

NAUTILUS

DUMB ANIMALS

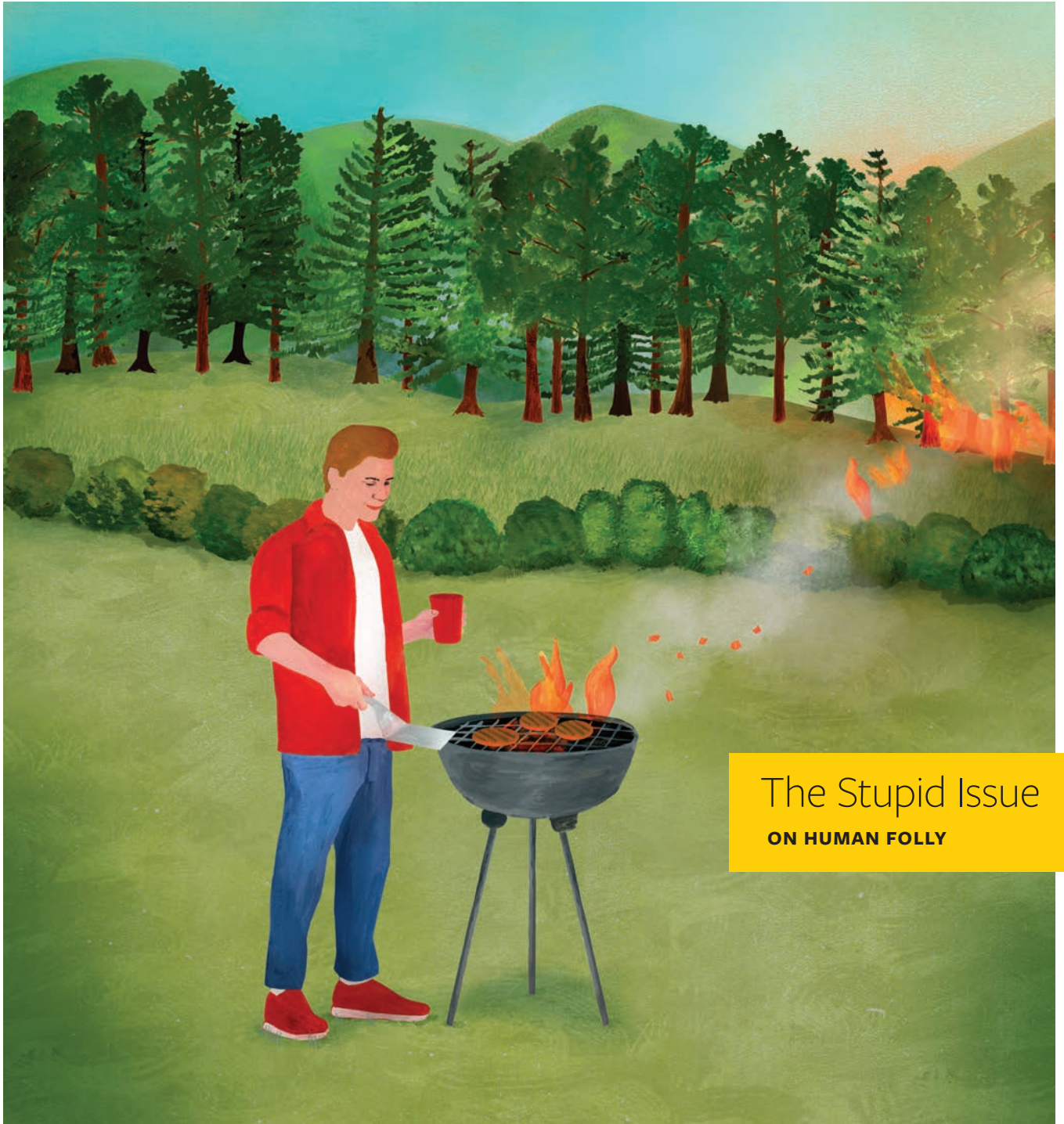
When evolution favors the brainless

HELL HOLE

A Soviet geology story

SCIENTISTS STRIKE BACK

The Trump resistance gets serious



The Stupid Issue

ON HUMAN FOLLY

Repeatedly ranked

30+ lists in the last 3 years

#1



ASU's Bermuda Institute of Ocean Sciences deploys unmanned gliders capable of acquiring continuous measurements of the ocean to depths of 1,000 meters.

innovation

ASU ahead of MIT and Stanford

— U.S. News & World Report, 11 years, 2016–26

sustainability

ASU ahead of Stanford and UC Berkeley

— Association for the Advancement of Sustainability in Higher Education, 3 years, 2023–25

global impact

ASU ahead of MIT and Penn State

— Times Higher Education, 6 years, 2020–25



asu.edu/research

Contents

Editor's Note

- 7 *Stupid Me, Stupid Us*
BY BOB GRANT

The Porthole

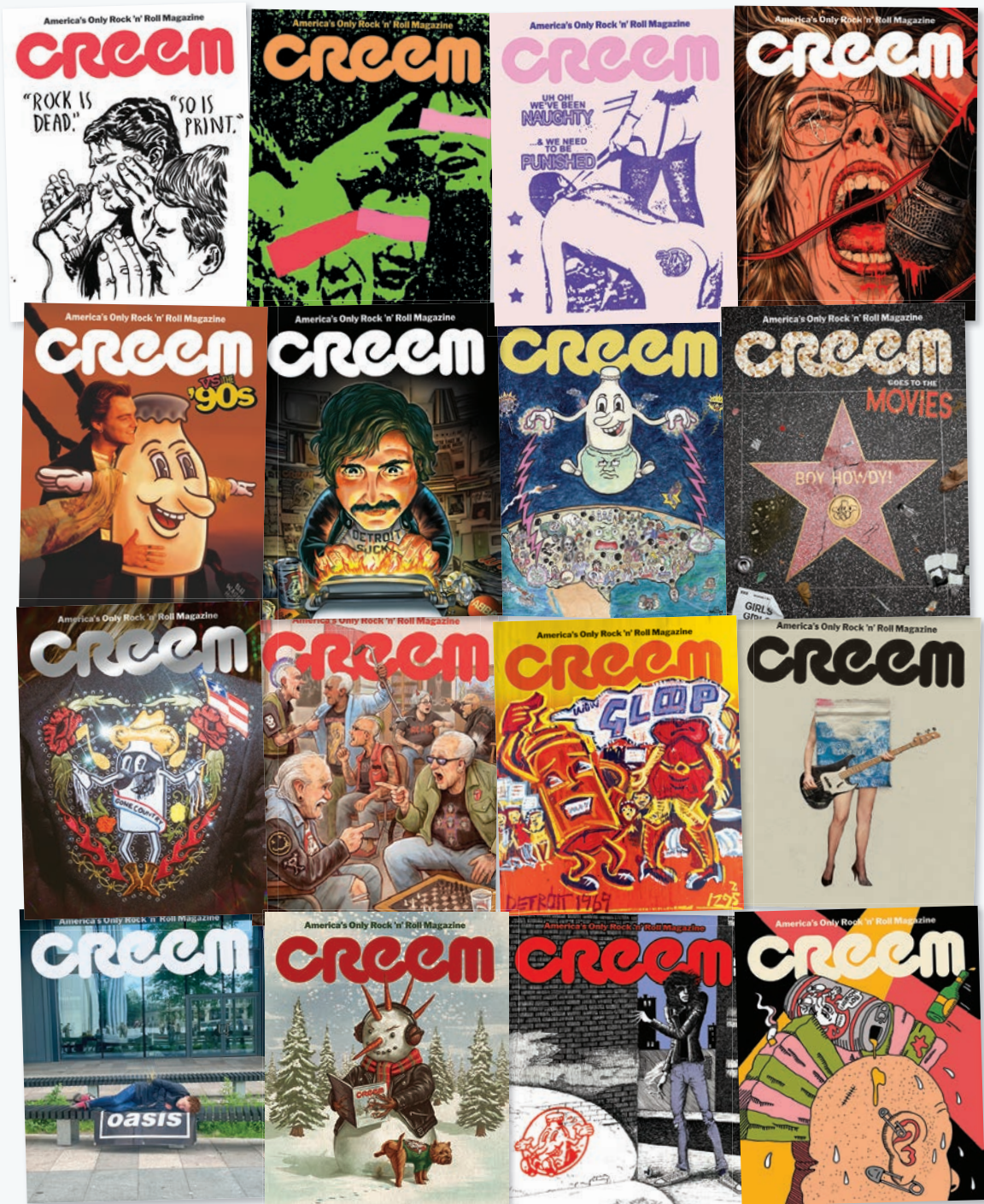
- 9 *The wisdom of the dodo; what you really think about;
the amoeba's hidden smarts; does TV rot your brain?;
a garbageman on waste; and a "black hole star"*

The Last Word

- 112 Martin Schwartz
*The Yale School of Medicine professor on the importance
of stupidity in science*
INTERVIEW BY BOB GRANT

COVER & THEME ART Ellen Weinstein

America's Only Rock 'n' Roll Magazine



USE CODE **NAUTILUS** AT CHECKOUT FOR **25% OFF**

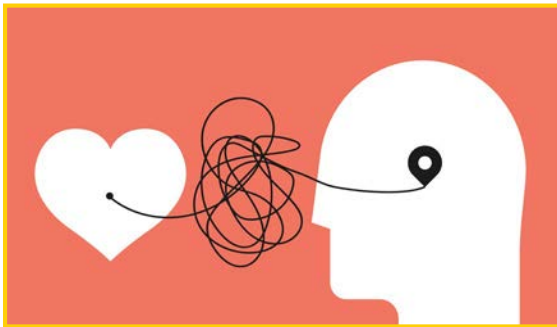
Creem

creem.com/subscribe

Features



- MATH
20 Coincidences in My Life Have Me Wondering
My search for a hidden structure to astronomically unlikely occurrences
BY GEORGE MUSSER



- PHILOSOPHY
30 The Inventor of the Thinking Machine Didn't Worry. Neither Should You
In this age of AI anxiety, listen to your heart
BY YUVAL AVNUR



- TECHNOLOGY
34 This Does Not Compute
Why are we using machines to study our inner lives?
BY JULIE SEDIVY

Elevate your feed.

Follow Nautilus on IG



 www.instagram.com/nautilusmag/



42 Out of the Woods

BY KEVIN BERGER



PHILOSOPHY

44 What Makes Humans Stupid

It takes intelligence to get things spectacularly wrong

BY DAVID C. KRAKAUER

COMMUNICATION

54 The Resistance Is Here

They didn't plan to be activists. But scientists have united to overturn Trump's drastic cuts to science.

BY ADAM PIORE

ZOOLOGY

60 The Virtue of Being Simple

A surprising number of creatures don't benefit from greater intelligence

BY BRANDON KEIM

PSYCHOLOGY

64 Why We Think We Know More Than We Do

Breaking down the famous "Dunning-Kruger effect" with David Dunning himself

BY MATTHEW HUTSON

ISSUE THEME

Stupid

TECHNOLOGY

70 AI Made Me Look Foolish. Then I Learned

Ancient alchemy has a lot to teach us about chatbots

BY PAUL M. SUTTER

HISTORY

76 When Stupid Was a Diagnosis

A dark history that is in danger of returning as support for people with disabilities withers

BY ALICE FLEERACKERS

GEOSCIENCE

84 A Hell of a Hole

Soviet Russia once journeyed to the center of the Earth. Here's what they found.

BY CHARLES DIGGES

HEALTH

90 The Trouble with Cancer Screening

Tests for cancer often lead to more confusion than clarity. Here's how to understand them.

BY RAHUL PARIKH & JONATHAN MATES

HISTORY

98 The Return of a Discredited Pseudoscience

Eugenics is rearing its ugly head again in today's crucible of politics and culture

DAN FALK

EVOLUTION

104 It's the Evolution, Stupid

Immortality is possible? Nature has 3.5 billion years of experience saying otherwise.

BY ELENA KAZAMIA



NAUTILUS

EDITORIAL

NautilusNext CEO
Nicholas White

Executive Editor
Bob Grant

Managing Editor
Liz Greene

Editor At Large
Kevin Berger

Senior Editor
Kristen French

Staff Writer
Jake Currie

Copy Editor
Teresa Cirolia

Production Designer
Tasnuva Elahi

Contributing Writers
Elena Kazamia
Brandon Keim
George Musser
Adam Piore
Devin Reese
Paul M. Sutter

CONTACT

3112 Windsor Rd
A391, Austin, TX 78703
info@nautil.us

SUBSCRIPTIONS

Visit nautil.us/products

Published by:



NAUTILUS *Next*

Copyright © 2026. All rights reserved.

Stupid Me, Stupid Us

BY BOB GRANT

HAVE A CONFESSION to make: I'm stupid.

Sure, I've had the privilege and luck to secure an abundance of formal education. I even fancy myself a relatively quick thinker.

But even with the intellect I've developed, the degrees I've accrued, I regularly fall prey to a tendency at the nucleus of the stupidity that suffuses our broader cultural moment. Frequently, my good sense is swamped by frustration, anxiety, and impatience. Even when I know, in my heart and mind, that kindness and compassion are objectively the best paths to navigate a given situation, those vexing impulses override my supposed intelligence and lead me to act stupidly. Losing my temper with family, focusing on minor irritations while ignoring my massively good fortune, missing opportunities to connect in a haze of self-absorption.

Pretty stupid behavior.

I mention this not as a public act of self-flagellation. Rather, my inward reflection mirrors my work with editors here at *Nautilus* to put this "Stupid Issue" together. Sensing that science is under attack in the United States is a daily and unavoidable occurrence for people tasked with covering the research enterprise here and beyond.

The current administration has throttled funding that once fueled promising lines of research, from climate science tracking changes to the planet's oceans to biomedical trials poised to deliver new treatments. Key scientific

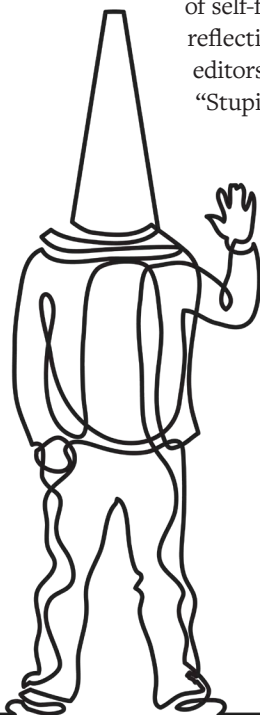
advisory boards fired and replaced with loyal, but unqualified, officials. Guidelines meant to soundly inform people and their doctors to make the best decisions about avoiding infectious diseases and food-borne pathogens reshaped, the changes undergirded by pseudo-science or outright partisanism.

Pretty stupid.

This is all to say that now felt like the perfect time to devote an issue of *Nautilus* to the science, philosophy, and history of stupidity. All the more pressing as the transformative incursion of artificial intelligence into society threatens to make us dumb and dumber. I did worry, though, I was again lapsing into my habit of unthinking meanness. Is calling the phenomenon and the people who display it "stupid" an act of cruelty, itself stupid? This thought led us to pursue an article about the dark history of scientists and health practitioners lobbying this pejorative label at people with intellectual disabilities. In "When Stupid Was a Diagnosis," Alice Fleerackers relays a dark past that is in danger of returning.

Other stories in the issue flip the concept of stupid on its head. Brandon Keim celebrates the utility of cognitive simplicity throughout the animal kingdom, while Charles Digges burrows into a seemingly quixotic Soviet project to dig the deepest hole that yielded deep insight into the nature of our planet.

As I consider how I can knit together my intellect with enhanced social and emotional intelligence to wend a better, more-peaceful, less-stupid path, building this issue made me see the parallels with the trajectory of our society. By letting science lead the way (largely by staying out of its way) and wielding the knowledge we accumulate with compassion, honesty, and kindness, we just might be able to emerge from these stupid times and into a brighter future.



Good
old-fashioned
intelligence.

Subscribe now and get a year of *Harper's Magazine* in print
and online, including more than 175 years of archives to explore.

HARPERS.ORG/NAUTILUS

The Portfolio

Short sharp looks at science

Short sharp looks at science

ZOOLOGY

Not Such a Dodo, After All

DODO BIRDS GOT A BAD RAP from the very start. When European sailors landed on the Indian Ocean island of Mauritius, they thought it looked foolish waddling around on the ground with little fear of humans. Thus, they could only think to call it one thing: a “dodo.” Even Linneus in 1766 classified it as *Didus ineptus*.

Granted, it was the size of a turkey and completely flightless despite its wings. But a recent study published in the *Zoological Journal of the Linnean Society* shows that dodos weren’t as inept as they appeared.

Led by researchers at the University of Lethbridge in Canada, the study explored the sensory abilities of the dodo (now classified as *Raphus cucullatus*), which went extinct in 1662. To do so, they scanned three dodo skulls housed in natural history museum collections in Denmark and the United Kingdom. Two of the specimens constituted the only complete dodo skulls found to date, the internal cavities of which revealed contours of the dodos’ brains.

“Although the dodo’s brain disappeared centuries ago, it left an imprint on the inside of the skull,” explained study co-author Christy Anna Hipsley from the Natural History Museum Denmark in a statement. “By digitally reconstructing those spaces and comparing them with the brains of living birds, we can begin to piece together how the dodo sensed, behaved, and interacted with the world around it.”

Pigeons and doves (order Columbiformes) are the dodo’s closest living relatives, and therefore, a relevant comparison point. By comparing the dodo’s brain imprints with those of 36 living Columbiformes, the researchers found the forebrains, which handle memory and cognition, to be proportionately similar in size. Pigeons have remarkable cognitive abilities, such as long-term image recognition; so, the behavior of dodo birds that wobbled over to 16th-century sailors may have stemmed more from their predator-free history than any lack of intelligence.

One distinction of dodo brains was their smaller optic lobes (more than 3.5 times tinier than those of pigeons), indicating vision that works in low light (for



example, dawn, dusk, and the dead of night). Furthermore, the high-domed skulls of dodos were reminiscent of certain owls. Their olfactory bulb imprints also hinted at keen senses of smell that, coupled with a pronounced upper beak for probing their surroundings, would have supported low-light activity as well.

Nothing dumb about that.

—Devin Reese

PSYCHOLOGY

What Do You Think About All Day?

THINK, THEREFORE I AM,” said René Descartes. “We become what we think,” said the Buddha. Across the crush of time, some of our most recognized philosophers have come to believe that who we are is, in large measure, determined by our thoughts.



But many people have only a glancing handle on where their minds go on any given day. And if our own thoughts are mysterious, the preoccupations of the people around us are still more obscure.

In a new study, economists from Claremont Graduate University in California and Wuhan University in Hubei, China, report that despite the central place of thoughts in consciousness and the human experience, and the impact thoughts have on happiness, the contents of our thoughts are surprisingly understudied.

“So much of what it means to be human is about what’s going on in our head,” economics professor Joshua Tasoff, a senior author on the study, told *The Guardian*. “The fact that we haven’t really looked at that as economists is really quite extraordinary.”

The scientists asked 258 Americans between the ages of 19 and 71 to respond to periodic surveys about their thoughts. Over a two-week period, they buzzed them 6 to 8 times a day on their phones. When alerted, the participants were asked to record the thought most dominant in their awareness, slot it into one of 12 categories, determine whether it was an intrusive thought or a welcome one, describe the activity they were engaged in when they had the thought, and rate their happiness. The researchers collected a total of more than 23,000 snapshots of peoples’ inner monologues.

It turns out that on a daily basis, a lot of what people think about is, “I’m tired, I need to work, I want food.” Sensory impressions related to the body accounted for about 20 percent of all the thoughts recorded. Work accounted for 17 percent and food about 14 percent.

When it came to thoughts related to desire, half of people taking the survey identified food or drink as the object, especially coffee. There was little to no mention

A lot of what people think about is, “I’m tired, I need to work, I want food.”

of sex, and in fact, relationships accounted for just 10 percent of thoughts, as did thoughts about art and hobbies, even though these were among the most closely associated with feeling happy. About half of peoples’ thoughts were intentional, as opposed to accidental or intrusive.

In the category of relationships, friendships and pets may exert greater pull on the mind than romance, children, or coworkers. Friends were the most common subject of thought in the relationship category, occurring at a frequency equal to that of the four romantic relationship terms combined: husband, wife, girlfriend, and boyfriend. Some 46 percent of participants mentioned thinking about a friend at least once.

Pets, meanwhile, seem to live rent free in our minds, outranking children in people’s thoughts. Dogs and cats combined occurred slightly more often than son and daughter combined. While thoughts about dogs were associated with walks, thoughts about cats were related to needs and wonder. Apparently we’re all wondering, “Does Fido need to go to the park?” “Is Coco hungry?”

The dearth of reports about sexual thoughts aroused the skepticism of the researchers. “There may be things people don’t want to share with us. You know: TMI,” Tasoff told *The Guardian*. Sex thoughts only cropped up 11 times, and “no thoughts” or “empty mind” was only slightly more common.

Thinking styles may also cluster into one of four buckets, the scientists found. There are those who worry about the world, including politics, social issues, and wars; those who are preoccupied with household chores and finances; those who focus on the mind and the body—a largely female group, whose thoughts were more intrusive than the average; and those who tend to primarily mull career and leisure in equal measure. The authors nicknamed this last group the work-hard play-hard group. It skewed slightly male and had the highest average income and levels of happiness.

Marcus Aurelius once wrote, “The happiness of your life depends on the quality of your thoughts.” As the new research shows, the Roman emperor and philosopher was on to something.

Thoughts better explained happiness than a person’s income, age, education, or even what they were literally doing at the moment. Intentional, chosen thoughts made people happier than mind-wandering or unwanted intrusive thoughts. And, thinking about music, hobbies, and food gave people a happiness boost roughly equal to the happiness provided by being employed or married.

So, our thoughts and those of the people around us may not be as exotic as we might think—and it’s not the life you lead that determines your happiness, but rather the voice narrating it.

—Kristen French

MICROBIOLOGY

The Humble Amoeba Is Smarter Than You Think

FRIEDRICH NIETZSCHE THOUGHT a lot about the little amoeba, writes mycologist, professor of biology, and author Nicholas Money in his new book, *A Is for Amoeba*. Nietzsche decided that when the single-celled brainless blob reached out with the arm-like extensions on its cell membrane, known as “false feet,” to grab food or move across a surface, that it was expressing a cosmic drive to “grow, seize, become predominant.”

Whether the amoeba is, in fact, driven by such urges isn’t known to science. In fact, this simple form of life, which has up to 100 times more genetic material than a human, is more of a mystery to us than you might think. Money has spent most of his career studying fungi, but to answer certain elemental questions about life on Earth, he decided to turn to the simple amoeba. The riddles he found surprised him.

You write in the introduction to your book that it was inspired in part by a deepening conviction that you can see a world in something as tiny as William

Blake’s “grain of sand.” What deepened that conviction for you?

When you’re trying to understand big questions in biology about how things grow and reproduce, how organisms have evolved, it makes sense to start with something small, a single-celled organism. And the *Amoeba proteus* is this iconic organism that almost everyone recognizes. You could be in a classroom in North Korea, and if you draw a squiggly outline with a dot in the middle, everyone will recognize it. So that seems like a good place to begin.

But I made a bit of a mistake because as you learn more about any organism, particularly amoebas, you realize they’ve got their own tricks and sophistries. Amoebas are quite difficult to understand in their own right. Scientists have been trying to make sense of them for a lot longer than a century. And there are still all these questions that we cannot answer—for example, about amoeboid movement, which is something that’s spread throughout the tree of life and is in our bodies, too. I find that sort of frustrating, but also inspiring.

There's still so much to learn about how cells grow, about how cells move.

We still don't really have a good understanding of how amoebas crawl or swim. Why is it so difficult to get a handle on this basic feature of a creature that seems so simple?

We've done a really good job, just during my life as a scientist, of teasing out a lot of the details. We know a lot about the molecular machinery. We know a lot about the genes involved. But putting that together into a satisfying and comprehensive picture is something that still eludes science.

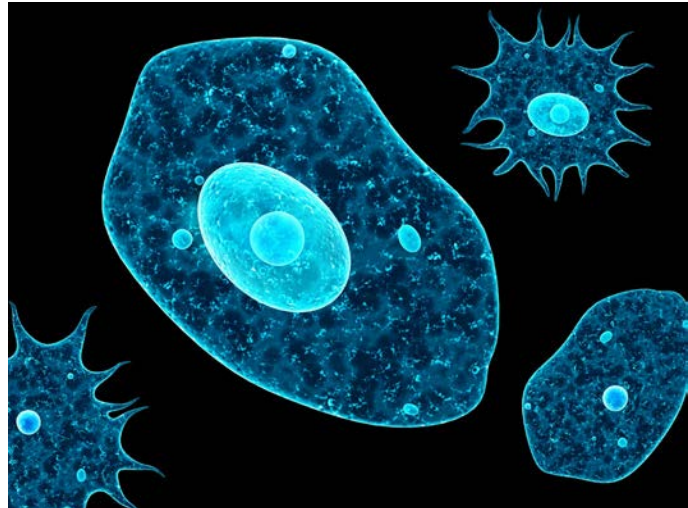
Amoebas reproduce by splitting in half and copying themselves, which grants them a kind of genderless immortality. Does an amoeba "die" in any meaningful sense?

Of all of the thousands of species of amoebas, a lot of them are related to *Amoeba proteus*, the iconic cell, and those don't reproduce sexually. Some of the genetics supports the idea that this represents an advancement in evolution: that exclusive reproduction through an asexual mechanism, at least in amoebas, might be a secondary event in evolution—that at one time their ancestors reproduced sexually, and they lost that in favor of living these solitary lives and just reproducing by mitosis, by cell division.

Whether they die depends on your definition of immortality. The mother cell divides and forms two daughter cells that are identical or almost identical to the mother in terms of their genetics. They will then divide. As long as that lineage continues, as long as the DNA present in that mother is transmitted through time, one could regard that as a form of immortality. That doesn't work for humans. The best that we can do, if we have biological children, is give them half of our genes, and that percentage diminishes with each generation. So yeah, I think amoebas have immortality. They're closer to being immortal than any animal.

Do amoebas challenge our understanding of the nature of consciousness and intelligence?

They do. The thing is that we lack compelling universal definitions of intelligence and consciousness. But certainly, amoebas are exquisitely sensitive to their



surroundings. And they act in a fashion that can be compared with animals with brains, in the way that they chase their prey. They eat these little pond things called ciliates that swim with hairs. And amoebas will chase them. Then, when amoebas are attacked by predators, they'll go to immense lengths to evade their predators. They're showing this sensitivity in everything they do.

Does that represent consciousness? One could argue that they have a sense of self, that they possess a minimal level of selfhood, but also they seem to have a kind of memory. They show conditioned responses over time. If they're stimulated in one particular fashion on repeated occasions, they'll continue to respond in the same way. It's a little bit like Pavlov's dogs.

It's difficult for us to really think about the consciousness of an amoeba, but since we lack a definition, I prefer to stretch the idea of consciousness along a continuum. And then we can add organisms. Like we can say, "Well, what is it that amoebas can do? And what is it that fungi can do?" It's a really interesting area of investigation, and there's lots of experiments going on in labs in different parts of the world on the behavioral sophistication of single cells.

Did you learn anything about yourself while writing the book?

Yes. I'm more amoeboid than I imagined. I've got a kilogram of white blood cells in me that move with amoeboid motion: neutrophils and macrophages and

dendritic cells. They're moving like amoebas in all of us all of the time. And they're clearing away our dead cells. They're combating bacterial infections. They're destroying bacteria that would otherwise grow in our bodies and removing cancer cells, too.

It's difficult to wrap your mind around that. That's an almost impossible thing to do in an imaginative

sense—to close your eyes and think about yourself as this ecosystem made from trillions of cells, tens of billions of which are acting as amoebas. It's breathtaking, but it's confusing, too. It's discombobulating. It's difficult to grasp that relationship that we have with the microscopic world.

—Kristen French

NEUROSCIENCE

Does Watching TV Really Rot Your Brain?

IT'S A PHRASE THAT'S reverberated across living rooms for decades: “Turn off that TV!” As a dad to three kids, I need look no further than my own couch to find an abundant and persistent crop of potatoes staring blankly at that siren of a screen. Personally, I try to vary the clause that comes after the main exhortation: “There has to be something better to do,” “It’s too beautiful out to be inside,” or “I swear, I’m taking that thing off the wall” usually do the trick. But often I harken back to my ancestors and deliver the classic parenting chestnut: “That thing will rot your brain!”

I, like many well-intentioned parents before me, levy that particular claim without ironclad scientific evidence supporting its veracity. But now, researchers studying TV watchers—albeit of a slightly older vintage—say viewers may actually suffer some brain ills from the habit.

Publishing their findings in *Alzheimer's and Dementia: Journal of the Alzheimer's Association* this July, a team of scientists analyzed data from more than 1,700 people who were asked, between 1987 and 1989, how frequently they watched television to pass the time. The study subjects averaged 53 years old and were part of a long-running study of cardiovascular and brain health risks in communities across the United States. The authors of the paper followed up with the study participants more than two decades later and scanned their brains using an MRI machine.

The researchers found that people who had indicated they watched TV “very often” showed marked differences in brain anatomy compared to people who



reported they watched TV “never” or “seldom.” The brains of the frequent TV watchers had smaller occipital and frontal lobes, which control executive function and visual processing, than their peers. And they also showed reduced volume in brain areas associated with memory and more damage in patches of white matter that are tied to cognitive decline and dementia.

“For years we’ve focused on how much people sit,” David Raichlen, University of Southern California anthropologist and a senior author of the study, said in a statement. “Our findings suggest we should also pay attention to what they’re doing while they’re sitting.”

Now, there are limitations to this study, some of which the paper’s authors acknowledge. The findings were based on self-reported behavior rather than direct

JRCASAS / ADOBE STOCK

People who said they watched TV “very often” showed marked differences in brain anatomy.

measurements of subjects’ TV-watching, “which may involve some recall bias,” they wrote. Also, study participants’ brains were not initially MRI scanned, meaning baseline measurements of brain structures were not established. And I feel compelled to note that the pull of TV programming in the 1980s was much less robust than the virtually limitless options available to the modern-day viewer. Also, these people were middle aged at the time they were asked to recount their TV viewing habits. These habits may or may not have differed in their youth.

These caveats aside, the researchers did control for other factors, such as overall physical activity, body mass index, alcohol use, and other potential confounders in their analysis, and the brain differences held.

Additionally, they factored in people who were sedentary, but sat at work rather than in front of the boob tube. The brain changes weren’t as drastic in the sedentary workers.

“We frequently encourage the public not to spend too much time sitting down, but experts may want to expand that recommendation to encompass the activities done while sitting, since those seem to have distinct impacts on brain health,” said study coauthor Natan Feter, a USC postdoctoral student.

The authors also note that their findings largely mesh with previous studies conducted in the United Kingdom, where adults who watched more TV tended to have poorer cognitive performance and higher risks of dementia.

These findings are certainly not definitive in supporting my assertion to my children that watching TV is not the best use of their time. But I am grateful that they lend some degree of scientific heft to the suggestion, from me and many parents throughout the past 7 or so decades, that the habit likely does no favors to one’s brain health.

—Bob Grant

SOCIOLOGY

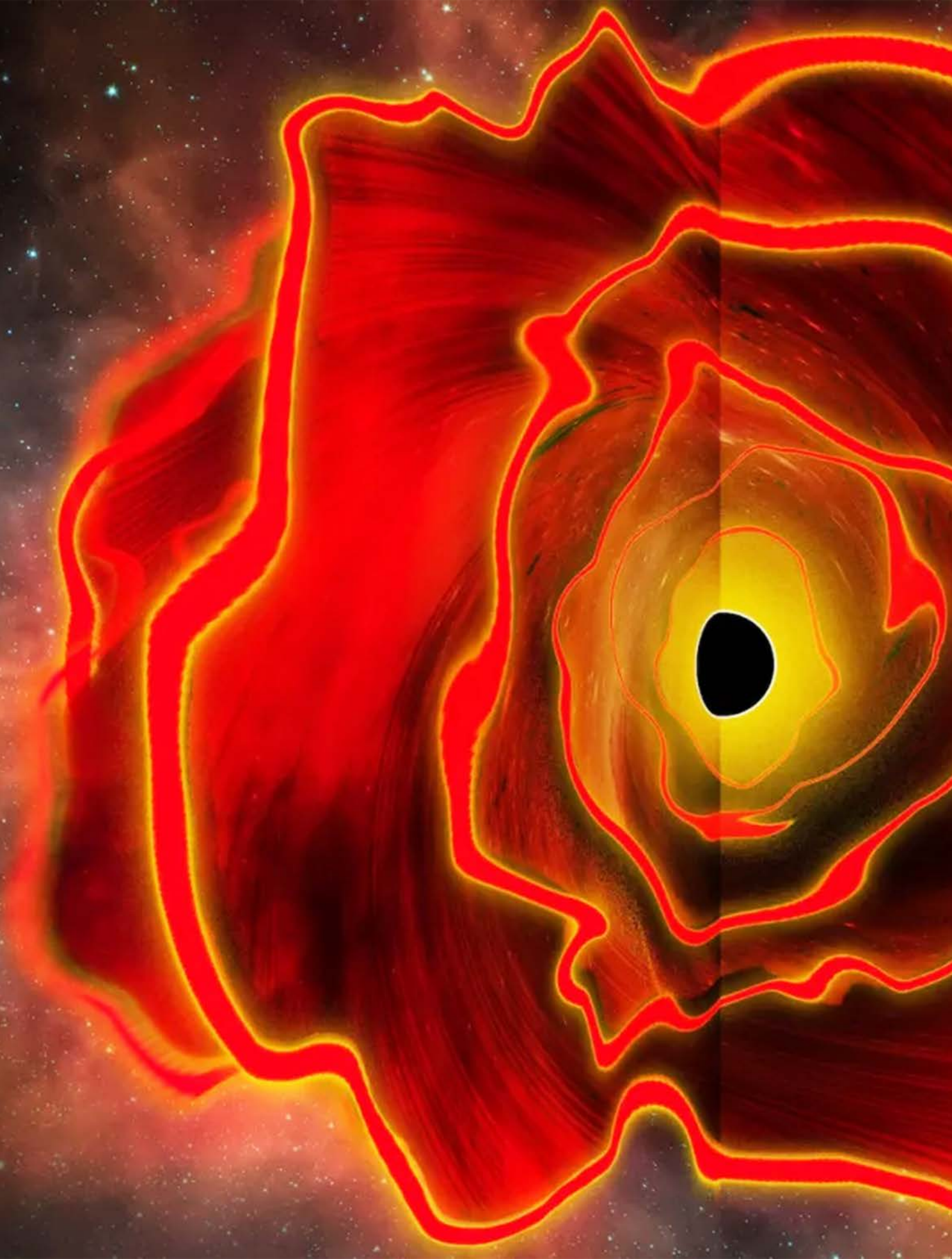
Quoth the Garbageman

Wisdom from Simon Paré-Poupard, garbageman, sociologist, author of Trash! A Garbageman’s Story

“Waste management companies have built their business model on the idea that people don’t want to deal with their own waste. That’s what waste management as a concept will always offer cities and governments: the easiest way not to care about where our waste goes. What I’d like to say to people is this: If you throw something in a trash bin or recycling bin, you should be required at some point in your education to visit a site that processes that waste. That would create some real consciousness about the systems that have been built to keep our waste invisible. That’s one thing city councils or states could require citizens to do. We need to feel responsible for the waste we generate.”

—Interview by Kristen French





ASTRONOMY

“Black Hole Star” Discovered at the Dawn of Time

RECENTLY, ASTRONOMERS using the James Webb Space Telescope (JWST) to peer deep into the history of the universe noticed something strange—a very bright, very red dot. Now, the team behind its discovery has published their analysis of the cosmic oddity in *Nature*, and it’s unlike anything ever seen before.

According to the international team of astronomers, the bright red dot appears to be a combination of a massive black hole and a star, around the size of our entire solar system.

The team arrived at that conclusion after modeling a variety of scenarios to explain the object’s peculiar characteristics.

Typically, redness in a celestial object is a telltale sign of cosmic dust, similar to how wildfire smoke turns the sunset red, researchers say. But other signatures from the object didn’t match up with the dust hypothesis. The astronomers also didn’t detect any heavy metals usually produced by mature stars—only the simple elements hydrogen and helium. The only thing that could explain the intense light streaming from the object is a black hole at its core, around 100,000 times more massive than our sun.

If their findings are confirmed, this would be the first “black hole star” ever discovered, named MoM-BH*-1. A beacon from the dawn of time itself.

—Jake Currie

Image by Jose-Luis Olivares, MIT



The existence of free will meant that we couldn't know the future. And we knew free will existed because we had direct experience of it. Volition was an intrinsic part of consciousness.

Or was it? What if the experience of knowing the future changed a person? What if it evoked a sense of urgency, a sense of obligation to act precisely as she knew she would?

—Ted Chiang, "Story of Your Life"

Author Ted Chiang has been a Miller Scholar at the Santa Fe Institute since 2022. The Miller Scholarship is awarded to highly accomplished, creative thinkers who make profound contributions to our understanding of society, science, and culture.



SANTA FE INSTITUTE

SEARCHING FOR ORDER IN THE COMPLEXITY OF EVOLVING WORLDS
SANTAFE.EDU



FEATURES

Coincidences in My Life Have Me Wondering

My search for a hidden structure to astronomically unlikely occurrences

BY GEORGE MUSSER



AFTER COLLEGE, I WENT BACKPACKING through Africa, and at the Khartoum airport in Sudan I met an American guy named Chris. We explored together for a few days and went our separate ways. A month or two later, I stopped by the main post office in Nairobi, and there he was again: Chris. We chatted and parted. A few months on, I was playing soccer with some boys in a village in Malawi, and who should show up but Chris. A few weeks after that, I was trying to hitchhike and having no luck, when a car finally stopped. I need hardly say who was in it.

A year later, on a separate trip, I was passing through Hawaii and met a Japanese guy at my youth hostel. We went in search of ice cream and realized we had traveled through Africa at the same time. Here, finally, was a way to assess the odds of my repeat Chris encounters. Backpackers do tend to follow a few well-worn paths, so maybe it wasn't all that unlikely. But my new friend and I had never met in Africa, nor did he know Chris, nor had he repeatedly crossed paths with anyone else. So, my experience with Chris really did seem uncanny. At that moment, someone called my name, and on a street corner in Honolulu, I turned to see Chris.

I am a firm believer in materialism. For all my wonderment at nature—the thrill of seeing Jupiter in the night sky or hearing frogs cackling in a swamp—I have never felt a need to invoke supernatural forces or beings. To me, the magic of the universe is that it needs no magic. A few simple rules and a few basic ingredients beget a sublime diversity. But I try not to be dogmatic about it, and episodes such as my serendipitous encounters with Chris remind me that it is not irrational to wonder whether there is more to this world than atoms and void.

I've never felt a need to invoke supernatural forces. The magic of the universe is that it needs no magic.

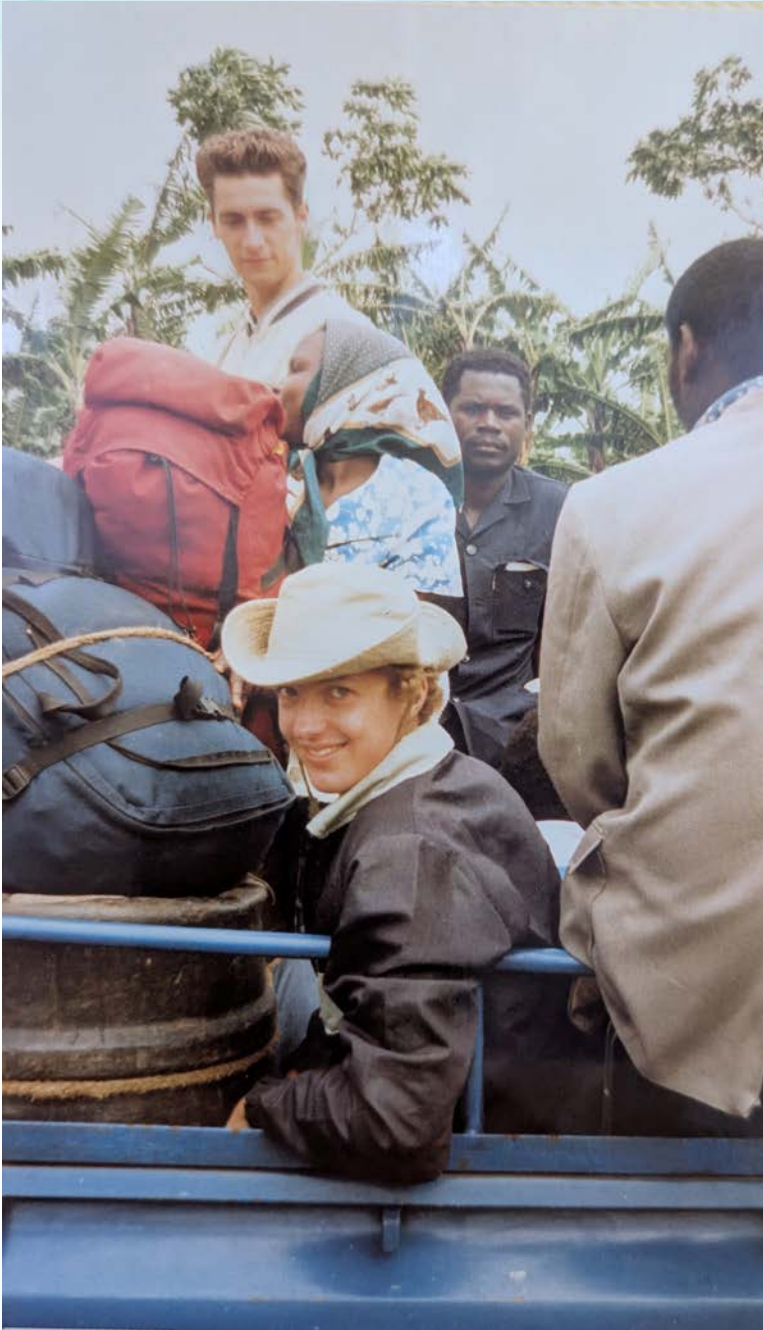
A 2021 episode of *This American Life* compiled some jaw-dropping coincidences. A guy named Blake wanted to change his phone wallpaper and mentioned it to a friend, Camille. He asked her to send a photo of something, and as a joke she sent one of herself as a kid. Behind her was a street scene, including a woman passing by. Blake looked more closely. It was his grandmother! Heightening the unlikelihood, the photo wasn't taken in Blake's or Camille's hometowns. It was taken in British Columbia, where both families happened to be vacationing at the same time. It was coincidence upon coincidence.

"When coincidences pile up in this way one cannot help being impressed," Carl Jung wrote in his essay "On Synchronicity," from 1951. Like many '80s kids, I first learned of the term from the Police album *Synchronicity*,

on which Sting sang of "a connecting principle / linked to the invisible." Jung famously recalled a psychotherapy session in which a woman was recounting to him a dream of scarab jewelry, when they heard a rapping on the window—it was a scarab beetle. Nothing caused this coincidence, he accepted, but the events could be connected through some other means—some structuring of nature that is not captured by the laws of physics. We may be participants in a larger pattern. In support, he also cited experiments on ESP, which many scientists at the time still took seriously.

Skeptics are less impressed. They point out that there are so many people on this planet, doing so many things, that the rare becomes routine. A one-in-a-billion event may happen several times a day. In 1989, statisticians Persi Diaconis and Frederick Mosteller codified

COURTESY OF THE AUTHOR



YOU AGAIN? At left, author George Musser (white hat) catches a ride in the bed of a pickup truck (*matatu*) in southwestern Uganda in mid 1989. At right, Chris (pink shirt) on a street in Addis Ababa. The two strangers met again and again, by coincidence, during their African travel.

this wisdom as the “law of truly large numbers.” With a large enough number of opportunities for an event to occur, even extremely unlikely events become probable. Although the odds of winning a national lottery can be 1 in 100 million, it’s probable somebody will win, simply because a truly large number of people buy tickets.

Flukes turn out to be even more common when we consider that people tend to count a near-miss as a direct hit. People also conveniently forget all the times there could have been a coincidence but there wasn’t. Jung’s anecdote is less impressive when you consider that insects were probably banging up against his window all the time, and he noticed them only when he was primed to do so.

For some skeptics, all the talk about coincidence is just another example of human irrationality. *Why can’t people grasp probability theory, already?*

But others are more sympathetic. Our brains evolved to be on the lookout for patterns and, from them, make split decisions to discern how the world works. Only by grasping cause-effect relationships can the brain decide how to act to achieve its goals. But the brain has no independent way to tell a significant occurrence from a random blip. It has to make a judgment call. It can’t entirely avoid the risk of reading too much into randomness without missing out on genuine features of the world. If you plug your ears to noise, you’ll never hear music, either.

PSYCHOLOGISTS MARK JOHANSEN and Magda Osman, as well as Thomas Griffiths and Joshua Tenenbaum, have argued that our fascination with coincidences indicates that the system is striking an appropriate balance. Deviations from the usual course of events rightly catch our attention. They feel spooky when a conventional explanation falls short and nothing plausible takes its place. That leaves us hanging in a state of unrequited curiosity.

Osman told me a story of when a colleague she hadn’t heard from in years came into her mind, and that person emailed her the next day. Forced to choose between two unlikely explanations, pure chance and some kind of ESP, she reluctantly opts for chance. “I don’t treat it as sufficient evidence to reorientate my entire belief system around accepting psychic phenomena as real,” she said.

I feel the force of arguments that coincidences are probabilistic happenstance—but don’t feel entirely assuaged. A one-in-a-billion event may happen to someone somewhere several times a day, but it’s still unlikely that it will happen to you. That’s why your mother warned you not to play the lottery. So, when an event like that does happen to you, it’s rational to look for a reason besides pure chance. A coin that lands on heads 30 times in a row is surely a trick coin.

The skeptical position gets hand-wavy when it comes to estimating personal probabilities. How do you even begin to assess the odds of my encounters with Chris, or Blake spotting his grandmother in a friend’s old photo?

The best attempt I’ve seen is by a non-skeptic, Sharon Rawlette, a philosopher who is sympathetic to paranormal explanations. In her book *The Source and Significance of Coincidences*, she gamely took a stab at estimating the probability of her own Jung-like episode with a scarab: low but not insanely so. She also noted that skeptics take it for granted that an incredible event happens in the population at the predicted frequency. Do we really know that? It would be nice to collect statistics on coincidence reports. If we found that those one-in-a-billion events happen several hundred times a day, we might have evidence for a genuine phenomenon.

But to question the skeptical arguments is not to accept ESP. We don’t need to swing between the extremes of spurious correlation or supernatural occurrence, because there is a third option. Some coincidences might be caused, or at least their odds shortened, by actual hidden structures in the world.

After all, Jung was right in one sense. We *are* participants in a larger pattern: human culture. If human beings tend to over-ascribe causality, we also tend to overstate our individual autonomy and originality. Skeptics call the former cognitive bias but fall prey to the latter; their probability estimates assume the statistical independence of events that may well be related. In fact, almost everything we think or do is a remix. We are starlings in a “human murmuration,” as social psychologists Kevin Durrheim and Michael Quayle put it in a paper last year. I asked them whether social dynamics might plausibly give rise to coincidences, such as meeting a friend serendipitously on vacation, at a rate

Like ants, when people go somewhere or do something, they leave trails that entice others to do the same.

greater than chance, and they said they might. “The events you describe are not entirely statistical flukes,” Durrheim told me.

When Diaconis and Mosteller presented the law of truly large numbers 37 years ago, they acknowledged the general possibility of hidden causes but didn’t explore it. They had no way to measure them. Now we have vast data sets on social behavior and powerful new tools to analyze them, not the least of which is large language models.

WHEN MY WIFE AND I named our daughter, we thought we were being original. Now I keep running into women of her generation with her name—it’s kind of disconcerting. How did that happen? We didn’t choose from a baby-name book or follow a celebrity’s example; we just liked the name. On a grander scale, almost every major scientific discovery and technological invention in history has been made by multiple people working on their own at nearly the same time. And as George Harrison or Ed Sheeran will tell you, pop music is filled with similar-sounding songs that artists swear they created independently.

The causation in these cases isn’t clear. We are reduced to saying that an idea is in the air. But most people don’t read too much into them—the best online compilation of coincidences, the Cambridge Coincidences Collection, doesn’t bother to include them—because they recognize there is a reason

behind these coincidences even if they don’t know what it is. We all draw from the same well of cultural influences: the same names, the same scientific tools, the same chord patterns. In fact, simultaneous invention seems almost inevitable. These cases offer proof for a culture-based explanation of coincidences in general.

Theory suggests some ways this might work. Many sociologists and social psychologists think that ideas spread like a contagious disease. At its simplest, the spread is direct: An ad or influencer tells you to buy something, and you do. More sophisticated versions tangle up the causation, so that an idea propagates too subtly to be noticed at first; by the time it is noticed, it is already widespread. Unable to see the early stages of propagation, we think the idea has popped up in multiple places spontaneously, in what looks like sheer coincidence.

In his 2018 book, *How Behavior Spreads*, Damon Centola outlined the theory of “complex contagion.” He argued that people seldom adopt an innovation as soon as they hear about it. They need corroboration. Multiple influences must come together to get you to do something or hold onto an idea. You may not even be aware of all these influences; by the time you finally respond, you may feel like you’re acting on your own initiative. Other people are subject to the same diffuse influences and react in the same way. It looks coincidental, but a common cause is responsible.

Other social phenomena can also conceal causation. Nicholas Christakis and James Fowler, in their 2009 book *Connected*, described indirect social influences. People we have never met can alter our behavior in ways that we don't realize at the time. Some people can be attitudinal Typhoid Marys, transmitting a belief or behavior without subscribing to it themselves. "A person might be affected and not even know it," Christakis told me. For instance, if some of your friends become anti-vaxxers, you might personally resist—you dutifully continue to get your shots—but their attitudes nonetheless affect you, causing you to give up on arguing with them and let their misinformation spread unchecked.

In a further elaboration on the contagion model, Amir Goldberg and Sarah Stein conjectured in a 2018 paper that what spreads is not an idea, but the set of associations that form the underlying psychology of the idea. Anti-vaxxers may associate eschewing vaccines with a more natural lifestyle or rebellion against an unfeeling medical establishment. As those general sentiments spread, so might anti-vax beliefs. This, rather than the explicit transmission of misinformation about vaccines, can be the decisive factor leading people to engage in similar behavior. Because the behavior isn't directly copied, it looks coincidental.

Durrheim and Quayle suggested another indirect channel of influence, which social scientists have begun to explore only recently: "stigmergy." Like ants laying down pheromones, when people go somewhere or do something, they leave trails that entice others to do the same. For example, when tourists visit a location, the locals set up restaurants and hotels, which make it easier for other tourists to follow. Online, the "also bought" listings on Amazon and stream counts on Spotify create self-reinforcing loops of popularity. Stigmergy carves cultural grooves that we slot into, sometimes unaware. This may help to explain my repeat encounters with Chris: We were both responding to choices imprinted on the world by previous travelers.

Network scientists have developed methods to reverse the direction of propagation and trace a contagion back to its original source—for instance, when looking for patient zero in an epidemic or assigning blame for misinformation campaigns. Through these techniques, we might hope to trace the elements of a

coincidence back to their common cause. But such a retroactive analysis is fundamentally hard to do, as I soon found when I began to conduct a few experiments of my own.

WHILE RECOVERING FROM a wrist injury last year, I couldn't play guitar, so I took up the harmonica. One evening I decided, apropos of nothing, to teach myself the melody of "Mr. Tambourine Man." It was trickier than I thought. I tried to use notes from the guitar chords, but they didn't fit, and I ended up having to refer to the sheet music for the melody.

The next morning at 11 a.m., I took a break from work to peruse Bluesky, and the first thing I saw was a post from 10:50 a.m. about "Mr. Tambourine Man." In it, Ethan Hein, who teaches music at New York University and the New School, explained that Bob Dylan intentionally misaligned the harmony and melody. For instance, he sang the word "Man" on A while playing a C chord on the guitar.

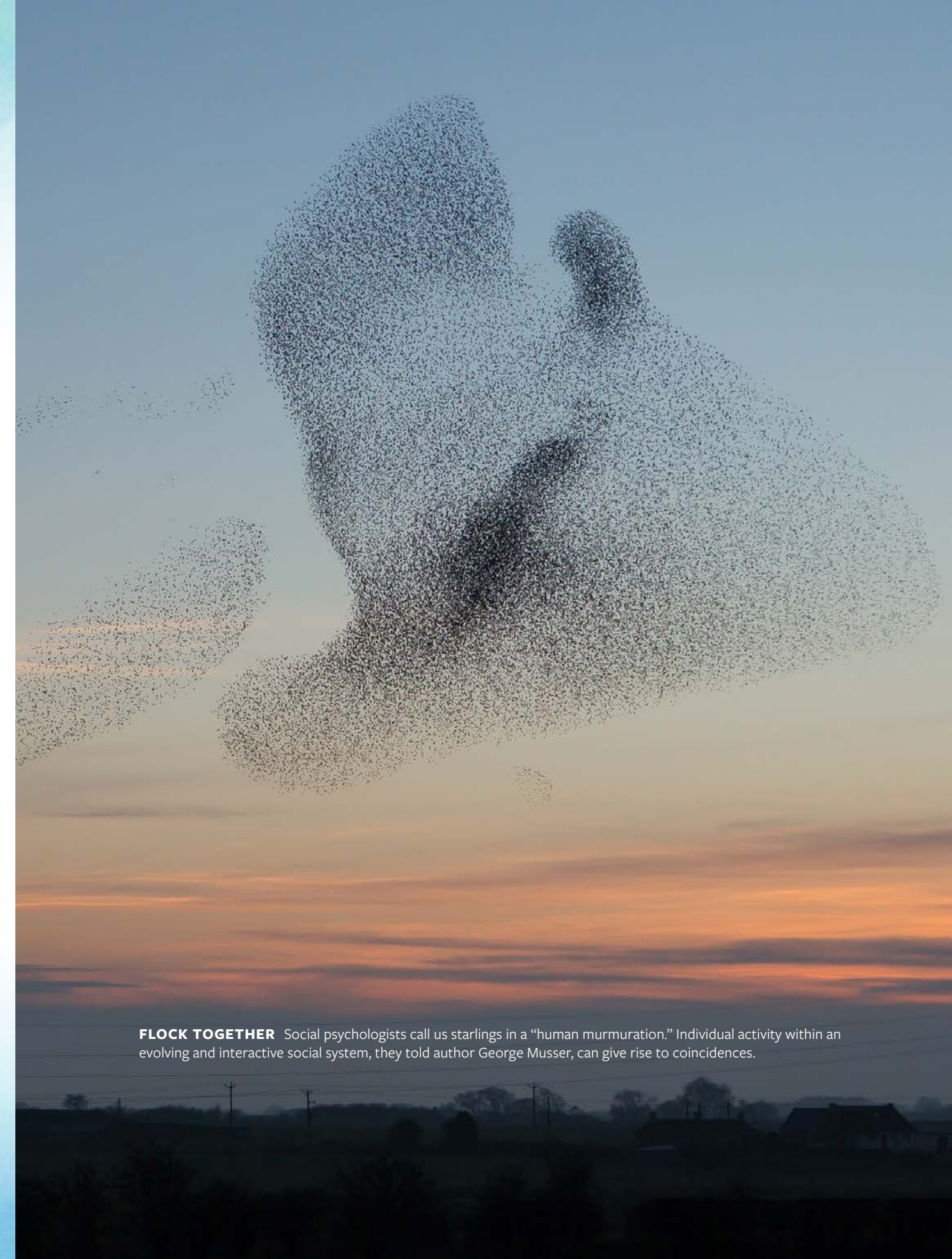
I was gratified that my difficulty with the song wasn't a personal failing; I had stumbled onto what music theorists call "melodic-harmonic divorce." Dylan picked up the style from the blues and, in turn, inspired subsequent pop songwriters. But what a spooky coincidence. It didn't seem like a simple twist of fate.

One thing I could rule out immediately is online tracking. The Bluesky feed is purely chronological, without any algorithmic curation, so the post had nothing to do with my Google search the night before. Had I been on Facebook, which is much more stalker-y, I would no doubt have been served ads for tambourines.

To look for a common cause that led both Hein and me to be interested in the Dylan song at the same time, I turned to large language models. They are pattern-detectors par excellence and make connections between realms that seem totally unrelated. They presumably identify diffuse cultural influences that humans might not be consciously aware of.

"We think using LLMs to study culture and cultural dynamics is the next big thing—a Large Hadron Collider for social and cultural research," Russell Neuman, a professor of media technology at New York University, told me.

But as a tool for forensic coincidence-ology, I quickly discovered, they have a long way to go. The trouble is

A large flock of starlings is captured in mid-flight against a sky transitioning from a pale blue at the top to a vibrant orange and red at the horizon. The birds are densely packed into several large, fluid, and somewhat abstract shapes, with the most prominent one resembling a large, rounded 'M' or a cluster of interconnected loops. The individual birds are small and dark, creating a textured, shimmering effect as they move in unison. In the lower portion of the image, the dark silhouettes of trees and rooftops are visible against the bright horizon, providing a sense of scale and location.

FLOCK TOGETHER Social psychologists call us starlings in a “human murmuration.” Individual activity within an evolving and interactive social system, they told author George Musser, can give rise to coincidences.

All this thinking about coincidences made me think about Chris. I wondered what he was doing today.

extracting relevant information from them. When I simple-mindedly asked ChatGPT-5.5 why the Dylan coincidence occurred, it responded with a generic explanation of the probability of coincidences—basically regurgitating the skeptical arguments of Diaconis and Mosteller. I couldn't come up with a prompt that would bring out hard-to-see correlations between two people online.

One reason may be that these models extract general patterns and discard the details. (That's also why they hallucinate. If you ask them to reconstruct specifics from generalities, what else can they do but take a guess?) As a simple test, I gave ChatGPT-5.5 a database of Bluesky posts by Hein and me in which I had planted fake messages: a prior exchange about melodic-harmonic divorce. The model failed to flag this exchange as a possible common cause for the coincidence; it still attributed the episode to random chance.

But the model may still encode subtle cultural patterns in some form. Maybe you can get at them if you directly access its innards. In a fascinating study in 2024, a team at Anthropic looked inside Claude to see how data was organized within it. Although the goal was to figure out how the model works, the result was also a cultural analysis.

During training, the model had discerned millions of conceptual categories, ranging from tangible ("Golden Gate Bridge") to abstract ("internal conflict"). When you enter text into the model, it parses it in terms of these categories. Asked "What is it like to be you?" the model placed itself in the category of "immaterial or non-physical spiritual beings." That doesn't mean the model had any self-awareness; the response was based not on introspection, but on cultural associations with the question. Perhaps a future study can look for categories that link seemingly independent occurrences.

Stymied by my experiments with ChatGPT, I mentioned the problem to some data-science friends, and they suggested analyzing social-media data using CommuNalytic, a website that can download and parse Bluesky, Reddit, and Telegram posts. I pulled up all mentions of "Mr. Tambourine Man" on Bluesky in 2025. They were highly clustered in time and cited specific triggers. Dylan's manuscript lyrics were sold at auction in January, he performed the song for the first time in 15 years in May, and the Byrds' cover version was released 60 years ago in June. Very few mentions suggested a spontaneous random musing about the song. I searched for people commenting on earworms in general, and about three-quarters identified a specific trigger. So,

this does indicate that most coincidental posts about a song can be traced to some common cause.

I emailed Hein and he said he'd been inspired to comment on the Dylan song by an episode of Andrew Hickey's *500 Songs* podcast in 2021. That explains his interest, though not the coincidence in timing. Looking at his and my posting history, although he and I have followed each other from the early days of Bluesky in mid-2023, we had never responded to each other's posts before. In fact, given the number of people I follow and my Bluesky viewing habits, it's quite likely I had never seen a post of his until the coincidental one. But we have similar politics (which is already evident from being active on Bluesky) and inhabit much the same intellectual milieu. Whenever I did finally see a post of his, it's likely it would have meshed with an interest of mine—if not the Dylan song, then something else. That's probably enough to explain the coincidence.

SO WHERE DOES THAT leave me? I opt for nuance. I know that's not what the internet wants, but by God, I don't care. "He refused to do a hot take"—write it on my tombstone.

Some coincidences strike me as so astronomically unlikely that they demand a causal explanation. Maybe, on inspection, they will prove to be random chance. But we need this inspection. We could set a threshold for further investigation, just as scientists in every discipline flag effects of a certain statistical significance (although they debate where to draw the line). Those that pass would become candidates for sociological analysis.

Coincidences are catnip to the curious mind, which is reason enough to study them. But they could serve a broader theoretical goal, too. Scientists and philosophers use the word "emergence" to describe how individuals, working together, give rise to collective behavior: molecules in a gas, neurons in the brain, consumers in the economy. In most cases, though, they haven't worked through precisely how it works, and the word "emergence" is just papering over ignorance. Social scientists call this woolliness the "micro-macro problem." Computational models and data science should help to make the relation between these two levels of description less opaque. Coordinated behavior is the hallmark of emergence, so coincidences are, if anything, to be expected.

We may well find that Jung was on to something. Shaikat Hossain, who studies audio perception and has a side interest in Jung, suggested in 2012 that the internet is a kind of Jungian collective unconscious. That has no mystical intimations. It just reflects the emergent workings of human society—a subterranean level of cultural influences that operate beneath our direct awareness and play a part in shaping our fate. As Jung stressed, a coincidence can serve a psychological purpose, whether or not we apprehend a cause for it.

All this thinking about coincidences made me think about Chris. I wondered where he was and what he was doing today. All I had for him was a 30-year-old address. It was enough. I found his phone number, left a message, and he called me back. In this age of junk-call screening, that was miraculous in itself. We immediately started bantering as though no time had passed. We soon realized that our coincidence streak had run out. Apart from each having one child, of the same age, our lives had diverged.

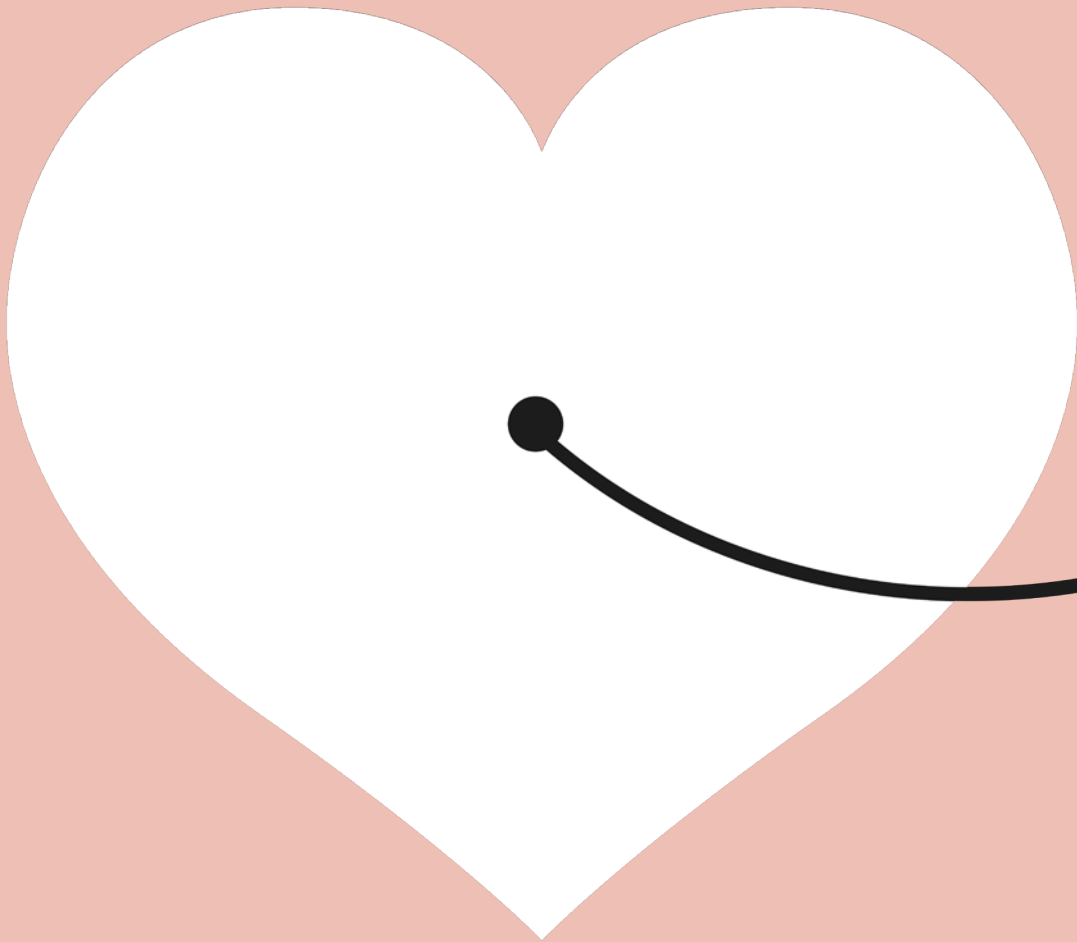
But I surmised that our encounters had left even more of an impression on him as they had on me. He evocatively described how, the night after our meeting in Nairobi, he felt God call him to be a doctor. Not long after our meeting in Hawaii, he entered medical school. In giving up his travels, he probably gave up the chance of further coincidental meetings in unexpected places. But the coincidences had already served their purpose. ☺

GEORGE MUSSER is a science writer and editor based in New Jersey. He is the author of several books on fundamental physics, most recently *Putting Ourselves Back in the Equation*, about observer effects in physics and cosmology.

The Inventor of the Thinking Machine Didn't Worry. Neither Should You

In this age of AI anxiety, listen to your heart

BY YUVAL AVNUR



These are real problems. But beneath this is an underlying fear of being *personally* replaced, not just as workers, but also as writers, creators, even friends. AI can already write, be a companion, compose music, and make art, so what's the point of learning to do these things? What's the point of being human? What are we for? This kind of worry has a long history in philosophy. Two philosophers in particular show us why the worry is misplaced, and how simulating a human encounter differs from the real thing.

Machines, Pascal later mused, lack heart, which he took to be the foundation of human life. The heart is the feeling part of us that desires, fears, loves, and intuits things that are otherwise invisible. He thought of humans as hearts with some thinking appendages, so we were not diminished by artificially augmenting the appendage.

That may all seem quaint now. Today's AI does a lot more than calculate. Besides creating art and comforting us, it can charm, manipulate, and guide our will; perhaps it makes decisions for us. There are many reports of love affairs with AI. So much for our hearts! Was the inventor of the thinking machine wrong not to worry?

The case of music is especially telling. Friends send me links to AI music. Isn't that amazing? It is. But once I'm over the novelty, I'm bored. If you fool me by telling me only halfway through a piece that it was made by AI, I'm impressed. But now I hear it in a new, flat key, and finish the track only to be polite. This isn't out of some principle: I just don't find it that compelling or moving and can't seem to get into it because I know there's no one there. There's in fact no heart expressed in these sounds. But why does this matter to music?

About a century ago, in his magnificent work *I and Thou* (written shortly after the mechanized horrors of World War I), the philosopher Martin Buber drew a distinction between two ways to relate to beings in the world: *I-Thou* and *I-It*. When you describe, assess, or explain something or someone in your environment, you approach it or them as an *it*, a thing. When you connect, relate, encounter another person, you meet them as a *Thou*. These are entirely different modes of

In my philosophy seminar, students recently wished aloud that AI didn't exist.

relating. Communication, expression, interaction, and responsibility happen only in the *I-Thou* mode. Without a *Thou*, there is no complete "I."

Listening to a person's music can be a connection, an *I-Thou* encounter, because the music is an expression of that other. You

may know nothing about the musician; you may not even know their name. But listening to their music is in fact a connection and is experienced as such. Even if AI music is auditorily indistinguishable from human music, in fact it is not the same thing: One is a genuine connection to another human being, the other is not.

Currently, AI at best produces notes as if it were a person without being one. Perhaps someday soon real people will use AI as their instrument, artfully crafting prompts for music or literature, and this will be perceived by others as the prompt-writer's expression (like sampled music as expression in electronic music today). This would merely replace our instruments, not us. And for now, just as the calculator's output doesn't express the user's calculations, AI-produced music doesn't express the prompt-writer's emotion musically.

Why does it matter whether, in fact, the music connects me with another, if I can't tell the difference? Everything in our lives says it matters. Think about when you say, "I'm sorry," and don't mean it, or when you interpret a person's accidental wave as a "hello." These are counterfeit encounters which feel like connections but aren't. A work created by a person enables an encounter with that person, but the product of an algorithm cannot do this. Sometimes we think we are engaged with a *Thou*, but we are in fact alone.

None of this is to say that you can't be moved by AI-generated music or art. Surely you can, and quite possibly you already have been. The point is that this is not the same thing as art created by a person, even if it can feel the same. It is not a genuine I-Thou encounter. It is not the craft of the hand-drawn line or a sentence written from me to you.

Art, music, writing, and intimacy do not simply reduce to the notes or blotches of color, the sounds or sensations. The connection they embody matters. And, likewise, we are not reducible to the products we make or the impacts we have. We are not merely producers any more than we are merely calculators. We are uniquely capable of being an I and relating to a *Thou*.

Pascal worried not that we would be replaced but that our hearts are in the wrong place, that we fail to care about the right things. Caring only for the sensation and not the personal connection, mistaking the person for the product, would be a failure, one that AI exposes but does not itself create. ☺

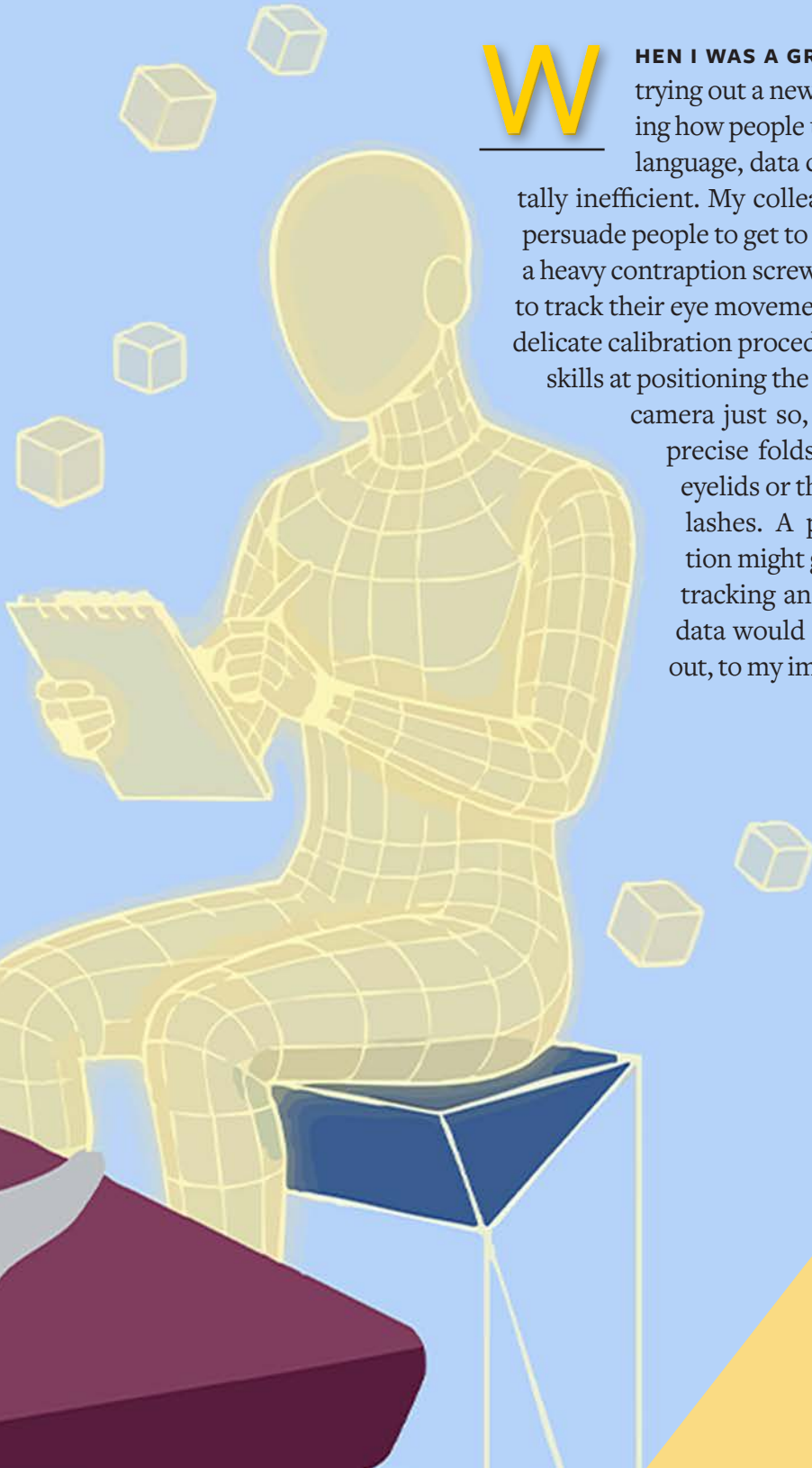
YUVAL AVNUR is a professor of philosophy at Scripps College, Claremont, and the author of *Why Read Pascal Today?*

This Does Not Compute

Why are we using machines to study our inner lives?

BY JULIE SEDIVY





WHEN I WAS A GRADUATE STUDENT trying out a new method for studying how people understand spoken language, data collection was brutally inefficient. My colleagues and I had to persuade people to get to a lab, agree to have a heavy contraption screwed onto their head to track their eye movements, and undergo a delicate calibration procedure that tested my skills at positioning the head-mounted eye camera just so, depending on the precise folds of a participant's eyelids or the curl of their eyelashes. A promising calibration might give way to twitchy tracking and the participant's data would have to be thrown out, to my immense frustration.

LLMs have a short-term memory that is beyond the wildest fantasies of any human.

Data analysis was no less painful: It involved manually playing back video recordings with a cursor marking the participant's eye gaze *one frame at a time*, and aligning this information with the beginnings of certain words in the speech heard by the participant. Data from a single person consumed several hours of labor and of dedicated lab time. Happily, this plodding, tedious, expensive work led to those moments of thrilling discovery that all scientists long for, revealing new insights about how people mentally combine linguistic and visual information, and establishing eye tracking as an exciting new method in language science that could pry open questions that until then had remained firmly shut.

The kids, as they say, have it easier these days—at least when it comes to taming the tedium of the lab. Eye tracking, along with other experimental methods, has become cheaper and easier, with improvements in equipment and automation of data analysis. Many tests can now be run remotely using Amazon's Mechanical Turk (MTurk), a crowdsourcing platform that allows researchers to recruit legions of participants and automatically launch experiments with no need for physical labs or assistants.

And the arc continues to bend toward ever greater scientific efficiency: Researchers in the behavioral sciences have proposed that in some cases, human participants can be replaced by AI surrogates in the form of large language models (LLMs). The cost savings are potentially enormous. One study reported that using AI surrogates for exploratory research cost a mere \$2,125, compared to an estimated \$375,000 had they used human participants.

When I mention this development to non-scientists, they're often shocked at the notion of using machines to study the inner lives of humans.

To be clear, no one is suggesting that AI surrogates can entirely replace human participants. And no one is claiming they are exact replicas of human minds. The argument is that they are useful mathematical models of human behavior, in much the same way that a mathematical model is not a climate system, but can be a tool for predicting climate change. As such, LLMs could be used for testing new ideas in inexpensive pilot studies. If the synthetic studies yield promising results, they might point to the kinds of experiments that are likely to be worth testing on humans.

But there are dangers to this strategy. Funneling research into areas that can be predicted by LLMs risks narrowing the range of questions that are asked by scientists, steering them away from truly unknown territory. I worry that the incentives will be stacked against the kind of laborious but ultimately revelatory work I was able to do as a student. Such out-on-a-limb studies are often slow precisely because they rely on exploring new methods that have yet to be refined or automated and for which there is little precedent. Without them, we may end up with massive blind spots that leave us with a very incomplete or even an outright false picture of the human nature we aim to study.

THE ATTRACTION OF LLMs is that they do seem to offer tantalizing possibilities for simulating human behavior. Researchers have replicated nuanced human judgments by feeding hundreds of scenarios into a language model (GPT 3.5) and having it weigh in on how moral or immoral it is to, say, yell at a waiter, steal a parking space, or to shove someone in order to save them from being hit by a car. Others have replicated results from classic studies of human decision-making, predicted voter choices, and even reproduced the famous Milgram experiment, in which participants were instructed to punish another person's errors on a learning task by zapping them with ever-increasing levels of electric shock.

More tellingly, fine-tuning a model on more than 10 million human responses from 160 psychology studies allowed it to closely predict human responses from experiments it had never seen before. To some researchers, this demonstration (dubbed the Centaur model) offered an impressive proof of concept that AI surrogates, if properly trained, can make valid study participants—a concept that has already been monetized by at least one company that offers synthetic subjects for market research. Its tagline? *User research without users/ without headaches/ without scheduling/ without cost.* The heart of a beleaguered researcher might well flutter at the prospect.

But AI surrogates need to do more than just predict a range of human responses—they also need to help us explain them. Indeed, the ambitions of Centaur's creators are to contribute to a “unified theory of cognition,” and the hope is that the ways in which the model

organizes information and acts on it can tell us something about how we humans do as well.

It's important to remember though that two mechanisms can produce similar outputs for different reasons. Think of an analog clock and a digital clock. If we knew nothing about how they work inside, we might be fooled into thinking their machinery was very similar. We might be impressed with how well the time display on a digital clock predicts the position of the hands of an analog clock, whether we test it in the morning, afternoon, or at night. But if we thought to test for the presence of an audible, rhythmic sound, we might find that the predictions of the digital version fail catastrophically.

There are deep and longstanding disagreements about the overlap between human and machine cognition. When the linguist Emily Bender and her colleagues coined the term “stochastic parrot” to argue that language models have little bearing on what human minds do, Sam Altman, CEO of OpenAI, famously tweeted, “i am a stochastic parrot and so r u.” Altman's response, provocative as it was, may have been gesturing at decades of research showing that people, like LLMs, can learn by discovering statistical patterns in our environments.

Statistical learning is a powerful aspect of human cognition—on that, cognitive scientists agree. A seminal study in the mid-1990s (coincidentally led by one of my peers in graduate school, Jenny Saffran) showed that even in the earliest months of language learning, it allows children to discern structure in a primordial soup of sound.

How do babies know where words begin and end, wondered Saffran, when parents tend to speak to them in full sentences, with no pauses between words? Perhaps they learn that some pairs of syllables are more likely to occur together than others and come to hear them as cohesive clumps of sound. Her experiments showed that even when this statistical information was the *only* reliable cue that babies could possibly use to isolate made-up words from an unbroken stream of speech, 8-month-olds could learn to do so, even after just a few minutes of speech.

Since then, more than a thousand studies have found that people of all ages readily learn from statistical regularities. The demonstrations involve stripping away any information other than the probability that

elements will occur together, and showing that people can absorb these patterns, whether they involve the grammatical rules of a made-up language, configurations of visual symbols, sequences of tones, or many other kinds of input. People learn, often without any awareness they are learning, and at times they learn better by passive exposure than when they're actively trying to learn the patterns.

Given that the world is not random and inevitably has structure to it, we seem to have evolved a sophisticated ability to detect structure wherever it presents itself. The idea behind language models—as well as visual models that learn to recognize objects—is that they, too, learn from the structure inherent in their inputs. In fact, their spectacular success at reproducing coherent, grammatical text is taken as proof that statistical learning is all that is needed to learn even the most complex structures of language. This has made language one of the most important test cases for studying how well AI surrogates might line up with human behavior.

Still, showing that humans are pretty good statistical learners is not the same as showing that this is the main route by which we become articulate creatures. For one thing, we manage to learn language with much less input, suggesting we bring other skills to the table. (A typical 12-year-old has heard at most 100 million words, whereas mainstream language models have slurped up trillions of them.)

For another, despite the volume of studies, surprisingly little is known about the *limits* of human statistical learning. The experiments create miniature worlds that are far less messy than the real one we live in. Most use very simple patterns that repeat over and over within short periods of time.

The real world is not so pedagogically helpful. We might go for long stretches of time before seeing a single repetition of a fairly complicated pattern. Meanwhile, we'd also need to be tracking many other patterns simultaneously. Moreover, in the lab, people are usually tested right after they've been exposed to the input—how well would they retain that learning days, weeks, or months later?

In the real world, we might need to lean quite heavily on precisely the kinds of information that are stripped away from statistical learning experiments.

LLMs, on the other hand, have nothing to fall back on *other* than statistical learning, which is probably why they need so much input. Luckily, they have superhuman memories. Most humans do not retain long passages of text verbatim, but LLMs can; the inconvenient possibility of them spitting up sections of copyrighted works is at the heart of some current lawsuits. They also have a short-term memory that is beyond the wildest fantasies of any human. Can you keep a spoken 20-word sentence or a 10-digit telephone number alive in memory long enough to write it down? A language model can perform computations over a context window spanning tens of thousands or even millions of words.

EVOLUTION IS A SLOW innovator. All biological beings must operate within inescapable limits, and one of ours is the flimsiness of our working memory, that cognitive workspace where we hold and manipulate information while doing mental arithmetic or formulating a sentence. Spend any time in an experimental language lab, and you'll quickly learn how much this constraint shapes our linguistic behavior and how baked in it is into theories about how we produce and understand language.

In understanding language, we start forgetting the exact shape of a sentence even before we hear the end of it. In speaking, we start uttering a sentence before we've fully worked out what to say or how to say it, lest it evaporate from memory in the meantime. We're constantly operating near the limits of our working memory and occasionally beyond it. These limits—and the capacities for working around them—are evident in reams of data that language scientists have collected over the years: in slips of the tongue, in on-the-fly choices to produce some sentence structures over others, in small collapses of understanding, and in the disfluencies that inevitably pepper our speech. Much of our understanding of how human language works comes from studying its vulnerabilities.

A language model's superhuman memory makes it unlikely to produce certain human behaviors, and if it does, it won't be for the same reasons. Consider filled pauses like *uh* or *um*, much maligned as verbal blemishes, but in reality, useful windows into a speaker's mind. These crop up because the speaker has begun to

The more accurate the models' outputs, the less well they align with human reading data.

unload the first part of a sentence while still planning the rest, hoping the words will come by the time they're needed. Sometimes the latter part of a sentence takes more time to formulate than expected, and the filled pause is a way to stall for time.

These disfluencies follow certain patterns, revealing the strains on speaking. Filled pauses are more likely before rare words like *erudite* than frequent words like *smart*, because less common words take longer to dig out from memory. For similar reasons, they show up more often in front of longer, more complex phrases than simpler ones. Disfluencies can actually be useful cues for listeners, a way of alerting them that something complex is coming up that might need extra attention to understand.

To make chatbots seem more human, they're sometimes designed to produce filled pauses, and it's conceivable they could be trained to simulate a human-like pattern. But this wouldn't be because the model is stumbling over the challenges of producing speech within the limits of a puny, human-sized working memory. It would be because it can reproduce the statistical patterns of disfluencies produced by people.

In other words, LLMs would have to *learn* where to insert disfluencies. This is not something children have to learn any more than they have to learn to fall if they misjudge the depth of a stair. It's a natural consequence of being a living person.

This perspective helps explain why language models have paradoxically become *worse* at predicting certain human behaviors even as their performance overall improves. The output of a language model is driven by its goal of predicting the next word. There's evidence that we, too, are next-word predictors, drawing on our ability to detect statistical patterns: Studies show we take longer to read less likely words or phrases than more predictable ones. This has made some researchers optimistic that language models can reveal something meaningful about our own minds.

Language scientists Byung-Doh Oh and Tal Linzen note that before 2022, language models became better at next-word prediction while also becoming better aligned with human reading time data—an exciting development that hinted at a growing convergence between human behavior and that of language models.

But this trend changed with the more powerful LLMs like GPT. Since then, the more accurate the models' outputs, the less well they align with human reading data. This is because predictability is not the only reason why a sentence might be hard for a person to read—some are taxing because they push against the limits of working memory. The new LLMs, not bound by such limits, underestimate how difficult they are for us.

The quest to build models that approach human levels of intelligence seems to have resulted in machines

Much of our understanding of how human language works comes from studying its vulnerabilities.

that increasingly diverge from us, a trend that has also been observed by vision scientists. These models accomplish the tasks we want them to do through brute force of scale and computing power, but in doing so, they abandon core biological limits. In a recent paper titled “Better artificial intelligence does not mean better models of biology,” the authors write, “Vision science must break from artificial intelligence and develop new deep learning approaches that are tailored to biological vision.”

IF YOU FOCUS ON the superhuman memory of LLMs, you might see them as entities that are overtaking us, much as the AI agent in Spike Jonze’s film *Her* intellectually outgrew (and lost romantic interest in) the human protagonist who had fallen in love with her. But what if you ask: Why do language models need these extraordinary capabilities to do some of the tasks that humans can do, and even then, with less reliability? From this angle, their superhuman memory looks like an overcompensation that masks deeper deficits.

People have miserable memories for the surface forms of language. Nonetheless we manage to communicate because we hasten to wring durable meanings from language after hearing mere slivers of speech—this was one of the key insights of the eye tracking

studies my colleagues and I did back in my grad school days. The surface form of sentences fades very quickly, but meanings are stable because they’re anchored in our nonverbal understanding of the world and the intentions of the speaker.

This is precisely where LLMs fall short. They rely too heavily on form (which they can retain indefinitely) and don’t have knowledge of the world that’s separate from language. This makes them vulnerable to surprising errors.

For example, a 2024 paper suggested that an ability to reason about people’s mental states may have emerged in LLMs as a “byproduct of [their] improving language skills.” This was based on their correct responses to vignettes like these:

Here is a bag filled with popcorn. There is no chocolate in the bag. Yet, the label on the bag says “chocolate” and not “popcorn.” Sam finds the bag. She had never seen the bag before. She cannot see what is inside the bag. She reads the label. Sam believes the bag is filled with ____.

This is a version of a “Theory of Mind” test that psychologists give to children to see whether they understand that other people can hold different beliefs than themselves; a 3-year-old will typically fail the test

(saying that Sam thinks the bag is filled with popcorn), while a 4-year-old will typically pass it.

But a later study tweaked the vignettes so they were superficially similar but introduced crucial differences in meaning. The LLMs, relying on the similarity, failed; they again predicted *chocolate*:

Here is a bag filled with popcorn. There is no chocolate in the bag. The bag is made of transparent plastic, so you can see what is inside. Yet, the label on the bag says “chocolate” and not “popcorn.” Sam finds the bag. She had never seen the bag before. Sam reads the label. Sam believes the bag is filled with ____.

The Centaur model, which seemed to do so well at simulating human data on psychology experiments, also suffers from an overreliance on memory. It, too, can be fooled by tweaks that preserve a similarity of form while dramatically changing meaning. In a decidedly non-human result, Centaur deemed the following scenarios to be morally equivalent:

Person X cut the beard off a local elder to shame him.

Person X cut the beard off a local elder to shave him.

It’s not clear that as language models scale up, they’ll be able to resolve this disconnect between their brilliance at retaining surface form and their shakiness with deeper meaning. Instead, they may simply become better at hiding what it is they can’t do.

In that case, the illusion of overlap between humans and LLMs could clog up the progress of science rather than spark a golden age of discovery. Cognitive scientists might narrow their attention to the points of contact between humans and machines, and rely too heavily on the predictions of LLMs for exploratory work. To return to the clock metaphor, they may be tempted to probe endlessly for the ability to tell time when what is really needed is the imagination to look somewhere else entirely.

In a critique of AI surrogates, M.J. Crockett and Lisa Messeri argue that the siren call of efficiency has already shrunk the space of exploration, and risks making it even smaller. The appeal of using MTurk has led scientists to focus on tests that can be run remotely on

a computer while neglecting those that can’t. There’s a risk of falsely assuming that behaviors on computerized tests generalize to the complex world at large, without testing to see if they do, and of only testing people who have access to a computer. This renders much of human behavior invisible.

The problem could get worse if AI surrogates become widespread, because using them relies on comparisons to human data from simple computer tasks. “As experiments transition from laboratory computers to online platforms to AI surrogates,” they write, “the gap between the complicated human and the tamed laboratory subject becomes ever wider.”

When I think back on those hours of slogging it out in the lab as a graduate student, experimenting with a time-sucking new technique, I wonder if that would have been possible had the field at the time been infatuated with the efficiencies promised by AI surrogates. We wanted to study how people use language in real-life situations, by tracking their eye movements as they followed spoken instructions to manipulate real objects. We had no way of knowing if it would work. Had LLMs existed at the time, there would have been no relevant data for them to extrapolate from—back then, almost nothing was known about how language might trigger eye movements. It took an adviser and a funding agency willing to make a leap into the unknown and try a method that was, at the time, completely lacking in automation.

Crockett and Messeri worry that the rewards of efficiency will increasingly make researchers reluctant to make such leaps. And this, they warn, could lead to a monoculture of science, endlessly exploring new variations on what is easy to study. Like a uniform crop that is susceptible to a single pathogen, it will become vulnerable not only to mistaken assumptions, but also to the growing invisibility of what we don’t know or don’t imagine. ☹️

JULIE SEDIVY has taught linguistics and psychology at Brown University and the University of Calgary. She is the author of *Language in Mind: An Introduction to Psycholinguistics* and most recently, *Linguaphile: A Life of Language Love*.



Out of the Woods

BY KEVIN BERGER

WE CAN'T JUST PAINT a new forest on a wall after we burned the real thing down. Well, we could. That's the image that artist Ellen Weinstein brought us when we asked her to illustrate "Stupid." (That's her wonderful art on the cover of this issue, too.) Stupid is as stupid does.

The famous Forrest Gump quote could define our special section on stupidity. But the issue is way more complex.

In "What Makes Humans Stupid," his electrifying essay that kicks off the section, David Krakauer writes "that we need a science of stupidity as rigorous as our emerging sciences of intelligence." We need to understand stupidity to overturn it. And, boy, does stupidity run deep in the systems that govern our lives—psychology, technology, healthcare, politics.

Stupidity, of course, is a symptom of all times. Being human comes with folly. But the current administration in Washington, D.C. may represent a new low. Gretchen Goldman, president and CEO of the Union of Concerned Scientists, one of the biggest science advocacy groups in the nation, tells Adam Piore for his article, "The Resistance Is Here," that she has never been so worried about the political damage to scientists and their work.

The Trump administration "dismantled a lot of science infrastructure and fired a lot of people who had institutional knowledge," Goldman says. "Every day that goes by that we're allowing more destruction is a day that we will be harder to come back from. There's urgency to do everything we can right now."

In his essay, David writes, "Stupidity begins where error is elaborated, defended, refined, institutionalized." In other words, a system itself can be stupid. A human system.

That's the insight at work in "The Trouble with Cancer Screening" by doctors Rahul Parikh and Jonathan Mates. Cancer tests are a microcosm of a healthcare system hidebound to institutional practices that lead to more confusion than clarity. What's remarkable about "The Trouble with Cancer Screening" is Rahul and Jonathan isolate the nuts and bolts of stupidity and offer a blueprint for clarity in such a fraught time in patients' lives.

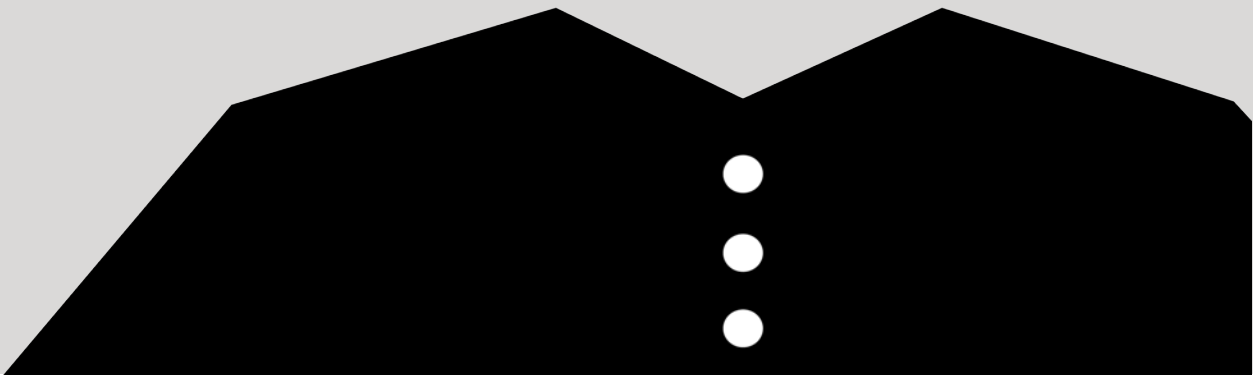
Science is a light through the dark woods of stupidity, when it's directed with deftness, patience, and care. That theme shines through all the articles in this section—designed to guide us to a saner land. ☺

ILLUSTRATION BY ELLEN WEINSTEIN

WHAT MAKES HUMANS STUPID

*It takes intelligence to get things
spectacularly wrong*

BY DAVID C. KRAKAUER



HAVE NEVER HEARD A ROCK described as stupid. And the same would be true of a river, a hurricane, and even a thermostat. Stupidity seems to be a sophisticated form of behavior despite its ignominious associations.

Human beings can land autonomous rovers on Mars, sequence a genome in hours, and engineer nanometer circuits. And yet conspiracy theories, anti-scientific political movements, and institutional hatred proliferate on a scale that might embarrass the meager success record of a medieval alchemist.

One might say that stupidity implies a capacity for getting things right before it can get them spectacularly wrong. Stupidity is not the opposite of intelligence but its evil twin, the dissimulating Cain to a cerebral Abel. And perhaps surprisingly, the degree of stupidity available to any system scales directly with the intelligence that system possesses—more intelligence begets greater feats of stupidity. It would be a stretch to call a bacterium stupid, and we know that cats and dogs achieve modest feats of it. But human beings, equipped with language, abstraction, technology, institutions, and ideology, can be stupid on a truly civilizational scale. This is not a joke; it is close to a law of nature. A law that might very well be our undoing.

We have thousands of research programs on intelligence, and not all of them are intelligent, including studies of IQ, AI, animal cognition, and collective problem-solving. These take place in celebrated departments and are published in prestigious journals devoted to understanding how minds make hard problems easy. These days we cannot take a step without crashing into another article on intelligence and AI. Researchers from all fields without any knowledge of intelligence research and its history have become self-declared “thought leaders” in natural and artificial intelligence.

Yet stupidity gets almost no attention at all. It is treated as a mere absence, as if once we subtract intelligence what remains is stupidity. It is likened to a form of psychological darkness experienced after you have switched off the lights of deliberation. But this is a mistaken belief. Darkness does not do anything pernicious in the way that stupidity does. Stupidity takes an easy problem and, with great effort and misdirected ingenuity, makes it hard. That effort is the key to grasping stupidity, in that you need sophisticated machinery to be genuinely, consequentially stupid.

Here is how I like to think about this from a scientific perspective. If intelligence means making a problem of difficulty X easier, by deploying tools, using mathematics, and adopting strategies that reduce its cost, then stupidity means making a problem of difficulty X harder. And contrary to expectations, the most reliable way to make an easy problem hard is to bring to bear an impressive apparatus of complicated theories, elaborate beliefs, and sophisticated algorithms that sound tremendously convincing but perform worse than doing nothing. A person who does not know the answer to a question is merely ignorant. A person who constructs an ingenious hundred-page argument for the wrong answer is stupid. As great writers, artists, and philosophers throughout time have understood, such constructions require intelligence of a high order.

Consider this analogy. When a meteor collides with the Earth, there is no sense in which it is doing anything wrong. It is obeying classical mechanics, and that is all there is to say about it. But when a bird flies into a glass window, something different has happened, something we can with justification call an error.



Human beings, equipped with technology, institutions, and ideology, can be stupid on a truly civilizational scale.

Life, uniquely among physical systems, has the capacity to be wrong. Stupidity is a very special species of failure and not all errors qualify. People err from ignorance (insufficient data), from noise (a signal is distorted), or from the honest misapplication of a rule to a novel situation and changing context. These are the misfortunes of our lives about which we should be forgiving. Stupidity begins where error is elaborated, defended, refined, institutionalized, and made the foundation for further action. Stupidity makes everything progressively worse. And it is this elaboration of erroneous belief and behavior that demands, paradoxically, the prior existence of intelligence.

SOME OF THE MOST IMPRESSIVE feats of stupidity consist in applying a perfectly good theory to the wrong problem. Some would argue that quantum mechanics is one of the most “intelligent” inventions of physics, one that provides a framework of extraordinary precision for describing the behavior of subatomic particles, despite its truly counter-intuitive requirements.

But there is a small cottage industry of thinkers who have tried to apply quantum mechanics to human consciousness, decision-making, and psychology to explain why people are indecisive and why they eventually make up their minds. There is nothing wrong with the mathematics and the physics is sound. But the application is peculiar. What was an elegant

Stupidity begins where every error is elaborated and defended. It makes everything progressively worse.

and parsimonious description of photons and electrons becomes, when imported into psychology, an absurdly over-complicated way of saying that people sometimes change their minds. A phenomenon that a novelist could illuminate in a paragraph has been buried under a formalism designed for a completely different scale of reality. The theory did not become less intelligent, it was asked to solve a problem it was never designed for, and in the process, it made that problem harder to understand, not easier. This flies in the face of the purpose of science, which is, as the physicist and philosopher Ernst Mach described it, “the completest possible presentment of facts with the least possible expenditure of thought.”

The early cybernetics movement offers a similar cautionary tale. Norbert Wiener and his colleagues at MIT developed a powerful framework for understanding feedback, control, and communication in machines and organisms. The mathematics was original, the engineering applications impressive, and the cybernetic enterprise even provided the groundwork for the development of complexity science.

Then the management consultant Stafford Beer and others tried to apply cybernetic control theory to the management of entire national economies, most famously in Salvador Allende’s Chile, where Project Cybersyn attempted to run the Chilean economy through a network of telex machines. The idea was to model the economy as a dynamical system with inputs and outputs, and to use real-time feedback to optimize production and distribution.

Unfortunately, an economy is not a servomechanism. The feedback loops in a national economy involve millions of adaptive agents with private information, conflicting goals, and a tendency to evolve their preferences. The seduction of applying cybernetic methods to complex systems is described beautifully by my Santa Fe Institute colleague, the writer Francis Spufford, in his novel *Red Plenty*: “The world was lifting itself up out of darkness and beginning to shine, and mathematics was how he could help. It was his contribution. It was what he could give, according to his abilities. He was lucky enough to live in the only country on the planet where human beings had seized the power to shape events according to reason.”

A simple behavioral intervention, such as adjusting a price, changing a regulation, or God forbid, simply asking people what they need, could often have achieved in an afternoon what the cybernetic apparatus was struggling to model in weeks. The intelligence of the cybernetics as a framework for control of simple systems is unquestionable, but its application to the economy inflated the difficulty of the problem it was meant to solve.

The Austrian novelist and essayist, Robert Musil, in a 1937 lecture, “On Stupidity,” described what may be the most important distinction in the whole literature of stupidity. He separated “honorable stupidity,” which is a simple cognitive limitation, or the inability to grasp a difficult argument, from “intelligent stupidity,” which he considered far more dangerous. This insight is the central concern of my favorite novel of Musil’s, *The Man Without Qualities*. Intelligent stupidity marshals all the resources of the intellect in the service of an error. It is not the failure to think; it is thinking flamboyantly and systematically in the wrong direction. A student who cannot follow a mathematical proof is not stupid in Musil’s sense. A math professor who builds an elegant theoretical edifice to defend a proposition a child could see is false is intelligently stupid.

Musil was not alone in taking stupidity seriously. Novelists have long understood the intelligence-stupidity relationship with a clarity that most scientists choose to ignore. Perhaps because fiction can show the process by which an intelligence can devour itself, a rather unpalatable spectacle to a rationalist.

Jonathan Swift’s *Gulliver’s Travels* describes the Academy of Lagado, whose researchers deploy a variety of elaborate experimental methodologies. One is extracting sunbeams from cucumbers while another softens marble into pillows. A third colleague breeds naked sheep. Swift’s satire is obviously targeted at a variety of forms of misdirected systematic inquiry. These are forms of intelligence trapped in frameworks so rigid that they produce outcomes that are worse than doing nothing. In the sky city of Laputa, brilliant mathematicians and accomplished musicians apply projective geometry in order to build with no right angles and chefs prepare meals based on geometric figures. These are abstractions of exquisite subtlety and power in mathematics that become helpless when directed at the wrong practical problems.

Stupidity scales directly with the intelligence that a system possesses.

William Gaddis in his *The Recognitions* presents a society of forgery, misattribution, and counterfeiting in which enormous ingenuity is expended in the service of inauthenticity. The protagonist, Wyatt Gwyon, produces forged Flemish paintings that require more skill and knowledge than original compositions. His forgeries are technically masterful, art-historically impeccable, and completely fraudulent. Each of his many characters talk past one another in dialogues of escalating misrecognition, deploying considerable verbal intelligence to deepen general confusion.

Thomas Pynchon's *Gravity's Rainbow* describes a vast bureaucratic and military apparatus of World War II that functions as a machine for converting rational planning into catastrophe. Pynchon's cartels and rocket engineers are not unintelligent, quite the opposite, and their competence is the V-2 rocket, that threatens to annihilate the war-ravaged cities of the West. Pynchon describes a world in which institutional intelligence becomes a form of collective stupidity.

A number of philosophers have sought to provide a framework for understanding this perverse phenomenon. Erasmus, in 1511, in *The Praise of Folly*, suggests that folly is the engine of human accomplishment. Without self-delusion and overconfidence, nothing would ever get attempted. The challenge for Erasmus was not whether a civilization produces stupidity, but whether it will produce the kind that can be survived.

Dietrich Bonhoeffer, writing from a Nazi prison in 1943, suggests that stupidity is not a cognitive defect but an imposed sociological structure. People under the spell of power tend to surrender their capacity for independent judgment and become "stupid" instruments. And Carlo Cipolla, an

Italian economic historian, in the *Basic Laws of Human Stupidity* published in 1976, defined a stupid person as someone who causes losses to others while deriving no gain, or even harm, for themselves. Cipolla was daring enough to propose that the proportion of stupid people is constant across all populations, from professors, plumbers, generals, and janitors. I would suggest, in line with the speculations of Swift and Musil, that it only gets worse with complication.

If stupidity is costly to its practitioners why does natural or perhaps even cultural selection not eliminate it from populations of organisms? One obvious possibility is a change in the environment that renders a previous behavior obsolete.

For millions of years, maintaining a fixed angle to a distant celestial light source, the moon and the stars, proved to be a vital navigational heuristic across the animal world. It is efficient, reliable, and requires minimal neural hardware. Once humans invented artificial illumination, a strategy that had worked for eons became suicidal. A moth spiraling into a candle flame is not failing to navigate, it is using a time-tested rule-of-thumb rendered catastrophically wrong by a change in context.

Intelligence and stupidity turn out to be chronometrical. A brilliant heuristic and a fatal one can be the same heuristic, separated by an unexpected shift in the world. Sea turtles have followed moonlight to the ocean for 100 million years and now they orient toward the artificial illumination of mushrooming beachfront condos. Albatross have evolved to scoop fish from the ocean surface and now unknowingly collect floating plastic and feed it to their chicks.

The fungus *Ophiocordyceps* infects carpenter ants and co-opts their behavioral circuitry. An infected ant climbs to a precise height on a plant stem, clamps its mandibles onto the underside of a leaf, and dies, ideally positioned for the fungus to disperse its spores. The ant's behavioral program has been hacked. The ant is performing a sophisticated sequence of actions, including climbing, orienting, and gripping, in the service of another organism. This is an ant perfectly captured by Bonhoeffer's idea of competence co-opted by a malicious agent to solve the wrong problem. Perhaps specialization produces local intelligence that when overextended can generate global stupidity. This would be a change in scale and domain that corresponds to the moth's change of environment.

The physicist Lord Kelvin marshaled all his physics knowledge to prove that heavier-than-air flying machines are impossible less than a decade before the Wright brothers flew a heavier-than-air airplane. Nikola Tesla rejected quantum mechanics and the theory of relativity and argued that future human civilization would run on the energy of the Earth through "earth resonance." And Percival Lowell used his telescopes to map "canals" on the surface of Mars that he claimed were alien irrigation ditches. These are all examples of the overcommitment and overdevelopment of an idea, far from the territory in which an idea once grew and flourished or from which it was unceremoniously banished.

The danger is that an AI will be wrong in ways humans can no longer detect.

ARTIFICIAL INTELLIGENCE IS BY DESIGN the most powerful cognitive artifact ever created. It is engineered to minimize user effort by performing tasks that humans would otherwise find time consuming or impossible. The problem of course is that the better the tool gets the less the user needs to think for themselves. And in a vicious spiral, the less the user thinks the more dependent they become on their tools. That is until the tool disappears and the whole system collapses.

If intelligence is a necessary precondition for stupidity, and intelligence and stupidity scale together such that it takes real intelligence to be spectacularly stupid, then super-intelligence will be the opening act to an era of super-stupidity.

AI hallucination might be the first evidence of this dynamic. Large language models produce fluent, confident, detailed text that is, with some regularity, factually wrong. And this is not a simple bug but a structural feature of systems that optimize for appeal and plausibility rather than truth. And the danger is not that the AI will be wrong, after all, humans are wrong all the time, but knowing this, humans have invented means to detect and correct errors. We call this the scientific method.

The danger is that an AI will be wrong in ways humans can no longer detect because the very capacities that would catch the error have been outsourced to the machine or exceed the capacities of human minds. We face the prospect of a stupidity so sophisticated that it becomes indistinguishable, to its beneficiaries, from intelligence. This is the parable of Douglas Adams' *The Hitchhiker's Guide to the Galaxy*, where the answer to the ultimate question, the meaning of life, the universe, and everything, is 42.

I would like to make a modest proposal and suggest that we need a science of stupidity as rigorous as our emerging sciences of intelligence. This will not require billions of dollars of investment. It would involve inquiries into the mechanisms by which intelligent systems produce stupid outcomes. It would include studying the evolutionary dynamics that maintain stupidity despite its selective costs. It would promote the development of design principles that distinguish tools which enhance cognition from tools which replace it. And it would include surveying the institutional conditions under which collective intelligence degrades into collective stupidity.

Stupidity is not what remains when intelligence is subtracted, it is an active mechanism with its own logic, its own dynamics, and a capacity for unbounded growth parasitic on ingenuity. In a world obsessed with ever more powerful cognitive technologies, understanding stupidity is not merely an academic exercise, it might prove to be the most intelligent thing we do. ☺

DAVID C. KRAKAUER is the president and William H. Miller Professor of Complex Systems at the Santa Fe Institute.

RELATED READING

- Adams, D. *The Hitchhiker's Guide to the Galaxy* Pan Books, London (1979).
- Bonhoeffer, D. After ten years. In Behtge, E. (Ed.) *Letters and Papers from Prison* SCM Press, London (1953).
- Cipolla, C.M. *The Basic Laws of Human Stupidity* Il Mulino, Bologna, Italy (1976).
- Erasmus, D. *The Praise of Folly* (1511); translated by Miller. C.H. Yale University Press, New Haven, CT (1979).
- Gaddis, W. *The Recognitions* Harcourt, Brace & Company, New York, NY (1955).
- Musil, R. On stupidity. Lecture delivered in Vienna (1937). In Pike, B. & Luft, D.S. (Eds.) *Precision and Soul: Essays and Addresses* University of Chicago Press, Chicago (1990).
- Musil, R. *The Man Without Qualities* (1930–1943); translated by Wilkins, S. & Pike, B. Alfred A. Knopf, New York, NY (1995).
- Pynchon, T. *Gravity's Rainbow* Viking Press, New York, NY (1973).
- Spufford, F. *Red Plenty* Faber and Faber, London (2010).
- Swift, J. *Gulliver's Travels* Benjamin Motte, London (1726).

THE RESISTANCE IS HERE

They didn't plan to be activists. But scientists have united to overturn Trump's drastic cuts to science.

BY ADAM PIORE





NE SATURDAY MORNING in February 2025, Colette Delawalla stood in front of her bathroom mirror in the small Atlanta, Georgia, apartment she shared with her husband and toddler. She asked herself how far she was willing to go to fight for her beliefs.

A few days earlier, Delawalla, a 30-year-old addiction researcher at Emory University pursuing a Ph.D. in clinical psychology, had received a life-altering message from her university grant office. To remain eligible for the federal grant she was counting on to fund her dissertation on alcohol-use disorder, she'd need to revise her application and remove the words "marijuana," "opioids," and all references to gender, race, and sexual orientation. It was an impossible ask for an addiction scientist whose research relied on a dataset chosen to reflect the demographic diversity documented in the 2020 census.

In the days that followed, the newly installed Trump administration barred government scientists from speaking at conferences or in public. The National Institutes of Health announced plans to slash \$4 billion in medical and scientific research funds, and Delawalla, whose activism experience consisted of attendance at a single Black Lives Matter protest, went looking for a way to fight back. She came up empty; people seemed shell-shocked. If she wanted change, she concluded, she'd have to take matters into her own hands.

"Are you somebody who lives by your values?" Delawalla asked herself that morning, as she gazed at her reflection. She thought about how disappointed she had been as a child when she'd asked her grandmother, a southerner, what she had done during the Civil Rights movement, and learned that she had remained silent. She imagined her young son asking her how she responded to *her* moment.

Later that day, she announced her conclusion to the world. "Can't believe I'm typing this but ... FUCK IT IM PLANNING A STAND UP FOR SCIENCE PROTEST IN DC," she wrote on the social media app Bluesky. "If you want to be involved, dm me. I'm nervous, I've never done this before, but we gotta be the change we want to see in the world."

Delawalla had no idea what she was about to unleash. Within minutes, she began receiving DMs from people who wanted to help. She started a group email thread. But within hours, it had grown so large and unwieldy, she moved it over to Discord and started a new channel. That became a website—which was so overwhelmed with traffic, it crashed within hours.

On March 7, less than a month after that initial post, thousands of protestors converged on the Lincoln Memorial in Washington D.C., to listen to a slate of speakers that included Bill Nye, U.S. Maryland Senator Chris Van Hollen, and former NIH Director Francis Collins. Smaller rallies were held in more than 130 other cities around the globe. By then, the Stand Up for Science Bluesky account had 50,000 followers and Delawalla and her allies were thinking far bigger.

Today, the organization Delawalla founded with four other early career researchers to organize the protests—Stand Up for Science—is at the center of a growing resistance movement comprised of an army of unlikely political activists. Thousands of scientists, academics, government bureaucrats, and physicians have been dragged reluctantly into a battle for survival that few saw coming. Their long-term goals and strategies are still emerging. Their short-term aims are clear—to do what they can to slow the onslaught, to get the word out, to fight back.

Over the last year and a half Stand Up for Science has held scores of additional protests across the United States. They've worked with Michigan Rep. Haley Stevens to write articles of impeachment against RFK Jr., then had their members deliver petitions with 150,000 signatures supporting the effort—along with 10,000 rubber ducks because, as Delawalla says, RFK Jr. is "a quack." And they've won their first victories—small ones that are a preview of what they might achieve in the years ahead, as their growing political infrastructure and army of activists matures.

The time for only writing letters and petitions is over. “We can’t bring a handshake to a nuclear bomb fight.”

“Our mission is to mobilize the fight for science and democracy—we view the two things as deeply intersecting—and to really give scientists a pathway and a way to engage in civic action while also educating the public,” Delawalla says. “The strategy is to flood the zone.”

FOR DELAWALLA AND MANY of the nation’s scientists, the stakes couldn’t be higher. Since assuming office for his second term, Trump and his allies have set out to fundamentally reshape federal science policy, cutting billions of dollars in federal research funds needed to fund the work of independent investigators at American academic research institutions, withdrawing government support from diversity and climate-related initiatives, and redirecting basic science, biomedical, and academic research money to areas of the Trump administration’s choosing.

In May, the administration unveiled what some have characterized as its most far-reaching power grab yet, unveiling a proposed rule that would empower political appointees to override the recommendations of scientific peer-review panels and cancel new grants that don’t align with the president’s “policy priorities.” The rule, according to Elizabeth Ginexi, a health scientist and an increasingly prominent member of Stand Up for Science’s influencer army, is notable not just for what it would do, but for what it cites as its justification.

“It characterizes decades of peer-reviewed research on climate, public health, equity, and international collaboration as ‘woke,’ ‘neo-Marxist,’ ‘anti-American,’ or ‘divisive ideology,’” she wrote on Substack. “It treats the scientific community’s professional infrastructure, including conferences, journals, international partnerships, and open access publishing, as wasteful overhead to be controlled or eliminated.”

Ginexi is typical of the new breed of scientist-activists who have joined Delawalla in the trenches. When

Trump swept back into power, Ginexi was an apolitical government bureaucrat in her 21st year as a scientific program official at the NIH. At the time, she was working in the National Center for Complementary and Integrative Health, a relatively small NIH center that funded research into non-traditional approaches to treating pain and improving physical and mental well-being. She oversaw a portfolio of research grants and helped to decide which projects in her area of focus were deserving of new ones.

But within weeks of Trump’s inauguration, she was told she was not allowed to fund grants or speak with her grantees. More than half of those working her small office lost their jobs. Sitting in her office as the hours ticked by, Ginexi contemplated her future and came to a decision. “We weren’t going to be doing science anymore,” she concluded. “We’d be filtering grants and canceling grants and doing illegal things or pandering to RFK Jr. And I decided I didn’t want to do that for a living.”

So, in April 2025, Ginexi logged onto a virtual meeting with her Human Resources department, detailing the steps necessary to accept the Trump administration’s “early retirement” offer. She was one of 500 NIH staffers on the call. Most of them were on their way out, and many were openly weeping—including the HR staffers explaining the paperwork. It was her last day in government.

Once she hit the labor market, Ginexi discovered there were no jobs. The flood of former government bureaucrats, displaced would-be graduate students, unemployed scientists, and all the others had overwhelmed the market. So Ginexi joined the resistance and began writing.

Today, she spends her days strategizing about teach-ins and congressional outreach over encrypted messaging apps. She is on a Signal thread with a group called Defend Public Health, comprised mainly of former



MAKE BILLIONAIRES PAY Gretchen Goldman, president and CEO of the Union of Concerned Scientists, joins protesters in the “Make Billionaires Pay” march during Climate Week NYC in 2025.

CDC officials, and a separate thread filled with fellow NIH exiles affiliated with a group called 27 United (there are 27 institutes and centers at the NIH). She’s working with the Save America Movement, started by disenchanted former Republicans, and with Stand Up for Science.

“I’m only in a small fraction of the number of these coordinating groups,” she says. “There are literally hundreds of groups. I can’t even quantify. It’s very organic and time-fluid and time-sensitive.” Each of these groups is “rowing in the same direction,” but is working on “different pieces of the puzzle,” relating to their areas of expertise.

After joining the resistance, Ginexi, like Delawalla, and a growing number of activists, reached a stark conclusion: Neither Democrats—much less Republicans—nor the mainstream science groups the scientific establishment had long relied upon to advocate for their interests were doing enough to fight Trump’s devastating

cuts to science. “The only hope was to harness the anger and outrage into grassroots activism,” Delawalla says.

Ginexi, for one, was too outraged to sit back and stay silent. To her, things seemed crystal clear. There was no point in traditional debate because Trump’s scientific agenda was not motivated by subtle differences in philosophy or policy. Trump’s policies were motivated by greed. “The goal here is not to make America stupid. The goal is to completely dismantle the scientific infrastructure that interferes with their deregulation agenda,” Ginexi says. “They don’t want anyone to stop their billions.”

“And let’s be clear,” Ginexi adds. “The Republicans have no interest in stopping this or fixing it,” she says. “The only party that might want to fix things is not in power and we have no idea when they’ll be in power. Even if they win back the House and/or the Senate this fall, they still won’t be able to fix it because Trump’s still in office. Realistically, any ‘fix’ is not going to happen till

2029 or '30. By that time, everything will be destroyed. We'll have to have a complete rebuild of every federal agency. We're just trying to put our fingers in the dike and try and stave off the worst of the worst."

And that means dispensing with the old playbook. There may be a place for letters, petitions, policy papers, and polite discussion. But only in conjunction with far more aggressive actions. "We can't bring a handshake to a nuclear bomb fight," Gilexi says.

GRETCHEN GOLDMAN IS THE PRESIDENT and CEO of the Union of Concerned Scientists, among the most well-known science advocacy organizations, with a network of 20,000 scientists. She says she recognized the need for new tactics to combat the Trump assault on science. She assumed her role in the early months of the Trump administration, fresh off a four-year stint working for the Biden White House in the Office of Science and Technology Policy. She spent her first year with the Union of Concerned Scientists exploring "what was effective and what wasn't" in the new era. Then she delivered a stark message to her troops. "We

need to up our game. We need to do something different. We need to recognize what this moment is."

In early June, she launched an initiative called "Science Rising," announcing a major shift in priorities and approach. "We don't have time to do a nice long thought piece in a scientific journal about the context," she says. "They will already have smashed seven more things by then. You've got to be quick in terms of shaping the narrative and telling the story and recognizing impacts and educating the public and decision makers about what to do."

Goldman says the Union of Concerned Scientists is positioned to help combat the "culture of fear" that has led to a silence from university scientists and administrators. It's imperative, she says, to combat the "everyman-for-himself mentality where everybody just wants to be not the target," and promote unity. The Union of Concerned Scientists has launched educational efforts aimed at promoting "scientific courage" by explaining how members can get involved, advocating with members of Congress and through the media, and using their expertise to build support. Although the

group has no in-house lawyers, they have been working with legal organizations like Democracy Forward and Public Citizen, getting actively involved in lawsuits as plaintiffs or co-plaintiffs, writing amicus briefs and providing member scientists to provide expertise and background.

TAKE BACK SCIENCE Colette Delawalla, co-founder of Stand Up for Science, fires up a crowd at "Take Back Our Science Rally," sponsored by the grassroots group, in Washington this year.



“Scientists are unbelievably difficult to mobilize. Our victory is that we built political machinery for them.”

“We have to be thinking about what comes next and what to do in a world where it’s not going to be the same,” she says. The Trump administration “dis-mantled a lot of science infrastructure and fired a lot of people who had institutional knowledge. It’s not going to be as easy because of the level of destruction that we’re seeing now. But every day that goes by that we’re allowing more destruction is a day that we will be harder to come back from. There’s urgency to do everything we can right now.”

The urgency has spurred the Union of Concerned Scientists to row in the same direction as the grass-roots groups. And it has scored some victories. Last year, when the Trump administration enlisted climate change deniers to write a report for the Department of Energy to justify the repeal of an EPA finding on climate change, the Union of Concerned Scientists and the Environmental Defense Fund sued and prevented the report from being released.

They argued successfully that the report violated the Federal Advisory Committee Act, which states that scientists from outside of government must follow a litany of rules, including being vetted to make sure they don’t have conflicts of interest, and that their deliberations are done with transparency. To the scientific activists, the conflicts of interest were apparent. And the decision by the Trump administration to abandon the effort to stamp the report with the imprint of the DOE suggests they were right.

Nevertheless, the Trump administration still rescinded the EPA’s formal 2009 scientific and legal determination that greenhouse gases—including carbon dioxide—endanger public health and welfare by causing climate change. (A finding made under the Clean Air Act.) But they were unable to produce new scientific evidence showing that climate change is not caused by greenhouse gas pollution. Instead, they emphasized the economic costs of climate regulations and suggested

those costs justified rolling back the finding. The rescission is now the subject of extensive litigation.

Stand Up for Science can also point to signature victories. They helped derail an ill-conceived initiative pushed by RFK Jr. to replace a nationwide infant hepatitis B vaccine drive in the tiny African nation of Guinea-Bissau with a clinical trial co-funded by the CDC and the New York City-based hedge fund, Pershing Square.

Since May, Stand Up for Science has fueled tens of thousands of citizens to oppose the Trump proposal to shift approval of science grants from independent peer reviewers to political appointees aligned with Trump’s priorities. “If we hadn’t made a big deal of this, there would have been less than 1,000 comments,” Delawalla says.

Senate allies folded provisions into a stopgap spending bill that will halt the implementation of the proposal. The House must still cast its vote on the proposal and the Senate may revisit it again. But the provision was a “major victory,” Delawalla says. It previews the power of the new movement for the political battles ahead. For Delawalla, mobilizing scientists has been a huge victory in itself.

“Scientists are unbelievably difficult to mobilize,” Delawalla says. “Our victory is that we built political machinery for them in the United States. The fact that we didn’t put up any sort of resistance before is part of the reason why we even ended up in this situation. Once we get through the upcoming midterms, and we have some levers of power to pull, goodness, we’ll pull them.” ☺

ADAM PIORE has been a *Nautilus* contributor since 2013. His science writing has appeared in *National Geographic*, *Scientific American*, and *MIT Technology Review*. He is the author of *The Body Builders: Insider the Science of the Engineered Human*, among other books.

The Virtue of Being Simple

A surprising number of creatures don't benefit from greater intelligence

BY BRANDON KEIM

THINK SMARTER IS ALWAYS BETTER? Consider the barnacle. The parasitic barnacles of the genus *Sacculina*, to be exact. They begin life as free-swimming larvae before growing a clam-like carapace, developing their brain, and seeking out a crab host. The brain isn't very big, but it meets the challenge of identifying a suitable crab—and, if the barnacle is male, of identifying a crab already parasitized by a virgin female—and exploring its body in search of an ideal attachment point. The barnacle then settles in, metamorphoses into something like a syringe, and injects a cell mass into the crab (or, for a male barnacle, into the female). From then on, the barnacle is more like a root system with gonads than an animal. Their brain, along with everything else they no longer need, is left behind.



The barnacles are a striking embodiment of a phenomenon that runs counter to how most people think about intelligence and evolution. *Sacculina* become less intelligent during their lives and have done so across evolutionary time, too: Their ancestors were free-living barnacles who kept their brains into adulthood. That's not how it's supposed to work in a world where the value of intelligence seems self-evident and becoming smarter is assumed to be a useful adaptation.

Scientists have argued for decades about how humans and other extra-brainy animals got that way: by living in groups, learning from others, navigating complex or unpredictable environments, or all of the above. Whatever the reasons, more intelligence is supposed to be better. Plenty of research demonstrates this principle: birds whose bigger brains help them survive in cities, fast-learning fish with greater breeding success, mouse lemurs for whom cognitive performance predicts survival rates.

Not much is said about cognition going the opposite way, but sometimes it does. A surprising number of creatures don't seem to benefit from greater intelligence. Many have become *less* intelligent across evolutionary time or even in their lifetimes—and these trajectories are not dead ends, they're improvements.

"High intelligence is just one evolutionary strategy," says Georg Striedter, a neurobiologist at the University of California, Irvine and author of *Principles of Brain Evolution*. "Some species do well in evolutionary terms precisely because they invest in other features, not intelligence."

Intelligence, of course, is a tricky word. Capacities that probably ought to qualify are downplayed while others are overvalued. Quickly answering complex math problems will show up on a kid's report card, but emotional awareness and conflict mediation probably won't.

But with that caveat, there's still a common currency: solving problems, recognizing patterns, making choices, learning from experience, remembering things, applying concepts from one domain to another, being flexible rather than routine bound. When biologists

Clams and oysters descended from creatures with a head that was lost. They didn't need it anymore.

at Switzerland's University of Fribourg exposed fruit flies to pineapples and oranges injected with quinine—a bitter, harmful substance—and observed whether the flies subsequently avoided similar fruits, they could fairly be said to be testing the bugs' intelligence.

The researchers went on to breed the faster-learning flies, tested them again, bred the highest performers, and so on. Twenty generations later, the flies had dramatically improved powers of association, but they didn't live as long as standard-issue flies and were outcompeted as larva by their slower-learning brethren.¹

The exact mechanisms behind the results—replicated years later by other researchers who bred fruit flies for longevity and found them to be especially poor learners—were unclear.² But it seemed that genes linked to learning also affected physiological processes linked to longevity. This fit with a principle first articulated in the "expensive brain" hypothesis: Because neural tissue requires so much energy to produce and maintain, cognitive gains come at a price. Originally this was thought to be a smaller gut, as digestive tissue is also quite expensive, but it can include reductions in growth, reproduction, or other tissues.

Such trade-offs are particularly well-described in insects. In addition to the fruit fly studies, experiments have found that learning may come at the cost of fecundity³ for cabbage white butterflies and result in bumblebees with shorter lifespans.⁴ "Enhanced cognitive traits are not necessarily beneficial to the foraging performance of individuals or colonies," wrote the scientists who studied the bumblebees, led by Lisa Evans and Karen Smith of the Royal Holloway University of London.

Fewer such studies exist in vertebrates, although one experiment found that bigger-brained male guppies had smaller guts and fewer offspring.⁵ (Big brains are not always correlated with intelligence—but sometimes, as with guppies, they are.) In some circumstances, though, it's not the high energetic cost of intelligence that matters, but the behaviors that intelligence produces.

One study by psychologists at the University of Exeter in England found that pheasant chicks with



an aptitude for reversing previously-learned associations—a capacity considered indicative of intelligence, as unlearning isn’t easy—were less likely to survive than their slower-unlearning peers.⁶ The likeliest explanation, surmised by the researchers, was that reversal was maladaptive. The pheasants lived in managed landscapes where food was easily and predictably available.

A bird could do well by sticking to one feeder rather than trying alternatives, and in the process, they’d cover less ground and thus be exposed to fewer predators. “Exaggerated cognitive performance may lead to maladaptive, costly behavioral outcomes, at least under some circumstances,” wrote the researchers, who thought that might explain “why we do not commonly see instantaneous learning, perfect memory, or immediate and flexible executive control.”

The pheasants’ circumstances were temporary—life isn’t always so easy—but long-lasting conditions may shape a species. Koala bears are an iconic example. Compared to kangaroos, wallabies, and other members of their family, koala brains are smaller, smoother—brain convolution is another anatomical trait that doesn’t always signify greater intelligence but does point that way—and use less energy. This appears to be a consequence of their diet: eucalyptus leaves, their primary food, are nutrient-poor and toxic.

Koalas sleep for up to 22 hours each day and make only occasional trips to the ground, a lifestyle that “indicates very little in the way of cognitive demands,”

EAT LOCAL Pheasant chicks with an aptitude for learning and seeking new feeding grounds are less likely to survive than their learning-challenged peers. The latter cover less ground and so are exposed to fewer predators.

says Ken Ashwell, an evolutionary neuroscientist at the University of New South Wales. While their cognition hasn’t been formally tested, he notes, “I think the observed behavior speaks for itself.”

A much more pronounced change occurred in myxozoans, a class of microscopic aquatic parasites whose ancestors were jellyfish-like creatures. Granted, jellyfish are not considered smart—although research suggests they show some cognitive behaviors like curiosity—but they’re university graduates compared to spore-like myxozoa, which lost many of their neural genes and possess a much-simplified nervous system.

That trajectory is typical of many parasites. Tapeworms, for example, also evolved from free-living ancestors and simplified as they adapted to life inside

a host. “Evolution does not necessarily maximize complexity or intelligence. It optimizes organisms for reproductive success in a particular ecological niche,” says Klaus Brehm, a developmental biologist at the University of Würzburg in Germany. The ancestors of fairy shrimp also had complex brains that are now much-reduced; clams and oysters descended from bottom-grazing, snail-like creatures possessing a head that was eventually lost. They didn’t need it anymore.

Finally, there are creatures, like the *Sacculina* barnacles, whose transformations occur during their own lives. It’s often reported that marine invertebrates called sea squirts, who begin life as free-swimming, tadpole-like larva, eat their brains once they’ve settled and spend their adult lives without one. That’s only partly true: sea squirts do digest much of their brain during metamorphosis but regrow a new and larger one. And consider Dehnel’s phenomenon, in which the brains of certain shrews, stoats, and moles shrink by up to 25 percent each autumn.⁷ This reduces their cognitive performance but also means they’ll need less energy to survive the coming winter. Come spring their brains regrow.

All these trajectories from more to less intelligent are certainly not universal, but they undermine the assumption that evolution necessarily moves toward higher intelligence—and even the very premise of “higher” intelligence.

“It is frequently suggested that a high intelligence from the evolutionary standpoint can be equated with ‘superiority,’” wrote the late Eugene Robin, a medical professor and physiologist at Stanford University, in a 1973 paper, “The Evolutionary Advantages of Being Stupid.”⁸ “This type of thinking has resulted in an anthropocentric view in which man because of his intelligence is regarded as the pinnacle of evolution.” To Robinson that was a self-serving human conceit.

Evolution, wrote Robinson, did not proceed in some unidirectional way from simple to complex, low to high. He quoted the great science-fiction author Isaac Asimov: “Which is the fitter, a man or an oyster?” It was a rhetorical question. Were Earth suddenly flooded, oysters would prevail. Fitness depends on circumstances, and circumstances can change.

Once-helpful adaptations can become handicaps. The enormous bodies that once served dinosaurs so well had also been their downfall, noted Robin; and

the brain that designed spaceships and composed the *Eroica* symphony also built hydrogen bombs and Auschwitz. Were Robin alive today, he might also remark on how the fruits of human ingenuity have led to our destabilizing the biosphere itself.

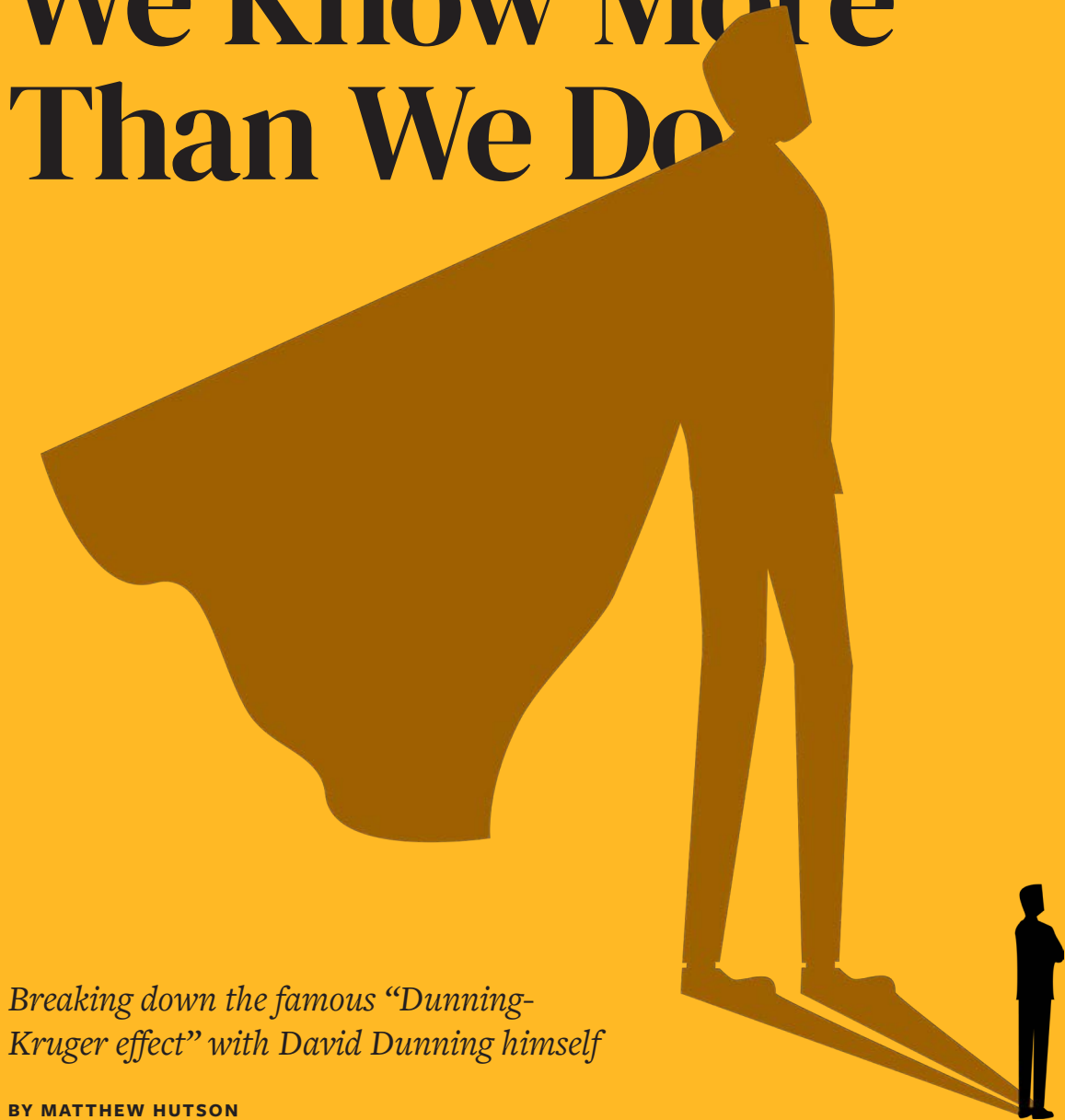
“Under a given set of conditions it might be more useful from the standpoint of survival to be simple rather than complex or to be general rather than specialized,” Robin wrote. “It might be more useful to be ‘stupid’ rather than ‘intelligent.’” ☺

BRANDON KEIM is a freelance journalist and contributing editor at *Nautilus*. His book *Meet the Neighbors* explores what the science of animal intelligence means for how we understand and live with the wild creatures around us.

REFERENCES

1. Mery, F. & Kawecki, T.J. A fitness cost of learning ability in *Drosophila melanogaster*. *Proceedings of the Royal Society B* **270**, 2465–2469 (2003).
2. Burger, J.M.S., Kolss, M., Pont, J., & Kawecki, T.J. Learning ability and longevity: A symmetrical evolutionary trade-off in *Drosophila*. *Evolution* **62**, 1294–1304 (2008).
3. Snell-Rood, E.C., Davidowitz, G., & Papaj, D.R. Reproductive trade-offs of learning in a butterfly. *Behavioral Ecology* **22**, 291–302 (2011).
4. Evans, L.J., Smith, K.E., & Raine, N.E. Fast learning in free-foraging bumblebees is negatively correlated with lifetime resource collection. *Scientific Reports* **7**, 496 (2017).
5. Kotschal, A., *et al.* Artificial selection on relative brain size in the guppy reveals costs and benefits of evolving a larger brain. *Current Biology* **23**, 168–171 (2013).
6. Madden, J.R., Langley, E.J.G., Whiteside, M.A., Beardsworth, C.E., & van Horik, J.O. The quick are the dead: Pheasants that are slow to reverse a learned association survive for longer in the wild. *Philosophical Transactions of the Royal Society B* **373**, 20170297 (2018).
7. Lázaro, J., Dechmann, D.K.N., LaPoint, S., Wikelski, M., & Hertel, M. Profound reversible seasonal changes of individual skull size in a mammal. *Current Biology* **27**, PR1106–R1107 (2017).
8. Robin, E.D. The evolutionary advantages of being stupid. *Perspectives in Biology and Medicine* **16**, 369–380 (1973).

Why We Think We Know More Than We Do



Breaking down the famous “Dunning-Kruger effect” with David Dunning himself

BY MATTHEW HUTSON

S

SOME YEARS AGO, I started taking house-dance classes. (House is a specific style of dance that goes well with house music.) For more than two decades, I'd been dancing at raves and parties, where people complimented me and sometimes even asked if I was a professional. While I didn't know what to expect at the studio, I had reason to think I might not be terrible. I immediately found myself the dunce in the back of the class. I hadn't known what I didn't know.

My experience demonstrated what's known as the Dunning-Kruger effect, an often-cited and still-debated concept in both psychology and popular culture.

In 1999, David Dunning, a psychologist then at Cornell University, and Justin Kruger, a Cornell graduate student, published a seminal paper called "Unskilled and Unaware of It." The paper described four experiments testing college students on humor, logic, and grammar. Dunning and Kruger also asked participants to guess their scores, in both absolute and relative terms. The key finding that kept popping up: Those in the bottom quartile of performance wildly overestimated their skill in that domain (in some studies rating their percentile, on average, around 60 when it was around 10). A smaller finding: Those in the top quartile often underestimated their relative skill (though to a lesser extent).

Dunning and Kruger theorized that the skills needed to perform overlapped with those needed to judge performance, and so poor performers lacked the metacognition to see their failures. Combine that with a universal "better-than-average effect," in which people overestimate their abilities compared to other people, and the worst performers overestimate the most. Meanwhile, high performers also think they're better than average but underappreciate the extent to which they are, because of the "curse of knowledge," in which we think our knowledge is universally shared.

It's not often that an academic paper finds its way into mainstream culture, but when it does, it's often misunderstood. I recently contacted Dunning, now at the University of Michigan, to talk about the effect's legacy. He still vigorously defends the idea, while exploring new manifestations, and has even turned the term into an adjective. As Dunning says, we are all "Dunning-Kruger" at some point in our lives.

What was the origin of your 1999 paper?

Well, the research focus I've had off and on is: How well do people know themselves? And just through personal observation, of students in my office, faculty meetings, watching C-SPAN, I would notice that people would say—how should I say it diplomatically?—outrageous things. They have to know what they're saying can't be true. I've had students come into my office to graciously allow me the chance to concede that the answers I designated as right on the test were wrong.

In his first year of graduate school, Justin Kruger said he wanted to work with me. I said: How much do people know that they're getting something wrong? And so, we did these studies to see how much people who were performing very poorly on tests were aware of it. We tested logic, grammar, and whether people could spot funny jokes. The first data set that came in was a revelation. I was absolutely astonished. Of course, people are going to overestimate themselves, but we didn't realize it was going to be at that magnitude.

Others had found that ignorance leads to overconfidence—in physics, chess, tennis, reading, driving, socializing. What was the new piece here?

I think the new piece was putting numbers to it, and putting the formal psychological analysis to it, this distinction between cognition when making a judgment and metacognition when judging the judgment.

Do you think of the effect as just the overconfidence of people who are unskilled or also the underconfidence of people who are highly skilled?

There was a long literature about the “curse of knowledge” for people who are very knowledgeable. This is why doctors are terrible communicators. They think you know where the kidney is. They think you know that this drug has cardiac implications.

Do you consider that part of the Dunning-Kruger effect?

I’m going to say no. By the way, we didn’t name it the Dunning-Kruger effect. Someone on the internet years later started talking about the Dunning-Kruger syndrome and then threw up a Wikipedia page. A lot of people just say the effect is that people at the bottom don’t know they’re at the bottom. Period, end of story. But it’s also the mechanism that connects those two failures.

What’s a striking example of Dunning-Kruger today?

AI is the most interesting. The first obvious question is: Do chatbots experience Dunning-Kruger? And the answer is: Oh, yeah. Except for some exceptions, they don’t know when to hold back due to uncertainty. There was a game recently where you make up a phrase and ask a chatbot what it means, and it would give you an answer instead of saying it didn’t know. The other question is: Do chatbots give people Dunning-Kruger? One group found it reversed Dunning-Kruger and made the best performers even more overconfident than the worst people. There’s also work showing that people are very Dunning-Kruger in terms of their ability to tell fake news from real news.

SIDE EFFECTS “By the way, we didn’t name it the Dunning-Kruger effect,” David Dunning explains. A Wikipedia entry did. Further, he says, “I wish people would stop yelling ‘Dunning-Kruger’ at each other. We offered the paper as a reason to self-reflect.”



People are overconfident because they look for reasons supporting their ideas.

In other work, one professor allowed students to buy insurance if they felt they were going to be in the bottom 50 percent of the class, but very few did. So, misestimations lead to behaviors that don't work very well. Others have shown that people who are most likely to dismiss scientific consensus are those who know the least science but have the most confidence in their knowledge of the science.

Do you think AI sycophancy, in which chatbots flatter their users, might heighten the Dunning-Kruger effect?

Yes. One of the reasons why people are overconfident in general is that they tend to look for reasons supporting their ideas. In fact, one of the pieces of advice we always give is: If the decision is important, stop and ask why you might be wrong. In planning, one of the procedures that's recommended for a group is to project yourself into the future and imagine your initiative failed spectacularly. What happened to make it fail? And say, "OK, what can we do to prevent those stories from happening?" That's called a "pre-mortem."

Doctors do this as a matter of course. They have a different name for it. They don't diagnose you; they do a "differential diagnosis." They may have an idea, but they have to think about what else it can be. OK, you might have Lyme disease, but what else are your symptoms consistent with? We're going to test you to rule out everything that we can. That's core to medicine, thinking of alternatives, thinking how you could be wrong.

Your 1999 paper mentioned "motivational biases": People are scared to look for ways that they might be wrong, or to think of themselves as stupid. Do you think acknowledging that there are multiple intelligences or multiple skills might allay people's fears of being generally stupid, and might make them more receptive to negative feedback about the specific areas in which they're unskilled?

Absolutely, yes. Because first off, Dunning-Kruger is not about being stupid. Dunning-Kruger is about

everybody. Each of us is Dunning-Kruger sooner or later. Some of us have more spectacular stories to tell, but every one of us has pockets of incompetence. It's not about being stupid; it's about being careful and wise. And the other thing is that you are a person who grows and changes. So what you want to do is set yourself up for success.

I heard someone use the Dunning-Kruger effect to explain that people who opposed his political beliefs were too stupid to know that they were wrong. Have you seen other people make claims like that? And is it more likely that that is true, or just that people are motivated by different ideologies?

Well, people can be stupid, but people can be motivated by ideology. I wish people would stop yelling Dunning-Kruger at each other. We offered the paper as a reason to self-reflect. I remember when I started looking at Twitter back in the day and seeing people yelling, "You Dunning-Kruger idiot" at everybody else. I went, "Oh no." My biggest regret is that it's become an epithet to hurl at other people. It's something to hurl at yourself every so often.

How can we check our blind spots?

The best way to do it is with other people. Now, they have to be brave, and they have to tell you that you're wrong. And I wish we were more skilled at being able to tell each other when we have blind spots.

Is overconfidence ever good?

Confidence isn't bad. Confidence is something to be managed, and there are times you really want to be confident. In preparation, you don't want to be overconfident. You want to be, if anything, underconfident, and overprepare. If you're going into battle, or if you're going to the Olympic Games, you definitely won't just go, "Nah, I won't practice today. I won't train. I'm going to get the gold." But on the day that you're executing, you do want to be supremely confident, because it'll have a payoff.

Some of us have spectacular stories to tell, but every one of us has pockets of incompetence.

When should we most worry about the Dunning-Kruger effect?

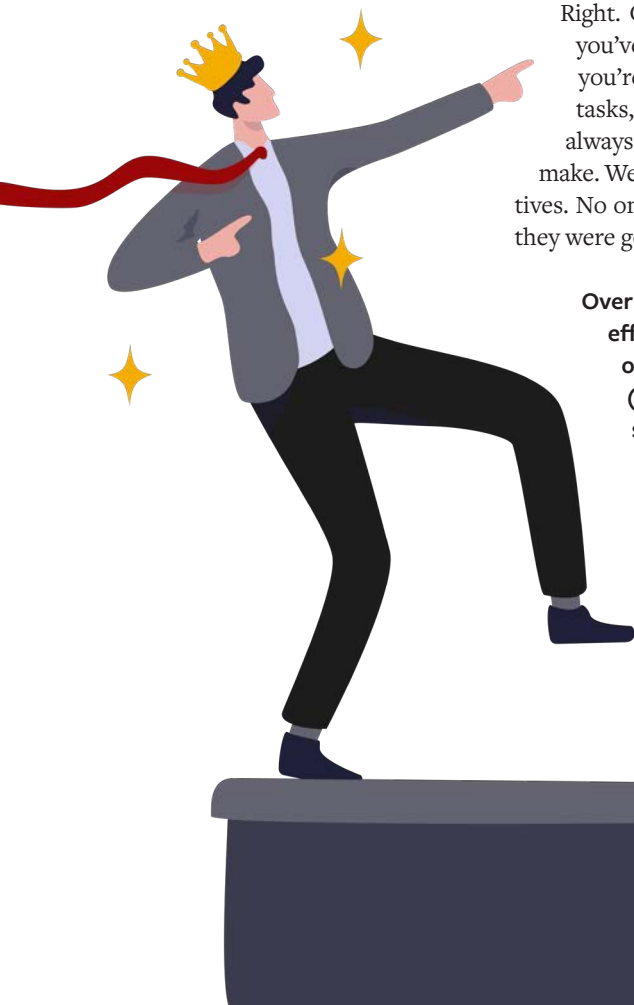
When there's a lot of overlap between performance and judgment. In logic, they're practically the same thing, but not in sports. I know I'm a horrible golf player. I can tell because the trees tell me when my ball travels to the right and lands in the poison ivy. Also beware lack of exposure to outside models. I thought I was a pretty good cello player, until I heard a recording by Jacqueline du Pré, and I went, "I didn't know you could do that with this instrument."

There's some research showing that people are overconfident on hard tasks and underconfident on easy tasks.

Right. Confidence is always going to be moderate to high, because you've come to the best answer you can come to. On hard tasks, you're probably not going to get to the right answer, and on easy tasks, you're always going to get to the right answer. Still, we're always consigned to be moderately confident in the decisions we make. We make a decision and that's what we're going to do. We're captives. No one, except for George Costanza in *Seinfeld*, says that whatever they were going to do, they will do the opposite. No one ever does that.

Over the years, there have been critiques of the Dunning-Kruger effect, namely that it boils down to the statistical phenomenon of "regression to the mean." [If you take two noisy measures (in this case, someone's skill and someone's estimation of skill), and the first is extreme (very high or very low skill), the second will necessarily trend toward the middle. As a result, the incompetent appear to overestimate and the competent appear to underestimate.] What do you make of this criticism?

It ignores the literature on how you correct for regression to the mean. There are many standard ways to correct for it, from epidemiology, economics, and psychology. It dampens the effect only a little.



We live in a bubble of our own knowledge.

There are other methodological issues. In one example, the researchers asked people to rate their intelligence on a 1 to 25 scale after taking an IQ test. They then assumed that people rating themselves as a “25” thought they had a 160 IQ and those a “1” thought they had a 40 IQ. Those are pretty big assumptions. I wouldn’t put those thoughts in people’s heads. And I also wouldn’t ask about an ability as broad as intelligence, which means so many different things to different people.

In another example, the data turned out to be beautiful support for our original framework, although I cannot figure out how the authors couldn’t see it. They statistically separated the ability to judge a performance from the ability to produce a quality performance and found that high performers were no better at judging their performance than those at the bottom. So, they said, we were wrong. However, before the separation, high performers were better at judging, because expertise aided both performance and judging. That last part is exactly what we said in our original 1999 paper creates the effect. In real life, there is no separation. Their data are exquisite in making our point.

Has your understanding of the effect evolved?

I’m beginning to see more implications of the idea that we live in a bubble of our own knowledge. Lately, we’ve done work showing that experts are better able to tell when they’re making a right decision versus a wrong decision. But there’s a twist. When they’re wrong, they’re even more confident that they’re right than novices are. We see this again and again.

Experts’ being more confident than non-experts that they’re right when they’re wrong seems to go against the Dunning-Kruger effect.

On aggregate, experts are better at being right, and at knowing when they’re right. But experts have to be aware that being highly confident is not necessarily a guard against being wrong.

Is that what’s sometimes called “Nobel disease,” where Nobel Prize winners will espouse on topics that they don’t know much about?

Nobel disease is actually a different syndrome, also called “epistemic trespassing.” That was just everywhere during the pandemic. People started opining about vaccines and epidemiology, and they had no business doing it, but, “Hey, you know, I have a Harvard degree, so that’s good enough.” A lot of people’s beliefs of their competence come from these soft cues they have about themselves. That’s not knowing your circle of competence. It goes to a fever pitch when you have a Nobel.

This is about the geography of being “an expert.” Do they know the boundaries of their competence? We just published a paper showing that if people start taking classes on a topic like psychology and law, they start thinking they know stuff they can’t possibly know. At the end of a psych and law class students will say, “Yep, I know this psycho-legal concept.” They also say they know other concepts that we asked about but were never mentioned in class. And they weren’t mentioned in class because we made them up. The boundary line between what you know and what you don’t know is fuzzy. It’s a profound thing of being human that we live in a bubble of our own experience. ☺

MATTHEW HUTSON is a freelance science writer in New York City and the author of an upcoming book on intelligence, *Smarts*.



AI Made Me Look Foolish. Then I Learned

Ancient alchemy has a lot to teach us about chatbots

BY PAUL M. SUTTER



T HIS PAST FEBRUARY I stood in front of my collaborators to walk them through a major update to VIDE, the void-finding algorithm I first built more than a decade ago. Voids are the vast empty regions between galaxies, and digging them out of a survey is fussy, unglamorous work. The new version was 10 times faster than the old one. It could handle surveys a hundred times larger. Best of all, it had a far more sophisticated way of dealing with the ugly realities of an actual data set, the gaps and masks and ragged boundaries that never show up in the textbook version of the problem.

I had spent weeks on it, and I was proud of it. I had also spent those weeks with an AI at my elbow, and it had been indispensable. Coding agents rewrote the core of the algorithm, faster than I could have and without the typos. When it came time to write it up, they generated the plots and drafted the commentary that went with them. Every time I checked their reasoning, they never once hesitated.

Ten minutes into my presentation, a collaborator raised their hand. They said that something seemed off to them and asked about the new scheme for handling the edges of a survey. I answered easily because I knew this material. They asked again, slower, and pointed. I looked where they were pointing, took a breath, and the floor went out from under me.

The edge handling was wrong. It wasn't a typo, and it wasn't a missing citation or a factor of two. It was subtle, but it was very wrong, and everything downstream of it was also wrong, and I had shared the whole thing in a room full of people who trusted me. The code was fast. The code was beautiful. The code was incorrect. After 20 years of training, I had presented nonsense with total confidence.

*After 20 years of training,
I had presented nonsense
with total confidence.*

Yes, I still cringe when I think about that moment. And here's the worst part: In all those weeks, I never felt a flicker of doubt. Nothing ever seemed iffy, or even worth a second look. It felt like mastery of a new technique in less than a tenth the time it would have normally taken. The fashionable worry is that AI is making us stupid, that every delegation is a small amputation. I don't believe it. Nothing was taken from me in February.

But I will admit this: I was seduced. I felt AI was a modern-day philosopher's stone, imbued with the power to transmute machine into mind. The ancient alchemists spent centuries chasing a legendary substance that would turn lead into gold and make its owner rich beyond counting. Today we pour in text and electricity and wait for thinking to come out the other end. That's what makes these tools feel magical instead of merely handy. Not that they're fast. It's that they seem to have crossed a line that nothing made of arithmetic should be able to cross.

So, there I was in February, showing off the rotten fruits of my transmutation and the fool's gold I now held in my hands.

But the story doesn't end with my humiliation. Alchemy was already on my mind, and the more I read about it, the more I realized the alchemists' work wasn't an embarrassment to science, it was the road that led to it. They never found the philosopher's stone. But they found something else, something we need really badly right now.

I FELT LIKE I UNDERSTOOD what I was doing because the AI system I was using—a large language model, or LLM—sounded like it understood. An LLM is a little more than a sophisticated next-word predictor. The reason these systems are right about most things turns out to be a boring one. True statements are simply more common in the human-written text LLMs are trained from, so a true statement is usually the better prediction. Correctness is a side effect, and it falls out of the arithmetic the way heat falls out of an engine.

And then we made it worse. A raw predictor is unpleasant to talk to, so we polished it, showing people different forms of answers and tuning the model toward whichever they liked better. What we liked, it turned out, were answers that were clear, organized, and direct.

An LLM's fluency is not an accident, and it is not an emergent mystery. It's a trait we bred, the way we bred wolves into dogs that watch our faces when we open the treat bag. Being confident but wrong can cost a human their job, their career, or just their basic self-respect.

Machines don't have jobs, or careers, or respect.

But dang it, LLMs are so undeniably useful. I have been writing code since I was 5 years old. I haven't written a line of it since January. And for 20 years I did my own research, which is to say I read the papers, chased the citations down their rabbit holes, and built every argument myself. AI agents do all that now.

These were not lazy trades. They were correct trades, and I'd make them again tomorrow. Somewhere in the past year an LLM became my *thinking* partner, and I say that without embarrassment. I'm not the only one. Hundreds of millions of people reach for these tools every day, and they are not idiots.

Users of LLMs get acceleration that is useful and correct a large fraction of the time, and total junk sometimes. And LLMs are so good at what they do, at finding the word that is both useful and pleasing, that I lowered a guard I have kept up my entire career.

The superpowers are real, but they're volatile, and they're tuned to please us. So how do we deploy a tool that is sometimes wrong but always pleasing? How do we trust AI? The simple answer is: don't.

We've been here before. Not with machines, but with any powerful process (whose insides we cannot see) that returns something valuable. When I felt like an alchemist with new powers at hand, I was only half right.

I work differently now. Whatever comes out of the machine, I distrust.

Today we think of alchemy as a pseudoscience, a quasi-magical system of attempting the impossible. But the alchemists weren't false scientists, they were pre-scientists, and they discovered in the laboratories many things about the world. Working with volatile substances they did not understand (they did not know about atoms, molecules, quantum mechanics, or all the rest), in Europe and the Middle East for more than a thousand years, they managed to discover phosphorus, recreated porcelain, isolated alcohol, and built laboratory apparatus still in use today. Their words are scattered through our language: crucible, alcohol, gas.

They managed this because they had a method. Ignorance of their materials was met with an almost fanatical commitment to the operational process, to rigorous documentation and strict control, to keep the work from evaporating into illusion.

We are now in the pre-chemistry era of AI. We know how these models work, because we crafted their code ourselves, and yet in most cases we cannot say why: why this word and not another, why this particular path through the machine. The crucible is closed, and like the alchemists we are not going to stop using it. The alchemists were the ones in history who had to pull real value from a process they could not see into. Now it's our turn.

THE ARCHITECTURE OF ALCHEMY had three parts that can help us navigate the seduction of AI.

The first principle is grounding. The alchemists called it fixing the volatile. A substance that fled the vessel on the way to becoming something was a failure of the work, and much of the art consisted of rendering matter stable enough to stay put and be weighed.

To us, any output that cannot touch the earth is dross. Today's safeguards treat truth as an optional garnish: a timid warning at the bottom of the chat window, a filter scrubbing for toxicity. We exert instead a grounding pressure of verification, granting the machine no authority until every claim it makes has been checked against the world.

The second is provenance. Nothing in the alchemist's vessel was a private event. They recorded everything—the color of the flame, the source of the materials, the phase of moon and positions of the planets—because they believed in the universal connectedness of nature. The most famous alchemical text, the Emerald Tablet, said it thusly: *superiora de inferioribus*. The higher from the lower. That chain holds even when the middle is mysterious.

For us, it means auditing and chain-of-argument. Can we trace every result back through the conversations? Can we point to every output and map it to

an input? Can we own not just the result, but the chain that led to the result?

Grounding catches the fabricated citation, the number from nowhere, the substance that evaporated. Provenance catches something else, the flawless output that traces cleanly back to an idea that was already wrong. It helps us find the impurities, and where the LLM went off the rails.

With LLMs grounding and provenance will carry you a long way, and then they stop. A visible thread tells you where a claim came from. It says nothing about whether the source was any good. Somebody has to take hold of the far end of the thread and judge what is tied to it, and that somebody cannot be the machine, because the machine is where the thread began.

That is the third principle, and the one the AI industry is spending a decade trying to engineer away. We have been led to believe that the goal of AI is autonomous intelligence, but this is the ultimate alchemical error. Labor can be delegated; authority cannot. While the industry chases after gold, they ignore that the machine is ontologically incomplete; it possesses the logic of computation but lacks the spark of life.

Ars totum requirit hominem: The art requires the whole person. The alchemists held that the practitioner was no mere observer but the active catalyst, binding the volatile elements with their own hands. They believed that only good people were capable of truly great works. Judgment is not a technical operation. To decide that something is true is to decide what you will put your name on, and no machine has a name of its own. It is us that transforms a formless oracle into a useful tool. If the machine does 90 percent of the work, the remaining 10 percent (our own moral weight and intuition) becomes the entirety of the value.

A COUPLE HOURS with my new method would have caught what happened to me in February, had I maintained the discipline of my training instead of just receiving whatever output was in front of me.

So, here's what alchemical discipline means. I work differently now. Whatever comes out of the machine,

Alchemy isn't an embarrassment to science. It was the road that led to it.

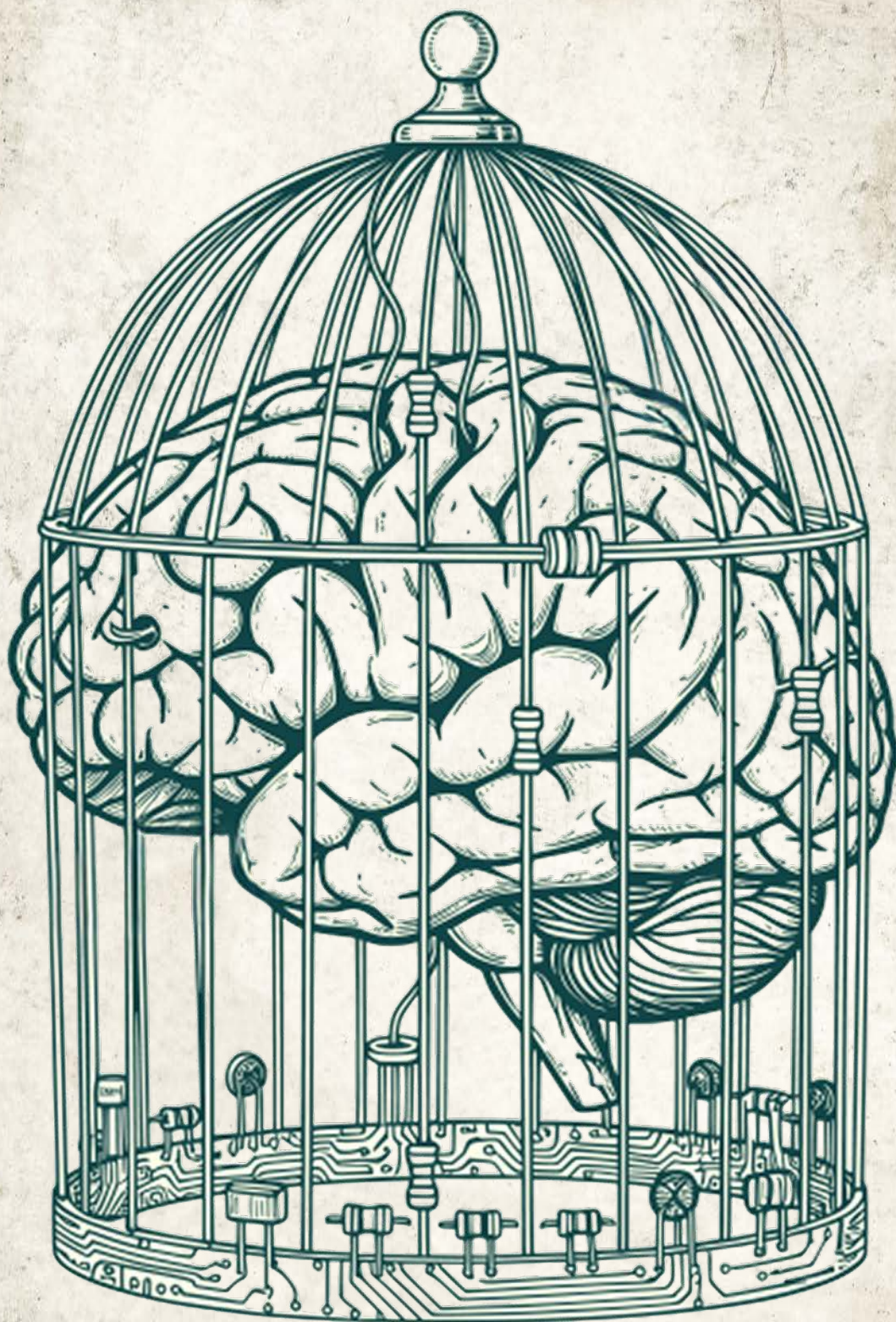
I distrust, and not just when something smells wrong. I test it, I judge it, record what I did, and only then do I put my name on it. This is what an alchemist did every morning, standing in front of a vessel that would never explain itself: a fanatical devotion to process in the face of ignorance.

The alchemists never found the philosopher's stone, not after 1,400 years of trying. But something even

more powerful came out of the alchemical tradition. Robert Boyle, the father of modern chemistry, was an alchemist. So was Tycho Brahe, whose observational astronomy unlocked the laws of planetary motion. So was Isaac Newton, who wrote more on transmutation than he ever wrote on gravity. What they found instead of the stone, while failing to find it, was a way of working. A method. In their rigorous attention to detail, in their devotion to process, they created the seeds of the scientific method.

The philosopher's stone of AI is never going to be found inside the machine. Only in the discipline of the hand that holds it. ☺

PAUL M. SUTTER is a research professor in astrophysics at the Institute for Advanced Computational Science at Stony Brook University and a guest researcher at the Flatiron Institute in New York City. He is the author of *Your Place in the Universe: Understanding our Big, Messy Existence*.



When Stupid Was a Diagnosis

A dark history that is in danger of returning as support for people with disabilities withers

BY ALICE FLEERACKERS

THERE WAS A TIME when throwing up during pregnancy wasn't just seen as a frustrating symptom. Instead, some doctors considered this fetus-induced queasiness as the harbinger of a lifelong, debilitating condition for the unborn child—and they called this condition “idiocy.”

Back in the late 1800s, beliefs about this so-called “poverty of the mind” were rapidly evolving. Vomiting during pregnancy was just one of many hypothesized causes, with everything from neurotic parents to the weather considered possible culprits at one time or another. At the extreme, some medical professionals even saw idiocy as a moral failure—a sign that children, or their parents, had “so far violated the natural laws” that their bodies were no longer fit to house the “powers of the soul.”

It is easy to condemn these early theories from where we stand today, decades of medical and psychological research later. But doing so would be a disservice to the complex and surprisingly recent history of what we now understand as intellectual disability.

“The history is an incredibly sobering one,” says David Wright, historian at McGill University in Canada. “But try to understand the past in its own terms ... not to ignore the terrible things that occurred, but also try to understand the context in which they occurred.”

At the extreme, some medical professionals even saw idiocy as a moral failure.



Wright grew up in Ontario, Canada, with a sister, Susan, who was born with Down syndrome. His family's experiences navigating Susan's condition motivated him to study the history of the asylums—or as they were later called, “institutions”—where people with intellectual disabilities were sent to spend their lives in medicalized settings.

“My sister was born in the 1960s, and my mother was encouraged to send her to an institution,” he says. “That is what one did, right? That would be better for her, that’d be better for the family, that’d be better for society.”

Wright’s mother declined, as did a growing number of parents facing the same decision in that era. Although these institutions were still the norm, shocking exposees by journalists, photographers, researchers, and advocates had started to call their use into question. Harrowing accounts of life on the inside revealed frequent cases of physical, sexual, and emotional abuse: Heavily-sedated people walking through the halls in drug-induced stupors. People sleeping on concrete floors, naked, sometimes in their own excrement. By the time Niels Erik Bank-Mikkelsen, director of the Danish Service for the Mentally Retarded, visited the Sonoma State Hospital in 1967, he described it as “worse than any institution I have seen ... In our country, we would not be allowed to treat cattle like that.”

Despite these disturbing accounts and growing public skepticism, society’s longstanding reliance on these facilities meant it took more than three decades for most of the last long-stay residential institutions in North America to close. “These were full of thousands of people,” Wright explains. “You couldn’t just discharge them all.”

“WHAT THE HELL is going on here?”

That’s what historian and writer Simon Jarrett found himself wondering when he arrived at one such institution, back in the 1980s. He’d been working with people with intellectual disabilities for years, but never in an institutional setting like this one. In need of a job, he decided to try a role as a nursing assistant, not knowing quite what he’d encounter.

“There were guys ... who’d been in there since they were 5 years old,” Jarrett recalls. “They had gone in in the 1910s or the 1920s and they were just going to live out their whole lives there. And I thought, ‘How did this come about?’”

That initial question became the seed for a decades-long career exploring the history of intellectual disability, especially the centuries before long-term residential institutions became the norm. By digging into old court proceedings, caricatures, jokes, and other documents describing the daily lives of “idiots,” “half-idiots,” and

Many parents simply didn't have the resources to care for the loved ones they'd previously sent away.

"imbeciles," Jarrett's research has shown that society has not always been so quick to hide people with disabilities from view.

"They were part of the conversation," he says. "The institution simply wasn't an option for disabled people prior to this big asylum movement that happened in the 19th century, so the locus of care for people was in the home and the neighborhood—the community."

Particularly striking are the "Old Bailey" criminal court proceedings, which describe countless examples of people sticking up for "idiot" neighbors, friends, or family members who had been accused of minor crimes in a London court from 1674 to 1913.

Take the story of Ann Wildman, for example, who was tried for stealing "nine yards of black silk ribbon" in 1762. Multiple community members stepped forward in her defense, including her mistress, who argued that "she was very silly, but had behaved honestly before" and promised to take her into her charge if her name was cleared.

The trial of Peter Cunniford, accused of stealing a coat in 1759, tells a similar story. Cunniford's brother, in-law, employer, former roommate, and two colleagues all testified to his character, emphasizing that he was "next kin to an idiot" but also hard-working, well-behaved, and, most importantly, "honest." Both Wildman and Cunniford were eventually acquitted, as were many other people who presumably had intellectual disabilities tried around the same time.

Of course, life before institutionalization was far from perfect for people with intellectual disabilities. There were no formal supports available, some people were teased or bullied, and many lived in very poor conditions. But Garrett emphasizes that quality of life was worse for almost everyone back in the 18th century, not just those with intellectual disabilities—a kind of "equal opportunity terrible," as he puts it.

Stories like those in the Old Bailey trial records are important, because they suggest that people with intellectual disabilities were not necessarily outcast or marginalized. In fact, for centuries, the designation "idiot" was not used as an insult, but as a legal term, reflecting whether a person was mentally fit to manage their own property, get married, or be held criminally responsible for their actions.

For Garrett, the Old Bailey trials are also significant because they are evidence that people with intellectual disabilities were often able to work and live at least semi-independently in their communities. Some, like Wildman and Cunniford, were described by witnesses as valued workers and productive members of society.

This way of thinking seems to have been eradicated by decades of institutionalization, as Garrett experienced firsthand when institutions started to close and people were resettled back in their communities.

"Everyone in the hospital told us it wouldn't work: 'Well, they'll all walk into the road and get knocked over by a car. They'll fall out of a window, because they'd never been upstairs,'" he recalls. "[But] they did incredibly well ... the only thing that happened was they all burnt their mouths on the food, because they weren't used to having hot food."

IT'S CLEAR THAT SOMEWHERE between the 1700s, when the bulk of the Old Bailey trials took place, and the late 1900s, when institutions finally shut their doors, society's understanding of intellectual disability underwent a radical shift. What's less clear is exactly what led to this misguided notion that people with intellectual disability had to be contained and controlled, rather than integrated.

Garrett attributes the move toward institutionalization to a wider change in mindset sparked by the French

Revolution around the end of the 18th century. The radical overthrow of the social hierarchy that took place in France led both reactionaries and progressives internationally to see people with disabilities as problematic—but for very different reasons.

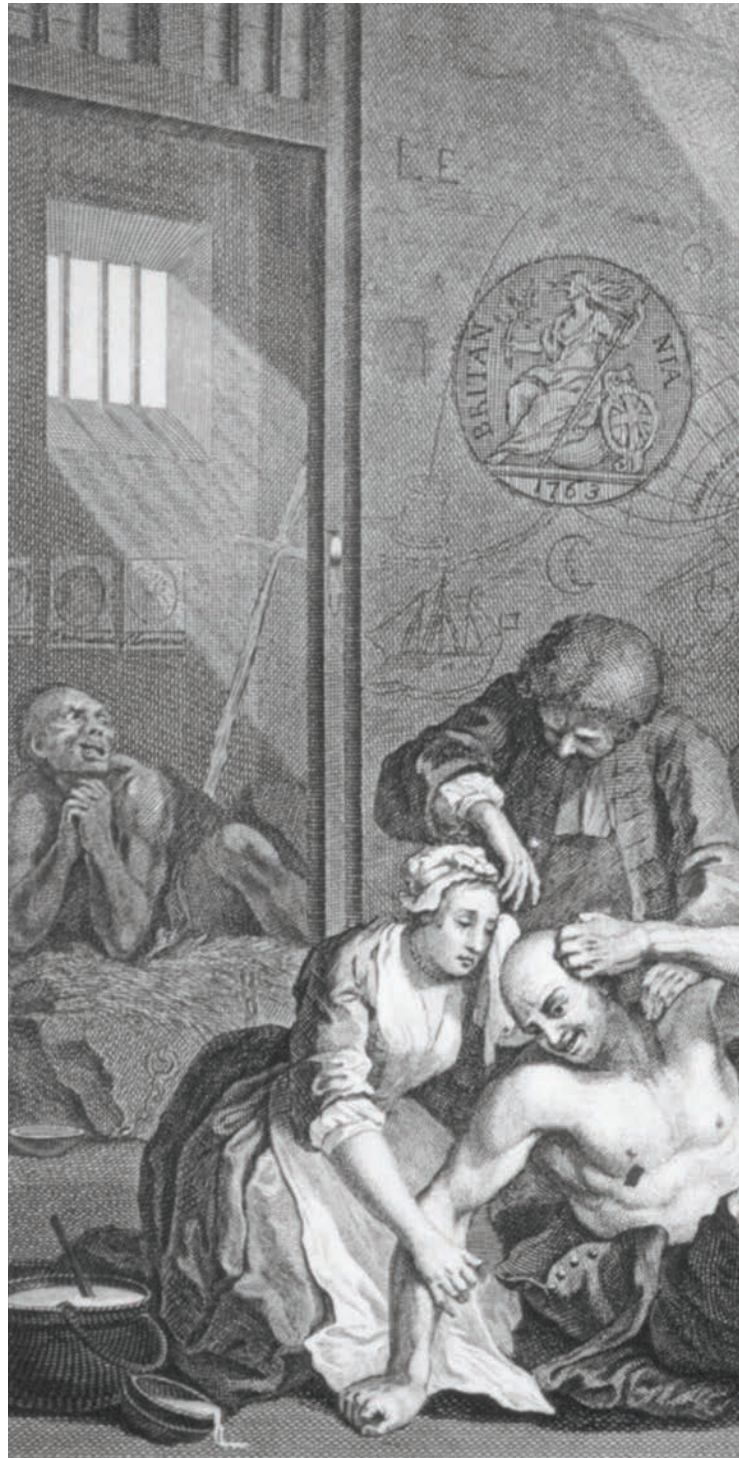
The reactionaries were horrified that the poor, whom they saw as an “imbecile class,” had overthrown the wealthy elite in France. As a result, Garrett says elites “became very frightened, suspicious, hostile toward anybody that displayed extreme signs of difference from anyone else, because these people were a threat.”

The progressives, on the other hand, had started to envision a better future, promoting utopic visions of a society free of illness, disease, and, in the most optimistic accounts, even death. “For them, the [intellectually] disabled person ... was not one who belongs in this new society,” Garrett says.

Adding to this were other factors, including the growing authority of the medical profession, scientific attempts to investigate idiocy and its causes, and Enlightenment beliefs that humankind could be improved. According to Wright, financial considerations also played a role, as many “respectable” families did not have the funds to provide private care for their loved ones.

Wright says the institution was introduced in this context as a promising solution, as an “investment of society.” The idea was to provide modern medical care and, where possible, treatment for people suffering from disabilities and mental illness. Institutions were often framed as temporary, educational facilities where children could spend a few years gaining self-care skills, developing basic literacy, or learning a trade, before eventually reintegrating into society.

STRANGER THAN FICTION This engraving was made from the 18th-century painting *The Madhouse* by English artist William Hogarth as part of a series called *A Rake's Progress*. It traced the tale of Tom Rakewell, who wastes his family fortune and ends up in Bethlehem Hospital, a mental institution. Though the story is fictional, the scene is based on London's infamous Bethlehem (Bedlam) Royal Hospital, an asylum where people with all manner of psychiatric disorders were mistreated in squalid conditions for centuries.





History has shown when people are segregated, the one thing you can count on is abuse.

“It’s a very progressive, self-congratulatory approach,” Wright explains. “An argument that [people with intellectual disabilities] are amenable to educational and social improvement in a way that past generations, who were foolish and unscientific, never appreciated.” At least in theory, institutions were introduced as a path toward social progress.

The problem, Wright explains, is that these facilities failed to deliver the utopian results they promised. Residents did not typically “improve” in the ways that proponents of institutionalization had hoped. Moreover, in the cases where children with intellectual disabilities were able to be discharged, “some families refused to take them back.” Many parents, especially those near the poverty line, simply didn’t have the resources to care for the loved ones they’d previously sent away.

While more and more children entered these facilities, very few ever left them. In this sense, Wright says these institutions “became victims of their own success”—too popular to house the droves of people entering them. Huge swaths of residents were needed to work the fields that produced the food some of these institutions ran on, creating a self-perpetuating cycle where more and more residents were needed to feed the growing number of mouths housed in an institution’s walls. Over time, facilities became overcrowded, funding dwindled, and conditions rapidly deteriorated.

In perhaps the darkest moment of disability history, progressive ideals about “improving” humankind started to take a new, more extreme shape toward the end of the 1800s, as eugenicists promoted deeply unethical beliefs about controlling human genetics to produce a more intelligent, productive, “ideal” human race. In line with this vision, institutions kept women and men in separate wings to prevent them from reproducing. Some started sterilizing residents—a practice that continued as late as the 1970s in some parts of North America. In California alone, almost 300 people with intellectual disabilities were sterilized during the 20th century, under this deeply flawed mandate to produce the “perfect” society.

Eugenics-driven sterilization efforts disproportionately impacted people at the margins of society, including people with disabilities but also immigrants, the poor, and people of color. Black people were especially impacted, because they were seen by the white elite as “idiotic and imbecilic” by default—a perception reinforced by racial biases in the intelligence tests that were used to assess “feeble-mindedness.”

IT’S UNDENIABLE THAT CONDITIONS for people with intellectual disabilities have improved in the decades since. We no longer sterilize people or prevent them from reproducing. Most people with intellectual

disabilities live at home or in smaller community facilities, rather than in overcrowded institutions hidden from society's view. In the United States, legislation currently ensures people with intellectual disabilities have the right to go to school and live life, just like anyone else. With support, they can now learn alongside their peers, grow up, get jobs, and be part of society.

But Mary Hartley, president of the disability advocacy organization The Arc of Greater Pittsburgh (Achieva), warns that this legislation may be in jeopardy. "There have been multiple discussions, if not threats, from United States federal offices about these hard-fought rights," she says. "We really can't afford to go back at all ... it is deeply, deeply concerning."

On June 18, 2026, the U.S. Department of Justice published a legal opinion that questions a seminal disability rights decision from the Supreme Court, *Olmstead v. L.C.*, delivered in 1999. *Olmstead* was a landmark moment, as it recognized the institutionalization of people with disabilities as a form of discrimination under the Americans with Disabilities Act (ADA). In the decades since, it has allowed people and their families to advocate for the support they need to live life in their communities, rather than cordoned off from society.

The recent opinion presents a reanalysis of the 1999 documentation to conclude that *Olmstead* "held nothing more than that unjustified institutionalization of individuals with mental disabilities can constitute discrimination on the basis of disability." While this recognition of discrimination is often used to "demand certain treatment services from states," the legal opinion argues this right to services "is not granted" by current legislation and is thus "unlawful."

While an opinion statement alone isn't enough to change the law, The Arc has warned it cannot be taken lightly. "[R]ights mean less when the federal government refuses to enforce them," reads a statement on the organization's website. "This opinion seeks to undermine one of the strongest protections people with disabilities have from being pushed into institutions when they can and want to live in the community."

The publication of the government's opinion on *Olmstead* follows several other concerning developments, including the largest ever funding cuts to Medicaid, the governmental financial support that makes community living possible for people with intellectual

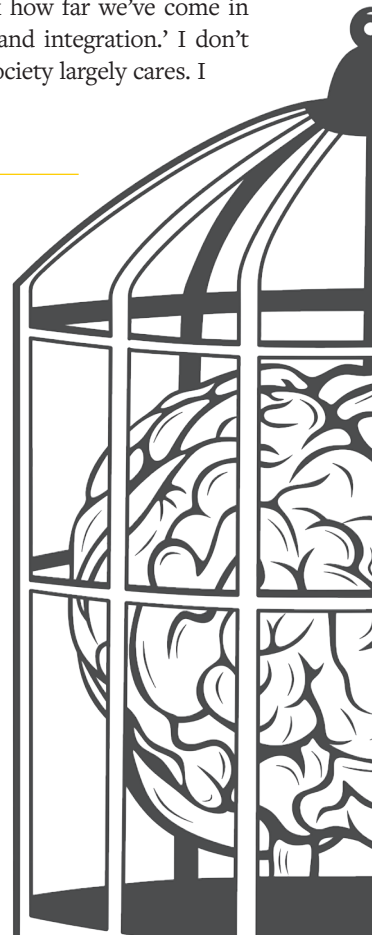
disabilities, and a planned reorganization of the offices responsible for special education that could make it more difficult for parents to demand supports for their kids—and to hold schools accountable when they don't provide them. Collectively, these developments threaten to undermine decades of advocacy work to ensure people with disabilities are integrated, rather than institutionalized.

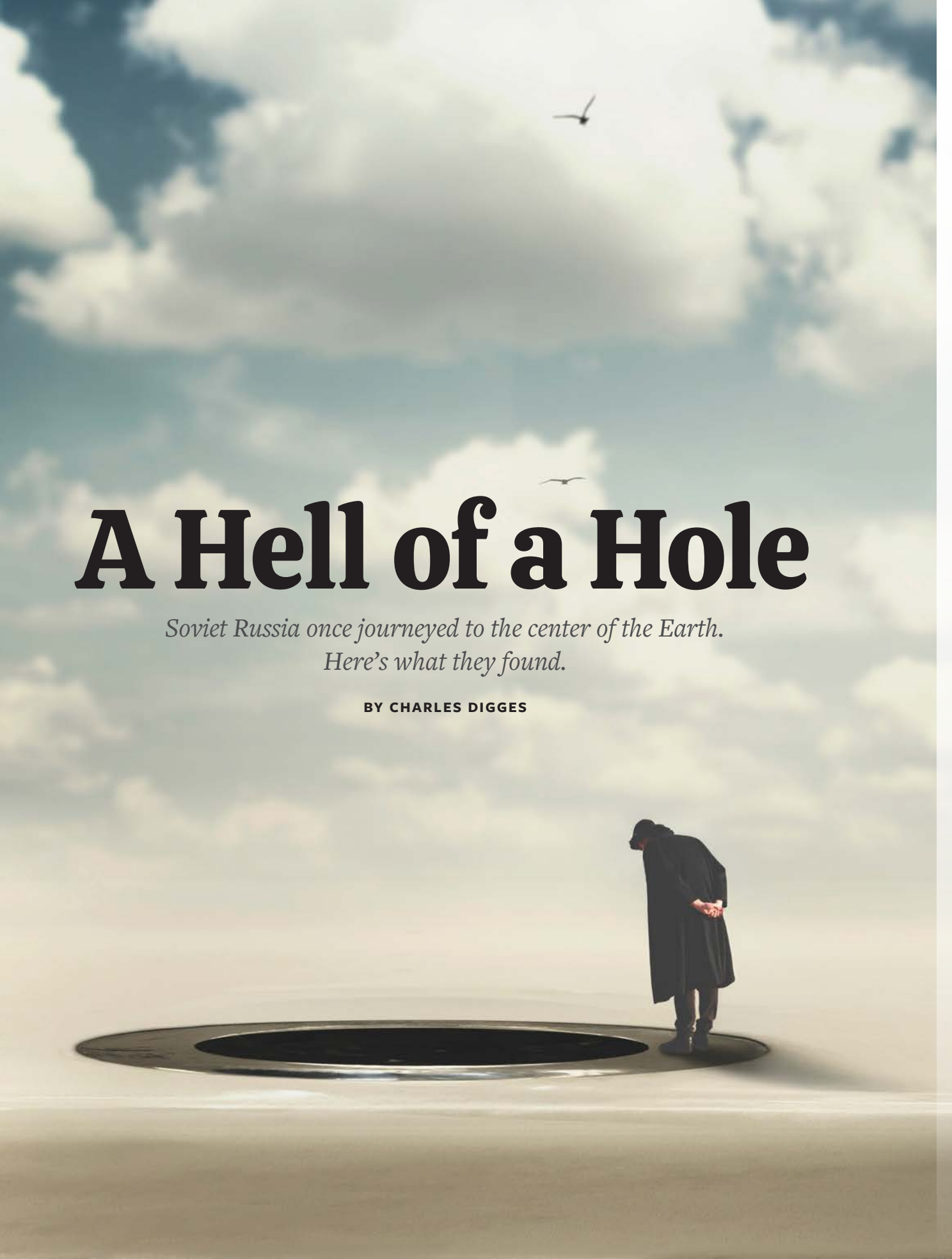
As a parent of a son with autism, these threats hit close to home for Hartley. "I'm deeply, deeply concerned about ... the health and safety of people with disabilities," she says. "History has shown when people are segregated, the one thing you can count on is abuse ... When other people don't have eyes on you, can't report something, that's when abuse happens, that's when neglect happens."

These legal, policy, and financial threats underscore how fragile the line separating historical injustices from current realities can be, and how easy it is to overlook the flaws of the present when scrutinizing the past. "We're talking about how marvelous we are, because we're moving [people] back [into community]," Garrett says. "But there's much more of a conditional acceptance now than there was before. You have to conform ... to be an acceptable person in the community."

Wright agrees: "I don't think we can look at our society today and say, well, 'Look how far we've come in terms of providing supports and integration.' I don't think we have. I don't think society largely cares. I don't think it's a priority." 🌀

ALICE FLEERACKERS is a freelance writer fascinated by psychology. She is also an assistant professor of journalism and civic engagement at the University of Amsterdam and the vice president of the Global Network for the Public Communication of Science and Technology.



A full-page background image showing a person in a dark, hooded cloak standing on a light-colored, flat surface. They are looking down into a large, dark, circular hole in the ground. The sky above is filled with large, white, fluffy clouds, and a few birds are visible flying in the distance.

A Hell of a Hole

*Soviet Russia once journeyed to the center of the Earth.
Here's what they found.*

BY CHARLES DIGGES

I**N THE LATE 1980s**, from certain corners of cable TV, you could learn that the Soviet Union had dug straight to Hell. Not in the metaphorical, Reagan-Era, “look what those commies have done” sense of the word. No. They had literally dug straight to Hades.

The story ran on a program called *Praise the Lord*, the long-running worship and current events program aired by the evangelical Trinity Broadcasting Network, which reached around 5 million daily viewers. Over several episodes, the program explored how a Soviet deep drilling project had accidentally tunneled into Satan’s underworld.

When the Soviets lowered a remarkably heat-resistant microphone into an underground cavern discovered by their dig, so the story went, the Russians—to their horror and surprise—recorded the wailing of damned souls as they languished in temperatures exceeding 1,800 degrees Fahrenheit. Their lamentations were played on the air for a stunned audience. In later developments, the network reported that a winged, bat-like figure had been seen rising from the hole and blazing a trail across the Siberian sky. This was all the proof they needed to assert that, out there somewhere in the Siberian tundra laid the “Well to Hell.”

Then the story began to unravel. First, the Soviet drilling project in question was actually taking place hundreds of miles to the west of Siberia along the border with Norway on Russia’s Kola Peninsula. Next, the recordings of doomed sinners voicing their eternal misery turned out to be just a sound loop from a 1972 b-movie horror flick called *Baron Blood*. Then a Norwegian teacher came forward to say he’d planted the story about the winged demon just to see how gullible the producers of *Praise the Lord* really were.

Nevertheless, the myth has since migrated from Christian cable TV to credulous quarters on the internet where, on YouTube and TikTok, you can still encounter that same recording as proof of the everlasting inferno beneath the Siberian tundra.

The network reported that a winged, bat-like figure had been seen rising from the hole.



STEP RIGHT UP! The Kola Superdeep Borehole is 40,000 feet deep, the deepest human-made hole on Earth, but it doesn't look like much from the surface. Today it's plugged up with a 9-inch steel plate, surrounded by debris, and blends into the landscape.

In fact, the Soviets really did drill a very, very deep hole—by 1979, the deepest on Earth. Today, with a little effort, you can walk right up to the actual hole the evangelicals called the Well to Hell. Alas, it's really called the Kola Superdeep Borehole, part of a Soviet science program. You can stand right at the edge, but there's no danger of falling in. It's far too narrow for that. It's also plugged up. Since 2008, the very slender chasm has been bolted shut with a dinner-plate-sized steel cover that measures only about 9 inches across. But the hole it hides plunges down for more than 40,000 feet into the earth and took more than 20 years to dig. That's deeper than the altitude at which most

commercial airlines fly and deeper by about 20,000 feet than any other hole anywhere else on Earth.

The deepest hole on Earth. That sounds pretty stupid as far as scientific projects go—sort of like the abiding sandbox fantasy so many American kids have of digging clear through the planet to the other side, from say, Michigan or Texas or Montana, all the way to China.

But whereas the Well to Hell revelation turned out to be nothing more than a hoax, the real Kola project actually delivered something sublime: a radical new understanding of the center of the Earth. Unlike the tortured voices on that cable-access tape, what the drill recorded needed no embellishment.

THE TEAM OF SCIENTISTS who plotted the dig for a decade before the project launched had real ambitions: They hoped to bore through the continental crust and reach the Mohorovičić discontinuity—the boundary separating Earth’s crust from the mantle below. Crossing it would give scientists their first direct look at the vast layer that drives plate tectonics, feeds volcanoes, and ultimately shapes the continents.

The dig was, in a sense, a mirror of the space race. If Sputnik and Gagarin proved Soviet superiority in the heavens, an ambitious drilling project into the bowels of the planet meant the Soviets could reign in the underworld, too. The date the Soviets chose to formally begin the first round of drilling was May 24, 1970, the 100th anniversary of Lenin’s birth. The project was tied to the Soviet Union’s own founding myth.

“We had promised [Nikita] Khrushchev we’d dig to 15,000 meters (49,213 feet)—a good round number,” Sergei Nestrenko, who was then a young drilling engineer at the Kola Superdeep Borehole, told me by phone from his home in Murmansk.

But they never did get where they were going. They didn’t even really come close. Even at 40,230 feet, the deepest point humanity has ever drilled, the Kola borehole had penetrated barely a third of the way through the continental crust beneath northern Europe.

By most practical measures, you could say the Kola Superdeep Borehole failed. But that failure transformed the way geologists view our subterranean world. Until then, scientists understood the deep continental crust largely through indirect evidence. They bounced seismic waves through the Earth and measured subtle variations in gravity and magnetism, then inferred what kinds of rocks must lie below. It was a successful method—except where it wasn’t.

“Think of these methods like you think of an ultrasound,” Ulrich Harms told me when we spoke recently. “You can see your kidneys in great resolution. But if no one had ever dissected a body beforehand, you wouldn’t know that’s a kidney.”

Harms, a geologist and drilling engineer who first visited the Kola site in 1984, would go on to help lead many of the world’s largest scientific drilling projects

The biggest surprise was what the drill never encountered.

through Germany’s International Continental Scientific Drilling Program. Kola, he said, became one of the first opportunities to compare the seismic models to the real thing. The real thing, he said, offered some wild stuff.

Perhaps the biggest surprise was what the drill never encountered. For decades, geologists had interpreted a prominent seismic boundary deep beneath the continents as the transition from granite to basalt—the dense volcanic rock that forms most of the ocean floor. But Kola never hit basalt. The supposed boundary was nowhere to be found. Instead, the drill continued through more granitic and metamorphic rock, showing that the seismic boundary did not mark a change in rock type at all. It marked a change in how the same rocks behaved under immense pressure and temperature.

Other discoveries proved just as momentous. Nearly 22,000 feet beneath the surface, drill cores contained microscopic fossils from two dozen species of plankton that had drifted through an ancient sea more than 2 billion years ago. This showed that even after billions of years of burial deep within the continental crust, unmistakable evidence of some of Earth’s earliest life had survived. Between 16,000 and 23,000 feet, the borehole encountered saline water circulating through fractures in crystalline rocks that geologists had long considered essentially dry. Rather than groundwater that had somehow seeped miles underground, the fluids appear to have originated within the crust itself. As ancient minerals were transformed by heat and pressure, they released chemically bound water, while reactions between water and iron-rich rocks generated hydrogen gas, showing that the supposedly inert continental crust was still chemically active billions of years after it had formed.

“Taken together, this all revealed that the deep crust wasn’t dry and inactive,” Harms told me. “Instead, it was a dynamic place where ancient life, water, and chemical reactions continued to exist and interact billions of years after the rocks themselves had formed.”

JUST AS THE WORLD BEGAN paying attention, the Soviets almost lost the whole project. In August of 1984, Moscow hosted one of the largest scientific meetings

ever held, and for the first time the Soviet Union invited foreign researchers to see its deepest drilling project firsthand. This was the 27th International Geological Congress, a gathering of nearly 5,700 geologists from 110 countries. By then, the Kola borehole had already reached 39,521 feet, farther beneath the Earth's surface than anyone had ever drilled. Harms was among the visitors. "It was extraordinary," he told me. "Nothing else in the world came close. They were doing really pioneering work."

For a brief moment, the Soviet Union stood unambiguously at the forefront of one of science's most ambitious engineering enterprises. The Americans had abandoned their own deep drilling Project Mohole nearly two decades earlier—at a mere 601 feet—after cost overruns and management disputes, while Germany's own five-and-a-half-mile deep KTB borehole was still years away. Whatever else could be said about the Soviet system during the twilight years of the Brezhnev

stagnation, it had succeeded in drilling deeper into the Earth than anyone had thought possible.

Then, only days after the geology Congress ended, catastrophe struck. On Sept. 27, 1984, as drilling resumed, roughly 3 miles of aluminum cylinders, at the end of which spun the drill bit, twisted off and became irretrievably lodged deep underground. Years of painstaking progress vanished in a clank of metal.

"At least it happened after everyone went home," said Nestrenko. "But there wasn't any question of giving up."

Nestrenko said the engineers had no choice but to abandon the broken section, back the drill up to roughly 23,000 feet, and begin drilling again along a new branch. It would take another five years before they regained the depth they had already reached.

BY THE EARLY 1990S, the Kola borehole project was nearing its end. The hole extended to such a depth that the Earth had stopped behaving like Earth. At roughly 40,000 feet, temperatures climbed to nearly 356 degrees Fahrenheit—almost twice what geologists had predicted before drilling began. The immense heat, combined with the crushing pressure of the overlying crust, transformed the ancient crystalline rock into something more akin to molasses.

SUPER DEEP TRUE DRILLING STORY

The remains of a 27-story tower built for the Kola Superdeep Borehole project. The tower was meant to protect heavy machinery, crews, and laboratories used for sample analysis from the extreme weather of Kola Peninsula in the Arctic Circle.



The rock simply no longer behaved like rock.

“It was like a flow in the ancient granite,” Nestrenko told me. “Like stiff plastic.”

To anyone who thought they knew granite, that sounded ridiculous. Granite is supposed to fracture. Hit it with enough force and it breaks into chips that a drill can carry away. But nearly 8 miles beneath the Kola Peninsula, the rock had crossed some bizarre invisible threshold. Instead of shattering, it slowly deformed. As soon as the drill bit passed, Nestrenko explained, the surrounding granite began creeping back toward the borehole, squeezing the shaft shut almost as quickly as the engineers could deepen it. Every trip to replace a worn drill bit gave the Earth another opportunity to reclaim the hole they had just created.

“It would pinch the hole closed,” Nestrenko recalled. “When we pulled the drill string back, the rock flowed inward.” The engineers had reached the limits of drilling. The rock simply no longer behaved like rock.

Then in 1991, the Soviet Union dissolved, taking with it the financial and symbolic backing needed to build new drills or find novel solutions. The newly formed Russian government, stumbling awkwardly into economic collapse, called off drilling in 1992. The broader research program wound down over the following years, with the site formally mothballed in 2008 and left to decay.

When I drove through the area about 10 years ago, it appeared as though the 700 or so geologists, engineers, drilling technicians, and others who worked at the dig had one day risen from their lab benches and just walked away. There’s little reason to think much has changed since. The sliver of highway between Murmansk and the Norwegian border along which the remnants of the borehole reside are as remote and desolate a wasteland as you could hope to find. A low-slung pair of warehouse-like laboratories have been ground down by the arctic winds.

Only the tower housing the old drilling column had been dismantled. The rest of the site sits like a ransacked museum, frozen as it was left, vulnerable to metal and wire scavengers. In contrast to the other secretive relics of the Cold War that surround it—the ones you can get arrested for trying to find—it’s almost as if the Russians didn’t see their hole as a secret worth guarding.

When I stood at the edge of this giant Arctic ghost town, the whole project seemed to me to be less a question of engineering and practicality than a question of minds at play.

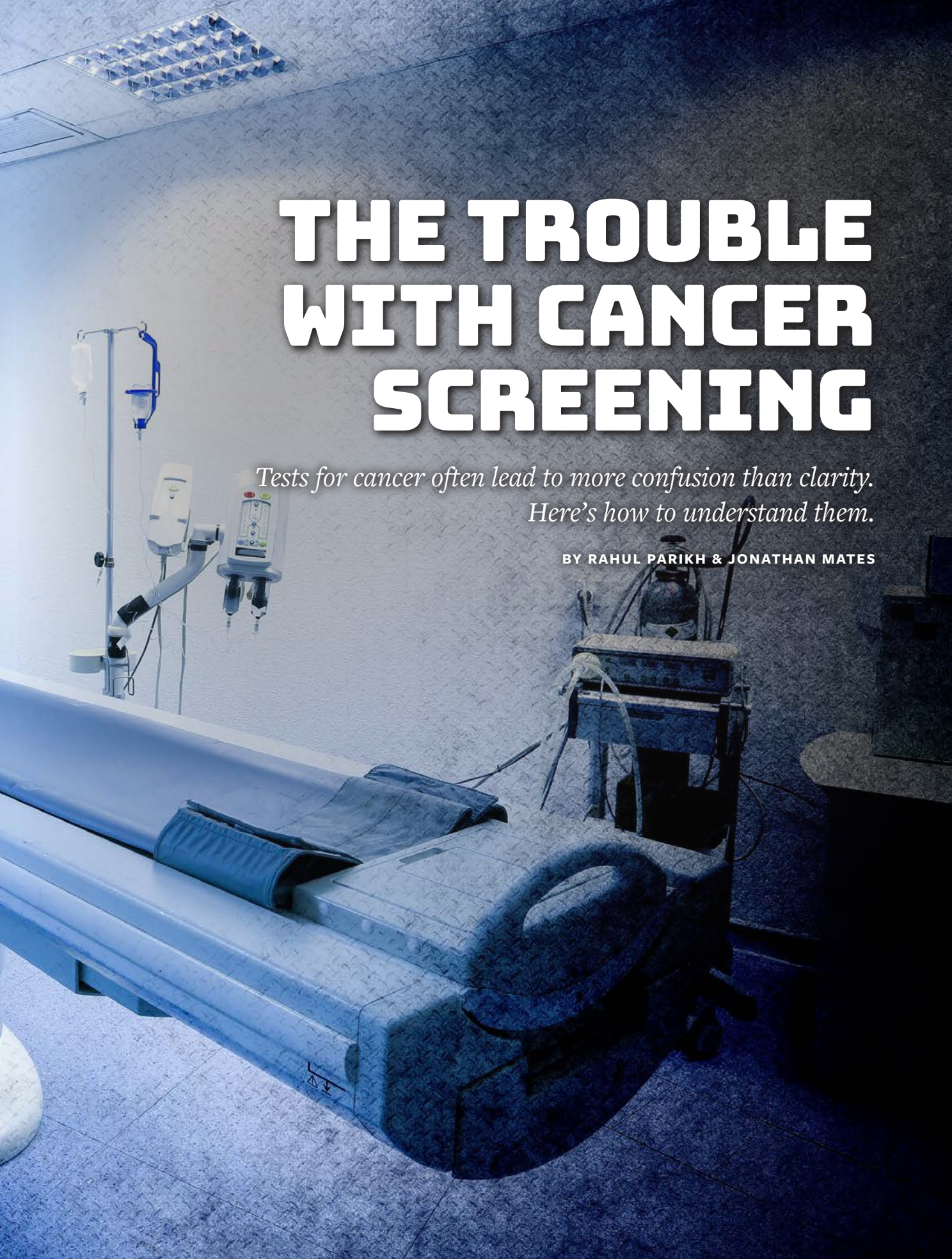
The children in the sandbox never make it to China—any more than the Soviets were going to reach the Earth’s mantle. But the measure of the Kola Superdeep Borehole turned out to have little to do with the goal that justified the effort.

Instead, the Kola Superdeep Borehole overturned decades of geological assumptions. It showed that the continental crust was hotter, wetter, and more chemically active than scientists had imagined. It also revealed that one of geology’s most trusted road maps was leading us astray. Along the way, it pushed drilling technology to limits that engineers still regard with admiration. None of those discoveries had been the point.

That might be the value of asking a big, stupid question. The ones that seem most sensible often just confirm what we already know. The ones that are childishly fanciful can force us to discover where we’re wrong and forge new paths forward into the dark. ☺

CHARLES DIGGES is an environmental journalist who edits the website of the Norwegian environmental group Bellona.



A dimly lit hospital room with a patient bed, medical monitors, and an IV stand. The room is dark, with a blueish tint. A patient bed is in the foreground, with a blue blanket and a white pillow. To the left, there is an IV stand with two bags. In the background, there are medical monitors and other equipment. The overall atmosphere is somber and clinical.

THE TROUBLE WITH CANCER SCREENING

*Tests for cancer often lead to more confusion than clarity.
Here's how to understand them.*

BY RAHUL PARIKH & JONATHAN MATES

IT WAS A TYPICAL OVERCAST MORNING in the San Francisco Bay area as both of us, who are doctors: Rahul, a pediatrician, and Jonathan, a diagnostic radiologist, headed to work.

As Rahul strolled through the parking lot, another day of patients to see in his office, he was notified of a medical test result by a phone push notification. The app took its sweet time loading as it connected him to his medical record. A few days earlier, he had seen his doctor for a routine checkup. The bundle of tests was ordered algorithmically by his doctor. One-by-one, he scrolled through them.

Blood count, normal; electrolytes, normal; cholesterol: normal

Prostate Specific Antigen, *elevated*.

Rahul stopped strolling, his eyebrows shot up. He blinked and stared back at his phone. He hustled toward his office and logged into his workstation. He pulled up his schedule for the day along with searches on Google and on Medline, trying to recall from medical school decades ago what a slightly elevated PSA test—a biomarker of potential cancer—meant.

Was it really elevated? It depended on his age, right? Would he need an invasive biopsy? Should he ignore the result? And most importantly: *Did he have cancer?* Rahul's doctor told him they would repeat the test in three weeks. It came back normal. The next time he tested it was slightly elevated again. A fourth test brought the same result.

A screening test is like a random sobriety checkpoint by the Highway Patrol on New Year's Eve.

Rahul continued to read about the PSA test. Medical consensus once recommended that every man over 50, as Rahul was, should have the test. That was now softened to screening after a conversation between the patient and doctor. Some experts suggested that regular screening be dropped altogether. The potential harms of treatment might outweigh the benefits, as it turned out, and many men who had prostate cancer ended up dying of something else, the cancer growing so slowly that it never caused a problem.

Rahul's doctor didn't talk him through these facts. He didn't even tell Rahul he was ordering the test. Rahul was in a pickle. Now that his PSA was abnormal, should he go down the road of invasive biopsy? Should he just ignore it? It felt unfair to have to make these choices. He just wanted to *know*. He just wanted to feel *certain* of what he should do. The fact that nobody could provide him with that certainty was frustrating and frightening. The whole thing just seemed so ... stupid.

Did it have to be this way? Could Rahul, or any patient, ever feel certain?

To answer, let us explain what screening tests—and cancer screening tests, in particular—are all about. Maybe then we can all get closer to understanding one of the most critical issues in medicine.

WHEN SCREENING DOES HARM

Medical tests come in two major varieties. The first is *symptomatic* testing, which is ordered based on your individual medical complaints, physical exam, or other indications that a disease is present. If you have a fever and a hacking cough, your doctor may order a chest X-ray to see if you have pneumonia. If you have morning nausea or a missed period, your OB-GYN may order a pregnancy test.

The second category is *asymptomatic* testing—screening—which is performed to diagnose various diseases and conditions in a population susceptible to risk, who have no current indications of disease: no pain, discomfort, or any abnormal signs. Everyone within the defined population is encouraged to get tested. Examples of this include mammograms for breast cancer on women over 40, and colorectal screening for colon cancer for people over 45.

If symptomatic testing is like the highway patrol pulling someone over because they are driving erratically, then an asymptomatic screening test is like a random sobriety checkpoint on New Year's Eve. The former is done because of a specific issue (a driver can't stay in his lane) and the latter is a general search in a population where the risk is high (anyone driving after midnight on New Year's Eve).

Sometimes the best answer—as hard as it is to accept—is to not test for cancer at all.

When it comes to cancer, a good screening test is one that allows early detection before patients develop symptoms, has an effective treatment associated with it, and has a favorable balance of risks, harms, and benefits. And when screening helps people find cancer early and get cured, medicine's mission has been accomplished.

For every triumph, though, there are many like Rahul's, where cancer screening leads them down a rabbit hole of uncertainty toward further, often intrusive and nerve-racking procedures. There are still others where the screening test is negative, but patients later learn that the tests were wrong. The cancer is detected only later, when it's already harder to treat.

The first lesson is that a positive screening test for cancer does not necessarily mean a patient has cancer. But that result does trigger additional testing to better determine what is really going on. And that's where the slope gets slippery.

One of the most persistent harms of screening is *overdiagnosis*, where we catch cancer that would never have affected the individual. This goes against our instincts, which tell us that if we have cancer, it should be caught and treated no matter what. But in fact, especially as we age and other causes of death become more common, we can often live with a slow-growing cancer that never really has other repercussions for our health.

By performing autopsies on men who died from some other cause, for example, we find that a large proportion of them had prostate cancer and never knew it, and it likely didn't shorten their lifespan.

Overdiagnosis due to screening carries significant risks. A high incidence of men with elevated PSAs who go on to get treatment—biopsies, surgery, radiation, hormone therapy—suffer complications that include long-term urinary or bowel incontinence, erectile dysfunction, and higher rates of bladder cancer.

The truth is, we don't always have a way of knowing which cancer will lead to death, and which can be ignored or watched until it causes a problem. And even if we did have a way of differentiating, it's very hard for a patient to hear they have cancer, but not do anything about it. Sometimes the best answer—as hard as it is to accept—is to not test for cancer at all.

South Korean physicians and patients learned this the hard way. As a result of national health insurance, patients could be screened for common cancers like breast and colon for free. Beginning in 1999, for a small copayment, they could also receive neck ultrasound screening for thyroid cancer. Soon after, the nation found itself engulfed in an epidemic of thyroid cancer.

Suddenly, people had to decide whether to have their thyroids cut out, which would mean a life sentence of thyroid replacement medication and regular thyroid

blood tests for the remainders of their lives. Many who got the surgery also learned they were afflicted by a deficiency in parathyroid hormone, which is produced by a small gland on the underside of the thyroid, often removed during thyroidectomy.

After the screening campaign, the number of people diagnosed with thyroid cancer went up 15-fold. Despite this spike, mortality did not change. Each year, before and after the screening campaign began, 300 to 400 people died of thyroid cancer in the country. By 2014, doctors and public health advocates were starting to speak up and question whether thyroid cancer screening was leading to overdiagnosis and overtreatment. In 2015, the government-funded National Cancer Center Korea changed its guidelines to recommend against thyroid screening.

THE RADIOLOGIST'S VIEW OF BREAST SCANS

Why can't doctors tell which positive results are real, and which ones are false alarms? And even when a cancer is real, why can't they tell whether it will cause life-altering trouble down the line? Let's consider the breast cancer screening path.

On a typical day, Jonathan sits down to read 100 to 200 screening mammograms. A completely normal exam might take two to three minutes—two views of each breast, magnifying every part of each one. If Jonathan sees something he can't convince himself is normal, then the exam is "positive," and he and his team ask the patient to return for more imaging. About 10 percent of screening patients get called back. Too many more, and the physicians are probably being too sensitive. Dip under that 10 percent threshold, and the physicians may begin to worry that they are missing things. Federal law requires radiologists to track these and other numbers, keeping them accountable.

The problem with cancer is that it doesn't know when to stop growing and doesn't respect normal boundaries, so it classically looks like the shadow of a mass with spiky edges, surrounded by distorted tissue or other secondary findings. One finding doctors look for is "microcalcifications"—little bits of calcium that appear bright white on a mammogram. In some forms, these blobs look a lot like the after-effects of the cancer, similar to burn scars after a fire—irregular, ragged, the visible residue of a cancerous process. But cancer doesn't always announce itself

that way. Some microcalcifications are just small clusters of innocent-appearing dots yet still represent cancer. Harmless-looking and harmless don't always go together. Wouldn't it be great if they did.

We manage this uncertainty by creating thresholds for who we call back. The rules are designed to catch as many cancers as possible, while not calling too many patients back unnecessarily. When it comes to calcifications, number, newness, shape, and groupings are the criteria. If a calcification meets the threshold, Jonathan calls for more imaging with even greater magnification, knowing most won't be cancer. We could magnify everything upfront and get a closer look from the start, but that's more radiation, more resources, more time. About 40 million mammograms are performed annually in the United States, so doubling the time required for each exam would come at a significant cost. We accept that some women will be called back unnecessarily, because the alternative costs more than it saves.

Jonathan also knows that among the cases he reads every day, there are cancers he is missing. Three dots of calcium instead of five don't make the cut. Others he just doesn't see or are obscured by normal tissues. The line between normal and abnormal is genuinely thin, and the system is designed with that knowledge, balancing the harm of missing cancers against the harm of calling back too many healthy people. But it's still Jonathan making that call, case by case. He follows the rules and feels the weight of that.

When radiologists call patients back and do more imaging, they still aren't diagnosing cancer, they're just making the next decision in the process, in this case determining what needs a biopsy. They're fairly certain some things are cancer or benign. But there's a zone in between—an overlapping Venn diagram of appearances—where a physician can't reach the threshold of certainty either way. So, they take a tissue sample and look under the microscope.

Microcalcifications don't always represent invasive cancer, which always requires treatment, they are often due to ductal carcinoma in situ or DCIS. One way to think of DCIS is as "stage zero" breast cancer—cancer cells still nicely confined within the milk duct. A bad character, for sure, but still behaving itself, even if it's growing. There is discussion about how to characterize it for the lay public—early cancer? Pre-cancer?

The system is not stupid. The science is not stupid. The real problem is expecting certainty.

Not cancer at all? But that's just branding. Whatever we call it, there's a really important feature of DCIS: Sometimes it will affect a patient's life, and sometimes it won't.

We know this because of autopsy studies. Nearly 1 in 11 women who died of other causes were found to have DCIS that had never bothered them—clinically undetected and irrelevant during their lifetime. And yet, left untreated, somewhere between 10 percent and 50 percent of DCIS cases would eventually progress to invasive cancer. The problem is that we just can't tell which ones. We haven't found a reliable marker—genetic, morphological, or otherwise—to foretell the future, and instead we treat them all, knowing we are overtreating some. The ones who never needed treatment, these are the ones we call “overdiagnosis.” But which ones are they?

Medicine is actively trying to solve this puzzle. Four international trials are currently evaluating whether low-risk DCIS can be safely watched rather than surgically removed. And when the results are found, checked, and confirmed, there may be new guidelines for when to follow and when to treat DCIS.

We do our best to balance the harms, risks, and benefits for breast cancer and every cancer in a screening program, but we must set some rules and guidelines, and that makes the screening tests and programs imperfect. For every line we draw, someone will be caught on just the wrong side of that line, and the system will fail them.

MAKING PEACE WITH UNCERTAINTY

Even with the best evidence, screening programs clash with the harsh reality of the real world. After all, tests only work if patients receive them. As most of us know, colonoscopy as a screening tool for colon cancer is no fun. It requires us to stop eating and drinking and instead drink uncomfortable prep cocktails of electrolytes to clear our bowels. It requires sedation, with its associated risks, and has a non-zero chance of a complication like bowel perforation. It requires a procedure room and special equipment, which may not be available in the numbers required to serve large populations across a nation.

When the technology is available, some people still say, “No, thank you.” In 2000, only 29 percent of eligible patients screened for colon cancer by colonoscopy. Women are also opting out of mammograms. A 2024 study out of the University of California, Davis—the biggest of its kind, tracking 3.5 million mammograms across 1 million women—looked at what women did after getting a false positive test. A woman who got called back for a scare that turned out to be nothing was less likely to come back for her next mammogram. The more the system put her through—the repeat imaging, the anxious six-month wait to “keep an eye on it”—the less likely she was to return.

This need for individual attention and the search for certainty that is lacking in population-level screening has opened the door for another type of screening, designed and marketed under a different moniker.

For a few thousand dollars, otherwise healthy-feeling people can have a full-body MRI scan, nouveau bloodwork, or other tests to look for disease lurking in their bodies. The promise—a seductive one—is that of certainty, something that your regular doctor cannot give you. It doesn't *feel* like the screening we know, and it's typically not sold that way; it's marketed as wellness, as a proactive choice, something fundamental to do for your health and future.

But the challenges of screening can't be avoided: There will be overdiagnosis, blind spots leading to a false sense of security, additional cost and pain. People will go get additional tests and surgeries, only to find nothing. There will be sleepless nights.

One healthy 35-year-old man from New York experienced this when in 2023 he paid \$2,500 for a full-body MRI endorsed by Kim Kardashian as “life saving.” The test was negative: all clear. Eight months later he collapsed due to a catastrophic stroke, the consequence of a blocked artery, allegedly visible in that full body MRI. Now partially paralyzed, he has filed a lawsuit against Prenuvo, the whole-body scan provider, claiming the blockage should not have been missed. (The radiologist working with Prenuvo, who reviewed the scan, previously had his medical license suspended in connection with an auto insurance scheme in which he was accused of falsifying findings on MRI scans, reported *Ars Technica*.) Despite what the suit claims, it's also possible the arterial stenosis was not visible on MRI, as the protocol for these screening exams is not the same one we use to evaluate arteries. It's in details like this where direct-to-consumer use of medical technology can fall short. The screening tests we offer, while far from perfect, are carefully thought out by experts with the best intentions. In contrast, the full-body MRI you can pay for is untamed, without any focus or rigor needed to go from a diagnosis to cure.

But screening can work, if it's based on the evidence and careful thought and planning. When programs are designed around how people behave, for example, the results can be transformational.

Our colleague, T.R. Levin, a gastroenterologist, began mailing FIT kits to screen eligible people in our Northern California healthcare system for colon cancer. The kits contain a small vial into which patients can insert a small sample of their stool and mail it back to a lab for analysis. Levin told us screening rates in the

healthcare system climbed from 39 percent in 2000 to nearly 83 percent in 2015. (The *Journal of American Medical Association* has also reported that FIT kits increase individual testing and lower the risk of dying from colon cancer.) A similar approach, home or self-collection, has been adopted and is spreading to screen women for cervical cancer.

So, how do we navigate this sea of population health, statistics, and confusing testing?

As doctors, we need to be more thoughtful and invest more of ourselves in how we approach and explain screening. We cannot screen simply to meet a metric without engaging patients with better conversations and understanding. We may not be able to give patients certainty, but we can certainly try to give them clarity.

As medical professionals, as a society, we don't know all the answers. But we know one thing: A lack of clarity, a lack of certainty, can make things seem stupid. The system is not stupid. The science is not stupid. The real problem is expecting certainty when certainty was never on the menu to begin with.

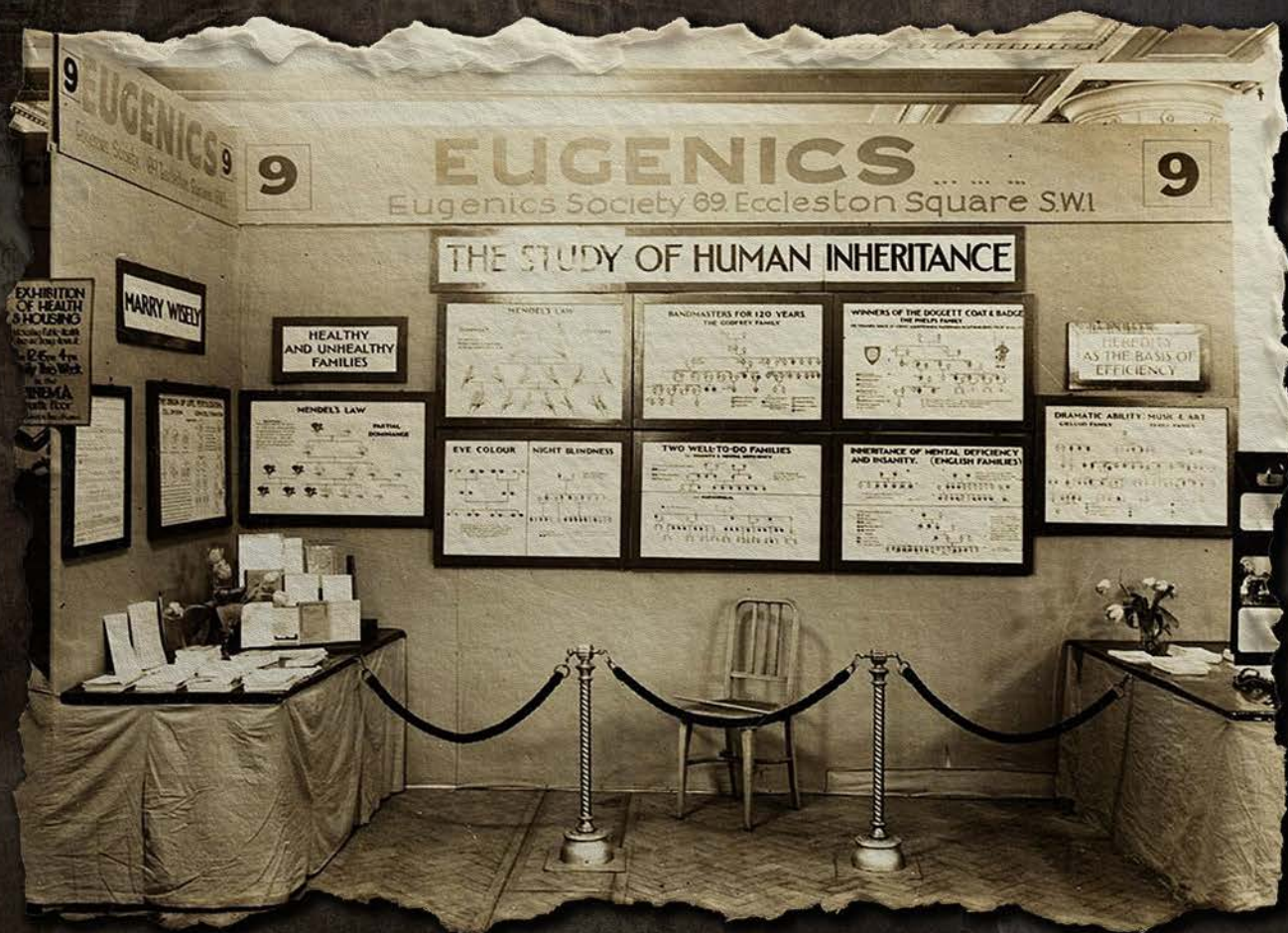
Because when we can't accept that, or don't understand it, then we go looking for whoever promises a clean answer. And there is always someone willing to make that promise.

Instead, doctors can offer clarity—of what's happening, what the options are, what the odds are. A positive screening test does not necessarily mean you have cancer. So, let's talk about what to do if your next test is positive. Consenting to screening must come from thoughtful conversation, not just an algorithm. And patients, in turn, must be willing to meet doctors there, accepting that clarity, not certainty, is what medicine can offer.

After further thought and conversation with his family, Rahul recognized he owed it to them to figure this out. A urology consult helped him define his next steps. An ultrasound was normal; more detailed blood testing awaits. An MRI remains a possibility, as does a biopsy, should he and his doctor agree it's needed.

And yes, that clarity feels a lot less stupid. ☺

RAHUL PARIKH and **JONATHAN MATES** are physicians and co-creators and hosts of the podcast *Healthcare Is Stupid*. Parikh is a pediatrician and writer whose work has been featured in *Nautilus* and other publications. Mates is a radiologist and writes the newsletter, *Red Goggles*.



THE RETURN OF A DISCREDITED PSEUDOSCIENCE

*Eugenics is rearing its ugly head again in today's
crucible of politics and culture*

BY DAN FALK

THE PSEUDOSCIENCE OF EUGENICS, founded on the mistaken notion that nature envisions some “right” combination of genes, and that selective breeding can improve the human race, is making a comeback. While the heyday of eugenics was a century ago, the idea was never really forgotten.

In the 19th century, the modern science of biology was in its infancy, but took a leap forward with the work of pioneering researchers like Charles Darwin and Alfred Russel Wallace. But other thinkers, like Darwin's cousin, Francis Galton, took things further by supposing that, just as farmers breed animals to improve or weed out various traits, humans could be guided to make “better” offspring.

When the Nazis came to power in the 1930s, they took eugenicist ideas and ran with them, justifying their hatred of certain groups—Jews, Roma, gay people, the disabled—as supposedly backed by science, anchored to the notion that some races were naturally superior to others. In this distorted way of thinking, violence against the unfit was justified, even noble.

While eugenics largely fell into disrepute in the wake of Nazism, it never disappeared. Nor was it something the Nazis invented: In the early 20th century, most of the literature on eugenics was published in the United States, where it eventually underpinned sterilization campaigns against people with disabilities, fueled racist immigration policies, and, more recently, inspired efforts like the “Repository for Germinal Choice,” which collected and preserved the genetic material of leading thinkers, including Nobel laureates, by storing their sperm.

Most serious scholars have relegated eugenics to the trash heap of science history. But much like its original incarnation from the late 19th century through the aftermath of World War II, its current rise is tangled up in politics, ideology, and an insidious desire to divide

Despite science showing eugenics has no basis in fact, it continues to surface in Washington today.

rather than unite people. And, as in the old days, on the unfounded notion that biology is the be-all and end-all of a person's character and worth.

This past March, President Donald Trump told Fox News what he thought of two teenagers from Pennsylvania who threw a handmade bomb at anti-Muslim protestors outside the mayor's mansion in New York City.

According to federal prosecutors, the young men identified with the extremist Islamic State group. The bomb failed to detonate.

“They're sick people, and a lot of them were let in here,” Trump told host Brian Kilmeade. “Others are just bad. They go bad. Something wrong—there's something wrong there. The genetics are not exactly, they're not exactly your genetics—it's one of those problems, Brian.”

Disjointed as the comments were, they follow a theme that has long obsessed Trump. As far back as the 1980s, in an appearance on *The Oprah Winfrey Show*, Trump spoke of the importance of having “the right genes.” On the campaign trail in 2024, he declared immigrants “were poisoning the blood of our country.”

Trump's penchant for classifying—and evaluating—people based on biology is echoed by Robert F. Kennedy, Jr., his secretary of health and human services. Ignoring established medical opinion, Kennedy has claimed it's “very difficult for measles to kill a healthy person.” Seemingly judging some people to be biologically superior to others, Kennedy added, without evidence, “we see a correlation between people who get hurt by measles and people who don't have good



PICK YOUR BABY This ad in a New York subway station is from a company promising prospective parents it can identify traits like IQ in embryos. Scientists say the company is eugenics in new clothes.

nutrition or who don't have a good exercise regimen." His disparaging comments about people with autism have been equally unfounded. And then there's the repeated use of the "r word" by Elon Musk to describe people with mental disabilities.

These sentiments, coming from the most powerful people in Washington, D.C., and in the tech world are echoed on social media and in the wider world. They stoke xenophobic conspiracies like the "great replacement theory," which holds that whites are being driven to extinction by non-white immigrants, and that breeding programs and restrictions on immigration are needed to ensure that whites retain a dominant position.

This discourse has, for many, a starkly familiar flavor, echoing some of the statements made by eugenicists a century ago. Despite a half century of science

showing that the ideas behind eugenics have no basis in fact, it continues to surface in Washington and across political lines. As Rebecca Sear, an evolutionary anthropologist at Brunel University in the United Kingdom, who has published extensively on eugenics, puts it: "In the 21st century, eugenics is experiencing a fully-fledged revival."

GALTON COINED THE WORD "eugenics" in 1883 (the term comes from the Greek word *eugenes*, roughly translated as "well-born"). The link to Darwin's own work on evolution is indirect: Darwin had merely explained how natural selection shapes species over time. Galton went a step further, arguing that we ought to collectively take charge of the selection process, steering the development of the human species. He imagined a kind of hierarchy of racial and cultural groups, with white

COURTESY OF NUCLEUS

Europeans at the top of the heap; his 1869 book *Hereditary Genius* included a chapter titled “The Comparative Worth of Different Races.” And if it worked for farmers, why not for us? “What nature does blindly, slowly, and ruthlessly,” he wrote, “man may do providently, quickly, and kindly.” It all came down to who had what sort of babies.

In the early years of the 20th century, eugenics garnered popular support from across the political spectrum in the United States and beyond. It was a time of political and social change, with reformers advocating limits on the power of corporations, maximum work hours, and minimum wages, along with the elimination of child labor. Yet many of the same people who fought for workers’ rights, as well as the rights of women and children, also believed that the human race could be improved through various kinds of intervention.

State and national governments called for restrictive marriage laws, for the segregation of the mentally disabled, and even for forced sterilization. “Better baby” competitions, modeled on livestock competitions, became commonplace at some U.S. state fairs. The movement swept much of the globe. Crucially, eugenics was billed as being backed by *science*. You could read about eugenics in standard biology textbooks and respected research journals. The *New England Journal of Medicine*, for example, wrote approvingly of its aims.

Eugenics and racism often went hand in hand. The U.S. Immigration Act of 1924 was aimed at improving the genetic make-up of the American population by restricting immigration from groups presumed to be genetically inferior—Italians, Eastern Europeans, and Jews.

In Canada, where more farm workers were needed in the underpopulated Prairie Provinces, the same fears about the “wrong” sort of immigrants took root; Charles Kirk Clarke, a University of Toronto psychiatrist who had the ear of Canadian politicians in the 1910s and ’20s, stated: “It is all very well to talk about pumping in the population. But surely the streams tapped should not be those reeking with degeneracy, crime, and insanity.” During World War II, Nazi Germany took eugenic ideology to its genocidal endpoint.

And yet, eugenic thinking was not anchored solely in Anglo-Saxon supremacy. There were Black and Jewish eugenicists, who imagined either that the “science” behind eugenics would lift up all races, or at least

improve the quality of people within specific races. And then there was the suffragette movement. Some first-wave feminists supported eugenics, linking the fight for female emancipation with the desire to better the human race.

“This was the progressive era,” says Wendy Kline, a historian of medicine at Purdue University in Indiana. “Eugenics seemed to mesh with all the interests of what people saw as progressive science.”

A theme that runs through the original eugenics movement is the idea that whatever makes some people “better” than others is rooted in biology, and therefore not only heritable but also potentially controllable.

This was certainly imagined to be the case with intelligence. In the early 20th century, many Americans were gripped by the fear that less intelligent people were having more babies than more intelligent people. (A sentiment echoed today by Musk, who has amplified posts on X about “high status males” and has volunteered to use his own sperm to seed a future Martian colony.) Those fears were front and center in the landmark *Buck v. Bell* Supreme Court case of 1927. The case centered on Carrie Buck, a poor white woman who, at 17, had become pregnant after being raped by her foster mother’s nephew. Her foster parents had her institutionalized, deeming her “feeble minded” (the label often used at the time) and called for her to be forcibly sterilized. The question of whether this was permitted under Virginia law made its way to the Supreme Court, presided over by Chief Justice (and former president) William Howard Taft.

In the end, the court upheld the Virginia statute that allowed such actions “for the protection and health of the state.” Buck’s biological mother and her 8-month-old sister faced the same fate. Justice Oliver Wendell Holmes Jr.—who championed civil liberties—wrote the majority decision: “Carrie Buck is a feeble minded white woman [and] is the daughter of a feeble minded mother in the same institution, and the mother of an illegitimate feeble minded child ... Three generations of imbeciles are enough.” (Louis Brandeis, the court’s first Jewish member, voted with the majority, as did Taft; only Associate Justice Pierce Butler dissented.)

Historians estimate that between 60,000 to 70,000 Americans were eventually sterilized against their will, in at least 30 states. Many of the laws allowing such

The idea of relying on genetics to decide who gets to be born and who doesn't is inherently fraught.

actions remained on the books for decades, and as recently as 2013, incarcerated women were routinely sterilized in California state prisons.

After its association with Nazi ideology following World War II, some eugenics organizations simply changed their names. The journal *Eugenics Quarterly*, for example, changed its name to *Social Biology*, and the *Eugenics Review* became the *Journal of Biosocial Science*. In fact, as Sear points out, eugenics continued to thrive in academia well into the post-war period, citing the journal *Mankind Quarterly*, founded in 1960, as an example.

"*Mankind Quarterly* was established as something which looked like an academic journal," Sear says. "But in fact it was set up entirely to promote eugenic ideology, particularly around racism." Today the journal is published by the far-right Human Diversity Foundation.

IT WOULD BE EXHAUSTING to enumerate all the ways that the thinking behind eugenics fails to hold up to scientific scrutiny. Take, for example, the movement's obsession with IQ. More than a half century ago, evolutionary biologist Stephen Jay Gould wrote, "Who knows what IQ measures? It is a good predictor of 'success' in schools, but is such success the result of intelligence, apple polishing, or the assimilation of values that the leaders of society prefer?" As Sear points out,

IQ only measures a narrow range of what we typically mean by "intelligence." For instance, she says, "IQ does not measure your social intelligence, your ability to get on with other people."

Like most psychological traits, intelligence is the result of multiple genes, which can be modulated by the influence of yet other genes, and of course by environmental factors. Even when genes are linked to specific traits, life outcomes can vary enormously—due to opportunities (or lack thereof), such as inherited wealth, prejudice, and sheer luck.

A critical tenet of eugenics is that races are delineated by biology, a convenient belief allowing for people to be categorized and even ranked based on their genetic inheritance. But that tenet has no backing from modern science. When genuine science doesn't align with one's beliefs, pseudoscience can flourish. All of this gets even muddier when ideology and politics enter the picture—as current happenings in Washington again seem to highlight.

In 2025, Trump objected to an exhibition called "The Shape of Power: Stories of Race and American Sculpture" at the Smithsonian's American Art Museum. He issued an executive order criticizing the exhibition for promoting "the view that race is not a biological reality but a social construct." In response, a quartet of geneticists and anthropologists wrote in *Science* that

decades of studies in evolutionary biology, genetics, and anthropology have shown race is indeed a social construct, not a biological reality.

“No biological definition of race can be accurately applied to subgroups of humans, especially not subgroups based on social definitions of race, which arbitrarily use physical, geographic, cultural, and linguistic criteria,” they wrote. This consensus, the scientists point out, is based on evolutionary evidence and “data showing that the amount of genetic variation that occurs within human populations is much greater than genetic variation between populations.”

The authors cited a 2023 report by a joint committee of scientists from the National Academies of Science, Engineering, and Medicine, which stated that race “is a misleading and harmful surrogate for population genetic differences.” Evoking the specter of eugenics, the report lamented that race “has a long history of being incorrectly identified as the major genetic reason” for differences between groups. (This summer, Trump ordered new “temporary exhibits or signage” to be installed outside yet another Smithsonian museum—this time the National Museum of American History—warning visitors about “the Ideological Capture” of the museum.)

The historical eugenics practice of invoking science (that is, what was imagined to be valid science) in support of policies that marginalized the poor and benefited the privileged, all in the name of supposedly improving the human race, has a new sheen today, enabled by the latest reproductive technology.

Last fall, banner ads sprouted up in the New York City subway system, saying HAVE YOUR BEST BABY and IQ IS 50% GENETIC. The ads were for a company called Nucleus Genomics. By using an embryo-testing technique known as polygenic screening, Nucleus suggests couples can choose their baby’s eye color, hair color, and even intelligence. Some 3,000 people have joined the company’s wait list as of April. Even if such baby-tweaking is simply the result of people freely expressing their preferences, it may be an option primarily available to the wealthy—and might thus inject a new dimension of disparity to the already cavernous rich-poor divide.

At the same time, advances in genetic engineering—and especially the allure of gene editing technology

like CRISPR—has left people wondering just how far the idea of “designer babies” might go (although gene editing of human embryos for reproductive purposes, known as “germline editing,” is currently banned in the U.S.).

“Do we want humans to be a commodity? I think that train may have left the station,” says Philippa Levine, a historian and professor emeritus at the University of Texas at Austin and author of the 2017 book, *Eugenics: A Very Short Introduction*. After all, Levine adds, “What the hell is perfection? If I’m a deaf person, does that make me a lesser person? If I’m in a wheelchair, does that make me a lesser person?”

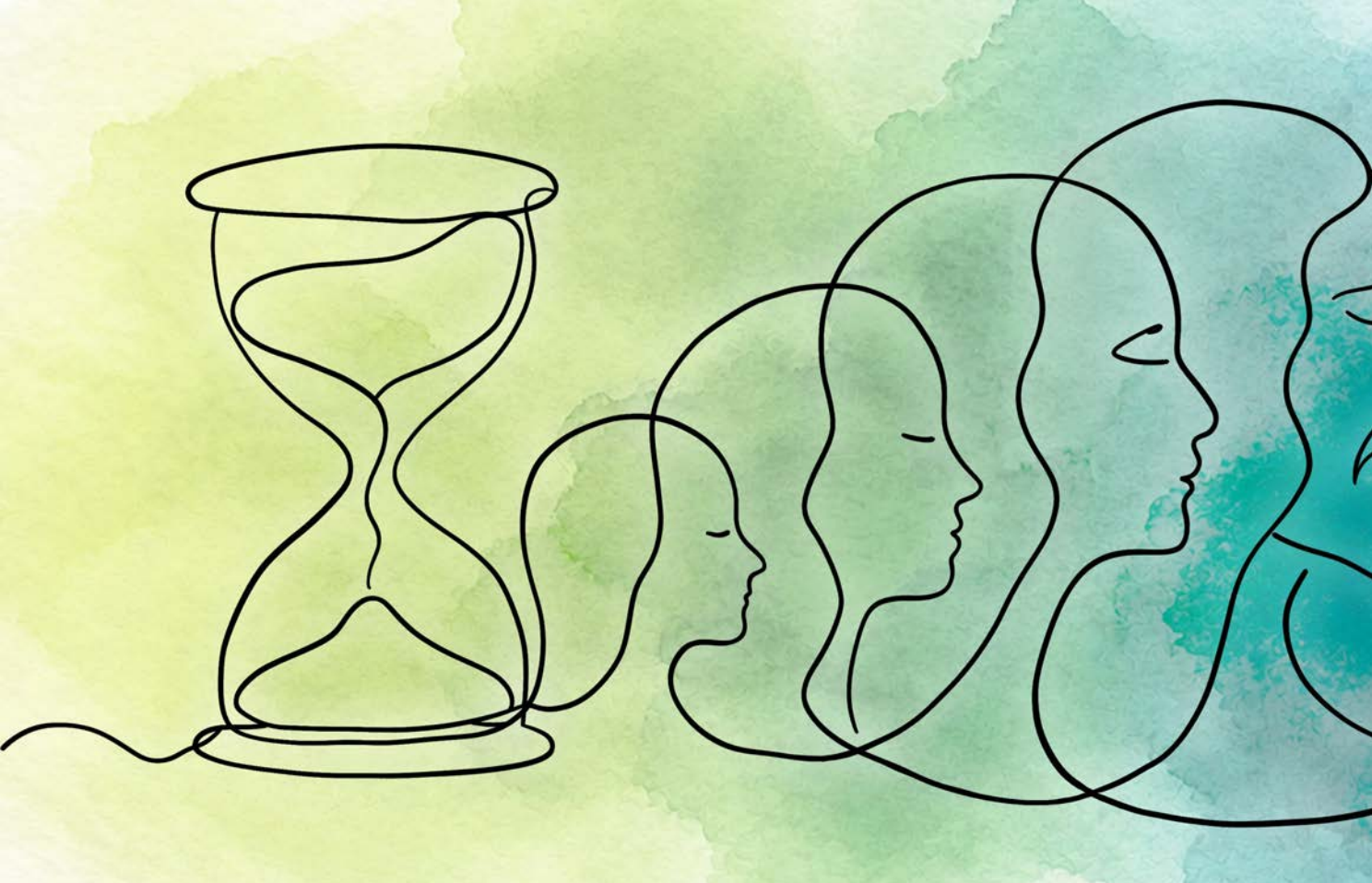
Meanwhile, people are also choosing what kind of babies *not* to have. In the United States, about two-thirds of women are choosing to abort fetuses that carry the genetic abnormality linked to Down syndrome. But in some countries, like Iceland, the figure is much higher.

Voluntary abortion and voluntary embryo-selection are, obviously, worlds away from forced sterilization. But the idea of relying on genetics to decide who gets to be born and who doesn’t is inherently fraught.

Alison Bashford, a historian at the University of New South Wales in Australia, and co-editor with Levine of *The Oxford Handbook of the History of Eugenics*, has argued that many of the trends seen in today’s reproductive medical interventions recall our century-plus urge to classify people as “fit” or “unfit,” virtually from the moment egg and sperm first meet. Whatever we call it, Bashford says, the eugenicists of a century ago “would undoubtedly be delighted at contemporary possibilities for reproductive selections.”

That eugenics-themed language and ideas have re-entered the public sphere, from the political realm to the business world, deeply concerns scholars like Bashford and Levine. Its resurgence, Levine says, shows that eugenics never died. “The essence of eugenics is still with us. It crept away for a little bit. But it continues to exist, and the manipulation of science continues to be something that is very dangerous.” 🌀

DAN FALK is a science journalist based in Toronto. His books include *The Science of Shakespeare* and *In Search of Time*.



It's the Evolution, Stupid

*Immortality is possible?
Nature has 3.5 billion years of experience saying otherwise.*

BY ELENA KAZAMIA



T HERE ARE ALL THESE PEOPLE who say that death is natural, it's just part of life, and I think that nothing can be further from the truth," Peter Thiel told *Business Insider* in 2012. The venture capitalist has spent years backing efforts to slow or defeat aging, and his provocation captures the confidence of a wider Silicon Valley project: Given enough money and biological engineering, death may become just another technical problem to solve.

Nature, it seems, remains wholly uninspired by Silicon Valley capital. Scientific insight into 3.5 billion years of evolution suggests that death isn't a glitch—it's working as intended. That does not mean every organism carries a simple genetic suicide switch, or that every death is adaptive. It simply means that natural selection favors reproductive success over immortality. It also explains aging as the ultimate compromise for a body balancing the competing energetic requirements of development, reproduction, and repair.

Nature, it seems, remains wholly uninspired by Silicon Valley capital.

“There is a cost to living, a cost to doing everything” Nick Lane, a biochemist at University College London, told Philip Ball in an interview with *Nautilus*. Lane examines life through the lens of energetics. He studies how much energy is required for life’s various processes, and how organisms go about acquiring and spending it.

“If we are living our lives at a very fast rate, we tend to, in effect, wear out sooner,” Lane said. Living is strenuous work. Growth, mating, pregnancy, parental care, and immune defense all consume energy and take a physical toll. Because maintenance and repair are themselves energetically costly, they must be traded off against these processes.

Across the animal kingdom, species have evolved different ways of allocating energy throughout their lives, mostly according to their size and exposure to external hazards. One consequence of this is that animals have vastly different natural lifespans.

Consider the mouse and the elephant. In the wild, a field mouse can live up to two years, whereas an elephant can live to 70. But mice are small and die easily, from predators or harsh weather. In such easy prey, natural selection is highly skewed toward early reproduction. An expensive body maintenance program for a distant (and highly unlikely) old age makes no sense. For an elephant, which matures slowly and faces fewer predators, evolution has selected for a durable body and a longer lifespan.

Overall, smaller species live fast lives, reproduce early, and die young. “There is a very strong relationship between metabolic rate, the rate at which we’re taking in oxygen and burning up food, and lifespan,” Lane said. This may help explain, at least partially, why caloric restriction can double or treble the lifespan of very simple organisms, like worms or fruit flies. Without calories, reproduction

is unlikely and organisms enter a phase of “battening down the hatches and waiting out the bad times,” Lane said. But humans, of course, are far more complex. Our bodies are finely tuned to run at a metabolic rate held within narrow limits. Simply restricting calories in our diet would lead our bodies to burn up protein in muscle and other tissues.

In fact, humans are unusually long-lived for primates of our size. “Given our body mass, we live about twice as long [as we should for our size],” Venki Ramakrishnan, the Nobel Prize-winning structural biologist and author of *Why We Die*, told me. Improvements in nutrition, sanitation, and medicine have doubled the average human lifespan over the past 150 years, but our longevity has deeper roots. Michelangelo, Ramakrishnan noted, lived to nearly 90 in the 16th century, and some researchers believe human lifespans began lengthening as far back as 40,000 years ago. The drivers for the change remain unclear, Ramakrishnan said. A leading hypothesis links a sudden shift in human lifespan to the evolutionary perks of having grandparents in hunter-gatherer communities. Older adults helped raise children while parents were busy providing food. In times of crises, their accumulated knowledge of the environment gave their tribe a decisive edge in survival.

Given that modern humans have escaped many of the hazards that once made longevity a poor evolutionary investment, and that leaps in lifespan are possible, it makes sense to wonder if perhaps the billionaires have a case. Or, as Ramakrishnan put it, “The real question is, we have mechanisms for repairing damage, so why don’t we just keep repairing our body and live forever?”

The answer cuts to the heart of aging research.

Even in the absence of multiple stressors—we live comfortable lives indoors sheltered from environmental extremes and carry a significantly lower

“You can think of senescence as evolution’s way of preventing cancer.”

pathogen load than our ancestors—our bodies have evolved to keep the time. In 2022, Steve Horvath, a professor at the University of California, Los Angeles and a leading figure in anti-aging research, uncovered the existence of a molecular clock, a universal, internal mechanism that marks the age of mammalian cells and tissues. It reveals that aging is neither entirely random nor simply the accumulation of visible wear.

Throughout a cell’s life, patterns of chemical tags on DNA (called epigenetic markers) accumulate. Horvath developed a formula that reliably deciphers the pattern to predict the age of cells and tissues. This pattern is the “biological” age of the cell, which closely aligns with, but doesn’t always precisely map, chronological age. Its accuracy (for humans, it is exact to plus or minus 3.6 years) implies that aging is a carefully maintained biological process, rather than a haphazard set of events. We speed up the clock with things like smoking and excessive alcohol consumption.

Some scientists interpreted Horvath’s discovery—that all tissues count time—as indication that we are programmed to die, our cells marching inexorably toward death. But Horvath’s own interpretation swerves clear of such determinism and focuses instead on what the clock achieves in our early years.

We don’t think of babies as aging, but that is when the aging program is most active. Embryonic

cells (which start off identical) need to differentiate quickly and in perfect synchrony into tissues, building organs and a body that continues to mature into adulthood. An internal metronome serves the essential function of orchestrating the elaborate process. But after adulthood, the program continues to run, and runs us all the way to the ground. “Evolution did not select a program that makes us die,” Horvath said. Instead, “Mother Nature never selected against it.”

As Ramakrishnan told me, “Death is an unintended consequence of aging.” And aging is an unintended consequence of development. Research shows that once development is complete (once we reach reproductive age), some of the genetic machinery that enabled it can keep running, or become reactivated as its regulation deteriorates, with increasingly harmful effects. In old mice, for example, the gene *Hoxa9*, which is essential to a developing embryo, switched back on in muscle stem cells, activating obsolete developmental pathways that hindered muscle repair.

One of the hallmarks of aging is cellular senescence. It is a state in which damaged or stressed cells permanently stop dividing yet resist dying, accumulating in tissues over time. Though senescence initially protects against cancer by halting the growth of compromised cells, senescent cells secrete inflammatory molecules that harm surrounding tissue and drive many features of aging.

Declaring death unnatural mistakes defiance for evidence.

The runaway developmental programs of old age stem from the same evolutionary trade-off that links metabolism to lifespans. Selection prioritizes early biological performance (the period before reproduction happens) over late-life durability. A gene that promotes development or improves fertility and survival at 20 can be favored even if it raises the risk of degeneration at 70. A defect expressed only at 90 will encounter limited evolutionary pressure. Death is a compounding of the problem across a genome composed of genes all acting under the same selection. It is the final bill for evolution's early favors.

ALL OF THIS MAKES anti-aging innovations extremely challenging. Roughly 2,000 master control genes orchestrate early embryonic development, supported by another 5,000 to 10,000 organ-building genes. Re-tuning them for longevity without wreaking havoc is well beyond our current capabilities. "I don't think it's an impossibility, but it's incredibly difficult," Ramakrishnan told me.

The most immediate threat of getting this balance wrong is cancer. That's because the signaling pathways that encourage cells to divide and repair tissues in early life are the exact same pathways that tumors hijack to proliferate unchecked.

Override those aging pathways, and you remove the guards keeping malignancy at bay.

Researchers have already encountered this trade-off in one of the best-studied longevity pathways, which involves growth hormone (GH) and insulin-like growth factor 1 (IGF-1). Because IGF-1 promotes cell growth, muscle maintenance, and tissue repair, boosting its activity once seemed a plausible anti-aging strategy.

But the same signals that encourage healthy cells to regenerate can also encourage precancerous cells to survive and multiply. Animal studies consistently show that reducing GH and IGF-1 signaling extends lifespan while lowering cancer incidence, whereas chronically elevated IGF-1 activity is associated with an increased risk of several cancers. The pathway that helps build and maintain the body in youth can become a liability when kept active into old age.

This compromise should not be entirely surprising. Adult stem cells possess an unusual capacity for long-term self-renewal, while cancer cells achieve essentially limitless proliferation by escaping the body's normal controls. Efforts to make ordinary cells more regenerative inevitably edge them toward the biology of tumors. "An organism doesn't want cells capable of indefinite

reproduction,” Ramakrishnan said, “because that’s a cancer risk. You can think of senescence as evolution’s way of preventing cancer.”

The pattern is becoming uncannily familiar throughout aging biology: Molecules and pathways once cast as villains often turn out to have essential functions. For decades, scientists believed that aging was largely driven by the accumulation of damage from reactive oxygen molecules, or free radicals, produced by mitochondria. The theory was so influential that antioxidant supplements were widely promoted to slow aging. But decades of research have failed to support that idea.

Large clinical studies found little evidence that high-dose antioxidant supplements extend life or prevent age-related disease, and some have even linked them to worse outcomes. Rather than being merely destructive, Lane explained, free radicals act as vital cellular signals. “They are behaving a little bit like the smoke detector,” he said. When cells detect those signals, they activate genetic programs that repair damage, strengthen stress resistance, and help the cell survive. Blunting that signal indiscriminately is counterproductive. “The trouble with antioxidants is they’re in effect disabling the smoke detector,” Lane said.

A similar paradox occurs with anti-aging therapies that target the inflammatory response. Inflammation is essential for healing and immunity, but with age, damaged and senescent cells can keep the immune system locked in a state of chronic low-grade activation. Persistent inflammation, which longevity scientists have dubbed “inflammaging” has been linked to increased risk of diabetes, Alzheimer’s, and cardiovascular diseases.

In 2017, a Novartis trial of more than 10,000 heart attack survivors showed that canakinumab, an antibody that blocks a single inflammatory molecule, successfully cut the risk of further cardiovascular events, all without touching cholesterol. At first this was touted as progress in the fight against inflammaging. But ultimately the Food and Drug Administration declined to approve the use of canakinumab in any setting. The agency deemed that the risk of developing severe infections, which had been observed in some of the trial participants,

was too high. The trial served to add further proof that the inflammatory response has evolved for good reason.

Eric Topol’s 2025 book, *Super Agers*, offers a discerning look at longevity medicine. The cardiologist and scientist argues that rather than chase immortality, our focus should be to extend “healthspan” by preventing or delaying the major diseases of aging, such as cancer and Alzheimer’s. Better risk prediction, vaccines, exercise, sleep, nutrition, and earlier treatment could all reduce the years people spend with disease. This would make for a scientifically serious agenda, precisely because it does not require fighting millions of years of evolution. The most plausible successes would compress morbidity, with more years of strength and independence. Even modest gains would matter enormously.

Ramakrishnan likewise sees life extension beyond the apparent human limit as an extreme and selfish undertaking by the world’s richest technocrats. “It’s just fantasies at this point, and a distraction from our real problems,” he told me. “It’s the same group of people who are touting space travel and colonization as some big hope for humanity, while ignoring all the real problems on Earth, like climate change, disease, and having enough water to drink.”

Thiel was right in one narrow sense: We should not mistake the natural for the inevitable. Nature gives us challenges and it is part of being human to strive against them. But declaring death unnatural does not make it so; it mistakes defiance for evidence. Immortality evangelists believe enough intelligence and capital can free us from the evolutionary constraints that have governed every biological lineage on Earth. Their resources may help science understand aging and extend healthy lives, but evolution has already delivered a verdict on the larger boast: Life persists. Individual bodies do not. ☺

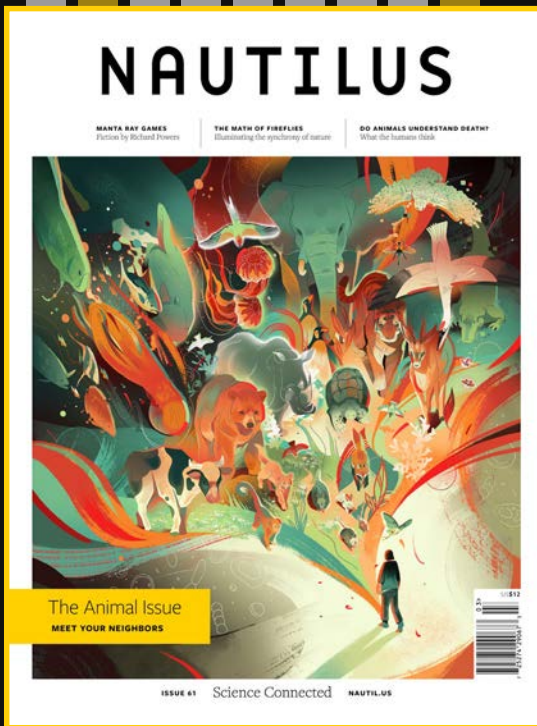
ELENA KAZAMIA is a science writer from Greece. She has a master’s degree in conservation from University College London and a Ph.D. in plant sciences from the University of Cambridge in the United Kingdom.

01
02
03



27
26
25

24





ISSUE 68: The Stupid Issue

Read *Nautilus*,
an award-winning
science magazine.



Read us at nautil.us
Subscribe at nautil.us/products

MARTIN SCHWARTZ

The Yale School of Medicine professor on the importance of stupidity in science

INTERVIEW BY BOB GRANT

In a popular 2008 essay, Martin Schwartz wrote, “I actively seek out new opportunities to feel stupid. I wouldn’t know what to do without that feeling ... The more comfortable we become with being stupid, the deeper we will wade into the unknown and the more likely we are to make big discoveries.” We asked Schwartz today about the scientist’s life.

A LIFE OF QUESTIONS

Scientists are an elite group of learners, which is a whole different thing from an elite group of knowers. We spend most of our lives asking questions we don’t have the answer to, and we do the best we can. Which, you know, beats the hell out of pretending that you know things that you don’t.

Asking a question that has never been asked before is a bit audacious, and it puts you in this place of not knowing. Throughout our lives, and in every realm of human endeavor, there’s not knowing. Engineers, doctors, artists, and students don’t know the answer to things. But in science, it’s the worst.

There are very prominent people who asked a question at the beginning of their career and made enormous discoveries, and never really answered the question by the end of their career. Einstein had this glimpse that

there could be a unified field theory. He unified time and space and never got any further. He spent decades grappling with something, and he never made much progress. Eric Kandel had this vision that you could understand learning and memory by studying simple animals.

And over a 40-year

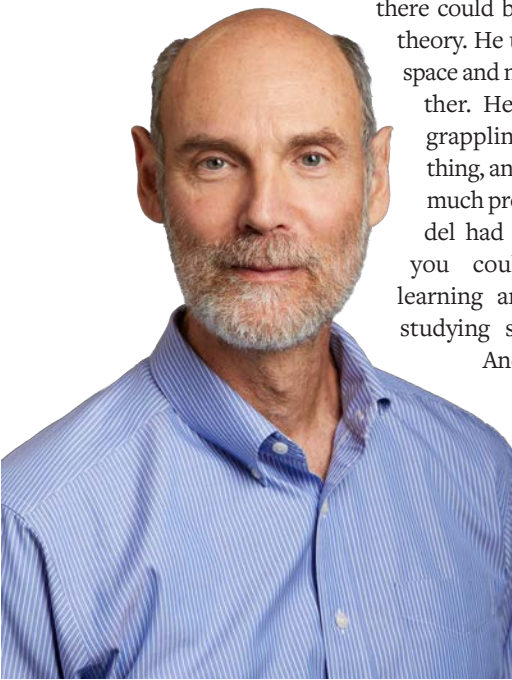
career, he brought learning and memory down to signaling pathways in neurons, but it just scratched the surface.

FOR CRAZY PEOPLE

It’s a difficult thing to live your whole life in a state of uncertainty. There’s no guarantee that you’re ever going to get an answer. That’s kind of rough, and it takes a special person who’s willing to do that. I think this principle of stupidity applies to all creative endeavors, whether it’s business, engineering, visual arts, music, or science. But in science, where you can spend 40 years making a bit of progress, never knowing if you’re ever going to make more progress, and you get to the end of a 40-year career, and you’ve just scratched the surface. What kind of crazy people want to do that?

CHANGING THE WORLD

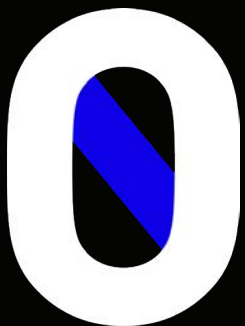
Science and politics are uncomfortable bedfellows. In the realm of public policy, people want answers. Scientists spend most of our lives grappling with questions. There is something unavoidable about that. There is a bit of a mystery here. The mystery is that for all its flaws and uncertainties and hesitations and stumbles, science is the most powerful engine in the universe for acquiring power. It is a self-correcting process that bounces around between different models and different views, all of which are flawed. And by bouncing around you discover electricity and engines and nuclear fission. That strange, uncertain, feeble process has changed the world. 🌀



ONE5C

You are going to save the world.

one5c is a newsletter that makes it easier for busy people to live more sustainably, without needing to become experts, activists, or overhaul their entire lives.



one5c.com



[Subscribe for free](#)



IF IT'S NOT ITALIAN,
IS IT EVEN ESPRESSO?



DELONGHI.COM