

Dear Chair,

Thank you again for inviting Google to provide oral evidence to the Science, Innovation and Technology Committee hearing on 24 March. I was grateful for the opportunity to update your Committee on my team's work over the past year since my colleague Amanda Storey gave evidence. I am writing to provide responses to the additional information you requested from us and recap on the broader context for this work.

My team continues to focus on building Google's mission - to organise the world's information and make it universally accessible and useful - into trusted, reliable products and features. The trust users put into those products and services is central to our business success, and is why we continually strive for the highest standards of product integrity and efficacy.

This work begins with robust policies, human and automated review systems, and proactive monitoring for abuse leveraging best in class technologies.

We empower our users to evaluate the content on our platforms by delivering reliable information and providing best-in-class tools that put them in control and able to make informed decisions, especially in contexts that are particularly susceptible to the rapid propagation of false information (such as breaking news events).

Misinformation on Google platforms: figures and analysis

While we develop built-in advanced protections that enable us to prevent, detect and respond to harmful and illegal content, automated detection tools face inherent challenges as norms around what constitutes misinformation are constantly evolving.

An objective indicator is the factual accuracy of our AI-powered search features, AI Overviews and AI Mode (referred to in the session as AIO/AIM).

These features run on Gemini 3, Google's most capable model to date, which was released in November 2025. This became the default model for these features globally in January 2026. Specific stats on Gemini 3's factual accuracy include:

- On SimpleQA Verified, a benchmark specifically measuring factual correctness, Gemini 3 Pro scores 72.1%.
- The model also demonstrates PhD-level reasoning, scoring 91.9% on GPQA Diamond and 37.5% on Humanity's Last Exam (an elite-level difficulty benchmark) without the use of any external tools.
- These scores are produced using the FACTS Leaderboard methodology, published in December 2025, which assesses factual accuracy across four dimensions: responses to image-based questions (FACTS Multimodal); world knowledge from training data alone (FACTS Parametric); accuracy when using web search (FACTS Search); and whether long-form responses are grounded in source documents (FACTS Grounding).

Further information is available in these materials: [Gemini 3 announcement](#); [AI Overviews and AI Mode update](#); [FACTS Leaderboard methodology](#).

We believe these benchmarks are a helpful guide to factual accuracy in AI-assisted search. We will continue to improve and publish progress with these efforts.

While these technical solutions play a key part, we also aim to empower users to navigate AI and the wider digital world confidently and safely. To that end, we also invest in scalable media literacy features and tools that help users everywhere understand what they are seeing online. We partner with UK civil society organisations to support targeted training in settings like schools and libraries, as we know this is not a challenge we will solve alone.

Improvements to internal governance systems

Our internal governance systems are designed to protect users and the integrity of our platforms, and we are continuously investing and evolving to meet new challenges as they emerge. These systems are built on three core pillars: robust policies enforced by people and technology, proactive threat disruption, and transparency and user empowerment.

We have a global team of tens of thousands of people who work to enforce our policies. Their work is augmented by our Threat Analysis Group (TAG) and Mandiant Intelligence, who continue to proactively identify, monitor, and disrupt emerging threats, including coordinated influence operations. We have continued to report on these actions in our [public quarterly TAG bulletin](#), which include security actions taken following investigations into coordinated influence operations throughout the period Q3 2024 to Q4 2025, including YouTube channels, Blogger blogs, ads and AdSense account terminations and also domain level enforcement. Reports can be found here: [Q3 2024](#), [Q1 2025](#), [Q2 2025](#), [Q3 2025](#), [Q4 2025](#).

To address emerging challenges posed by new technologies, we continue to invest heavily in developing and deploying new tools. SynthID is a state-of-the-art toolkit to apply an imperceptible, digital watermark and identify AI-generated content. We have recently made it easier for anyone using the Gemini app to verify if an image was created or edited by Google AI. We are a steering committee member of the Coalition for Content Provenance and Authenticity (C2PA) and have begun integrating its 'Content Credentials' technical standard into Google Search and YouTube to provide more context on the provenance of digital content. We also require prominent disclosures for realistic, AI-generated content in election ads and have introduced labels on YouTube to identify synthetic media for viewers.

These efforts represent a multi-layered and adaptive approach to governance, ensuring we can respond effectively to the dynamic information environment.

Content classification for YouTube Kids

We offer parents three age-appropriate content settings: Preschool (ages 4 and under), Younger (ages 5 to 8), and Older (ages 9 to 12). Content is curated through a combination of algorithmic filtering, human review and parent feedback. For parents seeking the highest level of assurance, an "Approved content only" setting allows them to hand-pick every video their child can access, bypassing the algorithm entirely. Since incorporating our quality principles into YouTube Kids recommendations, viewership of content meeting those criteria has increased by over 45% in the app.

As such, YouTube Kids does not carry a single equivalent to a BBFC rating. We believe a static classification system would not be an appropriate or effective framework for a platform of this kind.

YouTube Kids is a standalone video distribution app. Rather than applying a single post-production classification to content, our safety model is dynamic and continuous, designed to adapt as content and children's needs evolve.

Our approach is guided by our Youth Principles and a set of High Quality Content Principles developed in collaboration with independent child development experts, which are embedded directly into our recommendation systems to promote content that inspires curiosity and creativity while limiting content that is overly commercial, sensational or misleading.

Further information on our approach to child safety across our platforms is available at <https://protectingchildren.google>. We believe this layered, adaptive model offers stronger protections for children than any single classification rating could provide, and we remain open to further discussion with the Committee on this point.

Google's API for researchers

Google provides substantial and structured access to data for academic researchers studying our platforms, and is committed to supporting independent scrutiny of our products and their societal impact.

Our researcher support is organised around the following areas:

- **On data and tools**, we make available the Google Transparency Report, the Ads Transparency Center, Google Trends, Dataset Search, Data Commons, and a range of publications from Google Research on security, privacy and abuse prevention.
- **On formal researcher programmes**, we operate a Researcher Programme for academics at non-profit, EU-based institutions, providing access to data from Google Maps, Google Play, Google Search, Google Shopping and YouTube. The YouTube Researcher Program offers equivalent access for YouTube-focused research to eligible academics worldwide.
- **For Search** specifically, the Search Researcher Result API (SRR API) provides approved researchers with programmatic access to search results.
- **On outreach**, we support the research community through the Trust and Safety Research Awards, the Research Scholar Programme for early-career academics, and the Award for Inclusion Research, as well as researcher-focused events run through our Safety Engineering Centres in Munich, Dublin and Malaga.

Full details of all programmes, eligibility criteria and application processes are available on our [Transparency Hub](#). We are committed to supporting independent research into our platforms and welcome engagement from the academic community on these issues.

I hope this additional information is helpful. I remain at your disposal if you have any further questions and look forward to future opportunities to engage with you and your Committee.

Best wishes

Zoe Darmé