

SCIENTIFIC AMERICAN

OCTOBER 1997

\$4.95

**SPECIAL
ISSUE**

THE FUTURE OF **TRANSPORTATION**

Spaceflight Made Easy

Sideways Elevators

High-Speed Trains

Tiltrotor Planes

750-mph Cars

Microsubs

and more



THE FUTURE OF TRANSPORTATION



SCIENTIFIC AMERICAN

October 1997
Volume 277
Number 4

FROM THE EDITORS

6

LETTERS TO THE EDITORS

8

50, 100 AND 150 YEARS AGO

12

NEWS AND ANALYSIS

IN FOCUS

Tissue engineers try to grow organs
in the laboratory.

15

SCIENCE AND THE CITIZEN

Wallaby science.... A schizophrenia
virus?... Protein alchemists turn
sheets into coils.... Why Darwin
flunks with students.

20

PROFILE

Jane Goodall cares about science
but loves chimpanzees.

42

TECHNOLOGY AND BUSINESS

Short-circuiting the senses....
A consumer choice on energy....
Bye-bye, batteries.

46

CYBER VIEW

Masters of their domain (name)
find crowding on-line.

52

About the Cover
Image by Bryan Christie.



54

Transportation's Perennial Problems

W. Wayt Gibbs



58

The Past and Future of Global Mobility

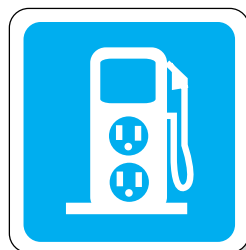
Andreas Schafer and David Victor



64

13 Vehicles That Went Nowhere

John Rennie



70

Hybrid Electric Vehicles

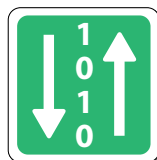
Victor Wouk



80

Automated Highways

James H. Rillings



86

Unjamming Traffic with Computers

Kenneth R. Howard



93

Now That Travel Can Be Virtual, Will Congestion Virtually Disappear?

Patricia L. Mokhtarian



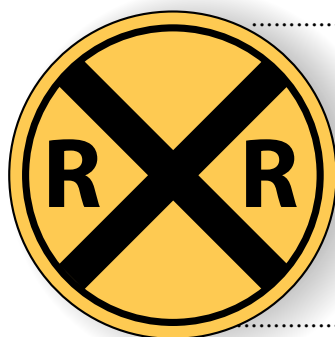
94

Driving to Mach 1

Gary Stix

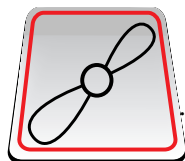


98 **Speed versus Need** Kristin Leutwyler



100 **How High-Speed Trains Make Tracks** Jean-Claude Raoul

106 **Fast Trains: Why the U.S. Lags** Anthony Perl and James A. Dunn, Jr.



109 **Maglev: Racing to Oblivion?** Gary Stix



110 **Straight Up into the Blue** Hans Mark



116 **The Lure of Icarus** Shawn Carlson



120 **A Simpler Ride into Space** T. K. Mattingly



126 **Faster Ships for the Future** David L. Giles



132 **Microsubs Go to Sea** Graham S. Hawkes

136 **Elevators on the Move** Miriam Labob

THE AMATEUR SCIENTIST Hear the beating of an unborn heart with an electronic stethoscope. 138

MATHEMATICAL RECREATIONS Jigsaw puzzles with more than one solution. 140

REVIEWS AND COMMENTARIES

Homosexuality under the
microscope.... Extraordinary
beauty in commonplace things....
Fleeing the DNA cops.

Wonders, by Philip Morrison
The cool secrets
of champion bicyclists.

Connections, by James Burke
Decimals, Descartes and dollars.

146

WORKING KNOWLEDGE How fish can climb ladders. 156

Scientific American (ISSN 0036-8733), published monthly by Scientific American, Inc., 415 Madison Avenue, New York, N.Y. 10017-1111. Copyright © 1997 by Scientific American, Inc. All rights reserved. No part of this issue may be reproduced by any mechanical, photographic or electronic process, or in the form of a phonographic recording, nor may it be stored in a retrieval system, transmitted or otherwise copied for public or private use without written permission of the publisher. Periodicals postage paid at New York, N.Y., and at additional mailing offices. Canada Post International Publications Mail (Canadian Distribution) Sales Agreement No. 242764. Canadian BN No. 127387652RT; QST No. Q1015332537. Subscription rates: one year \$34.97 (outside U.S. \$47). Institutional price: one year \$39.95 (outside U.S. \$50.95). Postmaster: Send address changes to Scientific American, Box 3187, Harlan, Iowa 51537. Reprints available: write Reprint Department, Scientific American, Inc., 415 Madison Avenue, New York, N.Y. 10017-1111; fax: (212) 355-0408 or send e-mail to info@sciam.com Subscription inquiries: U.S. and Canada (800) 333-1199; other (515) 247-7631.

Visit the SCIENTIFIC AMERICAN Web site (<http://www.sciam.com>) for more information on this issue's articles and other on-line features.

The Way to Go

Modern humans probably walked out of Africa about 100,000 years ago, then kept on going. First by foot, then on horseback, boat, wheels and wings, our kind has charged across the land and seas to every part of the globe. While one courageous minority invaded the depths of the oceans, another built rockets to visit the moon and near space. Not content to go places once, our entire civilization is bound up with the enterprise of getting to places again and again: more quickly, more easily, with more luxury or more cargo or less expense.

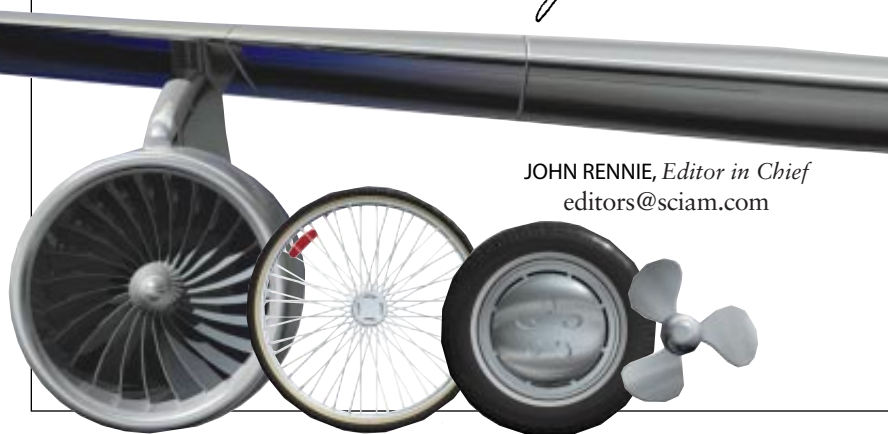
One striking point in most serious predictions is that modes of transportation in the next century will be, by and large, not too different from the ones we use now. (Well, there go my personal gyrocopter stocks.) Automotive technology will advance considerably, migrating away from so much reliance on polluting fossil fuels and toward use of electricity or other sources of power, yet the American love affair with the car will remain torrid. We may log proportionally more miles in aircraft or high-speed trains, but driving will still be our day-to-day first choice for most travel. Vastly more people around the world will be expressing the same preference, too, because they can afford to. Andreas Schafer and David Victor explain why that will be so in "The Past and Future of Global Mobility," beginning on page 58.

In aviation, the greatest changes may come in the numbers of aircraft, their safety, their efficiency and the transfer of advanced military technologies to the commercial sector. Average flight times may get shorter, not because new hypersonic aircraft will be making jaunts between Tokyo and New York in a few hours but largely because air-traffic management will be computerized and subsonic planes will get incrementally faster. Nevertheless, expect some novel vehicles, such as the vertical-takeoff planes described by Hans Mark (see page 110), to take to the skies.

In this issue, we have highlighted some of the more important trends and innovations that will shape transportation—over the land, through the air, across and under the oceans and into space—for the next few decades. Improvements even in low-glamour technologies, such as those for elevators and bicycles, can leave a big impression. But because travel and transportation are often fascinating for their own sake, we have also included a few ideas that lack something in practicality but make up for it in sheer fun. Human-powered planes, supersonic cars and microsubmarines are the perfect vehicles for chasing dreams. In your heart, do you know a better way to go?

John Rennie

JOHN RENNIE, *Editor in Chief*
editors@sciam.com



**SCIENTIFIC
AMERICAN®**
Established 1845

John Rennie, EDITOR IN CHIEF

Board of Editors

Michelle Press, MANAGING EDITOR
Philip M. Yam, NEWS EDITOR
Ricki L. Rusting, ASSOCIATE EDITOR
Timothy M. Beardsley, ASSOCIATE EDITOR
Gary Stix, ASSOCIATE EDITOR
Corey S. Powell, ELECTRONIC FEATURES EDITOR
W. Wayt Gibbs; Kristin Leutwyler; Madhusree Mukerjee;
Sasha Nemecek; David A. Schneider; Glenn Zorpette
Marguerite Holloway, CONTRIBUTING EDITOR
Paul Wallich, CONTRIBUTING EDITOR

Art

Edward Bell, ART DIRECTOR
Jana Brenning, SENIOR ASSOCIATE ART DIRECTOR
Johnny Johnson, ASSISTANT ART DIRECTOR
Jennifer C. Christiansen, ASSISTANT ART DIRECTOR
Bryan Christie, ASSISTANT ART DIRECTOR
Bridget Gerety, PHOTOGRAPHY EDITOR
Lisa Burnett, PRODUCTION EDITOR

Copy

Maria-Christina Keller, COPY CHIEF
Molly K. Frances; Daniel C. Schlenoff;
Terrance Dolan; Katherine A. Wong

Administration

Rob Gaines, EDITORIAL ADMINISTRATOR
Sonja Rosenzweig

Production

Richard Sasso, ASSOCIATE PUBLISHER/
VICE PRESIDENT, PRODUCTION
William Sherman, DIRECTOR, PRODUCTION
Janet Cermak, MANUFACTURING MANAGER
Tanya DeSilva, PREPRESS MANAGER
Silvia Di Placido, QUALITY CONTROL MANAGER
Carol Hansen, COMPOSITION MANAGER
Madelyn Keyes, SYSTEMS MANAGER
Carl Cherebin, AD TRAFFIC; Norma Jones

Circulation

Lorraine Leib Terlecki, ASSOCIATE PUBLISHER/
CIRCULATION DIRECTOR
Katherine Robold, CIRCULATION MANAGER
Joanne Guralnick, CIRCULATION PROMOTION MANAGER
Rosa Davis, FULFILLMENT MANAGER

Advertising

Kate Dobson, ASSOCIATE PUBLISHER/ADVERTISING DIRECTOR
OFFICES: NEW YORK:
Meryle Lowenthal, NEW YORK ADVERTISING MANAGER;
Kevin Gentzel; Thomas Potratz; Timothy Whiting.
DETROIT, CHICAGO: 3000 Town Center, Suite 1435,
Southfield, MI 48075;
Edward A. Bartley, DETROIT MANAGER; Randy James.
WEST COAST: 1554 S. Sepulveda Blvd., Suite 212,
Los Angeles, CA 90025;
Lisa K. Carden, WEST COAST MANAGER; Debra Silver.
225 Bush St., Suite 1453,
San Francisco, CA 94104
CANADA: Fenn Company, Inc. DALLAS: Griffith Group

Marketing Services

Laura Salant, MARKETING DIRECTOR
Diane Schube, PROMOTION MANAGER
Susan Spirakis, RESEARCH MANAGER
Nancy Mongelli, ASSISTANT MARKETING MANAGER

International

EUROPE: Roy Edwards, INTERNATIONAL ADVERTISING DIRECTOR,
London. HONG KONG: Stephen Hutton, Hutton Media Ltd.,
Wanchai. MIDDLE EAST: Peter Smith, Peter Smith Media and
Marketing, Devon, England. PARIS: Bill Cameron Ward,
Inflight Europe Ltd. PORTUGAL: Mariana Inverno,
Publicosmos Ltda., Parede. BRUSSELS: Reginald Hoe, Europa
S.A. SEOUL: Biscom, Inc. TOKYO: Nikkei International Ltd.

Business Administration

Joachim P. Rosler, PUBLISHER
Marie M. Beaumonte, GENERAL MANAGER
Alyson M. Lane, BUSINESS MANAGER
Constance Holmes, MANAGER, ADVERTISING ACCOUNTING
AND COORDINATION

Chairman and Chief Executive Officer

John J. Hanley

Corporate Officers

Robert L. Biewen, Frances Newburg,
Joachim P. Rosler, VICE PRESIDENTS
Anthony C. Degutis, CHIEF FINANCIAL OFFICER

Program Development **Electronic Publishing**
Linnéa C. Elliott, DIRECTOR Martin O. K. Paul, DIRECTOR

Ancillary Products

Diane McGarvey, DIRECTOR

SCIENTIFIC AMERICAN, INC.

415 Madison Avenue • New York, NY 10017-1111
(212) 754-0550

PRINTED IN U.S.A.

LETTERS TO THE EDITORS

SHARPER IMAGE

We read David Schneider's profile of Raymond V. Damadian ["Scanning the Horizon," News and Analysis, June] with interest. Damadian indeed performed an important early experiment, published in 1971, showing that excised samples had different magnetic resonance characteristics depending on whether they arose from normal or tumor tissue. It spurred on the development of magnetic resonance imaging (MRI), and he deserves recognition for that. But Schneider's article leaves the impression that MRI was single-handedly invented and developed by Damadian, and that view is plainly

wrong. The crucial contribution was made by Paul C. Lauterbur, who in that same year conceived the idea of using magnetic-field

gradients to obtain spatial information on the distribution of magnetic nuclei in a sample placed inside an NMR coil and thus was able to generate "pictures" that way.

WILLIAM J. LE NOBLE
CHARLES S. SPRINGER, JR.
State University of New York
at Stony Brook

Schneider replies:

My profile of Raymond V. Damadian indeed mentioned others' contributions to the development of MRI only in passing. Lauterbur clearly advanced the art significantly, and I should have noted that he jointly received the National Medal of Technology with Damadian. But Damadian needs to be credited with more than just measuring excised samples, as le Noble and Springer imply. Damadian realized that some method of localizing the signal would be needed to accomplish whole-body scanning, and he conceived of manipulating the magnetic field to do so in early 1971, some months before Lauterbur began his in-

vestigations. That Lauterbur's method proved technically superior to Damadian's technique is not in question. But in my view the first crucial step was Damadian's, even if the footwork was clumsy.

TREASURES AT DUNHUANG

I read with interest the article "China's Buddhist Treasures at Dunhuang," by Neville Agnew and Fan Jinshi [July]. I wonder, however, if the "foreign devils" who "began a systematic discovery and removal of the cultural heritage of the Silk Road" actually helped or hindered the preservation of this fascinating period in world history. Current efforts notwithstanding, can a case be made that the removed antiquities owe their very existence to the curatorship of these "foreign devils"? One can only speculate as to how the Buddhist treasures at Dunhuang would have fared at the hands of the agents of Mao's Cultural Revolution.

DARREL ZBAR
Hollywood, Fla.

GETTING A FIX ON NITROGEN

The potential environmental hazards posed by increased fixed nitrogen from anthropogenic sources are well stated by Vaclav Smil ["Global Population and the Nitrogen Cycle," July]. Yet his statement that lightning plays a minor role compared to bacteria in the global fixation of nitrogen may be premature. Research done by Carl J. Popp and myself (published in the *Journal of Geophysical Research* in 1989) suggests that lightning may be the major source of fixed nitrogen worldwide, supplying more than even human activities do. The implications of this possibility are far reaching and include a rethinking of much of atmospheric chemistry and the chemistry of global warming, environmental degradation and the origin of life.

EDWARD FRANZBLAU
Albuquerque, N.M.

Smil replies:

I am familiar with Franzblau's research in which he has estimated that a total of 100 million tons of nitrogen is fixed every year by lightning. And I agree that

there may be more reactive nitrogen fixed by lightning than is credited by many conservative estimates. But there is not enough nitrate (generated by the oxidation of nitrogen fixed by lightning) in the world's precipitation and dry deposition to balance this figure. Different studies constrain the amount of reactive nitrogen derived from lightning to between one and 20 million tons a year. Thus, a large uncertainty remains, but lightning is almost assuredly a less important source of reactive nitrogen than biofixation or synthesis of ammonia.

DECOHERENT STATE

Philip Yam's discussion in the June issue of the recent developments in the foundations of quantum physics ["Trends in Physics: Bringing Schrödinger's Cat to Life"] may leave readers with an impression that the phenomenon of decoherence is an ad hoc addition to quantum physics proper and that it allows the environment to determine the outcome of a measurement. Even though the role played by decoherence in the transition from quantum to classical mechanics has been recognized only recently, decoherence is, in fact, a consequence of quantum theory. It is essentially inevitable in macroscopic systems, which are all but impossible to isolate from the environment. The environment determines only which quantum states can stand such scrutiny and, therefore, will appear on a classical menu of the possibilities. In other words, dead or alive Schrödinger cats are okay, but their coherent superposition is not. This is why scientists with quite diverse interpretations of decoherence—such as Murray Gell-Mann, John A. Wheeler or one of its pioneers, H. Dieter Zeh—can agree on its consequences.

WOJCIECH H. ZUREK
Los Alamos National Laboratory

Letters to the editors should be sent by e-mail to editors@sciam.com or by post to Scientific American, 415 Madison Ave., New York, NY 10017. Letters may be edited for length and clarity. Because of the considerable volume of mail received, we cannot answer all correspondence.



Modern MRI machine

50, 100 AND 150 YEARS AGO



OCTOBER 1947

SYNTHETIC QUARTZ—"Quartz crystals, required in optical and electronic devices, and hitherto available only from scattered natural deposits, will be produced by the Naval Research Laboratories, Washington, D.C., as soon as equipment is installed for a new process of growing them. The method is based on techniques developed in Germany, and depends on the growth of a crystal from a seed placed in a solution of silica, sodium hydroxide or carbonate, and water, heated to 350 to 400 degrees Centigrade. Pressures generated may reach 2,000 to 3,000 pounds per square inch."

OCTOBER 1897

ARCTIC RESEARCH—"The latest Arctic adventure of Lieut. R. E. Peary, U.S.N., while devoid of sensational adventures and discoveries, was crowned with success from a scientific point of view. The great meteorite and the collections he gathered are worth all the expense and labor of the voyage. His vessel the *Hope* came into Sydney, Cape Breton, on September 20, nearly as deep in the water as when she left the port for the North—the great Cape York meteorite, the largest in the world, being in the hold embedded in tons of ballast. The meteorite is estimated to weigh up to 90 tons, and is composed of about 92 per cent iron and 8 per cent nickel." [Editors' note: *The meteorite is on display at the American Museum of Natural History in New York City.*]

PARASITES ON ANTS—"One of the most common parasites of the ants of the genus *Lasius* is an acarid, the *Antennophorus Uhlmanni*. This parasite does not move around in the formicary [ant nest], but lives constantly upon the body of the ants. As a general thing, an ant carries one acarid under the head and two to the right and left of the abdomen (*at left in illustration*). As soon as the *Antennophorus* has suc-

ceeded in creeping upon the ant, the latter, even in cases in which it is already carrying several of these parasites, struggles vigorously but soon resigns itself to the labor of carrying its new burden. Another common acarid parasite is *Discopoma comata* (*at right in illustration*)."

ARSENIC AND OLD WALLPAPER—"The fact that pigments containing arsenic are dangerous to health is widely known. It has been found that arsenical wallpaper, hung in damp rooms, has frequently caused chronic cases of poisoning in the occupants. Extensive researches have been made for the first time by Prof. Emmerling of the Berlin University. The results seem to confirm the correctness of the theory that the dust which becomes separated from the paper through wiping, as well as through expansion and contraction caused by changes in the temperature, is scattered about and enters the lungs of the occupants, thus giving rise to poisoning."

OCTOBER 1847

THERMAL TELESCOPE—"Professor Joseph Henry, of Princeton, N.J., communicated some interesting experiments with a Thermo Electrical apparatus, a very delicate instrument which will indicate 1/500th of a degree of a Fahrenheit thermometer. The apparatus was applied to form a Thermal Telescope: when turned to the heavens the coldest part was found to be directly over head. Experiments made upon the spots of the sun showed that they were colder than the surrounding parts; also, that the surface of that body was variously heated. The Thermo Electrical Telescope, when in a state of perfection, may reveal many new facts in astronomy, which thus far have only been opened to sight."

WATER AS FUEL—"This seemingly strange idea originated in a remark of Sir Humphrey Davy that, on the problematic exhaustion of coal, men will have recourse to the hydrogen of water, as a means of obtaining light and calefaction [heat]. As the gas used for lighting consists of hydrogen and a little carbon, it is only the latter which would have to be added, after the water had been decomposed into its elementary parts of hydrogen and oxygen."

FLOATING ROCKS—"The Association of American Geologists have just closed their annual meeting. Huge round rocks called *boulders*, found throughout different parts of our continent, have engaged a large share of their discussion, in accounting for their origin, where they have come from and by what means. It appears that the theory of their transportation is the 'age of Drifts'—that this continent was once the bed of the sea and that these boulders were brought from the North Pole by icebergs. This theory has a drift foundation."



Parasites on ants

NEWS AND ANALYSIS



20
SCIENCE
AND THE
CITIZEN



42 PROFILE *Jane Goodall*



24 IN BRIEF
38 ANTI GRAVITY
40 BY THE NUMBERS

46

TECHNOLOGY
AND
BUSINESS



52
CYBER VIEW

IN FOCUS

GROWING A NEW FIELD

Tissue engineering comes into its own

When Betty Shabazz suffered third-degree burns in a fire set by her grandson, doctors covered parts of her body with an artificially manufactured skin product. The widow of Malcolm X ultimately succumbed to her injuries. But the Shabazz case did serve to highlight the promise of tissue engineering: physicians have credited engineered skin with helping others survive severe burns with less extensive skin autografts from a patient's body or without the use of sometimes scarce cadaver skin. The nascent field promises to supply not only replacement skin but cartilage as well—and perhaps, one day, hearts, livers and other complex organs that substitute for transplants.

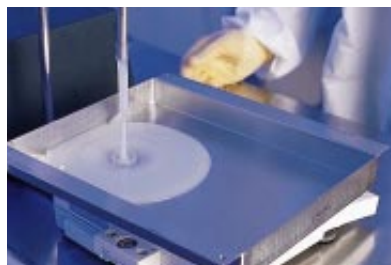
Since last year, the Food and Drug Administration has approved two artificial skin products for third-degree burns and is about to license cartilage replacement for damaged knees. Canadian regulators have given their sanction to a graft for skin ulcers. And U.S. clinical trials are under way for still more products, including cartilage and other engineered skin as well as cells encapsulated in polymers that deliver a nerve growth factor to the spinal columns of patients with amyotrophic lateral sclerosis (Lou Gehrig's disease). "We've moved from important laboratory discoveries in the 1980s to a number of real products," says Robert Langer, a professor of



BETTY SHABAZZ
received grafts of artificial skin but later died of her injuries.

chemical and biomedical engineering at the Massachusetts Institute of Technology who is a leading researcher in the field.

Integra, the artificial skin administered to Shabazz, consists of a porous matrix made of collagen (fibrous connective tissue from a cow) and a derivative of shark cartilage, materials that were tested for human biocompatibility. The size of the pores induces new connective tissue and blood vessels from tissue underneath the dermis (the inner skin layer) to grow into the biodegradable matrix. The manufactured dermis comes with a synthetic silicone covering, a substitute for the epidermis (the top layer). The synthetic must be replaced with a graft of the patient's own epidermis once the inner dermal cells have regenerated and the matrix has largely eroded. The



PETER MURPHY

TISSUE ENGINEERS AT INTEGRA LIFESCIENCES
make a component of artificial skin first by cleaning cow tendon (far left). Then they freeze it (center left), process it with other compounds, pour it for weighing (center right) and finally freeze-dry it into thin sheets (far right).

patient needs only a thin transplant of skin rather than a much thicker and potentially scar-inducing autograft.

Two other companies—Advanced Tissue Sciences (ATS) in La Jolla, Calif., which received FDA approval this spring, and Organogenesis in Canton, Mass., which now has Canadian licensing—grow new skin tissue from cells taken from the foreskins of newborns. The tissue generated then serves as either a temporary covering for burn patients (ATS) or a permanent graft for the treatment of skin ulcers (Organogenesis). At press time, another company, Genzyme Tissue Repair, expected FDA approval for a process called Carticel, which cultures the patient's own cartilage cells in vitro before injecting them into a damaged knee.

On the research front, universities and biotechnology companies have begun to develop concepts for bioengineering kidneys, bone, livers, hearts and—in one much publicized case—a human-shaped ear (implanted onto the back of a mouse). In late July a researcher from Harvard University, Dario Fauza, described how he and Harvard Medical School surgeon Anthony Atala had collaborated to grow replacement sections of organs from the tissue of prenatal lambs. At the conference of the British Association for Pediatric Surgeons, Fauza explained how cells harvested from the lambs were cultured on a polymer scaffolding that assumed the shape of a section of bladder. At birth, the lambs, which had surgically induced bladder malformations, received the contoured replacement bladder tissue. It functioned better than surgical repairs alone in a set of control lambs. Atala has plans in coming months to use a similar form of tissue engineering to rectify bladder abnormalities in children. And someday the method may replicate whole human organs: in the laboratory, Atala has created replacement bladders for adult beagles, a result that he expects to report at a conference of the American Academy of Pediatrics in October.

The promise of such an experiment cannot obscure daunting technical challenges. "People have made nice progress with transplanting cells into matrices," says Jeffrey Hubbell, a professor of biomedical engineering at the Swiss Federal Institute of Technology. "But there is a long way to go even for geometrically simple structures like skin and cartilage." Tissue designers face the difficult task of perfusing a blood supply into more voluminous parts—bone or liver, for instance—than the flat skin tissue. And an organ such as the heart (or even a whole hand or arm, one of tissue engineers' futuristic dreams) will need to be wired with nerve fibers.

A creative approach to the problem of ensuring an adequate vascular network for newly forming tissue came in a report from biomedical engineer Antonios G. Mikos of Rice University and his co-workers in the July issue of the *Journal*

of *Biomedical Materials Research*. Mikos's team took bone-forming cells from the marrow of a rat and transplanted them onto a porous polymer foam before culturing them in an incubator. They then sewed the cell-laden scaffolding into the rat's mesentery, the membrane that holds the intestine together. The bone tissue that grew on the scaffolding hooked up with blood vessels in the well-vascularized mesentery. Ultimately, this technique could serve as a novel means of cultivating new tissue for human bone replacement. The new bone produced, for example, in a vascularized membrane around the rib can be transferred to another site in the patient's body, an alternative to the painful harvesting of existing bone or the use of complication-laden synthetic bone.

Peripheral nerve tissue has drawn the attention of bioengineers because it does not regenerate easily. In rodents, researchers have sutured polymer or collagen tubes to the two severed ends of a disconnected nerve. The precise geometry of the cylinders promotes the reconnection of segments of up to a few centimeters in length. These nerve guidance channels can also be seeded with a type of cell that manages fiber regrowth. Integra LifeSciences, the artificial skin developer, has even begun a clinical trial on a collagen guidance channel in humans.

A nerve channel that can conduct an electric current may improve the growth of new nerve tissue. A report in the August 19 issue of the *Proceedings of the National Academy of Sciences* by an M.I.T.-Harvard team—Robert Langer, Joseph Vacanti, Christine E. Schmidt and Venkatram R. Shastri—demonstrated that a voltage applied through a conductive polymer, polypyrrole, produced an electrical field that induced nerve fibers from a rat to lengthen significantly more than those that did not receive the stimulus. Normally, nerve fibers do not grow well at all on the various polymers used to craft nerve guidance channels. Polypyrrole or other electrically conductive polymers may become candidates in the constant quest for new materials that can be used in tissue engineering. A scaffolding built of the right polymer might be used both to regenerate nerves and to grow other tissue types, a step toward the vision of building entire new limbs.

Prospects for tissue engineering have brightened as government research funding expands. Last spring the National Institutes of Health, for one, began soliciting proposals for a tissue-engineering grants program. Tissue engineering can even become a matter of civic pride. Since 1994 the Pittsburgh Tissue Engineering Initiative has brought together a research collaboration of area hospitals and universities. Discoveries related to this nascent technology, it is hoped, will eventually bring renewed life to the city's industrial base, a goal similar to tissue engineers' vision of reinvigorating an aging population.

—Gary Stix

RESEARCH MODELS

THE NEXT HOP

Can wallabies replace the lab rat?

Bare-handed, Craig Zaidman reaches into the pouch of a female tammar wallaby. At the neurobiologist's touch, this squirming, 18-inch-high cousin of the kangaroo becomes as docile as a milk cow, possibly because the hand feels like a young joey crawling in. After Zaidman separates the pouch entrance from the surrounding gray-brown fur, he plucks out a hairless, finger-length "pouch young" from a teat; it comes away from the nipple like a grape off a vine.

"This is what makes wallabies so great for study. It takes virtually no effort to hold what is essentially an embryo in the palm of my hand," says Zaidman, a visiting Fulbright scholar to the Australian National University (ANU) Research School of Biological

Because of this obstacle, researchers are usually forced to dissect dead specimens, examine cells in petri dishes or study such nonmammals as frogs. Some neurobiologists in Italy were able to make electrical recordings of embryonic brain activity in live rats, but the effort proved so difficult that no one has yet repeated the feat, Zaidman says. But by studying the wallaby (*Macropus eugenii*), scientists can make recordings in a live, intact animal well before visual activity begins, says Richard Mark, founder of the ANU's program.

Like all marsupials, wallabies are mammals; they have hair, produce milk and are warm-blooded. But unlike the rest of the mammal class, marsupials do not nourish their young in a placenta from conception to delivery. Instead their partially developed young spend only 28 days in the womb before crawling, sluglike, to the marsupium, or pouch, outside the mother's body. There they take another 180 days to suckle, differentiate and grow into fully formed joeys.

Meanwhile these developing pouch young are basically free-living, readily

establish a colony. "You can't just call up a biological supply house and say, 'I'd like 100 wallabies,'" points out Peter Janssens, co-author of *The Developing Marsupial: Models for Biomedical Research*. He adds that a lot of work had to be done on simple care and feeding, as well as on the applicability of wallabies to other mammals. "It took 15 years of background work before we could even get results," Mark adds.

Simply collecting the first animals was an adventure. The team had to improvise tools to capture the fast-hopping wallabies from an island off South Australia, where their numbers had become unnaturally high. The first nets often snapped from the force of the speeding marsupials. (They now use modified, oversized butterfly nets.)

There was also the obstacle of overcoming the bias against the use of marsupials as lab animals. Early 19th-century taxonomists thought Australian marsupials were a more primitive subcategory of mammals because they lacked a corpus callosum—the brain formation that enables the two hemispheres to communicate. It took decades before scientists discovered that marsupials did indeed have an equivalent structure, called the fasciculus aberrans. Other aspects of the marsupial brain also later proved to be similar to typical mammalian brains. "Contrary to early taxonomists, wallabies are not second-class mammals," Mark says, adding that "it's not the differences between wallabies and other mammals that make wallabies so interesting as a research model; it's the things that make them the same."

Wallaby studies have already paid dividends: by using these animals, Mark and his team found that optic axons do not randomly form connections with the superior colliculus, as previously thought. Instead axons target specific spots. Other workers in several research centers throughout Australia now use marsupials as lab animals, and in the U.S. the wallaby's South American cousin *Monodelphis domestica* (the gray, short-tailed opossum) has occasionally been imported for study. University of Melbourne's Marilyn Renfree, who has spent 30 years studying wallaby reproduction and development, sees this interest as long overdue. "But then I'm a marsupial chauvinist," she says.

—Dan Drollette in Canberra, Australia



DAN DROLLETTE

LIVE, 55-DAY-OLD WALLABY EMBRYO

taken from the pouch may be the ideal model in mammalian neurobiology.

Sciences in Canberra. "This one can be returned to the pouch, alive and well, for further monitoring," he adds, before weighing the rugged 55-day-old for his inquiry into how the developing eyeball makes connections with the brain.

Neurobiologists who use more traditional laboratory animals only dream of such easy access. Most of the brain's "hardwiring" occurs early in embryonic development, when access is difficult. By the time the young of popular lab animals such as rats or cats are available, their brains are already past the crucial stage when the onset of visual activity occurs.

accessible fetuses. Neither surgery nor anesthesia is required to get them, which eliminates a potential source of error.

An additional bonus is that maturation happens slowly inside the pouch; a developmental activity that takes 24 hours in rats takes three weeks in a wallaby. The drawn-out pace means that sequential events can be viewed distinctly in the embryonic brain: for instance, optic axons can be easily tracked as they extend from the back of the eyeball into the superior colliculus—the part of the brain controlling eye movement.

But for Mark and his team to even make such studies, they first had to es-

IN BRIEF

Hot Deals

It's the first agreement of its kind in the U.S.: Diversa Corporation in San Diego recently made a five-year "bioprospect-

ing" deal with Yellowstone National Park. The contract lets Diversa delve into the park's hot springs, geysers, fumaroles and boiling mud pots for extremophiles—microorganisms that live under extreme conditions and that, Diversa hopes, may make enzymes of commercial value. Scientists have identified fewer

than 1 percent of the fauna that thrive in the park's 10,000 thermal sites. Diversa gets the rights to any discoveries and products from them, and the park shares in the knowledge and royalties.

No Joking, Mr. Feynman

Physicists have at last seen—and heard—a phenomenon forecast long ago by Brian Josephson and the late Richard Feynman, among others. Josephson won the 1973 Nobel Prize in Physics for predicting what happens when a thin insulator joins two superconductors: the particles in each begin to oscillate back and forth. Now, James C. Davis and Richard Packer of the University of California at Berkeley have shown that when two containers of superfluid helium 3 are separated by a microscopic hole, the quantum liquid, which can flow without resistance, exhibits the same quirky trait. They report that the vibration of the particles, amplified more than a billion billion times, produced a high-pitched whistle.

Hey Diddley Ho, Neighbor

It may not be so surprising, but now it's official: People who trust the folks next door enjoy lower rates of violent crime. As part of the Project on Human Development in Chicago, researchers led by R. J. Sampson of the University of Chicago interviewed 8,872 residents in 343 city neighborhoods. In areas where families were willing to intervene on behalf of the common good, crime was far less frequent. In addition, the survey showed that social cohesion among neighbors was more effective at curbing crime than organized watches and other local services.

More "In Brief" on page 28

ECOLOGY

FIELD AND STREAM

A new way to identify the inhabitants of an ecosystem

Life is tough in the tundra. Most of the year snow covers the ground, and during summer the permafrost keeps many nutrients frozen below ground, unavailable to plants and animals. Without much to go around, few species thrive—making the tundra a relatively simple ecosystem. Which also makes it an ideal study site for researchers to tease apart some of the ecological processes that would be too dizzying to decipher in other, more diverse places.

By examining Arctic lakes and streams, Anne E. Hershey, Gretchen Gettel and their colleagues at the University of Minnesota appear to have uncovered a new way of determining species composition in an ecosystem. The idea, dubbed "a geomorphic-trophic hypothesis," could apply to other ecosystems. And it could eventually permit researchers to use remote sensing—aerial photography and radar—to determine species makeup, a potentially valuable tool for conservation.

The hypothesis brings together two fundamental ways of looking at ecosystems: who eats what, and how the physical terrain constrains the resident creatures. After years of studying aquatic food webs around the Toolik Field Station (a 22-year-old Arctic research site situated about 130 miles south of Prud-

hoe Bay and run by the University of Alaska-Fairbanks), Hershey and others mapped out how six species of fish set the stage for the entire biological composition of Arctic lakes, ponds and streams. Because fish are top predators, they control the zooplankton and the rest of the biota, explains Hershey, so "if we know what fish are present, we know what else is present."

The researchers then combined this trophic knowledge with geomorphic data: the physical characteristics of water bodies, including the gradient of the outflow from a lake, as well as the depth and area of the lake and connections to other lakes. Such features determine which species of fish are present. Trout, for example, cannot swim up steep slopes into high-gradient lakes, whereas grayling can navigate smaller waterfalls and steeper inclines.

Taken together, these approaches form the basis of the geomorphic-trophic hypothesis. The team found that lakes and ponds with very steep gradients have a diverse invertebrate community and no fish; those of moderate depth, and with somewhat gentler slopes, contain grayling, which eat the large invertebrates; lower gradient, deep lakes have trout, sculpin and grayling. In principle, such a complete picture of every organism could even come from satellite pictures and maps, which can measure lake depth and stream gradient. "The idea that these two forces interact to have an influence on the food web is a breakthrough," comments Gary A. Lamberti of the University of Notre Dame. "If [Hershey] can demonstrate that up in Alaska, the idea will catch fire."

First, though, the researchers had to



SAMPLING FOR INHABITANTS IN ALASKAN WATERS
this past July helped to confirm a new method for determining ecosystem makeup.

KEITH GUNNAR
Bruce Coleman Inc.

MARGUERITE HOLLOWAY

In Brief, continued from page 24

Monkeys Do, Scientists See

Schizophrenia has long been one of the most puzzling psychiatric conditions, but neurologists have a new model for studying the disorder. Robert Roth and his colleagues at Yale University recently reported that monkeys treated with phencyclidine (PCP) display the same immediate and long-term dysfunction as schizophrenic humans do. In particular, repeated PCP treatments rendered the prefrontal cortex less able to utilize the neurotransmitter dopamine. Giving the monkeys clozapine—a medication used to treat schizophrenia—improved their cognitive abilities.

Fat Tax

"Extra value meals" might become a thing of the past if Kelly Brownell, director of the Yale Center for Eating and Weight Disorders, has his way. Brownell wants to slap a tax on all fatty foods. He

notes that over the past 15 years, the prevalence of obesity has risen an alarming 25 percent in the U.S. Rather than

blame less-than-diligent dieters, Brownell targets a "toxic food environment," in which 7 percent of Americans eat at McDonald's on any given day, and the average child sees 10,000 food commercials on television a year. A fat tax, he adds, could subsidize more healthful foods and public exercise programs.

"Immortality" Gene Revealed

Two teams of scientists—from Geron Corporation, the University of Colorado at Boulder and the Whitehead Institute for Biomedical Research, among others—have cloned the gene for the human telomerase catalytic protein, the "holy grail" of aging research. This enzyme serves as a key of sorts for rewinding the cellular clock: cells that produce telomerase, such as cancer cells, are immortal. Those that lack the enzyme have a limited life span. The researchers hope that by having identified the enzyme, they will be able to screen for drugs that can inhibit or activate it. Inhibitors might prove to be highly specific and potent anticancer agents, whereas activators may well ameliorate diseases caused by cell death, including Alzheimer's.

More "In Brief" on page 32

explain one mystery: why certain fish appeared in places they shouldn't. For instance, trout were observed in some high-gradient lakes. Earlier this year Hershey conferred with a geologist and began to incorporate paleogeology into her lake profiles. The two found that "stream piracy" had occurred after the last glaciation. Lakes that in ancient times drained on a gentle slope in one direction would have permitted fish access. Over time, though, the lake may have broken through its banks to drain, say, down a steep slope into a different watershed. That event would have iso-

lated the trout in high-gradient lakes.

This past July, on a hot, sunny morning that gave way to a gray torrential downpour by late afternoon, a dozen or so biologists set out to see whether all the elements of the hypothesis, ancient and current, held together. Dropped by helicopter near a series of lakes, they spent the day carrying lightweight boats from one body of water to the next and sampling just about everything—fish, water, microorganisms and algae. It looked good. Every fish present was accounted for.

—Marguerite Holloway in Alaska

BIOCHEMISTRY

GOTTA KNOW WHEN TO FOLD 'EM

A scientific wager reveals details about how proteins fold

Thanks to some betting biochemists, proteins now belong right up there with poker and ponies. This past summer a team from Yale University collected on a \$1,000 bet that a certain type of protein couldn't be made. In addition to pocketing the cash, the Yale researchers also learned more about the way proteins work—knowledge that could one day improve our understanding of Alzheimer's and Creutzfeldt-Jakob disease.

The chains of amino acids that make up proteins exist as elaborate, three-dimensional structures, including combinations of corkscrewlike coils known as alpha-helices or extended flat surfaces called beta-sheets. Exactly how a se-

quence of amino acids assembles into its final conformation—called the protein-folding problem—is a topic of intense study. But researchers were sure of at least one thing: if two proteins have as little as 30 percent of their amino acid sequences in common, their structures would be very similar. In other words, making them different would require at least a 70 percent change in amino acid sequence.

Confident of this view, in 1994 George D. Rose of Johns Hopkins University and Trevor Creamer, now at the University of Kentucky, laid down a \$1,000 challenge in the journal *Proteins: Structure, Function, and Genetics*: take one protein structure (say, a beta-sheet) and transform it into another (say, an alpha-helix) by replacing no more than half the amino acids. According to Rose, "we thought it could not be done."

Enter Lynne Regan and her colleagues at Yale, Seema Dalal and Suganthi Balasubramanian. Last year, while chatting in the car on the way home from a conference, Regan suggested that two proteins being studied in the lab might just lend themselves to the modern-day alchemy required to win the wager.

The two proteins, called B-1 and Rop, had been looked at extensively in Regan's lab in an effort to understand how the proteins fold into their three-dimensional configurations. The protein B-1 is predominantly a beta-sheet, and Rop consists of several alpha-helices. Regan's group had been able to determine which amino acids in each protein controlled the formation of either a beta-sheet or an alpha-helix.



LYNNE REGAN

MODERN-DAY ALCHEMY
converts a beta-sheet protein (left)
to an alpha-helical structure (right).

In Brief, continued from page 28

Still Cloning Around

Scientists at ABS Global in Wisconsin have recently dispelled any lingering doubts about Dolly, the lamb cloned last spring by Keith Campbell of PPL Therapeutics and Ian Wilmut of the Roslin Institute in Scotland. The U.S.



team copied the earlier experiment and also copied Holstein cows (photograph). In the meantime Dolly's creators have made another lamb that has a human gene in each cell. Unlike Dolly, Polly, as the Poll Dorset newborn has been named, was cloned from skin cells, using a technique that appears to have many advantages over traditional genetic engineering. In particular, the method allows removal of genes from a cell. Thus, this type of cloning could be ideal for generating transgenic transplants; humans would most likely tolerate organs harvested from pigs cloned from cells that have had genes encoding rejection-causing proteins removed.

MOHREY GASH/AP PHOTO

Sun Sweat

Water on the sun? Peter F. Bernath of the University of Waterloo and his colleagues first suggested so in 1995, when they observed sunspot spectra resembling those from ordinary water molecules. It was possible. Although the sun's surface blazes at some 5,000 degrees Celsius, sunspots are generally 2,000 degrees cooler, which might permit water vapor to exist. For proof, the astronomers needed to calculate the wavelength patterns that H₂O molecules would emit at scorching temperatures. It hasn't proved a simple problem, requiring serious number crunching on a supercomputer. But now, two years later, their solutions exactly match their empirical data. The results should help scientists make better models of sundry planetary atmospheres. And closer to home, the finding may help satellites spot budding forest fires: burning trees probably release water molecules with similar chemical signatures.

—Kristin Leutwyler

So last summer the group experimented with the two structures, first on computer models, then on the real thing. The researchers removed small stretches of amino acids from B-1 that contributed to the formation of beta-sheets and replaced them with segments from Rop that could lead to alpha-helix formation. The result, published in the July issue of *Nature Structural Biology*, is a new protein. Janus, named for the two-headed Roman god, retains half of the amino acid sequence of B-1 but has the helical structure of Rop. In more recent work, the team created Janus II, which carries 61 percent of the B-1 sequence, meaning that the researchers had to substitute only 39 percent of the original amino acids.

One message of this work is "don't

treat all amino acids equally," according to Regan. Only certain amino acids actually dictate how the protein will fold into its final configuration, she says. Better knowledge of this specificity may eventually improve scientists' understanding of certain so-called protein-folding diseases. In conditions such as Alzheimer's or Creutzfeldt-Jakob disease (the human form of "mad cow" disease), researchers theorize that spontaneous alterations to a protein's structure can lead to the neural degeneration characteristic of these maladies.

In the meantime, Regan's group is still deciding what to do with the money. Rose, for his part, is pleased with the findings but laments taking such an expensive gamble: "Would that it had been a T-shirt." —Sasha Nemecek

FIELD NOTES

SCIENCE IN COURT

Reflections on science and truth in an asbestos trial

When I got the call I was startled, curious and perversely pleased that an editor at *Scientific American* had been selected as a juror in an asbestos trial. For the next five weeks, I spent my days inside the imposing New York State Supreme Court building, hearing testimony in the case of *Vincent Cangiane*

v. Westinghouse Electric and watching scientific evidence emerge bent, muffled, truncated—and ultimately, I hope, triumphant—in a high-stakes civil suit.

A few basic facts were undisputed. Asbestos exposure, especially with cigarette smoking, can cause lung cancer. The 64-year-old Cangiane was a heavy smoker but gave up cigarettes in 1967. Nevertheless, in 1993 and again in 1996 he developed cancers in his left lung. The key points of contention: Could Cangiane have been exposed to asbestos as a result of his work repairing subway cars that contained electrical components sold by Westinghouse? If so, did the asbestos contribute to his lung cancers?

Answering these questions seemed a straightforward matter of scientific investigation. But the courtroom is not a laboratory; we jurors know only what the lawyers and their witnesses are willing or able to show us. Under these circumstances, testing a hypothesis often becomes an exercise in reading facial expressions and inferring the subtext of the lawyers' questions.

Fibers and the associated asbestos bodies are few and far between even in someone who has had moderately severe asbestos exposure. And, in fact, none of the medical experts could find either of these in Cangiane's lungs. Years of fiber inhalation can also produce a scarring of the lung called asbestosis. But mild asbestosis appears as an almost imperceptible haziness on a chest x-ray,



JASON GOLTZ

CONEY ISLAND FACILITY
is where the plaintiff once repaired subway cars and claims he was exposed to asbestos.

undermining the defense argument that an absence of x-ray markings means an absence of asbestos exposure. One of the plaintiff's witnesses, Emanuel Rubin of Thomas Jefferson University, smartly dismissed the value of x-rays with a quip: "I don't believe in those shadows."

Then there was the matter of the asbestos source itself. Could an asbestos-impregnated arc chute (a molded sleeve that blocks electrical sparks from a high-voltage contact) release respirable fibers? Surely a simple bench test would tell. Only we learned of no such test; we had to rely on 25-year-old memories of job practices as recalled by witnesses who worked for Westinghouse and the New York City Transit Authority.

In the end we needed information from outside populations to put the medical evidence in perspective. Cancer risk from tobacco declines with time after a smoker quits; cancer risk from asbestos, in contrast, seems to peak many years after the initial exposure. In the absence of concrete proof, the statistical considerations proved critical, tipping the case to the plaintiff's side.

Since Galileo, quantification has been a hallmark of scientific method. But Galileo was timing balls rolling down inclined planes; we now had to determine the monetary value of a traumatized and shortened human life. It took a few hours of delicate, sometimes tense negotiation to reach a consensus number. Even then, several of us felt uneasy

as we considered the implications of multiple layers of conclusions based on a "preponderance of the evidence," in which 51 percent certainty is good enough.

One mystery remained: How *did* I end up on this jury? After the trial, I asked Jim Long, the lead plaintiff lawyer. "We ran out of challenges," he confessed with a relieved laugh. "One more, and you would have been off." I fleetingly considered how, in justice as in nature, small initial variations can lead to wildly disparate outcomes. One different juror, one different witness, and the outcome of the trial might well have changed. I reverted to the faith of a rationalist: truth somehow emerges from the chaos.

—Corey S. Powell

SCIENCE EDUCATION

WHAT ARE THEY THINKING?

Students' reasons for rejecting evolution go beyond the Bible

These days even the Pope will tell you that biological evolution is "more than a hypothesis," but nearly half of Americans still beg to differ. Poll after poll shows a country almost equally divided between those who accept and those who reject the theory that all the earth's flora and fauna descended from a common ancestor (in contrast, the scientific community has no doubts). In a country where the overwhelming majority professes some degree of religious faith, it

might seem logical to assume that those who discount evolution have simply taken the divine word over Darwin's. Harvard University researcher Brian J. Alters thinks there is more to it.

A veteran science educator, Alters has long sought to understand why so many students complete high school without coming to comprehend and accept one of biology's central tenets. Alters is particularly interested in pinpointing any nonreligious rationales. These, he argues, could appropriately be addressed in a public school setting.

With educational psychologist William B. Michael of the University of Southern California, Alters conducted interviews and administered surveys to pick the brains of more than 1,200 college freshmen at 10 different schools. In this unpublished study, he found that those who reject evolution (approxi-

mately 45 percent) tend more than their counterparts to hold specific misconceptions about evolutionary science. They are more likely to agree with statements such as "mutations are never beneficial to animals" and "the methods used to determine the age of fossils and rocks are not accurate." Indeed, nearly 40 percent of those skeptical of evolution believe the chance origin of life to be a statistical impossibility.

Having identified these and other erroneous beliefs, Alters says, the next step is to develop a curriculum that addresses them head-on. Although "the purpose of public school education is not to change people's religious beliefs," he notes, students' preconceptions about genetics, radiometric dating and statistical probability are certainly fair game.

Philip M. Sadler, the director of science education at the Harvard-Smithsonian Center for Astrophysics, has reviewed Alters's data and agrees that the type of curriculum that Alters envisions is crucial to the teaching of evolution and to science in general. Sadler concludes that for children "the process of learning science is a process of abandoning their own previous views." Until misconceptions are countered with specific evidence (a good explanation of how fossils are dated, say), "the ideas simply will not change," Sadler says.

Some physicists have begun to implement curricula that first address preconceptions, subsequently enabling students to "fly through" physics courses, Sadler comments. Perhaps with a similar approach in biology, educators could help students' understanding of Darwinism evolve as well.

—Rebecca Zacks



ELIZABETH CREWS/ImageWorks

MANY HIGH SCHOOL SCIENCE CLASSES
fail to correct misconceptions about the facts and methods of evolutionary biology.

MATTER OVER MIND

Do viruses cause severe mental illness?

Despite the enormous human and economic toll of schizophrenia and other psychoses, medical science has yet to provide a compelling account of what causes these mind-robbing disorders. Geneticists have found indications that heredity may play a part. But most researchers think other causes must be involved as well, mainly because when one member of a pair of identical twins has a psychotic illness, the other twin's chances of developing a similar affliction are very far from a sure thing.

One controversial theory, accepted still by only a minority of investigators, posits that an unrecognized infection by a virus or other agent might trigger at least some cases of schizophrenia or other psychoses. Several times over the past 20 years, researchers have reported that medicines used to treat schizophrenia or bipolar (manic-depressive) disorder may have antimicrobial effects. Moreover, physicians have occasionally noted that giving such drugs to a patient seemed to have a beneficial effect on a recognized viral infection. A recent study published in *Schizophrenia Research* puts these casual observations on a somewhat firmer footing.

Metabolic by-products of the antipsychotic drug clozapine, it turns out, inhibit the growth of HIV, the AIDS virus, in a standard cell-culture system. Although HIV does not cause schizophrenia or bipolar disorder, champions of the viral-causation theory note that other viruses might be similarly affected by antipsychotic medicines. Conceivably, they suggest, clozapine and some other antipsychotic drugs whose mode of action is uncertain might work by suppressing an unknown virus. "We believe this effect is not random," says Lorraine V. Jones-Brando of the Stanley Laboratory for the Study of Schizophrenia and Bipolar Disease at Johns Hopkins University, the lead author of the study. The new study does not mean that clozapine might become an anti-HIV drug, however: indications suggest existing therapies are better.

The most obvious objection to the viral schizophrenia theory is that nobody

ANTI GRAVITY

He Shoots, He Scars

The marathon known as the National Hockey League regular season is about to begin. Hundreds of robust young warriors will soon find themselves, at one time or another, writhing in agony. A recent report in the *American Journal of Sports Medicine*, "Predictors of Injury in Ice Hockey Players," notes that "injuries are attributed to collisions with players skating at speeds up to 30 mph, pucks traveling at 100 mph, sharp skates, and long sticks." Well, put *Lord of the Flies* on ice, and, yes, people are going to get hurt.

Sport entails risk. The collisions common to hockey and other contact sports often cause the temporary brain-scrambling known as concussion. A recent review in *Medicine & Science in Sports & Exercise* with the coy title "Were You Knocked Out?" provides a summary of concussion management. It includes a list of questions to be asked as a "post-concussion memory assessment," to help determine a player's wooziness coefficient. This list includes "Which team are we playing today?" and "How far into the quarter is it?" As a rule of dislocated thumb, trainers should note that a concussed New Yorker who responds to any question with "Who wants to know?" is totally coherent.

Speaking of concussions, boxers are obviously at great risk for becoming unconscious. The infamous Mike Tyson-Evander Holyfield rematch showed that boxing's risks now include rabies. Tyson, who felt he had been wronged by a Holyfield head butt, was perfectly free to take revenge by pummeling Holyfield in the face. Other sports discourage this form of retaliation, but in boxing, heck, it's the whole point. Tyson instead decided to attempt to bite off Holyfield's ears. Because repeated concussions can cause long-term brain damage, the possibility exists that any prior incidents may have taken their toll on Iron Mike's iron head.

Speaking of irons, even pastoral sports such as golf have their risks, some of which likewise include sticking things in your mouth. The journal *Gut* has reported that a 65-year-old retiree who golfed daily came down with hepatitis. Doctors searching for the

cause discovered that he licked his balls before putting. This habit exposed the golfer to Agent Orange, a pesticide used on the course, and made him the first proved victim of—deep breath now—Golf War Syndrome.

Lousy golfers face other hazards. A study published a couple of years back in the *New England Journal of Medicine* found that bad players in a Tennessee retirement community were more likely to get the tick-borne disease ehrlichiosis. Presumably, they spend more time in tick-ridden woods and high grass looking for errant tee shots. "What's your handicap, Arnie?" "Why, the fever and muscle aches, Jack!" (This reporter recently played a round of golf in which, for the first time, he didn't lose a single ball. Perhaps still



impaired from a baseball concussion some quarter of a century ago, however, he did finish minus a sand wedge.)

Golf is for the faint of heart compared with the rough-and-tumble action reported in a *Journal of the Royal Society of Medicine* article, "A Survey of Croquet Injuries." Although wrist, hand or forearm problems were not uncommon, croquet also leads to more serious harm. "Falling as a result of standing on a ball had the worst effects," the researcher notes. One player broke a foot bone "putting on a Wellington boot"; another "suffered a black eye from being struck on the head by a mallet."

The difference then between croquet and boxing? Mishaps of the Three Stooges variety in croquet are accidental. Tyson earned the sobriquet "Madman!" from *Sports Illustrated* for biting Holyfield. For administering a concussion, on the other hand, he would have been called "Champion!" Go figure.

—Steve Mirsky

MICHAEL CRAWFORD

has yet found a virus to fit the bill. On the other hand, notes E. Fuller Torrey of St. Elizabeth's Hospital in Washington, D.C., a longtime champion of the theory and a collaborator of Jones-Brando's, "almost nobody has looked" in psychotic patients for viruses other than the well-known types. "My own feeling is that if there's a virus it won't be one of the easily recognizable ones," says Robert H. Yolken of Johns Hopkins, who also worked on the HIV-clo-

zapine study. "The geneticists have not found a gene yet either, and we feel the same way about viruses." Yolken says he has been impressed by how many psychotic patients say their illness developed after signs of a viral infection.

A virus link no longer seems as outlandish as it once did: within the past five years Liv Bode of the Robert Koch Institute in Berlin has demonstrated that a virus originally found in horses, Borna virus, can cause depression or mood

disorders in humans. Yolken has failed so far to find evidence of Borna virus among patients with depression or psychosis. Still, some kind of virus-psychosis link is "becoming remarkably respectable," Torrey says. He and his associates are planning a study in which they would treat psychotic patients with antiviral drugs, probably anti-HIV protease inhibitors, to see whether they might somehow soothe tortured minds.

—Tim Beardsley in Washington, D.C.

BY THE NUMBERS

Chronic Obstructive Pulmonary Disease

Chronic obstructive pulmonary disease (COPD) is the term applied to several related conditions, of which the most serious are emphysema and chronic obstructive bronchitis. In emphysema the alveoli—the terminal sacs of the lung at which oxygen and carbon dioxide are exchanged with the blood—become permanently enlarged. In chronic obstructive bronchitis, which usually occurs with emphysema, the trachea and bronchial tubes become irreversibly inflamed, restricting airflow. Two other conditions often labeled as COPD have a better prognosis: simple chronic bronchitis with normal airflow and asthmatic bronchitis. Simple asthma, which is caused by hypersensitivity to allergens and other stimuli, is reversible and is not included in the definition of COPD.

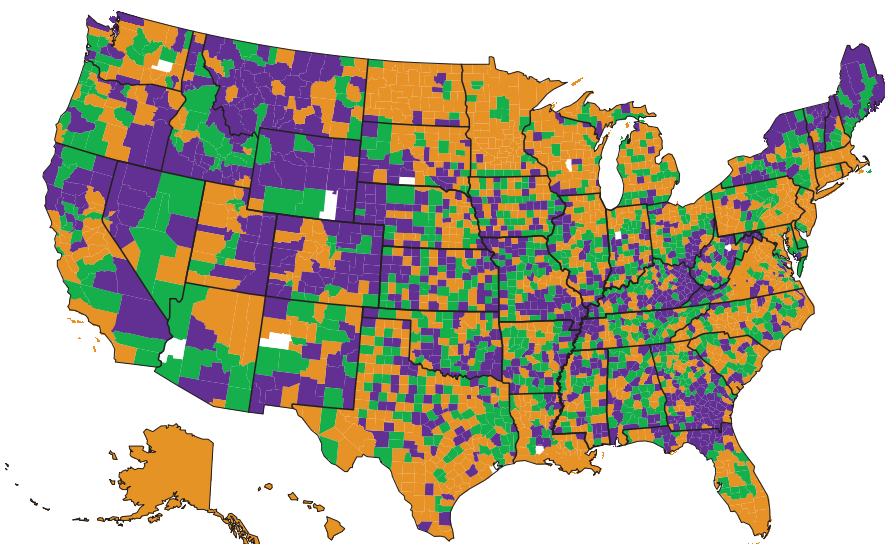
The chief symptoms of COPD are coughing, wheezing, expectoration and labored breathing. Unlike lung cancer, which kills its victims relatively quickly, COPD progresses slowly, gradually reducing the ability to breathe. Like lung cancer, it is caused primarily by cigarette smoking. Passive smoking and occupational exposure to dust and fumes play a part, and dust and sulfur dioxides outside the workplace may also be risk factors. In the normal healthy non-smoker, lung capacity gradually declines with age, but in those with COPD, capacity declines more rapidly, particularly among heavy smokers. Those who give up smoking do not regain lost lung capacity, but the rate of decline in capacity slows to that of nonsmokers. The prognosis in patients with mild airway obstruction is good, but for those with severe obstruction the prognosis is poor, particularly if the blood level of carbon dioxide is high. In most cases, death from COPD is precipitated by acute respiratory disease such as pneumonia or by other complications such as cardiac arrhythmia or pulmonary embolism.

About two million Americans have emphysema, and another 14 million have some form of

chronic bronchitis. About 105,000 died of COPD in 1996, making it the fourth leading cause of death in the U.S. after coronary heart disease, stroke and lung cancer. Nineteen out of 20 of those dying of COPD are 55 or older. Men are more likely to die of the disease than women.

The reasons for the regional differences in mortality are not clear, but it may be no accident that deaths from COPD and lung cancer are greater in the Southeast, where smoking is historically high. COPD mortality, unlike that of lung cancer, tends to increase with altitude, as illustrated by the high mortality rates in the mountain states. Altitude as a disease contributor has not been established but is biologically plausible. Those living in Denver, for example, get 15 percent less oxygen in the same volume of air as those living in a sea-level city such as Miami and so, if they have developed COPD, could be at higher risk of death. Poverty may also influence the pattern on the map: one of the highest concentrations of COPD is in eastern Kentucky, where poverty rates among whites are particularly high.

—Rodger Doyle



DEATHS PER 100,000 WHITE MALES 55 AND OVER, 1979–1994 (AGE-ADJUSTED)

ORANGE UNDER 220 GREEN 220 TO 259 PURPLE 260 AND OVER WHITE NO DATA

SOURCE: National Center for Health Statistics. County data for Alaska not available.

RODGER DOYLE

PROFILE: JANE GOODALL

Gombe's Famous Primate

She is standing on the porch of a wooden house in Washington, D.C., just under the thick branch of a tree and just to the side of a tangle of creepers that gives the carefully kept urban backyard a hint of the unkempt, of the vegetative wild, when she does it again. A loud, breathy, nonhuman crescendo silences the garden-party goers and the Goodall groupies, some of whom have driven hours to see her. It is the chimpanzee pant-hoot call, and it has become one of Jane Goodall's signatures. She punctuates most of her speeches and lectures with the wild cry, bringing Tanzanian forests to audiences who have never set foot in Africa and, at least for a few moments, eliminating whatever distinction her listeners were drawing between the scientist and her subjects.

Even as she makes the eerie sound—which is used to establish contact between far-flung members of a troop—Goodall manages to seem completely still. Thirty or so years of sitting quietly, observing the chimpanzees at the Gombe Stream Research Center, have left their mark. Goodall moves without seeming to move; she laughs and turns and gestures while giving the impression of utter calm and stasis. Which is something Goodall has needed a lot of in her dealings with people as well. Renowned and revered today, Goodall's approach to primatology was anything but standard when she started her work. Now that the researcher has moved out of the forest and onto the road, advocating for animal rights and raising money for chimpanzee sanctuaries, she has again met with controversy.

None of that conflict is in the air in this sloping, sunlit garden. Carrying copies of her books, including *In the Shadow of Man* and *Through a Window*, members of the rapt audience listen to Goodall review some of what she has learned about wild chimpanzees. The simian characters—Flo, Flint, Fifi, Pom, Passion—are as familiar to many as family or old acquaintances. Goodall talks about the importance of mothering styles in shaping chimp development, about how a four-year mother-daughter killing spree eliminated all but one newborn chimp and about how it

was Louis Leakey who pointed out that chimpanzees, with whom we share 98 percent genetic homology, provide a window into our distant past.

It was, of course, Leakey who sent Goodall out to peer through that frame. It is a famous story by now. Goodall, who was born in London in 1934 and who was always obsessed with animals and with stories of Dr. Doolittle, worked as a waitress and a secretary to raise enough money to get to Africa. Once in Kenya, Goodall called Leakey to say she wanted to work with animals. After informally testing her knowledge of wildlife during a tour of a game reserve, he took her on as assistant secretary and then, in 1960, sent her, untrained, into the field to observe chimpanzees.

Leakey's plan was to find young women—whom he felt would be patient observers and perhaps less threatening to their male subjects than men would be—to study each of the great apes. The other "trimates," Dian Fossey, who studied gorillas, and Birute Galdikas, who studies orangutans, followed soon behind Goodall. The legacy of the legendary paleontologist and his protégés has been far-reaching: primatology is one of the few scientific fields that has equal numbers of men and women. "Jane Goodall has had a profound effect as a role model. Thirty years ago she showed that it was okay for a woman to live in the jungle and watch wild animals," explains Meredith Small, an anthropolo-

gist at Cornell University. "I have several young women every year coming into my office telling me that they want to become an animal behaviorist like Goodall. She opened the door for women who dream of doing fieldwork."

Goodall herself initially went into the field accompanied by her mother, Vanne, because the remote forest on the banks of Lake Tanganyika was considered unsafe for an unescorted young woman. The chimps eluded Goodall at first, but months of patience paid off when she observed two previously unrecorded activities: meat eating and the use of long grass as a tool to pluck termites from a mound. By consistently following the apes, Goodall was able to observe their various interactions and to piece together the social structure of her group. She described strong and not so strong mother-infant bonds, sibling loyalty and rivalry, male displays and attacks and dominance, and sexual behavior—all in terms of individuals with humanlike personalities. Flo was a wonderful mother and a very sexually attractive female; her son, Flint, was overly attached and died of grief shortly after his mother died; Passion was cold-hearted, killing and eating the offspring of other females.

Such personal descriptions were not standard fare. "One of the things that was happening in primatology and in evolutionary biology in general as Jane was beginning to influence the field was that people were just beginning to look



MICHAEL NUEGBAUER/The Jane Goodall Institute

PANT HOOTS bring together family and friends.

at individuals. She was already doing that as a matter of temperament,” notes Sarah Blaffer Hrdy, an anthropologist at the University of California at Davis. “She was unabashed in her willingness to anthropomorphize and to allow her emotions to inform what she saw the animals doing.”

“In 1960 I shouldn’t have given the chimps names,” Goodall sardonically recalls, fingering the bone Maori talisman she wears as a necklace. “They didn’t have personalities, only humans did. I couldn’t have studied the chimp mind, because only humans had minds.” She goes on to explain in a voice simultaneously soft, hard, strong, calm and passionate that her first paper for *Nature* came back with the words “he” and “she” changed to “it.” “How they would even want to deprive them of their gender I can’t imagine. But that is what it was, animals were ‘it.’ Makes it a lot easier to torture them if they are an ‘it.’ Sometimes I wonder if the Nazis during the Holocaust referred to their prisoners as ‘its.’”

Goodall has written that missing a background in science allowed her to view animals in more human terms. Rather than thinking of them as other, she thought of stages of life and of emotion—childhood, adolescence, grief, attachment, rage, play—and because of that, saw animal behavior in new terms. Yet her lack of education could have been a liability as she tried to get her discoveries out into the world, and so Leakey arranged for her to study ethology at the University of Cambridge. Goodall received her doctorate in 1965, the same year that *National Geographic* introduced “Miss Goodall and the Wild Chimpanzees” to the world.

Fame and scientific imprimatur secure, Goodall continued her work at Gombe, training a stream of students. As the camp grew in size, however, so did the number of interactions between subjects and researchers. Some of the field observations have been criticized as difficult to interpret, such as fights for food. “By changing the environment and feeding them bananas, it skewed results,” maintains Robert Sussman of Washington University. “You can’t tease apart the effect of humans.”

Goodall regrets banana feeding—particularly as it made Leakey skeptical of all her subsequent observations—but she is neither sorry about intervening during a polio epidemic among the chimps nor sorry about threatening Passion and her daughter with a stick so Little Bee could escape with her newborn baby. “I wasn’t a scientist. I didn’t want to be a



CHRIS STEELE-PERKIN/MAGNUM PHOTOS

“I didn’t want to be a scientist,” Jane Goodall says. “I wanted to learn about chimpanzees.”

scientist, I wanted to learn about chimpanzees,” she says emphatically. “So there was this huge outcry: ‘You know you are interfering with nature!’ But, on the other side, there were all these scientists going out and shooting lots of their study population to examine their stomach contents. Is that not interfering with nature? It is so illogical.”

Part of her current work, she explains, is to talk to students about science, to correct the misapprehension that science has to be dispassionate. “I am often asked to talk about the softer kind of science as a way of bringing children back into realizing that it is not all about chopping things up and being totally objective and cold.”

Goodall describes this educational effort as her fourth phase of life. The first entailed preparation: reading and dreaming about getting to Africa. “Phase two was probably the most wonderful I will ever have in my life. I was so lucky I spent all of this time in paradise with the most fascinating animals you can possibly imagine.” Phase three was get-

ting the work into the scientific community. And her current stage came to her, she recounts, like the vision to St. Paul on the road to Damascus, during a conference in Chicago. “Everybody showed slides of what was happening in their area, and it was like a shock. Then we had a session where people showed videos secretly taken in some of the labs, where chimps are in medical research, and that was like a visit to Auschwitz for me. It was as simple as that. I thought: now it is the payback time.”

Payback means speaking out against the unnecessary use of animals in medical research and establishing sanctuaries for illegally captured chimpanzees. Goodall has been attacked for her activism, but she is careful to note that she supports certain uses, that her mother’s life was saved by a pig’s heart valve. Goodall has also been criticized for saving captured apes, rather than putting money into maintaining habitat in the few places where the estimated 250,000 remaining wild chimps live. Again, the individual is paramount, she says: How could she ignore the starving, bedraggled chimps she has met in markets all over Africa?

Although she spends all her time these days fund-raising, Goodall still ponders chimp behavior. She is particularly interested in female transfer: why some females leave their group and stay away, why others leave, become pregnant and come back. Findings continue to come out of Gombe as well. In an August issue of *Science*, Goodall and Anne Pusey and Jennifer Williams of the University of Minnesota describe the role of hierarchy in female reproductive success. Although female hierarchy is difficult to establish—it is not as blatant as male dominance—the researchers used submission calls recorded between 1970 and 1992 to determine social standing. They concluded that the offspring of high-ranking females have higher survival rates and that their daughters reach sexual maturity earlier.

Finished with her garden talk, Goodall stands on the porch, shaking people’s hands before she has to rush off to another talk in a vast, sold-out auditorium. The line is long, and it is filled with young women. —Marguerite Holloway

ELECTRICITY

CHANGE IN THE WIND

Utilities are starting to offer renewable energy—for a price

Most utilities offer as much choice in how your electricity is created as Henry Ford offered to those buying his Model T: you can have any color you want, as long as it is black. But as power companies face deregulation and the prospect of competing for customers, many are beginning to sell a second, distinctly greener stream of energy. The juice flowing from solar cells, windmills and biomass furnaces is still a mere trickle running into an ocean of fossil- and nuclear-fueled power. But pilot projects are revealing just how many people will pay more for electricity that pollutes less.

The tiny, city-owned utility that serves Traverse City, Mich., gambled that many of its customers would pay a 23

percent premium (typically about \$7.50 a month) to light their lamps with wind rather than coal. With a grant from the state and a subsidy from the U.S. Department of Energy, the electric company erected a giant, 600-kilowatt windmill with blades 44 meters (144 feet) in diameter—the largest such turbine in North America.

Some 145 residents and 20 businesses signed up; another 75 filled a waiting list. “That amounts to 3 percent of our 8,000 customers,” says Steve Smiley, who managed the project. Love of Mother Earth was not the only incentive for these people, he notes. “We also promised ‘green’ customers that we would not increase their rates in the future, since the fuel is free.”

Several years ago the Sacramento Municipal Utility District began installing small photovoltaic panels on the roofs of those willing to pay an extra \$4 a month. Thousands applied, but the panels cost about \$20,000 apiece, so the company has so far set up only 420, enough to generate 1.7 megawatts. In May the utility signed contracts to add 10 megawatts’ worth of solar cells over the next five years.

The company also kicked off a new green pricing program similar to Traverse City’s: for an extra cent per kilowatt-hour, subscribers will get all their electricity from new renewable sources. (Not literally: green customers still draw power from every oil- and gas-fired dynamo on the grid. But their checks pay for cleaner generators.)

Some 23 other companies have followed suit. Public Service Company of Colorado has begun enlisting buyers for a 10-megawatt wind farm. Wisconsin Electric signed up more than 7,000 volunteers for hydroelectric and biomass power.

The trend is encouraging, says Blair G. Swezey of the National Renewable Energy Lab, but should not be mistaken for a resurgence in renewables. In fact, utilities are adding renewable capacity at just one fifth the

rate they did a decade ago. Nonpolluting energy is closing in on the cost of coal and oil, but it is not there yet.

How close is close enough? In surveys, 40 to 60 percent say they would pay more for cleaner power. “But the story changes when people get their checkbooks out,” observes Terry Peterson of the Electric Power Research Institute in Palo Alto, Calif. Few green-power programs have enrolled more than 5 percent of ratepayers. To be sure, most were poorly advertised and asked for premiums of 20 percent or more.

But an exception may prove to be the rule. When Massachusetts let homeowners in four cities choose among nine power vendors last summer, 16 percent chose Working Assets Green Power, which buys no electricity from nuclear or coal plants. Although Working Assets’s rates were the highest of the nine competitors, they were still cheaper than the monopoly that customers were leaving. “For green pricing to make a real difference, you need to charge less than what people pay today,” says Laura Scher, who managed the project.

That will be difficult, Swezey argues, as long as utilities can bill customers separately for failed investments, such as prematurely closed nuclear reactors. If those costs were instead factored into the price of electricity, then wind and dam power would look like more of a bargain. Because they are not, Swezey wagers it will take several years of healthy competition before the renewable power industry starts seeing green.

—W. Wray Gibbs in San Francisco

MATERIALS SCIENCE

HEAVY METAL MEETS ITS MATCH

Two new materials strip pollutants from toxic wastes

Chemistry sometimes seems almost magical in its ability to transform a mundane substance, such as pencil lead, into a valuable one, such as diamond, simply by reorganizing its atoms. Recently chemists demonstrated an impressive new trick. Starting with silica, the stuff of



J. CARL GANTER

WIND TURBINE

in Traverse City, Mich., produces premium-priced energy for 145 homes.

sand and window glass, a team of chemists has created a spongelike material so effective at absorbing certain heavy metals that it can render hazardous wastewater clean enough to drink. Researchers believe the material may prove cheap and adaptable enough to use in agriculture, electronics, manufacturing and perhaps even medicine.

Scientists have known for five years now how to make mesoporous silica—a form that, like a microscopic honeycomb, is riddled with long corridors, each just nanometers wide. With all those internal walls, a three-gram chunk of this substance contains as much surface area as a football field. Such a structure could cram lots of chemical reactions into a very small space. Unfortunately, silica doesn't react with much—one reason there is so much of it at the bottom of the ocean.

But in May, Jun Liu and his colleagues at Pacific Northwest National Laboratory in Richland, Wash., published a recipe for coating the walls inside mesoporous silica with other chemicals that do handy things. Liu used sulfur compounds that lock up mercury, silver and lead—common industrial pollutants that if ingested can cause brain damage and worse. In tests on water and oil wastes similar to those produced at the Savannah River weapons facility, Liu reports, the sulfur-laced silica powder reduced toxic concentrations of heavy metals to well below federal drinking-water standards. Equally important, the new material does not react with other, less dangerous metals—such as sodium and zinc—that often clog conventional filters.

The trick to placing useful chemicals inside the silica sponges, Liu says, lies in getting just the right amount of water inside its tiny tubes. Liu first dries them, then adds water back, along with a solvent. With his recipe, he claims, “you can make these things in your kitchen. The process seems simple enough to scale to large quantities” and to adapt for other chemical reactions. Other silica specialists agree.

“I think the prospects for environmental applications of this are quite high,” comments Ilhan A. Aksay, a chemical engineer at Princeton University. Although the coated silica costs about 50 percent more per pound than commer-



TRAP FOR HEAVY METALS,
mesoporous silica is filled with channels (shown here in cross section). Each tunnel can be lined with chains bearing sulfur (yellow) to lock up mercury (blue).

cial filter materials, it absorbs metals 30 to 10,000 times more effectively, Liu reports. Once mercury or lead is inside, it does not appear to leach out, even at high temperatures. Yet strong acid will wash out the metals for recycling, leaving the silica intact and quite reusable. “We’ve had many calls from environmental and chemical companies who want to work with us,” Liu says, although he declines to name them.

Galen Stucky, a chemist at the University of California at Santa Barbara, claims to have pushed Liu’s work a step

further, making stable mesoporous silica with tunnels twice as wide. That should be plenty large enough to contain biological molecules. The agriculture department is reportedly interested in packing silica powders full of pheromones to make long-acting pesticides. Others, Stucky says, are lacing the material with enzymes.

For removing metals, mesoporous silica is a tough act to follow. But for filtering out organic pollutants such as dyes, it faces new competition. In August, DeQuan Li of Los Alamos National Laboratory announced that through another bit of chemical sleight of hand, he had created a spongelike material built from cyclodextrins, compounds in common

starch. Linked into polymers, the cyclodextrins bind organic toxins 100,000 times more tightly than does activated charcoal yet can be washed clean with alcohol. Or so Li claims; the research has yet to be peer-reviewed.

“In order to treat large amounts of waste or have a big industrial impact,” Liu concedes, “we will need ways to make these materials dirt cheap”—a trick that often fails to materialize. But, he adds quickly, “we have a few ideas” about how to pull that out of a hat.

—W. Wayt Gibbs in San Francisco

ELECTRONICS

CHARGING TO MARKET

Supercapacitors are set to give batteries a jolt

Batteries are great at holding electricity. It’s in the giving and receiving that they cause problems. Charge them too fast, and they die. Draining them quickly—to zoom from zero to 60 in your electric roadster, for example—is equally damaging and often impossible. Capacitors can pick up where batteries leave off, because they store power as static electricity rather than chemical energy. But despite their name, capacitors have offered only small capacities: enough zap to pop a flashbulb but not enough to accelerate a car. That is about to change.

Three companies have begun small-scale production of supercapacitors that can store 10 to 5,400 times as much electricity as conventional capacitors. PolyStor in Dublin, Calif., rolls sandwiches of plastic and electrolyte-soaked carbon to make supercapacitors the size of penlight batteries. The carbon is in the unusual form of an aerogel, a porous solid that is sometimes called frozen air. “There is no chemical reaction involved in their operation,” points out PolyStor president James L. Kaschmitter, so the devices can be charged and discharged thousands of times without wearing out. In portable phones, laptop computers and other machines that often need large pulses of power, Kaschmitter says, supercapacitors can make batteries’ lives smoother and thus longer, for only an extra dollar or two.

The Pinnacle Research Institute in Los Gatos, Calif., is manufacturing ceramics inside the supercapacitors that can dis-

charge even faster than carbon, claims D. Bruce Merrifield, chairman of the firm's parent company. "They can make NiCad and lithium ion cells last five times longer," he says. Pinnacle is also aiming its higher-voltage devices at hospital defibrillators and "smart" missiles as well as mobile phones.

Supercapacitors fill a much larger need than just these niches, argues Maurice E. P. Gunderson, a venture capitalist with Nth Power Technologies in San Francisco. Deregulation, he says, will soon force

electric utilities to compete on quality and on price. Large enough capacitors can reduce a utility's cost to power mundane equipment, such as lights, by filling in during brief interruptions. More important, the devices could flatten surges and sags in the power going to sensitive manufacturing equipment. "If power problems in a pharmaceutical plant ruin a reactor full of some drug, it can cost millions," Gunderson points out. "The same applies to microchips and even Oreo cookies." The market for

devices that can prevent such mishaps could ultimately run to \$2 billion a year, he projects.

Maxwell Technologies in San Diego, Calif., appears best positioned to grab those dollars. Its carbon-cloth supercapacitors are the biggest to hit the market, and in July it formed a joint venture with PacifiCorp, an electricity wholesaler. "It's too early to say which design is best," Gunderson hedges. But that isn't stopping anyone from thinking big.

—W. Wayt Gibbs in San Francisco

VIRTUAL REALITY

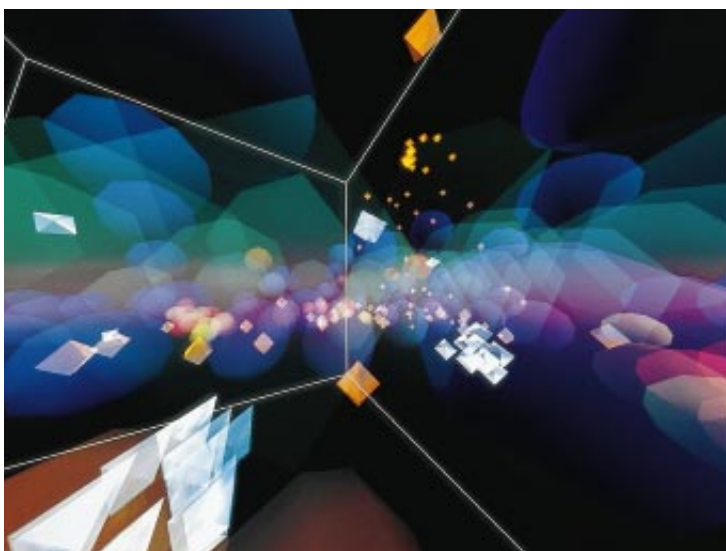
A Sense of Synesthesia

You are coming into the CAVE," Rita Addison begins. She is describing a virtual environment that she created to help people feel what it is like to have one's senses crossed, a phenomenon doctors call synesthesia (also the name of Addison's project). Five years ago a car accident scrambled sensory pathways in Addison's brain. Her vision clouded; the world seemed to zoom in and out, to spin. "Smells, absent at first, returned distorted," she recalls. "Sound wasn't heard but felt, like a push into my skin. With aphasia and vocabulary loss, frustration mounted whenever I tried to use words to explain what my world was like." So Addison instead turned her artistic skills to high-tech.

"The CAVE at the San Diego Supercomputer Center is a nine-foot cube; the walls are rear-projected video screens," she continues. "You are wearing a pair of liquid-crystal-shuttered glasses and a tracking device on top of your head. You are also carrying a little wand as a navigation tool. You are attired with an instrument that measures your chest's movement as you breathe.

"All around you there is a weblike image in pastels that have a subtle sheen [*below*]. When you start breathing, the web moves in and out with your breath.

"Now another person comes into the CAVE, wearing a band around his thumb to measure his heart rate. It creates ripples, moving the web up and down.



RYTA ADDISON

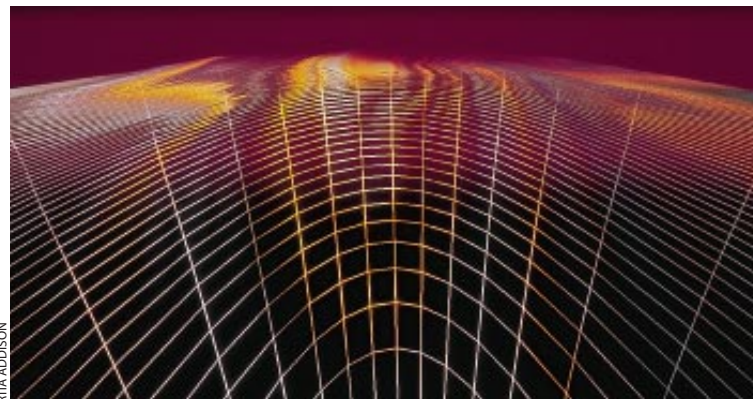
"We recorded the sound as blood flows from a big vessel to little vessels to capillaries. We also mixed in a recorded heart-beat. That sound is keyed to your heartbeat, the pace set by your own pulse. You are also making all of these wind sounds—we accentuate the swoosh of your breath.

"Now we change the environment on you [*above*]. Diamonds and spheres begin swirling around you. Your heartbeat presents itself in a new way, as a spurt of color rather than as a sound. If your breathing changes, the whole CAVE alters its flow patterns in response.

"*Synesthesia* is a linear experience," Addison explains; although participants can affect the environment, it still follows a script. "But it is a first step toward being able to have our physical senses as well as subtle physiological changes—skin conductivity, eye dilation, stuff that is pre-thalamus—notify the program that we are keener and modify the environment accordingly."

Although the project exhausted its funding last year, with more time (and money) Addison believes that virtual reality can be shaped into a potent tool for neuropsychological exploration.

—W. Wayt Gibbs in San Francisco



RYTA ADDISON

Master of Your Domain

What's in a name? On the Internet, it's your whole identity. Proposals to change the way Internet names are allocated have sparked arguments that expose the fact that behind the Net's apparent anarchy is a centralized structure. Controlling this structure is a relatively homogeneous group of engineers, lawyers and technical experts, a group that itself needs to be updated to match the radically changing character of the Net.

The current naming system was designed in 1983 as a human-friendly interface to the dotted clumps of numbers that routing computers understand. Each organization setting up on-line chooses what's called a domain name—like *Scientific American's* sciam.com. The name, along with the numbered address it represents, is added to the database for its top-level domain (the .com part), which in turn updates the world's routers. Besides .com, the other top-level domains in use in the U.S. are .edu, .gov, .net, .mil (military) and .org (nonprofits). Elsewhere, top-level domains are two-letter country codes, such as .fr for France, plus .int for international treaty organizations. Within those top-level domains, second-level identifiers distinguish types of organizations. This thoughtfully structured system has been stable through the stampede on-line except for one thing: almost everyone wants to be .com, which is short, memorable and easy to guess. This is partly snobbery. Businesses think .com sounds large, multinational and appealing to American customers. (Large American businesses, conversely, register names like microsoft.co.uk in Britain so they'll sound local.)

That has led to some conflict: a British consultancy uses prince.com, to the resentment of the American sports company. Had the U.S. followed the standard rules, this collision wouldn't have happened. American companies would sit in the disused .us domain, and .com would be reserved for multinationals. But the Net has traditionally rejected geographical divisions in favor of topics of interest. The underuse of .us is a

shame: acme.ithaca.ny.co.us is long but clear and leaves room for acme.ithaca.mn.co.us.

The current plan, developed by a group pulled together by old-time Net organizations—including the Internet Society, the Internet Assigned Numbers Authority and the Internet Engineering Task Force, plus the standards-setting International Telecommunication Union and the World Intellectual Property Organization—assumes that .us is lost. It introduces new top-level domains, widens the list of domain-name registrars and creates a council of registrars, to be established under the laws of Switzerland and overseen by two policy bodies appointed from the groups above.



DAVID SUTER

Opening up registration to competition is relatively uncontroversial. No one likes the present monopoly held by Virginia-based Network Solutions, now simultaneously floating an initial public offering and facing an antitrust investigation. People complain that the company, which charges \$50 a year per registration, mismanages its billing and other processes. Network Solutions's contract, awarded by the National Science Foundation, expires early in 1998.

There is less consensus about moving overall authority outside of the U.S., especially to appointed bodies with no commercial, education, government or consumer voices. Whereas some Americans believe the U.S. owns the Net (the Department of Defense paid only for the U.S. part, folks), and some call the plan an "attempted coup," the rest of the world wants Prince-style disputes to be settled in what they see as less partisan courts. "What this is really about is

not top-level domains but governance of the Internet," says Ivan Pope of Net-names UK, a British firm offering worldwide registration services. "Professionalizing governance is crucial—creating structures that are accountable and controllable by all interested parties."

But, judging from the comments I've seen, people hate the names: .firm, .store, .web, .arts, .rec, .info and .nom (for personal domains). "What is the problem we are trying to solve?" asks Donna Hoffman, an electronic commerce specialist at Vanderbilt University. If, she argues, we want to create more "good" names, this system fails because companies will register multiple names. If the goal is a directory structure, it fails again, because the names are confusing. "The categories should be mutually exclusive and exhaustive but also flexible enough to accommodate evolution," she says. She believes the Department of Commerce's call for public comments on the proposals is bringing the process to where it should have started: research.

Both Internet Society head Don Heath and Robert Shaw, an adviser at the International Telecommunication Union, laugh at the notion of significant opposition to their plan. They believe it will go through, with U.S. government support, by the end of the year, including the technical challenge of creating the shared registration database.

Other ideas, however, are worth considering. Domain-name dissident AlterNIC of Bremerton, Wash., promotes .xxx and .kids as easier ways to filter the Net than ratings systems. Or consider the logic of .radio, .air and .tv. No, wait, that last one is a country, the South Pacific island group Tuvalu. Would it sell? Tonga sells .to addresses via an automated Web server to all comers.

Hoffman is right: more research is needed. The right structure could solve a number of persistent problems if it took into account the changing nature of the Net, the fact that rules will always be broken, and the increasing value names and concepts acquire with use. The current plan does not do enough of the first two things, although it correctly says that domain names are a public trust, reflecting the human ability to create something valuable out of nothing.

—Wendy M. Grossman in London



Transportation's Perennial Problems

by W. Wayt Gibbs, *staff writer*

When the first issue of *Scientific American* hit the streets on August 28, 1845, its lead story excitedly touted “superbly splendid” new railroad cars able to “secure safety and convenience, and contribute ease and comfort to passengers, while flying at the rate of 30 or 40 miles per hour.” Half a century later this journal devoted almost an entire issue to innovations in bicycles, ships and the new steam-, electric- and gas-powered automobiles. “If there are faults” with cars, the editors concluded, “only time is wanted to make them disappear.... There is no mechanism more inoffensive, no means of transport more sure and safe.”

In hindsight, such blind faith that technology would solve the transportation woes of cities might seem quaint, even ironic. Now about half the travel on U.S. expressways slows to a crawl during the peak hours every day. Car crashes cause some three million injuries annually. According to the American Lung Association, roughly 100 million Americans live in cities where vehicle emissions regularly push ozone levels above federal standards. Hardly inoffensive, sure and safe.

But cars seemed a logical, progressive choice in 1899 because they helped to fulfill common human desires for mobility, space and status. They still do. For that reason, many developing nations are beginning to follow their rich peers down the asphalt path, with enormous consequences to their cities and environment. Also for that reason, attempts to reduce auto use have largely failed. Jane Holtz Kay argues in *Asphalt Nation* (Crown Publishers, 1997) that to solve the perennial problems of transportation “we must question why we travel at all.... We must alter our notions of mobility.” Many urban planners agree but caution that such funda-

mental changes typically require generations. In the meantime, technological advances may offer the most realistic means to take us from here to there faster, more safely and more cleanly.

Man versus Machine

At least, technology is what worked in the past—if only for a time. Consider safety. “All through the 19th century there were spectacular train wrecks: boiler explosions, fires. Head-on collisions were not unusual,” reports George M. Smerk, director of the Institute for Urban Transportation at Indiana University at Bloomington. To quell public outcry, rail and trolley lines installed steel cars, electric signals and air brakes. Accident rates fell. And then engineers responded by speeding up.

Drivers have shown the same tendency to adjust their behaviors to maintain a steady level of risk. Autos were initially safer than horses, says Clay McShane, a historian at Northeastern University: “Cars don’t run away on their own, they don’t bite, and they don’t kick.” In time, of course, drivers more than compensated for the predictability of their vehicle by stepping on the gas.

More recently, seat belt use has jumped from 11 percent in the early 1980s to about 68 percent now; air bags are making similar inroads. Perhaps predictably, drivers have begun traveling faster and following more closely, so the 41,798 highway fatalities in 1995 were down only about 2,400 from 1983. On the other hand, they were up just 11,750 from the automotive death toll in 1931, despite a fivefold increase in the number of cars on the road.

Drivers seem less interested in cleaner vehicles than in safer ones: “A third of the cars today are larger than any auto on the road in the 1950s,” McShane says. Yet here again, he recounts, “the



TRAFFIC almost always rises over time to exceed highway capacity. Building more roads can actually make congestion worse. New York City streets were as jammed in 1875 (*above*) and 1917 (*top right*) as they are today.

auto looked initially like an enormous improvement to the environment.” New York City in 1900 was buried under roughly four million pounds of manure every day. Horses had to be stabled away from their carriages, he states, “because their urine fumes were strong enough to blister paint. But the worst pollution problem was the air loaded with bacteria-carrying dust, through which respiratory diseases were transmitted.” When autos displaced horses in the 1920s, he says, tuberculosis rates plummeted.

“Many argue that the current air-quality problem in urban America will be ‘solved’ with cleaner vehicles,” notes Michael D. Meyer of the Georgia Institute of Technology. Indeed, hydrocarbon emissions fell 35 percent from 1984 to 1993, thanks to more efficient cars and cleaner gas. “However, the growth in vehicle miles traveled is expected to overwhelm any improvements that will likely occur in vehicle emissions,” Mey-



The congestion, accidents and pollution that plague modern travel are hardly new. History and recent research suggest they may remain intractable for generations to come



CORBIS BETTMANN



MICHAEL YAMASHITA Corbis

er adds. Odometers will spin ever faster so long as cities continue to spread out.

“Jam Yesterday ... Jam Tomorrow”

Like Alice at the Mad Hatter’s tea party, highway planners are caught in a vicious cycle, says Martin Wachs of the University of California Transportation

Center. “You can never build enough roads to keep up with congestion. Traffic always rises to exceed capacity.”

Part of the problem, operations engineer Dietrich Braess showed in 1968, is that adding new routes often makes congestion worse, not better. That paradox seems to have vexed every age. “Rush hours have always been a mess,” Smerk

says. “Traffic jams were so bad in Rome 2,000 years ago that the city banned chariot riding during peak hours.” In New York, McShane adds, “people complained about crowding on the horse cars 10 years after they began operation. Trolleys were overcrowded within five years of electrification. Mass auto-mobility comes in 1907; by 1914 you have traffic jams. The U.S. built the first interstate highways in the early 1920s, and they were already jammed by the end of the decade.”

More important than Braess’s paradox is the fact that with increased mobility people move not just around but away. “The horse car allowed city dwellers to move out to single-family homes,” McShane observes. “Then the laying of rails lowered fares to a nickel, allowing movement into the suburbs.” By the time autos appeared, cities had already begun to sprawl along the main rail lines.

Cars—especially when abetted by government-subsidized housing loans after 1945—kindled that spark into explosive suburban growth. It continues today: about 86 percent of the population growth in the U.S. since 1970 happened

in suburbs, Meyer reports. And for good reason, remarks Robert W. Burchell of the Center for Urban Policy Research at Rutgers University: "As you go farther out, your taxes fall, your housing generally costs less, your schools improve, you get increasing amounts of public recreation facilities, you are safer from crime, and you are more likely to be surrounded by people like yourself. Given its ability to deliver all that, it is no wonder the public loves sprawl."

Western European governments have showered fewer gifts and more auto taxes on their exurbanites. As a result, says John Pucher, an urban planner at Rutgers, their central cities typically have four times the population density of

and rail stations, they suggest, and people will drive less. Put businesses closer to homes, and citizens should reduce their travel altogether [see "Why Go Anywhere?" by Robert Cervero; *SCIENTIFIC AMERICAN*, September 1995].

Forward to the Past?

The New Urbanism movement, as it is called, has noble goals. But it faces tremendous practical obstacles. Razing and rebuilding entire suburbs is not feasible, so most neotraditional communities have been, and will be, built on cities' outskirts. Unfortunately, "there is no cost-effective way to build a transit system that serves beltway locations,"

towns tend to attract residents who already use public transit to get to work. Moreover, when Crane and his colleague Marlon G. Boarnet studied all 232 transit stations in southern California, they found that almost without exception, cities tend to put their stations near shopping centers and offices (which bring in jobs and taxes), not homes. "Transit-based housing will struggle," the two predicted, until cities begin chasing residents instead of businesses. "For the most part," they conclude, "that seems unlikely to happen."

In the interim, U.S. cities might find a different European strategy more effective: tolls. In 1991 Trondheim, Norway, placed electronic tollbooths on all routes

leading into the city, closing free access by road. It gave away radio tags; nearly all drivers now use them to pay without stopping at the gate. The city recouped its capital investment in six months, boasts Tore Hoven of the Trondheim Public Roads Administration. Tolls have since paid for new roads, sidewalks and buses. And because tolls rise during the morning rush hours (a technique called congestion pricing), many drivers switched to trains, boosting transit ridership 7 percent in a single year. When Stuttgart tested a similar system in 1995, it found that congestion pricing cut rush-hour travel by 12 percent.

"Is the American public ready for full pricing? I don't think so," Meyer comments. But that may change; there is nothing inherently un-American about tolls. Indeed, most of the first highways built in the U.S. were privately owned turnpikes. At least 2,000 companies maintained toll roads during the 19th century. The fashion may be returning;

private highways have recently opened in Dulles, Va., and Orange County, California. Houston, Tex., is also considering congestion pricing on one of its interstates.

States will be increasingly forced to squeeze more out of existing roads, Burchell says, because "the consequences of sprawl are costly. We just did a study for South Carolina that calculated their infrastructure tab for the next 20 years



PAUL CHESLEY/Doni Stone Images

CONGESTION IN BANGKOK fritters away 35 percent of the city's yearly economic output.

America's urban centers. Because stores and job sites are closer, Pucher adds, "Europeans make 40 to 50 percent of trips by walking or biking and about 10 percent by public transit. In contrast, 87 percent of trips in the U.S. are by car; only 3 percent involve transit."

Many urban planners in the U.S. now prescribe similar strictures to reduce traffic flows. Replace cul-de-sacs and parkways with old-fashioned street grids

McShane argues. Boston has tried to do this, Harvard University professor Jose A. Gomez-Ibanez points out in a recent article, and as a result its transit agency has faced budget crises every decade or so since 1961. It is due for another soon.

A recent microeconomic analysis by Randall Crane of the University of California at Irvine concluded that neotraditional designs may be good ideas but will not necessarily curb traffic. Such

as \$57 billion. That is \$1,000 a year for every person in the state for the rest of their lives. Increasing the gas tax by four cents would raise only \$56 million. But just by living differently, by setting growth boundaries around cities, doubling the amount of development inside the circle and halving the amount outside, you could save \$2.5 billion" in public infrastructure and services.

"Our best hope for easing sprawl" and the congestion it causes, Burchell contends, "is that we will run out of money. Sooner or later we will not be able to continue building so much infrastructure, because we can no longer afford to maintain it." Michigan and other states are already considering growth boundaries for that reason, he says.

On the other hand, McShane observes, "during a recession, highway building is a great way to inject money into the economy. If you had told me in 1988 that a city as environmentally conscious and transit-intensive as Boston would invest \$10 billion in downtown highways, I would have laughed at you. It happened."

The World Speeds Up

Recent work by Andreas Schafer of the Massachusetts Institute of Technology may help explain why, despite the well-known evils of automobiles, Americans—and, increasingly, Europeans—drive more miles year after year, often rearranging their communities to make that possible. Drawing on decades of travel surveys, Schafer found that city dwellers in the U.S., Europe, Russia, eastern Asia and even villages in Ghana share two important traits, which appear to have remained constant for at least 30 years. First, people in each location spend an average of 60 to 90 minutes traveling a day. And in every industrial country except Japan, people spend an average of 10 to 15 percent of their income doing it [see "The Past and Future of Global Mobility," by Andreas Schafer and David Victor, page 58].

As nations all over the world have grown richer, they have consistently used part of their wealth to buy speed. "Mobility is an underrated human right," Wachs declares. "You can never have enough of it."

If Schafer's trend holds true, it could have important implications for the developing world and those who share its atmosphere. Many Third World megacities already face huge transportation



ELECTRONIC TOLLBOOTHs in Trondheim, Norway, allow cars to zip through without stopping. Radio transceivers collect the fees, which rise during rush hours. Such congestion pricing might ease chronic traffic jams elsewhere.

snarls. Cars in Manila average seven miles (11 kilometers) per hour, reports Ralph Gakenheimer of M.I.T. A typical auto in Bangkok is stopped in gridlock the equivalent of 44 days each year; the congestion eats 35 percent of the city's gross annual output. New Delhi already loses six citizens a day on its highways, and air pollution harms many more.

Yet as incomes rise in Asia, so will the number of motor vehicles. "Around the world, one of the first things people buy when they can is a car," Pucher says. Gakenheimer points to a Chinese government survey that found citizens typically willing to spend up to two years' income on an automobile. (The average

American invests just six months' earnings.) Schafer estimates that if India follows the example of other nations, it will have 267 million cars on its roads by 2050. Rising car ownership, Gakenheimer predicts, will overwhelm developing cities, causing explosive sprawl. And thus the cycle begins again.

Meanwhile auto-saturated countries such as the U.S., finding it difficult to eke more speed out of their cars, are taking increasingly to the air. That has already begun to spawn a host of new traffic, safety and pollution problems. Will it ever end—will we ever finally quench our thirst for mobility? Think warp drive.

Further Reading

DOWN THE ASPHALT PATH: THE AUTOMOBILE AND THE AMERICAN CITY. Clay McShane. Columbia University Press, 1994.

URBAN PASSENGER TRANSPORT IN THE UNITED STATES AND EUROPE: A COMPARATIVE ANALYSIS OF PUBLIC POLICIES. J. Pucher in *Transport Reviews*, Vol. 15, No. 2, pages 99–117; 1995.

BIG-CITY TRANSIT RIDERSHIP, DEFICITS, AND POLITICS: AVOIDING REALITY IN BOSTON. Jose A. Gomez-Ibanez in *APA Journal*, Vol. 62, No. 1, pages 30–50; Winter 1996.



The Past and Future of Global Mobility

With growing wealth, people everywhere travel farther and faster. That trend inevitably brings a shift in the dominant transportation technologies

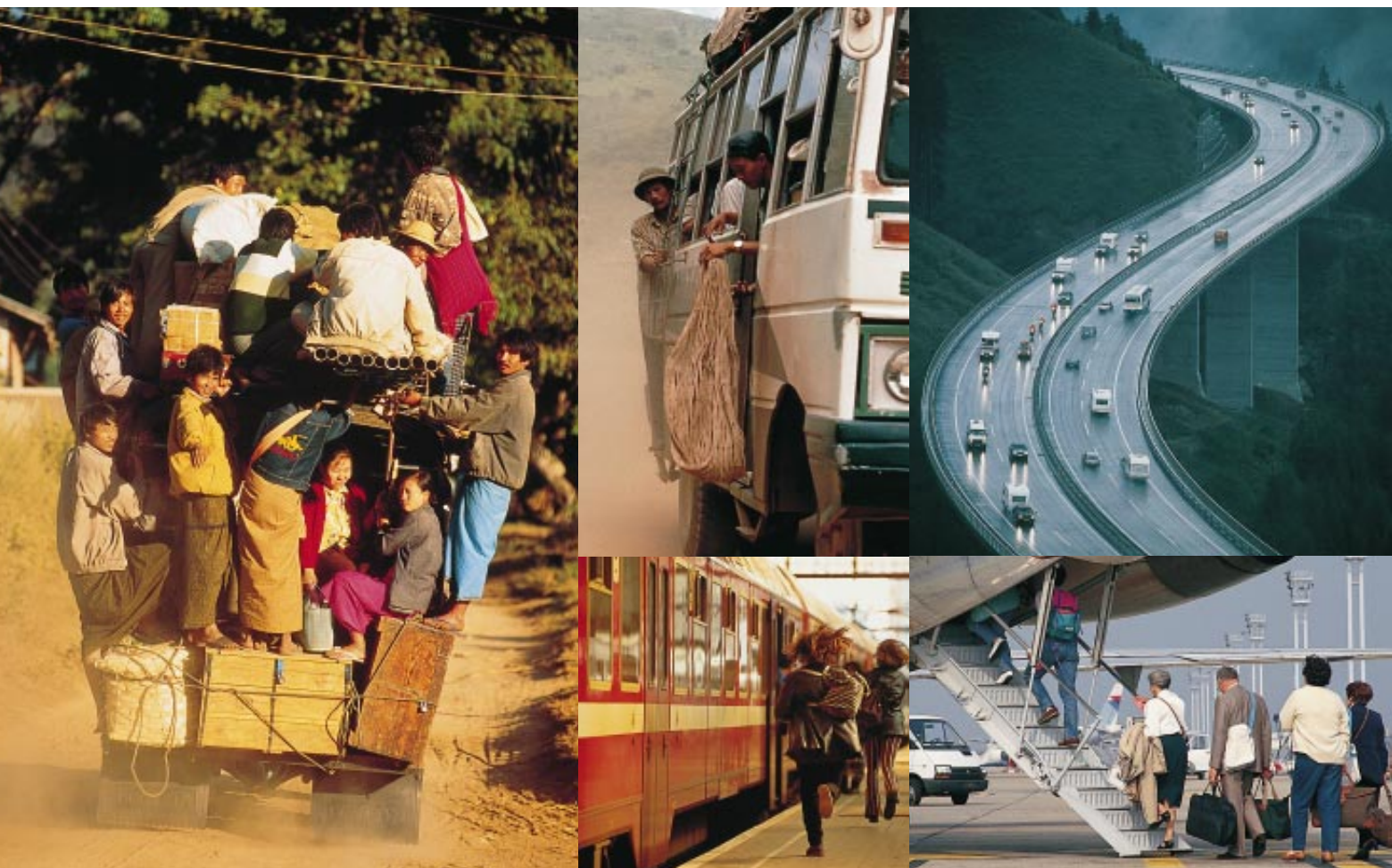
by Andreas Schafer and David Victor

How much will people travel in the future? Which modes of transport will they use? Where will traffic be most intense? The answers are critical for planning infrastructures and for assessing the consequences of mobility. They will help societies anticipate environmental problems such as regional acid rain and global warming, which are partially caused by transport emissions. These

questions also lie at the center of efforts to estimate the future size of markets for transportation hardware—aircraft, automobiles, buses and trains.

In our research, we have tried to answer these questions for 11 geographic regions specifically and more generally for the world. One of us (Schafer) compiled historical statistics for all four of the principal motorized modes of transportation—trains, buses, automobiles

and high-speed transport (aircraft and high-speed trains, which we place in a single category because both could eventually offer mobility at comparable quality and speed). Together we used that unique database to compose a scenario for the future volume of passenger travel, as well as the relative prevalence of different forms of transportation through the year 2050. Our perspective was both long term and large scale because trans-



port infrastructures evolve slowly, and the effects of mobility are increasingly global. The answers to those fundamental questions, we found, depend largely on only a few factors.

Historical data suggest that, throughout the world, personal income and traffic volume grow in tandem. As average income increases, the annual distance traveled per capita by car, bus, train or aircraft (termed motorized mobility, or traffic volume) rises by roughly the same proportion. The average North American earned \$9,600 and traveled 12,000 kilometers (7,460 miles) in 1960; by 1990 both per capita income and traffic volume had approximately doubled.

In developing countries the relation has been less tight. Between 1960 and 1990 the average income in China tripled, but motorized traffic volume rose 10-fold, to 630 kilometers. This discrepancy reflects, in part, the fact that growing wealth allows the poor to substitute motorized mobility, typically by bus or train, for nonmotorized forms such as walking and biking, for which the statistics are notoriously unreliable and so are excluded from our database.

The charted relation between income and traffic volume affirms a postulate by the late analyst Yacov Zahavi: on aver-

age, humans devote a roughly predictable fraction of their expenditures to transportation. This fraction is typically 3 to 5 percent in developing countries, where people rely predominantly on nonmotorized and public transportation. The fraction rises with automobile ownership, stabilizing at 10 to 15 percent at ownership levels of 0.2 car per capita (one car per family of five). Nearly all members of the Organization for Economic Cooperation and Development (OECD)—the rich industrial nations—have completed this “automobile transition.” Figures from the U.S., for example, show that this fraction remained nearly constant even during the two oil-price shocks of the 1970s; travelers compensated for higher operating costs by demanding less expensive (and more fuel-efficient) vehicles.

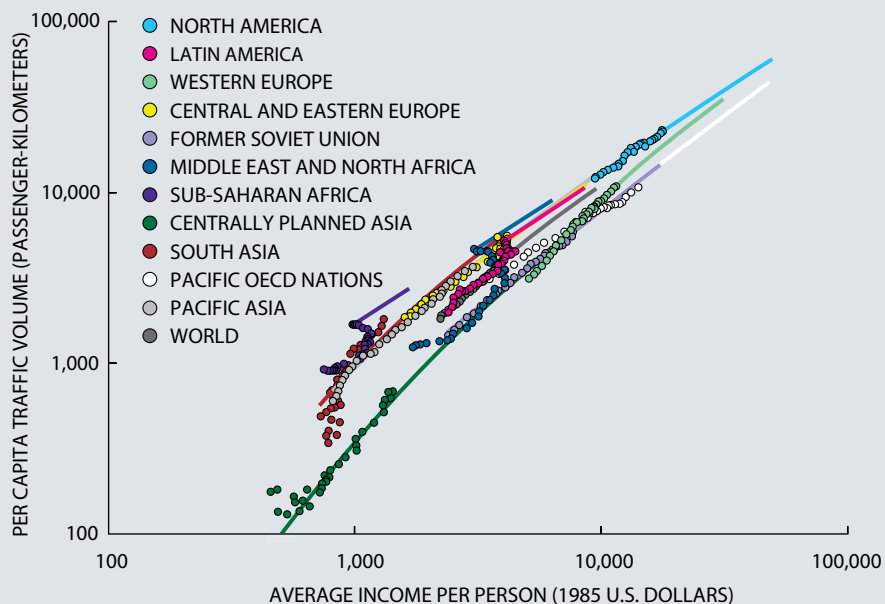
This predictable relation between income and transport spending allowed us to conjecture plausibly about the future. In the absence of major economic upsets, traffic volume should continue to rise with income, as in the past. Using reasonable assumptions for future income growth, we estimated that traffic volume in North America will rise to 58,000 passenger-kilometers a year in 2050. In China, annual motorized mo-

bility will reach 4,000 passenger-kilometers, which is comparable with western European levels in the mid-1960s. Developing countries will contribute a rising share to global traffic volume because, although their per capita mobility will stay lower, both their populations and their average incomes will grow faster than those of OECD nations. In 1960 the developing countries could claim only 22 percent of the world traffic volume, but by 2050, we estimate, they will account for about half of it—51 trillion passenger-kilometers.

Higher Incomes, Higher Speeds

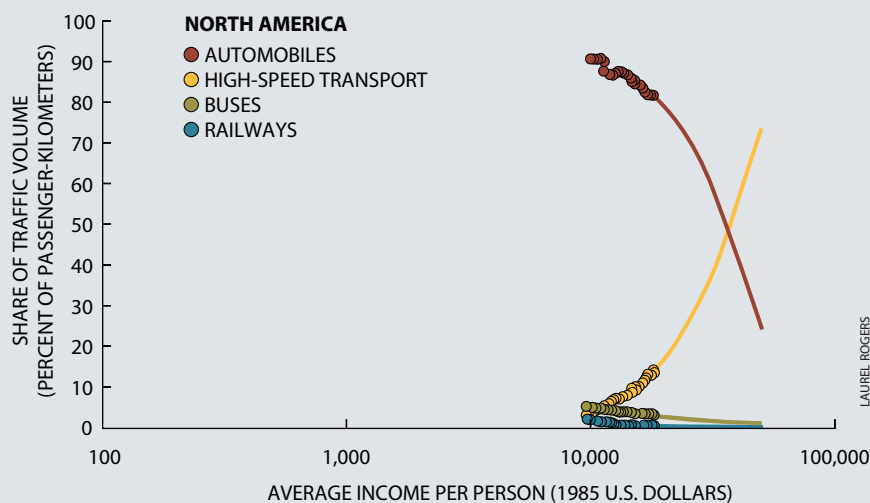
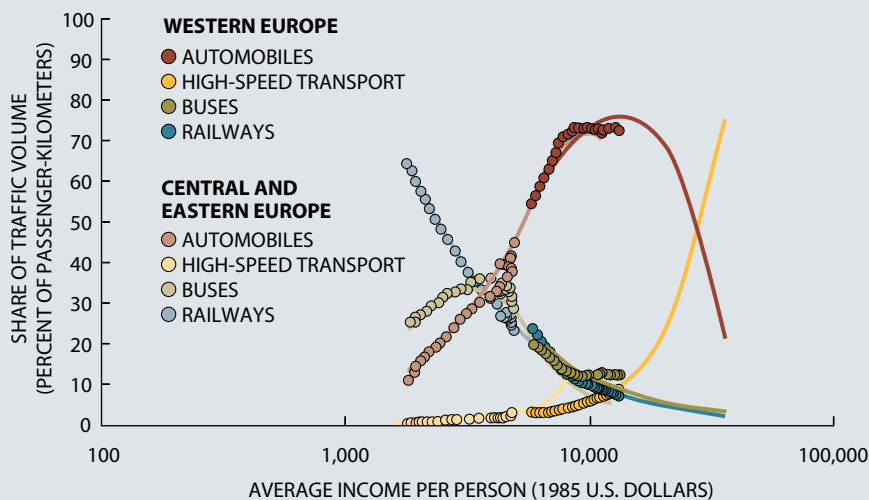
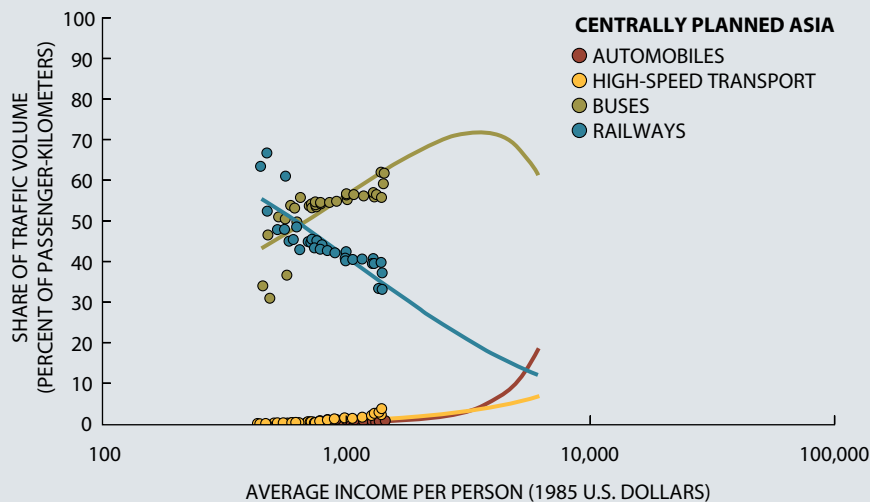
How will people satisfy their growing demand for mobility? We searched for patterns in how modes of transportation compete. Again, Zahavi offered a useful starting point: he argued that people devote on average a constant fraction of their daily time to travel—what he called the travel-time budget. All the reliable surveys that we have found support this hypothesis: the travel-time budget is typically between 1.0 and 1.5 hours per person per day in a wide variety of economic, social and geographic settings. Residents of African villages have a travel-time budget similar to those of Japan, Singapore, western Europe and North America. Small groups and individuals vary in their behavior, but at the level of aggregated populations, a person spends an average of 1.1 hours a day traveling.

If people hold their time for travel constant but also demand more mobility as their income rises, they must select faster modes of transport to cover more distance in the same time. Data from every region are consistent with that expectation. At low incomes (below \$5,000 per capita), motorized travel is dominated by buses and low-speed trains that, on average, move station-to-station at approximately 20 to 30 kilometers per hour. As income rises, slower public transport modes are replaced by automobiles, which typically operate door-to-door at 30 to 55 kph and offer greater flexibility. (These average speeds, which vary by region, are lower than the posted speed limits because of congestion and other inefficiencies.) The share of traffic volume supplied by automobiles peaks at approximately \$10,000 per capita. At higher incomes, aircraft and high-speed trains supplant slower modes. At present, aircraft supply 96 percent of all high-speed transport, fly-



SOURCE: Andreas Schafer and David Victor

MOTORIZED TRANSPORTATION takes many forms around the world, ranging from relatively slow public transit through private automobiles to high-speed planes (*opposite page*). Data from 11 regions collected between 1960 and 1990 generally demonstrate that as income rises, societies become more mobile (*above*). All income data are weighted for differences in local prices.



SOURCE: Andreas Schafer and David Victor

SHIFT FROM SLOW TO FAST MODES of transportation occurs with rise in income, as trends in several regions show. These curves represent historical data and future scenarios between 1960 and 2050. In centrally planned Asia (primarily China), buses are the preferred mode; trains are in decline, whereas cars and planes are of minor importance. In central and eastern Europe, cars are still on the rise, but in western Europe, a transition in favor of planes and high-speed trains is occurring. In North America, planes are already taking a share of traffic volume from cars.

ing airport-to-airport at about 600 kph.

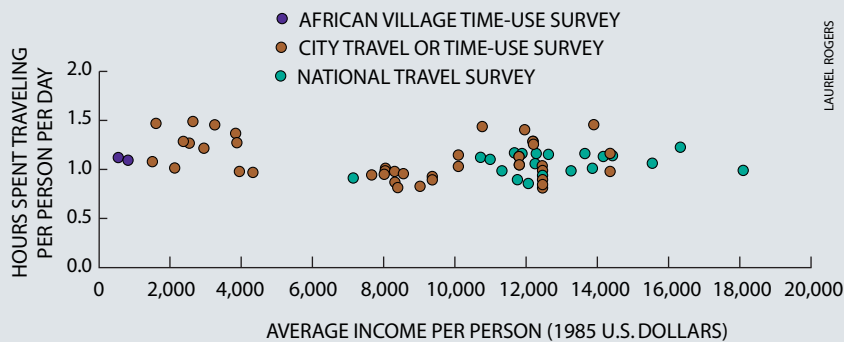
Although the constancy of the travel-time budget pushes people with rising incomes toward faster modes of transportation, the share of motorized mobility that each mode holds is strongly determined by geography. In the late 1950s, when Jack Kerouac extolled the open road in America, relatively few kilometers were motored by other means: by the 1960s, private automobiles delivered 90 percent of North American traffic volume because the continent had plenty of space and plenty of roads. In contrast, in more densely populated western Europe, the share of automobiles never climbed so high—it has been stagnant at about 70 percent and is poised to decline. Asia is even more compact, with an urban density three times that of western Europe. Accordingly, we expect that automobiles will peak at only 55 percent of the total traffic volume in the high-income Pacific OECD nations, which is primarily attributable to Japan. Public transport will continue to account for a higher share of mobility in Asia than in less densely populated regions.

In addition, the availability of roads, rail beds, airports and other essential infrastructures constrains the transport choices. Because transport infrastructures are expensive and long-lived, it typically takes six to seven decades to eliminate them (for example, canals) or to make new ones (for example, roads). New infrastructures could be built for a radically different transportation system by late in the next century, but transport choices for the next few dozen years will be limited by earlier investments.

On the Move in 2050

Assuming that a constant travel-time budget, geographic constraints and short-term infrastructure constraints persist as fundamental features of global mobility, what long-term results can one expect? In high-income regions, notably North America, our scenario suggests that the share of traffic volume supplied by buses and automobiles will decline as high-speed transport rises sharply. In developing countries, we anticipate the strongest increase to be in the shares first for buses and later for automobiles. Globally, these trends in bus and automobile transport are partially offsetting. From 1960 to 2050 the share of world traffic volume by buses will remain roughly constant, whereas

LAUREL ROGERS



SOURCE: Andreas Schafer and David Victor

TRAVEL-TIME BUDGET, the amount of time that people devote to travel, is consistently about 1.1 hours per person a day in all societies, according to surveys.

the automobile share will decline only gradually to 35 percent. High-speed transport should account for about 40 percent of all passenger-kilometers traveled in 2050. In all regions, the share of low-speed rail transport will probably continue its strongly evident decline.

Despite the sharply rising share of air travel, other types of vehicles, including

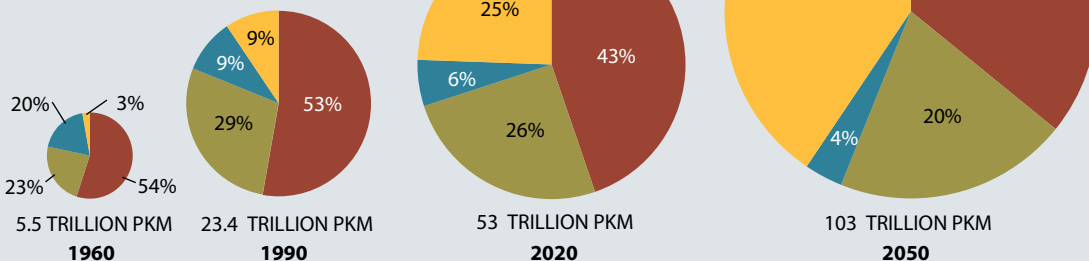
automobiles, will remain crucial parts of the transportation system. Even in North America, where we expect the relative decline of automobiles to be steepest, the absolute traffic volume supplied by cars will decline only after peaking at 22,000 passenger-kilometers per person in 2010. By 2050, automobiles will still supply 14,000 passenger-kilome-

ters per person, which means that North Americans will be driving as much as they did in 1970.

The allocation of travel time reflects the continuing importance of low-speed transport. We expect that throughout the period 1990–2050, the average North American will continue to devote most of his or her 1.1-hour travel-time budget to automobile travel. The very large demand for air travel (or high-speed rail travel) that will be manifest in 2050 works out to only 12 minutes per person a day; a little time goes a long way in the air. In several developing regions, most travel time in 2050 will still be devoted to nonmotorized modes. Buses will persist as the primary form of motorized transportation in developing countries for decades. No matter how important air travel becomes, buses, automobiles and even low-speed trains will surely go on serving vital niches. Some of the super-rich already commute and shop in aircraft, but average people will continue to spend most of their travel time on the ground. SA

WORLD TRAFFIC VOLUME, measured in passenger-kilometers (pkm), will continue to balloon, with higher-speed transport gaining market share. By 2050, automobiles will supply less than two fifths of global volume.

RAILWAYS
BUSES
AUTOMOBILES
HIGH-SPEED TRANSPORT



SOURCE: Andreas Schafer and David Victor

The Authors

ANDREAS SCHAFFER and DAVID VICTOR collaborate on long-term and large-scale models of transportation. Schafer, an aeronautical engineer, works at the Massachusetts Institute of Technology Center for Technology, Policy and Industrial Development. He does systems analysis on transportation and global change in the Cooperative Mobility Program and the Joint Program on the Science and Policy of Global Change. Victor, a political scientist with the Environmentally Compatible Energy Strategies Project at the International Institute for Applied Systems Analysis (IIASA), focuses on energy technology and international environmental governance.

Further Reading

PERSONAL TRAVEL BUDGETS. Edited by H. R. Kirby. Special issue of *Transportation Research*, Part A (Pergamon Press, U.K.), Vol. 15, No. 1; January 1981.
THE EVOLUTION OF TRANSPORT SYSTEMS: PAST AND FUTURE. Arnulf Gruebler and Nebojsa Nakicenovic. Research Report RR-91-008, International Institute for Applied Systems Analysis, Laxenburg, Austria, 1991.
ANTHROPOLOGICAL INVARIANTS IN TRAVEL BEHAVIOR. C. Marchetti in *Technological Forecasting and Social Change*, Vol. 47, No. 1, pages 75–88; September 1994.
THE FUTURE MOBILITY OF THE WORLD POPULATION. A. Schafer and D. G. Victor in *The Cooperative Mobility Program*. Discussion Paper 97-6-4, Center for Technology, Policy and Industrial Development, Massachusetts Institute of Technology, 1997.
THE GLOBAL DEMAND FOR MOTORIZED MOBILITY. Andreas Schafer in *Transportation Research*, Part A (in press).



13 Vehicles That Went Nowhere

Perhaps “nowhere” is too harsh. But all these transportation concepts—however brilliant or eccentric—fell far short of their enthusiasts’ great hopes. Some ran afoul of technical glitches or practical constraints. Some couldn’t compete with other transports. Some had bad luck. Some evolved into different types of vehicles. And some... well, maybe they weren’t very good to start with. In any case, they illustrate one of the most important lessons of transportation technology: it takes more than a bright idea to get somewhere.

—John Rennie, Editor in Chief



UPI/CORBIS-BETTMANN



UPI/CORBIS-BETTMANN

T. P. Hall’s flying auto (1948)

THE FLYING CAR

Background: Not just a car and not just a plane, but both, this fantasy has gripped inventors for as long as there have been both cars and planes. Why forsake the comforts of the family sedan while flying cross-country? Pilots wouldn’t need to hire a car at their destination. Perhaps these craft were meant to bring low-cost flying to the masses. But it’s hard not to think that the builders were inspired at least as much by a spirit of pure “because-we-can” intrepidity.

Quite a few “flying flivvers” were tried, including models in which the wings could be removed for driving. Henry Ford and major manufacturers such as Studebaker and Convair flirted with them. The Aerocar, featured in the TV comedy *Love That Bob*, was in production from 1946 to 1967; five were sold.

Problems: On the ground, car-plane hybrids were more cramped and fragile than ordinary cars; in the air, they handled worse than ordinary planes. They could be both expensive and unsafe. Imagine the air-traffic nightmares that would result from thousands of unscheduled takeoffs and landings on highways.

Status: Rest assured, you haven’t seen the last of these.



UPI/CORBIS-BETTMANN

Fuller’s three-wheeler (1933)

THE DYMAXION CAR

Background: The brainchild of inventor Buckminster Fuller, this 1933 automobile embodied for transportation the same principles of economic form and functionality that the geodesic dome brought to architecture. It had only three wheels—one that steered in back and two motorized drive wheels in front—which made it highly maneuverable. The 20-foot, 11-seater version could U-turn in less than its own length. Its raindrop contour was streamlined for fuel efficiency (about 30 miles per gallon). And because of its light weight, the car re-

portedly had a top speed of 120 miles (over 190 kilometers) per hour.

Problem: While racing in 1935, a Dymaxion car was involved in a fatal accident. (Ironically, the other car may have been at fault.) The resulting bad publicity scared away investors, scuttling the project.

Status: Although Bucky continued to refine the Dymaxion car, making it smaller and easier to steer, commercial interest had evaporated. It survives only as an inspiration to other designers of more economical, ecologically sound automobiles.

ROCKET BELTS, JET BELTS AND THE WASP

Background: Strap on an engine and take to the skies! These wonderful gadgets epitomized solo aerial freedom. Wendell F. Moore of Bell Aerosystems invented the rocket belt in 1953; it was little more than a steerable pair of chemical rockets worn like a backpack. Yet it captured the popular imagination at air shows, in commercials and in the James Bond movie *Thunderball*. Further refinements led Bell to build the jet belt, in which a high-thrust turbojet took the place of rockets. In 1970 Bell sold the rights to the jet belt to Williams Re-

search Corporation. In the subsequently developed Williams Wasp, instead of wearing the engine, the pilot stepped onto a platform that housed the vertically oriented turbojet.

Problems: Limited range was one restriction. The original rocket belt could carry only enough fuel to stay aloft for slightly more than 20 seconds—not much of a ride. The jet belt stretched that to about five minutes. Another problem common to both belts was that the pilot's legs had to serve as landing gear, so a misstep during takeoff or landing could be hazardous. The Wasp, however, overcame those obstacles, because it had its own legs and could carry more fuel.

What ultimately seems to have done in these devices was lack of a well-justified mission. The military had contracted for them, but it could not find enough reasons to send infantrymen into the air for short hops or for aerial reconnaissance that might be performed by conventional aircraft.

Status: Aside from occasional special appearances, such as at the opening ceremonies of the 1984 Olympics, these devices appear to be well-loved but idle historical pieces.

Mr. Bond's belt (1965)



EVERETT COLLECTION



SUPERSTOCK

Standing on air (1962)

HILLER FLYING PLATFORM

In the 1950s, long before the Williams Wasp, Hiller Corporation experimented with a manned platform that flew on the power of a giant, ducted fan. The pilot steered by leaning side to side. Its poor maneuverability and undefined utility discouraged further development.

PNEUMATIC TRAINS

Background: In 1870 Alfred Ely Beach, then editor of *Scientific American*, financed the construction of a prototype subway in New York City. Based on experimental European pneumatic trains, it consisted of a block-long stretch of tunnel through which a cylindrical car was pushed and pulled by a huge fan. Though popular, this system failed to win over the municipal authorities, who later built elevated trains instead.

But the idea of using air pressure to propel a train never lost its appeal. In the mid-1960s Lockheed and the Massachusetts Institute of Technology, in conjunction with the U.S. Department of Commerce, contemplated the feasibility of pneumatic trains connecting cities along the Boston-to-Washington corridor. Train cars would snugly fit into evacuated tubes hundreds of miles long. Opening and closing

valves would allow ambient air pressure to push the tube cars to their destination. For an added boost, the tubes would slope downward out of each station, creating a "gravitational pendulum" assist for the trains. Calculations suggested that on the run between Philadelphia and New York, for example, the average speed might be 390 miles per hour.

Problems: Boring tunnels to the required mechanical tolerances and then emptying them of air would have been expensive (to say the least). Any accident that compromised the vacuum or integrity of a tube at any point in its length would force a shutdown of the entire intercity line. Improving the highway, rail and air transit systems seemed like a better bet.



SCIENTIFIC AMERICAN

Beach's tube car (1870)

HOVERCRAFT

Background: Hovercraft, also known as air-cushion or ground-effect vehicles, float almost frictionlessly above a surface rather than rolling across it and so can move with equal ease over paved roads, dirt beds or lakes. Designs date back to the 1800s, but hovercraft did not become practical until after the 1950s with the invention of the inflatable skirt, which helps to trap the fan-driven air cushion underneath the vehicle.

Buoyed with enthusiasm (so to speak), some aficionados once believed that hovercraft might render conventional cars, trucks, boats and trains obsolete. Prototypes for hover rail systems between Paris and Orleans were tested. The military publisher Jane's looked forward to an era of hovering naval vessels as big as destroyers and traveling at 100 knots. Futurist Arthur C. Clarke speculated that once hovercraft blurred the distinctions between moving over land or water, the trade advantages of port cities would vanish; land-locked metropolises such as Oklahoma City might be the major crossroads of the 21st century.

Problems: The low-friction ride of hovercraft has a down side—it makes them hard to control. Above anything except a flat, evenly packed surface, they tend to slide downhill. (Hence, they are very stable on ice.) On rough seas, they lose maneuverability and can be blown off course. Moreover, the fans that generate the air cushion and thrust can be too loud for urban or residential areas and even for some military missions.

Status: Even without replacing cars or boats, hovercraft have carved out healthy niche businesses. They routinely serve as high-speed ferries across the English Channel and other bodies of water. In Canada, hovercraft make superlative ice breakers for shipping lanes, shattering the ice below them with shock waves rather than smashing through it. Navies are primarily interested in hovercraft as amphibious landing transports for troops and equipment, because they can quickly move from a carrier, across water and onto dry land. Hobbyists also continue to enjoy building and racing recreational hovercraft.

Princess Margaret hovers over the Thames



HULTON GETTY (ony Stone Images)

MOVING SIDEWALKS

Science-fiction writers used to imagine that cities of the future might have conveyor-beltlike sidewalks for speeding pedestrians on their way. But maintaining lengthy stretches of conveyor against the outside elements is an expensive proposition. Moving sidewalks have therefore found a more suitable home inside the sprawl of modern airports, where they are an efficient solution for bringing people and their luggage from point A to point B.



UPI/CORBIS-BETTMANN

London Underground Travolator (1960)



UPI/CORBIS-BETTMANN

Amphicar (1961)

THE CAR BOAT

Seagoing cousins to the car plane, amphibious roadsters have been reinvented many times. They do eliminate the headaches of towing a boat to its launch site—the boat can drive to the shore under its own power. But few people want to sacrifice the convenience of proper cars or boats for the dubious merits of Davy Jones's convertible.

THE ATOMIC-POWERED PLANE

Background: After the Manhattan Project, the U.S. Air Force and the Atomic Energy Commission collaborated to develop aircraft propelled by nuclear power. An onboard reactor would have provided thrust by superheating incoming air. In theory, a nuclear-powered bomber would have tremendous strategic advantages: it could jet along at high speeds, and its range was virtually unlimited because it never needed to refuel. It might fly for years without landing.

Problems: The twin bugaboos were weight and radiation. Building reac-

tors compact enough to fit on an aircraft was a challenge, although contractors did test some promising designs. But reactors need shielding, not only to protect the crew and the outside environment but their own critical systems. Adequate shielding raised the plane's weight prohibitively. In one early design, for example, the propulsion system would have reportedly weighed more than 80 tons, of which five tons was reactor and almost 50 tons was shielding.

Aside from these technical problems, the program was dogged by

poor organization and political unpopularity. Outsiders were understandably leery of flying a potential atomic disaster over populated areas. Ballistic-missile technology also raced forward faster than anticipated, which diminished the cold war need for a nuclear-powered bomber.

Status: President John F. Kennedy canceled the program in 1961. More than \$1 billion was spent on it over the years, and it never produced a working test aircraft.

THE ATOMIC CAR

The U.S. government briefly sponsored a project that had even less *raison d'être* than an atomic plane: an atomic car. Long after the research was defunct, auto designers continued to roll out fanciful chassis for futuristic atomic-powered cars, such as the 1958 Ford Nucleon. Just drop a small reactor in the back and drive 'er out of the showroom.



Ford Nucleon (1958)

UPI/CORBIS-BETTMA

ZEPPELINS

Background: The majestic passenger zeppelins that graced the skies before World War II were, hands down, the most dreamily luxurious craft that ever flew. Thousands of the well-to-do traversed the Atlantic on board the *Graf Zeppelin* and its successor, the *Hindenburg*. Reverence turned to horror, however, when the

latter burst into flames while landing at Lakehurst, N.J., on May 6, 1937. A *Graf Zeppelin II* followed, but it was dismantled by 1940, and the facilities in southern Germany that had constructed these craft were obliterated during the war.

Problems: The highly flammable hydrogen that filled the envelope of

the zeppelins posed an obvious danger. Yet, ironically, this past spring a new study by Addison Bain, formerly of the National Aeronautics and Space Administration, and Richard G. Van Treuren suggested that the real cause of the *Hindenburg's* fiery end was static igniting the envelope's chemically treated canvas. Safety aside, zeppelins also could not compete for passengers with airplanes, which were faster and less expensive.

Status: Zeppelins could be poised for a comeback of sorts. This past May, Luftschiffbau Zeppelin unveiled its new technology craft, a helium-filled airship 243 feet long that seats 12 passengers. Although they are unlikely to steal customers from commercial airlines—the cost and speed disadvantages remain—modern zeppelins are being ordered for tourism and for scientific applications. (Their buoyancy suits them for observational jobs that involve hovering in place for long periods.) With any luck, zeppelins will be among the few vehicles that ever traveled to oblivion—and made the trip back.

SA



Landing in Los Angeles (1929)

UPI/CORBIS-BETTMA



Hybrid Electric Vehicles

They will reduce pollution and conserve petroleum. But will people buy them, even if the vehicles have astounding fuel efficiency?

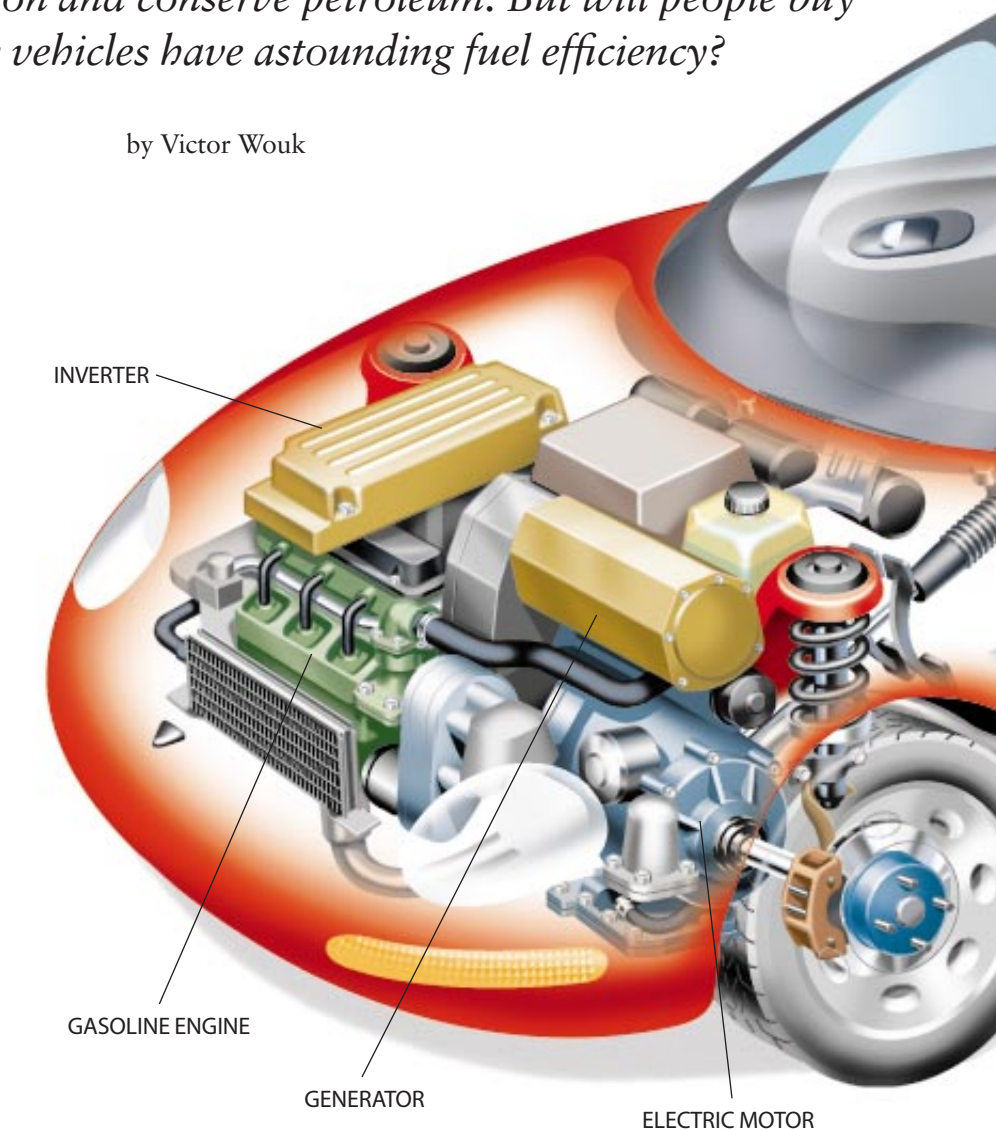
by Victor Wouk

To paraphrase Mark Twain: “Everybody talks about hybrids, but nobody does anything about them.” Okay, the assertion is a bit overstated. There are actually hundreds of engineers around the world who are working on hybrid electric vehicles. But almost a century after the hybrid was first conceived, more than 25 years after development work began on them in earnest, and after more than \$1 billion has been spent worldwide in recent years on development, not a single hybrid vehicle is being offered to the general public by a large automaker. In fact, not a single design is anywhere near volume production. In the U.S., where the government has spent about \$750 million since 1994 on almost frenzied efforts to advance the technology of hybrid electric vehicles (HEVs), the concept is still a political football rather than a commercial reality.

Why does this lack of progress matter? Because many experts believe the HEV could be—and, in fact, *should* be—the car of the near future. In simple terms, an HEV is an electric car that also has a small internal-combustion engine and an electric generator on board to charge the batteries, thereby extending the vehicle’s range. The batteries may be charged continuously or only when they become depleted to some level.

Thus, HEVs do not share an electric vehicle’s main drawback: limited range between chargings. The few thousand electric vehicles on the roads in the U.S. today can travel only about 130 kilometers (roughly 80 miles) before their batteries need recharging, which can take anywhere from three to eight hours. These facts mean that an HEV can have the best of both worlds: it can function as a pure electric vehicle for relatively short commutes while retaining the capability of a conventional automobile to make long trips.

The power of a hybrid’s internal-com-

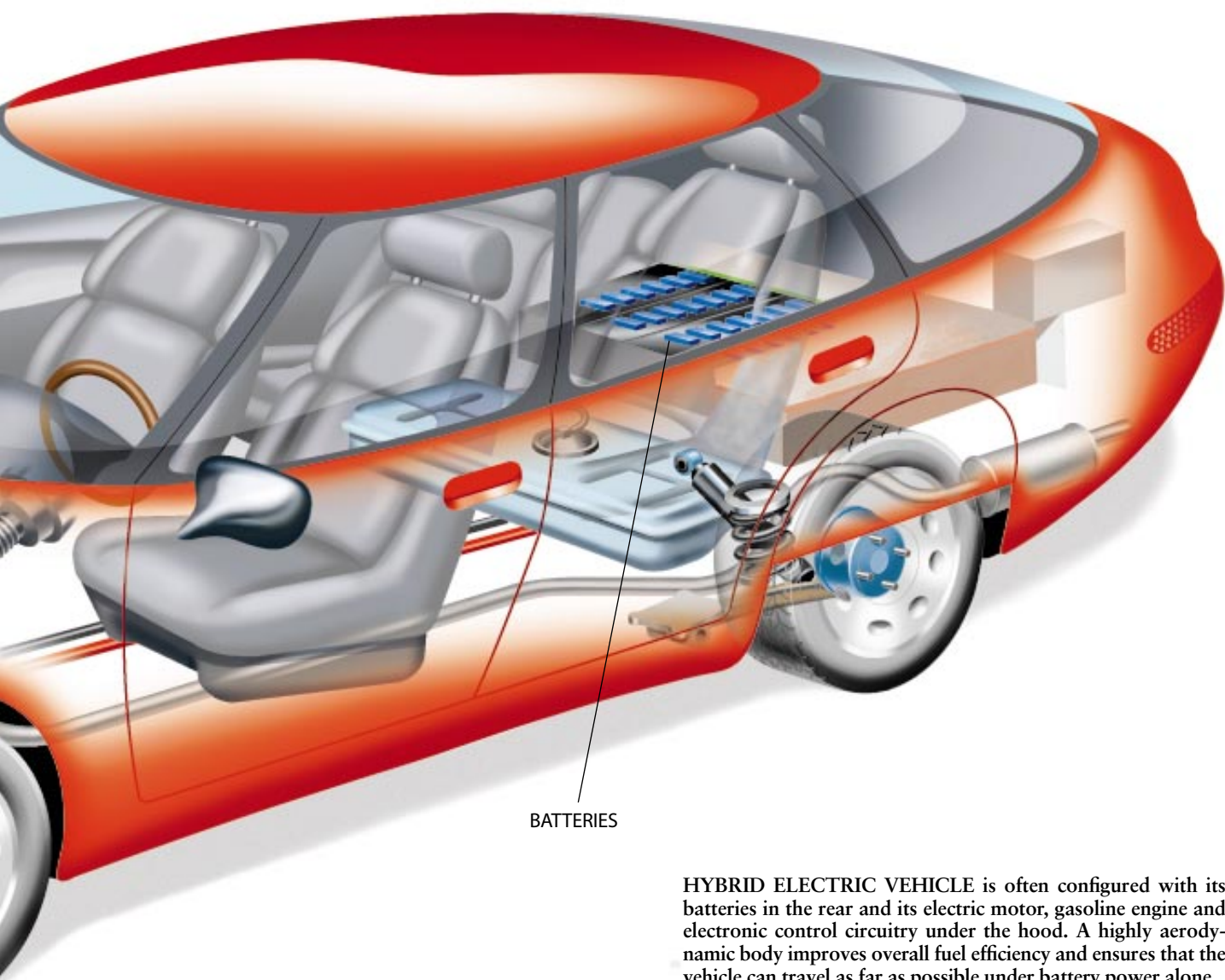


bustion engine generally ranges from one tenth to one quarter that of a conventional automobile’s. This engine can run continuously and efficiently, so although an HEV, when its internal-combustion engine is running, emits more pollutants than a pure electric, it is much cleaner than a conventional car. In fact, a hybrid can be made almost as “clean” as a pure electric. When pollution from the generating sources that charge its batteries is taken into account, an electric vehicle is about one tenth as “dirty” as a conventional car with a well-tuned engine. An HEV, in comparison, can be about one eighth as polluting. With good design, moreover, HEVs can achieve sev-

eral times the fuel efficiency of a gasoline-powered vehicle. Thus, if HEVs ever do become a success in the U.S., there could be benefits both for the environment and for the balance of trade: imported petroleum now accounts for almost half of the country’s consumption.

Here We Go Again

The HEV concept goes back to 1905. On November 23 of that year, American engineer H. Piper filed for a patent on a hybrid vehicle. Piper’s design called for an electric motor to augment a gasoline engine to let the vehicle accelerate to a rip-roaring 40 kilometers



GEORGE RETSECK

BATTERIES

HYBRID ELECTRIC VEHICLE is often configured with its batteries in the rear and its electric motor, gasoline engine and electronic control circuitry under the hood. A highly aerodynamic body improves overall fuel efficiency and ensures that the vehicle can travel as far as possible under battery power alone.

(25 miles) per hour in a mere 10 seconds, instead of the usual 30. But by the time the patent was issued, three and a half years later, engines had become powerful enough to achieve this kind of performance on their own. Nevertheless, a few hybrids were built during this period; there is one from around 1912, for example, in the Ford Museum in Dearborn, Mich.

The more powerful gasoline engines, along with equipment that allowed them to be started without cranks, contributed to the decline of the electric vehicle and of the nascent HEV between 1910 and 1920. In the early to mid-1970s, though, a brief flurry of interest and

funding, prompted by the oil crisis, led to the construction of several experimental HEVs in the U.S. and abroad.

During this time, I and a partner, Charles Rosen, built an HEV using my own funds and those of an investor. We outfitted the vehicle, a converted Buick Skylark, with eight heavy-duty police-car batteries, a 20-kilowatt direct-current electric motor and an RX-2 Mazda rotary engine. In 1974 it was tested at the Environmental Protection Agency's emissions-testing laboratories in Ann Arbor, Mich. The vehicle was optimized for low pollutant emissions, not for good fuel economy. Still, on the highway and with the batteries discharging,

the vehicle got nearly 13 kilometers per liter (30 miles per gallon)—more than twice the fuel economy of the vehicle before it was converted.

The vehicle's emission rates (per kilometer) of 1.53 grams of carbon monoxide, 0.5 gram of nitrogen oxides and 0.21 gram of hydrocarbons were only about 9 percent of those of a gas-powered car from that era. The project showed that a pair of determined individuals could use readily available and proved technologies to build quickly an HEV that met the requirements of the Clean Air Act of 1970. (As it happened, Detroit's conventional automobiles did not meet these requirements until 1986.)

My vehicle and others were described in numerous articles in the technical and lay press. In one report, I showed how with modest improvements my HEV could wring 21 kilometers out of a liter of fuel (50 miles per gallon). Nevertheless, interest in, and funding for, HEVs began to wane almost as soon as oil became plentiful again.

The dormancy went on until 1993, when the Clinton administration announced the formation of the Partnership for a New Generation of Vehicles (PNGV) consortium, which includes the "Big Three" automakers and about 350 smaller technical firms. Its members are spending about \$500 million a year—including \$250 million in federal funds—to develop a car that can travel 34 kilometers per liter (80 miles per gallon) of gasoline. Such a vehicle would be about three times as efficient as today's comparable, gas-fueled, midsize cars. Moreover, the efficiency is to be achieved without any sacrifices in performance or safety and in a vehicle that does not cost significantly more and emits perhaps one eighth of the pollutants.

The PNGV never specified that its supercar had to be an HEV that used an internal-combustion engine as its second power source. Indeed, the HEV is only one kind of hybrid; other possibilities include vehicles that have a fuel cell and a battery or that have an internal-combustion engine and a flywheel [see "Flywheels in Hybrid Vehicles," by Harold A. Rosen and Deborah R. Castleman, page 75]. Practically speaking, however, only the internal-combustion engine and battery combination has any chance of meeting the PNGV's stringent requirements in the near future.

Pick Your Configuration

Even settling on this type of HEV does not end the choices. The wide variety of possible engine-battery HEV configurations fall into two basic categories: series and parallel [see *illustration at right*]. In a series hybrid, the internal-combustion engine drives a generator that charges the batteries, which power the electric motor. Only this electric motor can directly turn the vehicle's driveshaft. In a parallel hybrid, on the other hand, either the engine or the motor can directly torque the driveshaft. A parallel HEV does not need a generator, because the motor serves this function. (When the engine turns the driveshaft, it also spins the motor's rotor when the

clutch is engaged. The motor thus becomes a generator, which can charge the batteries.)

Both the parallel and the series hybrid can be operated with propulsion power coming only from the battery (in an all-electric mode), with power supplied only by the internal-combustion engine (in a series hybrid, this power must still be applied through the generator and the electric motor), or with power from both sources. One advantage of the parallel scheme is that a smaller engine and motor can be used, because these two components can work together.

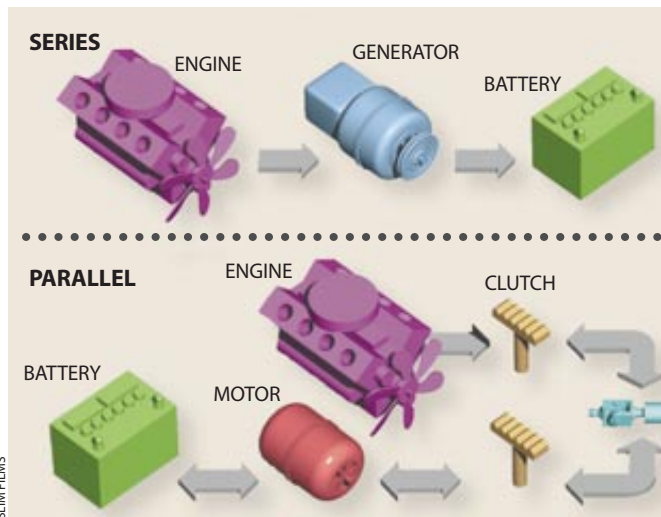
Disadvantages of the parallel configuration include the fact that the designer no longer has the luxury of putting the internal-combustion engine anywhere in the vehicle, because it must connect to the drivetrain. In addition, if a parallel hybrid is running electrically, the batteries cannot be charged at the same time, because there is no generator.

The distinct features of the two types of HEV suit them to different driving needs. Briefly: a series hybrid is generally more efficient but less powerful than a parallel HEV. So if the car is to be used for a daily commute of 35 kilometers or less each way, and perhaps the odd longer trip every now and then, a series HEV will do just fine. On the other hand, if the vehicle is to function—and feel—more like a conventional, gasoline-powered car, then a parallel hybrid may be necessary. Whereas a series hybrid might very well suffice for a mission involving many short trips a day (to make deliveries, say), a parallel hybrid would be preferable for heavy highway use, where bursts of speed are necessary for passing. A parallel HEV can also usually muster more speed on hills than a series hybrid, if both have been designed basically for moderate performance.

Yet power is not everything. For the average commute over terrain that is not too rugged, a series hybrid will get you there and back at higher efficiencies. In a series HEV the internal-combustion engine can be restricted so that it avoids rapid changes in speed and load, which cause surges in pollutant emissions. By running constantly and in a limited range, the engine can operate in its most

fuel efficient range, thereby using less fuel during a given driving mission [see *illustration on page 74*].

Another attractive feature of the series hybrid is that it can have a long range with a surprisingly small engine-generator set. In 1986 Roy A. Renner and Lawrence G. O'Connell of the Electric Power Research Institute in Palo Alto, Calif., did some calculations for a pure electric van with a 40-kilowatt electric motor and a bank of batteries capable of storing 34 kilowatt-hours. Renner and O'Connell found that the van's



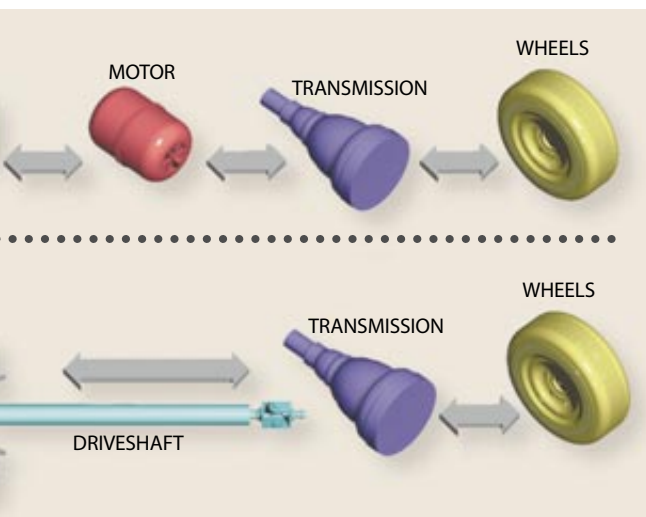
range of 100 kilometers could be doubled if the vehicle were converted into a series hybrid with a gas-powered engine-generator set capable of putting out a mere three kilowatts.

This doubling of range occurred with the engine-generator running continuously. When the vehicle was moving, all the generator current went to the motor, reducing the drain from the batteries. When the vehicle was standing still, the batteries charged slightly but not enough to replace the charge lost during starts and acceleration.

If the vehicle's battery bank could store only 17 kilowatt-hours instead of 34, then a six-kilowatt engine-generator set would be needed to achieve the same doubling of range, to 200 kilometers. More charging of the batteries would have to take place when the vehicle was not moving. In the extreme case, with a battery only big enough to help accelerate the vehicle, a 7.5-kilowatt (10 horsepower) engine-generator would be needed. In that situation, the batteries would never become depleted, because the 7.5-kilowatt engine-generator would be supplying the entire average load. In com-

parison, the engine of a conventional small car puts out about 75 kilowatts.

It should be noted that Renner and O'Connell obtained these figures by running computer simulations and putting experimental vehicles through a standard automotive test cycle of acceleration, cruising and stopping on a level surface. To climb a hill or pass another car on the highway (which requires a burst of power), the motor and engine-generator would need 50 to 100 percent more power than the 40 kilowatts and 7.5 kilowatts mentioned above.



TWO TYPES OF HYBRID are designated series and parallel. In the series type, a gasoline engine drives a generator that charges the batteries that power the electric motor, which turns the wheels. Only this motor can turn the wheels. In the parallel scheme, the gasoline engine or the electric motor—or both—can turn the wheels.

Despite these attractive features of the series configuration, all signs are that the PNGV consortium will choose a parallel HEV as its first prototype. The high degree of similarity to conventional cars that the PNGV is aiming for would be difficult to achieve in a series HEV. For a parallel HEV of about 1,000 kilograms, similar to one of today's midsize cars, a 100-kilowatt internal-combustion engine would be needed. Before such a vehicle could meet the desired specifications, though, major improvements in a host of "enabling technologies" will be necessary. These innovations include lighter-weight bodies, more efficient engines, better batteries and more efficient electric motors and generators. Almost a quarter century ago, I wrote in several papers and reports that all these improvements were necessary.

The PNGV's approach to technology

development has been centered on a three-year evaluation program. At the end of this year there is to be an announcement of technologies that will be supported further.

In a report released this past April the National Research Council (NRC) faulted the PNGV for not focusing its effort sharply enough on the most promising technologies needed to meet its goals. "Failure to address this issue may ultimately jeopardize the program," according to the NRC. Specifically, more attention and funds should be aimed at improving lithium-ion and nickel-metal-hydrate batteries, electronic systems and lightweight, low-cost diesel engines, the NRC committee wrote in the report.

In the U.S., serious development of HEVs is supported by joint industry and government funding, mainly through the PNGV. Ford, General Motors and Chrysler are designing both series and parallel hybrids, which the companies aim to have available for production in the reasonably near future. Success will depend on significant improvements in certain components, especially batteries. One of the few real success stories has been the use of nickel-metal-hydrate batteries, produced by GM Ovonic and Energy Conversion Devices in Troy, Mich., with funds from industry and the PNGV. A hybrid with nickel-metal-hydrate batteries will go twice as far, under battery power, as will an identical HEV with the same weight of lead-acid batteries.

A Very Tall Order

How likely is it that the PNGV consortium will meet its goals? Not very. By 2000 the PNGV is supposed to have test vehicles running; by 2004 the consortium must have a production-ready prototype that has a fuel efficiency of three liters per 100 kilometers (80 miles per gallon), low emissions and the same performance, cost and safety as a conventional car. Given the nearness of these deadlines, the consortium has little choice but to build this "supercar" without any intermediate steps, such as a vehicle with, for example, a fuel effi-

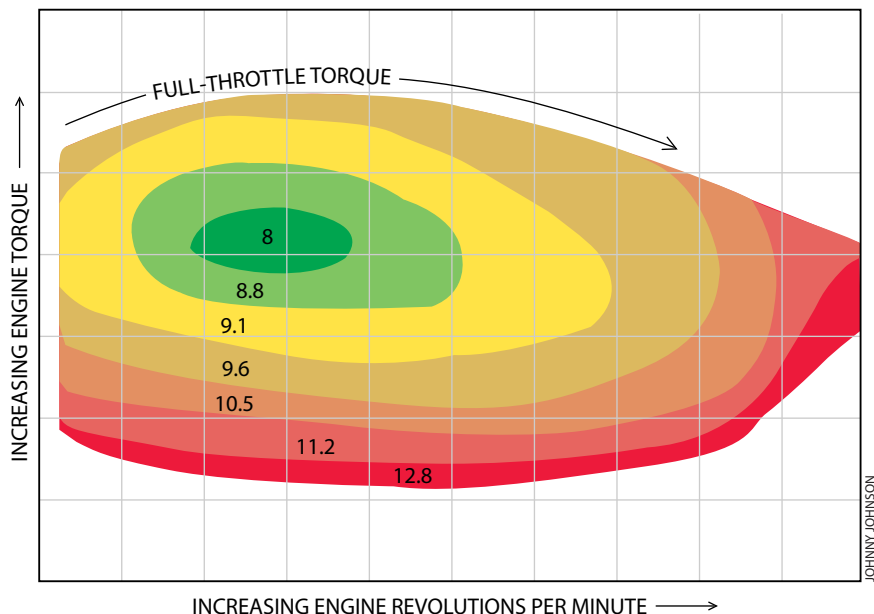
ciency of four or five liters per 100 kilometers. It would be as though the first manned space launch had to travel all the way to the moon and back.

Moreover, satisfactory procedures for testing an HEV still have not been devised. For example, one proposal would require that tests to determine the fuel economy of an HEV begin and end with the vehicle's batteries at the same state of charge. In calculating the fuel economy, the electricity that was used to recharge the batteries would be converted to an energy equivalent and added to the fuel consumed during the test. At first glance, the methodology seems reasonable. But by simply lumping the two energy amounts together, the calculation method basically ignores the shift in energy source, from onboard gasoline to electricity generated off board—which is the whole point of alternative vehicles.

If the PNGV seems to be going a bit off course, what about programs in other countries? Small HEV fleets are being demonstrated worldwide. A successful one in Japan, where HINO Motors has produced about a dozen HEV buses, is part of an effort to eliminate the particulate emissions that come from diesel engines during acceleration. The diesel is assisted by an electric motor and nickel-cadmium batteries during acceleration, eliminating the smoke. The batteries are charged during runs from stop to stop and by regenerative braking.

In general—and in stark contrast to the PNGV—European and Japanese HEV development is emphasizing existing or modestly improved technology. To a greater extent than their U.S. counterparts, the Europeans and Japanese are concentrating on ways of reducing production costs and making HEVs more marketable in the near term. Volkswagen, Mitsubishi and Toyota, among others, are developing HEVs with their own money. A two-year demonstration of 20 Volkswagen parallel HEVs in Zurich recently showed, again, that lower emissions and lower fuel consumption are simultaneously possible.

Regardless of the country in which they are built, whether or not HEVs (or, indeed, any alternative vehicles) succeed will depend on the relative costs of buying and operating them. And the operating cost will in turn depend on the price of gasoline. The formula is simple: the higher the price of gasoline, the more likely people will be to seek alternatives. Although it is true that there is



EFFICIENCY OF THE INTERNAL-COMBUSTION ENGINE is highest for a small range of values of torque and rotational speed—those in the darker green “sweet spot.” At this level of efficiency, if the engine were propelling a vehicle, it might burn eight liters of gasoline per 100 kilometers. A series-type hybrid can be designed so that its engine operates only in this highest-efficiency mode; a parallel hybrid can be designed so that its engine stays within the efficiency represented by the dark and light green regions.

virtually no history of HEV sales to analyze, the short, recent history of electric vehicle sales suggests that gasoline prices must go much higher indeed before people rent or buy these cars.

Drivers in Europe and Japan pay about three times as much for gasoline as do motorists in the U.S. Nevertheless, relatively few electric vehicles have been sold in those places. Despite generous government and manufacturers’ subsidies, sales of electric vehicles do not constitute even 1 percent of automobile sales anywhere in Europe or Japan. In France, plans by Peugeot and Renault to sell several thousand electric vehicles

in 1996 and 1997 have fallen short of those goals.

GM’s flashy and peppy EV1, introduced last December in southern California and Arizona, is meeting with more modest success than had been hoped. The few hundred vehicles on the road are being driven mainly by environmentally conscious people who have multiple vehicles and who might be called Greens with plenty of green.

It remains to be seen whether government mandates can do what subsidies and aggressive marketing have so far been unable to achieve. Specifically, in 1990 the California Air Resources Board

(CARB) mandated that by 1998, 2 percent of cars sold by the U.S. Big Three automakers and by Japan’s “Big Four” be so-called zero-emission vehicles. Electric vehicles were then, and still are, the only viable vehicle type that emits no pollutants as it is driven. Unfortunately, the batteries now available commercially do not provide the kind of range that the average consumer seems to demand, even from a second car. The HEV is an obvious alternative. Although CARB initially refused to consider HEVs, it now deems them acceptable, albeit with complicated rules governing the determination of their emission levels and fuel consumption.

It is even possible that in the future CARB or some other body might simply mandate that HEVs make up a certain percentage of vehicle sales by some date. Although the scheme was tried and not very well received for pure electrics, there would be a critical difference for HEVs: manufacturers could not reasonably complain that the public will not buy HEVs because of inadequate range or performance.

Will the HEV finally unite consumer acceptance, higher fuel economy and reduced emissions? It certainly will if political problems (another war in the Persian Gulf, for instance) or some other shock sends the cost of petroleum spiraling. But before an emergency forces us into a crash program, why don’t we try going about this in a rational way? Let us build reasonable—and mass-producible—HEVs that get at least 21 kilometers out of a liter of fuel (50 miles per gallon) but still drive like conventional cars. And let us not give up on the project until the cars can go 34 or more kilometers on a liter of fuel—and, of course, until people are buying them. SA

The Author

VICTOR WOUK is a New York City-based consultant on hybrid electric and electric vehicles. He holds 10 patents on electrical and electronic devices and systems, including ones related to the speed and braking control of electric vehicles. After founding two successful companies in the 1960s, he devoted himself full-time to hybrid and electric vehicles, building, among other things, an experimental hybrid electric vehicle in the early 1970s. He is currently the U.S. technical adviser to the International Electrotechnical Commission’s committee on electric road vehicles and is also active in the New York Academy of Sciences.

Further Reading

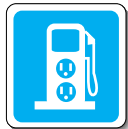
AN EXPERIMENTAL ICE/BATTERY-ELECTRIC HYBRID, WITH LOW EMISSIONS AND LOW FUEL CONSUMPTION CAPABILITY. Victor Wouk. Society of Automotive Engineers, Warrendale, Pa. (Congress and Exhibition, Detroit, Mich.) Publication SAE No. 760123, February 1976.

THE HYBRID CAR REVISITED. Roy A. Renner and Lawrence G. O’Connell in *Proceedings of the 8th International Electric Vehicle Symposium*, Washington, D.C., Oct. 21–23, 1986, pages 219–227. U.S. Department of Energy, Report No. CONF-8610122, 1986.

HISTORY OF THE ELECTRIC AUTOMOBILE: BATTERY-ONLY POWERED CARS. Earnest H. Wakefield. Society of Automotive Engineers, Warrendale, Pa., 1994.

EVs UNPLUGGED: HYBRID ELECTRIC VEHICLES BOOST RANGE, BUT NOT POLLUTION. Victor Wouk, Bradford Bates, Robert D. King, Kenneth B. Hafner, Lembit Salasoo and Rudolph A. Koegl in *IEEE Spectrum*, Vol. 32, No. 7, pages 16–31; July 1995.

POLICY IMPLICATIONS OF HYBRID-ELECTRIC VEHICLES. Final Report to NREL, Subcontract No. ACB-5-15337-01. J. S. Reuhl and P. J. Schuurmans. NEVCOR, Inc., Stanford, Calif., April 22, 1996. Available at <http://www.hev.doe.gov/papers/nevcor.pdf> on the Web.



Flywheels in Hybrid Vehicles

A rapidly spinning flywheel combines with a gas-turbine engine to power a novel hybrid electric vehicle

by Harold A. Rosen and Deborah R. Castleman

The search for an alternative to the internal-combustion engine used in today's cars is motivated by two societal concerns: the need to reduce fossil-fuel consumption and the need to reduce air pollution. Unfortunately, most car buyers do not make their purchases based on these criteria. Instead, when looking for a new automobile, most consumers consider issues such as cost, safety, performance and fuel efficiency. (This last factor does, of course, have an effect on fuel consumption and pollution, but it is rarely a car buyer's primary concern.)

In 1993 one of us (Rosen), along with his brother, Benjamin, founded Rosen Motors with the goal of producing a new type of powertrain for cars that would not only address concerns about pollution and fuel efficiency but would also be something that consumers would actually want to own.

Over the past four years, Rosen Motors has been developing a hybrid vehicle that incorporates a rather unusual technology—the flywheel. Although the concept of a flywheel is quite simple, the implementation has been difficult. The flywheel in our powertrain consists of a spinning cylinder made of a high-strength, carbon-fiber composite that can both store and generate energy. The faster the flywheel spins, the more energy it retains. Energy can be drawn off as needed by slowing the flywheel.

Like all hybrids, the automobile developed at Rosen Motors draws power from two separate sources. In our hybrid, we use a flywheel and a gas-turbine engine that is akin to a miniature jet engine. [For an overview of the technology behind hybrids, see "Hybrid

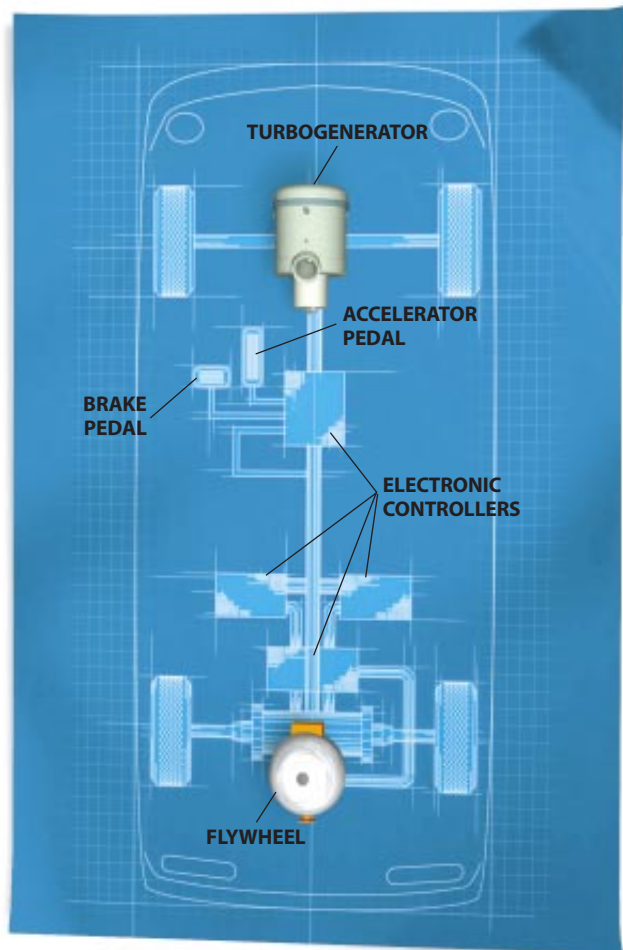
Electric Vehicles," by Victor Wouk, page 70.] Both the flywheel and the turbine have electric generators attached; we refer to the combination of the turbine and the generator as a turbogenerator.

These two power sources are better than one internal-combustion engine. High-power internal-combustion engines found in today's cars provide high acceleration but poor fuel economy, whereas low-power engines yield better fuel economy but poor acceleration. In addition, noxious emissions are an unavoidable by-product of operation.

In the hybrid electric powertrain developed at Rosen Motors, the turbogenerator propels the car while cruising, and it also recharges the flywheel, which we use to supply bursts of power for acceleration. In addition, the flywheel has been set up so that during braking it will recover energy that would otherwise be lost to friction.

The advantages of the flywheel lie mainly in its efficiency: chemical batteries that could generate and recapture the same power as the flywheel would weigh considerably more and would recover and reuse only half as much energy during stop-and-go driving. Furthermore, when the flywheel, rather than the combustion engine, is used to supply power for acceleration, the peak power required from the engine drops. As a result, the turbine engine can be smaller and lighter.

We selected a gas turbine because the system emits inherently low levels of pollutants; indeed, these emissions ap-



TURBINE-FLYWHEEL CAR responds when the driver steps on the accelerator or the brake, prompting electronic controllers to divert power from, or to, the flywheel or the turbogenerator. The flywheel supplies most of the power for acceleration and absorbs energy otherwise lost during braking; the turbogenerator is mostly used during cruising and for maintaining the flywheel at its optimum speed.

proach zero when catalytic combustion is used. The turbogenerator can be small and relatively simple and thus will have a long, reliable service life. The turbine runs on unleaded gasoline, so car owners can use existing gasoline stations.

On January 5 of this year, we watched the first successful test drive of a turbine-flywheel-powered automobile. We are now working on improved versions of the flywheel, turbogenerator and other components of the powertrain. In the near future we plan to operate the hybrid powertrain in a converted luxury sports sedan to demonstrate the acceleration, fuel economy and low emissions that are possible with the Rosen Motors turbine-flywheel-powered hybrid electric vehicle.

The Flywheel and How It Works

MAGNETIC BEARINGS The flywheel designed by Rosen Motors is made of a titanium hub and a high-strength, carbon-fiber-composite cylinder that can spin as fast as 60,000 revolutions per minute. To reduce friction at these speeds, the flywheel spins without touching anything. Magnetic bearings support the flywheel and preserve the tight clearances—as small as 0.005 of an inch (0.013 centimeter)—between the rotating and nonrotating parts of the assembly even as the car rides over bumps and potholes. The energy consumed by the magnetic bearings must be low enough so it does not discharge the vehicle's 12-volt batteries when the car is parked and the turbine is off. (These batteries supply power for accessories such as the radio and headlights.) To get energy in and out, the flywheel must include a motor generator; the motor rotor of the generator is attached to the central shaft of the flywheel cylinder.

UPPER MAGNETIC-BEARING SYSTEM

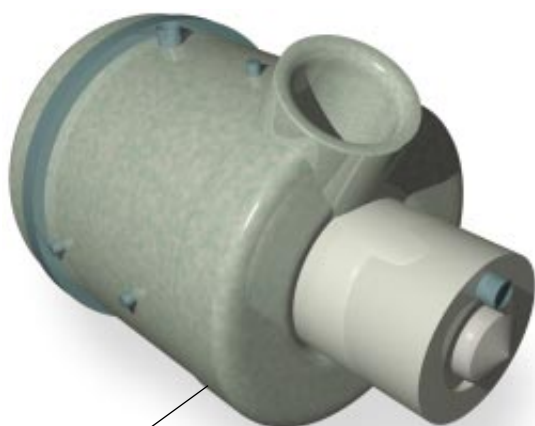
CENTRAL SHAFT

FIBER-COMPOSITE CYLINDER

MOTOR ROTOR

GIMBAL RING

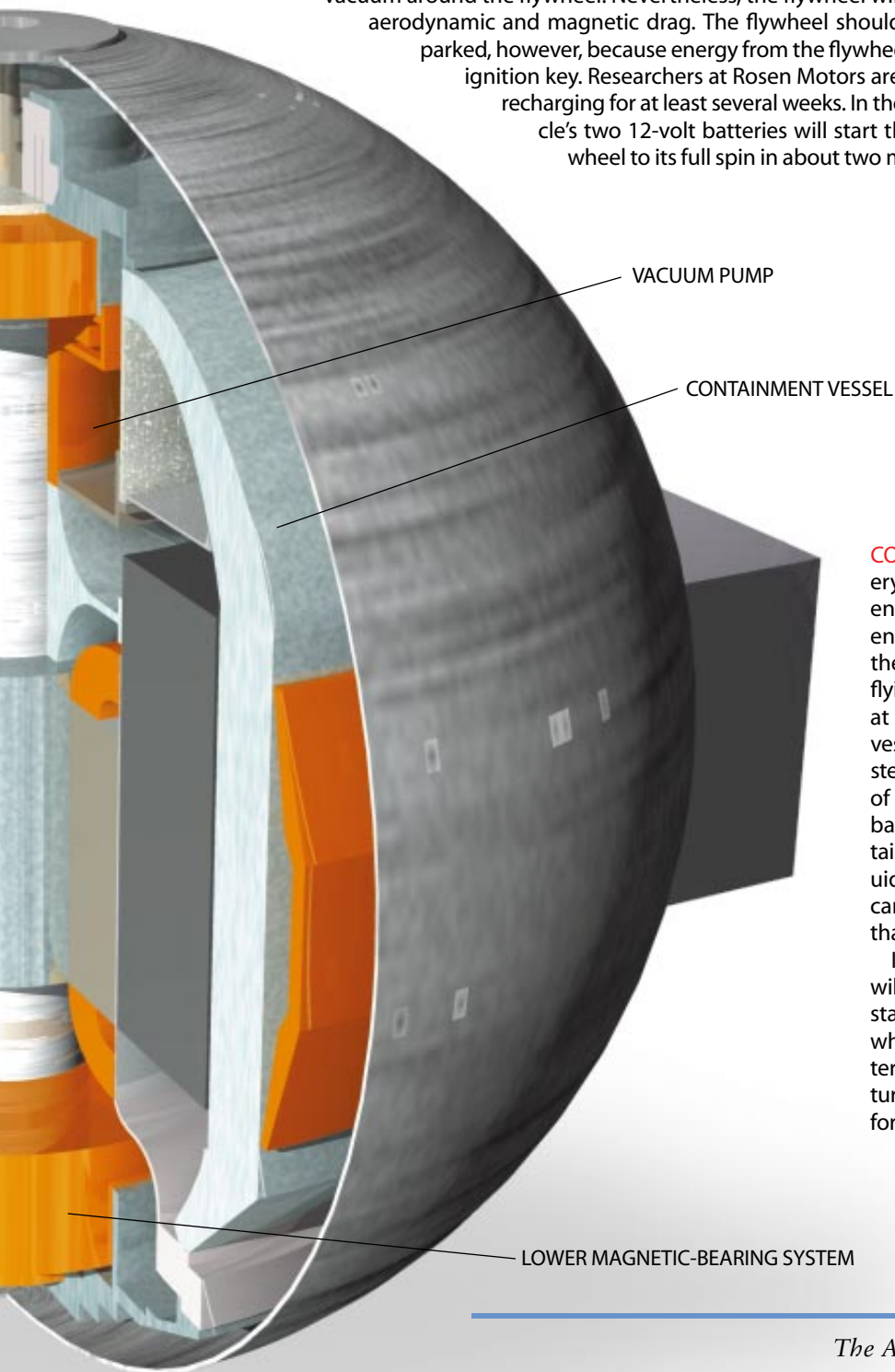
GIMBALS Theoretically, the rapid revolution of the flywheel could generate sufficient gyroscopic forces to interfere with the handling of the vehicle as well as to overload the magnetic bearings. A system of gimbals therefore cradles the flywheel assembly, isolating the spinning cylinder from the rotational motions of the vehicle.



TURBOGENERATOR

TURBOGENERATOR The turbogenerator, which is being developed by Capstone Turbine Corporation in Tarzana, Calif., consists of a clean-burning gas turbine (the type of engine used in jets) that drives an internal electric generator. Energy from the turbogenerator is used to keep the flywheel spinning at the proper speed. This turbogenerator is an advanced version of the turbogenerators now in production at Capstone for such applications as auxiliary power generators for buildings.

VACUUM PUMP Because aerodynamic drag can slow the flywheel and generate a considerable amount of heat, a vacuum pump consisting of a lightweight molecular drag pump and molecular sieves maintains a vacuum around the flywheel. Nevertheless, the flywheel will lose energy over time as a result of residual aerodynamic and magnetic drag. The flywheel should remain spinning even when the vehicle is parked, however, because energy from the flywheel starts the turbine when the driver turns the ignition key. Researchers at Rosen Motors are developing a flywheel that will run without recharging for at least several weeks. In the event the flywheel does run down, the vehicle's two 12-volt batteries will start the turbine, which will then recharge the flywheel to its full spin in about two minutes.



CONTAINMENT VESSEL For safety reasons, every high-speed rotating system, from huge jet engines to smaller flywheels, must be properly encased. Otherwise, in the unlikely event that the system breaks down, debris would be sent flying outward with considerable force. Workers at Rosen Motors have created a containment vessel of carbon-fiber-composite reinforced steel that will safely contain the flywheel in case of failure. Should such an event occur, the gimbal supports will pull away, allowing the containment vessel to spin to a stop in a cooling liquid that surrounds it. In that way, the flywheel can dissipate its energy relatively slowly, rather than imparting a sudden jerk to the vehicle.

In case of a crash, the containment structure will remain intact because it is designed to withstand the forces released if the flywheel breaks, which are much higher than the forces encountered during a collision. The containment structure itself is anchored to the car with Kevlar-reinforced, high-strength straps.

The Authors

HAROLD A. ROSEN and DEBORAH R. CASTLEMAN both work at Rosen Motors in Woodland Hills, Calif. Rosen is president, chief executive officer and co-founder (with his brother, Benjamin M. Rosen) of the company. He was a vice president at Hughes Aircraft Company and led the team that developed the first geostationary satellite. Castleman is a vice president of the company. She received a master's degree in electrical engineering from the California Institute of Technology and was formerly a deputy assistant secretary of defense for the Clinton administration.



Automated Highways

*Cars that drive themselves
in tight formation might alleviate
the congestion now plaguing urban freeways*

by James H. Rillings



HANDS-FREE DRIVING has become a realistic prospect because the electronic devices required for such automation—magnetometers, video cameras, radar, lasers and computers—are now sufficiently inexpensive. Although it currently takes teams of research engineers to outfit a passenger car to travel without constant human supervision on specially modified roads, much of the gear required for that capability may soon find its way into ordinary vehicles equipped with advanced cruise control, navigation aids or traffic-warning indicators.

Highway travel is the lifeblood of modern industrial nations. But in the U.S., as in many other places, the larger roads are sorely overburdened: around major cities, heavy usage slows most peak-hour travel on freeways to less than 56 kilometers (35 miles) per hour. In all, excessive traffic causes more than five billion hours of delay every year; it wastes countless gallons of fuel and needlessly multiplies exhaust emissions.

The answer, one might imagine, is simply to lay more asphalt. Yet building new highways is enormously expensive, particularly in urban areas. For instance, the reconstruction of an 11-kilometer stretch of the Central Artery in Boston will require approximately \$8 billion. With such costs involved, it is not economically feasible to expand urban thoroughfares on a large scale. So if highway transportation is to keep pace with the growth of urban areas, people must somehow learn to use existing roadways more efficiently.

One possibility is to develop an automated highway system, a lane or set of lanes where specially equipped cars, trucks and buses could travel together under computer control. That effort need not demand some giant central computer to direct the movement of all vehicles. Rather networks of small computers installed in vehicles and along the sides of certain roadways could coordinate the flow of traffic, increasing both efficiency and passenger safety.

Such automation may, in fact, be the least expensive way to boost highway capacity. A typical freeway lane can handle about 2,000 vehicles per hour, but a lane equipped to guide traffic automatically should be able to carry about 6,000, depending on the spacing of entrances and exits. The savings brought about by not having to build more roads or not having to widen existing ones should more than pay for the sophisticated electronic equipment needed for cars to drive themselves.

As visionary as this notion might appear, self-driving cars are not a new concept. Indeed, a working model of an automated highway was the hit of the General Motors pavilion at the 1939 World's Fair in New York City. During the late 1950s and early 1960s, researchers at General Motors went on to refine various driverless vehicles. They showed, for example, how robotic trucks could work in open-pit mines. Then, during the 1960s and early 1970s, Robert E.

Fenton of Ohio State University demonstrated that wire-guided cars could operate successfully on a test track.

Although these early attempts at automation were valuable research exercises, the results proved too crude to be truly workable. Yet by the late 1980s, advances in microprocessors, wireless communications and various electronic sensors prompted many people to rethink the idea of automated highways. One group, which originally called itself Mobility 2000, convened in 1988 to consider the possibilities. It subsequently formed the Intelligent Vehicle Highway Society of America (later named the Intelligent Transportation Society of America), which now has more than 1,000 organizations as members. Its mission is to foster the introduction of various "intelligent" transportation systems, including automated highways.

The U.S. government has also been working toward this end. In 1991 Congress called for a prototype system to be tested by 1997—an experiment that my organization, the National Automated Highway System Consortium, has just carried out on a stretch of California freeway. That demonstration showed how automation might allow existing highways to accommodate a larger number of vehicles, while ensuring a higher degree of safety.

Driving on Autopilot

What might driving on an automated highway be like? The answer depends on what kind of system is ultimately adopted. Two distinct types are on the drawing board. The first is a dedicated lane system, in which certain lanes are reserved for automated vehicles. The second is a mixed traffic system: fully automated vehicles would share the road with partially automated or manually driven cars. A dedicated lane system would require more extensive physical modifications to existing highways, but it promises the greatest gains in freeway capacity.

Under either scheme, the driver would specify the desired destination, furnishing this information to a computer in the car at the beginning of the trip or perhaps just before reaching the automated highway. If a mixed traffic system was in place, automated driving could begin whenever the driver was on suitably outfitted roads. If dedicated lanes were available, the car could enter them and join existing traffic in two different

ways. One method would use a special on-ramp. As the driver approached the point of entry for the highway, devices installed on the roadside would electronically interrogate the vehicle to determine its destination and to ascertain that it had the proper automation equipment in good working order. Assuming it passed such tests, the driver would then be guided through a gate and toward an automated lane. In this case, the transition from manual to automated control would take place on the entrance ramp.

An alternative technique could employ conventional lanes, which would be shared by automated and regular vehicles. The driver would steer onto the highway and move in normal fashion to a "transition" lane. The vehicle would then shift under computer control onto a lane reserved for automated traffic. (The limitation of these lanes to automated traffic would, presumably, be well respected, because all trespassers could be swiftly pinpointed by authorities.)

Either approach to joining a lane of automated traffic would harmonize the movement of newly entering vehicles with those already traveling. Automatic control here should allow for smooth merging, without the usual uncertainties and potential for accidents. And once a vehicle had settled into automated travel, the driver would be free to release the wheel, break open the morning paper or just relax.

The centralized part of the system that governs access to the highway need not remain in constant command of the vehicles operating under its control. The responsibility for sensing dangers and managing movements could be shared among the vehicles and the automated roadway. For example, every car might be individually required to detect obstacles and to control its steering, braking and acceleration so as to avoid collision. Communication between vehicles might serve only for sharing information about traffic. In more demanding situations, computers monitoring the roadways could take on a supervisory role, assigning speeds and directing the passage of vehicles between the automated highway and local roadways to keep the flow running smoothly at all times.

Indeed, orchestration of traffic on automated highways could take many forms. The two ends of the spectrum of possibilities now envisioned are known as free-agent vehicles and platooned vehicles. Free-agent vehicles would, as the

name implies, operate independently. Each would drive so that it would be able to stop without mishap even if the vehicle ahead applied maximum braking—perhaps because an obstacle suddenly appeared on the road. The spacing required between two cars would depend on the braking capabilities of both vehicles, the condition of the road surface and the electronic reaction time of the controlling equipment (which would be considerably less than the reaction time of any driver).

Platooned vehicles, at the other extreme, would operate in closely coordinated groups to maximize the capacity of the highway. These vehicles would be linked together with a wireless local communications network, which could continuously exchange information about speed, acceleration, braking, obstacles and the like. Such constant interchange would enable the vehicles of the platoon to become, in essence, an electronically coupled train. Platoons might contain 10 to 20 automated cars operating, most likely, within dedicated lanes. But unlike a railway train, these chains would be dynamic—forming, splitting and rejoining again as traffic conditions and individual destinations demanded. Precise control would permit the spacing between members to shrink to only a few meters.

With cars so near one another, collisions might conceivably ensue if one vehicle unexpectedly slowed, for example, because equipment malfunctioned. But such rare collisions would involve small relative velocities and therefore probably cause only minimal damage. The gap between adjacent platoons would be sufficiently large to prevent them from crashing together even if the lead platoon stopped abruptly.

Whether a trip down this 21st-century highway is made as a free agent or as part of a platoon, each vehicle would eventually have to leave the automated lanes and return to normal driving. Reasonably enough, the process of exiting an automated traffic stream would very likely follow in reverse the steps required to enter it. But there is a wrinkle. Because the car or truck may have been under automatic control for some time, the driver might not be prepared to resume control as the vehicle approaches its intended exit. The driver might be preoccupied, asleep—or even dead. The automated highway system would have to handle all such situations. To do so, it might signal the approach of the exit



ANDY RYAN



ANDY RYAN

HIGHWAY CONGESTION clogs the Central Artery in Boston (*top*). Yet the construction there to improve traffic flow (*bottom*) is exceedingly costly, amounting to about \$8 billion for 11 kilometers of expanded highway. In the future, other cities facing traffic overloads may prefer to invest in automation (*right*), which would increase capacity without the need for building new roads.

and observe the behavior of the vehicle while the driver assumed limited duties. If he or she acted appropriately, the computer would release control completely, and the driver would complete the exit. But if the driver was not responding properly, warnings would sound. Should these promptings fail to wake the driver, the automated system would notify authorities of an emergency and bring the vehicle to a safe stop in a nearby holding area.

Loaded with Options

At first glance, the automated roadway of the future might look very much like many of today's highways. Yet sundry modifications would be required to transform current roads into automated freeways. That process would probably begin with the conversion of part of an existing highway, along with the construction of special ramps, transition lanes and barriers. Provisions would also have to be made for check-in areas and the diversion of





NATIONAL AUTOMATED HIGHWAY SYSTEM CONSORTIUM/VRW

Some Recent Tryouts

On August 7, 1997, the National Automated Highway System Consortium (NAHSC) began a four-day demonstration of the technical feasibility of automating highways. The proving ground was an 11-kilometer stretch of Interstate 15, just north of San Diego, Calif. Conveniently, two lanes of this highway are physically isolated from the others. They normally carry cars with two or more occupants during rush hours; at other times they are closed to traffic and serve the California Department of Transportation for various tests.

In preparation, members of the NAHSC installed digital communications equipment at the roadside and magnets down the center of both lanes. The first demonstration involved two buses and three cars, all acting as free agents. They performed a series of automated lane changes and passing maneuvers at highway speeds, using side- and rear-looking sensors to check for traffic. These vehicles also showed how they were able to follow with constant headway and avoid obstacles cooperatively. (When a lead vehicle detected something blocking the roadway, it transmitted this information to the vehicles behind, which then also changed lanes to avoid a collision.)

The second demonstration showed how greater levels of cooperation could increase highway capacity. Eight passenger cars started from rest and accelerated to highway speed while maintaining a constant spacing of approximately four meters between them. This eight-car platoon then split by opening the space between two of the vehicles to allow one vehicle to leave the platoon. At the end of the test lanes, the re-formed platoon slowed and came to a complete stop, all the while maintaining its tight spacing.

A third exercise showed some possible intermediate steps in the evolution to full automation. Two manually driven vehicles warned their drivers when the cars began to drift from their lanes. These cars also demonstrated adaptive cruise control: they slowed to maintain a comfortable distance to the vehicle in front and automatically managed acceleration and braking in stop-and-go traffic. The same two vehicles then switched to full automation, whereby they made lane changes and avoided obstacles cooperatively. Other demonstrations carried out in California during this period showed how highway automation could operate heavy trucks and maintenance vehicles and how such automation might vary in style between urban and rural settings.

These efforts extended the capabilities achieved recently in Japan. Engineers there operated automated vehicles in September 1996 along a six-kilometer section of isolated freeway near Komoro City. That demonstration showed how cars could be fully controlled using magnets embedded in the roadway, video cameras mounted on the automobiles and a radio antenna running beside the test lanes.

In Europe, much of the research focuses on automating commercial vehicles. The European Commission is currently funding a program to modify heavy trucks. It hopes to develop an "electronic towbar" that will allow one or more tractor trailers to follow automatically a truck being driven manually. These trucks would form a platoon of heavy vehicles with a single human driver. Road testing should begin next year. —J.H.R.

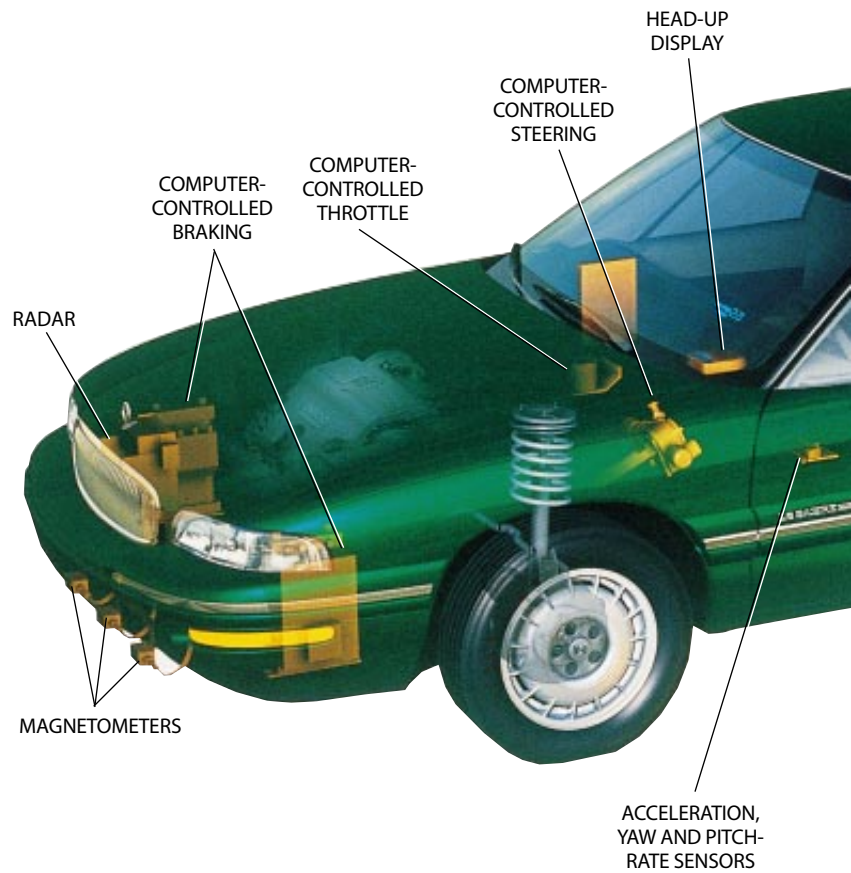


JAMES ARONOVSKY

AUTOMATED CRUISING in formation could be done without worry during trials that were recently conducted on a section of Interstate 15 in southern California.



ELECTRONIC EQUIPMENT required for automated operation includes magnetometers mounted below the bumpers (*top*), magnets embedded in the road (*bottom*) and an array of other sensors, actuators and computers installed inside the vehicle (*right*).



cars denied access. Finally, safe holding areas would have to be established near exits for the occasional driver who might be unable to resume manual control.

Although the roads themselves would have to contain a certain amount of specialized equipment, most of the new technology required for automated highway travel would be packed into future cars. For instance, a fully outfitted vehicle might employ a set of magnetometers designed to detect magnets embedded every meter or so at the center of each lane. In addition to providing a reference for steering, each magnet could convey one bit of information as the car passed over it (encoded by whether the north or south pole of the magnet was pointed upward). These digital data might, for example, inform an automated vehicle of its location or about upcoming curves in the road.

A forward-looking sensor, perhaps based on a millimeter-wavelength radar or an infrared laser, would detect dangerous obstacles and other vehicles ahead. Video cameras linked to computers that process images rapidly could also serve this function. Although the engineering required to develop these de-

VICES would be more challenging, such video equipment would also be able to track lane boundaries, eliminating the need for magnets and magnetometers.

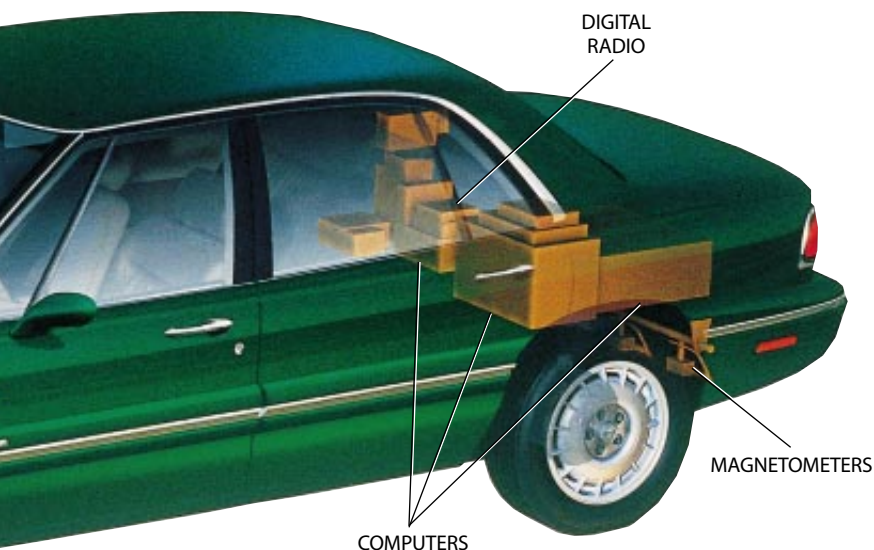
Accelerometers coupled to various actuators would manage steering, braking and throttle to maintain proper velocity and position. Digital radio equipment carried in each car would enable the computer on board to communicate with other vehicles in the vicinity and with supervisory computers monitoring the roadway. A display, perhaps projected onto the windshield in the manner of today's fighter aircraft, would give the driver information about the operation of the vehicle.

(Automated) Road Barriers?

Although such futuristic creations seem remote, most experts agree that the automation of highways is technically feasible, even with existing technology. Yet there are many other factors that might prevent such systems from ever becoming truly practical. Not the least of these concerns is the final price tag: the average car buyer must be able to afford the equipment for auto-

mation. Experience in the automotive industry shows that options for passenger cars costing more than about \$900 sell rather poorly. So this level is about the maximum one can expect consumers to pay for the special gear required for automated driving.

Yet such cost consciousness may not be as much of an impediment as it appears. The natural progression of automotive technology will undoubtedly introduce various forms of sophisticated electronics for improving safety or convenience. For example, video sensors that could control steering on an automated highway might first serve to alert a drowsy driver that the vehicle was leaving the lane or roadway—a condition responsible for the majority of fatalities on rural roads. And automakers might initially offer the forward-looking sensors needed to control the throttle and brakes in automated driving as part of an advanced cruise control (which is, in fact, being sold today in Japan) or as an early-warning indicator of potential collisions. Such equipment could even apply the brakes if the driver failed to do so in time. Safety experts hope such devices might lessen the dan-



BUICK MOTOR DIVISION

gers that arise through driver error and inattention, factors that now contribute to nine out of 10 crashes.

Two-way digital radios might also anticipate automated travel. They could notify drivers of hazardous road conditions, suggest optimum routing or, perhaps, automatically summon assistance during emergencies. Indeed, the "On-Star" navigation aid offered by General Motors and the "RESCU" radio system sold by Ford boast such features now.

Thus, the cost of the additional equipment for automation may not prove prohibitive, because new cars may slowly acquire most of the needed electronics anyway. But the success of automated highways will depend equally on an array of societal and institutional issues. It is unclear how willing automakers

will be to accept the potential liabilities involved in selling cars that drive themselves. Properly engineered, automated vehicles and roadways could be significantly safer than present-day highway travel, so overall liability costs should be reduced. But the proportion borne by various parties would shift: drivers' share of these costs would presumably decline, whereas the portion paid by manufacturers and highway agencies would rise.

Even if liability concerns could be addressed by new legislation and novel forms of insurance, one wonders whether people would want such a drastic change in driving. Certainly, no one desires that a large-scale system of automated highways be put into place overnight. Still, automated highways might develop incrementally, just as the auto-

motive technology involved slowly becomes more widespread.

A case study of such an evolution toward automated travel is, in fact, under way. At present, special lanes running down the center of the Katy Freeway in Houston are reserved for buses and passenger cars with two or more occupants. The public transit authorities in Houston are investigating adding automation so that platoons of buses could increase capacity on this freeway by acting as electronically coupled trains. A further step might allow properly equipped cars to join these platoons, perhaps after paying a premium for access to fast, effortless travel. No new roads would have to be built, yet highway usage could increase substantially. And no intermediate step in the process of conversion would require that a fully automated highway emerge in the end. Rather each advance would in itself make good technical and economic sense.

The automation of roadways in this manner could ultimately prove the easiest way to meet the growing demand for highway capacity around urban centers. But some people question whether society should necessarily accommodate that demand: Does the freedom to travel by personal automobile add more to the quality of life than arranging for convenient travel by rail or planning communities so that people can comfortably live where they work? The answer to this question will vary from place to place. And the recent efforts to demonstrate highway automation should help planners and government officials envision the full range of possibilities ahead as they make these important decisions. SA

A hyperlinked version of this article is available at <http://www.sciam.com/> on the SCIENTIFIC AMERICAN Web site.

The Author

JAMES H. RILLINGS is program manager of the National Automated Highway System Consortium, a group of government agencies and private companies collaborating to develop automated highways. Rillings received bachelor's and master's degrees in electrical engineering and a doctorate in systems engineering from the Rensselaer Polytechnic Institute. Before joining General Motors in 1970, he worked for the National Aeronautics and Space Administration for two years designing electrical power systems for spacecraft. At General Motors, he has done research on a wide range of advanced electronic equipment for automobiles, including devices for avoiding collisions and providing drivers with information about traffic.

Further Reading

REVIEW OF THE STATE OF DEVELOPMENT OF ADVANCED VEHICLE CONTROL SYSTEMS. Steven E. Shladover in *Vehicle System Dynamics: International Journal of Vehicle Mechanics and Mobility*, Vol. 24, Nos. 6-7, pages 551-595; July 1995.

LIFE IN THE FAST LANE: THE EVOLUTION OF AN ADAPTIVE VEHICLE CONTROL SYSTEM. Todd Jochem and Dean Pomerleau in *AI Magazine*, Vol. 17, No. 2, pages 11-50; Summer 1996.

THE AUTOMATED HIGHWAY. Terry Quinlan in *ITS Quarterly* (Intelligent Transportation Society of America), Vol. 5, No. 2, pages 7-16; Summer 1997.

ROBOT ROADS ARE JUST AROUND THE CORNER. Norman Martin in *Automotive Industries*, Vol. 177, No. 6, pages 65-66; June 1997.

DEMO '97: PROVING AHS WORKS. Editorial in *Public Roads*, Vol. 61, No. 1, pages 30-34; July-August 1997.



Unjamming Traffic with Computers

Insights gleaned from realistic simulations are already moving from computer screens to asphalt

by Kenneth R. Howard

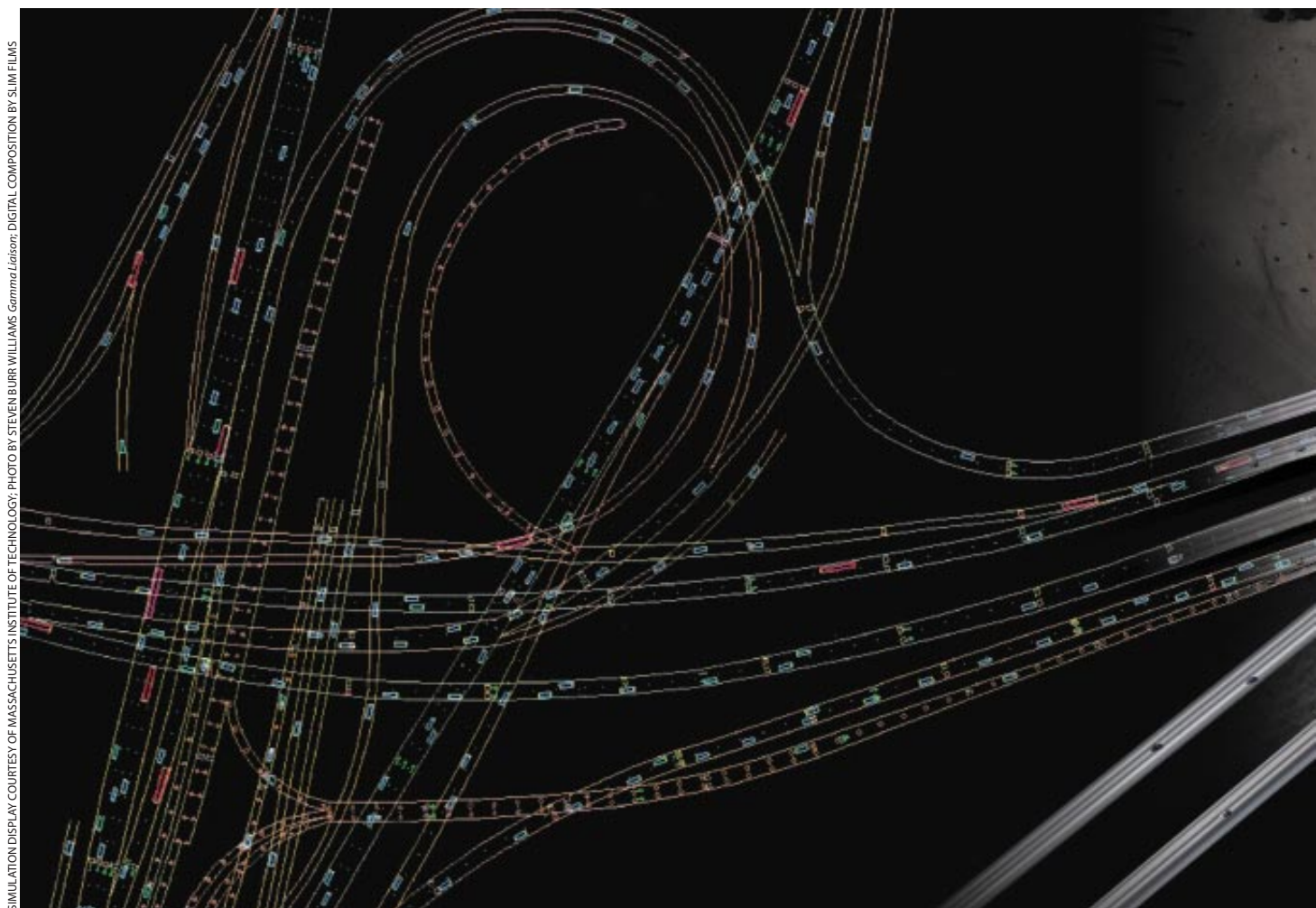
As any driver who has crawled in the “fast lane” of a crowded highway can attest, traffic has become a plague. Infrastructures strain at ever more cars arriving on the road. And although colossal sums of money are being spent on solutions—\$140 billion over five years by the federal gov-

ernment’s Highway Trust Fund alone—tangible results have been elusive, and road planning remains somewhat akin to gambling without knowing the odds.

Traffic problems can stem from cars slowing or stopping, sometimes in response to accidents. Yet the more usual cause is simply too many people want-

ing to be in the same place at the same time. Limitations in the road system or minor inconsistencies in the behavior of drivers can then compound to cause torturous slowdowns. Planners and scientists often state the problem as “too many people, not enough roadway.”

Past methods of predicting traffic pat-



SIMULATION DISPLAY COURTESY OF MASSACHUSETTS INSTITUTE OF TECHNOLOGY; PHOTO BY STEVEN BURR WILLIAMS Gamma Liaison; DIGITAL COMPOSITION BY SLIM FILMS

terns relied on statistical models that treated traffic as a homogeneous fluid, ignoring differences between individual drivers. Sections of transportation networks were often analyzed in a vacuum, without regard for the interactions between the components. Refinements in the mathematical treatment of complexity, however, coupled with hugely greater computing power, have revolutionized traffic analysis.

"Transportation is in a very cool spot between a social system and a physical system," explains Christopher L. Barrett, who studies traffic at Los Alamos National Laboratory. "Traffic lies in the middle. It goes beyond physics to a human scale." On the physical side are the stunning variety and number of vehicles and road systems, each contributing its own peculiarities, as well as the weather and other environmental factors. The social, behavioral side encompasses not only the individual preferences and second-by-second reactions of drivers but also actions taken by the rest of society—from corporate choices about

where to locate headquarters to sports teams' play strategies that influence attendance at games. To understand the forces acting on traffic flow, transportation planners must analyze the many possible outcomes from this snarled network of decisions.

As a further complication, seemingly logical solutions to traffic problems can have counterintuitive consequences. One instance is known as Braess's paradox, named for the German operations researcher Dietrich Braess, who in 1968 first noticed the phenomenon. He discovered that raising a network's traffic capacity can sometimes slow the average travel speed. "You have to be on the lookout for this problem when designing traffic networks," says Joel E. Cohen, mathematician and professor of populations at the Rockefeller University. Cohen explains that in road networks, adding a lane or route can increase driving time as unpredicted bottlenecks arise from what was thought to be a fix—for example, when too many drivers pile onto a new shortcut, causing gridlock.

As Steen Rasmussen of Los Alamos National Laboratory points out, "When you design roads you want to maximize throughput [traffic flow], but it turns out at the point of most throughput, predictability drops. This means as you go toward capacity, reliability of the traffic system breaks down." As variability within the transportation system explodes, more and more parts of it are pushed into a "critical regime," he says, where "small perturbations can cause the system to break down." The goal, then, is to design systems that function just under capacity.

Silicon Traffic

To that end, recent simulations have made important strides in mimicking the broad dynamics of transportation systems. Many traffic researchers now view transportation systems as what complexity theorists call "self-organizing systems"—entities that manifest a cohesive behavior even though they lack a central controller. According to Rasmussen, traffic can be looked at as a system of diverse elements widely distributed over space. "The elements that interact with one another are like biological systems. They are dynamical hierarchies with controls at many different levels, like organelles, cells, tissues, humans," he says. The challenge for designers of transportation simulations is isolating and modeling the different elements, then bringing them together to operate as a whole.

The tools that allow complexity researchers to unite the myriad components and organizational layers into a real-time or faster computer model of a traffic system are often cellular automata. Cellular automata are a type of computer simulation best known through the "Game of Life," invented by John Conway in 1970. Various agents, or elements with defined properties, are placed on a grid and assigned an initial state. (The states for car agents might be "moving" or "not moving," for ex-

SIMULATIONS OF TRAFFIC patterns as they arise on actual highways are becoming increasingly realistic, thanks to new computational approaches. With mathematical models and cellular automata, researchers can mimic the behavior of individual drivers under particular road conditions, then study the aggregate flow of cars that results.

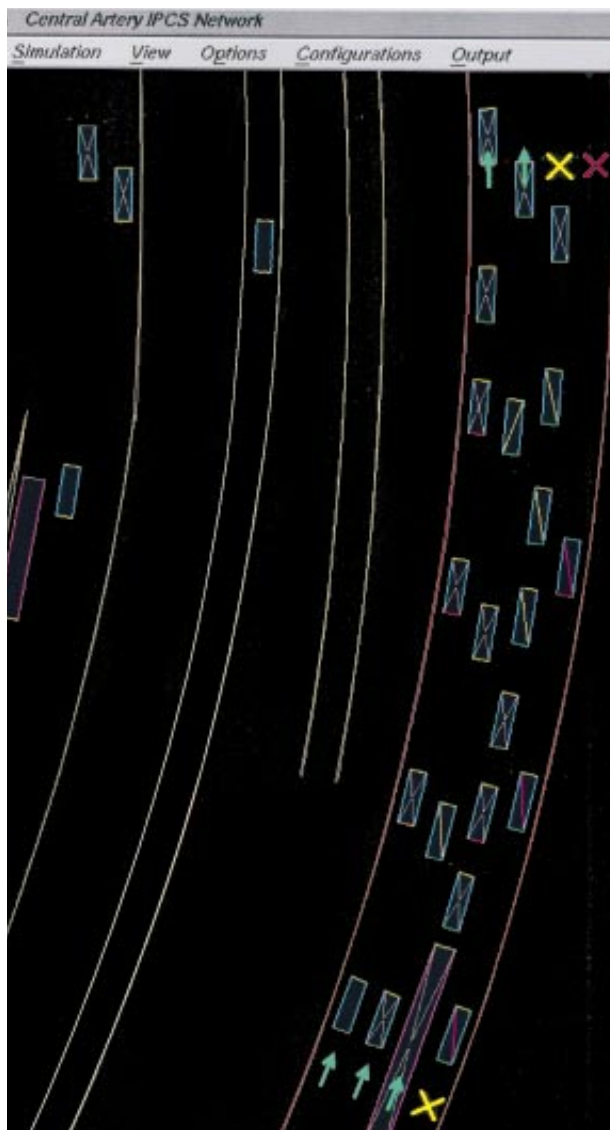


ample.) As time advances, each agent changes state in keeping with the rules of its own behavior and the current state of its neighbors. A typical rule could be that an agent is in motion if adjacent to two or fewer other agents and at rest if surrounded by more than two. According to Barrett, with each tick of the clock, the computer tries to update the state of every agent by looking at all its neighbors; from these interactions a global system emerges.

Barrett applied the supercomputers of Los Alamos to create a model of traffic scenarios called the Transportation Analysis Simulation System (TRANSIMS). John L. Casti, a complexity theorist at the Santa Fe Institute, describes TRANSIMS as "copying in silicon the traffic of a metropolitan area and putting it in real time as if you were in a helicopter watching the second-by-second movements."

The model creates a lab for testing traffic scenarios. Transportation planners can use it to predict, with reasonable accuracy, what the effects of building a bridge or adding a highway lane might be—options too costly to test in the real world. In this way, a Braess's paradox might be caught before a planning error was set in concrete. The simulation also brings the scientific method to traffic planning; an experimental situation can be precisely re-created.

TRANSIMS, which is sponsored by the U.S. Department of Transportation and the Environmental Protection Agency, was first used in 1993 to model the traffic system of Albuquerque, N.M. The simulations successfully mimicked the actual observed traffic patterns. In its second incarnation, which began in late 1995, it simulated traffic for Dallas/Fort Worth, Tex., an area of 3,600 square miles with 2.3 million residents. Beginning in 1998, work will commence on a more advanced simulation of Portland, Ore., that will attempt to incorporate realistic patterns of people changing traffic modes, such as driving to a train station, taking the train and then rid-



VIRTUAL VEHICLES with individual characteristics travel on Interstate 93 in this display from M.I.T.'s Intelligent Transportation Systems simulation. The green arrows and the red and yellow crosses represent electronic lane-use signs, one of many ideas being tested for real-world application. These signs could be used to guide drivers away from blocked lanes, smoothing the flow of traffic.

ing a bus to finish a commute to work.

TRANSIMS is still only in a research and development phase, but experts with knowledge of local highways can study the results of the simulations and offer insights into why certain traffic patterns emerge. Says Casti: "There is a parallel in evolutionary biology. It is hard to predict change, but hindsight can give a good explanation of why things turned out as they did."

Meanwhile transportation planners are learning from simulations by the Massachusetts Institute of Technology's Intelligent Transportation Systems Program. This system (which is not based on cellular automata) mathematically

models behavior down to individual driver habits—creating digital cars with a penchant for cutting off other cars, speeding down lanes and generally exhibiting traits seen every day on the highways in and around Boston. The program is being used by the city's \$8-billion Central Artery/Third Harbor Tunnel project. In addition to testing road-plan scenarios before construction, the simulations are helping with designs for traffic management systems (such as traffic signal algorithms and driver information systems) that will smooth the flow of traffic. The strategy, according to Moshe Ben-Akiva, professor of civil and environmental engineering at M.I.T., is to alleviate traffic congestion by designing a physical system that guides drivers toward better choices.

Using computer simulations to find the best solution to traffic problems ultimately calls for the inclusion of many factors beyond driver behavior, traffic density and the like. Possible fixes such as congestion pricing for tolls—with higher prices for peak-hour usage—and mass transit could be taken into account. Pollution analysis is another important consideration, one mandated by law and subject to paradoxical effects. (Shortening the distances that cars travel, for example, might seem like a good way to reduce emissions. But during a short trip, a car's engine and its catalytic converter

stay too cool to run efficiently and so proportionally emit more pollutants.) The ideal traffic simulators would consider aspects of air chemistry and construction patterns, because the geometry of buildings affects air movement.

The accumulating complexity of all these variables cannot yet be modeled or predicted easily. Planners will therefore have to wait for more complete computational tests. In the meantime, however, simulation is still likely to provide the insights necessary to keep traffic moving in the right direction. SA

KENNETH R. HOWARD is a writer based in New York City.



Now That Travel Can Be Virtual, Will Congestion Virtually Disappear?

by Patricia L. Mokhtarian

MARK LEWIS/Tony Stone Worldwide

The idea that telecommunications technology could substitute for travel dawned on people soon after the invention of the telephone. In the late 1870s letters and articles speculating on the potential of the telephone to replace face-to-face meetings appeared in various London newspapers. The science fiction of H. G. Wells ("When the Sleeper Wakes," 1899) and E. M. Forster ("The Machine Stops," 1909) described videoconferencing machines (or "kineto-tele-photographs," as Wells put it) that could accomplish the same goal. And an article in a *Scientific American* supplement from 1914 predicted that telecommunications would reduce transit congestion.

These ideas resurfaced in the 1960s and 1970s, as computing technology began to permeate society and the energy crises of the period prompted efforts to limit the burning of fossil fuels. But today, with fax machines and personal computers ubiquitous and videoconferencing almost mundane, congestion on the roads is worse than ever.

What's going on? Is the tidal wave of telecommuting still imminent, or are we anticipating something that is not likely to happen? When I started researching this question 15 years ago, I was optimistic about the power of telecommuting to reduce congestion, but now I'm more skeptical.

Commuting to work is the single most common trip people make and is a major contributor to overcrowding on the roads. Unlike trips to the grocery store or the doctor, the commute to the office can be more easily eliminated or reduced with telephones, faxes and e-mail. So if we hope to use communications technology to mitigate congestion, telecommuting is perhaps our best option. Conversely, if telecommuting does not do much good, then it is unlikely that teleshopping, teleconferencing, telemedicine, telebanking and other "telestuff" will have much impact.

Many people say they telecommute, but what really counts is how many are actually doing so on any given day. That number is the product of several factors, such as how many people can telecommute. Of those, how many want to? Of that group, how many actually do, and how often and for how long?

For many workers, telecommuting is simply not feasible. The job may be unsuited to it, or they lack the proper equipment. Some people are not aware they could telecommute; others face unwilling managers. I estimate that at present only about 16 percent of the entire workforce can actually consider telecommuting.

Not everyone who can telecommute even wants to, and not everyone who claims to want to actually does. Many people desire the professional and social interactions of the office. Others may be concerned about lack of self-discipline and domestic distractions. Many workers also consider the commute to and from their job a desirable buffer between home and office.

Although full-time telecommuting suits some, most people prefer doing it part-time—one or two days a week on average. In addition, several studies have noted that most people who try telecommuting do not do so forever: one half are likely to drop out within a year or so of starting. So what do all these statistics mean? Probably no more than 2 percent of the total workforce is telecommuting on any given day.

Now the question becomes whether these telecommuters actually have an effect on congestion. We do know that telecommuters drive less. This may seem obvious until you realize that there are several ways in which telecommuting could theoretically increase travel. People might decide to take more excursions to avoid cabin fever at home. Or telecommuters who once carpooled to work might decide to drive alone on the

days when they do go into the office.

The challenge arises in trying to extrapolate such findings to determine the overall effect of telecommuting on traffic on the roads. I estimate that telecommuting by 2 percent of the workforce translates to a reduction in the total number of miles driven in personal vehicles (cars and light-duty trucks) of 1 to 2 percent—an amount swamped by the increasing number of miles traveled by Americans in general.

Even this modest reduction will most likely decline over time. Today's telecommuters tend to live twice as far from work as the average employee. But when (or if) telecommuting moves into the mainstream, the distances saved on each occasion will fall closer to the average. Extra trips will probably increase as well. Early telecommuters may have been reluctant to take excursions from home because they were already traveling so far on the days they still had to commute. If telecommuters have a shorter drive to work, new trips for shopping or socializing on their telecommuting days may become more appealing.

The long-term effects of telecommuting, especially on where workers will choose to live, are not well understood at all. Telecommuting could potentially motivate some people to move even farther from work than they live now.

If telecommuting ever did reduce congestion noticeably, the excess capacity on the highways would almost certainly be quickly filled by changes in current travel patterns. For example, more people might decide to drive alone instead of using public transportation.

Historically, transportation and communications have been complements to each other, both increasing concurrently, rather than substitutes for each other. And we have no reason to expect that relationship to change. Yet under the right circumstances, telecommuting offers substantial benefits to employers, employees and society at large. Communications technology can increase productivity, reduce costs and give workers more personal flexibility. Sizeable reductions in travel will not be among these benefits, but telecommuting is still an idea worth promoting. SA

PATRICIA L. MOKHTARIAN, professor of civil and environmental engineering at the University of California, Davis, has studied the effects of telecommunications on travel behavior, land use and the environment since 1982.

750
mph

Driving to Mach 1



“Jetmobiles” try to go supersonic

by Gary Stix, staff writer

Before Chuck Yeager broke the sound barrier in 1947 in the *X-1* experimental plane, engineers had predicted that the buffeting produced by supersonic shock waves might tear apart his sleek craft. As drivers—one might call them pilots—of two custom-made supersonic cars recently prepared to punch through Mach 1, the engineering community voiced similar concerns, perhaps this time with more reason. “Anything that upsets a vehicle at 600 miles [around 965 kilometers] per hour or more puts it in a regime you don’t want to be in,” comments Make McDermott, a professor of mechanical engineering at Texas A&M University. “Aerodynamic forces make a ground vehicle not a ground vehicle. They make it want to fly.”

As this issue goes to press, the most serious attempt ever to break the speed of sound in a ground-based vehicle is scheduled to take place this year in September and early October in the Black Rock Desert northeast of Reno, Nev., the largest dry lake bed in America. The push toward Mach 1—the speed of sound, which is around 750 mph at the temperatures encountered at Black Rock—promises to be a dramatic face-off between two teams that have at different times claimed ownership of the title of fastest car on earth.

One contender is Craig Breedlove, the 60-year-old driver of a “jetmobile” called the *Spirit of America*. Breedlove captured the record five times between 1963 and 1970. The other team is headed by Richard Noble, the now 51-year-old driver of the British vehicle that achieved the current record of 633 mph in 1983. Although Noble is overseeing the effort, his car, *Thrust SSC*, will be driven by Royal Air Force fighter pilot Andy Green.

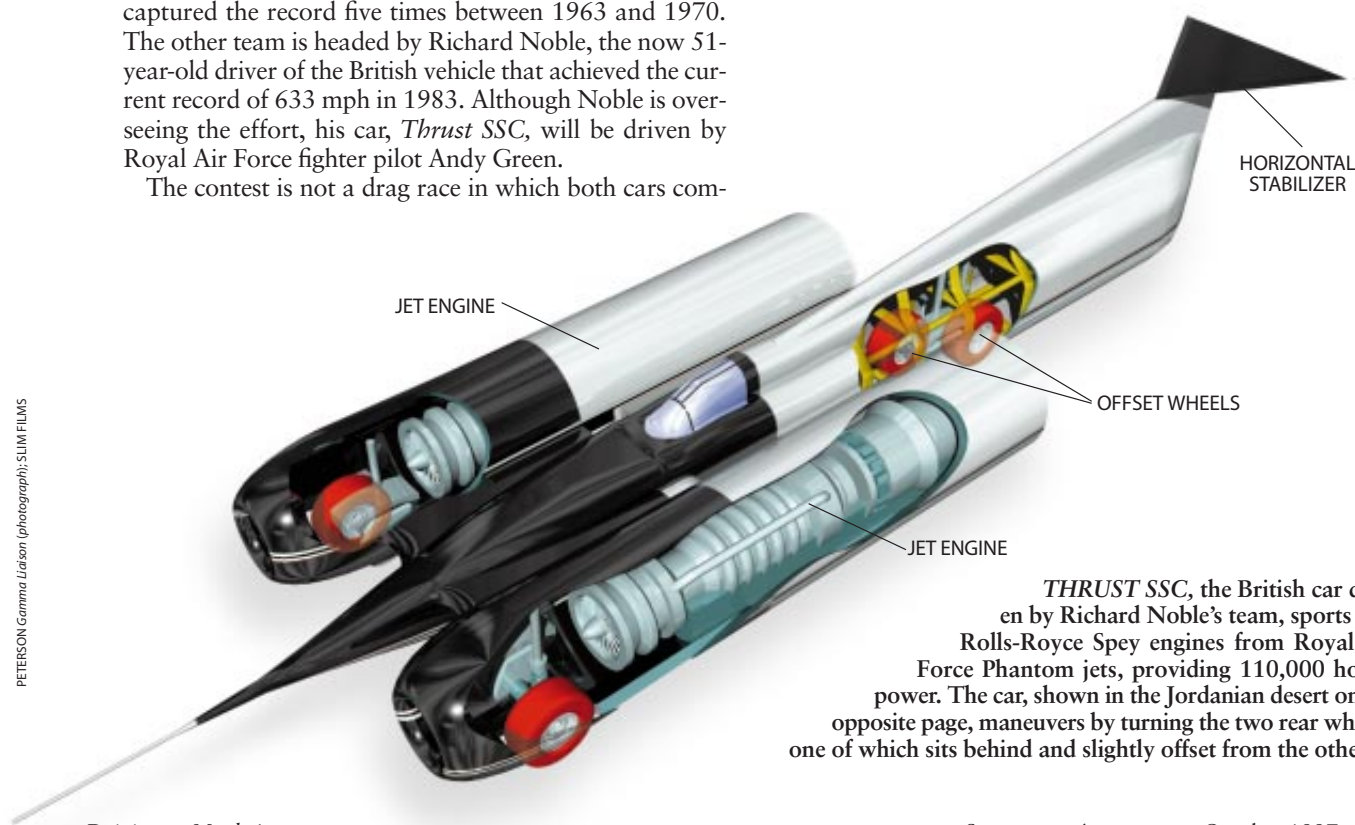
The contest is not a drag race in which both cars com-

pete simultaneously. The two teams will share the desert, making separate runs at gradually increasing speeds. Even if they do not break the sound barrier, they could still best Noble’s 1983 record or the 700-mph mark.

These teams are not the only ones in the world trying to break the 1983 record. But the intensive engineering and expense that have gone into both their vehicles make them the only candidates that can expect to approach anywhere near the speed of sound.

Unofficially, the sound barrier may already have been broken. In 1979 stuntman Stan Barrett claims to have piloted a rocket-powered car, the *Budweiser Rocket*, to a speed of nearly 740 mph.

But if the car did go that fast—which Breedlove and others hotly debate—it achieved that speed going only in one direction. The International Automobile Federation, the Paris-based organization that certifies these records, requires that a vehicle must average a record-breaking speed over a measured mile during two runs going in opposite directions, each drive within an hour of the other. In the Black Rock Desert, the cars will move along a 15-mile flat. They will accelerate for nearly five miles, move through a measured mile in about five seconds in the middle of the course, then slow down for five miles by cutting power and releasing parachutes before applying brakes at



THRUST SSC, the British car driven by Richard Noble’s team, sports two Rolls-Royce Spey engines from Royal Air Force Phantom jets, providing 110,000 horsepower. The car, shown in the Jordanian desert on the opposite page, maneuvers by turning the two rear wheels, one of which sits behind and slightly offset from the other.

speeds below 300 mph. Then they will turn around and go back the same way.

How does one build a car to drive to Mach 1? Both Breedlove's and Noble's teams have chosen jet engines originally used on fighter aircraft. But choosing to strap a driver to a jet engine turns out to be one of the simplest design decisions. Keeping the driver alive is another matter. Little is known about what happens to a car when it reaches the speed of sound. In an airplane the shock wave that occurs when the vehicle nears Mach 1 attenuates in the surrounding air, as Yeager learned. When a car approaches the speed of sound, the boundary between air flowing at supersonic and subsonic speeds creates shock waves between the vehicle and the ground that could initiate a potentially fatal back flip or side roll.

Jet cars have already demonstrated their perils. In an attempt to achieve a record last year in the *Spirit of America*, Breedlove veered out of control at an unofficial 677 mph and damaged the rear wheels. The British team has also experienced travails because of stresses on the car's frame. *Thrust SSC*

sustained damage this past July when a rear suspension bracket failed during a 540-mph-plus test run in Jordan's Al Jafr Desert.

Airplanes without Wings

The competitors have adopted divergent design approaches. Breedlove has tried to reduce the car's frontal cross section to lessen drag from an oncoming wind of 750 mph or more. The *Spirit of America* weighs 4.5 tons and is 44 feet long and 8.5 feet across at its widest point, the span between the back wheels. That width is almost four feet less than its British competitor. The elliptical shape of the front of the body is intended to let destabilizing shock waves underneath the car escape to the sides. In addition, the front part of the body sits only one inch off the ground to reduce the area in which pressure can build up. A larger clearance of 18 inches in the rear allows pressure waves to escape from the back.

The front wheels are three aluminum disks that spin on a single axle, each separated by a tenth of an inch, a configuration designed to increase the iner-

tial forces that prevent yaw—side-to-side movement. The wheels themselves are wound around their circumference with graphite fibers capped with fiberglass. These high-performance tires can protect the outer wheel rims, which may be subject to 35,000 times the force of gravity. "Should a wheel hit a rock when highly stressed, you don't have to be a rocket scientist to figure out that it could fracture from the outer periphery to the hub, and you'd have catastrophic failure," Breedlove says.

Each rear wheel extends out a few feet from the body of the car. This choice of design tends to move forward the center of mass (center of gravity), thereby enhancing stability. "It's like handles on a wheelbarrow," Breedlove says. "The longer you make the handles, the easier it is to pick up the load and the more weight goes on the front wheel of the wheelbarrow." The rear axles are encased in a flattened, horizontal winglike structure called a fairing. Fins are attached to the far edges of each fairing to guard against yaw forces.

Since his accident last year, Breedlove has refined the aerodynamic shape of



the fairing so that airflow speeds up on top while slowing down on the bottom. This design change is intended to prevent the air underneath from becoming supersonic and generating shock waves. Taking this step also required adding a set of flaps—wing surfaces that can be set to prevent the wing from lifting the car from the ground.

Stability First

Taking a different tack, the *Thrust* SSC team first determined the most stable design needed to reach Mach 1 safely. Only then did it decide on power and drag reduction measures. In contrast to Breedlove's penchant for trial-and-error development, team aerodynamicist Ronald F. Ayers, who once designed guided missiles, relied heavily on supercomputer simulation and supersonic tests with a two-foot-long model on a sled that was propelled by rocket fuel along a track.

From Ayers's technical analyses, the *Thrust* SSC emerged as a larger, bulkier vehicle than *Spirit of America*, weighing seven tons and measuring 54 feet in length. The team went to great pains to keep stable the vehicle's pitch—the slant of the nose above or below the vehicle's horizontal axis. "Too much nose up, and you take off like an airplane," Ayers says. "Too much nose down, and you bury yourself in the desert."

The difference between becoming a

Patriot missile or a miner's drill is an angle of only a degree or so. That angle also changes at supersonic speeds. The car incorporates an active suspension that can make the necessary adjustments in pitch as the vehicle nears the sound barrier. Strain gauges measure the load on the wheels and relay this information to an onboard computer. Hydraulic jacks at the rear of the car can then adjust attitude automatically. Between runs, the angle of a horizontal stabilizer at the back of the car can also be adjusted to ensure that the rear wheels remain firmly on the ground.

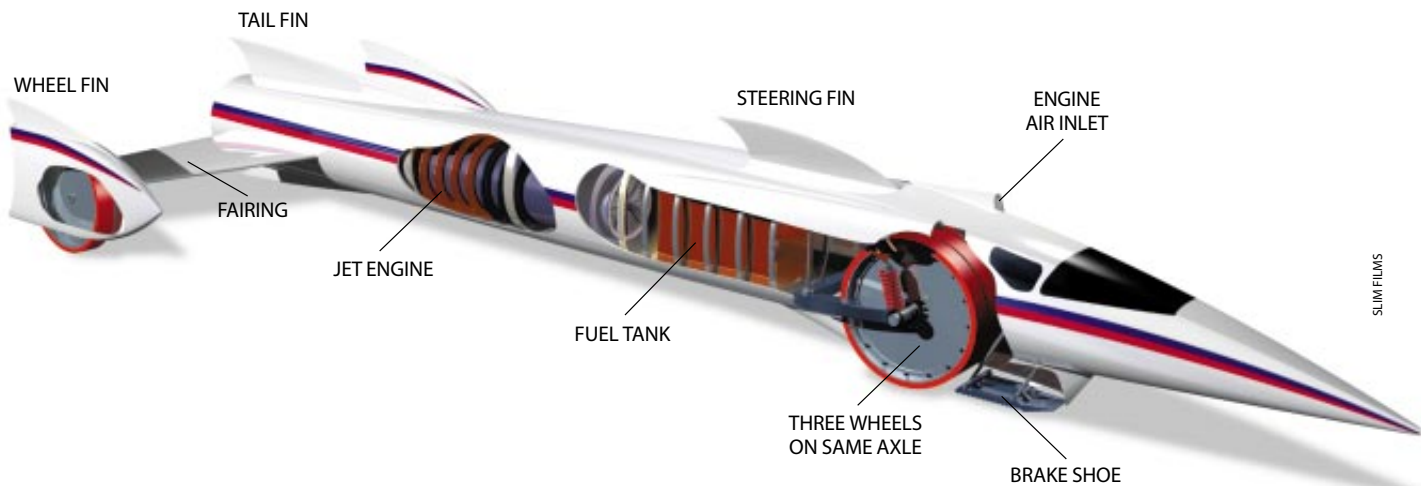
The *Thrust* SSC's front wheels hide inside the engine cowl, the covering that houses the engine, which reduces the cross-sectional area that faces into the wind. Placing the engines on the side and maintaining a wide front wheelbase moves the center of gravity farther forward than that of *Spirit of America*, a measure intended to keep the nose in position. Setting forward the center of gravity counteracts the tendency of the car to lurch into a spin. Increased thrust from two engines, Ayers says, compensates for added weight and the additional drag produced by the wider front profile. "There's no weight limit for this class of car," he quips. Unlike Breedlove, the *Thrust* SSC team uses forged aluminum wheels that turn without tires. The team hopes the desert surface will be soft enough to make up for the absence of tires on the wheels.

Driver Green sits in a cockpit placed at midsection between the two engines, allowing him to gain a better feel for side-to-side movement of the car. Green steers the two rear wheels, which are tucked toward the back of the body's underside to avoid interference with the jet exhausts. One rear wheel is placed slightly behind and to one side of the other, avoiding a drag-producing bulge in the rear section that would have occurred if the wheels had been placed parallel to each other. The front wheels are fixed in place: avoiding a front steering mechanism reduces drag.

Lessons learned in building a supersonic car may have scant value beyond an entry in the *Guinness Book of World Records* and thrills for those who make and drive the cars. "I am convinced that taking part in the world land-speed record is the most exciting thing you can do on God's earth," Noble says, expressing the missionary zeal that his team brings to the task. Practical spin-offs of running a car at these speeds are at best conjectural. Breedlove notes the possibilities for the tire technology. But when asked about where this type of graphite tire might be useful, Breedlove ponders for a moment and then replies, "I have no idea. My mission is to get the land-speed record. I'm not moved by much else. I think Kennedy wanted to go to the moon. He didn't care about spin-offs from it. He just wanted to beat the Russians."

5A

SPIRIT OF AMERICA, which rides with a General Electric J-79 jet engine that once powered a U.S. Navy F-4 Phantom fighter, supplies 48,000 horsepower. The engine is enclosed in the fuselage toward the rear of the vehicle. Driver Craig Breedlove steers by turning the three front wheels and by manipulating a steering fin atop the vehicle. The car, shown on the opposite page before an accident last year, has now been rebuilt.



SLIM FILMS



Speed versus Need

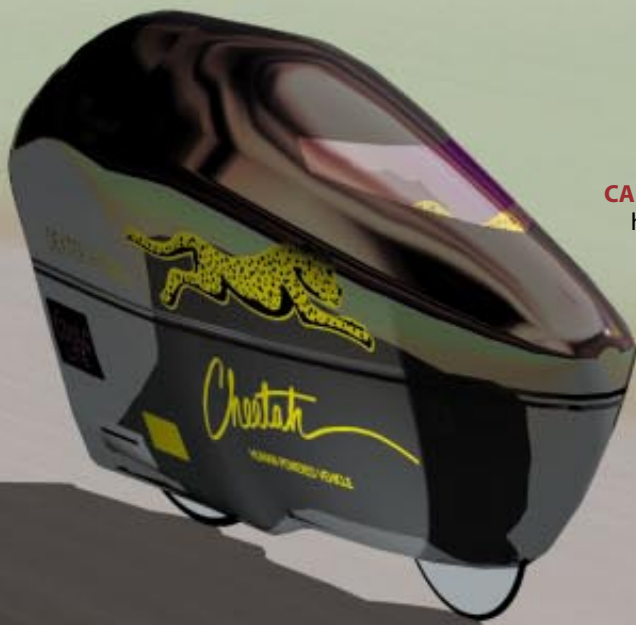
by Kristin Leutwyler, *staff writer*

Rugged mountain climbers, bamboo rigs built for two, three-speeds with banana seats—bicycles, in their many forms, exist the world over. We use them to make deliveries, to exercise, to explore and to get to work. And in recent years, creative engineers relying on sundry

lighter, stronger materials have come up with a variety of new models. Superfast cycles, such as the Cheetah (*below*), lie at one extreme. At the other are sturdier, more affordable two-wheelers, like the Kangaroo (*opposite page*). Here we highlight some of the main innovations in both designs.

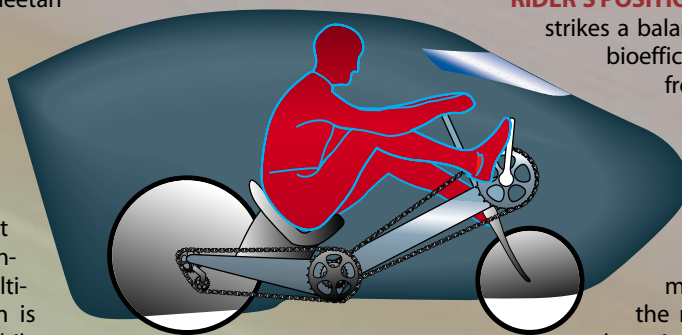
CHEETAH

This Human Powered Vehicle, as it is known, takes its name from the famed cat because, similarly, it is the fastest of its kind. Built by graduate students at the University of California at Berkeley, the Cheetah, which weighs in at a slight 29.5 pounds (a bit over 13 kilograms), set a world speed record on September 22, 1992. On that day, it reached an average speed of 68.73 miles (nearly 111 kilometers) per hour along a 656.2-foot (200-meter) roadway in Colorado's San Luis Valley.



CARBON-FIBER SHELL, called a fairing, makes the Cheetah highly aerodynamic. Because there are no openings through which the rider might extend his feet, his team members must help him balance until he gets under way, and they must catch him when he slows to a stop. Computer analysis helped designers optimize its shape to slip through the wind: the fairing is a mere 18 inches wide but more than nine feet long.

CUSTOM GEARING lets the Cheetah cruise near 70 mph, whereas conventional bikes reach top speeds of only 25 to 30 mph. A bike's speed depends on its gear ratio, defined by how many times the rear wheel turns for each turn of the pedals. The fastest traditional bikes use front chain rings of 53 teeth. Using an intermediate gear assembly to multiply the gear ratio, the Cheetah is equivalent to a conventional bike with a ground-scraping front ring of 117 teeth.



RIDER'S POSITION, described as semirecumbent, strikes a balance between aerodynamics and bioefficiency. Because his legs are out in front of him, the rider cuts a smaller cross section with the oncoming wind than if he were sitting upright, which minimizes resistance. Of course, this cross section would be even less were he lying down. To maximize pedaling power, though, the rider's heart must be above his legs. A semirecumbent position also allows for greater visibility and control.

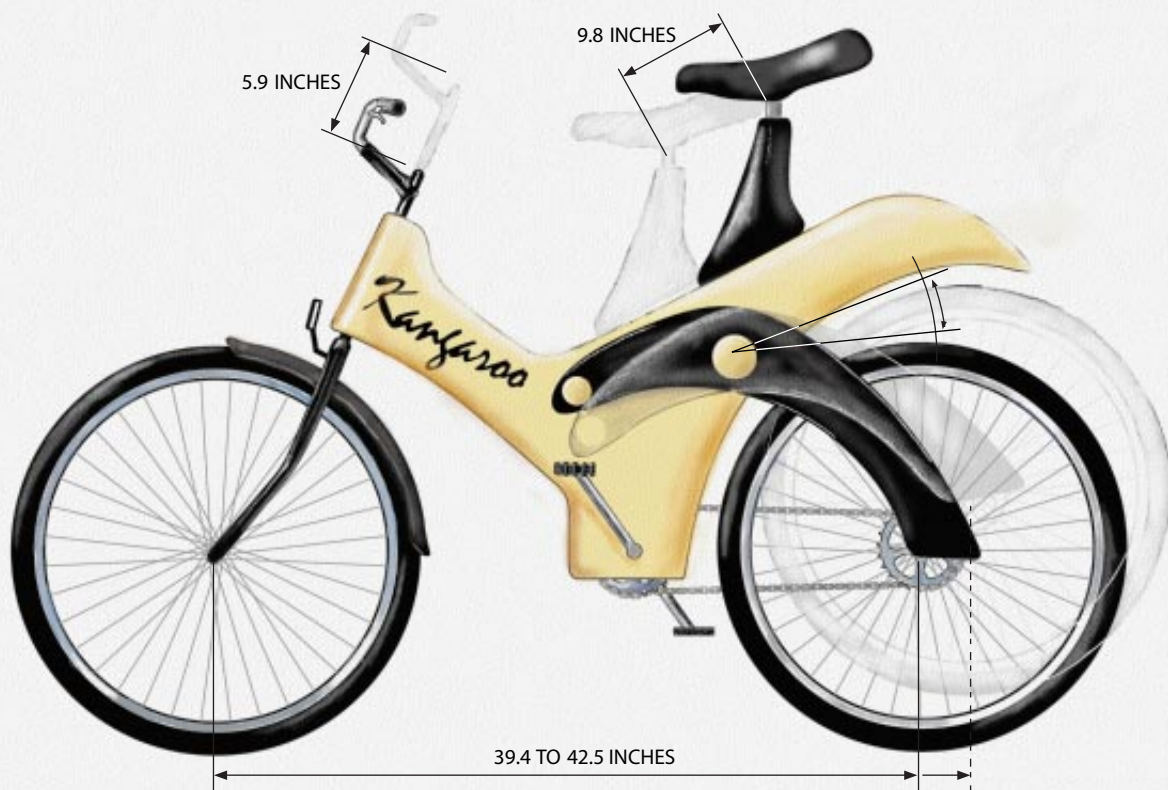


KANGAROO

This sturdy design, created by students from the University of São Paulo in Brazil, won a competition last year for the ideal “world bike”—defined as one that would be practical and affordable for some 80 percent of all people planetwide. Owens Corning sponsored the challenge in recognition of the fact that more than half the world’s population uses a bicycle as its primary means of transportation.

THIRD-PARTY PRODUCTION SYSTEM, in which bicycle companies need only assemble the Kangaroo’s premade parts, would demand little investment and should benefit from low labor costs. Using this model, the Kangaroo’s inventors figured that five million pieces could be made every year and that the finished product would cost \$82.

SHEET-MOLDING COMPOSITE of polyester resin and chopped-glass fibers—at roughly \$4 per pound—makes for an extremely cost-effective frame. But the parts must be carefully constructed to ensure their strength. All the Kangaroo’s main pieces have a low aspect ratio—that is, they are not much longer than they are wide or high.



ILLUSTRATIONS BY BRYAN CHRISTIE

VARIABLE GEOMETRY ensures that the Kangaroo is well suited to many users and uses. Fully 95 percent of adults can ride the Kangaroo comfortably because the position of the seat can be adjusted over a 9.8-inch (25-centimeter) range, and the height of the handlebars can be raised or lowered over a 5.9-inch range. Also, by shortening the wheelbase from 42.5 to 39.4 inches, users can adjust the Kangaroo’s handling characteristics from those of a cruising bike to those of a more agile city bike.



How High-Speed Trains Make Tracks

In Europe and Japan, train manufacturers are gearing up to achieve ultrafast speeds routinely, without relying on levitation

by Jean-Claude Raoul

Over the past 30 years or so, Japan and Europe have invested heavily in networks of high-speed trains to link major cities. They have turned to fast trains, exceeding 200 kilometers (roughly 125 miles) per hour, in part to relieve congestion on roads and at airports while minimizing operating costs and pollution.

Of course, for trains to live up to their financial and environmental promise, they must draw high numbers of paying passengers. The Japanese and European experience has shown that railways can often meet that demand if the rides

are comfortable, competitively priced and able to deposit travelers at their destinations about as quickly as an airplane would. Aircraft still go much faster than trains, often exceeding 600 kph, but long travel times to and from airports often cut significantly into time savings.

Engineers knew as early as the 1950s that simply by using more power they could force some conventional trains to reach 331 kph, much faster than the 130-kph top speed of many American long-distance trains today. But the higher speeds were deemed infeasible for commercial application because the fast-

moving vehicles damaged the tracks severely. High-speed trains, it seemed, would have demanded extensive, and thus prohibitively expensive, track maintenance efforts.

Nevertheless, Japanese and European innovators soon found ways to exploit existing technology to improve speeds to about 200 kph between some cities. For instance, without altering the trains themselves greatly, the Japanese designers achieved gains by such maneuvers as building tracks that avoided tight curves and steep grades. The huge popularity of their original Shinkansen, or

TRAIN À GRANDE VITESSE (TGV), shown in its double-decker (duplex) version, runs at up to 320 kilometers (almost

200 miles) per hour in France. The map displays the European Community's master plan for a high-speed train network.



bullet train, which began operation in 1964 between Tokyo and Osaka, sparked new interest in overcoming the technological obstacles to operating routinely at still higher speeds.

Those efforts have since resulted in a number of trains that go significantly faster than 200 kph. Among the best-known examples are the Train à Grande Vitesse (TGV) series in France, the InterCity Express (ICE) lines in Germany and the Eurostar trains (linking Paris and Brussels with London by way of the English Channel Tunnel ("Chunnel")). These trains and newer generations of the Shinkansen can all zoom at or near 300 kph on dedicated high-speed tracks (although they go more slowly on older tracks). And plans are under way at the French National Railway and at GEC Alsthom, respectively the owner and builder of the TGVs, to produce another series of trains—dubbed the "new generation"—able to cruise regularly at

360 kph. These vehicles are the product of an intensive research effort involving about 50 university laboratories, mostly in France but also in the U.S., Belgium and Sweden.

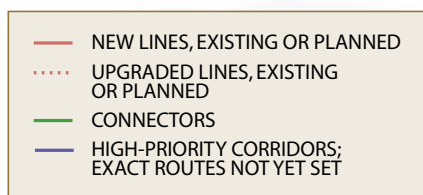
Stability Is Critical

Reaching these milestones has required innovation in all aspects of railroad engineering, including the design of tracks and signaling systems. For instance, as speeds rose, roadside signals became useless for the drivers; the cabs went by the signals too fast. The trains are now run with guidance from on-board computers that collate information emitted from monitoring and control equipment in the tracks and in the individual cars and from dispatching stations; the computers can also force the train to stop if critical safety commands go unheeded. But some of the most interesting inventions have altered

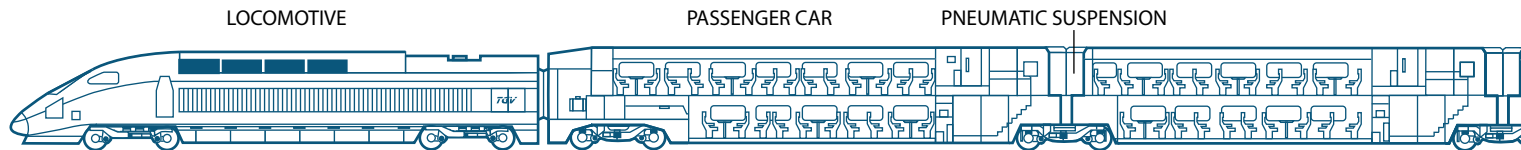
components of the trains themselves.

The design elements that the French have introduced for the TGVs offer an example of the kinds of technology that make wheel-on-rail travel at high speeds possible. Those solutions differ in some respects from those chosen in other countries, but they provide a sense of the work that has allowed speeds to increase steadily since the 1960s.

High-speed or not, most long-distance trains have certain features in common. They are moved by one or two locomotives, cars containing the power-generating equipment. This equipment converts energy from onboard fossil fuel or from an electrical feed into the specific form needed to move the train. In the U.S. many trains still run on diesel fuel. But in Europe most trains, and all high-speed trains, run on cleaner electrical power; this power is usually drawn from overhead lines, or catenaries, through a pantograph—a conducting rod—pro-



GEC ALSTHOM/JDD (photograph); LAURIE GRACE (map)



LAUREL ROGERS

truding from the top of the train. The energy that is produced allows traction motors under the locomotives to rotate the axles that join pairs of driving wheels—those that grip the track (use traction) to propel the train forward.

The driving wheels, as well as other wheels that simply carry the weight of the train and allow it to move smoothly along the tracks, reside in support structures known as bogies. Bogies, also called trucks, consist of two or more pairs of wheels and their axles connected by a frame that supports the cars above. A suspension system linking the bogies and the cars holds the cars in place and cushions riders from vibrations.

When train speeds rise, the vibrations produced by contact between the wheels and the rails increase dramatically. These vibrations can cause the bogies to become extremely sensitive to imperfections in the track, to sway from side to side and, ultimately, to jump the track. Moreover, as the 1950s tests showed, such rocking can damage the rails and incur huge maintenance costs. Hence, increasing ride stability became an early priority.

In the early 1970s, when scientists at the French National Railway and GEC Alsthom first began aiming for speeds above 200 kph, powerful computer simulation tools were not available. But experimentation and calculation indicated that increasing the distance between axles in the bogies to three meters from the 2.5 meters of conventional trains would maintain stability even at speeds in excess of 300 kph. Moreover, lengthening the distance would obviate the need for adding a great deal of vibration-dampening equipment that would have to be monitored constantly and replaced periodically.

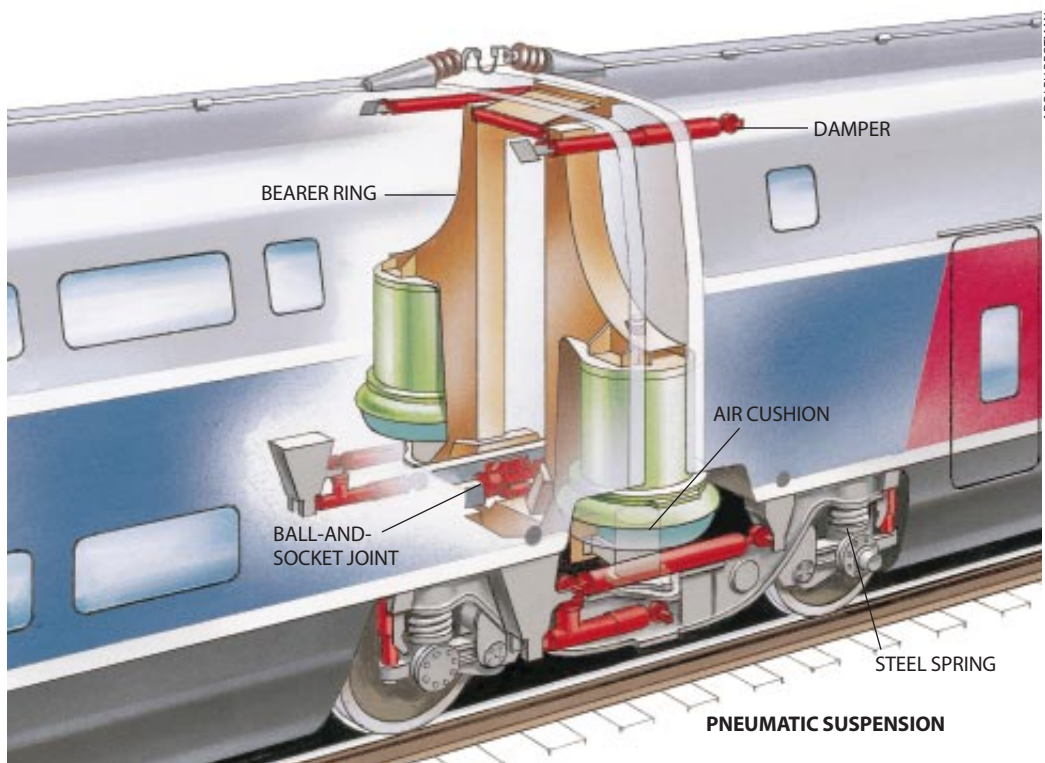
Suspending traction motors from the bottom of the locomotive or passenger cars, instead of mounting them as usual on the bogies, improved train stability as well. As bogies get heavier, the risk of

bogie instability and derailment increases. Moving the motors off the bogies lowered the weight of the trucks. Today researchers at GEC Alsthom, where I am technical director, continuously test new materials for the bogies—such as aluminum alloys or carbon fibers—looking for substances that will further re-

duce weight while retaining strength.

In a key departure from conventional construction, the designers of the TGVs also altered the placement of the bogies. Most trains allot two bogies to each car, setting them some distance in from the ends of the car. But with the exception of locomotives, TGV cars share bogies.

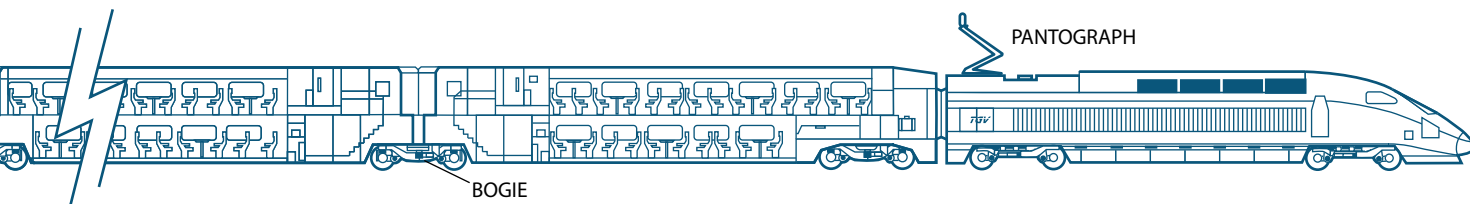
HIGH-SPEED TRAINS operating in Japan and Germany are known, respectively, as the Shinkansen, or bullet train (*left*), and the InterCity Express, or ICE (*right*). The commercial



ADOLPH BROTHMAN



KATO OSHIHARA Spira Press



DUPLEX TRAIN has several features that enable current TGVs to reach high speeds without destroying the tracks. They include aerodynamic styling; lightened materials throughout the train, including in the transformer (in the locomotive), the car frames and the bogies; shared bogies between passenger cars (instead of two bogies per car); use of a single pantograph (instead of the many used in other trains); and a pneumatic suspension (*detail at left*). The bearer ring in the suspension puts the weight of the cars on the air cushions, and the ball-and-socket joint links the cars. The dampers keep the cars aligned with one another and prevent them from rotating around their various axes.

We fit one bogie between each car, so that an individual car has a total of just one bogie (half a shared bogie at one end, plus half at the other end). These between-car trucks knit adjacent cars together semipermanently, preventing them from pivoting away from one another on curves.

This tight coupling of all cars limits managerial flexibility to an extent; cars cannot be added or removed readily to adapt to changing passenger loads throughout a day. But such changes are inadvisable in any case, because the computer systems that monitor and control every car on the train would have to be reprogrammed constantly to accommodate the rearrangements—a process that would require a great deal of labor and care.

The design of the suspension system also influences stability, and so investigators have tested several types. If stability were the sole concern, the ideal system would totally prevent cars from

swaying, but such a suspension would cause riders to feel every vibration underneath them. For the first generation of TGVs—running between Paris and Lyon—engineers settled on a steel-spring suspension, in which the vertical springs become stiffer as the frequency of the vibration increases. Those trains began operation in 1981 at 270 kph and later set a speed record when a test showed they could accelerate to 380 kph.

Later we switched to a pneumatic suspension: air cushions take the place of some of the steel springs and provide better insulation from vibrations. This new suspension, in addition to making for a more comfortable ride, helped the second generation of TGVs—the *Atlantique* trains, serving areas west of Paris—to set a world speed record of 515.3 kph in 1990 and to operate commercially at 300 kph.

In Germany, Sweden and other places, the problem of stability is being addressed somewhat differently than in

France. For instance, instead of altering the placement of the bogies, various manufacturers install tilt technology to cope with curves: the cars can pivot on the bogies and lean to balance the forces acting on the train and on passengers. Tilt technology has allowed trains to go as fast as 220 kph on upgraded older tracks, without forcing newer, straighter ones to be built.

Optimizing Shape and Weight

In addition to ensuring that high-speed trains will be stable, designers have to minimize the amount of fuel required to run the vehicles, both to limit pollution from the power plant that provides the electricity and to save on the costs of that electricity. To achieve the greatest speed for the lowest cost, the vehicles, above all, have to be aerodynamically designed to minimize the amount of drag that is produced when they race down the track. For that reason, high-speed trains as a group have smoother surfaces and fewer angles than standard trains do.

To reach speeds of 360 kph, more design changes will be needed. Diverse analytical tools—including sophisticated computer simulation programs, scale-model tests in wind and water tunnels, and analyses of wind flow around full-size trains on tracks—all show that most of the drag impeding the forward motion of current high-speed trains derives from the bogies and other equipment under the frame. Future generations of TGVs will therefore have smoother underframe contours.

Although some people might suspect that a train's weight would affect fuel consumption as much as its shape would, weight actually has little influence on that aspect of the operation of high-speed trains. But a heavy train stresses the tracks more than a lighter one does and consequently increases maintenance costs. Therefore, to protect the tracks, fast trains need to weigh as little as possible.

The novel arrangement of the bogies in TGVs helps to keep down the weight; by providing one bogie per passenger car instead of two, we almost halve the number of bogies in the train. We also

success of the first Shinkansen, which began running in 1964, stimulated the later development of still faster trains, including TGVs, ICEs and subsequent Shinkansen generations.



craft the cars from lighter material than has been incorporated into conventional trains. Use of such materials in the passenger cars has made it possible to produce double-decker (duplex) vehicles that weigh no more than the single-deck Atlantique trains, even though the duplexes boast seats for 45 percent more passengers. Thanks to aerodynamic styling, the duplexes also run as fast as their one-level counterparts but consume less energy.

The motors directly responsible for turning the driving wheels—the traction motors—have been lightened, too, without sacrificing power. The first TGVs were equipped with motors that each produced 535 kilowatts of power; the second generation uses motors that generate 1,100 kilowatts. Motors on the faster, new-generation trains will each put out 1,100 kilowatts but will be 40 percent lighter than the latest TGV motors. These weight reductions have been achieved by design changes as well as by using lighter materials.

On a per-seat basis, the TGVs are among the lightest trains in the world, but researchers continue to examine all parts of the train for other ways to reduce the load on the tracks. For example, transformers, which convert electricity of different power levels into voltages and frequencies required by the train's motors, are among the heaviest parts of the train. By building transformers from cobalt-alloyed steel and aluminum sheets instead of from copper wires, we have recently brought the mass of those devices down to 7.5 metric tons from 11 metric tons.

New-generation trains will carry the lighter transformers. They will also save on the weight of their electronic equipment, through use of a compact new device known as an insulated gate bipolar transistor. Such transistors will precisely control the electricity delivered to the traction motors. This is the first time this kind of transistor has been used to produce such high-power outputs. We have put a lot of work into the seats as

well. To save a few kilograms per seat, those in new-generation TGVs will be made of carbon fibers, magnesium and composites.

Stop Smoothly, Go Quietly

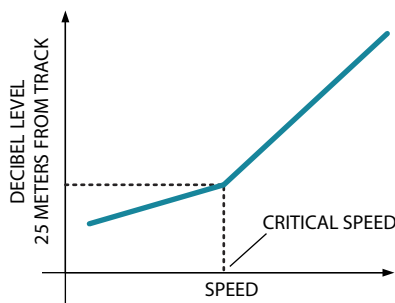
Innovations that encourage high speeds have to be accompanied by technologies that enable the train to stop efficiently without jolting passengers or derailling the train. The first generation of TGVs employed a disc-braking system resembling those found in racing cars. It was advanced for its time and quieter than conventional brakes, but it still relied on friction—that is, on something pressing on the discs (which themselves are on the axles) to dissipate kinetic energy and thus stop the rotation of the wheels. Operating such brakes consumes energy and also causes wear and tear on the braking components and the bogies.

To save on fuel and on maintenance expenses, the newer TGVs complement disc brakes with state-of-the-art “dynamic” braking systems. These help to stop the train by converting mechanical energy from the traction motors back into electricity. This electricity can then usually be recycled—passed to the overhead catenaries for use by moving trains up or down the line or perhaps to the air-conditioning or to other electrical components of the train. In the new-generation trains the braking system will dissipate some unneeded electricity by feeding it safely into the tracks as heat. More than 90 percent of the deceleration in the new-generation trains will be achieved through the dynamic braking systems.

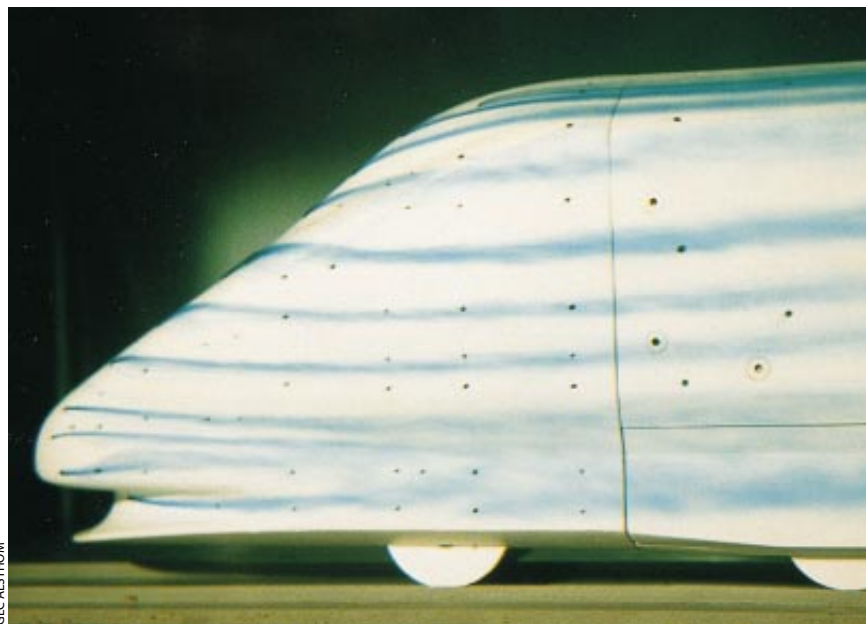
None of these innovations would be of any value if manufacturers failed to control the substantial noise generated by speeding trains. Most of the sound derives from the interaction of wheels and the rail and from wind passing over and under the train. At high speeds, the sound level increases exponentially. The rise caused by aerodynamic effects is especially huge, proportional to the sixth power of the speed.

The least noisy shape is the smoothest one, and so we strive to limit edges not only to reduce drag but to minimize the annoyance to passengers and people who live near train lines. But not all components can be made to have smooth contours, including the bogies. To cope with these realities, we shield underframe devices with aerodynamic deflectors that reduce wind resistance.

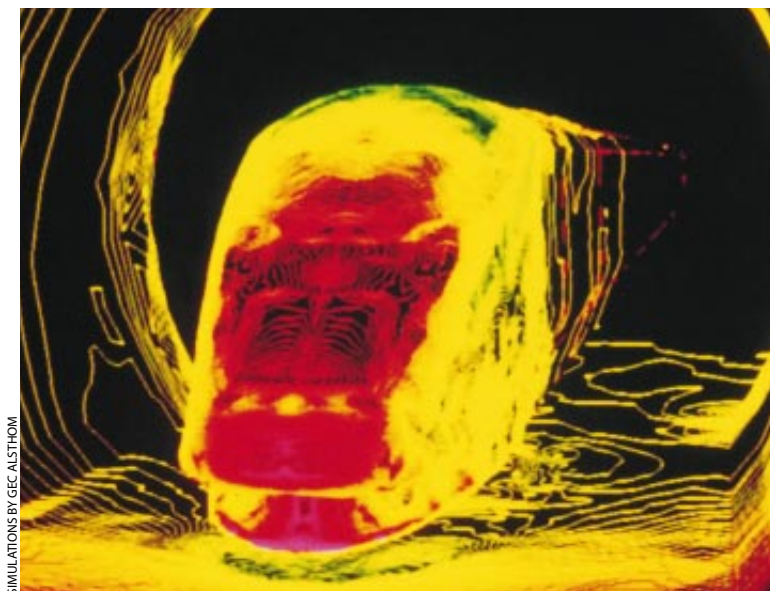
WATER-TUNNEL STUDIES of scale-model trains have informed efforts to optimize the aerodynamic properties of future TGVs. The relatively straight lines formed by green dye on the model below mean the model has a good configuration. Beyond enhancing speed, an aerodynamic design limits noise. Noise levels jump abruptly at some critical speed that varies for each train but is often in the neighborhood of 300 to 350 kph (*graph*).



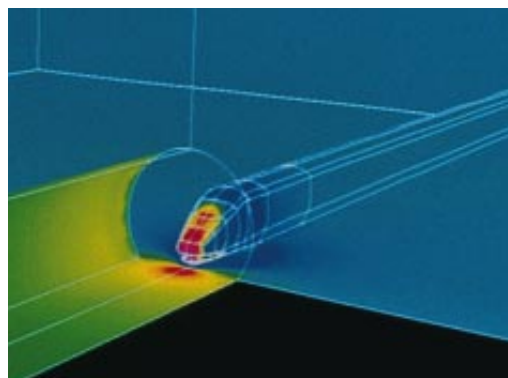
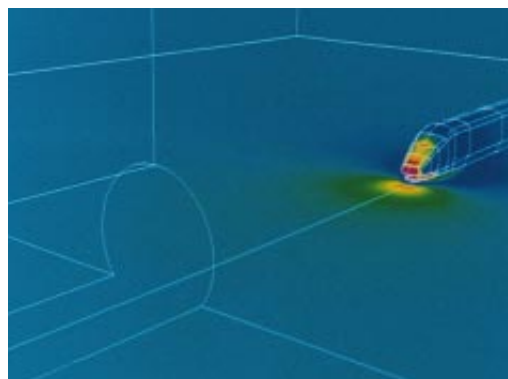
BRYAN CHRISTIE



GEC ALSTHOM



PRESSURE WAVES, which can cause pain in a passenger's ears, arise when a train enters a tunnel; the waves travel the length of the tunnel and back again. Such waves have been measured in computer simulations; red indicates the highest pressure, followed by yellow and green. The results of recent simulations indicate that long noses on trains can help minimize the waves, as can certain modifications in the shape of the tunnel itself.



We also use only one pantograph instead of the multiple pantographs found on more conventional trains. It is situated on the rear locomotive, and a cable brings electricity to the second power unit, at the front of the train. To reduce noise further, we have redesigned the pantographs for the new-generation TGVs, giving them fewer edges. The Japanese, however, have another solution for the pantographs: encasing them in aerodynamic chimneys.

Contact between the wheels and the rails contributes to the noise level by producing vibrations that are capable of exciting both elements. Guided by computer simulations, we have now subtly altered the design and, in places, the thickness of the wheels to reduce the noise without increasing their heft; the improved wheels, which have been test-

ed extensively, will roll on future TGVs.

In France, existing high-speed trains rarely pass through tunnels. But elsewhere in the world, railroads sometimes have to be routed through such passageways. As the vehicles enter tunnels, they create pressure waves that run the length of the tunnel and back again at the speed of sound. These waves created by high-speed trains can cause pain to eardrums and can potentially shatter glass.

Computer simulations and other experiments indicate that the intensity of the waves can be minimized by changing the shape of the trains, such as giving them a longer nose. Other ways to ensure passenger comfort include making trains airtight and controlling cabin pressure internally. Optimizing the shape of the tunnel can help as well.

Research on the new-generation trains

has demonstrated that speeds of up to 360 kph are technically and economically realistic. And we are already constructing a locomotive for testing in 1999 that will move a full train at 400 kph. Indeed, in anticipation of success, track systems under construction in France are already being built to handle equipment rolling at that higher rate. Even 400 kph could conceivably be bettered, although whether the fuel required to achieve significantly higher speeds will be worth the cost is an open question.

In Europe, financial considerations are now slowing the pace at which planned high-speed rail lines are being constructed. But building does continue. It seems reasonable to predict that speeds of 400 kph could be commonplace on the new tracks early in the next century. SA

The Author

JEAN-CLAUDE RAOUL, who joined GEC Alsthom in 1983, is technical director of GEC Alsthom Transport in Paris. He is also coordinator for industry and interoperability of the high-speed network in Europe and a member of the scientific councils of the Mechanical Laboratory of Lille and INRETS, the French National Research Institute for Transportation Systems.

Further Reading

RESEARCH DETERMINES SUPER-TGV FORMULA. François Lacôte in *Railway Gazette International*, Vol. 149, No. 3, pages 151–155; March 1993.
EUROPE'S HIGH-SPEED TRAINS: A STUDY IN GEO-ECONOMICS. Mitchell P. Strohl. Praeger, Westport, Conn., 1993.
SUPERTRAINS: SOLUTIONS TO AMERICA'S TRANSPORTATION GRIDLOCK. Joseph Vranich. St. Martin's Press, 1993.
THE 21ST CENTURY LIMITED: CELEBRATING A DECADE OF PROGRESS. High-Speed Rail/Maglev Association. Reichman Frankle, Englewood Cliffs, N.J., 1994.
HIGH-SPEED RAIL: ANOTHER GOLDEN AGE? Tony R. Eastham in *Scientific American*, Vol. 273, No. 3, pages 100–101; September 1995.



Fast Trains: Why the U.S. Lags

The reasons are more political than technological

by Anthony Perl and James A. Dunn, Jr.

Impressed by Japanese bullet trains, French TGVs and German ICEs, Americans inevitably ask, “Why don’t we have such high-speed trains here?” The trains, which roll at speeds in excess of 125 miles (about 200 kilometers) per hour in Japan and Europe, are missing in the U.S. mainly because of political and social factors, not a lack of technical know-how.

The nations that successfully pioneered high-speed rail had the financial and organizational means to pursue innovation in their train systems. They already owned the railroads and had national public enterprises in place for running them. With trains and tracks an entrenched public responsibility, governments faced with clogged roadways and strained airports were willing to make substantial long-term investments in railroad infrastructure and were prepared to provide the subsidies needed to offset operating deficits.

Moreover, the tradition of public funding for railways had maintained relatively strong enthusiasm for passenger train services and had nourished a vibrant manufacturing industry for passenger rail equipment. Consequently, bureaucrats, legislators, interest groups

and the general public all saw benefits from developing high-speed rail lines, and opposition remained feeble.

In contrast, 20th-century American railroads have not received strong, consistent support from the national government. America’s lawmakers have rejected public ownership of freight railroads and have authorized only relatively modest and sporadic capital investments in railroads of any kind. And they have excluded railroads from trust-fund mechanisms that have historically provided ongoing investment in the infrastructure for highways, airports, waterways and even mass transit.

Little to Go On

The lack of sustained federal commitment to railroads becomes even more evident when Congress’s relation with Amtrak is considered. Amtrak is the closest thing in America to a nationalized railway. It was created in 1971 as a “quasi-public, for-profit” corporation, meaning that the government subsidized it, required it to operate on certain routes and wanted it to turn a profit (although the “for-profit” part of its mandate was more of a politically imposed hope than a commercial possibility). Whether public financing should continue indefinitely was left ambiguous.

Amtrak has struggled to improve its bottom line while maintaining unpro-

fitable routes and reinvesting in equipment and infrastructure. It has faced a series of liquidity crises and repeatedly gone to Congress for supplemental appropriations. More recently, it has borrowed money to raise cash for major purchases and to meet its payroll.

Clearly, the federal government is unlikely to spearhead efforts to build a network of high-speed rail lines across America. What factors, then, could enable such a network to become a reality? A review of several past and ongoing high-speed rail projects offers some idea of the needed ingredients.

Amtrak has pursued the most promising initiative so far. In the early 1990s the corporation convinced Congress that improving speed and upgrading service in the densely populated Northeast corridor, from Washington, D.C., to Boston, could relieve pressures on existing roads and airports and improve Amtrak’s financial condition.

The plan couples federal funding for electrification of the New Haven-to-Boston stretch (the rest is already electrified) with an innovative financing package for European-designed high-speed equipment. High-speed trains and tracks would make Amtrak competitive with airlines in both the Washington–New York and New York–Boston markets. In March 1996 Amtrak chose a consortium that includes GEC Alstom (the maker of the French Train à Grande Vitesse, or TGV), and Bombardier of Canada to build the new high-speed trains, which are slated to run—at speeds of up to 150 mph—by 1999.

A number of state-level efforts have also been attempted. In the early 1980s, for example, Ohio developed plans for a high-speed train network connecting 13 Ohio cities and two out-of-state destinations—Detroit, Mich., and Pittsburgh, Pa. Most of the projected \$8 billion needed for the undertaking was to come from a specially dedicated one-cent statewide sales tax. In 1982, however, voters rejected the tax. The state’s high-speed rail prospects have languished in the planning stage ever since.

Ohio’s experience demonstrates that publicly funded high-speed rail is not easily sold to voters. Yet an attempt to



PUBLIC MONEY has long been poured into highway infrastructure but has never been channeled consistently to railroads. Lack of dedicated funding has been one of the major impediments to the development of high-speed train lines in the U.S. Raising private money for high-speed rail has also been problematic.

obtain fast trains for Texas through strictly private finance proved equally problematic. In 1989 the state legislature created the Texas High-Speed Rail Authority to oversee the development of a fast rail system in the "Texas Triangle" between Dallas, San Antonio and Houston; it also stipulated that the system be financed purely by the private sector. In 1991 the rail authority awarded a 50-year franchise to Texas TGV, a consortium that included GEC Alsthom. But the consortium initially had difficulty amassing the needed equity capital, so it was given a one-year extension.

In the meantime, public criticism and opposition to the plan mounted across the state. Southwest Airlines, with dozens of daily flights between the three cities, attacked the scheme and fought it in the courts. Farmers and other landowners turned out in droves to express their concerns at public hearings. A group calling itself DERAILED (Demanding Ethics, Responsibility and Accountability in Legislation) lobbied for the repeal of the High-Speed Rail Act. Finally, a major American partner in the consortium admitted that private funding could not cover 100 percent of the costs and withdrew from the arrangement. In August 1994 the rail authority formally canceled Texas TGV's franchise, ending the initiative.

Another state-level approach has sought to blend public and private funds to get high-speed rail off the drawing board. In 1996 the Florida Department of Transportation (FDOT) committed \$70 million (plus an inflation adjustment) annually for 40 years to a high-speed rail program authorized by the state legislature. These funds come from a part of Florida's gasoline tax earmarked for nonhighway expenditures. By contributing to the expensive environmental impact studies, engineering design work and legal permit costs that must be handled before any construction begins, Florida significantly reduced the risks to private investors.

In the same year, the FDOT selected the Florida Overland eXpress (FOX) group to design, build and operate a high-speed line between Miami and Tampa, via Orlando. This consortium

is made up of GEC Alsthom and Bombardier, as well as Fluor-Daniel, an American construction engineering firm. Current plans call for TGV-type trains to begin operating between Miami and Orlando in 2004, with service extended to Tampa by 2006. All-new high-speed track will be built, but the route will follow existing rail corridors for some 65 percent of its 320-mile length. The state of Florida will own all the infrastructure. FOX will purchase and operate the new trains as a private, unsubsidized business and pay the state a fee for the use of its track. FOX has promised to contribute nearly \$350 million in equity and loan guarantees to the project.

With the state's proposed contribution limited to approximately \$2.8 billion, an obvious financial gap has to be filled to reach the estimated price tag of \$5.3 billion. The FDOT has called federal financial participation "critical to the success of the project," and its financial plan calls for some \$300 million in direct federal funds, plus loan guarantees, to support the issuance of construction bonds. So far Congress has not authorized any construction funds or loan guarantees.

What Is Needed

These efforts and others suggest that, at a minimum, the following pieces must fall into place if the development of high-speed rail is to take off in the U.S. First, the federal government will have to make a substantial financial investment in the needed infrastructure.

Second, state governments will have to provide steady support for promising high-speed rail projects. Public representatives will have to cooperate to win approval from the public and to overcome the inevitable "not in my backyard" resistance that arises in response to virtually any major public works proposal. Legislatures will also have to commit significant levels of funding to attract federal matching support and to push efforts beyond the phase of paying for consultants and studies.

Third, private investors and managers will have to play a significant role both before and after train lines are established. In addition to providing needed capital, entrepreneurs can help in the planning stages to produce realistic projections of ridership and revenue and to hold down development costs. And private managers will be needed, at least outside of the Northeast corridor, to own and operate high-speed rail lines on publicly provided infrastructure (as is the case with private airline, bus and trucking companies).

Fourth, a high-speed rail line that is finally established will have to show it can attract large numbers of riders and make a profit. It was the economic success of the French TGV, which generated both an operating surplus and enough revenue to repay capital costs, that focused worldwide attention on high-speed rail as a potential solution to many transportation needs. If Amtrak's high-speed rail efforts in the Northeast corridor can generate significant operating profits, that feat will do more to legitimize high-speed rail in North America than any amount of technical wizardry or political lobbying.

Will these four demands be met in the next 10 to 20 years? We can imagine three plausible scenarios for the coming decades. The first and least optimistic can be labeled "Stagnation and More Studies." At the federal level, partisan bickering and balanced-budget politics result in a withdrawal or diminution of funds for Amtrak's high-speed rail program in the Northeast corridor. This setback for America's most advanced high-speed rail effort makes it more difficult for states to gather local support for their own initiatives. Without a successful example of American high-speed rail, and without realistic hope of federal funding for anything beyond planning studies, state efforts fail to attract other



HIGH-SPEED TRAINS resembling the computer-generated image here are expected to roll in Amtrak's Northeast corridor beginning in 1999. If Amtrak's service attracts large numbers of riders and turns a healthy profit, that success would encourage investment in other high-speed rail projects now being considered in several states across the nation.

Some Legislation to Watch

• **S. 436 (Intercity Passenger Rail Trust Fund Act of 1997) and H.R. 1437**, companion bills in the U.S. Senate and House, would create a rail fund financed by half a cent of the federal gas tax now deposited in the general fund. A guaranteed stream of federal money for rail infrastructure is needed to spark high-speed rail development in America.

• **S. 738 (Amtrak Reform and Accountability Act of 1997)** aims to modernize the structure of Amtrak and free it from many politically imposed burdens, moves that would help it succeed with its high-speed initiative.

• **S. 468 (National Economic Crossroads Transportation Efficiency Act, or NEXTEA) and S. 586 (ISTEA Reauthorization Act of 1997)** are competing versions of bills that would update ISTEA—the Intermodal Surface Transportation Efficiency Act of 1991. ISTEA governs all surface transportation programs and needs to be reauthorized this year. Rail supporters have been trying to incorporate language that would increase the ability of states to apply federal transportation aid to local needs, including high-speed rail projects.

forms of support and eventually stall.

In our second scenario, “Slow and Steady,” Congress agrees to provide a stream of public revenue to Amtrak to ensure that its national system remains in business and that its Northeast initiative has a chance to demonstrate the commercial viability of high-speed rail. Ridership and revenues in the first few years of the new service’s operation meet or exceed Amtrak’s projections. Proof that high-speed rail can lure passengers from airplanes and automobiles and pay its own operating expenses gives momentum to the most promising state initiatives currently under development. In particular, Florida’s FOX plan receives federal aid as a “project of national significance.” FOX is completed more or less on schedule. With Florida showing that organizers can successfully combine state funds, federal aid and private investment, other states scramble to attract private investors and federal money for their own endeavors.

Federal finances cannot handle more than one major new undertaking like Florida’s at a time, however. So states in other regions form new partnerships

with Amtrak or with other rail operators to provide less expensive, but still improved, train service running at up to 125 mph in various corridors—an incremental strategy known as accelerail. By 2010, these improvements have reached a critical mass, setting the stage for more ambitious investments.

The third and most optimistic scenario can be called “Breakthrough and Boom.” In the near future, the U.S. finds a way to integrate rail infrastructure into the existing financial and administrative framework governing assistance to all the other modes of surface transportation. Creation of an intercity passenger rail trust fund or its administrative equivalent is accomplished and allocates a small part of the federal gasoline tax (starting perhaps with half a cent or one cent per gallon) to the states. This becomes the key to a breakthrough and boom in high-speed rail development.

As with the highway-building explosion that followed passage of the 1956 Interstate Highway Act, this breakthrough allows, indeed it positively requires, multiple projects to proceed

quickly in several regions of the country. By 2010 half a dozen new high-speed rail lines are being run by private operators on public infrastructure, with more in the planning stages in most parts of the country. The world’s multinational manufacturers of high-speed rail equipment have all located important design and manufacturing operations in the U.S., which is becoming one of their lucrative markets, providing tens of thousands of new industrial and engineering jobs across the nation.

Yea or Nay?

We wish we knew which of these possibilities will come true. We suspect the prospects for the last scenario are dim. The priorities of Congress lean today toward balancing the budget and cutting taxes, not toward making major investments in railroad infrastructure. The “Slow and Steady” scenario has a chance of becoming a reality if Washington maintains adequate support for Amtrak in general and its Northeast corridor high-speed rail project in particular. Outside the Northeast, Congress will have to unlock the highway trust fund to match state high-speed rail spending before initiatives with real promise can get implemented. Otherwise, the “Stagnation and More Studies” scenario will become the future. Readers interested in making their own prognostications would do well to keep abreast of railroad-related bills that work their way through Congress.

Although the political environment does not now seem overly encouraging for high-speed rail in America, we would not discount the prospects forever. The major, well-documented benefits of the technology in such countries as France, Germany and Japan inspire hope that the climate will change one day and that fast trains will become a significant part of America’s 21st-century transportation system. SA

The Authors

ANTHONY PERL and JAMES A. DUNN, JR., have collaborated since 1993 on analyzing the factors that influence transportation policy in Europe and North America. Perl is director of the Research Unit for Public Policy Studies at the University of Calgary in Alberta, Canada. Dunn is associate professor of political science at Rutgers University in Camden, N.J.

Further Reading

POLICY NETWORKS AND INDUSTRIAL REVITALIZATION: HIGH-SPEED RAIL INITIATIVES IN FRANCE AND GERMANY. James A. Dunn and Anthony Perl in *Journal of Public Policy*, Vol. 14, No. 3, pages 311–343; July 1994.

HIGH-SPEED GROUND TRANSPORTATION FOR AMERICA: OVERVIEW REPORT. Federal Railroad Administration, U.S. Department of Transportation, August 1996.

REINVENTING AMTRAK: THE POLITICS OF SURVIVAL. Anthony Perl and James A. Dunn in *Journal of Policy Analysis and Management*, Vol. 16, No. 4, pages 598–614; Fall 1997.



Two years ago the world's only magnetically levitated train in commercial service shut down. It had carried riders for a 90-second trip between the airport in Birmingham, England, and a conventional rail line 600 meters (almost 2,000 feet) distant. But after 11 years in operation, the high-tech train, which was once hailed as a step into the future, was replaced by humble shuttle buses. The buses lack glamour, but when they break down, replacement parts for them can be readily found—a virtue pointedly lacking in their one-of-a-kind predecessor. The end of the line for the Birmingham maglev may prove a bleak harbinger for the few lingering efforts worldwide to bring to maturity a form of transportation that was long envisaged as tomorrow's high-speed, energy-efficient alternative to trains and to short-distance air travel.

Maglevs use high-strength magnets to lift and propel a vehicle that speeds no more than a few centimeters above a monorail guideway. Transportation visionaries have dreamed of levitated locomotion since the early part of this century, but enthusiasm has risen and waned.

Despite 30 years of development, no maglev has entered service for carrying passengers long distances, and only a few short-hop projects, such as the Birmingham connection, have reached completion. Germany and Japan, which have led the world in maglev enterprises since the U.S. dropped out in the 1970s, have sunk billions of dollars into research and development. But their efforts have yet to progress beyond test tracks that serve as futuristic showpieces. It seems increasingly unlikely that the technology will ever compete in any significant way with airplanes, cars or more conventional trains for trips of up to 800 kilometers.

Why is there so little to show for so many years of work? Any radically new technology comes with inherent cost, safety and mechanical risks, which lead governments and the private sector to choose the most conservative options. As engineers have attempted to perfect maglevs, high-speed forms of conventional rail technology have become an increasingly attractive alternative. Thirty years ago many transportation designers considered about 250 kilometers

Maglev: Racing to Oblivion?

by Gary Stix, *staff writer*



L. KATO Sipa Press

per hour to be the maximum travel speed for a wheeled train rolling down a steel track. But France's Train à Grande Vitesse (TGV) now reaches 300 kph in routine service—and higher speeds, up to 350 or even 400 kph, are under study [see “How High-Speed Trains Make Tracks,” by Jean-Claude Raoul, page 100]. The TGV, in fact, holds the world speed record for a train, having achieved 513 kph in a 1990 demonstration run.

The literal absence of any track record for maglev means that today any nation seeking a timesaving rapid link between cities may well choose to purchase the TGV or another high-speed train—and they are doing so in Europe and the Far East. In fact, new high-speed lines connecting major European cities have been established recently.

Maglev might more easily obtain high speeds for routine operation. The lack of friction between the vehicle and the guideway prevents the wear and tear experienced by a wheeled vehicle and the track underneath it. But over a typical 500-kilometer run, a maglev that attains 450 to 500 kph might save only 30 to 60 minutes compared with a high-speed train traveling from 300 to 350 kph. (A faster maglev would consume too much power because of the rapidly increasing aerodynamic drag.)

The dwindling advantage in speed may hamper deployment of an untried new technology that can be justified only on the few medium-length routes that can support high enough passenger traffic. “Is it really worth it? I think not,” says Tony R. Eastham, an expert on high-speed trains and a professor in the departments of civil and electrical engineering at the Hong Kong University of Science and Technology. “My gut feeling is that maglev will not be implemented in Germany or Japan—although it might take another couple of years for people to reach this conclusion.”

If maglev has any chance left at all, it will probably come in Germany during the next decade. The country is putting together a \$5.9-billion public-private financing package for a maglev, known as the Transrapid, to connect the 292-kilometer stretch between Hamburg and Berlin, set to begin in 2005. It wants the technology as a symbol that the reunited nation remains an innovator—and to quell questions about why it has not built a commercial maglev inside its own borders while trying to sell these flying trains abroad. Yet this past spring German federal officials noted that the undertaking's costs have risen by 10 percent above earlier estimates and that ridership and revenue projections had dropped substantially. The Transrapid, moreover, still faces opposition from one of Germany's state governments and from environmentalists who object to the cost of the project and who dispute the contention that the train has low energy requirements at elevated speeds.

A train that flies just off the ground may continue to hold a certain allure for technophiles. Recent proposals seem to highlight not an expedient push for higher speed but rather a preoccupation with technology for its own sake. Commissioners in Allegheny County, Pennsylvania, have approved bond guarantees for a train that would use superconducting magnets to take passengers from a Pittsburgh parking garage to nearby shops. “Building this one-half-mile system is like opening an air route to Latrobe [a nearby city] using a Boeing 737,” wrote one irate citizen to the *Pittsburgh Post-Gazette*. Like a World's Fair monorail, the main prospect for maglev's future, if any, may be as a high-tech tourist ride.

SA

Straight Up into the Blue

Tiltrotors, which take off like a helicopter but fly like an airplane, will soon make their military debut. Can civilian applications be far behind?

by Hans Mark

Let us suppose that a few years from now, a massive volcanic eruption in the northern Andes traps and endangers thousands of people, as one did as recently as 1985. We will further imagine that the U.S. Navy assault ship *Wasp*—which is carrying two dozen V-22 Osprey tiltrotor aircraft—is off the coast of Colombia in training maneuvers with the Colombian navy. The Ospreys are ordered to render whatever assistance they can.

The scene of the disaster is more than 950 kilometers (600 miles) from the *Wasp*—beyond helicopter range—but well within the range of the Ospreys. Taking off from the ship like helicopters, their rotors tilt forward 90 degrees once they are airborne, enabling them to fly like conventional turboprop airplanes. The Ospreys arrive only hours after the catastrophe and in the next few days fly more than 1,000 sorties and rescue tens of thousands of people. (In the 1985 eruption, 23,000 people were killed.)

Of course, helicopters too could have been used in this rescue, but not as efficiently as tiltrotors. Helicopters could not have made the 1,900-kilometer round-trip flight without landing to refuel. Moreover, helicopters, which generally have top speeds of about 325 kilometers (200 miles) per hour, are significantly slower than the V-22, which can cruise efficiently at 510 kph and reach top speeds of 560 kph. The higher speed, combined with a relatively large cargo capacity—the V-22 can carry 24 fully equipped marines—means that tiltrotors can be more productive than helicopters, in the sense that they can fly more sorties in the same time, delivering significantly more people or materiel.

The V-22, which has been under development since 1982, will be the first tiltrotor aircraft to go into large-scale production. In 1999 the U.S. Marine

Corps is scheduled to receive the first samples of an order that will eventually include more than 500 aircraft over 25 years. After decades of efforts to unite the two great classes of aircraft (and after political and administrative battles that turned out to be almost as challenging as the technical ones), the V-22 is about to deliver on the promise of tiltrotor aviation.

Moreover, just as intercontinental bomber aircraft, such as the Boeing B-47 and B-52, led to the development of large jet-propelled passenger aircraft in 1950s, there are indications that this kind of technology transfer will occur for tiltrotors. Last year the two top contractors for the V-22, Bell Helicopter Textron and Boeing Company, announced that they would develop and produce an executive/utility tiltrotor, the Bell-Boeing 609. This aircraft, which will carry up to nine passengers, will be produced through a private investment of more than \$2 billion. The first flight of the Bell-Boeing 609 is scheduled for 1999. It would seem that we are on the threshold of a revolution in civilian air transportation.

Many Configurations

The tiltrotor concept is almost as old as aviation. The very first proposal for a tiltrotor appears to have been the Baynes Heliplane, conceived in Britain in the 1930s but never built. After the war, in the 1950s and 1960s, many experimental aircraft were built and flight-tested. One of the first important proponents of the tiltrotor was Robert Lichten, who in 1950 founded a company to commercialize the idea. Lichten's first tiltrotor vehicle, the Transcendental Model 1-G, hovered several times in 1954 but never achieved full conversion from vertical to horizontal flight.

Researchers soon realized that any air-

craft that could fly both vertically and horizontally would have to pay a penalty in weight. The extra weight is a function of several factors: relatively powerful engines are needed for vertical flight, along with a more complex mechanical system to redirect the engines' thrust. And at a time when airframes could be made only out of metal, a frame sturdy enough to withstand the different stresses of both vertical and horizontal flight also added to the weight.

The key question in those days centered on which thrust configuration minimized the weight penalty while still providing the needed performance. Developers were considering many different configurations, which were known collectively as V/STOL, for vertical or short takeoff and landing. Some were jet-powered; others were propeller-driven. At present, the only V/STOL in service is British Aerospace's AV-8B Harrier, an attack and air-support jet aircraft.

Among the propeller-driven V/STOL configurations, some had separate engines for the rotor and the propellers; others used the same engines but redirected the thrust by reorienting the propeller-rotors (or "proprotors") after takeoff. This second category of so-called thrust-vectoring aircraft included both tiltrotors, in which only the engines and proprotors rotated to redirect the thrust, and tiltwings, in which the wings and proprotors rotated as a single unit.

Arguably, the most important propeller-driven V/STOL of this era was the XV-3 tiltrotor. Lichten designed it at the then Bell Helicopter Company in Fort Worth, Tex., which he joined in the mid-1950s after his own business folded. In 1958 the XV-3 became the first tiltrotor aircraft to achieve full conversion from vertical to horizontal flight.

Thrust-vectoring systems, researchers eventually found, had a smaller weight penalty than any of the other concepts.





V-22 TILTROTOR can ascend vertically, with its proprotors pointing straight up. An alternative, shown here, is a short takeoff mode, which enables the craft to be more heavily loaded. In this mode the proprotors are tilted forward, providing some horizontal thrust.



PHOTOGRAPHS BY ERIK HILDEBRANDT

Nevertheless, it would take decades before the means were found to direct the engines' thrust efficiently. Not until about 1970 did the vectored-thrust concept win out over the others.

By this time, various tiltwing and tiltrotor experimental aircraft had accumulated enough flight data to allow comparisons. The XV-3 tiltrotor was still flying in the mid-1960s, and its primary competition was the Vought XC-142A tiltwing. Crashes, linked to the XC-142A's mechanical complexity, marred this aircraft's extensive flight research program. The Bell XV-3 also had problems, all of them stemming from the fact that it was seriously underpowered.

By the early 1970s, however, technological advances were showing a way past this weakness. Turboprop engines—which are a form of jet engine in which a turbine drives a propeller rather than a compressor, as in a conventional jet engine—had been around since the 1950s. Advances in lightweight materials and better engine designs led to higher thrust-to-weight ratios, making it possible to outfit a tiltrotor aircraft with two engines—one on each wingtip—rather than putting just one in the fuselage. This configuration would eliminate the complex mechanical system in the XV-3 needed to transfer power out to each wingtip, thereby leading to a much simpler and more practical aircraft.

Administrative developments were also conducive. A new experimental aircraft program at the National Aeronau-

tics and Space Administration chose the tiltrotor concept for further development. The army was also interested in the tiltrotor for various applications, such as medical evacuation. Vertical takeoff and landing capability, combined with the speed and range of a tiltrotor aircraft, would make it possible to evacuate wounded soldiers directly from battlefields to base hospitals, making it unnecessary to have field hospitals.

Thus, a new tiltrotor development program was initiated in 1971 as a joint NASA-army program with a total budget of \$50 million. The contract to build this aircraft, which was known as the XV-15, was awarded to Bell Helicopter Textron. Designed around a conventional aluminum structure, the XV-15 became an effective test bed for addressing the unique problems in aero-

dynamics, control systems and propulsion that would have to be solved if tiltrotors were to succeed.

A Laboratory That Hovered

At the NASA Ames Research Center, where the project was headquartered for the federal government and where I was director at the time, we performed an extensive series of tests using Ames's full-scale wind tunnel. The XV-15 was small enough for us to put the whole vehicle in the tunnel and simulate the flight regime under many of the conditions it would encounter. In these tests, we turned up a fascinating aerodynamic problem, unique to tiltrotors. Specifically, when the XV-15's proprotors were tilted at certain angles with respect to the wings, a large vortex was

generated over each wing [see illustration on page 114]. This vortex caused the aircraft's tail to vibrate unacceptably. Our solution was to redesign the tail, bracing it and stiffening it so that it could withstand the forces.

Another major problem area was in the development of the aircraft's power transmission system. In rotary-wing aircraft, such as helicopters, this transmission has two functions: it lowers the high rotational speed of the engine to a speed appropriate for the rotor, and it translates the engine shaft's horizontal spinning into vertical rotation for the rotor. In a tiltrotor, on the other hand, no such translation is required, because the engine and propellers move together, from a horizontal to a vertical orientation. This fundamental difference meant that Bell could not rely on its proved helicopter-transmission designs and had to come up with a new transmission for the XV-15. Inevitably, there were kinks. For example, during early

tests, gear teeth broke, requiring that the gears be redesigned to better distribute mechanical stress.

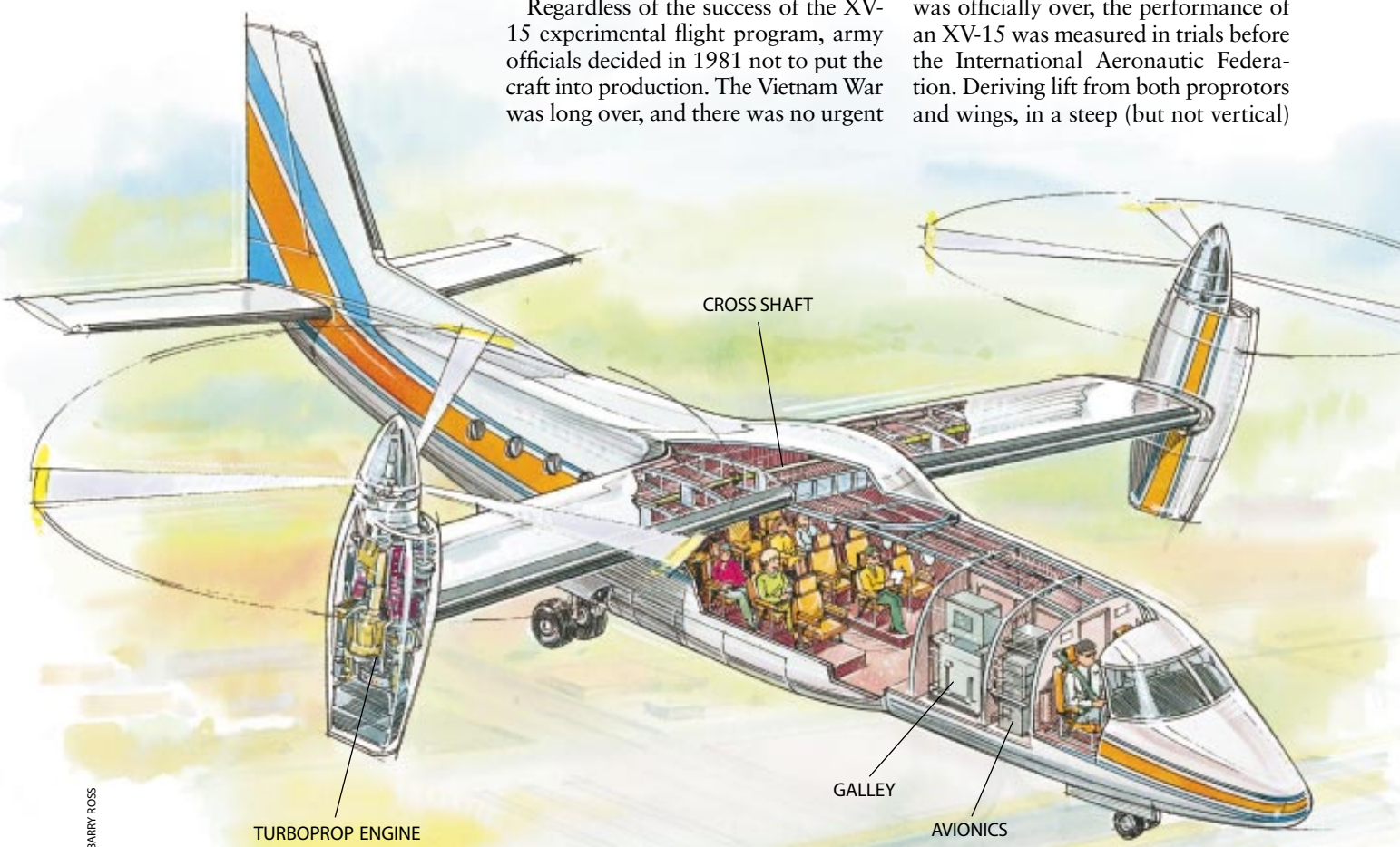
Two XV-15s were built, and the first flight took place on May 3, 1977. The flight-test program was unusually demanding because it sought to show whether the vehicles could meet both NASA's research requirements and the army's operational ones. During the program, conversion from vertical to horizontal flight was executed under many different circumstances, and the aircraft flew at speeds in excess of 550 kph.

One of the most crucial tests was unscheduled. The XV-15 was the first tiltrotor with a safety feature known as a cross-shaft system, which is a mechanical linkage that enables one engine to turn both rotors. This system proved itself, unexpectedly, when an engine suddenly failed during a test flight. The cross shaft behaved as anticipated, enabling the pilot to bring the craft down safely.

Regardless of the success of the XV-15 experimental flight program, army officials decided in 1981 not to put the craft into production. The Vietnam War was long over, and there was no urgent

need for a medical-evacuation aircraft. Although various segments of the army's aviation community still strongly supported development of a tiltrotor, opponents felt that all the army's requirements could be met with helicopters that were already in the inventory or close to it. This faction also pointed out that, unlike the tiltrotor, whose supporting role was to have been limited to medical evacuation, many of the new helicopters carried weapons. In this kind of situation, aircraft with firepower generally garner more support among the military leadership. In the end, the controversy over the XV-15 was an internal military matter and was therefore decided by rank: the undersecretary of the army, James Ambrose, was against the XV-15. Not much additional opposition was necessary.

Though destined to join the long list of practical, effective aircraft that never saw production, the XV-15 would have one more moment in the sun. On March 15, 1990, years after its test program was officially over, the performance of an XV-15 was measured in trials before the International Aeronautic Federation. Deriving lift from both propellers and wings, in a steep (but not vertical)



PASSENGER TILTROTOR for commercial, short-haul flights would have to be large enough to carry about 40 passengers.

ascent, the aircraft climbed to 3,000 meters in just over four minutes and 24 seconds and to 6,000 meters in just under eight minutes, 29 seconds.

Though it never went into production, the XV-15 attracted influential supporters in the 1980s, including navy secretary John Lehman, Senator Barry Goldwater of Arizona and the commandant of the Marine Corps, General P. X. Kelley. Kelley concluded that a tiltrotor would be the best choice for a new troop-transport aircraft to replace the Boeing Vertol CH-46 helicopter. A study by Bell Helicopter had shown that a tiltrotor's combination of range and payload would enable the ships involved in an amphibious assault to remain several hundred kilometers from the beach, while the tiltrotors delivered troops and materiel for the invasion. This advantage, known as increased standoff distance, is important because it makes it harder for the defending forces to attack the invasion fleet at sea.

The tiltrotor transport designed to meet these needs became known as the V-22 Osprey. It was designed and built by two lead contractors: Bell Helicopter Textron, the creator of the XV-15, and Boeing Company. Although the V-22 was in some ways a logical progression, adapting some of the findings and achievements of the XV-15 to a more practical transport aircraft, there were fundamental technical differences and challenges in the two programs.

Flying by Wire

The V-22, for example, would have an electronic fly-by-wire control system—something deemed too costly for the XV-15. In a traditional control system, a host of mechanical linkages physically connect cockpit controls to the actuators for control surfaces, such as the flaps and ailerons on the wings, which allow the direction of the aircraft to be changed. In a fly-by-wire system, all these mechanical connections are replaced with electronic circuits. Fly-by-wire systems eliminate the great weight of the many mechanical couplings in a complex aircraft and also permit more precise control that can improve handling. They are now used in the Boeing 777 and Airbus 340 passenger jets but were still fairly developmental in the early design days of the V-22.

The V-22's most important technical innovation, however, is its use of advanced composite materials throughout

its primary structure. Because of its size, the V-22 would have been too heavy and sluggish—and therefore unable to carry an adequate payload—if it had an aluminum structure, like the XV-15. Besides offering two or three times the strength and stiffness of aluminum in a sample that weighs 25 percent less, composites are much more resistant to corrosion and can be tailored, depending on their intended use, with different characteristics.

Composite materials consist of a plastic, such as epoxy, within which a tough, filamentary material has been embedded to enhance the strength. Most of the V-22's airframe, fuselage, tail and wings are fabricated with a type of composite known as a fiber-reinforced graphite-epoxy laminate (the proprotors are made of a similar composite). Because of the limited history of composites in the primary structure of aircraft, extensive tests had to be devised to ensure that the ways in which composites behave are well understood and are in conformance with Defense Department requirements for military aviation.

Advantageous as they are in many respects, composites do have drawbacks. One of them is that unlike metals, they offer extremely limited protection against lightning strikes. Engineers addressed this problem in the V-22 by laminating into the aircraft's outer surface a fine copper mesh, which disperses the charge of a lightning hit and adds only modestly to the weight.

Although the V-22's composite frame obviously meets all the requirements for military use, a transfer of this technology to the civilian sector will not be straightforward, because the Federal Aviation Administration may set much more rigorous criteria for approving composite structures for use in commercial aircraft. In any event, the materials and the manufacturing procedures that have been developed are important, and they will eventually see major applications in the aircraft-manufacturing industry. The applications include the precisely controlled manufacture of very large pieces of the aircraft, such as large pieces of the wing of the V-22.

The significant advance here is that the pieces are part of the primary—that is, weight-bearing—structure of the aircraft. In comparison, although 9 percent of the structural weight of the Boeing 777 consists of composite materials, none of them are in the weight-bearing structure of the aircraft. The hope and

expectation is that the experience the aviation community gains flying the V-22 (and the B-2 as well) for hundreds of thousands of hours on military missions eventually will persuade the FAA to certify a primary structure made of fiber-reinforced composite materials in civilian commercial service aircraft.

Another "Cancellation"

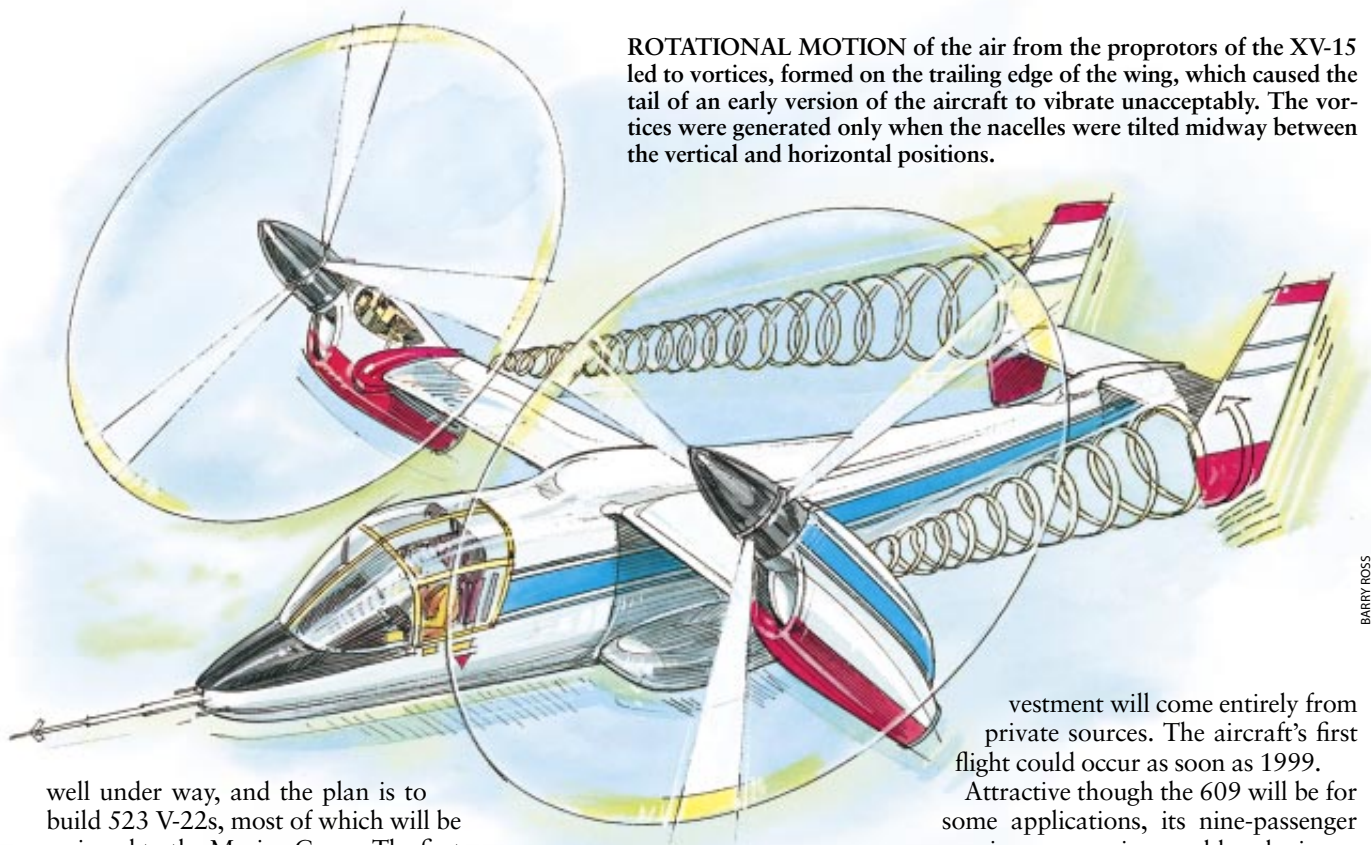
A V-22 first left the ground on March 19, 1989. The aircraft hovered for a short while and did little else during the brief flight. Even as the Osprey was at last showing that a practical tiltrotor transport could be built, a new development was threatening the program. Around the time of this first flight, Richard B. Cheney, President George Bush's secretary of defense, accepted the recommendation of an assistant that the V-22 program be canceled. Thanks, however, to intensive lobbying in Congress by the Marine Corps, the navy and, especially, Bell Helicopter Textron, the V-22 project did manage to survive.

A devastating blow to the program nearly did end it in 1992. On July 20 the fourth V-22 prototype crashed into the Potomac River, killing all seven people on board. An investigation revealed that the accident occurred because a flammable liquid (either fuel or hydraulic fluid) had leaked and collected in one of the engine cowlings while the aircraft was in airplane mode—with the cowlings horizontal. When the pilot changed the aircraft to helicopter mode in preparation for landing, the fluid was sucked into the engine's air intake and burst into flame, damaging not only the engine but also the cross shaft that would have transferred power from the other engine, enabling a safe landing. The flaw, which was unrelated to the viability of the tiltrotor principle, was corrected by putting a drain in the cowlings, so that pooling could not occur.

More than one lesson was learned from the tragedy. The Marine Corps was roundly criticized for putting what were essentially passengers on board an aircraft that was still experimental. Because of this lapse in judgment, five more casualties occurred than should have.

After the 1992 crash, the future of the V-22 program again seemed in doubt. Yet the project got another boost later that year after the election of Bill Clinton, whose first secretary of defense, William J. Perry, was a V-22 supporter. Currently the production program is

ROTATIONAL MOTION of the air from the propellers of the XV-15 led to vortices, formed on the trailing edge of the wing, which caused the tail of an early version of the aircraft to vibrate unacceptably. The vortices were generated only when the nacelles were tilted midway between the vertical and horizontal positions.



BARRY ROSS

well under way, and the plan is to build 523 V-22s, most of which will be assigned to the Marine Corps. The first four of these aircraft have already been built and are now undergoing testing; five others are under construction for delivery in 1999. Deliveries are to continue for as many as 25 years. The cost per aircraft has been estimated at \$30 million to \$40 million.

Technology Transfer

With their ability to take off vertically and cruise efficiently, tiltrotors have been found by many observers to have the potential to transform not only military aviation but also the commercial, short-haul commuter air market as well. This market, for trips of about 500 to 800 kilometers, is estimated to be growing at a rate of nearly 7 percent a year in the U.S., contributing to congestion and flight delays at most of the country's major airports. At seven of the top 10 U.S. airports, between 23 and 36 percent of all flights involve airplanes with 50 seats or fewer, traveling less than 800 kilometers.

Tiltrotors could help alleviate this congestion by shifting takeoffs and landings from airports to much smaller "vertiports" that are closer to city or suburban population centers. If, moreover, the flight originates or terminates at an airport in an uncongested area, the tiltrotor could operate like a turboprop airplane on a runway or, alternatively, at a

vertiport on the grounds of the airport, enabling passengers to make connections to other airlines.

These and other issues in nonmilitary tiltrotor aviation were analyzed in a recent report by the Civil Tilt-Rotor Development Advisory Committee, established in 1992 by Congress. Although the committee found that additional research and development would be needed before an airplane manufacturer could commit to producing a large commercial tiltrotor transport aircraft, it also suggested that the time was right for a small tiltrotor developed for the corporate and utility market.

In fact, this latter view already seems to have been vindicated. In November 1996 Bell Helicopter Textron and Boeing Company announced the establishment of a consortium that would produce a civil tiltrotor aircraft for such applications as shuttling executives among scattered corporate sites, transporting crews and equipment to offshore oil rigs, search and rescue operations, disaster relief and medical evacuation, and border patrol. This six- to nine-passenger aircraft, known as the Bell-Boeing 609, would be similar in size and shape to the XV-15 but would make use of the composites and several other advances achieved with the V-22. It is estimated that \$2 billion to \$3 billion would be required to develop the 609, and the in-

vestment will come entirely from private sources. The aircraft's first flight could occur as soon as 1999.

Attractive though the 609 will be for some applications, its nine-passenger maximum capacity would make it too small for use in the market for commercial, short-haul flights. The basic trade-off that influences the size of these aircraft balances economies of scale, which favor a larger aircraft that can generate more revenue per flight during peak hours, and efficiency, which demands a smaller aircraft that does not fly with dozens of empty seats too often. For the short-haul commuter market, a 40-passenger capacity has proved quite successful for numerous turboprop aircraft.

Entry into the commuter market would entail a number of formidable challenges. For example, the FAA's criteria for certifying aircraft used for commercial service are much more stringent than the ones applied to executive and utility aircraft. Just how rigorous they will be for tiltrotors is impossible to say. Yet for conventional aircraft—both fixed-wing and rotary-wing—the FAA requires that the vehicles be designed so that the pilot can maintain "positive control" of the aircraft under various failure conditions. It is reasonable to surmise that the FAA will apply this criterion in some form to tiltrotors.

The most likely anomalous condition that must be anticipated is engine failure. For a tiltrotor in airplane mode, it is possible that the positive-control requirement might be interpreted to mean that the aircraft must be able to glide in for a landing if both engines fail. In the

helicopter mode, a failure of a single engine can be handled by the cross-shaft system, which enables one engine to turn both rotors. For a double-engine failure in helicopter mode, the FAA might require that the rotors autorotate (spin without engine power) while the aircraft is descending, so that the pilot can execute a soft landing.

The most difficult issue the FAA will face in developing certification criteria for tiltrotor aircraft may very well concern the use of composite materials throughout the vehicles' frame, skin and structure. The problem is that there are several possible failure modes in these materials, and it is not clear that procedures can be developed to predict when such failures will occur. Such predictive methods are a basic part of the certification process for the structural metals used in aircraft. Should it not be possible to find ways to certify the composite structures, manufacturers would be left with the less desirable option of building commercial tiltrotor aircraft out of conventional aluminum parts.

The research program recommended by the Civil Tilt-Rotor Development Advisory Committee would eventually give the industry the kind of technical and economic confidence that seems to be necessary before a new, technically advanced commercial transport can be built. There is good reason to believe that the proposed research program will be adopted and that we will see a large commercial tiltrotor aircraft early in the next century, perhaps by 2020. Among the technical problems to be addressed are reducing noise and engine emissions. In the committee's judgment, both the noise and the emissions can be brought below the current federal standards.

Research in areas involving safety will, among other things, be directed toward the problem of how to deal with the loss of one or both engines. In this area also, there are several encouraging options. An intriguing one involves the use of

proprotors that would telescope to change their diameter: a larger diameter would provide the greater lift needed for vertical takeoff, whereas the smaller diameter would permit higher rotational speed and therefore even more efficiency in airplane mode. A larger diameter should also make it possible for the proprotors, in the event of a total loss of power during vertical flight, to autorotate. With autorotation, the spinning of the blades caused by the uprushing air slows the rate of descent, enabling the passengers (if not the craft) to survive. Finally, a research program will be initiated to look at avionics and flight-control systems that will make it possible to use the available airspace more safely and efficiently in bad weather.

Private-Sector Investments

The committee also performed extensive economic analyses designed to establish whether investments in a commercial tiltrotor air-transportation system by the private sector would be attractive in comparison with other transportation modes. The answers to economic questions such as these depend on various assumptions, so they must be taken with a grain of salt. The analyses did show, however, that for regions of high population density, a well-designed tiltrotor air-transportation system based on 40-passenger aircraft could divert more than 10 percent of passengers from conventional airports. In general, flight times would be shorter, because delays caused by congestion would be avoided and because tiltrotors could operate closer to population centers. But it is difficult to estimate how much shorter the flight times would be and how many more passengers would be willing to pay for this advantage. Flights on a tiltrotor would have to be more expensive because of the higher costs to an airline of buying and operating a tiltrotor. The committee's estimate was that

the airfare for some popular routes in the northeastern U.S. would be about 45 percent higher for a tiltrotor flight.

Although many of the technical specifics of a future tiltrotor transport are yet to be determined, one major point is certain: the aircraft will be designed and built in the U.S. Currently the U.S. is the only country seriously pursuing tiltrotor technology. In Europe a consortium called Eurofar was formed in 1986 to build a commercial tiltrotor transport. The consortium produced a tiltrotor design strongly resembling the V-22, but it was never built. In the late 1980s T. Ishida Aerospace Research, a Japanese-owned firm, was started in Texas. The company's design for a tiltwing vehicle was never built.

My personal judgment is that the process of introducing commercial service using tiltrotor transports will be an iterative one in which the various sectors of the civilian air-transportation system—the airport authorities, the airlines, the manufacturers and the federal regulators—will each take steps at some time to further the process. How long will this process take? We have only one historical precedent. The Boeing B-47 bomber, on which all subsequent large jet aircraft are based, flew for the first time in 1947. It took 11 years before the first commercial jet aircraft, the Boeing 707, was introduced into commercial service by Pan-American Airways. If the first flight of the V-22 in 1989 is taken as the analogous event, then we will see the first commercial tiltrotor flight shortly after the turn of the century.

The prediction may be overly optimistic. But it could be argued that behind any major technological breakthrough is a healthy dose of optimism. In any event, it would be a fitting tribute to Bob Lichten if the first commercial tiltrotor flight took place exactly half a century after he took a deep breath and lifted the first experimental tiltrotor very gingerly off the ground. SA

The Author

HANS MARK is professor of aerospace engineering at the University of Texas at Austin, where he also is the John J. McKetta Centennial Energy Chair in Engineering. He was previously chancellor of the University of Texas system (1984–1992), deputy administrator of the National Aeronautics and Space Administration (1981–1984) and secretary of the U.S. Air Force (1979–1981). In the 1970s he was director of the NASA Ames Research Center. Before joining NASA, he was associated with the University of California, Berkeley, first as a student and then as a faculty member.

Further Reading

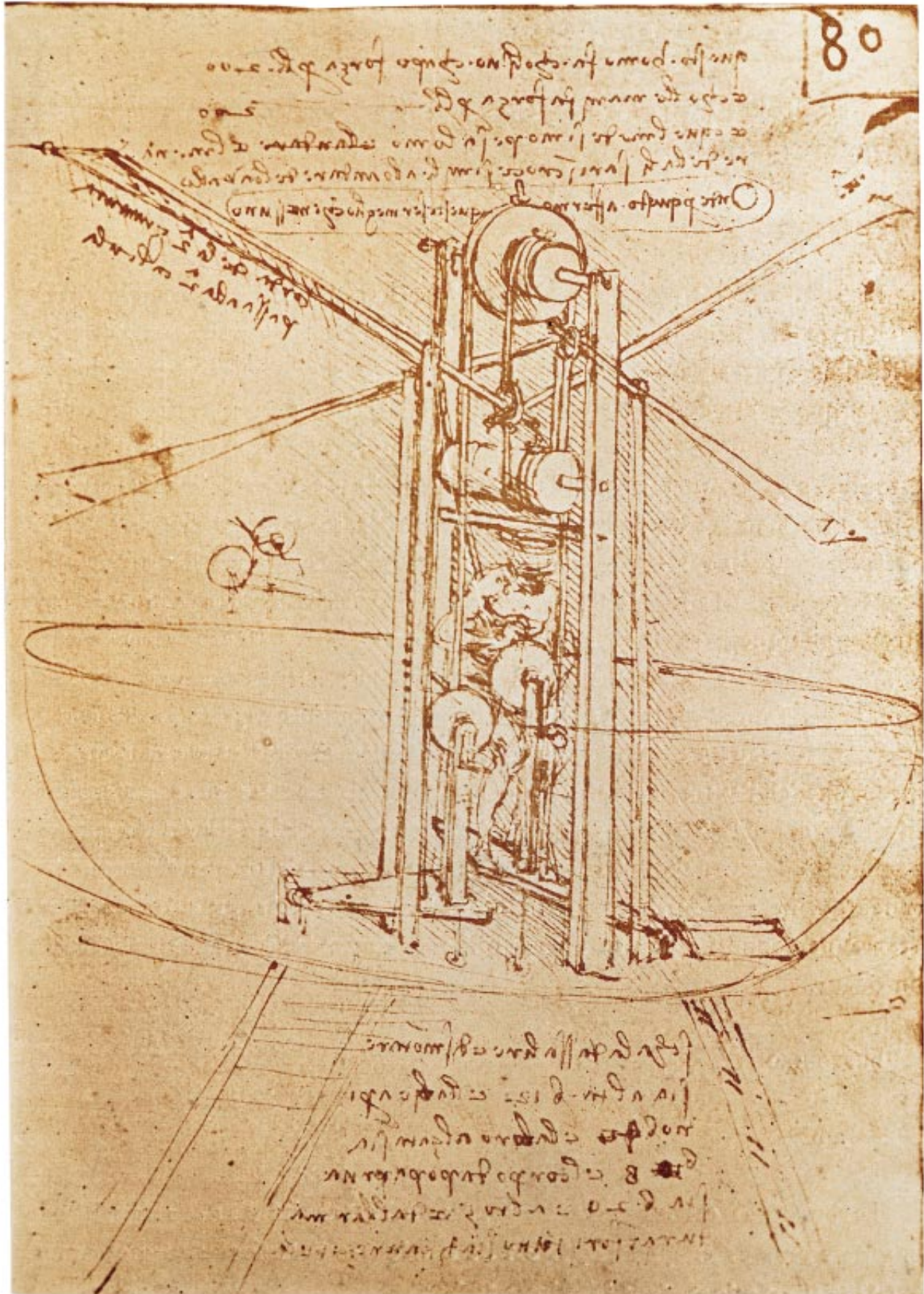
CIVIL TILTROTOR MISSIONS AND APPLICATIONS: A RESEARCH STUDY. Boeing Commercial Airplane Company, Bell Textron, Boeing Vertol and the National Aeronautics and Space Administration (NASA CR 177452), July 1987.

V-22 OSPREY IMPROVES MARKET PROSPECTS. M. E. Rhett Flater in *Vertiflite*, Vol. 41, No. 3, pages 28–32; May–June 1995.

CIVIL TILTROTOR DEVELOPMENT ADVISORY COMMITTEE: REPORT TO CONGRESS, Vols. 1 and 2. U.S. Department of Transportation, December 1995.



The Lure of Icarus



SCALA Art Resource

With new designs and materials, human-powered fliers challenge the distance record

by Shawn Carlson

A well-known parable warns against the dangers of letting excitement get the best of you. As the story goes, Daedalus, a brilliant Greek artisan, crafted two magnificent pairs of wings out of feathers and wax for himself and his brash young son, Icarus. Once airborne, Icarus became so enraptured by the thrill of flight that he ignored his father's warnings to stay close to the sea. When he soared too high, the sun melted his wings, and Icarus plummeted to his death. The parable has survived for 3,000 years, probably because it is so easy for us to put ourselves in Icarus's place. After all, who has not dreamed of flying like a bird?

For a lucky few, that dream is edging closer to reality. And although it may not be quite birdlike, human-powered flight has nonetheless arrived. Relying on ultra-lightweight yet incredibly strong space-age materials, modern-day pilot-powered craft fly with wingspans of 32 meters yet weigh only 34 kilograms. On small airfields all over the world, a handful of dedicated aeronautical engineers are eking ever better performance from their minimalist machines. Record flights have distanced over 115 kilometers (70 miles) and lasted as long as four hours. Though meager by the standards of commercial planes whose engines can deliver more than

LEONARDO DA VINCI drew this odd-looking flying machine. In all, he penned more than 500 drawings and 35,000 words on the topic of human flight. Although his concepts were imaginative, materials science and aeronautics were far too primitive in his day for success.



JUDY MACREADY



STEVE FINBERG

TWO INNOVATIVE DESIGNS won prizes sponsored by British industrialist Henry Kremer. The *Gossamer Condor* (top) was the first human-powered aircraft to demonstrate, in 1977, sustained flight and aerodynamic control. In 1984 *Monarch B* (bottom) set a speed record, completing a triangular, 1,500-meter course in just under three minutes.

10 million watts, these achievements are remarkable given that the engine-pilot can sustain only enough power to illuminate a few lightbulbs. This new generation of flying machines, moreover, has transcended the Icarus-like hubris that undermined earlier design and sometimes brought disaster on pilots.

Human-powered flight has been centuries in coming. All the earliest aeronautical innovators, including Leonardo da Vinci, focused exclusively on human power because no other source was available. When, in 1903, the Wright

brothers proved the potential of internal-combustion engines in aviation with their success at Kitty Hawk, N.C., engineers flocked to the challenge of building machine-powered airplanes. The lure of human-powered flight suddenly lost its luster.

A few well-heeled individuals, however, refused to give up the dream. The decades after Kitty Hawk saw several cash prizes offered for human-powered aircraft that could achieve modest feats of range and aerodynamic control. In 1933 Polytechnische Gesellschaft, a group in Frankfurt, offered 5,000 marks for the first human-powered airplane that could fly around two markers 500 meters apart. They upped the ante to 10,000 marks two years later. The offer did catch the attention of seasoned designers, but the circumstances were not then ripe for success. Similar prizes were offered in the U.S.S.R. and Italy, but all went unclaimed.

Then, in 1959, a visionary British industrialist named Henry Kremer offered £5,000 for the first human-powered aircraft that could demonstrate the same degree

of aerodynamic control as the early Wright fliers: trace a figure eight around two markers four fifths of a kilometer apart. This prize did lead to major advances, although it took another 18 years and a 10-fold increase in prize money before the necessary technology and engineering genius would come together.

They finally did in August 1977, when a 24-year-old cyclist, Bryan L. Allen, pedaled the revolutionary *Gossamer Condor* around the prescribed course and into history. Today the *Condor* is housed at the Smithsonian Institution's

National Air and Space Museum, but to get there the *Condor*'s design team, headed by famed technologist Paul B. MacCready, Jr., had to solve some vexing problems.

The first problem was power. The best athletes can deliver only about 400 watts for long periods. To fly on such little power, the *Condor* needed a wingspan of 29 meters, wider than a DC-9, yet it could weigh no more than a hang glider a third that size. Ailerons, movable flaps on the rear of the wings, proved too inconvenient to steer the plane. To solve the steering problem, the team rigged the wings to twist during turns (a trick also used by the Wright brothers) and mounted a small tiltable stabilizing wing a few meters in front of the pilot.

To enter a left turn, the pilot twisted the wings to increase lift on the left wing and decrease it on the right—the opposite of what is done with the ailerons on an ordinary airplane. But with the unusually shaped *Gossamer Condor*, practically just a flying wing, the greater drag on the left wing yawed the craft (that is, turned it sideways) to the left and simultaneously rolled the left wing down and the right wing up to bring the plane into a coordinated left turn. The small stabilizing wing in front acted like a hawk's tail feathers, forcing against the air to limit the degree of yaw.

Neither Kremer nor MacCready was through yet. Kremer knew that it was Louis Blériot's 1909 flight across the English Channel that ignited Europe's passion for planes. Hoping to generate similar excitement for human-powered aircraft, Kremer soon voiced his intention to sponsor a much bigger prize for the first pilot-powered airplane to duplicate Blériot's feat. This announcement set MacCready's team racing to develop a distance flier. The vehicle, dubbed the *Gossamer Albatross*, was a lean and elegant clone of the *Condor*; it incorporated advanced composite materials and improved streamlining but no new ideas. MacCready's team developed the flier so rapidly that by the time Kremer's £100,000 competition officially opened, the *Albatross* had already been flying for six months. On June 12, 1979, Allen once again flew a MacCready aircraft into history. Battling a head wind that added an unanticipated hour to the flight, he powered his way across the 35-kilometer channel in 169 minutes.

Unfortunately, Allen's Channel crossing did not spark the renaissance in human-powered flight for which Kremer had hoped. So in 1983 Kremer offered another prize, this time for speed—£20,000 for the first flight to complete a triangular 1,500-meter course in under three minutes. The winner would

have to average 32 kilometers per hour. This time the MacCready team's entry was trumped by *Monarch B*, the innovative creation of upstart students at the Massachusetts Institute of Technology.

More than £100,000 in Kremer prize money is still waiting to be won: a £50,000 purse for the first human-powered aircraft to fly a complicated marathon circuit of 40.5 kilometers; £10,000 for a human-powered seaplane; and £50,000 for an aircraft that can fly in minimal amounts of wind, instead of in the dead-still air usually required. These prizes continue to stimulate intense research. To date, nearly 100 human-powered aircraft have been built and flown all over the world.

The current distance record is held by Kanellos Kanellopoulos, who in 1988 nearly completed Icarus's mythic journey by flying in slightly less than four hours the 115 kilometers from the island of Crete to Santorini Island. He flew on average a mere five meters above the water. Like Icarus, Kanellopoulos also fell into the Aegean Sea, when a gust of wind snapped the craft's tail boom. He settled into the surf just 10 meters short of his destination.

Still, a new generation of adventurers is closing in on this record fast. Three hundred volunteer college students and 100 industry professionals worked this past summer in 19 integrated teams to assemble a human-powered plane called *Raven* that could be the most sophisticated aircraft of its type ever conceived. If things go well, sometime in the winter of 1998 *Raven* will obliterate Kanellopoulos's records by nearly 45 kilometers and over an hour of flight time.

Paul R. Illian and Heather A. Costantino of Boeing head the team of engineers who volunteered to design this

AIRGLOW is one of several vehicles that have flown purely for research. Its instruments are designed to unravel the aerodynamic peculiarities encountered by aircraft flying at extremely slow speeds.

The Lure of Icarus



JOHN MCINTYRE

RAVEN (*simulation at left*) will attempt this winter a continuous flight of 160 kilometers lasting about five hours. With a wingspan of 35 meters, it is the largest human-powered aircraft ever built. Its designers used extensive computer modeling (*below*) to achieve the best possible performance for an airplane that must fly on a mere 300 watts of power.

high-tech wonder. *Raven's* load-bearing structures are precision-machined, high-strength carbon graphite. The skin and propeller are fabricated from a ridged carbon-graphite mat and foam sandwich. Although its wings span an awesome 35 meters and have an area of 33 square meters, *Raven* tips the scales at a slight 34 kilograms, equivalent to the weight of all the pillows carried by a Boeing 777. Its specially engineered autopilot will steer the rudder and elevators, its only control surfaces. *Raven* should cruise at an altitude of about six meters at better than 32 kph.

While *Raven* is going for glory, other aircraft are being built for purely scientific reasons, to better understand how vehicles fly at such extremely slow speeds (physicists prefer to say low Reynolds numbers, which quantify how fluids flow over objects). John McIntyre of the University of Cambridge is a member of a small but highly dedicated group of researchers who are systematically working out every aspect of the aerodynamics of these machines. A longtime veteran of human-powered-flight research, his current plane, *Airglow*, is instrumented to analyze subtle aspects of the craft's performance. McIntyre flies it every chance he gets. He insists that it is easy to understand why some people are so passionate about human-powered aircraft.

"This is a world of extremes," McIntyre comments. "Things have to be so optimized. We have an airplane that can be picked up in one hand, flies on the

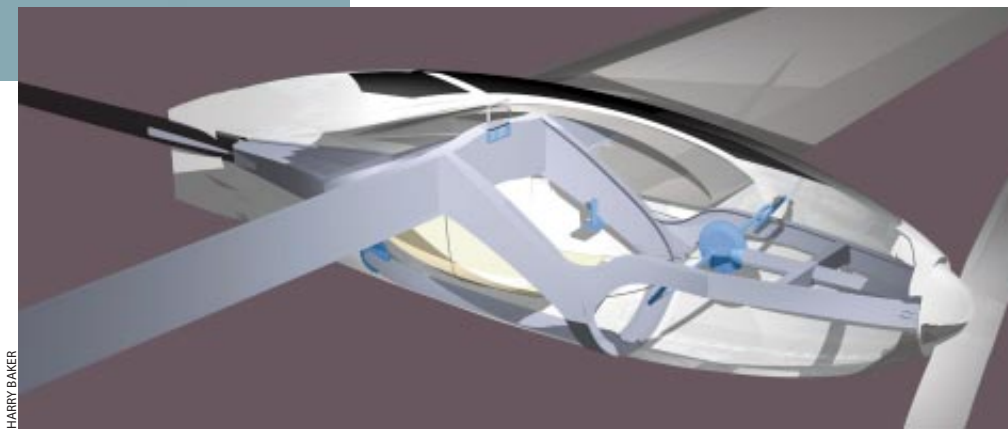
power needed to run a lightbulb and has a wingspan of a commercial aircraft."

Applications of such extreme vehicles have nearly arrived. Making use of the expertise gained on human-powered aircraft, MacCready's team has developed the unmanned and solar-powered *Pathfinder*. In July 1997 *Pathfinder* set a new altitude record for propeller-driven planes by reaching 21.8 kilometers (71,500 feet). The next solar-powered planes will extend *Pathfinder's* 30-meter wingspan to an astonishing 67 meters, larger than that of almost any other plane. These vehicles could stretch the altitude record to 30.5 kilometers (100,000 feet) and, at a somewhat lower altitude, serve as a kind of poor-man's satellite, spending months observing both the earth and the sky while roving the stratosphere on solar power by day and battery power by night.

Tantalizing as this possibility may be, 21st-century historians will most likely mark the legacy of human-powered flight more for what it gave us on the ground than in the air. While seeking

ways of storing energy on board a human-powered aircraft—by means of a battery charged by the pilot's pedaling—MacCready's team gained insights into making efficient use of very limited battery power. Practical electric cars are one example of the resulting technology. AeroVironment, MacCready's company in Pasadena, Calif., developed the all-electric Impact (now marketed as the EV1) for General Motors in 1989. MacCready traces his company's success in this field in no small part to the experience his team gained while running after his fragile flying machines.

The lure of Icarus will no doubt continue inspiring our future engineers for generations to come. The challenge of edging ever closer to the absolute limits of human performance bridges theory to harsh reality like no classroom can, forcing our engineers to think in revolutionary, and not evolutionary, ways. By struggling to bring us closer to Icarus's ancient dream, these young technologists are honing the skills they need to make a better life for us all. SA



HARRY BAKER

The Author

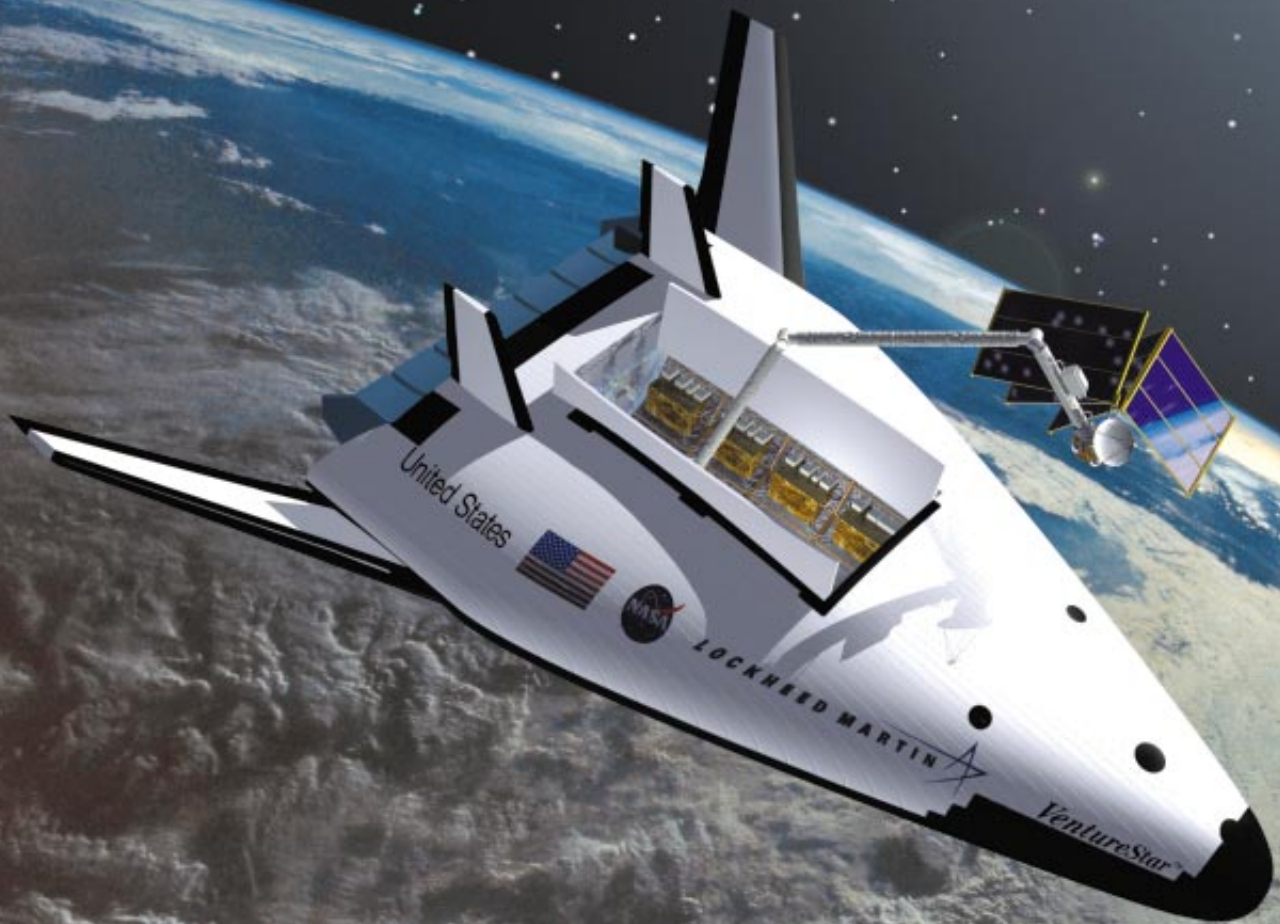
SHAWN CARLSON is founder and executive director of the Society for Amateur Scientists, a nonprofit organization that involves amateurs worldwide in scientific research. He received his Ph.D. in nuclear physics from the University of California, Los Angeles, in 1989 and is an adjunct professor of physics at San Diego State University. Carlson writes the Amateur Scientist column for this magazine and is the creator and principal author of *Physics: The Core* (Harcourt Brace, 1997) a college-level CD-ROM-based textbook.

Further Reading

MAN-POWERED AIRCRAFT. Don Dwiggins. Tab Books, 1979.
GOSSAMER ODYSSEY: THE TRIUMPH OF HUMAN-POWERED FLIGHT. Morton Grosser. Houghton Mifflin, 1981.
HUMAN-POWERED FLIGHT. Mark Drela and John S. Langford in *Scientific American*, Vol. 253, No. 5, pages 144–151; November 1985.
Project Raven information available at <http://mars.sonoma.edu/Raven> on the World Wide Web.
Japan International Birdman Rally information available at <http://www.nasg.com/bm=e.html> on the World Wide Web.



A Simpler Ride into Space



*Technological advances may allow rockets
of the next century to operate much as aircraft do today.
That change might cut the cost of reaching orbit by 10-fold*

by T. K. Mattingly

Forty years after Sputnik gained distinction as the first artificial satellite to circle the earth, activities in space have become ubiquitous. Today communications satellites continuously transmit messages around the globe; orbiting sensors beam down detailed measurements of the terrestrial surface; and robotic craft have started to explore the solar system.

Yet people are still far from realizing the full potential of spaceflight. Some visionaries have long foreseen large orbiting laboratories and expeditions to the surface of distant bodies in the solar system, for example. Others with an interest in commerce look forward to profitable space businesses that will provide services for voyagers as well as goods for earthbound customers. And inexpensive space-based communications and tracking networks could further improve quality of life on our planet.

The main reason current reality lags behind a much richer set of possibilities is that getting into space with the means now available—expendable rockets or the space shuttle—is a complex and hugely expensive undertaking. At going rates, sending a copy of this magazine into orbit would cost between \$1,000 and \$5,000, whereas airborne delivery to any country in the world demands no more than a few tens of dollars.

Several years ago the National Aeronautics and Space Administration studied various ways to lower the economic barriers to space and concluded that one of the most promising avenues was to develop a relatively simple and fully reusable launch vehicle. NASA then picked an industry team led by Lockheed Martin's "Skunk Works" (the division that designed the SR-71 spy plane and the F-117 stealth fighter, among others) to build a reduced-scale, suborbital rocket, designated the X-33, which could test many of the concepts needed to build a

reusable launch vehicle of this kind. That craft would take off vertically, attain orbit with a single set of rocket engines and fly back to the earth, landing horizontally.

If building and flying the X-33 encounters no insurmountable technical or economic hurdles, the team's engineers and technicians will construct a commercial space vehicle—dubbed the VentureStar—that uses the same principles to haul large payloads (which may include astronauts) into orbit. At that point, sometime early in the 21st century, putting people and cargo into orbit should demand, perhaps, only a tenth of what such missions cost today.

Why So Dear?

To understand how the X-33 will pave a less expensive road into space, one needs first to appreciate why current launch costs are so high. Some of the difficulty arises from the unalterable laws of physics. To place a satellite in orbit, a rocket must rise above almost all of the atmosphere and also give its payload sufficient horizontal velocity so that, when it falls back toward the earth, the earth's curved surface "falls" away at the same rate. For example, the horizontal speed required for a standard 100-nautical-mile-high orbit is about 17,000 miles per hour (nearly eight kilometers per second). Adding together the potential and kinetic energy needed to achieve this height and velocity with the inevitable losses that occur over a typical launch trajectory amounts to a sizable quantity of raw oomph. For instance, the energy expended, in little more than eight minutes, to place a space shuttle into orbit could power a typical automobile for millions of miles.

Yet the cost of the propellants carrying all this energy (usually in the form of liquid hydrogen or kerosene and liquid oxygen) is a relatively minor concern.

VENTURESTAR

will be a completely reusable launch vehicle for routine spaceflight in the next century. It will take off vertically, deliver a large payload of freight or passengers to orbit, reenter the atmosphere and glide to a horizontal landing. Unlike the space shuttle or the various expendable launch vehicles now employed, VentureStar will achieve orbit using a single stage of rocket engines.

PHOTOGRAPH BY NATIONAL AERONAUTICS AND SPACE ADMINISTRATION/DIGITAL COMPOSITION BY SLIM FILMS

WINGLESS X-24A, an experimental aircraft constructed during the 1960s, was designed as a lifting body. Destined only for 28 test flights, the X-24A helped to pioneer the fundamental aerodynamic design that will be used in VentureStar and its precursor, the X-33.



Much larger expenses result from the common practice today of building complex, high-performance launch vehicles that are used only once. The reason for this seemingly profligate strategy emerges from some of the basic principles of rocketry.

The payload that can be carried by a launch vehicle depends in large part on the performance of its engines and the ratio of propellant to structural weight. A rocket designer thus has two key tasks: to maximize propulsion efficiency and minimize the amount of mass to be ac-

celerated. But even the best efforts to improve efficiency and reduce mass have historically fallen short of what is needed in practice to attain orbital velocity with one set of rocket engines. That feat requires something like 90 percent of the weight of the vehicle to be allotted to

An Inside-Out Rocket

Traditional rocket engines employ a bell-shaped nozzle to hold in the expanding exhaust gases and direct their motion straight backward. For maximum thrust, those gases should exit the nozzle after they have expanded enough that their pres-

sure has dropped to match that of the surrounding atmosphere. If they leave the nozzle sooner (a), energy is wasted expanding the gas (yellow) far behind the vehicle.

Because a launch vehicle spends most of the time high above the earth, where air pressure is quite low, obtaining high efficiency requires a great deal of expansion to be harnessed—that is, rockets need large nozzles. But at sea level (b), such large bells would expand the exhaust gases so much that their pressure would drop to well below that of the surrounding atmosphere (arrows). The exhaust flow would then tend to detach from the walls of the nozzle, causing dangerous stresses to develop (red). Accordingly, each stage of a multistage launch vehicle has engine bells sized for the altitudes in which they operate.

To design a single engine that works safely and efficiently all the way from sea level to the vacuum of space, engineers at Rocketdyne during the 1960s devised a novel configuration. In essence, they removed half of the typical rocket nozzles and canted them inward, thereby forming a central ramp, or spike. Because the exhaust gases are exposed to the atmosphere, at low altitudes they expand only to ambient pressure as they shoot down the ramp, which is shaped to direct the exhaust straight backward (c).

In the vacuum of space, exhaust gases from the combustion chambers arrayed around the central spike expand to their natural limit and merge in the center, forming an “aerodynamic bubble” of trapped gas between them (d). The



TEST FIRING of a linear aerospike engine took place at Rocketdyne's Santa Susana Field Laboratory in California during the early 1970s.

propellant. Only by using two or more separate stages, each with its own engines and propellant, have designers been able to build practical launch vehicles.

Such “staging” works because it allows segments of the vehicle to be jettisoned en route. That capability provides a great advantage. Quickly lifting a launch vehicle off the ground and out of the thickest part of the atmosphere (so that horizontal speed can be built up without excessive atmospheric drag) requires high-thrust engines and large tanks of propellant to feed them. But such large engines and tanks are more massive than they need be for conditions higher up, where the thrust necessary to accelerate the craft at a tolerable rate is much less. By dropping these heavy components and using more appropriately sized counterparts in the upper stages of the vehicle, the mass that must be accelerated to orbit can be minimized.

Using separate stages has other advantages as well. It turns out that rocket engines are most efficient when their

exhaust gases exit the nozzle at the prevailing atmospheric pressure. At low altitudes, where pressure is high, this effect favors a short nozzle. But in the near vacuum of the upper atmosphere, a longer nozzle is more effective. Staging thus allows the use of nozzles that work reasonably well even as the craft climbs through progressively thinner air.

The introduction of multiple stages (which continue to be employed by all launch vehicles, including the space shuttle) allowed practical rockets to be constructed from the same materials then used to build aircraft. Staging, in essence, made space transportation feasible. Yet this approach often forces much costly hardware to be discarded with each flight. What is more, each mission for such an expendable launch vehicle must be the culmination of a labor-intensive process during which engineers check and recheck everything needed to fly.

Flawless performance is demanded of these rockets because to reduce unnecessary weight, their designs lack backup

systems for all but the lightest components. Also to save weight, many parts must function under stresses that are quite close to design limits. About half the cost of an expendable launch vehicle can be attributed to the many careful inspections and tests required to ensure that its one and only flight goes exactly as planned.

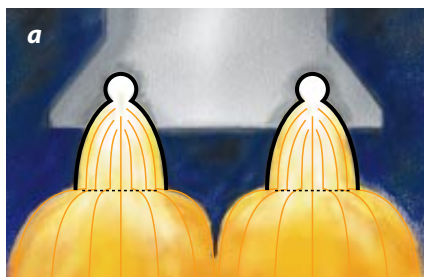
The construction of the space shuttle, with its many reusable components, was supposed to eliminate much of the expense incurred with expendable launch vehicles. But reusability in itself creates another set of problems: everything accelerated to orbital velocity has to be brought gently back to the earth. For the shuttle, the rigors of the return begin below about 65 miles altitude, where atmospheric drag progressively slows the orbiter as it glides toward its landing site. In the process, a great deal of thermal energy must be dissipated without melting or charring the materials on the outside of the vehicle. Also, to fly and land, the shuttle orbiter carries wings, control surfaces and landing gear.

And the shuttle, engineering marvel though it is, has to be minutely checked and reconditioned after every flight. The fuel and liquid-oxygen pumps that feed its main engines, which are severely stressed during a launch, need constant attention, and the ceramic tiles that insulate the craft from the heat of reentry require scrupulous inspection, frequent repairs and waterproofing. Each mission also demands manufacture of a new external fuel tank and retrieval of the solid rocket boosters from wherever they fall in the ocean. These boosters then require a recharge with solid propellant, careful reassembly and thorough checking. All this work entails thousands of skilled personnel.

Considering the technical complexities involved, the space shuttle and the various expendable launch vehicles now flying can be viewed as robust and efficient systems. But when compared with the fleets of airliners serving commercial aviation, these space vehicles appear fragile, inflexible and extraordinarily costly. Aircraft land, exchange cargo, refuel and return to flight in hours, whereas the space shuttle needs months to do so, and newly ordered expendable rockets require a year or so to fabricate. If aircraft were built and flown in this same fashion, air transportation would be too pricey for almost anyone. Until the process of sending rockets into orbit begins to mimic the routine opera-

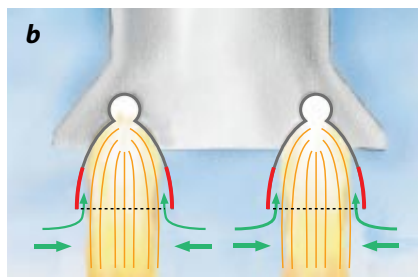
canted geometry ensures that the gases on the outside still go straight back. And the aerodynamic bubble (filled with a small amount of gas pumped in from the engine) serves the same function as a long, solid spike, which is why this combination of thrusters and ramps was dubbed an aerospike. —T.K.M.

HIGH ALTITUDE

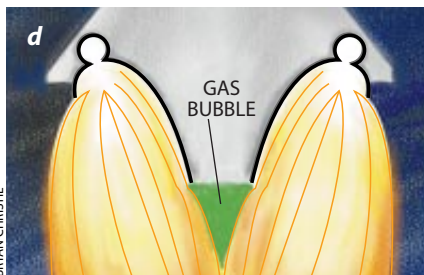


UNDEREXPANSION of exhaust gases in a bell nozzle is inefficient.

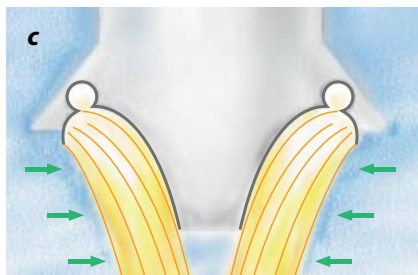
LOW ALTITUDE



OVEREXPANSION of exhaust gases in a bell nozzle may be dangerous.



HIGH-ALTITUDE OPERATION of an aerospike remains efficient because lateral expansion of the exhaust gases is harnessed by reaction against the ramp.



LOW-ALTITUDE OPERATION of an aerospike combustion chamber relies on the central ramp to deflect all exhaust gases directly backward.

BRYAN CHRISTIE

tion of aircraft, space transportation will remain prohibitively expensive for all but a few assignments.

Way to Go: SSTO

What can be done to lower the cost of reliable space transportation? An attractive solution is to develop a fully reusable single-stage-to-orbit (SSTO) launch vehicle that can be flown much like an airliner—in short, the VentureStar. Aerospace engineers can plan such an ambitious project today because improvements in propulsion efficiency and lightweight composite materials are just now bringing within reach a fully reusable single-stage spacecraft.

To maximize propulsive efficiency, the designers at Lockheed Martin have chosen to incorporate an unconventional engine configuration called a linear aerospike, which was pioneered by engineers at Rocketdyne during the 1960s. Unlike the rockets employed in existing launch vehicles, which use a bell-shaped nozzle to control the expansion of exhaust gases, a linear aerospike shoots its exhaust gases across a central ramp. Because they are not enclosed by a nozzle, the exhaust gases can expand to the prevailing atmospheric pressure while reacting against the ramp. This arrangement allows the engine to operate near maximum effectiveness at all altitudes [see box on pages 122 and 123].

Linear aerospike engines should also be advantageous for the VentureStar because they are arrayed broadly along the back of the vehicle. This geometry will allow a certain amount of steering

to be done by throttling individual engines to achieve varying amounts of thrust. In existing launch vehicles, such directional control requires the centrally mounted engines to be affixed to gimbals—heavy, costly and complex mechanisms that physically adjust the orientation of the engine. Eliminating the gimbals will decrease the thrust loads concentrated on a small part of the airframe and reduce weight.

The widespread adoption of composite materials in the VentureStar should also shave off weight. For instance, fabricating major structural components and propellant tanks of graphite-fiber composites instead of aluminum can, in principle, lower the empty weight of a vehicle by about 15 percent. Although this level of savings might seem modest, every pound of weight eliminated from the load taken to orbit automatically reduces the amount of propellant the rocket must carry by nearly eight pounds. Less propellant means the tanks can be made smaller, which in turn leads to a further savings in structural weight, and so on. In the end, the original pound of savings allows for a redesigned rocket that is about 40 pounds lighter at the launchpad.

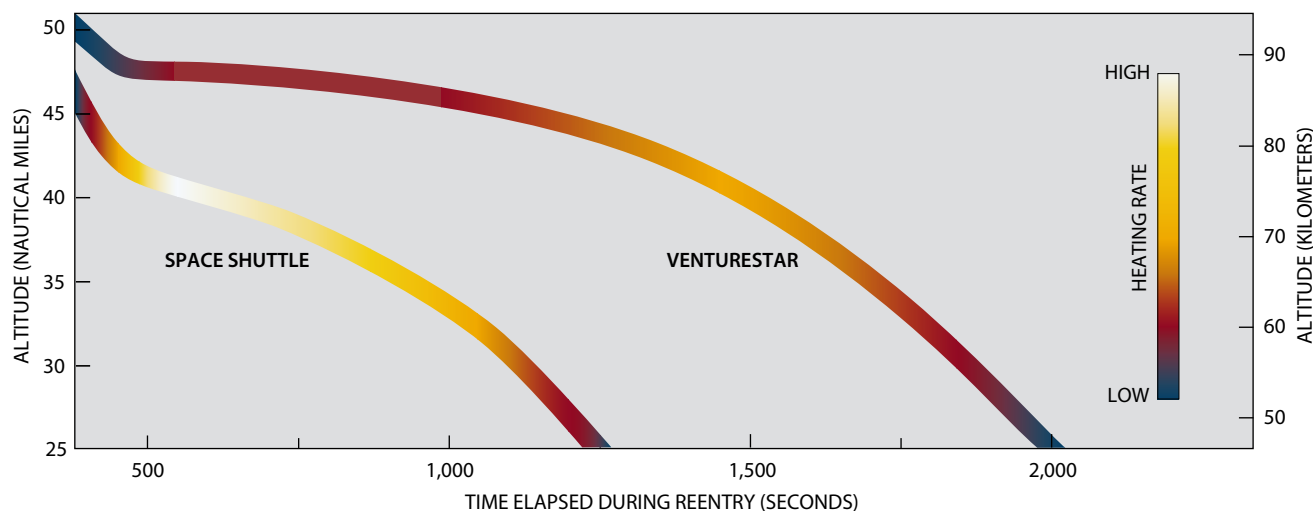
The VentureStar can also be made lighter than some other designs because its shape—called a lifting body—generates significant aerodynamic lift. A light craft can linger at relatively high altitudes as it begins to decelerate during reentry. The tenuous air at such altitudes heats the vehicle comparatively slowly. So the return from orbit would produce less sudden heating than occurs with a more

compact craft such as the space shuttle. The shape of the shock wave formed around a lifting body during reentry also works to minimize the surface area subjected to exceedingly high temperatures. Thus, the need for sheathing the airframe with protective insulation is lessened, which permits engineers to exploit heat-resistant materials that last longer and weigh less.

Achieving significant weight reductions without sacrificing strength or durability is a daunting technical challenge for designers of any single-stage launch vehicle. But success will bring rewards beyond just savings in cost and complexity: a single-stage spacecraft will have all its engines up and running before it takes off, providing a full functional test of each power plant on board before the craft begins to ascend. A single-stage vehicle should thus be inherently more reliable than a multistage rocket.

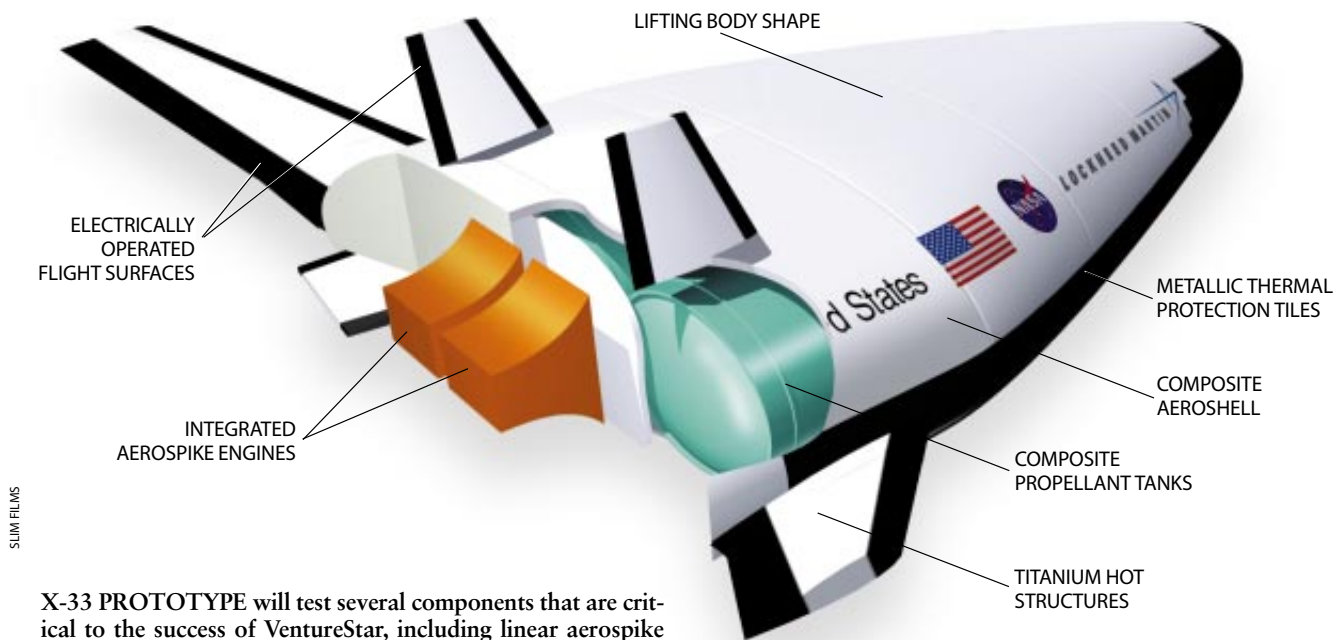
What is more, such a launch vehicle will necessarily throttle back its engines shortly after liftoff (otherwise the force of acceleration would become too great for the airframe and payload to handle). Most of the time, there would be a large excess in available engine capacity. For instance, at the moment of orbital insertion, the engines would be running at a very leisurely 30 percent of maximum thrust. So VentureStar could be expected to tolerate having one of its multiple engines shut down without mishap: others would simply be throttled up to compensate for the loss.

Although the overall philosophy behind VentureStar clearly holds great



REENTRY PROFILE of the VentureStar lifting body will be substantially different from that followed now by the space shuttle. The larger ratio of aerodynamic force to weight for Ven-

tureStar allows it to remain in the thin air of high altitudes for a comparatively long time. Reentry heating is thus more prolonged but much less severe than for the space shuttle.



X-33 PROTOTYPE will test several components that are critical to the success of VentureStar, including linear aerospike engines, composite propellant tanks, electric flight-control mechanisms and metallic thermal protection tiles.

promise for the future, there are many particulars of the design that will need to be worked out before such a craft can be built. Engineers can refine many of the details by modeling the performance of a newly designed device or structure with computers. But there is no better way to test ideas and software simulations than through the real-world experience of building and flying a prototype—the X-33.

Dress Rehearsals

The X-33 will be just 67 feet long, roughly half the size of the proposed VentureStar, and it will not carry a crew or payload. Its mission is only to test how well a lifting body driven by aerospike engines will perform. Sometime around the turn of the century, the X-33 will take off from Edwards Air Force Base in California and accelerate upward for several minutes before shutting down its two engines. It will then coast to a maximum height of about 40

nautical miles (73 kilometers) and fall back toward the earth, reentering the thickest part of the atmosphere before it glides to a landing at a suitable site hundreds of miles away. Although it will never reach orbit (it will attain only half the speed required), its abbreviated reentry will, in fact, be quite stressful because the craft penetrates the atmosphere at a steep angle.

For protection against the heat of reentry, the X-33 will sport heat-resistant metallic tiles, rather than the ceramic tiles now arrayed on the underside of the shuttle. Such metallic tiles (which have already been tested on certain parts of the shuttle) should require much less maintenance than ceramics. The X-33 will test a variety of titanium “hot structures”—components that do not experience the searing temperatures of the tiles but nonetheless must withstand substantial heating. The X-33 will also serve as a test bed for propellant tanks that are made of low-weight graphite composites and fashioned with multi-

ple lobes, a configuration that has never before flown in a rocket. Another innovation that will be tested on the X-33 is the use of high-voltage electrical actuators for the flight-control surfaces (which should be simpler to maintain than the hydraulic devices typically employed).

The X-33 is a fitting member of a long line of pioneering experimental vehicles that later ushered in a variety of operational aircraft and spacecraft. The outgrowth of the X-33—the VentureStar and its cousins of the 21st century—will quite likely follow a similar evolution. These spacecraft and their successors will eventually take their place alongside trains, cars and airplanes—everyday modes of transportation today that were each initially regarded as exotic, costly and of limited utility but that eventually transformed countless lives as they became commonplace. SA

A hyperlinked version of this article is available at <http://www.sciam.com/> on the SCIENTIFIC AMERICAN Web site.

The Author

T. K. MATTINGLY directed the reusable launch vehicle programs for Lockheed Martin during the first year of development of VentureStar. He received a bachelor's degree in engineering from Auburn University in 1958. Mattingly then became a naval aviator and advanced to the rank of Rear Admiral before retiring in 1989. As a Rear Admiral, he headed the Space Sensor Systems directorate of the Naval Warfare Systems Command. He participated in the development of the space shuttle and commanded two missions, in 1982 and 1985. Mattingly also piloted the command module of *Apollo XVI* during its flight to the moon in 1972.

Further Reading

INTERNATIONAL REFERENCE GUIDE TO SPACE LAUNCH SYSTEMS. Steven J. Isakowitz. AIAA Press, Washington, D.C., 1995.

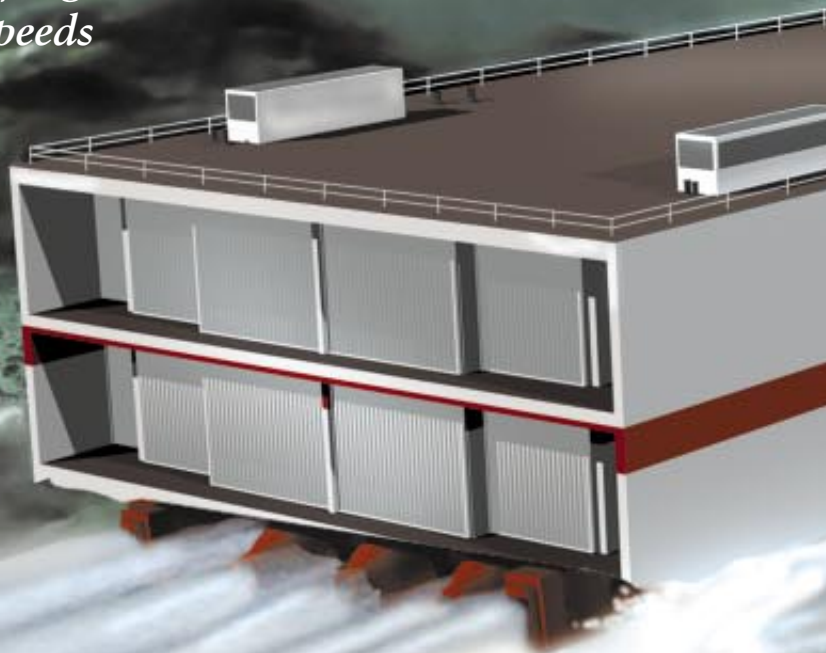
THE HISTORY OF DEVELOPING THE NATIONAL SPACE TRANSPORTATION SYSTEM. Second edition. Dennis R. Jenkins. D. R. Jenkins, Indian Harbor Beach, Fla., 1996. More on the X-33 is available at <http://stp.msfc.nasa.gov/stpweb/x33/x33home.html> on the World Wide Web.



Faster Ships for the Future

*New designs for oceangoing freighters
may soon double their speeds*

by David L. Giles



FASTSHIPS, such as the one rendered above, may well ferry cargo between the U.S. and Europe as soon as the year 2000. Thanks to an innovative hull design and high-powered propulsion system, FastShips can sail twice as fast as traditional freighters. As a result, valuable cargo should be able to cross the Atlantic Ocean in days rather than weeks.

GARDY MCGRATH



For many centuries, ships were the fastest vehicles on earth; they were also capable of carrying the greatest loads. As such, they enabled people from distant lands to exchange their ideas and wares; world trade grew and thrived thanks to Greek and Phoenician vessels, Viking ships and the clipper ship, among other designs. The ancient Greek historian Thucydides pointed out that whoever commands the sea commands, in essence, everything. And until modern times, his dictum stood. Indeed, only recently have vehicles swifter than ships—namely, railroads, trucks and aircraft—emerged.

During this century, international trade has become increasingly dependent on several modes of conveyance and on using them in combination. Many farsighted shippers view the coordination of different types of transportation as the last frontier in boosting productivity. They maintain that keeping pace with the burgeoning demands of international commerce will require a more effi-

cient, synchronized pipeline across the planet—one that, like a moving warehouse, can deliver vital goods within hours of when they are needed rather than within days or weeks.

At present, the weakest links in the supply chain are container ships—freighters that haul their cargo stacked in spacious metal containers. As did vessels built at the beginning of this century, they travel at speeds only a little faster than a running man. Although airplanes can also carry freight, sending cargo by air costs up to 10 times more than dispatching it over the water. And because of delays on land, it still takes three to six days for most airfreight to travel door-to-door between Europe and the U.S. Also, airplanes can transport only a minute fraction of all cargo. As a consequence, many perishable and other time-sensitive goods, on average worth \$10,000 per ton, waste some of their valuable shelf life in transit. Other concerns, too, argue against a substantial expansion of air transport: jets flying at high altitudes release nitrogen oxides, which can harm the environment.

For these reasons, there is a revitalized interest in improving shipping. An array of new technologies—many borrowed from computer science, the aerospace industry and even America's Cup sailboats—are helping naval architects

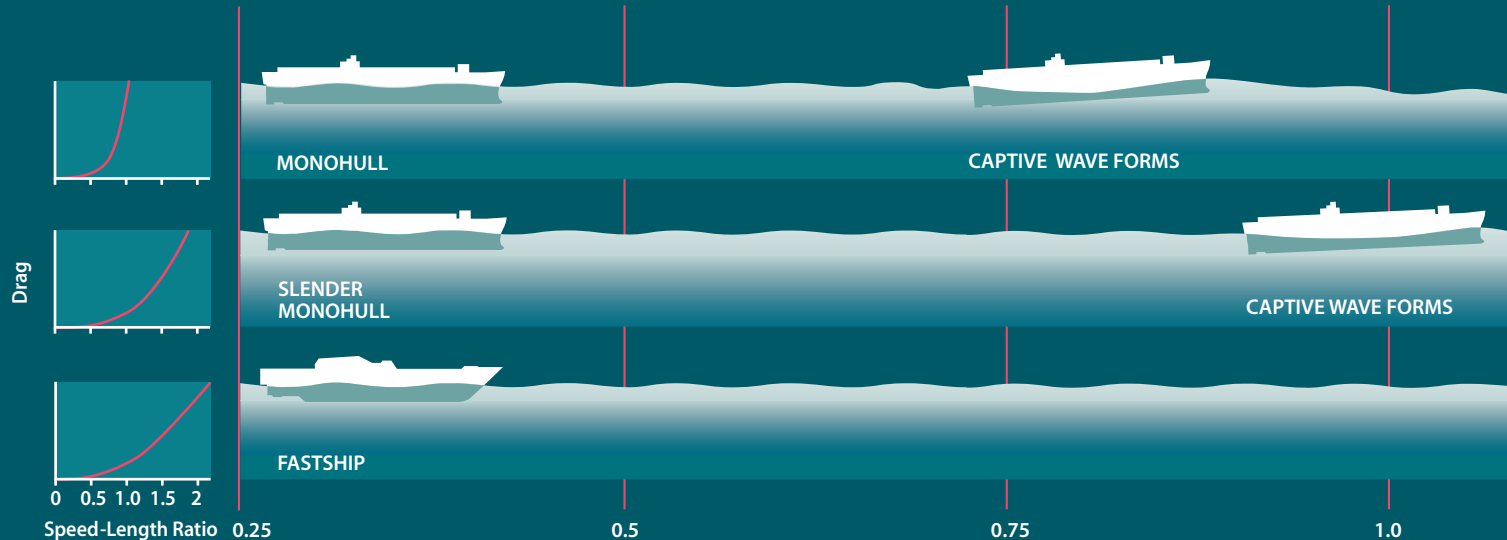
fashion fleets of faster, more reliable ships. Novel propulsion devices and hull designs may enable some of these craft to travel at twice the speed of current cargo carriers.

That shipbuilders have been reluctant to challenge the sea with new hull designs or propulsion systems is hardly surprising. The sea is one of the natural world's most powerful forces. A gentle breeze can excite small ripples on the water that, given time and distance, build into enormous, rolling cylinders of energy. Indeed, a typical ocean wave is nearly three stories tall, 600 feet (183 meters) long and moves as fast as a galloping horse. As one would expect, such waves bear tremendous force. Those that prevail in the oft-traveled North Atlantic during nearly half the winter can slow a container ship at full power by 20 to 30 percent, or four to six knots.

A Seaworthy Challenge

Even in fair weather, ships battle waves of their own making. As a vessel moves through the sea, it displaces the water around it and creates vortices, just as the wind does. The faster the ship travels, the larger these disturbances become, until they merge into a single wave, called a captive wave. Captive waves, like other large swells, can present serious problems. If a vessel increases its speed beyond that of its captive wave, the wave lengthens, and the stern of the ship sinks, or "squats," in the trough between crests. An elongated captive wave then places additional drag on a ship as it tries to climb the wave's back.

In the 1800s William Froude, a British



CAPTIVE WAVES form when any vessel moves through water. At a certain speed, these waves become as long as the ship itself. If the ship tries to go faster, the wave elongates, and the ship “squats.” The captive wave places so much drag on the hull that it cannot climb up the wave’s back and move ahead of it; the vessel sinks into the trough between crests. In fact, a ship’s maximum speed depends on its length and shape. In general, longer

ships can attain higher speeds. Among vessels of the same length, slender monohulls can go faster than conventional ones because they displace less water. In other words, they can achieve a higher speed-length ratio. FastShips are capable of an even higher speed-length ratio because of a concave hollow at their stern. This hydrodynamic curve produces a second wave that shortens the initial captive wave.

naval architect, deduced that the speed of a captive wave depends on the length and volume of the ship that produces it. This wave has the same broad characteristics as a typical wind-generated ocean wave, which varies in speed according to its length and size. For example, a 600-foot-long ocean wave has a height of 27 feet and moves at 31 knots, whereas a 900-foot-long wave is 38 feet tall and travels at about 38 knots. Indeed, the speed of a wave—or a ship—is approximately proportional to the square root of its length. But the maximum velocity of a ship also varies with the volume of the hull. For this reason, naval architects since Froude’s time have used the speed-length ratio (the speed in knots divided by the square root of the length in feet) to describe the relative

performance of ships of different size.

Froude noted several reasons why a ship’s speed-length ratio is typically somewhat less than that of an ocean wave, which has a ratio of 1.25. Most important, he noted, was the wave-making resistance (now called wave or pressure drag) imposed by the captive wave. Although in theory a modern cargo ship of 700 feet should be able to make 34 knots, its optimum speed in calm water is in fact only about 23 knots, which represents a speed-length ratio of 0.87.

In Froude’s view, the speed at which pressure drag from the captive wave becomes significant—in this case, 23 knots—represented the “ne plus ultra,” the point beyond which naval architects dared not venture. Today most designers of traditional ships similarly view

this mark as a fixed speed limit. In fact, a conventional ship cannot go faster without expending excessive amounts of power and fuel to fight the tremendous water resistance involved.

But traditional ships are not the only option. Centuries ago the Vikings discovered one way to reduce pressure drag and attain higher speeds. They simply made their ships longer and slimmer. These “longships” made smaller waves and so moved faster than shorter, wider vessels with the same hull volume and sail area. Longships were less stable than wider craft and less able to carry large loads. Nevertheless, designers continued drafting narrow hulls in later years for warships and passenger liners. In doing so, they found another speed limit on the high seas: above 30 knots, the

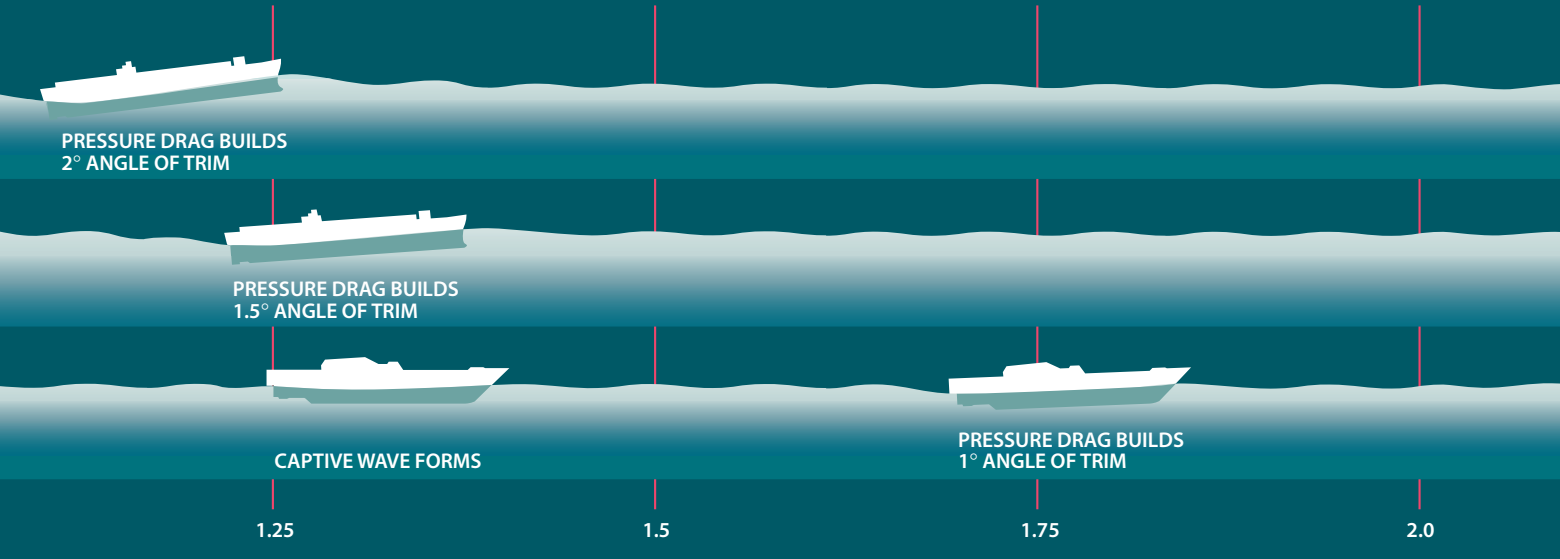
CATAMARAN Speed-length ratio of 2.5

FASTSHIP Speed-length ratio of 1.46

CONVENTIONAL MONOHULL Speed-length ratio of 0.87

COMPUTATIONAL FLUID DYNAMICS, a technique naval architects have borrowed from aeronautical engineers, reveals pressure differences along hulls. Less variation translates into less drag and higher speed-length ratios. Catamarans (*top*), which

are very slender and light, exhibit little pressure variation. The FastShip (*middle*) is much wider and heavier and shows moderate pressure variation. Conventional monohulls (*bottom*) exhibit great pressure variation.



propellers on large ships begin to cavitate—that is, the pressure on their forward surfaces becomes low enough to cause the water to boil, which induces powerful, hull-cracking vibrations.

Breaking the Speed Limit

Confronted with pressure drag and cavitation, naval architects have for many years simply accepted that container ships are bound to be slow. In fact, to balance large, heavy loads, freighters must have hulls that are fairly wide for their length. Thus, few marine engineers have contemplated creating faster hulls or developing novel means of propulsion for freighters. It has become an aphorism that because sea cargo cannot go fast, it has no need to go fast.

This impasse is comparable to one that faced aircraft designers during the 1950s. They found that airplanes experienced a dramatic increase in drag as they approached the speed of sound. Also, the efficiency of aircraft propellers declined at this point. Rather than accept defeat, however, the aerodynamicists labored to reduce the effect of pressure drag and to exploit advantages that might arise by venturing into unknown territory. The solution they came on—jet engines combined with new wing designs—proved ideal: such engines could function in the lower air density found at high altitudes, which rendered propellers and piston engines useless. In return, the thinner air passing over the new airfoils produced less pressure drag even as the aircraft approached the speed of sound.

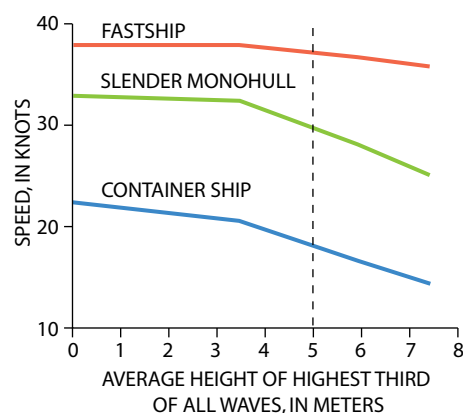
In the realm of modern maritime technology, there are two nascent developments that are the seagoing equivalent of the jet engine: gas turbines and water

jets. These innovations are already employed in many small passenger and car ferries, and it now seems feasible to scale them up to sizes at which they can drive big ships at high speeds. Companies such as General Electric and Rolls-Royce are now developing suitable gas turbines, adapted from aircraft, naval and power-generation applications. Compared with marine diesels of the same weight and volume, the latest gas turbines produce far greater amounts of power for no more fuel. In addition, gas turbines emit, per horsepower, only 4 percent of the sulfur oxides and 5 percent of the nitrogen oxides that diesel engines produce.

Water jets are modeled after the massive power turbines used in the hydroelectric industry. In turbines, flowing water drives a generator. In water jets, the process is reversed: a separate engine spins the turbine blades to produce powerful streams of water that propel the ship. Water jets are ideal for high-velocity cruising because, unlike propellers, their efficiency actually increases with speed. Also, cavitation does not occur at high velocity, because the pressure underneath the boat is sufficient to force water up into the jets and prevent air bubbles from forming on the propulsor blades.

Plunging into the realm of higher power is not worth much if the new technology only pushes traditional ships deeper into the water. So engineers are busy testing three alternative hull designs that, like modern aircraft wings, reduce drag sufficiently to benefit from jets. Some shipbuilders hope to enlarge catamarans, also known as multihulls, which have worked well ferrying cars and passengers in sheltered waters, operating at speed-length ratios around 2.5. These craft are in some ways the

THE EFFECT OF ADDED DRAG IN HEAD SEAS



SOURCE: Boston Marine Consultants/M.I.T. SWAN Codes

HEAD SEAS will slow any vessel, but FastShips fare much better than the rest. For instance, when enough typical waves in the Atlantic Ocean reach five meters in height (*dashed line*), a container ship loses six knots of speed and falls two days late (*blue line*). A slender monohull drops four knots and lags more than half a day behind schedule (*green line*). But a FastShip loses no more than 2 percent of its total speed, which makes it at most two hours late (*red line*).

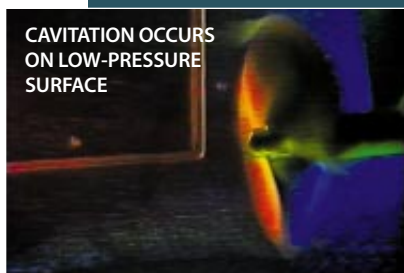
seaborne equivalent of the “flying wing,” which, by virtue of its smaller surface area and lower weight, experiences less drag than any other hull form.

Multihulls consist of two or more narrow hulls spanned by decks. The twin hulls provide increased stability, but they are also liable to split apart—particularly in the rough seas of the open ocean. And the limited buoyancy of their slim hulls means that these vessels must be light, a requirement that only further compromises their strength. For these reasons, it is unclear as yet whether multihulls can carry heavy cargo in high seas.

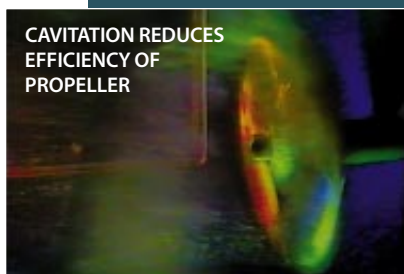
LAUREL ROGERS (top illustration and chart)



NORMAL TIP
VORTICES



CAVITATION OCCURS
ON LOW-PRESSURE
SURFACE

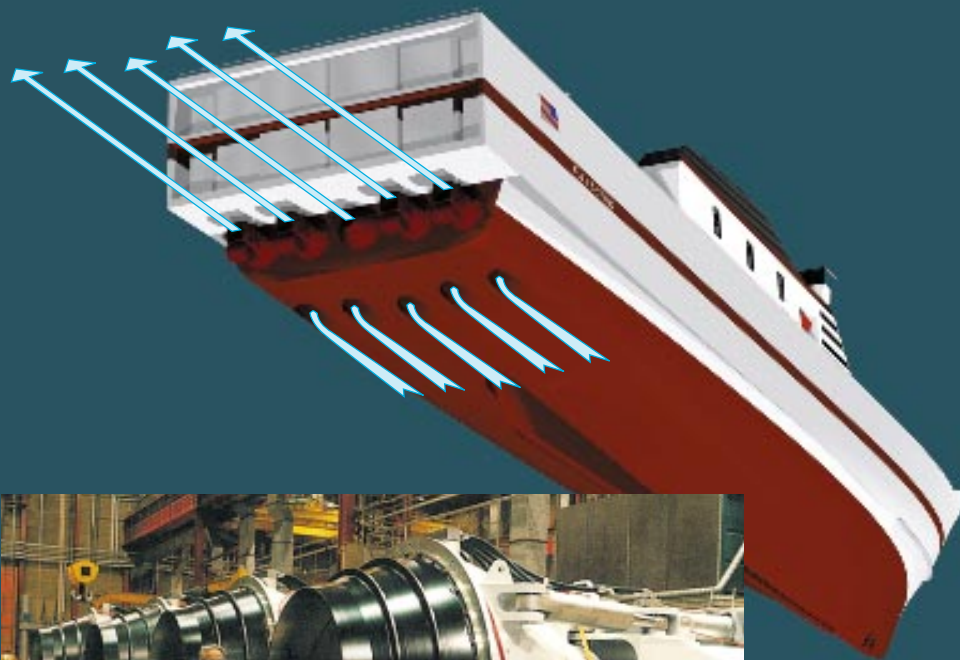


CAVITATION REDUCES
EFFICIENCY OF
PROPELLER



PROPELLER IS
ENTIRELY
INEFFICIENT

KAMEWA GROUP AB



CONVENTIONAL PROPELLERS cannot drive a ship at speeds above 30 knots, because they begin to cavitate: the pressure on their forward surfaces drops so low that the water boils and stirs up damaging vibrations. In contrast, water jets actually become more efficient with increasing speed. Current water jets (*photograph*) are about 60 percent the size of those required by a FastShip. Cavitation does not occur, because the increasing pressure under the hull forces water up into the jets, where high pressure is maintained, and keeps it from bubbling.

The second, more conventional approach for building swifter ships involves enlarging the traditional “destroyer” hull, also called a slender monohull. Being narrow and light, these structures can operate with a minimum of pressure drag. According to my calculations, an empty, 900-foot slender monohull in calm water could attain about 33 knots, representing an impressive speed-length ratio of 1.1. But because these vessels are so slim, they provide limited buoyancy and stability. Thus, many experts worry that slender monohulls might roll and yaw excessively—particularly if carrying tall stacks of containers through rough seas.

Indeed, slender monohulls are much beholden to the weather. Although these vessels reach relatively high speeds in still water, large waves can almost stop them in their tracks. To plow through 30-foot waves, they need more power than propellers can offer. Yet the speed

and resulting pressure under a slender monohull are barely sufficient to justify the use of water jets. For this reason, I believe the equivalent to the slender monohull, in aviation terms, must be airplanes powered by turbine-driven propeller engines, or so-called prop jets: both are too slow to benefit from true jets, and both are highly subject to the weather—making them somewhat unreliable for commercial purposes.

In the third new hull design for freighters—the semiplaning monohull, or “FastShip”—speed is actually the key to reliability in all conditions. One company, FastShip Atlantic, hopes to provide service between Europe and Philadelphia starting in 2000. In collaboration with the department of ocean engineering at the Massachusetts Institute of Technology, investigators are working out how well a cargo vessel with this hull shape will perform on the seas and in sales.

The basic design is not new. FastShip

Atlantic is licensing the patent from my firm, Thornycroft, Giles & Company, and we have tested the hull already in smaller naval and passenger craft and in several test tanks. A FastShip has a deep, V-shaped bow to cut through waves and a wide, shallow rear, with a concave, or slightly hollow, profile underwater. As the ship passes through a speed-length ratio of about 1, this hydrodynamic curve generates a second artificial crest at the stern, which helps to lift the back end of the vessel, thus reducing stern squat. The second wave also shortens the trough between the two waves, and so it keeps the pressure drag from the captive wave in check.

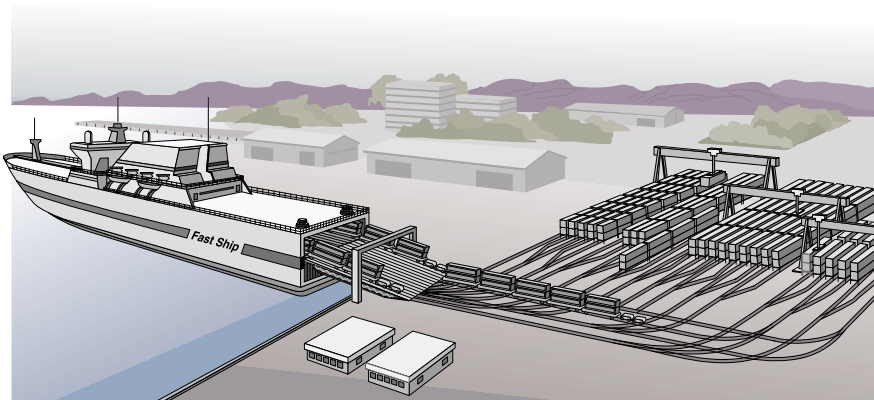
Because the second wave creates increased pressure below the hull, it also minimizes drag generated by heavy ocean waves. So, too, the dynamic lift generated by this hull shape renders FastShips ideal for water jets and affords these vessels increased stability as

they go faster. In contrast, traditional ships tend to pitch, roll, yaw and undergo various hull-slamming actions. In theory, FastShips should be able to maintain speed and stability, without slamming, up to speed-length ratios greater than 2. In practice, however, existing propulsion systems limit the maximum speed of a 750-foot FastShip to about 45 knots—or a speed-length ratio of about 1.5.

The 707 of the Seas

Because of its stability, a FastShip readily maintains speed even in terrible weather. Thus, I would respectfully suggest that it is the seafaring version of the Boeing 707, the airliner that ushered in the jet age by providing widespread, routine service unaffected by bad weather (because it had the power to fly above storms). Strong winds and high seas slow the average container ship from 23 knots to about 17 knots, thereby adding two days to the typical Atlantic crossing. In like conditions, a slender monohull drops its speed from 33 knots to about 29 knots and becomes more than half a day late. But a FastShip, ordinarily traveling at about 40 knots, would hardly slow at all. It should lose at most 2 percent of its speed to ocean waves and incur no more than two hours' delay.

Like the early 707s, FastShips will probably be expensive at first. They need a tremendous amount of power—at added cost—in order to move twice as fast as conventional ships. Yet passengers flocked to jet aircraft in the early days despite the extra expenditure, because they offered greater speed, reliability, frequency of service and comfort. These benefits attracted customers in such numbers that revenue soon covered



LOADING goods onto a FastShip should prove highly efficient. Because the gas turbines powering a FastShip are much smaller than the diesel engines on board traditional freighters, the propulsion system can be placed under the cargo decks, with exhausts up the side of the ship rather than in the middle. Thus, stevedores can roll cargo into the ship lengthwise along rails instead of relying on a crane to pack it in from the bottom up.

the extra investment. Fares ultimately fell below the prices predicted for tickets on slower prop jets. So, too, FastShip proponents believe the reliability, speed and earning capacity of these ships will offset their added expense. Just as in jetliners, improvements in engines, hull materials and other technology will steadily reduce costs.

To help secure competitive rates, however, FastShippers are taking several other steps: First, they are developing highly efficient loading and unloading systems, which should enable them to make more voyages and increase their earnings. Also, they will allow FastShips to call at only one port per voyage. As a result, all containers on board can be replaced at once, and delays at intermediate ports are avoided.

These improvements should reduce the total transit time between cities in Europe and the U.S. to a week or less (currently such cargo takes 14 to 35 days to reach its destination). In addition, shippers will be holding complementary services such as trucks and

railroads to tight schedules so that stevedores can use the same equipment to load exports and remove imports. Such a seamless transportation network might reduce the total door-to-door cost to levels approaching current container rates.

The improvements that new technology at sea and new practices on shore can bring to shipping should help the global economy develop to its full potential in the 21st century. Such advancements may well restore ships to their former status as the driving force behind world trade. Moreover, the benefits that the FastShip and its future cousins will very likely bring to moving cargo around the globe in the next 50 years should be no less dramatic than those that occurred over the past half century, when engineers developed novel technology for moving people and freight by air.

A hyperlinked version of this article is available at <http://www.sciam.com/> on the SCIENTIFIC AMERICAN Web site.

The Author

DAVID L. GILES is the inventor of FastShip and founder of Thornycroft, Giles & Company, which holds the U.S. and worldwide patents for large semiplaning monohulls. He received a master's degree in classics and English at New College, University of Oxford, in 1961. Later that same year he joined de Havilland Aircraft Company (later Hawker Siddeley Aviation and British Aerospace) and began studying aeronautical engineering at Hatfield College of Technology. In 1975 he went into partnership with the late Commander Peter Thornycroft. When Giles is not designing boats, he enjoys sailing them.

Further Reading

MODERN SHIP DESIGN. Second edition. Thomas C. Gillmer. U.S. Naval Institute Press, 1986.

THE LAST GRAIN RACE. Eric Newby. Picador/Macmillan, 1990 (1956).

SEALIFT OPTION FOR COMMERCIAL VIABILITY (SOCV). Available at <http://web.mit.edu/rhmeyer/www/sealift.html> on the World Wide Web.

A COMPUTATIONAL METHOD AS AN ADVANCED TOOL OF SHIP HYDRODYNAMIC DESIGN. Paul D. Sclavounos, David C. Kring, Yifeng Huang, Demetrios A. Mantzaris, Sungeun Kim and Yonghwan Kim. To be presented at the Society of Naval Architects and Marine Engineers (SNAME) '97 Annual Meeting, Ottawa, Canada, October 1997.

CFD IN SHIP DESIGN: PROSPECTS AND LIMITATIONS (18th Georg Weinblum Memorial Lecture). Lars Larsson in *Ship Technology Research (Schiffstechnik)*, published by Schiffahrts-Verlag HANSA, Hamburg) (in press).



Microsubs Go to Sea

Small, maneuverable, self-contained—these tiny submersibles may someday take a human to the bottom of the sea

by Graham S. Hawkes

Ocean covers two thirds of the earth and is home to much of the life on our planet. But humankind's freedom to enjoy this vast submerged habitat is sadly limited. Scuba divers can barely scratch the surface, reaching down to 50 meters (164 feet)— $\frac{1}{225}$ of the way to the bottom of the deepest ocean. About half a dozen aging submersible craft can carry people a little more than halfway to the deepest parts of the oceans; only a few robotic probes and cameras can go deeper. The piloted bathyscaphe *Trieste* plumbed the 11,275-meter-deep Marianas Trench once in 1960, but today only the new robotic Japanese vehicle *Kaiko* can attain these depths.

What makes deep-ocean exploration difficult and the realm so alien is a fundamental property of water: its high density. Pressure increases linearly with depth to a formidable 1,200 atmospheres (16,000 pounds per square inch) at the ocean's deepest point. For the *Trieste*, going to full ocean depth therefore required a very heavy, spherical, steel pressure hull, which in turn needed large

tanks filled with a light liquid to provide buoyancy. In addition, fluid-dynamic drag inhibits the movement of vehicles at the speeds needed to make submerged transportation over huge distances practical. In fact, current submersibles are so slow they take hours to sink or rise the few kilometers to and from depth, and they need to be transported, serviced and deployed from specially designed mother ships.

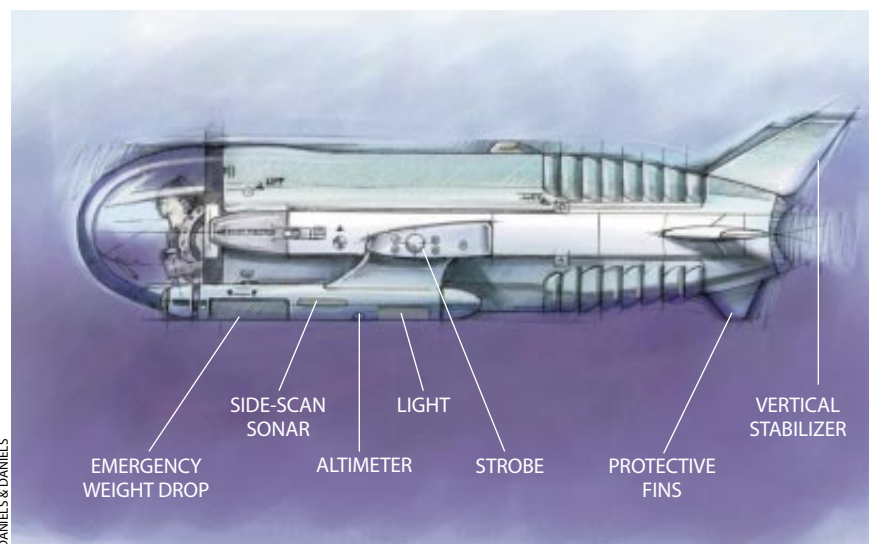
In an effort to circumvent these limitations, most scientists involved in exploration of the deep ocean have moved away from human-occupied submersibles and toward robotic craft. Tethered probes called remotely operated vehicles (ROVs) and small, computer-controlled, battery-powered vehicles (autonomous underwater vehicles, or AUVs) can be operated from any suitable ship. Furthermore, they are relatively inexpensive and, of course, carry no risk to a human operator. Indeed, ROVs have become so popular as tools for the offshore oil industry that economic forces could soon render conventional submersibles extinct.

For humans to lose altogether the ability to explore the ocean depths in person would be unfortunate. Quite aside from the question of whether some sub-sea jobs can best be done by someone on the scene, this loss would be a blow to the human spirit of adventure. For these reasons, it seems a worthwhile goal to develop a better class of deep-sea submersible—not to displace ROVs or AUVs but to offer a complementary in-situ capability for those who want it.

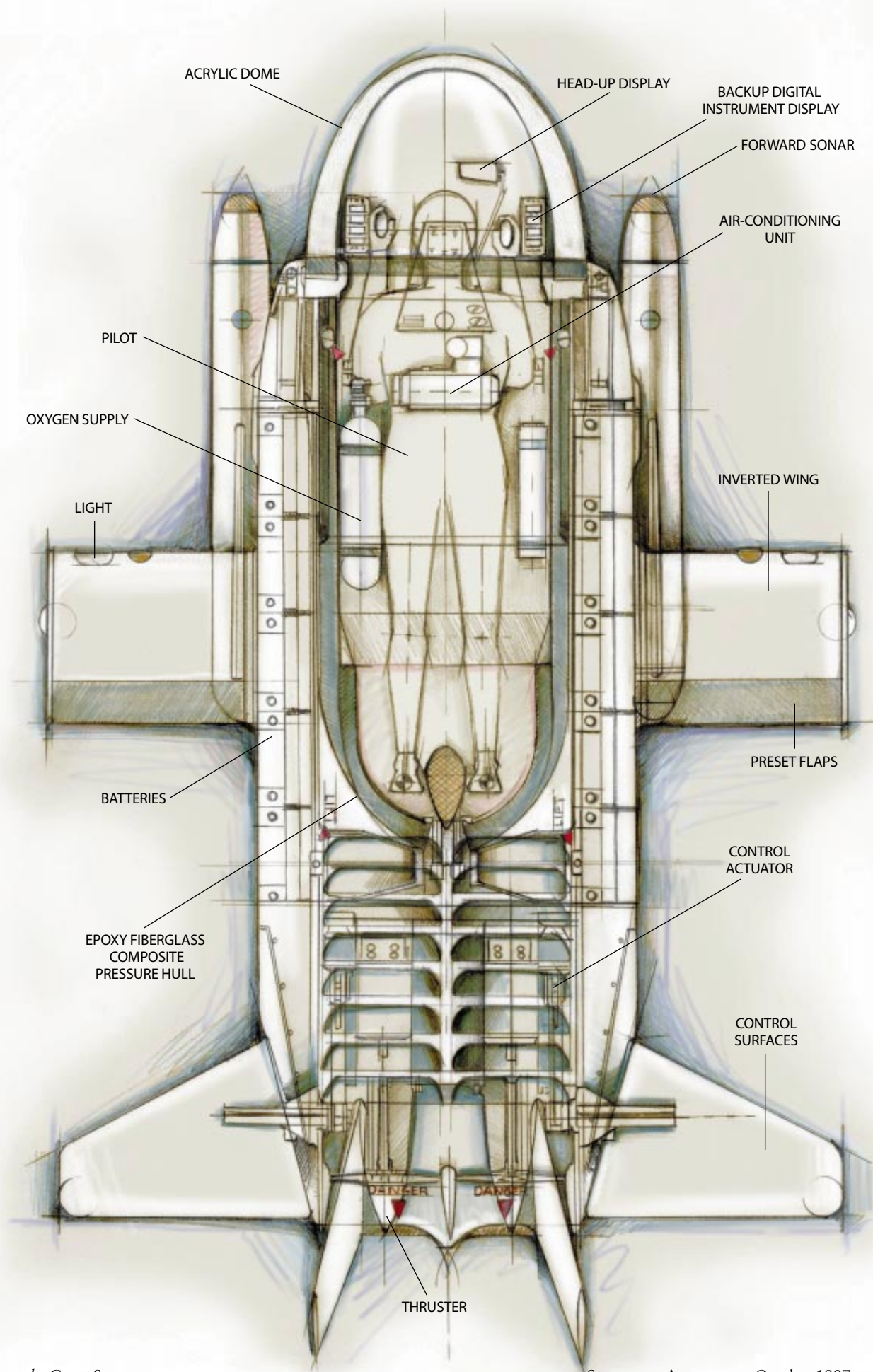
Deep Flight

The Deep Flight program is an attempt to move beyond existing constraints and develop a generation of lightweight, manned submersibles that could operate economically and independently from the research and commercial fleets. My colleagues and I built *Deep Flight I* purely as an experimental craft to evaluate the engineering concepts for improved hydrodynamic efficiency and to test other key systems that would shrink the bulk of the submersible down to a microsub. Guided by what we have learned from this prototype, we are now working on the design for *Deep Flight II*, a more practical craft that could conceivably take a person to the bottom of the deepest sea.

The dramatic difference in configuration between a Deep Flight craft and a conventional sub may look like a wild design inspiration, but it is really an engineering response to the need to move faster underwater. Because hydrody-



EXPERIMENTAL SUBMERSIBLE (*side view this page; top view facing page*) incorporates innovations as well as adaptations from other vehicles to reduce size and increase maneuverability. Designed by the author and dubbed *Deep Flight I*, the 3.5-meter-long sub can carry its pilot to depths of 1,000 meters (3,280 feet).





HAWKES OCEAN TECHNOLOGIES



AMOS NACHOUM PHOTOGRAPHY



AMOS NACHOUM PHOTOGRAPHY



AMOS NACHOUM PHOTOGRAPHY



CHUCK DAVIS

dynamic forces increase with the square, and power with the cube, of speed, raising a submersible's velocity from one knot to five knots increases the power requirements about 100-fold. With no near-term hope of improving battery power at those multiples, the speed gain has to come by reducing drag.

In many ways, *Deep Flight I*, which can descend to a maximum depth of 1,000 meters, is more like a deep-diving suit than a small submarine. It has a small frontal area and the inevitable streamlined form and wings (or fins) common to aircraft, birds, dolphins and whales. In designing *Deep Flight I*, I discarded the essential characteristic of bathyscaphes and submersibles, a variable buoyancy system that enables these vehicles to change their apparent weight in water and either sink to the bottom or float back to the surface. *Deep Flight I* remains slightly buoyant at all times. Once it is moving, it uses its wings (which are configured upside down with respect to those on an aircraft) to pull it down to depth. Movable aft wings provide control. The pilot "flies" the craft underwater with subtle movements of small joysticks that operate a fly-by-wire system. The clear acrylic nose cone extends past the pilot's peripheral vision—and with no structure visible, the effect when flying underwater is magical.

A major dilemma was how to shrink the instrumentation down to the very limited space available. But this difficulty was solved by the microprocessor revolution: just a few switches provide essentially unlimited control, and the mass of display instrumentation is reduced to various pages on a video screen. Old technologies die hard, however, and *Deep Flight I*, as a bridge to the past, has two digital display sets that provide basic instrumentation.

Perhaps the most obvious difference in our design is the one-person crew position: prone, face forward and strapped into a form-fitting body pan

instead of sitting upright in a chair. At first, this configuration appeared awkward, but it seemed acceptable for an experimental vehicle. What was not immediately obvious was that this position is precisely the one assumed by humans and other mammals swimming underwater. So it should be no surprise that, underwater, the posture feels natural and comfortable. After being in *Deep Flight I*, I immediately scrapped all my designs for other craft, which had been premised on a desire to offer a "proper" sitting position. As a concession to mental comfort in *Deep Flight II*, the basic body position was raised to a relaxing 30 degrees, which keeps the attitude within a non-alarming range for first-time crew during descents and ascents.

Even Deeper Flight

Will the hydrodynamic design that works so well in the relatively shallow waters traveled by *Deep Flight I* also work at the far greater depths for which *Deep Flight II* is intended? Happily, yes. The additional depth and pressure do not compound the drag, because water is, for all practical purposes, incompressible, so the dynamic behavior of a vehicle at full ocean depth is virtually the same as near the surface.

The mounting pressure on the hull, however, must be withstood. Switching to a standard steel hull would sacrifice the gains from lighter weight. Yet there is reason to think that a practical alternative for full-depth pressure hulls is possible. The U.S. Navy has successfully tested buoyant deep-submergence hulls made from new ceramic materials, which have the margin of safety required for human occupancy. Data about these ceramics have recently been declassified, which may facilitate their commercial development. Increased pressure would also require *Deep Flight II* to have more conventional view ports in place of the acrylic observation dome.

Constructed from the new ceramic materials and incorporating the innovations proved with *Deep Flight I*, *Deep Flight II* should be a remarkable and useful craft. It would operate within an envelope of what I like to call "intelligent autonomy" beyond unmanned probes: AUVs cannot match the intelligence (because they are robots), and ROVs cannot match the

DEEP FLIGHT I in its travel frame (1) is loaded onto a truck for transport to the launch site. The author enters the sub (2), which is hoisted by a crane (3) that places the craft in the water (4) for a test dive. Once underwater (5), the craft's acrylic dome offers unhampered views (6). Floating at the surface (7), the sub waits to be lifted from the water by a crane (8).

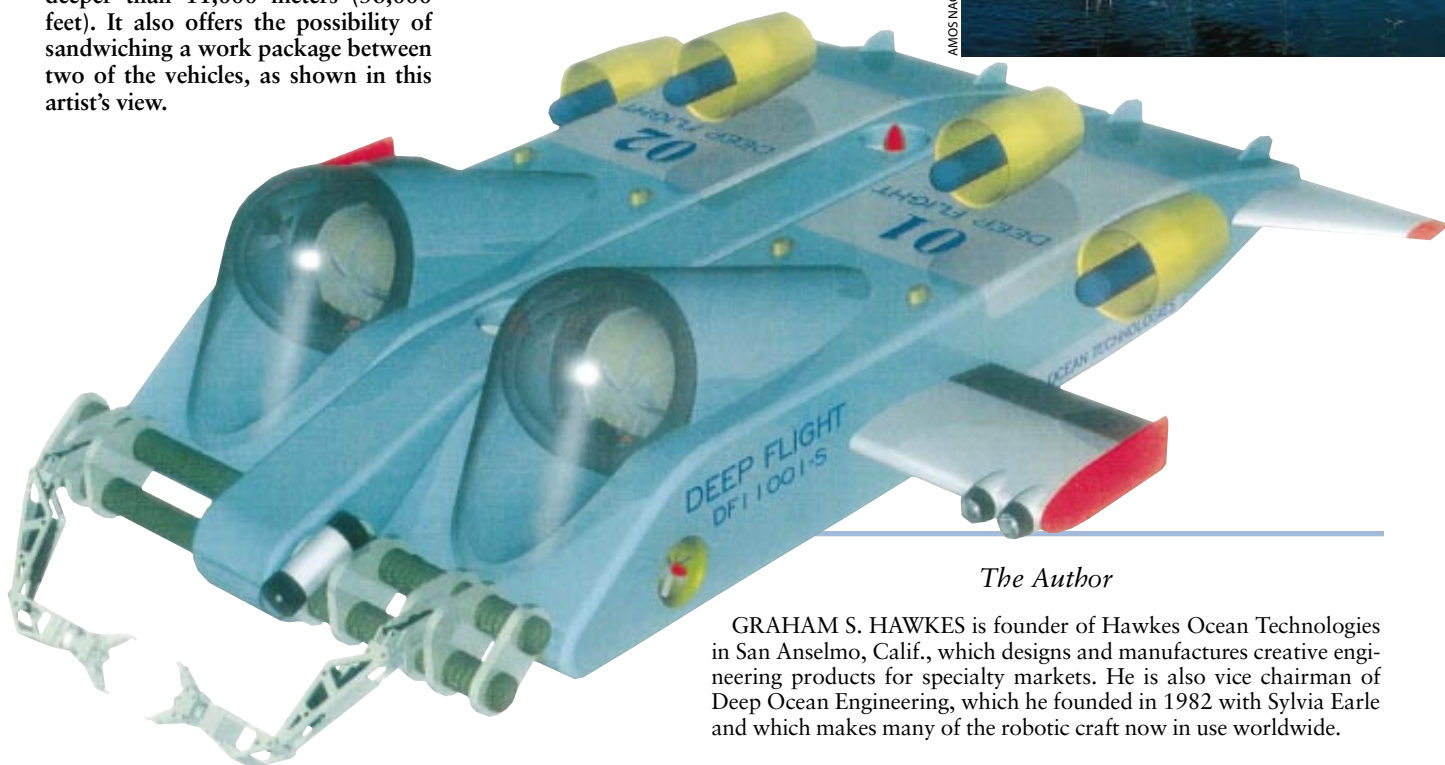
autonomy (because they are essentially tethered in place). Again, though, the goal is not to create a submersible that replaces AUVs and ROVs but to produce one that augments them by keeping open the chance for human travel to the depths.

The dearth of funding for new submersibles dictates that *Deep Flight II* will need to be an all-purpose craft, able to respond to widely variable requirements. It should be agile and stealthy for midwater biologists, for example, or be able to function as a heavy payload bulldozer for geologists. In response to various and conflicting needs, the basic design is modular and adaptable; the separate units could be quickly and easily reconfigured on a ship's deck into any of three basic modes: One version would be a single-person craft, with minimal weight and drag, for underwater surveys and exploration. A second configuration would provide a pair of single-person units joined together, for exploration

missions in which one of the passengers is a passive observer, not a pilot. Finally, *Deep Flight II* could be arranged as a more heavy-duty work vehicle consisting of two units with a work package sandwiched between them. These work packages would be equipped with vertical thrusters, enabling the craft to hover like a hummingbird around its work site.

Like all craft, submersibles will always have their limits, and for the foreseeable future they will remain strictly mission-oriented, such as for scientific exploration. The only practical transportation application would seem to be for shuttling personnel between sub-sea installations and the surface and for short-range submarine tours in shallower waters. Yet these craft do allow an appreciation of the oceans that robotic vehicles cannot. They keep the option open for those of us who want to wander and work freely and in person within the earth's largest habitat.

NEXT GENERATION of the Deep Flight microsubmersible promises to go to much greater depths, perhaps deeper than 11,000 meters (36,000 feet). It also offers the possibility of sandwiching a work package between two of the vehicles, as shown in this artist's view.



The Author

GRAHAM S. HAWKES is founder of Hawkes Ocean Technologies in San Anselmo, Calif., which designs and manufactures creative engineering products for specialty markets. He is also vice chairman of Deep Ocean Engineering, which he founded in 1982 with Sylvia Earle and which makes many of the robotic craft now in use worldwide.



In 1956 the visionary architect Frank Lloyd Wright sketched out plans for the Illinois, a mile-high skyscraper that would accommodate 100,000 people, parking for 15,000 cars and enough office space to house the entire state government. Sheathed in aluminum and stainless steel, Wright's 528-story edifice could have been built, he believed, with available technology. But a significant barrier stood in the way: the required elevators would have taken up too much space. This problem continues to limit the height of the world's tallest buildings to less than 1,800 feet (about 549 meters).

Conventional high-rise elevators are based on the age-old technology of winches, pulleys (called sheaves in elevator technology) and counterweights, which balance the weight of the cab and so reduce the amount of energy required to raise a load. For more than 140 years, elevators working on the same principles have proved remarkably efficient at raising and lowering people and freight. But certain drawbacks become particularly acute when a building reaches into the clouds.

In any size building, each elevator requires not only the cab, cables and counterweights but also a hoistway—the elevator shaft—and a hoisting machine, which is usually housed in a room of its own. In addition, tall buildings require a multitude of elevators. To provide efficient service for a densely populated building with a lot of traffic between floors, architects usually plan for roughly two elevators for every three floors; a 90-floor building might need about 60 elevators. The more lofty the construction, then, the more expansive and expensive the real estate taken up by elevators becomes.

One solution to this difficulty has been to create sky lobbies, such as those in the World Trade Center towers in New York City, where passengers traveling to the upper strata of the buildings must change elevators to reach higher floors. That strategy conserves space in the ground-floor lobby (by cutting down on the number of elevators traveling to the entry level). Still, the ultimate solution would be a ropeless elevator with a self-propelling drive system. "Getting rid of the counterweight and the ropes is the quantum leap everyone is looking for," says James W. Fortune of Lerch, Bates and Associates, an elevator systems consultancy.

Once unshackled from the heavy steel cables and attached counterweights that raise and lower the cab, the elevator of the future can become a far more versatile and efficient means of in-building transportation. More than one elevator will be able to travel in the same shaftway, thus saving on valuable building space. And elevators will no longer be limited to vertical movement. Companies worldwide are working to create elevators that go sideways as well as up and down. Joseph Bitar of Otis Elevator envisions elevators that can glide sideways to allow passing by another car or to transport passengers from one building to another. The mundane elevator car will thereby be transformed into a module that could take a passenger from a remote parking lot to the

Elevators on the Move

Elevator technology is taking off in new directions, including sideways

by Miriam Jacob

60th floor in about 90 seconds. This prospect is especially enticing for several huge building plans on the drawing board in Asia; those projects include many structures of different sizes and will require both vertical and horizontal conveyances.

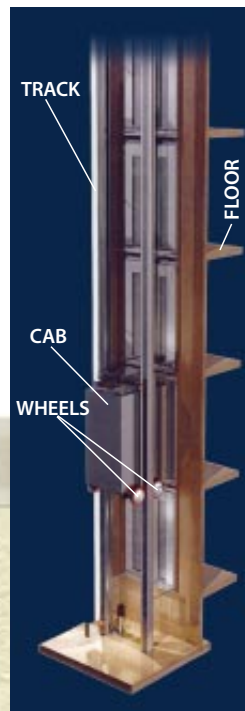
Alternatives to the conventional elevator have developed slowly, mainly because the established technology for moving the cars is reliable and remarkably energy-efficient. Among the alternative drive systems that

could ultimately lead to ropeless systems suitable for large buildings are linear induction motors. Such units are created from electric rotary motors by changing their geometric configuration—in effect, opening up and flattening the motor. Instead of producing torque that spins a rotor to pull a cable, the linear motor produces a longitudinal force that drives the elevator cab by magnetic repulsion.

A linear motor developed by Otis is in use now, although not yet in futuristic elevators. It is powering elevators in about 1,000 low-rise buildings in Japan. Its current incarnation still requires ropes, but the motor is incorporated into the elevator shaft, effectively becoming part of the counterweight that

is suspended at the end of the hoisting cables. This drive system does away with the machine room—a big advantage in Japan, where space is at a premium. Ultimately, though, flat linear motors might be installed on elevators themselves. This adjustment could eliminate cables and counterweights, potentially liberating the elevator from its solely vertical existence and even allowing multiple cars to operate in the same shaftway.

An innovative self-propelled elevator that is already run without ropes owes its strong traction to gearing principles first developed for a vehicle that explored the surface of the moon in the 1960s. Engineered by Schindler Elevator Corporation in Switzerland and dubbed the SchindlerMobile, this system is powered by a small motor attached to the bottom of the elevator cab. Melding automotive and elevator technology, the motor drives two wheels with special



SELF-PROPELLED ELEVATOR, the SchindlerMobile, rides up and down the tracks of two high-strength aluminum columns. A motor mounted under the cab moves the wheels.

polyurethane tires up the tracks of two high-strength aluminum columns. Each driven wheel is paired with an idler wheel not pushed by the motor. Constant pressure from a spring presses the wheel pairs against the track, providing the traction that keeps the vehicle suspended and moves it up and down the columns. The lightweight, aluminum-framed vehicle does use small counterweights to increase its efficiency; they are housed inside the aluminum columns. And the current version of the SchindlerMobile is small, fairly slow and designed for use in low-rise buildings. But Schindler officials say the concept may be applied to taller buildings and higher-capacity elevators in the future.

The high-capacity ropeless elevator awaits the creation of a drive system that can match the speed, comfort and energy efficiency of its conventional predecessor. Meanwhile various elevator manufacturers are exploring ways to liberate the elevator from its purely vertical shaftway with current technology. Otis, for instance, is developing an advanced elevator, named the Odyssey System, that combines horizontal and vertical people-moving technologies in a way that up to now has been exploited only in sophisticated amusement rides.

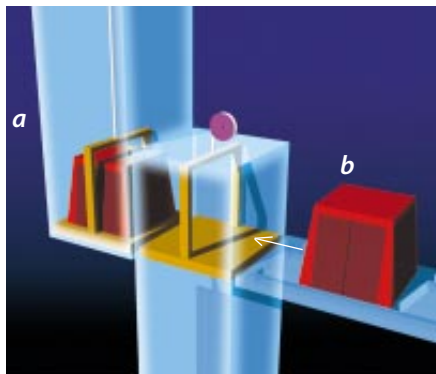
The key to the Odyssey System is eliminating direct contact between the elevator cab and the cables. This is accomplished by building a structure known as a frame, or platform, that fits around the cab. The platform holding the cab is tethered to the cables in the elevator shaft, and the cab itself becomes mo-

bile. This arrangement allows a platform to carry different cabs at different times—possibly including a passenger cab that has traveled horizontally from a different point in the building.

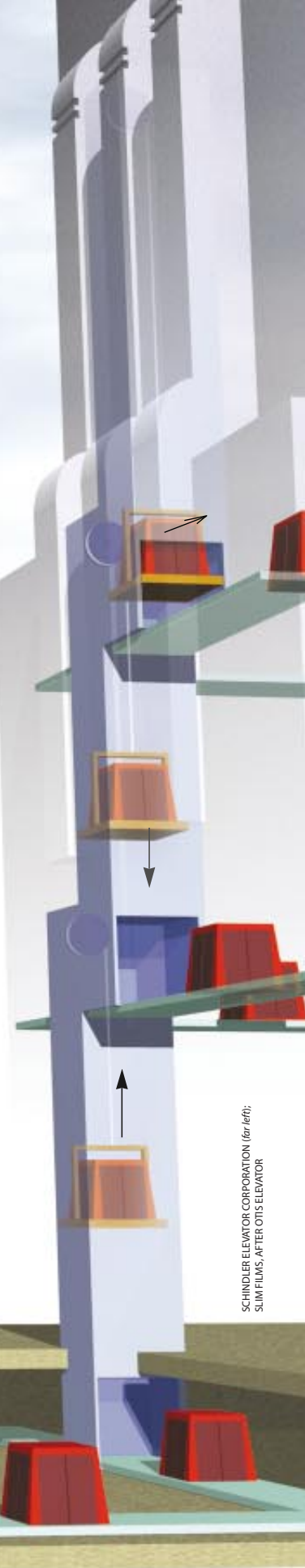
To transfer an elevator cab from a flatbed supporting horizontal movement to a platform for vertical movement, or vice versa, Otis has devised a special linear induction motor. The motor is in two sections such that one piece is on the platform or on the flatbed housing for horizontal modes, and the second is on the cab. The transfer of the cab is completed when the two pieces of the motor become locked together.

The Odyssey System and related projects foreshadow a far more complex role for the familiar elevator. No longer limited to vertical transportation, the tetherless elevator car of the future will travel far more widely and be considerably smarter. Without such a radically new vision for in-building transportation, the megabuilding projects being eyed for the future will be unlikely to materialize. SA

MIRIAM LACOB is a writer based in New York City.



ELEVATORS MOVE LEFT AND RIGHT as well as up and down in a scheme designed by Otis Elevator. Cabs moving sideways along a flatbed become vertically mobile by sliding into a frame, or "platform," which can be hoisted upward. Later (*detail above*), the cab may jog sideways into an adjacent platform to continue its journey in a different shaft (*a*). Or it may shift into a horizontal passageway (*b*) or into a loading area.



SCHINDLER ELEVATOR CORPORATION (far left); SLIM FILMS, AFTER OTIS ELEVATOR

THE AMATEUR SCIENTIST

by Shawn Carlson

Recording the Sounds of Life

Some months ago, in the predawn hours of a chilly winter morning, my life changed forever. At roughly 3 A.M., I was awakened from a deep sleep when a pair of strong hands seized me and started shaking my torso violently. I awoke in horror, expecting to confront a crazed killer who might be wielding a club or an ax. Instead I saw my wife standing over me, her face a hideous blend of terror and joy. I stared up at her, wondering for a moment whether she had gone suddenly insane. “Honey,” she blurted, frantically waving a short plastic wand in front of my confused eyes, “I’m pregnant!”

It was, of course, the best news of our lives. And, with both of us being scientists, it wasn’t long after sharing a celebratory cup of hot chocolate that we started thinking about the opportunities for discovery that Michelle’s pregnancy afforded. I wanted to find out how much entropy our growing baby will add to the universe by the time it is born. I even devised a simple experiment; it called for soaking Michelle in an insulated vat of tepid water repeatedly throughout her pregnancy, each time measuring how long it took the heat from her body to warm it. Alas, Michelle has made it quite clear to me that this fundamental number will have to remain a mystery (and a great amateur science project for

another expectant couple to take on).

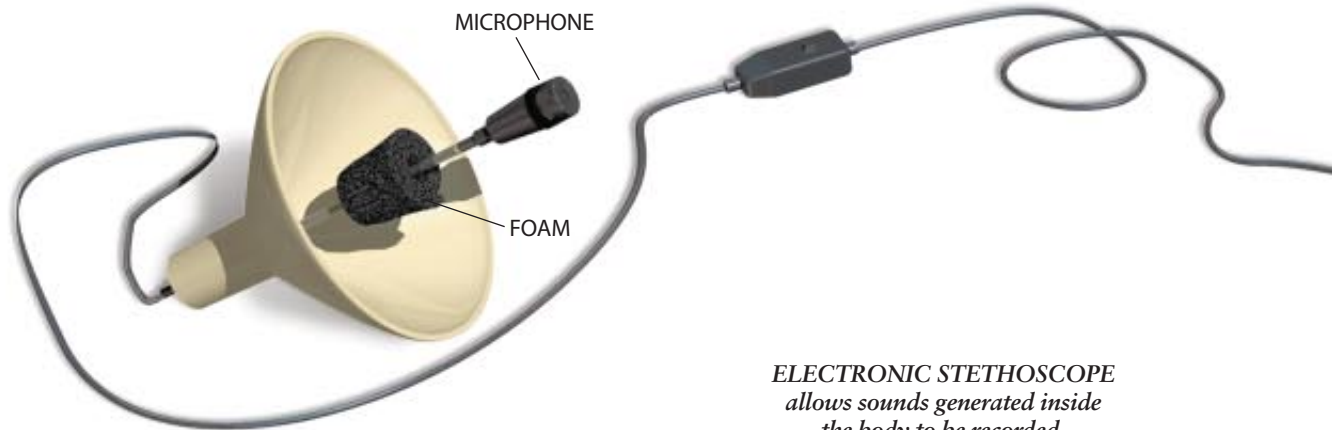
But my wife was happy to try other experiments with less demanding protocols. Just a few weeks ago I pieced together an electronic stethoscope that can detect all kinds of sounds produced inside the human body. Initially, I hoped to record the baby’s heartbeats and movements. But the apparatus worked better than I had anticipated. Michelle and I have now also recorded a myriad of sounds produced by our own hearts, lungs and gastrointestinal tracts—and a few truly odd gurgles that don’t seem to be emanating from any particular organ. You, too, may want to listen in on your own body or to record the internal sounds of your favorite cat or dog. Amateur scientists with an interest in marine creatures may want to adapt the apparatus for use underwater as a hydrophone. In each case, an ordinary tape recorder will serve to archive the sounds.

The device combines 19th-century and modern technologies. For more than 150 years, doctors have relied on a trick of geometry, not electronic circuitry, to amplify sounds within the body. Nearly all the sound energy that enters a stethoscope’s relatively large chest piece is channeled into a hollow tube, then directed through a headset and finally deposited onto the doctor’s eardrums. Focusing sound in this way increases the

intensity of the sound by roughly the same ratio as the area of the chest piece to the inside opening of the tube.

You can use the same technique to make a serviceable stethoscope quite easily using any small funnel. Just place the mouth of the funnel against a friend’s chest. When you press your ear over the neck of the funnel (something you should do very gently to avoid injuring your eardrum), you will hear your friend’s heart and lungs quite clearly. A small length of Tygon tubing, with one end pushed over the neck of the funnel and the other end delicately tucked just inside your ear canal, will let you listen to the noises created within your own body.

Once amplified by a funnel, these sounds can be captured with a small microphone and processed electronically. Today quality microphone transducers cost next to nothing, and sophisticated systems can be built from scratch for less than \$20. (You’ll need an electret-type condenser element, a low-noise op-amp, some shielded speaker wire and a few garden-variety resistors and capacitors. Die-hard do-it-yourselfers can consult the Society for Amateur Scientists’s Web page for details.) But it’s much easier simply to purchase a small lavalier microphone (also called a “tie clip” microphone). The Optimus omnidirectional microphone (Radio Shack catalogue number 33-3013), for example, costs less than \$25, and it outperformed all but my most extravagant cre-



ELECTRONIC STETHOSCOPE
*allows sounds generated inside
the body to be recorded.*

ations. The unit comes ready to be plugged into any conventional tape recorder that is compatible with an $\frac{1}{8}$ -inch plug. If your tape player has a different size jack, you'll also need to buy an adapter.

You can secure the microphone inside the funnel using a scrap of foam rubber or similar material. Begin by threading the microphone through the neck of the funnel, as depicted in the illustration on the opposite page. I used a short, thin strip of antistatic foam (Radio Shack catalogue number 276-2400) to hold it in place. Wrap the strip around the microphone a few times. Then secure this package into the neck of the funnel so that the microphone rests just at the apex of the cone.

To test the system, press the open end of the funnel firmly against your chest and switch on your recorder. If you're using a stereo tape deck, make sure to turn the volume on your stereo amplifier all the way down. If you try to record heart sounds and to listen to them through speakers at the same time, your entire neighborhood could be treated to an earsplitting sample of audio feedback. I first attempted to avoid this problem by listening through headphones. Big mistake. The microphone was so sensitive it picked up the faint sounds leaking from the headphones and fed them back into the amplifier. The result was an extremely painful high-frequency blast emitted directly into my ears, which abruptly ended the experiment.

Many tape recorders have a volume indicator that shows the amplitude of the signal being recorded. If yours does not have this feature, you'll have to set the overall amplification by adjusting the volume control, recording for a few seconds and then listening to how the newly recorded track sounds. Repeat the procedure until the signal is as loud as possible without being distorted.

Unfortunately, your microphone will not just register the sought-after body sounds; it will also pick up whatever extraneous noises may be polluting your local acoustic environment. To forestall problems, use a simple RC circuit as a low-pass filter to block any signal with a frequency greater than about 800 cycles per second [see diagram below]. The filter does not affect most body sounds, but it will help screen out chirping birds, honking horns and young neighbors' stereos. Although a single resistor-capacitor pair works, chaining two such pairs together, as shown, eliminates more noise, especially near the cutoff frequency. Make sure you use a shielded cable and that the electronics are housed in an all-metal and well-grounded project box.

My choice of 800 cycles per second for the cutoff frequency is completely arbitrary. Depending on your application, you may get better results by using a different limit. The cutoff frequency (in cycles per second) for any simple RC filter will just be the reciprocal of the product of the resistance (in ohms), the

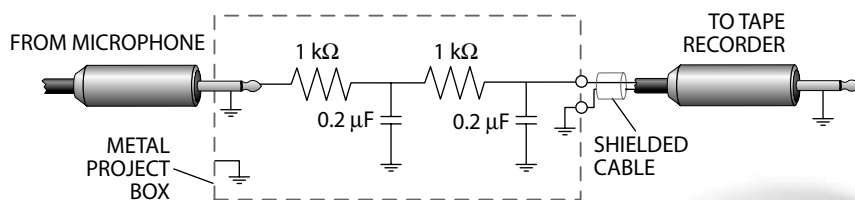
capacitance (in farads) and 2π (6.28).

Michelle and I have been regularly recording our baby's heartbeats since early July. We have noticed the sound getting steadily louder over the past few months and expect soon to observe the slowing of heart rate that happens as a baby develops. (In the fourth month of pregnancy, a baby's heart will beat typically at about 160 beats per minute; by the ninth month it normally drops below 140 beats per minute.)

Taking time out to listen in on our baby's internal doings has given us a special closeness with our unborn child. The emotion is not unlike that experienced by many scientists, professional and amateur alike, who develop a profound sense of intimacy with whatever they are examining. Often it is this personal connection that pushes such scientists onward in the pursuit of understanding. The motivation to undertake a program of careful observation is, of course, particularly strong when the subject is your own baby girl. (Body sounds don't reveal gender, but a routine ultrasound did.) Baby Katherine Joanne is due November 4.

SA

For information about this project or other activities for amateur scientists, write the Society for Amateur Scientists, 4735 Clairemont Square, Suite 179, San Diego, CA 92117. You can also visit the society's World Wide Web site at www.thesphere.com/SAS/, call (619) 239-8807 or leave a message at (800) 873-8767.



MATHEMATICAL RECREATIONS

by Ian Stewart

Two-Way Jigsaw Puzzles

Early in their careers the puzzlists Sam Loyd and Henry Ernest Dudeney—one American, one English—collaborated on a regular puzzle column for the magazine *Titbits*. Loyd wrote the puzzles, and Dudeney, under the pseudonym “Sphinx,” provided the commentary and awarded prizes. Collaboration soon turned into rivalry, however, and the two men went their separate ways. In so doing, they created a puzzle industry on both sides of the Atlantic, by formulating tantalizing mathematical questions within simple stories.

A typical example of their work is Loyd’s Sedan Chair Puzzle. The mathematical problem is to cut the sedan shape into as few pieces as possible and re-

assemble them to form a square; Loyd embeds it in a tale in which a young lady’s sedan chair folds up cunningly to protect its occupant from the rain.

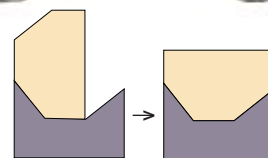
Puzzles of this kind are known as dissections. A wonderfully entertaining book on this time-honored theme is soon to be published; it is *Dissections: Plane and Fancy*, by Greg N. Frederickson (Cambridge University Press).

The basic mathematical concept that underlies all dissection puzzles is area. When a shape is cut up and the pieces are rearranged, the total area does not change. Some very deep mathematics indeed lies behind this apparently self-evident statement. Oddly, it is false in three dimensions if the “pieces” are allowed to be sufficiently complicated. In the celebrated Banach-Tarski Paradox, a solid sphere is “dissected” into six pieces, which can be reassembled to form two solid spheres, each the same size as the original. Polish mathematician Stefan Banach and his Polish-American collaborator Alfred Tarski proved this weird theorem in 1924. It is a logically valid result, but it seems so bizarre that “paradox” has stuck.

How can the volume double, just by

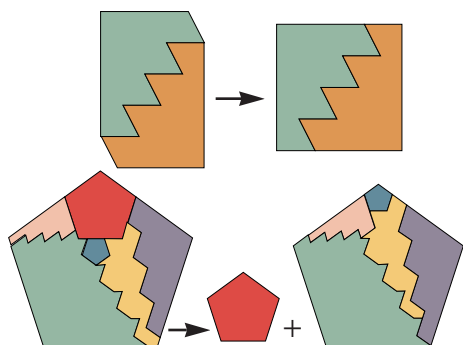
rearranging the pieces? The trick is to employ pieces that are so strange that they do not possess a well-defined volume, more like infinitely complex spherical dust clouds than single, connected objects. The ideas are summarized in my book *From Here to Infinity* (Oxford University Press, 1996) and described in all their gory glory in Stan Wagon’s *The Banach-Tarski Paradox* (Cambridge University Press, 1985).

There is, of course, no practical way to realize this dissection with a physical object—so traders in precious metals can breathe a sigh of relief—but it does demonstrate how subtle the concept of volume is. Curiously, no “paradoxical”

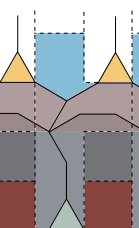
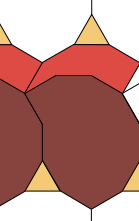
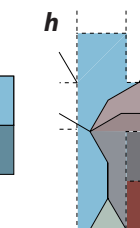
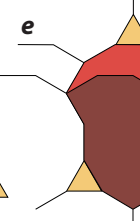
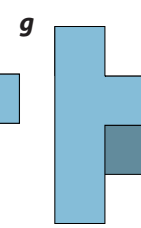
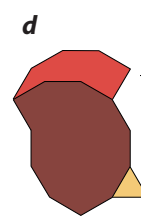
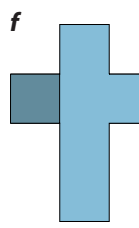
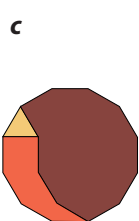
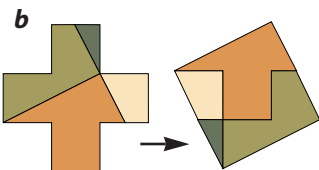
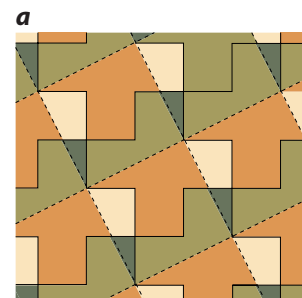


SEDAN CHAIR PUZZLE
and its solution.

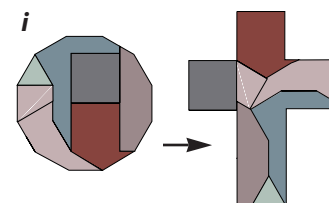
LAUREL ROGERS



STAIRCASE TRICK
is useful for finding dissections.



TILINGS
of the plane can allow dissections to be “read off,” as with the Greek cross to a square (a and b). More elaborately, a dodecagon (c) can be cut and rearranged into a shape (d) that tiles a plane (e), as can the cross (f and g). Overlaying the two tilings (h) leads to a dissection of the dodecagon into a cross (i).

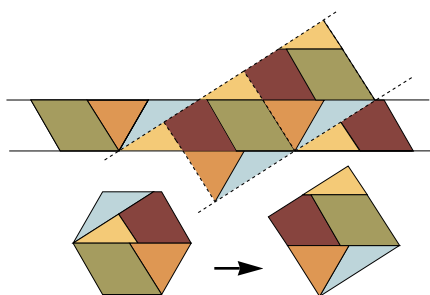


dissections, in which areas change, are possible in plane geometry, no matter how complex the pieces may be—as Tarski proved in 1925. (They are possible, however, on the surface of a sphere.)

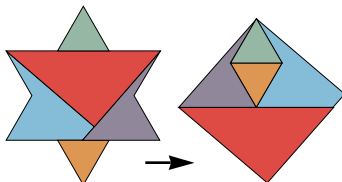
When the pieces into which the object is cut are nice enough to have well-defined areas or volumes, there are no intuition-bending constructions. Indeed, in 1833 P. Gerwein, a lieutenant in the Prussian army, answered a basic question about dissections raised by Hungarian mathematician Farkas Wolfgang Bolyai. Gerwein proved that given any two plane polygons of equal area, there is a finite set of identical polygonal pieces that can be assembled to form either shape. This result is called the Bolyai-Gerwein Theorem (although it seems to have first been proved by one William Wallace in 1807).

The Bolyai-Gerwein Theorem does not generalize to three dimensions. In 1900 the great German mathematician David Hilbert asked whether any two polyhedrons of equal volume were “equivalent by dissection”—that is, could be assembled from the same set of polyhedral pieces. One year later German-American topologist Max Dehn proved the startling result that a cube and a regular tetrahedron of equal volume are not equivalent by dissection.

The real fun comes in finding neat examples of shapes that are equivalent by dissection. You can make some proto-



INFINITE STRIPS
tiled by different shapes can lead to dissections, here of a hexagon to a square (above). A Star of David can be similarly dissected to a square (below).



LAUREL ROGERS

Mathematical Recreations

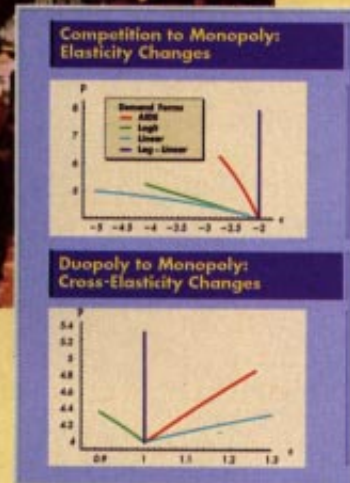
MATHEMATICA[®]

EMPOWERMENT

Trustworthy Models for Antitrust Investigation



Suppose major cosmetics companies X and Y have proposed a merger. How do antitrust investigators at the U.S. Department of Justice determine beforehand whether the merger would result in a sharp rise in the price of cosmetics?



Using software developed by Vanderbilt Professors Philip Crooke, Luke Froeb, and Steven Tschantz (shown here) and Justice Department Research Director Gregory Werden, antitrust enforcement agencies increasingly use *Mathematica* to simulate the effects of such horizontal mergers. The computational approach involves two steps: estimating a structural demand model and then computing or simulating the postmerger equilibrium to predict postmerger prices and quantities. *Mathematica* has been used to simulate mergers in products as diverse as breakfast cereal, bread, ski resorts, frozen fish, and tissue paper.

The Vanderbilt professors have also put the simulations onto a *Mathematica*-driven web site for use as a teaching tool for attorneys and MBA students. “For example,” says Dr. Froeb, “the cosmetics merger case turned on a low cross-elasticity of demand for the two companies’ mascara products. Using the web, I can tell the story of the case and then show how this affected the predicted anticompetitive price rise.”

“*Mathematica* makes it easy to prototype and simulate models,” says Froeb. “I think that *Mathematica* has the potential to radically change the field of economics, much as calculus did seventy years ago.”

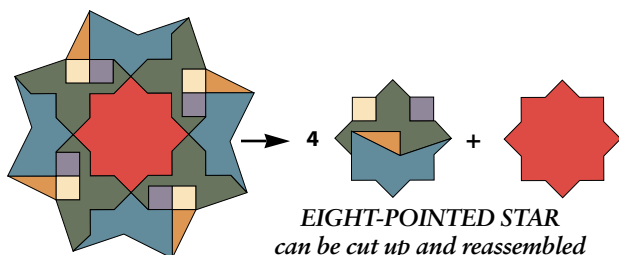
No matter what they’re using it for, researchers, scientists, engineers, hobbyists, and others all agree on one thing: *Mathematica* makes their lives easier and helps them accomplish more. *Mathematica* 3.0 introduces major new concepts in computation and presentation, with unprecedented ease of use and a revolutionary symbolic document interface. *Mathematica* 3.0 is available for Microsoft Windows, Macintosh, and over twenty Unix and other platforms. Purchase or upgrade on the web at <http://www.wolfram.com/orders>.

For more information on how you can use *Mathematica* for work or play, visit <http://www.wolfram.com/look/scg> or call toll free 1-800-416-8064.



WOLFRAM
RESEARCH

Wolfram Research, Inc.: <http://www.wolfram.com>; info@wolfram.com; +1-217-398-0700. Wolfram Research Europe Ltd.: <http://www.wolfram.co.uk>; info@wolfram.co.uk; +44-(0)1993-883400. Wolfram Research Asia Ltd.: <http://www.wolfram.co.jp>; info@wolfram.co.jp; +81-(0)3-5276-0506. © 1997 Wolfram Research, Inc. Mathematica is a registered trademark of Wolfram Research, Inc. and is not associated with Mathematica Policy Research, Inc. or MathTech, Inc.



EIGHT-POINTED STAR
can be cut up and reassembled
into five such stars.

ress by inspired trial and error, but only if you have a vivid spatial imagination. One of the virtues of *Dissections* is that it explains many of the general principles involved in finding them.

One principle is cutting a shape along a “staircase,” which can then be moved one step along to create a different shape. David Collison, a dissection enthusiast who was born in England and worked as a computer programmer and consultant in the U.S., devised elaborate dissections based on this principle. His pentagon dissection [see top left il-

lustration on page 140] offers dissection proof of the well-known “Pythagorean” fact $5^2 + 12^2 = 13^2$. Another general method is the so-called tessellation principle. Many shapes of interest can be embedded in tessellations—tiling patterns that cover the plane. If two different tessellations, each formed from tiles of the same area, are superposed, it often becomes possible to “read off” a dissection from one shape to the other. For instance, a Greek cross can thereby be dissected into a square.

A more elaborate use of the same basic idea—the dissection of a dodecagon to a Latin cross—is due Harry Lindgren, the author of *Geometric Dissections* (Van Nostrand, 1964). The first step, and the hardest, is to cut the dodecagon

FEEDBACK

Cows in the Maze,” the December 1996 column, was clearly a source of considerable amusement. It was about a self-referential maze that you had to solve by moving one or the other of two pencils, starting on boxes (1,7) and ending with at least one pointing to GOAL. The crux was the notorious “rule 60” in box 60, which suspended instructions in red text until further notice.

Before discussing the mailbag, I regret to report that the book *Supermazes*, which I promised, will not appear. The author, Robert Abbott, notified me of this fact by e-mail, but I was slow to read it and failed to correct the column in time.

Readers’ feedback caused me many moments of panic, as they wrote in with claims of shorter answers, better answers, errors in my answer and the like. Several claimed that I was wrong to state that any solution must involve getting to boxes (50,50) with rule 60 not in force. When I checked these attempted solutions, however, I found that in every case there was an error. I won’t mention names—I’m sure you’ll know who you are—but you deserve an explanation. I’ll use the notation of the column, with red font indicating the pencil to be moved and an asterisk showing that rule 60 is in force.

One reader’s attempt began (1,7) (1,26) (1,55) (1,15) (9,15) (35,15) (35,40).... But when one moves from position (35,15), the instruction in box 15 reads “Is the other pencil in a box whose number is evenly divisible by five?” The answer here is “yes,” and that leads us to (35,5), not (35,40).

Another error occurred in a claimed solution: (1,7) (2,7) (15,7) (15,26) (15,61) (40,61) (60,61) (25,61)* (7,61)* (26,61)* (61,61)* (1,61)* (2,61)* (15,61)* (40,61)* (65,61)* (75,1) (50,1) GOAL. Its author observed that “rule 60 is not canceled” as a result of the maneuver from (65,61)* to (75,1). There is a misunderstanding here. If you have arrived at (65,61)* and you choose to move pencil 61, then because rule 60 is in force, you must ignore the red text—which is all of box 61. This leads you to (65,1) because rule 60 tells you to use the “yes” exit for the chosen pencil. In order to get to (75,1), you must obey the red text in box 61, which tells you to move both pencils—but you can’t do this when rule 60 is in force.

Like all the other instructions, the rule in box 60 goes into effect only when you choose to move the pencil pointing to that box, not as soon as one of the pencils arrives at box 60. My solution involves a move from (26,60) to (55,60) with rule 60 not in force. Because I move pencil 26, the rule in box 60 is not activated at that time.

—I.S.

SCIENTIFIC AMERICAN

IN OTHER LANGUAGES

LE SCIENZE

LE SCIENZE

Piazza della Repubblica, 8, 20121 Milano, Italy

サイエンス

NIKKEI SCIENCE, INC.

1-9-5 Otemachi Chiyoda-ku, Tokyo 100-66, Japan

CIENCIA

PRENSA CIENTIFICA, S.A.

Muntaner, 339 pral. 1.a, 08021 Barcelona, Spain

SCIENCE

POUR LA SCIENCE, ÉDITIONS BELIN
8, Rue Férou, 75006 Paris, France

Spektrum

SPEKTRUM DER WISSENSCHAFT

Verlagsgesellschaft mbH

Vangerowstraße 20, 69115 Heidelberg, Germany

科学

KE XUE-Chongqing Branch

Institute of Scientific & Technical Information of China
P.O. Box 2104, Chongqing, Sichuan, Peoples Republic of China

العلوم

MAJALLAT AL-OLOOM

Kuwait Foundation for the Advancement of Sciences
P.O. Box 20856, Safat, 13069, Kuwait

SWIAT NAUKI

SWIAT NAUKI, Proszynski i Ska S.A.

ul. Garazowa 7, 02-651 Warszawa, Poland

CALLING ALL TEACHERS

Tired of the same old lesson plan?

SCIENTIFIC AMERICAN

TEACHER'S KIT

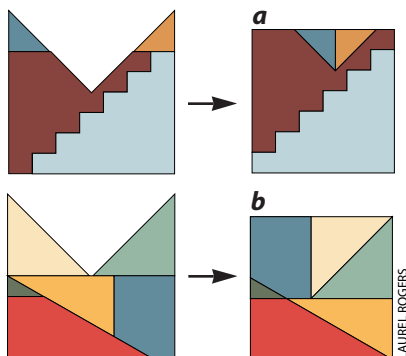
1-800-377-9414

OR FAX

1-212-355-0408

CALL for a FREE package of instructional materials and teaching ideas.

SCIENTIFIC AMERICAN will help your students get their HANDS and MINDS on SCIENCE.



DECEPTIVE DISSECTION
by Sam Loyd turned a miter into an apparent square (a). Unfortunately, the "square" is actually a rectangle; the correct dissection is b.

and rearrange it into a shape that tiles the plane. The Latin cross (of equal area to the dodecagon) can also be made to tile the plane. Comparing the two tilings leads to the dissection desired.

A third method is the strip principle. The two shapes are cut into pieces that between them can tile an infinitely long strip. If the strips are then overlapped, they determine a dissection. A dissection of a hexagon to a square, derived by Paul Busschop, comes from the strip method. (Busschop was a Belgian who wrote a book on peg solitaire, published posthumously in 1879 by his brother.) Harry Bradley, an American engineer who was an instructor at the Massachusetts Institute of Technology in 1897, discovered how to dissect a Star of David to a square by the same method.

Dissection puzzles are not confined to changing just one shape to another. Often an entire set of shapes must be cut up and reassembled, or vice versa.

Dissections are so compelling that occasionally the eye can mislead the head. In 1901 Loyd made a memorable error—described by Frederickson as "perhaps his biggest goof"—when he claimed to dissect a miter (a square with one quarter removed) into a square. Unfortunately, the apparent "square" is actually a rectangle whose sides are in the proportion 49:48. Ironically, Loyd called this the Smart Alec Puzzle. His great rival, Dudeney, pointed out the error in 1911 and gave the correct dissection shown. So if you want to look for your own dissections, take the advice given by many a parent to their offspring: "Have fun—but be careful."

SA



It may be small. But the Bose® Acoustic Wave® music system is definitely an overachiever. The unit features a compact disc player, an AM/FM radio, a handy remote control, and our patented acoustic waveguide speaker technology. And it produces a rich, natural sound quality comparable to audio systems costing thousands of dollars. We know that's hard to believe. So we're ready to prove it. Call or write now for our complimentary guide to this award-winning system. Because, like the system itself, it's available directly from Bose.

Call today. 1-800-898-BOSE, ext. A2266.

Mr./Mrs./Ms. _____ Daytime Telephone _____ Evening Telephone _____
Name (Please Print)
Address _____
City _____ State _____ Zip _____
Or mail to Bose Corporation, Dept. CDD-A2266, The Mountain, Framingham, MA 01701-9168.

BOSE®
Better sound through research®

SINGLES IN SCIENCE ...

are meeting via *Science Connection*, a North America-wide social network for single science professionals and others who enjoy science or nature.

Register on-line (or learn more) at:
<http://www.sciconnect.com/>

To receive our information package by mail, please contact us by phone or e-mail.



Science Connection

(800) 667-5179

sciconnect@compuserve.com

EVOLUTION EVOLUTION

SOUND BITES and INSIGHTS

0-9649118-1-7 © 1997 by Howard A. Royle
160pp 4 1/4" x 5 1/2" List @ \$6.50

A Popular, Easy-to-Understand,
One-of-a-Kind Book

Includes 400 classical and original
quotations and thoughts about Evolution.

Everything you ever wanted to know about
Evolution at your fingertips. Evolution from
A to Z—nothing else even comes close to
the concentrated knowledge in this book.

**Authorative, compact, easy to
read and comprehend.**

—a thorough grounding in evolutionary thought and practice.

A gem of a book for the scientifically curious
reader with a short attention span and an
energetic mind. A great source book for
students, teachers, writers and lecturers. From
the author of *SCIENCE - SCIENCE*.

To order send \$8.00 Check or M.O.
(Includes P & H.)

Canyon Publishing Company
P.O. Box 751 • Canyonville, OR 97417

REVIEWS AND COMMENTARIES

QUEER SCIENCE INDEED

Review by Tom Boellstorff and Lawrence Cohen

Queer Science: The Use and Abuse of Research into Homosexuality

BY SIMON LEVAY

MIT Press, Cambridge, Mass., 1996 (\$25)

A Natural History of Homosexuality

BY FRANCIS MARK MONDIMORE

Johns Hopkins University Press, Baltimore, Md., 1996 (\$15.95)

Despite findings that remain inconclusive, contested and sometimes irreproducible, the biology of human sexual variation has gradually become respectable. In these two ambitious books, which epitomize current thinking, neuroanatomist Simon LeVay and physician Francis Mark Mondimore review several decades of biological research on homosexuality and place it in its historical and political contexts. Both authors devote great space to summarizing the work of anthropologists and historians; their willingness as biologists to examine this research demonstrates an interest in dialogue with social scientists that is all too rare and seldom reciprocated.

Yet these efforts at integration meet with limited success. The two authors refer to the ethnographic and historical record without engaging it. Once he has mentioned the myriad ways that humans have coupled, for instance, LeVay glosses over them: "To try to take all this potential [cultural] diversity into account right at the beginning would be a recipe for paralysis."

This attitude is symptomatic of a more general malaise in the academy. The "two cultures" gap between science and the humanities has led researchers to believe that sexuality must be either biological or cultural. These books demonstrate how even the best of authors get drawn into turf wars. Instead of evaluating good science, they become less than critical boosters of any putatively scientific effort to de-

feat the "higher superstition" of those who suggest that culture plays a fundamental role in sexual orientation. In spite of their detailed historical and cultural discussions, both LeVay and Mondimore frequently reduce a broad spectrum of anthropological and historical work to a caricature of social constructivism—the notion that sexuality is entirely the product of cultural decisions.

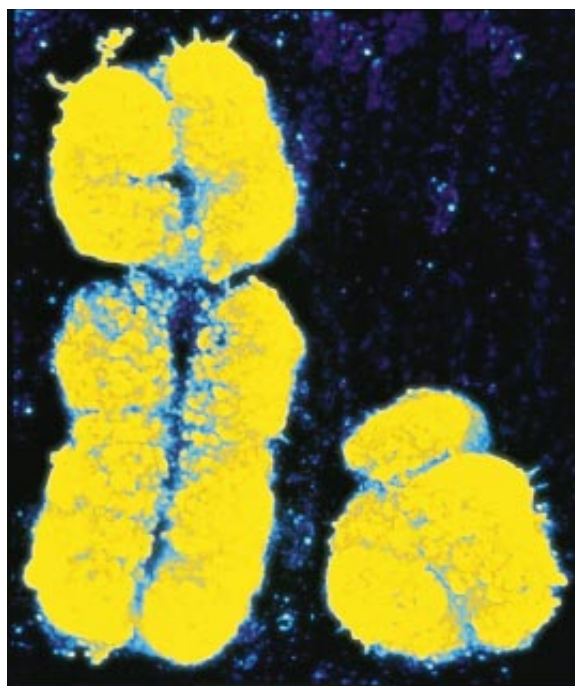
At times, the authors offer a more synthetic point of view—that humans have biological potentials that take a completed shape under specific personal and social circumstances, leading to great diversity in sexual experience and iden-

tity. Mondimore in particular focuses on the ways in which the brain is structured through interaction with the environment, an emphasis that puts him somewhat at odds with LeVay's enthusiasm for hormonal explanations. But LeVay's and Mondimore's "biological potentials" look suspiciously like the folk categories familiar to most Americans: straight and gay. Indeed, much of the biological research in this field seems to be based on the popular stereotype that gay men are feminized males and lesbians masculinized females.

The ethnographic record, however, documents behaviors that do not accord with these ostensible biological categories. For example, Tomas Almaguer of the University of Michigan has looked at the common Latin American split between people who penetrate ("active" men) and people who get penetrated (women and "passive" men). Men who penetrate other men are not marked as homosexual. Jonathan Marks, a biological anthropologist at Yale University, has said that the standards of validity for scientific research should be higher when the results appear to reinforce

folk assumptions than when they contradict the popular wisdom. Thus, we should be wary when LeVay claims, without any supporting evidence, that "we have a core identity, of which our sexual orientation is an important element, that radiates outward and richly informs and energizes our lives." Does this intuition point to a biologically determined universal or to a very American notion of individualism, choice and self-integrity? LeVay's formulation begs the question of why sexuality of any form is seen as a central aspect of one's identity. Understanding why people have felt this way in some cultural contexts but not in others would be of great use in formulating a biology of sexual orientation.

The worldwide intensification of media and consumer culture, along with the international battle against AIDS, has led to the



HUMAN SEX CHROMOSOMES
*show the familiar X (left) and Y (right) shapes.
But human sexuality is not so clearly defined.*

BIOPHOTO ASSOC. Photo Researchers, Inc.

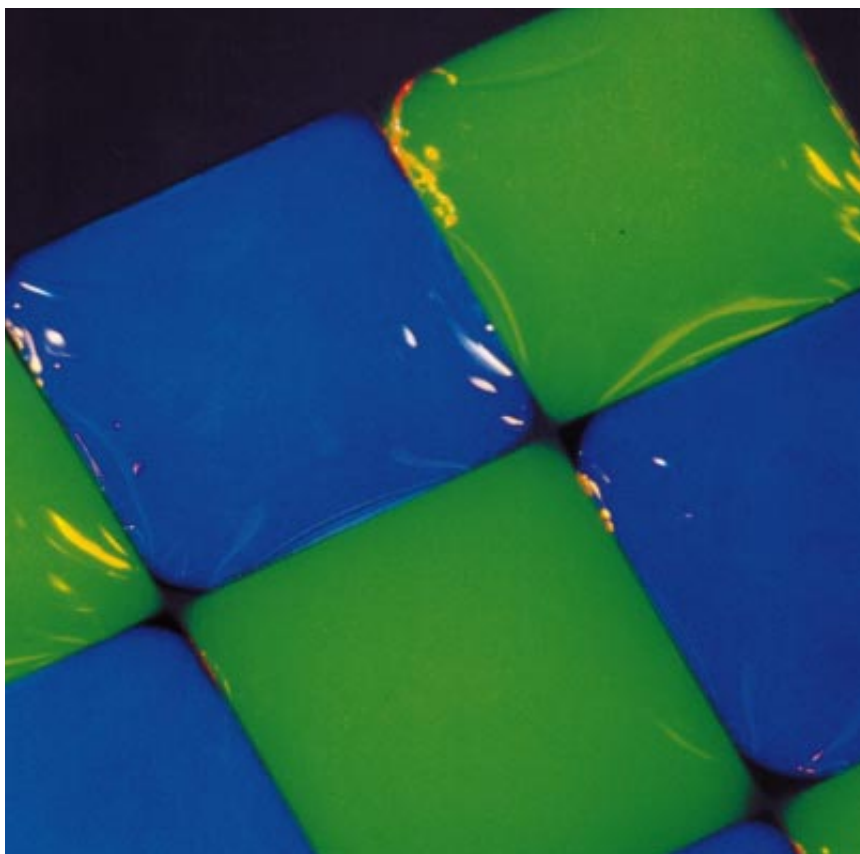
THE ILLUSTRATED PAGE

On the Surface of Things: Images of the Extraordinary in Science

BY FELICE FRANKEL AND
GEORGE M. WHITESIDES
Chronicle Books, San Francisco,
1997 (\$22.95)

Materials science bears an unfortunate reputation for dullness, dealing as it does with the stuff of everyday life. A ramble through the pages of this poetic volume, however, exposes the field's underlying luster. A checkerboard of water droplets (*shown at right*), a shard of broken glass or a swatch of plastic fabric reveal themselves as things of colorful, otherworldly beauty.

The words are no less remarkable, balancing weighty concepts from the laboratory with a literate tone as light and elegant as a spider's web. A wonderful achievement indeed. —Corey S. Powell



FELICE FRANKEL

This is what happens



Ovonic NiMH battery-powered car. **Available now.**
High energy for extended driving range.



The world won't be
the same with us.

global diffusion of the terms “lesbian,” “gay” and “straight.” But we should not infer from the spread of these categories that they reveal an underlying similarity. One of us (Boellstorff) has shown how Indonesians who use terms such as “gay” do not simply import them from the West but transform them in unexpected ways. Societal differences in both the classification and experience of sexuality continue to exist and may actually be increasing.

Such complications should be helpful to biologists who are interested in moving beyond general folk models of sexuality toward a science more grounded in the biological correlates of behavioral plasticity and its limits. They force us to be ever more precise in what we mean by a biological potential and to examine sources of bias in our methods and conclusions.

Mondimore and LeVay recognize the usefulness of social analysis, but both reach an impasse when they try to square a constructionist view of sexuality with notions of choice. They seem to believe that because most people do not choose their sexual orientation, it must be purely a matter of biology, immediate and precultural. LeVay claims that “sexual attraction is an aspect of consciousness; it is directly experienced, like hunger, thirst, seeing the color red, taking fright, loving one’s mother, and countless other aspects of our mental life.” This statement is indicative of the primary conceptual weakness of his reasoning. The empirical data simply do not support the notion that any aspect of consciousness—sexual attraction or love or even color vision—is directly experienced before the influence of culture. All humans grow up in a specific

culture, and *Homo sapiens sapiens* has evolved to be shaped by culture not only on the level of ideas and symbols but on neurological levels as well.

As a result, any analysis that omits specific cultural contexts in favor of biological “foundations” is inadequate. Only by beginning with all the data—cultural, genetic and neurological—can scientists undertake a rigorous study of the wide range of human sexuality. Furthermore, instead of assuming that similarities in such diversity are determining or “underlying” factors, we believe researchers should consider the diversity of causes that can lead to the most similar of results. Such a paradigm shift would have profound implications for the ways in which people think about and conduct research on sexuality.

Even where sameness appears across societies (the existence of plural grammatical forms, nurturing of children, same-sex behavior), it manifests itself only in the cultural context. It is to these contexts that we must turn to understand human actions and artifacts that seem to be cross-cultural. To start with the postulate that diversity underlies sameness recognizes the remarkably underdetermined nature of human genetics (as Mondimore acknowledges). Such a scientific methodology does not exclude biology, nor does it relegate biology to a subordinate position. Instead it recognizes that human beings have evolved biologically so that they need to live in specific cultures.

In a sense, we would argue that *H. sapiens sapiens* has evolved so that we have no sexual orientation without reference to a particular, historically located culture, just as we have no way of speaking without using a particular, historically located language. Such a framework does not discredit the work of researchers such as LeVay and Mondimore; rather it builds on their own stated desire to integrate biology and culture by offering a way to unite these apparently disparate domains of human existence in a way that subordinates neither.

TOM BOELLSTORFF is a graduate student in the department of anthropology at Stanford University.
LAWRENCE COHEN is assistant professor of anthropology at the University of California, Berkeley.

ON THE SCREEN

Fast, Cheap and Out of Control

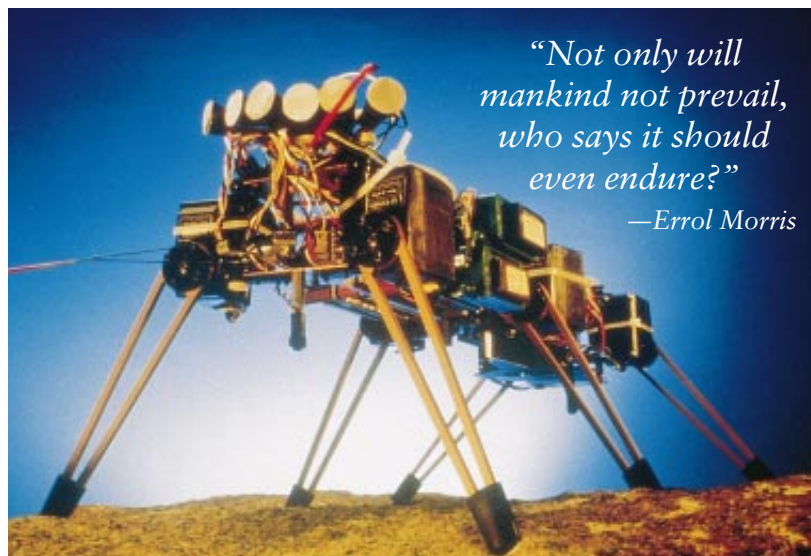
DIRECTED BY ERROL MORRIS

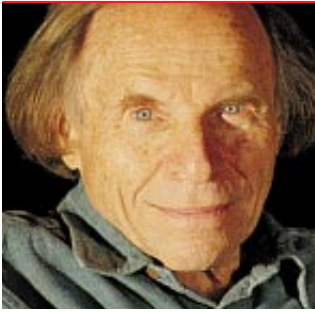
Sony Pictures Classics, 1997

Premieres in New York City and Los Angeles on October 3;
 opens nationwide throughout October

Ray Mendez is obsessed with naked mole rats. Rodney Brooks makes walking robots, like the one shown below. George Mendonça expresses himself through topiary. Dave Hoover thinks like a lion in order to tame a lion. Errol Morris (who directed the movie adaptation of *A Brief History of Time*) profiles these four men, overlapping and intercutting between them. The snippets seem superficial until larger patterns begin to emerge: an obsession with taming and controlling nature, a shared fascination with the animal world, a sense of our vulnerability to replacement. The gorgeous cinematography, by Robert Richardson, effectively complements the movie’s introspective mood.

—Corey S. Powell





WONDERS

by Philip Morrison

Air-Cooled

Fifty years ago power overloads and outages were the menace of winter. The power system controllers of any northern city grew anxious whenever late in the afternoon there came darkened clouds, cold winds, hints of sleet and snow. Most people hastened home at dusk in fear of worsening weather. Meanwhile the lights all glowed to brighten offices, shops and kitchens. Overcrowded packs of trolleys and commuter rail drew maximum power, while often frozen streams constricted hydro-power. The stage was set for peak power—and sometimes failure.

No longer. Now it is in May and not November that the newspapers in Bos-

ton and New York City carry warnings of brownouts ahead. The elements have not changed, but we power users have. It is on humid, oppressive afternoons that power dispatchers frown country-wide and check their reserves time and again. For now strange louvered boxes, air conditioners large and small, stud the walls and windows and crowd the roofs of every U.S. city. More than two thirds of our households and offices provide fan-blown air, refrigerator cooled and dried. As some hot day darkens into an all but intolerable evening, all those thermo-

stats call for the maximum. Come the wrong day, half a kilowatt or more by each compressor for many hours implies an increase large enough to boost the peak demand of 13 million New Englanders by three million kilowatts, or about 20 percent.

*No head wind, no cooling,
no sustained performance.*

We Americans cool ourselves with air-conditioning not only at home and at

work but on the wing, the road or the rail, even in the cab of the tractor pulling the plow. We go without power-assisted summer cooling mainly when we move by our own muscles.

Insight into air-cooling emerges from

when a bunch of socially responsible dreamers



Photovoltaic roof panels. **Available now.**
Flexible, thin-film solar cells for clean electric power from rooftops.



The world won't be
the same with us.

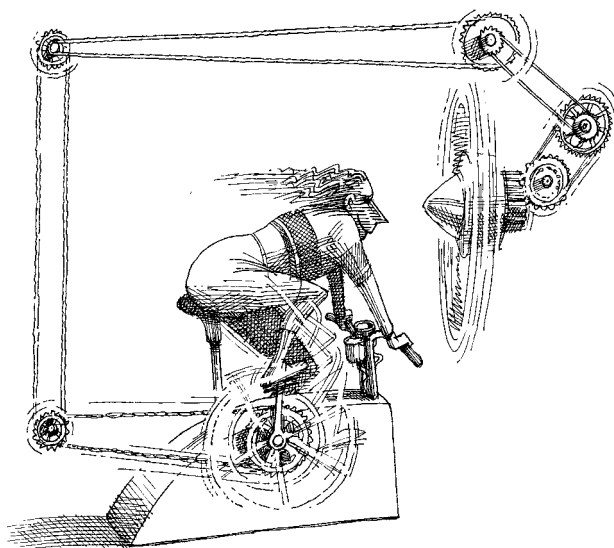
another phenomenon of summer, the Tour de France. The cyclists of that annual Gallic extravaganza are outstanding athletes. In 1997, 139 competitors finished the 2,500-mile road race, pros whose arcane team tactics are legendary. Even the slowest of them put on a sustained physical performance hard to match by any other well-documented feat in work or in sport. A marathon runner is justly hailed for physical effort over two long hours, but every day each bicyclist delivers muscular output more than double that of any marathoner. The contrast is a measure of the aptness of bicycles on a surfaced road; their flowing, rotary smoothness does not inflict the oscillatory muscular starts and stops and the pounding blows of foot against road that the runner must endure. Not many runners are eager to repeat a marathon two days running, but the cyclists come back for 22 days in all, with just one traditional day off.

There is one major hidden difference. An anecdote may make it plain. The Belgian racer Eddy Merckx was overall winner in the postwar Tour de France not once but five times, a paradigm of endurance. Curious physiologists besought Iron Eddy to show what he could do on an instrumented stationary bike. Known to be masterful over a full six hours up and down the most daunting Alpine passes, he began with élan, only to quit, tired, drenched in sweat and bitterly disappointed after an hour. What the lab bike did not provide was the 25-mile-an-hour head wind Eddy took with him everywhere he pedaled. Air drag was his chief adversary, for he had to push aside masses of air to reach speed, but his main ally, too, because that draft alone cooled him as energy income required. No head wind, no cooling, no sustained performance.

By the 1990s all this had been tightly documented. The racers neither gain nor lose weight over their weeks of work, although they eagerly scarf the calorie equivalent of eight square meals a day, where three such meals are enough for sedentary American males. The cyclists must balance that input by a matching

output of energy in work and as heat. They produce on average a kilowatt of muscular power, supplied by adenosine triphosphate (ATP) and its precursor, the biochemical fuel of aerobic dark-meat muscle fibers. A quarter of that goes into actual mechanical work against air drag and other losses; the other 75 percent is lost as heat during each six-hour performance. Getting rid of one kilowatt-hour of heat energy requires vaporization of about 1.6 quarts of water. The lion's share of the racer's heat loss is from this single evaporative process, driven to extreme values by the strong headwind his road speed generates.

The necessary coolant is his supply of



drinking water, as urgent as his food and even heavier. He is brought six or eight quarts while on the daily roll and plenty to drink at other times. Yet the power of evaporative cooling grows only slowly as speed increases. Even a minor draft is valuable. Many a concertgoer recalls genuine relief in a stifling hall from the slow "wind" made by one program sheet waved as a fan. But the power lost to air drag grows far more rapidly with air speed than the cooling does. Fluid-flow theory explains both results.

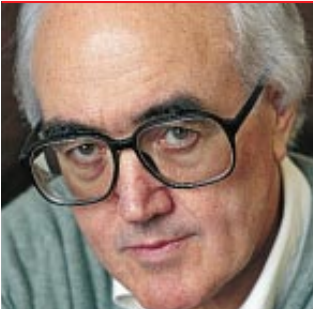
There is new interest today in the limits of muscular performance. A top Tour de France cyclist reaches a sustained metabolic rate about five times the minimal metabolism that accompanies bed rest. Among 50 vertebrate species measured, only a handful can do better than *Homo sapiens rotulans*, and none can beat the seven times minimum

racked up by a lab mouse mother nursing her 14 pups. Without laying claim to a final answer, one is satisfied to note that the old saw holds, if a little modified: it's not *only* the heat, it's the humidity. Indoor weather, like outdoor weather, is not a matter of air alone but of its watery content. On a bad day all those air conditioners dribble fluid water; you may say that the buildings sweat instead of their occupants, typically at around an ounce a minute from any ordinary room.

An old experience offers a helpful glimpse. Some may recall the practice of placing a canvas water bottle into the slipstream outside a moving auto. It produced deliciously cooled drinking water, well below air temperature, even as you sped down some broiling, bone-dry desert highway. Water turned to vapor takes out more energy by an order of magnitude than the heat it held while liquid. Consequently, much heat flows to generate the vapor, and the remaining fluid is cooled quite efficiently.

No sweat, we say, whenever a task is performed with little effort. The phrase needs a gloss. Overwork does induce sweating, and yet the visible loss of water drops from the skin is not the sign of cooling. Rather it warns of inadequate cooling. Losing bodily water as liquid means expending some without loading it with the latent heat that water vapor carries off. Called "insensible perspiration," it best fights human thermal stress.

Imagine a superhero cyclist who one day can reach a 10-fold increment of metabolic rate over the resting value. The increase in air drag would limit his speed increase to 50 percent, but his heat generation would grow threefold. Over an unchanged route he might start out at 40 miles an hour, but he could not even finish the first 100-mile stage. His core temperature would rise by about 10 degrees Fahrenheit during each hour at speed. Soon he would have to withdraw, as sagacious Eddy Merckx did long ago, or fall a fevered victim to distortion of almost all enzymatic reaction rates, in biochemical analogy to the wonderful one-hoss shay. Uncool. SA



CONNECTIONS

by James Burke

The Buck Stops Here

I was going Dutch the other day at lunch and handing over my share in U.S. dollar bills when I remembered that it was a 16th-century polymath from Holland who started all that decimal money stuff. Simon Stevin was his name. Unsung hero would be nearer the mark. His motto could have been that of any of the Scientific Revolution biggies who later eclipsed him: "There's always a rational explanation for what looks like magic."

Stevin was the engineering genius who first popularized an alternative to the mind-wrenching medieval practice of calculating everything in fractions. (For a flavor of the torture involved, try

$\frac{3}{144} \times \frac{2}{322} - \frac{1}{85} = ?$) He turned such gibberish into the decimals with which scientists, and innumerate like me, could more easily work. He even gave the treatise he wrote on the subject a user-friendly title: *The Tenth*.

In 1585 Stevin became quartermaster and commissioner of public works for Prince Maurice of Nassau, ruler of northern Holland at the time. Maurice was a bit of a military history freak and thought there was much to be learned from the disciplined way the Romans had fought. So he built an army that

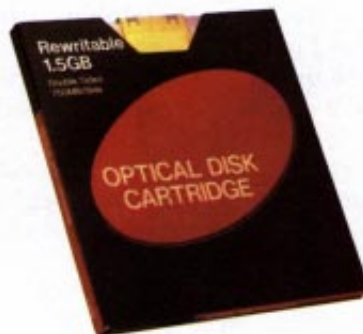
was, so to speak, all dressed up with nowhere to go. That is to say, he introduced new techniques (the use of cartridges, military drill, an instruction manual for firing muskets by numbers), any one of which could have given him

victory at any major battle he ever found himself involved in. But he never really got much more than a

large skirmish. This left all the glory to his contemporary, Scandinavian King Gustavus Adolphus, who a few years later gave his soldiers even more winning ways. One of his key improve-

*Europe's richest publisher
had a profitable sideline
in French underwear.*

decide they really don't care



Rewritable optical memory disks. **Available now.**
ECD's phase-change technology has 650 Mbyte removable capacity.
And soon, rewritable DVD.



The world won't be
the same with us.

ments was to put the musketeers in three rows, so that while the front row was firing, the second and third rows were reloading, getting ready to step forward and pull triggers. Thus, there was a constant stream of hot lead heading toward the opposition. As a result of this trick, Gustavus won every battle he fought (except the last, at which he was killed) and made Sweden a world power for all of 15 minutes.

The next Swedish ruler, Gustavus's daughter King Christina (that's not a typo; only the wife of a Swedish monarch was called "Queen"), cut her hair short, wore men's clothing, turned Catholic and abdicated in 1654 after only 10 years in the post. Whereupon she high-tailed it to Italy and (it is suspected) a long-term affair with one of the cardinals in a troubleshooting team of elite Curia bureaucrats known as the Flying Squad. Although some Swedes may disagree, there are many who believe we have much to thank Christina for. Such as Rome's first opera house and the successful careers of Bernini, Scarlatti and Corelli (all of whom she protected from various forms of Roman backstabbing). Less admirable, perhaps, was what she did to René Descartes. While she was still queen (sorry: king), she invited the eminent French philosopher to come and be thinker-in-residence and then obliged him to give her philosophy lessons at five in the morning. In Stockholm. In January. Surprise, surprise, he caught pneumonia and died.

Fortunately, however, not before he had produced (among other works) the *Discourse on Method*, which taught us all to think straight, as well as presented a fundamentally new view of the cosmos and a major piece on how the human body functioned like a machine. This last described how the brain worked by a system of tubes and valves controlling the distribution of a fluid "animal spirit" that made the different parts of the anatomy move. Which moved one Tom Willis, a well-to-do physician at the University of Oxford, to spend years preparing an opus on matters cerebral entitled *The Anatomy of the Brain*. Definitive for the

next 150 years, it contained the first reference to the autonomic responsibilities of the cerebellum. Willis's book became an international best-seller because it was the first to feature copious illustrations so detailed and accurate even the *Lancet* would have accepted them for its pages.

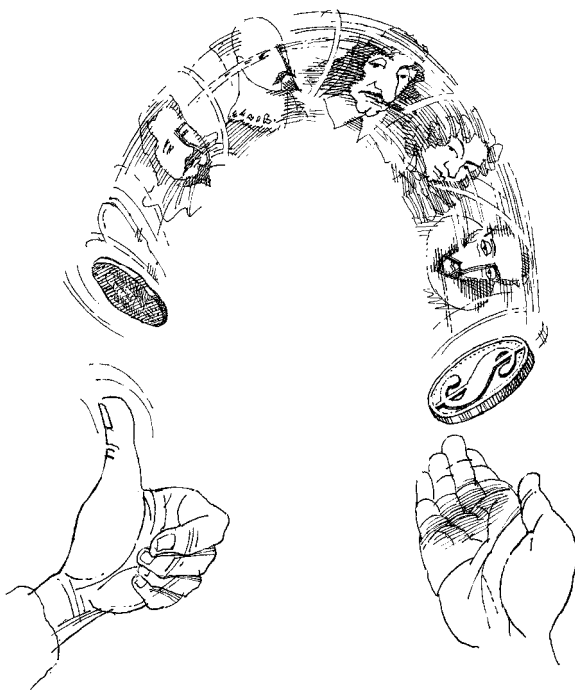
They were done by England's greatest draughtsman, architect, linguist, mathematician, weather forecaster, astronomer and general big-head—well, at the age of 36 you'd have to have a lot of chutzpah to apply for and get the contract to build St. Paul's Cathedral after the Great Fire of London. Christopher Wren was also a canny businessman (you're not surprised), being one of the first to get into the new stocks-and-shares game and becoming a director of that license-to-print-money known as the Hudson's Bay Company. Considering how much profit this organization made (and still does) for its backers, it seems a pity the bay's eponymous discoverer did all the hard work for so little reward. Like many early European navigators, Henry Hudson spent a lot of his life going nowhere. In particular, in 1609 he was commissioned by the Dutch East India Company to find an Arctic Northwest Passage over the top of Greenland and America, so the Dutch could get the spices, porcelain and tea they wanted out of the Far East without

being hassled by the Spanish and Portuguese, who had the southern routes sewn up. Well, after sailing up and down the Greenland coastline for months and repeatedly bumping into Spitsbergen or pack ice, Henry went back to Antwerp to give a piece of his mind to the so-called cartographer who had put him on this wild goose chase.

This unfortunate was theologian-turned-mapmaker Pieter Platvoet, who had learned all he knew (Hudson: "Not enough!") from the truly great cartographer Gerardus Mercator, whose fame spread rapidly when his work was printed by Christophe Plantin, Europe's richest publisher, who had a profitable sideline in French underwear.

Plantin made a fortune, after the Council of Trent decided to standardize worship, by churning out more than 40,000 identical liturgical texts for Philip II of Spain. Or he would have if Philip could have paid his bills on time. Philip's little financial problem was his father, who had got the job of Holy Roman Emperor by greasing the right palms with money he had borrowed from (and left Philip to pay back to) a German banker named Anton Fugger, the Rothschild of the day. By this time the Fugger family had been in the money game for over 100 years, and pretty much every crowned head in Europe was in hock to them. The monarchs all used mercenary armies and never had the ready cash to pay them off, so the Fuggers would helpfully provide the wherewithal, in return for property, or tax breaks, or concessions.

One such recoupment package included a mining franchise in the mountains of Bohemia: literally a hole in the ground that produced so much silver it became the official source of coinage for the entire Holy Roman Empire. The mine was in a valley called Joachimsthal, and the coins came to have the same name: "Joachimstalers." Over time this became shortened to "Talers." And over *more* time, the American pronunciation of the word became the name for the currency I was doling out at the end of that meal at the beginning of this column.



CLEAN GENES

Review by Philip Yam

Gattaca

WRITTEN AND DIRECTED

BY ANDREW NICCOL

Columbia Pictures, 1997

Gattaca joins *The Fly* and *The Island of Dr. Moreau* as the latest piece of science fiction to explore the dangers of tampering with the genetic code. Scheduled to open this month, the movie tells of a not too distant future in which almost everyone has been screened in a petri dish for the most desirable traits possible from their parents' DNA. The unfortunate few who are conceived naturally, the so-called In-Valids, are doomed to discrimination and an underclass life.

One In-Valid determined to beat the odds is Vincent Freeman (Ethan Hawke),

WITH A BORROWED GENOME
Vincent (Ethan Hawke) perpetuates a false identity in front of Irene (Uma Thurman) in Gattaca.

whose genes indicate he is likely to develop a fatal heart condition in his thirties. His dream to be an astronaut is thwarted, until he decides to buy the genetic code of Jerome Morrow (Jude Law), a genetic superior rendered a paraplegic in an accident. Armed with blood and urine samples that Jerome dutifully supplies every day for testing, Vincent rises through the ranks of Gattaca Aerospace Corporation, a space-launch firm. Just a week before Vincent is set to depart to Titan, a mission direc-

tor is brutally beaten to death. The subsequent investigation—with its sweeps for hair, flakes of skin and other traces of DNA—places Vincent's masquerade at risk. The resulting action has Vincent and Jerome trying to elude the detectives; to complicate matters, Vincent falls for co-worker Irene (Uma Thurman), who is assigned to help the detectives.

The film, written and directed by New Zealander Andrew Niccol in his directorial debut, achieves a terrific, sleek look, with its curved, stainless-steel of-



DARREN MICHAELS/Columbia TriStar

to be called dreamers.



Ovonic battery-powered car. Photovoltaic roof panels. Rewritable optical memory disks. **Available now.**

For more on how Energy Conversion Devices, Inc. is helping to change the way the world stores and generates energy and information, call 1-248-280-1900.

Write to us at Energy Conversion Devices, Inc., 1675 West Maple Road, Troy, Michigan 48064. Or visit our Web site: www.ovonic.com



The world won't be
the same with us.

SCIENTIFIC AMERICAN

COMING IN THE
NOVEMBER ISSUE...



EDWARD BELL

MERCURY, OUR UNEXPLORED NEIGHBOR

by Robert Nelson



MAKING RICE DISEASE-PROOF

by Pamela Ronald

Also in November...

Computer Viruses

Parasitic Wasps

Solving Fermat's
Last Theorem

Preventing Accidental
Nuclear War

The Great Zimbabwe

ON SALE OCTOBER 28

fice cubicles, unnaturally clean, reverberating spaces and astronaut-wear that refreshingly consists of slick, double-breasted suits rather than Star Trek-style pajamas. Add in a regimented workforce and random genetic checks that would make the American Civil Liberties Union apoplectic, and *Gattaca* sets an irresistibly Orwellian mood. The visuals, though, are undercut by punchless storytelling and flat characterizations. Grave sacrifices barely elicit a shrug, and the romance between Vincent and Irene does not even reach tepid.

But *Gattaca* has the hallmark of the best science fiction: it makes you consider "what if," using scientific ideas that are fairly accurate and well explained. Will genetic information be used to suppress a population? We already have hints of that today, as civil libertarians worry that insurance companies will use genetic information to screen applicants. As *Gattaca* reminds us, biology is not destiny, and the environment that shapes our motivations and passions is crucial to a fulfilled life. Yet it is easy to get carried away every time a report identifies a gene that is associated with a particular human condition.

An equally unsettling point explored by the film is the possibility of tracking one's every movement, a problem resulting at present from credit-card purchases, e-mail and other forms of communications technology rather than genetic engineering. If anything, *Gattaca* emphasizes the need to remain vigilant about issues of privacy as technology makes it easier to intrude.

Unfortunately, *Gattaca* undermines some of its worthier points when it resorts to what can be construed as science bashing. After the movie's final frame, a textual message states that the human genome will be completely mapped in a few years and warns that had that knowledge been available, important figures in history might never have been born because they had genetic diseases. Does Hollywood really need to be reminded that it is not the science that is dangerous—it is what we as a society choose to do with that science. With all the obvious talent that went into making this slick movie, the filmmakers should have known better.

PHILIP YAM is news editor of SCIENTIFIC AMERICAN.

BRIEFLY NOTED

DARWIN AMONG THE MACHINES, by George B. Dyson. *Helix Books/Addison Wesley, Reading, Mass.*, 1997 (\$25). Machines, like living organisms, evolve. They do so, however, in a peculiar ecosystem consisting largely of human designers and users. People's hands and minds are essential to breeding and reproduction in the mechanical world. George B. Dyson, son of physicist Freeman Dyson, demonstrates the complexity of talking about nonbiological evolution. His book provides useful information, regrettably immersed amid relentless prose and often gratuitous quotations.

PLANET QUEST, by Ken Croswell. *Free Press, New York*, 1997 (\$25). **WORLDS UNNUMBERED**, by Donald Goldsmith. *University Science Books, Sausalito, Calif.*, 1997 (\$28.50). **THE QUEST FOR ALIEN PLANETS**, by Paul Halpern. *Plenum Publishing Corporation, New York*, 1997 (\$27.95).

Two years ago astronomers discovered the first planets circling sunlike stars, proving that our solar system is not unique. These books offer fine but distinctive introductions to this mind-opening discovery. Ken Croswell takes a charming, historical approach, beginning with Giordano Bruno's vision of a multitude of worlds and continuing through the personalities and techniques involved in the latest findings. Donald Goldsmith omits some of the background details in favor of lively discussions about how planets form and whether any of the new bodies could support life. And Paul Halpern steers a middle course, adding a short discussion about the possible connection between planets and dark matter.

CARTOGRAPHIES OF DANGER, by Mark Monmonier. *University of Chicago Press*, 1997 (\$25).

Maps are powerful tools for understanding the distribution of risks all around us—ranging from radon to burglaries to tornadoes. Crime plots enable police to target their manpower more effectively; fault-zone maps guide architects building in earthquake-prone areas. The author calls on cartographic techniques to show the (often misunderstood) complexities of environmental hazards in the U.S., giving rigorous response to the victim's cry of "Why me?"

WORKING KNOWLEDGE

FISH LADDERS

by William S. Rainey

*Senior Fish Passage Engineer,
Northwest Region,
National Marine Fisheries Service,
Portland, Ore.*

A single man-made barrier can exterminate all migratory fish trying to move up a river. Just one dam bigger than a few feet tall will block salmon, for example, from ancestral spawning grounds, eliminating all traces of the fish from upstream waters.

Historically a highly valued food source, salmon have been getting help. They have been raised to higher water levels in special fish locks, pumped into tank trucks and driven upstream and, most effectively, provided with fishways or ladders that surmount a dam in steps the fish can handle.

Early ladders were pools separated by obstructions over which the fish could swim or jump, each pool typically about a foot higher than the previous one. But at times the ladder had insufficient flow, and fish were not attracted to it. In other instances, excessive flows prevented the fish from moving up the ladder and over the dam.

One type of contemporary fish ladder places narrow 12- to 15-inch-wide slots in the partitions between the four- to eight-foot-deep pools, which create favorable flow conditions for the fish. The pools also

dissipate the kinetic energy of the stream and provide resting zones, thereby allowing fish to swim easily from pool to pool.

Radio tracking studies show that a well-designed fish ladder can transform a killer dam into one with minimal effect on upstream migration. Unfortunately, the other half of the story remains grim. The small flow from the ladder does not lure the juveniles headed downstream. Instead screens may divert the young fish away from hydroelectric turbines or from being waylaid into irrigation canals. But engineers and biologists face a continuing challenge in trying to keep the tiny fish from harm.

FISH LADDER is often a 10-percent-graded flume interrupted with vertical, slotted partitions. The maximum one-foot drop in water level at each partition produces a flow that the salmon instinctively pursue. Setting the slots at an angle directs the flow exclusively into pools behind the partitions, so the dropping water never has more energy than the fish can resist. Changes in water level do not disrupt the ladder. Higher water increases the flow through the slots as well as the amount of energy-absorbing water in the pools.

ILLUSTRATIONS BY JACK UNRUH

SUPPLEMENTAL FLOW

ENTRANCE POOL

ENTRY GATE

ENTICING FISH TO ENTER

a ladder demands a stronger flow of water than the slots of the "steps" alone can provide. To supplement the volume coming down the ladder, water from the top (blue arrows) is piped directly into the lowest pool—the entrance pool—producing a faster flow through the entry gate. Engineers try to position the gate where fish tend to congregate—immediately adjacent to turbulent zones at the base of the dam.



SOCKEYE SALMON