

The Consequence Gap

An AI essay that slowly becomes an essay about ourselves.

Abby Buchanan

July 2026

ABSTRACT

Most conversations about AI ask whether it is good or bad. This memo argues that the more revealing question is not about AI itself, but about the limits of human perception. Drawing on evolutionary biology, history, organizational leadership, and personal experience, it explores the gap between the consequences humans evolved to perceive and the consequences modern technologies create. In doing so, it asks what happens when our instincts for adaptation and protection encounter systems operating across generations, continents, and entire ecosystems.

I've noticed that nearly every conversation about AI eventually becomes a conversation about who we believe we're protecting.

Jobs.

Innovation.

National security.

Shareholders.

Copyright.

Children.

The environment.

Future generations.

Underneath the arguments about technology is a quieter question:

Who belongs inside our circle of protection?

As I've sat with that question, I've noticed something about the way I think. When I encounter a complex problem, my mind rarely produces a single opinion. It assembles a panel.

I've come to trust that panel because, if I stay with the conversation long enough, a different question usually emerges.

The ethicist questions whether inevitability is ever a sufficient argument, particularly when humans are still the ones making the decisions. "It's inevitable" feels lazy at best and coercive at worst. The pragmatist assumes AI is here whether we're ready or not. The idealist looks for ways to steer the technology toward human flourishing. The realist notices AI creating extraordinary gains in productivity, analysis, and synthesis even as unintended consequences accumulate in its wake.

Then the researcher notices something different: AI doesn't eliminate work nearly as often as it redistributes it.

The steward thinks about the people responsible for leading organizations through this moment. They must balance the interests of employees, customers, investors, regulators, and boards whose priorities rarely align. Leadership begins to look less like choosing between right and wrong than choosing between competing responsibilities whose consequences unfold on different timelines. The danger lies both in moving too slowly and in moving too quickly.

I don't envy that position.

Then the pessimist quietly reminds everyone that we may simply be accelerating consequences we don't yet understand.

Finally, the historian clears her throat.

Human history is full of moments when technology expanded faster than our ability to understand its consequences.

Fire transformed survival.

Agriculture transformed civilization.

The printing press transformed knowledge.

Industrialization transformed labor.

Fossil fuels transformed our idea of prosperity.

The internet transformed communication.

Social media transformed attention.

Each expanded what humans could do.

Only later did we begin asking what those technologies were doing to us. History suggests that our greatest challenge isn't inventing new capabilities. It's understanding the systems those capabilities create.

The longer I sit with AI, the less interested I become in asking whether it's good or bad. The more interesting question is whether humans evolved to perceive consequences at the scale our technologies now create them.

What happens when our instincts evolved for a world of immediate, local consequences, but our technologies create consequences that are delayed, distributed, and largely invisible?

The biologist interjects.

Our ancestors didn't need to think fifty years ahead or halfway around the globe to survive. They needed to recognize immediate threats, protect the people closest to them, and secure enough food to make it through another season.

We are living proof that those instincts worked remarkably well.

A pizza or bowl of ice cream would have been an evolutionary jackpot. No wonder we struggle to find an off switch. We've changed our environment far more quickly than we've changed ourselves.

Evolution wired us to optimize for proximate success.

Modern systems require ultimate success.

The tension between those two ideas may define this century.

Once humans could reliably survive, the next order of business became acquiring and preserving knowledge. Knowledge reduced uncertainty. It increased the odds that our family, our tribe, and our community would survive.

But knowledge wasn't valuable simply because it answered questions.

It helped us adapt.

Humans who failed to adapt to changing environments rarely passed on their genes. We are descendants of the ones who learned, adjusted, experimented, and survived.

The promise of AI speaks directly to those instincts. It offers access to an almost unfathomable amount of rapidly synthesized information while promising to reduce effort, increase efficiency, and help us navigate an increasingly complex world.

But it isn't just that AI feels useful.

For many people, it feels necessary.

Companies are told they must adopt AI or be left behind. Workers are told they must learn AI or risk becoming obsolete. Some have already watched colleagues lose jobs as organizations restructure around AI. Others worry they may be next.

The same technology can feel like both the lifeboat and the storm.

Our ancestors adapted because survival depended on it. Today, the stakes are different, but economic uncertainty, social relevance, and employability can still feel deeply intertwined with survival.

No wonder AI is so compelling.

Eventually another voice joins the conversation.

The mother.

In my field, AI is already becoming part of the work, and I expect that trend to continue. Using these tools may help me remain employable. They may help me earn more, work more efficiently, and better provide for my children.

That feels like protection.

But I keep coming back to an uncomfortable question.

What if I'm confusing immediate security with long-term survival?

If I use AI to support my family today while participating—even in some tiny way—in systems that consume enormous amounts of energy and water, accelerate climate change, or widen inequality, have I actually protected my children? Or have I simply optimized for the next five years instead of the next fifty?

Human history is full of this paradox.

Wars are fought to protect nations, yet cities burn.

Forests are cleared to feed communities, yet ecosystems collapse.

Parents work longer hours to create security, only to lose the time they hoped to create with their children.

Again and again, humans act to protect the people they love while inadvertently changing the conditions those people ultimately depend upon.

That isn't necessarily because we're selfish or short-sighted. It may be because our instincts evolved for a world where cause and effect were visible, local, and immediate.

Perhaps every generation inherits technologies capable of producing consequences our instincts were never designed to perceive.

Technology has stretched cause and effect across continents and generations. Decisions made in one room may not reveal their consequences for decades, and those consequences may be borne by people we'll never meet. Our systems now operate at a scale our intuitions were never designed to perceive.

That also makes resistance more complicated than it's often portrayed.

Some people resist new technologies because they dislike change.

Some because their identity or livelihood is threatened.

Some because of deeply held moral, religious, or ideological convictions.

Some because they recognize harms others are minimizing.

Some because caution is itself an adaptation strategy.

And sometimes they're right.

History contains examples of every one of these, making it extraordinarily difficult, in the present moment, to distinguish resistance to progress from resistance to preventable harm.

That uncertainty is part of what keeps drawing me back—not because I think I have answers, but because I think we're asking the wrong questions.

Not, *Is AI good or bad?* but, *What consequences are we capable of perceiving—and which do our instincts systematically overlook?*

Maybe the question isn't simply how we build more powerful technologies. It's how we learn to perceive consequences beyond the horizon of our instincts.

Because what we perceive shapes what we choose to protect.

Evolution wired us to protect our own. That instinct is beautiful. But it's also incomplete.

Perhaps the next challenge isn't abandoning it, but expanding it. Can we widen the circle of who counts as our own? Can we place the most vulnerable—not at the edge of that circle, but at its center?

Children. Future generations. Communities with the least power to influence the systems shaping their lives. The ecosystems and non-human life that make our own survival possible.

Perhaps protecting the most vulnerable is another way of protecting ourselves—not tomorrow, but fifty years from now.

Every generation draws a circle around the lives it considers worthy of protection. Sometimes that circle expands. Sometimes it contracts. The technologies we build—and the systems we adopt—reveal where we've drawn it.

Perhaps every generation inherits the opportunity to draw that circle a little wider than the one before it. **I wonder if that's what this moment is asking of us.**