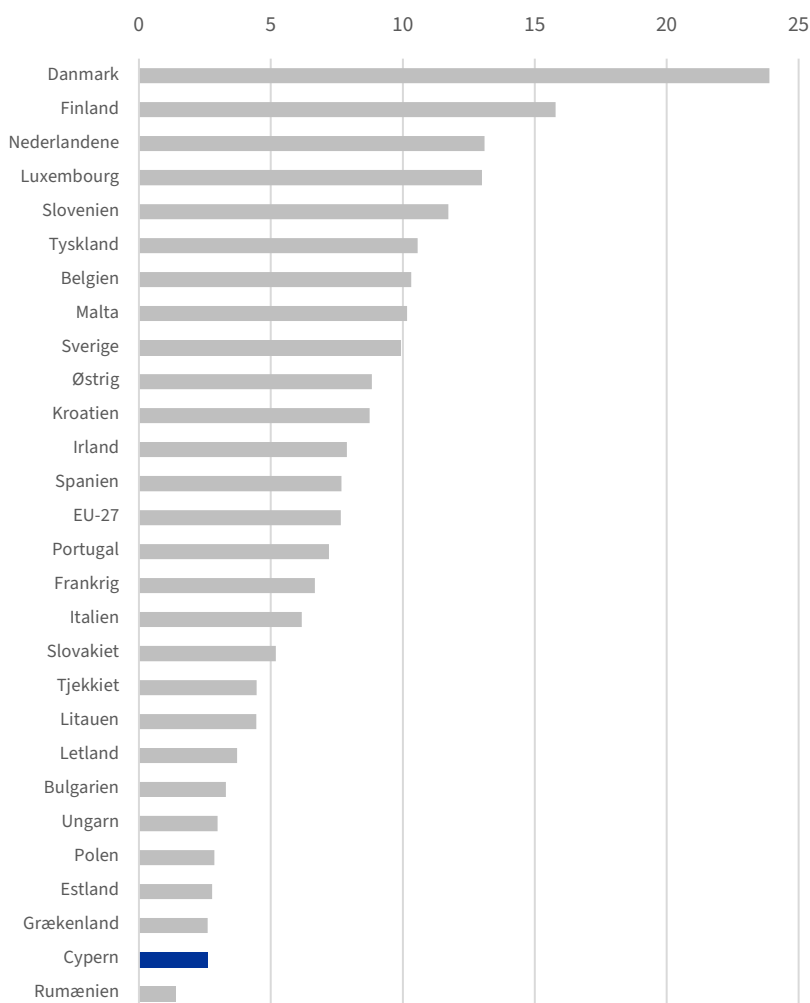
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

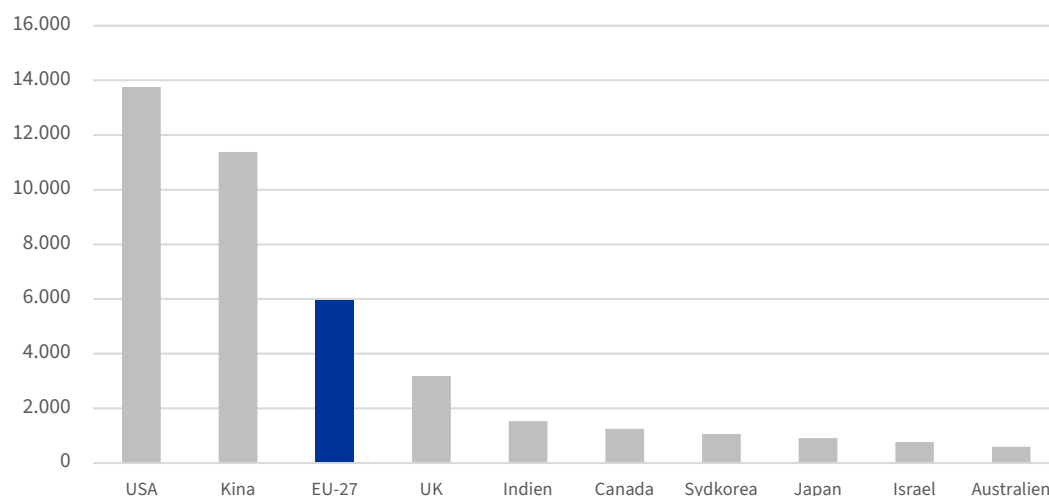
	Den Europæiske Union	Cypern
Befolkning	447 695 350	949 084
BNP pr. indbygger i €	38 150	32 720
Arbejdsløshedsprocent	6,1 %	5,8 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	4,7 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	25 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

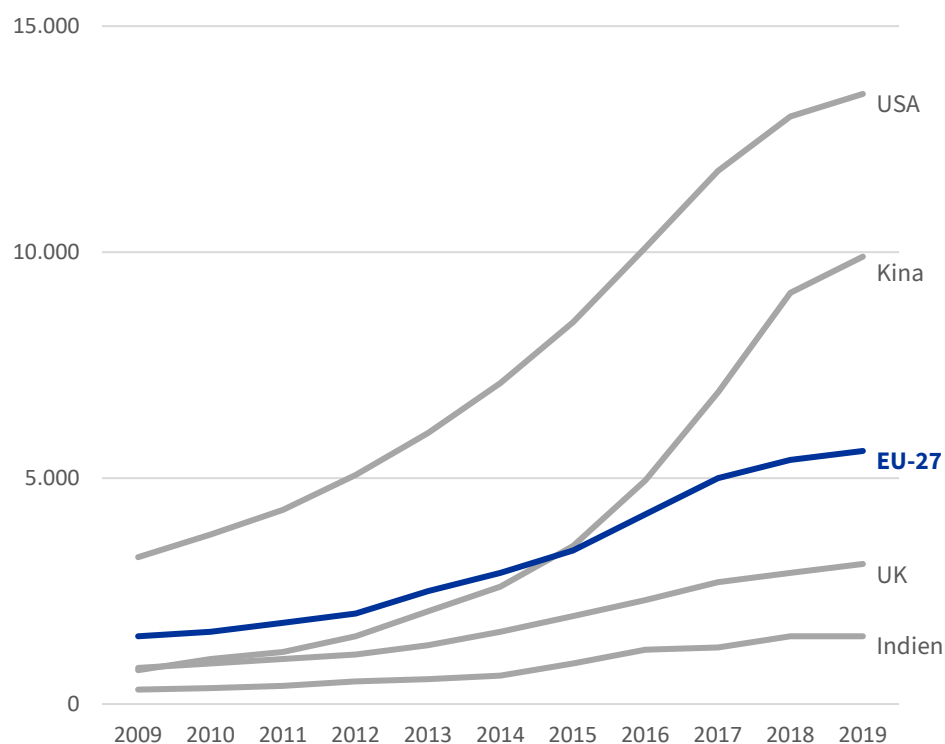


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

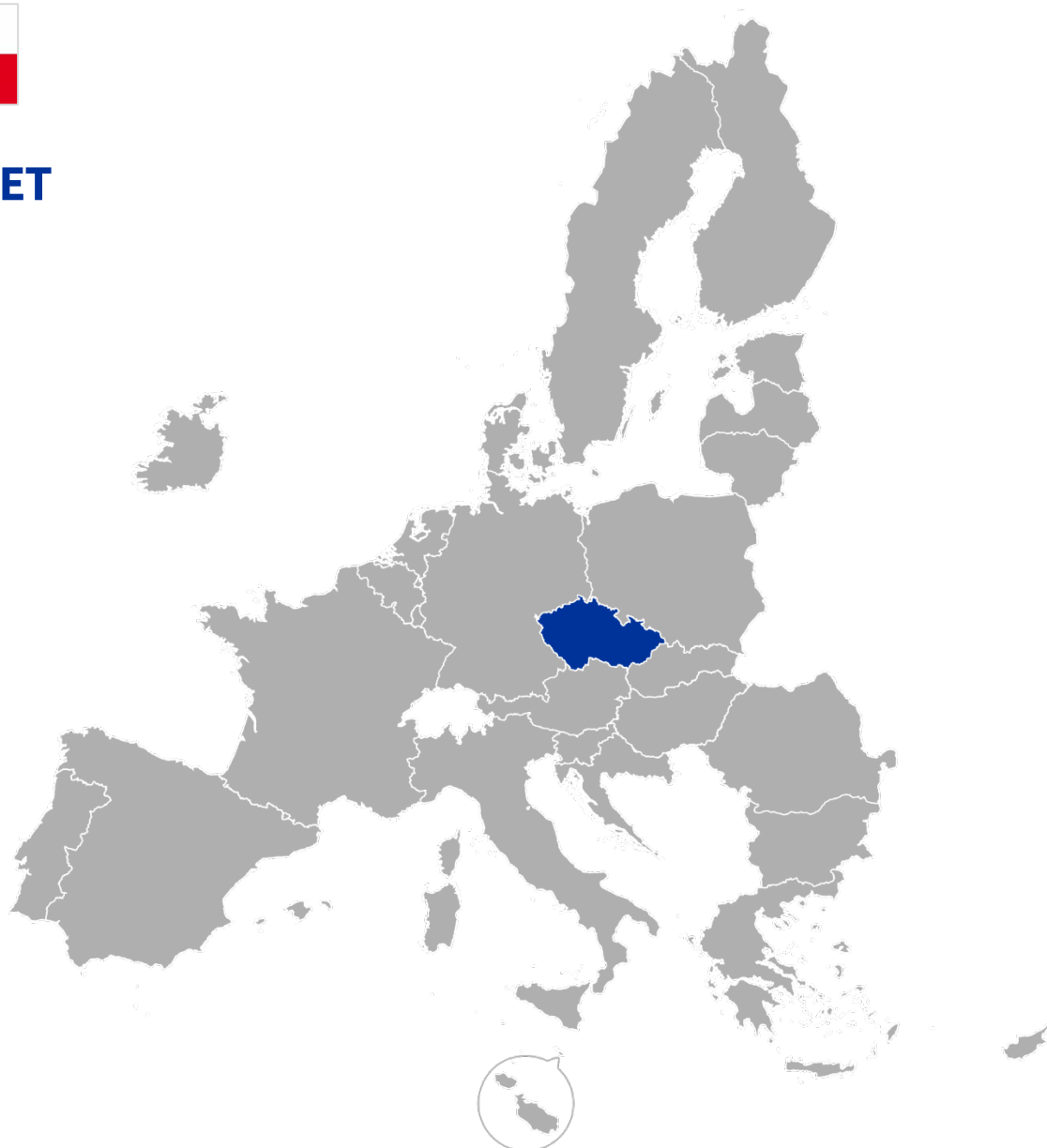
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltageres læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



TJEKKIET



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Estland, Letland, Litauen og Ungarn**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

_____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

_____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

_____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... **kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.**



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... **fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.**



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Din nationale regering ønsker at fremme AI-innovation. Selv om EU er stærkt med hensyn til forskning, ledes det globale marked af USA og Kina, som fastsætter spillets regler. Andre lande som Indien, Canada og UK gør også hurtige fremskridt.
- På denne baggrund mener du, at den logiske konklusion er, at EU fastlægger en afbalanceret lovgivningsmæssig ramme – som ikke er alt for rigid – for at fremme AI-innovation og positionere sig som en global AI-leder. EU's AI-forordning kan ændre spillets regler, hvis den vinder både investorernes og udviklernes tillid og støtte.
- Du støtter en kapabilitetsbaseret tilgang, da en formålsbaseret tilgang kan være i strid med eksisterende nationale sektorpolitikker og EU-sektorpolitikker og dermed skabe usikkerhed. Men det ville være for rigtigt at klassificere al AI til almen brug (GPAI) som højrisiko. Mange GPAI-systemer såsom AI, der anvendes til kundeservice eller oversættelse i realtid, er lav risiko. Overregulering kan lægge en dæmper på innovation og investeringer og styrke idéen om, at "EU har de strengeste AI-regler", hvilket i sidste ende svækker Europas konkurrenceevne.
- Hvis det er vanskeligt at finde et kompromis i dag, **er du åben over for, at der tilføjes en klausul om at revidere forordningen om fem år**, så det sikres, at den forbliver relevant i et AI-landskab i hastig udvikling.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Du er overbevist om, at etisk AI-udvikling skal gælde over hele linjen – ingen sektor og ingen udvikler bør undtages fra EU's AI-forordning. Retshåndhævelse, militær, forsvar og national sikkerhed er netop de områder, hvor risikoen for misbrug, bias og krænkelse af rettigheder er størst. Disse sektorer bør ikke behandles som undtagelser, men snarere omfattes af de højeste standarder.
- Hvis EU ønsker at være nummer et i verden inden for troværdig og ansvarlig AI, skal der foregå med et godt eksempel. Hvis der ses gennem fingre med visse applikationer, vil det svække troværdigheden og virkningen af EU's AI-forordning. Etiske principper skal være retningsgivende for alle anvendelser af AI – især dem, der har det største potentiale til at påvirke menneskers liv og friheder.
- Dine bemærkninger:

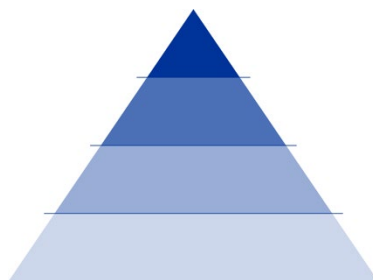
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet	kapabilitetsbaseret	fleksibel	revision om fem år	uden undtagelse	
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

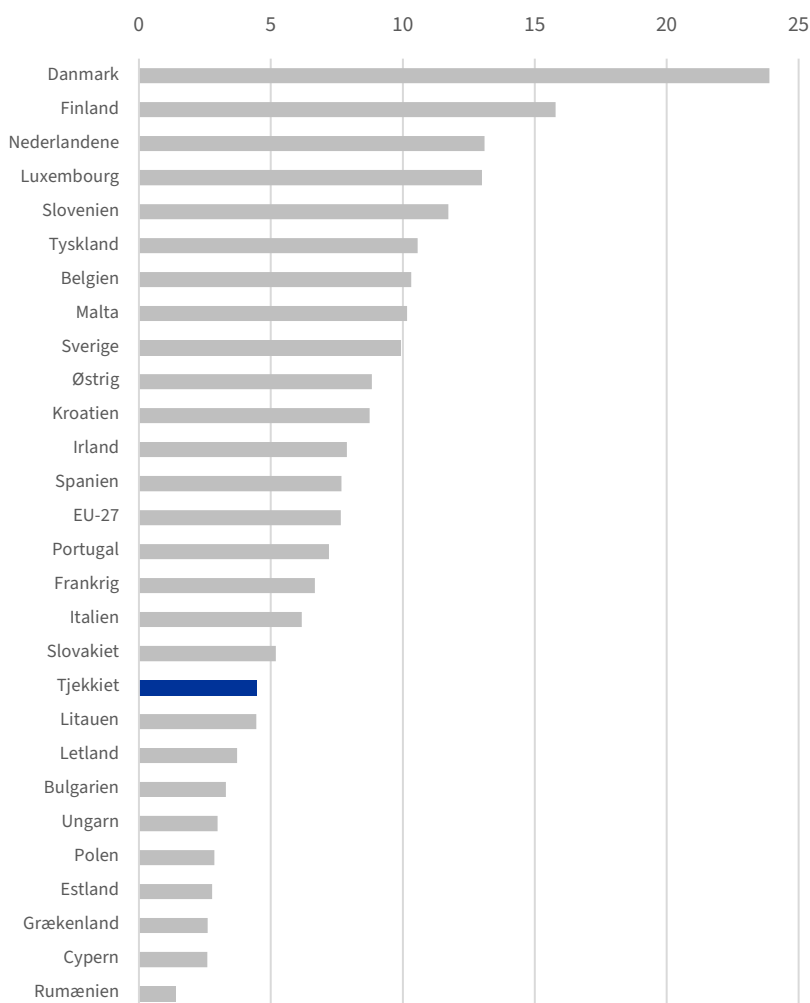
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

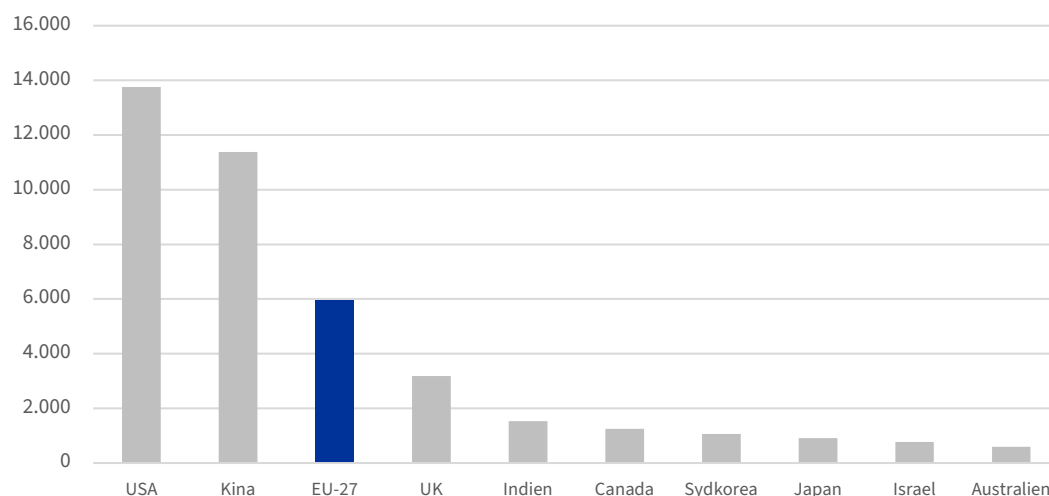
	Den Europæiske Union	Tjekkiet
Befolkning	447 695 350	10 827 529
BNP pr. indbygger i €	38 150	29 180
Arbejdsløshedsprocent	6,1 %	2,6 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	5,9 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	35,5 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

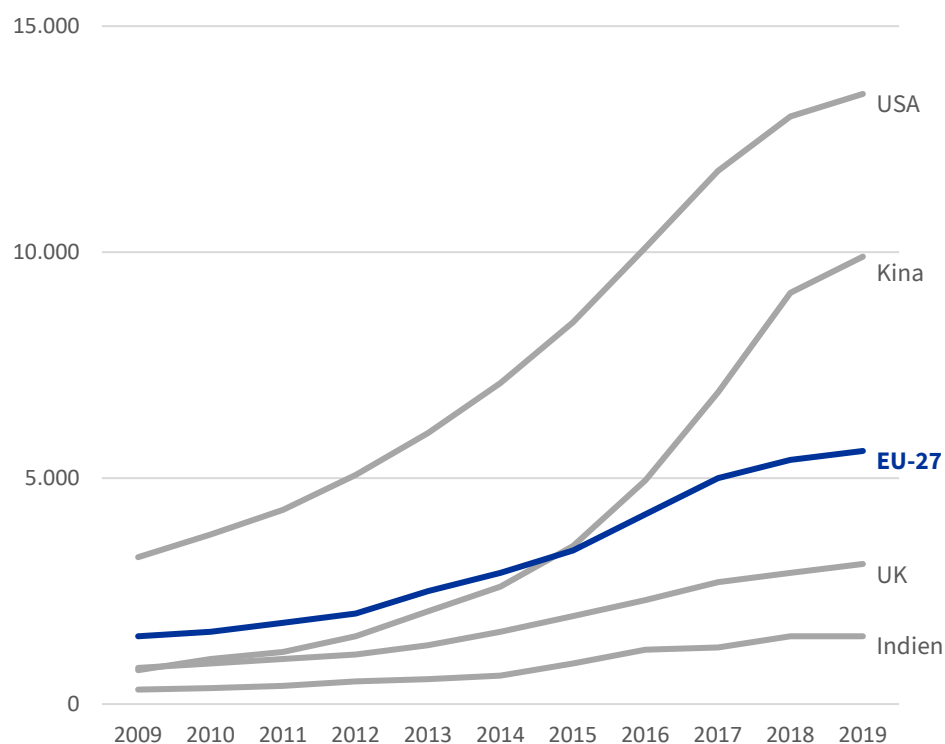


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

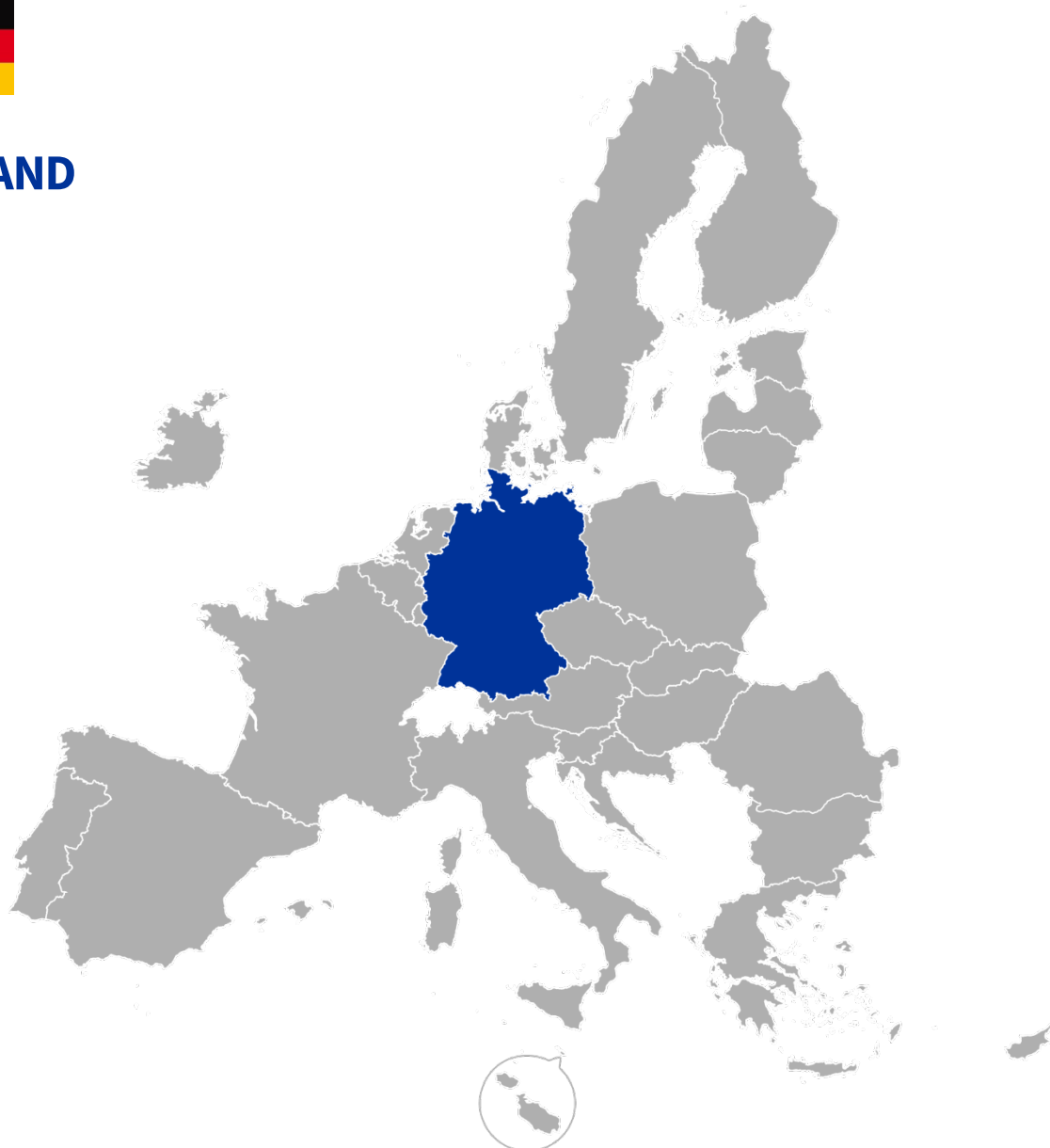
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



TYSKLAND



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



TYSKLAND

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Belgien, Cypern, Grækenland og Spanien**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Fremme af grundlæggende rettigheder er din topprioritet, og du anerkender det presserende behov for grænseoverskridende samarbejde for at imødegå både nationale og tværnationale trusler, herunder trusler fra digitale teknologier. Du har været vidne til, at AI kan give næring til misinformation og deepfake-kampagner, hvilket udgør en alvorlig risiko for ytringsfriheden og de demokratiske institutioner.
- Selv om AI-udviklingen skrider hurtigt frem og kommer til at spille en stadig større rolle i samfundet, mener du, at den skal designes til gavn for mennesker og miljø – ikke forårsage skade.
- Du går ind for strenge testkrav til højrisikomodeller, da et stort problem i AI-systemer er bias. Når AI trænes ved hjælp af datasæt, der ikke er forskelligartede eller ikke er repræsentative, kan den give diskriminerende resultater – som det var tilfældet med Amazons ansættelsesværktøj, som blev skrottet, fordi det favoriserede mandlige ansøgere på grundlag af tidligere ansættelsesmønstre. Dit endemål i dag er at arbejde på en ramme, der fjerner selv den mindste risiko for AI-genereret skade såsom misinformation. Du ved, at det er et ambitiøst mål – men du mener, at det er nødvendigt.
- Du støtter en kapabilitetsbaseret risikoklassificering: Ethvert AI-system med potentiale for misbrug eller skadelig indvirkning bør betragtes som højrisiko – ikke fordi AI i sig selv er farlig, men på grund af de måder, den kan bruges på.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Du er meget bekymret over at høre, at nogle lande ønsker at udelukke retshåndhævelse, militær, forsvar eller national sikkerhed fra anvendelsesområdet for EU's AI-forordning. Det er netop de områder, hvor AI kan forårsage alvorlig skade – gennem forskelsbehandling, falske positive og bias – og hvor der er størst behov for stærke sikkerhedsforanstaltninger. F.eks. har biaspræget ansigtsgenkendelse i forbindelse med politiarbejde allerede ført til uretmæssige anholdelser – af et uforholdsmæssigt stort antal racialiserede unge mænd – og alvorlige krænkelser af individuelle rettigheder.
- Du forstår, at mange lande oplever store effektivitetsgevinster ved at anvende AI i disse sektorer, og det respekterer du. Men der kan ikke tages let på menneskerettighederne. Du har fulgt AI-udviklingen i USA og Kina tæt, og du er dybt bekymret over manglen på stærke beskyttelsesmekanismer i deres systemer.
- Hele formålet med EU's AI-forordning er at fremme ansvarlig, menneskecentreret AI. Hvis visse sektorer udelukkes, vil det svække dette mål og underminere EU's ambitioner om at fastsætte en global standard for troværdig AI.
- Dine bemærkninger:

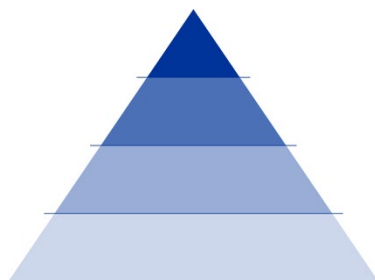
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko-AI-systemer**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland	kapabilitetsbaseret	forsigtigheds-		uden undtagelse	
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

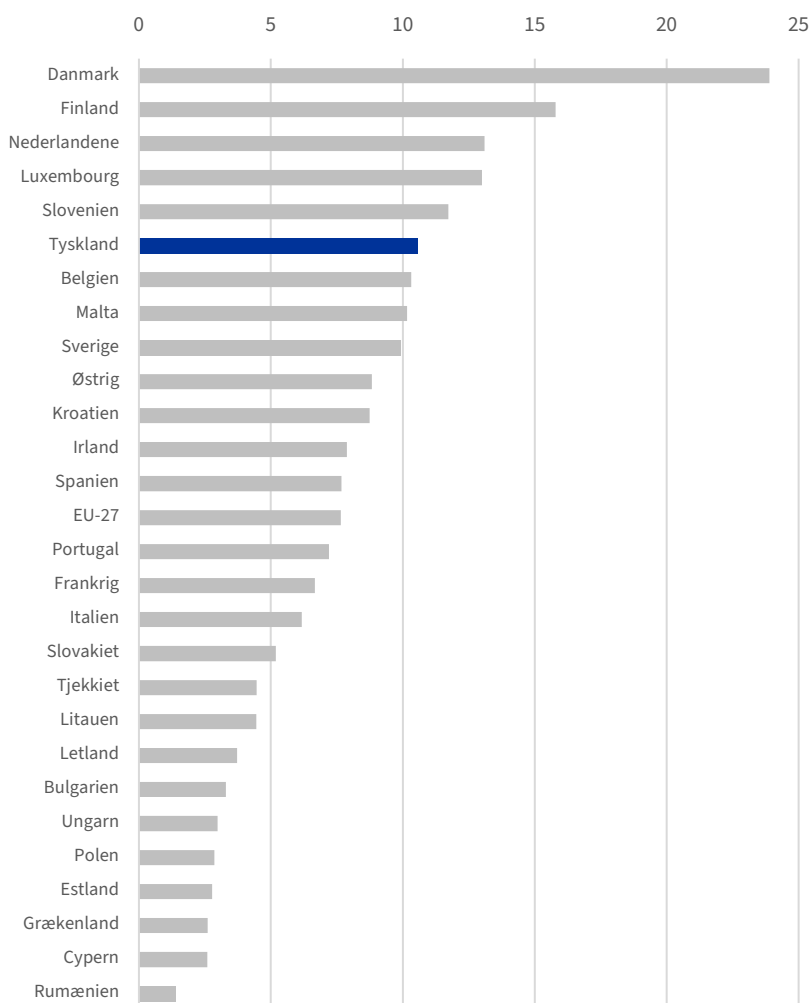
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

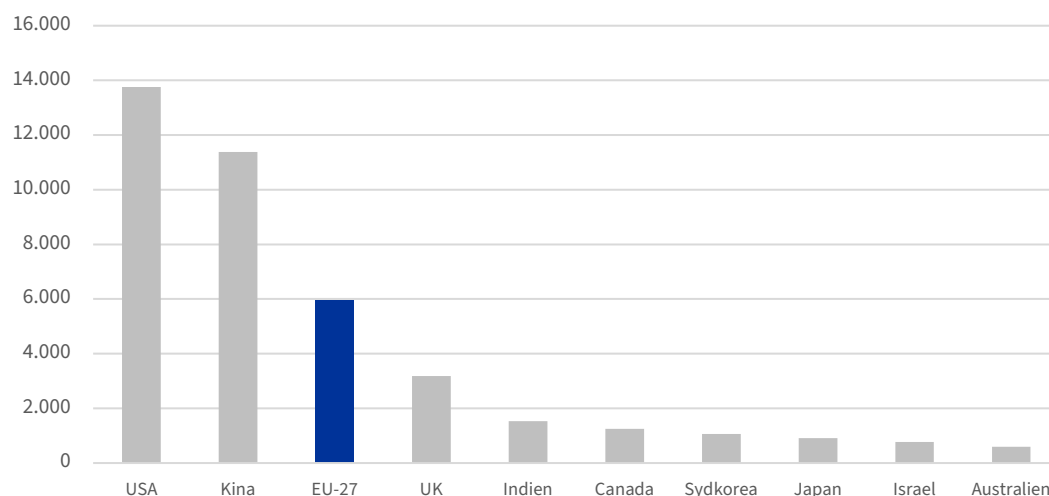
	Den Europæiske Union	Tyskland
Befolkning	447 695 350	83 118 501
BNP pr. indbygger i €	38 150	49 520
Arbejdsløshedsprocent	6,1 %	3,1 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	11,6 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	19,8 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

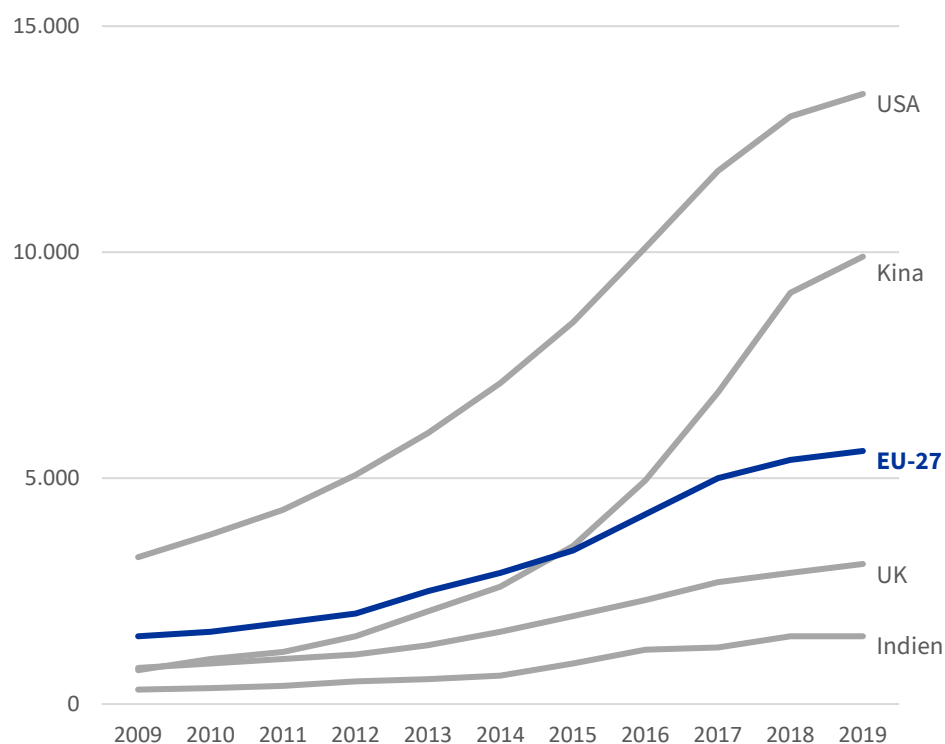


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

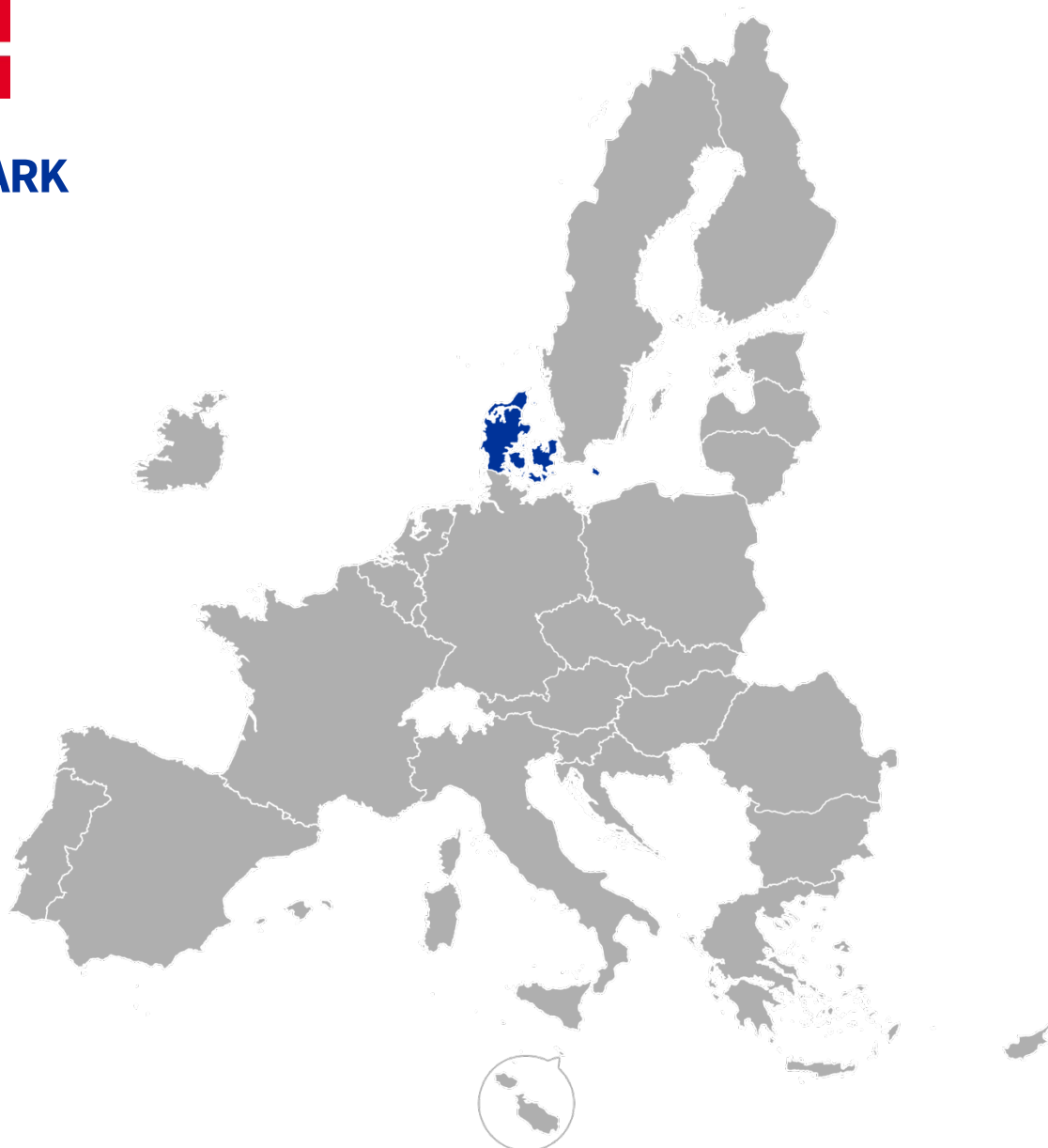
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltageres læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



DANMARK



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Frankrig, Polen, Rumænien og Østrig**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Du glæder dig over Kommissionens forslag, og du støtter den formålsbaserede tilgang som et konstruktivt første skridt til at imødegå udfordringerne i forbindelse med AI-udvikling. Du ønsker dog at inkludere uddannelsessektoren. Når skoler bruger AI til at overvåge adfærd, bedømme opgaver eller forudsige præstationer, kan eleverne føle, at de hele tiden bliver overvåget eller bedømt af en maskine. Det kan skabe et pres for altid at præstere på en bestemt måde, selv om eleverne bare har en dårlig dag. I modsætning til en menneskelig lærer mangler AI ofte empati eller forståelse for personlige udfordringer.
- AI muliggør nye måder at træffe beslutninger på, og det giver anledning til etiske betænkeligheder. Derfor går du ind for en ramme for forsigtighedstest, der fremmer en ansvarskultur, som gør det lettere at spore, hvor det er gået galt, når der opstår problemer.
- Desuden skal en fælles EU-ramme anerkende EU's sproglige mangfoldighed. De fleste store AI-modeller trænes ved hjælp af data fra meget udbredte sprog såsom engelsk, fransk eller tysk, hvilket kan føre til dårlige præstationer på mindre udbredte sprog som dansk.
- **Hvis EU beslutter at afsætte midler til AI-udvikling, vil det være rimeligt at prioritere støtte til mindre lande og mindre udbredte sprog.** Ellers kan AI-udvikling øge ulighederne og levere tjenester af lavere kvalitet på mindre udbredte sprog.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Det er afgørende for dig, at spørgsmål vedrørende national sikkerhed, militær eller forsvar ikke bliver omfattet af AI-forordningen. Medlemsstaterne er fast besluttet på at bevare kontrollen over beslutninger på disse områder, som under alle omstændigheder typisk falder uden for EU's jurisdiktion. Dette ses også af, at EU (endnu) ikke har en fælles EU-hær.
- Du argumenterer desuden for, at AI-systemer, der anvendes til indsamling af efterretninger, cyberforsvar eller slagmarksscener, ikke kan være underlagt de samme oplysningskrav eller reguleringsmæssige krav som civile systemer, fordi dette kunne kompromittere operationel hemmeligholdelse, nationale sikkerhedsinteresser og den strategiske fordel, som sådanne systemer er designet til at give.
- Hvis dine krav ikke kan opfyldes fuldt ud i dagens drøftelser, er du åben over for et kompromis, hvor formål, der vedrører militær og national sikkerhed, ikke undtages fra forordningen, men snarere underlægges de krav, der er anført under den fleksible testtilgang.
- Dine bemærkninger:

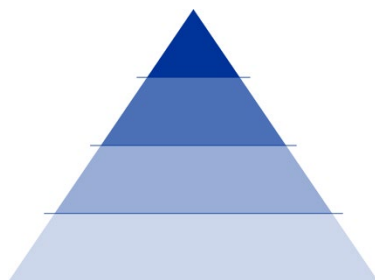
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyling for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark	formålsbaseret	forsigtigheds-	finansiel støtte	militær/forsvar + national sikkerhed	
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

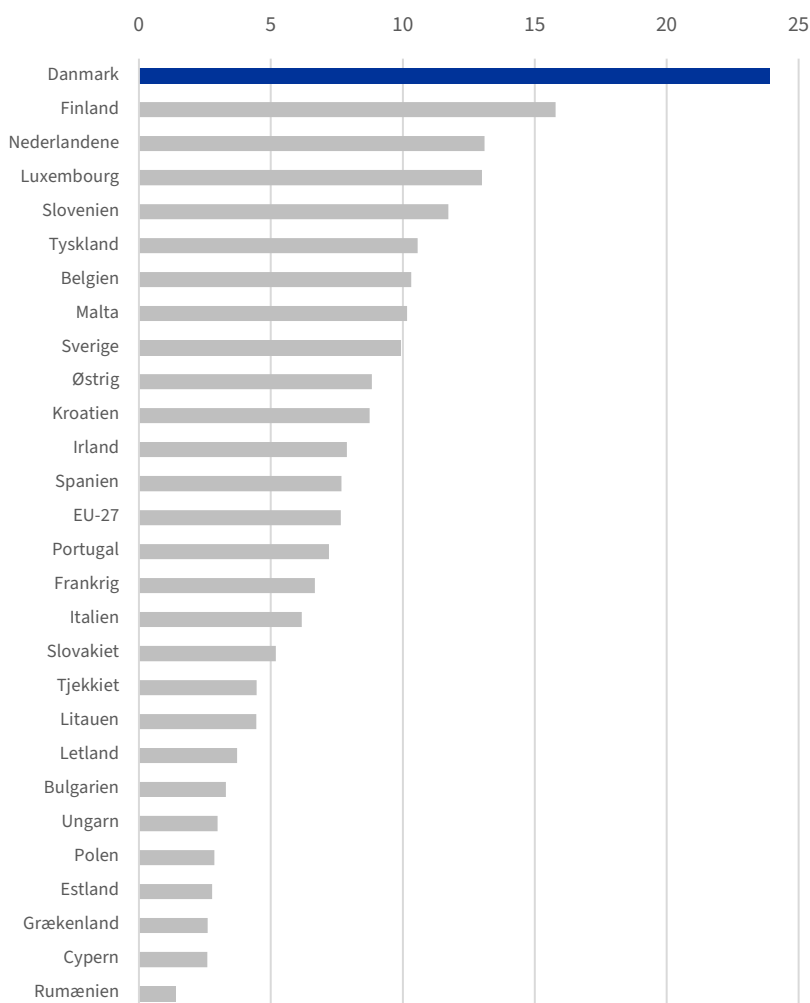
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

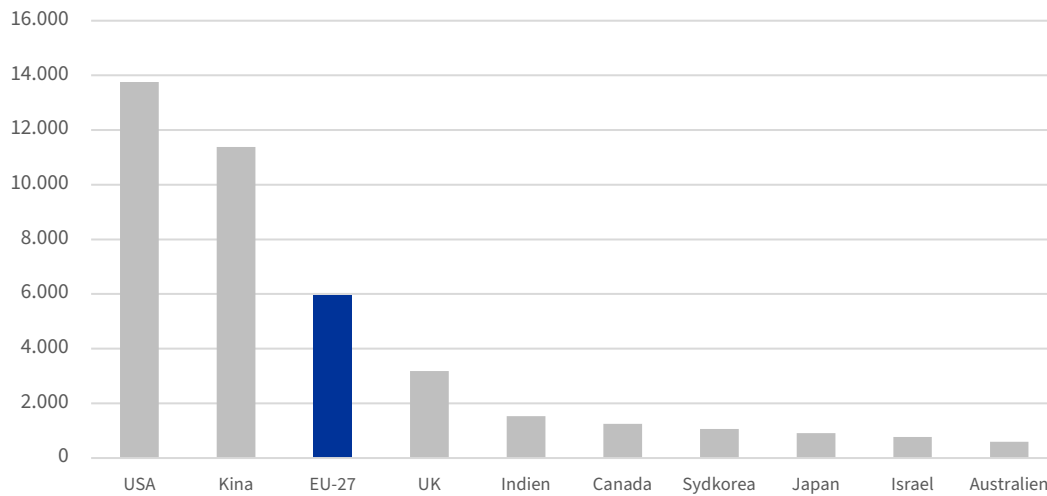
	Den Europæiske Union	Danmark
Befolkning	447 695 350	5 932 654
BNP pr. indbygger i €	38 150	63 290
Arbejdsløshedsprocent	6,1 %	5,1 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	15,2 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	39,4 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

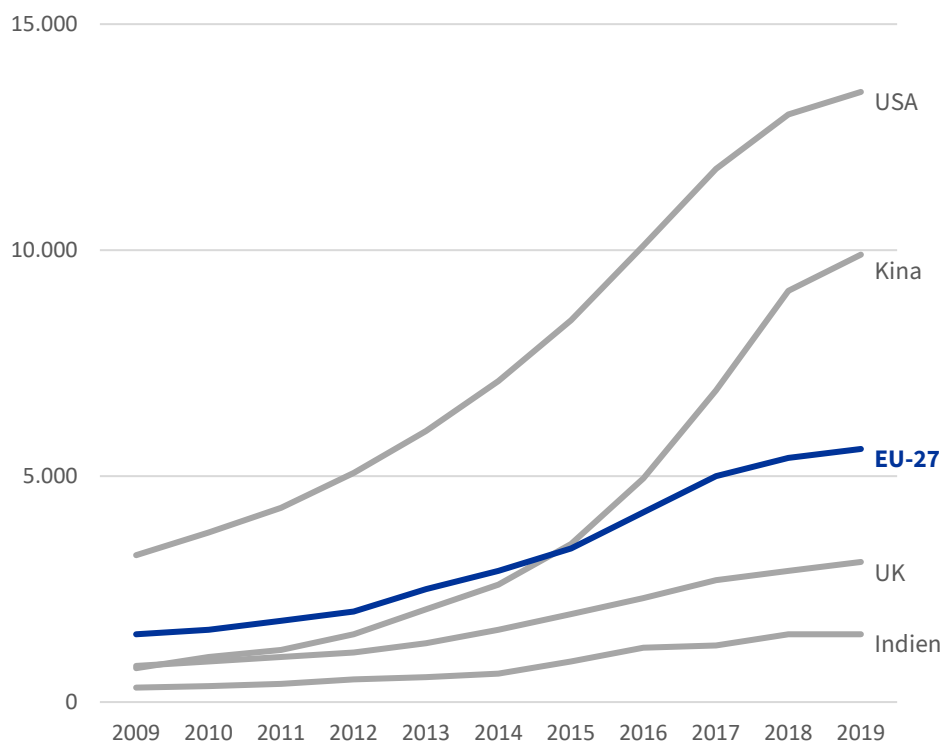


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

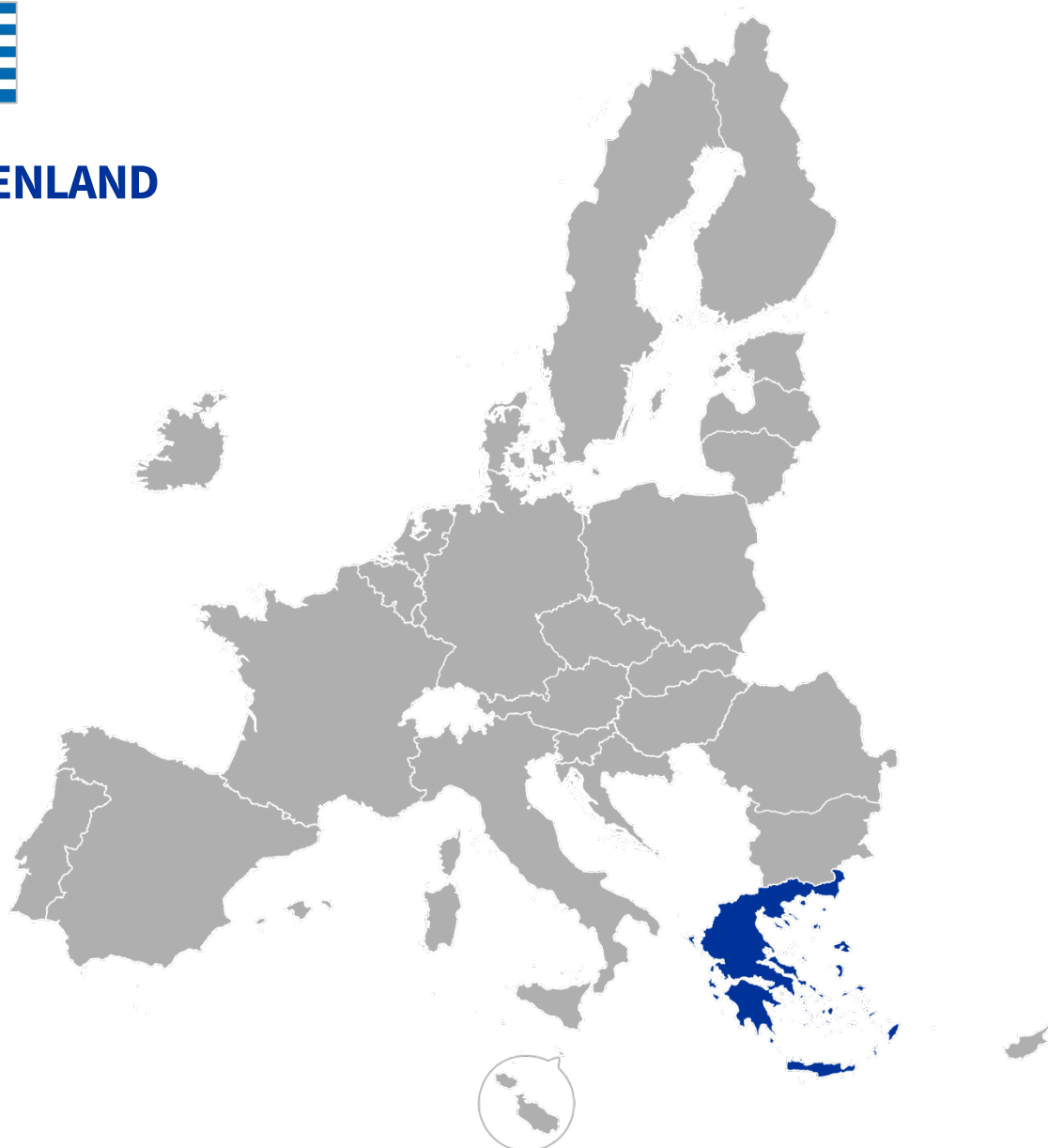
© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print	ISBN 978-92-850-0498-9	doi: 10.2860/0371762	QC-01-25-006-DA-C
PDF	ISBN 978-92-850-0497-2	doi: 10.2860/0996198	QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltageres læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



GRÆKENLAND



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



GRÆKENLAND

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Belgien, Cypern, Spanien og Tyskland**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.

Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... **kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.**



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... **forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).**

Dine argumenter

- Dit land sigter mod at skabe en teknologibaseret fremtid, der fremmer innovation og udvikling til gavn for alle og er baseret på Europas kerneværdier. For at forhindre, at AI har en negativ indvirkning på de demokratiske processer, støtter du den kapabilitetsbaserede tilgang, hvor ethvert AI-system anses for at være højrisiko, hvis det har kapabilitet til at "manipulere menneskelig adfærd" – også i forbindelse med valg.
- Du går stærkt ind for forsigtighedstesttilgangen. Det vigtigste element er efter din mening regulermæssige sandkasser, som tvinger udviklerne til at teste AI-systemer med små brugergrupper under nationalt tilsyn. Dette garanterer sikker eksperimentation og mindsker potentiel skade.
- Du er især bekymret over de potentielle farer ved biaspræget AI: Undersøgelser har vist, at store ansigtsgenkendelsessystemer har haft meget højere fejlprocenter, når de identificerer personer med mørkere hudfarve og kvinder, især sorte kvinder. Den slags fejl er efterfølgende blevet sporet tilbage til træningsdatasæt, der er stærkt skævvredne i retning af hvide mandlige ansigter, hvilket har ført til hyppig fejlidentifikation.
- Det kan være en udfordring at nå til enighed i dag. **Om nødvendigt kan det foreslås, at der som et kompromis tilføjes en klausul om revision om tre år.** Denne revision bør omfatte synspunkter fra civilsamfundet, den private sektor, industrien og de nationale regeringer.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Det er afgørende for dig, at spørgsmål vedrørende national sikkerhed, militær eller forsvar ikke bliver omfattet af AI-forordningen. Medlemsstaterne er fast besluttet på at bevare kontrollen over beslutninger på disse områder, som under alle omstændigheder typisk falder uden for EU's jurisdiktion. Som Schengengrænseland spiller Grækenland en afgørende rolle i forvaltningen af EU's ydre grænser. Men de nuværende systemer er forældede og ineffektive, hvilket lægger en tung operationel og humanitær byrde på de nationale myndigheder.
- Grækenland betragter derfor AI-modeller som et middel til at modernisere og strømline grænseoperationer – og dermed forbedre identifikations-, registrerings- og sagsforvaltningsprocedurerne og samtidig støtte hurtigere og mere præcis beslutningstagning. I betragtning af det presserende behov for at håndtere de mange ankomster har Grækenland brug for en fleksibel og rettidig ibrugtagning af AI-systemer, der kan støtte hurtigere og mere effektive operationer.
- Dine bemærkninger:

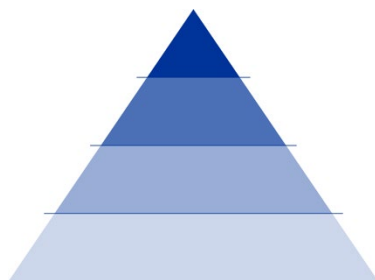
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland	kapabilitetsbaseret	forsigtigheds-	revision om tre år	militær/forsvar + national sikkerhed	
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

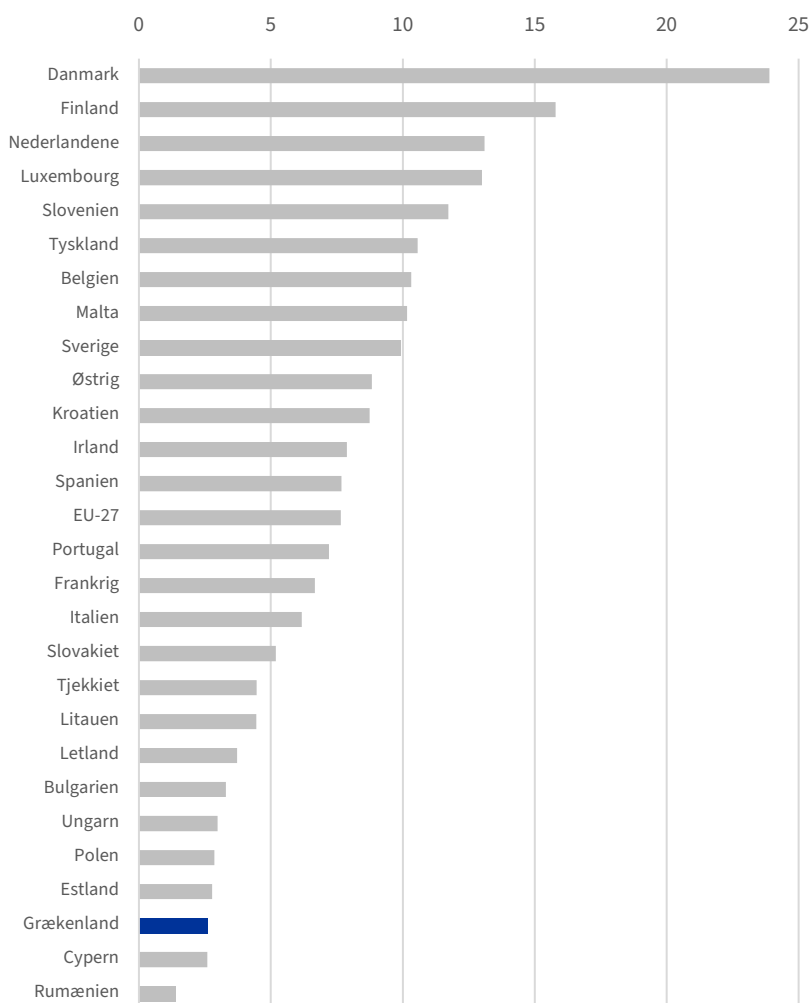
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

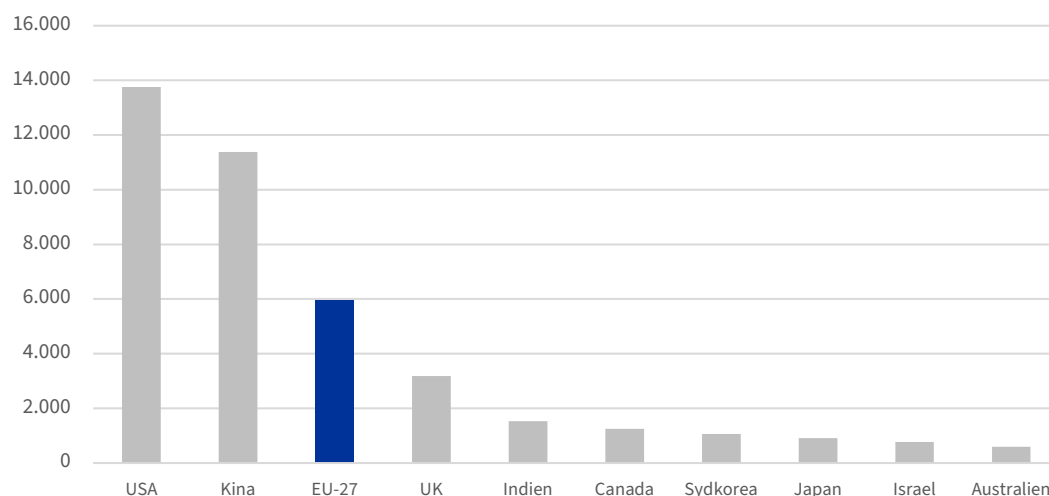
	Den Europæiske Union	Grækenland
Befolkning	447 695 350	10 413 982
BNP pr. indbygger i €	38 150	21 350
Arbejdsløshedsprocent	6,1 %	11,1 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	4 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	20 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

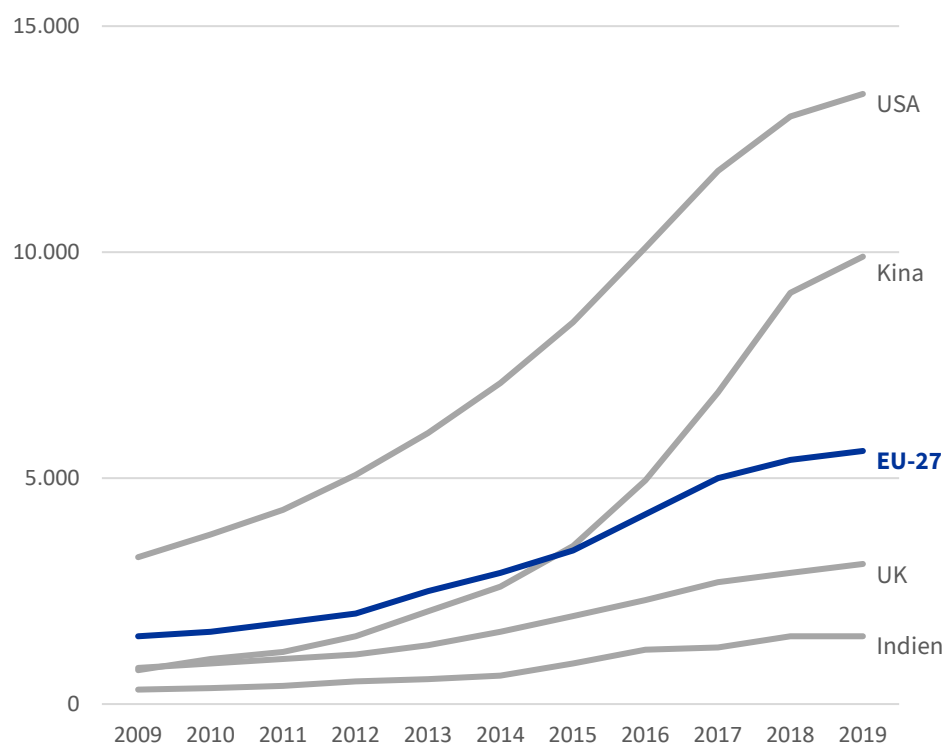


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

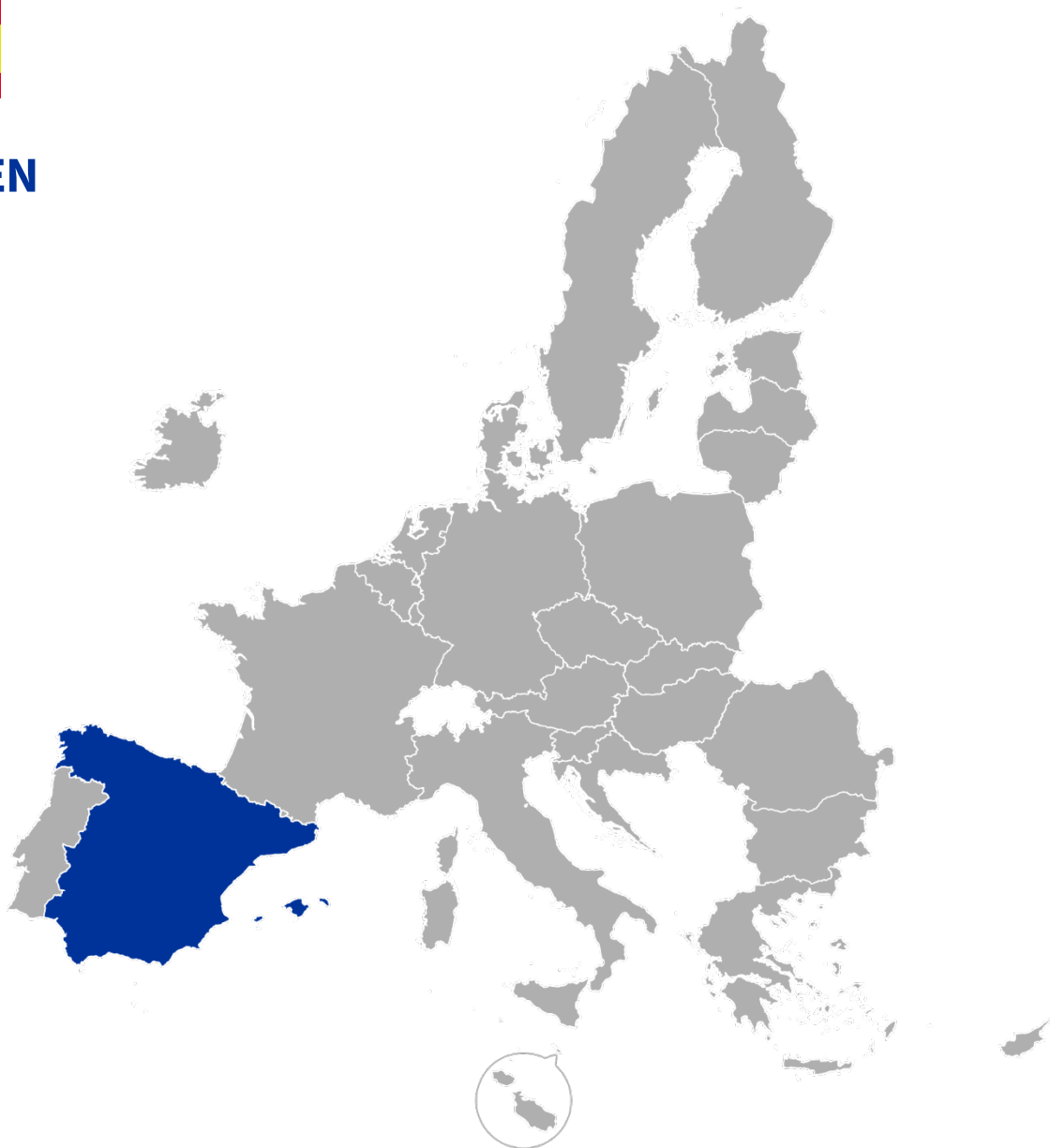
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



SPANIEN



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



SPANIEN

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Belgien, Cypern, Grækenland og Tyskland**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.

Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Dit land prioriterer menneskerettigheder og borgerlige frihedsrettigheder som grundlaget for et retfærdigt og sikkert samfund. Selv om du støtter virksomheders og universiteters nationale investeringer i AI, skal etiske hensyn fortsat være centrale. Dit endemål i dag er at arbejde på en ramme, der fjerner selv den mindste risiko for AI-genereret skade såsom misinformation. Du ved, at det er et ambitiøst mål – men du mener, at det er nødvendigt.
- Du er stærkt imod den formålsbaserede tilgang, som Kommissionen har foreslået. Hvis f.eks. et AI-system, der er designet til at skabe musik i den kreative sektor, også udsender lyde med skadelig styrke – som en lydkanon – kan det stadig forårsage skade, selv om det er mærket som lav risiko. Det er uansvarligt og farligt udelukkende at fokusere på den erklærede hensigt og se bort fra de tekniske kapabiliteter.
- Hvad angår udviklernes forpligtelser, mener du, at det er mest ansvarligt at indføre strenge testkrav, eftersom AI stadig befinder sig i en tidlig fase. Hvis udviklerne med tiden lærer af deres fejl, virksomhederne bliver mere opmærksomme på risici for bias, og teknologien modnes, er du åben over for at drøfte mulige undtagelser. Men indtil videre mener du, at det er afgørende at gå frem med forsigtighed og med stærke sikkerhedsforanstaltninger.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Din regering er meget engageret i borgerlige frihedsrettigheder. Samtidig støtter du en tosporet tilgang til AI-regulering – en tilgang, der sikrer stærk beskyttelse i det civile rum og samtidig giver mere fleksibilitet med hensyn til forsvar og national sikkerhed. National sikkerhed er afgørende for at bevare de friheder, som de borgerlige frihedsrettigheder bygger på, og AI kan spille en central rolle for beskyttelse af samfund. Derfor argumenterer du for, at anvendelser af AI i forbindelse med militær, forsvar og national sikkerhed udelukkes.
- Der bør imidlertid udvikles en særskilt, formålstjenlig lovgivningsmæssig tilgang for disse sektorer for at sikre, at AI udvikles ansvarligt – men indførelse af krav om fuld gennemsigtighed kan bringe operationel hemmeligholdelse i fare, nedbremse ibrugtagning og reducere effektiviteten i dynamiske miljøer eller højrisikomiljøer.
- Eksempler som Spaniens VioGén-system, der anvendes til at vurdere risikoen for vold i hjemmet, afslører farerne ved uigennemsigtig AI med begrænset tilsyn – tragisk nok blev nogle ofre fejlklassificeret som lav risiko og senere dræbt. Disse fejl understreger behovet for ansvarlig AI-udvikling, især inden for sikkerhed, men bør ikke retfærdiggøre overregulering. I stedet bør EU fremme en ansvarskultur, især i forbindelse med udvikling af AI til sikkerhedsapplikationer.
- Dine bemærkninger:

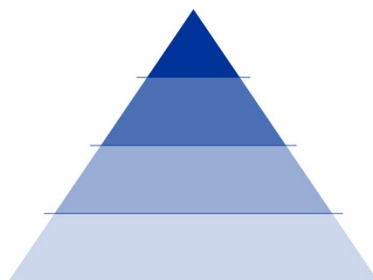
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien	kapabilitetsbaseret	forsigtigheds-		militær/forsvar + national sikkerhed	
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

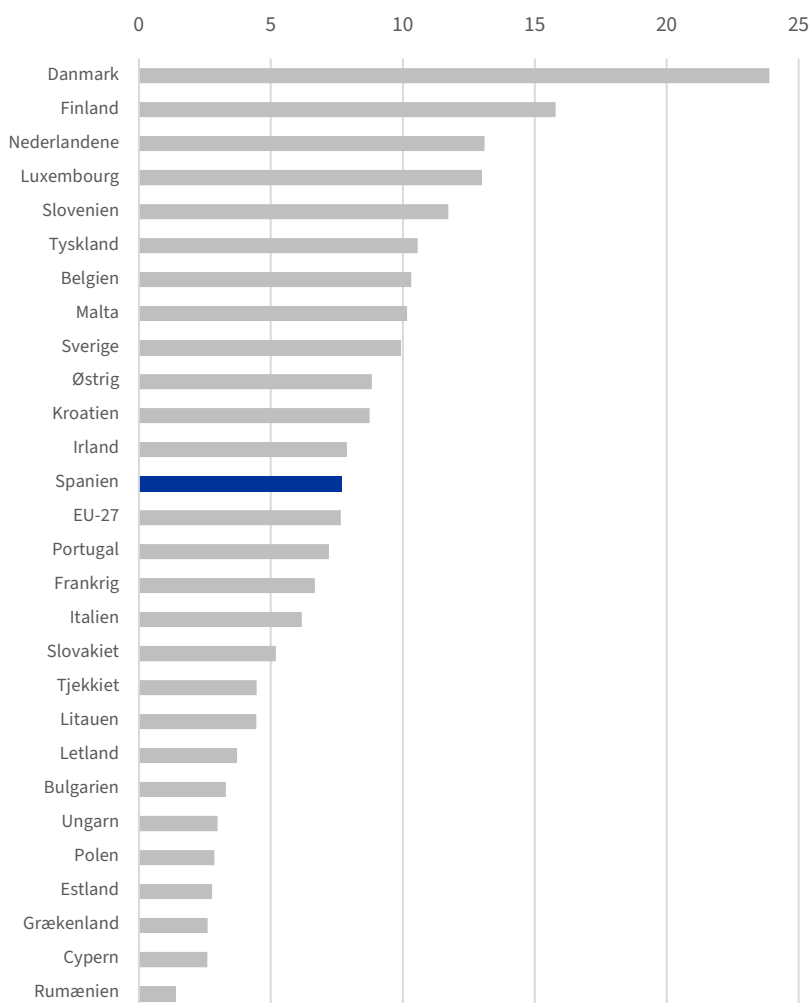
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

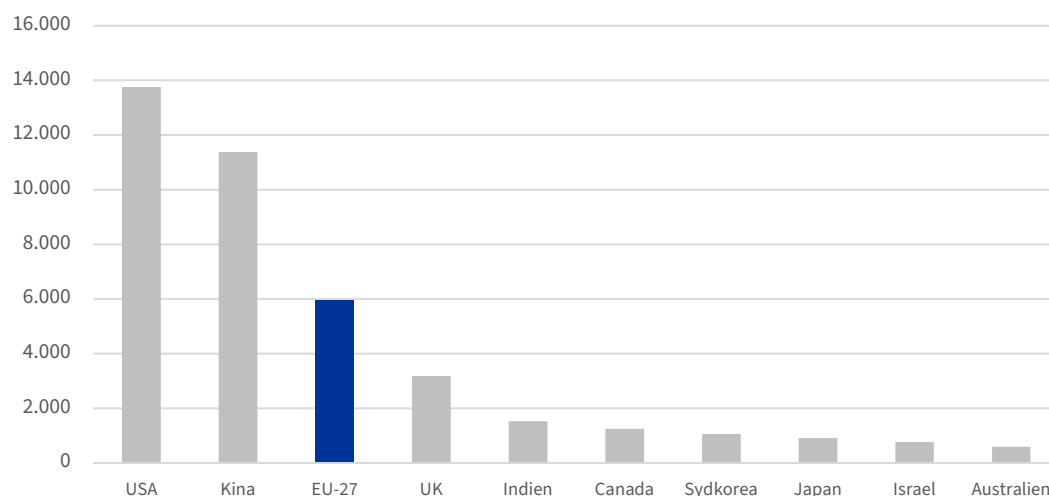
	Den Europæiske Union	Spanien
Befolkning	447 695 350	48 085 361
BNP pr. indbygger i €	38 150	30 970
Arbejdsløshedsprocent	6,1 %	12,2 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	9,2 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	38,7 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

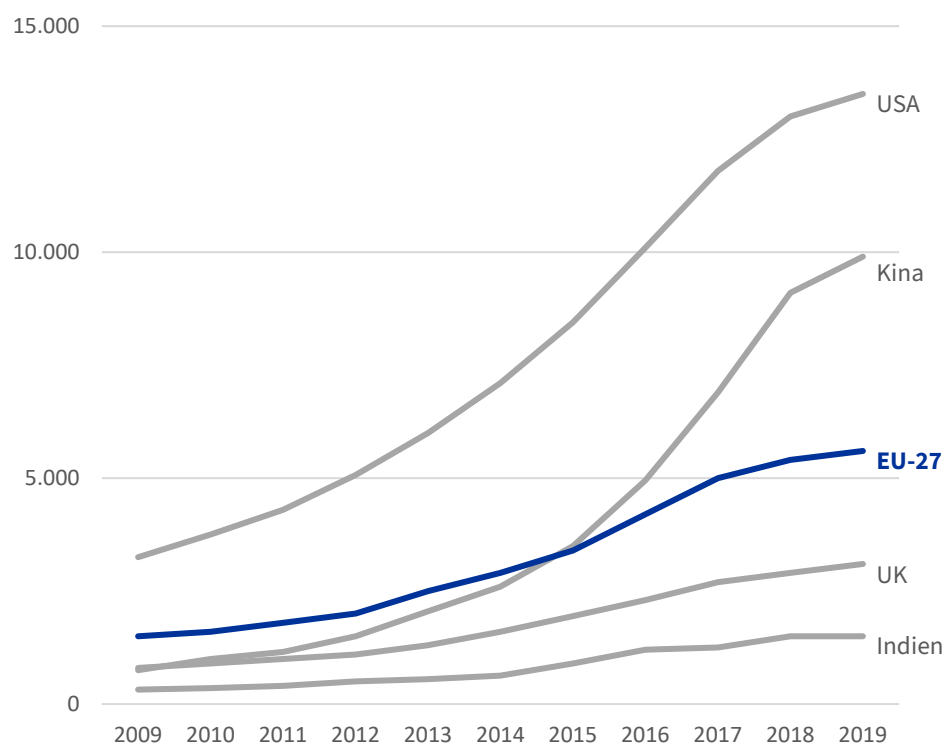


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

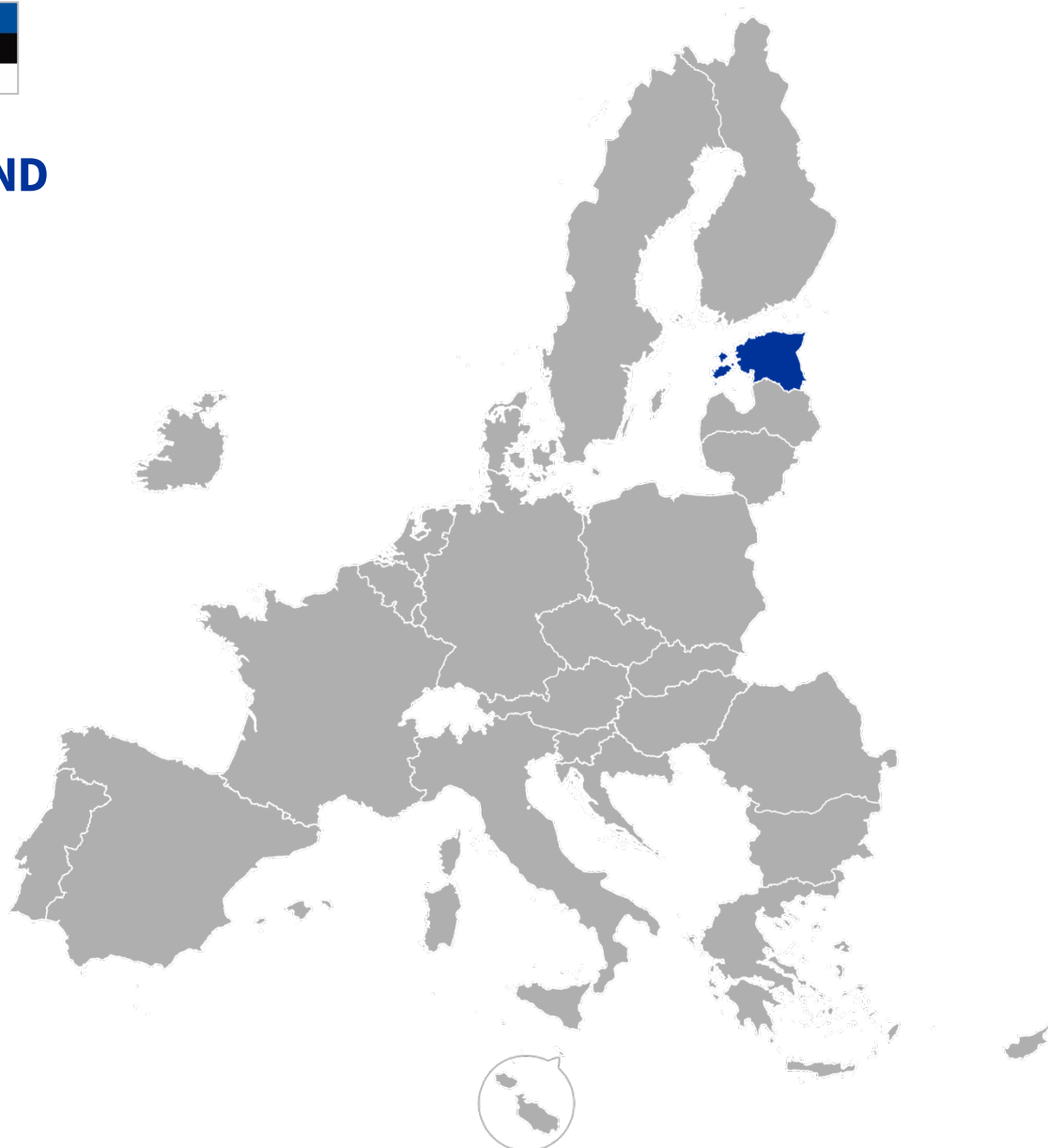
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltageres læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



ESTLAND



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



ESTLAND

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Letland, Litauen, Tjekkiet og Ungarn**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.

Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Dit land støtter en fælles AI-forordning i EU. Som en af EU's mest digitaliserede nationer har Estland en befolkning med et godt kendskab til digitale værktøjer – i modsætning til i nogle lande, hvor mange borgere (navnlig i de ældre aldersgrupper) er mindre fortrolige med dem og mere sårbare over for potentielle risici.
- Du mener, at Kommissionens forslag er problematisk: Nye AI-systemer vil snart have langt større kapabiliteter, og definitionen af "højrisiko" udelukkende baseret på de nuværende tekniske funktioner risikerer at blive forældet. Begrænsning af teknologi baseret på potentielle kapabiliteter kan også hindre innovation. I sektorer som f.eks. sundhedssektoren kan dette betyde, at værdifulde AI-løsninger udelukkes, før de overhovedet er blevet overvejet. Du går ind for en formålsbaseret tilgang, da den giver forskerne og udviklerne større frihed i lavrisikoområder, samtidig med at der skabes sikkerhed om testkrav og stærkere sikkerhedsforanstaltninger på områder, hvor risiciene er højere.
- Med hensyn til test støtter du en fleksibel model. Virksomheder sigter af egen drift mod at skabe sikre produkter, så overregulering er unødvendig. **Du mener, at det er vigtigt at fremme afprøvning under faktiske forhold af nye AI-modeller**, da sandkassmiljøer ofte ikke formår fuldt ud at replikere faktiske forhold.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Det er afgørende for dig, at spørgsmål vedrørende national sikkerhed, militær eller forsvar ikke bliver omfattet af AI-forordningen. Medlemsstaterne er fast besluttet på at bevare kontrollen over beslutninger på disse områder, som under alle omstændigheder typisk falder uden for EU's jurisdiktion. Dette ses også af, at EU ikke har et fælles EU-militær.
- Du argumenterer desuden for, at AI-systemer, der anvendes til indsamling af efterretninger, cyberforsvar eller slagmarksscenarier, ikke kan være underlagt de samme oplysningskrav eller reguleringsmæssige krav som civile systemer, fordi dette kunne kompromittere operationel hemmeligholdelse, nationale sikkerhedsinteresser og den strategiske fordel, som sådanne systemer er designet til at give.
- Dine bemærkninger:

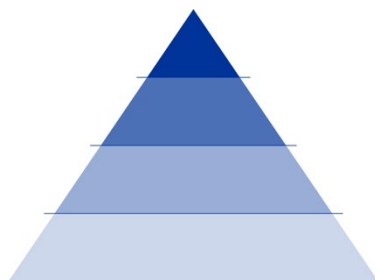
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland	formålsbaseret	fleksibel	afprøvning under faktiske forhold	militær/forsvar + national sikkerhed	
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

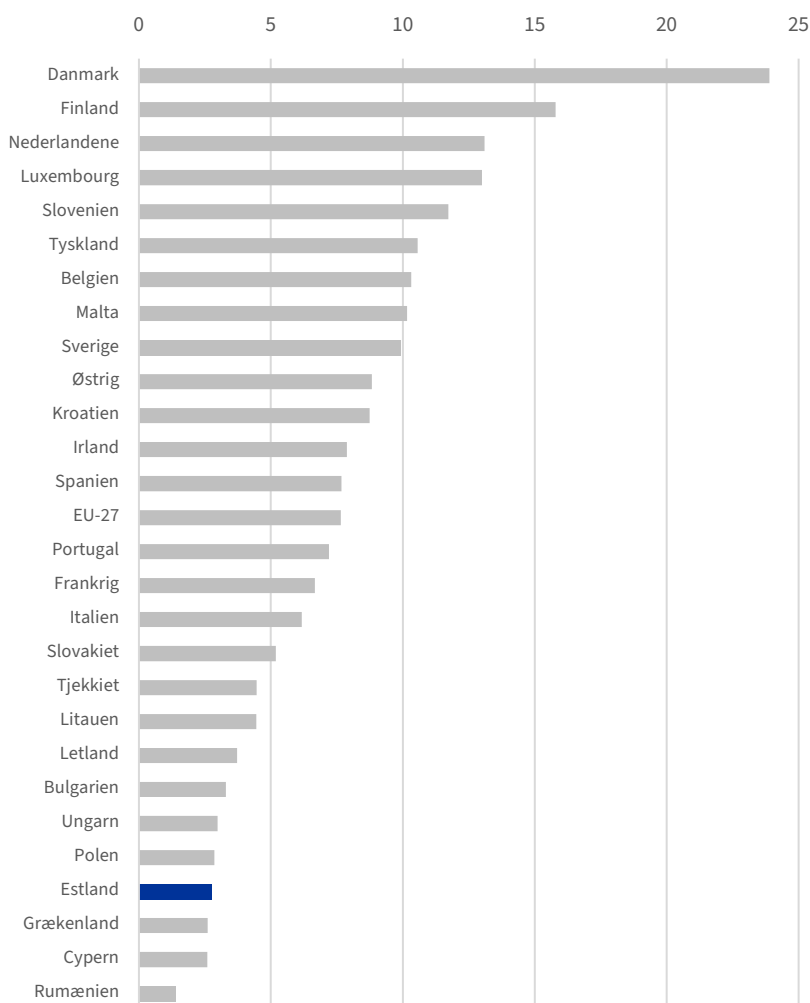
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

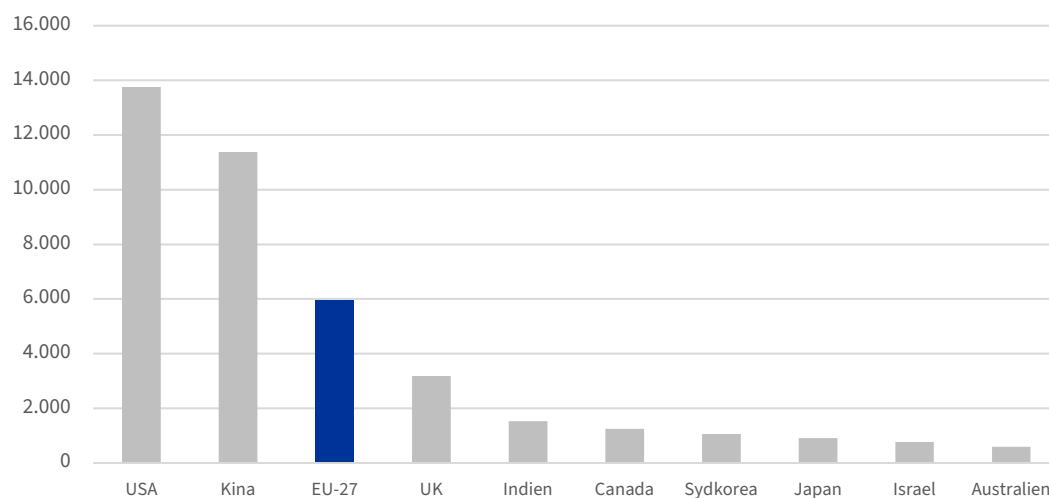
	Den Europæiske Union	Estland
Befolkning	447 695 350	1 365 884
BNP pr. indbygger i €	38 150	27 960
Arbejdsløshedsprocent	6,1 %	6,4 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	5,2 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	34,8 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

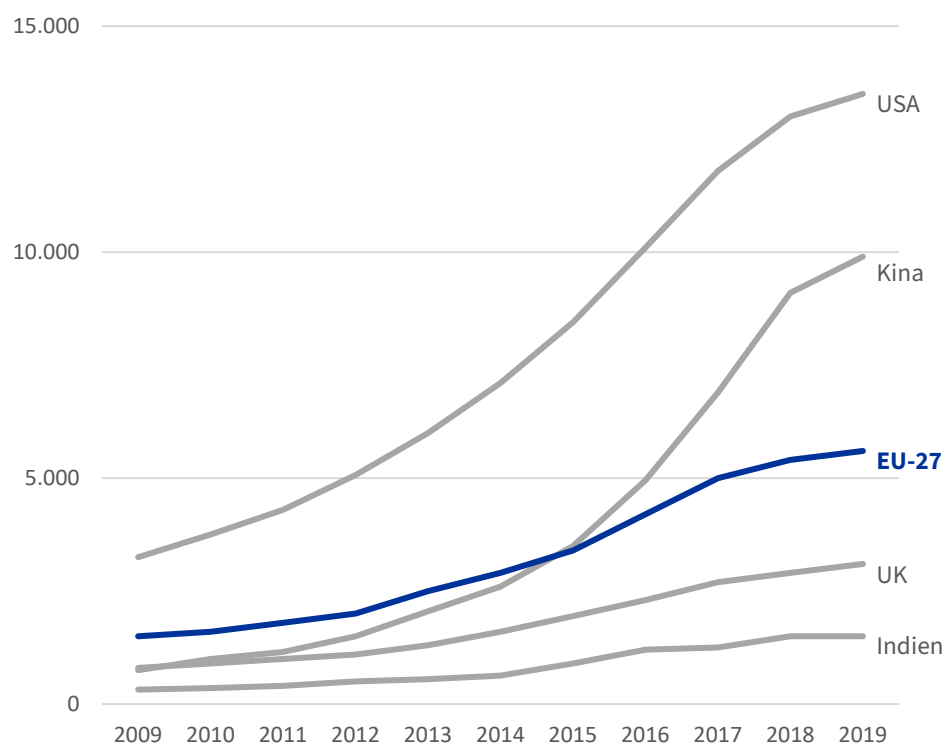


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

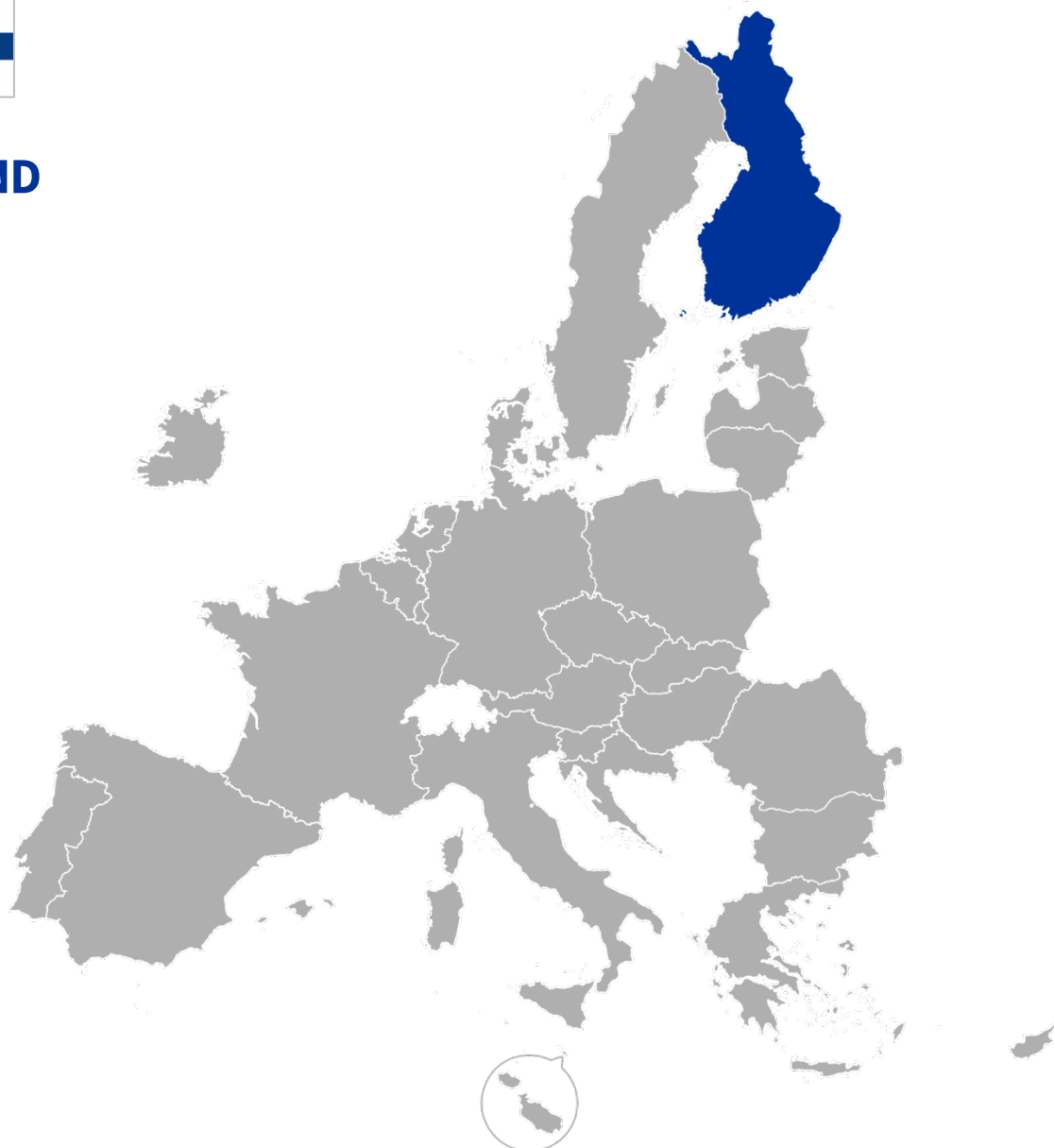
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltageres læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



FINLAND



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



FINLAND

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Irland, Portugal og Sverige**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Dit land er en af EU's AI-ledere med et nationalt mål om at blive en pålidelig, sikker og innovativ pioner i den digitale økonomi. Offentlighedens opbakning er stærk, bl.a. fordi der er offentlig undervisning i AI, som er frit tilgængelig for borgerne.
- Du støtter EU's formålsbaserede lovgivningsmæssige tilgang, som fokuserer på, hvordan AI anvendes, snarere end på dens tekniske funktioner. Derved undgås det, at AI stemples som i sig selv at være farlig, og virkningen under faktiske forhold prioriteres, hvilket efter din mening er den bedste måde at beskytte borgerne på.
- Du mener, at forslaget bør fremhæve EU's støtte til innovation gennem en fleksibel tilgang, bl.a. ved at give mulighed for AI-afprøvning under faktiske forhold under tilsyn af offentlige myndigheder. Det er ønskeligt at tilskynde medlemsstaterne til at oprette reguleringsmæssige sandkasser, men hvis de gøres obligatoriske, risikerer det at lægge en dæmper på udenlandske investeringer.
- Du håber, at der opnås enighed om den endelige tekst til EU's AI-forordning i dag som en ledetråd for troværdig AI-udvikling i Europa, men du er stærkt imod en indbygget revision om få år, da det vil underminere den lovgivningsmæssige stabilitet og investorernes tillid. **Hvis der stadig er usikkerhed, vil det være mere ansvarligt at udsætte afstemningen** end foregribende at planlægge at omskrive reglerne.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Indtil videre har EU haltet bagefter inden for global AI-udvikling. Udelukkelse af opstartsvirksomheder kan bidrage til at fostre et dynamisk økosystem. Opstartsvirksomheder mangler juridisk, finansiel og administrativ kapacitet til at overholde de samme regler som store techvirksomheder. Strenge regler kan kvæle innovation i de tidlige faser.
- Hvis du ikke kan få et tilstrækkeligt antal medlemsstater til at vedtage fleksible testkrav i artikel 1, ønsker du i stedet at sikre, at opstartsvirksomheder får nogle undtagelser fra disse krav. På den måde vil de ikke blive pålagt fuld opfyldelsesbyrde fra starten (hvis artikel 1 fører til det), men kan operere i et mere eksperimentelt rum.
- Dine bemærkninger:

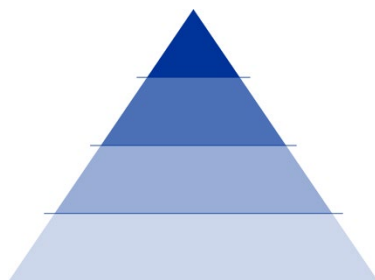
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland	formålsbaseret	fleksibel	udsætte afstemningen	opstartsvirksomheder	
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

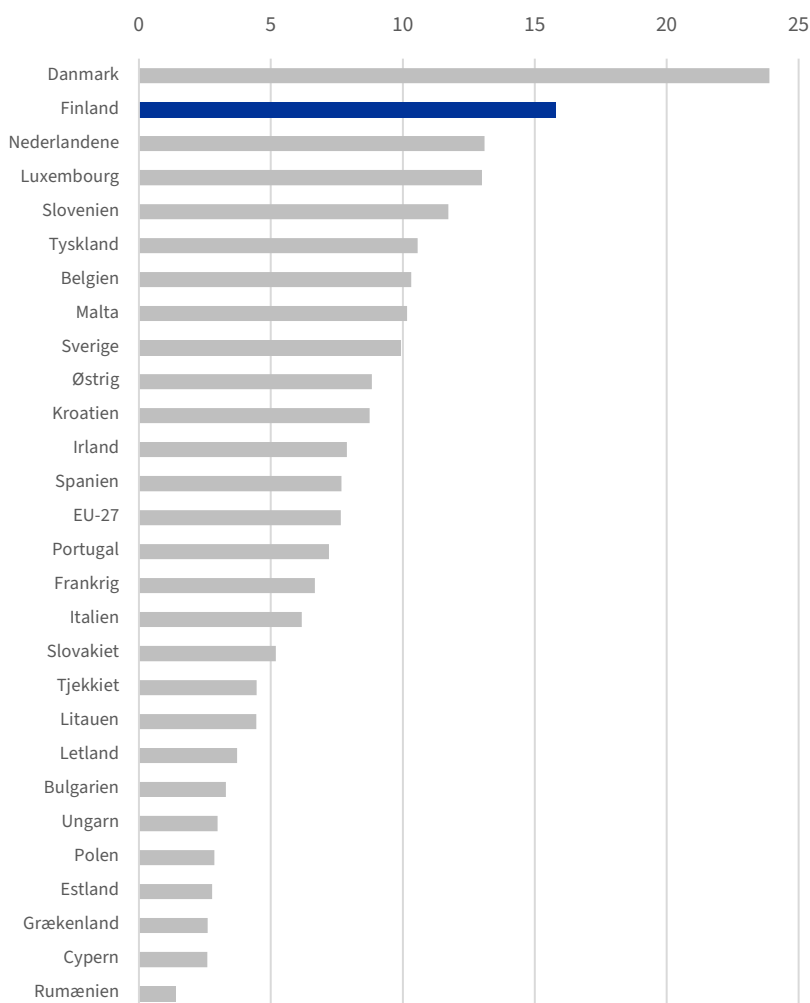
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

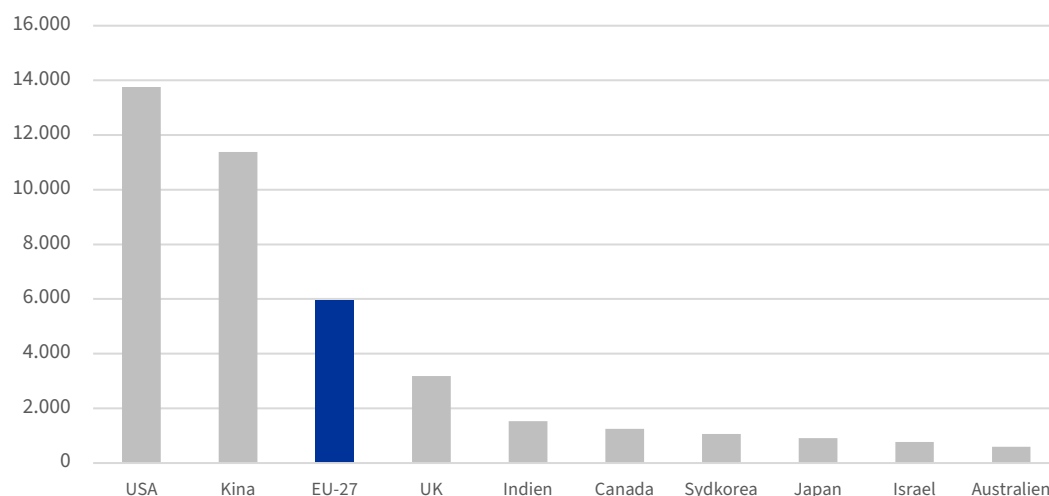
	Den Europæiske Union	Finland
Befolkning	447 695 350	5 563 970
BNP pr. indbygger i €	38 150	48 910
Arbejdsløshedsprocent	6,1 %	7,2 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	15,1 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	53,6 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

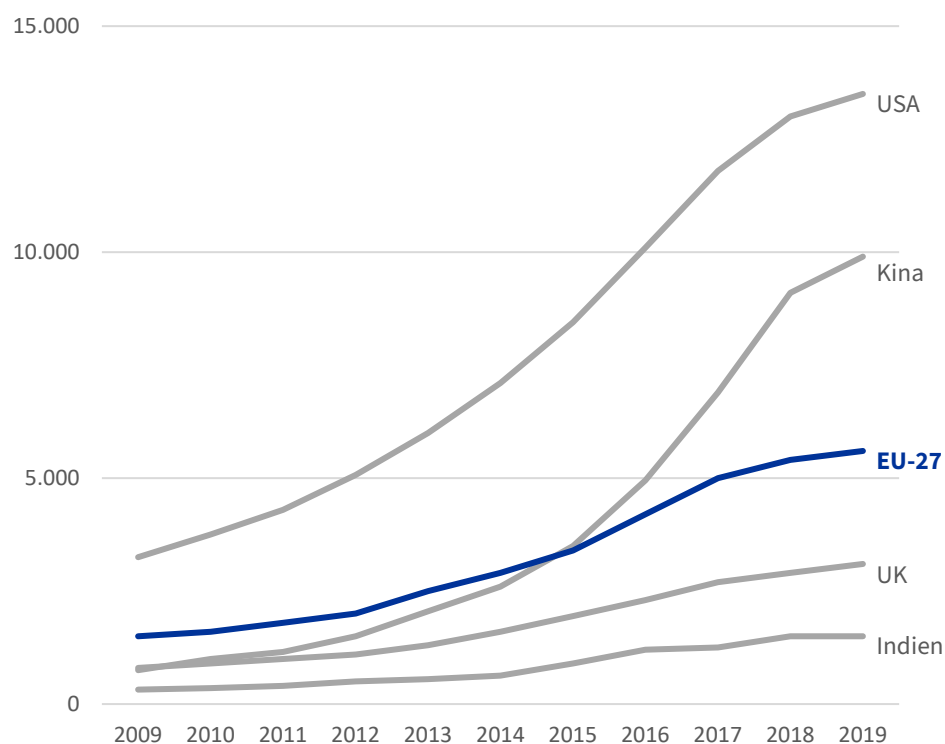


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

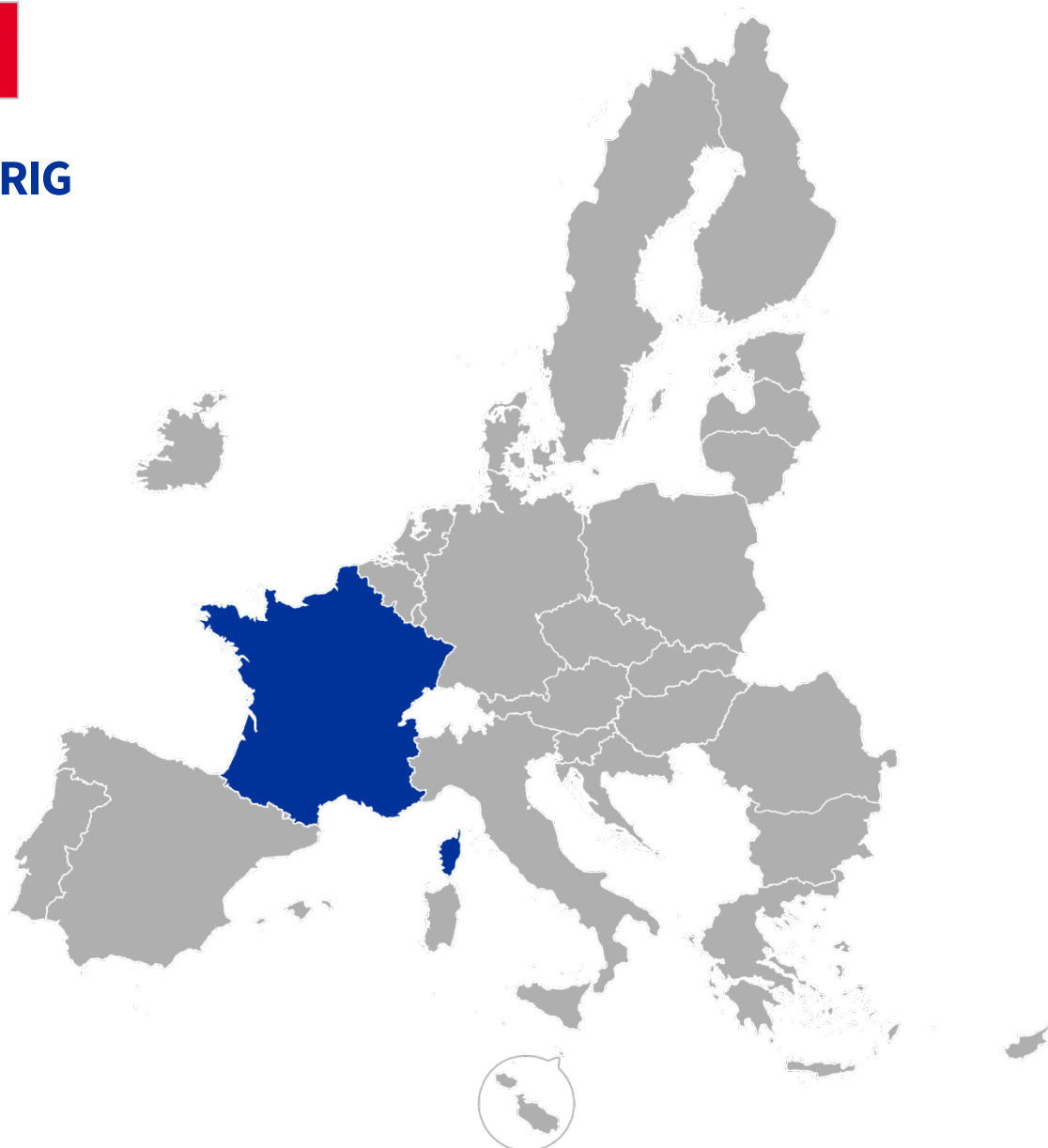
© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print	ISBN 978-92-850-0498-9	doi: 10.2860/0371762	QC-01-25-006-DA-C
PDF	ISBN 978-92-850-0497-2	doi: 10.2860/0996198	QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



FRANKRIG



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



FRANKRIG

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Danmark, Polen, Rumænien og Østrig**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Din nationale regering har allerede en national AI-strategi med titlen "AI for menneskeheden", der afspejler den centrale prioritet: menneskecentreret udvikling.
- Europa-Kommissionens forslag er betimeligt. En europæisk stemme i den globale AI-debat er stærkt påkrævet for at fremme etisk og ansvarlig udvikling. Indtil nu har AI i vid udstrækning udviklet sig ukontrolleret – men selv om der er risici, bør AI ikke dæmoniseres. Den er fortsat en stærk drivkraft for innovation og vækst.
- Du støtter den formålsbaserede tilgang, fordi den skaber klarhed for investorerne, **men listen bør udvides til at omfatte uddannelsessektoren**. AI-fejl i forbindelse med indplacering eller bedømmelse af elever og studerende kan føre til forskelsbehandling og styrke ulighed. Fair adgang til uddannelse er afgørende for at bekæmpe strukturel forskelsbehandling.
- Af erfaring ved du, at lovgivningsmæssige incitamenter er nyttige – men uden offentlig finansiering er der risiko for, at innovative idéer dør ud, hvilket fører til hjerneflugt og virksomhedsmigration. EU skal ikke kun indføre en ramme, men også forpligte sig til betydelige offentlige investeringer i AI-udvikling.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Din regering er meget engageret i borgerlige frihedsrettigheder. Samtidig støtter du en tosporet tilgang til AI-regulering – en tilgang, der sikrer stærk beskyttelse i det civile rum og samtidig giver mulighed for mere fleksible standarder med hensyn til forsvar og national sikkerhed. National sikkerhed er afgørende for at bevare de friheder, som de borgerlige frihedsrettigheder bygger på, og AI kan spille en central rolle for beskyttelse af demokratiske samfund mod nye trusler.
- Derfor argumenterer du for at udelukke applikationer til militær, forsvar og national sikkerhed fra denne forordning og udvikle en anden forordning om anvendelse af AI i disse sektorer. Selv om det stadig er afgørende, at AI udvikles ansvarligt, kan krav om fuld gennemsigtighed i forbindelse med forsvarsrelateret AI underminere operationel hemmeligholdelse, forsinke ibrugtagning og reducere effektiviteten i miljøer i hastig udvikling eller i fjendtlige miljøer.
- Desuden argumenterer du for, at det præciseres, at strenge testkrav til alle AI-systemer også gælder for AI, der importeres fra lande uden for EU. Du foreslår, at ordlyden ændres til "Forskrifterne fastsat i artikel 1 finder anvendelse på alle AI-systemer, der sættes i drift i EU, ...".
- Dine bemærkninger:

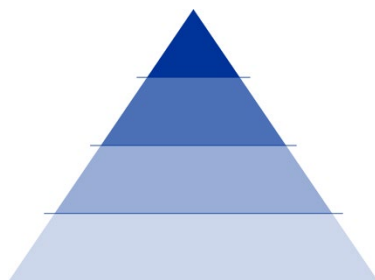
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig	formålsbaseret	forsigtigheds-	inkludere uddannelse	militær/forsvar + national sikkerhed	
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

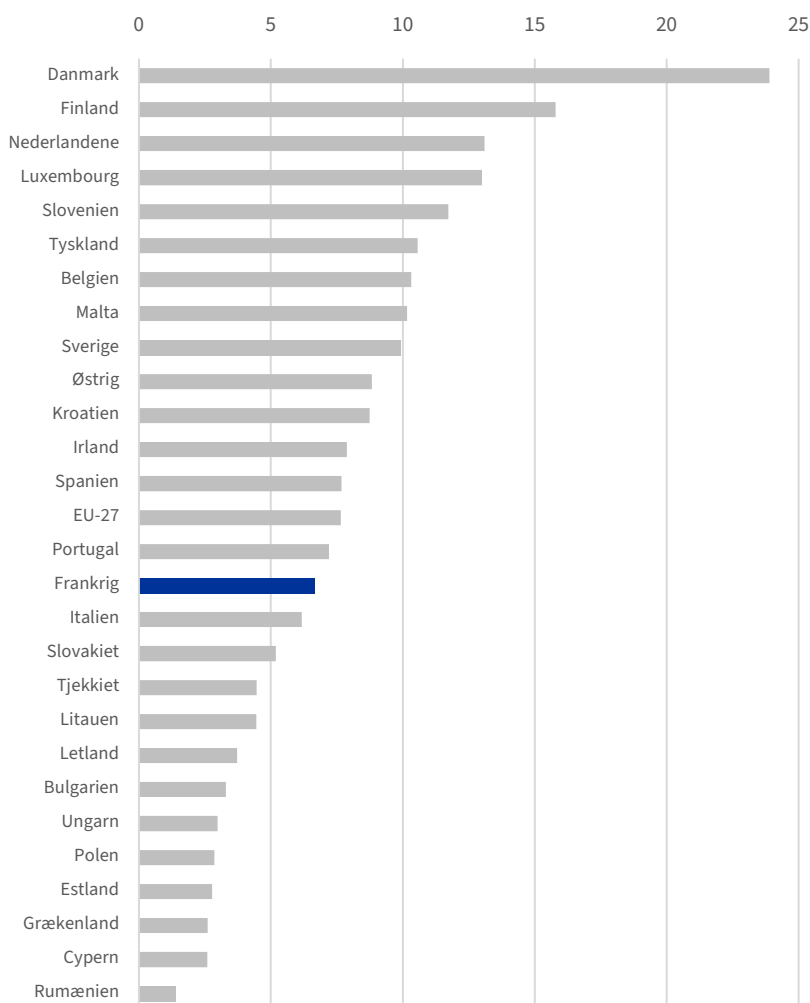
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

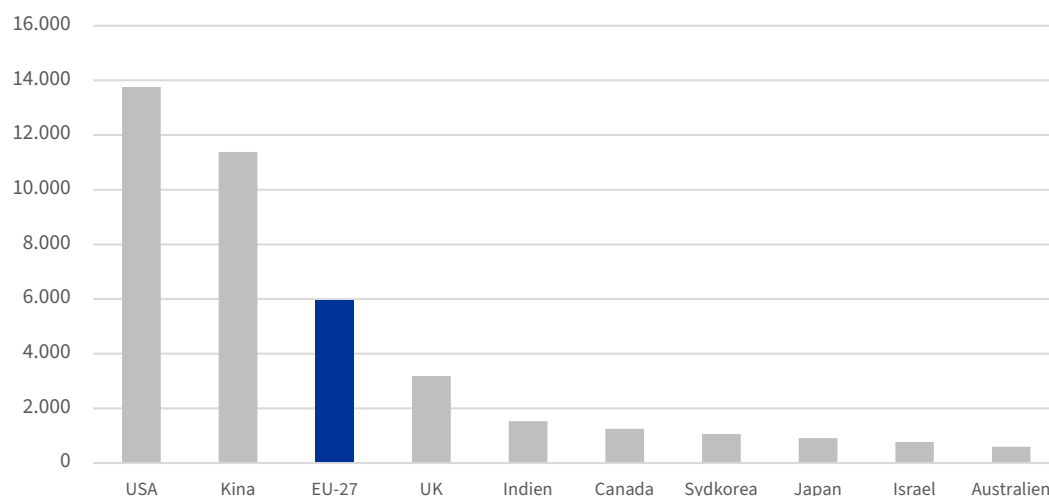
	Den Europæiske Union	Frankrig
Befolkning	447 695 350	68 277 210
BNP pr. indbygger i €	38 150	41 330
Arbejdsløshedsprocent	6,1 %	7,3 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	5,9 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	30,3 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

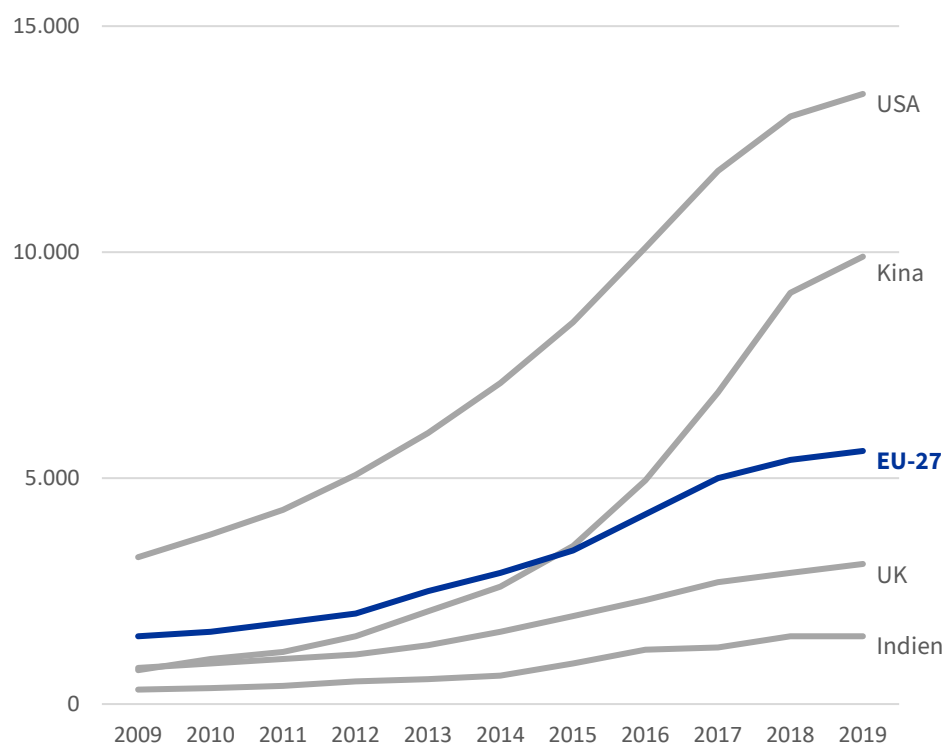


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

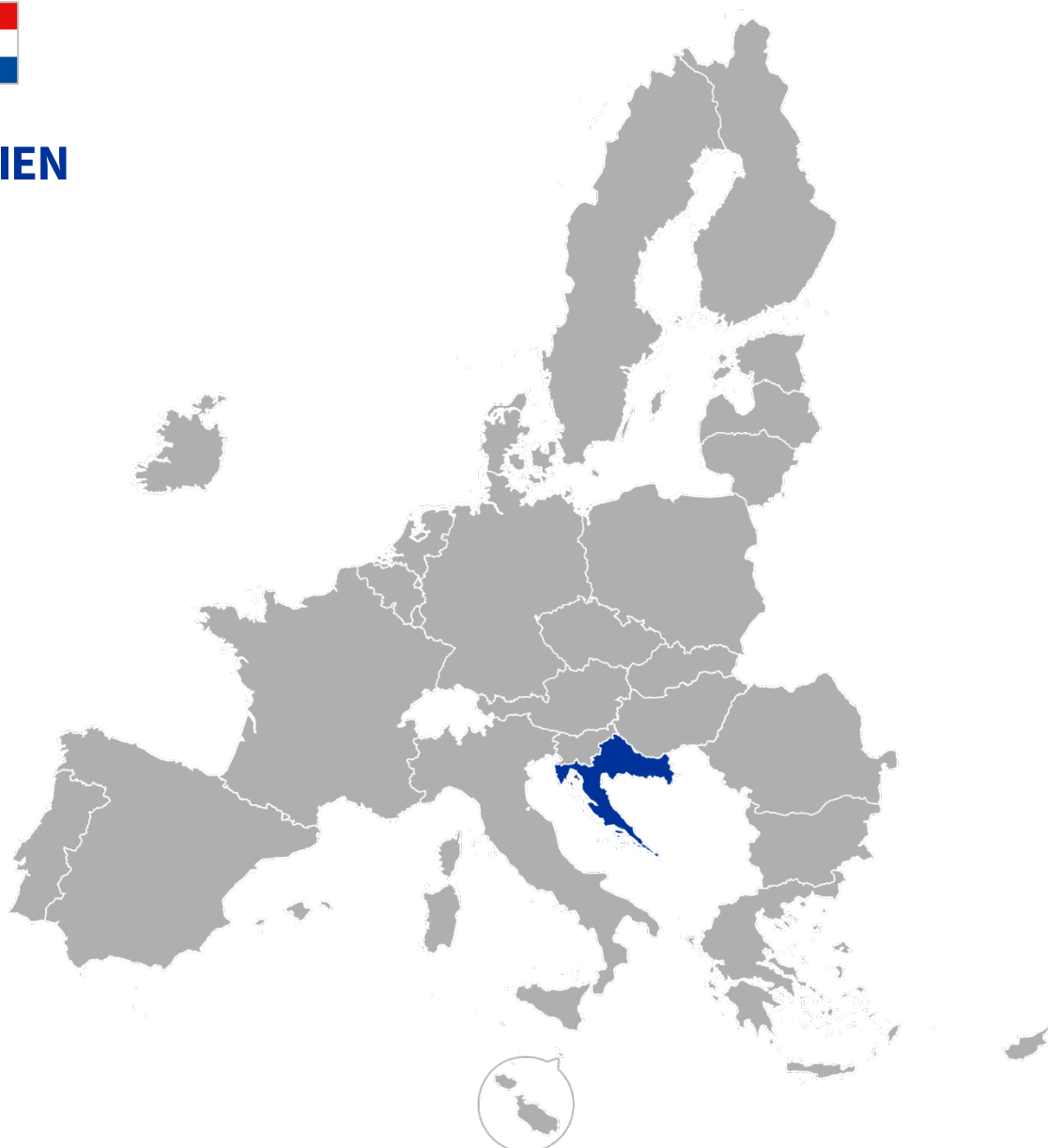
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



KROATIEN



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



KROATIEN

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Italien, Nederlandene og Slovenien**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Antallet af opstartsvirksomheder i Kroatien vokser hurtigt, og virksomhedskonsulenter skønner, at AI-drevet automatisering kan bidrage med op til 16 % af BNP i 2025, hovedsagelig gennem produktivitetsgevinster. Op til 52 % af arbejdstiden kan automatiseres. Selv om dette er lovende for økonomisk vækst, giver det også anledning til bekymring for potentiel arbejdsløshed i produktionssektoren.
- Du støtter Kommissionens formålsbaserede tilgang til AI-klassificering, fordi den afspejler kontekstspecifikke risici realistisk – f.eks. ansigtsgenkendelse for at låse en telefon op versus anvendelse heraf til masseovervågning.
- Men du mener, at gennemsigtighed er afgørende: Alle AI-systemer bør uanset risikoniveau informere brugerne klart om formålet med dem. Det sikrer ansvarlighed og beskytter individuelle rettigheder.
- Du støtter også stærke udviklingsstandarder, men mener, at de nuværende retningslinjer for risikovurdering er for vage. Du argumenterer for en obligatorisk menneskerettighedskonsekvensanalyse for højrisiko-AI, der er udviklet i samråd med de berørte samfund, for at beskytte sårbare grupper.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Efter din mening bør EU's AI-forordning fungere som katalysator for AI-udvikling i alle EU's medlemsstater. Nogle kapitalstærke lande har magtfulde AI-virksomheder, mens andre stadig forsøger at etablere deres opstartskosystemer. Hele idéen bag opstartsvirksomheder er at skabe noget disruptivt på kort tid og med begrænset adgang til kapital. Det er næsten umuligt, hvis der er en betydelig administrativ byrde.
- Derfor ser du gerne, at opstartsvirksomheder (med færre end ti ansatte og en årlig omsætning på under 2 mio. €) undtages fra EU's AI-forordning.
- Hvis du ikke kan få et tilstrækkeligt antal andre medlemsstater til at støtte din holdning, mener du, at supplerende EU-midler til opstartsvirksomheder er et godt kompromis. På den måde bliver de ikke hæmmet af de administrative og budgetmæssige begrænsninger, som forordningen forårsager.
- Dine bemærkninger:

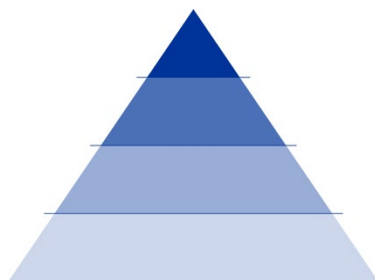
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien	formålsbaseret	forsigtigheds-		opstartsvirksomheder	
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

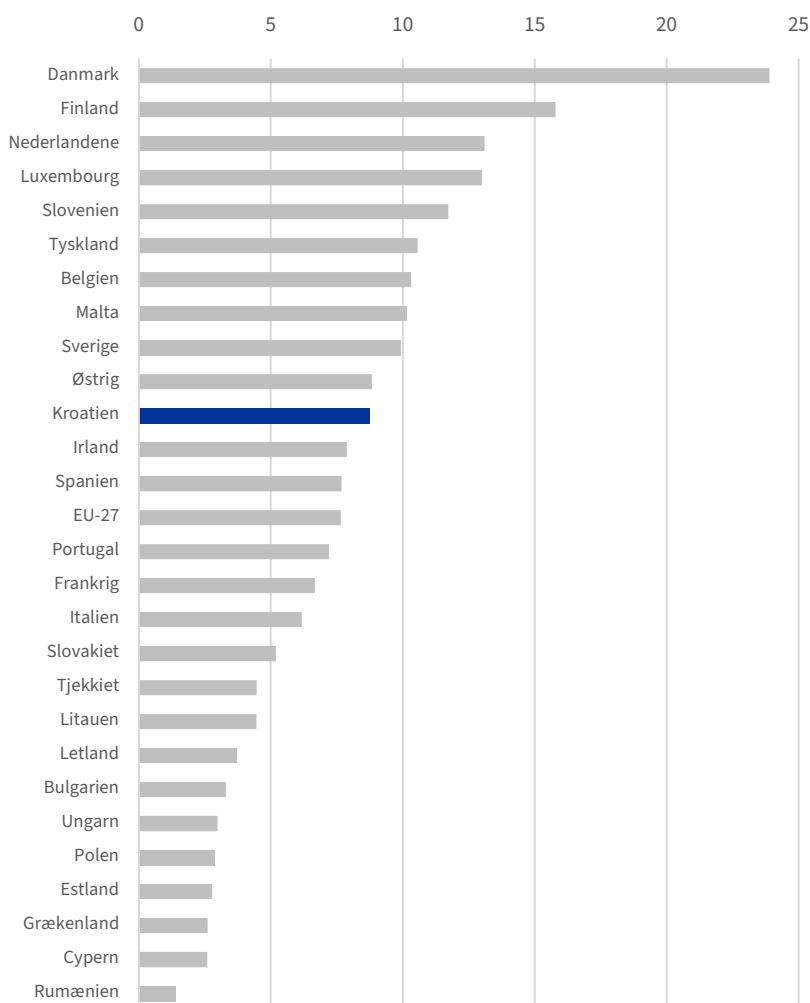
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

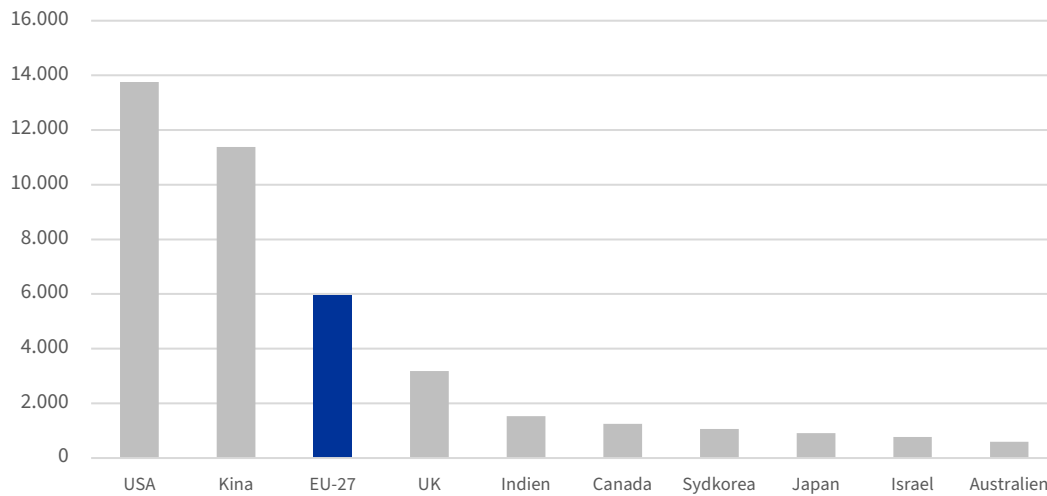
	Den Europæiske Union	Kroatien
Befolkning	447 695 350	3 850 894
BNP pr. indbygger i €	38 150	19 800
Arbejdsløshedsprocent	6,1 %	6,1 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	7,9 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	25 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

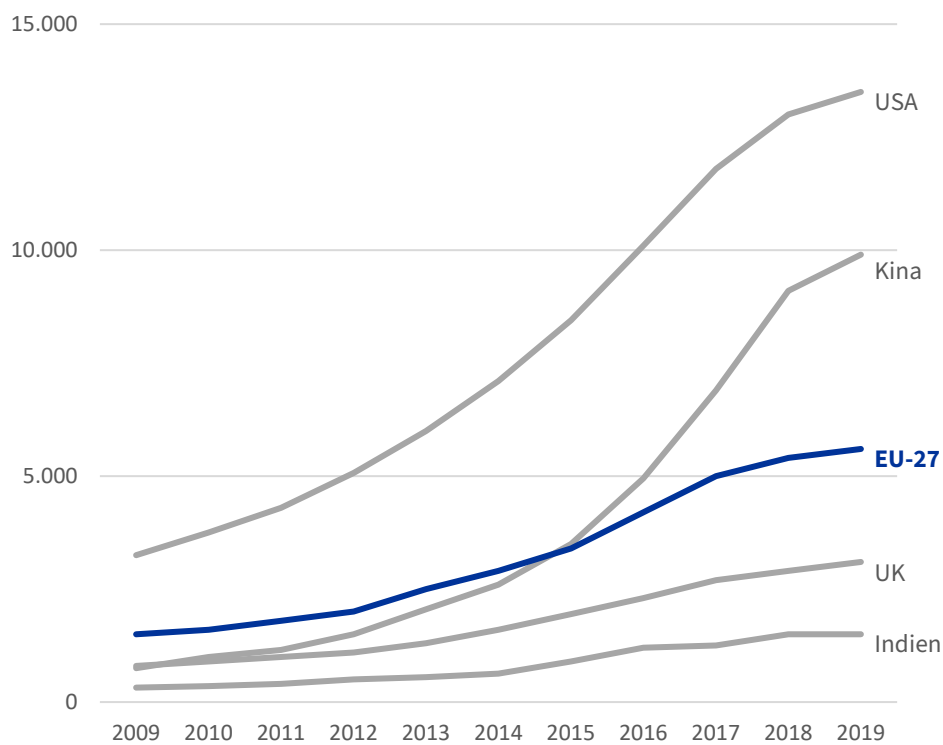


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

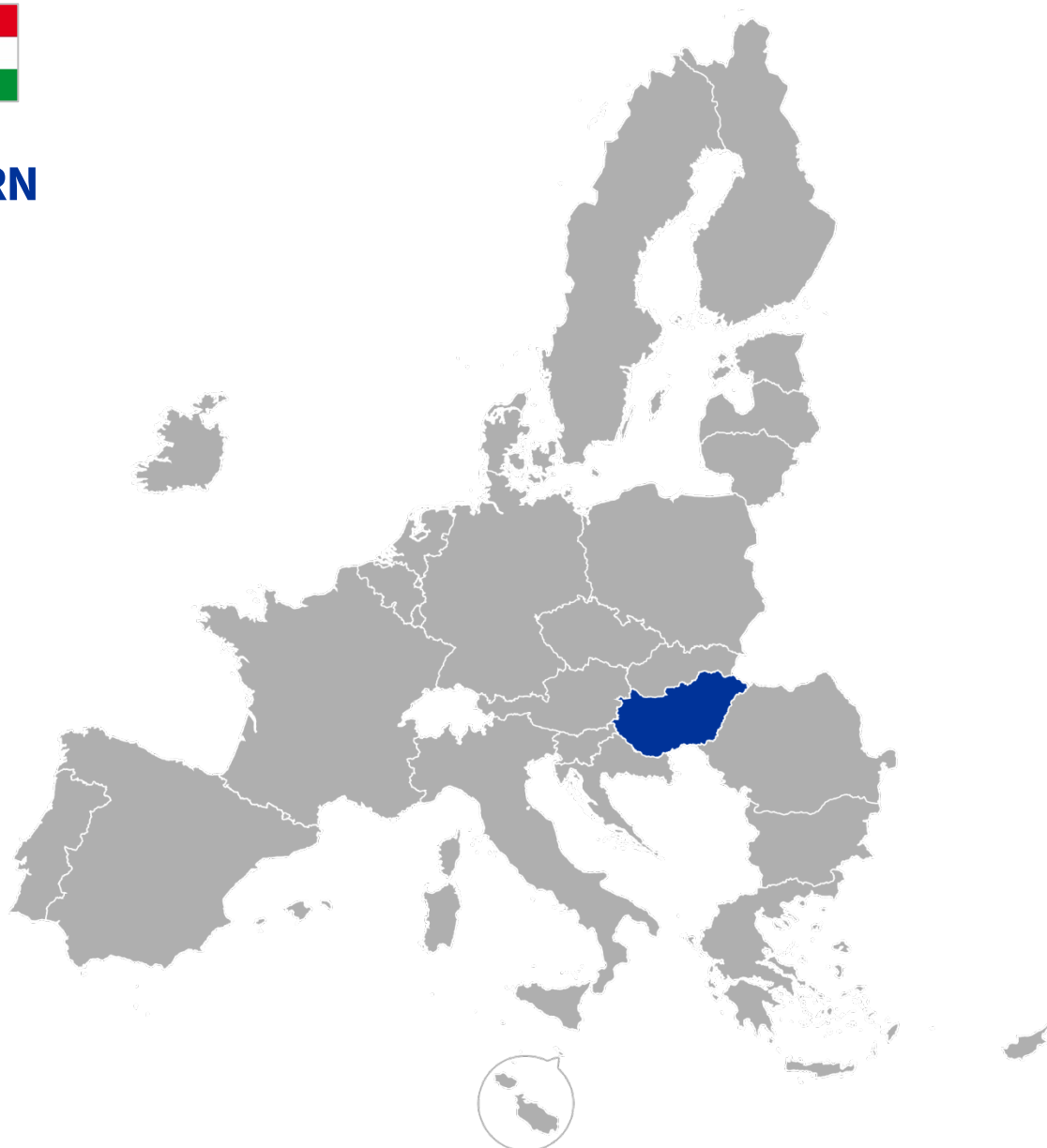
© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print	ISBN 978-92-850-0498-9	doi: 10.2860/0371762	QC-01-25-006-DA-C
PDF	ISBN 978-92-850-0497-2	doi: 10.2860/0996198	QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltageres læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



UNGARN



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



UNGARN

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Estland, Letland, Litauen og Tjekkiet**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Din regering har for nylig udviklet en national AI-strategi. En af dens søjler er at forberede samfundet på effektiv håndtering af de uundgåelige ændringer, som AI fører med sig, og at gøre det muligt for landets industri og universiteter at udnytte fordelene ved teknologien.
- Du mener, at øget AI-indførelse kan skabe betydelig merværdi for det nationale BNP – ressourcer, der er stærkt påkrævede til investeringer i offentlig infrastruktur.
- Med hensyn til klassificering af højrisiko-AI-modeller er den kapabilitetsbaserede tilgang mere velegnet end den formålsbaserede tilgang, som Kommissionen har foreslået. Hvis det samme AI-system betragtes som lav risiko i én sammenhæng og højrisiko i en anden, kan det skabe forvirring og retsikkerhed for brugere, udviklere og investorer.
- Du ser en stærk rolle for EU med hensyn til at sikre fair konkurrence, så både europæiske og ungarske virksomheder forbliver konkurrencedygtige på det indre marked og i de globale værdikæder. Ungarn sigter mod at være et attraktivt mål for udenlandske investeringer. Derfor støtter du en forordning, der kun fastsætter minimumskrav eller -standarder, for at undgå at pålægge udviklerne en stor omkostningsbyrde.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Du argumenterer for, at militære og forsvarsrelaterede applikationer udelukkes fra AI-forordningen, og understreger, at overregulering på dette område kan hindre hurtig ibrugtagning af AI-teknologier i konfliktområder eller i forbindelse med kriserespons. I forbindelse med den aktuelle regionale ustabilitet er fleksibilitet i anvendelsen af den slags teknologier afgørende.
- Som et land, der grænser op til Ukraine – en nation, der på nuværende tidspunkt er i krig med Rusland – er Ungarn dybt bekymret over den nationale sikkerhed og landets borgeres sikkerhed. Som Schengengrænseland står Ungarn desuden over for betydelige udfordringer med hensyn til at håndtere de øgede migrationsstrømme, navnlig langs landets sydlige grænse.
- I lyset af dette pres søger Ungarn større forståelse og støtte fra andre EU-medlemsstater, hvad angår landets unikke geopolitiske og sikkerhedsmæssige udfordringer.
- Selv om nogle lande er bekymrede over risiciene ved importerede AI-systemer – især fra Kina – mener du, at mange af disse bekymringer er politisk motiverede snarere end begrundet i reelle trusler. Ungarn har haft positive erfaringer med kinesiske teknologipartnerskaber, og efter din mening kan EU drage fordel af at løse sine handelsspændinger med Kina.
- Dine bemærkninger:

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

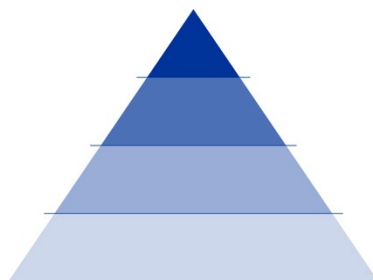
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn	kapabilitetsbaseret	fleksibel		militær/forsvar	
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

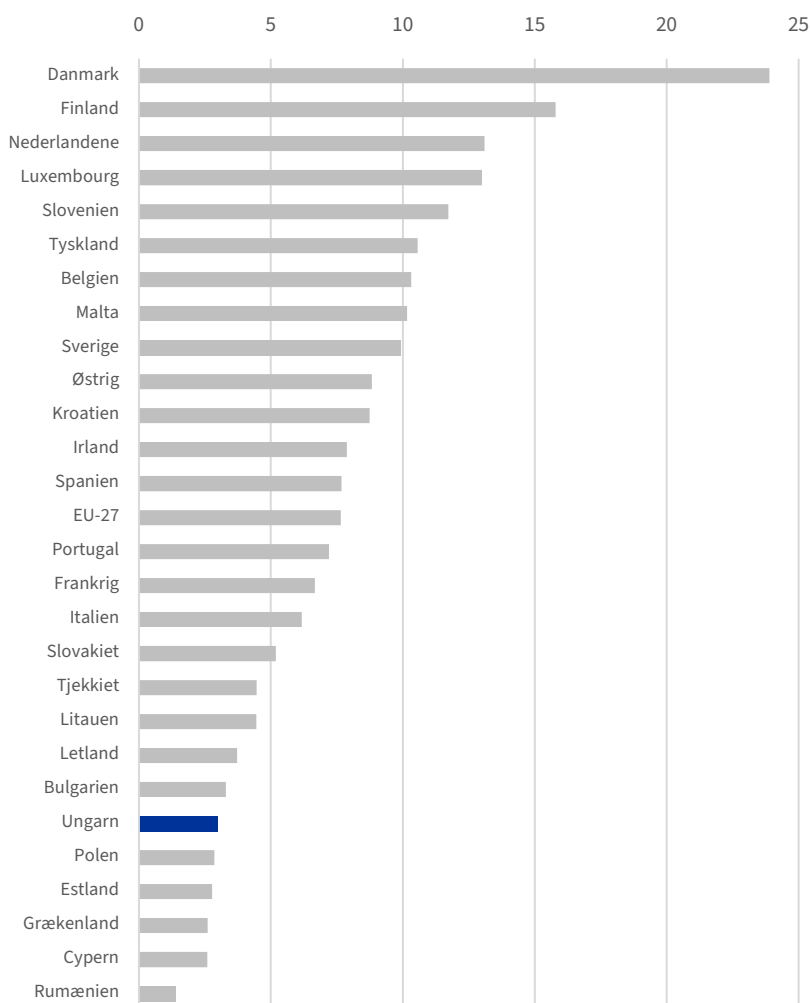
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

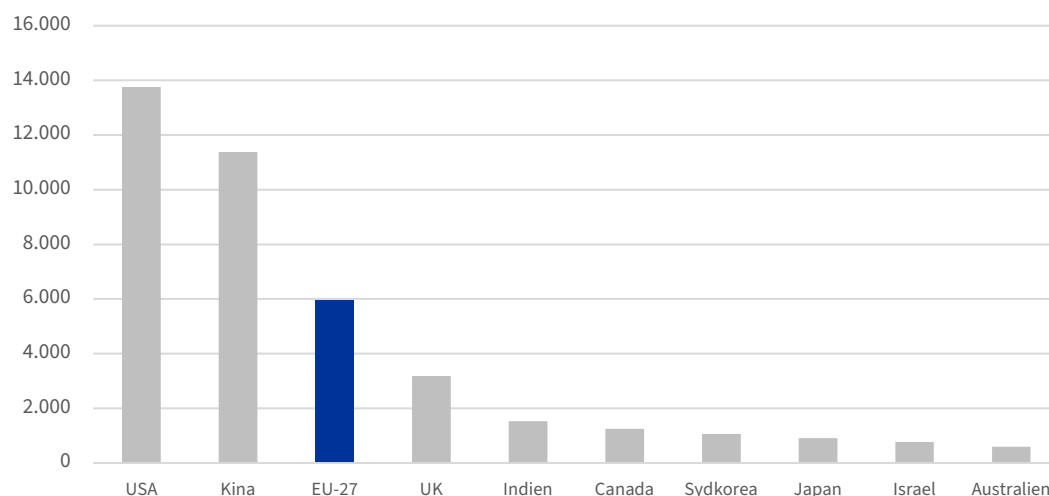
	Den Europæiske Union	Ungarn
Befolkning	447 695 350	9 599 744
BNP pr. indbygger i €	38 150	20 630
Arbejdsløshedsprocent	6,1 %	4,1 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	3,7 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	28,1 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

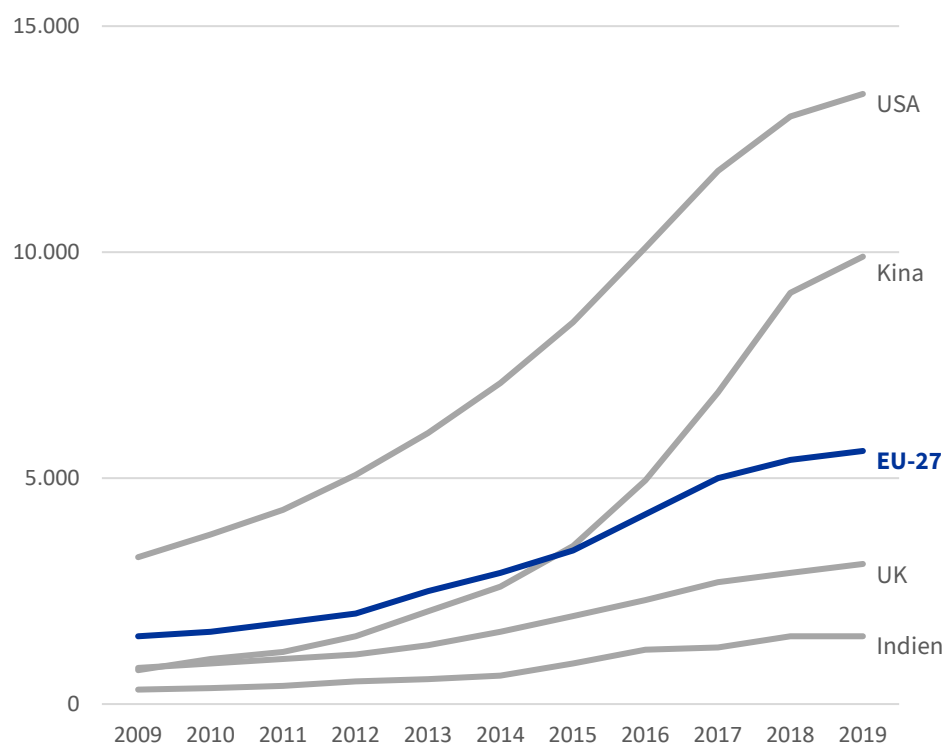


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9
PDF ISBN 978-92-850-0497-2

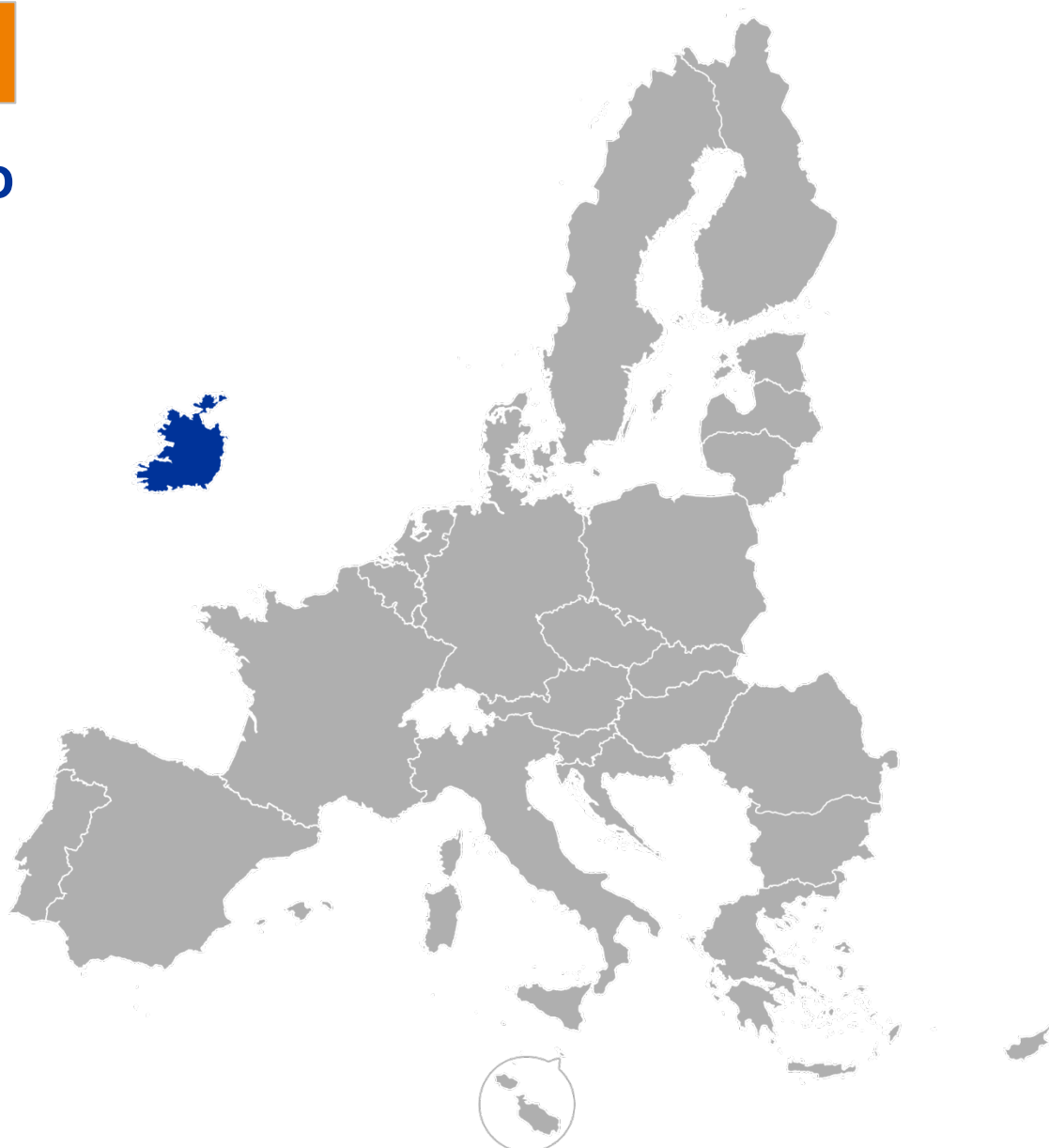
doi: 10.2860/0371762
doi: 10.2860/0996198

QC-01-25-006-DA-C
QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltageres læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



IRLAND



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



IRLAND

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Finland, Portugal og Sverige**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- OpenAI, Google og Meta har deres EU-hovedkvarter i dit land. At bibeholde disse virksomheder i Irland er i overensstemmelse med dine nationale interesser – og gavner EU som helhed.
- Du glæder dig over Kommissionens forslag om at indføre en fælles EU-ramme for AI. Du mener, det er vigtigt, da det sender det rigtige signal og viser verden, at EU er parat til at være et ansvarligt AI-kraftcenter. Derfor støtter du den formålsbaserede tilgang, som Kommissionen har foreslået. Med denne tilgang – sammenholdt med den kapabilitetsbaserede tilgang – er der mindre risiko for konflikt med eksisterende krav i andre sektorspecifikke forordninger.
- Hvad angår krav til udviklere, **understreger du behovet for fleksibilitet: Reglerne skal holde investorerne engagerede og samtidig sikre ansvarlighed.**
- Du fremhæver også, at mange AI-risici opstår efter ibrugtagning, så løbende overvågning er mere effektivt end udsættelse af markedsintroduktion på grundlag af hypotetiske risici.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Du mener, at AI og den digitale revolution også skal fremme ligestilling mellem kønnene i EU. Kvinder er underrepræsenteret på næsten alle områder og niveauer inden for forskning og innovation. Du understreger betydningen af aktivt at støtte deres deltagelse i udformningen af AI.
- Opstartsvirksomheder er en enestående mulighed. I modsætning til store techvirksomheder har de ofte mindre kønsforskelle og giver kvinder og genderqueers bedre muligheder inden for AI. Forskellige teams har også vist sig at opbygge mindre biaspræget og mere etisk AI – noget, EU bør tilskynde til.
- Det giver sig selv, at opstartsvirksomheder sigter mod at innovere hurtigt med begrænsede ressourcer. Tunge administrative byrder kan kvæle dette potentiale, navnlig for små teams, der forsøger at vinde fodfæste på området.
- Derfor støtter du, at små opstartsvirksomheder (under ti ansatte og 2 mio. € i omsætning) undtages fra den fulde indvirkning af EU's AI-forordning.
- Dine bemærkninger:

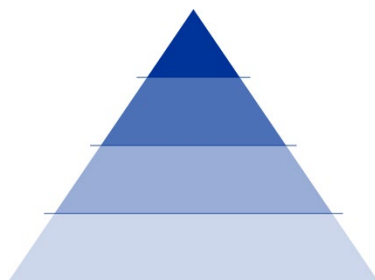
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland	formålsbaseret	fleksibel	fleksibilitet med hensyn til afprøvning	opstartsvirksomheder	

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

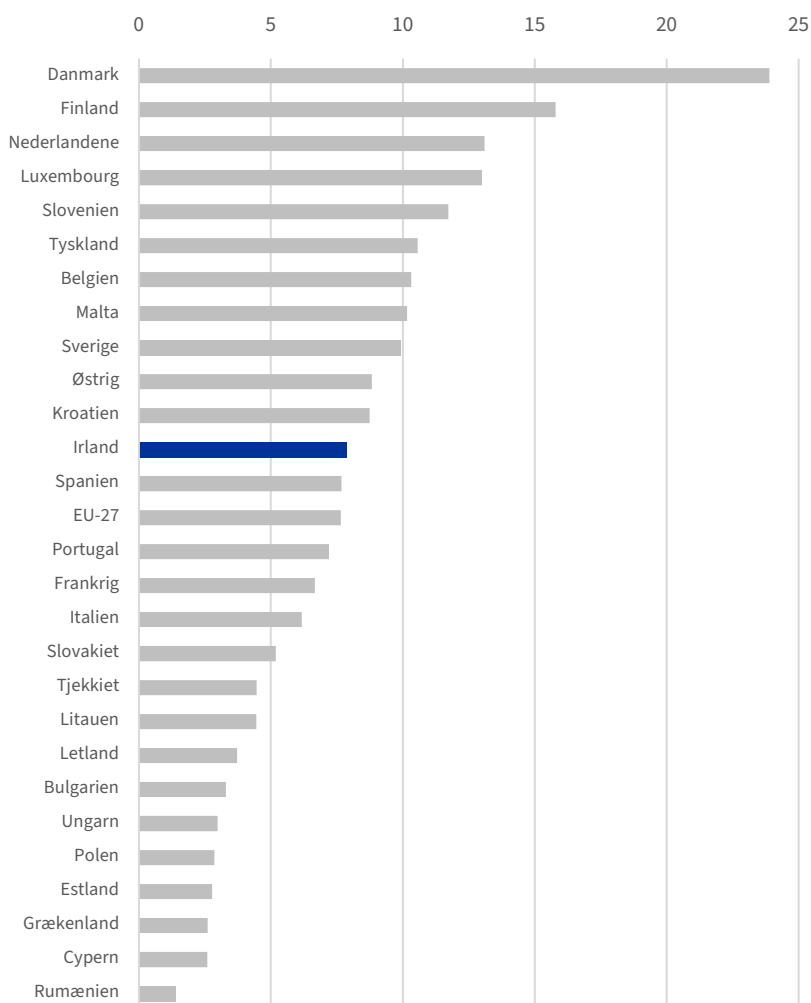
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

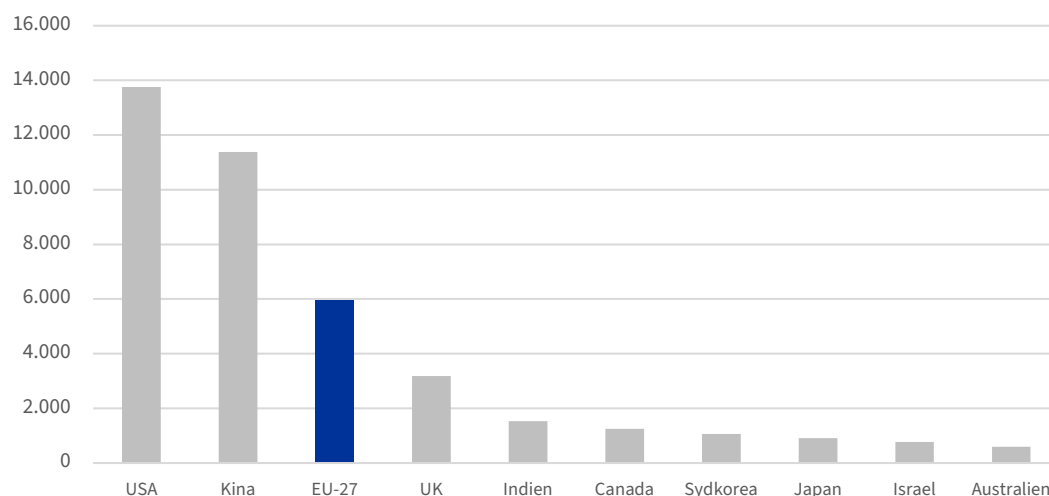
	Den Europæiske Union	Irland
Befolkning	447 695 350	5 271 395
BNP pr. indbygger i €	38 150	96 290
Arbejdsløshedsprocent	6,1 %	4,3 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	8 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	43,8 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

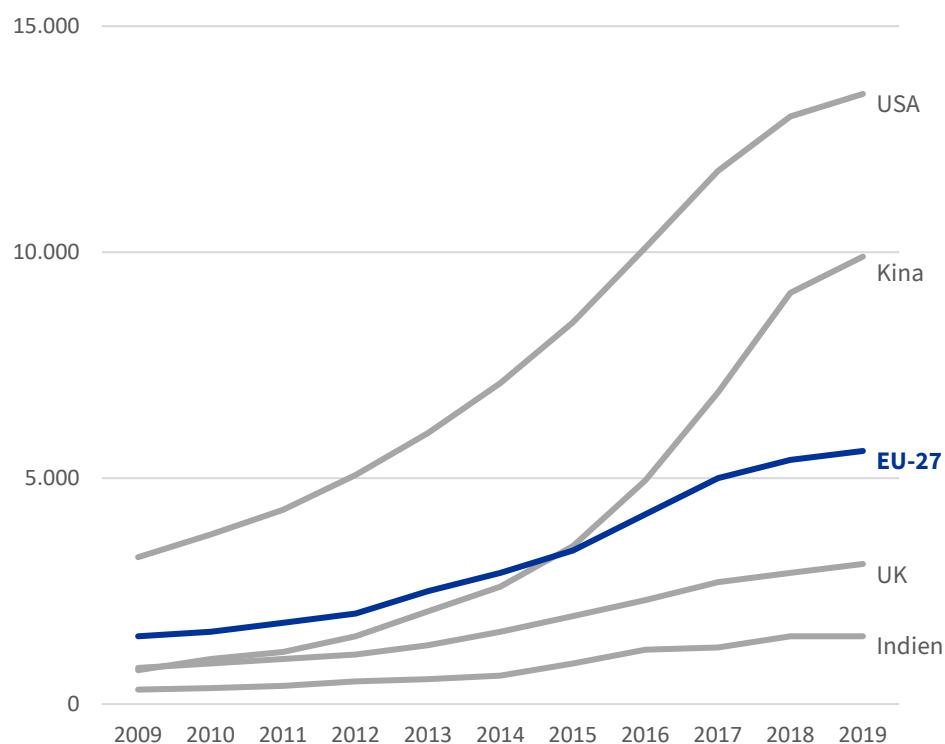


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

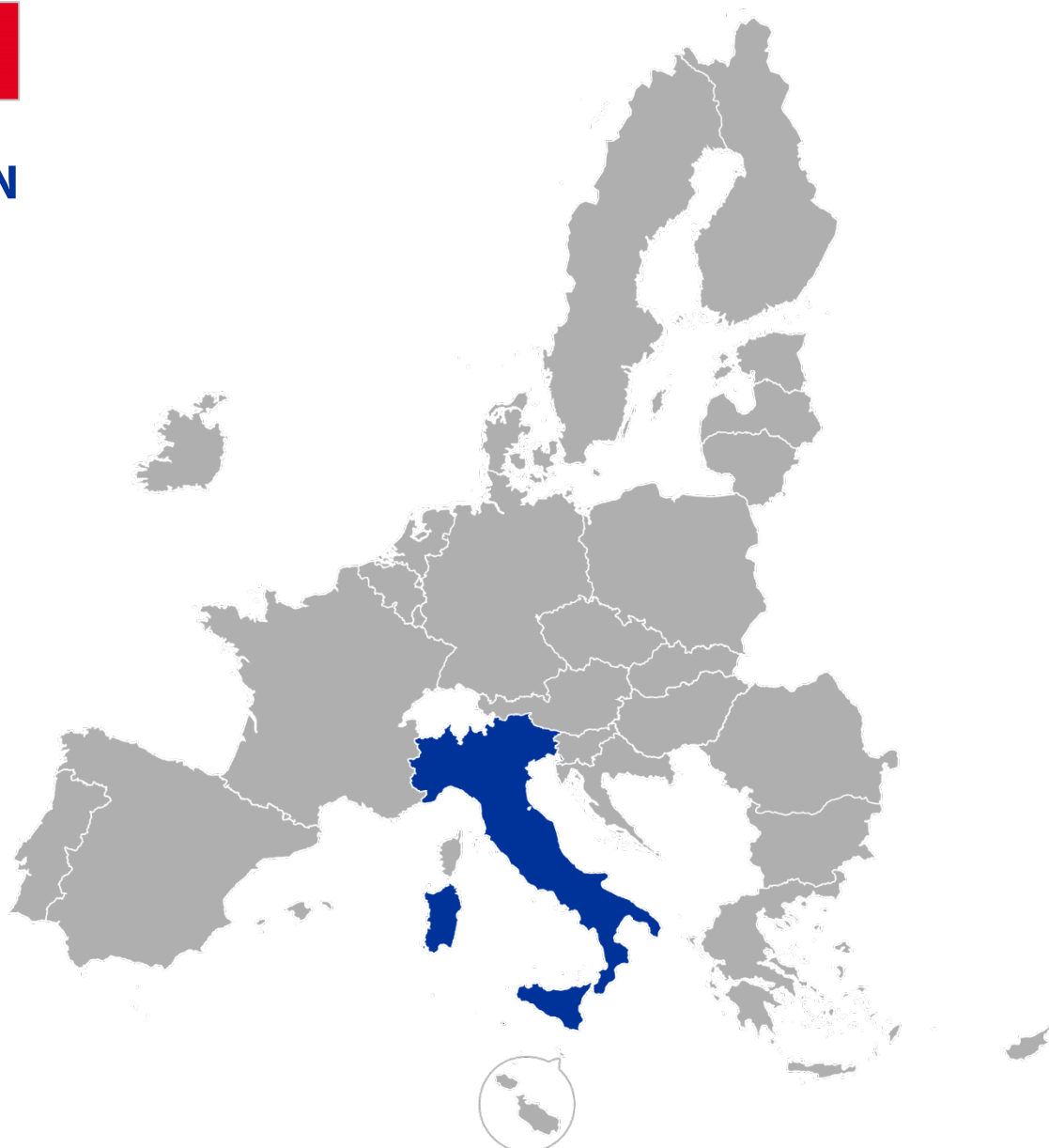
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



ITALIEN



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



ITALIEN

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Kroatien, Nederlandene og Slovenien**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Din nationale regering mener, at reguleringen af AI bør ske på europæisk plan, da nationale regler vil underminere EU's konkurrenceevne og strategiske autonomi.
- I overensstemmelse med artikel 3 i den italienske forfatning ("Alle borgere har samme sociale værdighed og er lige for loven") og FN's verdensmål for bæredygtig udvikling skal AI fremme inddragelse, bæredygtighed og menneskerettigheder – til gavn for både menneskene og kloden.
- Du mener, at den formålsbaserede tilgang er problematisk, fordi den kan skabe sektormæssige uoverensstemmelser: F.eks. klassificeres AI, der assisterer kirurgi, og AI, der tildeler hospitalsstuer, som højrisiko på trods af deres forskellige indvirkning. En kapabilitetsbaseret tilgang afspejler bedre farer i den virkelige verden, selv om du støtter regelmæssige ajourføringer af risikolisten for at afspejle nye AI-kapabiliteter.
- Du går også ind for en ramme for forsigtighedstest, ikke kun af sikkerhedshensyn, men også af bæredygtighedshensyn: AI-udvikling er arbejdskraft- og energiintensiv, og det kræver hundredvis af megawatttimer i elektricitet at træne store sprogmodeller. Derfor øger en ureguleret trial and error-tilgang miljøomkostningerne. En strengere ramme kan bremse udviklingen lidt, men fører i det lange løb til mere sikker, effektiv og troværdig AI.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Efter din mening bør EU's AI-forordning fungere som katalysator for AI-udvikling i alle EU's medlemsstater. Opstartsvirksomheder er en afgørende drivkraft for innovation, så AI-forordningen bør ikke pålægge AI-opstartsvirksomheder en uforholdsmæssigt stor byrde – strenge testkrav kan overstige en opstartsvirksomheds finansielle og administrative formåen.
- Hvis artikel 1 fører til en ramme for forsigtighedstest for højrisiko-AI-modeller – hvilket du støtter – foreslår du at undtage opstartsvirksomheder med færre end 20 ansatte og under 4 mio. € i omsætning. Medlemsstaterne bør i stedet vedtage den fleksible tilgang. Du ser dette som et afbalanceret kompromis og agter at lede drøftelsen om det.
- Dine bemærkninger:

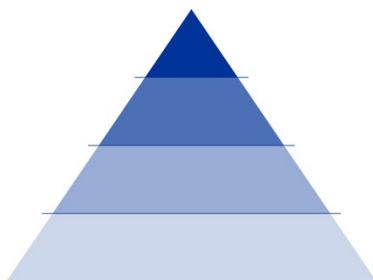
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien	kapabilitetsbaseret	forsigtigheds-		opstartsvirksomheder	
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

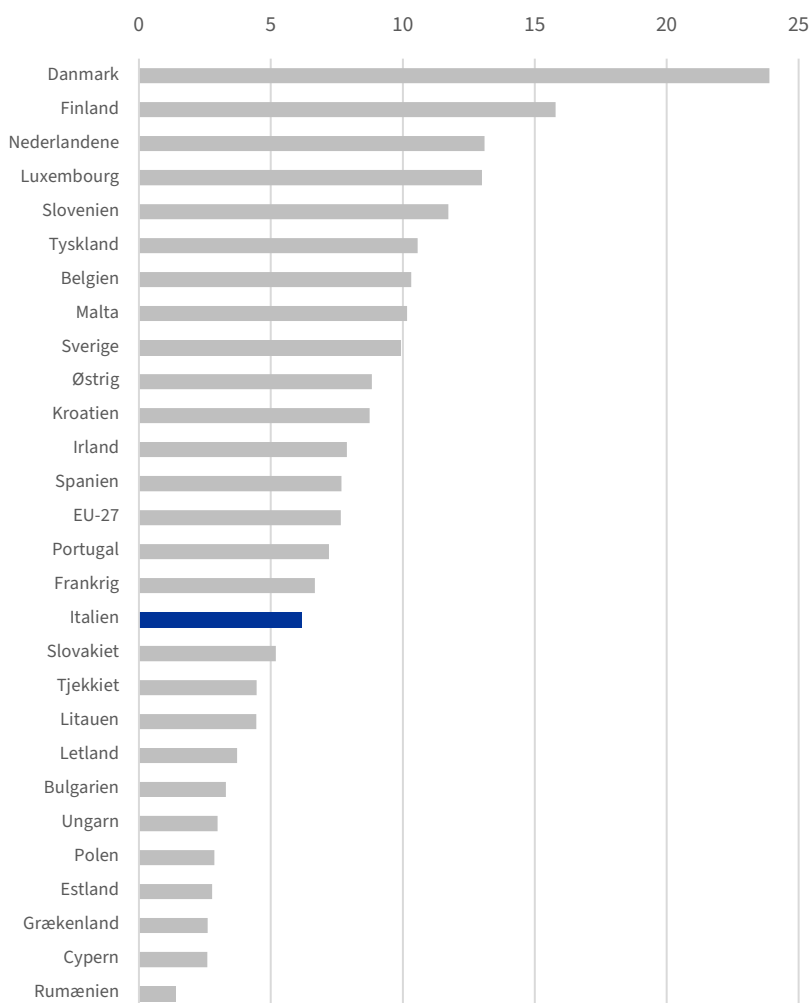
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

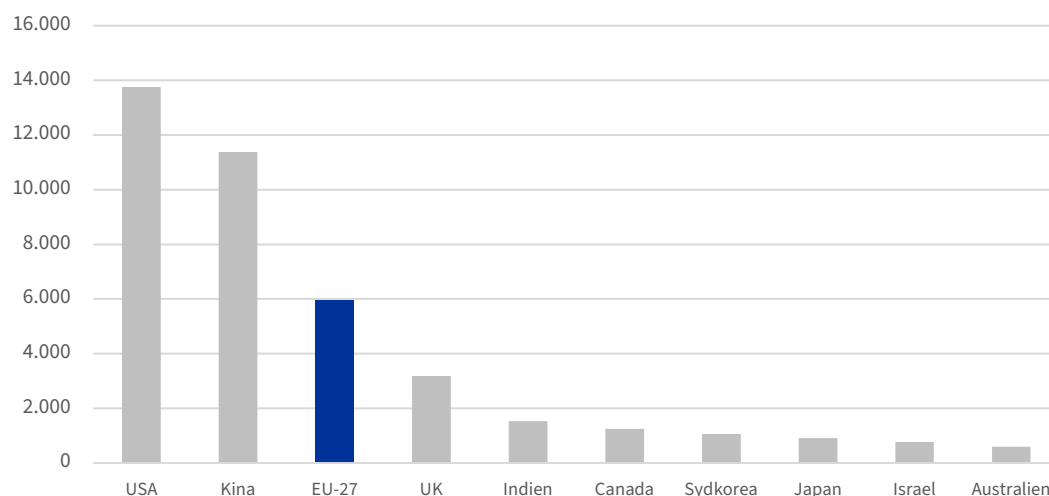
	Den Europæiske Union	Italien
Befolkning	447 695 350	58 997 201
BNP pr. indbygger i €	38 150	36 130
Arbejdsløshedsprocent	6,1 %	7,7 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	5,1 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	22,2 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

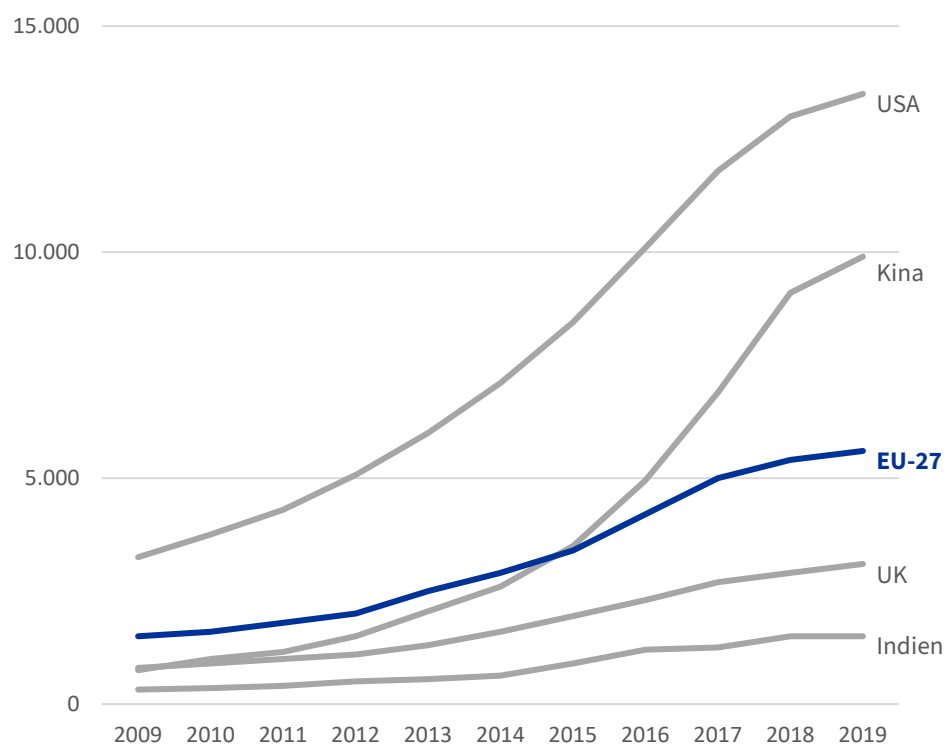


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

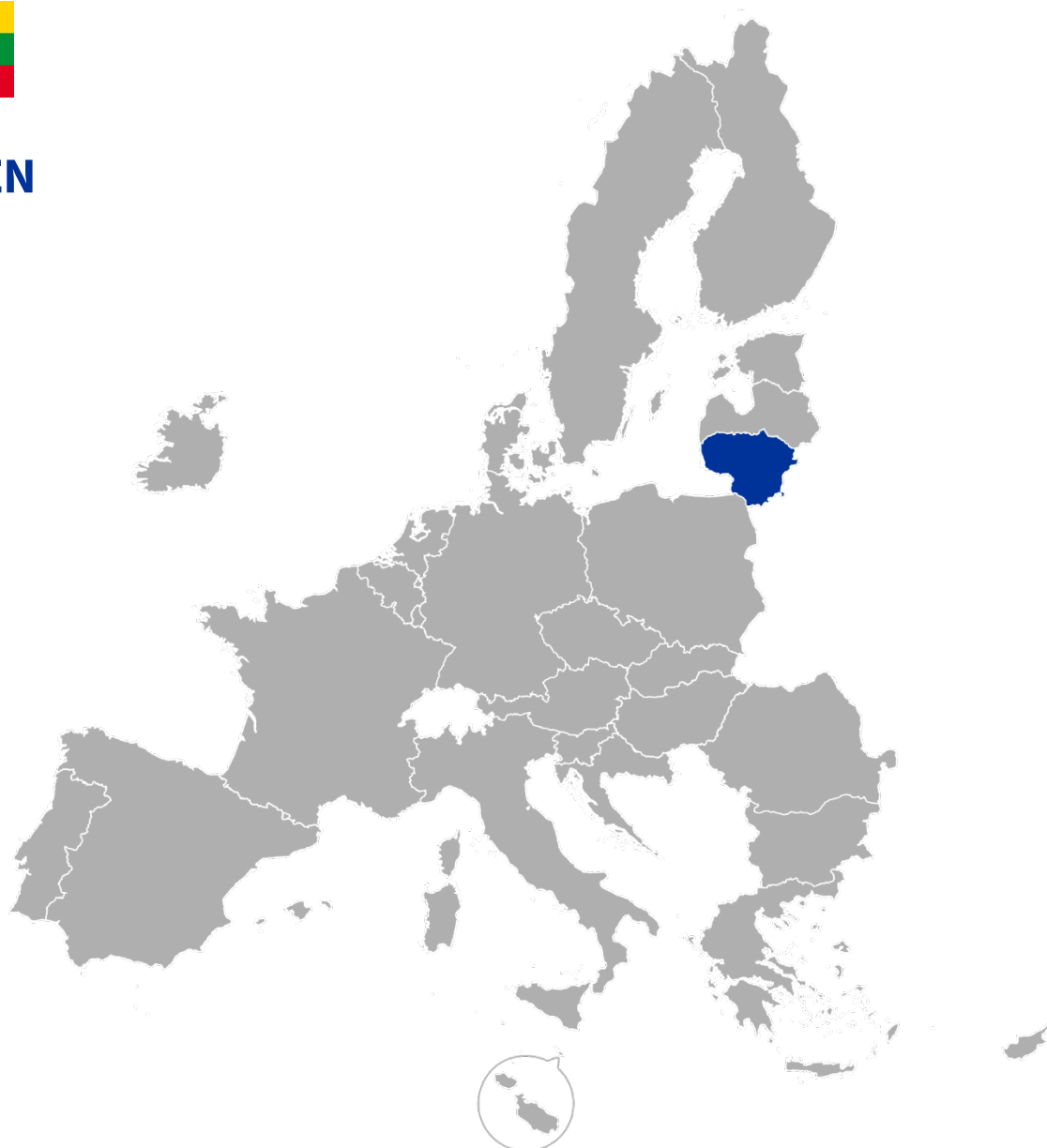
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltageres læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



LITAUEN



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



LITAUEN

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Estland, Letland, Tjekkiet og Ungarn**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Fremstillingsindustrien er den litauiske økonomis største sektor og tegner sig for 20,4 % af landets BNP. En af dens største udfordringer er lav arbejdsproduktivitet – et problem, som kunstig intelligens kan bidrage til at løse ved at automatisere rutineopgaver og optimere processer.
- Derfor argumenterer du for en ramme, der giver mulighed for fri innovation og eksperimenteren. Selv om utilsigtede konsekvenser skal undgås, understreger du, at selvevalueringer – selv med den fleksible tilgang – fortsat er obligatoriske, og udviklerne kan skræddersy dem, så de passer bedst til deres produkter.
- For meget regulering vil bremse innovationen og skade europæiske AI-opstartsvirksomheder, da obligatorisk ekstern revision og langvarige testfaser er dyre.
- Nogle medlemsstater presser på for strengere regler på grund af bekymringer for menneskerettighederne, men du argumenterer med, at AI også kan være et effektivt redskab til at beskytte disse rettigheder. Den kan bidrage til at mindske bias i forbindelse med ansættelse, politiarbejde og beslutningstagning, forsvare mod cyberangreb og beskytte personoplysninger. Og selv om AI kan generere misinformation eller deepfakes, kan den også trænes i at opdage og imødegå den slags fænomener. EU bør aktivt støtte innovation inden for disse beskyttende anvendelser af AI.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- I artikel 1 presser du på for færre restriktioner og mere frihed til at innovere – men alligevel er du på det rene med, at regler stadig er vigtige. Regulering giver mest mening, når den gælder for alle. Den giver udviklerne en klarere idé om, hvad der forventes, og hjælper EU med at vise en konsekvent tilgang til ansvarlig AI.
- Der er allerede advarselssignaler: I Kina er AI-baseret ansigts- og følelsesgenkendelsesteknologi blevet anvendt på det muslimske uighurmindretal, hvilket har ført til selvcensur og tilbageholdelse af mange i såkaldte genopdragelseslejre. I USA opsamlede Clearview AI milliarder af billeder fra sociale medieplatforme uden samtykke og solgte værktøjet til agenturer som FBI og ICE. I UK har sager, hvor Metropolitan Police har benyttet sig af ansigtsgenkendelse i realtid, ført til høje fejlprocenter, navnlig hvad angår sorte personer.
- Du mener, at disse problemer kunne være undgået med ordentlig afprøvning og regler, der gælder for alle uden undtagelse.
- Dine bemærkninger:

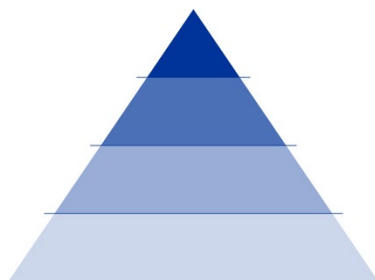
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen	kapabilitetsbaseret	fleksibel		uden undtagelse	
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

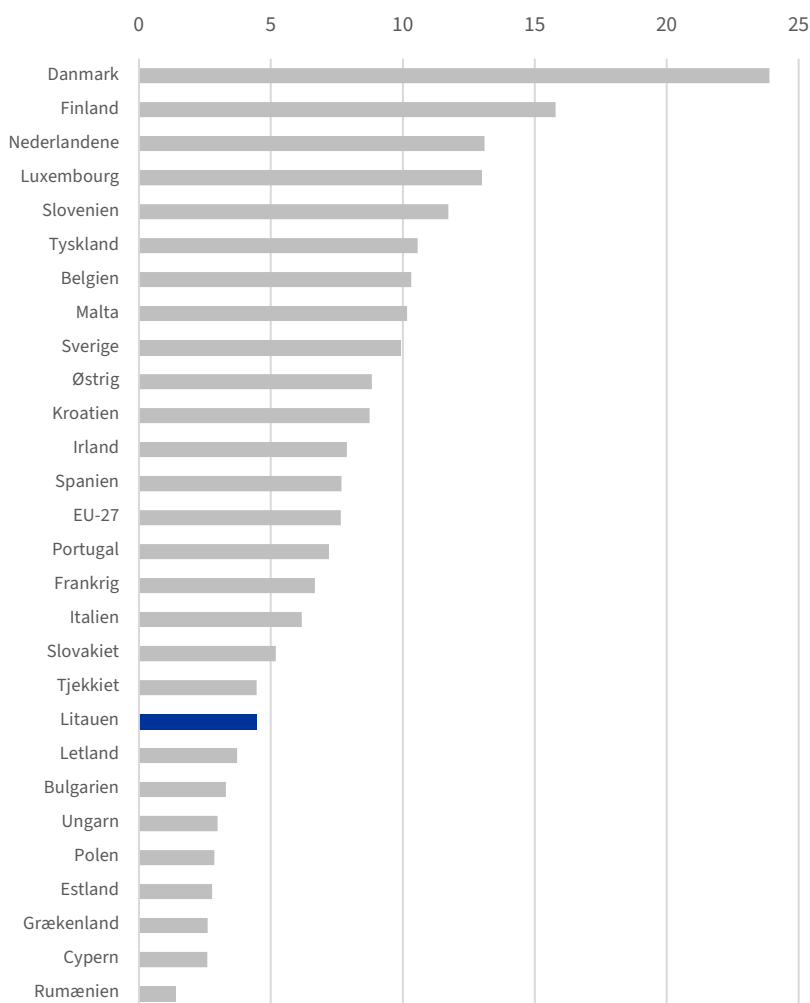
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

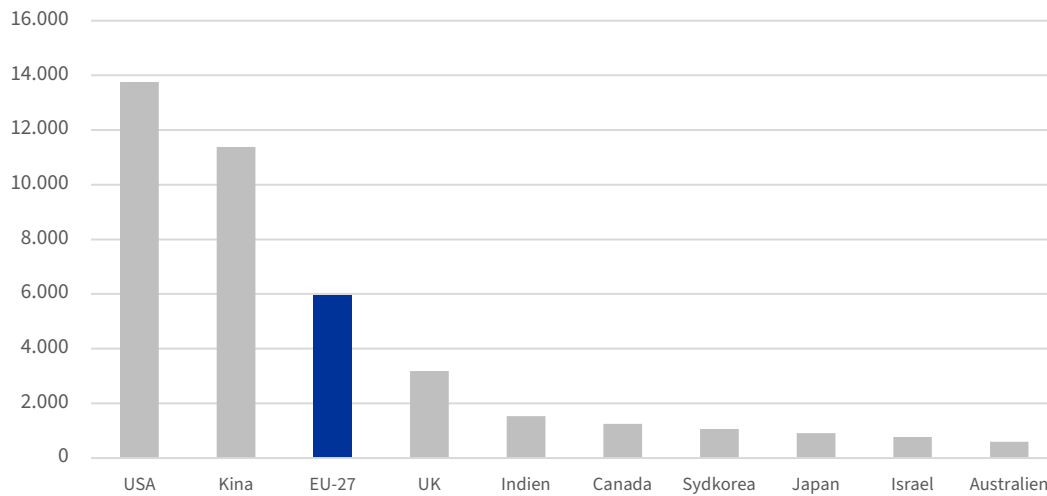
	Den Europæiske Union	Litauen
Befolkning	447 695 350	2 857 279
BNP pr. indbygger i €	38 150	25 700
Arbejdsløshedsprocent	6,1 %	6,9 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	4,9 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	25,9 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

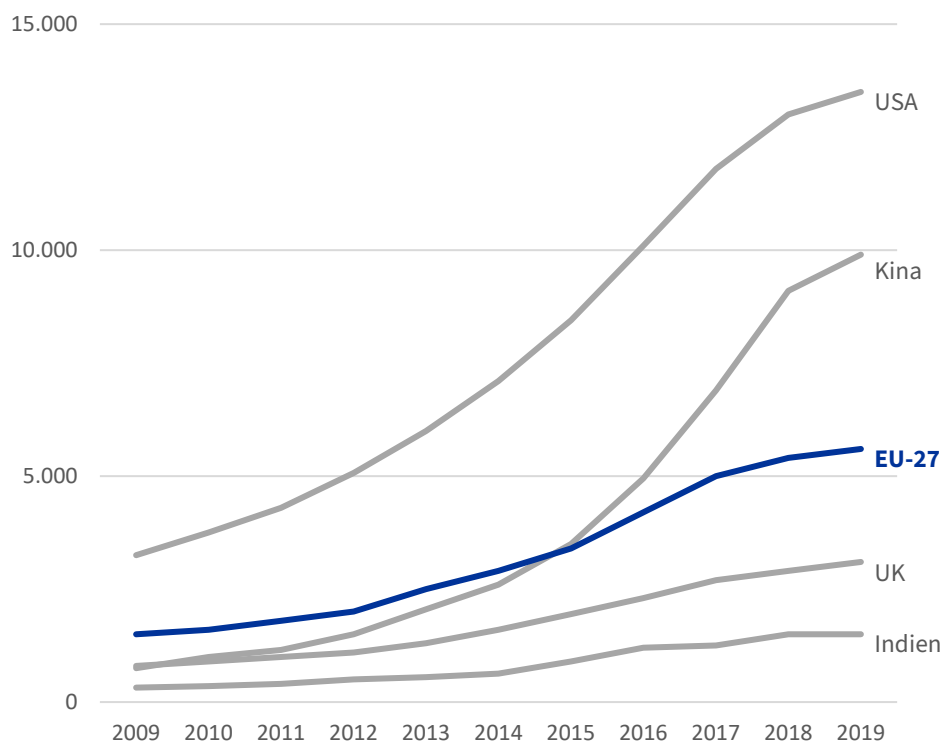


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print	ISBN 978-92-850-0498-9	doi: 10.2860/0371762	QC-01-25-006-DA-C
PDF	ISBN 978-92-850-0497-2	doi: 10.2860/0996198	QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltageres læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



LUXEMBOURG



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ **Indførelse af en AI-ramme**

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



LUXEMBOURG

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Bulgarien, Malta og Slovakiet**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.

Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Lande som Tyskland og Spanien sigter mod at fjerne selv de mindste AI-relaterede bias, hvilket du mener er urealistisk. Alle beslutninger involverer bias – Når AI udvikles ansvarligt, kan den rent faktisk mindske menneskelige bias. Du forsøger at styre debatten i en mere realistisk retning, der afspejler samtlige medlemmers behov. Hvis du bliver beskyldt for at bagatellisere AI-risici, protesterer du kraftigt: Du prioriterer bare de mest umiddelbare risici. Derfor er en formålsbaseret tilgang afgørende.
- F.eks. må AI ikke være til fare for el- eller internetsystemer, hvilket kan påvirke hospitaler og bringe liv på tværs af regioner i fare. Den kapabilitetsbaserede tilgang er for forvirrende og giver for mange smuthuller.
- Du går ind for at muliggøre hurtigere udvikling af AI, også i højrisikosektorerne. Men du er bange for, at dit land mangler den finansielle kapacitet til fuldt ud at udnytte AI-muligheder, der er underlagt rigide regler. Derfor argumenterer du for nogle undtagelser fra listen over krav (jf. bilag II), som giver mulighed for afprøvning under faktiske forhold, forudsat at ekstern revision forbliver obligatorisk. På den måde kan udviklerne ud fra deres finansielle kapacitet frit vælge, hvilke afprøvning de vil anvende.
- Hvis det er vanskeligt at opnå enighed om den formålsbaserede tilgang, **er du villig til at acceptere en revision om tre år.**

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Du mener, at AI og den digitale revolution også skal fremme ligestilling mellem kønnene i EU. Kvinder er underrepræsenteret på næsten alle områder og niveauer inden for forskning og innovation. Du understreger betydningen af aktivt at støtte deres deltagelse i udformningen af AI.
- Opstartsvirksomheder er en enestående mulighed. I modsætning til store techvirksomheder har de ofte mindre kønsforskelle og giver kvinder og genderqueers bedre muligheder inden for AI. Forskellige teams har også vist sig at opbygge mindre biaspræget og mere etisk AI – noget, EU bør tilskynde til.
- Det giver sig selv, at opstartsvirksomheder sigter mod at innovere hurtigt med begrænsede ressourcer. Tunge administrative byrder kan kvæle dette potentiale, navnlig for små teams, der forsøger at vinde fodfæste på området.
- Derfor støtter du, at små opstartsvirksomheder (under ti ansatte og 2 mio. € i omsætning) undtages fra den fulde indvirkning af EU's AI-forordning.
- Dine bemærkninger:

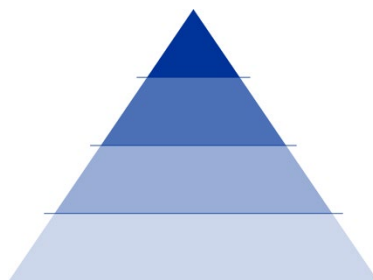
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg	formålsbaseret	fleksibel	revision om tre år	opstartsvirksomheder	
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

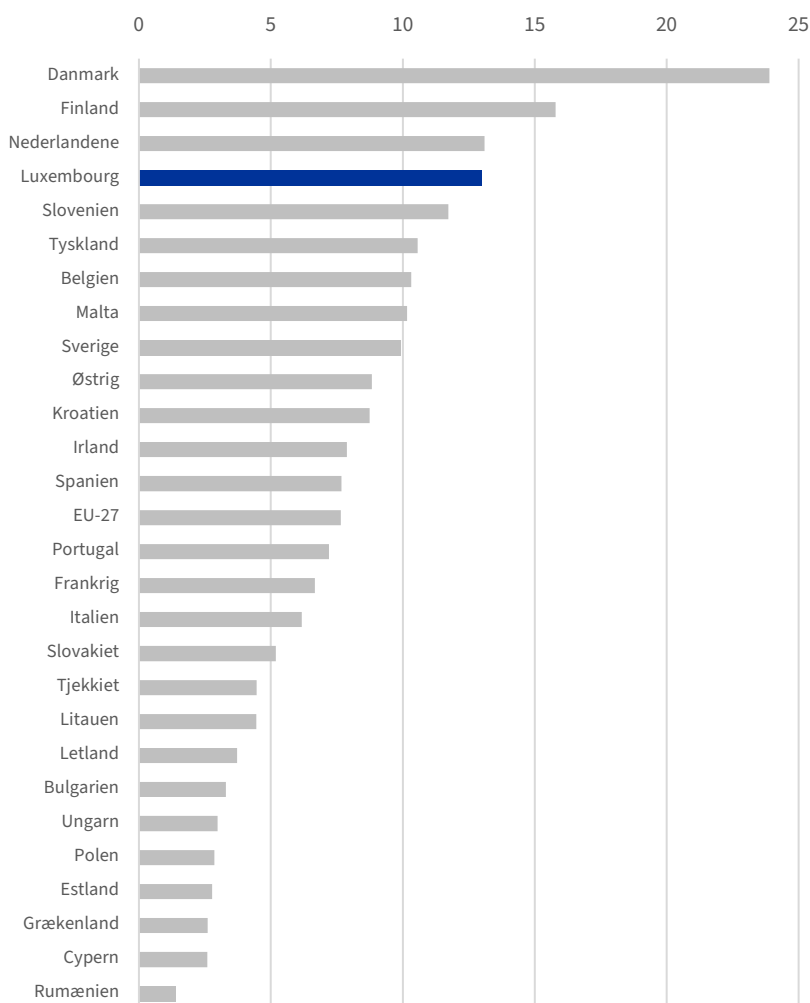
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

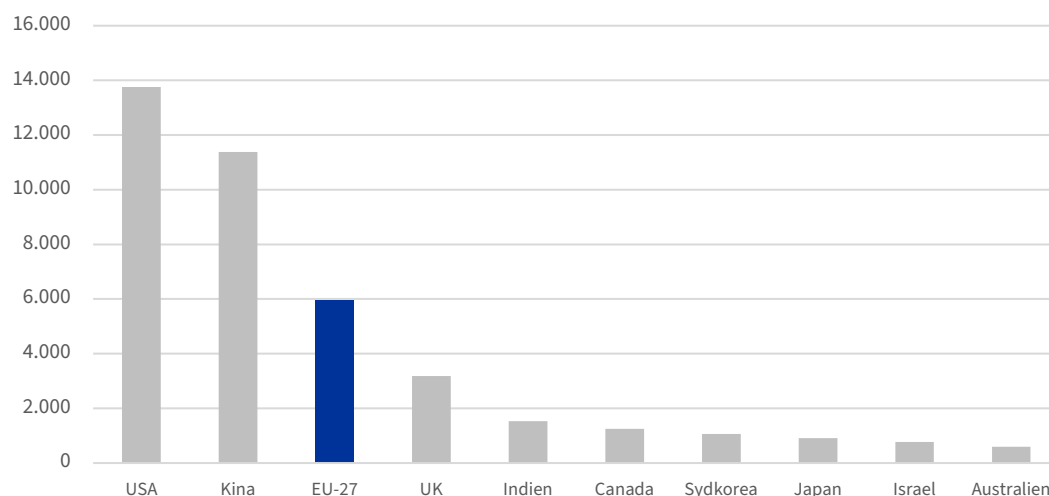
	Den Europæiske Union	Luxembourg
Befolkning	447 695 350	660 809
BNP pr. indbygger i €	38 150	121 290
Arbejdsløshedsprocent	6,1 %	5,2 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	14,5 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	27,9 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

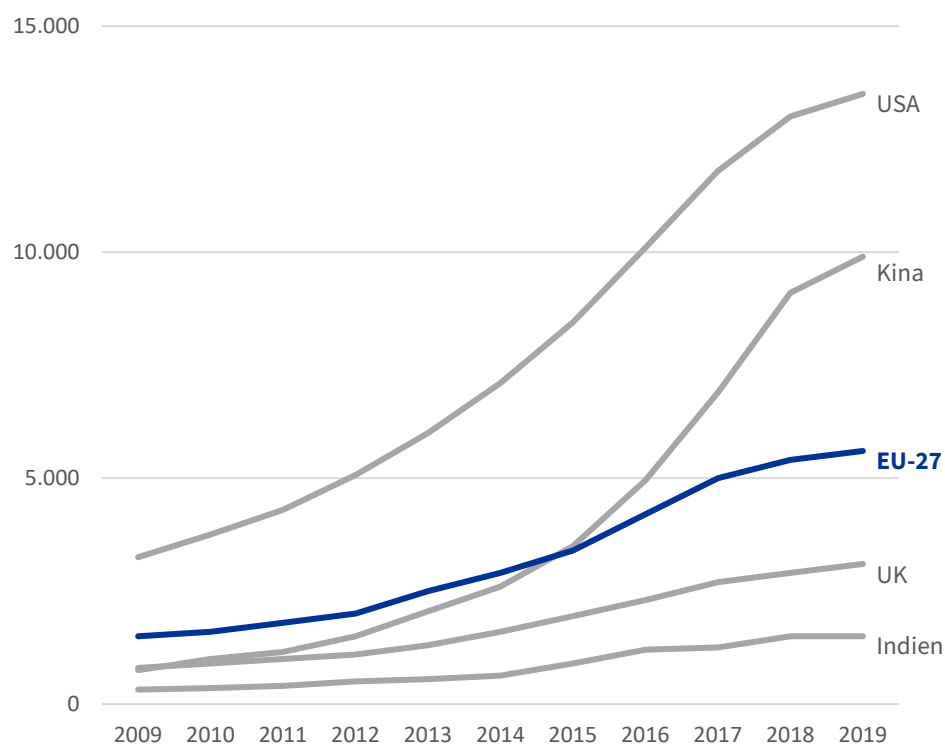


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

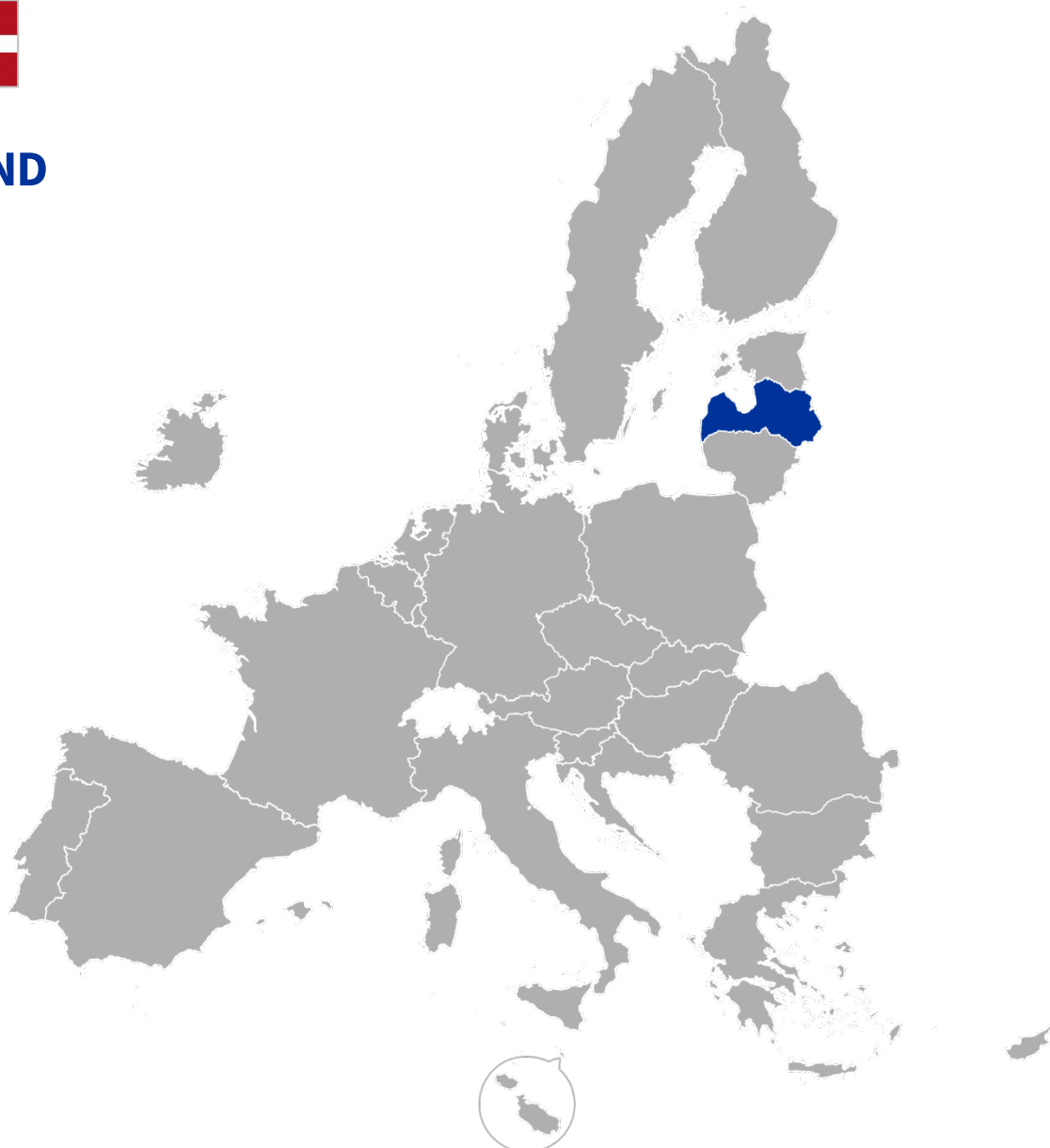
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



LETLAND



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



LETLAND

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Estland, Litauen, Tjekkiet og Ungarn**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.

Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____.

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____.

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____.

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____.

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... **kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.**



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... **fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.**



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Du er fast besluttet på at sikre, at der gennemføres digitale løsninger i Letland. Dette skaber nye krav om beskyttelse af følsomme oplysninger og digital infrastruktur, som er afgørende for at opretholde vitale samfundsmæssige funktioner.
- Med hensyn til klassificering af højrisiko-AI-modeller er den kapabilitetsbaserede tilgang mere velegnet end den formålsbaserede tilgang, som Kommissionen har foreslået. Hvis det samme AI-system betragtes som lav risiko i én sammenhæng og højrisiko i en anden, kan det skabe forvirring og retsikkerhed for brugere, udviklere og investorer.
- For at tiltrække de nødvendige investeringer i AI er det afgørende at have den korrekte ramme på plads. For dig er det en ramme, der beskytter de grundlæggende rettigheder grundigt og samtidig giver mulighed for frie og fleksible innovationsprocesser. Hvis EU vil konkurrere med AI-kraftcentre som USA og Kina, må det udvise ekspertise. En sådan ekspertise kan ikke udfolde sig under alt for restriktive regler, der begrænser eksperimenteren og udvikling. Af disse grunde vil du foreslå visse undtagelser fra listen over krav til højrisiko-AI-systemer (jf. bilag II). Du mener, at **det vil være bedst at fjerne "overvågning efter markedsføring" som et obligatorisk krav**, da det er uklart, hvad denne forpligtelse præcist indebærer.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Generelt er du enig med de lande, der argumenterer for, at der ikke skal være undtagelser, bl.a. din nabo og stærke allierede Litauen. Du synes, det er fornuftigt at støtte en fælles europæisk ramme for AI-regulering uden unødvendige fravigelser.
- Siden den russiske invasion af Ukraine har Letland og de andre baltiske stater imidlertid været i højt alarmberedskab af frygt for et potentielt russisk angreb på EU. Letland deler endda en grænse med Rusland, som Letland hver dag arbejder på at beskytte sammen med NATO-styrker.
- Derfor er du åben over for at drøfte begrænsede undtagelser for militære og forsvarsmæssige formål – men ikke for national sikkerhed i bredere forstand. Selv om national sikkerhed hører under medlemsstaternes kompetence, risikerer det at skabe smuthuller og lovgivningsmæssig fragmentering, hvis den holdes helt uden for EU's rammer. En minimal, grundlæggende EU-tilgang kan fremme konsekvens (og klarhed for udviklerne) og signalere en sammenhængende holdning til national sikkerhedsrelateret AI i hele Unionen.
- Dine bemærkninger:

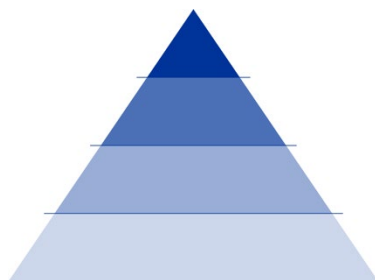
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland	kapabilitetsbaseret	fleksibel	fjerne overvågning efter markedsføring	militær/forsvar	
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

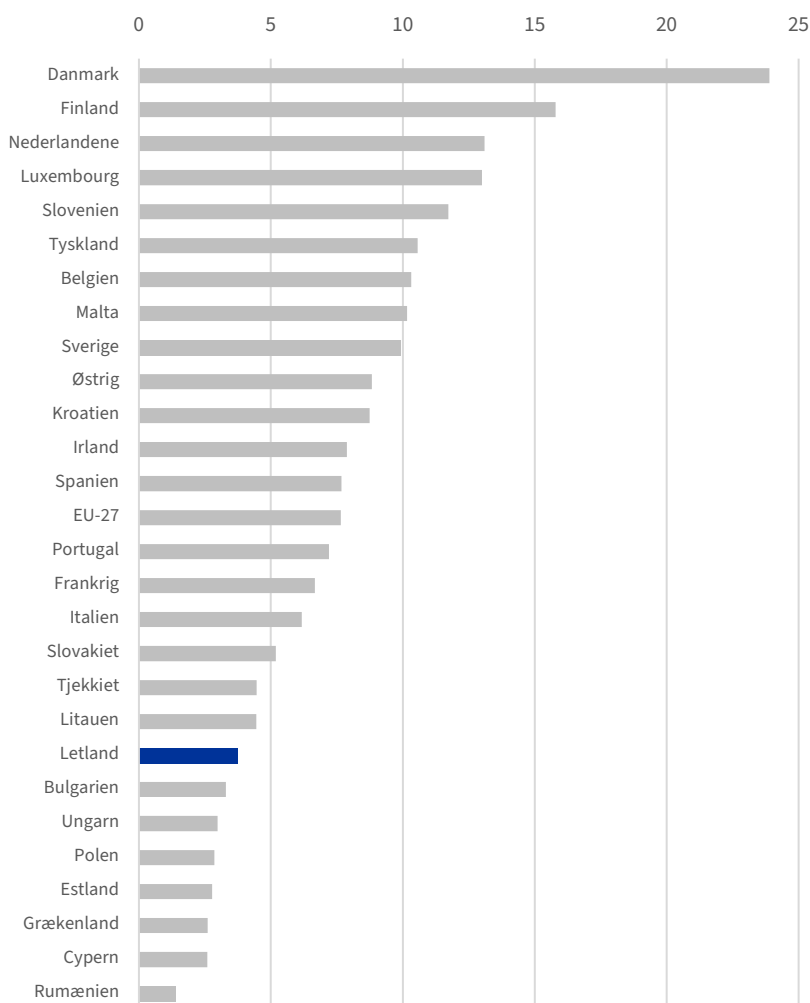
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

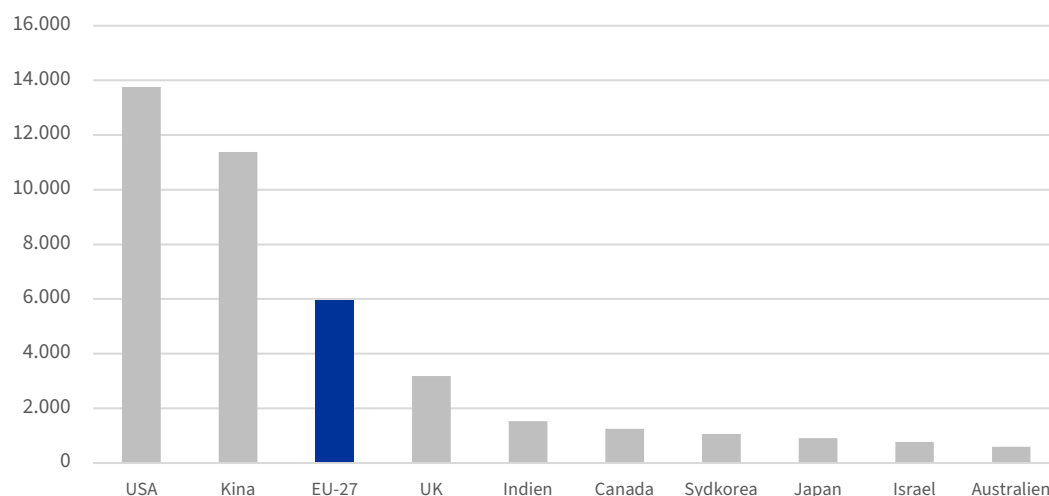
	Den Europæiske Union	Letland
Befolkning	447 695 350	1 883 008
BNP pr. indbygger i €	38 150	20 930
Arbejdsløshedsprocent	6,1 %	6,5 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	4,5 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	16,6 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

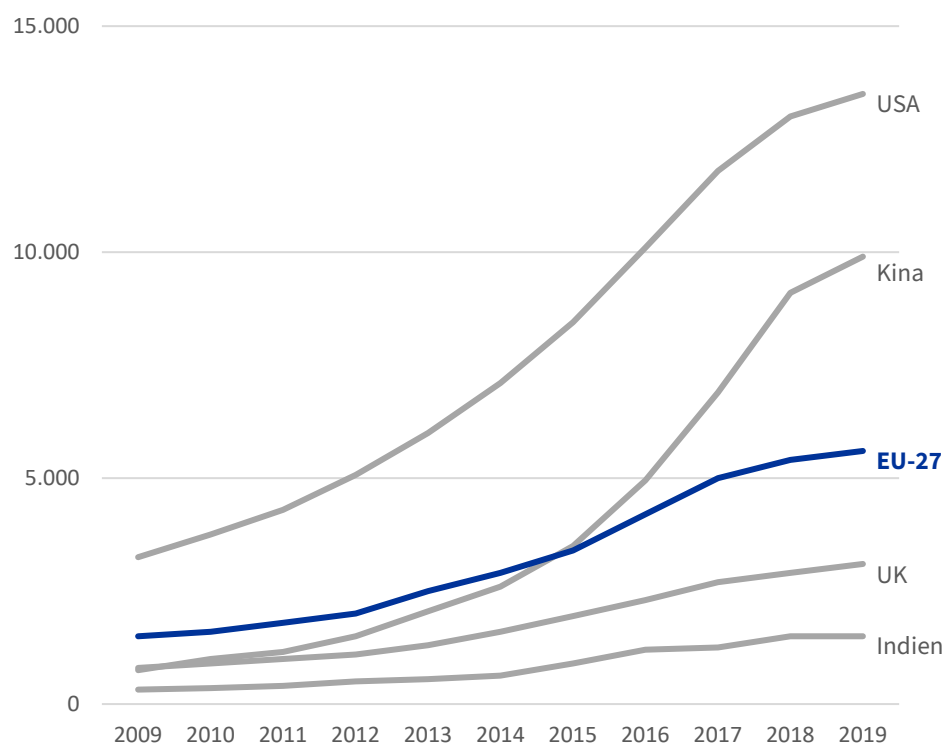


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



MALTA



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Bulgarien, Luxembourg og Slovakiet**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____.

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____.

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____.

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____.

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Som et af EU's mindste medlemmer med kun 550 000 indbyggere og som ønation er dit land stærkt engageret i bæredygtig økonomisk vækst. Du sigter mod at tiltrække udenlandske investeringer og kvalificerede talenter til at udvikle en robust AI-sektor med fokus på medicin og kritisk infrastruktur – områder, der ventes at give store fortjenester.
- Du støtter Kommissionens forslag til en horisontal AI-ramme og især dens formålsbaserede risikostyringstilgang, som du mener er logisk og hensigtsmæssigt, fordi det er lettere for reguleringsmyndighederne at vurdere specifikke applikationer end at evaluere en models abstrakte kapabiliteter.
- I dag er det afgørende at understrege, at teknologisk kreativitet kræver frihed – innovation bør omfavnes, ikke kvæles. For at støtte dette **er det vigtigt at fremme afprøvning under faktiske forhold af nye AI-modeller**, da sandkassemiljøer, selv om de er nyttige, ofte ikke formår fuldt ud at replikere faktiske forhold.
- Små lande som Malta står over for unikke udfordringer: begrænsede talentpuljer og færre ressourcer til dyre test. Derimod kan større, kapitalstærke stater lettere klare den slags byrder. Derfor er du stærkt imod alt for rigide testkrav, der ville lægge strukturelle hindringer i vejen for dine nationale ambitioner inden for AI-udvikling.

Artikel 2

Anvendelsesområde

Forskrifterne som fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Det er afgørende for dig, at spørgsmål vedrørende national sikkerhed, militær eller forsvar ikke bliver omfattet af AI-forordningen. Navnlig er national sikkerhed ikke en EU-kompetence. Det fremgår af artikel 4, stk. 2, i traktaten om Den Europæiske Union (TEU): "Den nationale sikkerhed forbliver den enkelte medlemsstats eneansvar." I lyset heraf giver det ikke mening at begrænse medlemsstaterne med hensyn til, hvordan de regulerer AI i nationale sikkerhedsapplikationer.
- På den anden side ser du ingen grund til at udelukke opstartsvirksomheder fra forordningen. Bare fordi en virksomhed er lille, betyder det ikke, at dens teknologi er uskadelig. Og hvis de undtages, vil store virksomheder måske bare overføre risikable projekter til dem for at undgå testkravene. Det er et smuthul, der bare venter på at blive udnyttet, og EU bør sørge for, at det ikke sker.
- Dine bemærkninger:

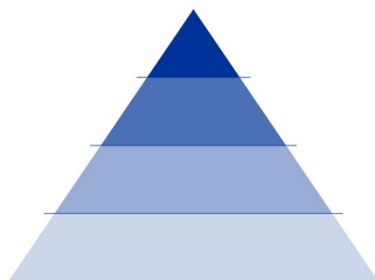
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta	formålsbaseret	fleksibel	afprøvning under faktiske forhold	national sikkerhed	
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

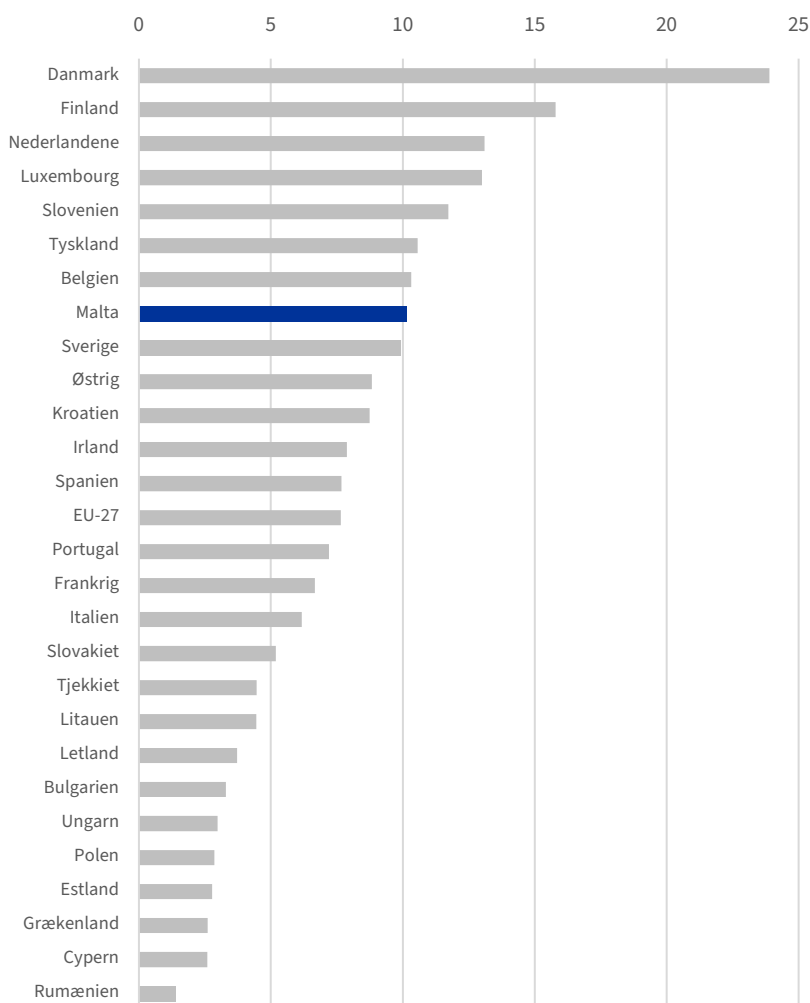
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

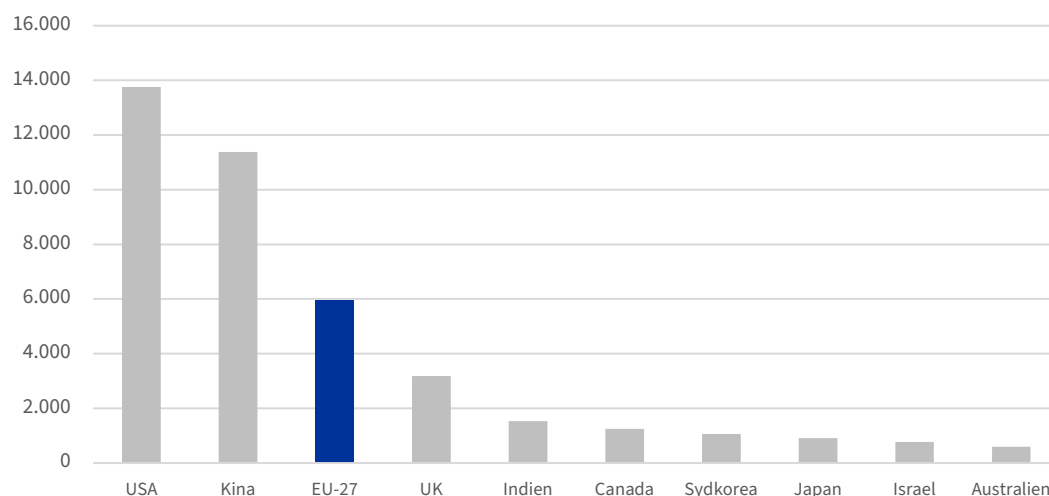
	Den Europæiske Union	Malta
Befolkning	447 695 350	542 051
BNP pr. indbygger i €	38 150	37 110
Arbejdsløshedsprocent	6,1 %	3,5 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	13,2 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	37 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

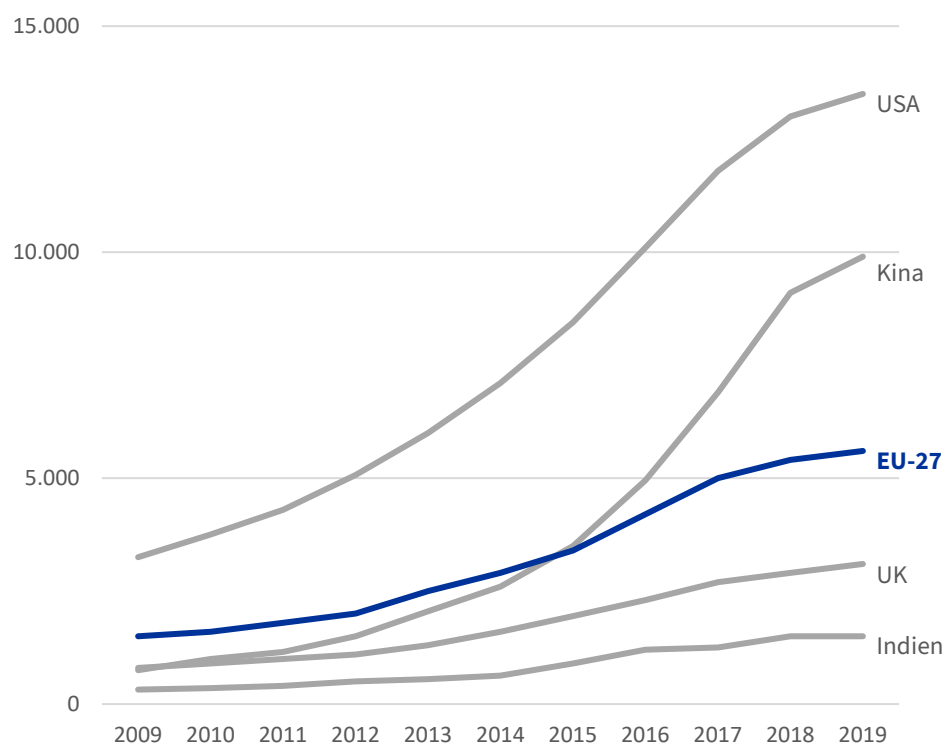


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

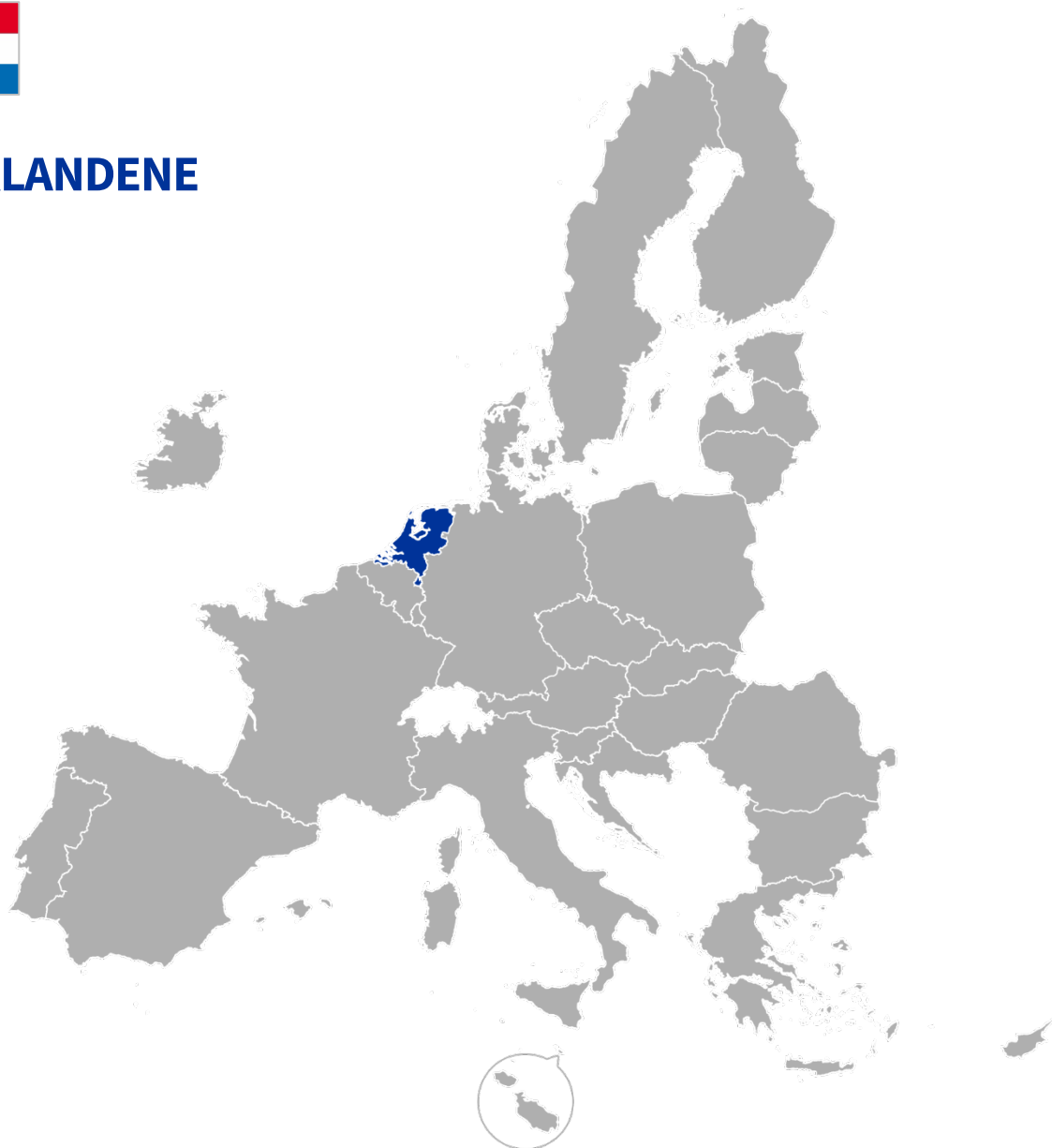
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltageres læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



NEDERLANDENE



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



NEDERLANDENE

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Italien, Kroatien og Slovenien**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Fastholdelse af nationale værdier og økonomisk velstand er en topprioritet for din regering. Ifølge IMF kan op til 60 % af alle job blive påvirket af AI, hvilket kan føre til fyringer og et stigende behov for medarbejdere med AI-kundskaber.
- Ukontrolleret arbejdsløshed vil være en alvorlig trussel for Nederlandenes økonomi, og derfor går du stærkt ind for en ramme for forsigtighedstest, så det sikres, at højrisikoteknologier ikke svigter og forstyrrer hele sektorer.
- Du argumenterer også for større AI-færdigheder i hele det nederlandske samfund, så folk kan genkende bias eller manipulation og ikke stoler blindt på automatiserede systemer. Den kapabilitetsbaserede tilgang er mere velegnet til dette formål, da den fokuserer på de faktiske risici for misbrug og manipulation: Visse kapabiliteter – f.eks. ansigtsgenkendelse – er tilbøjelige til at blive udsat for misbrug eller fejl uanset sammenhængen.
- Et kapabilitetsbaseret system undgår at underregulere den slags teknologier, bare fordi de anvendes i lavrisikosektorer. Et detail-AI-system, der sporer kundebevægelser, forudser adfærd og bruger ansigtsgenkendelse til at skræddersy reklamefremstød eller markere kunder med høj værdi, kan f.eks. gå fri af kontrol under en sektorbaseret model. Men de underliggende kapabiliteter giver anledning til alvorlig bekymring med hensyn til privatlivets fred, profilering og forskelsbehandling.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Du støtter en ramme for forsigtighedstest i forbindelse med artikel 1, men du går ind for undtagelser for formål, der vedrører national sikkerhed. National sikkerhed er ikke en EU-kompetence. Det fremgår af artikel 4, stk. 2, i traktaten om Den Europæiske Union (TEU): "Den nationale sikkerhed forbliver den enkelte medlemsstats eneansvar." I lyset heraf giver det ikke mening at begrænse medlemsstaterne med hensyn til, hvordan de regulerer AI i nationale sikkerhedsapplikationer.
- På den anden side ser du ingen grund til at udelukke opstartsvirksomheder fra forordningen. Bare fordi en virksomhed er lille, betyder det ikke, at dens teknologi er uskadelig. Og hvis de undtages, vil store virksomheder måske bare overføre risikable projekter til dem for at undgå testkravene. Det er et smuthul, der bare venter på at blive udnyttet, og EU bør sørge for, at det ikke sker.
- Dine bemærkninger:

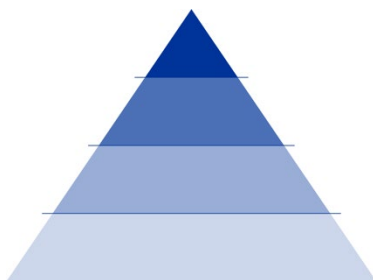
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene	kapabilitetsbaseret	forsigtigheds-		national sikkerhed	
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

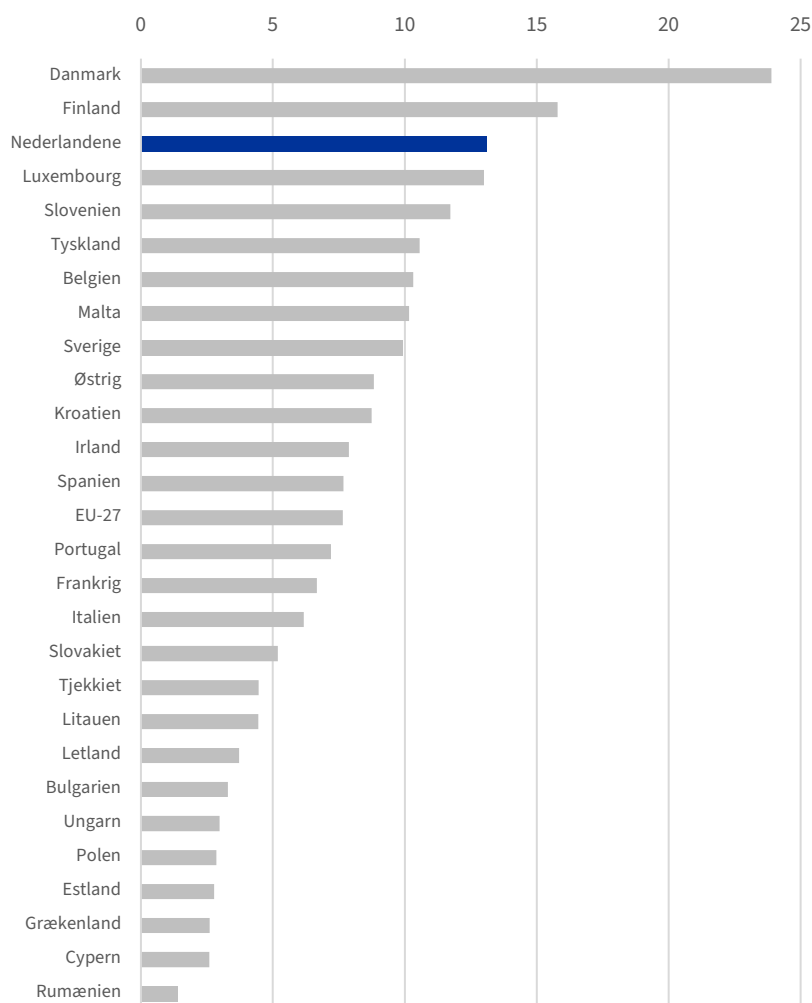
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

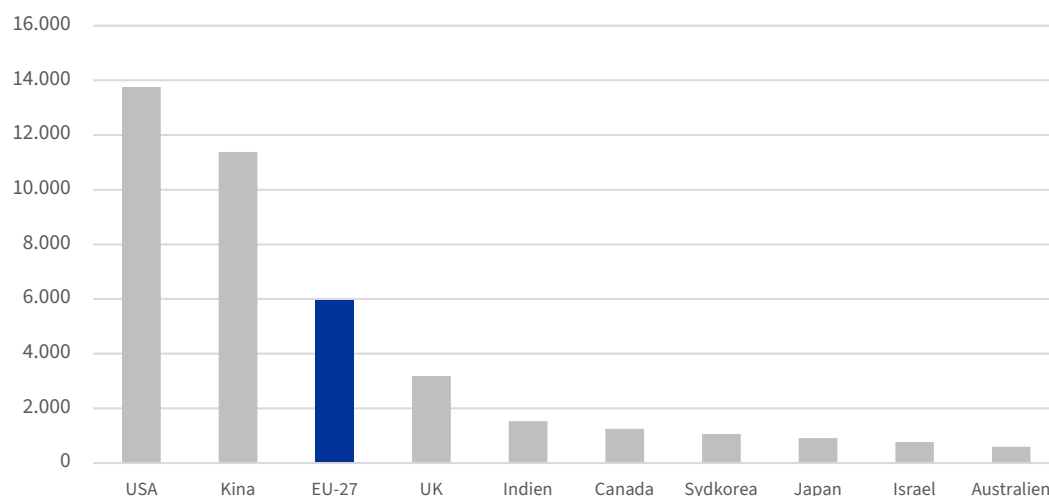
	Den Europæiske Union	Nederlandene
Befolkning	447 695 350	17 811 291
BNP pr. indbygger i €	38 150	59 720
Arbejdsløshedsprocent	6,1 %	3,6 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	14,1 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	54,5 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

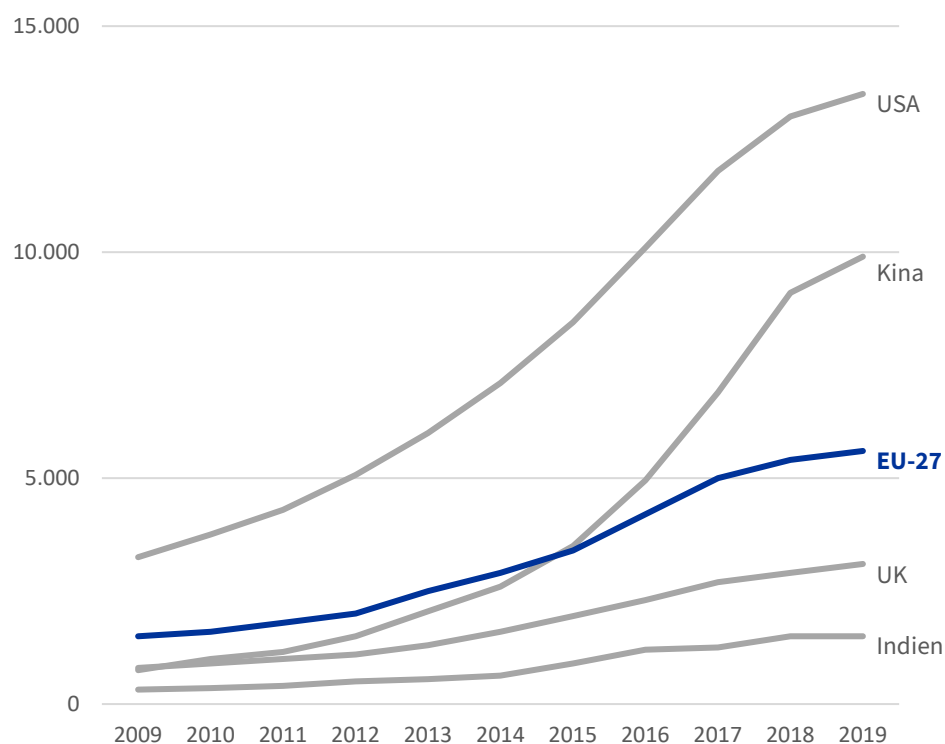


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

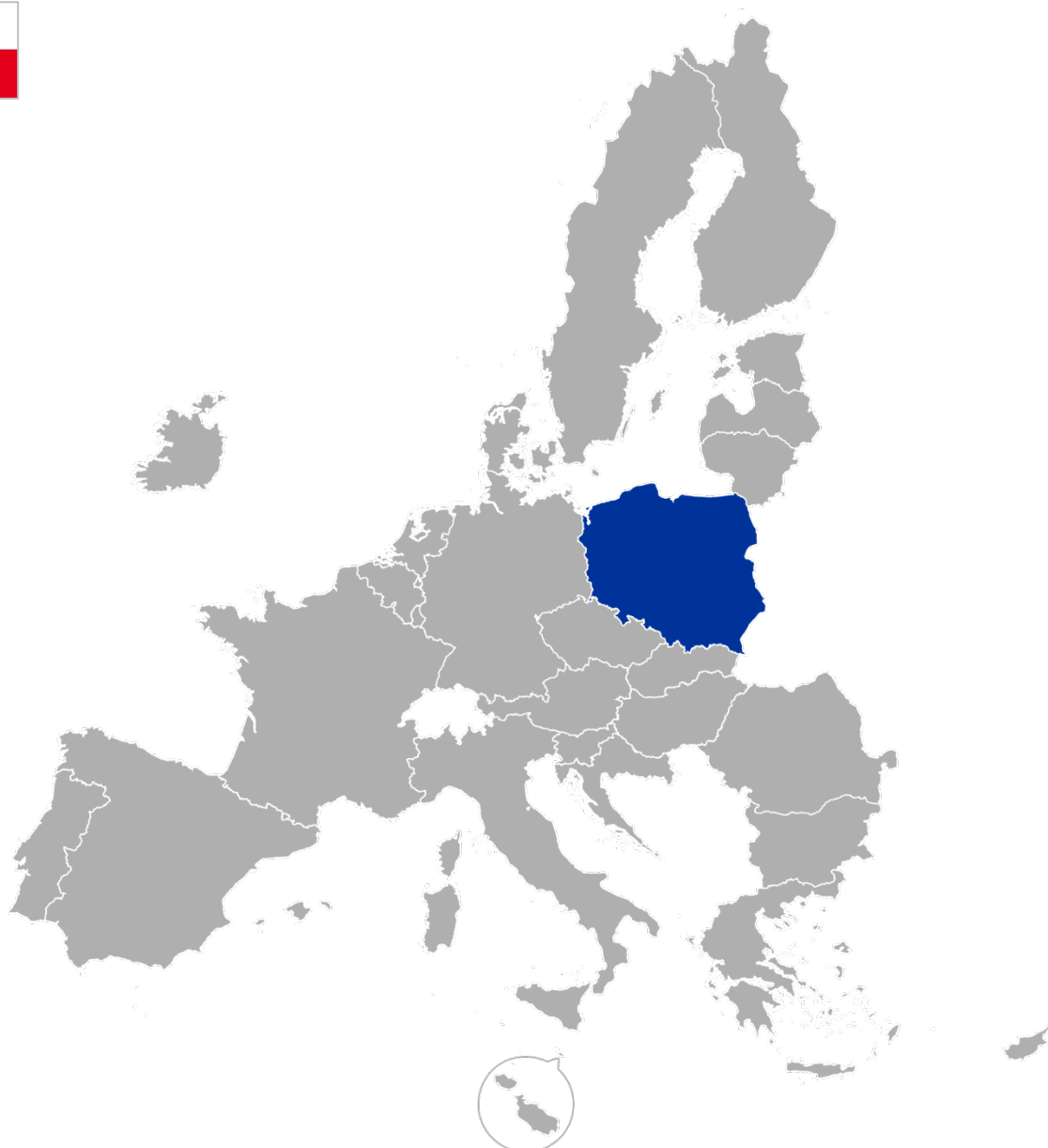
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



POLEN



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ **Indførelse af en AI-ramme**

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



POLEN

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Danmark, Frankrig, Rumænien og Østrig**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____.

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____.

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____.

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____.

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Polen har i øjeblikket kun få globale techgiganter, og det betyder, at AI giver små og mellemstore virksomheder store muligheder for at vokse. Landets økonomi – der bygger på energi, industri, landbrug og logistik – er godt rustet til at drage fordel af indførelsen af AI.
- Du støtter den formålsbaserede tilgang, som skaber klare incitament: Lavrisikosektorer får større fleksibilitet med hensyn til innovation, mens højrisikosektorer skal opfylde høje EU-standarder. Det vil styrke tilliden til europæisk produceret AI, som er internationalt kendt for sin kvalitet og sikkerhed.
- Du anerkender, at AI kan føre til bias, forskelsbehandling og risici for privatlivets fred, da intet datasæt er helt neutralt. En forsigtighedstilgang til test sikrer, at der gælder strengere regler, der hvor sandsynligheden for skade er størst, hvilket bidrager til at forebygge alvorlig indvirkning og beskytte de grundlæggende rettigheder.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Du argumenterer for, at militære og forsvarsrelaterede applikationer udelukkes fra AI-forordningen, og understreger, at overregulering på dette område kan hindre hurtig ibrugtagning af AI-teknologier i konfliktområder eller i forbindelse med kriserespons. I forbindelse med den aktuelle regionale ustabilitet er fleksibilitet i anvendelsen af den slags teknologier afgørende.
- Som et land, der grænser op til Rusland, Ukraine og Belarus, er Polen dybt bekymret over den nationale sikkerhed og landets borgeres sikkerhed. Desuden står Polen langs grænsen til Belarus over for betydelige udfordringer med hensyn til at håndtere de øgede flygtningebewægelser.
- I lyset af dette pres søger Polen større forståelse og støtte fra andre EU-medlemsstater, hvad angår landets unikke geopolitiske og sikkerhedsmæssige udfordringer.
- Dine bemærkninger:

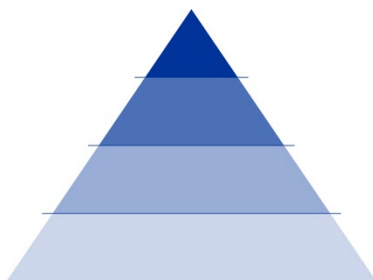
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen	formålsbaseret	forsigtigheds-		militær/forsvar	
Portugal					
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

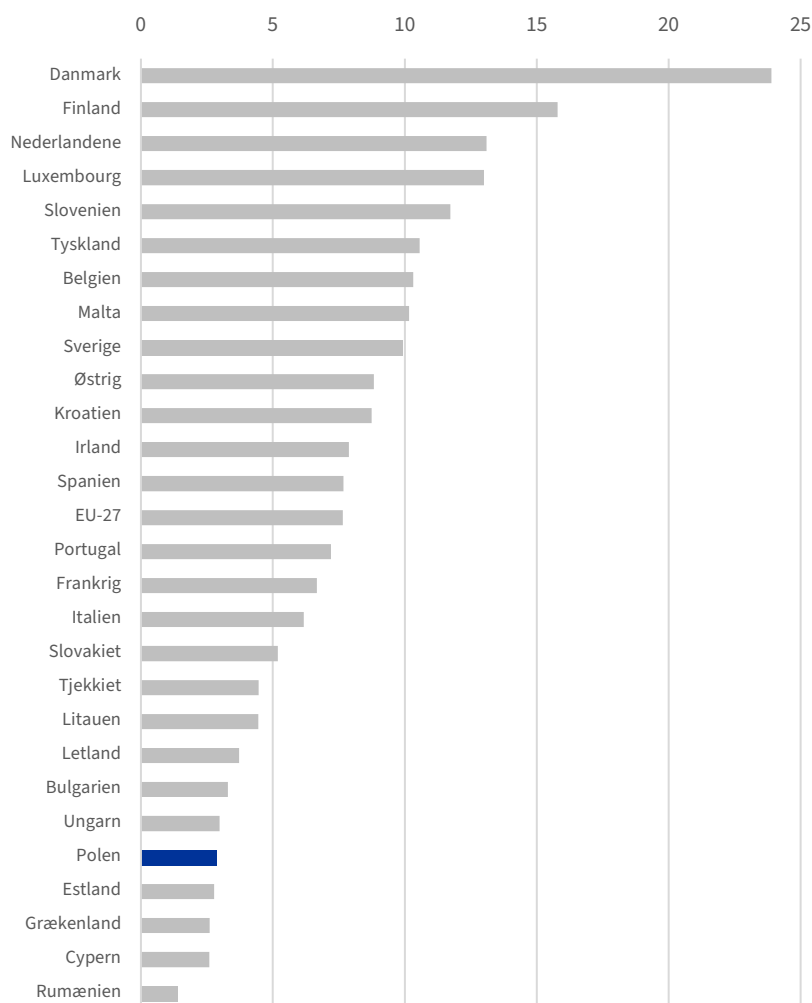
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

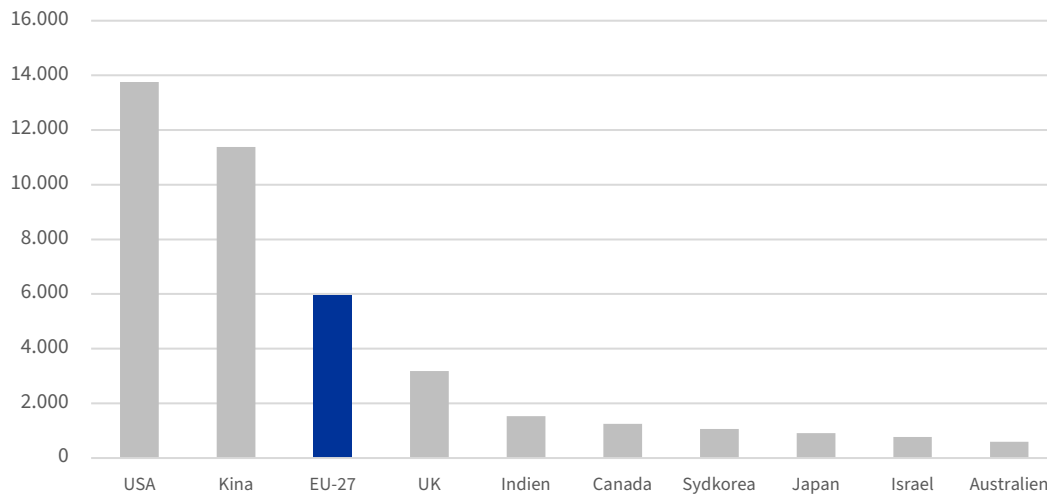
	Den Europæiske Union	Polen
Befolkning	447 695 350	36 753 736
BNP pr. indbygger i €	38 150	19 980
Arbejdsløshedsprocent	6,1 %	2,8 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	3,7 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	20,1 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

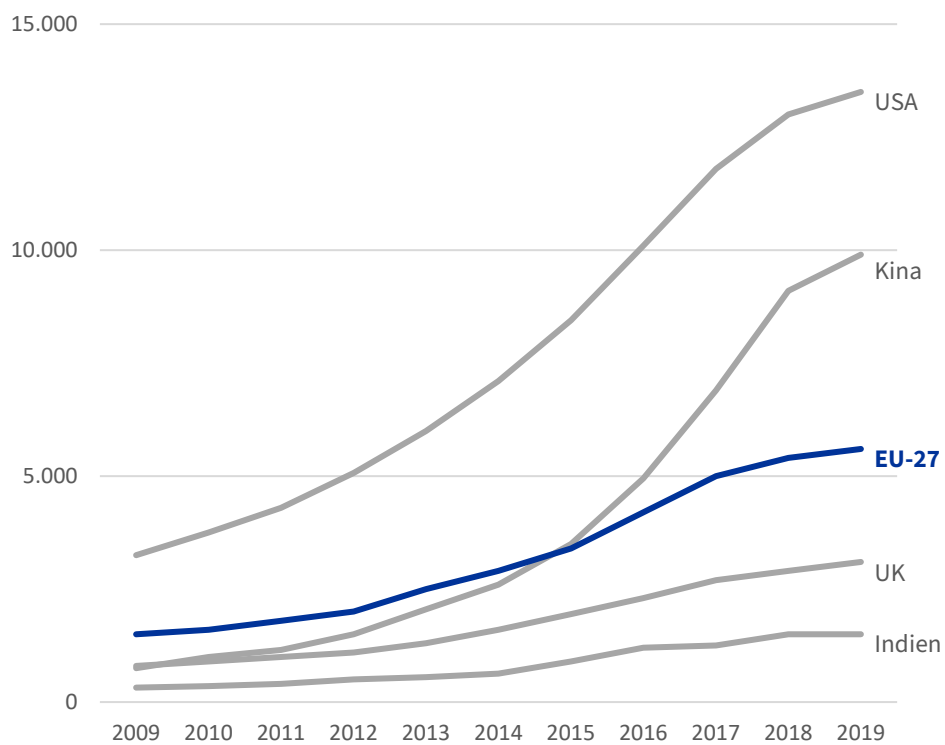


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

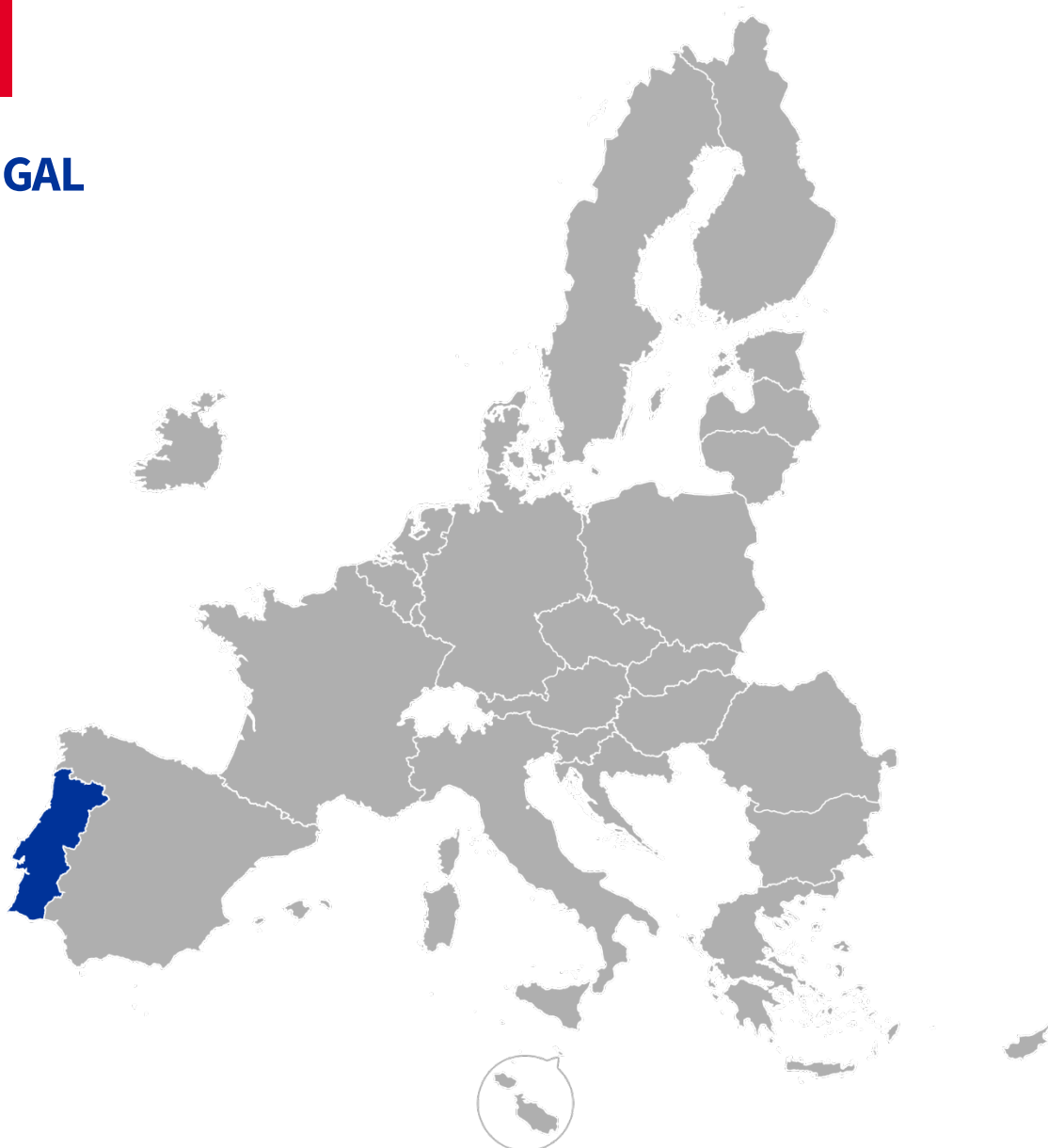
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltageres læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



PORTUGAL



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



PORTUGAL

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **delar nogle af argumenterne med lande som Finland, Irland og Sverige**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Portugal er et meget attraktivt mål for internationale hightechvirksomheder på grund af arbejdstagernes tekniske færdigheder. Et eksempel er TEKEVER, en virksomhed, der udvikler AI-modeller til autonom droneovervågning i både kommercielle og militære sammenhænge. Kommercielt anvendes dronerne til at overvåge rørledninger for lækager eller sabotagehandling og beskytter dermed kritisk infrastruktur. EU's Agentur for Søfartssikkerhed (EMSA) anvender også TEKEVER-systemer til at støtte kystvagtoperationer og miljøbeskyttelsesindsatser.
- Du går ind for den formålsbaserede tilgang, men du er bange for, at uklare sektorbeskrivelser vil få EU's AI-forordning til at virke svag og forvirrende. Du vil gerne have mere præcise beskrivelser af begrebet "kritisk infrastruktur", f.eks. ved at det gøres klart, at det omfatter elproduktion, telekommunikation, vandforsyning, landbrug, opvarmning og transport.
- Du er bange for, at en ramme for forsigtighedstest vil øge den finansielle og administrative byrde på AI-udviklerne.
- Du håber, at der opnås enighed om den endelige tekst til EU's AI-forordning i dag som en ledetråd for troværdig AI-udvikling i Europa, men du er stærkt imod en indbygget revision om få år, da det vil underminere investorernes tillid. Hvis der stadig er usikkerhed, **vil det være mere ansvarligt at udsætte afstemningen end foregribende at planlægge at omskrive reglerne.**

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Selv om du støtter fleksible bestemmelser i artikel 1, mener du stadig, at der bør gælde klare regler for alle sektorer – også nationale sikkerheds- og efterretningssektorer. Hvis de undtages, opstår der smuthuller, og der er større risiko for, at uetisk eller skadelig AI får frit spil.
- Dette handler ikke om at blokere AI i national sikkerhed, men om at sikre, at den er designet ansvarligt for at undgå bias eller forskelsbehandling.
- Den eneste undtagelse, du kan godtage, er open source-modeller: Disse modeller sætter gang i innovation, fordi de udvikles på en gennemsigtig måde. Udviklerne eksperimenterer og uddanner sig sammen og deler viden for at forbedre værktøjerne. Open source-AI er ofte mere inspicerbar og kontrollerbar end ejendomsretligt beskyttede modeller, hvilket er i tråd med retsaktens mål.
- Dine bemærkninger:

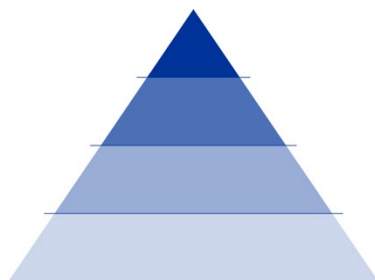
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal	formålsbaseret	fleksibel	udsætte afstemningen	open source	
Rumænien					
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

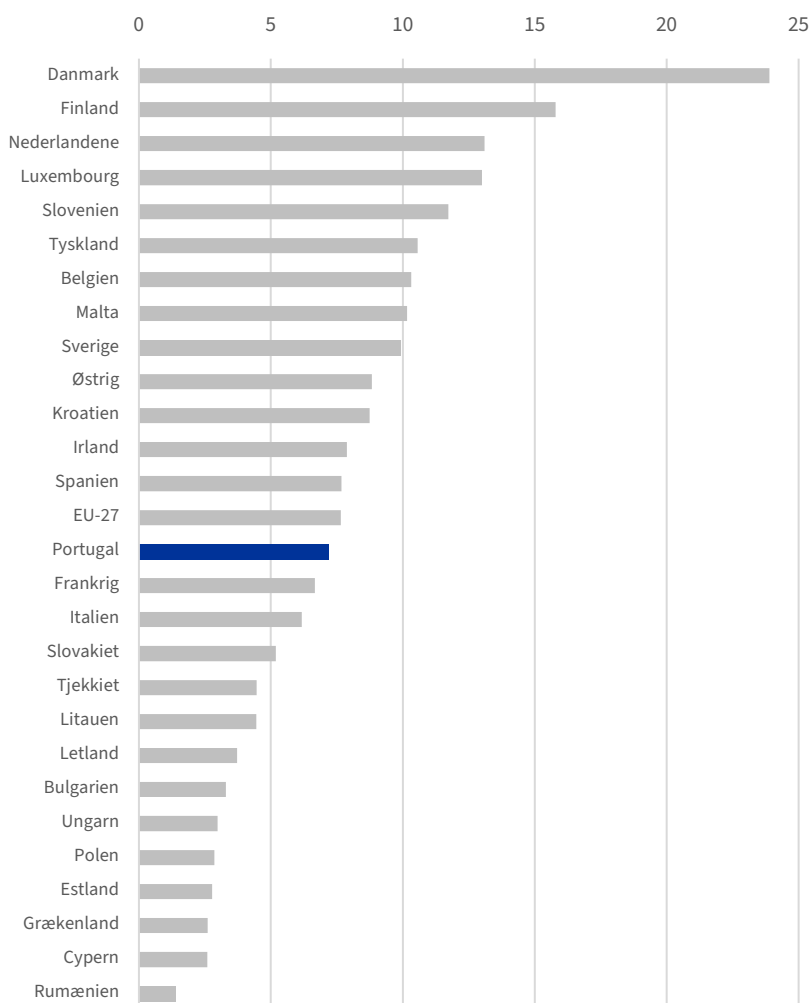
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

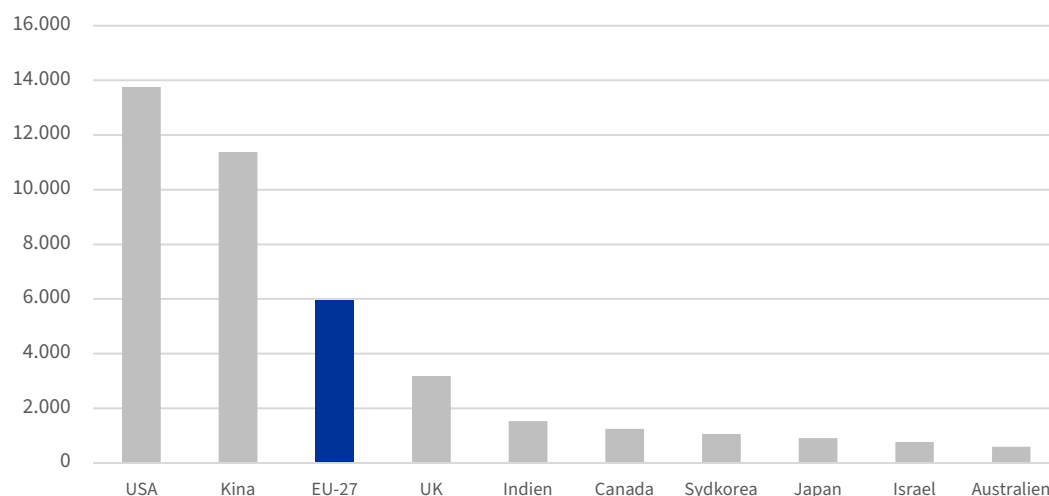
	Den Europæiske Union	Portugal
Befolkning	447 695 350	10 516 621
BNP pr. indbygger i €	38 150	25 330
Arbejdsløshedsprocent	6,1 %	6,5
Procentdel af virksomheder, der bruger AI-teknologi	8 %	7,9
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	29,9

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

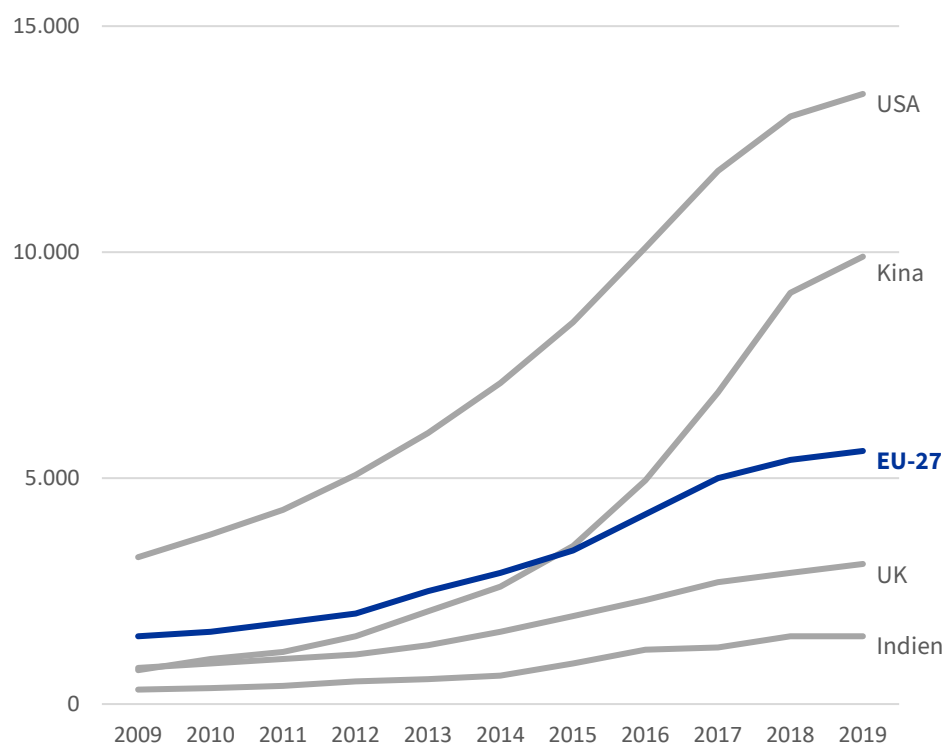


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

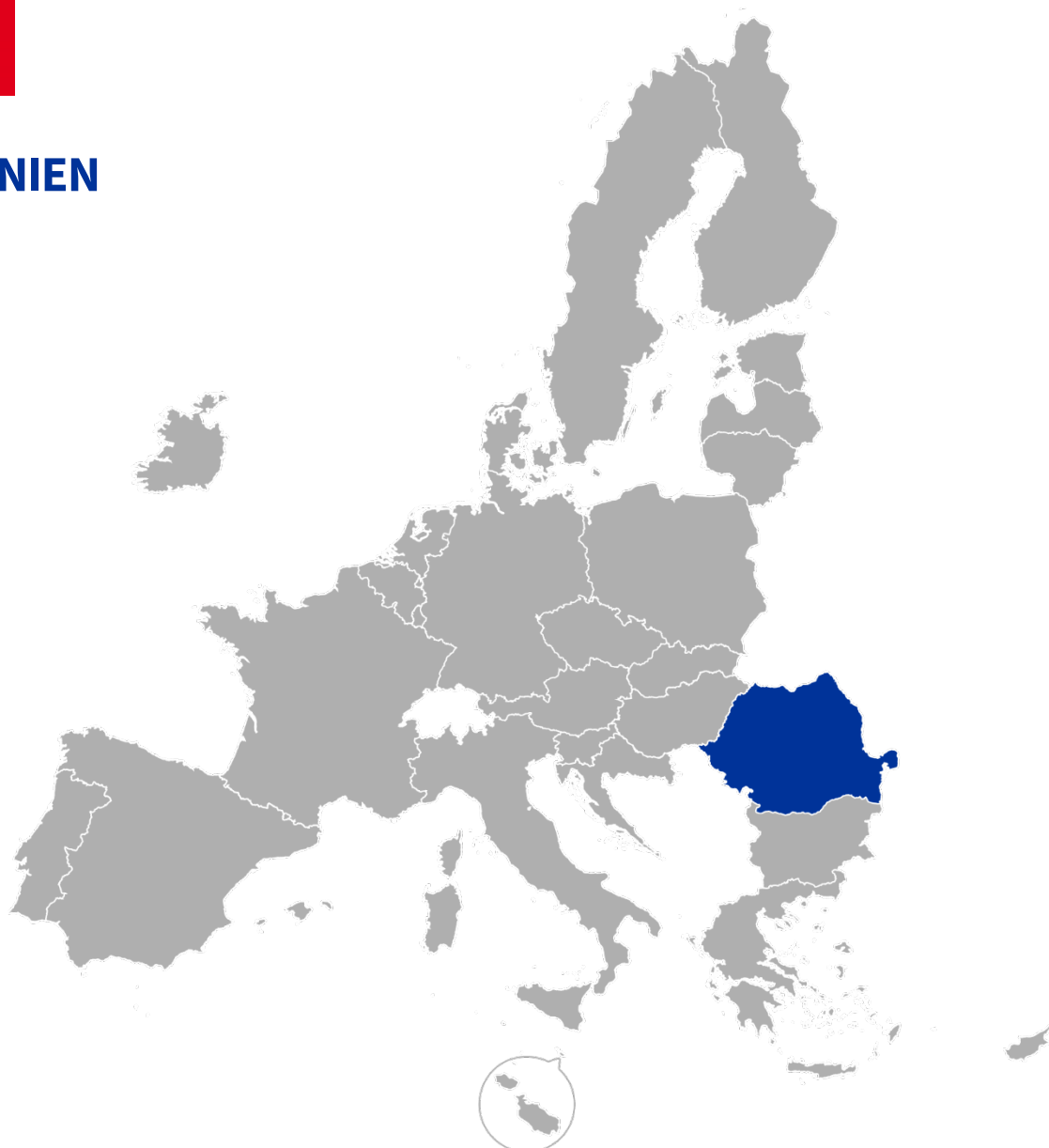
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltageres læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



RUMÆNIEN



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Danmark, Frankrig, Polen og Østrig**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Dit mål er at forbedre tilværelsen for de rumænske borgere og sikre, at AI-udviklingen har størst mulig positiv indvirkning på samfundet.
- Landbrug er en nøglesektor, hvor AI kan sætte betydeligt skub i den nationale økonomi, så du støtter lokal AI-udvikling for at fremme dette mål. Du går ind for en sektorbaseret tilgang, hvor vitale sektorer som landbrug undtages fra strengere regler, mens der anvendes strammere protokoller på mere følsomme områder.
- Men højrisiko-AI-udviklingskravene skal fortsat være hensigtsmæssige og sikre, at rammen giver retssikkerhed uden at pålægge virksomheder og forbrugere unødvendige byrder eller omkostninger. Derfor argumenterer du for nogle undtagelser fra kravene til udviklere.
- Du understreger vigtigheden af en AI-ramme, der støtter alle EU's medlemsstater, ikke kun velstående lande som Tyskland og Frankrig. Hvis drøftelserne i dag går i stå, og der ikke kan opnås et kompromis, er du parat til at acceptere en kapabilitetsbaseret tilgang. Men i så fald argumenterer du for, at det præciseres, at landbrug er en lavrisikosektor, så rigide testkrav i den sektor undgås.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Landbruget er den sektor, hvor du forventer de største fordele ved AI, men du regner også med betydelige fordele inden for retshåndhævelse. Du er imidlertid bekymret for, at man ved at underkaste de retshåndhævende myndigheder testkrav – hvor små de end måtte være – risikerer at udstille følsomme algoritmer eller operationelle detaljer for EU-tilsynsmyndighederne. Dette kan igen føre til potentielle lækager og kompromittere igangværende efterforskning.
- Derfor argumenterer du for, at nationale sikkerhedsapplikationer udelukkes fra forordningen, og understreger, at politi og nationale sikkerhedsagenturer skal kunne handle hurtigt og fleksibelt i deres indsats for at bekæmpe kriminalitet og terrorisme.
- Hvis dine krav ikke kan opfyldes fuldt ud i dagens drøftelser, er du åben over for et kompromis, hvor forordningen kan gennemgås og eventuelt revideres om et par år.
- Dine bemærkninger:

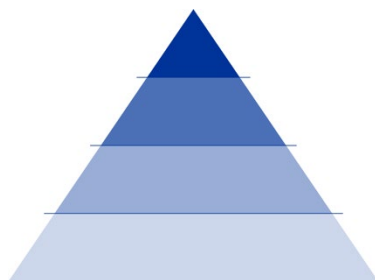
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien	formålsbaseret	fleksibel		national sikkerhed	
Slovakiet					
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

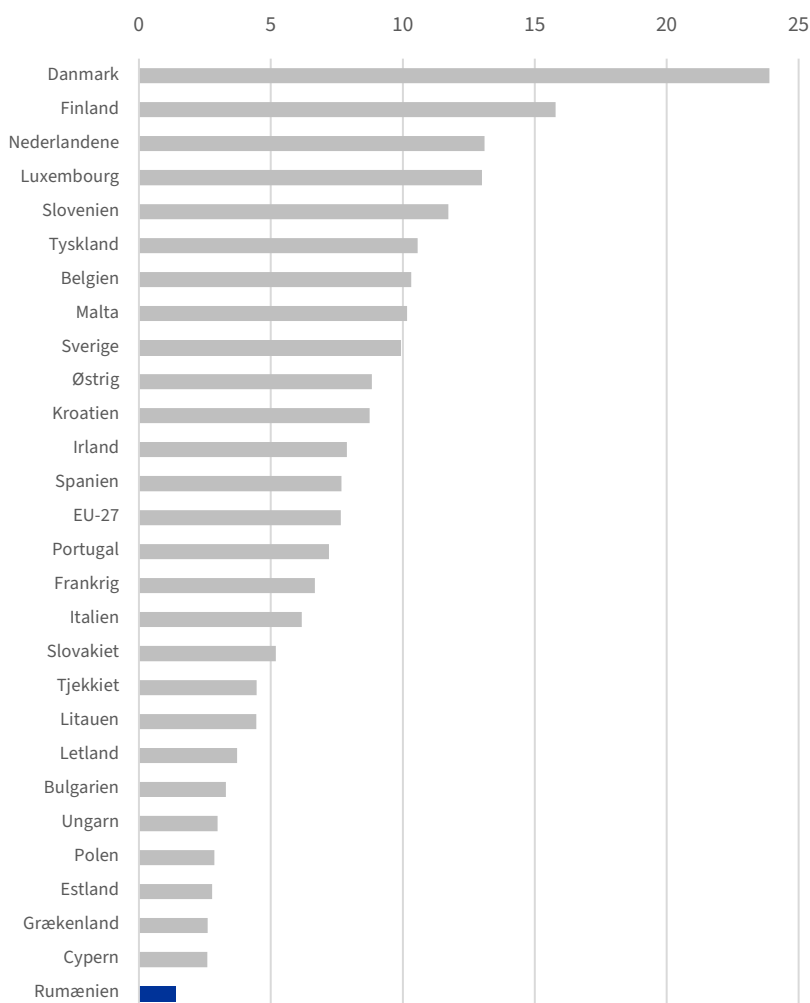
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

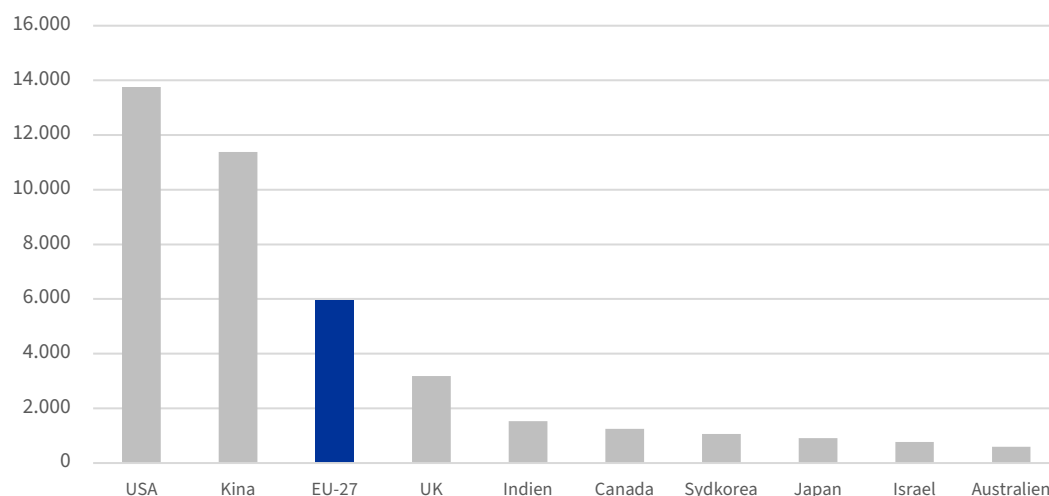
	Den Europæiske Union	Rumænien
Befolkning	447 695 350	19 054 548
BNP pr. indbygger i €	38 150	17 010
Arbejdsløshedsprocent	6,1 %	5,6 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	1,5 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	9 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

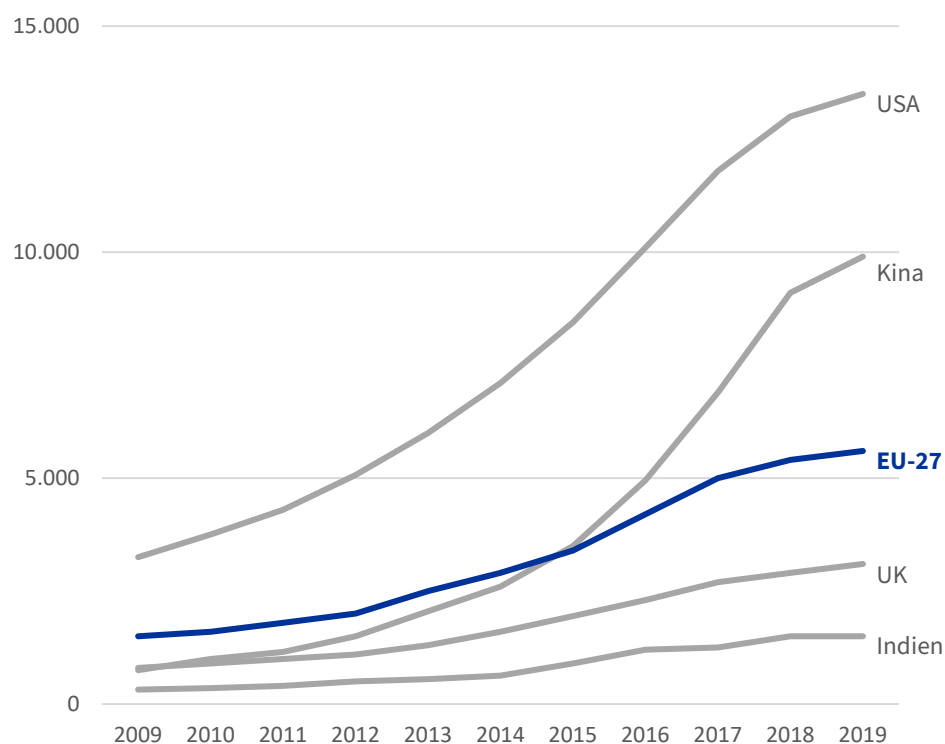


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

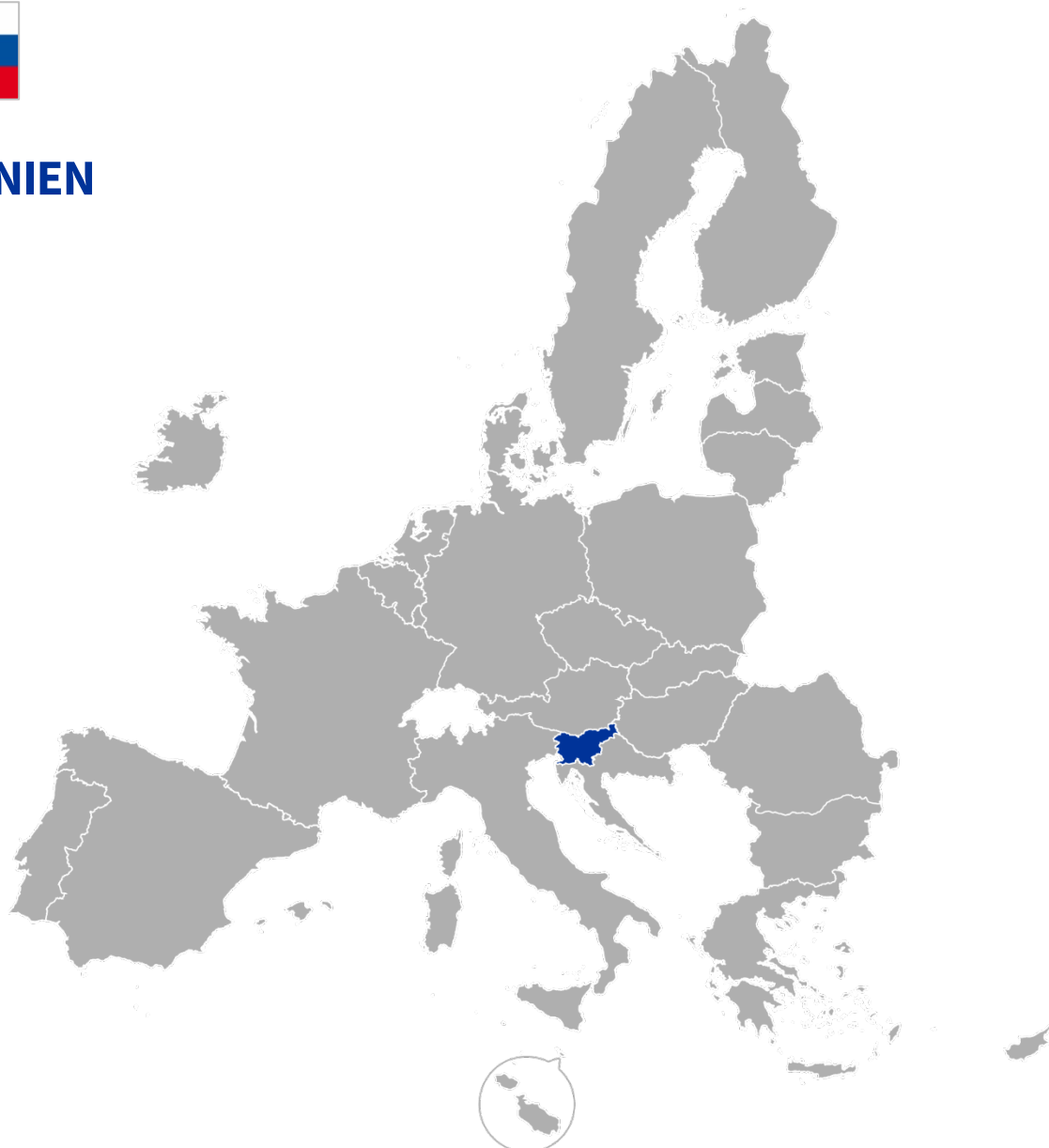
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



SLOVENIEN



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



SLOVENIEN

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Italien, Kroatien og Nederlandene**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____ .

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____ .

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____ .

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____ .

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Dit land er engageret i en omfattende og effektiv digital omstilling med udgangspunkt i livslang læring, samarbejde mellem generationerne og vækst i antallet af slovenske digitale talenter. Du ser AI som et redskab, der skal tjene menneskeheden.
- Du støtter den kapabilitetsbaserede tilgang, fordi den afspejler de faktiske risici bedre, navnlig hvad angår privatlivets fred. Du mener, at det er meget fornuftigt at klassificere alle AI'er, der er i stand til at "behandle, analysere og generere personoplysninger", som højrisiko, hvilket foreslås i den kapabilitetsbaserede tilgang.
- Under dagens drøftelser ønsker du at fremhæve potentielle risici for menneskerettigheder, retsvæsen og menneskelige aktiviteter. AI bliver betydeligt mere risikobetonet, når folk ikke har lært, hvordan den fungerer, hvad den kan, og hvilke begrænsninger den har. Uden grundlæggende AI-færdigheder kan folk stole blindt på automatiserede systemer, misse tilfælde af bias eller manipulation og være ude af stand til at drage udviklere eller institutioner til ansvar. En sådan manglende forståelse øger også sårbarheden over for misinformation, forværrer de sociale uligheder og kan føre til enten irrationel frygt eller ukritisk accept af AI-teknologier. For at sikre, at AI anvendes sikkert og retfærdigt, kræves der omfattende offentlig undervisning.
- Hvis betydningen af at undervise offentligheden understreges i AI-forordningen, er du åben over for, at der gives større fleksibilitet inden for AI-udvikling.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Efter din mening bør EU's AI-forordning fungere som katalysator for AI-udvikling i alle EU's medlemsstater. Nogle kapitalstærke lande har magtfulde AI-virksomheder, mens andre stadig forsøger at opbygge deres økosystemer med opstartsvirksomheder. Hele idéen bag opstartsvirksomheder er at skabe noget disruptivt på kort tid og med begrænset adgang til kapital. Det er næsten umuligt, hvis der er en betydelig administrativ byrde.
- Hvis artikel 1 fører til en ramme for forsigtighedstest for højrisiko-AI-modeller – hvilket du støtter – foreslår du at undtage opstartsvirksomheder med færre end 20 ansatte og under 4 mio. € i omsætning. De bør i stedet vedtage den fleksible tilgang. Du ser dette som et afbalanceret kompromis og agter at lede drøftelsen om det.
- Dine bemærkninger:

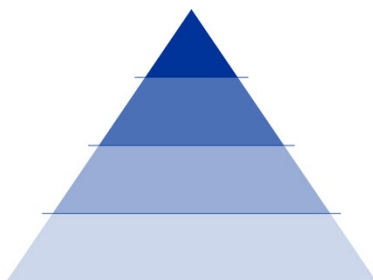
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet					
Slovenien	kapabilitetsbaseret	forsigtigheds-		opstartsvirksomheder	
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

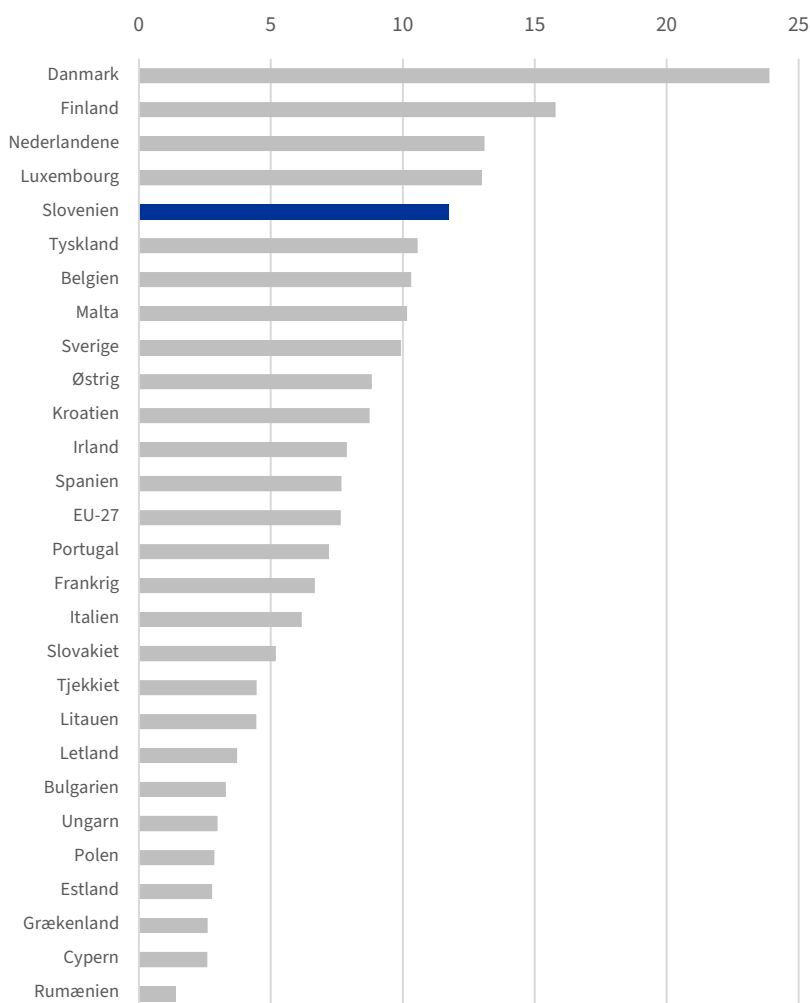
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

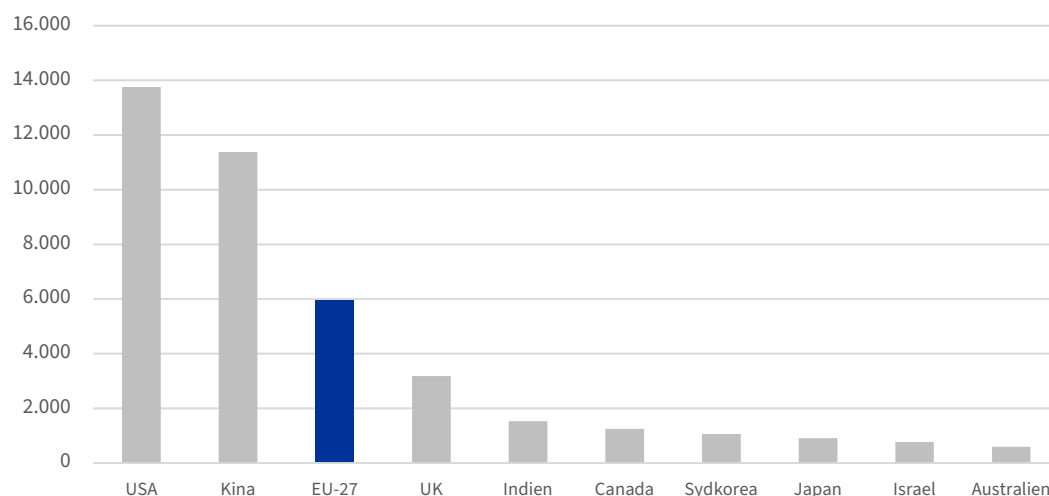
	Den Europæiske Union	Slovenien
Befolkning	447 695 350	2 116 972
BNP pr. indbygger i €	38 150	30 160
Arbejdsløshedsprocent	6,1 %	3,7 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	11,4 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	18,9 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

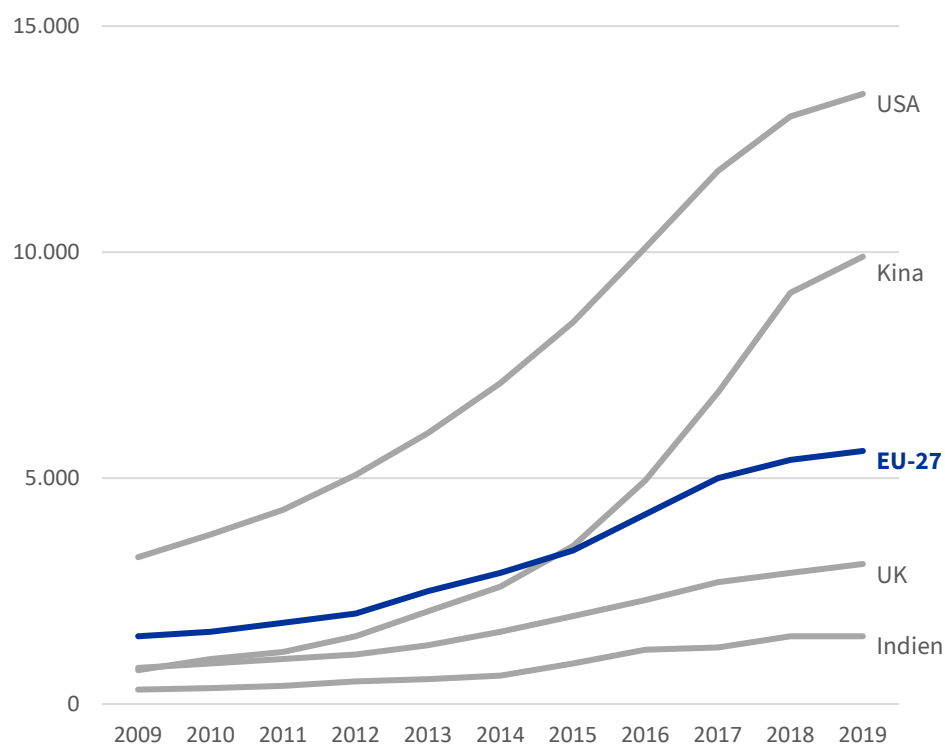


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

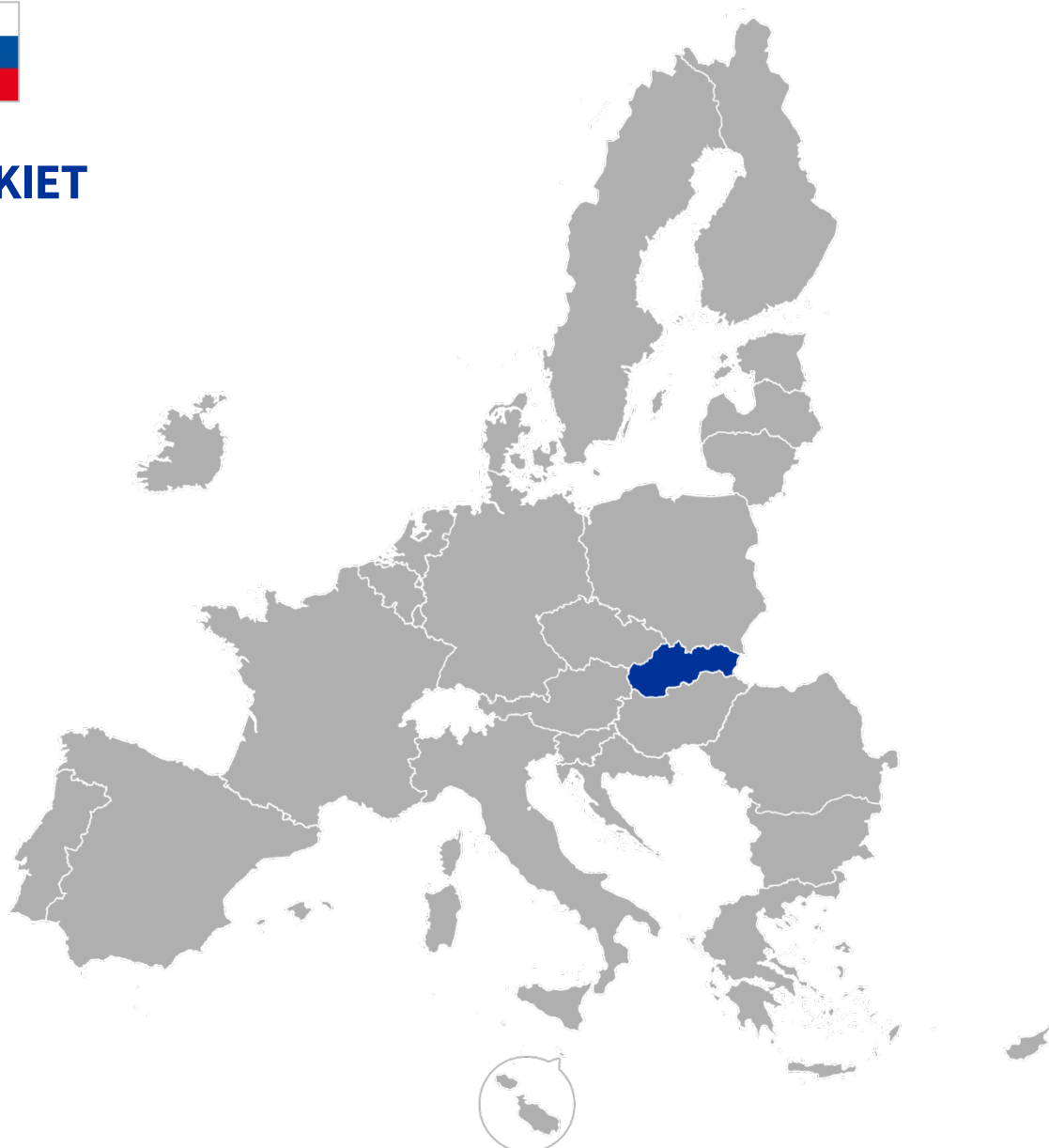
doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.



SLOVAKIET



MINISTER FOR EN DAG: NÅR 27 LANDE FORHANDLER

■ Indførelse af en AI-ramme

Rolleprofiler | **Avanceret udgave** | DA



Rådet for
Den Europæiske Union

Forslag til

EUROPA-PARLAMENTETS OG RÅDETS FORORDNING

om en retlig ramme for anvendelse af kunstig intelligens

BAGGRUND FOR FORSLAGET

Den Europæiske Union befinder sig på et kritisk tidspunkt i udformningen af sin digitale fremtid. Som led i strategien for det digitale årti sigter EU mod at føre an inden for kunstig intelligens (AI) og samtidig beskytte grundlæggende rettigheder, forbrugerbeskyttelse, markedskonkurrenceevne, innovation og teknologisk lederskab. Kernen i disse drøftelser er et lovgivningsforslag fra Europa-Kommissionen, som har til formål at regulere AI-udviklingen i EU. Efter flere forhandlingsrunder er følgende spørgsmål endnu ikke blevet behandlet.

Artikel 1

Klassificering af AI-systemer og risikostyring

- A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage **en formålsbaseret risikoklassificering: Risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.**
- B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge **en forsigtighedstilgang: overholde strenge standardiserede testkrav** (jf. bilag II).

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer, **undtagen når de er udviklet til militære eller forsvarsmæssige formål.**



SLOVAKIET

Dagsorden

På Rådets dagsorden i dag er Kommissionens forslag til en retlig ramme for anvendelse af kunstig intelligens. Du skal drøfte udkastet med repræsentanterne for de andre medlemsstater med henblik på at finde et kompromis.

Dit primære mål

En fælles holdning fra Rådet, der afspejler dine interesser. Du kan også overveje at støtte et kompromis, der delvis afspejler dine interesser.

Din opgave

1. **Læs grundigt dine gruppeoplysninger**, som beskriver dine interesser og holdninger.
2. Du **deler nogle af argumenterne med lande som Bulgarien, Luxembourg og Malta**. Tal med dem for at finde ud af, om I kan samarbejde om jeres fælles interesser, hvis I har nogen.
3. Forbered de **indledende bemærkninger** (maks. 1 minut) til den første runde officielle forhandlinger. Fremsæt din holdning til artikel 1 og 2. Du kan eventuelt hente inspiration i manuskriptet i tekstboksen.
4. Fremsæt dine indledende bemærkninger, når formandskabet (som leder mødet) giver dig ordet. Hvis du er enig i holdninger, som allerede er nævnt af andre lande, skal du gøre det klart.
5. Hør godt efter, når de andre lande fremsætter deres holdninger. **Udfyld tabellen på side 8 og 9** for at holde styr på deres holdninger. En stor del af forhandlingerne går ud på at lytte til hinanden og finde løsninger, der fungerer for alle.
6. **Fremlæg dine holdninger og argumenter** under de officielle forhandlinger. Hvis du har argumenter tilfælles med andre lande, kan det være en god idé at fremlægge jeres argumenter sammen for at give dem større vægt.
7. **Du bestemmer selvfølgelig selv, hvordan du vil stemme til sidst**. Vær opmærksom på, at du, når det er tid til den endelige afstemning om et kompromisforslag, skal beslutte, om du foretrækker et kompromis, som du måske ikke er 100 % tilfreds med, eller intet resultat.



Formand, kommissær, kolleger

Først og fremmest vil jeg gerne takke Kommissionen for forslaget og formandskabet for at sætte dette spørgsmål på dagsordenen i dag.

Forud for dagens møde har vi haft frugtbare drøftelser med _____.

Vi mener, at det er nødvendigt at vedtage en kapabilitets-/formålsbaseret risikoklassificering, fordi _____.

Med hensyn til artikel 1, litra B, mener vi, at udviklerne under udviklingen af højrisiko-AI-systemer skal anlægge en fleksibel tilgang/forsigtighedstilgang, fordi _____.

Med hensyn til artikel 2 vil visse undtagelser finde/ikke finde anvendelse, fordi vi går ind for _____.

Vi er overbevist om, at vi kan finde et brugbart kompromis i løbet af dagens møde.

Tak.

Artikel 1

Klassificering af AI-systemer og risikostyring

A) EU vedtager en fælles standard for risikostyring i forbindelse med udvikling, markedsføring, import og anvendelse af AI-systemer. Disse systemer klassificeres som uden risiko, lav risiko, høj risiko eller uacceptabel risiko (jf. bilag I). For at definere højrisiko-AI-systemer vil EU vedtage en



... kapabilitetsbaseret risikoklassificering: risikoen bestemmes af systemets tekniske formåen.



... formålsbaseret risikoklassificering: risikoen er baseret på den tilsigtede anvendelse af AI i følsomme sektorer.

På grundlag af denne risikoklassificering skal højrisiko-AI-systemer underkastes grundig afprøvning inden idriftsættelse.

B) Under udviklingen af højrisiko-AI-systemer skal udviklerne anlægge en



... fleksibel tilgang: foretage en selvevaluering af AI-modellen inden idriftsættelse.



... forsigtighedstilgang: overholde strenge standardiserede testkrav (jf. bilag II).

Dine argumenter

- Slovakiet er ivrigt efter at forstå, hvordan EU's AI-forordning udnyttes bedst muligt. Indtil videre er der ikke tildelt offentlige midler til national AI-forskning. Din regering er derfor afhængig af private investeringer eller støtte fra andre EU-medlemmer.
- Du støtter den formålsbaserede tilgang, fordi det normalt er sådan, EU regulerer teknologier og produkter, f.eks. medicinsk udstyr eller kemiske produkter. Men du mener, at den nuværende liste over højrisikosektorer bør begrænses til kritisk infrastruktur og sundhed, da disse sektorer har direkte indflydelse på liv og død. Andre sektorer kan, selv om de er vigtige, overvejes igen på et senere tidspunkt. Du er åben over for at revidere listen om tre år i samråd med civilsamfundet og industrien.
- Med hensyn til testkrav er du bekymret for, at alt for strenge regler kan udløse hjerneflugt og få europæiske AI-talenter til at søge mod mere fleksible miljøer i f.eks. USA. En fleksibel tilgang gavner også Slovakiet som en af de mindre EU-medlemsstater, der ikke har mange ressourcer til at investere i AI.
- Hvis drøftelserne i dag går i stå, og der ikke kan opnås et kompromis, er du parat til at acceptere lidt strengere testkrav. Men som led i et afbalanceret kompromis **bør der stilles finansiel støtte til rådighed for mindre velstående EU-medlemsstater** – for at hjælpe dem med at dække de øgede udviklingsomkostninger.

Artikel 2

Anvendelsesområde

Bestemmelsen fastsat i artikel 1 finder anvendelse på alle AI-systemer



uden undtagelse



uden undtagelse, medmindre systemet udvikles som open source



undtagen når de er udviklet af små virksomheder og opstartsvirksomheder (og)



undtagen når de er udviklet til militære eller forsvarsmæssige formål (og)



undtagen når de er udviklet med henblik på national sikkerhed, migration eller grænsekontrol

Dine argumenter

- Indtil videre har EU haltet bagefter inden for global AI-udvikling. Udelukkelse af opstartsvirksomheder kan bidrage til at fostre et dynamisk økosystem. Opstartsvirksomheder mangler juridisk, finansiel og administrativ kapacitet til at overholde de samme regler som store techvirksomheder. Strengt regler kan kvæle innovation i de tidlige faser.
- Hvis du ikke kan få et tilstrækkeligt antal medlemsstater til at vedtage fleksible testkrav i artikel 1, ønsker du i stedet at sikre, at opstartsvirksomheder får nogle undtagelser fra disse krav. På den måde vil de ikke blive pålagt byrden ved fuld opfyldelse fra starten (hvis artikel 1 fører til det), men vil snarere kunne operere i et mere eksperimentelt rum.
- Dine bemærkninger:

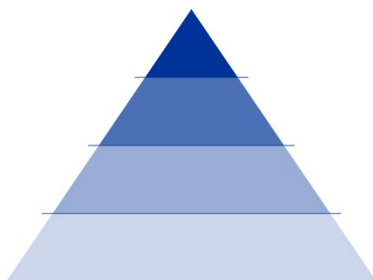
Bilag I: Risikoklassificering

Uacceptabel risiko

Høj risiko

Lav risiko

Ingen risiko



Totalforbud

Strengt krav (se bilag II)

Gennemsigtighed kræves

Ingen regulering

Uacceptabel risiko:

- **AI-baseret social scoring** – rangordning af borgere ud fra deres adfærd
- **AI-baseret masseovervågning** – biometrisk sporing uden samtykke
- **Følelsesgenkendelse på arbejdspladser/skoler** for at påvirke arbejdsresultater
- **Autonome dødbringende våben** – som kan dræbe uden menneskelig indgriben
- **AI-baseret subliminal manipulation** – påvirkning af folks beslutninger uden gennemsigtighed

Høj risiko



Kapabilitetsbaseret tilgang

Alle AI-systemer til almen brug og AI-systemer, der er i stand til:

- at træffe kritiske beslutninger om menneskers rettigheder, sikkerhed eller eksistensgrundlag – f.eks. AI-støttet kirurgi, automatiseret rekruttering af personale
- autonomt at påvirke eller manipulere menneskelig adfærd på følsomme områder, f.eks. politisk mikromålretning
- at behandle, analysere eller generere personoplysninger/biometriske data i stor skala
- at foretage komplekse forudsigelser eller risikovurderinger med samfundsmæssige konsekvenser, f.eks. klimamodellering
- at operere i miljøer, hvor fejl eller bias kan forårsage skade, f.eks. selvkørende biler

Formålsbaseret tilgang

Alle AI-systemer, der tilsigtes anvendt i forbindelse med

- kritisk infrastruktur
- sundhed
- arbejdspladsen (HR)
- offentlige bistandsydelser
- retslige institutioner

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Lav risiko

- **AI-systemer, der ikke er klassificeret som højrisiko**
- **Chatbots og virtuelle assistenter** (så længe brugerne ved, at de taler med AI)
- **Anbefalingsalgoritmer** – AI-baserede forslag på sociale medier
- **AI til forretningsautomatisering** – AI, der automatiserer grundlæggende administrative opgaver

Ingen risiko

- AI-baserede billed-, musik- eller kunstgeneratorer/-editorer
- AI i videnskabelig forskning
- AI i legetøjs- og underholdningsrobotter – AI-baseret legetøj uden risici i den virkelige verden



Bilag II: Krav til højrisiko-AI-systemer

Vægt på juridiske og etiske standarder:

- Risikovurdering inden ibrugtagning, identifikation af forudsigelige risici og potentielt misbrug.
- Menneskelig kontrol sikret for at forhindre fuld automatisering i beslutningstagningen.
- Datastyring for at sikre træningsdatasæt uden bias.
- Obligatorisk ekstern revision, herunder sandkasseafprøvning med menneskeligt tilsyn.
- Videregivelse af prøveresultater til EU's reguleringsmyndigheder.
- Overvågning efter markedsføring for at registrere ydeevne.



Definitioner

- **Algoritmer** – trinvis instruktion, som en computer følger for at løse et problem eller træffe en beslutning
- **Kunstig intelligens-systemer** (AI-systemer) – computerbaserede systemer, der anvender AI til at udføre opgaver, som typisk kræver menneskelig intelligens, f.eks. læring, ræsonnement, problemløsning, beslutningstagning, opfattelse og sprogforståelse
- **Hjerneflugt** – når faglærte arbejdere forlader deres land for at arbejde i udlandet
- **Virksomhedsmigration** – når en virksomhed flytter sine aktiviteter til et andet land, ofte for at opnå bedre vilkår eller lavere omkostninger
- **AI-baseret dyb læring** – et system, der lærer af store mængder data ved hjælp af neurale netværk, f.eks. hvordan den menneskelige hjerne fungerer. Det genkender mønstre og bliver bedre over tid, uden at der kræves eksplicit programmering
- **Kunstig intelligens til almen brug (GPAI)** – AI, der kan udføre mere end én opgave, f.eks. ChatGPT, Google Gemini
- **Open source** – software, hvis kode alle frit kan se, bruge og ændre
- **Regelbaseret AI** – et AI-system, der træffer beslutninger baseret på foruddefinerede regler og logik. F.eks. "Hvis X sker, så gør Y"
- **Sandkasseafprøvning** – afprøvning i et isoleret miljø, f.eks. afprøves en selvkvørende AI i en simuleret by, inden der køres på rigtige veje

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Østrig					
Belgien					
Bulgarien					
Kroatien					
Cypern					
Tjekkiet					
Danmark					
Estland					
Finland					
Frankrig					
Tyskland					
Grækenland					
Ungarn					
Irland					

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse.
Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

	Artikel 1A: klassificering	Artikel 1B: tilgang	Andre forslag/bemærkninger	Artikel 2: anvendelsesområde	Andre forslag/bemærkninger
Italien					
Letland					
Litauen					
Luxembourg					
Malta					
Nederlandene					
Polen					
Portugal					
Rumænien					
Slovakiet	formålsbaseret	fleksibel	finansiel støtte	opstartsvirksomheder	
Slovenien					
Spanien					
Sverige					
Kommissionen	formålsbaseret	forsigtigheds-		militær/forsvar	

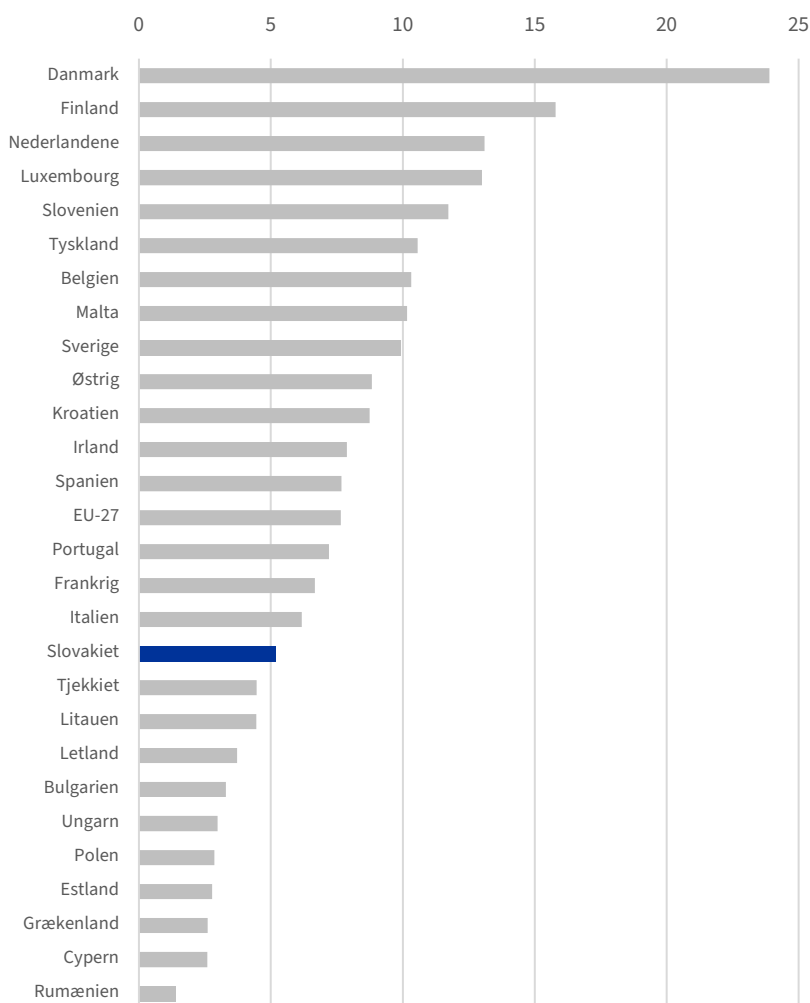
Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

Baggrundsinformation

Bemærk: **Forhandlingerne finder sted i 2023.**

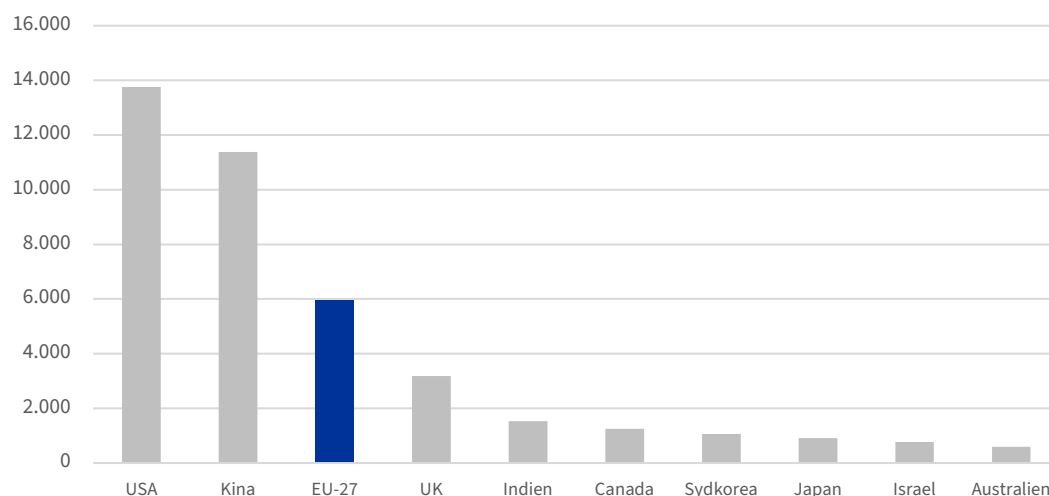
	Den Europæiske Union	Slovakiet
Befolkning	447 695 350	5 428 792
BNP pr. indbygger i €	38 150	22 690
Arbejdsløshedsprocent	6,1 %	5,8 %
Procentdel af virksomheder, der bruger AI-teknologi	8 %	7 %
Andel af befolkningen med avancerede digitale færdigheder	27,3 %	21,7 %

Procentdel af virksomheder med over ti ansatte,
der bruger mindst ét AI system (2021)

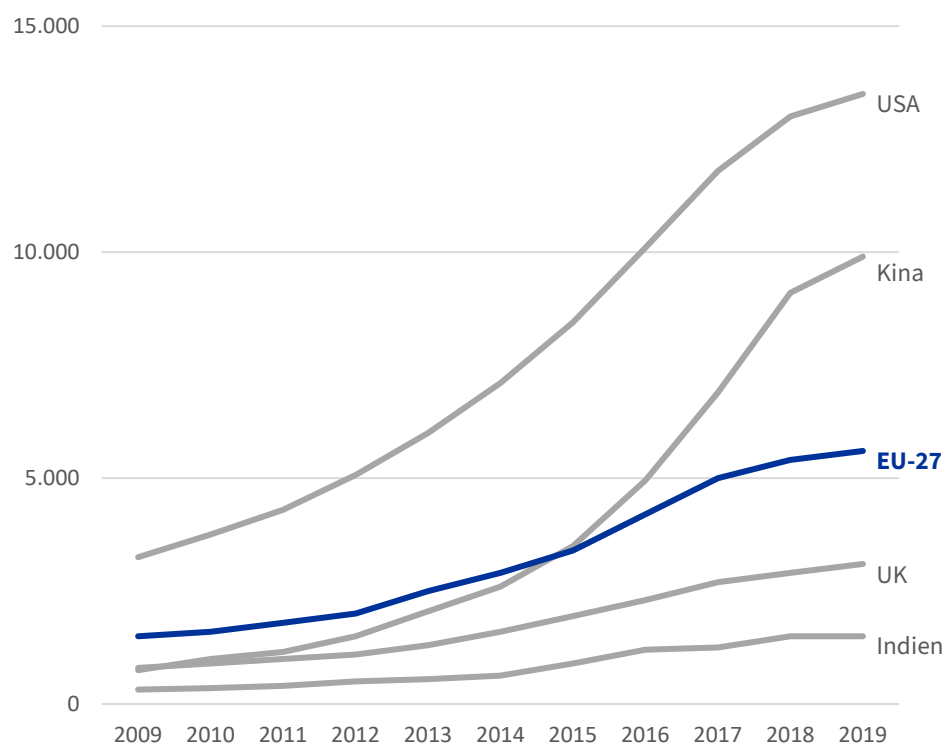


Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.

AI-virksomheder efter land (2020)



Vækst i antallet af AI-virksomheder



Datakilder: Eurostat, Det Fælles Forskningscenter

Denne publikation er udarbejdet af Generalsekretariatet for Rådet og er udelukkende til orientering. Den medfører ikke noget ansvar for EU-institutionerne eller medlemsstaterne. Du kan finde yderligere oplysninger om Det Europæiske Råd og Rådet på webstedet www.consilium.europa.eu.

© Den Europæiske Union, 2025 | Gengivelse tilladt med kildeangivelse

Print ISBN 978-92-850-0498-9

doi: 10.2860/0371762

QC-01-25-006-DA-C

PDF ISBN 978-92-850-0497-2

doi: 10.2860/0996198

QC-01-25-006-DA-N

Bemærk: Indholdet er undervisningsmæssigt udarbejdet for at sikre spillets forløb og deltagernes læringsoplevelse. Det repræsenterer ikke nødvendigvis virkeligt indhold eller virkelige politiske holdninger.