

Code of Conduct on Disinformation – Report of TikTok for the period 1 January 2026 - 30 June 2026



Table of Contents

Executive summary.....	12
Ensuring the Integrity of Services.....	14
Empowering Users.....	15
Empowering the Fact-Checking Community.....	16
Looking forward.....	16
II. Scrutiny of Ad Placements.....	17
Commitment 1.....	18
Measure 1.1.....	18
Measure 1.2.....	19
Measure 1.3.....	19
QRE 1.3.1.....	20
Measure 1.4.....	20
Measure 1.5.....	21
Measure 1.6.....	21
Commitment 2.....	22
Measure 2.1.....	22
Measure 2.2.....	25
Measure 2.3.....	28
Measure 2.4.....	34
Commitment 3.....	36
Measure 3.1.....	37
Measure 3.2.....	37
Measure 3.3.....	38
III. Political Advertising.....	39
Commitment 4.....	40
Measure 4.1.....	40
Measure 4.2.....	40
Commitment 5.....	41
Measure 5.1.....	41
Commitment 6.....	42
Measure 6.1.....	42
Measure 6.2.....	42
Measure 6.3.....	43
Measure 6.4.....	43
Measure 6.5.....	43
Commitment 7.....	44
Measure 7.1.....	44
Measure 7.2.....	45
Measure 7.3.....	45
Measure 7.4.....	45
Commitment 8.....	46
Measure 8.1.....	46

Measure 8.2.....	46
Commitment 9.....	47
Measure 9.1.....	47
Measure 9.2.....	47
Commitment 10.....	48
Measure 10.1.....	48
Measure 10.2.....	48
Commitment 11.....	49
Measure 11.1.....	49
Measure 11.2.....	49
Measure 11.3.....	49
Measure 11.4.....	49
Commitment 12.....	50
Measure 12.1.....	50
Measure 12.2.....	50
Measure 12.3.....	51
Commitment 13.....	51
Measure 13.1.....	51
Measure 13.2.....	51
Measure 13.3.....	52
IV. Integrity of Services.....	53
Commitment 14.....	54
Measure 14.1.....	55
Measure 14.2.....	63
Measure 14.3.....	92
Commitment 15.....	93
Measure 15.1.....	93
Measure 15.2.....	95
Commitment 16.....	96
Measure 16.1.....	96
Measure 16.2.....	97
V. Empowering Users.....	98
Commitment 17.....	99
Measure 17.1.....	101
Measure 17.2.....	107
Measure 17.3.....	127
Commitment 18.....	128
Measure 18.1.....	128
Measure 18.2.....	129
Measure 18.3.....	138
Commitment 19.....	138
Measure 19.1.....	139
Measure 19.2.....	141
Commitment 20.....	144

Measure 20.1.....	145
Measure 20.2.....	145
Commitment 21.....	145
Measure 21.1.....	147
Measure 21.2.....	157
Measure 21.3.....	158
Commitment 22.....	158
Measure 22.1.....	159
Measure 22.2.....	159
Measure 22.3.....	159
Measure 22.4.....	160
Measure 22.5.....	160
Measure 22.6.....	160
Measure 22.7.....	161
Commitment 23.....	162
Measure 23.1.....	162
Measure 23.2.....	163
Commitment 24.....	165
Measure 24.1.....	165
Commitment 25.....	175
Measure 25.1.....	176
Measure 25.2.....	176
VI. Empowering the research community.....	177
Commitment 26.....	178
Measure 26.1.....	178
Measure 26.2.....	179
Measure 26.3.....	183
Commitment 27.....	183
Measure 27.1.....	184
Measure 27.2.....	184
Measure 27.3.....	184
Measure 27.4.....	185
Commitment 28.....	185
Measure 28.1.....	185
Measure 28.2.....	186
Measure 28.3.....	187
Measure 28.4.....	187
Commitment 29.....	188
Measure 29.1.....	188
Measure 29.2.....	189
Measure 29.3.....	189
VII. Empowering the fact-checking community.....	190
Commitment 30.....	191
Measure 30.1.....	191

Measure 30.2.....	196
Measure 30.3.....	197
Measure 30.4.....	197
Commitment 31.....	197
Measure 31.1.....	198
Measure 31.2.....	198
Measure 31.3.....	205
Measure 31.4.....	205
Commitment 32.....	206
Measure 32.1.....	206
Measure 32.2.....	206
Measure 32.3.....	207
Commitment 33.....	207
Measure 33.1.....	208
VIII. Transparency Centre.....	209
Commitment 34.....	210
Measure 34.1.....	210
Measure 34.2.....	210
Measure 34.3.....	210
Measure 34.4.....	210
Measure 34.5.....	211
Commitment 35.....	211
Measure 35.1.....	211
Measure 35.2.....	211
Measure 35.3.....	212
Measure 35.4.....	212
Measure 35.5.....	212
Measure 35.6.....	212
Commitment 36.....	212
Measure 36.1.....	213
Measure 36.2.....	213
Measure 36.3.....	213
IX. Permanent Task-Force.....	214
Commitment 37.....	215
Measure 37.1.....	215
Measure 37.2.....	215
Measure 37.3.....	215
Measure 37.4.....	216
Measure 37.5.....	216
Measure 37.6.....	216
X. Monitoring of Code.....	217
Commitment 38.....	218
Measure 38.1.....	218
Commitment 39.....	219



Commitment 40.....	220
Measure 40.1.....	220
Measure 40.2.....	221
Measure 40.3.....	221
Measure 40.4.....	221
Measure 40.5.....	221
Measure 40.6.....	221
Commitment 41.....	221
Measure 41.1.....	222
Measure 41.2.....	222
Measure 41.3.....	222
Commitment 42.....	222
Commitment 43.....	223
War of aggression by Russia on Ukraine.....	225
Israel-Hamas Conflict.....	235
H1 2026 Elections.....	246

Commitments	Measures	Subscribed
II. Scrutiny of Ad Placements		
1	Measure 1.1	☐
	Measure 1.2	☐
	Measure 1.3	☒
	Measure 1.4	☐
	Measure 1.5	☐
	Measure 1.6	☐
2	Measure 2.1	☒
	Measure 2.2	☒
	Measure 2.3	☒
	Measure 2.4	☒
3	Measure 3.1	☒
	Measure 3.2	☒
	Measure 3.3	☒
III. Political advertising		
4	Measure 4.1	☐
	Measure 4.2	☐
5	Measure 5.1	☐
6	Measure 6.1	☐
	Measure 6.2	☐
	Measure 6.3	☐
	Measure 6.4	☐
	Measure 6.5	☐
7	Measure 7.1	☐
	Measure 7.2	☐
	Measure 7.3	☐
	Measure 7.4	☐
8	Measure 8.1	☐

	Measure 8.2	☐
9	Measure 9.1	☐
	Measure 9.2	☐
10	Measure 10.1	☐
	Measure 10.2	☐
11	Measure 11.1	☐
	Measure 11.2	☐
	Measure 11.3	☐
	Measure 11.4	☐
12	Measure 12.1	☐
	Measure 12.2	☐
	Measure 12.3	☐
13	Measure 13.1	☐
	Measure 13.2	☐
	Measure 13.3	☐
IV. Integrity of services		
14	Measure 14.1	☒
	Measure 14.2	☒
	Measure 14.3	☒
15	Measure 15.1	☒
	Measure 15.2	☒
16	Measure 16.1	☒
	Measure 16.2	☒
V. Empowering users		
17	Measure 17.1	☒
	Measure 17.2	☒
	Measure 17.3	☒
18	Measure 18.1	☐
	Measure 18.2	☒
	Measure 18.3	☐

19	Measure 19.1	☒
	Measure 19.2	☒
20	Measure 20.1	☐
	Measure 20.2	☐
21	Measure 21.1	☒
	Measure 21.2	☐
	Measure 21.3	☐
22	Measure 22.1	☐
	Measure 22.2	☐
	Measure 22.3	☐
	Measure 22.4	☐
	Measure 22.5	☐
	Measure 22.6	☐
	Measure 22.7	☒
23	Measure 23.1	☒
	Measure 23.2	☒
24	Measure 24.1	☒
25	Measure 25.1	☐
	Measure 25.2	☐
VI. Empowering the research community		
26	Measure 26.1	☒
	Measure 26.2	☒
	Measure 26.3	☒
27	Measure 27.1	☐
	Measure 27.2	☐
	Measure 27.3	☐
	Measure 27.4	☐
28	Measure 28.1	☒
	Measure 28.2	☒

	Measure 28.3	☒
	Measure 28.4	☐
29	Measure 29.1	☐
	Measure 29.2	☐
	Measure 29.3	☐
VII. Empowering the fact-checking community		
30	Measure 30.1	☒
	Measure 30.2	☒
	Measure 30.3	☒
	Measure 30.4	☒
31	Measure 31.1	☐
	Measure 31.2	☒
	Measure 31.3	☐
	Measure 31.4	☐
32	Measure 32.1	☒
	Measure 32.2	☒
	Measure 32.3	☒
33	Measure 33.1	☐
VIII. Transparency centre		
34	Measure 34.1	☒
	Measure 34.2	☒
	Measure 34.3	☒
	Measure 34.4	☒
	Measure 34.5	☒
35	Measure 35.1	☒
	Measure 35.2	☒
	Measure 35.3	☒
	Measure 35.4	☒
	Measure 35.5	☒

	Measure 35.6	☒
36	Measure 36.1	☒
	Measure 36.2	☒
	Measure 36.3	☒
IX. Permanent Task-Force		
37	Measure 37.1	☒
	Measure 37.2	☒
	Measure 37.3	☒
	Measure 37.4	☒
	Measure 37.5	☒
	Measure 37.6	☒
X. Monitoring of the Code		
38	Measure 38.1	☒
39	-	☐
40	Measure 40.1	☒
	Measure 40.2	☒
	Measure 40.3	☒
	Measure 40.4	☒
	Measure 40.5	☒
	Measure 40.6	☒
41	Measure 41.1	☒
	Measure 41.2	☒
	Measure 41.3	☒
42	-	☒
43	-	☒



Executive summary

About TikTok

TikTok's mission is to inspire creativity and bring joy. With more than 200 million people across Europe coming to TikTok every month, including 189 million in the EU, it's natural for people to hold different opinions. That's why we focus on a shared set of facts when it comes to issues that affect people's safety. A safe, authentic, and trustworthy experience is essential to achieving our goals. Transparency plays a key role in building that trust, allowing online communities and society to assess how TikTok meets its regulatory obligations. As a signatory to the Code of Conduct on Disinformation (the Code), TikTok is committed to sharing clear insights into the actions we take.

TikTok takes disinformation extremely seriously. We are committed to preventing its spread, promoting authoritative information, and supporting media literacy initiatives that strengthen community resilience.

We prioritise proactive content moderation, with the vast majority of violative content removed before it is reported. In H1 2026, more than 99% of videos violating our Integrity and Authenticity policies were removed proactively worldwide.

We continue to address emerging behaviours and risks through our Digital Services Act (DSA) compliance efforts, which the Code has operated under since July 2025.

Our actions under the Code demonstrate TikTok's strong commitment to combating disinformation while ensuring transparency and accountability to our community and regulators.

Our eighth report under the Code - 1 January to 30 June 2026

TikTok has been an active participant in the Code since 2020. We remain engaged in the Code's Taskforce, participate in all relevant working group discussions, and co-chair the working group on Elections (Crisis Response).

This eighth report provides detailed insights into the measures we implement to combat disinformation. Of note during this reporting period:

- We ran 14 temporary media literacy campaigns in countries with elections, whilst maintaining 14 media literacy and critical thinking campaigns across Europe.
- We removed more than 700,000 videos in the EEA for breaking our rules against misinformation.
- We approved more than 180 applications to use our Research Tools and study our platform.
- We have continued to invest in combating harmful AI-Generated Content (AIGC), consulting with experts, and partnering with peers on shared solutions. Several measures extend beyond EEA member states, including EU candidate countries and nations with significant diaspora communities, thereby helping limit disinformation in Europe.

TikTok remains committed to protecting our community from disinformation, strengthening resilience against misinformation, and upholding transparency and accountability throughout electoral and crisis contexts.

We are reporting on seven elections, reflecting the high volume of elections that took place across Europe in H1 2026, and with that, elevated risks of inauthentic behaviour and attempts to mislead people or our systems in order to influence public discussion. This report details our efforts to safeguard users and protect platform integrity during elections, including:

- Enforcing robust policies around misinformation, harmfully misleading AI-Generated Content, deceptive behaviour, and political advertising;
- Elevating reliable information from authoritative sources;



- Collaborating with our 13 IFCN-accredited fact-checking partners and other external experts to strengthen our approach.

During this period, we devoted multiple resources to election integrity, including:

- Operating Mission Control Centres for dedicated monitoring;
- Participating in the Code's Rapid Response System to streamline coordination with civil society, fact-checkers, and platforms;
- Launching dedicated Election Centres with localised media literacy campaigns;
- Launching up-to-date, detailed coverage of elections across Europe on our [Global Elections Integrity Hub](#) within the Trust & Safety Center.

We published seven specific post-election chapters in this report documenting our approach to the Portugal presidential election, Slovenia parliamentary election, Denmark general election, Bulgaria parliamentary election, Hungary parliamentary election, Cyprus parliamentary election, and Malta general election. We continue to provide dedicated crisis reports on the ongoing conflicts in Israel/Gaza and Russia/Ukraine, outlining our responses and safeguards.

Our policies

Our Integrity and Authenticity policies are designed to help promote a trustworthy, authentic experience for our users. Our policies cover specific types of misinformation and deceptive behaviours, as well as misleading AI-generated content, conspiracy theories, Covert Influence Operations (CIO) networks, and public safety events like natural disasters. We do not allow false or misleading content that may cause significant harm to individuals or society, regardless of intent. We do not allow misinformation or content about civic and electoral processes that may result in voter interference, disrupt the peaceful transfer of power, or lead to off-platform violence. Our policies are thoughtfully crafted to cover a broad range of content and the constantly changing nature of misinformation trends, often based on what's happening in the world. We also tackle deceptive behaviour by removing accounts that seek to mislead people or engage in platform manipulation.

Enforcing our policies

Disinformation presents unique challenges. It is highly complex, evolves quickly, and often requires context. Our specialised Integrity and Authenticity moderation team receives training to assess, confirm, and take action on harmful misinformation. They also have access to our [independent fact-checking partners](#) to help evaluate content accuracy. We place considerable emphasis on proactive detection and automated moderation technology to action violative content. In H1 2026, 93% of the violative videos we removed globally were taken down through automated and AI-based technology. We use machine learning models to help detect potential misinformation, and we rely on automated moderation when our systems have a high degree of confidence that content is violative. To further support proactive moderation at scale, we continued testing large language models (LLMs). Because LLMs can comprehend human language and perform highly specific, complex tasks, we are better able to moderate nuanced areas like misinformation by extracting specific misinformation "claims" from videos for our moderators to assess directly or route to our fact-checking partners.

We are transparent with our community about [how we moderate](#) content and [what moderation actions we take](#). This includes details about what content we make [ineligible for the For You Feed](#). We disclose the underlying metrics in this report.

We also address disinformation by removing accounts that repeatedly post misinformation that violates our policies, and have expert teams who continuously monitor potential disinformation campaigns, inauthentic behavior, and influence operations.



Transparency and Scrutiny of Advertising

Under our misinformation advertising policy, we are committed to maintaining a safe and trustworthy environment by taking action against misinformation, such as false, misleading, or unsubstantiated content, and manipulated content that misleads viewers and could harm individuals or society, regardless of intent.

Like all users of our platform, participants in content monetisation programs must adhere to our Community Guidelines, including our Integrity and Authenticity policies. Those policies make clear that we do not allow activities that may undermine the integrity of our platform or the authenticity of our user accounts. They also make clear that we remove content or accounts, including those of creators, which contain misleading information that causes significant harm or engage in deceptive behaviours. In certain scenarios, we may remove a creator's access to a creator monetisation feature. Our [Creator Code of Conduct](#) outlines the standards we expect creators involved in TikTok programs, features, events and campaigns to uphold, both on and off-platform, including in relation to misinformation-related activities.

Prohibiting Political Ads

TikTok is first and foremost an entertainment platform, and we're proud to be a place that brings people together through creative, engaging content. While sharing political beliefs and participating in political conversation is allowed as organic content on TikTok, our policies prohibit our community, including politicians and political party accounts, from placing [political ads](#) or posting political [branded content](#). We also prevent [governments, politicians and political party accounts](#) from accessing our monetisation features and campaign fundraising (with a limited exception for government-run public service announcements such as health campaigns).

We have continued to focus on the enforcement of our political advertising prohibition since the EU Regulation on the Transparency and Targeting of Political Advertising came into full application late last year.

By prohibiting political advertising, campaign fundraising, and limiting access to certain monetisation features, we're aiming to strike a balance between enabling people to discuss the issues that are relevant to their lives while also protecting the creative, entertaining platform that our community wants.

Ensuring the Integrity of Services

Our [Integrity and Authenticity policies](#) strictly prohibit deceptive behaviors, and we are committed to combating manipulative practices. We remain highly vigilant about the evolving disinformation tactics, techniques, and procedures employed by malicious actors, as outlined in detail in Commitment 14. To effectively combat these threats, we continuously assess and refine our policies, ensuring they remain robust and responsive to emerging challenges in the information ecosystem. Through proactive monitoring and enforcement, we aim to safeguard the integrity of our platform and protect users from harmful influence operations.

We continue to fight against covert influence operations (CIO), and we do not allow attempts to sway public opinion while misleading our platform's systems or community about the identity, origin, operating location, popularity, or purpose of the account. As shown in our H1 2026 reporting data, TikTok remains highly vigilant in detecting and disrupting CIOs that seek to manipulate public discourse through coordinated deceptive behaviour.



To provide granular transparency, TikTok maintains a [dedicated CIO Transparency Report within our Trust & Safety Centre](#), detailing accounts removed for attempting to mislead the community or platform systems. During this period, our expert teams—comprising specialists in threat intelligence, data science, and law enforcement—focused on behavioural signals and technical linkages to identify coordinated efforts.

We've continued to strengthen our transparency work in Artificial Intelligence (AI) through our Edited Media and AI-Generated Content (AIGC) policy, which addresses the use of content created or modified by AI on our platform. To support authentic and transparent experiences for our community, we require creators to label content that has been either completely generated or significantly edited by AI and disclose when their content shows realistic scenes. We detect AI-generated content through a combination of proactive technologies, alerts from experts and fact-checking partners, searches for clips or keywords related to known AI-generated content, and user reports. Since joining the [Coalition for Content Provenance and Authenticity \(C2PA\)](#) and the [Content Authenticity Initiative \(CAI\)](#) in May 2024, we've continued to collaborate on driving industry adoption of Content Credentials, a technology that helps platforms more easily label AI-generated content.

TikTok is also proud to actively participate in the Code's Rapid Response System, which streamlines the exchange of information among civil society organisations, fact-checkers, and online platforms.

Empowering Users

We continuously work to deter misinformation proactively by empowering our community with resources that help them recognise misinformation, critically assess content, and file reports about potentially violative content. We offer our community in Europe easy-to-use in-app and online reporting tools so they can alert us to content or accounts they believe may violate our Community Guidelines or break the law. Content that is reported for being illegal will be reviewed against our policies, and where a violation is detected, the content may be removed globally.

We continue to dedicate resources to expanding our in-app measures that show users additional context on certain content (e.g., natural disasters and rapidly unfolding events), and to redirect them to authoritative information. We make these tools available in 23 EU official languages, and Norwegian and Icelandic for EEA users. For example, during this reporting period, we launched temporary in-app search guides to help users navigate sensitive events with authoritative information, including public health crises such as the outbreak of Ebola in Sub-Saharan Africa (Mayotte) and the outbreak of Hantavirus. These guides connected users to TikTok's Safety Center and authoritative third-party resources on aid and relief support.

We maintained 14 ongoing media literacy and critical thinking campaigns across Europe, in Germany, Romania, Poland, Denmark, Finland, France, Georgia, Ireland, Italy, Moldova, Portugal, Spain, Sweden, and the Netherlands, some in collaboration with fact-checking and media literacy partners.

We ran 14 temporary media literacy election integrity campaigns, all with in-app Election Centers or Search Guide and Details Page, in advance of elections. For the elections in Portugal, Slovenia, Denmark, Bulgaria, Hungary, Cyprus, and Malta, our search interventions were seen by users more than 9.4 million times during the reporting period. And, to better inform us about our approach to upcoming elections, we hosted Election Speaker Series, inviting external experts, including those from the fact-checking community, to share their insights and market expertise with our internal teams.

We continue to offer greater transparency to users about our systems and integrity and authenticity efforts. Our Trust & Safety Centre provides a wealth of information about how we [counter deceptive behaviour](#), protect the integrity of specific [global elections](#), disrupt [Covert Influence Operations](#). We also provide regular updates about broader trust and safety work through the [TikTok Transparency Center blog](#).

Empowering the Research Community

We recognise the important role of researchers in helping to identify disinformation trends and practices. We maintain, and continue to iterate, our [Research API](#) (providing researchers in Europe with access to public data on content and accounts from our platform) as well as our [Commercial Content API](#) (bringing transparency to paid advertising and other content that is commercial in nature on TikTok) and [Commercial Content Library](#) (a publicly searchable EU ads database with information about paid ads and ad metadata). We also continue to refine the [Virtual Compute Environment](#) (VCE), which offers broader access to user data to qualifying not-for-profit researchers for querying and analyzing applicable data while ensuring robust security and privacy protections.

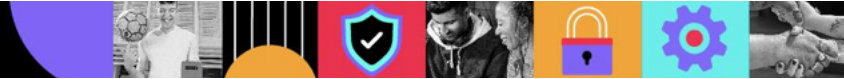
Empowering the Fact-Checking Community

TikTok recognises the important contribution of our fact-checking partners in the fight against disinformation. We work closely with 13 [JECN-accredited](#) fact-checking organisations across the EU, EEA, and wider Europe, which have technical training, resources, and industry-wide insights to impartially assess online misinformation.

Our fact-checking programme incorporates fact-checker input into our broader content moderation efforts. Fact-checkers do not moderate content directly on TikTok, but assess whether a claim is true, false, or unsubstantiated. They also share proactive insight reports that help us detect harmful misinformation and anticipate misinformation trends. This feedback from fact-checkers is relayed to TikTok's moderation teams so that they can ensure it is factored into their moderation work and take the relevant action based on our Community Guidelines. This approach effectively produces a force multiplier to the underlying fact-checking output, ensuring that the disinformation content or trend is more comprehensively and broadly addressed. Our expanding fact-checking repository ensures that our teams and systems fully utilise the insights provided by our fact-checking partners on TikTok.

Looking forward

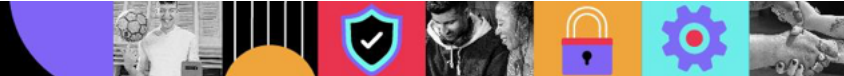
TikTok remains fully committed to the Code and to meaningful collaboration with industry peers, civil society, and EU authorities. By continuously sharing expertise, strengthening our policies, and enhancing enforcement mechanisms, we aim to stay ahead of the evolving tactics of malicious actors. TikTok will continue to iterate on its automated systems to improve the accuracy and coverage of misinformation moderation, ensuring the platform remains a safe and authentic space for creative expression.



II. Scrutiny of Ad Placements Commitments 1 - 3



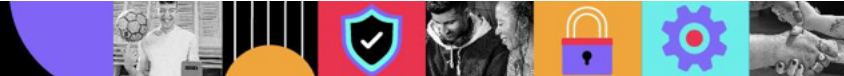
II. Scrutiny of Ad Placements	
Commitment 1	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points]. Guidance to support the identification of policies. Improving identification. Improvement of the enforcement of the policies themselves (not the policy wording).	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	No
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 1.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 1.1.1	N/A



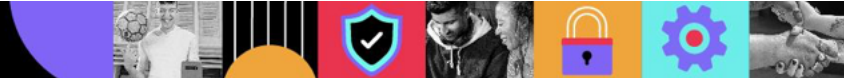
SLI 1.1.1 – Numbers by actions enforcing policies above (specify if at page and/or domain level)	N/A		
SLI 1.1.2 -	N/A		

Measure 1.2	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .			
QRE 1.2.1	N/A			
SLI 1.2.1	N/A			
Member States	N/A	N/A	N/A	N/A
Total EU				
Total EEA				

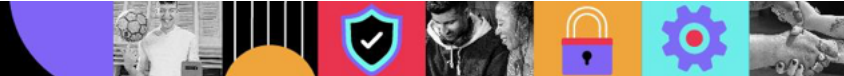
Measure 1.3	
-------------	--



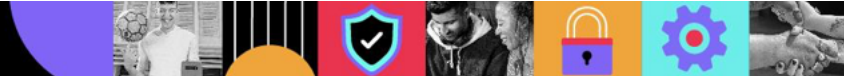
QRE 1.3.1	We partner with a number of industry leaders to provide a number of controls and transparency tools to advertising buyers with regard to the placement of ads: Controls: We offer pre-campaign solutions to advertisers so they can put additional safeguards in place before their campaign goes live to mitigate the risk of their advertising being displayed adjacent to certain types of user-generated content. These measures are in addition to the Community Guidelines, which provide overarching rules around the types of content that can appear on TikTok and are eligible for the For You feed: • TikTok Inventory Filter: This is our proprietary system, which enables advertisers to choose the profile of content they want their ads to run adjacent to. We expanded our Inventory Filter, which is available in various EEA markets designated as fully or de minimis monetised on the TikTok app, and is embedded directly in TikTok Ads Manager, the main system through which advertisers purchase ads. More details can be found here. The Inventory Filter is informed by industry standards and policies, which include topics that may be susceptible to disinformation. Additionally, this enables advertisers to: ○ Selectively exclude unwanted or videos that do not align with their brand safety requirements from appearing next to their ads through TikTok's Video Exclusion List solution. ○ Exclude specific profile pages from serving their Profile Feed ads through TikTok's Profile Feed Exclusion List. • TikTok Pre-bid Brand Safety Solution by Integral Ad Science ("IAS"): Advertisers can filter content based on industry-standard frameworks with all levels of risk (available in France and Germany). Some misinformation content may be captured and filtered out by these industry standard categories, such as "Sensitive Social Issues". Transparency: We have partnered with third parties to offer post-campaign solutions that enable advertisers to assess the suitability of user content that ran immediately adjacent to their ad in the For You feed against their chosen brand suitability parameters: • Zefr: Through our partnership with Zefr, advertisers can obtain campaign insights into brand suitability and safety on the platform. Zefr aligns with industry standards. • IAS: Advertisers can measure brand safety, viewability, and invalid traffic on the platform with the IAS Signal platform. As with IAS's pre-bid solution covered above, this aligns with industry standards. • DoubleVerify: We are partnering with DoubleVerify to provide advertisers with media quality measurement for ads. DoubleVerify is working actively with us to expand its suite of brand suitability and media quality solutions on the platform.
Measure 1.4	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document.



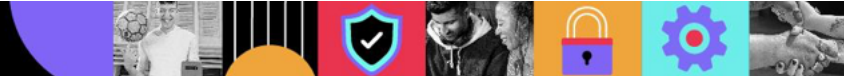
QRE 1.4.1	N/A
Measure 1.5	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 1.5.1	N/A
QRE 1.5.2	N/A
Measure 1.6	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 1.6.1	N/A
QRE 1.6.2	N/A
QRE 1.6.3	N/A
QRE 1.6.4	N/A
SLI 1.6.1	N/A



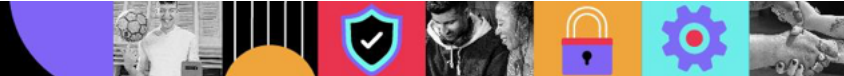
II. Scrutiny of Ad Placements	
Commitment 2	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	Yes
If yes, which further implementation measures do you plan to put in place in the next 6 months?	We continue to focus on improving the accuracy and coverage of our automated misinformation moderation systems for advertising.
Measure 2.1	
QRE 2.1.1	We continue to implement and enforce our granular advertising policies for misinformation in the EEA, which advertisers need to comply with. These policies were iterated in H2 2025 and currently cover: • Health Misinformation • Environment/Climate Misinformation • Public Safety & Trust Misinformation • Election Misinformation • Other Misinformation



	The policies provide clearer categorisation of misinformation types and build on the principles and enforcement experience of the policies previously set out in the H1 2025 report, enabling more consistent and targeted enforcement in line with evolving risks. Our automated detection supports the enforcement of our new misinformation advertising policies, and we continue to develop our models to optimise operationalisation of the misinformation advertising policies. We provide users with a simple and intuitive way to report advertisements in-app for breach of our misinformation advertising policies in the EEA. Our advertiser account policies expressly prohibit deceptive behaviours, including prohibiting advertisers from circumventing, evading, or interfering with our advertising systems and processes.	
SLI 2.1.1 – Numbers by actions enforcing policies above	Methodology of data measurement: We have set out the number of ads that have been removed from our platform for violation of our granular misinformation advertising policies on Health Misinformation, Environment/Climate Misinformation, Public Safety & Trust Misinformation, Election Misinformation, and Other Misinformation. We launched these iterated misinformation policies in August 2025. These policies were developed to provide clearer categorisation and more targeted, risk-based enforcement. We are pleased to be able to report on the ads removed for breach of our granular misinformation advertising policies. We have provided the political advertising enforcement metrics in the Elections Chapter of this Report. Note that numbers have only been provided for monetised markets and are based on where the ads were displayed.	
	Number of ads removals under the granular misinformation ad policies	
Member States		

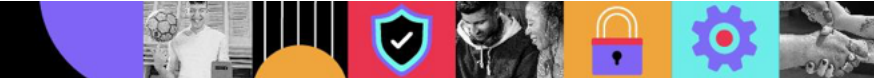


Austria	19	
Belgium	20	
Bulgaria	1	
Croatia	2	
Cyprus	0	
Czechia	7	
Denmark	3	
Estonia	0	
Finland	1	
France	37	
Germany	12	
Greece	0	
Hungary	4	
Ireland	0	
Italy	19	
Latvia	0	
Lithuania	0	
Luxembourg	0	
Malta	0	



Netherlands	8	
Poland	3	
Portugal	1	
Romania	4	
Slovakia	0	
Slovenia	0	
Spain	8	
Sweden	2	
Iceland	0	
Liechtenstein	0	
Norway	0	
Total EU	151	
Total EEA	151	

Measure 2.2	
QRE 2.2.1	TikTok places considerable emphasis on proactive moderation of advertisements. Advertisements are reviewed against our Advertising Policies through a combination of automated and human moderation. Our granular misinformation advertising policies currently cover: • Health Misinformation • Environment/Climate Misinformation



- Public Safety & Trust Misinformation
- Election Misinformation
- Other Misinformation

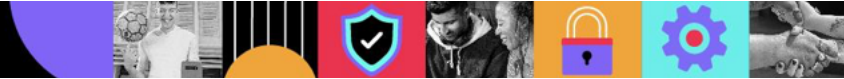
Our advertiser account policies expressly prohibit deceptive behaviours, including prohibiting advertisers circumventing, evading, or interfering with our advertising systems and processes.

We provide users with a simple and intuitive way to report advertisements in-app for breach of our Advertising Policies including for misinformation in the EEA.

There are two main ways to report an advertisement on TikTok, either:

- By ‘long-pressing’ (e.g., clicking and holding for 3 seconds) on the advertisement and selecting the “Report” option.
- By selecting the “Share” button available on the right-hand side of the advertisement and then selecting the “Report” option.

The user is then shown categories of reporting reasons from which to select. This feature has a specific “Misinformation” category and allows users to select a sub-category to report the reason with increased granularity.



Submit

×

Spam

>

Offensive

>

Sexually inappropriate

>

Misinformation

>

Intellectual property violation

>

Prohibited or dangerous

>

Fraud or scam

>

Report content that is harmful to children

>

Report illegal content

>

Political ads

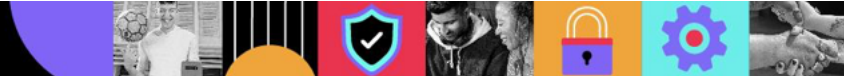
>

Others

>



	< Select report reason × Health misinformation > Environment/climate change misinformation > Election misinformation > Public safety & trust misinformation > Misleading AI-generated content & identity misuse > Other false or misleading claims >
Measure 2.3	



QRE 2.3.1

TikTok places considerable emphasis on proactive moderation of advertisements. Advertisements are reviewed against our Advertising Policies through a combination of automated and human moderation.

Our granular misinformation advertising policies launched currently cover:

- Health Misinformation
- Environment/Climate Misinformation
- Public Safety & Trust Misinformation
- Election Misinformation
- Other Misinformation

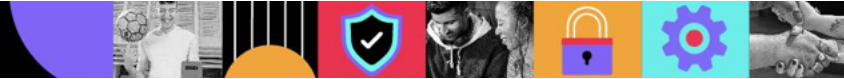
Our advertiser account policies expressly prohibit deceptive behaviours, including prohibiting advertisers circumventing, evading, or interfering with our advertising systems and processes.

We provide users with a simple and intuitive way to report advertisements in-app for breach of our Advertising Policies including for misinformation in the EEA.

There are two main ways to report an advertisement on TikTok, either:

- By ‘long-pressing’ (e.g., clicking and holding for 3 seconds) on the advertisement and selecting the “Report” option.
- By selecting the “Share” button available on the right-hand side of the advertisement and then selecting the “Report” option.

The user is then shown categories of reporting reasons from which to select. This feature has a specific “Misinformation” category and allows users to select a sub-category to report the reason with increased granularity.



Submit

×

Spam

>

Offensive

>

Sexually inappropriate

>

Misinformation

>

Intellectual property violation

>

Prohibited or dangerous

>

Fraud or scam

>

Report content that is harmful to children

>

Report illegal content

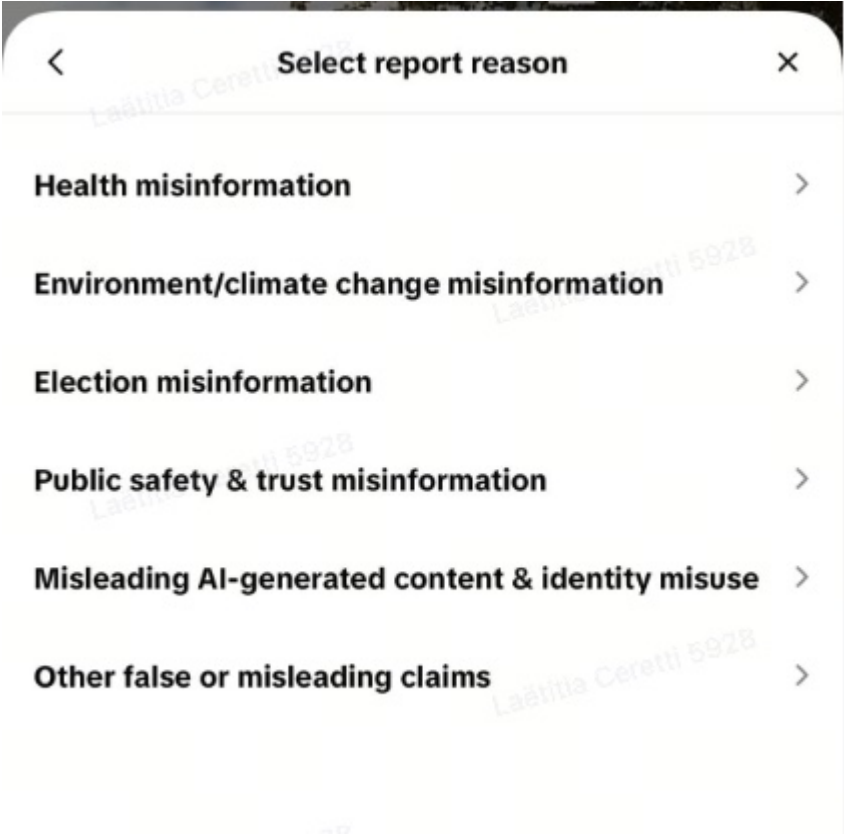
>

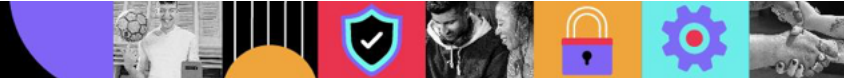
Political ads

>

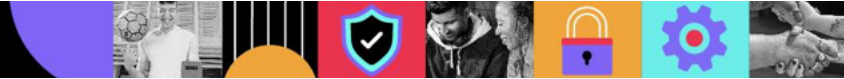
Others

>

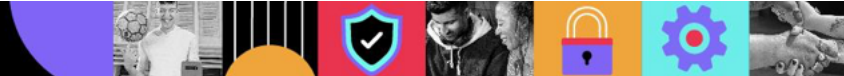
	
SLI 2.3.1	We are pleased to be able to report on the ads removed for breach of our granular misinformation advertising policies, including the impressions of those ads in this report. We have provided the political advertising enforcement metrics in the Elections Chapter of this Report.



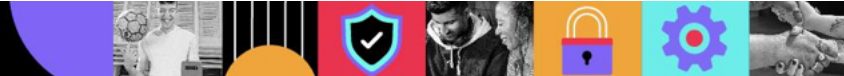
	Number of ads removals under the misinformation policies	Number of impressions for ads removed under the misinformation policies		
Member States				
Austria	19	0		
Belgium	20	0		
Bulgaria	1	0		
Croatia	2	0		
Cyprus	0	0		
Czechia	7	5576		
Denmark	3	0		
Estonia	0	0		
Finland	1	0		
France	37	4461453		
Germany	12	0		
Greece	0	0		
Hungary	4	0		
Ireland	0	0		
Italy	19	10870		



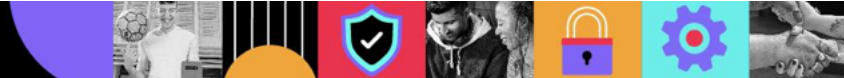
Latvia	0	0		
Lithuania	0	0		
Luxembourg	0	0		
Malta	0	0		
Netherlands	8	38022		
Poland	3	0		
Portugal	1	0		
Romania	4	0		
Slovakia	0	0		
Slovenia	0	0		
Spain	8	105691		
Sweden	2	0		
Iceland	0	0		
Liechtenstein	0	0		
Norway	0	0		
Total EU	151	4621612		
Total EEA	151	4621612		



Measure 2.4				
QRE 2.4.1	We are clear with advertisers that their ads must comply with our strict ad policies (see TikTok Business Help Centre). We explain that all ads are reviewed before being uploaded on our platform - usually within 24 hours. Ads already on TikTok may go through an additional stage of review if they are reported or if certain conditions are met (e.g. reaching certain impression thresholds). Where an advertiser has violated an ad policy, they are informed by way of a notification. This is visible in their TikTok Ads Manager account and/or sent by email (if they have provided a valid email address), or where an advertiser has booked their ad through a TikTok representative, then the representative will inform the advertiser of any violations. Advertisers are able to make use of the functionality to appeal rejections of their ads. Transparency is an important part of our overarching DSA compliance efforts. Notifications of restrictions include the restriction itself, reason for restriction, whether we made that decision by automated means, how we came to detect the violation (e.g. as a result of a user report or proactive TikTok initiatives) and what their rights of redress are. Advertisers can access an online functionality to appeal restrictions on their account or ads. These appeals are then also reviewed against our ad policies and additional information could be provided to advertisers to help them understand the violation and what to do about it.			
SLI 2.4.1	We are pleased to be able to share the number of appeals for ads removed under our granular misinformation advertising policies as well as the number of overturns.			
	Number of appeals for ads removed under the granular misinformation ad policies	Number of overturns of appeals under the granular misinformation ad policies		
Member States				
Austria	1	1		
Belgium	1	0		
Bulgaria	0	0		

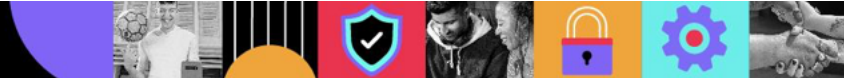


Croatia	0	0		
Cyprus	0	0		
Czechia	1	1		
Denmark	0	0		
Estonia	0	0		
Finland	0	0		
France	0	0		
Germany	1	0		
Greece	0	0		
Hungary	0	0		
Ireland	0	0		
Italy	0	0		
Latvia	0	0		
Lithuania	0	0		
Luxembourg	0	0		
Malta	0	0		
Netherlands	0	0		
Poland	1	0		
Portugal	0	0		

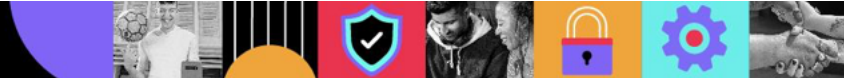


Romania	0	0		
Slovakia	0	0		
Slovenia	0	0		
Spain	0	0		
Sweden	0	0		
Iceland	0	0		
Liechtenstein	0	0		
Norway	0	0		
Total EU	5	2		
Total EEA	5	2		

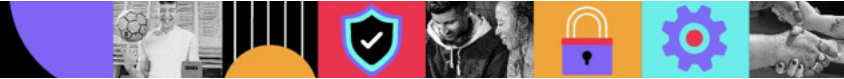
II. Scrutiny of Ad Placements	
Commitment 3	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	N/A



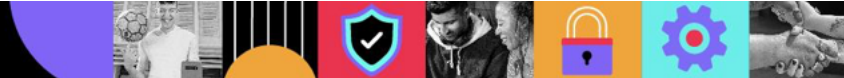
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	Yes
If yes, which further implementation measures do you plan to put in place in the next 6 months?	We will continue to enhance our misinformation detection capabilities through two key initiatives: • Optimising our collaboration framework with third-party fact-checking organisation in relation to advertising (i.e. Science Feedback); and • Continuing to enhance detection within the advertising ecosystem through signal-sharing to improve our internal databases.
Measure 3.1	
QRE 3.1.1	We cooperate with third party organisations and participate in conferences to facilitate the flow of information that may be relevant for tackling purveyors of harmful misinformation. This information is shared internally to help ensure consistency of approach across our platform. In this reporting period, we continued to partner with third-party fact-checking organisation Science Feedback. Science Feedback provides verified misinformation claims and signals which are integrated into our moderation workflows for ads. We also source claims and signals from across the platform to further enhance our moderation. We remain supportive of close collaboration with industry to ensure alignment and clarity on the reporting of these code requirements.
Measure 3.2	
QRE 3.2.1	We work closely with organisations such as TAG in the EEA and globally. TikTok participates in TAG audits on an annual basis, where our brand safety processes are reviewed and verified. We continue to share relevant insights and metrics within our quarterly transparency reports, which aim to inform industry



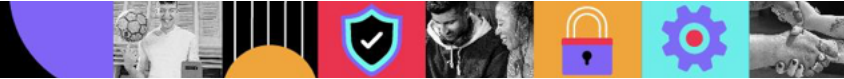
	peers and the research community. We remain supportive of close industry collaboration to ensure alignment and clarity in reporting against the Code's requirements.
Measure 3.3	
QRE 3.3.1	Our continued partnership with third-party fact-checking organisation Science Feedback provides verified claims that are integrated into our moderation workflows for ads. We also continue to work closely with organisations such as TAG in the EEA and globally.



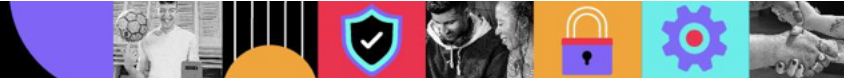
III. Political Advertising Commitments 4 - 13



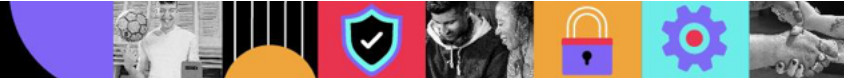
III. Political Advertising	
Commitment 4	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 4.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
Measure 4.2	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 4.1.1 (for measures 4.1 and 4.2)	N/A
QRE 4.1.2 (for measures 4.1 and 4.2)	N/A



III. Political Advertising	
Commitment 5	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 5.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 5.1.1	N/A

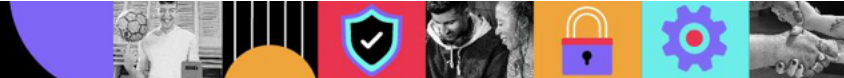


III. Political Advertising	
Commitment 6	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 6.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 6.1.1	N/A
Measure 6.2	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 6.2.1	N/A
QRE 6.2.2	N/A

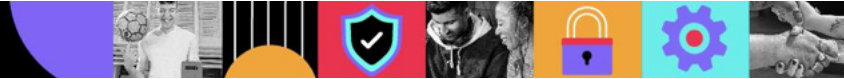


SLI 6.2.1 – numbers for actions enforcing policies above	N/A		
Member States	N/A	N/A	N/A
Total EU	N/A	N/A	N/A
Total EEA	N/A	N/A	N/A

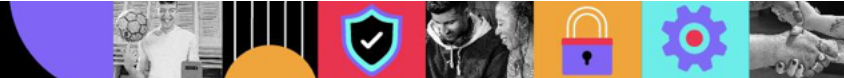
Measure 6.3	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 6.3.1	N/A
Measure 6.4	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 6.4.1	N/A
Measure 6.5	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 6.5.1	N/A



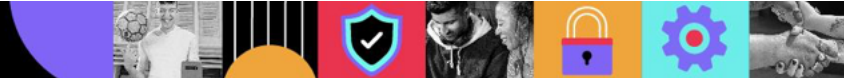
III. Political Advertising	
Commitment 7	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 7.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 7.1.1	N/A
SLI 7.1.1 – numbers for actions enforcing policies above (comparable metrics as for SLI 6.2.1)	N/A



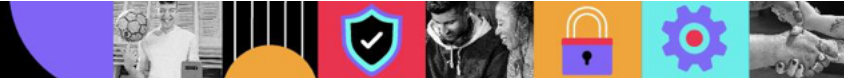
Measure 7.2	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 7.2.1	N/A
QRE 7.2.2	N/A
Measure 7.3	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 7.3.1	N/A
QRE 7.3.2	N/A
Measure 7.4	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 7.4.1	N/A



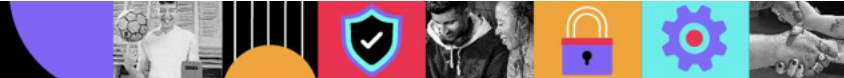
III. Political Advertising	
Commitment 8	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 8.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
Measure 8.2	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 8.2.1 (for measures 8.1 and 8.2)	N/A



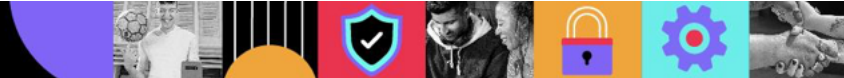
III. Political Advertising	
Commitment 9	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 9.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
Measure 9.2	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 9.2.1 (for measures 9.1 and 9.2)	N/A



III. Political Advertising	
Commitment 10	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 10.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
Measure 10.2	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 10.2.1 (for measures 10.1 and 10.2)	N/A

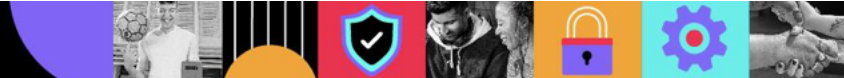


III. Political Advertising	
Commitment 11	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 11.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
Measure 11.2	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
Measure 11.3	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
Measure 11.4	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .



QRE 11.1.1 (for measures 11.1-11.4)	N/A
QRE 11.4.1	N/A

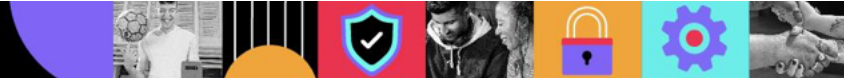
III. Political Advertising	
Commitment 12	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 12.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
Measure 12.2	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .



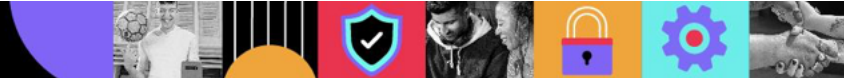
Measure 12.3	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 12.1.1 (for measures 12.1-12.3)	N/A

III. Political Advertising	
Commitment 13	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 13.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
Measure 13.2	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .

Measure 13.3	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 13.1.1 (for measures 13.1-13.3)	N/A



IV. Integrity of Services Commitments 14 - 16



IV. Integrity of Services

Commitment 14

In order to limit impermissible manipulative behaviours and practices across their services, Relevant Signatories commit to put in place or further bolster policies to address both misinformation and disinformation across their services, and to agree on a cross-service understanding of manipulative behaviours, actors and practices not permitted on their services. Such behaviours and practices, which should periodically be reviewed in light with the latest evidence on the conducts and TTPs employed by malicious actors, such as the AMITT Disinformation Tactics, Techniques and Procedures Framework, include:

The following TTPs pertain to the creation of assets for the purpose of a disinformation campaign, and to ways to make these assets seem credible:

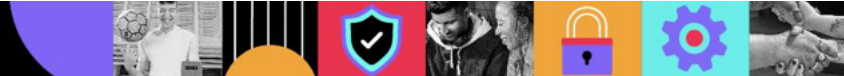
- 1. Creation of inauthentic accounts or botnets (which may include automated, partially automated, or non-automated accounts)
- 2. Use of fake / inauthentic reactions (e.g. likes, up votes, comments)
- 3. Use of fake followers or subscribers
- 4. Creation of inauthentic pages, groups, chat groups, fora, or domains
- 5. Account hijacking or impersonation

The following TTPs pertain to the dissemination of content created in the context of a disinformation campaign, which may or may not include some forms of targeting or attempting to silence opposing views. Relevant TTPs include:

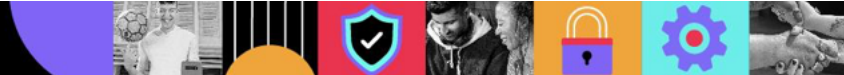
- 6. Deliberately targeting vulnerable recipients (e.g. location spoofing or obfuscation)
- 7. Deploy deceptive manipulated media (e.g. “deep fakes”, “cheap fakes”...)
- 8. Use “hack and leak” operation (which may or may not include doctored content)
- 9. Inauthentic coordination of content creation or amplification, including attempts to deceive/manipulate platforms algorithms (e.g. keyword stuffing or inauthentic posting/reposting designed to mislead people about popularity of content)-
- 10. Use of deceptive practices to deceive/manipulate platform algorithms, such as to create, amplify or hijack hashtags, data voids, filter bubbles, or echo chambers
- 11. Non-transparent compensated messages or promotions by influencers
- 12. Coordinated mass reporting of non-violative opposing content or accounts

In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]

Yes



If yes, list these implementation measures here [short bullet points].	• Our 2025 Community Guidelines update went live on September 13, 2025, ensuring that they remain aligned with our internal policies and were in force during the reporting period. ◦ Our Harmful Misinformation policies are referenced under the hack and leak section. They have all been refined in H2 2025, and they continue to drive our work in combating harmful misinformation. • We have enhanced our ability to detect covert influence operations, as detailed in our dedicated Covert Influence Operations Transparency Report available in our Trust & Safety Centre. • We updated and refined our policies around Covert Influence Operations in order to stay agile to changing behaviours and tactics on the platform and to ensure more granular detail is enshrined in our policy rationales.
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 14.1	



QRE 14.1.1

Our Integrity and Authenticity policies in our Community Guidelines expressly prohibit deceptive behaviours and relate to the TTPs as follows:

TTPs which pertain to the creation of assets for the purpose of a disinformation campaign, and the ways to make these assets seem credible:

Creation of inauthentic accounts or botnets (which may include automated, partially automated, or non-automated accounts)

Our Integrity and Authenticity policies expressly prohibit account behaviours that may spam or mislead our community. These include:

- Operating large networks of accounts controlled by a single entity, or through automation;
- Bulk distribution of a high volume of spam; and
- Manipulation of engagement signals to amplify the reach of certain content, or buying and selling followers, particularly for financial purposes

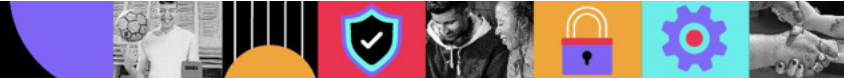
We also do not allow impersonation, including:

- Accounts that pose as another real person or entity without disclosing that they are a fan or parody account in the account name, such as using someone's name, biographical details, content, or image without disclosing it
- Presenting as a person or entity that does not exist (a fake persona) with a demonstrated intent to mislead others on the platform

Use of fake / inauthentic reactions (e.g. likes, up votes, comments) and use of fake followers or subscribers

Our Integrity and Authenticity policies do not allow the trade or marketing of services that attempt to artificially increase engagement or deceive TikTok's recommendation system. We do not allow our users to:

- facilitate the trade or marketing of services that artificially increase engagement, such as selling followers or likes; or
- provide instructions on how to artificially increase engagement on TikTok.

**Creation of inauthentic pages, groups, chat groups, fora, or domains**

TikTok does not have pages, fora, or domains. TikTok has invite-only private group chats. Therefore, this TTP is not materially relevant and - in any case - is generally addressed through broader integrity measures such as our Integrity and Authenticity policies.

Account hijacking or Impersonation

As described above, our Integrity and Authenticity policies prohibit **impersonation**, which refers to accounts that pose as another real person or entity or present as a person or entity that does not exist (a fake persona) with a demonstrated intent to mislead others on the platform. Our users are not allowed to use someone else's name, biographical details, or profile picture in a misleading manner.

In order to protect freedom of expression, we do allow accounts that are clearly parody, commentary, or fan-based, such as where the account name indicates that it is a fan, commentary, or parody account and not affiliated with the subject of the account. We continue to develop our policies to ensure that impersonation of entities (such as businesses or educational institutions, for example) is prohibited and that accounts which impersonate people or entities who are not on the platform are also prohibited.

Our Privacy and Security policies under our Community Guidelines address account hijacking. We expressly prohibit users from providing access to their account credentials to others or enabling others to conduct activities against our Community Guidelines. We do not allow access to any part of TikTok through unauthorised methods; attempts to obtain sensitive, confidential, commercial, or personal information; or any abuse of the security, integrity, or reliability of our platform. We also provide practical [guidance](#) to users if they have concerns that their account may have been hacked.

Our TikTok Shop Content Policy also prohibits impersonating other individuals or organisations to deceive or mislead users or make false representations. In addition, under TikTok Shop's Artificial Intelligence Generated Content (AIGC) policy, creators must not use AI-generated content to impersonate celebrities, medical professionals, political figures, public figures or official authorities for product promotion.



TTPs which pertain to the dissemination of content created in the context of a disinformation campaign, which may or may not include some forms of targeting or attempting to silence opposing views:

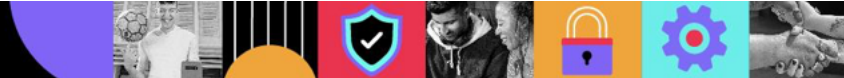
Deliberately targeting vulnerable recipients (e.g. location spoofing or obfuscation), inauthentic coordination of content creation or amplification, including attempts to deceive/manipulate platforms algorithms (e.g. keyword stuffing or inauthentic posting/reposting designed to mislead people about popularity of content, including by influencers), use of deceptive practices to deceive/manipulate platform algorithms, and coordinated mass reporting of non-violative opposing content or accounts.

We address covert influence operations (CIOs) through our Covert Influence Operations (CIO) policy, which prohibits attempts to sway public opinion while also misleading our systems or users about the identity, origin, approximate location, popularity or overall purpose.

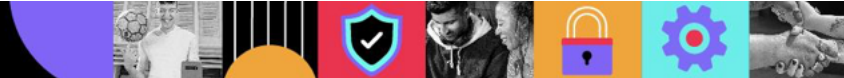
We focus on behavioural signals and linkages between accounts and techniques to determine if actors are engaging in a coordinated effort to mislead our systems or community. We take continuous action against these attempts, including banning accounts found to be linked to previously disrupted networks. Our approach is supported by ongoing research into complex deceptive behaviours and iterative updates to our product and policy solutions to address emerging disinformation risks. We also publish regular transparency reports detailing the CIO networks we detect and remove, which are available in our Trust & Safety Centre [here](#). For advertising-related CIO measures, please refer to Chapter 2.

TikTok Shop's Creator Fraud, Abuse and Misconduct Policy prohibits:

- Fake interactions, being activities that involve directly manipulating platform metrics or algorithms to artificially boost performance.
- Traffic manipulation, which involves generating traffic, exposure, or engagement through improper, deceptive, or non-genuine means.
- Inauthentic engagement behaviours, which involve deceptive or manipulative tactics intended to influence buyer behavior or artificially drive engagement.



	TikTok Shop's Content Policy also prohibits certain forms of artificial amplification and coordinated distribution of e-commerce content. This includes creating multiple accounts to distribute the same e-commerce content, coordinated mass posting across multiple accounts, and the use of malicious software or modified code to artificially increase views, likes, followers, shares or comments. Use “hack and leak” operation (which may or may not include doctored content) We have a number of policies that address hack-and-leak related threats, including: • Our hack and leak policy aims to further reduce the harms inflicted by the unauthorised disclosure of hacked materials on the individuals, communities, and organisations that may be implicated or exposed by such disclosures. • Our CIO policy addresses use of leaked documents to sway public opinion as part of a wider operation. • Our Edited Media and AI-Generated Content (AIGC) policy captures materials that have been digitally altered without an appropriate disclosure. • Our misinformation policies prohibit: ○ Misinformation that poses a risk to public safety or incites panic, including falsely presenting past crisis events as recent or claiming that critical resources are unavailable during emergencies. ○ Health misinformation that could cause significant harm, such as promoting unproven treatments that may be fatal, discouraging professional care for life-threatening conditions (e.g., vaccine effectiveness), or spreading false information about how such conditions are transmitted. ○ Misinformation that denies the existence of climate change, misrepresents its causes, or contradicts its established environmental impact.
--	--



	○ Conspiracy theories or hoaxes that could cause significant harm, such as those that make a violent call to action or have links to previous violence. Deceptive manipulated media (e.g. “deep fakes”, “cheap fakes”...) Our ‘Edited Media and AI-Generated Content (AIGC)’ policy outlines our existing prohibitions on AIGC. In accordance with our policy, we prohibit AIGC which features: ● Using the likeness of private figures without consent ● Sexualized, fetishized, or victimizing depictions ● AI-created likenesses made to bully or harass ● Accounts focused on AI images of youth in clothing suited for adults, or sexualized poses or facial expressions ● AIGC or significantly edited content that misleads about a matter of public importance, such as: ○ Content made to look like it comes from a real news source ○ A crisis event, like a natural disaster or conflict ○ A public figure being degraded, harassed, or linked to criminal behavior ○ A public figure taking political stances, supporting products, or commenting on public issues they haven't actually addressed ○ A political endorsement or condemnation that never happened ● Any content that breaks our Community Guidelines, including those on impersonation, misinformation, and hate speech, even if it's AI-generated TikTok Shop's AIGC policy also prohibits the use of AI-generated content that impersonates celebrities, medical professionals, political figures, public figures or official authorities for product promotion, including AI-generated deepfakes and other manipulated media.
--	--



Non-transparent compensated messages or promotions by influencers

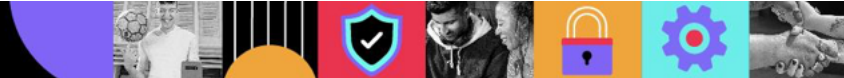
Under our [Terms of Service](#) and Commercial Disclosure and Paid Marketing section of our Community Guidelines, users posting about a **brand or product in return for any payment or other incentive** must disclose their content by enabling the Commercial Disclosure Toggle, which we make available for users. To support enforcement, we provide a reporting functionality for suspected undisclosed branded content, which triggers user prompts and encourages corrective action.

For TikTok Shop, creator commissions are paid through TikTok itself, meaning TikTok is aware whenever content contains a shoppable product link and the creator is remunerated. As a result, such content does not rely on user self-disclosure: the Commercial Disclosure Toggle is automatically enabled when a product link is added.

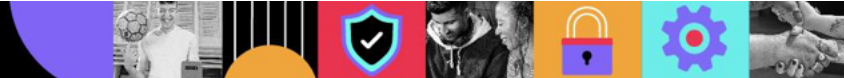
Our Political Advertising policy prohibits political advertising, including political branded content.

With regard to Advertising Policies, please refer to Chapter 2 concerning Misinformation Advertising Policies.

In addition, our CIO policy also applies to non-transparent compensated messages or promotions by influencers where it is found that those messages or promotions formed part of a covert influence campaign.



QRE 14.1.2	We proactively detect TTPs and related practices through a combination of automated systems, behavioural analytics and specialised investigative teams. We have created and used detection models and rule engines that: • Prevent inauthentic accounts from being created based on malicious patterns; and • Remove registered accounts based on certain signals (i.e., uncommon behaviour on the platform). We also manually monitor user reports of inauthentic accounts in order to detect larger clusters or similar inauthentic behaviours. Given the complex nature of the TTPs, human moderation is critical to success in assessing and addressing identified violations. We provide our moderation teams with detailed guidance on how to apply the Integrity and Authenticity policies in our Community Guidelines, allowing them to route new or evolving content to our fact-checking partners for assessment. In addition, where content reaches certain popularity levels in terms of the number of video views, it will be flagged for further review given the increase in potential harm if the content is found to be in breach of our Community Guidelines, including our Integrity and Authenticity policies. Furthermore, during this reporting period, we improved automated detection and enforcement of our 'Edited Media and AI-Generated Content (AIGC)' policy, effectively increasing the number of videos removed for policy violations. This also decreased the number of views per video over the reporting period, demonstrating an effective control strategy as the scope of enforcement increased. We've built international trust and safety teams with specialised expertise across threat intelligence, security, law enforcement, and data science to work on influence operations. These teams are trained to continuously pursue and analyse on-platform technical signals as well as leads from external sources to detect and investigate potential CIOs. They also collaborate with external intelligence vendors to support specific investigations on a case-by-case basis.
--------------------------	--



Accounts that engage in influence operations often avoid posting content that would be violative of platforms' guidelines by itself. That's why we focus on accounts' behaviour and technical linkages, specifically looking for evidence that:

- They are coordinating with each other. For example, they are operated by the same entity, share technical similarities like using the same devices, or work together to spread the same narrative.
- They are misleading our systems or users. For example, they are trying to conceal their actual location or use fake personas to pose as someone they're not.
- They are attempting to manipulate or corrupt public debate to impact the decision-making, beliefs, and opinions of a community. For example, they are attempting to shape discourse around an election or conflict.

These criteria are aligned with industry standards and guidance from the experts we regularly consult with. They're particularly important to help us distinguish malicious, inauthentic coordination from authentic interactions that are part of healthy and open communities.

Measure 14.2

**QRE 14.2.1**

The implementation of our policies is ensured by different means, including specifically-designed tools (such as toggles to disclose branded content - see QRE 14.1.1) or human investigations to detect deceptive behaviours (for CIO activities - see QRE 14.1.2).

In further details:

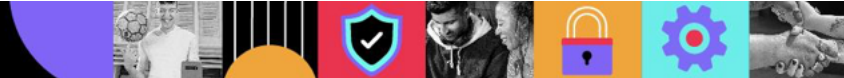
- **Creation of inauthentic accounts or botnets (which may include automated, partially automated, or non-automated accounts)**
 - If we determine someone has engaged in any deceptive account behaviours, we will ban the account, and may ban any new accounts that are created. We also issue warnings to users of suspected impersonation accounts and do not recommend those accounts on our For You Feed.
- **Use of fake / inauthentic reactions (e.g. likes, up votes, comments) and use of fake followers or subscribers**
 - If we become aware of accounts or content with inauthentically inflated metrics, we will remove the associated fake followers or likes. Content that tricks or manipulates others as a way to increase engagement metrics, such as “like-for-like” promises and false incentives for engaging with content (to increase gifts, followers, likes, views, or other engagement metrics) is ineligible for our For You feed.
- **Account hijacking or Impersonation**
 - If we determine someone has engaged in any account hijacking or impersonation, we will ban the account, and may ban any new accounts that are created. We also issue warnings to users of suspected impersonation accounts and do not recommend those accounts on our For You Feed.



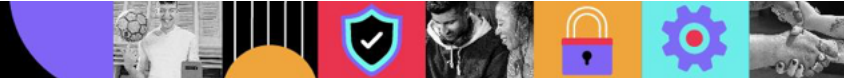
- **Deliberately targeting vulnerable recipients (e.g. location spoofing or obfuscation), inauthentic coordination of content creation or amplification, including attempts to deceive/manipulate platforms algorithms (e.g. keyword stuffing or inauthentic posting/reposting designed to mislead people about popularity of content, including by influencers), use of deceptive practices to deceive/manipulate platform algorithms, and coordinated mass reporting of non-violative opposing content or accounts.**
 - Targeting vulnerable recipients via deceptive practices in order to manipulate the information environment is covered by our policy on CIO. We use in-depth analysis to assess for technical signals such as location obfuscation and inauthentic coordination in these investigations. Where our teams have the necessary high degree of confidence that an account is engaged in CIO or is connected to networks we took down in the past as part of a CIO, it is removed from our Platform.
- **Engaging in a “hack and leak” operation (which may or may not include doctored content)**
 - If we determine with high confidence that someone has posted private, sensitive, or confidential information that was obtained without consent, we will remove the content from the platform. We will further ban accounts that posted the material if it is part of a wider covert influence operation. However, we allow limited discussion or distribution of such materials if there's a clear public interest and the content follows journalistic best practices.
- **Deceptive manipulated media (e.g. “deep fakes”, “cheap fakes”...)**
 - We remove content that was generated or significantly edited by AI or AI-assisted tools to falsely or misleadingly depict authoritative information, critical events, elections, and public figures. If we determine that the account has repeatedly posted such content, we will ban the account as well.



	○ If we determine that e-commerce content does not comply with TikTok Shop's Artificial Intelligence Generated Content (AIGC) policy, we take enforcement action, including the removal of product anchors from content and other account-level measures depending on the nature and severity of the violation. ● Non-transparent compensated messages or promotions by influencers ○ For undisclosed branded content, we remind the user who posted the suspected undisclosed branded content of our requirements and prompt them to turn the Commercial Disclosure Toggle on if required. ○ Where TikTok suspects content is political branded content and TikTok has high confidence that an individual was paid to post political content, TikTok removes the content. Where TikTok has only medium confidence, the content is restricted from appearing in the For You Feed. ○ For TikTok Shop, creator commissions are paid through TikTok itself, meaning TikTok is aware whenever content contains a shoppable product link and the creator is remunerated. As a result, such content does not rely on user self-disclosure: the Commercial Disclosure Toggle is automatically enabled when a product link is added. The implementation of our policies is also ensured through enforcement measures applied in all EEA countries. Similarly, where our teams have a high degree of confidence that specific content violates one of our TTPs-related policies (See QRE 14.1.1), such content is removed from TikTok.



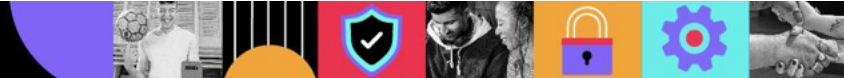
SLI 14.2.1 – SLI 14.2.4								
TTP OR ACTION1	TTP No. 1: Creation of inauthentic accounts or botnets (which may include automated, partially automated, or non-automated accounts) Methodology of data measurement We have based the number on: (i) fake accounts removed; and (ii) followers of the fake accounts (identified at the time of removal of the fake account), in the country the fake account was last active in. For SLI 14.2.4, we provide a ratio of the monthly average of fake accounts removed over monthly active users in the EU, based on the latest publication of monthly active users by TikTok under the DSA, in order to better reflect TTPs related content in relation to overall content on the service.							
	SLI 14.2.1	SLI 14.2.2	SLI 14.2.3			SLI 14.2.4		
	Number of actions taken by type	Interaction/ engagement before action				TTPs related content in relation to overall content on the service		
List actions per member states (see example table above)	Number of fake accounts removed	Number of followers of fake accounts identified at the time of removal				Ratio of monthly average of Fake accounts over monthly active users		



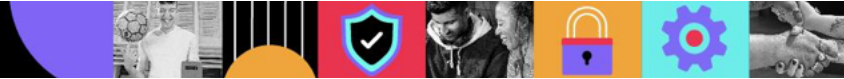
Member States								
Austria	970509	121661						
Belgium	1182975	175652						
Bulgaria	689877	65405						
Croatia	694966	47332						
Cyprus	356703	45024						
Czechia	647844	125046						
Denmark	1009026	65251						
Estonia	761476	41967						
Finland	1251547	199189						
France	4754983	2723599						
Germany	6651150	3087831						
Greece	1269946	207504						
Hungary	239605	63867						
Ireland	443832	99793						
Italy	1352923	571400						
Latvia	201734	40532						
Lithuania	163058	63406						
Luxembourg	410838	57750						



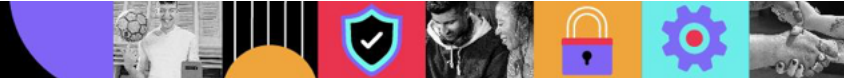
Malta	269526	14785						
Netherlands	1311942	815914						
Poland	1956516	575722						
Portugal	350462	386145						
Romania	960855	180313						
Slovakia	465397	49613						
Slovenia	376600	66800						
Spain	3413344	811709						
Sweden	2063601	239303						
Iceland	49563	21062						
Liechtenstein	30540	4757						
Norway	2311190	64496						
Total EU	34221235	10942513					3.01%	
Total EEA	36612528	11032828						



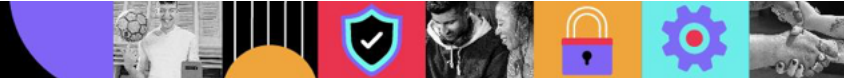
TTP OR ACTION 2	TTP no. 2: Use of fake / inauthentic reactions (e.g. likes, up votes, comments)			
	Methodology of data measurement: We based the number of fake likes that we removed on the country of registration of the user. We also based the number of fake likes prevented on the country of registration of the user. <i>H1 2026 figures for this metric reflect certain data limitations experienced earlier in the reporting period, impacting data quality and completeness. As such, the totals presented may not fully capture the overall scope of activity under this TTP.</i>			
	SLI 14.2.1	SLI 14.2.2	SLI 14.2.3	SLI 14.2.4
	Number of actions taken by type	Interaction/ engagement before action		
List actions per member states (see example table above)	Number of fake likes removed	Number of fake likes prevented		
Austria	17057524	8663536		
Belgium	33605557	15915836		
Bulgaria	8793507	17878056		
Croatia	6315458	4672220		



Cyprus	21419126	1947853		
Czechia	6906117	7512112		
Denmark	7251754	4643146		
Estonia	2478309	2201824		
Finland	15179538	5786710		
France	203055875	125365075		
Germany	222641814	276744482		
Greece	13148363	10549929		
Hungary	7627999	8971004		
Ireland	25037665	7587545		
Italy	84279565	69958143		
Latvia	7957686	3215055		

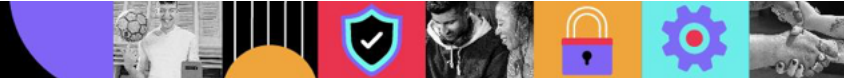


Lithuania	4287887	4004511		
Luxembourg	1348163	1821560		
Malta	5716023	1107908		
Netherlands	103826313	44523502		
Poland	28439166	43801134		
Portugal	11714277	13147308		
Romania	34964776	29599787		
Slovakia	2819307	5504841		
Slovenia	1587706	3738400		
Spain	72291234	59864715		
Sweden	27053458	51504755		
Iceland	691841	413196		

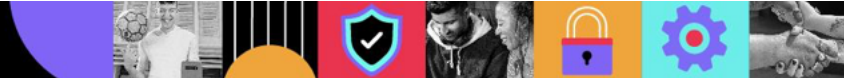


Liechtenstein	2201447	33424		
Norway	18254336	4840856		
Total EU	976804167	830230947		
Total EEA	997951791	835518423		

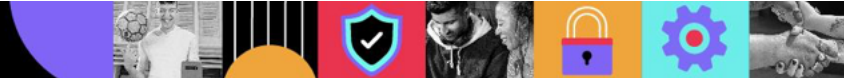
TTP OR ACTION 3	TTP No. 3: Use of fake followers or subscribers Methodology of data measurement: We based the number of fake followers that we removed on the country of registration of the user. We also based the number of fake followers prevented on the country of registration of the user. <i>H1 2026 figures for this metric reflect certain data limitations experienced earlier in the reporting period, impacting data quality and completeness. As such, the totals presented may not fully capture the overall scope of activity under this TTP.</i>			
	SLI 14.2.1	SLI 14.2.2	SLI 14.2.3	SLI 14.2.4
	Number of actions taken by type	Interaction/ engagement before action		



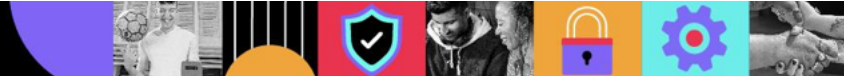
List actions per member states (see example table above)	Number of fake followers removed	Number of fake follows prevented		
Member States				
Austria	11390859	3749861		
Belgium	22137682	6413789		
Bulgaria	3813327	6509936		
Croatia	3128259	1394627		
Cyprus	3746433	643893		
Czechia	8180464	2416099		
Denmark	5817279	3017345		
Estonia	1470272	967121		



Finland	4549158	1951041		
France	169273428	98736923		
Germany	236112589	108719977		
Greece	7150537	7479380		
Hungary	5744976	4567699		
Ireland	7076834	18470877		
Italy	53195455	87949421		
Latvia	2274410	1437015		
Lithuania	5452681	1238838		
Luxembourg	1682553	1133166		
Malta	885786	215705		
Netherlands	40994677	25049518		

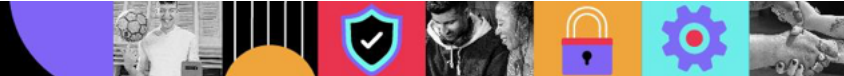


Poland	22161275	16669449		
Portugal	6973708	9044549		
Romania	20094090	26393859		
Slovakia	1976213	1644686		
Slovenia	1367979	6146667		
Spain	52866819	34466332		
Sweden	15038610	6068181		
Iceland	527905	456623		
Liechtenstein	61053	7324		
Norway	6628315	2507450		
Total EU	714556353	482495954		
Total EEA	721773626	485467351		

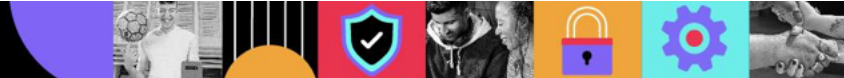


TTP OR ACTION 4	TTP No. 4: Creation of inauthentic pages, groups, chat groups, fora, or domains TikTok does not have pages, fora, or domains. TikTok has invite-only private group chats. Therefore, this TTP is not materially relevant and - in any case - is generally addressed through broader integrity measures such as our Integrity and Authenticity policies.
------------------------	--

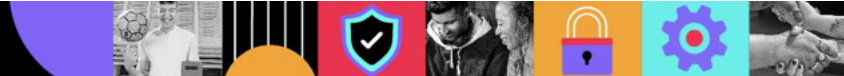
TTP OR ACTION 5	TTP No. 5: Account hijacking or impersonation Methodology of data measurement: The number of accounts removed under our impersonation policy is based on the approximate location of the users. We have updated our methodology to report the ratio of monthly average impersonation accounts banned over monthly active users, based on the latest publication of monthly active users, in order to better reflect TTPs related content in relation to overall content on the service.			
	SLI 14.2.1	SLI 14.2.2	SLI 14.2.3	SLI 14.2.4
	Number of actions taken by type			TTPs related content in relation to overall content on the service
Member States	Number of account banned under impersonation policy			Impersonation accounts over monthly active users
List actions per member states (see example table above)				
Austria	850			
Belgium	1441			



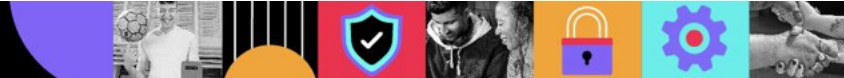
Bulgaria	827			
Croatia	316			
Cyprus	211			
Czechia	754			
Denmark	628			
Estonia	178			
Finland	692			
France	10362			
Germany	10732			
Greece	807			
Hungary	629			
Ireland	1088			
Italy	5227			
Latvia	273			
Lithuania	378			
Luxembourg	112			
Malta	0			
Netherlands	4678			
Poland	5140			



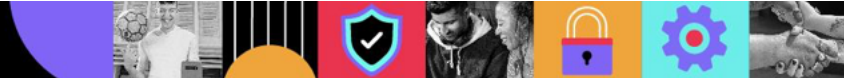
Portugal	1095			
Romania	3046			
Slovakia	471			
Slovenia	175			
Spain	5382			
Sweden	1554			
Iceland	68			
Liechtenstein	0			
Norway	901			
Total EU	57046			0.005%
Total EEA	58015			



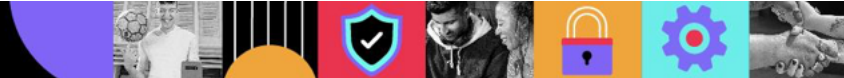
TTP OR ACTION 6	TTP No. 6. Deliberately targeting vulnerable recipients (e.g. via personalised advertising, location spoofing or obfuscation)					
	Methodology of data measurement: The number of new CIO network discoveries found to be targeting European audiences relates to our public disclosures for the period January 1 to June 30 2026. We have categorised disrupted CIO networks by the country we assess that the network targeted. We have included any network which we assess to have targeted one or more European markets, or have operated from an EU market. We publish details of the CIO networks we identify and remove within our transparency reports here. CIO networks identified and removed are detailed below, including the assessed geographic location of network operation and the assessed target audience of the network, which we assess via technical and behavioural evidence from proprietary and open sources. The number of followers of CIO networks has been based on the number of accounts that followed any account within a network as of the date of that network's removal.					
	SLI 14.2.1	SLI 14.2.2			SLI 14.2.3	SLI 14.2.4
Assessed Network Operating Location	Number of instances of identified TTPs and actions taken by type	Interaction/ engagement before action	Views/ impressions after action	Interaction/ engagement after action	Trends on targeted audiences	
January 1- June 30 2026						



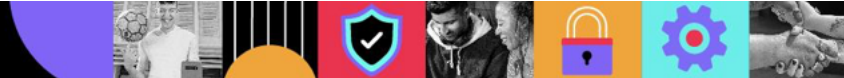
Hungary	107 Accounts	37,070 Followers	N/A	N/A	We assess that this network operated from Hungary and targeted a Hungarian audience. The individuals behind this network created inauthentic accounts in order to artificially amplify narratives critical of the Tisza political party. The network was found to post AI-generated content to reinforce its messaging.	
Ukraine	45 Accounts	47,114 Followers	N/A	N/A	We assess that this network operated from Ukraine and targeted European audiences. The individuals behind this network created inauthentic accounts in order to undermine certain European political figures such as Hungarian Prime Minister Viktor Orbán. The network was found to create accounts which it presented as news accounts.	



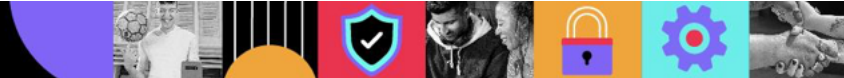
Hungary	91 Accounts	246 Followers	N/A	N/A	We assess that this network operated from Hungary and targeted a Hungarian audience. The individuals behind this network created inauthentic accounts in order to artificially amplify narratives critical of the Tisza political party. The network was found to create fictitious personas using stock or unoriginal imagery as profile pictures.	
Hungary	62 Accounts	1,452 Followers	N/A	N/A	We assess that this network operated from Hungary and targeted a Hungarian audience. The individuals behind this network created inauthentic accounts in order to artificially amplify narratives critical of the Fidesz political party. The network was found to create fictitious personas using profile pictures which we assess to be AI-generated.	



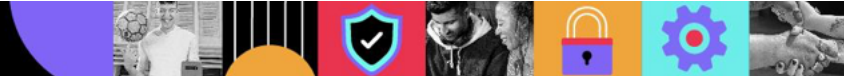
Bulgaria	34 Accounts	66,763 Followers	N/A	N/A	We assess that this network operated from Bulgaria and targeted a Bulgarian audience. The individuals behind this network created inauthentic accounts in order to artificially amplify narratives in support of the DPS-NN political party, within the context of the April 2026 Bulgarian parliamentary elections. The network was found to coordinate across multiple online platforms.	
Malta	9 Accounts	433 Followers	N/A	N/A	We assess that this network operated from Malta and targeted a Maltese audience. The individuals behind this network created inauthentic accounts in order to artificially amplify narratives favorable to the Labor Party within the context of the 2026 Maltese general elections. The network was found to create accounts that it presented as news accounts.	



China	41 Accounts	195,566 Followers	N/A	N/A	We assess that this network operated from China and targeted a global audience. The individuals behind this network created inauthentic accounts in order to artificially undermine public perception of US foreign policy decisions. The network was found to create fictitious personas that catered to each market using English and Chinese language videos.	
Germany	51 Accounts	46,388 Followers	N/A	N/A	We assess that this network operated from Germany and targeted a German audience. The individuals behind this network created inauthentic accounts in order to amplify narratives supporting the Alternative for Germany (AfD) political party. The network was found to create duplicative profiles in order to artificially amplify its narratives.	

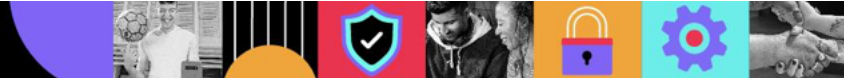


France	47 Accounts	7,281 Followers	N/A	N/A	We assess that this network operated from France and targeted audiences in Armenia and Azerbaijan. The individuals behind this network created inauthentic accounts in order to amplify anti-Azerbaijan narratives. The network was found to post in multiple languages, including Russian, Armenian and Georgian.	
Pakistan	26 Accounts	31,880 Followers	N/A	N/A	We assess that this network operated from Pakistan and targeted the Pakistani and Indian diasporas globally. The individuals behind this network created inauthentic accounts in order to artificially amplify narratives favorable to Pakistan, while criticizing India and the United Kingdom. We assess that the network used hashtags not commonly used within Pakistan in order to target a more global audience.	

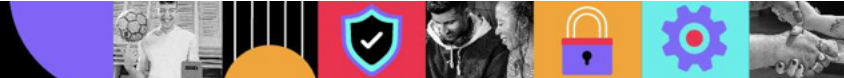


Kazakhstan	6 Accounts	492,739 Followers	N/A	N/A	We assess that this network operated from Kazakhstan and targeted a global English-speaking audience. The individuals behind this network created inauthentic accounts in order to amplify narratives critical of U.S. foreign policy. The network was found to repurpose accounts with existing high engagement in order to post content in support of its strategic aims.	
------------	------------	-------------------	-----	-----	---	--

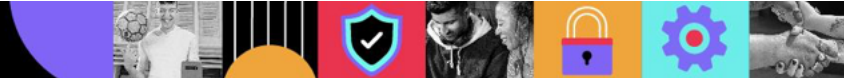
TTP OR ACTION 7	TTP No. 7: Deploy deceptive manipulated media (e.g. “deep fakes”, “cheap fakes”...) We have based the following numbers on the country in which the video was posted: videos removed because of violations of the Edited Media and AI-Generated Content (AIGC) policy. The number of views of videos removed because of violation of each of these policies is based on the approximate location of the user.				
Member States	Number of videos removed because of violation of Edited Media and AI-Generated Content (AIGC) policy	Number of views of videos removed because of Edited Media and AI-Generated Content (AIGC) policy	Number of unique videos labelled with AIGC tag of "Creator labeled as AI-generated"	Number of unique videos labelled with AIGC tag of "AI-generated"	
Austria	11719	3734469	242193	1352933	
Belgium	10754	3425892	270806	2128005	



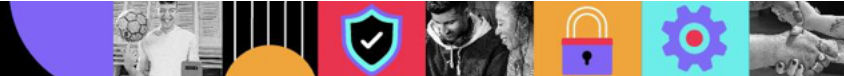
Bulgaria	42241	1708891	111688	3068025
Croatia	16366	722675	41838	448236
Cyprus	4826	912258	105076	395260
Czechia	7177	10261960	92420	1463723
Denmark	15176	8686782	101929	664029
Estonia	4085	2056073	31497	223889
Finland	5372	4735509	120414	645357
France	63056	94228457	2897475	13544507
Germany	99159	128170986	3698070	15869550
Greece	5441	867743	301057	2293565



Hungary	14592	3097025	225171	2683504
Ireland	6948	1363165	71710	575280
Italy	37133	13472997	1833428	12282164
Latvia	4805	5250952	63249	546291
Lithuania	4861	3304486	112851	1717055
Luxembourg	2078	23604	20863	130628
Malta	910	1778433	17863	124796
Netherlands	85596	33693554	640312	2910540
Poland	42850	30661634	455100	5781369
Portugal	11637	2861651	248864	2708704

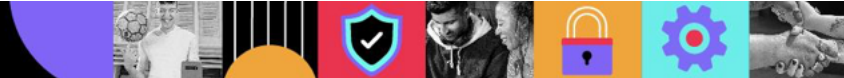


Romania	27051	6105146	533887	9301075
Slovakia	4077	146257	38502	1095414
Slovenia	1570	273450	24673	208580
Spain	45863	35355202	1960427	12916276
Sweden	10924	26469283	291759	1386206
Iceland	568	339717	4817	39335
Liechtenstein	63	0	436	2001
Norway	27016	1668462	118845	612409
Total EU	586267	423368534	14553122	96464961
Total EEA	613914	425376713	14677220	97,118,706



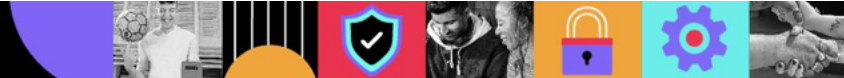
TTP OR ACTION 8 Member	TTP No. 8: Use “hack and leak” operation (which may or may not include doctored content) We have provided data on the CIO networks that we have disrupted in the reporting period under TTP No. 6. We have also provided data on violations of our Edited Media and AI-Generated Content (AIGC) policy under TTP No. 7. Our hack and leak policy was launched in H1 2024, but we do not have meaningful metrics under this policy to report for H1 2026.
TTP OR ACTION 9	TTP No. 9: Inauthentic coordination of content creation or amplification, including attempts to deceive/manipulate platforms algorithms (e.g. keyword stuffing or inauthentic posting/reposting designed to mislead people about popularity of content, including by influencers) In H2 2025, we launched the Evasive Techniques policy, which combats methods designed to evade moderation systems. We have provided data on the CIO networks that we have disrupted in the reporting period under TTP No. 6.
TTP OR ACTION 10	TTP No. 10: Use of deceptive practices to deceive/manipulate platform algorithms, such as to create, amplify or hijack hashtags, data voids, filter bubbles, or echo chambers We have provided data on the CIO networks that we have disrupted in the reporting period under TTP No. 6.

TTP OR ACTION 11	TTP No. 11. Non-transparent compensated messages or promotions by influencers			
	Methodology of data measurement: We are unable to provide this metric due to insufficient data available for the reporting period.			
	SLI 14.2.1	SLI 14.2.2	SLI 14.2.3	SLI 14.2.4
	Number of actions taken by type			
Member States				
List actions per member states (see example table above)				

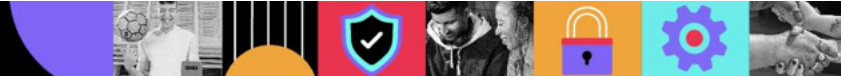


TTP OR ACTION 12	TTP No. 12: Coordinated mass reporting of non-violative opposing content or accounts We have provided data on the CIO networks that we have disrupted in the reporting period under TTP No. 6.			
	SLI 14.2.1	SLI 14.2.2	SLI 14.2.3	SLI 14.2.4
Member States				
List actions per member states (see example table above)				

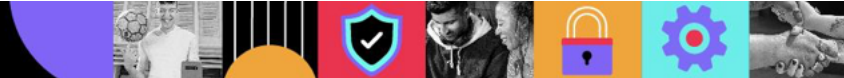
Measure 14.3	
QRE 14.3.1	We collaborated as part of the Integrity of Services working group to set up the first list of TTPs. We continue to provide updates on observed TTPs through our regular CIO transparency reporting , including observations on novel and emerging tradecraft.



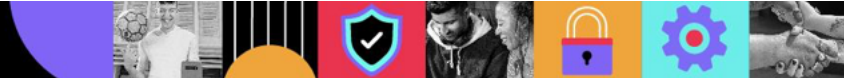
IV. Integrity of Services	
Commitment 15	
Relevant Signatories that develop or operate AI systems and that disseminate AI-generated and manipulated content through their services (e.g. deep fakes) commit to take into consideration the transparency obligations and the list of manipulative practices prohibited under the proposal for Artificial Intelligence Act.	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	Yes
If yes, list these implementation measures here [short bullet points].	• We provided extensive training for moderators and risk containment agents to help them better detect and remove deceptive AIGC more quickly. We also conducted a thorough assessment of the effectiveness of our AI policies and provided guidance to reduce systemic error. • We published our Responsible AI Principles.
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 15.1	



QRE 15.1.1	Our Edited Media and AI-Generated Content (AIGC) policy outlines our existing prohibitions on AIGC showing fake authoritative sources or crisis events, or falsely showing public figures in specific contexts. As AI evolves, we continue to invest in combating harmful AIGC by evolving our proactive detection models, consulting with experts, and partnering with peers on shared solutions. In line with our policy, users must proactively disclose when their content is AI-generated or manipulated but shows realistic scenes (i.e. fake people, places, or events that look like they are real). Our AI toggle allows users to self-disclose AI-generated content when posting. When this has been turned on, a tag “Creator labelled as AI-generated” is displayed to users. Alternatively, this can be done through the use of a sticker or caption, such as ‘synthetic’, ‘fake’, ‘not real’, or ‘altered’. We also automatically label content made with TikTok effects if they use AI. TikTok may automatically apply the “AI-generated” label to content we identify as completely generated or significantly edited with AI. This may happen when a creator uses TikTok AI effects or uploads AI-generated content that has Content Credentials attached, a technology from the Coalition for Content Provenance and Authenticity (C2PA). Content Credentials attach metadata to content that we can use to recognize and label AIGC instantly. Once content is labeled as AI-generated with an auto-label, users are unable to remove the label from the post. Under our AIGC policy, we do not allow: • Using the likeness of private figures without consent • Sexualized, fetishized, or victimizing depictions • AI-created likenesses made to bully or harass • Accounts focused on AI images of youth in clothing suited for adults, or sexualized poses or facial expressions • AIGC or significantly edited content that misleads about a matter of public importance, such as: ◦ Content made to look like it comes from a real news source ◦ A crisis event, like a natural disaster or conflict ◦ A public figure being degraded, harassed, or linked to criminal behavior ◦ A public figure taking political stances, supporting products, or commenting on public issues they haven't actually addressed ◦ A political endorsement or condemnation that never happened • Any content that breaks our Community Guidelines, including those on impersonation, misinformation, and hate speech, even if it's AI-generated
--------------------------	---



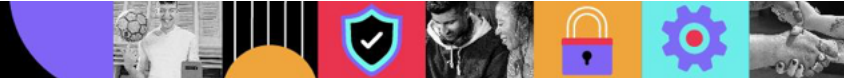
	In H1 2026 we updated our Enforcement Guidance for the Edited Media and AI-Generated Content (AIGC) policy to account for the latest AIGC trends on our platform. TikTok Shop operates a separate dedicated AIGC policy applicable to e-commerce videos and livestreams. Under this policy, creators are required to disclose AI-generated content, including through TikTok's AI-generated content disclosure tools where applicable. Under the TikTok Shop AIGC policy, the following are prohibited: • AI-generated content that impersonates or misrepresents celebrities, medical professionals, political figures, public figures, or official authorities for product promotion; • AI-generated content that exaggerates, fabricates or misrepresents product functions, efficacy, attributes or performance; • AI-generated content that uses fear-based health imagery or otherwise presents misleading health-related outcomes; and • Certain forms of high-volume, repetitive AI-generated promotional content. In addition, in relation to product listings, sellers are required to ensure that the product image depicts the actual product being offered for sale.
Measure 15.2	
QRE 15.2.1	We have a number of measures to ensure the algorithms we develop for detection moderation and sanctioning uphold the principles of fairness and comply with applicable laws. To that end: • We have in place internal guidelines on Algorithmic Fairness that are developed with adherence to our commitment to human rights as outlined here: https://www.tiktok.com/transparency/en/upholding-human-rights • We have continued to take a risk-based approach to monitoring our AI content moderation systems for fairness, and conducting technical fairness assessments where they meet certain risk-based thresholds.



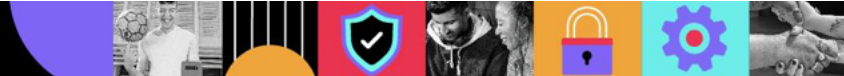
IV. Integrity of Services	
Commitment 16	
Relevant Signatories commit to operate channels of exchange between their relevant teams in order to proactively share information about cross-platform influence operations, foreign interference in information space and relevant incidents that emerge on their respective services, with the aim of preventing dissemination and resurgence on other services, in full compliance with privacy legislation and with due consideration for security and human rights risks.	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 16.1	



QRE 16.1.1	In addition to continuously enhancing our in-house capabilities, we proactively engage in comprehensive reviews of our peers' publicly disclosed findings in relation to CIO and swiftly implement necessary actions in alignment with our policies. To provide more regular and detailed updates about the CIO we disrupt to others, we have a dedicated report on covert influence operations, available in TikTok's Trust & Safety Centre. The insights and metrics in this report aim to inform industry peers and the research community. We also review relevant insights and metrics from other industry peers to cross-compare for any similar behaviour on TikTok. We continue to engage in the subgroups set up for insights sharing between signatories and the Commission, including: • Cross-industry forums such as EU elections roundtables in markets including Denmark and Slovenia; As we have detailed in other chapters to this report, we have robust advertising policies in place and have established joint operating procedures between specialist CIO investigations teams and Business Integrity teams to work on joint investigations of CIOs involving ads.		
SLI 16.1.1 Numbers of actions as a result of information sharing	N/A		
Data			
Measure 16.2			
QRE 16.2.1	We publish details of the CIO networks we identify and remove within our transparency reports here. As new deceptive behaviours emerge, we'll continue to evolve our response, strengthen enforcement capabilities, and publish our findings.		



V. Empowering Users Commitments 17 - 25



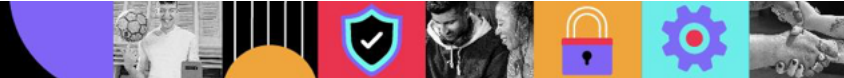
V. Empowering Users

Commitment 17

In light of the European Commission's initiatives in the area of media literacy, including the new Digital Education Action Plan, Relevant Signatories commit to continue and strengthen their efforts in the area of media literacy and critical thinking, also with the aim to include vulnerable groups.

In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]

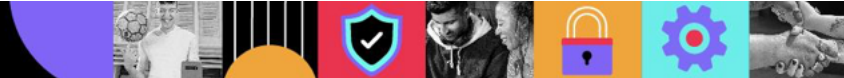
Yes



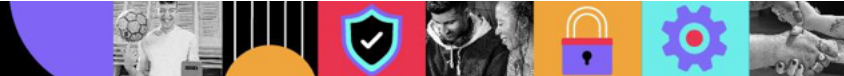
If yes, list these implementation measures here [short bullet points].

We have 14 ongoing media literacy and critical thinking skills campaigns in Europe (12 in EU/EEA—Denmark, Finland, France, Germany, Ireland, Italy, Romania, Spain, Sweden, Netherlands, Poland and Portugal; 2 in wider European countries—Georgia and Moldova).

- We ran 14 temporary media literacy election integrity campaigns in advance of elections, some in collaboration with our fact-checking and media literacy partners:
 - 14 in the EU
 - Portugal (Presidential Election) - Polígrafo
 - Aragon (Spain Regional Election)
 - Baden-Württemberg and Rheinland-Pfalz (Germany State Elections) - Deutsche Presse-Agentur (dpa)
 - Castile and Leon (Spain Regional Election)
 - France (Municipal Elections) - Agence France-Presse (AFP)
 - Netherlands (Municipal Elections)
 - Slovenia (Parliamentary Election)
 - Denmark (General Election)
 - Hungary (Parliamentary Election)
 - Bulgaria (Parliamentary Election)
 - Andalusia (Spain Regional Election) - Newtral
 - Cyprus (Parliamentary Election)
 - Malta (General Election)
 - New Caledonia (France Regional Election)
- Following the outbreak of Ebola in Sub-saharan Africa, we launched an in-app guide to provide users with guidance on verifying information using reliable sources. This guide was launched in Mayotte (France) where we linked to info.gouv.fr.
- Following the outbreak of Hantavirus, we launched an in-app guide and a video tag to provide users with guidance on verifying information using reliable sources. This guide was launched in all EU markets, where we linked to the World Health Organisation (WHO) and other government health advice websites (IT, DE, NL, FR).
- We continued to support mental well-being awareness and literacy and to combat misinformation with reliable content through the [WHO's Fides](#) network, a diverse community of trusted healthcare professionals and content creators in a number of countries, including France.

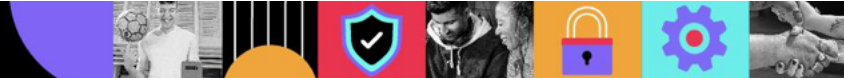


	We brought greater transparency about our systems and our integrity and authenticity efforts to our community by sharing regular insights and updates in our Global Elections Integrity Hub, including dedicated coverage of elections across Europe, the Middle East, and Africa. The Hub outlines our policies, product features, and moderation practices that help protect platform integrity during elections. Throughout this reporting period, we regularly updated the Hub with information on our safety efforts in markets with active elections, including Croatia, Germany, Netherlands, Portugal, Poland and Ireland.
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 17.1	

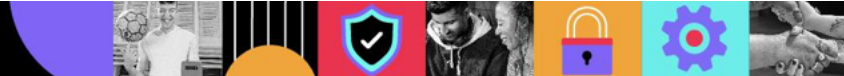


QRE 17.1.1	We continue to expand our in-app measures for front-end interventions that we outline below, which show users additional context on certain content (e.g., natural disasters and rapidly unfolding events) and redirect them to trusted information. We make these tools available in relevant EU languages, which includes 22 EU official languages (plus, for EEA users, Norwegian and Icelandic), as applicable. We work with external experts to combat harmful misinformation. For example, we work with our global fact-checking partners, taking into account their feedback to continually identify new topics and consider which tools may be best suited for raising awareness around that topic. We also deploy a combination of in-app user intervention tools on topical issues such as elections, Holocaust Education, and the War in Ukraine. These tools include: • Video notice tags: An information bar at the bottom of a video which is automatically applied to a specific word or hashtag (or set of hashtags). The information bar is clickable and invites users to “<i>Learn more about</i> [the topic]”. Users will be directed to an in-app guide, or reliable third party resource, as appropriate. • Search intervention: If users search for terms associated with a topic, they will be presented with a banner encouraging them to verify the facts and providing a link to a trusted source of information. Search interventions are not deployed for search terms that violate our Community Guidelines, which are actioned according to our policies. TikTok Shop also provides educational resources for sellers and creators through TikTok Shop Academy and the Policy Centre. These resources provide information on platform policies, compliance requirements and responsible content practices relevant to TikTok Shop. We have updated our methodology for this COCD report to include an additional Holocaust misinformation campaign in the data we provide below.
-------------------	--

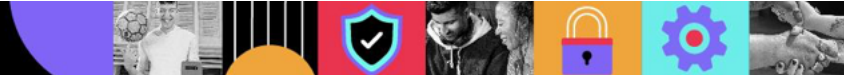
	Number of impressions of Video Notice Tag covered by Intervention (Holocaust Misinformation/Denial)	Number of clicks of Video Notice Tag covered by Intervention (Holocaust Misinformation/Denial)	Click Through Rate of Video Notice Tag covered by Intervention (Holocaust Misinformation/Denial)
Member States			
Austria	6945549	6479	0.2%
Belgium	4678337	6455	0.2%



Bulgaria	1752413	3521	0.4%
Croatia	2868651	4168	0.3%
Cyprus	629559	905	0.3%
Czechia	5765230	3434	0.2%
Denmark	3099745	2465	0.3%
Estonia	812547	827	0.2%
Finland	5084612	4947	0.3%
France	2651762	24841	0.2%
Germany	69318976	34666	0.2%
Greece	4752972	9943	0.3%
Hungary	5801809	5835	0.3%
Ireland	6082223	3330	0.2%
Italy	3479548	18142	0.3%
Latvia	1088268	1076	0.3%
Lithuania	1725812	2254	0.3%
Luxembourg	369517	435	0.2%
Malta	358985	277	0.2%
Netherlands	16224082	13614	0.2%
Poland	44471692	20533	0.2%

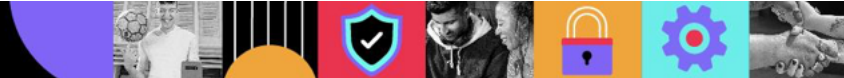


Portugal	3630727	3243	0.3%
Romania	7901518	9348	0.3%
Slovakia	3357329	2023	0.3%
Slovenia	1398765	1436	0.2%
Spain	13997758	50136	0.2%
Sweden	8885363	7200	0.3%
Iceland	457096	346	0.2%
Liechtenstein	23171	11	0.2%
Norway	4316912	3295	0.3%
Total EU	227,133,749	241533	0.2%
Total EEA	231930928	245185	0.2%
	Number of impressions of topic covered by video Intervention (Election)	Number of clicks by video Intervention (Election)	Click Through Rate by video Intervention (Election)
Member States			
Bulgaria	22,719,523	34,796	0.15%
Cyprus	130,980	278	0.21%
Denmark	15,825,969	22,143	0.14%
Hungary	142,693,830	169,863	0.12%
Malta	7,208,871	5,241	0.07%

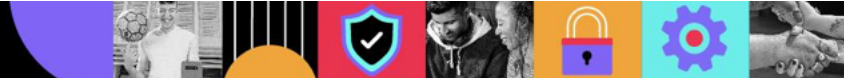


Portugal	1,467,705	2,096	0.14%
Slovenia	6,236,637	8,522	0.14%
Total EU	196,283,515	242,939	0.12%
Total EEA	196,283,515	242,939	0.12%

	Number of impressions of Search interventions (Holocaust Misinformation/Denial)	Number of clicks of Search interventions (Holocaust Misinformation/Denial)	Click Through Rate of Search interventions (Holocaust Misinformation/Denial)
Member States			
Austria	866216	6479	0.7%
Belgium	847731	6455	0.8%
Bulgaria	461877	3521	0.8%
Croatia	425219	4168	1.0%
Cyprus	137288	905	0.7%
Czechia	561558	3434	0.6%
Denmark	391749	2465	0.6%
Estonia	91227	827	0.9%
Finland	668526	4947	0.7%
France	4087849	24841	0.6%
Germany	5617788	34666	0.6%



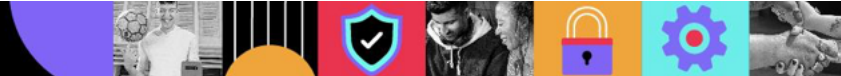
Greece	1652885	9943	0.6%
Hungary	518803	5835	1.1%
Ireland	605633	3330	0.5%
Italy	3481247	18142	0.5%
Latvia	122713	1076	0.9%
Lithuania	306055	2254	0.7%
Luxembourg	53953	435	0.8%
Malta	35479	277	0.8%
Netherlands	2475438	13614	0.5%
Poland	3460694	20533	0.6%
Portugal	779606	3243	0.4%
Romania	1073109	9348	0.9%
Slovakia	252424	2023	0.8%
Slovenia	130858	1436	1.1%
Spain	3367952	50136	1.5%
Sweden	1187509	7200	0.6%
Iceland	35346	346	1.0%
Liechtenstein	2521	11	0.4%
Norway	535322	3295	0.6%



Total EU	33661386	241533	0.7%
Total EEA	34234575	245185	0.7%

	Number of impressions of Search interventions (Election)	Number of clicks of Search interventions (Election)	Click Through Rate of Search interventions (Election)
Member States			
Bulgaria	210252	989	0.5%
Cyprus	36528	173	0.5%
Denmark	309385	2139	0.7%
Hungary	7469203	27530	0.4%
Malta	51903	489	0.9%
Portugal	1109063	3050	0.3%
Slovenia	219677	1358	0.6%

Measure 17.2	
--------------	--

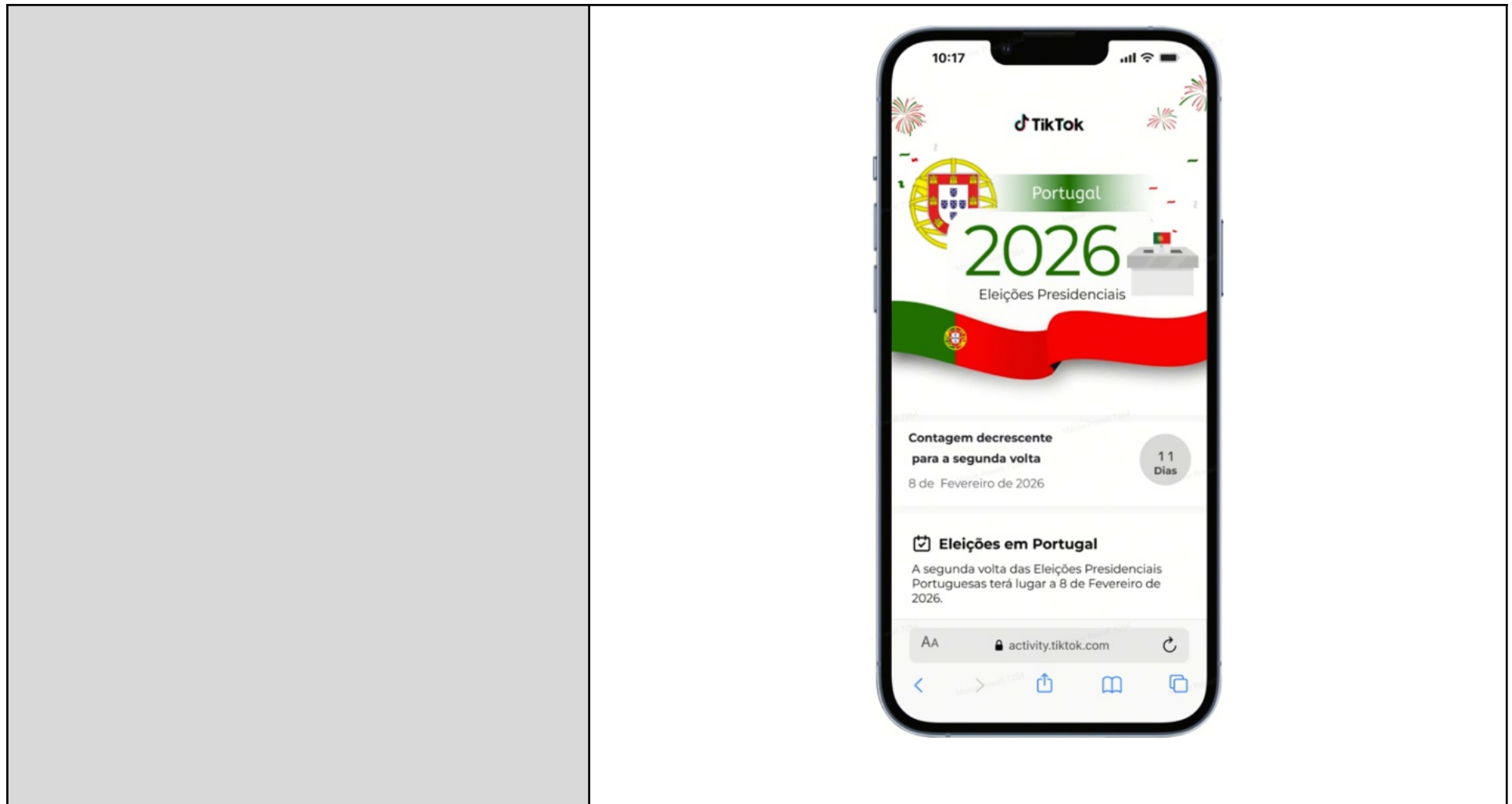
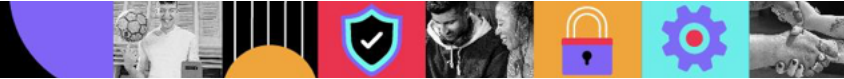
**QRE 17.2.1**

We run a variety of media literacy campaigns adapting our approach to the topic. We localise certain campaigns (e.g. for elections), meaning we collaborate with national partners to develop an approach that best resonates with the local audience. For other issues such as the War in Ukraine, our priority is to connect users to accurate and trusted resources.

Below are examples of the campaigns we have most recently run in-app which have leveraged a number of the intervention tools outlined above.

(I) Promoting election integrity. As well as the election integrity pages on TikTok's [Safety Center](#) and [Transparency Center](#), and the dedicated [Global Elections Hub](#), we launched media literacy campaigns in advance of several elections in the EU and wider Europe.

- **Portugal Presidential Election 2026:** From 9 Dec 2025, we launched an in-app [Election Centre](#) to provide users with up-to-date information about the 2026 Portugal presidential election, which contained a section about spotting misinformation, which included videos created in partnership with the fact-checking organisation [Polígrafo](#).



- **Aragon (Spain) Regional Election:** From 2 Feb 2026, we launched an in-app Search Guide and Details Page to provide users with up-to-date information about the Aragon local election, which contained a section about following our Community Guidelines, with a link to the local election office website.



- **German State Elections in Baden-Württemberg and Rheinland-Pfalz 2026:** From 9 Feb 2026, we launched an in-app [Election Centre](#) to provide users with up-to-date information about the German state elections taking place in March 2026, which contained a section providing tips for spotting misinformation, which included videos created in partnership with the fact-checking organisation [dpa](#).



- **Castile and Leon (Spain) Regional Election:** From 23 Feb 2026, we launched an in-app Search Guide and Details Page to provide users with up-to-date information about the Castile and Leon local election, which contained a section about following our Community Guidelines, with a link to the local election office website.



- **France Municipal Elections 2026:** From 4 Feb 2026, we launched an in-app [Election Centre](#) to provide users with up-to-date information about the France municipal elections, which contained a section providing tips for spotting misinformation, including videos created in partnership with the fact-checking organisation [Agence France-Presse \(AFP\)](#).



- **Netherlands Municipal Election:** From 18 Feb 2026, we launched an in-app Search Guide and Details Page to provide users with up-to-date information about the Netherlands' municipal elections, which contained a section about following our Community Guidelines, with a link to the local election office website and to isdatechtzo.nl for digital literacy resources.



- **Slovenia Parliamentary Election 2026:** From 19 Feb 2026, we launched an in-app [Election Centre](#) to provide users with up-to-date information about the Slovenia parliamentary election, which contained a section providing tips for spotting misinformation.



- **Denmark General Election 2026:** From 11 March 2026, we launched an in-app [Election Centre](#) to provide users with up-to-date information about the Denmark general election. The centre contained a section providing tips for spotting misinformation.



- **Hungary Parliamentary Election 2026:** From 10 March 2026, we launched an in-app [Election Centre](#) to provide users with up-to-date information about the Hungary parliamentary election, which contained a section providing tips for spotting misinformation.



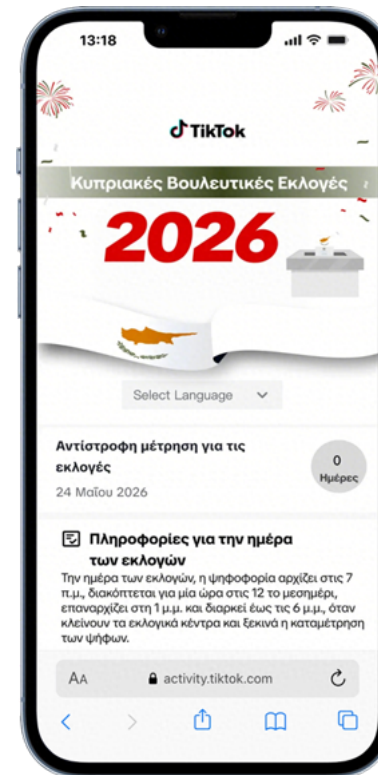
- **Bulgaria Parliamentary Election 2026:** From 19 April 2026, we launched an in-app [Election Centre](#) to provide users with up-to-date information about the Bulgaria parliamentary election, which contained a section providing tips for spotting misinformation.



- **Andalusia (Spain) Regional Election:** From 28 April 2026, we launched an in-app [Election Centre](#) to provide users with up-to-date information about the Andalusia election, which contained a section providing tips for spotting misinformation, including videos created in partnership with the fact-checking organisation [Newtral](#).



- **Cyprus Parliamentary Election 2026:** From 8 May 2026, we launched an in-app [Election Centre](#) to provide users with up-to-date information about the Cyprus parliamentary election, which contained a section providing tips for spotting misinformation.



- **Malta General Election 2026:** From 8 May 2026, we launched an in-app [Election Centre](#) to provide users with up-to-date information about the Malta general election, which contained a section providing tips for spotting misinformation.



- **New Caledonia (France) Election:** From 17 June 2026, we launched an in-app Search Guide and Details Page to provide users with up-to-date information about the elections in New Caledonia, which contained a section about following our Community Guidelines, with a link to the local election office website.



(II) Media literacy (General). We continue our ongoing **general media literacy and critical thinking skills campaigns** in the EU in collaboration with our fact-checking and media literacy partners, spanning 14 countries (Denmark, Finland, France, Georgia, Germany, Ireland, Italy, Romania, Spain, Sweden, Moldova, Netherlands, Poland, and Portugal).



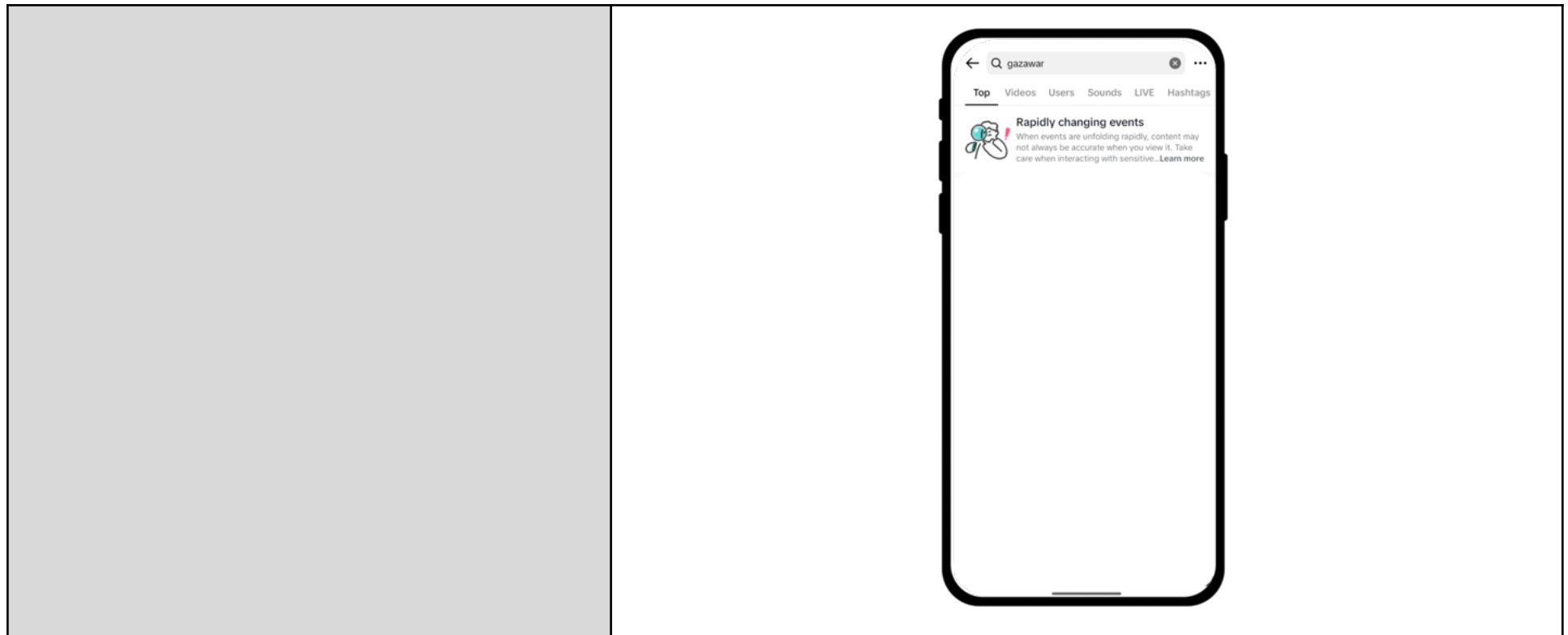
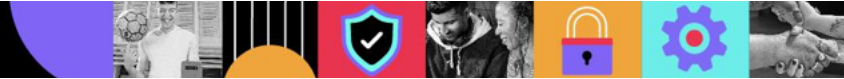
(III) Media literacy (War in Ukraine). We continue to serve 17 localised media literacy campaigns specific to the war in Ukraine in: Ukraine, Romania, Slovakia, Hungary, Latvia, Estonia, Lithuania, Czechia, Poland, Croatia, Slovenia, Bulgaria, Germany, Austria, Bosnia, Montenegro, and Serbia.

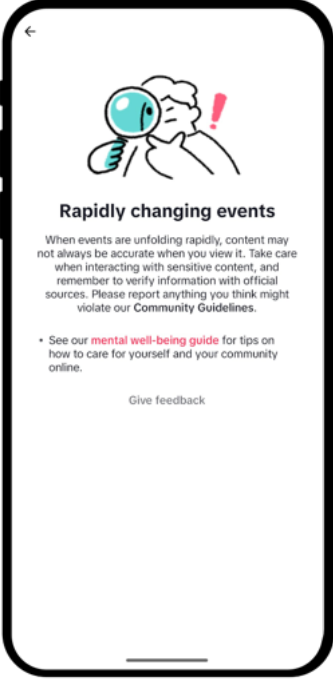
- Partnered with Lead Stories: Ukraine, Romania, Slovakia, Hungary, Latvia, Estonia, Lithuania.
- Partnered with fakenews.pl: Poland.
- Partnered with Correctiv: Germany, Austria.

Through these campaigns, users searching for keywords relating to the war in Ukraine on TikTok are directed to tips prepared in partnership with local media literacy bodies and our trusted fact-checking partners, to help them identify misinformation and prevent its spread on the platform.

(IV) Israel-Hamas conflict. To help raise awareness and to protect our users, we have search interventions that are triggered when users search for neutral terms related to this topic (e.g. Israel, Palestine). These search interventions remind users to pause and check their sources and also direct them to well-being resources.

- On 1 March 2026, we expanded our Rapidly Changing Events search guide to include keywords related to the conflict in Iran.
- The Rapidly Changing Events search intervention that related to the Israel/Hamas conflict, and subsequently expanded to include the Iran conflict, was ended on 11 June 2026.



				
SLI 17.2.1 - actions enforcing policies above	We are pleased to report metrics on the 12 general media literacy and critical thinking skills campaigns that ran in the EEA through the reporting period in Germany, Romania, Poland, Denmark, Finland, France, Ireland, Italy,, Portugal, Spain, Sweden, and the Netherlands.			



Member States	Total number of impressions of the H5 Page (Views generated between January 1 and June 30, 2026.	Number of impressions of the search intervention	Number of clicks on the search intervention	Click through rate of the search intervention
France (in partnership with AFP)	62,921	27,832,706	83,877	0.3%
Portugal (in partnership with Poligrafo)	9,004	5,389,311	19,453	0.36%
Denmark (in partnership with Logically Facts)	764	242,468	1,087	0.45%
The Netherlands (in partnership with Nieuwscheckers)	12,572	3,347,885	19,401	0.58%
Ireland (in partnership with The journal.ie)	1,064	296,704	2,003	0.68%
Finland (in partnership with Logically Facts)	1,109	186,749	2,051	1.1%
Sweden (in partnership with Logically Facts)	1,806	383,190	3,279	0.86%
Spain (in partnership with Maldita)	16,222	12,288,223	31,930	0.26%
Italy (in partnership with Facta)	1,666	513,261	2,706	0.53%
Germany (in partnership with Correctiv)	18,392	8,592,576	32,445	0.38%
Poland	10,384	13,917,257	100,745	0.72%
Romania	8,786	489,946	2,512	0.51%



Measure 17.3	
QRE 17.3.1	We work with fact-checking partners and media literacy bodies to develop campaigns that educate users and redirect them to authoritative resources. Specific examples of partnerships within the campaigns and projects set out in QRE 17.2.1 are: (I) Promoting election integrity. We partner with various media organisations and fact-checkers to promote election integrity on TikTok. For more detail about the input our fact-checking partners provide please refer to QRE 30.1.3. During this reporting period, we worked with European fact-checkers and media literacy organisations on 4 temporary media literacy election integrity campaigns, in advance of elections, through our in-app Election Centers: • Portugal (local election): Polígrafo • Baden-Württemberg and Rheinland-Pfalz (Germany State Elections) - Deutsche Presse-Agentur (dpa) • France (Municipal Elections) - Agence France-Presse (AFP) • Andalusia (Spain Regional Election) - Newtral (II) War in Ukraine. We continue to run our media literacy campaigns about the war in Ukraine, developed in partnership with our media literacy partners Correctiv in Austria and Germany, Fakenews.pl in Poland and Lead Stories in Ukraine, Romania, Slovakia, Hungary, Latvia, Estonia, Lithuania. This campaign is also active in Czechia, Croatia, Slovenia, Bulgaria.



V. Empowering Users

Commitment 18

Relevant Signatories commit to minimise the risks of viral propagation of Disinformation by adopting safe design practices as they develop their systems, policies, and features.

In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	Yes
If yes, list these implementation measures here [short bullet points].	- Improved the accuracy of, and overall coverage provided by, our machine learning detection models to better identify disinformation. - Continued testing large language models (LLMs) to further support proactive moderation at scale. Because LLMs can comprehend human language and perform highly specific, complex tasks, we are better able to moderate nuanced areas like misinformation by extracting specific misinformation "claims" from videos for moderators to assess directly or route to our fact-checking partners. - In June 2026, TikTok both sponsored and sent a delegation of attendees to the Global Fact conference in Vilnius, Lithuania.
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 18.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .



QRE 18.1.1	N/A			
QRE 18.1.2	N/A			
QRE 18.1.3	N/A			
SLI 18.1.1 - actions proving effectiveness of measures and policies	N/A			
Measure 18.2				
QRE 18.2.1	Our Terms of Service and Integrity and Authenticity policies under our Community Guidelines are the first line of defence in combating harmful misinformation and (as outlined in more detail in QRE 14.1.1) deceptive behaviours on our platform. These rules make clear to our users what content we remove or make ineligible for the For You feed when they pose a risk of harm to our users and our community. Specifically, our policies do not allow: Misinformation - Misinformation that poses a risk to public safety or incites panic about a crisis event or emergency, including falsely presenting past crisis events as recent or claiming that critical resources are unavailable during emergencies - Health misinformation that could cause significant harm, such as promoting unproven treatments that may be fatal, discouraging professional care for			



life-threatening conditions (e.g., vaccine effectiveness), or spreading false information about how such conditions are transmitted

- Misinformation that denies the existence of climate change, misrepresents its causes, or contradicts its established environmental impact
- Conspiracy theories or hoaxes that could cause significant harm, such as those that make a violent call to action or have links to previous violence

- **Civic and Election Integrity**

- Election misinformation, including:
 - How, when, and where to vote or register to vote
 - Voter eligibility or candidate qualifications
 - Laws or procedures for elections, referendums, ballot initiatives, or censuses
 - Election results

- **Edited Media and AI-Generated Content (AIGC)**

- Using the likeness of private figures without consent
- Sexualized, fetishized, or victimizing depictions
- AI-created likenesses made to bully or harass
- Accounts focused on AI images of youth in clothing suited for adults, or sexualized poses or facial expressions
- AIGC or significantly edited content that misleads about a matter of public importance, such as:
 - Content made to look like it comes from a real news source
 - A crisis event, like a natural disaster or conflict
 - A public figure being degraded, harassed, or linked to criminal behavior
 - A public figure taking political stances, supporting products, or commenting on public issues they haven't actually addressed
 - A political endorsement or condemnation that never happened
- Any content that breaks our Community Guidelines, including those on impersonation, misinformation, and hate speech, even if it's AI-generated

- **Fake Engagement**

- Facilitating the trade or marketing of services that artificially increase engagement, such as selling followers or likes.
- Providing instructions on how to artificially increase engagement on TikTok.



We have made clear to our users [here](#) that the following content is ineligible for the For You feed:

- **Misinformation**

- Conspiracy theories that are unfounded and claim that certain events or situations are carried out by covert or powerful groups, such as "the government" or a "secret society".
- Moderate harm health misinformation, such as an unproven recommendation for how to treat a minor illness.
- Repurposed media, such as showing a crowd at a music concert and suggesting it is a political protest.
- Misrepresenting authoritative sources, such as selectively referencing certain scientific data to support a conclusion that is counter to the findings of the study.
- Unverified claims related to an emergency or unfolding event.
- Potential high-harm misinformation while it is undergoing a fact-checking review.

- **Civic and Election Integrity**

- Unverified claims about an election, such as a premature claim that all ballots have been counted or tallied.
- Statements that significantly misrepresent authoritative civic information, such as a false claim about the text of a parliamentary bill.

- **Fake Engagement**

- Content that tricks or manipulates others as a way to increase gifts, or engagement metrics, such as "like-for-like" promises or other false incentives for engaging with content.

As outlined under Commitment 14, we also remove accounts that seek to mislead people or use TikTok to deceptively sway public opinion. These activities range from inauthentic or fake account creation, to more sophisticated efforts to undermine public trust.

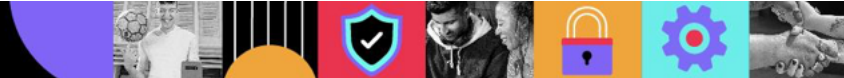
TikTok Shop. We maintain policies to prohibit harmful false or misleading information on TikTok Shop. These include the TikTok Shop Artificial Intelligence Generated Content (AIGC) policy and Misleading Functionality and Effect policies. These policies are described in more detail in Measure 14. Under these policies, creators and sellers must not use AI-generated or manipulated content, editing techniques, special effects or other representations that mislead users as to a product's functionality, attributes, effects, performance or expected results. TikTok Shop also



	prohibits certain misleading health-related claims, including claims that products can cure, treat, prevent or eliminate medical conditions or symptoms We have policy experts within our Trust and Safety team dedicated to the topic of integrity and authenticity. They continually keep these policies under review and collaborate with external partners and experts to understand whether updates or new policies are required and ensure they are informed by a diversity of perspectives, expertise, and lived experiences. Enforcing our policies. We remove content – including video, audio, livestream, images, comments, links, or other text – that violates our Integrity and Authenticity policies. Individuals are notified of our decisions and can appeal them if they believe no violation has occurred. We also make clear in our Community Guidelines that we will temporarily or permanently ban accounts and/or users that are involved in serious or repeated violations, including violations of our Integrity and Authenticity policies. For TikTok Shop content, enforcement measures may also include limiting recommendation or distribution of content through Feeds, Shop Tab, or Search, removing product anchors, and/or restricting creators' e-commerce permissions. As with our other policies, enforcement is applied proportionately, taking into account the severity and pattern of non-compliance. We are pleased to include in this report the number of videos made ineligible for the For You feed under the relevant Integrity and Authenticity policies.
SLI 18.2.1	Methodology of data measurement: We have based the following numbers on the country in which the video was posted: videos removed because of violations of our Misinformation, Civic and Election Integrity and Edited media and AIGC policies. The number of views of videos removed because of violation of each of these policies is based on the approximate location of the user. We also updated the methodology on the number of videos made ineligible for the For You feed under our Misinformation policy. For this report, we are also including TikTok Shop metrics under COCD for the first time. Under SLI 18.2.1, we are sharing the number of actions taken on TikTok Shop under the relevant AIGC and misinformation policies.



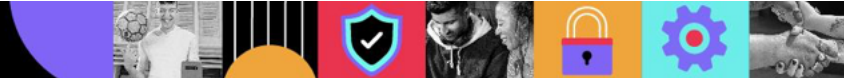
	On 15 June 2026, we launched TikTok Shop in eight additional EU member states: Poland, Netherlands, Belgium, Czech Republic, Austria, Greece, Portugal and Hungary. The metrics in this report include data for the TikTok Shop market launches covering the period from launch to the end of the reporting period.				
	Total no of violations	Metric 1: indicating the impact of the action taken		Total no of violations	Metric 1: indicating the impact of the action taken
List actions per member states and languages (see example table above)	Number of videos removed because of violation of Misinformation policy	Number of views of videos removed because of violation of Misinformation policy		Number of videos made ineligible for the For You feed under the Misinformation policy.	
Member States					
Austria	11308	3218475		4077	
Belgium	15105	3976317		7162	
Bulgaria	42197	4714923		5870	
Croatia	9243	1958426		1210	
Cyprus	3222	1172909		1752	
Czechia	10470	2537518		4305	
Denmark	15725	1389364		2593	



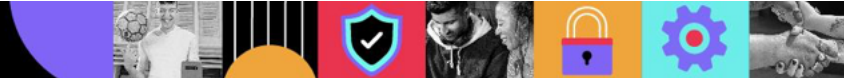
Estonia	1963	175119		826	
Finland	5307	16740320		2373	
France	100530	123033306		50124	
Germany	113943	94571527		57421	
Greece	8404	6095813		7711	
Hungary	16159	10995548		3802	
Ireland	9383	402371		4076	
Italy	41780	14807032		53760	
Latvia	3648	1357460		1427	
Lithuania	3146	7321042		1533	
Luxembourg	897	372475		354	
Malta	501	26668		458	
Netherlands	80988	11906044		42591	
Poland	63967	31092484		23880	
Portugal	11667	1374918		5248	
Romania	47144	10399917		25016	



Slovakia	4076	2402576		1894	
Slovenia	1273	2094126		561	
Spain	56000	26353206		43237	
Sweden	13090	3654101		4836	
Iceland	345	128459		218	
Liechtenstein	5	0		7	
Norway	17075	2120896		2801	
Total EU	691136	384143985		358097	
Total EEA	708561	386393340		361123	
List actions per member states and languages (see example table above)	Number of videos removed because of violation of Civic and Election Integrity policy	Number of views of videos removed because of violation of Civic and Election Integrity policy		Number of videos removed because of violation of Edited Media and AI-Generated Content (AIGC)	Number of views of videos removed because of violation of Edited Media and AI-Generated Content (AIGC)
Member States					
Austria	188	421509		11719	3734469
Belgium	328	149648		10754	3425892
Bulgaria	229	656676		42241	1708891
Croatia	73	3802		16366	722675



Cyprus	83	207		4826	912258	
Czechia	150	29797		7177	10261960	
Denmark	157	1309		15176	8686782	
Estonia	55	13465		4085	2056073	
Finland	105	193		5372	4735509	
France	1093	2595543		63056	94228457	
Germany	2443	6223320		99159	128170986	
Greece	210	3938		5441	867743	
Hungary	194	1125008		14592	3097025	
Ireland	206	2636		6948	1363165	
Italy	1154	755531		37133	13472997	
Latvia	62	2802		4805	5250952	
Lithuania	106	19690		4861	3304486	
Luxembourg	30	49		2078	23604	
Malta	22	24709		910	1778433	
Netherlands	817	55117		85596	33693554	
Poland	880	3333956		42850	30661634	
Portugal	257	612567		11637	2861651	
Romania	3243	48434		27051	6105146	

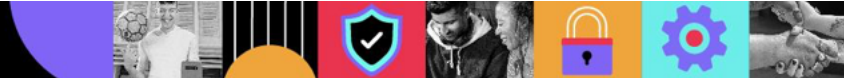


Slovakia	74	6274		4077	146257	
Slovenia	38	248		1570	273450	
Spain	920	52512		45863	35355202	
Sweden	210	219528		10924	26469283	
Iceland	9	205		568	339717	
Liechtenstein	0	0		63	0	
Norway	172	160		27016	1668462	
Total EU	13327	16358468		586267	423368534	
Total EEA	13508	16358833		613914	425376713	
List actions per member states and languages (see example table above)	Number of actions taken on TikTok Shop for violation of Edited Media and AI-Generated Content (AIGC) policy	Number of actions taken on TikTok Shop for violation of Misleading Functionality and Effect policy policy				
Member States						
Austria	0	2				
Belgium	0	2				
Czechia	1	8				
France	9212	3458				



Germany	14772	2331			
Ireland	18	57			
Italy	5674	1374			
Netherlands	2	21			
Poland	3	37			
Portugal	11	11			
Spain	9441	2435			
Total EU	39,134	9736			
Total EEA	39,134	9736			
Measure 18.3	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .				
QRE 18.3.1	N/A				

V. Empowering Users
Commitment 19 Relevant Signatories using recommender systems commit to make them transparent to the recipients regarding the main criteria and parameters used for prioritising or deprioritising information, and provide options to users about recommender systems, and make available information on those options.



In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 19.1	
QRE 19.1.1	The For You feed is the interface users first see when they open TikTok. It's central to the TikTok experience and where most of our users spend their time exploring the platform. We make clear to users in our Terms of Service (section 4) and our Help Center article (and in our and Safety Center guide) that each account holder's For You feed is based on a personalised recommendation system. As well as removing harmful misinformation content that violates our Community Guidelines, we take steps to avoid recommending certain categories of content that may not be appropriate for a broad audience, including general conspiracy theories and unverified information related to an emergency or unfolding event. We may also make some of this content harder to find in search. Main parameters. The system recommends content by ranking content based on a combination of factors, including: • User interactions (e.g. content users like, share, comment on, and watch in full or skip, as well as accounts of followers that users follow back);



- Content information (e.g. sounds, hashtags, number of views, and the country the content was published in); and
- User information (e.g. device settings, language preferences, location, time zone and day, and device type).

The main parameters help us make predictions on the content users are likely to be interested in. Different factors can play a larger or smaller role in what's recommended, and the importance – or weighting – of a factor can change over time. For many users, the time spent watching a specific video is generally weighted more heavily than other factors. These predictions are also influenced by the interactions of other people on TikTok who appear to have similar interests.

Users can access the “Why this video” feature, which allows them to see on any particular video that appears in their For You feed factors that influenced why it appeared. The feature essentially explains to users how past interactions on the platform have impacted the video they have been recommended.

TikTok Shop. TikTok Shop uses the same core recommendation parameters described above, namely user interactions, content information and user information. These factors are used to personalise recommendations across TikTok Shop features, including Feeds, Shop Tab and Search. Further information regarding TikTok Shop's recommendation systems is available in the TikTok Help Center page titled '[How TikTok Shop recommends content](#)', which is linked from the [TikTok Shop Terms of Use](#).

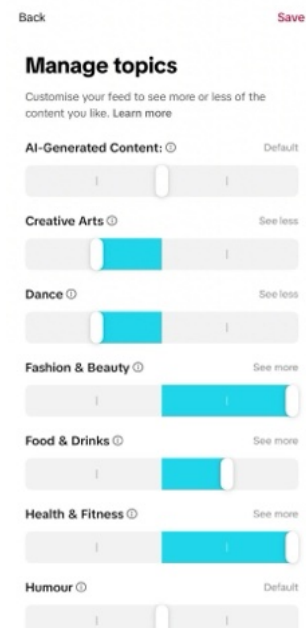
User preferences. We also empower our users to customise their experience to their preferences and comfort.

For example, in the For You feed:

- Users can click on any video and select “not interested” to indicate that they do not want to see similar content.
- Users are able to automatically filter out specific words or hashtags from the content recommended to them (see [here](#)).
- Users are able to [refresh their For You feed](#) if they no longer feel like recommendations are relevant to them or are too similar. When the For You feed is refreshed, users view a number of new videos which include popular videos, and their interaction with these new videos will inform future recommendations.



- Users can also personalise their "For You" page through our **Manage Topics** feature. This allows users to adjust the frequency of content they see related to particular topics. The settings don't eliminate topics entirely but can influence how often they're recommended as peoples' interests evolve over time.

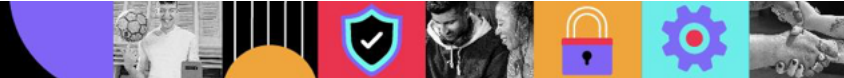


- As part of our obligations under the **DSA (Article 38)**, we introduced non-personalized feeds on our platform, which enables our EU users to turn off personalisation so that feeds show non-personalised content (including TikTok Shop content). The For You feed will instead show popular videos in their region and internationally. See [here](#).

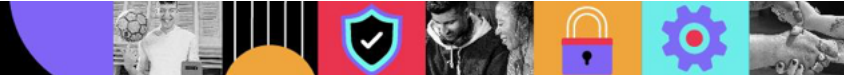
Measure 19.2



SLI 19.2.1 – user settings	Methodology of data measurement: The number of users who have filtered hashtags or a keyword to set preferences for For You feed, the number of times users clicked “not interested” in relation to the For You feed, and the number of times users clicked on the For You Feed Refresh are all based on the approximate location of the users that engaged with these tools. The number for videos tagged with AIGC label includes both automatic and creator-generated labeling. AIGC labelling is a tool that helps inform users about the content that is recommended to them. The following metrics reflect how users interact with tools that TikTok makes available for users to influence and customise recommendations.			
	No of times users actively engaged with these settings	No of times users actively engaged with these settings		
List actions per member states and languages (see example table above)	Number of users that filtered hashtags or words	Number of users that clicked on "not interested"	Number of times users clicked on the For You Feed Refresh	Number of Videos tagged with AIGC label
Member States				
Austria	79581	1423257	52640	1595126
Belgium	122761	1901101	83694	2398811
Bulgaria	64907	1282340	47197	3179713
Croatia	38331	547873	22887	490074
Cyprus	19877	292189	15614	500336
Czechia	66773	1279910	37427	1556143



Denmark	47932	830450	30642	765958
Estonia	20407	262317	12441	255386
Finland	73960	909321	51941	765771
France	724990	11462004	449590	16441982
Germany	860636	12080787	573382	19567620
Greece	107543	1685674	89778	2594622
Hungary	65667	1497804	42982	2908675
Ireland	86060	1241931	73537	646990
Italy	461075	8683480	301800	14115592
Latvia	27123	578124	23813	609540
Lithuania	34860	704421	28890	1829906
Luxembourg	7491	116236	4294	151491
Malta	8206	153729	6534	142659
Netherlands	262002	4963767	193401	3550852
Poland	308148	5010960	219510	6236469
Portugal	105726	1386953	63010	2957568
Romania	172035	3715213	192428	9834962
Slovakia	29416	452807	17069	1133916
Slovenia	14875	266430	12339	233253



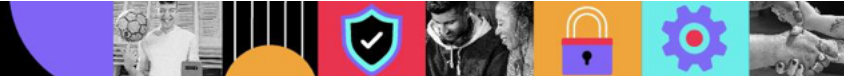
Spain	551153	9994444	349474	14876703
Sweden	126193	1983104	91940	1677965
Iceland	6683	78376	3943	44152
Liechtenstein	184	8264	278	2437
Norway	68484	930172	44812	731254
Total EU	4487728	74706626	3088254	111018083
Total EEA	4563079	75723438	3137287	111795926

V. Empowering Users

Commitment 20

Relevant Signatories commit to empower users with tools to assess the provenance and edit history or authenticity or accuracy of digital content.

In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A



If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 20.1	
QRE 20.1.1	N/A
Measure 20.2	
QRE 20.2.1	N/A

V. Empowering Users

Commitment 21

Relevant Signatories commit to strengthen their efforts to better equip users to identify Disinformation. In particular, in order to enable users to navigate services in an informed way, Relevant Signatories commit to facilitate, across all Member States languages in which their services are provided, user access to tools for assessing the factual accuracy of sources through fact-checks from fact-checking organisations that have flagged potential Disinformation, as well as warning labels from other authoritative sources.

In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	Yes
--	-----



If yes, list these implementation measures here [short bullet points].

- We ran 14 temporary **media literacy election integrity campaigns** in advance of regional elections, some in collaboration with our fact-checking and media literacy partners:
 - 14 in the EU
 - Portugal (Presidential Election) - Polígrafo
 - Aragon (Spain Regional Election)
 - Baden-Württemberg and Rheinland-Pfalz (Germany State Elections) - Deutsche Presse-Agentur (dpa)
 - Castile and Leon (Spain Regional Election)
 - France (Municipal Elections) - Agence France-Presse (AFP)
 - Netherlands (Municipal Elections)
 - Slovenia (Parliamentary Election)
 - Denmark (General Election)
 - Hungary (Parliamentary Election)
 - Bulgaria (Parliamentary Election)
 - Andalusia (Spain Regional Election) - Newtral
 - Cyprus (Parliamentary Election)
 - Malta (General Election)
 - New Caledonia (France Regional Election)
- Following the outbreak of Ebola in Sub-saharan Africa, we launched an in-app guide to provide users with guidance on verifying information using reliable sources. This guide was launched in Mayotte (France) where we linked to info.gouv.fr.
- Following the outbreak of Hantavirus, we launched an in-app guide and a video tag to provide users with guidance on verifying information using reliable sources. This guide was launched in all EU markets, where we linked to the World Health Organisation (WHO) and other government health advice websites (IT, DE, NL, FR).

Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]

N/A



If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 21.1	
QRE 21.1.1	We currently have 13 IFCN accredited fact-checking partners across the EU, EEA, and wider Europe: 1. Agence France-Presse (AFP) 2. Deutsche Presse-Agentur (dpa) 3. Demagog.pl 4. Facta 5. Faktograf 6. Internews Kosova 7. Lead Stories 8. Newtral 9. Poligrafo 10. Reuters 11. Teyit 12. Science Feedback - For advertising-related fact-checking partnerships, please refer to Chapter 2. 13. Geofacts These partners provide fact-checking coverage in 23 official EEA languages, including at least one official language of each EU Member States, plus Georgian, Russian, Turkish, and Ukrainian. We partner with fact checking organisations in the following ways: • Enforcement of misinformation policies. Our fact-checking partners play a critical role in helping us enforce our misinformation policies. As part of our review process, our human moderators can provide content suspected to be misinformation to our expert fact-checking partners for further evaluation. Where fact-checking partners



advise that content is false, our moderators take measures to assess and remove it from our platform. Our response to QRE 31.1.1 provides further insight into the way in which fact-checking partners are involved in this process.

- **Unverified content labelling.** Sometimes, our fact-checking partners determine that the accuracy of content cannot be confirmed or checks are inconclusive (especially during unfolding events). Where our fact-checking partners provide us with a rating that demonstrates the claim cannot yet be verified, we may use our unverified content label to inform viewers [via a banner](#) that a video contains unverified content, to raise user awareness about content credibility. In these circumstances, the content creator is also notified that their video was flagged as unsubstantiated content and the video will become ineligible for recommendation in the For You feed.
- **In-app tools related to specific topics:**
 - **Election integrity.** We have launched campaigns in advance of several major elections aimed at educating the public about the voting process, which encourage users to fact-check information with our fact-checking partners. For example, the election integrity campaign we rolled out in advance of Portugal's presidential elections in January and February 2026 included an in-app [Election Centre](#). The centre contained a section about spotting misinformation, which included videos created in partnership with fact-checking organisation [Polígrafo](#). In total, during the reporting period, we ran 14 temporary **media literacy election integrity campaigns** in advance of regional elections.
 - **Rapidly changing events:**
 - Following the outbreak of Ebola in Sub-saharan Africa, we launched an in-app guide to provide users with guidance on verifying information using reliable sources. This guide was launched in Mayotte (France) where we linked to [info.gouv.fr](#).
 - Following the outbreak of Hantavirus, we launched an in-app guide and a video tag to provide users with guidance on verifying information using reliable sources. This guide was launched in all EU markets, where we linked to the World Health Organisation (WHO) and other government health advice websites (IT, DE, NL, FR).



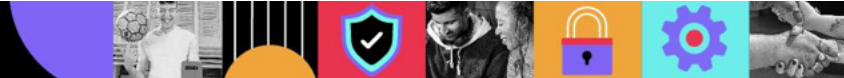
	User awareness of our fact-checking partnerships and labels. We have created pages on our new Trust & Safety Center to raise users' awareness about our fact-checking program, labels and to support the work of our safety partners.				
SLI 21.1.1 - actions taken under measure 21.1	Methodology of data measurement: The share of removals under our harmful misinformation policy, share of proactive removals, share of removals before any views and share of the removals within 24h are relative to the total removals of each policy. The share cancel rate (%) following the unverified content label share warning pop-up indicates the percentage of users who do not share a video after seeing the label pop up. This metric is based on the approximate location of the users that engaged with these tools.				
	Reach of labels/ fact-checkers and other authoritative sources	Other pertinent metric	Other pertinent metric	Other pertinent metric	Other pertinent metric
List actions per member states and languages (see example table above)	Share cancel rate (%) following the unverified content label share warning pop-up (users who do not share the video after seeing the pop up)	Share of removals under misinformation policy	Share of proactive removals under misinformation policy	Share of video removals before any views under misinformation policy	Share of video removals within 24h by misinformation policy
Member States					
Austria	30.5%	52.2%	99.7%	90.5%	96.5%
Belgium	31.0%	61.5%	99.6%	91.4%	96.0%



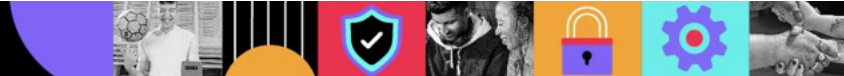
Bulgaria	30.0%	49.8%	99.9%	96.9%	99.3%
Croatia	28.2%	36.2%	99.7%	94.4%	98.7%
Cyprus	31.5%	39.9%	99.6%	94.4%	97.8%
Czechia	29.2%	50.3%	99.5%	80.0%	96.3%
Denmark	31.2%	54.1%	99.9%	92.7%	95.5%
Estonia	29.5%	31.6%	99.7%	84.2%	97.5%
Finland	33.0%	48.7%	99.5%	83.5%	96.9%
France	30.6%	64.3%	99.4%	82.2%	94.3%
Germany	33.6%	55.3%	99.4%	85.5%	95.6%
Greece	30.1%	60.9%	99.6%	79.0%	96.7%
Hungary	28.1%	53.0%	99.7%	95.6%	98.6%
Ireland	33.0%	57.3%	99.8%	92.0%	97.7%
Italy	34.0%	55.0%	99.6%	91.4%	96.5%
Latvia	32.6%	42.7%	99.8%	87.9%	98.7%
Lithuania	30.2%	38.4%	99.1%	86.2%	96.9%
Luxembourg	26.5%	31.1%	99.9%	92.0%	96.3%
Malta	29.1%	36.0%	99.6%	93.6%	96.4%
Netherlands	29.7%	62.5%	99.7%	96.2%	98.3%
Poland	27.5%	63.1%	99.5%	77.3%	97.6%



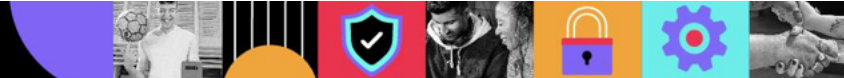
Portugal	27.1%	54.1%	99.8%	93.0%	98.0%
Romania	29.9%	63.3%	99.7%	70.0%	96.5%
Slovakia	28.2%	53.5%	99.6%	76.7%	96.9%
Slovenia	27.5%	44.6%	98.0%	73.8%	91.2%
Spain	30.6%	57.8%	99.7%	92.8%	97.5%
Sweden	31.8%	54.9%	99.7%	84.8%	96.4%
Iceland	25.9%	37.1%	99.7%	91.0%	95.7%
Liechtenstein	38.5%	7.0%	100.0%	100.0%	100.0%
Norway	29.7%	38.6%	99.9%	96.4%	98.9%
Total EU	30.8%	57.1%	99.6%	86.8%	96.7%
Total EEA	30.8%	56.4%	99.6%	87.1%	96.7%
Member States		Share of video removals under Civic and Election Integrity policy	Share of proactive video removals under Civic and Election Integrity policy	Share of video removals before any views under Civic and Election Integrity policy	Share of video removals within 24h under Civic and Election Integrity policy
Austria		0.9%	99.5%	91.0%	97.5%
Belgium		1.3%	98.8%	90.5%	93.0%
Bulgaria		0.3%	95.6%	78.2%	99.4%



Croatia		0.3%	100.0%	84.9%	99.4%
Cyprus		1.0%	100.0%	95.2%	98.2%
Czechia		0.7%	100.0%	86.7%	97.8%
Denmark		0.5%	100.0%	94.3%	98.1%
Estonia		0.9%	98.2%	94.5%	98.8%
Finland		1.0%	100.0%	94.3%	96.9%
France		0.7%	97.7%	81.2%	93.2%
Germany		1.2%	98.3%	84.2%	94.8%
Greece		1.5%	100.0%	84.8%	94.0%
Hungary		0.6%	88.7%	73.2%	98.6%
Ireland		1.3%	100.0%	87.9%	96.6%
Italy		1.5%	98.4%	86.4%	95.0%
Latvia		0.7%	100.0%	91.9%	98.3%
Lithuania		1.3%	100.0%	97.2%	97.3%
Luxembourg		1.0%	100.0%	76.7%	98.2%
Malta		1.6%	100.0%	68.2%	95.3%
Netherlands		0.6%	98.5%	92.5%	98.5%
Poland		0.9%	98.6%	72.7%	97.7%
Portugal		1.2%	98.4%	88.3%	96.8%



Romania		4.4%	99.9%	74.6%	96.1%
Slovakia		1.0%	100.0%	86.5%	98.4%
Slovenia		1.3%	100.0%	84.2%	96.0%
Spain		0.9%	98.5%	76.1%	94.4%
Sweden		0.9%	98.6%	84.8%	95.3%
Iceland		1.0%	100.0%	77.8%	97.0%
Liechtenstein		0.0%	0.0%	0.0%	100.0%
Norway		0.4%	100.0%	93.0%	99.2%
Total EU		1.1%	98.7%	81.6%	96.3%
Total EEA		1.1%	98.7%	81.7%	96.5%
Member States		% video removals under Edited Media and AI-Generated Content (AIGC) policy	% proactive video removals under Edited Media and AI-Generated Content (AIGC) policy	% video removals before any views under Edited Media and AI-Generated Content (AIGC) policy	% video removals within 24h under Edited Media and AI-Generated Content (AIGC) policy
Austria		54.1%	99.7%	95.6%	97.5%
Belgium		43.8%	98.5%	90.5%	93.0%
Bulgaria		49.9%	100.0%	98.9%	99.4%
Croatia		64.0%	99.9%	98.5%	99.4%



Cyprus		59.8%	99.6%	96.6%	98.2%
Czechia		34.5%	99.5%	95.1%	97.8%
Denmark		52.2%	99.9%	97.1%	98.1%
Estonia		65.8%	99.7%	97.3%	98.8%
Finland		49.3%	99.5%	95.1%	96.9%
France		40.3%	98.9%	90.4%	93.2%
Germany		48.1%	98.7%	91.7%	94.8%
Greece		39.4%	99.3%	88.9%	94.0%
Hungary		47.9%	99.4%	98.1%	98.6%
Ireland		42.4%	99.7%	93.7%	96.6%
Italy		48.9%	99.3%	92.8%	95.0%
Latvia		56.3%	99.6%	96.6%	98.3%
Lithuania		59.4%	99.2%	94.5%	97.3%
Luxembourg		72.0%	100.0%	96.0%	98.2%
Malta		65.3%	99.7%	93.0%	95.3%
Netherlands		66.1%	99.3%	97.5%	98.5%
Poland		42.2%	99.2%	93.7%	97.7%
Portugal		53.9%	99.5%	95.3%	96.8%
Romania		36.3%	99.5%	92.8%	96.1%



Slovakia		53.5%	99.8%	94.6%	98.4%
Slovenia		55.0%	99.6%	92.7%	96.0%
Spain		47.3%	99.6%	92.0%	94.4%
Sweden		45.8%	99.4%	91.9%	95.3%
Iceland		61.1%	98.2%	94.4%	97.0%
Liechtenstein		88.7%	100.0%	100.0%	100.0%
Norway		61.1%	99.9%	98.7%	99.2%
Total EU		48.4%	99.3%	94.1%	96.3%
Total EEA		48.9%	99.3%	94.3%	96.5%

SLI 21.1.2 - actions taken under measure 21.1

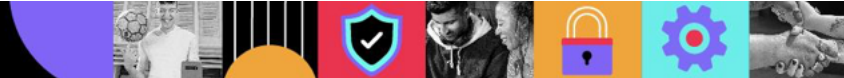
Methodology of data measurement:

The number of videos tagged with the unverified content label is based on the country in which the video was posted.

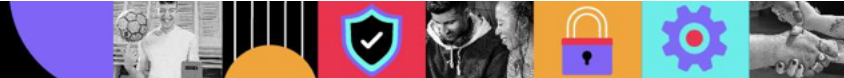
The share cancel rate (%) following the unverified content label share warning pop-up indicates the percentage of users who do not share a video after seeing the label pop-up. This metric is based on the approximate location of the users that engaged with these tools.

Number of labels applied to content, such as on the basis of such articles

Meaningful metrics such as the impact of 21.1. measures on user interactions with, or user re-shares of, content fact-checked as false or misleading



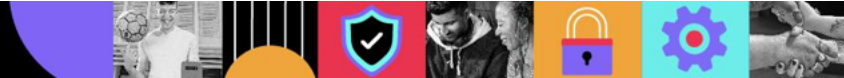
List actions per member states and languages (see example table above)		Number of videos tagged with the unverified content label		Share cancel rate (%) following the unverified content label share warning pop-up (users who do not share the video after seeing the pop up)
Member States				
Austria		48		30.5%
Belgium		99		31.0%
Bulgaria		99		30.0%
Croatia		15		28.2%
Cyprus		6		31.5%
Czechia		165		29.2%
Denmark		195		31.2%
Estonia		15		29.5%
Finland		33		33.0%
France		913		30.6%
Germany		826		33.6%
Greece		121		30.1%
Hungary		18		28.1%
Ireland		23		33.0%
Italy		481		34.0%



Latvia		23		32.6%
Lithuania		13		30.2%
Luxembourg		2		26.5%
Malta		0		29.1%
Netherlands		131		29.7%
Poland		207		27.5%
Portugal		44		27.1%
Romania		211		29.9%
Slovakia		78		28.2%
Slovenia		14		27.5%
Spain		309		30.6%
Sweden		75		31.8%
Iceland		0		25.9%
Liechtenstein		0		38.5%
Norway		53		29.7%
Total EU		4164		30.8%
Total EEA		4217		30.8%

Measure 21.2

TikTok did not subscribe to this measure as outlined in the January 2025 [Subscription Document](#).



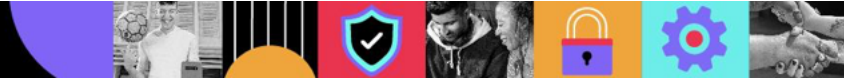
QRE 21.2.1	N/A
Measure 21.3	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 21.3.1	N/A

V. Empowering Users	
Commitment 22	
Relevant Signatories commit to provide users with tools to help them make more informed decisions when they encounter online information that may be false or misleading, and to facilitate user access to tools and information to assess the trustworthiness of information sources, such as indicators of trustworthiness for informed online navigation, particularly relating to societal issues or debates of general interest.	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	N/A.
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A



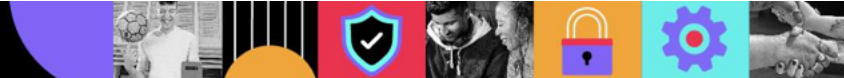
Measure 22.1	
QRE 22.1.1	N/A
SLI 22.1.1 - actions enforcing policies above	N/A
	N/A

Measure 22.2	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 22.2.1	N/A
Measure 22.3	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 22.3.1	N/A

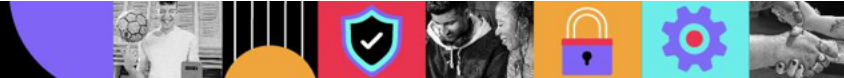


Measure 22.4	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 22.4.1	N/A
SLI 22.4.1 - actions enforcing policies above	N/A
	N/A
Data	N/A
Measure 22.5	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 22.5.1	N/A
SLI 22.5.1 - actions enforcing policies above	
	N/A

SLI 22.5.2 - actions enforcing policies above	N/A
	N/A
Data	
Measure 22.6	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .



QRE 22.6.1	N/A			
SLI 22.6.1 - actions enforcing policies above	N/A			
	N/A			
Data				
Measure 22.7				
QRE 22.7.1	As explained in further detail in our response to QRE 17.1.1, we deploy a range of in-app tools that proactively direct users to trusted sources. These include labels and search interventions, which surface links to trusted information providers when users engage with relevant content or queries. These tools are available in relevant EU member states and in 23 official EU languages (plus, for EEA users, Norwegian and Icelandic), as applicable. We also run targeted, localised campaigns to further promote authoritative information on specific topics.			
SLI 22.7.1 - actions enforcing policies above	N/A			

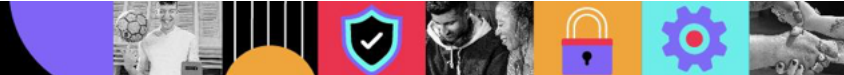


V. Empowering Users

Commitment 23

Relevant Signatories commit to provide users with the functionality to flag harmful false and/or misleading information that violates Signatories policies or terms of service.

In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 23.1	
QRE 23.1.1	We provide users with simple, intuitive ways to report/flag content in-app for any breach of our Terms of Service or Community Guidelines including for harmful, false and/or misleading information in each EU Member State and in an official language of the European Union.



	Users can access the reporting tool by: • Pressing and holding (i.e. tapping for 3 seconds) the video content and selecting the “Report” option. • Selecting the “Share” button available on the right-hand side of the video content and then selecting the “Report” option. The user is then shown categories of reporting reasons from which to select (which align with the harms our Community Guidelines seek to address), including - Misinformation. TikTok Shop. Users can also report TikTok Shop content where it breaches policies or applicable law, including for harmful, false and/or misleading information. For advertisements and commission-based content containing a TikTok Shop product anchor, users are able to select a specific “misinformation” reporting category. In line with our DSA requirements, we also continued to provide a dedicated reporting channel, and appeals process for our community in the European Union to ‘Report Illegal Content,’ enabling users to alert us to content they believe breaches the law.
Measure 23.2	
QRE 23.2.1	Reporting system To ensure the integrity of our reporting system, we deploy a combination of advanced moderation technologies and teams of human safety experts. Reported videos are initially reviewed by our automated moderation technology, which aims to identify content that violates our Community Guidelines. If a potential violation of our Community Guidelines is found, the automated review system will either pass it on to our moderation teams for further review or, if there is a high degree of confidence that the content violates our Community Guidelines, remove it automatically. Automated removal is only applied when violations are clear-cut, such as where the content contains nudity or pertains to youth safety.



To support the fair and consistent review of potentially violative content, where violations are less clear-cut, content will be passed to our human moderation teams for further review.

We have sought to make our Community Guidelines as clear and complete as possible and have put in place robust Quality Assurance processes (including steps such as review of moderation cases, flows, appeals and undertaking Root Cause Analyses).

Where content is reported under our DSA “Report Illegal Content” channel, TikTok will review the content against our Community Guidelines and where a violation is detected, the content may be removed globally. If it is not removed, our illegal content moderation team will further review the content to assess whether it is unlawful in the relevant jurisdiction - this assessment is undertaken by human moderators. If it is, access to that content will be restricted in that country.

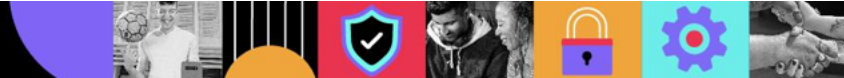
Appeals system

We are transparent with users in relation to appeals. We set out [the options](#) that may be available both to the user who reported the content and the creator of the affected content, where they disagree with the decision we have taken.

The integrity of our appeals systems is reinforced by the involvement of our dedicated moderation processes that take context and nuance into consideration when deciding whether content violates our Community Guidelines or is illegal.

To ensure consistency within this process and its overall integrity, we have sought to make our policies as clear and complete as possible and have put in place robust Quality Assurance processes (including steps such as auditing appeals and undertaking Root Cause Analyses).

If users who have submitted an appeal are still not satisfied with our decision, they can share feedback with us via the [webform](#) on TikTok.com. We continuously take user feedback into consideration to identify areas of improvement, including within the appeals process.

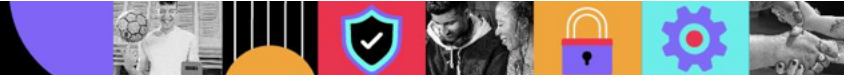


V. Empowering Users

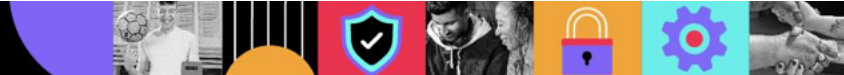
Commitment 24

Relevant Signatories commit to inform users whose content or accounts has been subject to enforcement actions (content/accounts labelled, demoted or otherwise enforced on) taken on the basis of violation of policies relevant to this section (as outlined in Measure 18.2), and provide them with the possibility to appeal against the enforcement action at issue and to handle complaints in a timely, diligent, transparent, and objective manner and to reverse the action without undue delay where the complaint is deemed to be founded.

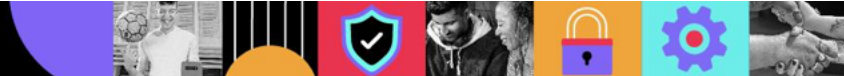
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 24.1	
QRE 24.1.1	Users in all EU member states are notified by an in-app notification in their relevant local language where the following action is taken: Removal or otherwise restriction of access to their content;



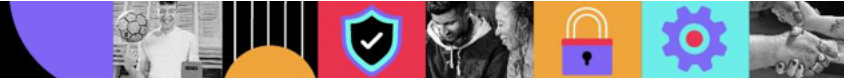
	• A ban of the account; • Restriction of their access to a feature (such as LIVE or TikTok Shop); or • Restriction of their ability to monetise. Such notifications are provided in near real time after action has been taken (i.e. generally within several seconds or up to a few minutes). Where we have taken any of these decisions, an in-app inbox notification sets out the violation deemed to have taken place, along with an option for users to “disagree” and submit an appeal. Users can submit appeals within 180 days of being notified of the decision they want to appeal. Further information, including about how to appeal a decision and other redress possibilities is set out here. All such appeals raised will be queued for review by our dedicated moderation process so as to ensure that context is adequately taken into account in reaching a determination. As mentioned above, our users have the ability to share feedback with us to the extent that they don't agree with the result of their appeal. They can do so by using the in-app function which allows them to "report a problem". We are continuously taking user feedback into consideration in order to identify areas of improvement within the appeals process.			
SLI 24.1.1 - enforcement actions	Methodology of data measurement: The number of appeals/overturns is based on the country in which the video being appealed/overtured was posted. These numbers are only related to our Misinformation, Civic and Election Integrity and Edited Media and AI-Generated Content (AIGC) policies.			
	Number of actions appealed	Metrics on results of appeals	Number of actions appealed	
List actions per member states and languages (see example table above)	Number of Appeals of videos removed for violation of misinformation policy	Number of overturns of appeals for violation of misinformation policy	Appeal success rate of videos removed for violation of misinformation policy	



Member States				
Austria	1178	916	77.8%	
Belgium	2478	2135	86.2%	
Bulgaria	1099	887	80.7%	
Croatia	518	470	90.7%	
Cyprus	287	231	80.5%	
Czechia	1928	1726	89.5%	
Denmark	867	758	87.4%	
Estonia	481	426	88.6%	
Finland	962	854	88.8%	
France	16207	13511	83.4%	
Germany	22230	17306	77.8%	
Greece	1636	1448	88.5%	
Hungary	709	637	89.8%	
Ireland	1597	1240	77.6%	
Italy	7364	6418	87.2%	



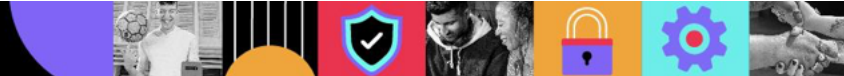
Latvia	733	666	90.9%	
Lithuania	495	434	87.7%	
Luxembourg	76	61	80.3%	
Malta	90	76	84.4%	
Netherlands	7025	6060	86.3%	
Poland	9081	7740	85.2%	
Portugal	1699	1550	91.2%	
Romania	8181	6699	81.9%	
Slovakia	739	660	89.3%	
Slovenia	392	361	92.1%	
Spain	9110	7101	77.9%	
Sweden	2695	2463	91.4%	
Iceland	66	56	84.8%	
Liechtenstein	3	2	66.7%	
Norway	906	750	82.8%	
Total EU	99,857	82,834	83.0%	



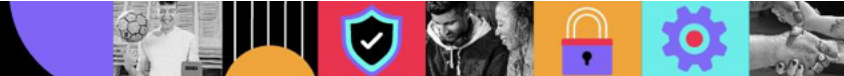
Total EEA	100,832	83,642	83.0%	
List actions per member states and languages (see example table above)	Number of appeals of videos removed for violation of Civic and Election Integrity policy	Number of overturns of appeals for violation of Civic and Election Integrity policy	Appeal success rate of videos removed for violation of Civic and Election Integrity policy	
Member States				
Austria	20	17	85%	
Belgium	59	52	88%	
Bulgaria	18	12	67%	
Croatia	7	7	100%	
Cyprus	12	9	75%	
Czechia	26	20	77%	
Denmark	13	13	100%	
Estonia	14	11	79%	
Finland	20	14	70%	
France	141	119	84%	
Germany	543	460	85%	
Greece	39	30	77%	



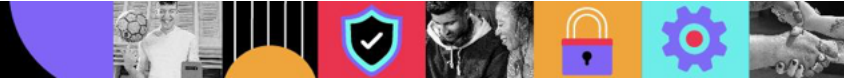
Hungary	20	17	85%	
Ireland	42	34	81%	
Italy	223	207	93%	
Latvia	17	12	71%	
Lithuania	14	11	79%	
Luxembourg	1	1	100%	
Malta	3	3	100%	
Netherlands	154	137	89%	
Poland	154	140	91%	
Portugal	66	63	95%	
Romania	833	707	85%	
Slovakia	17	13	76%	
Slovenia	7	5	71%	
Spain	144	130	90%	
Sweden	38	33	87%	
Iceland	2	1	50%	



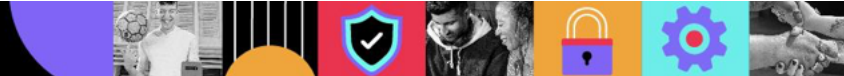
Liechtenstein	0	0	0%	
Norway	23	18	78%	
Total EU	2645	2277	86%	
Total EEA	2670	2296	86%	
List actions per member states and languages (see example table above)	Number of appeals of videos removed for violation of Edited Media and AI-Generated Content (AIGC)	Number of overturns of appeals for violation of Edited Media and AI-Generated Content (AIGC)	Appeal success rate of videos removed for violation of Edited Media and AI-Generated Content (AIGC)	
Member States				
Austria	971	859	88%	
Belgium	1411	1267	90%	
Bulgaria	496	434	88%	
Croatia	486	443	91%	
Cyprus	403	344	85%	
Czechia	737	645	88%	
Denmark	535	488	91%	
Estonia	724	625	86%	



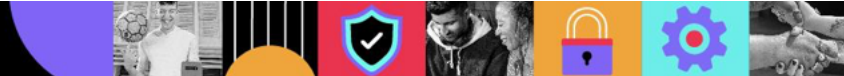
Finland	701	623	89%	
France	7398	6594	89%	
Germany	16437	14112	86%	
Greece	777	706	91%	
Hungary	428	378	88%	
Ireland	735	670	91%	
Italy	5151	4675	91%	
Latvia	950	805	85%	
Lithuania	783	674	86%	
Luxembourg	119	116	97%	
Malta	132	116	88%	
Netherlands	5261	4773	91%	
Poland	4496	3710	83%	
Portugal	1038	961	93%	
Romania	4837	4456	92%	
Slovakia	584	516	88%	



Slovenia	288	257	89%	
Spain	7276	6639	91%	
Sweden	1518	1362	90%	
Iceland	91	79	87%	
Liechtenstein	5	5	100%	
Norway	534	452	85%	
Total EU	64672	57248	89%	
Total EEA	65302	57784	88%	
List actions per member states and languages (see example table above)	Number of appeals against actions taken on TikTok Shop for violation of Misleading Functionality and Effect policy policy	Number of successful appeals against actions taken on TikTok Shop for violation of Misleading Functionality and Effect policy policy	Appeal success rate following appeals against actions taken on TikTok Shop for violation of Misleading Functionality and Effect policy policy	
Member States				
Austria	0	0	N/A	
Belgium	1	0	0%	
Czechia	0	0	N/A	
France	1,417	198	13.97%	



Germany	1,460	355	24.32%	
Ireland	39	8	20.51%	
Italy	775	137	17.68%	
Netherlands	10	2	20.00%	
Poland	14	3	21.43%	
Portugal	8	0	0.00%	
Spain	1,793	854	47.63%	
Total EU	5,517	1,557	28.22%	
Total EEA	5,517	1,557	28.22%	
List actions per member states and languages (see example table above)	Number of appeals against actions taken on TikTok Shop for violation of Edited Media and AI-Generated Content (AIGC) policy	Number of successful appeals against actions taken on TikTok Shop for violation of Edited Media and AI-Generated Content (AIGC) policy	Appeal success rate following appeals against actions taken on TikTok Shop for violation of Edited Media and AI-Generated Content (AIGC) policy	
Member States				
Austria	0	0	N/A	
Belgium	0	0	N/A	
Czechia	0	0	N/A	



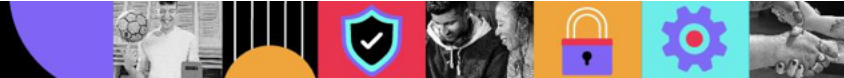
France	2,813	295	10.49%	
Germany	6,406	1,951	30.46%	
Ireland	15	0	0.00%	
Italy	1,859	198	10.65%	
Netherlands	0	0	N/A	
Poland	2	1	50.00%	
Portugal	6	0	0.00%	
Spain	3,922	581	14.81%	
Total EU	15,022	3,026	20.14%	
Total EEA	15,022	3,026	20.14%	

V. Empowering Users

Commitment 25

In order to help users of private messaging services to identify possible disinformation disseminated through such services, Relevant Signatories that provide messaging applications commit to continue to build and implement features or initiatives that empower users to think critically about information they receive and help them to determine whether it is accurate, without any weakening of encryption and with due regard to the protection of privacy.

In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A



Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 25.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 25.1.1	N/A
SLI 25.1.1	N/A
	N/A
Data	
Measure 25.2	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 25.2.1	N/A
SLI 25.2.1 - use of select tools	N/A
	N/A
Data	



VI. Empowering the research community

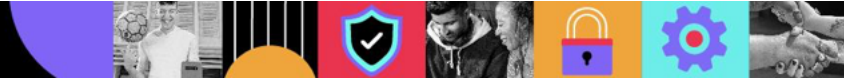
Commitments 26 - 29

VI. Empowering the research community

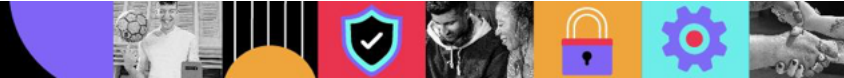
Commitment 26

Relevant Signatories commit to provide access, wherever safe and practicable, to continuous, real-time or near real-time, searchable stable access to non-personal data and anonymised, aggregated, or manifestly-made public data for research purposes on Disinformation through automated means such as APIs or other open and accessible technical solutions allowing the analysis of said data.

In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 26.1	
QRE 26.1.1	Our dedicated TikTok for Developers website hosts our Research Tools and Commercial Content APIs, which provide access to non-personal data and public data pertinent to research on disinformation (detailed in QRE 26.2 below).



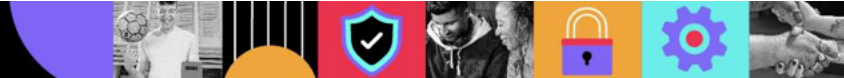
QRE 26.1.2	Information about the available data points can be found on our Research Tools page and Codebook .
SLI 26.1.1	
Data	
Measure 26.2	
QRE 26.2.1	(I) Research API Through our Research API, academic researchers from non-profit academic institutions in the US and Europe can apply to study public data about TikTok content and accounts. This public data includes comments, captions, subtitles, number of comments, shares, likes, followers and following lists, and favourites that a video receives on our platform. More information is available here. (II) Virtual Compute Environment (VCE) Through our VCE, qualifying not-for-profit researchers and academic researchers from non-profit academic institutions in the EU, intending to study youth safety can query and analyse TikTok's public data. To protect the security and privacy of our users the VCE is designed to ensure that TikTok data is processed within confined parameters. TikTok only reviews the results to ensure that there is no identifiable individual information extracted out of the platform. All aggregated results will be shared as a downloadable link to the approved primary researcher's email. (III) Commercial Content API Through our Commercial Content API, qualifying researchers and professionals, who can be located in any country, can request public data about searches on ads and other commercial contents including ads, ad and advertiser metadata, and targeting information. (IV) Commercial Content Library



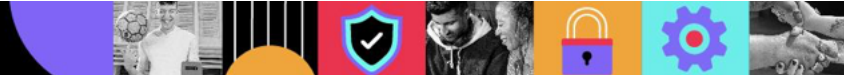
	TikTok's Commercial Content Library (CCL) is a repository of ads and other commercial content posted on TikTok. There are two main sub-libraries within the CCL: • Ad Library: This library features ads that we're paid to display to people, including those that aren't currently active or have been paused by the advertisers. • Other commercial content: This library features content that we're not paid to display, including content that promotes a brand, product, or service. The CCL currently includes information on ads available to users in the European Economic Area (EEA), Switzerland, the U.K. and Turkey.
QRE 26.2.2	See our response to QRE 26.2.1 which sets out the scope of public data available through each tool.
QRE 26.2.3	We make detailed information available to applicants about our Research Tools (Research API and VCE) and Commercial Content API, through our dedicated TikTok for Developers website, including on what data is made available and how to apply for access. Researchers who wish to utilise the Research API or VCE must submit an application form with information about their eligibility and research project. The application criteria for our Research Tools (Research API and VCE) and Commercial Content API is research topic agnostic and clearly set out in our dedicated TikTok for Developers website. Once an application has been approved for access to our Research Tools, we provide step-by-step instructions for researchers on how to access research data, how to comply with the security steps, and how to run queries on the data, through our Getting Started guides for Research API and VCE Similarly with the Commercial Content API, we provide participants with detailed information on how to query ad data and fetch public advertiser data.



SLI 26.2.1	Research Tools, Commercial Content API, and the Commercial Content Library During this reporting period we received: 282 applications to access TikTok's Research Tools (Research API and VCE) from researchers in the EU and EEA.					
	Number of applications received for Research Tools	Number of applications accepted for Research Tools	Number of applications rejected for Research Tools	Number of applications received for TikTok Commercial Content API	Number of applications accepted for TikTok Commercial Content API	Number of applications rejected for TikTok Commercial Content API
Austria	10	8	2	2	2	0
Belgium	9	7	2	12	12	0
Bulgaria	2	1	0	0	0	0
Croatia	0	0	0	1	1	0
Cyprus	2	1	1	1	1	0
Czechia	3	1	2	6	5	0
Denmark	7	7	0	82	82	0
Estonia	0	0	0	1	1	0
Finland	2	2	0	3	3	0
France	36	17	18	37	37	0
Germany	64	41	17	20	16	0
Greece	2	1	0	0	0	0



Hungary	6	2	4	3	3	0
Ireland	8	7	1	2	2	0
Italy	25	19	3	15	15	0
Latvia	1	1	0	2	2	0
Lithuania	1	1	0	2	2	0
Luxembourg	1	1	0	1	1	0
Malta	0	0	0	0	0	0
Netherlands	18	9	7	8	8	0
Poland	14	11	3	8	8	0
Portugal	7	4	3	2	2	0
Romania	11	7	2	6	6	0
Slovakia	1	0	1	4	4	0
Slovenia	0	0	0	1	1	0
Spain	37	25	11	17	17	0
Sweden	8	7	1	6	6	0
Iceland	0	0	0	0	0	0
Liechtenstein	0	0	0	0	0	0
Norway	7	6	0	2	2	0
EU Level	275	180	78	222	222	0

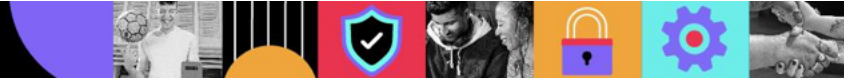


EEA Level	282	186	78	244	244	0
Measure 26.3						
QRE 26.3.1	We have a dedicated support form where researchers can provide feedback about their experience. In H1 2026, we refreshed our changelog, which is available on TikTok for Developers. TikTok provides technical support via a support form on the TikTok for Developers site as well as over email. Any reported bugs or outages are evaluated by our technical teams. We also launched a quarterly Research Tools Office Hours session. This is targeted at researchers who are already onboarded to the Research Tools and who reach out with technical support questions (by invitation). Our Office Hours session is in addition to our ongoing work to conduct demos and engage with researchers at conferences.					

VI. Empowering the research community	
Commitment 27	
Relevant Signatories commit to provide vetted researchers with access to data necessary to undertake research on Disinformation by developing, funding, and cooperating with an independent, third-party body that can vet researchers and research proposals.	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A

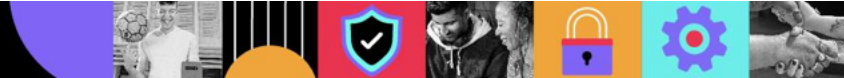


Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 27.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 27.1.1	N/A
Measure 27.2	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 27.2.1	N/A
Measure 27.3	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 27.3.1	N/A
SLI 27.3.1 - research projects vetted by the independent third-party body	N/A
	N/A
Data	N/A

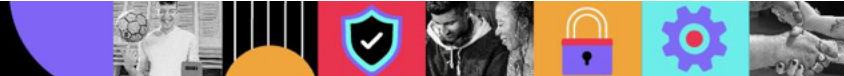


Measure 27.4	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 27.4.1	N/A

VI. Empowering the research community	
Commitment 28	
Relevant Signatories commit to support good faith research into Disinformation that involves their services.	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 28.1	



QRE 28.1.1	As set out above, TikTok is committed to facilitating research through our Research Tools, Commercial Content APIs and Commercial Content Library, full details of which are available on our TikTok for Developers and Commercial Content Library websites. TikTok teams and personnel also regularly participate in research-focused events. During this reporting period, we conducted office hours and demos. The transferable nature of the skills provided are such that they could be broadly applied to disinformation research. • We launched a quarterly series of Research Tools Office Hours for researchers who are onboarded to our Tools and require technical support. We hosted two sessions in February and June. • We conducted the demos for the following researcher groups: Johns Hopkins University (January), University of Lusofona (March), University of Coimbra (April), York University (April), University of Pennsylvania (April), and invited researchers in Slovakia (June).
Measure 28.2	
QRE 28.2.1	We have a dedicated TikTok for Developers website which hosts our Research Tools and Commercial Content APIs. With the Research API, researchers can access: • Public account data, such as user profiles, followers and following lists, liked videos, pinned videos and reposted videos. • Public content data, such as comments, captions, subtitles, and number of comments, shares and likes that a video receives. Through the VCE, qualifying not-for-profit researchers in the EU can access and analyse TikTok's public data, including public U18 data, in a secure environment that is subject to strict security controls. TikTok makes information about ads and other commercial content available through its Commercial Content Library (CCL). The CCL is a publicly accessible, searchable repository that



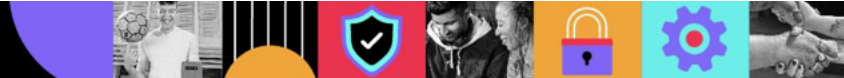
	allows users to view information about paid advertisements and other commercial content on TikTok, including relevant advertiser and ad information, such as the advertising creative, dates the ad ran, main parameters used for targeting (e.g. age, gender), number of people who were served the ad, and more. Separately, the Commercial Content API provides programmatic access to public data available through the CCL, allowing users to query and analyse this information through an API.
Measure 28.3	
QRE 28.3.1	As of April 2026, over 80 successful publications were published using data from our Research Tools.
Measure 28.4	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 28.4.1	N/A

VI. Empowering the research community

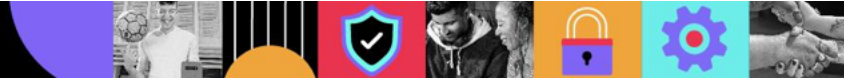
Commitment 29

Relevant Signatories commit to conduct research based on transparent methodology and ethical standards, as well as to share datasets, research findings and methodologies with relevant audiences.

In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 29.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 29.1.1	N/A
QRE 29.1.2	N/A
QRE 29.1.3	N/A



SLI 29.1.1 - reach of stakeholders or citizens informed about the outcome of research projects	N/A
	N/A
Data	N/A
Measure 29.2	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 29.2.1	N/A
QRE 29.2.2	N/A
QRE 29.2.3	N/A
SLI 29.2.1	N/A
	N/A
Data	N/A
Measure 29.3	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 29.3.1	N/A
SLI 29.3.1 - reach of stakeholders or citizens informed about the outcome of research projects	N/A
	N/A
Data	N/A



VII. Empowering the fact-checking community

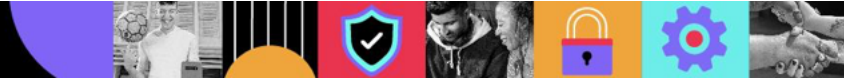
Commitments 30 - 33

VII. Empowering the fact-checking community

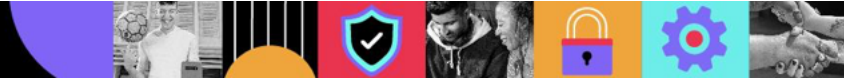
Commitment 30

Relevant Signatories commit to establish a framework for transparent, structured, open, financially sustainable, and non-discriminatory cooperation between them and the EU fact-checking community regarding resources and support made available to fact-checkers

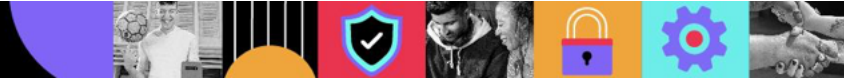
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 30.1	
QRE 30.1.1	Within Europe, we work with 13 fact-checking partners who provide fact-checking coverage in 23 EEA languages, including at least one official language of every EU Member State, and additional languages including Georgian, Russian, Turkish, Ukrainian, Albanian and Serbian. Our partners have teams of fact-checkers who review and verify reported content. Our Integrity and Authenticity moderators then use that independent feedback to take action and where



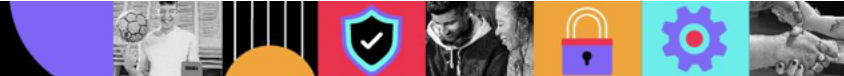
	appropriate, remove or make ineligible for recommendation false or misleading content or label unverified content. Our agreements with our partners are standardised, meaning the agreements are based on our template master services agreements and consistent with common standards and conditions. We reviewed and updated our template standard agreements as part of our annual contract renewal process. The terms of the agreements describe: • The service the fact-checking partner will provide, namely, that their team of fact checkers review, assess and rate video content uploaded to their fact-checking queue, and will provide regular pro-active Insights Reports about general misinformation trends observed on our platform and across the industry generally, including new/changing industry or market trends, events or topics that generate particular misinformation or disinformation. • The expected results e.g. the fact-checkers' advice on whether the content may be or contain misinformation and rate it using our classification categories. • Notifications of proactive flagging of potentially harmful misinformation from our partners. • The languages in which they will provide fact-checking services. • The ability to request temporary coverage regarding additional languages or support on ad hoc additional projects. • All other key terms including the applicable term and fees and payment arrangements.
QRE 30.1.2	We currently have agreements with 13 IFCN accredited fact-checking partners across the EU, EEA, and wider Europe: 1. Agence France-Presse (AFP) 2. Deutsche Presse-Agentur (dpa) 3. Demagog 4. Facta 5. Geofacts 6. Faktograf 7. Internews Kosova (Kallxo) 8. Lead Stories 9. Newtral



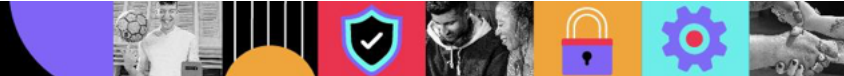
	10. Poligrafo 11. Reuters 12. Science Feedback - For advertising-related fact-checking partnerships, please refer to Chapter 2. 13. Teyit We put in place temporary agreements with these fact-checking partners to provide additional EU language coverage during high risk events like elections or an unfolding crisis. During this period, we also engaged an existing fact-checking partner to provide temporary fact-checking coverage during the Malta general election in both Maltese and English. Globally, we have 20 IFCN-accredited fact-checking partners and we keep users updated here.
QRE 30.1.3	We have fact-checking coverage in 23 official EEA languages: Bulgarian, Croatian, Czech, Danish, Dutch, English, Estonian, Finnish, French, German, Greek, Hungarian, Italian, Latvian, Lithuanian, Norwegian, Polish, Portuguese, Romanian, Slovak, Slovenian, Spanish and Swedish. We have fact-checking coverage in a number of other European languages or languages used in Europe which affect European users, including Georgian, Russian, Turkish, and Ukrainian. In terms of global fact-checking initiatives, we currently cover more than 60 languages and 130 markets across the world, thereby improving the overall integrity of the service and benefiting European users. In order to effectively scale the feedback provided by our fact-checkers globally, we have implemented the measures listed below. • Insights reports. Our fact-checking partners provide regular reports identifying general misinformation trends observed on our platform and across the industry generally, including new/changing industry or market trends, events or topics that generated particular misinformation or disinformation. • Proactive detection by our fact-checking partners. Our fact-checking partners are authorised to proactively identify content that may constitute harmful misinformation on our platform, which our moderators assess against our Community Guidelines, and



	suggest prominent misinformation that is circulating online that may benefit from verification. • Moderation guidelines. Where relevant, we create guidelines and trending topic reminders for our moderators which are informed by previous fact checking assessments. This helps our teams leverage the insights from our fact-checking partners and supports swift and accurate decisions on flagged content regardless of the language in which the original claim was made.
SLI 30.1.1 - Member States and languages covered by agreements with the fact-checking organisations	
Austria	Fact-checking coverage implemented
Belgium	Fact-checking coverage implemented
Bulgaria	Fact-checking coverage implemented
Croatia	Fact-checking coverage implemented
Cyprus	Fact-checking coverage implemented
Czechia	Fact-checking coverage implemented
Denmark	Fact-checking coverage implemented
Estonia	Fact-checking coverage implemented
Finland	Fact-checking coverage implemented
France	Fact-checking coverage implemented

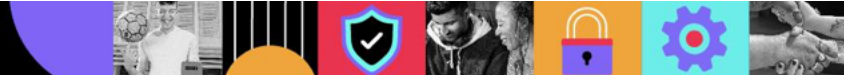


Germany	Fact-checking coverage implemented
Greece	Fact-checking coverage implemented
Hungary	Fact-checking coverage implemented
Ireland	Fact-checking coverage implemented
Italy	Fact-checking coverage implemented
Latvia	Fact-checking coverage implemented
Lithuania	Fact-checking coverage implemented
Luxembourg	Fact-checking coverage implemented
Malta	No permanent fact-checking coverage. We can, and have, put in place temporary agreements with fact-checking partners to provide additional EU language coverage during high risk events like elections or an unfolding crisis.
Netherlands	Fact-checking coverage implemented
Poland	Fact-checking coverage implemented
Portugal	Fact-checking coverage implemented
Romania	Fact-checking coverage implemented
Slovakia	Fact-checking coverage implemented
Slovenia	Fact-checking coverage implemented



Spain	Fact-checking coverage implemented
Sweden	Fact-checking coverage implemented
Iceland	No permanent fact-checking coverage. We can, and have, put in place temporary agreements with fact-checking partners to provide additional EU language coverage during high risk events like elections or an unfolding crisis.
Liechtenstein	Fact-checking coverage implemented
Norway	Fact-checking coverage implemented
Total EU	22 languages
Total EEA	23 languages

Measure 30.2	
QRE 30.2.1	Our agreements with our fact-checking partners are standardised, meaning the agreements are based on our template master services agreements and consistent with common standards and conditions. These agreements, as with all of our agreements, must meet the ethical and professional standards we set internally including containing anti-bribery and corruption provisions. Our partners are compensated in a fair, transparent way based on the work done by them using standardised rates. Our fact-checking partners then invoice us on a monthly basis based on work done. All of our fact-checking partners are independent organisations, which are certified through the non-partisan IFCN. Our agreements with them explicitly state that the fact-checkers are non-exclusive, independent contractors of TikTok who retain editorial independence in relation to the fact-checking, and that the services shall be performed in a professional manner and in line with the highest standards in the industry. Our processes are also set up to ensure our fact-checking partners' independence. Our partners access flagged content through a tool

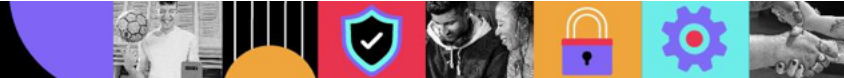


	dedicated for their use and provide their assessment of the accuracy of the content by providing a rating. Fact-checkers will do so independently from us, and their review may include calling sources, consulting public data or authenticating videos and images. To facilitate transparency and openness with our fact-checking partners, we regularly meet with them to gather feedback.
QRE 30.2.2	We meet regularly with our fact-checking partners and have an ongoing dialogue with them about how our partnership is working and evolving.
QRE 30.2.3	This provision is not relevant to TikTok, only to fact-checking organisations.
Measure 30.3	
QRE 30.3.1	Given our fact-checking partners are all IFCN-accredited, our fact-checking partners already engage in some informal cross-border collaboration through that network. In June 2026, TikTok both sponsored and sent a delegation of attendees to the Global Fact conference in Vilnius, Lithuania.
Measure 30.4	
QRE 30.4.1	TikTok is establishing a communication channel with EDMO to maintain regular dialogue. We also maintain dialogue with EFCSN on these and other issues.

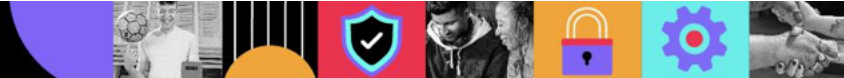
VII. Empowering the fact-checking community

Commitment 31

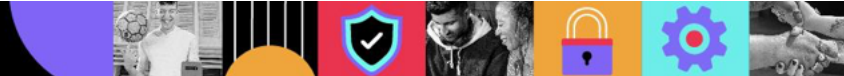
Relevant Signatories commit to integrate, showcase, or otherwise consistently use fact-checkers' work in their platforms' services, processes, and contents; with full coverage of all Member States and languages.



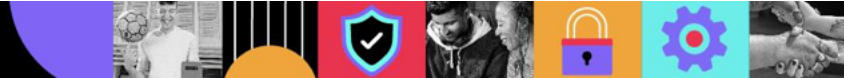
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 31.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
Measure 31.2	
QRE 31.1.1	We place considerable emphasis on proactive detection and automated moderation technology to action violative content. For example, "multi-modal LLMs" can perform complex, highly specific tasks related to visual content. We can use this technology to make misinformation moderation easier by extracting specific misinformation "claims" from videos for moderators to assess directly or route to our fact-checking partners. Our Integrity and Authenticity moderators receive direct access to our fact-checking partners who help assess the accuracy of content. We also use fact-checking feedback to provide additional context to users about certain content. As mentioned, when our fact checking partners conclude that the fact-check is inconclusive or content is not able to be confirmed, we inform viewers via a banner when we identify a video with unverified content in an effort to raise users' awareness about the credibility of the content and to reduce sharing. The video may also



	become ineligible for recommendation into anyone's For You feed to limit the spread of potentially misleading information.			
SLI 31.1.1 - use of fact-checks	Methodology of data measurement: The number of fact checked videos is based on the number of videos that have been reviewed by one of our fact-checking partners in the relevant territory.			
	Number of fact-checked articles published			
List actions per member states and languages (see example table above)	Number of fact checked videos (tasks)			
Member States				
Austria	40			
Belgium	78			
Bulgaria	140			
Croatia	4			
Cyprus	14			
Czechia	101			
Denmark	51			
Estonia	100			
Finland	34			

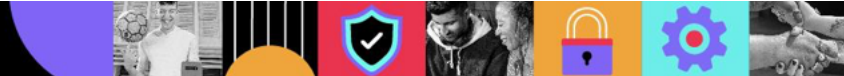


France	1659			
Germany	318			
Greece	51			
Hungary	12			
Ireland	36			
Italy	662			
Latvia	5			
Lithuania	3			
Luxembourg	1			
Malta	4			
Netherlands	343			
Poland	1926			
Portugal	263			
Romania	73			
Slovakia	108			
Slovenia	6			
Spain	726			
Sweden	30			
Iceland	1			

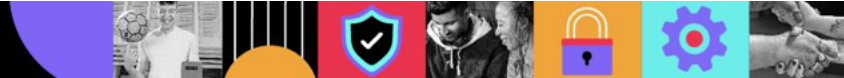


Liechtenstein	0			
Norway	44			
Total EU	6788			
Total EEA	6833			

SLI 31.1.2 - impact of actions taken	Methodology of data measurement: The number of videos removed as a result of a fact-checking assessment and the number of videos removed because of policy guidelines and known misinformation trends. These metrics correspond to the numbers of removals under the misinformation policy since all of its enforcement are based on the policy guidelines and known misinformation trends.		
	N/A		
List actions per member states and languages (see example table above)	Number of videos removed as a result of a fact checking assessment	Number of videos removed under Misinformation policy	
Member States			
Austria	13	11308	
Belgium	22	15105	
Bulgaria	11	42197	
Croatia	14	9243	
Cyprus	3	3222	
Czechia	6	10470	

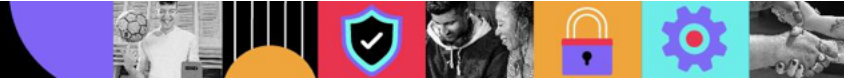


Denmark	3	15725	
Estonia	6	1963	
Finland	6	5307	
France	251	100530	
Germany	122	113943	
Greece	2	8404	
Hungary	2	16159	
Ireland	7	9383	
Italy	142	41780	
Latvia	0	3648	
Lithuania	1	3146	
Luxembourg	0	897	
Malta	0	501	
Netherlands	54	80988	
Poland	392	63967	
Portugal	21	11667	
Romania	10	47144	
Slovakia	18	4076	
Slovenia	1	1273	

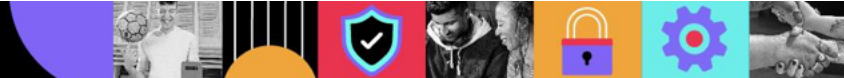


Spain	97	56000	
Sweden	0	13090	
Iceland	1	345	
Liechtenstein	0	5	
Norway	4	17075	
Total EU	1204	691136	
Total EEA	1209	708561	

SLI 31.1.3 – Quantitative information used for contextualisation for the SLIs 31.1.1 / 31.1.2	Methodology of data measurement: The metric we have provided demonstrates the % of videos which have been removed as a result of the fact-checking assessment, in comparison to the total number of videos removed because of violation of our harmful misinformation policy.
List actions per member states and languages (see example table above)	Videos removed as a result of a fact checking assessment as a percentage of total number of videos removed due to violation of harmful misinformation policy
Austria	0.11%
Belgium	0.15%
Bulgaria	0.03%



Croatia	0.15%
Cyprus	0.09%
Czechia	0.06%
Denmark	0.02%
Estonia	0.31%
Finland	0.11%
France	0.25%
Germany	0.11%
Greece	0.02%
Hungary	0.01%
Ireland	0.07%
Italy	0.34%
Latvia	0.00%
Lithuania	0.03%
Luxembourg	0.00%
Malta	0.00%
Netherlands	0.07%
Poland	0.61%
Portugal	0.18%



Romania	0.02%
Slovakia	0.44%
Slovenia	0.08%
Spain	0.17%
Sweden	0.00%
Iceland	0.29%
Liechtenstein	0.00%
Norway	0.02%
Total EU	0.17%
Total EEA	0.17%

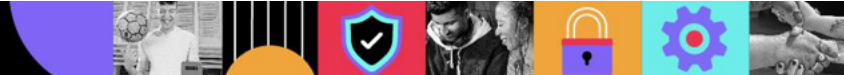
Measure 31.3	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 31.3.1	N/A
Measure 31.4	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document .
QRE 31.4.1	N/A

VII. Empowering the fact-checking community

Commitment 32

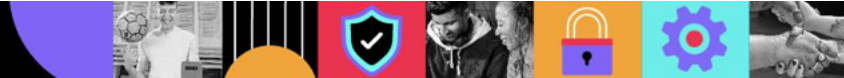
Relevant Signatories commit to provide fact-checkers with prompt, and whenever possible automated, access to information that is pertinent to help them to maximise the quality and impact of fact-checking, as defined in a framework to be designed in coordination with EDMO and an elected body representative of the independent European fact-checking organisations.

In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 32.1	
Measure 32.2	
QRE 32.1.1	Our fact-checking partners access content which has been flagged for review through a content review tool made available for their exclusive use. The dashboard shows our fact-checkers certain quantitative information about the services they provide, including the number of videos



	queued for assessment at any one time, as well as the time the review has taken. Fact-checkers can also use the dashboard to see the rating they applied to videos they have previously assessed.		
SLI 32.1.1 - use of the interfaces and other tools	Methodology of data measurement: N/A. As mentioned in our response to QRE 32.1.1, the dashboard we currently share with our partners only contains high level quantitative information about the services they provide, including the number of videos queued for assessment at any one time, as well as the time the review has taken. We are continuing to work with our fact checking partners to understand what further data it would be helpful for us to share with them.		
Data			
Measure 32.3			
QRE 32.3.1	We remain a participant in the Code Taskforce. Separately, we maintain regular dialogue with our fact-checking partners regarding the operation and evolution of our partnerships.		

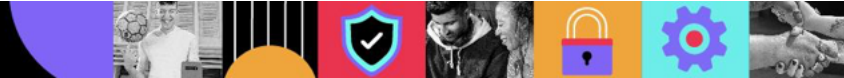
VII. Empowering the fact-checking community
Commitment 33 Relevant Signatories (i.e. fact-checking organisations) commit to operate on the basis of strict ethical and transparency rules, and to protect their independence.



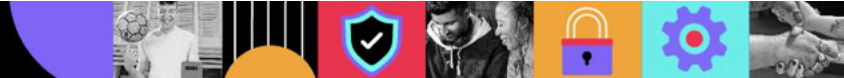
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 33.1	TikTok did not subscribe to this measure as outlined in the January 2025 Subscription Document
QRE 33.1.1	N/A
SLI 33.1.1 - number of European fact-checkers that are IFCN-certified	N/A



VIII. Transparency Centre Commitments 34 - 36

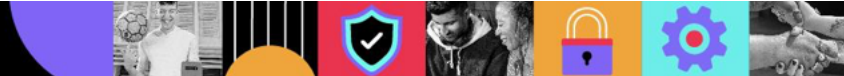


VIII. Transparency Centre	
Commitment 34 To ensure transparency and accountability around the implementation of this Code, Relevant Signatories commit to set up and maintain a publicly available common Transparency Centre website.	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 34.1	
Measure 34.2	
Measure 34.3	
Measure 34.4	



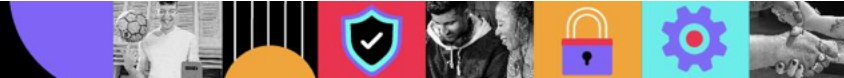
Measure 34.5	
--------------	--

VIII. Transparency Centre	
Commitment 35	
Signatories commit to ensure that the Transparency Centre contains all the relevant information related to the implementation of the Code's Commitments and Measures and that this information is presented in an easy-to-understand manner, per service, and is easily searchable.	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 35.1	
Measure 35.2	



Measure 35.3	
Measure 35.4	
Measure 35.5	
Measure 35.6	

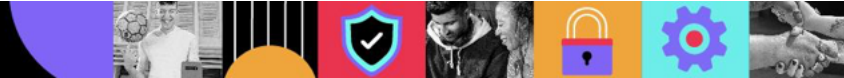
VIII. Transparency Centre	
Commitment 36 Signatories commit to updating the relevant information contained in the Transparency Centre in a timely and complete manner.	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	Yes
If yes, which further implementation measures do you plan to put in place in the next 6 months?	In line with the commitments set out in the Code, we plan to upload reports and relevant information to the Transparency Centre.



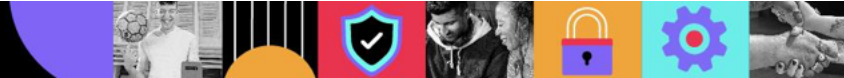
Measure 36.1	
Measure 36.2	
Measure 36.3	
QRE 36.1.1 (for the Commitments 34-36)	N/A
QRE 36.1.2 (for the Commitments 34-36)	No changes since last report. VOST maintains the Transparency Center and we continue to support this.
SLI 36.1.1	We worked with the administrator of the Transparency Centre to develop the below metrics for this SLI.
Data	Between January 1 and June 30 2026, our signatory profile was visited 2,585 times, and our signatory reports were downloaded 9,422 times. The Transparency Centre Webpage was visited 46,705 times overall.



IX. Permanent Task-Force Commitment 37



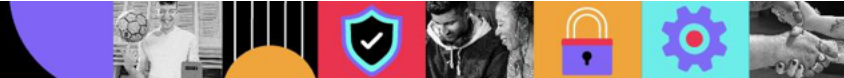
IX. Permanent Task-Force	
Commitment 37 Signatories commit to participate in the permanent Task-force. The Task-force includes the Signatories of the Code and representatives from EDMO and ERGA. It is chaired by the European Commission, and includes representatives of the European External Action Service (EEAS). The Task-force can also invite relevant experts as observers to support its work. Decisions of the Task-force are made by consensus.	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 37.1	We have attended all Plenary meetings and continue to participate in the Task-force, Plenaries and working groups.
Measure 37.2	We continue to participate in all relevant workstreams of the Task-force.
Measure 37.3	We continue to participate in all relevant workstreams of the Task-force.



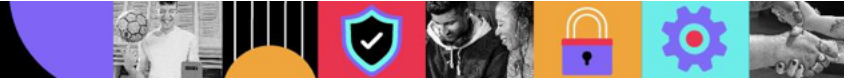
Measure 37.4	We continue to participate in all relevant workstreams of the Task-force, including the relevant subgroups.
Measure 37.5	We continue to participate in all relevant workstreams of the Task-force, including attending meetings and relevant events as required by the Code.
Measure 37.6	
QRE 37.6.1	We have meaningfully engaged in the Task-force and all of its working groups by attending and participating in meetings and engaging in any relevant discussions, in particular regarding elections and further developing/activating the Rapid Response System (RRS). We will continue to engage in the Task-force and all of its working groups and Election subgroups.



X. Monitoring of Code Commitment 38 - 43

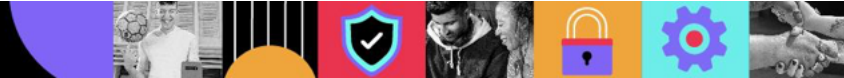


X. Monitoring of Code	
Commitment 38 The Signatories commit to dedicate adequate financial and human resources and put in place appropriate internal processes to ensure the implementation of their commitments under the Code.	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 38.1	



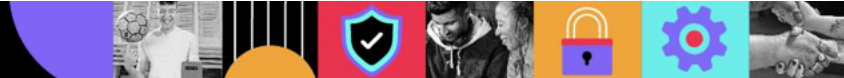
QRE 38.1.1	TikTok will continue to have appropriate resources in place to meet our commitments and compliance. Given the breadth of the Code and the commitments therein, our work spans multiple teams, including Trust and Safety, Legal, Business Integrity, Governance and Experience, Product and Public Policy. Teams across the globe are deployed to ensure that we meet our commitments and compliance, including achieving full coverage across the Member States and languages of the EU.
------------	---

X. Monitoring of Code	
Commitment 39 Signatories commit to provide to the European Commission, within 1 month after the end of the implementation period (6 months after this Code's signature) the baseline reports as set out in the Preamble..	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document .
If yes, list these implementation measures here [short bullet points].	N/A
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A



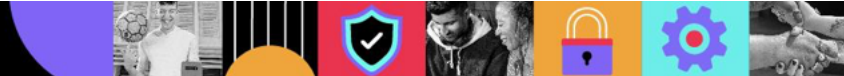
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
---	-----

X. Monitoring of Code	
Commitment 40 Signatories commit to provide regular reporting on Service Level Indicators (SLIs) and Qualitative Reporting Elements (QREs). The reports and data provided should allow for a thorough assessment of the extent of the implementation of the Code's Commitments and Measures by each Signatory, service and at Member State level.	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	Yes
If yes, list these implementation measures here [short bullet points].	We have reported on the SLIs and QREs relevant to the Commitments we signed-up to within this report.
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 40.1	We publish a report, detailing the implementation of the commitments and measures (including QREs and SLIs) we have signed up to under the Code, every 6 months.



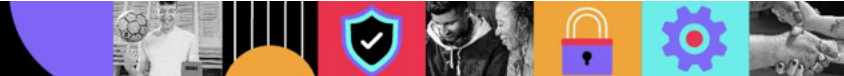
Measure 40.2	We publish a report, detailing the implementation of the commitments and measures (including QREs and SLIs) we have signed up to under the Code, every 6 months.
Measure 40.3	We publish our reports online. All reports published under the Code can be accessed through the Transparency Center .
Measure 40.4	We continue to work with the Taskforce, as applicable.
Measure 40.5	We continue to engage in the work of the Taskforce, as applicable.
Measure 40.6	We continue to work and cooperate with the EC, as applicable.

X. Monitoring of Code	
Commitment 41 Signatories commit to work within the Task-force towards developing Structural Indicators, and publish a first set of them within 9 months from the signature of this Code; and to publish an initial measurement alongside their first full report. To achieve this goal, Signatories commit to support their implementation, including the testing and adapting of the initial set of Structural Indicators agreed in this Code. This, in order to assess the effectiveness of the Code in reducing the spread of online disinformation for each of the relevant Signatories, and for the entire online ecosystem in the EU and at Member State level. Signatories will collaborate with relevant actors in that regard, including ERGA and EDMO.	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	No, pending further updates from the Commission.
If yes, list these implementation measures here [short bullet points].	N/A



Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A
Measure 41.1	N/A
Measure 41.2	N/A
Measure 41.3	N/A

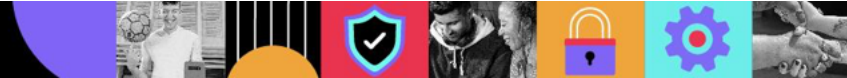
X. Monitoring of Code	
Commitment 42 Relevant Signatories commit to provide, in special situations like elections or crisis, upon request of the European Commission, proportionate and appropriate information and data, including ad-hoc specific reports and specific chapters within the regular monitoring, in accordance with the rapid response system established by the Taskforce..	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	Yes
If yes, list these implementation measures here [short bullet points].	We have been an active participant in the Crisis Response working group, which resulted in the implementation of the Rapid Response System being developed/ activated for elections. We



	have also published Crisis Reports specific to the War in Ukraine, the Israel-Hamas conflict, and 2026 elections reports on the Portugal Presidential Election, Slovenia Parliamentary Election, Denmark General Election, Bulgaria Parliamentary Election, Hungary Parliamentary Election, Cyprus Parliamentary Election, and Malta General Election along with this report.
Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A

X. Monitoring of Code	
Commitment 43 Signatories commit to produce reports and provide data following the harmonised reporting templates and refined methodology for reporting and data disclosure, as agreed in the Task-force.	
In line with this commitment, did you deploy new implementation measures (e.g. changes to your terms of service, new tools, new policies, etc)? [Yes/No]	Yes
If yes, list these implementation measures here [short bullet points].	• Participated in the monitoring and reporting working group. • Published transparency report in March 2026.

Do you plan to put further implementation measures in place in the next 6 months to substantially improve the maturity of the implementation of this commitment? [Yes/No]	N/A
If yes, which further implementation measures do you plan to put in place in the next 6 months?	N/A



Reporting on the service's response during a crisis

War of aggression by Russia on Ukraine

Threats observed or anticipated at time of reporting: [suggested character limit 2000 characters].

Since the start of the war of aggression by Russia on Ukraine in February 2022 (the "War in Ukraine"), we have observed false or unverified claims about specific attacks and events, the development or use of weapons, the involvement of particular countries, and military activities such as troop movements. We have also seen misleadingly repurposed footage, including clips from video games, AI-generated content, or unrelated past events presented as current. We have remained alert to the spread of harmful misinformation and covert influence operations ("CIO").

Mitigations in place at time of reporting: [suggested character limit: 2000 characters].

(I) Upholding TikTok's Community Guidelines

We are continuing to enforce our [policies](#) against [violence](#), [hate](#), and [harmful misinformation](#) by taking action to remove violative content and accounts. We use a combination of advanced moderation technologies and teams of human safety experts to identify, review, and action content that violates our policies.

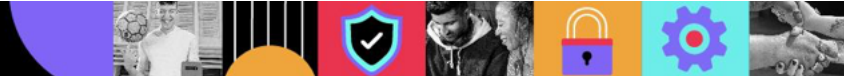
Automated Review

We place considerable emphasis on proactive detection to remove violative content and reduce exposure to potentially distressing content for our human moderators. Before content is posted to our platform, it's reviewed by automated moderation technologies, which identify content or behaviour that may violate our policies or For You feed eligibility standards, or that may require age-restriction or other actions. While undergoing this review, the content is visible only to the uploader.

If our automated moderation technology identifies content that is a potential violation, it will either take action against the content or flag it for further review by our human moderation teams. In line with our safeguards to help ensure accurate decisions are made, automated removal is applied when violations are the most clear-cut.

Some of the methods and technologies that support these efforts include:

- **Vision-based:** Computer vision models can identify objects that violate our Community Guidelines, such as weapons or hate symbols.
- **Audio-based:** Audio clips are reviewed for violations of our policies, supported by a dedicated audio bank and "classifiers" that help us detect audios that are similar or modified to previous violations.
- **Text-based:** Detection models review written content like comments or hashtags,. Artificial Intelligence (AI) that can interpret the context surrounding content—helps us identify violations that are context-dependent, such as words that can be used in a hateful way but may not violate



our policies by themselves.

- **Similarity-based:** "Similarity detection systems" enable us to not only catch identical or highly similar versions of violative content, but other types of content that share key contextual similarities and may require additional review.
- **Activity-based:** Technologies that look at how accounts are being operated help us disrupt deceptive activities like bot accounts, spam, or attempts to artificially inflate engagement through fake likes or follow attempts.
- **LLMs:** We use multimodal LLMs to help moderate content faster and more consistently at scale, from taking automated action on activity like fake engagement, to empowering teams with better moderation tools and risk insights.
- We work with external groups, for example [Tech Against Terrorism](#) in the context of [violent extremist](#) content, who help us to more quickly detect and remove violative content that has already been identified off the platform.

Scaling human expertise

Human insight plays a crucial role in the content moderation process, from our community or external experts, to our own safety professionals. Our global teams of human safety experts speak more than 60 languages and dialects, including Russian and Ukrainian. We strive to promote a caring working environment for all TikTok employees, and especially for trust and safety professionals. We use an evidence-based approach to develop programmes and resources that support their psychological well-being, including for Trust & Safety personnel working on mis & disinformation.

In H1 2026, we removed 9,967 videos in relation to the War in Ukraine, which violated our misinformation policies.

(II) Leveraging our Global Fact-Checking Program

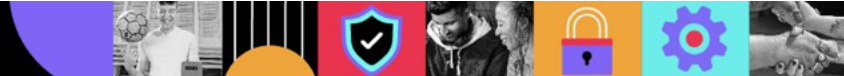
We use a layered approach to detect harmful misinformation that violates our Community Guidelines, with our Global Fact-Checking Programme playing a key role. We assess the accuracy of harmful or hard-to-verify claims by partnering with 20 [IFCN-accredited](#) fact-checking organisations who support over 60 languages on TikTok, including Russian, Ukrainian, and Belarusian. We also collaborate with certain fact-checking partners to receive advance warning of emerging misinformation narratives. This helps facilitate proactive responses against high-harm trends and ensures that our Integrity and Authenticity moderators have up-to-date guidance.

To limit the spread of potentially misleading information, we apply [warning labels](#) and prompt users to reconsider sharing content about unfolding or emergency events that have been reviewed by fact-checkers but cannot be verified—referred to as “unverified content.” Recognising that the situation around the Conflict can change rapidly, we have put in place a process allowing our fact-checking partners to quickly update us if claims previously marked as “unverified” are later verified or clarified with additional context.

(III) Disruption of CIOs

TikTok’s [integrity and authenticity policies](#) do not allow deceptive behaviour that may cause harm to our community or society at large. We have specifically-trained teams on high alert to investigate, disrupt and remove CIO networks from our platform and we provide regular updates in our dedicated [CIO transparency reports](#). For advertising-related CIO measures, please refer to Chapter 2.

Between January and June 2026, we took action to remove a total of four CIO networks targeting discourse related to the War in Ukraine.



(IV) Spread of harmful misinformation

TikTok takes a multi-faceted approach to tackling the spread of harmful misinformation, regardless of intent. This includes our [Integrity and Authenticity policies](#), as well as our products, operational practices, and external partnerships with fact-checkers, media literacy organisations, and researchers.

We support our Integrity and Authenticity moderators with detailed misinformation policy guidance, enhanced training, and direct access to our IFCN-accredited fact-checking partners, who help assess the accuracy of content.

We continue to take swift action against misinformation, conspiracy theories, fake engagement, and fake accounts relating to the War in Ukraine.

(V) Mitigating the risk of monetisation of harmful misinformation

Political advertising has been prohibited on our platform for many years, but as an additional risk mitigation measure against the risk of profiteering from the War in Ukraine we prohibit Russian-based advertisers from outbound targeting of EU markets. We also suspended TikTok in the Donetsk and Luhansk regions.

(VI) Localised media literacy campaigns

Proactive measures aimed at improving our users' digital literacy are vital, and we recognise the importance of increasing the prominence of authoritative information. We have 17 localised media literacy campaigns addressing disinformation related to the War in Ukraine in Austria, Bosnia, Bulgaria, Czechia, Croatia, Estonia, Germany, Hungary, Latvia, Lithuania, Montenegro, Poland, Romania, Serbia, Slovakia, Slovenia, and Ukraine, in close collaboration with our fact-checking partners. Users searching for keywords relating to the War in Ukraine are directed to tips, prepared in partnership with our fact-checking partners, to help users identify misinformation and prevent its spread on the platform.

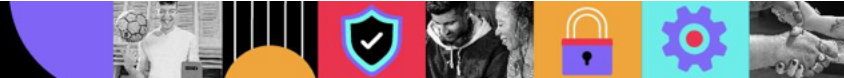
(VII) Adding opt-in screens over content that could be shocking or graphic

We recognise that some content that may otherwise break our rules can be in the public interest, and we allow this content to remain on the platform for documentary, educational, and counterspeech purposes. As we continue to make [public interest exceptions](#) for some content, we provide opt-in screens to help prevent people from unexpectedly viewing shocking or graphic content.

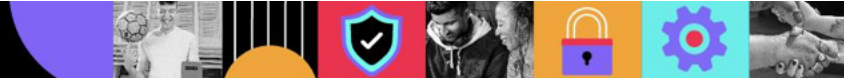
(VIII) External engagement

We are committed to engaging with experts across the industry and civil society, and cooperating with law enforcement agencies globally in line with our [Law Enforcement Guidelines](#), to further safeguard and secure our platform during times of conflict.

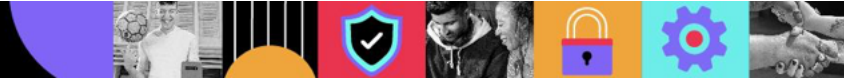
[Note: Signatories are requested to provide information relevant to their particular response to the threats and challenges they observed on their service(s). They ensure that the information below provides an accurate and complete report of their relevant actions. As operational responses to crisis/election



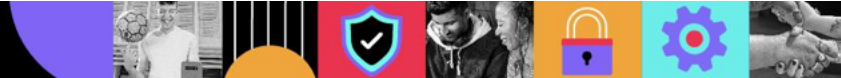
situations can vary from service to service, an absence of information should not be considered a priori a shortfall in the way a particular service has responded. Impact metrics are accurate to the best of signatories' abilities to measure them].		
Policies and Terms and Conditions		
Outline any changes to your policies		
Policy	Changes (such as newly introduced policies, edits, adaptation in scope or implementation)	Rationale
N/A No relevant updates in the reporting period.	N/A In a crisis, we keep under review our policies and to ensure moderation teams have supplementary guidance.	During the reporting period, no crisis-specific policy changes were implemented.
Scrutiny of Ads Placements		
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.		
Content moderation (Commitment 2, Measure 2.2)	Description of intervention TikTok places considerable emphasis on proactive moderation of advertisements. Advertisements are reviewed against our Advertising Policies through a combination of automated and human moderation. These policies prohibit, among other things, ad content and landing pages from displaying language or imagery which promotes, glorifies, advocates for any form of war, armed conflict, hostilities or uprising, as well as prohibiting attempts to profit from ongoing wars or armed conflicts. The policy can allow advocacy for stopping wars or ceasefires and raising awareness of war victims, as well as images or references to soldiers only if the presentation is respectful, memorial in tone and not political, provocative, disparaging or polemical.	



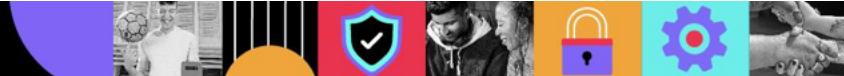
	<i>Indication of impact (at beginning of action: expected impact) including relevant metrics when available</i> Given the range of potential policy violations that could be engaged, we are currently unable to provide metrics specific to this issue.
Political Advertising	
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	
Political Advertising	<i>Description of intervention</i> TikTok did not subscribe to this Chapter as outlined in the January 2025 Subscription Document
	<i>Indication of impact (at beginning of action: expected impact) including relevant metrics when available</i> N/A
Integrity of Services	
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	
Identifying and removing CIO networks <i>(Commitment 14, Measure 14.1)</i>	<i>Description of intervention</i> Our Integrity and Authenticity policies prohibit attempts to manipulate public opinion while misleading our systems or users about identity, origin, approximate location, popularity, or purpose. Dedicated teams monitor and investigate CIO networks and have removed networks targeting discourse related to the War in Ukraine in line with these policies. We continually seek to strengthen our policies and enforcement actions in order to protect our community against new types of harmful misinformation and inauthentic behaviours.



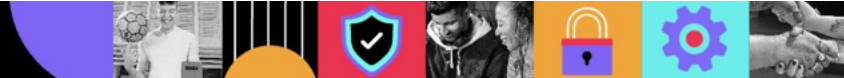
	<i>Indication of impact (at beginning of action: expected impact) including relevant metrics when available</i> Between January and June 2026, we took action to remove the following 4 networks (consisting of 170 accounts in total) that were found to be involved in coordinated attempts to influence public opinion about the War in Ukraine and mislead our community: 1. Network Origin: Ukraine Description: We assessed that this network operated from Ukraine and targeted a Russian audience. Accounts Removed: 27 Followers: 60,918 2. Network Origin: Russia Description: We assessed that this network operated from Russia and targeted a Ukrainian. Accounts Removed: 38 Followers: 124,008 3. Network Origin: Colombia Description: We assessed that this network operated from Colombia and targeted a Colombian audience (With attempts to shift public opinion about the war in Ukraine). Accounts Removed: 16 Followers: 22,773 4. Network Origin: Russia Description: We assessed that this network operated from Russia and targeted a Russian audience. Accounts Removed: 89 Followers: 577 We publish CIO networks we identify and remove, including those relating to the War in Ukraine, in our dedicated CIO transparency report.
	<i>Description of intervention</i>



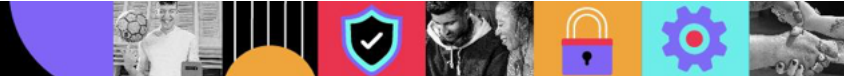
Tackling Edited Media and AI-Generated Content (AIGC) <i>(Commitments 14 and 15, Measures 14.1, 15.1 and 15.2).</i>	Our Edited Media and AI-Generated Content (AIGC) policy makes it clear that we do not want our users to be misled about crisis events. For the purposes of our policy, AIGC refers to content created or modified by AI technology or machine-learning processes. It includes images of real people and may show highly realistic-looking scenes. We do not allow misleading AIGC or edited media that falsely shows: • Content made to seem as if it comes from an authoritative source, such as a reputable news organization, scientific or medical society, or government entity providing critical services; • A critical event, such as an election, natural disaster, or a mass casualty incident; • Matters of public importance, including debates about significant and challenging policy issues; • A public figure who is: ◦ Being degraded or harassed, or engaging in criminal or anti-social behavior; ◦ Taking a position on a political issue, commercial product, or a matter of public importance (such as an election); ◦ Spreading misinformation about matters of public importance. In addition, all AI-generated or significantly edited content that shows realistic-looking scenes or people is not allowed. We have an AI-generated content label for users to easily inform their community when they post AIGC. The label can be applied to any content that has been completely generated or significantly edited by AI, which makes it easier to comply with the obligation to disclose AIGC that shows realistic scenes. Creators can do this through this label or through other types of disclosures, like a sticker, watermark, or caption. TikTok has invested in labeling technologies and tools, including the implementation of Content Credentials technology from the Coalition for Content Provenance and Authenticity (C2PA), which enables the automatic recognition and labeling of AIGC, including AIGC created on some other platforms. AI-generated content. This is complemented by a TikTok-developed tool that allows creators to easily label AI-generated content, which was used to label more than 14 million pieces of content in the EEA during the reporting period. TikTok's commitment to AIGC transparency ensures a safe environment for users, who can easily identify synthetic content and understand its context. <i>Indication of impact (at beginning of action: expected impact) including relevant metrics when available</i> Our efforts support transparent and responsible content creation practices, both in the context of the War in Ukraine and more broadly on our platform.
--	---



Removing harmful misinformation from our platform (Commitment 14, Measure 14.1)	Description of intervention The vast majority of violative content is proactively removed before it is viewed or reported. In H1 2026, more than 99% of videos violating our Integrity and Authenticity policies were removed proactively worldwide. We take action to remove accounts or content that contain inaccurate, misleading, or false information that may cause significant harm to individuals or society, regardless of intent. In conflict environments, such information may include content that is repurposed from past conflicts, content that makes false and harmful claims about specific events, or incites panic. In certain circumstances, we may reduce the prominence of such content. Indication of impact (at beginning of action: expected impact) including relevant metrics when available In the context of the crisis, we have proactively removed 9,825 videos in H1 containing harmful misinformation related to the War in Ukraine. We carry out targeted sweeps of certain types of content as well as working closely with our fact-checking partners and responding to emerging trends they identify. <i>Relevant metrics:</i> • Number of videos removed because of violation of misinformation policy with a proxy related to the War in Ukraine - 9,967 • Number of videos not recommended because of violation of misinformation policy with a proxy (only focusing on RU/UA) - 3,534 • Number of proactive removals of videos removed because of violation of misinformation policy with a proxy related to the War in Ukraine - 9,825
Empowering Users	
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	
Creating localised media literacy campaigns (Commitment 17, Measures 17.2 and 17.3)	Description of intervention We have localised media literacy campaigns related to the crisis to raise awareness amongst our users. We promoted the campaign through a combination of our in-app intervention tools to ensure that authoritative information is promoted to our users. Users searching for keywords related to the War in Ukraine are directed to tips, prepared in partnership with our fact-checking partners. These tips help users identify misinformation and prevent its spread on the platform.



	<i>Indication of impact (at beginning of action: expected impact) including relevant metrics when available</i> Working with our fact-checking partners, we have 17 localised media literacy campaigns addressing disinformation related to the War in Ukraine in Austria, Bosnia, Bulgaria, Czechia, Croatia, Estonia, Germany, Hungary, Latvia, Lithuania, Montenegro, Poland, Romania, Serbia, Slovakia, Slovenia, and Ukraine. <i>Relevant metrics for the media literacy campaigns (EEA total numbers, in countries where campaigns are active):</i> • Total Number of impressions of the search intervention - 26,592,133 • Total Number of clicks on the search intervention - 158,841 • Click through rate of the search intervention - 0.60%
Empowering the Research Community	
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	
Measures taken to support research into crisis related misinformation and disinformation <i>(Commitment 26, Measure 26.1 and 26.2)</i>	<i>Description of intervention</i> Through our Research API, academic researchers from non-profit universities in the US and Europe can apply to study public data about TikTok content and accounts. This public data includes comments, captions, subtitles, and number of comments, shares, likes, and favourites that a video receives from our platform. More information is available here. <i>Indication of impact (at beginning of action: expected impact) including relevant metrics when available</i> Number of Research API applications related to the War in Ukraine that have been approved from January - June 2026: 2
Empowering the Fact-Checking Community	
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	



Applying our unverified content label and making content ineligible for recommendation (Commitment 21, Measure 21.2)	Description of intervention Where our Integrity & Authenticity moderators or fact-checking partners determine that content is not able to be verified at the given time (which is common during an unfolding event), we apply our unverified content label to the content to encourage users to consider the reliability or source of the content. The application of the label will also result in the content becoming ineligible for recommendation in order to limit the spread of potentially misleading information. Our unverified content label is available to users in 23 EU official languages (plus, for EEA users, Norwegian and Icelandic). Indication of impact (at beginning of action: expected impact) including relevant metrics when available We share metrics relating to the unverified content label in the relevant section of the main report (SLI 21.1.2).
Ensuring fact-checking coverage (Commitment 30, Measure 30.1)	Description of intervention We work with 13 fact-checking partners in Europe, providing coverage in 23 official EEA languages, including at least one official language of each EU Member States, and additional languages including Georgian, Russian, Turkish, Ukrainian, Belarusian, Albanian and Serbian. Indication of impact (at beginning of action: expected impact) including relevant metrics when available • Number of fact-checked videos with a proxy related to the War in Ukraine - 665 • Number of videos removed as a result of a fact-checking assessment with words related to the War in Ukraine - 78 • Number of videos not recommended in the For Your Feed as a result of a fact-checking assessment with words related to the War in Ukraine - 147

Reporting on the service's response during a crisis

Israel-Hamas Conflict

Threats observed or anticipated at time of reporting: [suggested character limit 2000 characters].

TikTok continues to moderate violative content at scale, while respecting and protecting the fundamental rights and freedoms of European users. We remain committed to supporting freedom of expression, upholding our commitment to [human rights](#), and maintaining the safety and integrity of our platform during the Israel–Hamas conflict (referred to as the “Conflict” in this chapter). The main threats, both observed and anticipated in relation to the Conflict during the reporting period were the spread of harmful misinformation and Covert Influence Operations (“CIO”).

Mitigations in place at time of reporting: [suggested character limit: 2000 characters].

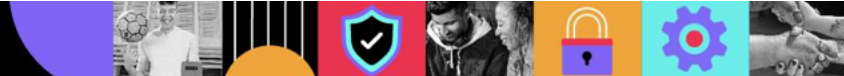
Since the beginning of the Conflict, we are:

Upholding TikTok's Community Guidelines

Continuing to enforce our [policies](#) against [violence](#), [hate](#), and [harmful misinformation](#) by taking action to remove violative content and accounts. For example, we remove content that promotes Hamas, or otherwise supports the attacks or mocks victims affected by the violence. We do not tolerate attempts to incite violence or spread hateful ideologies. We have a zero-tolerance policy for content praising violent and hateful organisations and individuals, and those organisations and individuals aren't allowed on our platform. We also block hashtags that promote violence or otherwise break our rules. We use a combination of advanced moderation technologies and teams of human safety experts to identify, review, and action content that violates our policies.

Automated Review

We place considerable emphasis on proactive detection to remove violative content and reduce exposure to potentially distressing content for our human moderators. Before content is posted to our platform, it's reviewed by automated moderation technologies which identify content or behaviour that may violate our policies or For You feed eligibility standards, or that may require age-restriction or other actions. While undergoing this review, the content is visible only to the uploader.



If our automated moderation technology identifies content that is a potential violation, it will either take action against the content or flag it for further review by our human moderation teams. In line with our safeguards to help ensure accurate decisions are made, automated removal is applied when violations are the most clear-cut.

Some of the methods and technologies that support these efforts include:

- **Vision-based:** Computer vision models can identify objects that violate our Community Guidelines, such as weapons or hate symbols.
- **Audio-based:** Audio clips are reviewed for violations of our policies, supported by a dedicated audio bank and "classifiers" that help us detect audios that are similar or modified to previous violations.
- **Text-based:** Detection models review written content like comments or hashtags, using foundational keyword lists to find variations of violative text. Artificial Intelligence (AI) that can interpret the context surrounding content—helps us identify violations that are context-dependent, such as words that can be used in a hateful way but may not violate our policies by themselves. We also work with various external experts, like our [fact-checking partners](#), to inform our keyword lists.
- **Similarity-based:** "Similarity detection systems" enable us to not only catch identical or highly similar versions of violative content, but other types of content that share key contextual similarities and may require additional review.
- **Activity-based:** Technologies that look at how accounts are being operated help us disrupt deceptive activities like bot accounts, spam, or attempts to artificially inflate engagement through fake likes or follow attempts.
- **LLMs:** We use multimodal LLMs to help moderate content faster and more consistently at scale, from taking automated action on activity like fake engagement, to empowering teams with better moderation tools and risk insights.
- We work with external groups, for example [Tech Against Terrorism](#) in the context of [violent extremist](#) content, who help us to more quickly detect and remove violative content that has already been identified off the platform.

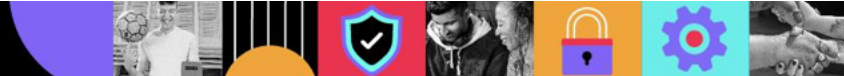
Scaling human expertise

Human insight plays a crucial role in the content moderation process, from our community or external experts, to our own safety professionals. TikTok has Arabic and Hebrew speaking content moderators who review content and assist with Conflict-related translations. We continue to focus on moderator care through the provision of internal training and well-being resources for T&S personnel working on mis & disinformation.

In H1 2026, we have removed 24,411 videos in relation to the Conflict, which violated our misinformation policies.

Leveraging our Global Fact-Checking Program

We use a layered approach to detect harmful misinformation that violates our Community Guidelines, with our Global Fact-Checking Program playing a key role. We assess the accuracy of harmful or hard-to-verify claims by partnering with 20 [IFCN-accredited](#) fact-checking organizations who support over 60 languages on TikTok, including Arabic and Hebrew. We also collaborate with certain fact-checking partners to receive advance warning of emerging



misinformation narratives. This helps facilitate proactive responses against high-harm trends and ensures that our Integrity and Authenticity moderators have up-to-date guidance.

To limit the spread of potentially misleading information, we apply [warning labels](#) and prompt users to reconsider sharing content about unfolding or emergency events that have been reviewed by fact-checkers but cannot be verified—referred to as “unverified content.” Recognising that the situation around the Conflict can change rapidly, we have put in place a process allowing our fact-checking partners to quickly update us if claims previously marked as “unverified” are later verified or clarified with additional context.

Disruption of CIOs

TikTok's [Integrity and Authenticity policies](#) do not allow deceptive behaviour that may cause harm to our community or society at large. We have specifically-trained teams on high alert to investigate, disrupt and remove CIO networks from our platform and we provide regular updates in our dedicated [CIO transparency reports](#). For advertising-related CIO measures, please refer to Chapter 2.

Between January and June 2026, we took action to remove a total of two CIO networks targeting discourse related to Israel and Palestine.

Spread of harmful misinformation

TikTok takes a multi-faceted approach to tackling the spread of harmful misinformation, regardless of intent. This includes our [Integrity and Authenticity policies](#), as well as our products, operational practices, and external partnerships with fact-checkers, media literacy organisations, and researchers. We support our Integrity and Authenticity moderators with detailed misinformation policy guidance, enhanced training, and direct access to our IFCN-accredited fact-checking partners, who help assess the accuracy of content.

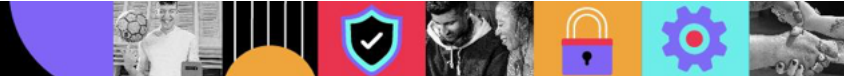
We continue to take swift action against misinformation, conspiracy theories, fake engagement, and fake accounts relating to the Conflict.

Deploying search interventions to raise awareness of potential misinformation

To help raise awareness and to protect our users, we provide in-app search interventions that are triggered when users search for non-violating terms related to the Conflict (e.g., Israel, Palestine). These search interventions remind users to pause and check their sources.

Adding opt-in screens over content that could be shocking or graphic

We recognise that some content that may otherwise break our rules can be of public interest, and we allow this content to remain on the platform for documentary, educational, and counterspeech purposes. As we continue to make [public interest exceptions](#) for some content, we provide opt-in screens to help prevent people from unexpectedly viewing shocking or graphic content.



External engagement

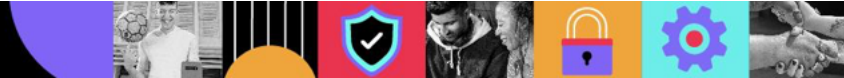
We are committed to engaging with experts across the industry and civil society, such as [Tech Against Terrorism](#), and cooperating with law enforcement agencies globally in line with our [Law Enforcement Guidelines](#), to further safeguard and secure our platform during times of conflict.

[Note: Signatories are requested to provide information relevant to their particular response to the threats and challenges they observed on their service(s). They ensure that the information below provides an accurate and complete report of their relevant actions. As operational responses to crisis/election situations can vary from service to service, an absence of information should not be considered a priori a shortfall in the way a particular service has responded. Impact metrics are accurate to the best of signatories' abilities to measure them].

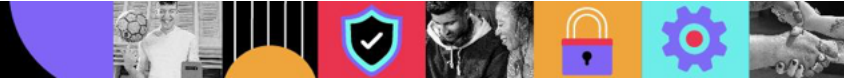
Policies and Terms and Conditions

Outline any changes to your policies

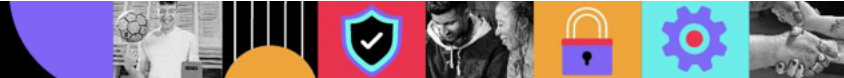
Policy	Changes (such as newly introduced policies, edits, adaptation in scope or implementation)	Rationale



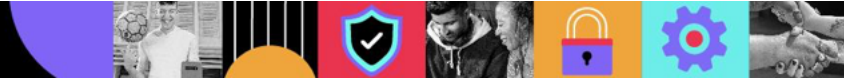
N/A No relevant updates in the reporting period	N/A In a crisis, we keep our policies under review to ensure moderation teams have supplementary guidance.	During the reporting period, no crisis-specific policy changes were implemented.
Scrutiny of Ads Placements		
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.		
Content moderation (Commitment 2, Measure 2.2)	Description of intervention TikTok places considerable emphasis on proactive moderation of advertisements. Advertisements are reviewed against our Advertising Policies through a combination of automated and human moderation. These policies prohibit, among other things, ad content and landing pages from displaying language or imagery which promotes, glorifies, advocates for any form of war, armed conflict, hostilities or uprising, as well as prohibiting attempts to profit from ongoing wars or armed conflicts. The policy can allow advocacy for stopping wars or ceasefires and raising awareness of war victims, as well as images or references to soldiers only if the presentation is respectful, memorial in tone and not political, provocative, disparaging or polemical.	
	Indication of impact (at beginning of action: expected impact) including relevant metrics when available Given the range of potential policy violations that could be engaged, we are currently unable to provide metrics specific to this issue.	
Political Advertising		



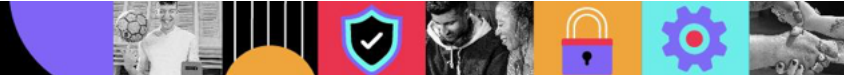
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	
Prohibition on Political Advertising	Description of intervention
	TikTok did not subscribe to this Chapter as outlined in the January 2025 Subscription Document
	Indication of impact (at beginning of action: expected impact) including relevant metrics when available
	N/A
Integrity of Services	
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	
Identifying and removing CIO networks (Commitment 14, Measure 14.1)	Description of intervention
	Our Integrity and Authenticity policies prohibit attempts to manipulate public opinion while misleading our systems or users about identity, origin, approximate location, popularity, or purpose. Dedicated teams monitor and investigate CIO networks and have removed networks targeting discourse related to Israel and Palestine in line with these policies. We know that CIO will continue to evolve in response to our detection and networks may attempt to reestablish a presence on our platform, which is why we continually seek to strengthen our policies and enforcement actions in order to protect our community against new types of harmful misinformation and inauthentic behaviours.
	Indication of impact (at beginning of action: expected impact) including relevant metrics when available
	Between January and June 2026, we took action to remove the following two networks (consisting of 32 accounts in total) that were found to be related to the Conflict:



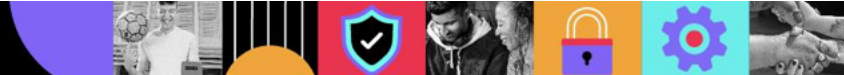
	1. Network Origin: Iran Description: We assess that this network operated from Iran and targeted American audiences. Accounts in network: 27 Followers of network: 110346 2. Network Origin: Iran Description: We assess that this network targeted an American and Israeli audience. Accounts in network: 5 Followers of network: 401 We publish details of the CIO networks we identify and remove, including those relating to the Conflict, in our dedicated CIO transparency report.
Tackling Edited Media and AI-Generated Content (AIGC) <i>(Commitment 15, Measures 15.1 and 15.2)</i>	Description of intervention Our Edited Media and AI-Generated Content (AIGC) policy makes it clear that we do not want our users to be misled about crisis events. For the purposes of our policy, AIGC refers to content created or modified by AI technology or machine-learning processes. It includes images of real people and may show highly realistic-looking scenes. We do not allow misleading AIGC or edited media that falsely shows: • Content made to seem as if it comes from an authoritative source, such as a reputable news organization, scientific or medical society, or government entity providing critical services; • A critical event, such as an election, natural disaster, or a mass casualty incident; • A public figure who is: ○ Being degraded or harassed, or engaging in criminal or anti-social behaviour; ○ Taking a position on a political issue, commercial product, or a matter of public importance (such as an election); ○ Spreading misinformation about matters of public importance. In addition, AI-generated or significantly edited content that shows realistic-looking scenes or people is not allowed. We have an AI-generated content label for users to easily inform their community when they post AIGC. The label can be applied to any content that has been completely generated or significantly edited by AI, which makes it easier to comply with the



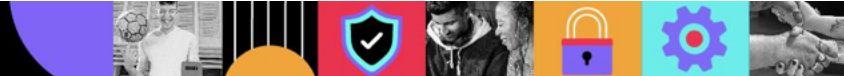
	obligation to disclose AIGC that shows realistic scenes. Creators can do this through this label or through other types of disclosures, like a sticker, watermark, or caption. TikTok has invested in labeling technologies and tools, including the implementation of Content Credentials technology from the Coalition for Content Provenance and Authenticity (C2PA), which enables the automatic recognition and labeling of AIGC, including AIGC created on some other platforms. This is complemented by a TikTok-developed tool that allows creators to easily label AI-generated content, which was used to label more than 14 million pieces of content in the EEA during the reporting period. TikTok's commitment to AIGC transparency ensures a safe environment for users, who can easily identify synthetic content and understand its context. <i>Indication of impact (at beginning of action: expected impact) including relevant metrics when available</i> Our efforts support transparent and responsible content creation practices, which are relevant both in the context of the Conflict and more broadly on our platform.
Removing harmful misinformation from our platform <i>(Commitment 14, Measure 14.1)</i>	<i>Description of intervention</i> We take proactive measures to remove accounts or content that contain inaccurate, misleading, or false information which may cause significant harm to individuals or society, regardless of intent. In conflict environments, this includes content repurposed from previous conflicts, false or harmful claims about specific events, or material that incites panic. In some cases, we may also reduce the visibility of such content. To ensure users can trust the information on our platform, we remove misleading AI-generated content that could cause harm. Misinformation is considered harmful when it is likely to directly lead to violence, fuel tensions, encourage harmful actions, or provoke public panic. Additionally, we restrict misinformation that undermines public trust or distorts public understanding on important matters, even if it does not directly result in violence. We prioritise proactive content moderation, with the vast majority of violative content removed before it is viewed or reported. In H1 2026, more than 99% of videos violating our Integrity and Authenticity policies were removed proactively worldwide. <i>Indication of impact (at beginning of action: expected impact) including relevant metrics when available</i> In the context of the crisis, we have proactively removed 24,177 videos in H1 containing harmful misinformation related to the Conflict. We carry out targeted sweeps of certain types of content (e.g. hashtags/sensitive keyword lists) as well as working closely with our fact-checking partners and responding to emerging trends they identify.



	We have Arabic and Hebrew speaking content moderation as we recognise the importance of language and cultural context in the misinformation moderation process. <i>Relevant metrics:</i> • Number of videos removed because of violation of misinformation policy with a proxy (IL-Hamas) - 24,411 • Number of videos not recommended because of violation of misinformation policy with a proxy (IL-Hamas) - 22,185 • Number of proactive removals of videos removed because of violation of misinformation policy with a proxy (IL/Hamas): 24,177
Empowering Users	
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	
Deploying search interventions to raise awareness of potential misinformation <i>(Commitment 21, Measure 21.1)</i>	Description of intervention
	Indication of impact (at beginning of action: expected impact) including relevant metrics when available These search interventions remind users to pause and check their sources and also direct them to well-being resources.
Empowering the Research Community	
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	
Measures taken to support research into Conflict- related misinformation and disinformation	Description of intervention Through our Research API, academic researchers from non-profit universities in the US and Europe can apply to study public data about TikTok content and accounts. This public data includes comments, captions, subtitles, and number of comments, shares, likes, and favourites that a video receives from our platform. More information is available here.



(Commitment 26, Measure 26.1 and 26.2)	Indication of impact (at beginning of action: expected impact) including relevant metrics when available Number of Research API applications related to the Conflict that have been approved from January - June 2026: 4
Empowering the Fact-Checking Community	
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	
Applying our unverified content label to make content ineligible for recommendation (Commitment 21, Measure 21.2)	Description of intervention Where our Integrity and Authenticity moderators or fact-checking partners determine that content cannot be verified at the given time (which is common during an emergency or unfolding event), we apply our unverified content label to the content to encourage users to consider the reliability or source of the content. The application of the label will also result in the content becoming ineligible for recommendation in order to limit the spread of potentially misleading information. Our unverified content label is available to users in 23 EU official languages (plus, for EEA users, Norwegian and Icelandic). Indication of impact (at beginning of action: expected impact) including relevant metrics when available We share metrics relating to the unverified content label in the relevant section of the main report (SLI 21.1.2).
Ensuring fact-checking coverage (Commitment 30, Measure 30.1)	Description of intervention As part of our fact-checking program, TikTok works with 20 IFCN-accredited fact-checking organisations that support more than 60 languages, including Hebrew and Arabic, to help assess the accuracy of content in this rapidly-changing environment. In the context of the Conflict, our independent fact-checking partners are following our standard practice, whereby they do not moderate content directly on TikTok, but assess whether a claim is true, false, or unsubstantiated so that our moderators can take action based on our Community Guidelines. Fact-checker input is then incorporated into our broader content moderation efforts in a number of different ways, as further outlined in the ‘indication of impact’ section below. Indication of impact (at beginning of action: expected impact) including relevant metrics when available We see harmful misinformation as different from other content issues. Context and fact-checking are critical to consistently and



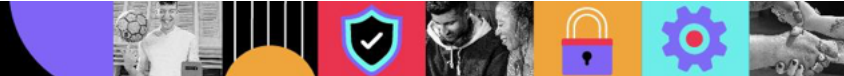
accurately enforcing our harmful misinformation policies, which is why we have ensured that, in the context of the Conflict, our fact-checking programme covers Arabic and Hebrew.

As noted above, we also incorporate fact-checker input into our broader content moderation efforts in different ways:

- Proactive insight reports that flag new and evolving claims they're seeing across the internet. This helps us detect harmful misinformation and anticipate misinformation trends on our platform.
- Collaborating with our fact-checking partners to receive advance warning of emerging misinformation narratives has facilitated proactive responses against high-harm trends and has helped to ensure that our Integrity and Authenticity moderators have up-to-date guidance.

Relevant metrics:

- **Number of fact checked tasks related to IL/Hamas - 851**
- **Number of videos removed as a result of a fact checking assessment with words related to IL/Hamas - 186**
- **Number of videos demoted (NR) as a result of a fact checking assessment with words related to IL/Hamas - 396**



Reporting on the signatory's response during an election

H1 2026 Elections

Threats observed during the electoral period: [suggested character limit 2000 characters].

We have comprehensive measures in place to anticipate and address the risks associated with electoral processes, including the risks associated with election misinformation, in the context of the following elections held during the reporting period:

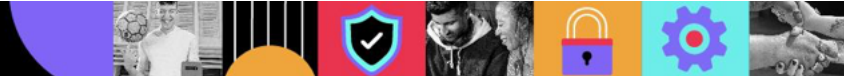
1. Portugal Presidential Election held on 18 January 2026.
2. Slovenia Parliamentary Election held on 22 March 2026.
3. Denmark General Election held on 24 March 2026.
4. Hungary Parliamentary Election held on 12 April 2026
5. Bulgaria Parliamentary Election held on 19 April 2026.
6. Cyprus Parliamentary Election held on 24 May 2026
7. Malta General Election held on 30 May 2026.

Collectively referred to in this chapter as the “**Elections**”.

In advance of the Elections, a core election Task-Force was formed, and consultations between cross-functional teams helped to identify and design response strategies.

Through the Elections, we monitored for and actioned inauthentic behaviour and removed content that violated our Community Guidelines.

Please see Identifying and removing CIO networks (Commitment 14, Measure 14.1) below for details on the covert influence operations we removed.



Mitigations in place during the electoral period: [suggested character limit: 2000 characters].

Investing in media literacy

We invest in media literacy campaigns as a counter-misinformation strategy.

Please see **Rolling out Media literacy campaigns** (Commitment 17, Measure 17.2) below for details on each Election Centre.

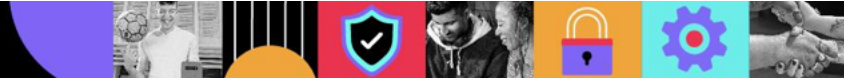
Additionally, we published Newsroom posts relating to the elections below:

- [Portugal Presidential Election](#)
- [Slovenia Parliamentary Election](#)
- [Denmark General Election](#)
- [Hungary Parliamentary Election](#)
- [Bulgaria Parliamentary Election](#)
- [Cyprus Parliamentary Election and Malta General Election](#)

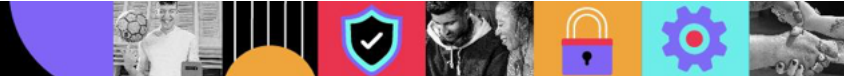
External engagement at the national and EU levels

- **Rapid Response System: external collaboration with COCD Signatories**
 - The COCD Election Rapid Response System (“RRS”) was utilised to exchange information among civil society organisations, fact-checkers, and online platforms. TikTok received 33 RRS during the relevant period. Throughout the election period for each of the Elections, the team maintained consistent prioritisation of RRS requests and ensured timely, accurate support for cross-functional partners.
- **Engagement with local experts**
 - To further promote election integrity, and inform our approach to the Elections, we organised an Election Speaker Series with local fact-checking partners who shared their insights and market expertise with our internal teams.
 - Portugal: Speaker Series was held with Polígrafo on 13 January 2026.
 - Slovenia: Speaker Series was held with LeadStories on 12 March 2026.
 - Hungary: Speaker Series was held with LeadStories on 1 April 2026.

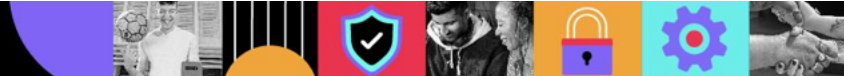
[Note: Signatories are requested to provide information relevant to their particular response to the threats and challenges they observed on their service(s). They ensure that the information below provides an accurate and complete report of their relevant actions. As operational responses to crisis/election



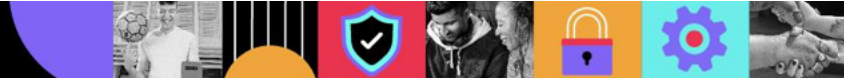
situations can vary from service to service, an absence of information should not be considered a priori a shortfall in the way a particular service has responded. Impact metrics are accurate to the best of signatories' abilities to measure them].		
Policies and Terms and Conditions		
Outline any changes to your policies		
Policy	Changes (such as newly introduced policies, edits, adaptation in scope or implementation)	Rationale
N/A	N/A	N/A
Scrutiny of Ads Placements		
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.		
Scrutiny of Ad Placements <i>(Commitment 2 and Measure 2.1 and Measure 2.3)</i>	Description of intervention We implemented specific granular misinformation policies that provide comprehensive coverage to address harmful misinformation in advertising. In particular, election-related misinformation is explicitly addressed within this policy framework under the Election Misinformation Policy. In addition, we are pleased to be able to report on the advertisements removed for breach of our Political Advertising policy in H1 2026, including the impressions associated with those advertisements. This information is set out in the "Political Advertising Data H1 2026" section at the end of this document.	
	Indication of impact (at beginning of action: expected impact) including relevant metrics when available By prohibiting political advertising, we help ensure our community can have a creative and authentic TikTok experience, and it is one way that we can reduce the risk of our platform being used to advertise and amplify narratives that may be divisive or false.	



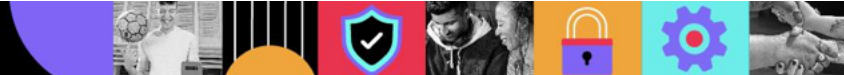
	• Number of ads removed for our political advertising policies during the 4 weeks leading up to and including the day of the Portugal Presidential Election, as well as the additional 3 weeks leading up to and including the day of the second round election (22 Dec, 2025, and 8 Feb, 2026): 1,666 • Number of ads removed for our political advertising policies during the 4 weeks leading up to and including the days of the Slovenia Parliamentary Election (23 Feb, 2026, and 22 Mar, 2026): 676 • Number of ads removed for our political advertising policies during the 4 weeks leading up to and including the days of the Denmark General Election (23 Feb, 2026, and 29 Mar, 2026): 1,025 • Number of ads removed for our political advertising policies during the 4 weeks leading up to and including the days of the Hungary Parliamentary Election (16 Mar, 2026, and 12 Apr, 2026): 2,666 • Number of ads removed for our political advertising policies during the 4 weeks leading up to and including the days of the Bulgaria Parliamentary Election (23 Mar, 2026, and 19 Apr, 2026): 937 • Number of ads removed for our political advertising policies during the 4 weeks leading up to and including the days of the Cyprus Parliamentary Election (27 Apr, 2026, and 24 May, 2026): 1,047 • Number of ads removed for our political advertising policies during the 4 weeks leading up to and including the days of the Malta General Election (27 Apr, 2026, and 31 May, 2026): 0
Political Advertising	
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	
Political Advertising	Description of intervention TikTok did not subscribe to this commitment as outlined in the January 2025 Subscription Document
	Indication of impact (at beginning of action: expected impact) including relevant metrics when available N/A
Integrity of Services	
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	



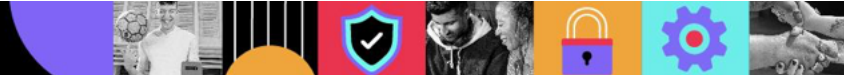
Identifying and removing CIO networks <i>(Commitment 14, Measure 14.1)</i>	Description of intervention During the Elections, we detected the following Covert Influence Operations: • Hungary: ○ In February 2026, we disrupted two networks targeting political discourse. ■ The first network, which we assessed to be operating from Hungary and targeting a Hungarian audience, included 107 accounts with 37,070 followers. ■ The second network, which we assessed to be operating from Ukraine and targeting a European audience, included 45 accounts with 47,114 followers. • The individuals behind this network created inauthentic accounts in order to undermine certain European political figures such as Hungarian Prime Minister Viktor Orbán. ○ In March 2026, we disrupted 2 networks targeting political discourse. ■ The first network, which we assessed to be operating from Hungary and targeting a Hungarian audience, included 91 accounts with 246 followers. ■ The second network, which we assessed to be operating from Hungary and targeting a Hungarian audience, included 62 accounts with 1,452 followers. • Bulgaria: In March 2026, we disrupted a network targeting political discourse. The network, which we assessed to be operating from Bulgaria and targeting a Bulgarian audience, included 34 accounts with 66,763 followers. • Malta: In May 2026, we disrupted a network targeting political discourse. The network, which we assessed to be operating from Malta and targeting a Maltese audience, included 9 accounts with 433 followers. We publish details of the CIO networks we identify and remove in our dedicated CIO transparency reports.
Tackling misleading AIGC and edited media <i>(Commitment 15, Measures 15.1 and 15.2)</i>	Description of intervention Our Edited Media and AI-Generated Content (AIGC) policy makes it clear that we do not want our users to be misled about crisis events. For the purposes of our policy, AIGC refers to content created or modified by AI technology or machine-learning processes. It includes images of real people and may show highly realistic-looking scenes. We do not allow misleading AIGC or edited media that falsely shows:



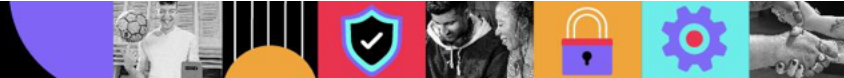
	• Content made to seem as if it comes from an authoritative source, such as a reputable news organisation, scientific or medical society, or government entity providing critical services; • A critical event, such as an election, natural disaster, or a mass casualty incident; • Matters of public importance, including debates about significant and challenging policy issues; • A public figure who is: ○ Being degraded or harassed, or engaging in criminal or anti-social behavior; ○ Taking a position on a political issue, commercial product, or a matter of public importance (such as an election); ○ Spreading misinformation about matters of public importance. In addition, AI-generated or significantly edited content that shows realistic-looking scenes or people is not allowed. We have an AI-generated content label for users to easily inform their community when they post AIGC. The label can be applied to any content that has been completely generated or significantly edited by AI, which makes it easier to comply with the obligation to disclose AIGC that shows realistic scenes. Creators can do this through this label or through other types of disclosures, like a sticker, watermark, or caption. TikTok has invested in labeling technologies and tools, including the implementation of Content Credentials technology from the Coalition for Content Provenance and Authenticity (C2PA), which enables the automatic recognition and labeling of AIGC, including AIGC created on some other platforms. AI-generated content. This is complemented by a TikTok-developed tool that allows creators to easily label AI-generated content, which was used to label more than 14 million pieces of content in the EEA during the reporting period. TikTok's commitment to AIGC transparency ensures a safe environment for users, who can easily identify synthetic content and understand its context. TikTok is a member of the Content Authenticity Initiative and the Coalition for Content Provenance and Authenticity, and was the first video sharing platform to put Content Credentials into practice. We have the ability to read Content Credentials that attach metadata to content, which we can use to instantly recognise and label AIGC. This helped us to expand auto-labelling to AIGC created on some other platforms. Indication of impact (at beginning of action: expected impact) including relevant metrics when available Number of videos removed for violating our Edited Media and AI-Generated Content (AIGC) policy during the Election periods: • Number of removals under this policy during the 4 weeks leading up to and including the day of the Portugal Presidential Election, as well as the additional 3 weeks leading up to and including the day of the second round election (22 Dec, 2025, and 8 Feb, 2026): 2,401
--	--



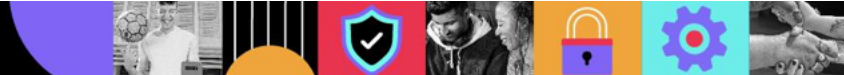
	• Number of removals under this policy during the 4 weeks leading up to and including the days of the Slovenia Parliamentary Election (23 Feb, 2026, and 22 Mar, 2026): 196 • Number of removals under this policy during the 4 weeks leading up to and including the days of the Denmark General Election (23 Feb, 2026, and 29 Mar, 2026): 1,950 • Number of removals under this policy during the 4 weeks leading up to and including the days of the Hungary Parliamentary Election (16 Mar, 2026, and 12 Apr, 2026): 621 • Number of removals under this policy during the 4 weeks leading up to and including the days of the Bulgaria Parliamentary Election (23 Mar, 2026, and 19 Apr, 2026): 911 • Number of removals under this policy during the 4 weeks leading up to and including the days of the Cyprus Parliamentary Election (27 Apr, 2026, and 24 May, 2026): 321 • Number of removals under this policy during the 4 weeks leading up to and including the days of the Malta General Election (27 Apr, 2026, and 31 May, 2026): 143
Empowering Users	
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	
Rolling out Media literacy campaigns <i>(Commitment 17, Measure 17.2)</i>	Description of intervention We launched in-app Election Centres to provide users with up-to-date election information, which contained a section providing tips for spotting misinformation. Below are the launch dates and links for each Election Centre: 1. 9 Dec 2025: Election Centre for the Portugal presidential election. Videos were created in partnership with the fact-checking organisation Polígrafo. 2. 19 Feb 2026: Election Centre for the Slovenia parliamentary election. 3. 11 March 2026: Election Centre for the Denmark general election. 4. 10 March 2026: Election Centre for the Hungary parliamentary election. 5. 19 April 2026: Election Centre for the Bulgaria parliamentary election. 6. 8 May 2026: Election Centre for the Cyprus parliamentary election. 7. 8 May 2026: Election Centre for the Malta general election. We directed people to these Election Centres through prompts on videos, LIVES and searches related to elections.



	Indication of impact (at beginning of action: expected impact) including relevant metrics when available The Election Centres launched before each of the Elections were visited as follows: 1. Slovenia Election Centre was visited 18,024 times. 2. Hungary Election Centre was visited 280,310 times. 3. Bulgaria Election Centre was visited 43,513 times. 4. Cyprus Election Centre was visited 2,338 times. 5. Denmark Election Centre was visited 54,185 times. 6. Malta Election Centre was visited 10,448 times. 7. Portugal Election Centre was visited 85,209 times.
Engagement with local and regional experts (Commitment 17, Measure 17.2)	Description of intervention To further promote election integrity, and inform our approach to the Elections, we engaged with the following partners: Portugal • Held a Speaker Series with our fact-checking partner, Polígrafo, on 13 January 2026. Denmark • Met with six Danish ministries to raise awareness of our election integrity efforts. • Handled one government escalation on election day regarding a Danish artist who was using TikTok as a means of offering valuable artworks in exchange for votes for a specific party. This was escalated and quickly removed to great satisfaction. Slovenia • Participated in a roundtable organized by Agency for Communication on Networks and Services of the Republic of Slovenia (Slovenian DSC) on 3 March 2026. Attendees included representatives of the national authorities, online platforms, and the European Commission. Bulgaria • Visited Bulgarian authorities in March 2026, including the security agency, the Digital Services Coordinator, the Ministry of Interior, and presented our election integrity efforts. Hungary • Met with Law Enforcement to discuss our process and support during elections. • Attended a roundtable organised by the Hungarian DSC National Media and Infocommunications Authority with VLOPs and national authorities on 5 March 2026.

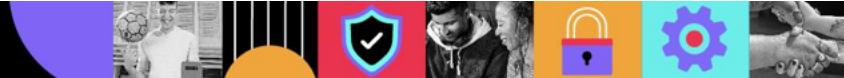


	Held a Speaker Series with our fact-checking partner, LeadStories, on 1 April 2026.
	Indication of impact (at beginning of action: expected impact) including relevant metrics when available This engagement with external regional and local experts allowed us to inform our country-level approach to the Elections.
Empowering the Research Community	
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	
Providing access to our Research API <i>(Commitment 26 and Measures 26.1 and 26.2)</i>	Description of intervention Through our Research API, academic researchers from non-profit universities in the US and Europe can apply to study public data about TikTok content and accounts. This public data includes comments, captions, subtitles, and number of comments, shares, likes, and favourites that a video receives, and comments from our platform. More information is available here .
	Indication of impact (at beginning of action: expected impact) including relevant metrics when available Number of Research API applications related to the Elections that have been approved in H1 2026: 0
Empowering the Fact-Checking Community	
Outline approaches pertinent to this chapter, highlighting similarities/commonalities and differences with regular enforcement.	
Ensuring fact-checking coverage <i>(Commitment 30, Measure 30.1)</i>	Description of intervention - Lead Stories serves as the fact-checking partner for Slovenia and provided fact-checking coverage throughout the election period. - Lead Stories serves as the fact-checking partner for Bulgaria and provided fact-checking coverage throughout the



	election period. • Lead Stories serves as the fact-checking partner for Hungary and provided fact-checking coverage throughout the election period. • Reuters serves as the fact-checking partner for Denmark and provided fact-checking coverage throughout the election period. • Polígrafo serves as the fact-checking partner for Portugal and provided fact-checking coverage throughout the election period. • Lead Stories provided temporary fact-checking coverage throughout the election period for Malta. • AFP serves as the fact-checking partner for Cyprus and provided fact-checking coverage throughout the election period.
	Indication of impact (at beginning of action: expected impact) including relevant metrics when available
	Please refer to Chapter 7 - Empowering the Fact-Checking Community for metrics.

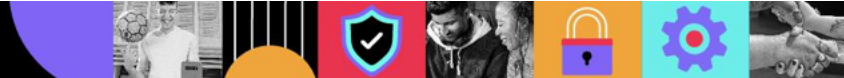
Political Advertising Data H1 2026				
	We are pleased to be able to report on the advertisements removed for breach of our political advertising policies in H1 2026. We also report on the impressions of those advertisements, the number of appeals for advertisements removed under our political advertising policies, and the number of overturns of appeals under political advertising policies in this report.			
	Number of ad removals under the political advertising policy	Number of impressions for ads removed under the political advertising policy	Number of appeals for ads removed under the political advertising policy	Number of overturns of appeals under political advertising policy
Member States				



Austria	6585	1560550	239	78
Belgium	9129	12170473	475	181
Bulgaria	23786	195744904	262	52
Croatia	2521	2940599	20	9
Cyprus	4239	3767090	18	9
Czechia	47999	722127	425	183
Denmark	5432	8209369	144	61
Estonia	2280	226403	91	22
Finland	3965	2916046	144	67
France	74136	93186358	3380	1511
Germany	36619	42474578	2593	943
Greece	9207	10730422	715	290
Hungary	9016	4171626	554	123



Ireland	54744	62038544	1429	186
Italy	27795	16475510	2096	741
Latvia	2620	1471333	165	46
Lithuania	27237	104688	134	50
Luxembourg	1234	80	36	17
Malta	0	0	0	0
Netherlands	39562	13193507	1320	444
Poland	39160	114582098	1174	436
Portugal	20768	3771063	692	256
Romania	20052	10774056	1439	522
Slovakia	2562	1582776	14	4
Slovenia	2660	1979542	60	25
Spain	29492	31093208	2080	783



Sweden	6836	15423231	405	163
Iceland	0	0	0	0
Liechtenstein	0	0	0	0
Norway	5523	1678818	299	108
Total EU	509,636	651,310,181	20,104	7,202
Total EEA	515,159	652,988,999	20,403	7,310