

Commentary

“More Dangerous than Splitting the Atom”

Learning from Previous Technology Panics



David Osimo

A fierce global debate has erupted over whether artificial intelligence poses an existential threat to humanity and whether we must pause before it is too late, as Anthropic Founder Dario Amodei urgently argued in a recent [blog post](#).

While this dispute may appear entirely novel, its core arguments closely echo the genetically modified organisms (GMO) debates of the 1980s, an analogy highlighted by [The Pessimist Archive](#). There is much to learn from it.

Both AI and GMOs share a "Frankensteinian" name that fuses artificial intervention with nature. GMOs promised dramatic yield increases and reduced chemical pesticide usage, yet they triggered intense public anxiety very similar to AI. Activists warned of unstoppable "superweeds," sterile "terminator seeds" and human health catastrophes due to antibiotic resistance. Press outlets fixated on the unique risks of gene-splicing, warning that unlike localised chemical spills, living engineered organisms reproduce, mutate and migrate uncontrollably across ecosystems. High-profile critics like Jeremy Rifkin argued in 1983 that altering DNA posed a threat *"potentially far more dangerous than our decision to split the atom,"* predicting ecological collapse and moral degradation. Corporate incumbents did not fight regulation; they actively [lobbied](#) for it, not only to build public trust, but to raise compliance costs for smaller startups.

Governments responded by erecting regulatory hurdles. Europe took the most rigid approach under directive 2001/18/EC, institutionalising a process-based "precautionary principle". Rather than evaluating the actual biological risk of the final plant, the European Union regulated the technological method used to modify DNA. In contrast, forward-looking jurisdictions like Canada focused on the novelty of the plant, regardless of the underlying breeding process.

The Visible Benefits and the Invisible Toll of "Missed Use"

30 years later, none of the existential risks materialised and GMOs delivered on their promises: as the [2014 meta-analysis](#) summarises: "GM technology adoption has reduced chemical pesticide use by 37%, increased crop yields by 22% and increased farmer profits by 68%." These numbers are even greater for smallholders in developing countries who saw a [15 to 20% decrease in food insecurity](#), as they were able to achieve high yields even without the capital and equipment of large-scale farmers.

Admittedly, the market remained an oligopoly with high barriers to entry due to regulation and corporates enjoyed tremendous profit guaranteed by very strong intellectual property provision. But at the same time, GMOs delivered life-changing benefits, especially for smallholders in developing countries and the urban poor.

On the other hand, Western precautionary bias was exported globally with devastating human consequences. A landmark [study](#) by economists Justus Wesseler and David Zilberman quantified the human cost of delaying Golden Rice—a crop engineered to synthesise Vitamin A. They calculated that a decade-long regulatory delay in India alone cost over 1.4 million disability-adjusted life years (DALYs) and resulted in 600,000 to 1.2 million preventable cases of childhood blindness.

By 2021, the European Commission formally acknowledged that its framework was no longer fit for purpose. In classic technocratic prose, the Commission [observed](#) that *"the use of new genomic techniques raises ethical concerns, but so does missing opportunities as a result of not using them."* One is left to ask why this basic trade-off is typically ignored when designing regulations.

What does this mean for artificial intelligence?

We can't extrapolate the conclusions from GMOs to declare the existential risks of AI a simple moral panic. Mr Amodi is most certainly sincere when he voices fears and concerns. His article is not fully convincing per se, as it refers to two well-known issues (the self-improving character of AI and the combination of misalignment with stronger capabilities) that, in his judgement, are becoming too much. It's an important and authoritative statement as it comes from a leading provider, but it does not provide evidence to help cut through the noise. It adds to it.

In this confusing context, there are valuable lessons to be learned from GMOs.

1. Technological progress is one of the most important ways to leapfrog human conditions and solve issues at large scale. Technological progress can at the same time benefit those most in need and make jaw-dropping profits for corporations. Then, of course, how the pie is divided is a matter for governments to regulate. But everyone is better off than they would be without technological advancements.
2. Certain innovations, in particular those that are close to the concept of life and intelligence, generate widespread fears well beyond what is reasonable. Any sufficiently disruptive innovation has always looked like an act of arrogance, of humans playing god, and becomes a source of moral panic ("this is too much!").
3. In a fear-driven debate, it is very difficult to find trusted voices. Almost any position can be traced back to a specific interest and incumbents can have an interest in raising barriers to entry through safety provisions.
4. Slowing down technology development and adoption through the precautionary principle is not always the safest option, as it can carry high costs, in particular for the most vulnerable. A well-documented psychological bias leads society to judge active harm ("misuse") far more harshly than passive, unseen loss ("missed use"). The moral panic risks concealing the very real opportunity cost of stalling AI breakthroughs in oncology, structural biology, grid optimisation and clean energy materials.

The need for a new vision

In this heated debate, today's political responses to AI risks are deeply fractured. The Trump administration has embraced the pro-tech vision, offering an apparently consequential combination of export controls with a free-market stance that commercial incentives alone will take care of safety externalities. The rest of us are left with strong ex ante precautionary principles such as the AI act and vague statements of principles. Ursula von der Leyen's assertion that *"AI has the potential to benefit humanity, if we get it right"* or Barack Obama's framing that outcomes *"depend on the choices that we make right now"* are certainly correct but offer no tangible guiding framework.

It is our duty to move beyond this recurrent stalemate between extreme free-market and rigid precautionary principles by articulating a new vision driven by opportunity, not fear. There are already initiatives in this sense, such as performance-based regulation, public AI, sandboxes and strengthened institutional capacity. As Nobel Prize Winner Jennifer Doudna stated about CRISPR, the gene-editing technology she pioneered: *"Few technologies are inherently good or bad; what matters is how we use them... Weighing the dangers inherent in a technology like CRISPR against the responsibility to use its power for the benefit of humanity and our planet will be a test like no other. Yet it's one that we must pass."*

The GMOs debate not only shows the damages of panic-driven regulation, but also the pillars for a constructive debate: scientific evidence, an open and inclusive debate and a constant focus on concrete benefits for individuals, in particular those who need it most. This debate should not embrace techno-optimism instead of techno-pessimism, but evidence instead of fear, the many instead of the few.

theForum

The Forum is a semi-weekly newsletter of cutting-edge blogs devoted to policy proposals and up-to-date political analysis. It is hosted by the Lisbon Council, a global think tank based in Brussels. For more commentary visit: <https://lisboncouncil.net/category/blog/>