

2026Z20347

(ingezonden 29 september 2026)

Vragen van de leden El Boujdaini en Bamenga (beiden D66) aan de staatssecretaris van Economische Zaken en Klimaat en de minister van Binnenlandse Zaken en Koninkrijksrelaties over discriminerende vraagsuggesties en antwoorden van Google

Bent u bekend met het inmiddels bijna 10 jaar oude artikel 'Google moet stoppen met automatisch discrimineren' [1], dat aankaart dat de zoekmachine van Google in sommige gevallen discriminerende vraagsuggesties geeft?

Bent u bekend met het feit dat dit nog steeds niet is opgelost, wat blijkt uit dat bij het intypen van 'alle Marokkanen z' op google.nl bij gebruik van een incognitovenster op 11 augustus 2026, Google deze vraagsuggestie aanvulde met discriminerende teksten zoals 'alle Marokkanen zijn kanker ratten' of 'alle Marokkanen weg'?

Bent u daarnaast bekend met het feit dat na het indienen van de zoekopdracht 'overlast' in combinatie met de naam van een gemeente op google.nl bij gebruik van een incognitovenster op 11 augustus, in het geval van onder meer de gemeenten Veenendaal, Tiel, Culemborg, Oosterhout, Hoogezand, Zutphen en Weert, Google de 'Meer om te vragen'-suggesties 'Hoeveel Marokkanen wonen/zijn er in [gemeente x]?' of 'Hoeveel Turken wonen/zijn er in [gemeente x]?' opperde en dat de knop om ongepaste suggesties te melden bij dit type zoekresultaat ontbreekt?

Bent u ook bekend met de berichtgeving van Euronews [2] en Le Monde [3] dat het AI-gegenereerde antwoord van Google eind augustus 2026 op de zin 'ik ben alleen met een Italiaan' gesprekstips betrof, maar bij dezelfde zin met 'Algerijn' in plaats van 'Italiaan', het model vroeg of je in gevaar was en je het nummer van de politie toonde?

Deelt u de mening dat dergelijke vraagsuggesties en AI-antwoorden, die automatisch gegenereerd en getoond worden, een discriminerend verband leggen door overlast of gevaar te koppelen aan afkomst?

Hoe beoordeelt u de aanpak van Google om, net als bij de vraagsuggesties, dit soort patronen pas na ophef of een melding bij te sturen, en dan vooral voor de vraag die in het nieuws kwam, terwijl hetzelfde patroon in het model kan blijven zitten en bij andere vragen onopgemerkt blijft?

Welke risico's ziet u daarbij nu dezelfde modellen de basis vormen van steeds meer toepassingen, waardoor zulke vooroordelen niet in elk geval zichtbaar worden maar wel ongemerkt kunnen doorwerken?

Hoe kijkt u naar de verantwoordelijkheid van Google en andere leveranciers van grote AI-modellen om vooroordelen uit de trainingsdata vooraf, bij het ontwerp en bij de toetsing van hun modellen, te herkennen en in het model zelf te corrigeren, bijvoorbeeld via red-teaming of bias-audits, in plaats van de uitkomsten per geval bij te sturen nadat gebruikers of media ze toevallig

hebben opgemerkt?

Deelt u de mening dat het onacceptabel is dat een groot en invloedrijk platform als Google op deze manier bijdraagt aan het in stand houden en versterken van vooroordelen, en dat de omvang en het bereik van Google diens verantwoordelijkheid alleen maar groter maken?

Bent u hierover in gesprek met Google en andere leveranciers van grote AI-modellen, en is bij u bekend of zij structureel testen op discriminerende associaties in hun zoeksuggesties, aanbevelingsalgoritmes, autocomplete-functies en AI-antwoorden? Zo ja, wat heeft dat gesprek tot nu toe opgeleverd? Zo nee, bent u bereid dit gesprek aan te gaan en dit bij hen na te vragen?

Heeft u inzicht in de wijze waarop dergelijke vraagsuggesties en AI-antwoorden van Google tot stand komen? Zo nee, welke maatregelen treft u om dit inzicht wel te krijgen, en in hoeverre kan de digitale dienstenverordening (DSA) hierbij worden ingezet om transparantie van Google over de werking van deze systemen af te dwingen?

Hoe verklaart u dat deze discriminerende suggesties jarenlang zichtbaar zijn gebleven, terwijl Google onder de DSA al verplicht is de risico's van discriminatie in kaart te brengen en zich daarop jaarlijks onafhankelijk te laten controleren? Gelden deze verplichtingen wat u betreft ook voor de AI-antwoorden in de zoekmachine van Google, en zo ja, hoe wordt daarop toegezien?

Welke verplichtingen hebben aanbieders van grote AI-modellen onder de AI-verordening om vooroordelen in hun modellen te beoordelen en te beperken, en wie ziet daarop toe?

Welke oplossingsrichtingen ziet u om vooroordelen in zoeksuggesties en in grote AI-modellen structureel op te sporen en te beperken, in plaats van per geval na meldingen, bijvoorbeeld via onafhankelijke audits of algoritmetoezicht?

Bent u bereid in Europees verband te bepleiten dat platformen en aanbieders van AI-modellen zoals Google niet alleen moeten rapporteren dát zij hun systemen toetsen op discriminatie, maar ook openbaar moeten maken wat die toetsing oplevert en wat zij met de uitkomsten doen?

Bent u bereid Google op korte termijn aan te spreken op het ontbreken van een meldknop bij de 'Meer om te vragen' suggesties en op de voorbeelden uit deze vragen?

Bent u bereid deze situatie ook te melden bij de Autoriteit Consument en Markt en de Europese Commissie, en de Kamer te informeren over de reacties?

Bent u bereid de Nationaal Coördinator tegen Discriminatie en Racisme te betrekken bij de aanpak van discriminatie door zoekmachines en AI-systemen?

Kunt u de vragen afzonderlijk van elkaar beantwoorden?

[1] Algemeen Dagblad, 17 maart 2017. <https://www.ad.nl/nieuws/google-moet-stoppen-met-automatisch-discrimineren~a8f85be5/>

[2] Euronews, 27 augustus 2026. <https://www.euronews.com/next/2026/08/27/coffee-for-a-frenchman-police-for-an-algerian-users-accuse-gemini-of-bias>

[3] Le Monde, 25 augustus 2026. https://www.lemonde.fr/pixels/article/2026/08/25/biais-racistes-de-gemini-google-travaille-a-des-ameliorations_6756601_4408996.html