

2026Z18543

(ingezonden 10 september 2026)

Vragen van het lid Stoffer (SGP) aan de staatssecretaris van Economische Zaken en Klimaat over de veiligheidsrisico's van geavanceerde kunstmatige intelligentie.

Bent u bekend met het bericht 'Anthropic-onderzoeker neemt ontslag: OpenAI en Anthropic handelen onverantwoordelijk en gokken met ons leven'[1] en het bericht 'Anthropic-medewerker: kans bestaat dat AI de mensheid uitroeit'[2], waarin onderzoekers van vooraanstaande AI-bedrijven waarschuwen voor ernstige veiligheidsrisico's van steeds krachtigere AI-systemen?

Hoe beoordeelt u de waarschuwingen van deze onderzoekers dat de ontwikkeling van zeer geavanceerde AI-systemen sneller verloopt dan de ontwikkeling van adequate veiligheidsmaatregelen en beheersingsmechanismen? Welke conclusies verbindt u daaraan?

Heeft het kabinet scenarioanalyses laten uitvoeren naar de gevolgen van geavanceerde AI voor de nationale veiligheid, waaronder cyberaanvallen, beïnvloeding van democratische processen, sabotage van vitale infrastructuur, militaire toepassingen en andere hybride dreigingen? Zo ja, kunt u de belangrijkste bevindingen met de Kamer delen?

In hoeverre acht u het bestaande nationale, Europese en internationale normatieve kader, waaronder de AI Act, toereikend voor het beheersen van risico's die voortvloeien uit de ontwikkeling van zeer krachtige algemene AI-modellen en toekomstige frontier AI-systemen?

Welke aanvullende maatregelen acht u noodzakelijk om de digitale weerbaarheid van Nederland te versterken tegen risico's die samenhangen met geavanceerde AI-systemen, in het bijzonder waar het gaat om vitale infrastructuur, overheidsorganisaties, defensie en de bescherming van burgers?

Hoe geeft het kabinet, mede in het licht van het streven naar een digitaal weerbare samenleving, invulling aan het voorzorgsbeginsel bij de ontwikkeling en toepassing van zeer geavanceerde AI-systemen waarvan de maatschappelijke gevolgen mogelijk verstrekkend en onomkeerbaar zijn?

Deelt u de zorg dat een internationale concurrentiestrijd tussen staten en marktpartijen kan leiden tot een situatie waarin veiligheidsbelangen ondergeschikt raken aan technologische en economische belangen? Zo ja, hoe wordt dit risico in het kabinetsbeleid geadresseerd?

Herkent u de zorgen die leven rondom superintelligente AI en het risico dat de mens niet meer in staat is deze ontwikkeling te stoppen? Deelt u de mening dat een verbod hierop wenselijk is?

Bent u bereid binnen de Europese Unie en internationale gremia te pleiten voor aanvullende afspraken over de ontwikkeling van geavanceerde AI, een verbod op superintelligente AI, en de mogelijkheid om tijdelijk een pas op de plaats te maken bij de ontwikkeling van krachtige AI-

systemen zonder dat de risico's voldoende helder zijn?

Kunt u de Kamer voor het einde van dit jaar informeren over de kabinetsinzet en de mogelijkheden en wenselijkheid van aanvullende waarborgen voor geavanceerde, superintelligente AI, waaronder de mogelijkheid van een tijdelijke ontwikkelpauze voor de krachtigste AI-systemen?

[1] Anthropic-onderzoeker neemt ontslag: 'OpenAI en Anthropic handelen onverantwoordelijk en gokken met ons leven', Volkskrant d.d. 9 september 2026; Anthropic-onderzoeker neemt ontslag: 'OpenAI en Anthropic handelen onverantwoordelijk en gokken met ons leven' | de Volkskrant.

[2] Anthropic medewerker: kans bestaat dat AI de mensheid uitroeit', Quotenet d.d. 9 september 2026; Anthropic-medewerker: 'Kans bestaat dat AI de mensheid uitroeit'.