

SCHRIFTELIJKE INBRENG RONDETAfelGESPREK INZAKE 'SOCIALE MEDIA EN INMENGING' 4 MAART 2026

1) Hoe Meta zich voorbereidt op verkiezingen en de integriteit van verkiezingen beschermt

Meta heeft door de jaren heen een uitgebreide aanpak ontwikkeld met betrekking tot verkiezingen op haar platforms. Bij de voltrekking van de Nederlandse Tweede Kamerverkiezing op 29 oktober 2025, waren deze beleidsmaatregelen en procedures ook van kracht en hebben verschillende teams binnen Meta maandenlang voorbereidingen getroffen om de belangrijkste platform-risico's rond de verkiezingen te identificeren en te mitigeren.

Sinds 2016 heeft Meta meer dan \$30 miljard geïnvesteerd in veiligheid (*safety*) en beveiliging (*security*). Elke verkiezing is uniek, maar door onze ervaring van over 200 verkiezingen wereldwijd te gebruiken, hebben wij een robuust programma gebouwd met processen, tools en beleid om de vrijheid van meningsuiting te beschermen en de integriteit van verkiezingen te waarborgen. We verbeteren deze processen en maatregelen voortdurend zodat ze responsief blijven voor opkomende risico's en blijven deze inspanningen versterken - ook in het licht van de regelgevende kaders van de Digital Services Act (DSA), de Election Guidelines en de relevante verplichtingen uit de EU Code of Practice on Disinformation.

Ter voorbereiding op de afgelopen Nederlandse Tweede Kamerverkiezing heeft Meta een speciaal team samengesteld dat verantwoordelijk was voor deze bedrijfsoverkoepelende inspanningen. Dit team bestond uit experts van onze *intelligence*-, *data science*-, *product*- en *engineering*-, *research*-, *operations*-, *content*-, *mensenrechten*-, *public policy*- en juridische teams.

Meta Platforms Ireland Limited

Registered Office:

Merrion Road

Dublin 4, D04 X2K5

Registered in Ireland as a private company limited by shares

Company No. 462932

Directors: David Harris, Majella Goss, Yvonne Cunnane, Anne O'Leary, Genevieve Hughes

www.meta.com

Dit multidisciplinaire team combineert regionale en taalkundige expertise: specialisten die in verschillende landen werken richten zich op verkiezings-gerelateerde onderwerpen en beoordelen samen de specifieke risico's van elke verkiezing.

Deze voorbereidingen stellen ons in staat om een op-maat-benadering op te stellen en te zorgen dat we voldoende middelen en mitigaties hebben om de geïdentificeerde risico's aan te pakken. Het team blijft mogelijke nieuwe risico's monitoren en beoordelen, zodat we relevante mitigaties in *real-time* kunnen toepassen.

Bij de aanpak en de voorbereiding van verkiezingen zetten wij o.a. in op het verwijderen van content die onze beleidsregels overtreedt, het veilig houden van mensen op onze platforms, het verstoren van vijandige netwerken, het beheersen van AI-risico's en kiezers voorzien van betrouwbare informatie. Wij handhaven ook [ons beleid dat kiezers- of census- beïnvloeding tegengaat](#).

We hebben gespecialiseerde onderzoeksteams om gemanipuleerde campagnes te stoppen en opkomende bedreigingen te identificeren.

- We [bestrijden](#) buitenlandse en binnenlandse inmenging bij verkiezingen en hebben sinds 2017 meer dan 200 netwerken van [Gecoördineerd Onecht Gedrag](#) (CIB) wereldwijd verwijderd. Onze bevindingen over CIB-netwerken delen we via regelmatige dreigingsrapportages.
- We nemen proactieve maatregelen om organisaties te detecteren die niet op onze platforms mogen zijn en deze te verwijderen.
- We plaatsen [groepscontent](#) lager van leden die ons beleid tegen beïnvloeding of andere Community Standards hebben overtreden, om het bereik van deze personen te beperken. We beperken ook hun mogelijkheden om te posten, te reageren, nieuwe leden toe te voegen of nieuwe groepen te maken. Meer informatie vindt u [hier](#).

We verwijderen schadelijke content

- We verwijderen content die direct kan bijdragen aan een risico op inmenging, zoals de mogelijkheid van mensen om deel te nemen aan verkiezingen (dit omvat bijvoorbeeld content die valse informatie verschaft over hoe of waar te stemmen). We hebben specifieke beleidsregels die misinformatie over data, locaties, tijden, stemmethoden of wie mag stemmen verbieden (het beleid is [hier](#) verder uiteengezet).



We hebben onze aanpak aangepast om het opkomende gebruik van AI bij het creëren van schadelijke of misleidende content te adresseren.

- We verwijderen content die ons beleid schendt, ongeacht of deze door AI of een persoon is gemaakt.
- Onze onafhankelijke factchecking-partners beoordelen en classificeren virale misinformatie, waarbij een label “Aangepast” voor content die is bewerkt op een manier die mensen kan misleiden, ook met AI. We verbieden advertenties als ze als *False, Altered, Partly False of Missing Context* zijn beoordeeld.

We verbinden mensen met betrouwbare informatie over stemmen en bestrijden misinformatie in verschillende talen

- Er zijn momenteel geen wijzigingen in ons third-party factchecking-programma in Europa; we blijven samenwerken met lokale, onafhankelijke factcheck-organisaties om virale misinformatie en hoaxes te ontkrachten. In Nederland werken we nauw samen met dpa-Faktencheck en AFP.
- We blijven gebruikers voorzien van betrouwbare verkiezingsinformatie en bestrijden misinformatie in verschillende talen. Daarom verbinden we gebruikers proactief met details over de verkiezingen (in hun lidstaat) via in-app notificaties.
- We verwijzen gebruikers naar betrouwbare informatie over het verkiezingsproces via in-app ‘Voter Information Units’ en ‘Election Day Information’. Dit zijn meldingen bovenaan de feed op Facebook en Instagram, die gebruikers herinneren aan de verkiezingsdag en hen doorverwijzen naar lokale, gezaghebbende bronnen over hoe en waar te stemmen. In Nederland is verwezen naar de [website van de Rijksoverheid](#).

2) Bestrijding van desinformatie

We vertrouwen op geautomatiseerde systemen die nepaccounts detecteren en verwijderen om onechte volgers of nep-engagement te bestrijden. Niet-authentiek of onecht gedrag verwijst naar complexe vormen van misleiding, uitgevoerd door een netwerk van neppe *assets* die door dezelfde persoon of personen worden beheerd, met als doel Meta of onze community te misleiden of handhaving van de Community Standards te ontwijken. Als deze gebruikers niet authentiek zijn (dus geen echte mensen vertegenwoordigen), worden ze door onze systemen verwijderd. We hebben ook automatische verdedigingsmechanismen via ons spambeleid tegen *gescripte* accountcreatie of grootschalige gecoördineerde acties. We handhaven niet op authentieke accounts, ongeacht diens interacties met politieke pagina's of profielen.



Let wel, de correlatie tussen engagement en content over politieke partijen kan het gevolg zijn van authentieke politieke voorkeuren van gebruikers. We beschermen de uiting van politieke voorkeuren en stemmen op onze platforms en faciliteren civiele discussie en engagement tussen authentieke accounts.

We hebben beleid en geautomatiseerde waarborgen tegen onechtheid en identiteit, met als doel zoveel mogelijk nepaccounts op Facebook en Instagram te verwijderen. Onze detectietechnologie blokkeert dagelijks miljoenen pogingen om nepaccounts aan te maken en detecteert er nog miljoenen meer, vaak binnen enkele minuten na aanmaak. Tussen juli en september 2025 hebben we 698 miljoen nep-accounts verwijderd, waarvan 99.9% door ons werd gevonden en aangepakt voordat gebruikers ze konden melden.

We blijven [onderzoek doen](#) en heimelijke beïnvloedingsnetwerken - zowel binnenlands als buitenlands - verwijderen via ons [Coordinated Inauthentic Behavior beleid \(CIB\)](#). Dit beleid is ontworpen om bijzonder geavanceerde vormen van onecht gedrag te voorkomen, zoals het gebruik van valse identiteiten of vijandige methoden om detectie te ontwijken of authentiek te lijken. Sinds 2017 hebben we meer dan 200 CIB-netwerken aangepakt, afkomstig uit meer dan 70 landen. We schalen onze CIB-handhaving verder op door de inzichten uit de wereldwijde onderzoeken terug te koppelen aan geautomatiseerde detectiesystemen die overtreders en recidivisten opsporen.

3) Aanbevelingssystemen

De afgelopen jaren hebben we ingezien dat gebruikers het niet altijd eenvoudig vinden om door de grote hoeveelheid communities en content op onze platforms te navigeren en te vinden wat voor hen waardevol is. We doen aanbevelingen aan gebruikers om hen te helpen om nieuwe communities en content te ontdekken. Zowel Facebook als Instagram kunnen in dit verband aanbevelingen doen waarvan wij denken dat mensen die waardevol vinden. De content die iemand ziet in de Facebook- en Instagram-feed aanbevelingen wordt door een AI-systeem geselecteerd en gerangschikt. Binnen een AI-systeem werken meerdere machine learning-modellen samen om deze ervaring te leveren. Deze modellen en hun input signalen, zijn dynamisch en veranderen regelmatig naarmate het systeem leert en verbetert.

Hoe ranking- en aanbevelingen algoritmen werken:

We gebruiken geavanceerde machine learning-modellen om te bepalen welke content aan gebruikers wordt aanbevolen en in welke volgorde deze content in de feed verschijnt. Samengevat is dit een [proces in 4 stappen](#):



- **Inventariseren:** Het systeem verzamelt alle potentieel interessante posts van vrienden, pagina's die iemand volgt en groepen waar iemand lid van is. Dit sluit posts uit die zijn verwijderd wegens schending van onze [Community Standards](#) of posts van lage kwaliteit of aanstootgevende posts die niet mogen worden aanbevolen volgens onze [Recommendation Guidelines](#).
- **Signalen benutten:** Vervolgens bekijkt het AI-systeem verschillende signalen over elke post, zoals wie de post heeft gemaakt, eerdere interacties, het type post (foto, video, link), en hoeveel vrienden de post leuk vonden.
- **Voorspellingen doen:** De signalen leiden tot tientallen voorspellingen over de kans dat een content-item waardevol is voor iemand.
- **Posts rangschikken op score:** Deze voorspellingen worden gewogen om een waardescore te creëren, die bepaalt in welke volgorde posts verschijnen wanneer iemand Facebook of Instagram opent.

System cards helpen mensen en het publiek onze algoritmen te begrijpen:

We hebben [22 system cards](#) gepubliceerd die uitleggen hoe de AI-gestuurde aanbevelingssystemen werken en meer details geven over onze signalen en voorspellingen. Deze *system cards* bieden transparantie over hoe we content rangschikken en aanbevelen op verschillende plekken in onze apps (bijvoorbeeld Facebook Stories of IG Reels).

Keuzes van gebruikers voeden onze algoritmen:

We gebruiken een combinatie van signalen om te bepalen wat in iemands feed verschijnt (zie hierboven), en bieden verschillende controles waarmee gebruikers expliciet feedback kunnen aangeven en voorkeur kunnen aanbrengen.

- **Gebruikers kunnen hun algoritme aanpassen:** Mensen kunnen bijvoorbeeld een lijst met “verborgen woorden” maken die niet meer aan hen worden aanbevolen. Ze kunnen ook aangeven of ze “geïnteresseerd” of “niet geïnteresseerd” zijn in aanbevolen content. Op Instagram hebben we [aangekondigd](#) dat gebruikers met de functie “Jouw Algoritme” direct kunnen aangeven welke soorten Reels ze meer of minder willen zien. Dit is sinds December 2025 beschikbaar voor alle gebruikers met de Engelse taalinstelling. We hopen dit ook binnenkort voor andere taalinstellingen mogelijk te maken.
- **Gebruikers kunnen kiezen voor een niet-algoritmisch gerangschikte ervaring:** Op Facebook en Instagram bieden we verschillende feedopties waarmee mensen kunnen kiezen hoe ze content zien in hun [feed](#). Bijvoorbeeld, door te schakelen naar de “Following”-feed krijgen gebruikers een feed die niet algoritmisch is



gerangschikt, maar in omgekeerde chronologische volgorde wordt weergegeven. Daarnaast kunnen gebruikers een “Favorieten”-feed bekijken, een lijst van accounts die ze zelf selecteren om bovenaan hun feed te verschijnen.

- Sinds 2024 kunnen gebruikers op Facebook, Instagram en Threads ook zelf bepalen hoeveel politieke content ze zien op onze platforms.

Overtredende content wordt van onze platforms verwijderd, terwijl content van lage kwaliteit niet mag worden aanbevolen.

- De veiligheid van gebruikers op onze platforms is onze prioriteit. De Community Standards van Meta beschrijven welke soorten content we volledig verwijderen, met behulp van (AI-gedreven) automatische en menselijke beoordeling.
- We stellen hoge eisen aan de content die we aanbevelen om kwaliteit en waarde te waarborgen. Daarom hebben we Recommendation Guidelines op [Facebook](#) en [Instagram](#) die bepaalde content uitsluiten van aanbevelingen, ook als deze niet de Community Standards schendt. Dit kan bijvoorbeeld aanstootgevende of gevoelige content zijn, zoals seksueel suggestieve content.