

scientific misconduct.

Ramirez, who once worked at Bell Labs, points out that, though researchers there have published prolifically on organics over the past two years, they spent several years before that assembling the necessary expertise. Establishing a cohesive group of workers with disparate backgrounds is often key to rapid success, Ramirez says.

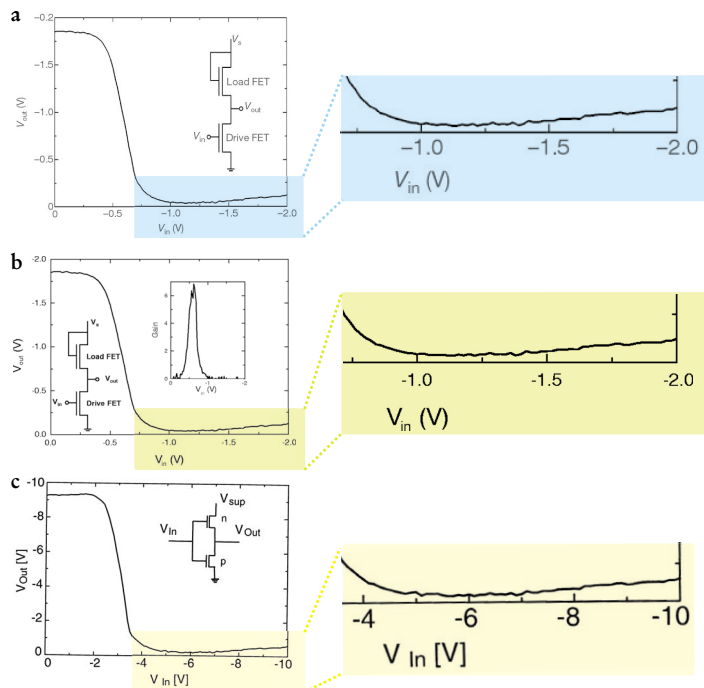
Two factors have been especially critical in the Bell Labs work: the purity of the organic crystals and the quality of the oxide layer. For holes and electrons to have high mobilities, organic crystals must be free of defects that can trap the excess charges.

Ramirez told us that Bell Labs' Christian Kloc, who has grown most of the labs' organic crystals, had been working on growing very pure crystals as far back as 1997.

The aluminum oxide layers for the experiments in question were laid down by Schön at the University of Konstanz in Germany, where Schön did his graduate work. He began applying oxide layers to the Bell Labs crystals in Konstanz while waiting for his US visa, and he later continued to use the same sputtering machine because of his familiarity with it. Those trying to reproduce the Bell Labs achievements have asked how the layers in the Bell Labs samples have been able to withstand the reported high voltages without breakdown. High voltages have been necessary to achieve the doping levels required for superconducting behavior, but not for all the observed phenomena.

The inquiry

Serving with Beasley on the investi-



CURVES LOOK IDENTICAL at the high voltage end in all three panels, taken from different papers, and at the low voltage end for the top two panels. All three give the input-output characteristics of a field-effect transistor, but for different materials: (a) a self-assembled monolayer of undiluted 4,4'-biphenyl-dithiol, (b) a self-assembled monolayer of the same molecule but diluted with non-conducting molecules, and (c) pentacene. (Adapted from refs. 1, 2, and 4.)

gatory committee are Herbert Kroemer of the University of California, Santa Barbara; Supriyo Datta of Purdue University; Herwig Kogelnik of Bell Labs; and Donald Monroe of Agere Systems, a spinoff of Lucent. Beasley had a hand in picking his committee. When asked whether it was appropriate to include members from Bell Labs or related institutions, he answered that there are precedents for such members: The benefit of including them is to help other

committee members understand the institution under investigation and to lend legitimacy for the researchers from that institution.

As for the scope of the inquiry, Beasley said that is not yet defined. "We have to go with what we find," he said. Committee members are receiving lots of information, from Bell Labs and from outside. Their inquiry is being guided by the federal policy on research misconduct, though it technically does not apply to work that does not receive federal funding. Many people are anxious to learn the outcome but, says Beasley, his committee will proceed only as fast "as fairness and thoroughness dictate."

In the meantime, Murray told us, Bell Labs has offered its complete cooperation. It will open its doors and its books to the investigators. She expressed her gratitude that the committee members, whom she described as "blue ribbon panelists," had agreed on such short notice to serve. They plan to make the committee's report public.

Murray said the researchers are currently employed, trying to reproduce their results; all are cooperating with the committee. Schön has said he stands behind his work.

BARBARA GOSS LEVI

References

1. J. H. Schön, H. Meng, Z. Bao, *Nature* **413**, 713 (2001).
2. J. H. Schön, H. Meng, Z. Bao, *Science* **294**, 2138 (2001).
3. Correction in *Science* **296**, 1400 (2002).
4. J. H. Schön, S. Berg, Ch. Kloc, B. Batlogg, *Science* **287**, 1022 (2000).
5. See, for example, J. H. Schön et al., *Phys. Rev. B* **58**, 12952 (1998).

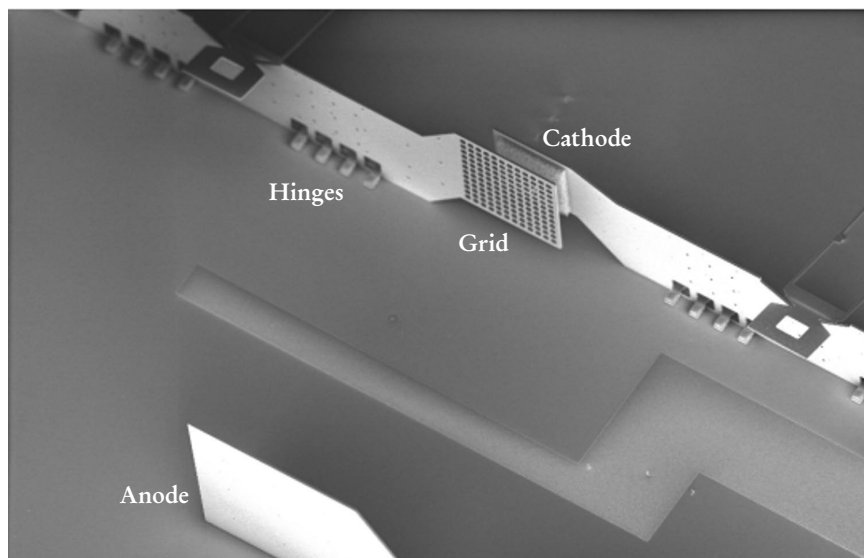
Researchers Combine Carbon Nanotubes with MEMS Technology to Make a Tiny Triode

Though superseded for most applications by the solid-state transistor, the venerable vacuum tube boasts one big advantage over its younger usurper: power. Free of stuff between its electrodes, a vacuum tube can operate at currents that would ohmically toast a transistor's semiconducting

▶ Miniature on-chip vacuum tubes could power efficient wireless communication devices.

innards. Vacuum tubes remain the technology of choice for high-power amplifiers, magnetrons, and klystrons.

Now, a team led by Wei Zhu of Agere Systems in Murray Hill, New Jersey, has built a device that could propel the vacuum tube out of its traditional niches.¹ The Agere device, a 100- μm -scale triode, exploits two of the hottest technologies of the past decade: carbon nanotubes and micro-



electromechanical systems (MEMS). And because the new triode sits on a silicon substrate, it could readily plug into an integrated circuit. Agere developed the triode as a proof of concept. Devices like it could end up in compact wireless transmitters.

Hot and cold cathodes

All vacuum tubes, including the Agere triode, share the same basic components. A cathode emits electrons, which shoot across an evacuated gap toward an anode. Triodes include an additional electrode—the grid—between the cathode and anode. Because it's placed close to the cathode, the grid can attract or repel electrons with modest changes in voltage, thereby controlling the current across the device.

In traditional vacuum tubes, ohmic heating provides electrons with the energy they need to escape the cathode surface. Though simple, hot cathodes are far from ideal. They can't turn on rapidly because they have to reach a temperature of at least 800°C. The operating temperature is so high that hot cathodes burn out when the rest of the device could keep going. And because of the heat, you can't place the grid or anode arbitrarily close to the cathode. As a result, hot cathode devices can't be easily miniaturized. And the smaller the gap between the electrodes, the higher the frequency of the radiation produced when running the vacuum tube as an oscillator.

The Agere triode uses field-emitting cathodes, which rely on a strong electric field, rather than heat, to release the electrons. As Ralph Fowler and Lothar Nordheim explained in 1928, the field reduces the tunneling barrier so that electrons can easily tunnel out

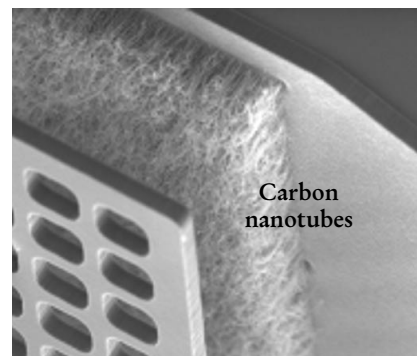
to the cathode surface.² Because they're cold, field-emitting cathodes avoid the pitfalls of their hotter cousins, but getting field emitters to work has its own set of challenges. At 1 gigavolt per meter, the required field to coax emission from most materials is so great that the emitters must be extremely narrow. And to produce enough current, you need a lot of them.

With those aims in sight, in 1970 Capp Spindt of SRI International invented a way to fabricate an array of pointy metallic field emitters using advanced lithography. But, despite 30 years of development, Spindt-type cold cathodes have attained somewhat limited success. The problem lies in creating an array of uniform emitters. Millions of individual emitters make up a typical Spindt-type cathode, but only the 1000 to 10 000 sharpest of them emit. Because these few emitters carry all the current, they tend to fail—sometimes explosively. Another disadvantage is the steep cost of the high-precision lithography needed to fashion the tips.

Nanotubes are cool

Eight years ago, Walt de Heer realized that carbon nanotubes would make superb field emitters. Although nanotube carbon has a high work function, nanotubes have such tiny tips that a field of just tens of volts per micrometer will liberate electrons. De Heer, who's now at Georgia Tech, and his collaborators, André Châtelain (Ecole Polytechnique Fédérale de Lausanne, Switzerland) and Daniel Uzgate (Brazilian National Synchrotron Light Laboratory), built a high-intensity electron gun in 1995 from a dense array of aligned nanotubes.³

To create their nanotube arrays, de Heer and company drew a suspension



THIS MINIATURIZED TRIODE (left) has an anode-to-grid distance of about 200 μm and the grid-to-cathode distance is about 20 μm . The hinges used in the pop-up fabrication are visible at the base of both the grid and the cathode. Shown above is the array of carbon nanotube emitters. Each emitter is 8 μm long and 10 nm in diameter. (Courtesy of Wei Zhu, Agere Systems.)

of ready-made tubes through a ceramic sieve and then transferred the aligned and upright tubes to a substrate. Since that pioneering work, new methods have emerged for growing aligned nanotubes directly on the cathode. The Agere group used a technique they'd developed two years ago.⁴ First, a thin layer of iron is deposited onto what will eventually be the cathode. Iron serves to nucleate the nascent nanotubes, which form when the gaseous reactants, ammonia and acetylene, are vaporized by a strong microwave field. Under the influence of the field, the nanotubes grow quickly and evenly in thick, dense arrays like a bamboo grove. Turning off the field halts the growth and controls the height of the tubes.

Pop-up triode

To build their triode, the Agere researchers exploited a MEMS technique called surface micromachining. The technique's practitioners assemble devices in a series of patterned layers of various materials, usually forms of silicon, such as polysilicon, silicon dioxide, and silicon nitride. Micromachinists etch away oxide layers to leave behind intricate three-dimensional structures on the substrate, just as a sculptor breaks apart a plaster mold to reveal the cast metal sculpture inside. (For more on MEMS, see the article on page 38 of the October 2001 PHYSICS TODAY.)

But building vertical structures in layers takes time. As a short cut, MEMS engineers create pop-up structures. Pioneered 10 years ago by Kris Pister of the University of California,

Berkeley, pop-up fabrication involves micromachining structures that resemble the hatches on ships. First, the structure, lying flat and replete with hinges, is patterned. Next, the layer beneath the structure and around the hinges is etched away. Pop-up structures can be engineered to spring upright spontaneously once the underlying layer disappears. But for their triode, shown in the accompanying figure, the Agere researchers chose instead to raise the three electrodes by hand under a microscope.

Other ideas

The Agere team developed the triode as a proof of concept, rather than as a production prototype. The triode shows the expected field-emitting behavior and amplifies the grid current by a significant, but modest, factor of four. However, the team couldn't run the triode in AC mode as a

microwave generator because too much current is lost each cycle to capacitance that strays into the circuit from the base of the device. With an optimized choice of materials, a triode of about the same dimensions should, the Agere team calculates, operate close to 200 MHz.

In the vacuum tube marketplace, pentodes and induction-output amplifiers, not triodes, are the biggest money earners. The Agere team has already built a pentode, but the biggest challenge remains: to build a microwave device that can operate at the 1–2 GHz frequencies used in wireless communications.

The basic concept behind the Agere triode, controlling the flight of ballistic electrons in minuscule settings, offers other possibilities. David Garner of University College London envisions creating a tiny UV light source

by sending a beam of electrons through a low density gas, such as nitrogen. Ionized by the electrons, the atoms would deexcite by emitting UV photons, just like a fluorescent light bulb. Dan Nicolaescu of Romania's National Institute for Research and Development in Microtechnologies has proposed using electrons from a cold cathode to measure magnetic fields. Deflected by the Lorentz force, field-emitted electrons could hit one of several anodes, depending on the magnitude of the field. **CHARLES DAY**

References

1. C. Bower et al., *Appl. Phys. Lett.* **80**, 3820 (2002).
2. R. H. Fowler, L. W. Nordheim, *Proc. R. Soc. London* **A119**, 173 (1928).
3. W. A. de Heer, A. Châtelain, D. Ugarte, *Science* **270**, 1179 (1995).
4. C. Bower et al., *Appl. Phys. Lett.* **77**, 830 (2000); C. Bower et al., *Appl. Phys. Lett.* **77**, 2767 (2000).

X-Ray Spectrum Challenges Models of Gamma-Ray Bursts

Astrophysicists have been studying gamma-ray bursts (GRBs) for more than 30 years, but they still don't fully understand the cataclysmic cosmic processes that give rise to these brief showers of energetic gamma rays.¹ One technique for learning about the explosions (see also page 24 in this issue) is to study the emissions of the x-ray, optical, or radio afterglows that follow the GRBs: Afterglows can reveal details of the temperature, ionization, composition, and other features of the material illuminated by the bursts.

This past April, James Reeves and colleagues at the University of Leicester, UK, presented an unusually detailed emission spectrum² of the x-ray afterglow following the gamma-ray burst GRB011211, so named because it was observed on 11 December 2001. The paper has generated a lot of questions, according to Reeves, with scientists puzzling over how to reconcile the data with their favored theories of GRB formation.

The first GRB was detected on 2 July 1967 by US surveillance satellites built to ensure that the Soviet Union was not testing nuclear weapons in space in violation of the Nuclear Test Ban Treaty. Thirty years later, the Italian–Dutch satellite BeppoSAX recorded a GRB with a redshift of about 0.8, confirming that the bursts were of cosmological origin, not confined to our galaxy (see *PHYSICS TODAY*, July 1997, page 17, and the article by Neil Gehrels and

▶ The debate is heating up: Does the progenitor of these powerful explosions collapse in one step or two?

Jacques Paul in *PHYSICS TODAY*, February 1998, page 26).

A GRB releases a staggering amount of energy, perhaps as much as 10^{44} – 10^{45} joules. It is generally agreed that GRB energy is released in a pair of jets. Because the bursts are jetted, we see only a fraction of the GRBs emitted in the universe on any given day. To estimate how often GRBs occur, one needs to know the solid angle subtended by the jets. Typical theoretical models give a value of 0.01–0.1 steradians. Combining this value with observed occurrences of GRBs, one concludes that there are, roughly speaking, hundreds of bursts every day.

Modeling bursts

The duration of the prompt gamma-ray luminous phase of a GRB can range from 0.001 to 1000 seconds. Most of the bursts are “long,” with durations of more than 2 seconds. The long bursts are the only ones for which afterglows have been observed. Short bursts display qualitatively different energy spectra with relatively more high-energy gamma rays. The spectral differences between the short and long bursts, and the different timescales associated with them, hint that they may originate from different physical mechanisms.

Models for gamma-ray bursts fall

into two main sets. One set posits that GRBs are generated by the coalescence of two compact objects, such as two neutron stars or a neutron star and a black hole. In the second set of models, the progenitor whose catastrophic collapse leads to a GRB is a single massive object.

For about five years, a consensus has been growing that neutron-star-binary mergers or similar processes are not the cause of long-duration GRBs. As early as 1997, the Hubble Space Telescope showed long-duration GRBs occurring near the optical disks of galaxies. Such observations argue against merger scenarios if, as many believe, the gradual decay of the binary orbits occurs over billions of years. Over that long span, the binary system should drift far from the galactic plane. Because of the drift, coalescing neutron-star binaries would emit GRBs in a region of space with interstellar matter too diffuse to allow for x-ray afterglow emission. Thus, other strikes against the coalescence picture are the x-ray spectrum Reeves and colleagues observed and iron x-ray fluorescence earlier researchers saw. Optical afterglows have generally been observed in relatively young, star-forming regions of galaxies, a further argument against coalescence models.

The timescale associated with the catastrophic final phase of binary merging is much shorter than the several-second timescale of long-duration GRBs and single progenitor models.