

Parkfield Quakes Skip a Beat

Seismologists' first official earthquake forecast has failed, ushering in an era of heightened uncertainty and more modest ambitions

As recently as the mid-1980s, predicting earthquakes years in the future was a promising concept. The basic idea was as simple as the ticking of a clock. At least some faults seemed to accumulate and release stress over a regular cycle; if you knew how often quakes had ruptured a given fault in the past, you could calculate—at least roughly—when the next quake was due. Seismologists agreed that one of the best places to test the idea was a 25-kilometer section of the San Andreas fault that runs by the tiny village of Parkfield in central California. The fault segment had broken in moderate, magnitude 6 quakes every 22 years or so since 1857. Allow 5 years of leeway to account for the inevitable irregularities of natural systems, the prediction went, and the next Parkfield earthquake should strike around 1988—or almost surely by the end of 1992.

That forecast was offered in 1984 by three California seismologists—Allan Lindh and William Bakun of the U.S. Geological Survey (USGS) in Menlo Park and Thomas McEvilly of the University of California, Berkeley. Not long afterward, it was endorsed by the National Earthquake Prediction Evaluation Council (NEPEC), and the San An-

dreas near Parkfield became the most closely watched fault segment in the world. Besides offering a test of seismologists' understanding of the earthquake cycle, it seemed an ideal place to watch for signs that an earthquake was imminent—clues that might one day make it possible to issue short-term earthquake warnings (see box).

That, at any rate, was the theory. Reality proved different. On 1 January, the prediction window closed—without the much anticipated temblor. And that nonevent has sent a tremor through the seismological community. Some suspect that the style of prediction epitomized by Parkfield might be resuscitated by adding in additional factors that could account for departures from a regular cycle. Lindh, for example, argues that the prediction's failure "says the game is harder than we hoped. It doesn't mean it upsets the whole apple cart."

Lindh's USGS colleague Paul Reasenberg disagrees; he thinks the idea underlying earthquake forecasting—simple repetitive cycles—is wrong. "The closing of the Parkfield window has to be a significant and symbolic event in a changing paradigm.... It was going to be like clockwork, but that is all falling apart."

Reasenberg and others contend that the earth's crust is such a remarkably complex system that chaos overwhelms predictability.

No fault is an island. The failure of the Parkfield prediction may have underscored doubts about earthquake forecasting, but they had started to set in soon after the forecast was issued. A fundamental assumption of the prediction—that a fault segment can be considered as an isolated system—quickly started to look shaky as seismologists studied the aftermath of the magnitude 6.5 Coalinga earthquake of 1983, which hit 30 kilometers off the San Andreas to the northeast of Parkfield. Coalinga released enough stress to reverse the slow creep of the San Andreas north of Parkfield for a year and a half. Within a few years of Coalinga, Robert Simpson of the USGS in Menlo Park and Terry Tullis of Brown University had independently calculated that that stress release probably also relieved some of the tension in the "spring" of the Parkfield clock, delaying the next quake by as much as several years.

Coalinga might have been written off as a rare stroke of bad luck, but Harvard University earthquake modelers Yehuda Ben-Zion, James Rice, and Renata Dmowska think an-

Short-Term Prediction Takes Its Knocks, Too

Researchers hoping to test their skill at long-term earthquake forecasting aren't the only group disappointed by the quake vigil at Parkfield (see main text). The earthquake's failure to appear in the expected time frame has disappointed specialists in short-term prediction as well. They had hoped that weeks or days before the quake materialized, their instruments would pick up premonitory stirrings that would enable seismologists to issue public warnings prior to future quakes.

Researchers were counting on the Parkfield Earthquake Prediction Experiment—with its closely packed web of seismometers, strainmeters, and other instruments—to help them change the discouraging record of short-term earthquake prediction. After years of scrutinizing records of other earthquakes, researchers have yet to nail down a single reliable earthquake precursor. In both the Loma Prieta earthquake of 1989 and last year's Landers earthquake, for example, the earth gave no warning that was



Taking the measure of Parkfield. A laser distance-measuring device watches for precursory crustal motion.

recognized as such before the quake let loose. The result, says Robert Wesson, head of the USGS's Office of Earthquakes, Volcanoes, and Engineering in Reston, Virginia, is that "geophysicists are sobered by the difficulties of short-term earthquake prediction."

What encourages them to keep looking for warning signs, says Wesson, is that at least in hindsight, "these big earthquakes don't just happen out of the blue. There is a process that has some symptoms" months before the quakes strike. In retrospect, seismologists have recognized smaller earthquakes recorded up to a year before Loma Prieta and Landers as foreshocks or "preshocks." But at the time, there was little to distinguish them from other earthquakes. As seismologist Thomas Heaton of the USGS in Pasadena notes, "There's no way we know of now that we could have definitely identified them as precursors"—until the mainshock arrived and proved that they were.

other external influence—this one from farther down the San Andreas—also affects the timing of earthquakes at Parkfield. Ben-Zion and his colleagues point out that the Parkfield clock is in part wound from the south, by stress emanating from great earthquakes there and stored in the “spring” of the ductile deep crust. A great earthquake that stretched south of Parkfield in 1857 gave the clock its last such winding. Since then, say Ben-Zion and colleagues, the clock should have been slowing down. Though the slowing has been hidden until now in the earthquake record’s fluctuations, by now it should be enough to delay the next quake by 4 to 7 years.

Further complicating the picture of fault segments as isolated systems has been the realization that their endpoints can vary, affecting the size, if not the timing, of the earthquakes they spawn. The magnitude 6 Parkfield forecast assumed that the earthquakes there are always confined to the same stretch of fault. That’s been true for Parkfield over the past 120 years, but other faults haven’t acted so constrained: When Wayne Thatcher of the USGS in Menlo Park looked at 10 places around the Pacific rim where two or more fault segments had broken together to produce great earthquakes, he found not a single case in which the same combination of segments broke twice in a row.

Drawing on past studies and their own work, tectonophysicists Evelyn Roeloffs of the USGS in Vancouver, Washington, and Ruth Harris of the USGS in Menlo Park recently emphasized the potential for the same kind of variable behavior at Parkfield. The next rupture, they say, could break

through what many researchers had taken to be a barrier at the southern end of the segment—a 1-kilometer jog in the surface track of the fault—to produce a magnitude 7.1 event, more than 10 times more powerful than was forecasted.

If earthquake forecasters only had to contend with external factors that might affect earthquake timing or size, they wouldn’t be so worried about the prospects for long-term prediction. Seismologists have become increasingly aware, however, of another complication—the underlying complexity of faults themselves. The image of a fault as a simple system that gradually stores stress, then releases it at a specific threshold is crumbling as seismologists try to reproduce faults and the forces acting on them in computer models.

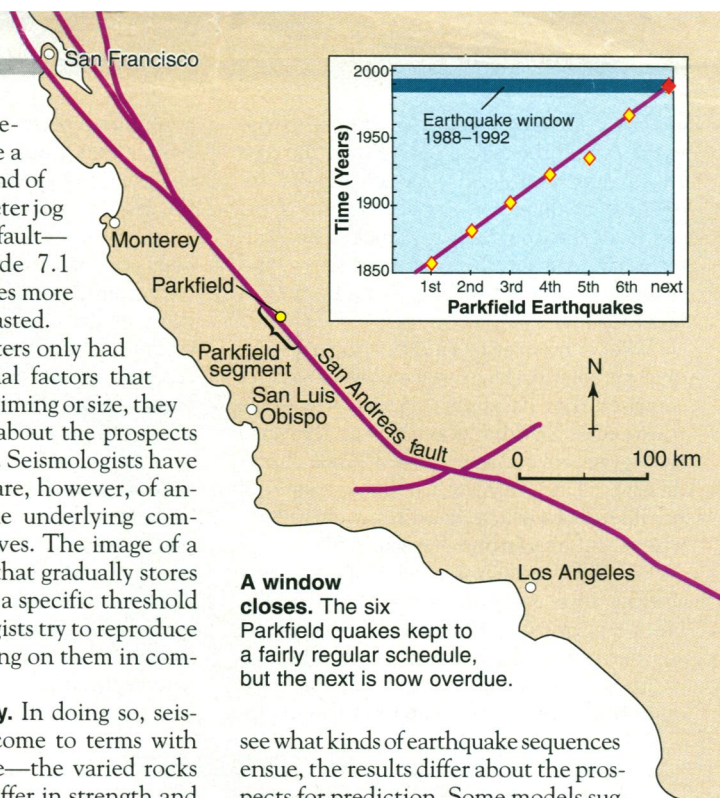
Creeping complexity. In doing so, seismologists have had to come to terms with faults’ complex structure—the varied rocks that line them, which differ in strength and other mechanical properties, resulting in a mosaic of weaker and stronger patches along the fault. “The basic problem is that the crust is massively fractured and complicated,” says Ben-Zion. “And complex systems, like the stock market, just don’t behave regularly. If you think in terms of a simple clock, you have no hope of predicting earthquakes.”

Just how badly fault complexity is mucking up the Parkfield clock, or any other earthquake cycle, scientists are hard pressed to say. When researchers run their fault models for decades, even centuries, of simulated time to

A window closes. The six Parkfield quakes kept to a fairly regular schedule, but the next is now overdue.

see what kinds of earthquake sequences ensue, the results differ about the prospects for prediction. Some models suggest that fault behavior is likely to be so chaotic that long-term prediction of the next quake in a sequence would be highly imprecise if not impossible. Others offer more hope. After all, as daily forecasts of the weather show, there can be some predictability even in chaotic systems.

The computer models developed by Tullis and Rice, for example, differ in the properties assigned to the rock along the fault, and those disagreements lead to quite different results. Tullis’ model, the more uniform of the two, produces a fairly regular sequence of quakes



Other hoped for precursors just haven’t appeared at all—precursory deformation of the crust, for example. According to seismological theory, failure of a patch of fault that is locked tight should begin around its edges. As the failure progresses inward, the crust around this “preparation zone” should deform at an increasing rate until, days or weeks later, a full-blown rupture occurs. But even though researchers had their instruments close enough to the Loma Prieta epicenter to detect any preparation zone larger than a couple of kilometers across, they saw nothing, according to Malcolm Johnston of the USGS in Menlo Park.

Preparation zones may be so tiny because the patches of fault on which earthquakes, even large ones, get their start are tiny as well. Rachel Abercrombie of the University of Southern California and James Mori of the USGS in Pasadena compared records of the first few seconds of the Landers earthquake with the start of a magnitude 5.5 Landers aftershock. Could they tell “which one was going to grow into a large earthquake? I think you’d be hard-pressed to say,” concedes Abercrombie.

If big quakes are simply little ones that run away, says Abercrombie, “you may have to predict 5’s if you want to predict a 7, and how do you tell which 5 will become a 7?” Seismologists have tried to lay the groundwork for doing so by dividing a fault into segments and gauging how likely each is to be triggered by a moderate quake on an adjacent segment. But Landers didn’t bode well for that effort either. “Landers broke all the standard rules of

fault segments,” says Heaton. The rupture jumped between faults through a previously unmapped fault, jumped from mid-fault to mid-fault, and stopped in the middle of a fault segment. “I think it would have been very difficult to predict this earthquake,” he says.

Does all that mean that the Parkfield experiment is a waste of resources? An answer to that question will be forthcoming soon from a panel of experts chaired by Bradford Hager of the Massachusetts Institute of Technology, which is scheduled to deliver its report to the National Earthquake Prediction Council.

Seismologists say the report stresses the experiment’s moderate cost—less than \$2 million per year, out of the \$50 million spent by the USGS on studying earthquakes and reducing quake hazards. The state of California, they add, is pleased by the smoothly operating public warning system developed for Parkfield, which was exercised recently following a possible foreshock (*Science*, 30 October 1992, p. 742). And most researchers feel that the unequaled density of proven instrumentation in the experiment is providing their first reliable, detailed views of the frequent twitchings of the fault, some of which may yet turn out to be earthquake precursors.

“I’m extremely enthusiastic about Parkfield,” says Harvard University quake modeler Yehuda Ben-Zion. “My studies say you have to work on the details” to get any understanding of fault behavior, he says, “and Parkfield is essential to that.”

—R.A.K.

at Parkfield, more or less in line with the variability of the actual record over the past 140 years. Rice's can produce regular sequences of quakes too, but only for short spurts. In the long run, this model's variability is twice what the relatively short, and therefore possibly misleading, historical record shows.

With complications multiplying in earthquake prediction, many seismologists weren't surprised that Parkfield failed to live up to the forecast. But that doesn't mean that they have given up on more modest kinds of predictions. They've recently had a couple of notable successes in forecasting where—if not when—future earthquakes will strike.

On 15 January, just as seismologists were settling into the post-Parkfield prediction blues, a magnitude 5.1 earthquake struck the southern end of the Calaveras fault, a branch of the San Andreas running east of San Francisco Bay. Much to the delight of David Op-

penheimer of the USGS in Menlo Park, Bakun, and Lindh, the quake fulfilled a partial prediction they had developed in 1989 (*Science*, 21 April 1989, p. 286). They had divided the southern 60 kilometers of the fault into six segments, based on patterns of background microseismicity, and found that three of the segments had been broken by recent quakes. After checking the historical record to see which segments might be due to break again, they concluded that the two segments at either end of the study area were "the most likely sites for the next [greater than magnitude] 5 earthquakes."

Oppenheimer and company were right on the money when the January quake ruptured the designated southern segment. And in another successful prediction of a quake's location and magnitude, a magnitude 7.6 shock struck offshore of Nicaragua last September at a location seismologists had predicted more than a decade earlier

(*Science*, 30 October 1992, p. 743).

Such successes, along with the historical record of frequent Parkfield quakes, convince researchers that the 1985 forecast is still likely to be right on a very basic level: Sooner rather than later, a good-sized quake will strike Parkfield. Eventually, the \$1.8 million being spent each year on the Parkfield Earthquake Prediction Experiment will pay off, most researchers say, and probably sooner than if the experiment were moved to any other fault segment. "No matter how you cut it," says McEvelly, "we're in the last quarter of a magnitude 6 cycle...and the experiment is heading for a conclusion." Better late than never.

—Richard A. Kerr

Additional Reading

Abstracts for the sessions "Parkfield as the Prediction Window Closes I and II," *EOS Trans. AGU* **73**, Fall Meeting Supplement, 396, 406 (1992).

MATERIALS SCIENCE

Shiny Molecules Gain New Luster

San Jose was teeming with thousands of points of light on 1 and 2 February as researchers gathered at a meeting sponsored by the Society for Imaging Science and Technology and the International Society for Optical Engineering. Much of that sparkle came from a symposium devoted to light-emitting organic molecules—"the first worldwide meeting of people working in this field," says Milan Stolka, a meeting chairman and research supervisor at Xerox Corp.'s Webster Research Center in Rochester. These substances, which give off light when sandwiched between electrodes, are brightening prospects for new kinds of computer monitors, area lighting, even TVs.

Stolka and his colleagues have been following this guiding light for years, but the meeting was marked by a new optimism. For one, it was energized by the announcement, during a talk by Shogo Saito of Kyushu University in Fukuoka, that colleagues at the Pioneer Electronic Corp. in Saitama had built an experimental electroluminescent (EL) patch that shone, at least for a short time, with an eye-shocking brightness about 2000 times that of a computer monitor. And while the Pioneer device, like most organic EL devices so far, is a short-lived parfait made up of molecular layers, researchers at the meeting were also talking about a newer approach that could be more practical: coaxing light from durable sheets of polymers.

Although the research is only now coalescing into a field, its roots go back to the discovery, some 30 years old, that certain organic molecules emit light when oppositely charged particles recombine within them, says Ching Tang of the Eastman Kodak Co. in

Rochester. That discovery acquired a new significance with the rise of laptop computers. Liquid crystal displays, a leading technology, quickly drain laptop batteries because they make inefficient use of light: They often work by blocking or reflecting light coming from fluorescent tubes behind the liquid crystal layer.

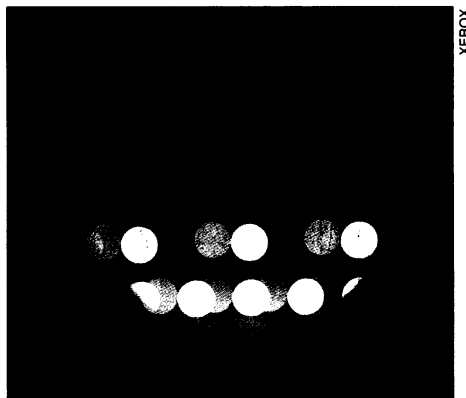
The alternative may turn out to be a multidecker sandwich centered on a thin layer of organic molecules. Electrodes on the outer layers of the sandwich supply negative charges, carried by electrons, and positive

In a standard recipe for these illuminating sandwiches, the light-emitting layer is a thin film of small organic or organometallic molecules such as the quinolium-aluminum complex (which consists of an aluminum atom flanked by a trio of quinolium molecules). Several groups, including Tang's group and the Pioneer researchers, are striving to increase the intensity or efficiency of the light emission by chemically modifying the emitting layer. Many of the designs rely on additives—such as electron-friendly oxadiazole or hole-hosting aromatic amines—that act as social directors, encouraging more excitonic meetings between opposite charges.

For all their bright promise, however, organic films are handicapped by their need for brittle components, such as transparent indium-tin-oxide electrodes, that restrict their flexibility and by chemical instabilities that limit their lifetimes. That's why many researchers, including Xerox's Gordon Johnson and Kathleen McGrane and Richard Friend and his colleagues at Cambridge University, are replacing the fragile layer of discrete organic molecules with more durable and flexible polymers such as poly-(phenylene vinylene)s. "These are more mechanically and thermally stable than the molecular thin films," Tang admits. But so far, he adds, they are only about one-tenth as efficient.

Everyone agrees that it is too early to place bets on which strategy will emerge in the commercial arena. But plenty of established companies are now earnestly investigating both approaches. And in another measure of the vogue for shiny molecules, it has already spawned at least two startup companies, one in Cambridge, Massachusetts, and another in Santa Barbara.

—Ivan Amato



Bright stuff. An electroluminescent polymer.

charges, which are called holes. The electrons and holes migrate inward; when they encounter each other they join in a transient liaison called an exciton. The organic layer has the rare property of allowing excitons to persist awhile, and instead of recombining in a lightless process, as would happen in most organic materials, the opposite charges in some excitons collapse in light-emitting unions.