

Google DeepMind and Isomorphic Labs:

Our approach to bioresilience

Google DeepMind has pioneered the application of AI to grand challenges in biology, from [protein structure prediction](#) to [gene expression](#). Isomorphic Labs is using AI to accelerate drug discovery. Our joint AI bioresilience program brings these capabilities together. Today, we are sharing an update on the program, which seeks to prevent threat actors from misusing our AI models, while ensuring that experts can use the models to prevent, detect, and respond to future outbreaks.

Biosecurity risks are growing. High rates of deforestation, urbanization, and travel increase the risk of a destructive natural outbreak. Fast-improving AI models and agents may erode some of the knowledge barriers that help prevent threat actors from developing and deploying bioweapons.

We need the scientific advances that frontier AI models will enable to make society more resilient to these biosecurity risks. But we need to achieve this while also putting in place stringent measures to prevent threat actors from misusing AI. This is the dual mandate of Google DeepMind and Isomorphic Labs' AI bioresilience program.

The program has three pillars:

- Prevent threat actors from misusing our models.
- Detect new outbreaks quickly.
- Respond decisively to these outbreaks.

Over the past 12 months, we have advanced more than 15 partnerships with government bodies, biosecurity organizations, and research groups across these three areas. We will advance these partnerships further in the next 6-12 months. If you are interested in partnering, please reach out [here](#).

1. [Prevent threat actors from misusing our models](#)

Frontier AI models like Gemini have an increasingly sophisticated knowledge of biology. As they are integrated into agents like [Antigravity](#), and paired with specialized AI biology models and third-party tools and databases, their capabilities will become stronger.

To ensure that Gemini is safe, but still useful to scientists and experts, we use **threat modeling** to understand which actors are most likely to be willing and able to carry out an attack. We draw on various **evaluation methods**—from expert red-teaming to randomized controlled trials—to help judge whether Gemini could help threat actors overcome the primary bottlenecks they face.

We deploy various **mitigations** to counter these risks. Our [post-training](#) methods teach the model to identify and refuse to assist queries with harmful intent, while aiming to avoid the over-refusal of beneficial science queries. Our [classifiers and probes](#) detect and block risky user activity, while our targeted [analysis of user logs](#) helps detect more subtle patterns of misuse and ensures that our mitigations are working effectively. These mitigation efforts are complemented by robust [security infrastructure](#) and privacy controls.

We partner closely with government bodies, external experts, and peers across this risk reduction process. In the next 6-12 months, our priorities include threat intelligence, evaluation methods for AI agents, and jailbreak mitigations. We are also working with our peers at the [Frontier Model Forum](#) to align on best practices for difficult questions, such as how to treat riskier datasets—for example in virology.

To prevent attacks, there are also specific risks that we are exploring further. For example, we want to help ensure that AI models do not make it easier for threat actors to access DNA. Leading DNA synthesis companies, like those in the [International Gene Synthesis Consortium](#), use lists of harmful pathogens and toxins, alongside screening algorithms, to help detect potentially risky orders. However, this approach is starting to fray, as AI makes it possible to design different DNA sequences with similar functionality.

To help DNA companies identify AI-generated sequences, **we are exploring how to adapt our watermarking technology, SynthID**, which has become [an industry standard](#), to biological data. In the future, we also hope to help partners use AI to develop a new kind of DNA synthesis screening that could [help predict the function](#) of a DNA sequence and whether it is likely to be toxic or pathogenic, irrespective of whether it resembles a known pathogen or toxin—[a major, open technical challenge](#).

2. [Detect new outbreaks quickly](#)

To detect new outbreaks, **the world needs ubiquitous biosurveillance** that can sequence genomic data—from wastewater, the air or patients—and characterize the pathogens present in it, in close to real-time. Traditional approaches detect a small number of known pathogens. Metagenomic sequencing can detect novel or less common outbreaks by sequencing all microorganisms in a sample.

To scale metagenomic sequencing to the parts of the world where it is most needed, it needs to become much more cost-effective. A recent [collaboration](#) between Google and Pacific Biosciences used our [AlphaEvolve](#) coding agent to improve the accuracy of sequencing and **we are now exploring other opportunities, from optimizing the algorithms used to analyze sequencing data to informing hardware design**. We are also exploring how technologies like [AlphaGenome](#) could be used to help detect and characterize pathogens from sequence data.

3. [Respond decisively to these outbreaks](#)

A large number of known pathogens have no licensed diagnostic, vaccine, or treatment available, leaving the world highly vulnerable to new outbreaks. Researchers are using our models to help close this medical countermeasure gap.

Over the past five years, **AlphaFold has been cited in more than 10,000 publications on infectious diseases**. Researchers have used it to help [better understand tuberculosis](#) and [malaria transmission](#), and to map vaccine and drug targets for threats like [Mpox](#) and [Nipah](#). We also recently partnered with the bioresilience program at [Lawrence Livermore National Lab](#), which will use AlphaFold 3 to accelerate broad-spectrum antibody design, such as the pan-filovirus antibody. To enable this kind of research, we have worked with our partners to steadily add relevant [protein structures](#) and [complexes](#) to the [AlphaFold Protein Structure Database](#). We plan to add more this year, with a focus on targets for medical countermeasure development.

We are also granting leading scientists targeted access to our most recent AI agent systems, such as Co-Scientist, which helps researchers develop compelling hypotheses. This includes researchers in the US Department of Energy's National Laboratories, who we are collaborating with as part of the US government's [Genesis Mission](#). These scientific partners are pursuing a range of goals, such as identifying novel drug combinations for pathogens that are resistant to current treatments. To scale these impacts, we [recently pledged](#) US\$7 million in support for Health for Human Potential, led by the Philanthropy Asia Alliance, for infectious disease research in Asia.

The complexity of biology means that it is impossible to predict, with perfect precision, the future outbreaks or attacks that may occur. To support government bodies and non-profit organizations during novel outbreaks, **Isomorphic Labs has established a focused unit to rapidly deploy its [drug design engine](#) to help develop medical countermeasures** that could address both naturally occurring pandemics and potential risks arising from the misuse of advanced AI. To ensure real-world impact, we would deploy such capabilities in collaboration with leading governments and national research centers,

such as Lawrence Livermore National Lab, the [UK AI Security Institute](#), [CEPI](#), and the [Francis Crick Institute](#).

Our recommendations for policymakers

We are partnering with government bodies around the world across these three pillars. We also make the following recommendations to US policymakers:

1. Prevent threats by passing a national frontier AI safety framework and modernizing DNA screening

Establish a federal [framework](#) to address biosecurity and other national security risks from frontier AI models, including guidance on how to judge risk-benefit trade-offs. Pass the [AI-Ready Bio-Data Standards Act](#) (H.R. 7907) to ensure that biology training data is safe, standardized, and suitable for training AI models. We also recommend [making DNA synthesis screening mandatory](#), via the [US Biosecurity Modernization and Innovation Act](#) (S. 3741), and directing NIH and NSF to carry out research to evolve screening towards more robust structure- and [function-based methods](#). We also recommend passing the [SCALE Biology Act](#) (H.R. 8981) to help develop the measurement standards, guidelines, and best practices needed to support safe biotechnology.

2. Detect biosecurity incidents quickly by expanding early-warning infrastructure

Identify and characterize new outbreaks before they spread more widely, by aggressively implementing [metagenomic sequencing](#) across key [locations](#) like [transit hubs](#) and high-density areas. Support this by passing the [America's Living Library Act](#) (S. 4023) to sequence and catalog genomic data across public lands, creating a database for biosurveillance. Agencies like DARPA and HHS should also scale up funding for cutting-edge research into new kinds of AI-enabled [early warning](#) and attribution systems.

3. Respond decisively by accelerating medical countermeasure pipelines

Invest in adaptable diagnostic and therapeutic platform technologies that can be quickly customized to counter new biological threats. Enable this by supporting the creation of secure, centralized, high-quality biology datasets, via the proposed Web of Biological Data Act ([H.R. 9307](#) / [S.4770](#)). Complement this with investment in the [infrastructure](#) needed to quickly test, scale up, and deploy novel countermeasures, for example by “warm-basing” and [evolving](#) manufacturing facilities so that they are operationally ready, pre-establishing clinical trial networks and creating accelerated pathways for regulatory approval.
