

Briselē, 21.4.2021.
SWD(2021) 85 final

KOMISIJAS DIENESTU DARBA DOKUMENTS
IETEKMES NOVĒRTĒJUMA KOPSAVILKUMA ZIŅOJUMS

Pavaddokuments dokumentam

[Mandatory element]

[Mandatory element]

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

Kopsavilkums
Mākslīgā intelekta tiesiskā regulējuma ietekmes novērtējums
A. Rīcības nepieciešamība
Problēmas būtība un nozīme ES mērogā
Mākslīgais intelekts (AI, arī MI) ir jauna vispārējas nozīmes tehnoloģija — datorprogrammēšanas tehnoloģiju kopums ar ļoti plašām iespējām. AI sistēmu ieviešanai ir liels potenciāls dot ievērojamu ieguvumu sabiedrībai, sekmēt ekonomikas izaugsmi un Eiropas Savienībā uzlabot inovāciju un ES spēju konkurēt pasaulē. Tomēr dažās jomās AI sistēmu izmantošana var radīt problēmas. Dažu AI sistēmu īpašās iezīmes var radīt jaunus riskus, kas saistīti ar 1) drošumu un drošību un 2) pamattiesībām, kā arī palielināt esošo risku iespējamību un intensitāti. AI sistēmas arī 3) apgrūtina tiesībaizsardzības iestāžu darbu pastāvošo noteikumu ievērošanas pārbaudīšanā un izpildes panākšanā. Minētie problemātiskie aspekti savukārt 4) rada juridisko nenoteiktību uzņēmumiem un 5) nepietiekamas uzticēšanās dēļ var kavēt AI tehnoloģiju ieviešanu uzņēmumos un sabiedrībā; turklāt 6), reaģējot uz tiem, valstu iestādes pieņem regulatīvos pasākumus, kuru mērķis ir mazināt iespējamās eksternalitātes, bet kuri draud sadrumstalot iekšējo tirgu.
Sasniedzamie mērķi
Tiesiskā regulējuma mērķis ir risināt minētās problēmas, lai nodrošinātu pareizu vienotā tirgus darbību, radot piemērotus apstākļus uzticama AI izstrādei un izmantošanai Savienībā. Specifiskie mērķi ir šādi: 1) nodrošināt, ka tirgū laistās un izmantotās AI sistēmas ir drošas un atbilst pašreizējiem tiesību aktiem par pamattiesībām un Savienības vērtībām; 2) nodrošināt juridisko noteiktību, lai veicinātu ieguldījumus AI un ar to saistītu inovāciju; 3) uzlabot pārvaldi un to pastāvošo AI sistēmām piemērojamo tiesību aktu efektīvu izpildi, kuros noteiktas pamattiesības un drošuma prasības; 4) veicināt vienota tirgus attīstību likumīgām, drošām un uzticamām AI sistēmām un novērst tirgus sadrumstalotību.
ES līmeņa rīcības pievienotā vērtība (subsidiaritāte)
Tā kā AI lietojumiem bieži tiek izmantoti lielapjoma dati un datu kopas, kam piemīt pārrobežu raksturs, iniciatīvas mērķus nevar sasniegt dalībvalstis, rīkojoties atsevišķi. Eiropas uzticama AI tiesiskā regulējuma mērķis ir noteikt saskaņotus noteikumus par tādu produktu un pakalpojumu izstrādi, laišanu tirgū un izmantošanu, kuros iegultas AI tehnoloģijas, kā arī par savrupu AI lietojumu izstrādi, laišanu tirgū un izmantošanu Savienībā. Tā nolūks ir nodrošināt vienādus konkurences apstākļus un aizsargāt visus Eiropas iedzīvotājus, vienlaikus stiprinot Eiropas konkurētspēju un rūpniecisko bāzi mākslīgā intelekta jomā. ES rīcība AI jomā stimulēs iekšējo tirgu, un šādi ir lielas iespējas sniegt Eiropas nozares dalībniekiem konkurences priekšrocību globālā mērogā, pateicoties apjomradītiem ietaupījumiem, kurus dalībvalstis nevarētu panākt, rīkojoties atsevišķi.
B. Risinājumi
Risinājumu varianti izvirzīto mērķu sasniegšanai. Vēlamais variants (ja ir), iemesli (ja nav)
Tika izskatīti šādi varianti: 1. variants – ES likumdošanas instruments, kas izveidotu brīvprātīgu marķējuma sistēmu; 2. variants – speciāla sektoriāla pieeja; 3. variants – horizontāls ES likumdošanas instruments, kas noteiktu obligātas prasības augsta riska AI lietojumiem; 3.+ variants – tas pats 3. variants, bet ar brīvprātīgiem rīcības kodeksiem attiecībā uz AI lietojumiem, kuri nerada augstu risku; 4. variants – horizontāls ES likumdošanas instruments, kas noteiktu obligātas prasības visiem AI lietojumiem. Vēlamākais ir 3.+ variants, jo tas nodrošina samērīgus aizsardzības pasākumus pret AI radītajiem riskiem ar minimalizētām administratīvajām un atbilstības nodrošināšanas izmaksām. Konkrētajam jautājumam par atbildību par AI lietojumiem tiks veltīti atsevišķi noteikumi, kas tiks pieņemti nākotnē, tādēļ tas šajos variantos nav aplūkots.
Dažādu ieinteresēto personu viedokļi. Atbalsts kādam no variantiem
Gan uzņēmumi, gan publiskā sektora iestādes, gan akadēmiskie darbinieki un nevalstiskās organizācijas piekrīt, ka pastāv regulējuma nepilnības vai ka ir nepieciešami jauni tiesību akti, taču uzņēmumu vidū šādi

domājošo vairākums ir mazāks. Nozares dalībnieki un publiskā sektora iestādes piekrīt, ka obligātās prasības jāattiecinā tikai uz augsta riska AI lietojumiem. Pilsoņi un pilsoniskā sabiedrība biežāk nepiekrīt obligāto prasību attiecināšanai vienīgi uz augsta riska lietojumiem.
C. Vēlamā varianta ietekme
Ieguvumi no vēlamā varianta (ja tāda nav, tad no galvenajiem variantiem)
No pilsoņu viedokļa vēlamais variants mazinās riskus, kas apdraud to drošību un pamattiesības. AI sagādātājiem šis variants nesīs juridisko noteiktību un novērsīs šķēršļus ar AI saistīto pakalpojumu un produktu pārrobežu aprītei. Uzņēmumiem, kuri izmanto AI, tas vairo klientu uzticēšanos. No valsts publiskās pārvaldes iestāžu viedokļa — tas vairo sabiedrības uzticēšanos AI izmantošanai un stiprinās izpildes panākšanas mehānismus (ieviešot Eiropas koordinācijas mehānismu, paredzot pienācīgas spējas un vienkāršojot AI sistēmu revīzijas ar jaunām prasībām par dokumentāciju, izsekojamību un pārredzamību).
Vēlamā varianta (ja tāda nav, tad galveno variantu) izmaksas
Uzņēmumiem vai publiskā sektora iestādēm, kas izstrādā vai izmanto AI lietojumus, kuri rada augstu risku iedzīvotāju drošībai vai pamattiesībām, būs jāievēro konkrētas horizontālās prasības un pienākumi, kas tiks noteikti saskaņotos tehniskajos standartos. Aplēsts, ka atbilstības nodrošināšanas summārās izmaksas laikposmā līdz 2025. gadam būs no 100 līdz 500 miljoniem euro; tas atbilst ne vairāk kā 4–5 % no ieguldījumiem augsta riska mākslīgā intelekta risinājumos (tādi ir aptuveni 5–15 % no visiem AI lietojumiem). Pārbaudes izmaksu apmērs varētu būt vēl 2–5 % no ieguldījumiem augsta riska mākslīgā intelekta risinājumos. Uzņēmumiem un publiskā sektora iestādēm, kas izstrādā vai izmanto AI lietojumus, kuri nav klasificēti kā augsta riska lietojumi, nekādas izmaksas nerastos. Tomēr tie varētu izvēlēties ievērot brīvprātīgos rīcības kodeksus, lai izpildītu attiecīgas prasības un nodrošinātu, ka to AI ir uzticams. Šādā gadījumā maksimālās izmaksas varētu būt tikpat lielas kā saistībā ar augsta riska lietojumiem, taču, visticamāk, tās būtu zemākas.
Ietekme uz MVU un konkurētspēju
No lielākas vispārējās uzticēšanās mākslīgajam intelektam MVU gūs lielāku labumu nekā lieli uzņēmumi, kuri var paļauties arī uz sava zīmola tēlu. MVU, kas izstrādā lietojumus, kuri klasificēti kā augsta riska lietojumi, rastos līdzīgas izmaksas kā lieliem uzņēmumiem. Pateicoties digitālo tehnoloģiju lielajām mērogojamības iespējām, mazie un vidējie uzņēmumi par spīti nelielajam izmēram var sasniegt plašu klientu loku, potenciāli ietekmējot miljoniem cilvēku. Tādēļ mākslīgā intelekta risinājumus piedāvājošo MVU izslēgšana no tiesiskā regulējuma piemērošanas attiecībā uz augsta riska lietojumiem varētu būtiski apgrūtināt uzticēšanās palielināšanas mērķa sasniegšanu. Regulējumā tomēr būs paredzēti īpaši pasākumi, tostarp “regulatīvās smilškastes” vai digitālās inovācijas centru sniegta palīdzība, ar mērķi atbalstīt MVU jauno noteikumu izpildē, ņemot vērā to īpašās vajadzības.
Nozīmīga ietekme uz valsts budžetiem un pārvaldes iestādēm
Dalībvalstīm būtu jānorīko uzraudzības iestādes, kas atbild par regulatīvo prasību īstenošanu. To pienākumu izpilde uzraudzības jomā varētu būt balstīta jau esošā kārtībā, piemēram, saistībā ar atbilstības novērtēšanas struktūrām vai tirgus uzraudzību, taču tai būtu nepieciešama pietiekama tehnoloģiskā zinātība un resursi. Atkarā no katrā dalībvalstī jau izveidotajām struktūrām tam varētu būt nepieciešams 1 līdz 25 pilnslodzes ekvivalenti uz vienu dalībvalsti.
Cita nozīmīga ietekme
Vēlamais variants būtiski mazinātu riskus, kas apdraud pilsoņu pamattiesības un plašākas Savienības vērtības, un uzlabotu tādu noteiktu produktu un pakalpojumu drošumu, kuros iegultas AI tehnoloģijas, kā arī savrupu AI lietojumu drošumu.
Proporcionalitāte
Priekšlikums ir proporcionāls un nepieciešams mērķu sasniegšanai, jo atbilst uz risku balstītai pieejai un paredz regulatīvo slogu tikai gadījumos, kad AI sistēmas var radīt augstu risku pamattiesībām vai drošībai. Pārējos gadījumos noteiktas tikai minimālas pārredzamības prasības, īpaši attiecībā uz informācijas sniegšanu situācijās, kad jānorāda uz AI sistēmas izmantošanu saskarsmē ar cilvēkiem vai dziļviltojumu

izmantošanu, ja tos neizmanto likumīgos nolūkos. Lai palīdzētu sagādātājiem un lietotājiem ievērot prasības un samazināt izmaksas, tiks izdoti saskaņoti standarti, kā arī atbalsta norādījumi un atbilstības nodrošināšanas rīki.
D. Turpmākā rīcība
Politikas pārskatīšanas termiņš
Komisija publicēs ziņojumu par regulējuma izvērtēšanu un pārskatīšanu piecus gadus pēc dienas, kad tas kļūs piemērojams.