

Our AI Principles

Our approach to developing and harnessing the potential of AI is grounded in our founding mission — to organize the world's information and make it universally accessible and useful.

We believe our approach to AI must be both bold and responsible. Bold in rapidly innovating and deploying AI in groundbreaking products used by and benefiting people everywhere, contributing to scientific advances that deepen our understanding of the world, and helping humanity address its most pressing challenges and opportunities. And responsible in developing and deploying AI that addresses both user needs and broader responsibilities, while safeguarding user safety, security, and privacy.

We approach this work together, by collaborating with a broad range of partners to make breakthroughs and maximize the broad benefits of AI, while empowering others to build their own bold and responsible solutions.

Our approach to AI is grounded in these three principles:

1. **Bold** innovation

We develop AI that assists, empowers, and inspires people in almost every field of human endeavor; drives economic progress; and improves lives, enables scientific breakthroughs, and helps address humanity's biggest challenges.

a. Developing and deploying models and applications where the likely overall benefits substantially outweigh the foreseeable risks.

- b. Advancing the frontier of AI research and innovation through rigorous application of the scientific method, rapid iteration, and open inquiry.
- c. Using AI to accelerate scientific discovery and breakthroughs in areas like biology, medicine, chemistry, physics, and mathematics.
- d. Focusing on solving real world problems, measuring the tangible outcomes of our work, and making breakthroughs broadly available, enabling humanity to achieve its most ambitious and beneficial goals.

2. Responsible development and deployment

Because we understand that AI, as a still-emerging transformative technology, poses evolving complexities and risks, we pursue AI responsibly throughout the AI development and deployment lifecycle, from design to testing to deployment to iteration, learning as AI advances and uses evolve.

- a. Implementing appropriate human oversight, due diligence, and feedback mechanisms to align with user goals, social responsibility, and widely accepted principles of international law and human rights.
- b. Investing in industry-leading approaches to advance safety and security research and benchmarks, pioneering technical solutions to address risks, and sharing our learnings with the ecosystem.
- c. Employing rigorous design, testing, monitoring, and safeguards to mitigate unintended or harmful outcomes and avoid unfair bias.
- d. Promoting privacy and security, and respecting intellectual property rights.

3. Collaborative progress, together

We make tools that empower others to harness AI for individual and collective benefit.

- a. Developing AI as a foundational technology capable of driving creativity, productivity, and innovation across a wide array of fields, and also as a tool that enables others to innovate boldly.
- b. Collaborating with researchers across industry and academia to make breakthroughs in AI, while engaging with governments and civil society to address

challenges that can't be solved by any single stakeholder.

Google AI

c. Fostering and enabling a vibrant ecosystem that empowers others to build innovative tools and solutions and contribute to progress in the field.

Other languages

Download in English 

Auf Deutsch herunterladen 

Download dalam bahasa Indonesia 

Baixe em Portugues 

Descargar en Español (España) 

Télécharger en Français 

□□□□□□□□□□ 

Descargar en Español (Latinoamérica) 

हिन्दी में डाउनलोड करें 

□□□□ □□□□ 

Our AI Principles in action

Guided by our AI Principles, our AI governance is operationalized through a comprehensive and multi-layered approach that spans the entire model lifecycle—from responsible model development and deployment to post-launch monitoring and remediation. We identify and assess AI risks through research, expert input, and comprehensive pre- and post-launch testing—which inform our policies and frameworks. We implement robust mitigations to address these risks, enabling our

safety measures to continuously adapt to the complexities of a rapidly advancing technological landscape.

Google AI

Responsible AI Progress Report

PUBLISHED IN FEBRUARY 2026



Learn more about our approach to responsible AI

Learn more [↗](#)

See past AI Progress reports

[2025](#)

[2024](#)

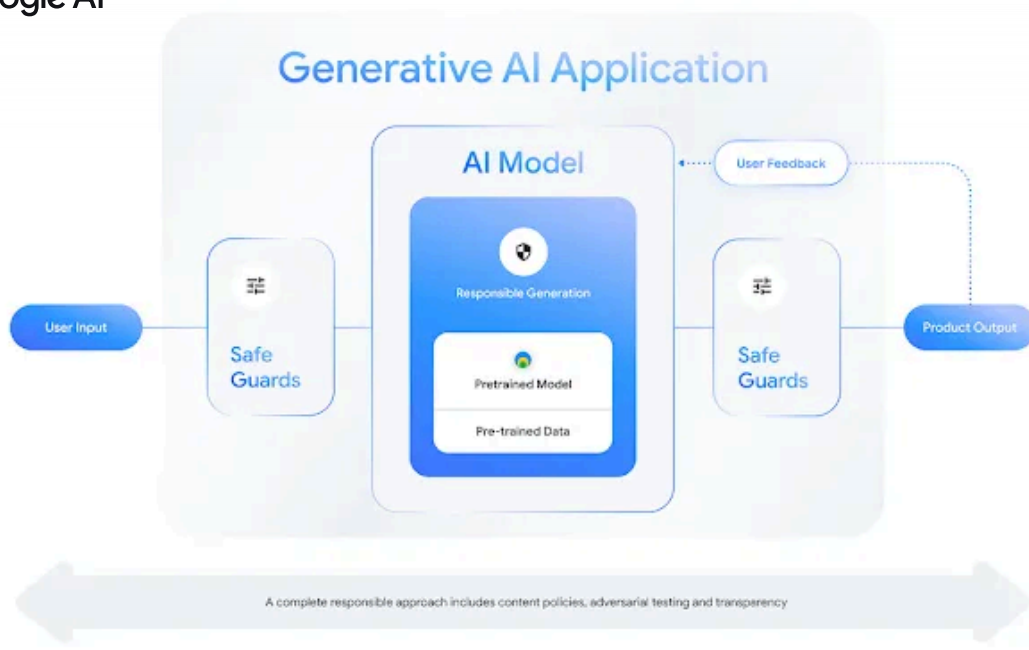
[2023](#)

[2022](#)

[2021](#)

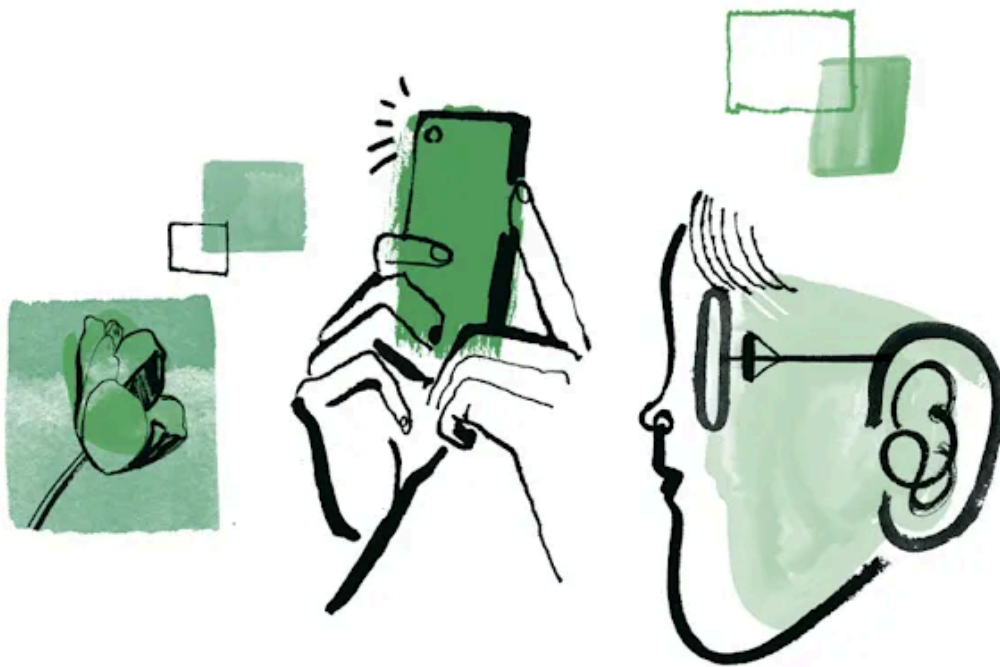
[2020](#)

[2019](#)



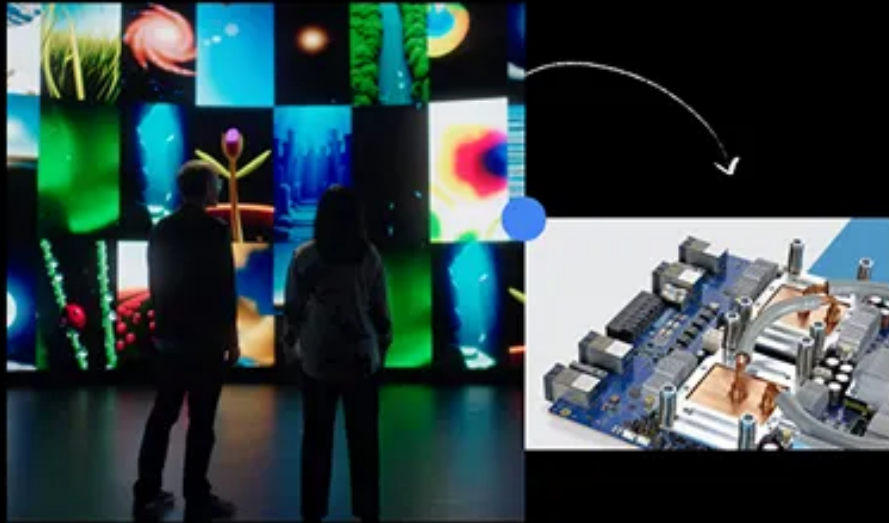
Responsible Generative AI toolkit

Learn more [🔗](#)



[Learn more](#) 

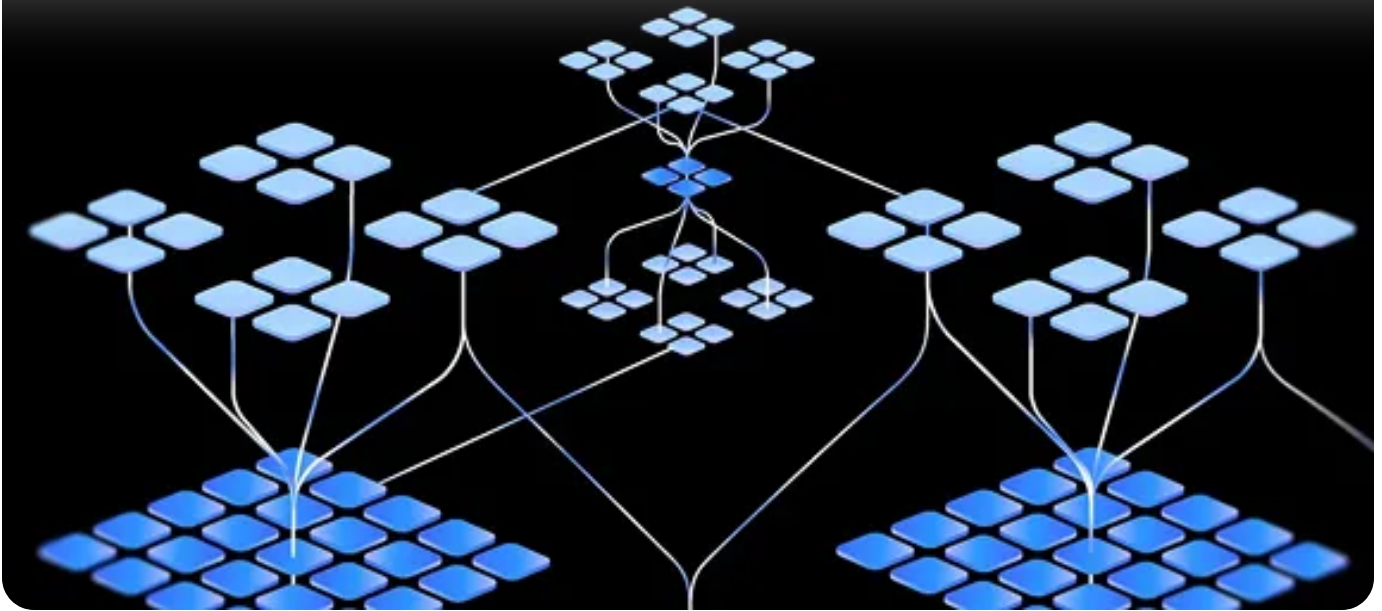
Google AI



Open Research on Responsible AI

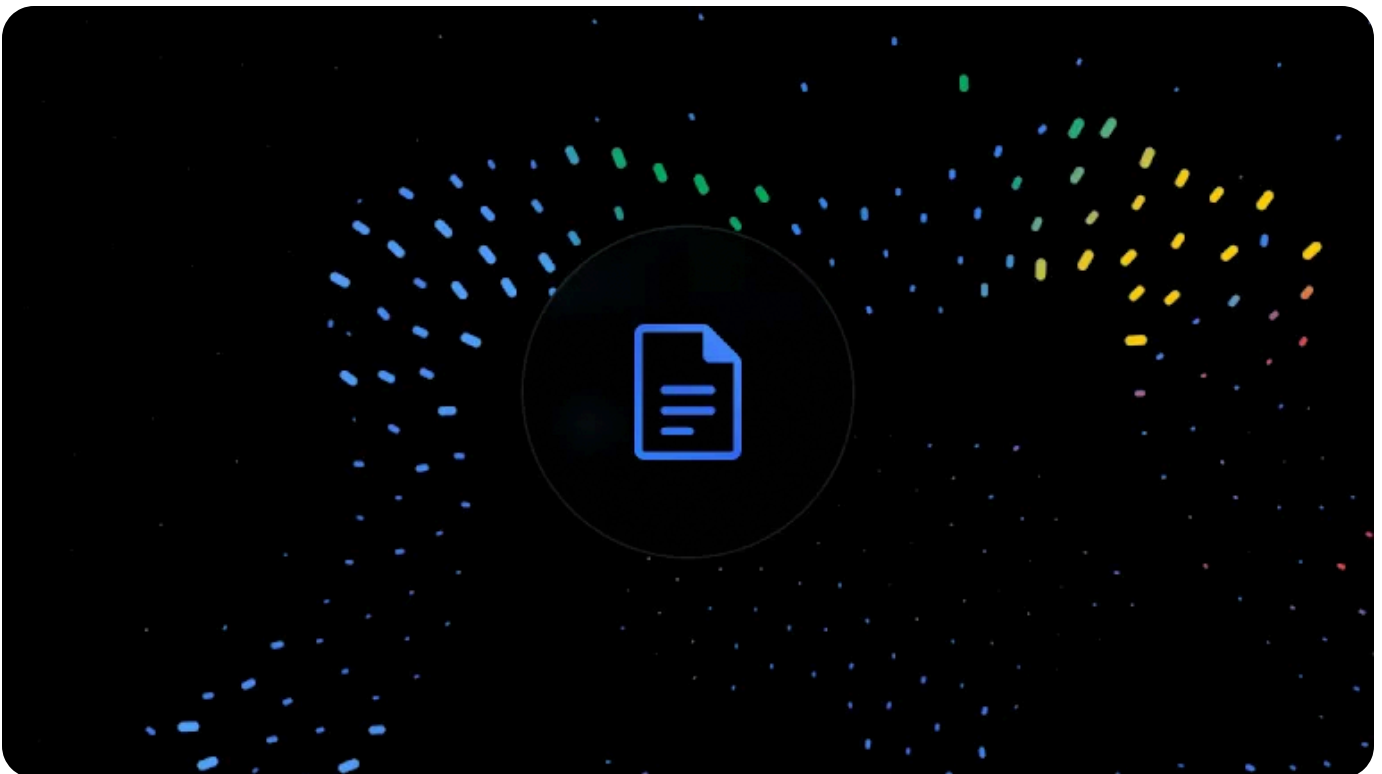
[Learn more](#) 

Google AI



Google's AI products and services: Guided by our AI Principles

Learn more [↗](#)



Simple, structured model information

Learn more 

Google AI

Follow us



Making AI helpful for everyone

Products

Discover how AI can be helpful, from work to everyday life

For knowledge

For creativity

For productivity

For students

For experimenting

Explore products

Research

Tackling the most challenging problems in computer science

Health

Science

Sustainability

Explore more research

About

We're committed to improving the lives of as many people as possible

Why AI

Our AI Journey

AI Principles

Build

Get started building with cutting-edge AI models and tools

Start building

Code with AI assistance

Leverage frameworks and tools

Build with Google AI

Responsibility

We're building and deploying AI responsibly

Responsibility and safety

Policy

Building for everyone

Societal impact

For organizations

Learn AI Skills

Google AI

Google

About Google

Google Products

Privacy

Terms